
Cours de Mathématiques

Classe de MPSI

Auteur : Marc LORENZI

Mis à jour le 3 juin 2025

Lycée Camille Guérin

Année 2024-2025

Table des matières

1	Logique	9
1	Propositions	10
2	Connecteurs logiques	11
3	Quantificateurs	17
2	Ensembles	21
1	Notion d'ensemble	22
2	Opérations sur les ensembles	24
3	Entiers naturels	31
1	L'ensemble des entiers naturels	32
2	Raisonnement par récurrence	32
3	Extensions du principe de récurrence	39
4	Suites définies par récurrence	41
5	Annexe : suites récurrentes	41
4	Applications et relations	43
1	Applications	44
2	Relations d'ordre	50
3	Relations d'équivalence	52
5	Nombres complexes	57
1	Introduction	58
2	Nombres complexes et trigonométrie	62
3	Résolution d'équations algébriques	70
4	Interprétations géométriques	75
6	Nombres réels	79
1	Le corps des réels	80
2	Borne supérieure	83

3	Intervalles	90
7	Suites réelles et complexes	93
1	Limite d'une suite	94
2	Théorèmes d'existence de limites	104
3	Suites complexes	110
4	Réurrences linéaires	112
8	Comparaison des suites	117
1	Suites équivalentes	118
2	Négligeabilité, domination	123
3	Puissances, exponentielles, logarithmes, factorielles	126
9	Groupes, anneaux, corps	129
1	Lois de composition interne	130
2	Groupes	132
3	Anneaux et corps	139
4	Espaces vectoriels	145
10	Fonctions	147
1	Opérations algébriques	148
2	Une relation d'ordre	148
3	Fonctions monotones	150
4	Extrema	150
5	Parité, périodicité	151
11	Limites - Continuité	155
1	Étude locale d'une fonction	156
2	Propriétés des fonctions continues	164
3	Continuité uniforme	169
4	Fonctions à valeurs complexes	171
12	Dérivation	173
1	Notion de dérivée	174
2	Accroissements finis	183
3	Fonctions à valeurs complexes	188
13	Fonctions Usuelles	191
1	Logarithmes et exponentielles	192
2	Puissances	197
3	Fonctions circulaires	200
4	Fonctions hyperboliques	204
14	Fonctions convexes	209
1	Intervalles de $\mathbb{R}$	210

2	Notion de fonction convexe	213
3	Fonctions convexes dérivables	217
4	Annexe - p -normes	222
15	Primitives et Équations différentielles	229
1	Calculs de primitives	230
2	Notion d'équation différentielle	236
3	Équations linéaires du premier ordre	238
4	Résolution approchée	242
5	Équations du second ordre	243
16	Arithmétique	249
1	Divisibilité dans $\mathbb{Z}$	250
2	PGCD, PPCM	252
3	Nombres premiers	263
4	Congruences	267
5	Annexe : les anneaux $\mathbb{Z}/n\mathbb{Z}$	268
17	Calcul matriciel	273
1	Notion de matrice	274
2	Systèmes linéaires	279
3	L'anneau des matrices carrées	285
4	matrices triangulaires	288
18	Polynômes	293
1	L'algèbre des polynômes	294
2	L'anneau des polynômes	299
3	Fonctions polynômes	300
4	Dérivation	303
5	Factorisation des polynômes	304
6	Pgcd, Ppcm	309
7	Interpolation de Lagrange	314
19	Fractions rationnelles	317
1	Le corps des fractions rationnelles	318
2	Éléments simples	320
3	Primitives des fractions rationnelles	326
20	Analyse asymptotique	329
1	Comparaison des fonctions au voisinage d'un point	330
2	Développements limités	335
3	Opérations sur les développements limités	342
4	Développements asymptotiques	349
21	Étude des Fonctions	355

1	Ensemble de définition	356
2	Réduction de l'ensemble d'étude	356
3	Régularité de la fonction	356
4	Variations	356
5	Points remarquables réels	357
6	Étude à l'infini	357
7	Concavité	358
8	Tracé	359
9	Un exemple complet	359
10	Un autre exemple	362
22	Espaces vectoriels	367
1	Généralités	368
2	Applications linéaires	370
3	Sous-espaces vectoriels	374
4	Opérations sur les s.e.v	377
5	Sous-espaces supplémentaires	379
6	Familles remarquables de vecteurs	382
23	Dimension finie	387
1	Dimension	388
2	Sev d'un espace de dimension finie	391
3	Dimensions classiques	392
4	Notion de rang	395
5	Dual d'un espace vectoriel	399
24	Matrices	405
1	Vecteurs, applications linéaires et matrices	406
2	Changements de base	411
3	Rang	412
4	Trace	415
5	Similitude	416
25	Groupes Symétriques	419
1	Compléments sur les groupes finis	420
2	Notion de permutation	422
3	Orbites	424
4	Signature	428
5	Annexe : Inversions	433
26	Déterminants	435
1	Applications multilinéaires	436
2	Déterminant d'une famille de vecteurs	438
3	Calcul des déterminants	443

4	Déterminant d'un endomorphisme	449
5	Inversion des matrices	454
27	Systèmes linéaires	457
1	Sous-espaces affines d'un espace vectoriel	458
2	Rappels sur les systèmes linéaires	464
3	Systèmes de Cramer	466
4	L'algorithme du Pivot de Gauss	468
5	Compléments	471
28	Intégration	475
1	Intégration des fonctions en escalier	476
2	Construction de l'intégrale	480
3	Propriétés de l'intégrale	484
4	Approximations de l'intégrale	490
5	Primitives	495
6	Formules de Taylor	499
7	Fonctions à valeurs complexes	501
29	Dénombrements	503
1	Notion d'ensemble fini	504
2	Opérations sur les ensembles finis	506
3	Applications entre ensembles finis	509
4	Parties d'un ensemble fini	511
5	Annexe : la formule du crible	515
30	Probabilités	519
1	Espaces probabilisables	520
2	Espaces probabilisés	523
3	Urnes	529
4	Probabilités conditionnelles	532
31	Variables aléatoires	541
1	Notion de variable aléatoire	542
2	Lois usuelles	545
3	Couples de variables aléatoires	547
4	Indépendance	551
32	Espérance, variance et covariance	557
1	Espérance	558
2	Variance	567
3	Covariance	574
33	Espaces préhilbertiens	581
1	Produits scalaires	582

2	Orthogonalité	587
3	Projecteurs orthogonaux - Symétries orthogonales	593
34	Endomorphismes orthogonaux	597
1	Généralités	598
2	Le groupe orthogonal du plan	603
3	Produit vectoriel	607
4	Le groupe orthogonal de l'espace	611
5	Compléments	615
35	Séries	621
1	Généralités	622
2	Séries à termes positifs	627
3	Comparaison entre séries et intégrales	630
4	Convergence absolue	635
5	Appendice - Représentations p-adiques des réels	639
36	Familles Sommables	647
1	Familles sommables de réels positifs	648
2	Familles sommables de nombres complexes	656
3	Produit de Cauchy	664
4	Complément - Le cardinal des familles sommables	666
37	Fonctions de deux variables	669
1	Un peu de topologie	670
2	Continuité	672
3	Dérivées partielles	675
4	Composition	680
	Index	685

Chapitre 1

Logique

1 Propositions

1.1 Le langage des mathématiques

Faire des mathématiques, c'est écrire des propositions vraies. Mais qu'est ce qu'une proposition ? Et qu'est ce que la vérité ? Ceci n'est pas un cours de logique, et nous nous contenterons de quelques idées très vagues sur la question. Une proposition est une assemblage de symboles, formé à partir de règles strictes. Par exemple, « $1 + 1 = 2$ » est une proposition. Un autre exemple de proposition est « tous les nombres entiers sont pairs », ou encore « tout corps fini est commutatif ». Où sont les symboles dans ces propositions ? Eh bien ils sont cachés dans les mots que le mathématicien a définis. Par exemple, « n est pair » veut dire « $\exists p \in \mathbb{N}, n = 2 \times p$ ». Et les symboles $\mathbb{N}$ et $\times$ sont eux mêmes des résumés pour des assemblages d'autres symboles encore plus primitifs. En réalité, le seul symbole primitif de la théorie dans laquelle nous faisons des mathématiques (la théorie des ensembles) est le symbole d'appartenance $\in$.

Étant donnée une proposition, celle-ci est vraie ou fausse, mais jamais les deux à la fois (nous ne chercherons pas à définir les mots VRAI et FAUX). La véracité ou la fausseté d'une proposition est ce que l'on appelle parfois sa *valeur de vérité*. L'activité principale d'un mathématicien consiste à écrire, en utilisant un certain nombre de règles, des propositions vraies. Nous conviendrons que, dorénavant, toutes les propositions que nous écrivons sont vraies. Ceci ressemble à une évidence, mais elle ne l'est pas complètement. En effet, nous pourrions dorénavant écrire $1 + 1 = 2$, au lieu de « la proposition « $1 + 1 = 2$ » est vraie ».

Nous serons cependant parfois assez souples sur l'emploi des guillemets. Par exemple, si P est une proposition nous n'écrirons pas « P » entre guillemets sous prétexte que nous savons pas si P est vraie ou fausse.

Le terme générique « proposition » convient pour toutes les phrases mathématiques vraies. On utilise parfois d'autres mots, pour insister sur l'objectif de la proposition. Citons entre autres :

- Théorème : une proposition d'une grande importance, parfois l'aboutissement d'années de recherches.
- Lemme : une proposition de moindre importance, souvent nécessaire à la démonstration d'un théorème. Cela ne signifie pas forcément qu'un lemme est simple à démontrer, ni qu'un lemme n'est pas important. Tout est relatif.
- Corollaire : une proposition qui est une conséquence (souvent) immédiate d'une proposition qui vient juste d'être établie.
- Axiome : une proposition arbitrairement décidée comme vraie dans la théorie mathématique dans laquelle on se place. Seuls les grands mathématiciens ont le droit de changer les axiomes des mathématiques.

1.2 Équivalence de deux propositions

Deux propositions peuvent avoir l'air très différentes, et être pourtant simultanément vraies ou simultanément fausses. C'est ce qui fait toute la difficulté des mathématiques

Définition 1. On dit que deux propositions P et Q sont équivalentes, et on note $P \iff Q$, lorsque P et Q sont simultanément vraies ou simultanément fausses.

Remarque 1. Comment montrer une équivalence ? Nous verrons que pour des propositions très simples, ce que l'on appelle une table de vérité suffit. Mais bien entendu, cela ne va pas durer. Une démonstration d'équivalence peut être très difficile. Nous allons en gros passer le reste de l'année à démontrer des équivalences.

Remarque 2. Étant données deux propositions P et Q , $P \iff Q$ est en fait une nouvelle proposition, soit vraie, soit fausse. Le symbole $\iff$ fait partie de la famille des connecteurs, dont nous allons maintenant étudier quelques exemples.

Remarquons que l'on a la proposition suivante.

Proposition 1.

- Soit P une proposition. Alors, $P \iff P$.
- Soient P et Q deux propositions. Si $P \iff Q$, alors $Q \iff P$.
- Soient P, Q, R trois propositions. Si $P \iff Q$ et $Q \iff R$, alors $P \iff R$.

Démonstration. Nous réservons la preuve de cette proposition pour un peu plus tard.
□

2 Connecteurs logiques

Le discours mathématique a tendance à relier, à connecter des propositions entre-elles, ceci afin de fabriquer de nouvelles propositions. La question est de savoir quand ces nouvelles propositions sont vraies.

2.1 Le connecteur « ET »

Définition 2. Étant données deux propositions P et Q , le connecteur $\wedge$ (ET) permet de fabriquer la nouvelle proposition $P \wedge Q$. Cette nouvelle proposition est vraie exactement lorsque les deux propositions P et Q sont vraies.

Voici la table de vérité du « ET ».

P	Q	$P \wedge Q$
0	0	0
0	1	0
1	0	0
1	1	1

FAUX et VRAI sont représentés par 0 et 1. Par exemple, la ligne 0 1 0 nous dit que $P \wedge Q$ est FAUX lorsque P est FAUX et Q est VRAI.

Proposition 2. Soient P, Q, R des propositions. On a

- $P \wedge P \iff P$ (idempotence)
- $P \wedge Q \iff Q \wedge P$ (commutativité)
- $(P \wedge Q) \wedge R \iff P \wedge (Q \wedge R)$ (associativité)

Démonstration. Ceci se prouve en faisant des *tables de vérité*. Montrons par exemple l'associativité.

P	Q	R	$P \wedge Q$	$Q \wedge R$	$(P \wedge Q) \wedge R$	$P \wedge (Q \wedge R)$
0	0	0	0	0	0	0
0	0	1	0	0	0	0
0	1	0	0	0	0	0
0	1	1	0	1	0	0
1	0	0	0	0	0	0
1	0	1	0	0	0	0
1	1	0	1	0	0	0
1	1	1	1	1	1	1

On constate que les deux dernières colonnes de ce tableau sont identiques, ce qui veut dire que les propositions correspondantes sont équivalentes. $\square$

Remarque 3. En tant que mathématicien, le reste de votre vie va être consacré à démontrer des propositions. La rédaction d'une démonstration est un acte parfaitement codifié. Pour chaque type de proposition, il existe un ou plusieurs types de démonstrations. Il s'agit pour vous de parfaitement connaître la façon dont on démontre chaque type de proposition, puisque votre vie d'étudiant en mathématiques consistera à être noté sur vos démonstrations.

Comment montrer un « ET » ? Eh bien, on montre d'abord l'une des deux propositions, PUIS on montre l'autre. Histoire d'être bien clair, pour montrer un « ET », il faut faire deux démonstrations distinctes.

2.2 Le connecteur « OU »

Définition 3. Étant données deux propositions P et Q , le connecteur $\vee$ (OU) permet de fabriquer la nouvelle proposition $P \vee Q$. Cette nouvelle proposition est vraie exactement lorsque l'une au moins des deux propositions P et Q est vraie.

Remarque 4. Pour être précis, nous venons de définir ici ce que l'on appelle le *ou inclusif*. Il existe un autre type de « ou », que l'on appelle le *ou exclusif* et que l'on nomme aussi « ou bien ».

Voici la table de vérité du « OU » (inclusif).

P	Q	$P \vee Q$
0	0	0
0	1	1
1	0	1
1	1	1

Pour le ou exclusif, la dernière ligne de la table serait 1 1 0.

Remarque 5. Comment montrer un « OU » ? Ce n'est pas si évident que cela. Il va falloir attendre de connaître l'implication.

Voici quelques propriétés du connecteur « OU » :

Proposition 3. Soient P, Q, R des propositions. On a

- $P \vee P \iff P$
- $P \vee Q \iff Q \vee P$
- $(P \vee Q) \vee R \iff P \vee (Q \vee R)$

Proposition 4. Soient P, Q, R des propositions. On a

- $P \vee (Q \wedge R) \iff (P \vee Q) \wedge (P \vee R)$ (distributivité du OU par rapport au ET)
- $P \wedge (Q \vee R) \iff (P \wedge Q) \vee (P \wedge R)$ (distributivité du ET par rapport au OU)

Démonstration. Ces deux propositions se prouvent en faisant des *tables de vérité*. $\square$

2.3 Le connecteur « NON »

Définition 4. Étant donnée une proposition P , le connecteur $\neg$ (NON) permet de fabriquer la nouvelle proposition $\neg P$. Cette nouvelle proposition est vraie exactement lorsque la proposition P est fausse.

Voici la table de vérité du « NON ».

P	$\neg P$
0	1
1	0

Proposition 5. Soient P, Q des propositions. On a

- $\neg\neg P \iff P$
- $\neg(P \wedge \neg P)$ (non contradiction)
- $P \vee \neg P$ (tiers exclu)
- $\neg(P \wedge Q) \iff \neg P \vee \neg Q$
- $\neg(P \vee Q) \iff \neg P \wedge \neg Q$ (lois de Morgan)

Démonstration. Encore une fois, des tables de vérité permettent de tout montrer. Mais histoire de voir que nous avons avancé un peu, voici une autre preuve du dernier point. Écrivons l'avant dernier point pour les propositions $\neg P$ et $\neg Q$:

$$\neg(\neg P \wedge \neg Q) \iff (\neg\neg P) \vee (\neg\neg Q)$$

Par le premier point, on a donc

$$\neg(\neg P \wedge \neg Q) \iff P \vee Q$$

Si deux propositions sont équivalentes, il est clair que leurs négations le sont aussi. Ainsi,

$$\neg\neg(\neg P \wedge \neg Q) \iff \neg(P \vee Q)$$

et à nouveau par le premier point,

$$\neg P \wedge \neg Q \iff \neg(P \vee Q)$$

2.4 Le connecteur « IMPLIQUE »

□

Définition 5. Étant données deux propositions P et Q , on crée une nouvelle proposition $P \implies Q$. Cette proposition est fausse exactement dans le cas où P est vraie et Q est fausse.

Voici la table de vérité de « IMPLIQUE ».

P	Q	$P \implies Q$
0	0	1
0	1	1
1	0	0
1	1	1

Proposition 6. Soient P, Q, R des propositions. On a

- $((P \implies Q) \wedge (Q \implies R)) \implies (P \implies R)$ (transitivité de l'implication)
- $(P \implies Q) \iff (\neg P \vee Q)$
- $(P \vee Q) \iff (\neg P \implies Q)$
- $\neg(P \implies Q) \iff (P \wedge \neg Q)$

Exemple. L'exemple qui suit est destiné à nous faire comprendre la nature du connecteur implique. Donnons nous un entier naturel n . Je vais partir du principe que tout le monde sera d'accord lorsque j'affirme que si n est multiple de 4, alors n est pair. Prenons quelques valeurs de n pour voir ...

Commençons par $n = 8$. Alors n est multiple de 4 et n est également pair. Donc, si l'on veut que la table de vérité du connecteur IMPLIQUE ne dépende que des valeurs de vérité des propositions, et pas de la nature de ces propositions, on doit avoir que VRAI implique VRAI.

Maintenant, attention ! Prenons $n = 2$. Alors, n n'est pas multiple de 4, et pourtant il est pair. Donc, FAUX implique VRAI.

Pour finir, prenons $n = 3$. Cette fois ci, FAUX implique FAUX.

Un seul cas n'a pas été envisagé : a-t-on aussi VRAI implique FAUX ? Si c'était le cas, toutes les implications seraient vraies. Donc, si l'on veut que le connecteur IMPLIQUE ne soit pas un connecteur trivial, on doit poser que VRAI implique FAUX est FAUX.

Remarque 6. Comment montrer que $P \implies Q$? La proposition P est soit vraie soit fausse. Si P est fausse, alors $P \implies Q$ sera vraie, ceci quelle que soit Q . On n'a donc pas à s'occuper de ce cas, et on peut donc supposer que P est vraie. Puis, il s'agit de vérifier que Q est vraie. Pour résumer :

Pour montrer $P \implies Q$:

- On suppose P .
- On montre Q .

S'il y a une seule chose à retenir dans ce chapitre, c'est celle-ci. Nous en ferons un usage quotidien.

Remarque 7. On a vu qu'il existe un lien entre « OU » et « IMPLIQUE ». Ainsi : Pour montrer $P \vee Q$,

- On suppose que P est fausse.
- On montre Q .

Proposition 7. PRINCIPE DE CONTRAPOSITION

Soient P et Q deux propositions. On a

$$P \implies Q \iff \neg Q \implies \neg P$$

Ce principe nous donne une seconde méthode pour prouver une implication.

- On suppose que le membre de droite est faux
- On montre que celui de gauche l'est aussi.

Proposition 8. PRINCIPE DE DÉMONSTRATION PAR L'ABSURDE

Soit P une proposition. Soit F une proposition fausse. On a

$$(\neg P \implies F) \implies P$$

Ce principe nous dit que pour montrer que P est vraie, on *peut* procéder comme suit :

- On suppose que P est fausse.
- On arrive à une contradiction.

Il ne faut pas abuser de ce genre de démonstrations, car elles sont souvent difficiles à lire (pour ne pas dire incompréhensibles), et paraissent artificielles. Moralité, on ne fait une démonstration par l'absurde qu'après que les autres techniques de démonstration aient échoué, *ou alors parce qu'on a une bonne raison pour le faire.*

2.5 Équivalence logique

Nous revenons ici à l'étude du connecteur d'équivalence. Rappelons la

Définition 6. Étant données deux propositions P et Q , on crée une nouvelle proposition $P \iff Q$. Cette proposition est vraie exactement dans le cas où P et Q ont la même valeur de vérité.

Voici la table de vérité de « EQUIVAUT ».

P	Q	$P \iff Q$
0	0	1
0	1	0
1	0	0
1	1	1

Nous sommes maintenant en mesure de prouver la proposition 1. Le lecteur est invité à faire une table de vérité.

Proposition 9. *Soient P et Q deux propositions. Alors,*

$$(P \iff Q) \iff ((P \implies Q) \wedge (Q \implies P))$$

Ainsi, pour montrer que P et Q sont équivalentes,

- On montre $P \implies Q$.
- On montre $Q \implies P$.

Pour montrer une équivalence il faut donc faire **DEUX** démonstrations.

3 Quantificateurs

3.1 Quantificateurs universel et existentiel

Soit $P(x)$ une proposition « dépendant » d'une variable x (on dit aussi un *prédicat*), cette variable appartenant à un certain ensemble E .

Définition 7. Le symbole $\forall$ est appelé quantificateur universel. Il permet de former la nouvelle proposition

$$\forall x \in E, P(x)$$

Cette proposition est vraie si et seulement si la proposition $P(x)$ est vraie pour tous les éléments x de l'ensemble E .

Définition 8. Le symbole $\exists$ est appelé quantificateur existentiel. Il permet de former la nouvelle proposition

$$\exists x \in E, P(x)$$

Cette proposition est vraie si et seulement si la proposition $P(x)$ est vraie pour au moins un élément x de l'ensemble E .

Remarque 8. Lorsque la proposition $P(x)$ est vraie pour exactement un élément x de l'ensemble E , on note

$$\exists! x \in E, P(x)$$

Nous reviendrons sans arrêt tout au long de l'année sur ces deux questions de *l'existence* et de *l'unicité*.

3.2 Négation d'une phrase quantifiée

Proposition 10. Soit $P(x)$ un prédicat, la variable x appartenant à un ensemble E . On a

- $\neg(\exists x \in E, P(x)) \iff (\forall x \in E, \neg P(x))$
- $\neg(\forall x \in E, P(x)) \iff (\exists x \in E, \neg P(x))$

Démonstration. Supposons $\neg(\exists x \in E, P(x))$. Soit $x \in E$. D'après l'hypothèse, on n'a pas $P(x)$. Donc on a $\neg P(x)$.

Inversement, supposons $\forall x \in E, \neg P(x)$. Tous les $x \in E$ vérifient $\neg P(x)$, ce qui veut dire qu'aucun $x \in E$ ne vérifie $P(x)$. $\square$

Exercice. Prouver la deuxième équivalence en prenant la négation de la première équivalence.

3.3 Échange de quantificateurs

Il arrive fréquemment que l'on ait à écrire des propositions comportant une succession de quantificateurs. Soit $P(x, y)$ une proposition dépendant de deux variables $x \in E$ et $y \in F$. Considérons les propositions

$$A = \forall x \in E, \forall y \in F, P(x, y)$$

et

$$B = \forall y \in F, \forall x \in E, P(x, y)$$

Ces deux propositions sont équivalentes, et sont aussi équivalentes à

$$C = \forall (x, y) \in E \times F, P(x, y)$$

(voir plus loin pour le produit cartésien). Lorsque $E = F$, on commet l'abus d'écrire

$$\forall x, y \in E, P(x, y)$$

Il en va de même si l'on met des quantificateurs existentiels à la place des quantificateurs universels.

Considérons maintenant les deux propositions

- $A = \forall x \in E, \exists y \in F, P(x, y)$
- $B = \exists y \in F, \forall x \in E, P(x, y)$

Proposition 11. *On a $B \implies A$. La réciproque est fausse en général.*

Démonstration. Supposons B . Il existe donc un élément y_0 de F tel que que $\forall x \in E, P(x, y_0)$. Soit maintenant $x \in E$. On a $P(x, y_0)$ donc $\exists y \in F, P(x, y)$ (prendre $y = y_0$). On a montré A .

Pour voir que la réciproque est fausse, prenons $E = F = \mathbb{N}$ et

$$P(x, y) \iff y = x + 1$$

On a bien A :

$$\forall x \in \mathbb{N}, \exists y \in \mathbb{N}, y = x + 1$$

On a même $\exists! y \in \mathbb{N}, y = x + 1$.

En revanche, B est fausse (aucun entier n'est le successeur de TOUS les entiers). Le puriste voudra une démonstration, alors raisonnons par l'absurde. Supposons donc

$$\exists y \in \mathbb{N}, \forall x \in \mathbb{N}, y = x + 1$$

Prenons un tel y et appelons-le y_0 . On a donc

$$\forall x \in \mathbb{N}, y_0 = x + 1$$

En particulier, comme $y_0 \in \mathbb{N}$,

$$y_0 = y_0 + 1$$

et donc, en soustrayant y_0 ,

$$0 = 1$$

Nous verrons dans le chapitre sur $\mathbb{N}$ que ceci est faux :-). $\square$

Chapitre 2

Ensembles

1 Notion d'ensemble

1.1 Ensembles, éléments

Un ensemble est, naïvement, une collection d'objets unis par une propriété commune. Un objet appartient à l'ensemble lorsqu'il possède cette propriété. Nous considérons la notion d'ensemble comme une notion « première », que nous ne chercherons pas à définir.

Notation. Étant donné un ensemble E et un objet x , on écrira $x \in E$ lorsque x est élément de E (ou appartient à E), et $x \notin E$ dans le cas contraire.

On peut décrire un ensemble de plusieurs façons, les plus courantes étant les suivantes :

- En listant ses éléments : $A = \{1, 2, 3\}$ ou $B = \{0, 1, 4, 9, 16 \dots\}$.
- En donnant une propriété caractérisant ses éléments : B est l'ensemble des nombres pairs, ce que l'on peut écrire $B = \{n \in \mathbb{N}, \exists p \in \mathbb{N}, n = p^2\}$.
- En regardant ses éléments comme les images des éléments d'un autre ensemble par une fonction : $B = \{n^2, n \in \mathbb{N}\}$.

Nous reviendrons sur ces descriptions dans la suite du chapitre.

1.2 Les difficultés du concept d'ensemble

Le concept naïf d'ensemble mène assez facilement à des contradictions. Ces contradictions, mises en évidence vers la fin du XIXe siècle, ont conduit les mathématiciens à concevoir la *Théorie des ensembles*. Donnons un exemple de telle contradiction :

Étant donné un ensemble X , la question de se demander si $X \in X$ a un sens. Par exemple, si $X = \{1, 2, 3\}$, alors $X \notin X$. En effet, X a trois éléments qui sont 1, 2 et 3. Et aucun de ces éléments n'est égal à $\{1, 2, 3\}$! En fait, il est extrêmement difficile de trouver un ensemble qui appartienne à lui-même (en fait l'un des axiomes de la théorie des ensembles « classique » dit que c'est impossible).

Considérons l'ensemble de tous les ensembles qui n'appartiennent pas à eux-mêmes :

$$\mathcal{E} = \{X, X \notin X\}$$

On a donc, pour tout ensemble X :

$$X \in \mathcal{E} \iff X \notin X$$

Posons-nous la question suivante : $\mathcal{E}$ appartient-il à $\mathcal{E}$? Remplaçant X par $\mathcal{E}$ dans la propriété ci-dessus, on trouve $\mathcal{E} \in \mathcal{E} \iff \mathcal{E} \notin \mathcal{E}$, ce qui est clairement contradictoire. Conclusion : $\mathcal{E}$ n'est pas un ensemble.

Il est bon de savoir qu'un tel souci ne peut arriver à condition de prendre des précautions :

Axiome 1. Soit E un ensemble. Soit $\mathcal{P}$ une propriété. Il existe alors un ensemble A tel que, pour tout objet x , $x \in A \iff x \in E$ et $\mathcal{P}(x)$. Cet ensemble A est noté

$$A = \{x \in E, \mathcal{P}(x)\}$$

Remarque 1. L'ensemble de tous les ensembles n'existe pas. En effet, supposons son existence. Appelons le $\mathcal{U}$. Considérons la propriété

$$\mathcal{P}(X) = (X \notin X)$$

Alors, par l'axiome 1, l'ensemble

$$E = \{X \in \mathcal{U}, X \notin X\}$$

existe. Contradiction, par le paradoxe de Russell.

1.3 Égalité d'ensembles

Axiome 2. Deux ensembles sont égaux si et seulement si ils ont les mêmes éléments.

Remarque 2. Ainsi, soient A et B deux ensembles. Pour montrer que $A = B$, on prouve les deux propriétés :

- $\forall x \in A, x \in B$.
- $\forall x \in B, x \in A$.

1.4 Inclusion

Définition 1. Soient A et B deux ensembles. On dit que A est inclus dans B , et on note $A \subset B$, lorsque

$$\forall x \in A, x \in B$$

Proposition 3. On a, pour tous ensembles A , B et C :

- $A \subset A$ (réflexivité)
- $A \subset B$ et $B \subset A \implies A = B$ (antisymétrie)
- $A \subset B$ et $B \subset C \implies A \subset C$ (transitivité)

Remarque 3. La relation d'inclusion est ce que l'on appelle une *relation d'ordre*. Nous reviendrons sur les relations d'ordre dans un chapitre ultérieur.

Démonstration. La réflexivité est évidente. L'antisymétrie résulte de l'axiome sur l'égalité de deux ensembles. Supposons maintenant $A \subset B$ et $B \subset C$. Soit $x \in A$. On a $x \in B$ puisque $A \subset B$. Donc $x \in C$ puisque $B \subset C$. Ainsi, $A \subset C$. $\square$

1.5 Ensemble vide

Axiome 4. Il existe un ensemble qui n'a pas d'élément. On l'appelle l'ensemble vide, et on le note $\emptyset$ (ou $\{\}$).

Remarque 4. Il existe un unique ensemble vide. Soient en effet $\emptyset_1$ et $\emptyset_2$ deux ensembles vides. Alors, pour tout $x \in \emptyset_1$, on a $x \in \emptyset_2$. Donc, $\emptyset_1 \subset \emptyset_2$. De même, $\emptyset_2 \subset \emptyset_1$. D'où $\emptyset_1 = \emptyset_2$.

Remarque 5. Attention, $\{\emptyset\}$ ne désigne pas l'ensemble vide, mais un ensemble qui possède un unique élément, cet élément étant l'ensemble vide.

1.6 Parties d'un ensemble

Axiome 5. Soit E un ensemble. Il existe un ensemble, noté $\mathcal{P}(E)$, tel que pour tout objet X ,

$$X \in \mathcal{P}(E) \iff X \subset E$$

L'ensemble $\mathcal{P}(E)$ est l'ensemble des *parties* de E .

Exemple.

- Soit $E = \{1, 2\}$. Alors, $\mathcal{P}(E) = \{\emptyset, \{1\}, \{2\}, \{1, 2\}\}$.
- Soit $E = \{a, b, c\}$. Alors, $\mathcal{P}(E) = \{\emptyset, \{a\}, \{b\}, \{c\}, \{a, b\}, \{a, c\}, \{b, c\}, \{a, b, c\}\}$.

Exemple.

- $\mathcal{P}(\emptyset) = \{\emptyset\}$. L'ensemble $\mathcal{P}(\emptyset)$ possède donc un élément (cet élément est l'ensemble vide).
- $\mathcal{P}(\mathcal{P}(\emptyset)) = \{\emptyset, \{\emptyset\}\}$.
- $\mathcal{P}(\mathcal{P}(\mathcal{P}(\emptyset))) = \{\emptyset, \{\emptyset\}, \{\{\emptyset\}\}, \{\emptyset, \{\emptyset\}\}\}$.

Exercice. Calculer $\mathcal{P}(\mathcal{P}(\mathcal{P}(\mathcal{P}(\emptyset))))$.

2 Opérations sur les ensembles

2.1 Réunion

Définition 2. Soient A et B deux ensembles. On appelle réunion de A et B l'ensemble

$$A \cup B = \{x, x \in A \text{ ou } x \in B\}$$

Remarque 6. L'un des axiomes de la théorie des ensembles affirme que $A \cup B$ est bien un ensemble.

Définition 3. Plus généralement, étant donnée une *famille* d'ensembles $(A_i)_{i \in I}$ (c'est à dire des ensembles A_i indexés par un indice i « variant » lui-même dans un ensemble I), on appelle réunion des A_i l'ensemble

$$\bigcup_{i \in I} A_i = \{x, \exists i \in I, x \in A_i\}$$

Exemple. On a

$$\bigcup_{n \geq 1} \left[0, 1 - \frac{1}{n}\right] = [0, 1[$$

Posons $A = \bigcup_{n \geq 1} [0, 1 - \frac{1}{n}]$.

Soit $x \in A$. Il existe $n \geq 1$ tel que $0 \leq x \leq 1 - \frac{1}{n}$. Or, $1 - \frac{1}{n} < 1$, donc $0 \leq x < 1$. Donc $x \in [0, 1[$.

Inversement, soit $x \in [0, 1[$. Soit n un entier non nul vérifiant $n \geq \frac{1}{1-x}$. Alors, $x \geq 1 - \frac{1}{n}$ et donc $x \in [0, 1 - \frac{1}{n}]$. Le réel x est donc un élément de la réunion.

Proposition 6. Pour tous ensembles A, B, C, D ,

- $A \cup B = B \cup A$ (commutativité)
- $(A \cup B) \cup C = A \cup (B \cup C)$ (associativité)
- $A \cup \emptyset = A$ ($\emptyset$ est élément neutre pour la réunion)
- $A \cup A = A$ (idempotence)
- $A \cup B = A \iff B \subset A$.
- $A \subset B$ et $C \subset D \implies A \cup C \subset B \cup D$.

Démonstration. Montrons la commutativité. Soit x un objet. On a $x \in A \cup B$ si et seulement si $x \in A$ OU $x \in B$. Mais OU est commutatif. Ceci équivaut donc à $x \in B$ OU $x \in A$, c'est à dire à $x \in B \cup A$. Les autres propriétés sont laissées en exercice. $\square$

2.2 Intersection

Définition 4. Soient A et B deux ensembles. On appelle intersection de A et B l'ensemble

$$A \cap B = \{x, x \in A \text{ et } x \in B\}$$

Remarque 7. Pas besoin d'axiome pour affirmer que $A \cap B$ est un ensemble, puisque $A \cap B = \{x \in A, x \in B\}$.

Définition 5. Plus, généralement, étant donnée une famille d'ensembles $(A_i)_{i \in I}$, on appelle intersection des A_i l'ensemble

$$\bigcap_{i \in I} A_i = \{x, \forall i \in I, x \in A_i\}$$

Exercice. Montrer que $\bigcap_{n \geq 1} \left[0, 1 + \frac{1}{n}\right] = [0, 1]$.

Proposition 7. Pour tous ensembles A, B, C, D ,

- $A \cap B = B \cap A$.
- $(A \cap B) \cap C = A \cap (B \cap C)$.
- $A \cap \emptyset = \emptyset$.
- $A \cap A = A$.
- $A \cap B = A \iff A \subset B$.
- $A \subset B \text{ et } C \subset D \implies A \cap C \subset B \cap D$.

Démonstration. Exercice. $\square$

Proposition 8. Pour tous ensembles A, B, C ,

- $A \cap (B \cup C) = (A \cap B) \cup (A \cap C)$ (distributivité de $\cap$ par rapport à $\cup$)
- $A \cup (B \cap C) = (A \cup B) \cap (A \cup C)$ (distributivité de $\cup$ par rapport à $\cap$)

Démonstration. Soit x un objet. On a $x \in A \cap (B \cup C)$ si et seulement si $(x \in A) \wedge ((x \in B) \vee (x \in C))$. On termine en utilisant la distributivité de $\wedge$ par rapport à $\vee$. $\square$

2.3 Différence, Complémentaire

Définition 6. Soient A et B deux ensembles. On appelle différence de A et B l'ensemble

$$A \setminus B = \{x \in A, x \notin B\}$$

Notation. Lorsque $B \subset A$, l'ensemble $A \setminus B$ est appelé le *complémentaire* de B (dans A). Lorsqu'il n'y a pas d'ambiguïté sur A , on le note aussi B^c ou $\overline{B}$.

Proposition 9. LOIS DE MORGAN

Pour tous ensembles A, B, C ,

- $A \setminus (B \cup C) = (A \setminus B) \cap (A \setminus C)$.
- $A \setminus (B \cap C) = (A \setminus B) \cup (A \setminus C)$.

Démonstration. Soit x un objet. On a $x \in A \setminus (B \cup C)$ si et seulement si $(x \in A) \wedge \neg(x \in B \vee x \in C)$. On termine en utilisant les lois de Morgan du chapitre de logique. $\square$

2.4 Couples, produit cartésien

De façon vague, un couple (a, b) est composé de deux objets a et b . Ce qui est important, c'est que deux couples (a, b) et (c, d) sont égaux si et seulement si $a = c$ et $b = d$. Il y a plusieurs façons de définir le concept de couple. Peu importe comment on définit les couples, ce qui importe c'est à quelle condition deux couples sont égaux. Nous allons ci-dessous donner une construction possible, due à von Neumann. Puis nous l'oublierons.

Définition 7. Soient a et b deux objets. On appelle couple de première composante a et de seconde composante b l'ensemble

$$(a, b) = \{\{a\}, \{a, b\}\}$$

Proposition 10. Deux couples (a, b) et (c, d) sont égaux si et seulement si $a = c$ et $b = d$.

Démonstration. Dans un sens, c'est évident. Supposons maintenant que $(a, b) = (c, d)$, c'est à dire

$$(1) \quad \{\{a\}, \{a, b\}\} = \{\{c\}, \{c, d\}\}$$

L'ensemble $\{a, b\}$ est un élément de l'ensemble $\{\{c\}, \{c, d\}\}$. Deux cas se présentent.

- Cas 1, $\{a, b\} = \{c\}$. Cela signifie que $a = b = c$. (1) devient

$$(1) \quad \{\{a\}\} = \{\{a\}, \{a, d\}\}$$

On a donc $\{a, d\} \in \{\{a\}\}$, donc $\{a, d\} = \{a\}$, d'où $a = d$. Ainsi, $a = b = c = d$.

- Cas 2, $\{a, b\} \neq \{c\}$. Comme $\{c\} \in \{\{a\}, \{a, b\}\}$, on a donc $\{c\} = \{a\}$, d'où $c = a$. L'hypothèse faite nous dit donc que

$$\{a, b\} \neq \{a\}$$

et donc $b \neq a$. De plus, par (1),

$$\{a, b\} \in \{\{c\}, \{c, d\}\} = \{\{a\}, \{a, d\}\}$$

Il en résulte que

$$\{a, b\} = \{a, d\}$$

On a donc $b \in \{a, d\}$, donc $b = a$ ou $b = d$. Comme $b \neq a$, $b = d$.

□

Remarque 8. Il faut maintenant s'empresser d'oublier la définition des couples pour ne retenir que leur propriété caractéristique !

Remarque 9. On peut généraliser cette notation et considérer des triplets (a, b, c) , des quadruplets (a, b, c, d) , ..., des n -uplets $(a_1, a_2, \dots, a_n)$.

Définition 8. Soient A et B deux ensembles. On appelle produit cartésien de A et B l'ensemble

$$A \times B = \{(a, b), a \in A, b \in B\}$$

On généralise au produit cartésien de 3, 4, ... ensembles. Si A est un ensemble et $n \geq 1$, on note A^n le produit cartésien de A n fois avec lui-même.

$$A^n = A \times A \times \dots \times A$$

2.5 Familles

Soit E un ensemble. Le paragraphe précédent nous permet de parler d'un couple d'éléments de E . On peut définir de la même façon la notion de triplet, quadruplet, ..., de n -uplet d'éléments de E . On peut aller encore plus loin et définir une *famille* d'éléments de E .

La construction de la notion de famille n'est pas difficile, mais elle fait appel au concept d'application, que nous verrons plus loin (en fait une famille est une application, ou plutôt, pour les puristes, le *graphe* d'une application). Nous dirons simplement qu'étant donné un ensemble I , on peut considérer des objets $\mathcal{F}$ que nous noterons

$$\mathcal{F} = (x_i)_{i \in I}$$

et qui sont tels que pour tout $i \in I$, $x_i \in E$. De plus, et c'est là la propriété importante,

$$(x_i)_{i \in I} = (y_i)_{i \in I} \iff \forall i \in I, x_i = y_i$$

Nous noterons E^I l'ensemble des familles d'éléments de E indexées par l'ensemble I .

En prenant $I = 1, 2$ on retrouve la notion de couple.

En prenant $I = \mathbb{N}$, on obtient la notion de suite. Ainsi, $\mathbb{R}^{\mathbb{N}}$ est l'ensemble des suites réelles.

Rien ne nous empêche (et nous le ferons) de prendre $I = \mathbb{R}$ ou $I = \mathbb{C}$, ou tout ce qui nous fera plaisir.

Chapitre 3

Entiers naturels

1 L'ensemble des entiers naturels

On admet les principales propriétés des entiers naturels : l'ensemble $\mathbb{N} = \{0, 1, 2, \dots\}$ est un ensemble infini, muni d'une addition et d'une multiplication possédant les propriétés bien connues de tous. Cet ensemble est muni de deux relations d'ordre :

- L'ordre naturel : $a \leq b$ lorsqu'il existe un entier naturel c tel que $b = a + c$. Cet ordre est total, c'est à dire que deux entiers sont toujours comparables. On note alors $b - a$ l'unique entier c ainsi défini.
- La divisibilité : a divise b (on note $a \mid b$) lorsqu'il existe un entier naturel c tel que $b = ac$. Cet ordre est partiel : par exemple, $2 \nmid 3$ et $3 \nmid 2$. Si $a \neq 0$, l'entier c ainsi défini est unique, et on le note $c = \frac{b}{a}$.

Une propriété des entiers naturels fréquemment utilisée est la suivante.

Proposition 1. Soient $a, b \in \mathbb{N}$. On a

$$a < b \iff a + 1 \leq b \iff a \leq b - 1$$

2 Raisonnement par récurrence

2.1 Le théorème de démonstration par récurrence

Axiome 2. Toute partie non vide de $\mathbb{N}$ possède un plus petit élément.

Proposition 3. Toute partie non vide ET majorée de $\mathbb{N}$ possède un plus grand élément.

Démonstration. Soit A une partie de $\mathbb{N}$ non vide et majorée. Considérons l'ensemble B des majorants de A . L'ensemble B est une partie non vide de $\mathbb{N}$, et possède donc un plus petit élément M . Montrons que M est le plus grand élément de A .

- Traitons tout d'abord le cas où $M = 0$: dans ce cas, $A = \{0\}$ et possède donc un plus grand élément, qui est 0.
- Supposons maintenant $M \geq 1$. $M - 1$ ne majore pas A , puisque M est le plus petit majorant de A . Donc, il existe $a \in A$ tel que $M - 1 < a \leq M$. D'où $a = M$ et donc $M \in A$.

□

Proposition 4. Soit A une partie de $\mathbb{N}$. On suppose :

- $0 \in A$.
- $\forall n \in \mathbb{N}, n \in A \implies n + 1 \in A$.

Alors, $A = \mathbb{N}$.

Démonstration. Supposons que $A \neq \mathbb{N}$, et considérons $B = \mathbb{N} \setminus A$. L'ensemble B est une partie non vide de $\mathbb{N}$, et admet donc un plus petit élément n_0 .

On a $n_0 \neq 0$, puisque $n_0 \in A$. Considérons alors $m = n_0 - 1$. Cet entier n'est pas dans B . Donc il est dans A , donc $n_0 = m + 1 \in A$. Contradiction. $\square$

Proposition 5. Soit $\mathcal{P}(n)$ une propriété dépendant de l'entier n . On suppose

- $\mathcal{P}(0)$
- $\forall n \in \mathbb{N}, \mathcal{P}(n) \implies \mathcal{P}(n + 1)$

Alors $\forall n \in \mathbb{N}, \mathcal{P}(n)$.

Démonstration. On considère $A = \{n \in \mathbb{N}, \mathcal{P}(n)\}$ et on applique la proposition précédente. $\square$

Remarque 1. On peut bien sûr « commencer la récurrence » à un rang différent de 0. Il faut alors adapter la conclusion du théorème.

2.2 Exemples

Les exemples ci-dessous sont essentiels. Ils seront utilisés tout au long de l'année.

Proposition 6. Pour tout $n \in \mathbb{N}$,

$$\sum_{k=0}^n k = \frac{n(n+1)}{2}$$

Démonstration. Posons, pour tout $n \in \mathbb{N}$,

$$S_n = \sum_{k=0}^n k$$

Ceci est notre première récurrence, nous allons particulièrement soigner sa rédaction.

Tout d'abord on annonce ce que l'on va montrer.

Montrons par récurrence sur n que pour tout $n \in \mathbb{N}$, $S_n = \frac{n(n+1)}{2}$.

On montre ensuite que les hypothèses du théorème de démonstration par récurrence sont vérifiées.

- On a $S_0 = 0 = \frac{0(0+1)}{2}$.
- Soit $n \in \mathbb{N}$. Supposons $S_n = \frac{n(n+1)}{2}$. On a alors

$$\begin{aligned}
 S_{n+1} &= \sum_{k=0}^{n+1} k \\
 &= \sum_{k=0}^n k + (n+1) \\
 &= S_n + (n+1) \\
 &= \frac{n(n+1)}{2} + (n+1) \quad (\text{HR}) \\
 &= \frac{(n+1)(n+2)}{2}
 \end{aligned}$$

Enfin, on conclut.

Par le théorème de démonstration par récurrence, la propriété est établie. $\square$

Proposition 7. Pour tout $n \in \mathbb{N}$,

$$\sum_{k=0}^n k^2 = \frac{n(n+1)(2n+1)}{6}$$

Démonstration. Posons, pour tout $n \in \mathbb{N}$,

$$S_n = \sum_{k=0}^n k^2$$

Montrons par récurrence sur n que pour tout $n \in \mathbb{N}$, $S_n = \frac{n(n+1)(2n+1)}{6}$.

- On a $S_0 = 0 = \frac{0(0+1)(2 \times 0 + 1)}{6}$.

- Soit $n \in \mathbb{N}$. Supposons $S_n = \frac{n(n+1)(2n+1)}{6}$. On a alors

$$\begin{aligned}
 S_{n+1} &= \sum_{k=0}^{n+1} k^2 \\
 &= \sum_{k=0}^n k^2 + (n+1)^2 \\
 &= S_n + (n+1)^2 \\
 &= \frac{n(n+1)(2n+1)}{6} + (n+1)^2 \quad (\text{HR}) \\
 &= \frac{(n+1)(n+2)(2n+3)}{6}
 \end{aligned}$$

Par le théorème de démonstration par récurrence, la propriété est établie. $\square$

Exercice. Montrer que pour tout entier $n \geq 2$, 10 divise $2^{2^n} - 6$.

Faisons une récurrence sur n .

- $2^{2^2} - 6 = 10$ qui est bien divisible par 10.
- Soit $n \geq 2$. Supposons que $10 \mid 2^{2^n} - 6$. On a alors

$$\begin{aligned}
 2^{2^{n+1}} - 6 &= 2^{2^n \times 2} - 6 \\
 &= (2^{2^n})^2 - 6 \\
 &= ((2^{2^n} - 6) + 6)^2 - 6
 \end{aligned}$$

Par l'hypothèse de récurrence, il existe $k \in \mathbb{N}$ tel que $2^{2^n} - 6 = 10k$. On a donc

$$\begin{aligned}
 2^{2^{n+1}} - 6 &= (10k + 6)^2 - 6 \\
 &= 100k^2 + 120k + 30 \\
 &= 10(10k^2 + 12k + 3)
 \end{aligned}$$

Ainsi, 10 divise $2^{2^{n+1}} - 6$.

Par le théorème de démonstration par récurrence, la propriété est démontrée.

Exercice. Montrer qu'il existe 4 réels a, b, c, d tels que, pour tout entier naturel n ,

$$\sum_{k=0}^n k^3 = an^4 + bn^3 + cn^2 + d^n$$

Calculer les entiers a, b, c, d .

Proposition 8. Soit $n \in \mathbb{N}$. Soit $x \in \mathbb{C}$.

- Si $x = 1$,

$$\sum_{k=0}^n x^k = n + 1$$

- Si $x \neq 1$,

$$\sum_{k=0}^n x^k = \frac{x^{n+1} - 1}{x - 1}$$

Démonstration. Le cas $x = 1$ est laissé au lecteur. Supposons donc $x \neq 1$ et faisons une récurrence sur n . Pour tout $n \in \mathbb{N}$, posons

$$S_n = \sum_{k=0}^n x^k$$

Montrons par récurrence sur n que pour tout $n \in \mathbb{N}$, $S_n = \frac{x^{n+1}-1}{x-1}$.

- On a $S_0 = 1 = \frac{x^1-1}{x-1}$.
- Soit $n \in \mathbb{N}$. Supposons que $S_n = \frac{x^{n+1}-1}{x-1}$. On a alors

$$\begin{aligned} S_{n+1} &= \sum_{k=0}^{n+1} x^k \\ &= \sum_{k=0}^n x^k + x^{n+1} \\ &= S_n + x^{n+1} \\ &= \frac{x^{n+1} - 1}{x - 1} + x^{n+1} \quad (\text{HR}) \\ &= \frac{x^{n+1} - 1 + x^{n+1}(x - 1)}{x - 1} \\ &= \frac{x^{n+2} - 1}{x - 1} \end{aligned}$$

La propriété est ainsi vérifiée pour tout entier n . $\square$

Corollaire 9. Soient $a, b \in \mathbb{C}$. Soit $n \in \mathbb{N}$. On a

$$a^{n+1} - b^{n+1} = (a - b) \sum_{k=0}^n a^k b^{n-k}$$

Démonstration. Si $b = 0$ ou $a = b$, c'est évident. Sinon,

$$S = \sum_{k=0}^n a^k b^{n-k} = b^n \sum_{k=0}^n \left(\frac{a}{b}\right)^k = b^n \frac{\left(\frac{a}{b}\right)^{n+1} - 1}{\frac{a}{b} - 1}$$

On conclut en mettant le tout au même dénominateur. $\square$

2.3 La formule du binôme

Nous reviendrons très longuement tout au long de l'année sur cette formule essentielle. Son usage est permanent.

Définition 1. Soit $n \in \mathbb{N}^*$. On appelle factorielle de n l'entier

$$n! = \prod_{k=1}^n k = 1 \times 2 \times \dots \times n$$

On convient également de poser $0! = 1$.

Définition 2. Soient $n, k \in \mathbb{N}$. Lorsque $k \leq n$, on pose

$$\binom{n}{k} = \frac{n!}{k!(n-k)!}$$

On convient également de poser $\binom{n}{k} = 0$ lorsque $k > n$ ou $k < 0$. Les nombres $\binom{n}{k}$ sont appelés les *coefficients binomiaux*.

Les coefficients binomiaux possèdent un très grand nombre de propriétés remarquables. Nous en reparlerons en fin d'année. Pour l'instant, nous nous limitons à un petit nombre de formules importantes.

Proposition 10. On a

- $\forall n \in \mathbb{N}, \binom{n}{0} = \binom{n}{n} = 1$
- $\forall n, k \in \mathbb{N}, k \leq n, \binom{n}{k} = \binom{n}{n-k}$.
- $\forall n, k \in \mathbb{N}, \binom{n}{k} + \binom{n}{k+1} = \binom{n+1}{k+1}$.

Démonstration. Il suffit d'appliquer la définition des coefficients binomiaux. Pour le troisième point, il convient de distinguer les deux cas $k < n$ et $k = n$. $\square$

Remarque 2. La dernière formule est ce que l'on appelle une relation de récurrence. Elle permet de calculer les coefficients binomiaux pour l'entier $n + 1$ lorsqu'on les connaît pour

l'entier n . Ceci peut être résumé dans un tableau appelé le triangle de Pascal, qui contient à la ligne n et à la colonne k le coefficient $\binom{n}{k}$.

	0	1	2	3	4	5	6	7
0	1	0	0	0	0	0	0	0
1	1	1	0	0	0	0	0	0
2	1	2	1	0	0	0	0	0
3	1	3	3	1	0	0	0	0
4	1	4	6	4	1	0	0	0
5	1	5	10	10	5	1	0	0
6	1	6	15	20	15	6	1	0
7	1	7	21	35	35	21	7	1

Voici maintenant l'une des formules les plus importantes de l'année. Il s'agit de la formule du binôme.

Proposition 11. Soient $a, b \in \mathbb{C}$. Soit $n \in \mathbb{N}$. On a

$$(a + b)^n = \sum_{k=0}^n \binom{n}{k} a^k b^{n-k}$$

Démonstration. On fait une récurrence sur n .

- Pour $n = 0$, c'est évident.
- Supposons que la formule soit vraie pour un certain entier n . On a alors

$$\begin{aligned} X &= (a + b)^{n+1} = (a + b)(a + b)^n \\ &= (a + b) \sum_{k=0}^n \binom{n}{k} a^k b^{n-k} \\ &= \sum_{k=0}^n \binom{n}{k} a^{k+1} b^{n-k} + \sum_{k=0}^n \binom{n}{k} a^k b^{n-k+1} \end{aligned}$$

On pose dans la première somme $k' = k + 1$. Il vient

$$\begin{aligned} X &= \sum_{k=1}^{n+1} \binom{n}{k-1} a^k b^{n+1-k} + \sum_{k=0}^n \binom{n}{k} a^k b^{n+1-k} \\ &= a^{n+1} + b^{n+1} + \sum_{k=1}^n \left(\binom{n}{k-1} + \binom{n}{k} \right) a^k b^{n+1-k} \end{aligned}$$

Il ne reste qu'à appliquer la formule sur les coefficients binomiaux vue un peu plus haut.

□

Exercice. Soit $n \in \mathbb{N}$. Appliquer judicieusement la formule du binôme pour montrer que

$$\sum_{k=0}^n \binom{n}{k} = 2^n$$

Calculer de même

$$\sum_{k=0}^n (-1)^k \binom{n}{k}$$

3 Extensions du principe de récurrence

3.1 Récurrence forte

Proposition 12. Soit $\mathcal{P}(n)$ une propriété dépendant de l'entier n . On suppose

- $\mathcal{P}(0)$
- $\forall n \in \mathbb{N}, [\forall k \leq n, \mathcal{P}(k)] \implies \mathcal{P}(n+1)$

Alors $\forall n \in \mathbb{N}, \mathcal{P}(n)$.

Démonstration. On pose $Q(n) = \ll \forall k \leq n, \mathcal{P}(k) \gg$ et on montre par récurrence simple sur n que $Q(n)$ est vraie pour tous les entiers naturels n . $\mathcal{P}(n)$ est donc a fortiori vraie. □

Exemple. Montrons que tout entier naturel supérieur ou égal à 2 possède un diviseur premier.

C'est vrai pour $n = 2$.

Soit maintenant $n \geq 2$. Supposons que tous les entiers entre 2 et n ont un diviseur premier. Si $n + 1$ est premier, il a clairement un diviseur premier : lui-même. Sinon, $n + 1 = ab$ où a et b sont deux entiers entre 2 et n . Par l'hypothèse de récurrence, a possède un diviseur premier. Ce nombre premier divise également $n + 1$.

3.2 Récurrence à 2 termes

Proposition 13. Soit $\mathcal{P}(n)$ une propriété dépendant de l'entier n . On suppose

- $\mathcal{P}(0)$
- $\mathcal{P}(1)$

$$\bullet \forall n \in \mathbb{N}, (P(n) \text{ et } P(n+1)) \implies P(n+2)$$

Alors $\forall n \in \mathbb{N}, P(n)$.

Démonstration. On le montre par récurrence forte sur n . $\square$

Exemple. Soit $(F_n)_{n \geq 0}$ la suite d'entiers définie par $F_0 = 0$, $F_1 = 1$, $F_{n+2} = F_{n+1} + F_n$. Soient φ et $\bar{\varphi}$ les racines de l'équation $x^2 = x + 1$. On a alors pour tout entier n ,

$$F_n = \frac{\varphi^n - \bar{\varphi}^n}{\varphi - \bar{\varphi}}$$

Montrons-le par récurrence à deux termes sur n .

Tout d'abord,

$$F_0 = 0 = \frac{\varphi^0 - \bar{\varphi}^0}{\varphi - \bar{\varphi}}$$

$$F_1 = 1 = \frac{\varphi^1 - \bar{\varphi}^1}{\varphi - \bar{\varphi}}$$

Soit $n \in \mathbb{N}$. Supposons la formule vraie pour n et $n+1$. On a alors

$$\begin{aligned} F_{n+2} &= F_{n+1} + F_n \\ &= \frac{\varphi^{n+1} - \bar{\varphi}^{n+1}}{\varphi - \bar{\varphi}} + \frac{\varphi^n - \bar{\varphi}^n}{\varphi - \bar{\varphi}} \quad (\text{HR}) \\ &= \frac{(\varphi^{n+1} + \varphi^n) - (\bar{\varphi}^{n+1} + \bar{\varphi}^n)}{\varphi - \bar{\varphi}} \\ &= \frac{\varphi^n(\varphi + 1) - \bar{\varphi}^n(\bar{\varphi} + 1)}{\varphi - \bar{\varphi}} \end{aligned}$$

Comme φ et $\bar{\varphi}$ sont racines de l'équation $x^2 = x + 1$, on a $\varphi + 1 = \varphi^2$ et de même pour $\bar{\varphi}$. Ainsi,

$$\begin{aligned} F_{n+2} &= \frac{\varphi^n \varphi^2 - \bar{\varphi}^n \bar{\varphi}^2}{\varphi - \bar{\varphi}} \\ &= \frac{\varphi^{n+2} - \bar{\varphi}^{n+2}}{\varphi - \bar{\varphi}} \end{aligned}$$

Remarque 3. On peut bien sûr avoir également des récurrences à 3, 4, ... termes.

4 Suites définies par récurrence

Proposition 14. Soit $f : E \rightarrow E$ une application. Soit $a \in E$. Il existe une unique suite $(u_n)_{n \geq 0}$ à valeurs dans E , telle que

$$\begin{cases} u_0 = a \\ \forall n \in \mathbb{N}, u_{n+1} = f(u_n) \end{cases}$$

Remarque 4. L'unicité, laissée en exercice, est simple à montrer. Supposons qu'il existe deux suites $(u_n)_{n \geq 0}$ et $(v_n)_{n \geq 0}$ vérifiant ces conditions. Montrons par récurrence sur n que pour tout $n \in \mathbb{N}$, $u_n = v_n$.

- Tout d'abord, $u_0 = a = v_0$.
- Soit $n \in \mathbb{N}$. Supposons $u_n = v_n$. On a alors

$$u_{n+1} = f(u_n) \stackrel{(\text{HR})}{=} f(v_n) = v_{n+1}$$

C'est l'existence qui est difficile. On pourra trouver la preuve en annexe.

Remarque 5. En utilisant les corollaires du principe de récurrence, on peut également définir des suites récurrentes à deux, trois, ..., n termes. Nous l'avons déjà fait un peu plus haut avec l'exemple des nombres de Fibonacci.

Exercice. Soit $(u_n)_{n \geq 0}$ la suite définie par récurrence forte par $u_0 = 1$ et, pour tout $n \in \mathbb{N}$,

$$u_{n+1} = \sum_{k=0}^n u_k$$

Calculer, pour tout $n \in \mathbb{N}$, u_n en fonction de n .

5 Annexe : suites récurrentes

Rappelons le théorème sur les suites récurrentes. La démonstration de l'existence est non triviale et nécessite quelques lemmes.

Théorème 15. Soit E un ensemble non vide. Soit $a \in E$. Soit $f : E \rightarrow E$. Il existe une unique suite $u \in E^{\mathbb{N}}$ telle que $u_0 = a$ et pour tout entier n , $u_{n+1} = f(u_n)$.

Démonstration. On pose pour tout entier n ,

$$\mathcal{E}_n = \{s \in E^{\mathbb{N}}, s_0 = a, \forall k < n, s_{k+1} = f(s_k)\}$$

Lemme 16. $\mathcal{E}_0 \supset \mathcal{E}_1 \supset \mathcal{E}_2 \supset \dots$

Démonstration. C'est évident. $\square$

Lemme 17. Soit $n \in \mathbb{N}$. Soient $s, s' \in \mathcal{E}_n$. On a $\forall k \leq n, s_k = s'_k$.

Démonstration. C'est une récurrence facile sur k . $\square$

Lemme 18. $\forall n \in \mathbb{N}, \mathcal{E}_n \neq \emptyset$.

Démonstration. Récurrence sur n . La suite $(a, a, a, \dots)$ est dans $\mathcal{E}_0$ qui est donc non vide. Soit $n \in \mathbb{N}$. Supposons que $\mathcal{E}_n$ est non vide. Soit $s \in \mathcal{E}_n$. La suite

$$s' = (s_0, s_1, \dots, s_n, f(s_n), a, a, \dots)$$

est bien dans $\mathcal{E}_{n+1}$ qui est ainsi non vide. $\square$

Lemme 19. Soit $n \in \mathbb{N}$. Soit $s \in \mathcal{E}_n$, soit $s' \in \mathcal{E}_{n+1}$. On a alors $\forall k \leq n, s_k = s'_k$.

Démonstration. Soient $s \in \mathcal{E}_0, s' \in \mathcal{E}_1$. On a $s_0 = a = s'_0$. La propriété est donc vraie pour $n = 0$. Soit $n \in \mathbb{N}$. Supposons la propriété vraie pour n . Soient $s \in \mathcal{E}_{n+1}, s' \in \mathcal{E}_{n+2}$. Alors, $s \in \mathcal{E}_n, s' \in \mathcal{E}_{n+1}$, donc par l'hypothèse de récurrence $\forall k \leq n, s_k = s'_k$. Puis $s_{n+1} = f(s_n) = f(s'_n) = s'_{n+1}$. $\square$

Nous pouvons maintenant finir la démonstration du théorème.

Pour tout entier n , choisissons $s^n \in \mathcal{E}_n$. Soit $u : \mathbb{N} \rightarrow E$ définie par $u_n = s^n_n$. On a $u_0 = s^0_0 = a$. Et pour tout entier n ,

$$u_{n+1} = s^{n+1}_{n+1} = f(s^{n+1}_n) = f(s^n_n) = f(u_n)$$

La suite u répond bien à la question.

L'unicité est quant à elle facile. Si deux suites u et v répondent à la question, on a $u_0 = a = v_0$. Et si $u_n = v_n$ pour un entier n , alors $u_{n+1} = f(u_n) = f(v_n) = v_{n+1}$. On a donc $u = v$.

$\square$

Chapitre 4

Applications et relations

1 Applications

1.1 Définitions

Définition 1. Soient E et F deux ensembles. On appelle *application* de E vers F tout triplet $f = (E, F, \mathcal{G})$ où $\mathcal{G}$ est une partie de $E \times F$, vérifiant de plus :

$$\forall x \in E, \exists ! y \in F, (x, y) \in \mathcal{G}$$

- E est l'ensemble de départ de f .
- F est l'ensemble d'arrivée de f .
- $\mathcal{G}$ est le graphe de f .

Définition 2. Soit $f = (E, F, \mathcal{G})$ une application. Pour tout $x \in E$, l'unique $y \in F$ tel que $(x, y) \in \mathcal{G}$ est appelé l'image de x par f et est noté $f(x)$.

Lorsqu'on se donne une application f de E vers F , on note très souvent

$$f : E \rightarrow F$$

Lorsqu'on a envie de préciser quelle est l'image des éléments de l'ensemble de départ, on note sur une deuxième ligne quelque chose du genre $x \mapsto f(x)$. Ainsi, par exemple,

$$\begin{array}{ccc} f & : & \mathbb{R} \longrightarrow \mathbb{R} \\ & & x \longmapsto x^2 + x - 1 \end{array}$$

Notation. On note $\mathcal{F}(E, F)$ ou encore F^E l'ensemble des applications de E dans F .

Remarque 1. Une suite à valeurs dans E est une application $u : \mathbb{N} \rightarrow E$. Tout le monde sait que l'image de l'entier n par la suite u est en général notée u_n et pas $u(n)$. Usuellement, on se donne la suite $(u_n)_{n \in \mathbb{N}}$, et pas la suite $u : \mathbb{N} \rightarrow E$. C'est juste une question de notation et de vocabulaire. Plus généralement :

Définition 3. Soient I et E deux ensembles. On appelle famille d'éléments de E indexée par I (le graphe de) toute application

$$\begin{array}{ccc} \mathcal{F} & : & I \longrightarrow E \\ & & i \longmapsto x_i \end{array}$$

On note de façon plus compacte

$$\mathcal{F} = (x_i)_{i \in I}$$

1.2 Restrictions, prolongements

La donnée d'une application entraîne la donnée de ses ensembles de départ et d'arrivée. Il arrive que l'on veuille modifier ceux-ci, en les élargissant ou en les restreignant.

Définition 4. Soit $f : E \rightarrow F$. Soit $E' \subset E$. On appelle restriction de f à E' l'application $f|_{E'} : E' \rightarrow F$ définie par

$$\forall x \in E' \quad f|_{E'}(x) = f(x)$$

Définition 5. Soit $f : E \rightarrow F$. Soit $E' \supset E$. On appelle UN prolongement de f à E' toute application $g : E' \rightarrow F$ telle que $g|_E = f$.

On peut également s'occuper de l'ensemble d'arrivée, et parler de co-restriction ou de co-prolongement.

1.3 Injections, surjections, bijections

Définition 6. Soit $f : E \rightarrow F$. Soit $y \in F$. On appelle UN antécédent de y par f tout élément x de E tel que $f(x) = y$. Un élément de F peut avoir zéro, un ou plusieurs (voire une infinité) d'antécédents. On ne note pas de façon particulière un antécédent.

Évidemment, il serait plus simple que tout élément de F ait un et un seul antécédent par f . Ceci suggère trois définitions :

Définition 7. Soit $f : E \rightarrow F$. On dit que f est :

- injective lorsque tout élément de F a au plus un antécédent par f .
- surjective lorsque tout élément de F a au moins un antécédent par f .
- bijective lorsque tout élément de F a exactement un antécédent par f .

Ainsi, une application est bijective lorsqu'elle est à la fois injective et surjective.

Proposition 1. Soit $f : E \rightarrow F$. f est injective si et seulement si

$$\forall x, x' \in E \quad f(x) = f(x') \implies x = x'$$

Démonstration. Supposons que f a la propriété ci-dessus. Soit $y \in F$. Supposons que y a deux antécédents x et x' par f . On a alors $y = f(x) = f(x')$ d'où $x = x'$. y a donc au plus un antécédent, et f est injective.

Supposons inversement que f est injective. Soient x et x' deux éléments de E tels que $f(x) = f(x')$. Soit y la valeur commune des images. y a au plus un antécédent par f , or x et x' sont des antécédents de y . Donc $x = x'$. $\square$

Remarque 2. L'injectivité peut aussi être formulée en

$$x \neq x' \implies f(x) \neq f(x')$$

mais il vaut mieux si possible utiliser la caractérisation de la proposition avec les égalités.

1.4 Composition

Définition 8. Soient $f : E \rightarrow F$ et $g : F \rightarrow G$. On appelle composée de f et g l'application $g \circ f : E \rightarrow G$ définie par

$$\forall x \in E \quad g \circ f(x) = g(f(x))$$

Proposition 2. *La composition des applications est associative.*

Démonstration. Soient $f : E \rightarrow F$, $g : F \rightarrow G$ et $h : G \rightarrow H$ trois applications. Les applications $(h \circ g) \circ f$ et $h \circ (g \circ f)$ ont mêmes ensembles de départ et d'arrivée, à savoir E et H . Soit $x \in E$. On a

$$\begin{aligned} (h \circ (g \circ f))(x) &= h((g \circ f)(x)) \\ &= h(g(f(x))) \\ &= (h \circ g)(f(x)) \\ &= ((h \circ g) \circ f)(x) \end{aligned}$$

$\square$

Proposition 3. *Soit*

$$\begin{array}{ccc} id_E & : & E \longrightarrow E \\ & & x \longmapsto x \end{array}$$

l'application « identité de E ». On a pour toute application $f : E \rightarrow F$,

$$f \circ id_E = f \text{ et } id_F \circ f = f$$

Remarque 3. La composition des applications n'est pas commutative. Par exemple, soient les deux applications de $\mathbb{R}$ vers $\mathbb{R}$ $f : x \mapsto x^2$ et $g : x \mapsto x + 1$. On a $(g \circ f)(1) = 2$ alors que $(f \circ g)(1) = 4$. Les applications $f \circ g$ et $g \circ f$ sont donc différentes.

Proposition 4. *La composée de deux injections (resp : surjections, bijections) est une injection (resp : surjection, bijection).*

Démonstration. Soient $f : E \rightarrow F$ et $g : F \rightarrow G$.

- Supposons f et g injectives. Soient $x, x' \in E$. Supposons que $(g \circ f)(x) = (g \circ f)(x')$. On a donc

$$g(f(x)) = g(f(x'))$$

Comme g est injective, il en résulte que $f(x) = f(x')$. Puis, par injectivité de f , $x = x'$.

- Supposons f et g surjectives. Soit $z \in G$. L'application g est surjective, il existe donc $y \in F$ tel que $z = g(y)$. Par la surjectivité de f , il existe $x \in E$ tel que $y = f(x)$. Ainsi,

$$z = g(y) = g(f(x)) = (g \circ f)(x)$$

- Pour la composition des bijections il suffit d'utiliser les deux premiers points.

□

1.5 Application réciproque

Définition 9. Soit $f : E \rightarrow F$. Soit $g : F \rightarrow E$. On dit que

- g est un inverse à gauche de f lorsque $g \circ f = id_E$.
- g est un inverse à droite de f lorsque $f \circ g = id_F$.
- g est un inverse (ou une réciproque) de f lorsque g est un inverse à gauche et à droite de f .

Une application est dite inversible à gauche (resp : à droite) lorsqu'elle admet un inverse à gauche (resp : à droite). Elle est dite inversible lorsqu'elle admet un inverse.

Proposition 5. *Une application possède au plus un inverse.*

Démonstration. Soit $f : E \rightarrow F$. Supposons que f a deux inverses g et h . On a $g \circ f = id_E$. En composant à droite par h et en utilisant l'associativité de $\circ$, on obtient

$$g \circ (f \circ h) = id_E \circ h$$

ou encore

$$g \circ id_F = id_E \circ h$$

c'est à dire $g = h$. $\square$

Quelles sont les applications inversibles à gauches ? Inversibles à droite ? Inversibles ?

Proposition 6. *Soit $f : E \rightarrow F$ où E est un ensemble non vide. L'application f est injective si et seulement si elle est inversible à gauche.*

Démonstration. Supposons f injective. Soit $x_0 \in E$. Soit $g : F \rightarrow E$ définie comme suit : pour tout $y \in F$, si y a un antécédent x par f (forcément unique), $g(y) = x$. Sinon, $g(y) = x_0$. On a bien $g \circ f = id_E$. Inversement, supposons l'existence d'un tel g . Soient $x, x' \in E$ tels que $f(x) = f(x')$. En composant par g , on obtient $x = x'$. $\square$

Proposition 7. *Soit $f : E \rightarrow F$. L'application f est surjective si et seulement si elle est inversible à droite.*

Démonstration. Supposons f surjective. Soit $h : F \rightarrow E$ définie comme suit : pour tout $y \in F$, $h(y)$ = un antécédent de y par f (il en existe au moins un puisque f est surjective). On a bien $f \circ h = id_F$. Supposons inversement l'existence d'un tel h . Soit $y \in F$. On a $f(h(y)) = y$, donc y a un antécédent par f : $h(y)$. f est bien surjective. $\square$

Proposition 8.

- Une application est bijective si et seulement si elle est inversible.
- Une bijection a un unique inverse à gauche et un unique inverse à droite, et ceux ci sont égaux à son inverse.

Démonstration. L'équivalence entre bijectivité et existence des inverses est immédiate d'après la proposition précédente. Supposons maintenant que la bijection f a deux inverses à gauche g et g' . On a $id_E = g \circ f = g' \circ f$. En composant à droite par un inverse à droite h , on obtient $g = g' = h$. D'où l'unicité de l'inverse à gauche, et son égalité avec tout inverse à droite (et donc l'unicité de l'inverse à droite). $\square$

Notation. Soit $f : E \rightarrow F$ une bijection. La réciproque de f est notée f^{-1} .

Remarque 4. Soit $f : E \rightarrow F$ une bijection. La réciproque de f est l'unique application $g : F \rightarrow E$ vérifiant

$$\begin{cases} g \circ f &= id_E \\ f \circ g &= id_F \end{cases}$$

Proposition 9. Soit $f : E \rightarrow F$ une bijection. La réciproque de f est bijective, et

$$(f^{-1})^{-1} = f$$

Soient $f : E \rightarrow F$ et $g : F \rightarrow G$ deux bijections. L'application $g \circ f$ est bijective, et

$$(g \circ f)^{-1} = f^{-1} \circ g^{-1}$$

Démonstration. En composant f^{-1} par f des deux côtés, on obtient l'identité. Donc f est bien la réciproque de f^{-1} . De même, en composant $g \circ f$ par $f^{-1} \circ g^{-1}$, on obtient bien l'identité. $\square$

1.6 Images directes, Images réciproques

Définition 10. Soit $f : E \rightarrow F$.

Pour tout ensemble $A \subset E$, on note

$$f(A) = \{y \in F, \exists x \in E, y = f(x)\}$$

Cet ensemble est appelé l'image directe de A par f .

Pour tout ensemble $B \subset F$, on note

$$f^{-1}(B) = \{x \in E, f(x) \in B\}$$

Cet ensemble est appelé l'image réciproque de B par f .

Remarque 5. Prendre bien garde au fait que les notations sont ambiguës. En particulier, f^{-1} ne désigne PAS la réciproque de f , qui peut d'ailleurs très bien ne pas exister.

Exemple. On a $\sin \mathbb{R} = [-1, 1]$, $\sin^{-1}\{0\} = 2\pi\mathbb{Z}$, $\exp \mathbb{R}_+ = [1, +\infty[$

1.7 Fonctions indicatrices

Définition 11. Soit E un ensemble. Soit $A \subset E$. On appelle fonction indicatrice (ou parfois fonction caractéristique) de A l'application $\mathbb{1}_A : E \rightarrow \{0, 1\}$ définie par

$$\mathbb{1}_A(x) = 1 \text{ si } x \in A \text{ et } \mathbb{1}_A(x) = 0 \text{ si } x \in E \setminus A$$

Exemple. Prenons par exemple $E = \mathbb{R}$. La fonction $\mathbb{1}_{\mathbb{R}}$ est la fonction constante égale à 1. La fonction $\mathbb{1}_{\emptyset}$ est la fonction nulle. La fonction $\mathbb{1}_{[0,1]}$ est la fonction qui vaut 1 sur le segment $[0, 1]$ et qui est nulle partout ailleurs.

Proposition 10. Soit E un ensemble. On a, pour toutes parties A et B de E :

- $\mathbb{1}_{A \cap B} = \mathbb{1}_A \mathbb{1}_B$.
- $\mathbb{1}_{A \cup B} = \mathbb{1}_A + \mathbb{1}_B - \mathbb{1}_A \mathbb{1}_B$.
- $\mathbb{1}_{A \Delta B} = \mathbb{1}_A + \mathbb{1}_B - 2\mathbb{1}_A \mathbb{1}_B$.
- $\mathbb{1}_{E \setminus B} = 1 - \mathbb{1}_B$.
- $\mathbb{1}_{A \setminus B} = \mathbb{1}_A(1 - \mathbb{1}_B)$.

Démonstration. Montrons par exemple la formule pour la réunion. Soit $x \in E$. Si x n'est ni dans A , ni dans B , les deux membres de l'égalité sont nuls. Si x est dans A mais pas dans B , les deux membres de l'égalité valent 1. On fait de même dans les deux cas restants. $\square$

2 Relations d'ordre

2.1 Relations binaires

Définition 12. Soit E un ensemble. On appelle relation (binaire) sur E tout couple $\mathcal{R} = (E, \mathcal{G})$ où $\mathcal{G} \subset E \times E$. L'ensemble $\mathcal{G}$ est appelé le graphe de la relation $\mathcal{R}$.

Étant donnés deux éléments x et y de E on peut avoir ou pas $(x, y) \in \mathcal{G}$. Lorsque c'est le cas, on dit que x et y sont en relation par la relation $\mathcal{R}$, et on note $x\mathcal{R}y$.

2.2 Ensembles ordonnés

Définition 13. Soit $\mathcal{R}$ une relation sur un ensemble E . On dit que $\mathcal{R}$ est une relation d'ordre sur E lorsque :

- $\mathcal{R}$ est réflexive : $\forall x \in E, x\mathcal{R}x$
- $\mathcal{R}$ est antisymétrique : $\forall x, y \in E, x\mathcal{R}y \text{ et } y\mathcal{R}x \implies x = y$
- $\mathcal{R}$ est transitive : $\forall x, y, z \in E, x\mathcal{R}y \text{ et } y\mathcal{R}z \implies x\mathcal{R}z$

Un couple $(E, \mathcal{R})$ où $\mathcal{R}$ est une relation d'ordre sur E est appelé un ensemble ordonné.

Exemple. $\mathbb{R}$ et ses sous ensembles sont ordonnés par l'ordre naturel. L'ensemble $\mathbb{C}$, en revanche, n'a pas d'ordre naturel. Il peut bien sûr être ordonné, par exemple par l'ordre

lexicographique, mais cet ordre n'est pas compatible avec la multiplication.

Si E et F sont deux ensembles, et que F est ordonné, l'ensemble F^E est ordonné en posant

$$f \leq g \iff \forall x \in E, f(x) \leq g(x)$$

L'ensemble $\mathcal{P}(E)$ est ordonné par l'inclusion $\subset$.

Remarque 6. La notion de relation d'ordre que l'on vient de définir est celle d'ordre large. Il existe, pour chaque ordre large $\leq$, un ordre strict associé, défini par

$$x < y \iff x \leq y \text{ et } x \neq y$$

Un ordre strict n'est pas un ordre au sens de notre définition !

Définition 14. Soit $(E, \leq)$ un ensemble ordonné. On dit que l'ordre est total lorsque

$$\forall x, y \in E \quad x \leq y \text{ ou } y \leq x$$

Sinon, on dit que l'ordre est partiel.

Exemple. L'ordre naturel sur $\mathbb{R}$ est total. L'ordre défini ci-dessus sur $\mathbb{R}^{\mathbb{R}}$ est partiel. Dès que l'ensemble E a au moins deux éléments, la relation d'inclusion sur $\mathcal{P}(E)$ est un ordre partiel.

2.3 Majorants, Minorants

Définition 15. Soit $(E, \leq)$ un ensemble ordonné. Soit $A \subset E$. Soit $x \in E$. On dit que x est un majorant de A lorsque

$$\forall a \in A \quad a \leq x$$

De même, on dit que x est un minorant de A lorsque

$$\forall a \in A \quad x \leq a$$

Exemple. Dans $\mathbb{R}$ muni de l'ordre naturel, le réel 7 majore l'ensemble $[0, 1]$. Ce n'est bien sûr pas le seul, il y a aussi, π , 53, et $\dots$, qui a l'air « mieux » que les autres.

Définition 16. Une partie A de l'ensemble ordonné E admet un plus grand élément lorsqu'il existe $x \in A$ majorant de A . On définit de même la notion de plus petit élément.

Proposition 11. *Le plus petit (resp : plus grand) élément de l'ensemble A , s'il existe, est unique. On le note $\min A$ (resp : $\max A$).*

Démonstration. Si x et y sont deux plus petits éléments de A , on a $x \leq y$ et $y \leq x$, d'où $x = y$ par antisymétrie. $\square$

Exemple. Dans $\mathbb{R}$ muni de l'ordre naturel, le plus grand élément de $[0, 1]$ est 1. En revanche, $[0, 1[$ n'a pas de plus grand élément.

Remarque 7. Si l'on reprend l'exemple ci-dessus, on a dit que 1 était un majorant de $[0, 1]$ mieux que les autres. Si l'on se demande pourquoi, on s'aperçoit qu'il y a deux raisons. La première, c'est que $1 \in [0, 1]$. Cela nous donne la notion de plus grand élément. La seconde raison, c'est que 1 est le plus petit des majorants de $[0, 1]$. Pour cet ensemble cela n'est pas très intéressant. En revanche, on s'aperçoit que, bien que $[0, 1[$ n'ait pas de plus grand élément, il possède un plus petit majorant. Cela nous conduira à la notion de borne supérieure, que nous aborderons dans le chapitre sur les nombres réels.

3 Relations d'équivalence

3.1 Notion de relation d'équivalence

Définition 17. Soit $\mathcal{R}$ une relation sur un ensemble E . On dit que $\mathcal{R}$ est une relation d'équivalence sur E lorsque :

- $\mathcal{R}$ est réflexive : $\forall x \in E, x\mathcal{R}x$
- $\mathcal{R}$ est symétrique : $\forall x, y \in E, x\mathcal{R}y \implies y\mathcal{R}x$
- $\mathcal{R}$ est transitive : $\forall x, y, z \in E, x\mathcal{R}y \text{ et } y\mathcal{R}z \implies x\mathcal{R}z$

3.2 Exemples simples

- L'égalité est une relation d'équivalence sur tout ensemble.
- Pour tout ensemble E , la relation triviale $\mathcal{T}$ définie par $\forall x, y \in E, x\mathcal{T}y$ est une relation d'équivalence.
- Soit $f : E \rightarrow F$ une application. La relation $\mathcal{R}$ sur E définie par $\forall x, y \in E, x\mathcal{R}y \iff f(x) = f(y)$ est une relation d'équivalence. L'identité et la relation triviale sur E en sont deux cas particuliers.

3.3 Congruences sur $\mathbb{Z}$

Prenons $E = \mathbb{Z}$, l'ensemble des entiers relatifs. Soit $n \in \mathbb{Z}$.

Définition 18. Étant donnés $a, b \in \mathbb{Z}$, on dit que a est congru à b modulo n , et on note $a \equiv b[n]$ lorsque

$$\exists k \in \mathbb{Z}, b = a + kn$$

Remarque 8. Pour $n = 0$, on obtient la relation d'égalité. Pour $n = 1$, on obtient la relation triviale. Par ailleurs, la relation obtenue pour l'entier $-n$ est la même que celle obtenue pour l'entier n . Nous supposons donc dorénavant que $n \in \mathbb{N}$.

Remarque 9. Notons $n\mathbb{Z} = \{kn, k \in \mathbb{Z}\}$ l'ensemble des multiples de n . On a alors

$$a \equiv b[n] \iff b - a \in n\mathbb{Z}$$

Proposition 12. Soit $n \in \mathbb{N}$. La relation de congruence modulo n est une relation d'équivalence sur $\mathbb{Z}$.

Démonstration. L'ensemble $n\mathbb{Z}$ vérifie les propriétés suivantes : il contient le nombre 0, il est stable par passage à l'opposé, et il est stable pour l'addition. Nous résumerons plus tard ces trois propriétés en disant que $n\mathbb{Z}$ est un sous-groupe de $\mathbb{Z}$. Montrons par exemple la transitivité. Soient $a, b, c \in \mathbb{Z}$. Supposons que $a \equiv b[n]$ et $b \equiv c[n]$. On a alors $c - a = (c - b) + (b - a)$. Mais $c - b$ et $b - a$ sont dans $n\mathbb{Z}$, qui est stable pour l'addition. Donc $c - a \in n\mathbb{Z}$ et donc $a \equiv c[n]$. $\square$

3.4 Congruences sur $\mathbb{R}$

Prenons $E = \mathbb{R}$, l'ensemble des nombres réels. Soit $\alpha \in \mathbb{R}$.

Définition 19. Étant donnés $a, b \in \mathbb{R}$, on dit que a est congru à b modulo α , et on note $a \equiv b[\alpha]$ lorsque

$$\exists k \in \mathbb{Z}, b = a + k\alpha$$

Remarque 10. Mêmes remarques que sur $\mathbb{Z}$. Nous utiliserons en particulier cette notion avec $\alpha = \pi$ ou $\alpha = 2\pi$, en relation avec la trigonométrie et les nombres complexes.

Remarque 11. Notons $\alpha\mathbb{Z} = \{k\alpha, k \in \mathbb{Z}\}$ l'ensemble des multiples de α . On a alors

$$a \equiv b[\alpha] \iff b - a \in \alpha\mathbb{Z}$$

Proposition 13. Soit $\alpha \in \mathbb{R}$. La relation de congruence modulo α est une relation d'équivalence sur $\mathbb{R}$.

Démonstration. Même démonstration que sur $\mathbb{Z}$. $\square$

3.5 Classes d'équivalences

Définition 20. Soit E un ensemble. Soit $\mathcal{R}$ une relation d'équivalence sur E . Soit $x \in E$. On appelle classe de x modulo $\mathcal{R}$ et on note $\bar{x}$ l'ensemble

$$\bar{x} = \{y \in E, x\mathcal{R}y\}$$

Notation. On note $E/\mathcal{R}$ l'ensemble des classes modulo $\mathcal{R}$.

Proposition 14. Avec les notations ci-dessus :

- Les classes sont non vides : $\forall \mathcal{C} \in E/\mathcal{R}, \mathcal{C} \neq \emptyset$.
- Les classes sont disjointes : $\forall \mathcal{C}, \mathcal{C}' \in E/\mathcal{R}, \mathcal{C} \neq \mathcal{C}' \implies \mathcal{C} \cap \mathcal{C}' = \emptyset$.
- Les classes recouvrent E : $\cup_{\mathcal{C} \in E/\mathcal{R}} \mathcal{C} = E$.

On dit que les classes modulo $\mathcal{R}$ forment une partition de E .

Démonstration. La réflexivité de $\mathcal{R}$ montre que $\forall x \in E, x \in \bar{x}$. Cela prouve les points 1 et 3. Soient maintenant deux classes $\mathcal{C} = \bar{x}$ et $\mathcal{C}' = \bar{x}'$. Supposons que $\mathcal{C} \cap \mathcal{C}' \neq \emptyset$. Il existe donc $z \in \mathcal{C} \cap \mathcal{C}'$. Montrons que $\mathcal{C} \subset \mathcal{C}'$, la démonstration de l'autre inclusion étant identique. Soit $y \in \mathcal{C}$. On a $x\mathcal{R}y$ et $x\mathcal{R}z$. Donc, par symétrie et transitivité, $z\mathcal{R}y$. Mais on a aussi $x'\mathcal{R}z$ donc, toujours par transitivité, $x'\mathcal{R}y$, et $y \in \mathcal{C}'$. $\square$

Proposition 15. Soit $n \in \mathbb{N}^*$. La relation sur $\mathbb{Z}$ de congruence modulo n possède exactement n classes, qui sont (par exemple) $\bar{0}, \bar{1}, \dots, \overline{n-1}$.

Remarque 12. Au lieu de $\mathbb{Z}/\equiv$, on note $\mathbb{Z}/n\mathbb{Z}$ l'ensemble des classes modulo n .

Démonstration. Nous admettons provisoirement le résultat suivant, qui sera prouvé dans le chapitre sur les entiers relatifs : Étant donné $a \in \mathbb{Z}$, il existe un unique $q \in \mathbb{Z}$ et un unique $r \in \llbracket 0, n-1 \rrbracket$ tels que $a = qn + r$. Il s'agit de ce que l'on appelle la division euclidienne de a par n . Cette égalité montre que

$$\forall a \in \mathbb{Z}, \exists r \in \llbracket 0, n-1 \rrbracket, a \equiv r[n]$$

ou encore $a \in \bar{r}$. Ainsi,

$$\mathbb{Z}/n\mathbb{Z} \subset \{\bar{0}, \bar{1}, \dots, \overline{n-1}\}$$

L'autre inclusion est évidente. Il y a donc au plus n classes dans $\mathbb{Z}/n\mathbb{Z}$. Attention, il reste maintenant à montrer que ces classes sont bien distinctes ! Soient donc $0 \leq i < j < n$. On a $0 < j - i < n$, donc $j - i \notin n\mathbb{Z}$. Ainsi, $\bar{i} \neq \bar{j}$. $\square$

Remarque 13.

- Pour $n = 0$ on a une infinité de classes, toutes constituées d'un singleton.
- Soit $\alpha \in \mathbb{R}, \alpha > 0$. La relation sur $\mathbb{R}$ de congruence modulo α comporte une infinité de classes. On peut par exemple prendre les classes des $x \in [0, \alpha[$, ou des $x \in]-\frac{\alpha}{2}, \frac{\alpha}{2}]$.
- Les relations de congruence sur $\mathbb{Z}$ possèdent bien d'autres propriétés, que nous verrons en temps utile. En particulier, elles sont compatibles avec l'addition et la multiplication (on peut multiplier et ajouter des congruences, comme s'il s'agissait d'égalités).

3.6 Une réciproque

Proposition 16. *Soit E un ensemble. Soit $\mathcal{P}$ une partition de E . Il existe une et une seule relation d'équivalence $\mathcal{R}$ sur E telle que $E/\mathcal{R} = \mathcal{P}$.*

Démonstration. Soit $\mathcal{R}$ une telle relation. Soient $x, y \in E$. Supposons que $x\mathcal{R}y$. Alors, $\exists \mathcal{C} \in \mathcal{P}, x, y \in \mathcal{C}$, puisque x et y sont par définition dans la même classe modulo $\mathcal{R}$. Inversement, supposons $\exists \mathcal{C} \in \mathcal{P}, x, y \in \mathcal{C}$. x et y sont donc dans la même classe modulo $\mathcal{R}$, et ainsi $x\mathcal{R}y$. En conclusion, si $\mathcal{R}$ est une relation d'équivalence sur E telle que $E/\mathcal{R} = \mathcal{P}$, alors

$$\forall x, y \in E, x\mathcal{R}y \iff \exists \mathcal{C} \in \mathcal{P}, x, y \in \mathcal{C}$$

Nous venons donc de montrer qu'il existe au plus une telle relation. Inversement, on vérifie sans peine que la relation définie ci-dessus est bien une relation d'équivalence, et que ses classes sont justement les éléments de la partition $\mathcal{P}$. $\square$

Chapitre 5

Nombres complexes

1 Introduction

Nous verrons dans un chapitre ultérieur la notion de corps. Sans entrer pour l'instant dans les détails, un corps est un ensemble $\mathbb{K}$ muni de deux opérations : une addition (+) et une multiplication ($\times$). Ces opérations vérifient un certain nombre de propriétés, comme la commutativité, l'associativité, la distributivité de la multiplication par rapport à l'addition, l'existence d'éléments neutres, l'inversibilité des éléments non nuls pour la multiplication, etc.

1.1 Notion de nombre complexe

Théorème 1. *Il existe un ensemble $\mathbb{C}$ muni de deux opérations $+$ et $\times$ vérifiant les propriétés suivantes :*

- $(\mathbb{C}, +, \times)$ est un corps
- $\mathbb{R} \subset \mathbb{C}$
- Il existe $i \in \mathbb{C}$ tel que $i^2 = -1$
- Tout élément $z \in \mathbb{C}$ s'écrit de façon unique $z = x + iy$ avec x et y réels.

Un tel corps est « unique ».

Démonstration. Nous ne détaillons pas ici la construction de $\mathbb{C}$. Il en existe de très nombreuses, à partir de couples de réels, de matrices, de polynômes, etc. $\square$

Remarquons simplement que les propriétés de $\mathbb{C}$ obligent les opérations à s'écrire comme suit :

$$\begin{aligned}(x + iy) + (x' + iy') &= (x + x') + i(y + y') \\ (x + iy) \times (x' + iy') &= (xx' - yy') + i(xy' + x'y)\end{aligned}$$

Remarque 1. L'opposé du nombre complexe z est noté $-z$. Bien évidemment,

$$-(x + iy) = -x + i(-y)$$

et on dispose de l'opération de soustraction, définie par $z - z' = z + (-z')$.

L'inverse d'un nombre complexe non nul étant noté z^{-1} , on dispose alors de l'opération de division, définie par

$$\frac{z}{z'} = zz'^{-1}$$

Bien entendu, pour tout nombre complexe non nul z , on a $z^{-1} = \frac{1}{z}$ puisque

$$\frac{1}{z} = 1 \times z^{-1} = z^{-1}$$

Définition 1. Soit $z \in \mathbb{C}$. Soit (x, y) l'unique couple de réels tel que $z = x + iy$. x est appelé la partie réelle de z , et y est la partie imaginaire de z . On note $x = \operatorname{Re} z$ et $y = \operatorname{Im} z$.

Les nombres de la forme iy , avec y réel sont dits *imaginaires purs*. On note $\mathbb{R}i$ (ou $i\mathbb{R}$) l'ensemble des imaginaires purs.

Remarque 2. Sauf mention contraire, lorsqu'on écrira « $z = x + iy$ » dans ce qui suit, on supposera que x et y sont réels. Dans une rédaction complète, il faudrait le préciser à chaque fois.

1.2 Affixe, image

Notons $\mathcal{P} = \mathbb{R}^2$ l'ensemble des points du plan (bref, le plan!).

Considérons les deux applications $\mathcal{A} : \mathcal{P} \rightarrow \mathbb{C}$ et $\mathcal{I} : \mathbb{C} \rightarrow \mathcal{P}$ définies comme suit :

- si $M = (x, y)$ est un point du plan, $\mathcal{A}(M) = x + iy$.
- Si $z = x + iy \in \mathbb{C}$, $\mathcal{I}(z) = (x, y) \in \mathcal{P}$.

On a bien évidemment

$$\forall M \in \mathcal{P}, \mathcal{I}(\mathcal{A}(M)) = M$$

et

$$\forall z \in \mathbb{C}, \mathcal{A}(\mathcal{I}(z)) = z$$

Les fonctions $\mathcal{A}$ et $\mathcal{I}$ sont des bijections, et chacune est la réciproque de l'autre.

Définition 2.

- Pour tout point $M \in \mathcal{P}$, le nombre complexe $z = \mathcal{A}(M)$ est appelé l'affixe de M . On notera parfois $M(z)$ pour indiquer que M est le point d'affixe z .
- Pour tout $z \in \mathbb{C}$, le point $M = \mathcal{I}(z)$ est appelé l'image de z .

Remarque 3. Au lieu de voir $\mathcal{P} = \mathbb{R}^2$ comme un ensemble de points, on peut aussi le voir comme un ensemble de vecteurs et parler ainsi du vecteur $\overrightarrow{u}$ d'affixe z ou du vecteur $\overrightarrow{u}$ image du nombre complexe z .

Remarque 4. Il nous arrivera parfois de confondre abusivement un nombre complexe avec son affixe. Ainsi, nous parlerons de nombres complexes alignés, ou sur un même cercle, etc.

1.3 Conjugué

Définition 3. Soit $z = x + iy$ un nombre complexe. Le conjugué de z est

$$\bar{z} = x - iy$$

Remarque 5. Géométriquement, l'image du conjugué de z est le symétrique de l'image z par rapport à l'axe Ox .

Proposition 2. Soit $z \in \mathbb{C}$. On a :

- $Re\ z = \frac{z + \bar{z}}{2}$ et $Im\ z = \frac{z - \bar{z}}{2i}$.
- $\bar{\bar{z}} = z$.
- $z \in \mathbb{R} \iff \bar{z} = z$.
- $z \in i\mathbb{R} \iff \bar{z} = -z$.

Démonstration. Il suffit de poser $z = x + iy$ avec $x, y \in \mathbb{R}$. $\square$

Proposition 3.

- *Le conjugué d'une somme est la somme des conjugués.*
- *Le conjugué d'un produit est le produit des conjugués.*
- *Le conjugué d'un quotient est le quotient des conjugués.*

Démonstration. Il suffit de poser $z = x + iy$, $z' = x' + iy'$ avec $x, x', y, y' \in \mathbb{R}$. $\square$

Exercice. Déterminons tous les nombres complexes z tels que $\frac{z+i}{z-i} \in \mathbb{R}i$. Soit $z \in \mathbb{C}$, $z \neq i$. z est solution du problème si et seulement si

$$\frac{z+i}{z-i} = -\frac{\overline{z+i}}{\overline{z-i}} = -\frac{\bar{z}-i}{\bar{z}+i}$$

On réduit au même dénominateur :

$$(z+i)(\bar{z}+i) = -(z-i)(\bar{z}-i)$$

qui se simplifie en $z\bar{z} = 1$. Les solutions du problème sont les nombres complexes de module 1, i excepté.

Remarque 6. Soit $\varphi : \mathbb{C} \rightarrow \mathbb{C}$ une application vérifiant les propriétés suivantes :

- Pour tous $z, z' \in \mathbb{C}$, $\varphi(z + z') = \varphi(z) + \varphi(z')$.
- Pour tous $z, z' \in \mathbb{C}$, $\varphi(zz') = \varphi(z)\varphi(z')$.
- Pour tous $x \in \mathbb{R}$, $\varphi(x) = x$.

Soit $z = x + iy \in \mathbb{C}$. On a alors

$$\varphi(z) = \varphi(x) + \varphi(i)\varphi(y) = x + \varphi(i)y$$

Mais $i^2 = -1$, donc $\varphi(i)^2 = \varphi(-1) = -1$. On en déduit que $\varphi(i) = \pm i$.

Si $\varphi(i) = i$, alors, pour tout complexe z , $\varphi(z) = z$: φ est l'identité de $\mathbb{C}$.

Si, au contraire, $\varphi(i) = -i$, alors, pour tout complexe z , $\varphi(z) = \bar{z}$: φ est la conjugaison.

En résumé, la conjugaison est l'unique endomorphisme non trivial de $\mathbb{C}$ laissant les réels invariants. C'est donc une application extrêmement précieuse : si vous la perdez, vous n'en aurez pas d'autre.

1.4 Module

Proposition 4. Soit $z \in \mathbb{C}$. Alors $z\bar{z}$ est un réel positif.

Démonstration. Posons $z = x + iy$, où $x, y \in \mathbb{R}$. Alors $z\bar{z} = x^2 + y^2 \in \mathbb{R}_+$. $\square$

Définition 4. On appelle module du complexe z le réel positif

$$|z| = \sqrt{z\bar{z}}$$

Proposition 5. Pour tout complexe z , on a

- $|z| = 0 \iff z = 0$
- $|z| = |\bar{z}|$
- $\operatorname{Re} z \leq |\operatorname{Re} z| \leq |z|$
- $\operatorname{Im} z \leq |\operatorname{Im} z| \leq |z|$

Démonstration. On a $|z|^2 = z\bar{z}$ donc $|z| = 0$ si et seulement si $z = 0$ ou $\bar{z} = 0$, c'est à dire si et seulement si $z = 0$. On a $|\bar{z}|^2 = \bar{z}\bar{\bar{z}} = \bar{z}z = |z|^2$.

On pose $z = x + iy$ où $x, y \in \mathbb{R}$. On a $|z|^2 = x^2 + y^2 \geq x^2$. Donc, en passant à la racine carrée, $|x| \leq |z|$. Même preuve pour la partie imaginaire. $\square$

Proposition 6. Pour tout complexe z non nul, on a

- $\frac{1}{z} = \frac{\bar{z}}{|z|^2}$.
- $|z| = 1 \iff \frac{1}{z} = \bar{z}$.

Démonstration. On a

$$\frac{1}{z} = \frac{\bar{z}}{z\bar{z}} = \frac{\bar{z}}{|z|^2}$$

Et aussi, $|z| = 1$ si et seulement si $z\bar{z} = 1$ ou encore $\bar{z} = \frac{1}{z}$. $\square$

Proposition 7. Pour tous complexes z_1 et z_2 , on a

- $|z_1 z_2| = |z_1| |z_2|$.
- Si $z_2 \neq 0$, $|\frac{z_1}{z_2}| = \frac{|z_1|}{|z_2|}$.
- $|z_1 + z_2| \leq |z_1| + |z_2|$ (inégalité triangulaire)
- $||z_1| - |z_2|| \leq |z_1 - z_2|$ (inégalité triangulaire 2)

Démonstration. On a

$$|z_1 z_2|^2 = (z_1 z_2) \overline{z_1 z_2} = z_1 \bar{z}_1 z_2 \bar{z}_2 = |z_1|^2 |z_2|^2$$

L'inégalité triangulaire est évidente si $z_2 = 0$. Sinon, on pose $u = \frac{z_1}{z_2}$. Il suffit de prouver que

$$|1 + u| \leq 1 + |u|$$

Or,

$$|1 + u|^2 - (1 + |u|)^2 = 2(|u| - \Re u) \geq 0$$

La dernière inégalité se prouve à l'aide de la précédente : $z_1 = z_1 - z_2 + z_2$, donc $|z_1| \leq |z_1 - z_2| + |z_2|$. Ainsi, $|z_1| - |z_2| \leq |z_1 - z_2|$. On obtient de même que $|z_2| - |z_1| \leq |z_1 - z_2|$. D'où le résultat. $\square$

Remarque 7. Il y a égalité dans l'inégalité triangulaire si et seulement si $z_2 = 0$ ou $z_1 = \lambda z_2$, avec $\lambda \in \mathbb{R}_+$.

2 Nombres complexes et trigonométrie

2.1 Fonctions cosinus, sinus et tangente

Nous rappelons ici les principaux résultats sur les fonctions trigonométriques : Les fonctions sin et cos sont dérivables sur $\mathbb{R}$, et de période 2π . La fonction sinus est impaire sa dérivée est cos. La fonction cosinus est paire, sa dérivée est sin. Nous n'énumérerons pas les différentes symétries de ces fonctions : $\sin(\pi - x) = \sin x$, $\sin(\pi + x) = -\sin x$, etc. Elles se retrouvent facilement à l'aide du cercle trigonométrique. Les formules suivantes sont à connaître. Les lettres a, b, p, q, x désignent des réels quelconques.

$$\begin{aligned}
\cos^2 x + \sin^2 x &= 1 \\
\sin(a+b) &= \sin a \cos b + \cos a \sin b \\
\sin(a-b) &= \sin a \cos b - \cos a \sin b \\
\cos(a+b) &= \cos a \cos b - \sin a \sin b \\
\cos(a-b) &= \cos a \cos b + \sin a \sin b \\
\sin 2x &= 2 \sin x \cos x \\
\cos 2x &= \cos^2 x - \sin^2 x = 2 \cos^2 x - 1 = 1 - 2 \sin^2 x \\
\sin a \cos b &= \frac{1}{2}(\sin(a+b) + \sin(a-b)) \\
\cos a \cos b &= \frac{1}{2}(\cos(a+b) + \cos(a-b)) \\
\sin a \sin b &= -\frac{1}{2}(\cos(a+b) - \cos(a-b)) \\
\sin p - \sin q &= 2 \sin \frac{p-q}{2} \cos \frac{p+q}{2} \\
\cos p - \cos q &= -2 \sin \frac{p-q}{2} \sin \frac{p+q}{2}
\end{aligned}$$

La fonction tangente est quant à elle définie par

$$\tan x = \frac{\sin x}{\cos x}$$

Elle est impaire, π -périodique. Son ensemble de définition est

$$\mathcal{D} = \mathbb{R} \setminus \left\{ \frac{\pi}{2} + k\pi, k \in \mathbb{Z} \right\}$$

Elle est dérivable sur son ensemble de définition. On a

$$\forall x \in \mathcal{D}, \tan' x = 1 + \tan^2 x = \frac{1}{\cos^2 x}$$

Les formules ci-dessous sont aussi à connaître. Pour tous réels a, b, x tels que les expressions ci-dessous aient un sens, on a

$$\begin{aligned}
\tan(a+b) &= \frac{\tan a + \tan b}{1 - \tan a \tan b} \\
\tan(a-b) &= \frac{\tan a - \tan b}{1 + \tan a \tan b} \\
\tan 2x &= \frac{2 \tan x}{1 - \tan^2 x} \\
\tan\left(\frac{\pi}{2} + x\right) &= -\frac{1}{\tan x} \\
\tan\left(\frac{\pi}{2} - x\right) &= \frac{1}{\tan x}
\end{aligned}$$

Exercice. Pour quelles valeurs de a, b, x ces expressions ont-elles un sens ?

Proposition 8. Soit $x \in \mathbb{R} \setminus \{\pi + 2k\pi, k \in \mathbb{Z}\}$, Soit $t = \tan \frac{x}{2}$. On a

$$\begin{aligned}\cos x &= \frac{1 - t^2}{1 + t^2} \\ \sin x &= \frac{2t}{1 + t^2} \\ \tan x &= \frac{2t}{1 - t^2}\end{aligned}$$

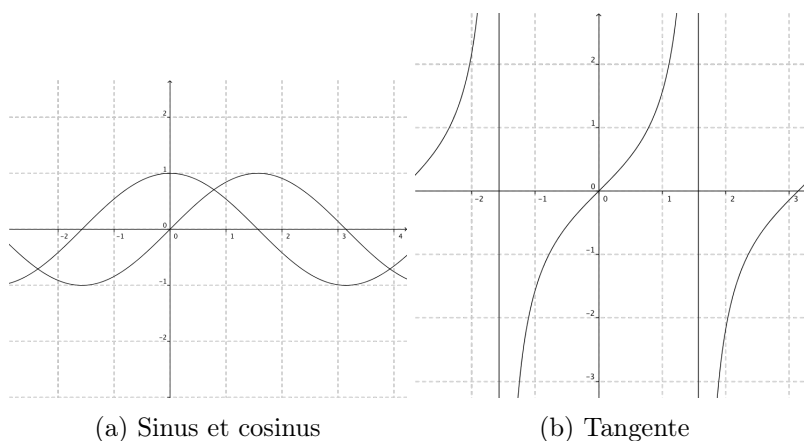


FIGURE 5.1 – Sinus, cosinus, tangente

2.2 Nombres complexes de module 1

Proposition 9. L'ensemble $\mathcal{U}$ des nombres complexes de module 1 est un sous-groupe de $(\mathbb{C}^*, \times)$.

Démonstration. Ici, petite incursion dans la théorie des groupes, avant d'en avoir parlé en cours. Pas d'inquiétude... Un élément de $\mathcal{U}$ est non nul car de module 1. Donc $\mathcal{U} \subset \mathbb{C}^*$. Le produit de deux complexes de module 1, l'inverse d'un complexe de module 1, sont encore de module 1. Nous verrons que cela prouve exactement que $\mathcal{U}$ est un sous-groupe de $\mathbb{C}^*$. $\square$

Définition 5. Pour tout réel θ , on pose $e^{i\theta} = \cos \theta + i \sin \theta$.

Proposition 10. On a $\mathcal{U} = \{e^{i\theta}, \theta \in \mathbb{R}\}$.

Démonstration. On a $|e^{i\theta}| = 1$ d'où l'inclusion réciproque. Inversement, soit $z = x + iy \in \mathcal{U}$. On a $x^2 + y^2 = 1$. Il existe donc un réel θ tel que $z = e^{i\theta}$ d'où l'inclusion. $\square$

Proposition 11. FORMULES D'EULER

Pour $\theta \in \mathbb{R}$, on a

$$\cos \theta = \frac{e^{i\theta} + e^{-i\theta}}{2} \text{ et } \sin \theta = \frac{e^{i\theta} - e^{-i\theta}}{2i}$$

Démonstration. C'est immédiat. On utilise la parité du cosinus et l'imparité du sinus. $\square$

Proposition 12. On a, pour θ et φ réels :

- $e^{i(\theta+\varphi)} = e^{i\theta}e^{i\varphi}$.
- $e^{-i\theta} = 1/e^{i\theta}$
- $e^{i(\theta-\varphi)} = e^{i\theta}/e^{i\varphi}$.
- $e^{i\theta} = 1 \iff \theta \equiv 0[2\pi]$.
- $e^{i\theta} = e^{i\varphi} \iff \theta \equiv \varphi[2\pi]$.

Démonstration. La première égalité se vérifie par un petit calcul. La deuxième résulte de la première. La troisième résulte des deux premières. Soit maintenant un réel θ . On a $e^{i\theta} = 1$ si et seulement si $\cos \theta = 1$ et $\sin \theta = 0$, c'est à dire si et seulement si θ est multiple de 2π . La dernière propriété est alors immédiate : $e^{i\theta} = e^{i\varphi} \iff e^{i(\theta-\varphi)} = 1 \iff \theta - \varphi \equiv 0[2\pi]$. $\square$

Remarque 8. On vient de prouver que l'application $\theta \mapsto e^{i\theta}$ est un morphisme surjectif du groupe $(\mathbb{R}, +)$ vers le groupe $(\mathcal{U}, \times)$ dont le noyau est $2\pi\mathbb{Z}$.

Proposition 13. FORMULE DE MOIVRE

Pour $\theta \in \mathbb{R}$ et $n \in \mathbb{Z}$, on a

$$(\cos \theta + i \sin \theta)^n = \cos n\theta + i \sin n\theta$$

Démonstration. C'est immédiat en utilisant le fait que l'application $\theta \mapsto e^{i\theta}$ est un morphisme. $\square$

Exercice. Soit x un réel et n un entier. On se propose de calculer

$$S_n = \sum_{k=0}^n \cos kx$$

et

$$T_n = \sum_{k=0}^n \sin kx$$

Pour cela, on forme

$$S_n + iT_n = \sum_{k=0}^n e^{ikx}$$

On reconnaît la somme des termes d'une suite géométrique de raison e^{ix} .

- Si $x \equiv 0[2\pi]$, alors $e^{ix} = 1$, donc $S_n + iT_n = n + 1$, d'où $S_n = n + 1$ et $T_n = 0$.
- Si, au contraire, $x \not\equiv 0[2\pi]$, alors $e^{ix} \neq 1$, donc

$$S_n + iT_n = \frac{e^{i(n+1)x} - 1}{e^{ix} - 1}$$

On factorise au numérateur et au dénominateur de cette fraction la quantité $e^{ix/2}$.
Il vient alors

$$S_n + iT_n = \frac{e^{i(n+1/2)x} - e^{-ix/2}}{e^{ix/2} - e^{-ix/2}} = \frac{e^{i(n+1/2)x} - e^{-ix/2}}{2i \sin(x/2)}$$

Il vient alors facilement

$$S_n = \frac{1}{2} + \frac{\sin(n + \frac{1}{2})x}{2 \sin \frac{x}{2}} \quad \text{et} \quad T_n = \frac{\cos \frac{x}{2} - \cos(n + \frac{1}{2})x}{2 \sin \frac{x}{2}}$$

2.3 Arguments d'un nombre complexe non nul

Soit z un nombre complexe non nul. Le complexe $\frac{z}{|z|}$ est alors de module 1, donc il existe $\theta \in \mathbb{R}$ tel que $\frac{z}{|z|} = e^{i\theta}$. Ainsi,

$$z = re^{i\theta} \text{ où } r = |z|$$

Définition 6. Soit z un complexe non nul. On appelle argument de z tout réel θ tel que $z = |z|e^{i\theta}$.

Proposition 14. Soit z un complexe non nul. Soit θ_0 un argument de z . L'ensemble des arguments de z est

$$\{\theta_0 + 2k\pi, k \in \mathbb{Z}\}$$

Si z est un complexe et θ est un argument de z , on note $\arg z \equiv \theta[2\pi]$. Souvent, on écrit abusivement $\arg z = \theta$.

Remarque 9.

- Si z est sous la forme $z = ae^{i\theta}$ avec a et θ réels, il convient de discuter sur le signe de a pour connaître son module et un de ses arguments.
- Si l'on astreint l'argument à demeurer dans un intervalle d'amplitude 2π , tel que $[0, 2\pi[$ ou $]-\pi, \pi]$, on a alors unicité de l'argument.

Exemple. Trouvons le module et un argument de $z = \sqrt{3} + i$. On a $|z|^2 = (\sqrt{3})^2 + 1 = 4$, donc $|z| = 2$. Maintenant, $\frac{z}{|z|} = \frac{\sqrt{3}}{2} + \frac{1}{2}i$. Un argument de z est donc un réel θ tel que $\cos \theta = \frac{\sqrt{3}}{2}$ et $\sin \theta = \frac{1}{2}$. Par exemple, $\theta = \frac{\pi}{3}$.

2.4 Forme trigonométrique

Soit $z \in \mathbb{C}$. On dit que z est mis sous forme trigonométrique lorsqu'on a écrit $z = re^{i\theta}$ avec $r \in \mathbb{R}_+$ et $\theta \in \mathbb{R}$. La forme trigonométrique est particulièrement adaptée aux calculs de produits, quotients, puissances, et pas du tout adaptée aux calculs de sommes.

Proposition 15. Si $z_1 = r_1 e^{i\theta_1}$ et $z_2 = r_2 e^{i\theta_2}$, alors :

$$z_1 z_2 = r_1 r_2 e^{i(\theta_1 + \theta_2)}$$

$$\frac{z_1}{z_2} = \frac{r_1}{r_2} e^{i(\theta_1 - \theta_2)}$$

Corollaire 16. Soient z_1 et z_2 deux complexes non nuls. On a :

$$\arg(z_1 z_2) \equiv \arg z_1 + \arg z_2 [2\pi]$$

$$\arg \frac{z_1}{z_2} \equiv \arg z_1 - \arg z_2 [2\pi]$$

Corollaire 17. Soient z un complexe non nul et n un entier relatif. On a :

$$\arg z^n \equiv n \arg z [2\pi]$$

Exercice. Soient $a, b, x \in \mathbb{R}$. On se propose de factoriser l'expression $a \cos x + b \sin x$.

Soit $z = a + ib$. Écrivons z sous forme trigonométrique : $z = re^{i\theta}$, où $r \geq 0$ et $\theta \in \mathbb{R}$. Si jamais $z = 0$, prendre $r = 0$ et θ quelconque. On a alors

$$\begin{aligned} a \cos x + b \sin x &= r \cos \theta \cos x + r \sin \theta \sin x \\ &= r \cos(x - \theta) \end{aligned}$$

Par exemple,

$$\begin{aligned} \sqrt{3} \cos x + \sin x &= 2 \left(\frac{\sqrt{3}}{2} \cos x + \frac{1}{2} \sin x \right) \\ &= 2 \left(\cos \frac{\pi}{6} \cos x + \sin \frac{\pi}{6} \sin x \right) \\ &= 2 \cos \left(x - \frac{\pi}{6} \right) \end{aligned}$$

Exercice. Soient $\varphi, \psi \in \mathbb{R}$. On se propose de factoriser l'expression $e^{i\theta} + e^{i\varphi}$.

Posons $\psi = \frac{1}{2}(\theta + \varphi)$ et $\delta = \frac{1}{2}(\theta - \varphi)$. On a alors

$$\begin{aligned} e^{i\theta} + e^{i\varphi} &= e^{i(\psi+\delta)} + e^{i(\psi-\delta)} \\ &= e^{i\psi} e^{i\delta} + e^{i\psi} e^{-i\delta} \\ &= e^{i\psi} (e^{i\delta} + e^{-i\delta}) \\ &= 2 \cos \delta e^{i\psi} \end{aligned}$$

Ainsi,

$$e^{i\theta} + e^{i\varphi} = 2 \cos \frac{\theta - \varphi}{2} \exp \left(i \frac{\theta + \varphi}{2} \right)$$

Exercice. Soient $\varphi, \psi \in \mathbb{R}$. Factoriser l'expression $e^{i\theta} - e^{i\varphi}$.

2.5 Exponentielle complexe

On suppose connues les propriétés de la fonction exponentielle sur $\mathbb{R}$.

Définition 7. Soit $z = x + iy$ un complexe. On appelle exponentielle de z le complexe

$$e^z = e^x e^{iy}$$

Proposition 18. Soit z un complexe. On a

- $e^{\bar{z}} = \overline{e^z}$
- $|e^z| = e^{\operatorname{Re} z}$
- $\arg e^z \equiv \operatorname{Im} z [2\pi]$

Proposition 19. Soient $z, z' \in \mathbb{C}$. On a :

- $e^{z+z'} = e^z e^{z'}$.
- $\frac{1}{e^z} = e^{-z}$.
- $e^{z-z'} = \frac{e^z}{e^{z'}}$.

Proposition 20. Soit $z \in \mathbb{C}$. On a $e^z = 1 \iff z \in 2i\pi\mathbb{Z}$.

Démonstration. Posons $z = x + iy$. Supposons $e^z = 1$. On a $e^x e^{iy} = 1$. En passant au module, on obtient $e^x = 1$ d'où $x = 0$ puisque x est réel. On reporte, on en déduit $e^{iy} = 1$, donc $y \in 2\pi\mathbb{Z}$, comme vu plus haut. On a donc bien $z \in 2i\pi\mathbb{Z}$. La réciproque est immédiate. $\square$

Corollaire 21. Soient $z, z' \in \mathbb{C}$. On a $e^z = e^{z'} \iff z' - z \in 2i\pi\mathbb{Z}$.

Démonstration. C'est évident, puisque $e^z = e^{z'}$ si et seulement si $e^{z-z'} = 1$. $\square$

Proposition 22. Tout nombre complexe non nul a est de la forme e^{z_0} , avec $z_0 \in \mathbb{C}$. Ses antécédents sont alors les $z_0 + 2ik\pi$, avec $k \in \mathbb{Z}$.

Démonstration. On écrit $a = |a|e^{i\theta}$ où θ est un argument de a . Soit $z_0 = \ln|a| + i\theta$. Alors $e^{z_0} = a$. De plus,

$$e^z = a \iff e^z = e^{z_0} \iff e^{z-z_0} = 1$$

c'est à dire $z - z_0 \in 2i\pi\mathbb{Z}$. $\square$

2.6 Application à la linéarisation de sinus et cosinus

Soient m un entier naturel et θ un réel. On écrit

$$\cos^m \theta = \left(\frac{e^{i\theta} + e^{-i\theta}}{2} \right)^m$$

et on développe par la formule du binôme de Newton :

$$\cos^m \theta = \frac{1}{2^m} \sum_{k=0}^m \binom{m}{k} e^{ik\theta} e^{-i(m-k)\theta} = \frac{1}{2^m} \sum_{k=0}^m \binom{m}{k} e^{i(2k-m)\theta}$$

On sait par ailleurs que cette quantité est réelle. Elle est donc égale à sa partie réelle. Ainsi,

$$\cos^m \theta = \frac{1}{2^m} \sum_{k=0}^m \binom{m}{k} \cos(2k - m)\theta$$

La même opération peut bien entendu être réalisée avec un sinus. Il faut alors distinguer suivant la parité de m .

2.7 Opération inverse : délinéarisation ?

On écrit cette fois-ci la formule de Moivre :

$$\cos m\theta + i \sin m\theta = (\cos \theta + i \sin \theta)^m = \sum_{k=0}^m \binom{m}{k} i^k \cos^k \theta \sin^{m-k} \theta$$

Si l'on veut par exemple $\cos m\theta$, il suffit de prendre la partie réelle :

$$\cos m\theta = \sum_{0 \leq 2p \leq m} \binom{m}{2p} (-1)^p \cos^{2p} \theta \sin^{m-2p} \theta$$

De même :

$$\sin m\theta = \sum_{0 \leq 2p+1 \leq m} \binom{m}{2p+1} (-1)^p \cos^{2p+1} \theta \sin^{m-2p-1} \theta$$

3 Résolution d'équations algébriques

3.1 Racine carrée d'un nombre complexe

Définition 8. On appelle racine carrée du complexe a tout complexe z tel que $z^2 = a$.

Proposition 23. Tout nombre complexe non nul possède exactement deux racines carrées.

Exemple. Les racines carrées de i sont $\pm \frac{\sqrt{2}}{2}(1 + i)$.

Démonstration. Soit $a = re^{i\theta}$, $r > 0$. Soit $z = \rho e^{i\varphi}$, $\rho > 0$. Alors, $z^2 = a$ si et seulement si $\rho^2 = r$ et $2\varphi \equiv \theta[2\pi]$, c'est-à-dire $\rho = \sqrt{r}$ et $\varphi \equiv \frac{\theta}{2}[\pi]$. On trouve donc deux racines carrées de a qui sont $\pm \sqrt{r}e^{i\frac{\theta}{2}}$. $\square$

3.2 Méthode algébrique

Soit $a = x + iy$, où $x, y \in \mathbb{R}$. Nous allons calculer explicitement les racines carrées de a , sans passer par les fonctions trigonométriques.

Si $y = 0$, a est réel. Si $a > 0$, les racines carrées de a sont $\pm\sqrt{a}$. Si, au contraire, $a < 0$, les racines carrées de a sont $\pm i\sqrt{-a}$.

Supposons maintenant $y \neq 0$. Soit $z = X + iY$. Alors $z^2 = a$ si et seulement si

$$(1) \begin{cases} X^2 - Y^2 &= x \\ 2XY &= y \end{cases}$$

Supposons (1). On a $z^2 = a$, donc aussi $|z|^2 = |a|$. De plus, XY est du signe de y . On a donc

$$(2) \begin{cases} X^2 - Y^2 &= x \\ X^2 + Y^2 &= |a| \\ XY &\text{du signe de } y \end{cases}$$

Inversement, supposons (2). En additionnant et soustrayant les deux égalités, on obtient

$$X^2 = \frac{1}{2}(|a| + x) \text{ et } Y^2 = \frac{1}{2}(|a| - x)$$

Ainsi,

$$X = \varepsilon \sqrt{\frac{1}{2}(|a| + x)} \text{ et } Y = \varepsilon' \sqrt{\frac{1}{2}(|a| - x)}$$

où $\varepsilon, \varepsilon' \in \{-1, 1\}$. De plus, la condition de signe impose que $\varepsilon\varepsilon'$ est du signe de y . On en déduit que

$$\begin{aligned} 2XY &= 2\varepsilon\varepsilon' \sqrt{\frac{1}{2}(|a| + x)} \sqrt{\frac{1}{2}(|a| - x)} \\ &= 2\varepsilon\varepsilon' \sqrt{\frac{1}{4}(|a|^2 - x^2)} \\ &= 2\varepsilon\varepsilon' \sqrt{\frac{1}{4}y^2} \\ &= \varepsilon\varepsilon'|y| \\ &= y \end{aligned}$$

Ainsi, si on a (2) alors on a (1). En conclusion, (2) est une condition nécessaire et suffisante pour que $z^2 = a$.

Exercice. Calculons les racines carrées de $1+i$ par la méthode algébrique. Soit $z = X + iY$. On a $z^2 = 1 + i$ si et seulement si

$$(2) \begin{cases} X^2 - Y^2 &= 1 \\ X^2 + Y^2 &= \sqrt{2} \\ XY &> 0 \end{cases}$$

On en tire $X^2 = \frac{1}{2}(\sqrt{2} + 1)$ et $Y^2 = \frac{1}{2}(\sqrt{2} - 1)$. Avec les conditions de signe, on en déduit

$$z = \pm \left(\sqrt{\frac{1}{2}(\sqrt{2} + 1)} + i\sqrt{\frac{1}{2}(\sqrt{2} - 1)} \right)$$

3.3 Équation du second degré

Proposition 24. Soient a, b, c trois complexes avec $a \neq 0$. On considère l'équation

$$(E) \quad az^2 + bz + c = 0$$

et son discriminant $\Delta = b^2 - 4ac$. Si $\Delta = 0$, l'équation (E) a pour unique racine racine $z = \frac{-b}{2a}$. Sinon, en appelant δ une racine carrée de Δ , l'équation (E) a deux racines distinctes qui sont $\frac{-b \pm \delta}{2a}$.

Démonstration. On écrit

$$az^2 + bz + c = a\left(z + \frac{b}{2a}\right)^2 - \frac{\Delta}{4a}$$

Ainsi, z est solution de E si et seulement si

$$\left(z + \frac{b}{2a}\right)^2 = \left(\frac{\delta}{2a}\right)^2$$

ou encore

$$z + \frac{b}{2a} = \pm \frac{\delta}{2a}$$

□

Exemple. Résolvons l'équation $z^2 - (1 + 2i)z + i - 1 = 0$. Son discriminant est $\Delta = 1$. Nous avons de la chance, le discriminant aurait pu être non réel. Une racine carrée de Δ est $\delta = 1$. Les solutions de l'équation sont donc $\frac{1}{2}((1 + 2i) + 1) = 1 + i$ et $\frac{1}{2}((1 + 2i) - 1) = i$.

Corollaire 25. Avec les mêmes notations que dans la proposition précédente, le produit des racines de (E) est $\frac{c}{a}$ et la somme des racines est $\frac{-b}{a}$.

Démonstration. Il suffit de prendre les expressions des racines données par la proposition précédente. □

Remarque 10. Lorsque $\Delta = 0$, ce que l'on appelle **les** racines est en réalité **la** racine comptée deux fois. nous parlerons dans cette situation de *racine double*.

Proposition 26. Soient $z, z', s, p \in \mathbb{C}$. On a $z + z' = s$ et $zz' = p$ si et seulement si z et z' sont les racines de l'équation $X^2 - sX + p = 0$.

Démonstration. Soient α et β les racines de cette équation. On a $\alpha + \beta = s$ et $\alpha\beta = p$ d'où l'un des deux sens de l'équivalence. Inversement, supposons $z + z' = s$ et $zz' = p$. z et z' sont évidemment les racines de l'équation $(X - z)(X - z') = 0$. Mais justement, $(X - z)(X - z') = X^2 - sX + p$. $\square$

3.4 Racines nièmes de l'unité

Définition 9. Soit z un complexe et n un entier non nul. On appelle racine nième de z tout complexe Z tel que $Z^n = z$.

Les racines nièmes de 1 sont appelées les racines nièmes de l'unité.

Proposition 27. L'ensemble, noté $\mathcal{U}_n$, des racines nièmes de l'unité est un sous-groupe du groupe $(\mathcal{U}, \times)$ des nombres complexes de module 1.

Démonstration. Si $z^n = 1$, alors $|z|^n = |1| = 1$ donc $|z| = 1$. Ainsi, $\mathcal{U}_n \subset \mathcal{U}$. Le reste de la preuve est facile : stabilité par produit car $(zz')^n = z^n z'^n$ et stabilité par inverse car $(\frac{1}{z})^n = \frac{1}{z^n}$. $\square$

Proposition 28. Il y a exactement n racines nièmes de l'unité. Précisément,

$$\mathcal{U}_n = \left\{ e^{\frac{2ik\pi}{n}}, k \in \llbracket 0, n-1 \rrbracket \right\}$$

Démonstration. Soit $z = e^{i\theta} \in \mathcal{U}$. Alors $z \in \mathcal{U}_n$ si et seulement si $z^n = e^{in\theta} = 1$ c'est à dire $n\theta \in 2\pi\mathbb{Z}$ ou encore $\theta = \frac{2k\pi}{n}$ où $k \in \mathbb{Z}$. Ainsi,

$$\mathcal{U}_n = \{\xi^k, k \in \mathbb{Z}\}$$

où $\xi = e^{\frac{2i\pi}{n}}$. Soit maintenant $k \in \mathbb{Z}$. On effectue la division euclidienne de k par n : $k = nq + r$ où $0 \leq r < n$. On a alors $\xi^k = (\xi^n)^q \xi^r = \xi^r$. Donc,

$$\mathcal{U}_n = \{\xi^k, 0 \leq k < n\}$$

Enfin, si $0 \leq i < j < n$, alors $0 < i - j < n$ donc

$$0 < \frac{2\pi(i-j)}{n} < 2\pi$$

donc $\xi^i \neq \xi^j$. L'ensemble $\mathcal{U}_n$ contient bien exactement n éléments. $\square$

Remarque 11. La somme des racines n èmes de l'unité est

$$1 + \xi + \xi^2 + \cdots + \xi^{n-1} = \frac{\xi^n - 1}{\xi - 1} = 0$$

puisque $\xi^n = 1$.

Exemple. On a $\mathcal{U}_1 = \{1\}$, $\mathcal{U}_2 = \{-1, 1\}$, $\mathcal{U}_4 = \{-1, 1, -i, i\}$. On a $\mathcal{U}_3 = \{1, j, j^2\}$ où $j = e^{2i\pi/3}$. Le nombre j apparaît dans un grand nombre de calculs. Il faut absolument se souvenir que $\frac{1}{j} = \bar{j} = j^2$, $1 + j + j^2 = 0$.

Plus généralement, si l'on dessine $\mathcal{U}_n$ dans le plan (en identifiant le complexe $x + iy$ et le point (x, y)), on obtient le polygone régulier à n côtés inscrit dans le cercle unité, dont l'un des sommets est le point $(1, 0)$.

Exercice. Dessiner $\mathcal{U}_6$.

3.5 Racines n èmes d'un nombre complexe

Proposition 29. Soit z un complexe non nul. Si u est une racine n ème de z , alors l'ensemble des racines n èmes de z est

$$\{u\xi, \xi \in \mathcal{U}_n\}$$

Démonstration. On a $w^n = z$ si et seulement si $w^n = u^n$, c'est à dire $\frac{w}{u} \in \mathcal{U}_n$. $\square$

Corollaire 30. Soit z un complexe non nul de module r et d'argument θ . Les racines n èmes de z sont les complexes

$$Z_k = \sqrt[n]{r} e^{i(\frac{\theta}{n} + \frac{2k\pi}{n})}, k \in [0, n-1]$$

Démonstration. C'est clair, puisque $\sqrt[n]{r} e^{i\theta/n}$ est clairement une racine n ème de z . $\square$

Exercice. Calculer les racines cubiques de $2i$.

On a $2i = 2e^{i\pi/2}$. Une racine cubique de $2i$ est donc

$$u = \sqrt[3]{2} e^{i\pi/6}$$

Les racines cubiques de $2i$ sont donc

$$u, ju, j^2u$$

où $j = e^{2i\pi/3}$.

4 Interprétations géométriques

4.1 Points, vecteurs

Nous avons déjà parlé de l'afixe d'un *point* du plan, qui est un nombre complexe. Nous pouvons également parler de l'afixe d'un *vecteur* du plan. Si $A(a)$ et $B(b)$ sont deux points du plan d'affixes $a, b \in \mathbb{C}$, l'afixe du vecteur $\overrightarrow{AB}$ est le nombre complexe $b - a$. Nous noterons éventuellement $\overrightarrow{u}(c)$ lorsque $c \in \mathbb{C}$ est l'afixe du vecteur $\overrightarrow{u}$.

4.2 Somme, produit par un réel

Ces opérations sur les complexes correspondent dans le plan euclidien aux opérations usuelles sur les vecteurs. Si $\overrightarrow{u}$ et $\overrightarrow{v}$ ont pour affixes z et z' , alors $\overrightarrow{u} + \overrightarrow{v}$ a pour affixe $z + z'$. Si, de plus, $\lambda \in \mathbb{R}$, alors $\lambda \overrightarrow{u}$ a pour affixe λz .

Ces opérations fournissent d'importantes transformations du plan.

Définition 10. Soit $\overrightarrow{u} \in \mathbb{R}^2$. L'application $\tau_{\overrightarrow{u}} : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ définie par

$$M \mapsto M + \overrightarrow{u}$$

est appelée la translation de vecteur u .

En termes d'affixes, si $\overrightarrow{u}(z_0)$, cette application est décrite par

$$M(z) \mapsto M(z + z_0)$$

Définition 11. Soit $\Omega \in \mathbb{R}^2$. Soit $\lambda \in \mathbb{R}$. L'application $h : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ définie par

$$M \mapsto \Omega + \lambda \overrightarrow{\Omega M}$$

est appelée l'homothétie de centre Ω et de rapport λ .

En termes d'affixes, si $\Omega(z_0)$, cette application est décrite par

$$M(z) \mapsto M(z_0 + \lambda(z - z_0))$$

4.3 Module, Argument

Module et argument représentent respectivement des distances et des angles dans le plan euclidien. Ainsi, soient $A(a)$ et $B(b)$ deux points du plan. On a

$$d(A, B) = |b - a|$$

Si $A(a)$, $B(b)$ et $C(c)$ sont trois points distincts, l'angle orienté entre les vecteurs $\overrightarrow{AB}$ et $\overrightarrow{AC}$ est donné par

$$\theta = \arg \frac{c - a}{b - a}$$

4.4 Produit

C'est l'opération la plus intéressante. Commençons par le produit d'un nombre complexe par un complexe de module 1 : Soit $z \in \mathbb{C}$ et $\theta \in \mathbb{R}$. On a alors :

- $|ze^{i\theta}| = |z|$
- $\arg(ze^{i\theta}) \equiv \arg z + \theta \pmod{2\pi}$

Définition 12. Soient $\Omega(z_0)$ un point du plan et θ un réel. L'application

$$M(z) \mapsto M(z_0 + (z - z_0)e^{i\theta})$$

est appelée la rotation de centre Ω et d'angle θ .

Plus, généralement, la multiplication par un nombre complexe non nul est la composée des deux opérations suivantes :

- Produit par un nombre complexe de module 1 (rotation)
- Produit par un réel non nul (homothétie)

Définition 13. Soient $\Omega(z_0)$ un point du plan, λ un réel non nul, et θ un réel. L'application

$$M(z) \mapsto M(z_0 + \lambda(z - z_0)e^{i\theta})$$

est appelée la similitude (directe) de centre Ω , de rapport λ et d'angle θ .

Plus généralement, on appelle similitude (directe) du plan toute application

$$M(z) \mapsto M(az + b)$$

où $a, b \in \mathbb{C}, a \neq 0$.

Proposition 31. *Les similitudes directes sont*

- *Les translations.*
- *Les similitudes « à centre ».*

Démonstration. Pour alléger les notations, nous confondons dans ce qui suit les points et leurs affixes.

Soient $a, b \in \mathbb{C}$, $a \neq 0$. Soit

$$f : z \mapsto az + b$$

la similitude correspondante. Il y a deux cas à considérer.

Si $a = 1$, l'application f est une translation.

Supposons maintenant $a \neq 1$. Soit $\omega \in \mathbb{C}$. On a $f(\omega) = \omega$ si et seulement si

$$\omega = \frac{b}{1-a}$$

L'application f a donc un unique point fixe. Mais alors, pour tout $z \in \mathbb{C}$,

$$f(z) - \omega = f(z) - f(\omega) = a(z - \omega)$$

f est donc la similitude de centre ω , de rapport $|a|$ et d'angle $\arg a$. $\square$

Proposition 32. *Les similitudes conservent les angles.*

Démonstration. On pose $f(z) = az + b$, avec $a \neq 0$. Soient $u, v, w \in \mathbb{C}$ distincts. On a alors

$$\frac{f(w) - f(u)}{f(v) - f(u)} = \frac{aw + b - (au + b)}{av + b - (au + b)} = \frac{w - u}{v - u}$$

Les arguments de ces quantités sont donc égaux, et f conserve les angles. $\square$

On peut en fait prouver le résultat plus général suivant :

Proposition 33. *Soit $f : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ une application injective. Alors, f conserve les angles si et seulement si f est une similitude directe.*

Démonstration. Seul le sens direct est à prouver. Supposons donc que f est une injection qui conserve les angles. Posons

$$g(z) = \frac{f(z) - f(0)}{f(1) - f(0)}$$

L'application g est encore une injection, et conserve toujours les angles. De plus, $g(0) = 0$ et $g(1) = 1$.

Soit $z \in \mathbb{C} \setminus \mathbb{R}$. On considère l'angle formé par 1, 0 et z . Cet angle étant conservé par g , on a donc

$$\arg \frac{g(z) - 0}{1 - 0} = \arg \frac{z - 0}{1 - 0}$$

Donc, $g(z) = \lambda z$, où $\lambda \in \mathbb{R}_+^*$. De même, l'angle formé par 0, 1 et z est conservé, donc

$$\arg \frac{g(z) - 1}{0 - 1} = \arg \frac{z - 1}{0 - 1}$$

Il existe donc $\mu \in \mathbb{R}_+^*$ tel que $1 - g(z) = \mu(1 - z)$. D'où

$$\begin{aligned} 1 - \lambda z &= \mu(1 - z) \\ 1 - \mu &= (\lambda - \mu)z \end{aligned}$$

Mais comme on a supposé z non réel, on a nécessairement $\lambda = \mu = 1$, d'où $g(z) = z$.

Si z est réel, on recommence avec $0, i$ et z (on sait maintenant que $g(i) = i$). Finalement, on a $g(z) = z$ pour tout complexe z . Donc, $f(z) = az + b$, où $a = f(1) - f(0) \neq 0$ et $b = f(0)$. L'application f est donc une similitude directe. $\square$

Chapitre 6

Nombres réels

1 Le corps des réels

1.1 Corps ordonnés

Définition 1. Soit $\mathbb{K}$ un corps. On suppose que $\mathbb{K}$ est partitionné en trois parties disjointes : $\mathbb{K} = \mathcal{P} \cup \mathcal{N} \cup \{0\}$. On dit que $\mathbb{K}$ est un corps ordonné lorsque

- $\mathcal{P}$ est stable pour l'addition et la multiplication.
- Pour tout élément x de $\mathbb{K}$, on a $x \in \mathcal{P}$ si et seulement si $-x \in \mathcal{N}$.

Exemple. $\mathbb{Q}$, $\mathbb{R}$, $\{a + b\sqrt{2}, a, b \in \mathbb{Q}\}$ sont des corps ordonnés lorsqu'on prend pour $\mathcal{P}$ l'ensemble des éléments strictement positifs du corps et pour $\mathcal{N}$ l'ensemble des éléments strictement négatifs.

On suppose jusqu'à la fin de cette section que $\mathbb{K}$ est un corps ordonné.

Proposition 1. Soient $x, y \in \mathbb{K}$. On a

- $x, y \in \mathcal{P} \implies xy \in \mathcal{P}$
- $x, y \in \mathcal{N} \implies xy \in \mathcal{P}$
- $x \in \mathcal{P}, y \in \mathcal{N} \implies xy \in \mathcal{N}$

Démonstration. Par exemple, si $x \in \mathcal{P}$ et $y \in \mathcal{N}$, on a $-(xy) = x(-y) \in \mathcal{P}$ donc $xy \in \mathcal{N}$. $\square$

Proposition 2. On a

- $\forall x \in \mathbb{K}, x^2 \in \mathcal{P}$
- $1 \in \mathcal{P}, -1 \in \mathcal{N}$
- $\forall x \in \mathbb{K}, x \in \mathcal{P} \iff \frac{1}{x} \in \mathcal{P}$

Démonstration. Si $x \in \mathcal{P}$, on a $x^2 = xx \in \mathcal{P}$. De même si $x \in \mathcal{N}$. Pour le second point, $1 = 1^2 \in \mathcal{P}$. Pour le dernier point, on a $x \frac{1}{x} = 1 \in \mathcal{P}$ donc x et $\frac{1}{x}$ sont tous les deux dans $\mathcal{P}$ ou tous les deux dans $\mathcal{N}$. $\square$

Proposition 3. On a $\forall x, y \in \mathcal{N}, x + y \in \mathcal{N}$

Démonstration. Car $x + y = -((-x) + (-y))$. $\square$

Définition 2. Pour $x, y \in \mathbb{K}$, on pose

$$x < y \iff y - x \in \mathcal{P}$$

et

$$x \leq y \iff x < y \text{ ou } x = y$$

On définit de même $x > y$ et $x \geq y$.

Remarque 1. Pour tout élément x de $\mathbb{K}$, on a donc $x > 0$ si et seulement si $x \in \mathcal{P}$, et $x < 0$ si et seulement si $x \in \mathcal{N}$. Les propositions ci-dessus nous indiquent donc que la règle des signes est vraie dans $\mathbb{K}$, que tout carré est positif, que l'inverse d'un élément positif est positif, etc.

Proposition 4. La relation $\leq$ est un ordre total sur $\mathbb{K}$.

Démonstration.

- La réflexivité est évidente.
- Si $x < y$ alors $y - x \in \mathcal{P}$ donc $x - y \notin \mathcal{P}$. Ainsi, si $x \leq y$ et $y \leq x$, on a $x = y$ d'où l'antisymétrie.
- Si $x < y$ et $y < z$ alors $z - x = (z - y) + (y - x) \in \mathcal{P}$ donc $x < z$ d'où la transitivité.
- Soient enfin $x, y \in \mathbb{K}$. $y - x$ est soit dans $\mathcal{N}$ auquel cas $y < x$, soit dans $\mathcal{P}$ auquel cas $x < y$ soit nul auquel cas $x = y$. L'ordre est total.

□

Proposition 5. Pour tous éléments x, y, z de $\mathbb{K}$, on a

- $x < y \iff x + z < y + z$
- Si $z > 0$, alors $x < y \iff xz < yz$
- Si $z < 0$, alors $x < y \iff xz > yz$

Démonstration. Supposons $x < y$. Alors

$$(y + z) - (x + z) = y - x \in \mathcal{P}$$

donc $x + z < y + z$. La réciproque et les autres points se démontrent de façon identique. □

Remarque 2. On ne peut pas munir $\mathbb{C}$ d'une structure de corps ordonné. En effet $-1 = i^2$, or $-1 < 0$ et tout carré est positif.

1.2 Valeur absolue

Définition 3. Pour tout $x \in \mathbb{K}$, on appelle valeur absolue de x , et on note $|x|$, l'élément de $\mathbb{K}$

$$|x| = \max(-x, x)$$

Remarque 3. On bien sûr $|x| = x$ si $x \geq 0$ et $|x| = -x$ si $x \leq 0$.

Proposition 6. La valeur absolue vérifie les propriétés suivantes :

- $\forall x \in \mathbb{K}, |x| \geq 0$.
- $\forall x \in \mathbb{K}, |x| = 0 \iff x = 0$.
- $\forall x, y \in \mathbb{K}, |xy| = |x||y|$.
- $\forall x, y \in \mathbb{K}, |x + y| \leq |x| + |y|$.
- $\forall x, y \in \mathbb{K}, ||x| - |y|| \leq |x - y|$.

Démonstration. Seules les inégalités sont non triviales. Soient $x, y \in \mathbb{K}$. On a

$$-|x| \leq x \leq |x| \text{ et } -|y| \leq y \leq |y|$$

En ajoutant on obtient

$$-|x| - |y| \leq x + y \leq |x| + |y|$$

Passant à l'opposé dans ces inégalités, on a aussi

$$-|x| - |y| \leq -(x + y) \leq |x| + |y|$$

Mais $|x + y|$ est l'un des deux réels $x + y$, $-(x + y)$. Donc

$$|x + y| \leq |x| + |y|$$

Passons à l'autre inégalité. On écrit $x = (x - y) + y$. De là,

$$|x| = |(x - y) + y| \leq |x - y| + |y|$$

d'où

$$|x| - |y| \leq |x - y|$$

En échangeant les rôles de x et y , on a aussi

$$|y| - |x| \leq |y - x| = |x - y|$$

Comme $||x| - |y||$ est l'un des deux réels $|x| - |y|$, $|y| - |x|$, on en déduit

$$||x| - |y|| \leq |x - y|$$

□

2 Borne supérieure

2.1 Notion de borne supérieure

Définition 4. Soit $(E, \leq)$ un ensemble ordonné. Soit $A \subset E$. On appelle borne supérieure de A le plus petit élément, lorsqu'il existe, de l'ensemble des majorants de A .

On définit de même la notion de borne inférieure.

En d'autres termes, la borne supérieure M de l'ensemble A est caractérisée par

- $\forall x \in A, x \leq M$.
- $\forall y < M, \exists x \in A, y < x$.

Notation. On note A^+ (resp. A^-) l'ensemble des majorants (resp. minorants) de A . On note $\sup A$ la borne supérieure de A , lorsqu'elle existe. De même on note $\inf A$ sa borne inférieure.

Lorsque la borne supérieure de A existe, on a donc

$$\sup A = \min A^+$$

De même, si la borne inférieure de A existe, alors

$$\inf A = \max A^-$$

Proposition 7. Si A possède un plus grand élément, alors A possède une borne supérieure, et $\sup A = \max A$. La réciproque est fausse.

Démonstration. Soit $a = \max A$. Soit A^+ l'ensemble des majorants de A . Soit $y \in E$. Supposons $a \leq y$. Alors, pour tout $x \in A$, $x \leq a \leq y$, donc $x \leq y$. Ainsi, $y \in A^+$. Inversement, soit $y \in A^+$. Comme $a \in A$, on a $a \leq y$. Ainsi, $y \in A^+$ si et seulement si $a \leq y$. Donc, A^+ a un plus petit élément : a . $\square$

Exemple.

- On se place sur $\mathbb{R}$ muni de l'ordre usuel. On a

$$\sup[0, 1] = 1$$

$$\sup[0, 1[= 1$$

$\mathbb{R}_+$ n'a pas de borne supérieure.

- On se place dans $\mathbb{R}^{\mathbb{R}}$ muni de l'ordre usuel. Soient $f, g \in \mathbb{R}^{\mathbb{R}}$, et $A = \{f, g\}$. L'ensemble A n'a pas, en général, de plus grand élément (prendre par exemple $f(x) = \sin x$ et $g(x) = \cos x$). En revanche, A a une borne supérieure. C'est la fonction h définie par $h(x) = \max(f(x), g(x))$ pour tout réel x . Cette fonction est notée $\sup(f, g)$.

Remarque 4. Soit A une partie de l'ensemble $\mathbb{R}$. Supposons que A possède une borne supérieure. Alors :

- A est évidemment majorée puisque, par définition de la borne sup, l'ensemble des majorants de A possède un plus petit élément (et donc, possède un élément!).
- A est non vide. En effet, tout réel majore l'ensemble vide, et $\mathbb{R}$ n'a pas de plus petit élément.

Ainsi, si A possède une borne supérieure, alors A est non vide et majorée. Nous verrons que dans le cas des réels la réciproque est vraie.

Exemple. On se place dans $\mathbb{Q}$. Soit

$$A = \{x \in \mathbb{Q}_+, x^2 \leq 2\}$$

Montrons que A est non vide et majorée, mais que A ne possède pas de borne supérieure. L'ensemble A est clairement non vide puisque $1 \in A$. De même, pour tout $x \in A$, on a

$$x^2 \leq 2 \leq 2^2$$

donc $x \leq 2$ et A est majoré. Supposons que A possède une borne supérieure α . Nous allons montrer que $\alpha^2 = 2$, ce qui est impossible car il n'existe pas de rationnel dont le carré vaut 2.

- Tout d'abord, supposons $\alpha^2 < 2$. Soit h un rationnel tel que

$$0 < h < \frac{2 - \alpha^2}{2\alpha + 1}$$

On a

$$(\alpha + h)^2 = \alpha^2 + h(2\alpha + h) < \alpha^2 + h(2\alpha + 1) < 2$$

Ainsi, $\alpha + h \in A$ ce qui n'est pas possible puisque α majore A . Donc, $\alpha^2 \geq 2$.

- Supposons maintenant que $\alpha^2 > 2$. Soit h un rationnel vérifiant

$$0 < h < \frac{\alpha^2 - 2}{2\alpha}$$

On a

$$(\alpha - h)^2 = \alpha^2 - 2\alpha h + h^2 > \alpha^2 - 2\alpha h > \alpha^2 - (\alpha^2 - 2) = 2$$

Pour tout $x \in A$, $x^2 \leq 2 \leq (\alpha - h)^2$, donc $x \leq \alpha - h$. Ceci n'est pas possible. En effet $\alpha - h$ ne majore pas A puisque α est le plus petit majorant de A .

Remarque 5. Dans tout ce qui suit, on ne parlera que de borne sup. Il existe évidemment des propriétés identiques pour la borne inférieure, que l'on obtient par exemple en passant à l'opposé.

2.2 Le Théorème fondamental

Théorème 8. *Il existe un corps $\mathbb{K}$ totalement ordonné contenant $\mathbb{Q}$, et tel que toute partie non vide et majorée de $\mathbb{K}$ possède une borne supérieure. Un tel corps est unique à isomorphisme près.*

On choisit un tel corps, et on l'appelle $\mathbb{R}$. Lequel ? Cela n'a pas d'importance, puisque tous ces corps sont isomorphes.

Remarque 6. Reprenons l'exemple du paragraphe précédent, mais dans $\mathbb{R}$ au lieu de $\mathbb{Q}$. Soit donc $A = \{x \in \mathbb{R}_+, x^2 \leq 2\}$. A est non vide et majorée, donc A possède une borne supérieure α . La même démonstration que ci-dessus nous prouve que $\alpha^2 = 2$. Nous avons donc prouvé l'existence d'un réel $\alpha > 0$ tel que $\alpha^2 = 2$. Si l'on remplace 2 par un réel $t > 0$ quelconque une démonstration identique nous dira que tout réel positif possède une racine carrée positive.

Exercice. Montrer l'unicité de la racine carrée positive d'un réel positif.

2.3 La droite réelle achevée

Pour simplifier certaines discussions dans le cours, on rajoute ici à $\mathbb{R}$ deux éléments, notés $+\infty$ et $-\infty$, pour fabriquer ce que l'on appelle la droite réelle achevée, notée $\overline{\mathbb{R}}$. On va voir que toute partie de $\overline{\mathbb{R}}$ possède une borne supérieure. En contrepartie, les opérations ne sont plus que partiellement définies.

Définition 5. On appelle droite numérique achevée l'ensemble $\overline{\mathbb{R}} = \mathbb{R} \cup \{-\infty, +\infty\}$ où $-\infty$ et $+\infty$ sont deux objets sans signification.

On étend ensuite les opérations et l'ordre sur $\mathbb{R}$ à $\overline{\mathbb{R}}$.

Définition 6. On étend l'ordre de $\mathbb{R}$ à $\overline{\mathbb{R}}$ en posant, pour tout réel x , $-\infty < x < +\infty$.

Définition 7. On étend partiellement l'addition de $\mathbb{R}$ à $\overline{\mathbb{R}}$ par le tableau suivant :

+	$-\infty$	$y \in \mathbb{R}$	$+\infty$
$-\infty$	$-\infty$	$-\infty$	<i>N.D.</i>
$x \in \mathbb{R}$	$-\infty$	$x + y$	$+\infty$
$+\infty$	<i>N.D.</i>	$+\infty$	$+\infty$

Définition 8. On étend partiellement la multiplication de $\mathbb{R}$ à $\overline{\mathbb{R}}$ par le tableau suivant :

$\times$	$-\infty$	$y \in \mathbb{R}_-^*$	0	$y \in \mathbb{R}_+^*$	$+\infty$
$-\infty$	$+\infty$	$+\infty$	<i>N.D.</i>	$-\infty$	$-\infty$
$x \in \mathbb{R}_-^*$	$+\infty$	xy	0	xy	$-\infty$
0	<i>N.D.</i>	0	0	0	<i>N.D.</i>
$x \in \mathbb{R}_+^*$	$-\infty$	xyy	0	xy	$+\infty$
$+\infty$	$-\infty$	$-\infty$	<i>N.D.</i>	$+\infty$	$+\infty$

Remarque 7. On peut également convenir que l'inverse des infinis est 0. En revanche l'inverse de 0 pose problème d'un point de vue algébrique à cause d'une ambiguïté de signe.

Proposition 9. Toute partie de $\overline{\mathbb{R}}$ possède une borne supérieure.

Démonstration. Soit $A \subset \overline{\mathbb{R}}$.

- Si $A = \emptyset$, alors $\sup A = -\infty$. Sinon :
- Si $+\infty \in A$, alors A possède un plus grand élément, qui est aussi sa borne supérieure : $+\infty$.
- Si $A = \{-\infty\}$, alors $\sup A = -\infty$.
- Sinon $A \setminus \{-\infty\}$ est une partie non vide de $\mathbb{R}$. Si cette partie est majorée dans $\mathbb{R}$ elle possède une borne supérieure réelle. Sinon, $\sup A = +\infty$.

□

Les cas intéressants pour nous sont ceux où A est une partie non vide de $\mathbb{R}$: Si A est majorée, alors A possède une borne supérieure réelle. Si A n'est pas majorée, alors $\sup A = +\infty$.

Remarque 8. Soit A une partie non vide de $\overline{\mathbb{R}}$. Soit $a \in A$. Alors, les majorants de A sont plus grands que a , les minorants de A sont plus petits que a , et donc

$$\inf A \leq \sup A$$

2.4 Propriété d'Archimède

Proposition 10. Soient a et b deux réels positifs, avec $b \neq 0$. Il existe alors un entier naturel n tel que $nb > a$.

Démonstration. On suppose le contraire. L'ensemble

$$E = \{nb, n \in \mathbb{N}\}$$

est alors une partie de $\mathbb{R}$, non vide, et majorée par a . Elle possède une borne supérieure M . Mais alors, $M - b$ n'est pas un majorant de E . On peut donc trouver un entier n tel que $nb > M - b$. D'où $(n + 1)b > M$ ce qui contredit le fait que M majore E . $\square$

Corollaire 11. Soient a et b deux réels, $b > 0$. Il existe un unique entier relatif n tel que

$$nb \leq a < (n + 1)b$$

Démonstration.

- Supposons d'abord $a \geq 0$. Soit E l'ensemble des entiers naturels n tels que $nb > a$. C'est une partie de $\mathbb{N}$, non vide d'après la proposition précédente. Donc E a un plus petit élément m . L'entier $n = m - 1$ répond à la question.
- Supposons maintenant $a < 0$. Alors, $-a > 0$ et il existe donc un entier n tel que $nb \leq -a < (n + 1)b$. Si $a \neq nb$, l'entier $-n - 1$ répond à la question. Sinon, l'entier $-n$ répond à la question. Pour l'unicité, on suppose que deux entiers m et n conviennent. On a $nb \leq a < (m + 1)b$ d'où $n < m + 1$ et donc $n \leq m$. De même $m \leq n$ et donc $m = n$.

$\square$

Une application importante de ce théorème est obtenue avec $b = 1$.

2.5 Partie entière, Approximations décimales

Définition 9. Soit $a \in \mathbb{R}$. On appelle partie entière de a l'unique entier relatif n tel que $n \leq a < n + 1$. On note $E(a)$ ou $[a]$ la partie entière de a .

Remarque 9. La partie entière de a est caractérisée par $[a] \in \mathbb{Z}$ et $[a] \leq a < [a] + 1$.

Remarque 10. On peut de la même façon définir le plafond de a , le plus petit entier supérieur ou égal à a . Il est caractérisé par $\lceil a \rceil \in \mathbb{Z}$ et $\lceil a \rceil < a \leq \lceil a \rceil + 1$.

Proposition 12. Soit a un nombre réel. Il existe un unique entier relatif p tel que

$$\frac{p}{10^n} \leq a < \frac{p}{10^n} + \frac{1}{10^n}$$

Démonstration. On a $[10^n a] \leq 10^n a < [10^n a] + 1$. En posant $p = E(10^n a)$, on a le résultat. $\square$

Remarque 11. L'ensemble

$$\mathcal{D} = \left\{ \frac{p}{10^n}, p \in \mathbb{Z} \right\}$$

est appelé *l'ensemble des nombres décimaux*. On vérifie facilement que la somme et le produit de deux nombres décimaux sont encore des nombres décimaux. En revanche, l'inverse d'un nombre décimal non nul n'est pas forcément décimal.

Définition 10. Soit a un nombre réel et $n \in \mathbb{N}$. On appelle valeur approchée de a à 10^{-n} près par défaut tout réel x tel que $x \leq a \leq x + 10^{-n}$.

On définit de même une valeur approchée de a à 10^{-n} près par excès comme tout réel x tel que $x - 10^{-n} \leq a \leq x$.

Si on réinterprète le résultat de la proposition précédente en termes de valeurs approchées, on obtient donc

Proposition 13. Le nombre décimal $\frac{\lfloor 10^n a \rfloor}{10^n}$ est une valeur approchée de a à 10^{-n} près par défaut. Le nombre décimal $\frac{\lfloor 10^n a \rfloor}{10^n} + \frac{1}{10^n}$ est une valeur approchée de a à 10^{-n} près par excès.

Exemple.

- 3.14 est une valeur approchée de π à 10^{-2} près par défaut.
- 1.415 est une valeur approchée de $\sqrt{2}$ à 10^{-3} près par excès.

2.6 Densité des rationnels et des irrationnels

Définition 11. Soit $A \subset \mathbb{R}$. On dit que A est *dense* dans $\mathbb{R}$ lorsque pour tous $x, y \in \mathbb{R}$ tels que $x < y$, il existe $a \in A$ tel que $x < a < y$.

Remarque 12. Soit A une partie dense de $\mathbb{R}$. Soient $x, y \in \mathbb{R}$ tels que $x < y$. Il existe $a_1 \in A$ tel que $x < a_1 < y$. On peut réutiliser la densité avec x et a_1 : il existe $a_2 \in A$ tel que $x < a_2 < a_1 < y$. Etc. Entre deux réels distincts, il existe donc une infinité d'éléments de A .

Proposition 14.

- $\mathbb{Q}$ est dense dans $\mathbb{R}$.
- $\mathbb{R} \setminus \mathbb{Q}$ est dense dans $\mathbb{R}$.

Démonstration. Soient $x, y \in \mathbb{R}$ tels que $x < y$.

- Par la propriété d'Archimède, il existe $q \in \mathbb{N}^*$ tel que $q > \frac{1}{y-x}$, d'où

$$0 < \frac{1}{q} < y - x$$

Toujours d'après Archimède, il existe $p \in \mathbb{Z}$ tel que

$$\frac{p-1}{q} \leq x < \frac{p}{q}$$

La première inégalité s'écrit aussi

$$\frac{p}{q} \leq x + \frac{1}{q} < x + (y - x) = y$$

Donc $x < \frac{p}{q} < y$.

- Appliquons ce qui précède à $x + \sqrt{2}$ et $y + \sqrt{2}$. Il existe $a \in \mathbb{Q}$ tel que

$$x + \sqrt{2} < a < y + \sqrt{2}$$

et donc

$$x < a - \sqrt{2} < y$$

Or, $a - \sqrt{2} \in \mathbb{R} \setminus \mathbb{Q}$.

□

Remarquons pour terminer que la propriété de densité peut s'écrire de bien des manières.

Proposition 15. Soit $A \subset \mathbb{R}$. On a équivalence entre

- (i) A est dense dans $\mathbb{R}$.
- (ii) Pour tout $x \in \mathbb{R}$, pour tout $\varepsilon > 0$, il existe $a \in A$ tel que $|x - a| \leq \varepsilon$.
- (iii) Tout réel est la limite d'une suite d'éléments de A .

Démonstration. Supposons (i). Soient $x \in \mathbb{R}$ et $\varepsilon > 0$. Par la densité de A , il existe $a \in A$ tel que $x - \varepsilon < a < x + \varepsilon$. D'où (ii).

Supposons (ii). Pour tout $n \geq 1$, il existe $a_n \in A$ tel que $|x - a_n| \leq \frac{1}{n}$. Nous verrons dans le chapitre sur les suites que cela entraîne la convergence de a_n vers x lorsque n tend vers l'infini. D'où (iii).

Supposons (iii). Soient $x, y \in \mathbb{R}$ tels que $x < y$. Il existe une suite $(a_n)_{n \in \mathbb{N}}$ d'éléments de A qui converge vers $z = \frac{x+y}{2}$. Nous verrons dans le chapitre sur les suites que cela entraîne que pour tout entier n suffisamment grand,

$$|a_n - z| < \frac{y - x}{2}$$

Pour un tel entier n , on a alors

$$x = \frac{x+y}{2} - \frac{y-x}{2} < a_n < \frac{x+y}{2} + \frac{y-x}{2} = y$$

D'où (i). $\square$

Exercice. Montrer que l'ensemble des nombres décimaux est dense dans $\mathbb{R}$.

3 Intervalles

3.1 Notion d'intervalle

On distingue diverses familles d'intervalles de $\mathbb{R}$:

- Les intervalles bornés, qui peuvent être de 4 sortes : $[a, b]$, $[a, b[$, $]a, b]$, ou $]a, b[$, où a et b sont deux réels. L'ensemble vide est donc a priori un intervalle. Les singletons aussi. Un cas important est celui des intervalles de la forme $[a, b]$, appelés aussi segments.
- Les intervalles non bornés, qui peuvent être de 5 types : $] - \infty, b]$, $] - \infty, b[$, $]a, +\infty[$, $[a, +\infty[$ et $] - \infty, +\infty[= \mathbb{R}$.

Parmi les intervalles, certains sont ouverts, d'autres sont fermés. Précisément, $\emptyset$ et $\mathbb{R}$ sont à la fois ouverts et fermés. À part ces deux cas bizarres, sont ouverts les intervalles du type $]a, b[$, $] - \infty, b[$, $]a, +\infty[$. Sont fermés les intervalles du type $[a, b]$, $] - \infty, b]$, $[a, +\infty[$.

Les intervalles du type $[a, b]$, où $-\infty < a \leq b < +\infty$ jouent un rôle particulier. On les appelle des *segments*.

3.2 Parties convexes

Définition 12. Soit A une partie de $\mathbb{R}$. On dit que A est convexe lorsque pour tous éléments x et y de A , avec $x \leq y$, le segment $[x, y]$ est inclus dans A .

Proposition 16. Soit A une partie de $\mathbb{R}$. Alors A est convexe si et seulement si A est un intervalle.

Démonstration. Il est évident que tout intervalle est convexe (10 cas).

Inversement, soit A une partie convexe de $\mathbb{R}$. Si A est vide, c'est bien un intervalle. Supposons donc A non vide. Prenons par exemple le cas où A est minorée et pas majorée. Alors A possède une borne inférieure, que l'on va noter a . On va prouver que $A =]a, +\infty[$, ouvert ou fermé en a .

- Montrons que $]a, +\infty[\subset A$: soit $z > a$. Par maximalité de a , z n'est donc pas un minorant de A . Il existe ainsi $x \in A$ tel que $x < z$. z n'est pas non plus majorant de A , vu que A n'est pas majorée. Donc, il existe $y \in A$ tel que $z < y$. Mais A est convexe, donc $[x, y] \subset A$. Or, $z \in [x, y]$, donc $z \in A$.
- Montrons que $A \subset [a, +\infty[$: a minore A , donc pour tout z de A , on a $a \leq z$. D'où le résultat.

□

Chapitre 7

Suites réelles et complexes

Une suite à valeurs dans l'ensemble E est une application $u : \mathbb{N} \rightarrow E$. Dans ce chapitre nous considérerons tout d'abord des suites à valeurs dans $\mathbb{R}$. Puis nous regarderons si les résultats obtenus peuvent être étendus aux suites à valeurs dans $\mathbb{C}$.

L'image de l'entier n par la suite u est appelée le n ième terme de la suite. Cette image est notée très souvent u_n au lieu de $u(n)$.

La suite u est aussi notée $(u_n)_{n \in \mathbb{N}}$.

On considère parfois des suites « démarrant au rang n_0 », c'est à dire des applications $u : [n_0, +\infty[\rightarrow E$.

1 Limite d'une suite

1.1 Limite réelle

Définition 1. Soit u une suite réelle. Soit $\ell \in \mathbb{R}$. On dit que la suite u tend vers ℓ lorsque

$$\forall \varepsilon > 0, \exists N \in \mathbb{N}, \forall n \in \mathbb{N}, n \geq N \implies |u_n - \ell| \leq \varepsilon$$

Le réel ℓ est alors appelé une limite de la suite u .

Notation. On note $u \longrightarrow \ell$ ou, de façon plus compliquée, $u_n \xrightarrow[n \rightarrow \infty]{} \ell$.

Exemple. La suite de terme général

$$u_n = \frac{3n - 1}{n^2 + 5}$$

tend vers 0. Soit pour cela $\varepsilon > 0$. Alors, $|u_n| \leq \varepsilon$ dès que $\frac{3}{n} \leq \varepsilon$, puisque $3n - 1 \leq 3n$ et $n^2 + 5 \geq n^2$. Donc, en posant

$$N = \left\lfloor \frac{3}{\varepsilon} \right\rfloor + 1$$

on a pour tout $n \geq N$ l'inégalité $|u_n| \leq \varepsilon$.

Exemple. Montrons que la suite de terme général

$$u_n = \frac{2n - 1}{3n + 5}$$

tend vers $\frac{2}{3}$.

Pour cela, remarquons que pour tout $n \in \mathbb{N}$,

$$u_n - \frac{2}{3} = \frac{2n - 1}{3n + 5} - \frac{2}{3} = -\frac{13}{3(3n + 5)}$$

Ainsi,

$$\left| u_n - \frac{2}{3} \right| = \frac{13}{3(3n+5)} \leq \frac{18}{9n} = \frac{2}{n}$$

Soit $\varepsilon > 0$. Soit $N \in \mathbb{N}$ tel que $N \geq \frac{2}{\varepsilon}$. On a alors pour tout $n \geq N$,

$$\left| u_n - \frac{2}{3} \right| \leq \frac{2}{n} \leq \frac{2}{N} \leq \varepsilon$$

Voici une suite d'une grande importance.

Proposition 1. *Soit q un réel vérifiant $|q| < 1$. Alors, q^n tend vers 0 lorsque n tend vers l'infini.*

Démonstration. Soit $\varepsilon > 0$. Posons $\frac{1}{|q|} = 1 + h$, où $h > 0$. Alors,

$$\frac{1}{|q|^n} = (1 + h)^n \geq 1 + nh$$

Donc, $|q^n| \leq \varepsilon$ dès que

$$\frac{1}{|q|^n} \geq \frac{1}{\varepsilon}$$

c'est à dire dès que

$$n \geq \frac{\frac{1}{\varepsilon} - 1}{h}$$

□

Proposition 2. *Une suite possède au plus une limite.*

Démonstration. Soit u une suite. Supposons que u tend vers ℓ et aussi vers ℓ' . Soit $\varepsilon > 0$. Il existe un entier N' tel que pour tout $n \geq N'$,

$$|u_n - \ell| \leq \varepsilon/2$$

et un entier N'' tel que pour tout $n \geq N''$,

$$|u_n - \ell'| \leq \varepsilon/2$$

Soit $N = \max(N', N'')$. On a

$$|\ell - \ell'| \leq |\ell - u_N| + |u_N - \ell'| \leq 2\varepsilon/2 = \varepsilon$$

Ceci est vrai pour tout $\varepsilon > 0$. On en déduit que $|\ell - \ell'| = 0$, c'est à dire $\ell = \ell'$. □

Nous avons dans la démonstration ci-dessus utilisé un argument très fréquent en analyse. Si x est un réel positif ou nul qui est plus petit que tous les réels strictement positifs, alors $x = 0$. Plus mathématiquement,

$$(\forall \varepsilon > 0, 0 \leq x \leq \varepsilon) \implies x = 0$$

En effet, en supposant un court instant $x > 0$ et en prenant $\varepsilon = \frac{x}{2}$, on obtient $x \leq \frac{x}{2}$, d'où $1 \leq \frac{1}{2}$, contradiction.

Remarque 1. Si une suite u tend vers un réel ℓ , alors ℓ est **LA** limite de la suite u . On peut alors noter

$$\ell = \lim_{n \rightarrow \infty} u_n$$

ou, plus simplement, $\ell = \lim u$.

Définition 2. Une suite est dite convergente lorsqu'elle a une limite réelle. Sinon, on dit qu'elle est divergente.

1.2 Suites bornées, valeurs absolues

Rappelons que la suite réelle u est bornée si et seulement si il existe $m, M \in \mathbb{R}$ tels que

$$\forall n \in \mathbb{N}, m \leq u_n \leq M$$

L'inconvénient de cette définition est qu'elle nous oblige à montrer deux inégalités. Pour simplifier le travail, nous avons le résultat suivant.

Proposition 3. Soit u une suite réelle. La suite u est bornée si et seulement si il existe $M \in \mathbb{R}$ tel que

$$\forall n \in \mathbb{N}, |u_n| \leq M$$

Démonstration. Laissée en exercice. $\square$

Proposition 4. Toute suite convergente est bornée. La réciproque est fausse.

Démonstration. Supposons que $u \rightarrow \ell \in \mathbb{R}$. Il existe un entier N tel que pour tout $n \in \mathbb{N}$ on ait

$$n \geq N \implies |u_n - \ell| \leq 1$$

et donc

$$|u_n| \leq 1 + |\ell|$$

On en déduit que pour TOUT entier n ,

$$|u_n| \leq \max(|u_0|, |u_1|, \dots, |u_{N-1}|, 1 + |\ell|)$$

Il est facile de trouver des suites bornées divergentes, comme par exemple la suite $((-1)^n)_{n \geq 0}$.
□

Proposition 5.

- Soient u une suite bornée et v une suite qui tend vers 0. Alors, la suite uv tend vers 0.
- Soit u une suite qui tend vers $\ell \in \mathbb{R}$. Alors, $|u|$ tend vers $|\ell|$.
- Une suite u tend vers 0 si et seulement si sa valeur absolue tend vers 0.

Démonstration.

- Soient u une suite bornée et v une suite qui tend vers 0. Il existe donc un réel $M > 0$ tel que pour tout entier n , $|u_n| \leq M$. Soit $\varepsilon > 0$. Il existe un entier N tel que pour tout $n \geq N$, $|v_n| \leq \varepsilon/M$. Mais alors

$$|u_n v_n| \leq M \frac{\varepsilon}{M} = \varepsilon$$

- Supposons que u tend vers ℓ . On a pour tout entier n

$$||u_n| - |\ell|| \leq |u_n - \ell|$$

donc $|u|$ tend vers $|\ell|$.

- Enfin, $||u_n| - |0|| = |u_n - 0|$ d'où l'équivalence lorsque la suite tend vers 0.

□

Proposition 6. Soient u et v deux suites. On suppose que $v \rightarrow 0$, et que pour tout n assez grand, $|u_n| \leq v_n$. Alors $u \rightarrow 0$.

Démonstration. Soit $\varepsilon > 0$. Il existe N tel que pour tout $n \geq N$, on ait $|v_n| \leq \varepsilon$. Soit N' tel que pour tout $n \geq N'$ on ait $|u_n| \leq v_n$. Alors, pour tout $n \geq \max(N, N')$, on a $|u_n| \leq \varepsilon$. Donc, u tend vers 0. □

1.3 Limite infinie

Définition 3. Soit u une suite réelle. On dit que la suite tend vers $+\infty$ lorsque

$$\forall M \in \mathbb{R}, \exists N \in \mathbb{N}, \forall n \in \mathbb{N}, n \geq N \implies u_n \geq M$$

On définit de même une suite tendant vers $-\infty$.

Exemple. La suite de terme général $u_n = \frac{n^2+5}{3n-1}$ tend vers $+\infty$. On remarque pour cela que pour tout $n \geq 1$,

$$u_n \geq \frac{n^2}{3n} = \frac{n}{3}$$

donc $u_n \geq M$ dès que $n \geq 3M$.

Le cas des suites géométriques est à citer.

Proposition 7. Soit q un réel, $q > 1$. Alors q^n tend vers $+\infty$.

Démonstration. Posons en effet $q = 1 + h$, où $h > 0$. Alors,

$$q^n = (1 + h)^n \geq 1 + nh$$

Étant donné un réel M , on a donc $q^n \geq M$ dès que $n \geq \frac{M-1}{h}$. $\square$

Proposition 8. Une suite possède au plus une limite, finie ou infinie.

Démonstration. Une suite convergente est bornée, alors qu'une suite qui tend vers l'infini ne l'est pas. Une suite ne peut pas tendre à la fois vers l'infini, et vers une limite finie. Une suite qui tend vers $+\infty$ a son terme général supérieur à 1 pour n assez grand. Une telle suite ne peut donc pas tendre vers $-\infty$. Enfin, on a déjà vu qu'une suite ne peut pas tendre vers deux réels distincts. $\square$

1.4 Opérations sur les limites

Proposition 9. Soient u et v deux suites réelles. On suppose que $u_n \rightarrow \ell$ et $v_n \rightarrow \ell'$ lorsque $n \rightarrow \infty$. Soit $\lambda \in \mathbb{R}$. Alors :

- $u + v \rightarrow \ell + \ell'$.
- $uv \rightarrow \ell\ell'$.
- $\lambda u \rightarrow \lambda\ell$.

Démonstration.

- On a

$$|(u_n + v_n) - (\ell + \ell')| \leq |u_n - \ell| + |v_n - \ell'| \leq \varepsilon$$

dès que

$$|u_n - \ell| \leq \varepsilon/2$$

et

$$|v_n - \ell'| \leq \varepsilon/2$$

ce qui est le cas pour n assez grand.

- Plus subtil,

$$|u_n v_n - \ell \ell'| = |u_n(v_n - \ell') + \ell'(u_n - \ell)| \leq |u_n||v_n - \ell'| + |\ell'||u_n - \ell|$$

La suite u converge, donc la suite $|u|$ est majorée par un réel $M > 0$. On a pour tout $n \in \mathbb{N}$,

$$|u_n v_n - \ell \ell'| \leq \varepsilon$$

dès que

$$|u_n - \ell| \leq \frac{\varepsilon}{2(|\ell'| + 1)}$$

et

$$|v_n - \ell'| \leq \frac{\varepsilon}{2M}$$

ce qui est le cas pour n assez grand.

- En ce qui concerne le produit par λ , preuve laissée en exercice.

□

Remarque 2. Notons $\mathbb{R}_c^{\mathbb{N}}$ l'ensemble des suites réelles convergentes. Le résultat ci-dessus nous apprend des choses sur la structure algébrique de cet ensemble. Il est stable pour l'addition, le produit, et le produit par un nombre réel. Nous avons là ce que l'on appelle une $\mathbb{R}$ -algèbre. De plus, l'application

$$\lim : \mathbb{R}_c^{\mathbb{N}} \rightarrow \mathbb{R}$$

vérifie :

- $\forall u, v \in \mathbb{R}_c^{\mathbb{N}}, \lim(u + v) = \lim u + \lim v$.
- $\forall u, v \in \mathbb{R}_c^{\mathbb{N}}, \lim(uv) = \lim u \lim v$.
- $\forall u \in \mathbb{R}_c^{\mathbb{N}}, \forall \lambda \in \mathbb{R}, \lim(\lambda u) = \lambda \lim u$.

L'application $\lim$ est un *morphisme* d'algèbres.

Proposition 10. *Les résultats ci-dessus restent vrais pour des limites infinies, à condition de ne pas être dans un cas d'indétermination.*

Remarque 3. Les indéterminations en question sont $\infty - \infty$ et $0 \times \infty$. Se reporter aux tableaux des opérations sur la droite réelle achevée.

Proposition 11. *Soit u une suite tendant vers $\ell \in \overline{\mathbb{R}}_+^*$. Alors :*

- Il existe un réel $m > 0$ et un entier n_0 tels que $\forall n \geq n_0, u_n \geq m$.
- $\frac{1}{u}$ tend vers $\frac{1}{\ell}$ lorsque n tend vers l'infini.

Démonstration. Faisons la preuve pour $\ell \in \mathbb{R}_+^*$. On applique la définition de limite avec $\varepsilon = \frac{\ell}{2}$: pour tout n assez grand, on a

$$\frac{\ell}{2} \leq u_n \leq \frac{3}{2}\ell$$

D'où le premier point en posant

$$m = \frac{\ell}{2}$$

Puis, pour n assez grand,

$$\left| \frac{1}{u_n} - \frac{1}{\ell} \right| = \frac{|u_n - \ell|}{|u_n|\ell} \leq \frac{|u_n - \ell|}{m\ell}$$

qui est inférieur à tout $\varepsilon > 0$ pourvu que

$$|u_n - \ell| \leq m\ell\varepsilon$$

□

Proposition 12. Soit u une suite strictement positive tendant vers 0. Alors, $\frac{1}{u}$ tend vers $+\infty$ lorsque n tend vers l'infini.

Démonstration. Exercice. □

Proposition 13. Soit u une suite tendant vers $\ell \in \overline{\mathbb{R}}$. Soit f une fonction définie au voisinage de ℓ , et telle que $f(x) \rightarrow \ell' \in \overline{\mathbb{R}}$ lorsque x tend vers ℓ . Alors, $f(u_n)$ tend vers ℓ' lorsque n tend vers l'infini.

Nous admettons pour l'instant ce théorème, puisque la notion de limite pour les fonctions n'a pas encore été abordée de façon théorique. On l'utilisera sur des cas simples.

Exemple. Si $u_n \rightarrow -\infty$, alors $e^{u_n} \rightarrow 0$ puisque la limite de l'exponentielle en $-\infty$ est nulle.

Exemple. Soit u la suite définie par récurrence par $u_0 = 0$, et

$$\forall n \geq 0, u_{n+1} = \sqrt{1 + u_n}$$

Supposons que u converge vers un réel ℓ . Alors, $\sqrt{1 + u_n}$ converge vers $\sqrt{1 + \ell}$. Nous verrons un peu plus loin (suites extraites) que $(u_{n+1})_{n \geq 0}$ converge aussi vers ℓ . Par unicité de la limite, il vient

$$\ell = \sqrt{1 + \ell}$$

Ainsi, $\ell \geq 0$ et

$$\ell^2 - \ell - 1 = 0$$

Le seul réel $\ell \geq 0$ vérifiant cette égalité est le nombre d'or $\frac{1+\sqrt{5}}{2}$. Ainsi, SI la suite converge, c'est forcément vers $\frac{1+\sqrt{5}}{2}$.

Bien entendu, il resterait à prouver que cette suite converge. Mais ceci est une autre histoire.

1.5 Limites et ordre

Proposition 14. *Soient u et v deux suites tendant respectivement vers les réels achevés ℓ et ℓ' . On suppose que pour tout entier n assez grand, on a $u_n \leq v_n$. Alors, $\ell \leq \ell'$.*

Démonstration. Soit $\varepsilon > 0$. Il existe trois entiers N, N', N'' tels que

$$\forall n \geq N, |u_n - \ell| \leq \varepsilon$$

$$\forall n \geq N', |v_n - \ell'| \leq \varepsilon$$

et

$$\forall n \geq N'', u_n \leq v_n$$

On a donc pour tout $n \geq \max(N, N', N'')$,

$$\ell - \varepsilon \leq u_n \leq v_n \leq \ell' + \varepsilon$$

Ainsi,

$$\ell \leq \ell' + 2\varepsilon$$

ceci pour tout $\varepsilon > 0$. D'où $\ell \leq \ell'$. $\square$

Ne pas confondre ce résultat (passage à la limite dans une inégalité) avec le suivant (théorème d'encadrement)

Théorème 15. *Soient u, v, w trois suites réelles. On suppose que u et w convergent vers une même limite réelle ℓ et que, de plus, la double inégalité $u_n \leq v_n \leq w_n$ est vérifiée pour tout n assez grand. Alors, la suite v converge vers ℓ .*

Démonstration. Soit $\varepsilon > 0$. On a pour tout n assez grand

$$\ell - \varepsilon \leq u_n \leq v_n \leq w_n \leq \ell + \varepsilon$$

d'où le résultat. $\square$

Théorème 16. Soient u, v deux suites réelles. On suppose que u tend vers $+\infty$ et que, de plus, l'inégalité $u_n \leq v_n$ est vérifiée pour tout n assez grand. Alors, la suite v tend vers $+\infty$.

Démonstration. Soit $M \in \mathbb{R}$. On a pour tout n assez grand $M \leq u_n \leq v_n$. $\square$

Exemple. Soit

$$u_n = \sum_{k=1}^{2^n} \frac{1}{k}$$

On a

$$u_{n+1} - u_n = \sum_{k=2^n+1}^{2^{n+1}} \frac{1}{k} \geq \frac{1}{2}$$

Minorer chaque terme de la somme par $\frac{1}{2^{n+1}}$ et remarquer que cette somme comporte 2^n termes. Il en résulte, par exemple par récurrence sur n , que

$$u_n \geq \frac{n}{2}$$

Ainsi u_n tend vers $+\infty$ lorsque n tend vers l'infini.

1.6 Suites extraites

Définition 4. Soit u une suite. On appelle suite extraite de u toute suite v pour laquelle pour tout n $v_n = u_{\varphi(n)}$, où φ est une fonction strictement croissante de $\mathbb{N}$ vers $\mathbb{N}$.

Lemme 17. Soit $\varphi : \mathbb{N} \rightarrow \mathbb{N}$ une fonction strictement croissante. On a

$$\forall n \in \mathbb{N}, \varphi(n) \geq n$$

Démonstration. On fait une récurrence sur n . Tout d'abord, $\varphi(0) \in \mathbb{N}$ et donc $\varphi(0) \geq 0$. Soit $n \in \mathbb{N}$. Supposons $\varphi(n) \geq n$. Comme $n+1 > n$, on a $\varphi(n+1) > \varphi(n)$ et donc $\varphi(n+1) > n$. Mais n et $\varphi(n+1)$ sont entiers, donc $\varphi(n+1) \geq n+1$. $\square$

Proposition 18. Toute suite extraite d'une suite tendant vers $\ell \in \overline{\mathbb{R}}$ tend également vers ℓ .

Démonstration. Faisons par exemple la preuve pour ℓ réel. Étant donné $\varepsilon > 0$, il existe N tel que pour tout $n \geq N$,

$$|u_n - \ell| \leq \varepsilon$$

De là, pour tout $n \geq N$,

$$\varphi(n) \geq n \geq N$$

donc

$$|u_{\varphi(n)} - \ell| \leq \varepsilon$$

Donc, $v_n \rightarrow \ell$. $\square$

Cette proposition sert essentiellement à montrer qu'une suite est divergente, ou à trouver des conditions pour qu'elle converge. Prenons deux exemples.

Exemple. La suite de terme général $u_n = (-1)^n$ n'a pas de limite. En effet, u_{2n} tend vers 1 et u_{2n+1} tend vers -1 .

Exemple. Soit $x \in \mathbb{R}$. Soit u la suite définie par

$$\forall n \in \mathbb{N}, u_n = \cos nx$$

Supposons que cette suite converge vers un réel ℓ . Alors, pour tout entier n ,

$$u_{2n} = 2 \cos^2 nx - 1 = 2u_n^2 - 1$$

Comme la suite $(u_{2n})_{n \in \mathbb{N}}$ est extraite de u , elle tend aussi vers ℓ . Mais elle tend aussi vers $2\ell^2 - 1$. Donc, par unicité de la limite, $\ell = 2\ell^2 - 1$, d'où $\ell = 0$ ou $\ell = -\frac{1}{2}$.

De la même façon,

$$u_{3n} = 4u_n^3 - 3u_n$$

donc $\ell = 4\ell^3 - 3\ell$ et donc $\ell = 0, 1$ ou -1 .

En réunissant les deux, on en déduit que $\ell = 1$. Mais alors, $\sin^2 nx = 1 - u_n^2$ tend vers 0. Donc, $\sin nx$ tend aussi vers 0.

On considère enfin

$$u_{n+1} = \cos nx \cos x + \sin nx \sin x$$

qui tend vers $\cos x$. Mais $(u_{n+1})_{n \in \mathbb{N}}$ est aussi une suite extraite de u , elle tend donc vers 1. Ainsi, $\cos x = 1$ donc $x \in 2\pi\mathbb{Z}$. Bilan : si x n'est pas un multiple de 2π , la suite diverge.

Enfin, si $x \in 2\pi\mathbb{Z}$, la suite est constante égale à 1, donc converge.

Ce résultat possède un certain nombre de « réciproques » utiles. Citons la suivante :

Proposition 19. Soit u une suite réelle. On suppose que les deux suites $(u_{2n})_{n \geq 0}$ et $(u_{2n+1})_{n \geq 0}$ tendent vers une même limite ℓ lorsque n tend vers l'infini. Alors, la suite u tend vers ℓ .

Démonstration. Soit $\varepsilon > 0$. Il existe deux entiers N et N' tels que, pour tout $n \geq N$, on ait

$$|u_{2n} - \ell| \leq \varepsilon$$

et pour tout $n \geq N'$ on ait

$$|u_{2n+1} - \ell| \leq \varepsilon$$

Alors, pour tout $n \geq \max(2N, 2N' + 1)$, on a

$$|u_n - \ell| \leq \varepsilon$$

□

Exemple. Considérons la suite de terme général

$$u_n = \sum_{k=1}^n \frac{(-1)^{k-1}}{k}$$

Posons, pour tout $n \geq 1$, $v_n = u_{2n}$ et $w_n = u_{2n+1}$. On montre alors l'une de ces deux suites croît, l'autre décroît, et leur différence tend vers 0. Nous verrons plus loin (suites adjacentes) que cela entraîne que v et w convergent vers une même limite. Ainsi, la suite u converge. Sa limite est $\ln 2$, mais ceci est une autre histoire.

2 Théorèmes d'existence de limites

2.1 Suites monotones

Définition 5. Soit $u \in \mathbb{R}^{\mathbb{N}}$. La suite u est dite

- croissante lorsque

$$\forall n, n' \in \mathbb{N}, n \leq n' \implies u_n \leq u_{n'}$$

- strictement croissante lorsque

$$\forall n, n' \in \mathbb{N}, n < n' \implies u_n < u_{n'}$$

On définit de même la décroissance et la décroissance stricte en renversant les inégalités.

La suite u est dite monotone lorsqu'elle est croissante ou décroissante. On définit de même la stricte monotonie.

Proposition 20. La suite réelle u est croissante si et seulement si

$$\forall n \in \mathbb{N}, u_n \leq u_{n+1}$$

Démonstration. Supposons u est croissante. On a clairement la propriété demandée. Supposons inversement que pour tout entier n , $u_n \leq u_{n+1}$. Une récurrence facile montre que pour tous entiers n et p , on a $u_n \leq u_{n+p}$, d'où la croissance de u .

On a évidemment des résultats analogues pour la décroissance, la croissance stricte et la décroissance stricte. $\square$

Proposition 21. *Soit u une suite croissante de réels.*

- *Si u est majorée, alors u converge vers $\ell = \sup\{u_n, n \geq 0\}$.*
- *Si u n'est pas majorée, alors u diverge vers $+\infty$.*

Démonstration.

- Supposons u majorée. L'ensemble $E = \{u_n, n \in \mathbb{N}\}$ est une partie de $\mathbb{R}$ non vide et majorée, qui possède donc une borne supérieure ℓ . Soit $\varepsilon > 0$. Le réel $\ell - \varepsilon$ ne majore pas E , il existe donc N tel que $u_N > \ell - \varepsilon$. Mais alors, pour tout $n \geq N$, on a

$$\ell - \varepsilon \leq u_N \leq u_n \leq \ell \leq \ell + \varepsilon$$

ou encore

$$|u_n - \ell| \leq \varepsilon$$

- Supposons u non majorée. Soit $M \in \mathbb{R}$. Le réel M ne majore pas u , il existe donc un entier N tel que $u_N > M$. Soit $n \geq N$. Par la croissance de u ,

$$u_n \geq u_N \geq M$$

Ainsi, u tend vers $+\infty$.

$\square$

2.2 Suites adjacentes

Définition 6. Soient u et v deux suites réelles. On dit qu'elles sont adjacentes lorsque

- L'une croît, l'autre décroît.
- Leur différence tend vers 0.

Proposition 22. *Soient u et v deux suites adjacentes. Pour fixer les idées, on suppose que u croît et v décroît. Alors :*

- *Les suites u et v sont convergentes.*
- *Elles ont la même limite.*

- En supposant par exemple u croissante et v décroissante, et en appelant ℓ la limite commune des deux suites, on a pour tous $m, n \in \mathbb{N}$,

$$u_m \leq \ell \leq v_n$$

Démonstration. La suite $v - u$ est décroissante et tend vers 0. On a donc par le théorème sur les suites monotones

$$\inf\{v_n - u_n, n \in \mathbb{N}\} = 0$$

Donc, pour tout $n \in \mathbb{N}$, $v_n - u_n \geq 0$, c'est à dire $u_n \leq v_n$. Comme v décroît, $v_n \leq v_0$. Ainsi,

$$\forall n \in \mathbb{N}, u_n \leq v_0$$

La suite u est donc croissante et majorée. Elle converge ainsi vers un réel ℓ .

De même, v est décroissante et minorée par u_0 , elle converge donc vers un réel ℓ' . Comme $v - u$ tend vers 0, on conclut que $\ell = \ell'$.

Passons au dernier point. Puisque u croît et tend vers ℓ ,

$$\ell = \sup\{u_m, m \in \mathbb{N}\}$$

Ainsi,

$$\forall m \in \mathbb{N}, u_m \leq \ell$$

De même, comme v décroît et tend vers ℓ ,

$$\forall m \in \mathbb{N}, \ell \leq v_n$$

□

Remarque 4. En reprenant les notations ci-dessus, on a en particulier pour tout $n \in \mathbb{N}$,

$$u_n \leq \ell \leq v_n$$

La précision de l'encadrement, $\varepsilon_n = v_n - u_n$, tend vers 0 lorsque n tend vers l'infini.

Exercice. On pose pour tout entier $n \geq 1$,

$$u_n = \sum_{k=1}^n \frac{1}{k} - \ln n$$

et

$$v_n = \sum_{k=1}^n \frac{1}{k} - \ln(n+1)$$

Montrons que ces deux suites sont adjacentes.

- Soit $n \geq 1$. On a

$$u_{n+1} - u_n = \frac{1}{n+1} - \ln(n+1) + \ln n$$

Soit $f : \mathbb{R}_+^* \rightarrow \mathbb{R}$ définie par

$$f(x) = \frac{1}{x+1} - \ln(x+1) + \ln x$$

La fonction f est dérivable et on a pour tout $x > 0$,

$$f'(x) = -\frac{1}{(x+1)^2} - \frac{1}{x+1} + \frac{1}{x} = \frac{1}{x(x+1)^2} > 0$$

Le fonction f est donc strictement croissante. De plus, on a pour tout $x > 0$

$$f(x) = \frac{1}{x+1} - \ln\left(1 + \frac{1}{x}\right)$$

et donc f tend vers 0 en $+\infty$. Ainsi, $f < 0$ sur $\mathbb{R}_+^*$. La suite u est donc décroissante.

- On montre de même (exercice) que la suite v est croissante.
- Enfin, on a pour tout $n \geq 1$,

$$u_n - v_n = \ln(n+1) - \ln n = \ln\left(1 + \frac{1}{n}\right) \xrightarrow{n \rightarrow \infty} 0$$

Ces deux suites sont donc adjacentes. Leur limite commune, $\gamma \simeq 0.577215$, est appelée la constante d'Euler. Remarquons que l'on a pour tout $n \geq 1$,

$$v_n \leq \gamma \leq u_n$$

et que l'erreur commise en approchant γ par un tel encadrement est $\ln\left(1 + \frac{1}{n}\right)$. Nous verrons plus tard que cette erreur est de l'ordre de $\frac{1}{n}$.

Proposition 23. THÉORÈME DES SEGMENTS EMBOÎTÉS.

Soit $(I_n)_{n \geq 0}$ une suite de segments. On suppose que

- Les segments sont emboîtés : $\forall n \in \mathbb{N}, I_n \supset I_{n+1}$.
- La longueur $L(I_n)$ de I_n tend vers 0 lorsque n tend vers l'infini.

Alors, $\bigcap_{n \in \mathbb{N}} I_n$ est un singleton.

Démonstration. Posons, pour tout $n \in \mathbb{N}$, $I_n = [a_n, b_n]$, où $a_n, b_n \in \mathbb{R}$ et $a_n \leq b_n$. Les hypothèses du théorème nous disent que (a_n) croît, (b_n) décroît, et $L(I_n) = b_n - a_n$ tend vers 0. Les suites a et b sont donc adjacentes, et convergent ainsi vers une même limite ℓ .

Posons $I = \bigcap_{n \in \mathbb{N}} I_n$. On a pour tout $n \in \mathbb{N}$

$$a_n \leq \ell \leq b_n$$

Ainsi, $\ell \in I$ et donc $\{\ell\} \subset I$.

Inversement, soit $x \in I$. On a alors pour tout $n \in \mathbb{N}$,

$$a_n \leq x \leq b_n$$

$$-b_n \leq -\ell \leq -a_n$$

d'où, en ajoutant,

$$a_n - b_n \leq x - \ell \leq b_n - a_n$$

c'est à dire

$$|x - \ell| \leq b_n - a_n$$

D'où, en passant à la limite dans l'inégalité, $x = \ell$. $\square$

2.3 Théorème de Bolzano-Weierstrass

Théorème 24. [BOLZANO-WEIERSTRASS]

De toute suite bornée de réels on peut extraire une suite convergente.

Nous allons démontrer ce théorème en plusieurs étapes. Donnons-nous une suite réelle u . Supposons que u est bornée.

Définition 7. Si $A \subseteq \mathbb{R}$, notons $\mathcal{P}(A)$ la propriété « A contient une infinité de termes de la suite u ».

Précisément, on a $\mathcal{P}(A)$ si et seulement si l'ensemble

$$u^{-1}(A) = \{n \in \mathbb{N} : u_n \in A\}$$

est un ensemble infini.

Remarque. Si $A, B, C \subseteq \mathbb{R}$ et $B \cup C = A$, alors

$$\mathcal{P}(A) \iff \mathcal{P}(B) \text{ ou } \mathcal{P}(C)$$

En effet,

$$u^{-1}(A) = u^{-1}(B) \cup u^{-1}(C)$$

et une réunion finie d'ensembles est infinie si et seulement si l'un au moins des ensembles est infini.

Étape 1. *Construisons par récurrence sur n deux suites réelles $(a_n)_{n \in \mathbb{N}}$ et $(b_n)_{n \in \mathbb{N}}$ telles que pour tout $n \in \mathbb{N}$, on ait $\mathcal{P}([a_n, b_n])$.*

- Rappelons que la suite u est bornée. Appelons a_0 un minorant de la suite u et b_0 un majorant de u . Remarquons que le segment $[a_0, b_0]$ contient tous les termes de la suite u . On a donc $\mathcal{P}([a_0, b_0])$.
- Soit $n \in \mathbb{N}$. Supposons construits a_n et b_n tels que l'on ait $\mathcal{P}([a_n, b_n])$. Soit $c_n = \frac{1}{2}(a_n + b_n)$. Comme $[a_n, c_n] \cup [c_n, b_n] = [a_n, b_n]$, on a $\mathcal{P}([a_n, c_n])$ ou $\mathcal{P}([c_n, b_n])$.
 - (1) Si $\mathcal{P}([a_n, c_n])$ est vérifiée, on pose $a_{n+1} = a_n$ et $b_{n+1} = c_n$.
 - (2) Sinon, par la remarque faite un peu plus haut, on a $\mathcal{P}([c_n, b_n])$. On pose alors $a_{n+1} = c_n$ et $b_{n+1} = b_n$.

Lemme 25. *Les suites a et b sont adjacentes.*

Démonstration. Soit $n \in \mathbb{N}$. Si on est dans le cas (1) ci-dessus, alors $a_{n+1} = a_n$ et $b_{n+1} = c_n \leq b_n$. Si on est dans le cas (2), alors $a_{n+1} = c_n \geq a_n$ et $b_{n+1} = b_n$. Ainsi, la suite a croît et la suite b décroît.

Il reste à montrer que la suite $b - a$ tend vers 0. Remarquons que dans le cas (1),

$$b_{n+1} - a_{n+1} = c_n - a_n = \frac{1}{2}(a_n + b_n) - a_n = \frac{1}{2}(b_n - a_n)$$

et dans le cas (2),

$$b_{n+1} - a_{n+1} = b_n - c_n = b_n - \frac{1}{2}(a_n + b_n) = \frac{1}{2}(b_n - a_n)$$

La suite $b - a$ est donc une suite géométrique de raison $\frac{1}{2}$. Ainsi, $b - a \longrightarrow 0$. $\square$

Les suites adjacentes a et b convergent ainsi vers une même limite que nous noterons ℓ .

Étape 2. *Construisons maintenant par récurrence sur n une suite $(\varphi(n))_{n \in \mathbb{N}}$ strictement croissante d'entiers naturels telle que pour tout $n \in \mathbb{N}$, $u_{\varphi(n)} \in [a_n, b_n]$.*

- Rappelons que $u^{-1}([a_0, b_0]) = \mathbb{N}$. Posons donc, par exemple, $\varphi(0) = 0$.
- Soit $n \in \mathbb{N}$. Supposons construit un entier $\varphi(n)$ tel que $u_{\varphi(n)} \in [a_n, b_n]$. On a $\mathcal{P}([a_{n+1}, b_{n+1}])$, c'est à dire qu'il existe une infinité d'entiers m tels que $u_m \in [a_{n+1}, b_{n+1}]$. On choisit un tel entier $m > \varphi(n)$ et on pose $\varphi(n+1) = m$.

La suite φ définit une suite extraite de u . C'est la suite v définie pour tout $n \in \mathbb{N}$ par $v_n = u_{\varphi(n)}$.

Lemme 26. *La suite v converge vers ℓ .*

Démonstration. On a pour tout $n \in \mathbb{N}$, $u_{\varphi(n)} \in [a_n, b_n]$, c'est à dire

$$a_n \leq v_n \leq b_n$$

Les suites a et b convergent vers ℓ . Par le théorème d'encadrement, la suite v converge aussi vers ℓ . $\square$

3 Suites complexes

3.1 Introduction

Soit $u \in \mathbb{C}^{\mathbb{N}}$ une suite à valeurs complexes. On peut fabriquer à partir de u un certain nombre de suites, comme par exemple :

- $(\operatorname{Re} u_n)_{n \geq 0}$, $(\operatorname{Im} u_n)_{n \geq 0}$, et $(|u_n|)_{n \geq 0}$, qui sont des suites réelles.
- $(\overline{u_n})_{n \geq 0}$, qui est une suite complexe.

Définition 8. La suite u est dite bornée lorsque la suite réelle $(|u_n|)_{n \geq 0}$ est bornée.

Proposition 27. Une suite complexe est bornée si et seulement si sa partie réelle et sa partie imaginaire sont bornées.

Démonstration. Ceci résulte des inégalités

$$|\operatorname{Re} z| \leq |z| \leq |\operatorname{Re} z| + |\operatorname{Im} z| \quad \text{et} \quad |\operatorname{Im} z| \leq |z| \leq |\operatorname{Re} z| + |\operatorname{Im} z|$$

valables pour tout nombre complexe z . $\square$

3.2 Limites

La définition de limite dans $\mathbb{C}$ reste inchangée pour les suites complexes, le module remplaçant la valeur absolue. En revanche, on ne parle pas de limite infinie pour une suite complexe.

Définition 9. Soit $u \in \mathbb{C}^{\mathbb{N}}$. Soit $\ell \in \mathbb{C}$. On dit que u tend vers ℓ lorsque

$$\forall \varepsilon > 0, \exists N \in \mathbb{N}, \forall n \geq N, |u_n - \ell| \leq \varepsilon$$

Le théorème d'unicité de la limite reste vérifié (pourquoi ?). On emploie le même vocabulaire que pour les suites réelles : « ℓ est la limite de la suite u », « la suite u est convergente », etc.

Proposition 28. Soit $u \in \mathbb{C}^{\mathbb{N}}$. La suite u est convergente si et seulement si sa partie réelle et sa partie imaginaire sont convergentes. Plus précisément, soit $\ell \in \mathbb{C}$. On a

$$u \rightarrow \ell \iff (\operatorname{Re} u \rightarrow \operatorname{Re} \ell \quad \text{et} \quad \operatorname{Im} u \rightarrow \operatorname{Im} \ell)$$

Démonstration. Supposons que $\operatorname{Re} u \rightarrow \operatorname{Re} \ell$ et $\operatorname{Im} u \rightarrow \operatorname{Im} \ell$. On a pour tout $n \in \mathbb{N}$,

$$|u_n - \ell| \leq |\operatorname{Re} u_n - \operatorname{Re} \ell| + |\operatorname{Im} u_n - \operatorname{Im} \ell|$$

On en déduit facilement en revenant à la définition de limite que u tend vers ℓ .

Inversement, supposons que $u \rightarrow \ell$. on a pour tout $n \in \mathbb{N}$,

$$|\operatorname{Re} u_n - \operatorname{Re} \ell| \leq |u_n - \ell|$$

d'où, en appliquant la définition de limite, la convergence de $\operatorname{Re} u$ vers $\operatorname{Re} \ell$. On fait de même pour la partie imaginaire. $\square$

Restent vrais tous les théorèmes « ayant un sens », c'est à dire ne faisant pas appel à des comparaisons de termes de la suite (il n'y a pas de relation d'ordre dans $\mathbb{C}$!) : opérations sur les limites, théorèmes sur les suites extraites, relations de comparaison, Bolzano-Weierstrass.

Terminons cette section par l'exemple important des suites géométriques.

Proposition 29. Soit $q \in \mathbb{C}$.

- Si $|q| < 1$, la suite $(q^n)_{n \geq 0}$ tend vers 0.
- Si $|q| > 1$, la suite $(q^n)_{n \geq 0}$ diverge. Plus précisément, $|q^n|$ tend vers l'infini.
- Si $|q| = 1$, et $q \neq 1$ la suite $(q^n)_{n \geq 0}$ diverge.
- Si $q = 1$ la suite $(q^n)_{n \geq 0}$ est constante égale à 1 (et donc converge vers 1).

Démonstration.

- Supposons $|q| < 1$. Nous avons vu dans la partie du chapitre sur les suites réelles que $|q^n| = |q|^n$ tend vers 0 lorsque n tend vers l'infini. Donc, q^n tend aussi vers 0.
- Supposons $|q| > 1$. Alors $|q^n|$ tend vers $+\infty$ lorsque n tend vers l'infini, et donc $(q^n)_{n \geq 0}$ diverge.
- Supposons $|q| = 1$ et $q \neq 1$. Il existe donc $x \in \mathbb{R} \setminus 2\pi\mathbb{Z}$ tel que $q = e^{ix}$. Supposons un instant que $q^n = e^{inx}$ vers $\ell \in \mathbb{C}$ lorsque n tend vers l'infini.

Or, $(q^{n+1})_{n \geq 0}$ est une suite extraite de la suite $(q^n)_{n \geq 0}$. Cette suite tend donc aussi vers ℓ . Mais on a aussi, pour tout $n \in \mathbb{N}$,

$$q^{n+1} = e^{ix} q^n$$

et cette quantité tend vers $e^{ix}\ell$. Par unicité de la limite,

$$\ell(1 - e^{ix}) = 0$$

Comme $e^{ix} \neq 1$, il en résulte que $\ell = 0$. Mais ceci est absurde, puisque $|u_n| = 1$ pour tout n et donc, toujours par passage à la limite, $|\ell| = 1$.

- Le dernier point est évident.

$\square$

4 Récurrences linéaires

4.1 Suites arithmétiques et géométriques

Définition 10. La suite complexe u est dite arithmétique lorsqu'il existe $b \in \mathbb{C}$ tel que

$$\forall n \in \mathbb{N}, u_{n+1} = b + u_n$$

Le nombre b est appelé la raison de la suite.

Proposition 30. Soit u une suite arithmétique de raison b . On a pour tout entier n , $u_n = u_0 + bn$.

Démonstration. Évident par récurrence sur n . $\square$

Proposition 31. Soit u une suite arithmétique de raison b . Alors,

$$\sum_{k=0}^n u_k = (n+1)u_0 + \frac{n(n+1)}{2}b$$

Démonstration. Utiliser la propriété précédente et l'égalité

$$\sum_{k=0}^n k = \frac{1}{2}n(n+1)$$

$\square$

Définition 11. La suite complexe u est dite géométrique lorsqu'il existe $a \in \mathbb{C}$ tel que

$$\forall n \in \mathbb{N}, u_{n+1} = au_n$$

Le nombre a est appelé la raison de la suite.

Proposition 32. Soit u une suite géométrique de raison a . On a pour tout entier n , $u_n = u_0 a^n$.

Démonstration. Récurrence facile. $\square$

Proposition 33. *Soit u une suite géométrique de raison a .*

- Si $a \neq 1$,

$$\sum_{k=0}^n u_k = u_0 \frac{a^{n+1} - 1}{a - 1}$$

- Si $a = 1$,

$$\sum_{k=0}^n u_k = (n + 1)u_0$$

Démonstration. Cette identité remarquable a déjà été prouvée dans le chapitre sur les entiers naturels. $\square$

4.2 Suites « arithmético-géométriques »

Soient $a, b \in \mathbb{C}$. Soit $u \in \mathbb{C}^{\mathbb{N}}$ la suite définie par son terme initial u_0 et la relation

$$\forall n \in \mathbb{N}, u_{n+1} = au_n + b$$

Calculons u_n en fonction de n . Si $a = 1$, la suite est une suite arithmétique, et on sait faire. Supposons donc $a \neq 1$.

Soit ω l'unique complexe vérifiant $\omega = a\omega + b$ (c'est à dire $\frac{b}{1-a}$). On a alors pour tout entier n ,

$$u_{n+1} - \omega = a(u_n - \omega)$$

On retrouve une suite géométrique, d'où pour tout n ,

$$u_n = \omega + a^n(u_0 - \omega)$$

Remarquons que si $|a| < 1$, la suite u converge vers ω .

4.3 Récurrences linéaires à deux termes

On se donne deux nombres complexes a et b . On s'intéresse aux suites u vérifiant la relation de récurrence

$$(\mathcal{R}) \quad \forall n \geq 0, u_{n+2} = au_{n+1} + bu_n$$

Soit $q \in \mathbb{C}$. À quelle condition sur q la suite géométrique $(q^n)_{n \geq 0}$ vérifie-t-elle $(\mathcal{R})$? q convient si et seulement si pour tout $n \geq 0$,

$$q^{n+2} = aq^{n+1} + bq^n$$

ce qui équivaut à

$$(E) \quad q^2 = aq + b$$

Remarquons qu'une petite discussion, laissée au lecteur, s'impose dans le cas où $q = 0$ est solution du problème.

L'équation (E) est appelée *l'équation caractéristique* de la récurrence. Deux cas surviennent.

Cas 1. L'équation caractéristique a deux racines distinctes (a priori complexes) q et q' .

Proposition 34. Dans le cas 1, une suite u vérifie $\mathcal{R}$ si et seulement si existe $\alpha, \beta \in \mathbb{C}$ tels que pour tout $n \in \mathbb{N}$,

$$u_n = \alpha q^n + \beta q'^n$$

Démonstration. Supposons l'existence d'un tel α et d'un tel β . Ces deux nombres vérifient

$$\begin{cases} \alpha + \beta &= u_0 \\ \alpha q + \beta q' &= u_1 \end{cases}$$

On vérifie facilement que ces deux égalités sont vérifiées par un unique couple (α, β) . On montre ensuite par une récurrence à deux termes sur n que pour tout $n \in \mathbb{N}$, $u_n = \alpha q^n + \beta q'^n$. $\square$

Exemple. La suite de Fibonacci $(F_n)_{n \geq 0}$ est définie par $F_0 = 0$, $F_1 = 1$ et pour tout entier $n \in \mathbb{N}$,

$$F_{n+2} = F_{n+1} + F_n$$

L'équation caractéristique de cette récurrence est

$$(E) \quad q^2 = q + 1$$

Cette équation possède deux racines réelles φ et $\hat{\varphi}$. Les valeurs de ces racines importent peu. Remarquons simplement que

$$\varphi + \hat{\varphi} = 1 \text{ et } \varphi \hat{\varphi} = -1$$

Il existe $\alpha, \beta \in \mathbb{R}$ tels que pour tout $n \in \mathbb{N}$,

$$F_n = \alpha \varphi^n + \beta \hat{\varphi}^n$$

Ces deux nombres α et β sont les solutions du système

$$\begin{cases} \alpha + \beta &= 0 \\ \alpha \varphi + \beta \hat{\varphi} &= 1 \end{cases}$$

d'où, facilement, $\alpha = \frac{1}{\varphi - \hat{\varphi}}$, $\beta = \frac{-1}{\varphi - \hat{\varphi}}$, et donc, pour tout $n \in \mathbb{N}$,

$$F_n = \frac{\varphi^n - \hat{\varphi}^n}{\varphi - \hat{\varphi}}$$

Cas 2. L'équation caractéristique a une racine double q .

Proposition 35. Dans le cas 2, et dans l'hypothèse où $q \neq 0$, une suite u vérifie $\mathcal{R}$ si et seulement si existe $\alpha, \beta \in \mathbb{C}$ tels que pour tout n ,

$$u_n = \alpha q^n + n\beta q^n$$

Démonstration. Remarquons que l'hypothèse $q \neq 0$ n'est pas très restrictive. En effet, si on est dans le cas 2 et que $q = 0$, alors $a = b = 0$ et la relation de récurrence est

$$(\mathcal{R}) \quad u_{n+2} = 0$$

Prenons donc $q \neq 0$. Supposons l'existence d'un tel α et d'un tel β . Ces deux nombres vérifient

$$\begin{cases} \alpha &= u_0 \\ (\alpha + \beta)q &= u_1 \end{cases}$$

Il existe un unique couple (α, β) vérifiant ces deux égalités. On montre ensuite par une récurrence à deux termes sur n que pour tout $n \in \mathbb{N}$, $u_n = \alpha q^n + n\beta q^n$. $\square$

Exemple. Considérons la suite u définie par $u_0 = 2$, $u_1 = 1$ et pour tout $n \in \mathbb{N}$,

$$u_{n+2} = 4u_{n+1} - 4u_n$$

Son équation caractéristique est

$$(E) \quad q^2 = 4q - 4$$

Cette équation possède une unique racine qui est 2. Il existe donc $\alpha, \beta \in \mathbb{R}$ tels que pour tout $n \in \mathbb{N}$,

$$u_n = \alpha 2^n + n\beta 2^n$$

Ces deux nombres sont les solutions du système

$$\begin{cases} \alpha &= 2 \\ 2(\alpha + \beta) &= 1 \end{cases}$$

d'où $\alpha = 2$, $\beta = -\frac{3}{2}$, et pour tout $n \in \mathbb{N}$,

$$u_n = 2^{n+1} - 3n2^{n-1}$$

Chapitre 8

Comparaison des suites

Lorsqu'on recherche la limite d'une suite, on est fréquemment confronté à des indéterminations, dont les plus courantes sont $\infty - \infty$, $0 \times \infty$, $\frac{0}{0}$, $\frac{\infty}{\infty}$, 1^∞ et ∞^0 . En étudiant certaines relations entre les suites (équivalence, domination, négligeabilité), nous parviendrons à « lever » certaines indéterminations.

1 Suites équivalentes

1.1 C'est quoi ?

Définition 1. Soient u et v deux suites réelles. On dit que u est équivalente à v lorsqu'il existe une suite $(\lambda_n)_{n \geq 0}$ qui tend vers 1 telle que pour tout n assez grand $u_n = \lambda_n v_n$.

On note alors $u \sim v$ ou $u_n \underset{n \rightarrow \infty}{\sim} v_n$, ou un peu abusivement $u_n \sim v_n$.

Proposition 1. Soient u et v deux suites réelles. On suppose que pour tout n assez grand $v_n \neq 0$. Alors

$$u_n \sim v_n \iff \frac{u_n}{v_n} \underset{n \rightarrow \infty}{\longrightarrow} 1$$

Démonstration. C'est immédiat. $\square$

Exemple.

- $2n^2 + 3n - 1 \sim 2n^2$.
- Plus généralement, soient $d \in \mathbb{N}^*$ et $a_0, \dots, a_d \in \mathbb{R}$ tels que $a_d \neq 0$. On a

$$\sum_{k=0}^d a_k n^k \sim a_d n^d$$

Proposition 2. L'équivalence des suites est une relation d'équivalence.

Démonstration. Preuve facile. $\square$

1.2 Conservation de propriétés par équivalence

Proposition 3. Soient u, v deux suites réelles. Soit $\ell \in \overline{\mathbb{R}}$. On suppose que $u \sim v$ et que $u \longrightarrow \ell$. Alors $v \longrightarrow \ell$.

Démonstration. Pour tout n assez grand, $v_n = \lambda_n u_n$ où $\lambda_n \rightarrow 1$, d'où le résultat. $\square$

Remarque 1. Si deux suites u et v ont une même limite ℓ , sont-elles équivalentes ? La réponse est **NON**, sauf dans les cas triviaux.

- Si $\ell \in \mathbb{R}^*$, la réponse est **OUI**, puisque $\frac{u}{v} \rightarrow \frac{\ell}{\ell} = 1$. Mais ceci est sans intérêt, car nous sommes dans un cas où le calcul des équivalents ne sert à rien, vu qu'il n'y a pas d'indétermination.
- Si $\ell = 0, \pm\infty$, la réponse est **NON**, comme tout un tas d'exemple peuvent le montrer. Par exemple, $\frac{1}{n}$ et $\frac{1}{n^2}$, ou n et n^2 .

En bref, si u et v tendent vers 0 ou vers ∞ , on ne peut pas se prononcer sur leur équivalence. *Le but de ce qui suit est justement de pouvoir décider.*

Proposition 4. Soient u et v deux suites équivalentes. Alors, pour tout n assez grand,

- $u_n = 0 \iff v_n = 0$.
- $u_n > 0 \iff v_n > 0$.
- $u_n < 0 \iff v_n < 0$.

Démonstration. Pour tout n assez grand, $v_n = \lambda_n u_n$ où $\lambda_n \rightarrow 1$ lorsque n tend vers l'infini. Appliquons la définition de limite à la suite (λ_n) , en prenant $\varepsilon = \frac{1}{2}$. Il en résulte que pour tout n assez grand, $\lambda_n \geq \frac{1}{2}$ et donc $\lambda_n > 0$. D'où le résultat. $\square$

1.3 Opérations sur les équivalents

Proposition 5. Soient u, u', v, v' des suites réelles. On suppose $u \sim u'$ et $v \sim v'$. Alors

- $uv \sim u'v'$.
- Si v_n (et donc v'_n) est non nul pour n assez grand, $\frac{u}{v} \sim \frac{u'}{v'}$.
- Pour tout réel α , $u^\alpha \sim u'^\alpha$.
- En revanche, **on ne peut rien dire en ce qui concerne $u + v$ et $u' + v'$.**

Démonstration.

- Le produit et le quotient de deux suites qui tendent vers 1 tend encore vers 1.
- Si la suite λ tend vers 1 et si $\alpha \in \mathbb{R}$, alors λ^α tend encore vers 1.
- L'addition pose problème. Par exemple, on a $n^2 \sim n^2 + \sqrt{n}$ et $-n^2 + n \sim -n^2$. Mais on n'a PAS $n \sim \sqrt{n}$.

$\square$

1.4 Équivalents classiques

Proposition 6. Soit u une suite réelle **tendant vers 0**. Soit α un nombre réel. On a alors les équivalents suivants :

$e^{u_n} - 1$	$\sim$	u_n
$\ln(1 + u_n)$	$\sim$	u_n
$(1 + u_n)^\alpha - 1$	$\sim$	αu_n
$\sin u_n$	$\sim$	u_n
$1 - \cos u_n$	$\sim$	$\frac{u_n^2}{2}$
$\tan u_n$	$\sim$	u_n

Nous disposerons d'autres équivalents lorsque nous aurons abordé le cours sur les fonctions usuelles.

Démonstration. Faisons la démonstration en supposant que la suite u ne s'annule pas. La plupart de ces équivalents se montrent de la même façon, en prouvant qu'un taux d'accroissement tend vers une certaine dérivée. Par exemple,

$$\frac{e^{u_n} - 1}{u_n} = \frac{e^{u_n} - e^0}{u_n - 0}$$

tend vers la dérivée de l'exponentielle en 0, c'est à dire 1. De même,

$$\frac{(1 + u_n)^\alpha - 1}{u_n} = \frac{f(u_n) - f(0)}{u_n - 0}$$

où f est définie pour tout x voisin de 0 par $f(x) = (1 + x)^\alpha$. Cette quantité tend vers $f'(0) = \alpha$ lorsque n tend vers l'infini.

Seul l'équivalent pour le cosinus se traite différemment. On a

$$\frac{1 - \cos u_n}{u_n^2} = \frac{2 \sin^2 \frac{u_n}{2}}{u_n^2} \sim \frac{2 \frac{u_n^2}{4}}{u_n^2} = \frac{1}{2} \longrightarrow \frac{1}{2}$$

Donc,

$$1 - \cos u_n \sim \frac{u_n^2}{2}$$

□

Deux mots sur les puissances . . . Pour tout réel $x > 0$ et tout réel α , on pose

$$x^\alpha = \exp(\alpha \ln x)$$

Nous reviendrons sur les fonctions puissances dans le chapitre sur les fonctions usuelles. Tout ce que nous aurons besoin de savoir dans ce chapitre est que si α est un entier relatif on retrouve la valeur classique, que les formules usuelles $x^{\alpha+\beta} = x^\alpha x^\beta$, $(x^\alpha)^\beta = x^{\alpha\beta}$ et $(xy)^\alpha = x^\alpha y^\alpha$ restent valides, et que la dérivée de $x \mapsto x^\alpha$ est $x \mapsto \alpha x^{\alpha-1}$.

1.5 Exemples

- Quelle est la limite éventuelle de $u_n = \left(1 + \frac{1}{n}\right)^n$ lorsque n tend vers l'infini ? On a

$$\ln u_n = n \ln \left(1 + \frac{1}{n}\right) \sim n \frac{1}{n} = 1$$

et ainsi, $\ln u_n \longrightarrow 1$. Donc $u_n \longrightarrow e^1 = e$.

- Soit u une suite qui tend vers 0. Quelle est la limite éventuelle de

$$v_n = \frac{(1 - e^{u_n}) \sin u_n}{\ln(1 + 3u_n^2)}?$$

On a $\sin u_n \sim u_n$, $1 - e^{u_n} \sim -u_n$ et $\ln(1 + 3u_n^2) \sim 3u_n^2$. Ainsi,

$$v_n \sim \frac{-u_n^2}{3u_n^2} = -\frac{1}{3}$$

et donc $v \longrightarrow -\frac{1}{3}$.

- Soit u une suite qui tend vers 0. Quelle est la limite éventuelle de

$$v_n = \frac{\ln \cos(3u_n)}{\sin^2(2u_n)}?$$

On a

$$\ln \cos(3u_n) = \ln(1 + w_n)$$

où $w_n = \cos(3u_n) - 1 \longrightarrow 0$. Ainsi,

$$\ln \cos(3u_n) \sim w_n = \cos(3u_n) - 1 \sim -\frac{3}{2}u_n^2$$

De là,

$$v_n \sim \frac{-\frac{3}{2}u_n^2}{4u_n^2} = -\frac{3}{8}$$

et donc $v \longrightarrow -\frac{3}{8}$.

- Soit a un réel strictement positif différent de 1. Soit u une suite qui tend vers a et ne prend pas la valeur a . Déterminons un équivalent de

$$v_n = \log_a u_n - \log_{u_n} a$$

Posons tout d'abord $u_n = a + h_n$, de sorte que h_n tend vers 0. On a

$$v_n = \frac{\ln u_n}{\ln a} - \frac{\ln a}{\ln u_n} = \frac{(\ln u_n - \ln a)(\ln u_n + \ln a)}{\ln a \ln u_n}$$

$\ln u_n + \ln a$ tend vers $2 \ln a \neq 0$ et est donc équivalent à cette quantité. De même, $\ln u_n \ln a \rightarrow \ln^2 a \sim \ln^2 a$. Il reste à examiner dernier facteur :

$$\ln u_n - \ln a = \ln \frac{a + h_n}{a} = \ln \left(1 + \frac{h_n}{a} \right) \sim \frac{h_n}{a}$$

Finalement,

$$v_n \sim \frac{h_n 2 \ln a}{a \ln^2 a} = \frac{2}{a \ln a} (u_n - a)$$

1.6 Logarithmes et exponentielles d'équivalents

Proposition 7. Soient u et v deux suites réelles strictement positives. On suppose que $u \sim v$ et $u \rightarrow \ell \in \overline{\mathbb{R}}_+$ (et donc v aussi).

- Si $\ell \neq 1$ alors $\ln u_n \sim \ln v_n$.
- Si $\ell = 1$ on ne peut rien dire.

Démonstration.

- Supposons $\ell \neq 0, 1, +\infty$. On a $\ln u_n \rightarrow \ln \ell$ et, de même, $\ln v_n \rightarrow \ln \ell$. Comme $\ln \ell$ est un réel non nul, $\frac{\ln u_n}{\ln v_n} \rightarrow \frac{\ln \ell}{\ln \ell} = 1$ par les opérations sur les limites.
- Supposons $\ell = +\infty$. On a

$$\frac{\ln u_n}{\ln v_n} = \frac{\ln \frac{u_n}{v_n} + \ln v_n}{\ln v_n} = \frac{\ln \frac{u_n}{v_n}}{\ln v_n} + 1$$

Le numérateur de la fraction tend vers $\ln 1 = 0$ et son dénominateur tend vers $+\infty$, celle-ci tend donc vers 0.

- Supposons $\ell = 0$. $\frac{1}{u_n}$ et $\frac{1}{v_n}$ tendent alors vers $+\infty$, donc, par le point précédent,

$$-\ln u_n = \ln \frac{1}{u_n} \sim \ln \frac{1}{v_n} = -\ln v_n$$

d'où le résultat en passant à l'opposé.

- Lorsque $\ell = 1$, on ne peut rien dire. Par exemple, $u_n = 1 + \frac{1}{n^2}$ et $v_n = 1 + \frac{1}{n}$ tendent vers 1. Pourtant, $\ln u_n \sim \frac{1}{n^2}$ et $\ln v_n \sim \frac{1}{n}$ ne sont pas équivalents.

□

Proposition 8. Soient u et v deux suites réelles. On suppose que $u \sim v$ et $u \rightarrow \ell \in \overline{\mathbb{R}}$ (et donc v aussi).

- Si $\ell \neq \pm\infty$ alors $e^u \sim e^v$.
- Si $\ell = \pm\infty$ on ne peut rien dire.

Démonstration.

- Supposons $\ell \neq \pm\infty$. Alors $e^{u_n} \rightarrow e^\ell$ et $e^{v_n} \rightarrow e^\ell$. De là, $\frac{e^{u_n}}{e^{v_n}} \rightarrow 1$.
- Lorsque $\ell = \pm\infty$, tout peut arriver. Par exemple, $n \sim n+1$ et pourtant

$$e^{n+1} = ee^n \not\sim e^n$$

□

Moralité : On ne peut pas passer aux exponentielles dans les équivalents, sauf dans les cas triviaux.

2 Négligeabilité, domination

2.1 C'est quoi ?

Définition 2. Soient u et v deux suites réelles. On dit que

- u est dominée par v lorsque pour tout n assez grand, $u_n = w_n v_n$, où la suite w est bornée. On note alors $u = O(v)$ ou $u_n = O(v_n)$.
- u est négligeable devant v lorsque pour tout n assez grand, $u_n = \varepsilon_n v_n$, où ε_n tend vers 0 lorsque n tend vers l'infini. On note alors $u = o(v)$ ou $u_n = o(v_n)$ ou $u \ll v$ ou encore $u_n \ll v_n$.

Proposition 9. Soient u et v deux suites réelles. On suppose que pour tout n assez grand $v_n \neq 0$. Alors,

- $u = O(v)$ si et seulement si la suite $\left(\frac{u_n}{v_n}\right)$ est bornée.
- $u = o(v)$ si et seulement si $\frac{u_n}{v_n} \rightarrow 0$ lorsque n tend vers l'infini.

Démonstration. C'est immédiat. □

Exemple.

- On a $3n^2 + 8n - 1 = O(n^2)$, ou aussi $O(n^3)$ ou $O(n^4) \dots$
- Plus généralement, soient $d \in \mathbb{N}$ et $a_0, \dots, a_d \in \mathbb{R}$. On a

$$\sum_{k=0}^d a_k n^k = O(n^d)$$

- Soient $\alpha, \beta \in \mathbb{R}$. On a $n^\alpha = o(n^\beta)$ si et seulement si $\alpha < \beta$.
- Soient $\alpha, \beta \in \mathbb{R}$. On a $\alpha^n = o(\beta^n)$ si et seulement si $\alpha < \beta$.

Les notations « petit o » et « grand O » sont très pratiques. Voici un exemple de manipulation de ces notations.

Proposition 10. *Soit v une suite de réels non nuls. Alors*

- Pour tout réel a , $o(av_n) = o(v_n)$.
- Pour tout réel a non nul, $o(v_n) = o(av_n)$.
- $o(v_n) + o(v_n) = o(v_n)$.
- $o(v_n) \times o(v_n) = o(v_n^2)$.

Démonstration. Avec des notations évidentes, on a

- $o(av_n) = \varepsilon_n av_n = \varepsilon'_n v_n = o(v_n)$.
- $o(v_n) = \frac{1}{a} \varepsilon_n av_n = \varepsilon'_n (av_n) = o(av_n)$.
- $o(v_n) + o(v_n) = (\varepsilon_n + \varepsilon'_n) v_n = \varepsilon''_n v_n$.
- $o(v_n) \times o(v_n) = \varepsilon_n \varepsilon'_n v_n^2 = \varepsilon''_n v_n^2 = o(v_n^2)$.

Le but n'est pas de *retenir* ce genre d'égalités, mais de les *comprendre*. En cas de doute, il suffit de revenir à $o(v_n) = \varepsilon_n v_n$ où ε_n tend vers 0. $\square$

Remarque 2. Dans la proposition précédente, on a l'impression que $1 + 1$ ça ne fait pas 2 et que l'égalité n'est pas une relation symétrique. Il faut effectivement faire attention lorsqu'on manipule des o , car $o(v_n)$ représente UNE suite négligeable devant v_n . Donc, par exemple lorsqu'on écrit $o(v_n) + o(v_n)$ on additionne deux suites négligeables devant v , mais ces deux suites sont a priori différentes. Encore une fois, en cas de doute il n'y a qu'à revenir à la définition.

Exercice. Montrer que la proposition ci-dessus reste vraie en remplaçant les petits o par des grands O .

2.2 Équivalents et petits o

Proposition 11. *Soient u et v deux suites de réels non nuls. On a*

$$u \sim v \iff u = v + o(v)$$

Démonstration. Tout simplement parce qu'une suite λ tend vers 1 si et seulement si on peut écrire $\lambda_n = 1 + \varepsilon_n$ où $\varepsilon_n \longrightarrow 0$. $\square$

On peut donc toujours remplacer un équivalent par une égalité, *en n'oubliant pas, bien entendu, de rajouter un petit o .*

2.3 Bilan

Lorsque v_n est non nul pour tout n assez grand, nous pouvons résumer tout ce qui précède dans un tableau.

u_n/v_n	
$\rightarrow 0$	$u_n = o(v_n)$
$\rightarrow 1$	$u_n \sim v_n$
est bornée	$u_n = O(v_n)$

2.4 Un exemple

Soit a un réel strictement positif. Soit u une suite qui tend vers a . Déterminons un équivalent de

$$v_n = a^{u_n} - u_n^a$$

Posons $u_n = a + h_n$, de sorte que h_n tend vers 0. On a

$$v_n = a^{a+h_n} - (a+h_n)^a = a^a \left(e^{h_n \ln a} - \left(1 + \frac{h_n}{a}\right)^a \right)$$

On a $e^{h_n \ln a} - 1 \sim h_n \ln a$, donc

$$e^{h_n \ln a} = 1 + h_n \ln a + o(h_n)$$

De même,

$$\left(1 + \frac{h_n}{a}\right)^a - 1 \sim a \frac{h_n}{a} = h_n$$

et donc

$$\left(1 + \frac{h_n}{a}\right)^a = 1 + h_n + o(h_n)$$

Regroupons le tout :

$$\begin{aligned} v_n &= a^a ((1 + h_n \ln a + o(h_n)) - (1 + h_n + o(h_n))) \\ &= a^a (\ln a - 1) h_n + o(h_n) \end{aligned}$$

Si $a \neq e$, on a $\ln a \neq 1$. Dans ce cas,

$$v_n \sim a^a (\ln a - 1)(u_n - a)$$

Si $a = e$, on a $v_n = o(u_n - a)$ et nos calculs ne nous permettent pas d'obtenir un équivalent de v_n . Nous verrons lors de l'étude des développements limités que des outils plus puissants que ceux de ce chapitre permettent d'obtenir un tel équivalent.

3 Puissances, exponentielles, logarithmes, factorielles

3.1 Croissances comparées

On prouve dans cette dernière section un résultat qui permet de comparer logarithmes, puissances, exponentielles et factorielle de suites tendant vers l'infini. Ce résultat est parfois appelé le « théorème des croissances comparées ». Noter les guillemets, les suites que nous allons considérer n'ayant pas de raison d'être croissantes. Un terme plus approprié aurait été « théorème de comparaison des vitesses auxquelles les logarithmes, puissances, exponentielles et factorielles tendent vers l'infini », mais il faut avouer que c'est un peu long comme titre.

Lemme 12. *Soit $a > 1$. Soit u une suite qui tend vers $+\infty$. On a*

$$\ln u_n \ll u_n \ll a^{u_n}$$

Démonstration. *La démonstration générale repose sur des propriétés que nous verrons dans le chapitre sur les fonctions. Dans ce qui suit nous allons uniquement prouver la proposition lorsque $u_n = n$.*

- Soit $n \geq 1$. On a

$$\begin{aligned} \frac{\ln(n+1)}{n+1} - \frac{\ln n}{n} &= \frac{n \ln(n+1) - (n+1) \ln n}{n(n+1)} \\ &= \frac{n(\ln n + \ln(1 + \frac{1}{n})) - (n+1) \ln n}{n(n+1)} \\ &= \frac{n \ln(1 + \frac{1}{n}) - \ln n}{n(n+1)} \\ &\sim -\frac{\ln n}{n(n+1)} \leq 0 \end{aligned}$$

Ainsi, la suite $(\frac{\ln n}{n})_{n \geq 1}$ est décroissante pour n assez grand. Comme elle est minorée par 0, elle converge vers $\ell \geq 0$. De plus

$$\frac{\ln(n^2)}{n^2} = \frac{2 \ln n}{n} \frac{1}{n}$$

Passons à la limite (suite extraite!). Il vient $\ell = 0$ d'où $\ell = 0$.

- Soit $n \geq 1$. On a

$$\frac{a^{n+1}}{n+1} - \frac{a^n}{n} = \frac{a^n}{n(n+1)}((a-1)n - 1)$$

Cette quantité est positive dès que $n \geq \frac{1}{a-1}$. La suite $(\frac{a^n}{n})_{n \geq 1}$ est donc croissante à partir d'un certain rang. Ainsi, elle tend vers $\ell \in \overline{\mathbb{R}}_+^*$. Supposons un court instant ℓ réel. On a

$$\frac{a^{n+1}}{n+1} = a \frac{a^n}{n} \frac{n}{n+1}$$

Passons à la limite (suite extraite!). Il vient

$$\ell = a\ell$$

d'où $\ell = 0$. Contradiction. Ainsi, $\ell = +\infty$.

□

Proposition 13. Soit $a > 1$. Soient $\alpha, \beta > 0$. Soit u une suite qui tend vers $+\infty$. On a

$$\ln^\beta u_n \ll u_n^\alpha \ll a^{u_n}$$

Démonstration.

- On a

$$\frac{\ln^\beta u_n}{u_n^\alpha} = \left(\frac{\ln u_n}{u_n^{\alpha/\beta}} \right)^\beta$$

Posons $v_n = u_n^{\alpha/\beta}$. Il vient

$$\frac{\ln^\beta u_n}{u_n^\alpha} = \left(\frac{\ln v_n^{\beta/\alpha}}{v_n} \right)^\beta = \left(\frac{\beta}{\alpha} \right)^\beta \left(\frac{\ln v_n}{v_n} \right)^\beta$$

Par le lemme, cette quantité tend vers 0 lorsque n tend vers l'infini.

- On a

$$\frac{u_n^\alpha}{a^{u_n}} = \left(\frac{u_n}{b^{u_n}} \right)^\alpha$$

où $b = a^{1/\alpha}$. Par le lemme, cette quantité tend vers 0 lorsque n tend vers l'infini.

□

Proposition 14. Soit $a > 1$. Soit u une suite d'entiers naturels tendant vers $+\infty$. On a

$$a^{u_n} \ll u_n!$$

Démonstration. Faisons la démonstration pour $u_n = n$. Il existe un entier N tel que pour tout $n \geq N$, on ait

$$\frac{a}{n} \leq \frac{1}{2}$$

Pour tout $n \geq N$, écrivons

$$\frac{a^n}{n!} = \frac{a}{1} \frac{a}{2} \cdots \frac{a}{n} = K \frac{a}{N} \frac{a}{N+1} \cdots \frac{a}{n}$$

où

$$K = \frac{a}{1} \frac{a}{2} \cdots \frac{a}{N-1}$$

On a alors

$$0 \leq \frac{a^n}{n!} \leq K \left(\frac{1}{2}\right)^{n-N+1} = \frac{K'}{2^n}$$

où $K' = K2^{N-1}$. On conclut par le théorème d'encadrement des limites que $\frac{a^n}{n!}$ tend vers 0 lorsque n tend vers l'infini. $\square$

3.2 Un exemple

Soit u une suite qui tend vers $+\infty$. Déterminer un équivalent de

$$v_n = u_n^{\frac{1}{u_n}} - u_n$$

Nous avons là une terrible indétermination du type $\infty^{\infty^0} - \infty$. Factorisons u_n :

$$v_n = u_n \left(u_n^{\frac{1}{u_n} - 1} - 1 \right) = u_n \left(e^{(u_n^{\frac{1}{u_n}} - 1) \ln u_n} - 1 \right)$$

Posons $w_n = u_n^{\frac{1}{u_n}} - 1$. On a

$$w_n = e^{\ln u_n / u_n} - 1$$

Comme $\ln u_n \ll u_n$, le quotient $\frac{\ln u_n}{u_n}$ tend vers 0, et donc

$$w_n \sim \frac{\ln u_n}{u_n}$$

et

$$w_n \ln u_n \sim \frac{\ln^2 u_n}{u_n}$$

Par le théorème des croissances comparées, $w_n \ln u_n$ tend vers 0 lorsque n tend vers l'infini. De là,

$$v_n = u_n \left(e^{w_n \ln u_n} - 1 \right) \sim u_n w_n \ln u_n \sim \ln^2 u_n$$

Chapitre 9

Groupes, anneaux, corps

1 Lois de composition interne

1.1 Définition

Définition 1. Soit E un ensemble non vide. On appelle loi de composition interne (ou plus simplement opération) sur E toute application $\star : E \times E \rightarrow E$.

L'usage est de noter $x \star y$ l'image du couple (x, y) par l'opération $\star$. La plupart des lois classiques sont notées additivement $(+)$ ou multiplicativement $(\times, \text{ ou } \cdot)$ mais ce n'est pas obligé (penser à l'intersection et la réunion des ensembles $(\cap, \cup)$, à la composition des applications $(\circ)$, à l'opération $(x, y) \mapsto x^y, \dots$).

1.2 Commutativité

Définition 2. Soient E un ensemble non vide et $\star$ une opération sur E . On dit que $\star$ est commutative lorsque pour tous éléments x et y de E , on a :

$$x \star y = y \star x$$

Exemple. L'addition et la multiplication sont commutatives dans $\mathbb{R}$. La composition n'est pas commutative dans l'ensemble $\mathbb{R}^{\mathbb{R}}$ des fonctions de $\mathbb{R}$ vers $\mathbb{R}$. La soustraction n'est pas commutative dans $\mathbb{Z}$. La réunion est commutative dans l'ensemble $\mathcal{P}(A)$ des parties d'un ensemble A .

1.3 Associativité

Définition 3. Soient E un ensemble non vide et $\star$ une opération sur E . On dit que $\star$ est associative lorsque pour tous éléments x, y et z de E , on a :

$$(x \star y) \star z = x \star (y \star z)$$

Remarque 1. Les parenthèses deviennent alors inutiles, et on peut alors écrire plus simplement $x \star y \star z$.

Exemple. L'addition et la multiplication sont associatives dans $\mathbb{R}$. La composition est associative dans l'ensemble $\mathbb{R}^{\mathbb{R}}$ des fonctions de $\mathbb{R}$ vers $\mathbb{R}$. La soustraction n'est pas associative dans $\mathbb{Z}$. La réunion et l'intersection sont associatives dans l'ensemble $\mathcal{P}(A)$ des parties d'un ensemble A .

1.4 Élément neutre

Définition 4. Soient E un ensemble non vide et $\star$ une opération sur E . Soit e un élément de E . On dit que e est neutre pour l'opération $\star$ lorsque pour tout x dans E on a :

$$x \star e = e \star x = x$$

Exemple. 0 est neutre pour l'addition dans $\mathbb{C}$. 1 est neutre pour la multiplication dans $\mathbb{Q}$. $id_{\mathbb{R}}$ est neutre pour la composition dans $\mathbb{R}^{\mathbb{R}}$. $\emptyset$ est neutre pour la réunion dans $\mathcal{P}(A)$. La soustraction dans $\mathbb{R}$ ne possède pas d'élément neutre : il n'existe aucun réel a tel que pour tout réel x , on ait $x - a = a - x = x$.

Proposition 1. *L'élément neutre, lorsqu'il existe, est unique.*

Démonstration. Soient e et e' deux neutres pour $\star$. On a $e \star e' = e$ puisque e' est neutre, mais on a aussi $e \star e' = e'$ puisque e est neutre. Donc, $e = e'$. $\square$

1.5 Élément absorbant

Définition 5. Soient E un ensemble non vide et $\star$ une opération sur E . Soit ζ un élément de E . On dit que ζ est absorbant pour l'opération $\star$ lorsque pour tout x dans E on a :

$$x \star \zeta = \zeta \star x = \zeta$$

Exercice. L'élément absorbant, s'il existe, est-il forcément unique ?

1.6 Élément inversible

Définition 6. Soient E un ensemble non vide muni d'une opération $\star$ possédant un neutre e . Soit x un élément de E . On dit que x est inversible pour l'opération $\star$ lorsque il existe un élément x' de E tel que :

$$x \star x' = x' \star x = e$$

Un tel élément x' est appelé symétrique (ou inverse) de x pour l'opération $\star$.

Proposition 2. *Si l'opération est associative, le symétrique, lorsqu'il existe, est unique.*

Démonstration. Soit x un élément de E . Soient x' et x'' deux symétriques de x . On a $x \star x' = e$ d'où $x'' \star (x \star x') = x'' \star e = x''$. Mais l'opération $\star$ est associative : $(x'' \star x) \star x' = e \star x' = x' = x''$. $\square$

Remarque 2. Le symétrique de x est couramment noté x^{-1} . Pour certaines lois (comme l'addition), le symétrique est appelé *l'opposé*, et le symétrique de x est noté $-x$.

1.7 Élément régulier

Définition 7. Soit E un ensemble non vide muni d'une opération $\star$. Soit x un élément de E . On dit que x est régulier pour l'opération $\star$ lorsque, pour tous éléments y et z de E , on a

$$x \star y = x \star z \Rightarrow y = z \quad \text{et} \quad y \star x = z \star x \Rightarrow y = z$$

Proposition 3. *Si l'opération est associative, un élément inversible est régulier. La réciproque est fausse.*

Démonstration. Soit x inversible, d'inverse x' . Soient y et z tels que $x \star y = x \star z$. On a alors $y = x' \star x \star y = x' \star x \star z = z$. L'inversibilité entraîne la régularité. En revanche, Dans l'ensemble $\mathbb{Z}$ muni de la multiplication, le nombre 2 est régulier ($2y = 2z \Rightarrow y = z$) mais pas inversible : il n'existe pas d'entier x tel que $2x = 1$. $\square$

1.8 Distributivité

Définition 8. Soit E un ensemble non vide muni de deux opérations $\star$ et $\perp$. On dit que $\star$ est distributive par rapport à $\perp$ lorsque pour tous éléments x, y, z de E , on a :

$$x \star (y \perp z) = (x \star y) \perp (x \star z) \quad \text{et} \quad (y \perp z) \star x = (y \star x) \perp (z \star x)$$

2 Groupes

2.1 Définitions

Définition 9. Soient G un ensemble non vide, et $\star$ une opération sur G . On dit que $(G, \star)$ est un groupe lorsque

- $\star$ est associative.
- $\star$ possède un élément neutre.
- Tout élément de G est inversible pour la loi $\star$.

Si, de plus, $\star$ est commutative, on dit que le groupe G est commutatif (ou abélien¹).

Remarque 3. Dans la suite, lorsqu'on travaillera sur un groupe G « abstrait », on notera multiplicativement l'opération sur ce groupe. On notera également e le neutre G , sauf mention contraire (et si le groupe est G' , on note son neutre e' , etc).

Exercice. Déterminer tous les groupes à 1, 2, 3 et 4 éléments.

Exercice. Notons $\mathfrak{S}_3$ l'ensemble des bijections de l'ensemble $\llbracket 1, 3 \rrbracket$ dans lui-même. Montrer que $(\mathfrak{S}_3, \circ)$ est un groupe, et donner sa table.

2.2 Puissances d'un élément

Définition 10. Soit G un groupe de neutre e . On définit pour $x \in G$ et $n \in \mathbb{Z}$, l'élément $x^n \in G$ par

- $x^0 = e$
- $\forall n \in \mathbb{N}, x^{n+1} = x^n x$
- $\forall n \in \mathbb{Z}_-, x^n = (x^{-n})^{-1}$

Proposition 4. Soit G un groupe. On a pour tous $x, y \in G$,

- $(x^{-1})^{-1} = x$
- $(xy)^{-1} = y^{-1}x^{-1}$

Démonstration. C'est évident. $\square$

Proposition 5. Soit G un groupe. Soient $x, y \in G$ tels que $xy = yx$. On a pour tous entiers $m, n \in \mathbb{Z}$ $x^m y^n = y^n x^m$.

Démonstration. On montre d'abord par récurrence sur n que pour tout $n \in \mathbb{N}$, $x^n y = y^n x$.

C'est clair pour $n = 0, 1$. Soit $n \in \mathbb{N}$. Supposons que $x^n y = y x^n$. On a alors

$$x^{n+1} y = x^n x y = x^n y x = y x^n x = y x^{n+1}$$

Soit maintenant $n < 0$. On a

$$x^{-n} y = y x^{-n}$$

puisque $-n \in \mathbb{N}$. En multipliant par x^n à gauche et à droite de cette égalité, on obtient $y x^n = x^n y$. Ainsi, pour tout $n \in \mathbb{Z}$, $x^n y = y x^n$.

Soient maintenant $m, n \in \mathbb{Z}$. Par ce qui précède, x^m et y commutent. Puis, à nouveau par ce qui précède, x^m et y^n commutent. $\square$

Lemme 6. Soit G un groupe. Pour tout $x \in G$ et tout $n \in \mathbb{Z}$, on a

$$x^{n+1} = x^n x$$

Démonstration. Pour $n \in \mathbb{N}$ c'est vrai par définition des puissances. Supposons donc $n < 0$. On a donc $n + 1 \leq 0$. De là,

$$\begin{aligned} x^{n+1} x^{-1} &= (x^{-n-1})^{-1} x^{-1} \\ &= (x x^{-n-1})^{-1} \\ &= (x^{1-n-1})^{-1} \\ &= (x^{-n})^{-1} \\ &= x^n \end{aligned}$$

Il reste à multiplier les deux côtés de l'égalité par x pour obtenir le résultat. $\square$

Proposition 7. Soit G un groupe. On a

- $\forall x, y \in G, \forall n \in \mathbb{Z}, (xy)^n = x^n y^n$ à condition que $xy = yx$.
- $\forall x \in G, \forall m, n \in \mathbb{Z}, x^{m+n} = x^m x^n$.
- $\forall x \in G, \forall m, n \in \mathbb{Z}, x^{mn} = (x^m)^n$.

Démonstration.

- Démontrons la première propriété. On la montre d'abord pour $n \in \mathbb{N}$ par récurrence sur n . C'est évident pour $n = 0, 1$. Supposons que cela soit vrai pour un certain entier n . On a alors

$$(xy)^{n+1} = (xy)^n xy = x^n y^n xy = x^n x y^n y = x^{n+1} y^{n+1}$$

Montrons maintenant la propriété pour $n < 0$. On a

$$(xy)^n = ((xy)^{-n})^{-1} = (x^{-n}y^{-n})^{-1} = (y^{-n})^{-1}(x^{-n})^{-1} = y^n x^n$$

Il reste à remarquer que comme x et y commutent, x^n et y^n commutent aussi.

- Démontrons la deuxième propriété.

Montrons tout d'abord par récurrence sur n que pour tout $n \in \mathbb{N}$, pour tout $m \in \mathbb{Z}$, $x^{m+n} = x^m x^n$. Pour $n = 0$, $x^{m+0} = x^m = x^m e = x^m x^0$. Soit maintenant $n \in \mathbb{N}$. Supposons que pour tout $m \in \mathbb{Z}$, $x^{m+n} = x^m x^n$. Alors

$$x^{m+n+1} = x^{m+n} x = x^m x^n x = x^m x^{n+1}$$

Soit maintenant $n < 0$. On a, en utilisant le fait que $-n \in \mathbb{N}$,

$$x^m = x^{(m+n)-n} = x^{m+n} x^{-n}$$

En multipliant à droite par x^n , on obtient

$$x^m x^n = x^{m+n} x^{-n} x^n = x^{m+n} e = x^{m+n}$$

- La troisième propriété est laissée en exercice.

□

Remarque 4. Conséquence importante de ces formules, les puissances d'un élément x de G commutent entre-elles.

Remarque 5. Pour certains groupes abéliens G , l'opération est notée « additivement ». On ne parle pas de puissances, mais de multiples. Le neutre est la plupart du temps noté 0, et ce qui précède devient :

Définition 11. Soit $(G, +)$ un groupe de neutre 0. On définit pour $x \in G$ et $n \in \mathbb{Z}$, l'élément $nx \in G$ par :

- $0x = 0$
- $\forall n \in \mathbb{N}, (n+1)x = nx + x$
- $\forall n \in \mathbb{Z}_-, nx = -(-nx)$

Proposition 8. Soit $(G, +)$ un groupe abélien. On a

- $\forall x \in G, \forall m, n \in \mathbb{Z}, (m+n)x = mx + nx$.
- $\forall x \in G, \forall m, n \in \mathbb{Z}, (mn)x = m(nx)$.
- $\forall x, y \in G, \forall n \in \mathbb{Z}, n(x+y) = nx + ny$.

2.3 Sous-groupes

Définition 12. Soit $(G, \star)$ un groupe. Soit H une partie de G . On dit que H est un sous-groupe de G lorsque $(H, \star)$ est lui-même un groupe.

Proposition 9. *Le neutre de H est le même que celui de G , et l'inverse d'un élément de H est le même que celui de cet élément, vu comme un élément de G .*

Démonstration. On a $e_H e_H = e_H$. En multipliant par l'inverse de e_H dans G , on obtient $e_H = e_G$. On le note simplement e dans la suite. Soit maintenant $x \in H$. On dispose de x' , inverse de x en tant qu'élément de H , et x'' , l'inverse de x dans G . On a $xx' = e$. On multiplie par x'' à gauche, il vient $x''xx' = x' = ex'' = x''$. $\square$

Proposition 10. *Soit G un groupe. Soit $H \subset G$. Alors, H est un sous-groupe de G si et seulement si*

- $H \neq \emptyset$
- Pour tous éléments x et y de H , $xy \in H$
- Pour tout élément x de H , $x^{-1} \in H$.

Démonstration. Dans un sens, c'est clair. Un sous groupe de G contient le neutre, donc il est non vide. De plus, il est stable par multiplication, et on a vu que l'inverse d'un élément de H était l'inverse de cet élément dans G . Inversement, soit H non vide, stable par multiplication et inversion. L'associativité est automatique puisque H est inclus dans G . Soit $x \in H$. Alors $x^{-1} \in H$ donc $e = xx^{-1} \in H$. Donc H a bien un neutre (celui de G , évidemment). $\square$

Remarque 6. Les deux dernières assertions sont équivalentes à : pour tous x et y de H , $xy^{-1} \in H$. Preuve laissée en exercice.

2.4 Morphismes de groupes

Définition 13. Soient $(G, \star)$ et $(G', *)$ deux groupes. Soit $f : G \rightarrow G'$. On dit que f est un morphisme (de groupes) lorsque

$$\forall x, y \in G, f(x \star y) = f(x) * f(y)$$

Proposition 11. Soit $f : G \rightarrow G'$ un morphisme de groupes. On a :

- $f(e) = e'$
- $\forall x \in G, f(x^{-1}) = f(x)^{-1}$
- $\forall x \in G, \forall n \in \mathbb{Z}, f(x^n) = f(x)^n$

Démonstration.

- On a $ee = e$ donc $f(e)f(e) = f(e)$. On multiplie des deux côtés par l'inverse de $f(e)$ et on obtient $f(e) = e'$.
- Soit $x \in G$. On a $xx^{-1} = e$ donc $f(x)f(x^{-1}) = f(e) = e'$. On multiplie des deux côtés par l'inverse de $f(x)$ et on obtient $f(x^{-1}) = f(x)^{-1}$.
- Pour la dernière propriété, prenons d'abord $n \in \mathbb{N}$ et faisons une récurrence. Pour $n = 0$, c'est clair puisque $f(x^0) = f(e) = e' = f(x)^0$. Soit maintenant $n \in \mathbb{N}$. Supposons que $f(x^n) = f(x)^n$. Alors

$$f(x^{n+1}) = f(x^n x) = f(x^n)f(x) = f(x)^n f(x) = f(x)^{n+1}$$

Soit enfin $n < 0$. On a $n = -p$ où p est un entier naturel. Alors,

$$f(x^n) = f((x^p)^{-1}) = f(x^p)^{-1} = (f(x)^p)^{-1} = f(x)^n$$

□

Remarque 7. Il convient bien entendu de traduire la proposition précédente lorsque l'une des lois de groupe (voire les deux) est une loi additive. Par exemple, l'exponentielle est un morphisme de $(\mathbb{R}, +)$ vers $(\mathbb{R}^*, \times)$, puisque $e^{x+y} = e^x e^y$. Donc, on a $e^{nx} = (e^x)^n$ (l'exponentielle d'un multiple est une puissance de l'exponentielle).

Définition 14. Soit $f : G \rightarrow G'$ est un morphisme de groupes, on dit que :

- f est un isomorphisme lorsque f est bijectif.
- f est un endomorphisme lorsque $G = G'$ (et que les opérations sur G et G' sont identiques).
- f est un automorphisme lorsque f est un endomorphisme bijectif.

Définition 15. Deux groupes sont isomorphes lorsqu'il existe un isomorphisme de l'un vers l'autre.

2.5 Noyau et image d'un morphisme

Proposition 12. *L'image directe et l'image réciproque d'un sous-groupe par un morphisme de groupes est encore un sous-groupe.*

Démonstration. Soit $f : G \rightarrow G'$ un morphisme de groupes.

Soit H un sous-groupe de G . Soit $H' = f(H)$. H est non vide, donc H' aussi. Soient $y, y' \in H'$. Il existe $x, x' \in G$ tels que $y = f(x)$ et $y' = f(x')$. D'où

$$yy' = f(x)f(x') = f(xx') \in H'$$

Et

$$y^{-1} = f(x)^{-1} = f(x^{-1}) \in H'$$

H' est bien un sous-groupe de G' .

Soit maintenant H' un sous-groupe de G' . Posons $H = f^{-1}(H')$. On a $f(e) = e' \in H'$ donc $e \in H$ et H est non vide. Soient $x, x' \in H$. On a

$$f(xx') = f(x)f(x') \in H'$$

donc $xx' \in H$. Enfin,

$$f(x^{-1}) = f(x)^{-1} \in H'$$

donc $x^{-1} \in H$. H est bien un sous-groupe de G . $\square$

Définition 16. Soit $f : G \rightarrow G'$ un morphisme de groupes. On appelle

- Noyau de f , et on note $\ker f$ l'ensemble des éléments de G dont l'image par f est le neutre de G' :

$$\ker f = f^{-1}(\{e'\}) = \{x \in G, f(x) = e'\}$$

- Image de f , et on note $\operatorname{Im} f$, l'ensemble des images de tous les éléments de G :

$$\operatorname{Im} f = f(G) = \{f(x), x \in G\}$$

Proposition 13. *Soit $f : G \rightarrow G'$ un morphisme de groupes.*

- *Le noyau de f est un sous-groupe de G .*
- *L'image de f est un sous-groupe de G' .*

Démonstration. En effet, $\ker f$ est l'image réciproque par f d'un sous-groupe de G' et $\operatorname{Im} f$ est l'image directe par f d'un sous-groupe de G . $\square$

Proposition 14. *Soit $f : G \rightarrow G'$ un morphisme de groupes. Alors :*

- *f est injectif si et seulement si $\ker f = \{e\}$.*
- *f est surjectif si et seulement si $\operatorname{Im} f = G'$.*

Démonstration. Supposons f injectif. Soit $x \in G$. On a $x \in \ker f$ si et seulement si $f(x) = e' = f(e)$, si et seulement si $x = e$ par l'injectivité de f .

Supposons maintenant que $\ker f$ ne contient que e . Soient $x, x' \in G$ tels que $f(x) = f(x')$. Alors

$$f(x)f(x')^{-1} = e'$$

ou encore

$$f(xx'^{-1}) = e'$$

Donc $xx'^{-1} \in \ker f$ d'où $xx'^{-1} = e$ et on a bien $x = x'$.

La caractérisation de la surjectivité est évidente. $\square$

2.6 Composition de morphismes

Proposition 15.

- *La composée de deux morphismes de groupes est encore un morphisme de groupes.*
- *La réciproque d'un isomorphisme de groupes est encore un isomorphisme de groupes.*

Démonstration. Exercice. $\square$

Proposition 16. *Soit G un groupe. L'ensemble $\operatorname{Aut} G$ des automorphismes du groupe G est un groupe pour la composition des applications.*

Démonstration. Exercice. $\square$

3 Anneaux et corps

3.1 Définitions

Définition 17. Soit A un ensemble muni de deux opérations $+$ et $\times$ (notées additivement et multiplicativement pour simplifier). On dit que $(A, +, \times)$ est un anneau lorsque

- $(A, +)$ est un groupe abélien (dont on notera le neutre 0).
- A possède un élément neutre pour la multiplication (que l'on notera 1).
- La multiplication dans A est associative.
- La multiplication dans A est distributive par rapport à l'addition.

Si, de plus, la multiplication dans A est commutative, on dit que A est un anneau commutatif.

Exemple. $\mathbb{Z}$, $\mathbb{Q}$, $\mathbb{R}$, $\mathbb{C}$ avec les opérations usuelles sont des anneaux commutatifs. $\mathbb{N}$ n'en est pas un. Si E est un ensemble, l'ensemble $\mathbb{R}^E$ des applications de E dans $\mathbb{R}$ est un anneau, en posant $(f+g)(x) = f(x) + g(x)$ et $(fg)(x) = f(x)g(x)$ pour toutes $f, g \in \mathbb{R}^E$ et tout $x \in E$. Cela fonctionne toujours si, à la place de $\mathbb{R}$ on prend un anneau A quelconque.

Remarque 8. L'ensemble $\{0\}$, muni des opérations $0+0 = 0 \times 0 = 0$ est un anneau, appelé l'anneau nul. Dans cet anneau, on a $1 = 0$. Cette situation est tout à fait exceptionnelle. En effet, soit A un anneau dans lequel $1 = 0$. Soit $x \in A$. On a alors $x.1 = x$, et $x.1 = x.0 = 0$. Donc, $x = 0$. Ainsi, A est l'anneau nul. En conclusion, dans un anneau autre que l'anneau nul, on a toujours $1 \neq 0$.

Définition 18. Soit $(A, +, \times)$ un anneau commutatif. On dit que A est un corps lorsque

- A est différent de l'anneau nul.
- Tout élément non nul de A est inversible pour la multiplication.

Exemple. $\mathbb{Z}$ n'est pas un corps. En revanche, $\mathbb{Q}$, $\mathbb{R}$ et $\mathbb{C}$ en sont. L'anneau $\mathbb{R}^{\mathbb{R}}$ n'est pas non plus un corps. Par exemple, la fonction $x \mapsto x$ n'a pas d'inverse pour la multiplication.

3.2 Sommes et produits

Soit $(A, +, \times)$ un anneau. Étant donnés des éléments $a_1, a_2, \dots, a_n$ de l'anneau A , on définit par récurrence sur n

$$\sum_{k=1}^n a_k \text{ et } \prod_{k=1}^n a_k$$

en convenant que $\sum_{k=1}^0 a_k = 0$ et $\prod_{k=1}^0 a_k = 1$. Plus généralement, si I est un ensemble fini, et $(a_i)_{i \in I}$ est une famille finie d'éléments de A , on définit par récurrence sur le cardinal de I

$$\sum_{i \in I} a_i \text{ et } \prod_{i \in I} a_i$$

en convenant que si $I = \emptyset$, la somme vaut 0 et le produit vaut 1. Attention, pour définir le produit, on a besoin de savoir que les a_i commutent (ce qui sera bien sûr le cas si l'anneau est commutatif).

3.3 Puissances et multiples

Dans un anneau $(A, +, \times)$ on dispose à la fois de la notion de multiple d'un élément (pour l'opération $+$) et de puissance **positive** d'un élément (pour l'opération $\times$). Seules les puissances positives sont autorisées pour un élément quelconque, les puissances négatives n'ayant de sens que lorsque l'élément en question est inversible (pour la multiplication).

3.4 Sous-Anneaux

Définition 19. Soit $(A, +, \times)$ un anneau. Soit $B \subset A$. On dit que B est un sous-anneau de A lorsque

- $(B, +, \times)$ est lui-même un anneau.
- Le neutre pour la multiplication dans B est le même que le neutre pour la multiplication dans A .

Un sous-corps $\mathbb{L}$ du corps $\mathbb{K}$ est un sous-anneau $\mathbb{L}$ de $\mathbb{K}$ tel que tout élément non nul de $\mathbb{L}$ possède un inverse **dans** $\mathbb{L}$.

Proposition 17. Soit $(A, +, \times)$ un anneau de neutre 1 pour la multiplication. Soit $B \subset A$. Alors, B est un sous-anneau de A si et seulement si

- $1 \in B$.
- $\forall x, y \in B, x - y \in B$.
- $\forall x, y \in B, xy \in B$.

B est un sous-corps de A si et seulement si B est un sous-anneau de A et si, de plus,

$$\forall x \in B \setminus \{0\}, x^{-1} \in B$$

Démonstration. La preuve est laissée en exercice. $\square$

3.5 Éléments inversibles d'un anneau

Proposition 18. Soit A un anneau. L'ensemble A^* des éléments inversibles de A , muni de la multiplication, est un groupe.

Démonstration. La multiplication est associative. Il suffit de vérifier que le produit de deux inversibles est inversible, et que tout inversible a un inverse, qui est lui-même inversible. Mais on le sait déjà. $\square$

3.6 Morphismes d'anneaux

Définition 20. Soit $f : A \rightarrow A'$ une application d'un anneau A vers un anneau A' (les lois des deux anneaux sont notées ici $+$ et $\times$ pour simplifier). On dit que f est un morphisme d'anneaux lorsque

- $\forall x, y \in A, f(x + y) = f(x) + f(y)$.
- $f(1_A) = 1_{A'}$.
- $\forall x, y \in A, f(xy) = f(x)f(y)$.

Remarque 9. On peut parler d'image et de noyau d'un morphisme d'anneaux. Si $f : A \rightarrow A'$ est un morphisme d'anneaux, $\ker f = f^{-1}(\{0\})$ et $\operatorname{Im} f = f(A)$. Un morphisme d'anneaux étant avant tout un morphisme de groupes additifs, on a que f est injectif si et seulement si son noyau est réduit à $\{0\}$ et f est surjectif si et seulement si $\operatorname{Im} f = A'$. Mais attention, le noyau de f n'est pas en général un sous-anneau de A (c'est ce que l'on appelle un idéal de A , on verra apparaître cette notion dans les exercices).

Exercice. Quels sont les endomorphismes de $\mathbb{Z}$? De $\mathbb{Q}$? De $\mathbb{R}$? Citer deux endomorphismes de $\mathbb{C}$.

3.7 Identités remarquables

Proposition 19. Soient a et b deux éléments d'un anneau A vérifiant $ab = ba$. Soit n un entier naturel non nul. Alors :

$$b^n - a^n = (b - a) \sum_{k=0}^{n-1} a^k b^{n-1-k}$$

Démonstration. On a

$$X = (b - a) \sum_{k=0}^{n-1} a^k b^{n-1-k} = \sum_{k=0}^{n-1} a^k b^{n-k} - \sum_{k=0}^{n-1} a^{k+1} b^{n-1-k}$$

En posant $k' = k + 1$ dans la deuxième somme, on voit que celle-ci est égale à $\sum_{k=1}^n a^k b^{n-k}$. Ainsi,

$$X = \sum_{k=0}^{n-1} a^k b^{n-k} - \sum_{k=1}^n a^k b^{n-k}$$

Tous les termes s'éliminent deux à deux, sauf le terme d'indice $k = 0$ dans la première somme et le terme d'indice $k = n$ dans la deuxième somme, qui valent respectivement b^n et a^n . $\square$

Un cas particulier important est le cas $b = 1$:

Proposition 20. *Soit a un élément d'un anneau A . Soit n un entier naturel non nul. Alors :*

$$1 - a^n = (1 - a) \sum_{k=0}^{n-1} a^k$$

Remarque 10. Si $1 - a$ est inversible, ce qui arrive par exemple lorsque $A = \mathbb{C}$ et $a \neq 1$, on peut diviser l'identité précédente par $1 - a$, ce qui donne la formule bien connue

$$\sum_{k=0}^{n-1} a^k = \frac{1 - a^n}{1 - a}$$

Bien entendu, si $a = 1$, alors $\sum_{k=0}^{n-1} a^k = \sum_{k=0}^{n-1} 1 = n$.

Proposition 21. *Soient a et b deux éléments d'un anneau A vérifiant $ab = ba$. Soit n un entier naturel. Alors :*

$$(a + b)^n = \sum_{k=0}^n \binom{n}{k} a^k b^{n-k}$$

Il s'agit de la formule du « binôme de Newton ». Les quantités $\binom{n}{k}$ sont les coefficients binomiaux : $\binom{n}{k} = \frac{n!}{k!(n-k)!}$.

Démonstration. Nous avons déjà démontré la formule du binôme pour a et b complexes. Il n'y a aucune différence, la preuve se fait par récurrence sur n . $\square$

3.8 Anneaux intègres

Nous nous contentons dans ce paragraphe de quelques définitions (élément régulier, diviseur de zéro, anneau intègre) et de propriétés faciles.

Définition 21. Soit A un anneau. Soit $x \in A$.

- On dit que x est régulier lorsque pour tous $y, z \in A$,

$$xy = xz \implies y = z \quad \text{et} \quad yx = zx \implies y = z$$

- On dit que x est un diviseur de zéro lorsqu'il existe $y \in A \setminus \{0\}$ tel que $xy = 0$ ou $yx = 0$.

Proposition 22. *Un élément x de l'anneau A est régulier si et seulement si x n'est pas un diviseur de zéro.*

Démonstration.

($\Rightarrow$) Supposons que x est un diviseur de zéro. Il existe donc $y \neq 0$ tel que, par exemple, $xy = 0$. On a donc $xy = x0$, avec $y \neq 0$. x n'est donc pas régulier.

($\Leftarrow$) Supposons que x n'est pas régulier. Il existe donc $y, z \in A$ distincts tels que, par exemple, $xz = yz$. De là, $x(y - z) = 0$ et $y - z \neq 0$. Ainsi, x est un diviseur de zéro. $\square$

Exemple. Considérons l'anneau $\mathbb{R}^{\mathbb{R}}$ des fonctions de $\mathbb{R}$ vers $\mathbb{R}$. La somme des deux fonctions f et g est la fonction $f + g$ définie pour tout réel x par

$$(f + g)(x) = f(x) + g(x)$$

et de même pour le produit fg . Une fonction f est un diviseur de zéro dans A si et seulement si elle s'annule. Elle est régulière si et seulement si elle est inversible si et seulement si elle ne s'annule pas.

Proposition 23. *Tout élément inversible d'un anneau A est régulier. La réciproque est fausse.*

Démonstration. Soit x un élément inversible de l'anneau A . Soient $y, z \in A$. Supposons que $xy = xz$. En multipliant cette égalité à gauche par x^{-1} , on obtient $y = z$. De même pour l'implication $yx = zx \implies y = z$.

Pour la réciproque, il suffit de considérer l'anneau $\mathbb{Z}$. Tous les entiers non nuls sont réguliers, alors que les seuls éléments inversibles sont -1 et 1 . $\square$

Définition 22. Soit A un anneau. On dit que A est intègre lorsque A est différent de l'anneau nul et que tous les éléments non nuls de A sont réguliers.

Par exemple, $\mathbb{Z}$ est un anneau intègre. Les corps sont aussi des anneaux intègres.

Les anneaux commutatifs intègres sont « presque » des corps. On peut montrer que tout anneau commutatif intègre A peut être en un certain sens plongé dans un corps K dont les éléments sont les « fractions » d'éléments de A . Le corps K est appelé le corps des fractions de l'anneau A .

Par exemple, le corps des fractions de $\mathbb{Z}$ est le corps $\mathbb{Q}$ des rationnels. Un peu plus loin dans ce cours, nous étudierons l'anneau $\mathbb{K}[X]$ des polynômes à une indéterminée sur un corps $\mathbb{K}$. L'anneau $\mathbb{K}[X]$ est un anneau commutatif intègre, et son corps des fractions, noté $\mathbb{K}(X)$ est le corps des fractions rationnelles à une indéterminée sur le corps $\mathbb{K}$.

4 Espaces vectoriels

Nous nous contenterons ici de deux définitions. L'étude des espaces vectoriels, appelée *l'algèbre linéaire*, sera faite plus tard dans l'année et représente une part importante du programme.

Définition 23. Soient E un ensemble et $\mathbb{K}$ un corps. On suppose E muni d'une loi de composition interne $+$ et d'une « loi de composition externe » $\cdot : \mathbb{K} \times E \rightarrow E$. On dit que $(E, +, \cdot)$ est un $\mathbb{K}$ -espace vectoriel lorsque

- $(E, +)$ est un groupe abélien.
- Pour tout $x \in E$, $1x = x$.
- Pour tous $x, y \in E$ et tout $\lambda \in \mathbb{K}$, $\lambda(x + y) = \lambda x + \lambda y$.
- Pour tout $x \in E$ et tous $\lambda, \mu \in \mathbb{K}$, $(\lambda + \mu)x = \lambda x + \mu x$.
- Pour tout $x \in E$ et tous $\lambda, \mu \in \mathbb{K}$, $\lambda(\mu x) = (\lambda\mu)x$.

Les éléments de E sont appelés des *vecteurs*. Les éléments de $\mathbb{K}$ sont appelés des *scalaires*.

Par exemple,

- $\mathbb{K}^n$ muni de l'addition des n -uplets terme à terme et du produit par λ terme à terme est un $\mathbb{K}$ -espace vectoriel.
- L'ensemble $\mathbb{R}^{\mathbb{R}}$ des fonctions de $\mathbb{R}$ vers $\mathbb{R}$, muni de l'addition des fonctions et du produit d'une fonction par un réel, est un $\mathbb{R}$ -espace vectoriel.

Une opération importante dans les espaces vectoriels est la *combinaison linéaire*. Si $\lambda_1, \dots, \lambda_n \in \mathbb{K}$ et $x_1, \dots, x_n \in E$, on peut former le vecteur

$$\sum_{k=1}^n \lambda_k x_k = \lambda_1 x_1 + \dots + \lambda_n x_n$$

On peut montrer qu'une partie F d'un espace vectoriel E est lui-même un espace vectoriel si et seulement si F est non vide et stable par combinaisons linéaires. Une telle partie est appelée un *sous-espace vectoriel* de E .

Les morphismes d'espaces vectoriels sont appelés des *applications linéaires*.

Définition 24. Soient E et F deux $\mathbb{K}$ -espaces vectoriels. Soit $f : E \rightarrow F$. On dit que f est une application linéaire lorsque

- Pour tous $x, y \in E$, $f(x + y) = f(x) + f(y)$.
- Pour tout $x \in E$ et tout $\lambda \in \mathbb{K}$, $f(\lambda x) = \lambda f(x)$.

Chapitre 10

Fonctions

Dans ce court chapitre, nous nous intéressons aux propriétés algébriques des fonctions à valeurs réelles. En particulier, nous mettons en place du vocabulaire qui sera ensuite couramment utilisé.

Dans ce chapitre, X désigne un ensemble quelconque. On s'intéresse aux propriétés algébriques de l'ensemble $\mathbb{R}^X$ des fonctions de X vers $\mathbb{R}$. Deux cas sont particulièrement intéressants : celui où X est un intervalle de $\mathbb{R}$ (c'est ce qui se passe dans le cours d'Analyse) et celui où $X = \mathbb{N}$ (ou une partie de $\mathbb{N}$) dans le cours sur les suites à valeurs réelles.

1 Opérations algébriques

Soient $f, g \in \mathbb{R}^X$. On peut fabriquer à partir de f et g de nouvelles fonctions :

- La fonction $f + g : X \rightarrow \mathbb{R}$ est définie, pour tout $x \in X$, par

$$(f + g)(x) = f(x) + g(x)$$

- La fonction $fg : X \rightarrow \mathbb{R}$ est définie, pour tout $x \in X$, par

$$(fg)(x) = f(x)g(x)$$

- Si $\lambda \in \mathbb{R}$, la fonction $\lambda f : X \rightarrow \mathbb{R}$ est définie, pour tout $x \in X$, par

$$(\lambda f)(x) = \lambda f(x)$$

Proposition 1. *L'ensemble $\mathbb{R}^X$, muni des opérations ci-dessus est un espace vectoriel et un anneau commutatif (on dit aussi une algèbre).*

Démonstration. Ce sont de simples vérifications. $\square$

Remarque 1. Une fonction est inversible pour la multiplication si et seulement si elle ne s'annule pas. Or, pour un ensemble X ayant au moins deux éléments, il existe des fonctions $X \rightarrow \mathbb{R}$ différentes de la fonction nulle, et qui pourtant s'annulent. L'anneau $\mathbb{R}^X$ n'est donc pas un corps.

2 Une relation d'ordre

2.1 Un ordre partiel

Définition 1. Pour $f, g \in \mathbb{R}^X$, on dit que $f \leq g$ lorsque

$$\forall x \in X, f(x) \leq g(x)$$

Proposition 2. *Si X a au moins deux éléments, cette relation est une relation d'ordre partiel.*

Démonstration. Laissé en exercice. Pour l'ordre partiel, prendre par exemple la fonction $x \mapsto x$ et $x \mapsto -x$. Aucune de ces fonctions n'est plus petite que l'autre. $\square$

Définition 2. Pour $f, g \in \mathbb{R}^X$, on appelle $\sup(f, g)$ la fonction définie par

$$\sup(f, g)(x) = \max(f(x), g(x))$$

On définit de même $\inf(f, g)$.

Remarque 2. On a

$$\sup(f, g) = \frac{f + g}{2} + \left| \frac{f - g}{2} \right|$$

et

$$\inf(f, g) = \frac{f + g}{2} - \left| \frac{f - g}{2} \right|$$

Exemple. On pose $f^+ = \sup(f, 0)$ et $f^- = -\inf(f, 0)$. Alors, f^+ et f^- sont positives, et on a $f = f^+ - f^-$ et $|f| = f^+ + f^-$.

2.2 Fonctions bornées

Définition 3. Soit $f : X \longrightarrow \mathbb{R}$. On dit que f est

- majorée lorsque

$$\exists M \in \mathbb{R}, \forall x \in X, f(x) \leq M$$

- minorée lorsque

$$\exists m \in \mathbb{R}, \forall x \in X, f(x) \geq m$$

- bornée lorsqu'elle est majorée et minorée, ce qui peut s'écrire :

$$\exists M \in \mathbb{R}, \forall x \in X, |f(x)| \leq M$$

Exercice. Montrer que si $f, g : X \longrightarrow \mathbb{R}$ sont bornées, alors $f + g$, fg et, pour tout réel λ , λf , sont encore des fonctions bornées.

3 Fonctions monotones

Dans cette section, on suppose que X est une partie de $\mathbb{R}$.

Définition 4. Soit $f : X \longrightarrow \mathbb{R}$. On dit que

- f est croissante sur X lorsque

$$\forall x, y \in X, x \leq y \implies f(x) \leq f(y)$$

- f est strictement croissante sur X lorsque

$$\forall x, y \in X, x < y \implies f(x) < f(y)$$

On définit de même (strictement) décroissante. Enfin, f est (strictement) monotone lorsqu'elle est croissante OU décroissante (strictement).

Remarque 3. La somme de deux fonctions croissantes est croissante. On ne peut en revanche « rien » dire de la différence de deux fonctions croissantes.

Proposition 3. La composée de deux fonctions monotones est monotone. Le sens de monotonie obéit à la règle des signes.

Démonstration. Laissé en exercice. $\square$

Exemple. La fonction $x \mapsto \exp\left(-\frac{1}{x^2}\right)$ est décroissante sur $\mathbb{R}_+^*$.

4 Extrema

Définition 5. Soit I un intervalle de $\mathbb{R}$. Soit $f : I \longrightarrow \mathbb{R}$. Soit $a \in I$. On dit que f admet un maximum (global) en a lorsque

$$\forall x \in I, f(x) \leq f(a)$$

Le réel $f(a)$ est le *maximum* de f sur I . On définit de même la notion de minimum.

Définition 6. Soit I un intervalle de $\mathbb{R}$. Soit $f : I \longrightarrow \mathbb{R}$. Soit $a \in I$. On dit que f

admet un maximum local en a lorsque

$$\exists \delta > 0, \forall x \in I, |x - a| \leq \delta \implies f(x) \leq f(a)$$

On définit de même la notion de minimum local.

Remarque 4. Si une fonction admet un extremum global en un point, c'est aussi bien entendu un extremum local. La réciproque est fausse. Nous verrons plus tard des conditions d'existence d'extrema locaux ou globaux, liées à la nature de l'intervalle I (segment) et à la régularité de f (continuité, dérivabilité).

5 Parité, périodicité

5.1 Fonctions paires, fonctions impaires

Définition 7. Soit $X \subset \mathbb{R}$. Soit $f : X \longrightarrow \mathbb{R}$, où X est une partie de $\mathbb{R}$ symétrique par rapport à 0 : $\forall x \in X, -x \in X$.

- On dit que f est paire lorsque

$$\forall x \in X, f(-x) = f(x)$$

- On dit que f est impaire lorsque

$$\forall x \in X, f(-x) = -f(x)$$

Proposition 4. Toute fonction $X \longrightarrow \mathbb{R}$ s'écrit de façon unique comme somme d'une fonction paire et d'une fonction impaire.

Démonstration. Soit $f : X \longrightarrow \mathbb{R}$. Supposons que $f = g + h$ où g est paire et h est impaire. On a pour tout $x \in X$ les deux égalités

$$f(x) = g(x) + h(x) \text{ et } f(-x) = g(x) - h(x)$$

On en déduit

$$g(x) = \frac{f(x) + f(-x)}{2} \text{ et } h(x) = \frac{f(x) - f(-x)}{2}$$

d'où l'unicité. Inversement, on vérifie que ces deux fonctions sont bien respectivement paire et impaire, et que leur somme est f , d'où l'existence. $\square$

Exemple. On a $e^x = \cosh x + \sinh x$ où $\cosh x = \frac{e^x + e^{-x}}{2}$ et $\sinh x = \frac{e^x - e^{-x}}{2}$. Ces deux fonctions sont les fonctions cosinus et sinus hyperboliques. Nous les étudierons plus loin dans ce chapitre.

Proposition 5. Soit $f : X \rightarrow \mathbb{R}$. Alors,

- f est paire si et seulement si son graphe est symétrique par rapport à l'axe Oy .
- f est impaire si et seulement si son graphe est symétrique par rapport au point O .

Démonstration. Supposons f paire. Notons $\mathcal{G}$ le graphe de f . Rappelons que

$$\mathcal{G} = \{(x, f(x)), x \in X\}$$

Soit $s : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ définie par

$$s(x, y) = (-x, y)$$

L'application s est la *symétrie orthogonale* par rapport à l'axe Oy . Il s'agit de montrer que

$$s(\mathcal{G}) = \mathcal{G}$$

Soit $(a, b) \in s(\mathcal{G})$. Il existe donc $(x, y) \in \mathcal{G}$ tel que $(a, b) = s(x, y)$. On a $x \in X$, $y = f(x)$, $a = -x$ et $b = y$. Comme X est symétrique par rapport à O , $a \in X$. De plus,

$$b = y = f(x) = f(-x) = f(a)$$

Ainsi, $(a, b) \in \mathcal{G}$. Nous avons donc prouvé que $s(\mathcal{G}) \subset \mathcal{G}$.

Appliquons s à l'inclusion précédente. Il vient

$$s \circ s(\mathcal{G}) \subset s(\mathcal{G})$$

Or, $s \circ s = \text{id}_{\mathbb{R}^2}$. Ainsi, $\mathcal{G} \subset s(\mathcal{G})$.

Inversement, supposons que $s(\mathcal{G}) = \mathcal{G}$. Soit $x \in X$. Soit $y = f(x)$. On a $(x, y) \in \mathcal{G}$, donc $s(x, y) = (-x, y) \in \mathcal{G}$. On en déduit que $-x \in X$ et $y = f(-x)$, d'où $f(-x) = f(x)$. La fonction f est paire.

Le cas où f est impaire est laissé en exercice. $\square$

5.2 Fonctions périodiques

Définition 8. Soit $X \subset \mathbb{R}$. Soit $f : X \rightarrow \mathbb{R}$. Soit $T \in \mathbb{R}$. On dit que T est une période de f lorsque

- $\forall x \in X, x + T \in X$.
- $\forall x \in X, f(x + T) = f(x)$.

La fonction f est dite périodique si elle admet au moins une période non nulle.

On note $\mathcal{T}(f)$ l'ensemble des périodes de la fonction f .

Proposition 6. *Soit $f : X \rightarrow \mathbb{R}$. L'ensemble $\mathcal{T}(f)$ est un groupe pour l'addition : il contient 0, et il est stable pour l'addition et par passage à l'opposé.*

Démonstration. Il est clair que 0 est une période de f .

Soit T une période de f . On a pour tout $x \in X$ $f(x - T) = f(x - T + T) = f(x)$. Donc $-T$ est une période de f .

Enfin, si T_1 et T_2 sont deux périodes de f , on a pour tout $x \in X$ $f(x + T_1 + T_2) = f(x + T_1) = f(x)$ donc $T_1 + T_2$ est une période de f . $\square$

Remarque 5. On peut montrer que si G est un sous-groupe de $\mathbb{R}$, alors on est dans un (et un seul) des trois cas ci-dessous :

- $G = \{0\}$.
- G possède un plus petit élément strictement positif : dans ce cas, il existe un réel $\alpha > 0$ tel que $G = \alpha\mathbb{Z}$.
- G ne possède pas de plus petit élément strictement positif : G est alors dense dans $\mathbb{R}$.

Ceci signifie qu'il existe deux types de fonctions périodiques :

- Les fonctions dont le groupe des périodes est dense dans $\mathbb{R}$. Ce sont des fonctions très compliquées.
- Les fonctions dont le groupe des périodes est de la forme $T_0\mathbb{Z}$ avec $T_0 > 0$. On peut montrer que c'est par exemple le cas lorsque la fonction est continue en au moins 1 point. C'est le cas simple, la fonction possède alors une plus petite période strictement positive que l'on appelle SA période.

Proposition 7. *Soit T un réel. L'ensemble des fonctions T -périodiques sur l'ensemble X est une algèbre.*

Démonstration. Cette proposition nous dit essentiellement que la somme, le produit de deux fonctions T -périodiques, le produit d'une fonction T -périodique par un réel, sont encore des fonctions T -périodiques. $\square$

Chapitre 1 1

Limites - Continuité

1 Étude locale d'une fonction

1.1 Un peu de topologie

Définition 1. Soit $a \in \overline{\mathbb{R}}$.

- Si $a \in \mathbb{R}$, un voisinage de a est un segment $[a - \varepsilon, a + \varepsilon]$ où $\varepsilon > 0$.
- Si $a = +\infty$, un voisinage de a est un intervalle $[M, +\infty[$, où $M \in \mathbb{R}$.
- Si $a = -\infty$, un voisinage de a est un intervalle $] -\infty, M]$, où $M \in \mathbb{R}$.

Les trois propositions ci-dessous sont faciles à montrer. Leur preuve est laissée en exercice.

Proposition 1. Soit $a \in \overline{\mathbb{R}}$. L'intersection de deux voisinages de a est encore un voisinage de a .

Proposition 2. Soit $a \in \overline{\mathbb{R}}$. Si $a \in \mathbb{R}$, l'intersection des voisinages de a est le singleton $\{a\}$. Si $a = \pm\infty$, l'intersection des voisinages de a est vide.

Proposition 3. Soient $a, b \in \overline{\mathbb{R}}$, $a \neq b$. Il existe un voisinage V de a et un voisinage W de b tels que $V \cap W = \emptyset$.

Notation. On note $\mathcal{V}(a)$ l'ensemble des voisinages de a .

Définition 2. Soit A une partie de $\mathbb{R}$. Soit $a \in \overline{\mathbb{R}}$. On dit que a est *adhérent* à A lorsque tout voisinage de a rencontre A .

Exemple. Les points adhérents à un intervalle sont les points de l'intervalle et ses bornes.

Notation. On note $\overline{A}$ l'ensemble des points adhérents à A . Cet ensemble est appelé l'adhérence de A .

Exemple. L'adhérence d'un intervalle est l'intervalle fermé correspondant.

Définition 3. Soit $\mathcal{P}(x)$ une propriété du réel x . Soit $a \in \overline{\mathbb{R}}$. On dit que la propriété $\mathcal{P}$ est vraie au voisinage de a lorsque $\exists V \in \mathcal{V}(a), \forall x \in V, \mathcal{P}(x)$.

1.2 Limite et continuité en un point

Définition 4. Soit A une partie de $\mathbb{R}$. Soit $f : A \rightarrow \mathbb{R}$. Soit $a \in \overline{A}$. Soit $\ell \in \overline{\mathbb{R}}$. On dit que f tend vers ℓ en a lorsque

$$\forall V \in \mathcal{V}(\ell), \exists W \in \mathcal{V}(a), \forall x \in W \cap A, f(x) \in V$$

Notation. On note $f(x) \xrightarrow{x \rightarrow a} \ell$.

Remarque 1. Les ensembles A et les objets a que nous considérerons usuellement seront de l'une des formes ci-dessous :

- A est un intervalle et a est un point de A .
- A est un intervalle et a est une borne de A , $a \notin A$.
- A est un intervalle privé d'un point a . On parle alors quelquefois d'intervalle épointé.

Proposition 4. La limite d'une fonction en un point, si elle existe, est unique.

Démonstration. Supposons que f a deux limites ℓ et ℓ' en a . Soient V un voisinage de ℓ et V' un voisinage de ℓ' . Il existe W, W' voisinages de a tels que

$$f(W \cap A) \subset V \text{ et } f(W' \cap A) \subset V'$$

$W \cap W'$ est encore un voisinage de a et

$$f(W \cap W' \cap A) \subset f(W \cap A) \subset V$$

De même,

$$f(W \cap W' \cap A) \subset V'$$

Ainsi, $V \cap V' \neq \emptyset$. Ceci étant vrai pour tous les voisinages de ℓ et ℓ' , on a donc $\ell = \ell'$. $\square$

Remarque 2. On parle donc de LA limite de f en a (si elle existe!). L'unicité de la limite justifie la notation $\ell = \lim_a f$ ou notations semblables.

Proposition 5. Si une fonction a une limite finie en un point a , cette fonction est bornée au voisinage de a .

Démonstration. Supposons que f tend vers $\ell \in \mathbb{R}$. Pour tout x assez proche de a , on a donc $|f(x) - \ell| \leq 1$, donc $\ell - 1 \leq f(x) \leq \ell + 1$. La fonction f est donc bornée au voisinage de a . $\square$

Proposition 6. Soit $f : A \longrightarrow \mathbb{R}$. Soit $a \in A$. Si f a une limite en a , ce ne peut être que $f(a)$.

Démonstration. Supposons $f(x) \longrightarrow \ell$ lorsque $x \longrightarrow a$. Pour tout voisinage V de ℓ , il existe un voisinage W de a tel que $f(W \cap A) \subset V$. Mais a appartient à $W \cap A$, donc, $f(a) \in V$. Ainsi, $f(a)$ est dans tous les voisinages de ℓ . Donc, ℓ est réel et $\ell = f(a)$. $\square$

Remarque 3. Ce que nous venons de montrer peut se révéler gênant dans certains cas. Ainsi, il arrive parfois que l'on « retire » a de l'ensemble de définition de f et que l'on considère la limite de $f(x)$ lorsque x tend vers a , $x \neq a$.

Définition 5. Soit $f : A \longrightarrow \mathbb{R}$. Soit $a \in A$. On dit que f est continue en a lorsque f a une limite en a (qui est donc forcément $f(a)$).

Remarque 4. f est continue en a si et seulement si

$$\forall \varepsilon > 0, \exists \alpha > 0, \forall x \in A, |x - a| \leq \alpha \Rightarrow |f(x) - f(a)| \leq \varepsilon$$

Définition 6. Soit $f : A \longrightarrow \mathbb{R}$. On dit que f est continue sur A lorsque f est continue en tout point de A .

Notation. Soit $A \subset \mathbb{R}$. On note $\mathcal{C}^0(A)$ l'ensemble des fonctions continues sur A .

1.3 Reformulations de la notion de limite

La définition de limite par les voisinages a l'avantage d'être valable dans tous les cas, en un point fini ou pas, pour une limite finie ou pas. C'est très bien pour faire des démonstrations très générales, mais cela l'est beaucoup moins dans les situations concrètes. On se donne $f : A \longrightarrow \mathbb{R}$. Soit $a \in \overline{A}$. Soit $\ell \in \mathbb{R}$. On réécrit selon la finitude ou non de a et ℓ la définition de limite.

Limite réelle en un point réel

Proposition 7. $f(x)$ tend vers ℓ lorsque x tend vers a si et seulement si

$$\forall \varepsilon > 0, \exists \alpha > 0, \forall x \in A, |x - a| \leq \alpha \implies |f(x) - \ell| \leq \varepsilon$$

Limite finie à l'infini

Proposition 8. $f(x)$ tend vers ℓ lorsque x tend vers $+\infty$ si et seulement si

$$\forall \varepsilon > 0, \exists M \in \mathbb{R}, \forall x \in A, x \geq M \implies |f(x) - \ell| \leq \varepsilon$$

Remarque 5. On définit de même la notion de limite réelle en $-\infty$.

Limite nulle

Les fonctions ayant une limite nulle en un point donné possèdent des propriétés remarquables.

Proposition 9. Soit I un intervalle de $\mathbb{R}$. Soit $a \in \bar{I}$.

- Une fonction f a une limite nulle en a si et seulement si $|f|$ a une limite nulle en a .
- Si f a une limite nulle en a , et g est bornée au voisinage de a , alors fg a une limite nulle en a .

Démonstration. Démonstration analogue à celle qui a été faite pour les suites. $\square$

Limite infinie en un point réel

Définition 7. $f(x)$ tend vers $+\infty$ lorsque x tend vers $a \in \mathbb{R}$ si et seulement si

$$\forall M \in \mathbb{R}, \exists \alpha > 0, \forall x \in A, |x - a| \leq \alpha \implies f(x) \geq M$$

Limite infinie à l'infini

Définition 8. $f(x)$ tend vers $+\infty$ lorsque x tend vers $+\infty$ si et seulement si

$$\forall M \in \mathbb{R}, \exists M' \in \mathbb{R}, \forall x \in A, x \geq M' \implies f(x) \geq M$$

On définit de même des limites en $-\infty$, égales à $\pm\infty$. Il suffit de changer le sens des inégalités.

1.4 Opérations sur les limites

Proposition 10. Soient $f, g : A \rightarrow \mathbb{R}$. Soit $a \in \overline{A}$. Soit $\lambda \in \mathbb{R}$. Soient $\ell, \ell' \in \mathbb{R}$. On suppose que $f(x) \rightarrow \ell$ et $g(x) \rightarrow \ell'$ lorsque $x \rightarrow a$. Alors :

- $f(x) + g(x) \rightarrow \ell + \ell'$ lorsque $x \rightarrow a$.
- $f(x)g(x) \rightarrow \ell\ell'$ lorsque $x \rightarrow a$.
- $\lambda f(x) \rightarrow \lambda\ell$ lorsque $x \rightarrow a$.

Démonstration. Démonstration analogue à celle qui a été faite pour les suites. $\square$

Remarque 6. Les résultats précédents subsistent lorsque $\ell, \ell' \in \overline{\mathbb{R}}$, sauf dans les cas d'indétermination « $\infty - \infty$ » et « $0 \times \infty$ ».

Proposition 11. Soit $g : A \rightarrow \mathbb{R}$. Soit $a \in \overline{A}$. Soit $\ell \in \overline{\mathbb{R}}_+^*$. On suppose que $g(x)$ tend vers ℓ lorsque x tend vers a . Alors :

- Il existe un réel $m > 0$ tel que $g(x) \geq m$ au voisinage de a .
- $\frac{1}{g(x)}$ tend vers $\frac{1}{\ell}$ lorsque x tend vers a .

Démonstration. Démonstration analogue à celle qui a été faite pour les suites. $\square$

Remarque 7. Le résultat précédent s'adapte évidemment au cas où $\ell < 0$. Il convient en revanche de faire attention au cas où $\ell = 0$. Il faut dans ce cas *supposer* que la fonction g a un signe constant au voisinage de a pour conclure que g tend vers l'infini en a .

Lemme 12. Soient A et B deux parties de $\mathbb{R}$. Soit $f : A \rightarrow B$. Soit $a \in \overline{A}$. Soit $b \in \overline{B}$. On suppose que $f(x) \rightarrow b$ lorsque $x \rightarrow a$. Alors b appartient à $\overline{B}$.

Démonstration. Soit V un voisinage de b . Il existe un voisinage W de a tel que $f(W \cap A) \subset V$. Comme a est adhérent à A , $W \cap A$ est non vide, donc $f(W \cap A)$ est non vide aussi. Mais $f(W \cap A) \subset B$, donc $V \cap B$ est non vide. Ceci étant vrai pour tout voisinage V de b , on a bien $b \in \overline{B}$. $\square$

Proposition 13. THÉORÈME DE COMPOSITION DES LIMITES.

Soient $f : A \rightarrow B$ et $g : B \rightarrow \mathbb{R}$. Soit $a \in \overline{A}$, $b \in \overline{B}$, $\ell \in \overline{\mathbb{R}}$. On suppose que $f(x) \rightarrow b$ lorsque $x \rightarrow a$, et $g(y) \rightarrow \ell$ lorsque $y \rightarrow b$. Alors $g \circ f(x) \rightarrow \ell$ lorsque $x \rightarrow a$.

Démonstration. Soit V un voisinage de ℓ . Il existe un voisinage W de b tel que $g(W \cap B) \subset V$. Il existe ensuite un voisinage X de a tel que $f(X \cap A) \subset W$. On a alors $g \circ f(X \cap A) \subset g(W \cap B) \subset V$. La fonction $g \circ f$ tend bien vers ℓ en a . $\square$

1.5 Limites de fonctions, limites de suites

Proposition 14. CARACTÉRISATION SÉQUENTIELLE DES LIMITES

Soit $f : A \rightarrow \mathbb{R}$. Soit $a \in \overline{A}$. Soit $\ell \in \overline{\mathbb{R}}$. On a $f(x) \rightarrow \ell$ lorsque $x \rightarrow a$ si et seulement si pour toute suite $(x_n)_{n \geq 0}$ de points de A qui tend vers a , la suite $(f(x_n))_{n \geq 0}$ tend vers ℓ .

Démonstration. On fait la démonstration pour a et ℓ réels. Supposons que $f(x) \rightarrow \ell$ lorsque $x \rightarrow a$. Soit (x_n) une suite de points de A qui converge vers a . Soit $\varepsilon > 0$. Il existe $\delta > 0$ tel que

$$\forall x \in A, |x - a| \leq \delta \Rightarrow |f(x) - \ell| \leq \varepsilon$$

Soit $N \in \mathbb{N}$ tel que pour tout $n \geq N$, $|x_n - a| \leq \delta$. On a alors $|f(x_n) - \ell| \leq \varepsilon$. La suite $(f(x_n))$ tend donc vers ℓ .

Supposons maintenant que $f(x) \not\rightarrow \ell$ lorsque x tend vers a . Il existe donc $\varepsilon > 0$ tel que pour tout $\delta > 0$ il existe $x \in A$ tel que $|x - a| \leq \delta$ et $|f(x) - \ell| > \varepsilon$. Pour tout entier $n \geq 1$, prenons $\delta_n = \frac{1}{n}$. Soit $x_n \in A$ tel que $|x_n - a| \leq \frac{1}{n}$ et $|f(x_n) - \ell| > \varepsilon$. On a $x_n \rightarrow a$ lorsque $n \rightarrow \infty$. Mais si l'on suppose que $f(x_n) \rightarrow \ell$, on obtient, par passage à la limite, $0 \geq \varepsilon$, contradiction. La suite $(f(x_n))$ ne tend donc pas vers ℓ . $\square$

Exemple. La fonction $f : x \mapsto \sin \frac{1}{x}$ n'a pas de limite en 0. En effet, supposons que f a une limite ℓ en 0. Alors $0 = f\left(\frac{1}{2n\pi}\right) \rightarrow \ell$ donc $\ell = 0$. Mais aussi, $1 = f\left(\frac{1}{2n\pi + \frac{\pi}{2}}\right) \rightarrow \ell$ donc $\ell = 1$, contradiction.

1.6 limites et ordre

Proposition 15. PASSAGE À LA LIMITE DANS UNE INÉGALITÉ

Soient $f, g : A \rightarrow \mathbb{R}$. Soit $a \in \overline{A}$. On suppose :

- $f(x) \rightarrow \ell$ lorsque $x \rightarrow a$.
- $g(x) \rightarrow \ell'$ lorsque $x \rightarrow a$.
- $f \leq g$ au voisinage de a .

Alors, $\ell \leq \ell'$.

Démonstration. Démonstration analogue à celle du théorème correspondant sur les suites. $\square$

Proposition 16. THÉORÈME D'ENCADREMENT

Soient $f, g, h : A \rightarrow \mathbb{R}$. Soit $a \in \overline{A}$. Soit $\ell \in \mathbb{R}$. On suppose :

- $f \leq g \leq h$ au voisinage de a .
- $f(x) \rightarrow \ell$ lorsque $x \rightarrow a$.
- $h(x) \rightarrow \ell$ lorsque $x \rightarrow a$.

Alors, $g(x) \rightarrow \ell$ lorsque $x \rightarrow a$.

Démonstration. Démonstration analogue à celle du théorème correspondant sur les suites. $\square$

Proposition 17. THÉORÈME D'ENCADREMENT

Soient $f, g : A \rightarrow \mathbb{R}$. Soit $a \in \overline{A}$. On suppose :

- $f \leq g$ au voisinage de a .
- $f(x) \rightarrow +\infty$ lorsque $x \rightarrow a$.

Alors, $g(x) \rightarrow +\infty$ lorsque $x \rightarrow a$.

Démonstration. Démonstration analogue à celle du théorème correspondant sur les suites. $\square$

1.7 Limites dans une direction

Définition 9. Soit $f : I \rightarrow \mathbb{R}$, où I est un intervalle de $\mathbb{R}$. Soit $a \in I$. On suppose que a n'est pas la borne inférieure de I . On définit la limite à gauche de f en a comme la limite (si elle existe !) de la restriction de f à $I \cap]-\infty, a[$, en a . On note cette limite $\lim_{x \rightarrow a, x < a} f(x)$ ou encore $f(a^-)$. On définit de même la limite à droite de f en a .

Remarque 8. Si a est la borne inférieure de I , la notion de limite à droite de $f(x)$ lorsque $x \rightarrow a$ et celle de limite de $f(x)$ lorsque $x \rightarrow a, x \neq a$ coïncident. Même remarque pour la limite à gauche lorsque a est la borne supérieure de I .

Proposition 18. Soit $f : I \rightarrow \mathbb{R}$. Soit $a \in I$, a n'étant pas la borne inférieure de I . La fonction f admet une limite en a si et seulement si elle y admet une limite à gauche et une limite à droite, et si celles-ci sont égales à $f(a)$. On a alors

$$\lim_a f = f(a^-) = f(a^+) = f(a)$$

Démonstration. Laissée en exercice. $\square$

Remarque 9. Si f est définie sur l'intervalle épointé $I \setminus \{a\}$, il convient d'adapter la proposition précédente en y enlevant les références à $f(a)$.

1.8 Fonctions monotones

Dans ce paragraphe, les fonctions sont définies sur des intervalles.

On se donne une fonction $f : I \rightarrow \mathbb{R}$ croissante. Soit également $a \in \bar{I}$.

Cas où a n'est pas une borne de I

Proposition 19. f admet une limite à gauche et une limite à droite réelles en a . De plus,

- $f(a^-) = \sup\{f(x), x < a\}$.
- $f(a^+) = \inf\{f(x), x > a\}$.
- $f(a^-) \leq f(a) \leq f(a^+)$.

Démonstration. Soit $E = \{f(x), x \in I, x < a\} = f(I \cap]-\infty, a[)$. L'ensemble E est non vide car a n'est pas la borne inférieure de I . De plus, E est majoré par $f(a)$ puisque f est croissante. Donc, E possède une borne supérieure $\ell \leq f(a)$. Soit $\varepsilon > 0$. Comme $\ell - \varepsilon < \ell$, il existe $y_0 \in E$ tel que $y_0 > \ell - \varepsilon$. Autrement dit, il existe $x_0 \in I, x_0 < a$ tel que $f(x_0) > \ell - \varepsilon$. Posons $\delta = a - x_0$. On a alors pour tout $x \in I, x < a, |x - a| = a - x \leq \delta \Rightarrow x_0 \geq x < a \Rightarrow \ell - \varepsilon \leq f(x_0) \leq f(x)$. De plus, $f(x) \in E$ donc $f(x) \leq \ell \leq \ell + \varepsilon$. Bref, $|f(x) - \ell| \leq \varepsilon$. Ainsi, $f(a^-)$ existe et vaut ℓ . Même démonstration pour la limite à droite. $\square$

Cas où a est la borne supérieure de I , et $a \in I$

Proposition 20. f admet une limite à gauche réelle en a . De plus,

- $f(a^-) = \sup\{f(x), x < a\}$.
- $f(a^-) \leq f(a)$.

Démonstration. Voir le cas précédent. $\square$

Cas où a est la borne supérieure de I , mais $a \notin I$

Proposition 21. f admet une limite à gauche réelle ou égale à $+\infty$ en a . De plus,

$$f(a^-) = \sup\{f(x), x < a\}$$

Si f est majorée sur I , ce sup est réel. Sinon, il est égal à $+\infty$.

Démonstration. La seule différence, c'est que E n'est plus nécessairement majoré. S'il l'est, f a une limite réelle en a . S'il ne l'est pas, f tend vers $+\infty$ en a . $\square$

Le cas de la borne inférieure de I , et celui des fonctions décroissantes, se traitent évidemment de façon identique.

2 Propriétés des fonctions continues

2.1 Opérations sur les fonctions continues

Proposition 22. *Soit $A \subset \mathbb{R}$. L'ensemble $\mathcal{C}^0(A)$ des fonctions continues sur A est une algèbre.*

Démonstration. Ce sont tout simplement les opérations sur les limites qui nous le disent. $\square$

Proposition 23. *La composée de deux fonctions continues est encore une fonction continue.*

Démonstration. Théorème de composition des limites. $\square$

Remarque 10. La fonction identité $x \mapsto x$ est clairement continue sur $\mathbb{R}$. Il s'ensuit que toute fonction polynôme

$$P(x) = \sum_{k=0}^n a_k x^k$$

est continu sur $\mathbb{R}$. De là, les fractions rationnelles

$$F(x) = \frac{P(x)}{Q(x)}$$

sont continues en tout point où Q ne s'annule pas.

Nous verrons plus tard que la fonction exponentielle, les fonctions sinus et cosinus, sont continues sur $\mathbb{R}$. La fonction logarithme, les fonctions $x \mapsto x^\alpha$ (où $\alpha \in \mathbb{R}$) sont continues sur $\mathbb{R}_+^*$.

2.2 Prolongement par continuité

Étant donnée une fonction $f :]a, \dots[\rightarrow \mathbb{R}$, on cherche souvent à prolonger f à $[a, \dots[$. Un tel prolongement est particulièrement intéressant lorsqu'il est continu.

Définition 10. Soit $a \in \mathbb{R}$. Soit $f :]a, \dots | \rightarrow \mathbb{R}$ continue sur $]a, \dots |$. On dit que f est prolongeable par continuité en a lorsqu'il existe un prolongement de f à $[a, \dots |$ qui soit continu.

Proposition 24. *Un tel prolongement existe si et seulement si f admet une limite finie en a .*

Démonstration. Supposons l'existence d'un prolongement g . On a $g(x) \rightarrow g(a)$ lorsque $x \rightarrow a$. Mais hormis en a , f et g coïncident. donc $f(x) \rightarrow g(a)$ lorsque $x \rightarrow a$. Inversement, supposons que $f(x) \rightarrow \ell \in \mathbb{R}$ lorsque $x \rightarrow a$. Le prolongement g de f en a défini par $g(a) = \ell$ est alors continu. $\square$

2.3 Image d'un intervalle

Proposition 25. THÉORÈME DES VALEURS INTERMÉDIAIRES

Soient $a, b \in \mathbb{R}$ tels que $a < b$. Soit $f : [a, b] \rightarrow \mathbb{R}$, continue sur $[a, b]$. On suppose $f(a) < 0$ et $f(b) > 0$. Alors, il existe $c \in [a, b]$ tel que $f(c) = 0$.

Démonstration. On définit par récurrence sur n deux suites $(a_n)_{n \geq 0}$ et $(b_n)_{n \geq 0}$ de points de $[a, b]$. Tout d'abord, $a_0 = a$ et $b_0 = b$. Pour tout entier $n \in \mathbb{N}$, soit $c_n = \frac{1}{2}(a_n + b_n)$.

- Si $f(c_n) < 0$, on pose $a_{n+1} = c_n$ et $b_{n+1} = b_n$.
- Sinon, on pose $a_{n+1} = a_n$ et $b_{n+1} = c_n$.

On montre alors facilement par récurrence sur n que pour tout $n \in \mathbb{N}$,

- $a \leq a_n \leq a_{n+1} < b_{n+1} \leq b_n \leq b$.
- $f(a_n) < 0$, $f(b_n) \geq 0$.
- $b_n - a_n = \frac{1}{2^n}(b - a)$.

Les suites $(a_n)_{n \geq 0}$ et $(b_n)_{n \geq 0}$ sont donc des suites adjacentes de points de $[a, b]$. Ainsi, elles convergent vers une limite commune $c \in [a, b]$.

Comme $c \in [a, b]$, f est continue en c . Par le théorème de caractérisation séquentielle des limites, $f(a_n) \rightarrow f(c)$ lorsque n tend vers l'infini. Or, $f(a_n) < 0$. Par passage à la limite dans les inégalités larges, on en déduit que $f(c) \leq 0$. De même, $f(b_n) \rightarrow f(c)$ lorsque n tend vers l'infini et $f(b_n) \geq 0$, donc $f(c) \geq 0$. En conclusion, $f(c) = 0$. $\square$

Remarque 11. La preuve ci-dessus donne une méthode *effective* pour calculer des approximations d'une racine de f . La fonction Python ci-dessous fait le travail en question. Elle prend en paramètres une fonction f , deux réels $a < b$ tels que $f(a) < 0$ et $f(b) \geq 0$,

et un réel $\varepsilon > 0$. Elle renvoie un couple (x, y) tel que $x < y$, $[x, y]$ contient une racine de f , et $y - x \leq \varepsilon$.

```
def racine(f, a, b, eps):
    while b - a > eps:
        c = (a + b) / 2
        if f(c) < 0: a = c
        else: b = c
    return (a, b)
```

Exemple. Soit $f : [0, 1] \rightarrow [0, 1]$ continue. Il existe $x \in [0, 1]$ tel que $f(x) = x$. Il suffit de considérer la fonction $g : x \mapsto f(x) - x$. Cette fonction est continue, $g(0) \geq 0$ et $g(1) \leq 0$.

Corollaire 26. THÉORÈME DES VALEURS INTERMÉDIAIRES

Soit $f : [a, b] \rightarrow \mathbb{R}$ continue. Soient $\alpha = f(a), \beta = f(b)$. Soit γ un réel compris entre α et β . Il existe $c \in [a, b]$ tel que $f(c) = \gamma$.

Démonstration. Supposons par exemple $\alpha \leq \beta$. Considérons la fonction $g : x \mapsto f(x) - \gamma$. Cette fonction est continue, négative en a et positive en b . Donc, elle s'annule. $\square$

Corollaire 27. THÉORÈME DES VALEURS INTERMÉDIAIRES

L'image d'un intervalle par une fonction continue est un intervalle.

Démonstration. Soit I un intervalle de $\mathbb{R}$. Soit $f : I \rightarrow \mathbb{R}$ continue. Soit $J = f(I)$. Montrons que J est convexe. Ce sera donc un intervalle. Soient $\alpha, \beta \in J$. Il existe donc $a, b \in I$ tels que $f(a) = \alpha$ et $f(b) = \beta$. Soit γ un réel entre α et β . f étant continue sur $[a, b]$, il existe donc $c \in [a, b]$ tel que $f(c) = \gamma$. Mais I est convexe puisque c'est un intervalle. Donc $c \in I$ et $\gamma = f(c) \in J$. Ainsi, J est convexe. J est bien un intervalle. $\square$

2.4 Image d'un segment

Proposition 28. *Soit $f : [a, b] \rightarrow \mathbb{R}$, continue sur $[a, b]$. Alors, f admet un maximum et un minimum sur le segment $[a, b]$.*

Démonstration. Considérons l'ensemble $E = f([a, b])$. Cet ensemble est non vide, et possède donc une borne supérieure $M \in \mathbb{R}$, éventuellement égale à $+\infty$. Soit (x_n) une suite de points de $[a, b]$ telle que $f(x_n) \rightarrow M$. La suite (x_n) est bornée, elle admet donc d'après le théorème de Bolzano-Weierstrass une suite extraite $(x_{\varphi(n)})$ convergeant vers un réel c . Comme on a pour tout n $a \leq x_{\varphi(n)} \leq b$, on a par passage à la limite $a \leq c \leq b$. La fonction f est donc continue en c , et ainsi, $f(x_{\varphi(n)}) \rightarrow f(c)$. Mais on a aussi que $f(x_{\varphi(n)}) \rightarrow M$.

Donc $M = f(c)$. On en déduit que M est réel, et aussi que M est une valeur prise par f : $M = \max E$. On fait de même pour le minimum. $\square$

Corollaire 29. *L'image d'un segment par une fonction continue est un segment.*

Démonstration. L'image d'un segment par une fonction continue est un intervalle (théorème des valeurs intermédiaires) possédant un plus petit et un plus grand élément. Pas de doute, c'est un segment. $\square$

2.5 Fonctions continues strictement monotones

Proposition 30. *Soit $f : I \longrightarrow \mathbb{R}$ continue. Si f est injective, alors f est strictement monotone.*

Démonstration. Supposons f continue et injective. Soit $T = \{(x, y) \in I^2, x < y\}$. Soit $F : T \longrightarrow \mathbb{R}$ définie par

$$F(x, y) = \frac{f(y) - f(x)}{y - x}$$

Remarquons que pour tout $(x, y) \in T$, on a $x \neq y$. En conséquence, F est bien définie et F ne s'annule pas car f est injective. Il s'agit de montrer que F est de signe constant sur T . Pour cela, soient (a, b) et (x, y) deux éléments de T . On vérifie facilement que pour tout $t \in [0, 1]$,

$$((1 - t)a + tx, (1 - t)b + ty) \in T$$

Soit $\varphi : [0, 1] \longrightarrow \mathbb{R}$ définie par

$$\varphi(t) = F((1 - t)a + tx, (1 - t)b + ty)$$

La fonction φ est clairement continue sur $[0, 1]$ et ne s'annule pas. Par le théorème des valeurs intermédiaires, φ est de signe constant. Ainsi :

Pour tout $(x, y) \in T$, le signe de $F(x, y)$ est le même que celui de $F(a, b)$. La fonction F est bien de signe constant. $\square$

Il reste à examiner la réciproque.

Proposition 31. *Soit $f : I \longrightarrow \mathbb{R}$ continue, strictement croissante. Alors :*

- $J = f(I)$ est un intervalle de même « type » que I (i.e. fermé ou ouvert aux mêmes endroits).
- f est une bijection de I sur J .
- $f^{-1} : J \longrightarrow I$ est continue, strictement croissante.

Démonstration. Prenons par exemple $I = [a, b[$, où $-\infty < a < b \leq +\infty$. On sait que J est un intervalle puisque f est continue. Mieux, pour tout $x \in I$, $f(a) \leq f(x) \leq f(b^-)$. $f(a)$ est dans J , $f(b^-)$ est la borne supérieure de J (limite d'une fonction monotone), donc $J = [f(a), f(b^-)[$ ou alors $J = [f(a), f(b^-)]$.

Supposons que $f(b^-)$ soit dans J . Il existe alors $x < b$ tel que $f(x) = f(b^-)$. Mais alors, par croissance de f , f est constante sur $[x, b[$. Ce n'est pas possible puisque f est *strictement* croissante. f est évidemment injective, donc f est une bijection de I sur J . Notons g sa réciproque.

On voit facilement que g est strictement croissante. Il s'agit de prouver qu'elle est continue sur J .

Soit $y \in J$. Traitons le cas où $y > f(a)$. Posons $y = f(x)$ où $x \in I$ et $x > a$. Par ce que nous venons de voir appliqué à la restriction de f à $[a, x[$, on a $f([a, x[) = [f(a), f(x^-)[$. Mais $x \in I$, donc f est continue en x . Ainsi, $f(x^-) = f(x) = y$ et donc $f([a, x[) = [f(a), y[$. De là

$$g([f(a), y[) = g(f([a, x[)) = [a, x[$$

La fonction g est croissante, donc $g(y^-)$ existe. De plus,

$$g(y^-) = \sup g([f(a), y[) = \sup g(f([a, x[)) = \sup [a, x[= x = g(y)$$

De même, $g(y^+) = g(y)$. Ainsi, g est continue en y . $\square$

Exemple. Soit $n \in \mathbb{N}, n \geq 1$. La fonction $f : x \mapsto x^n$ est continue sur $\mathbb{R}_+$. Elle est strictement croissante. Donc $f(\mathbb{R}_+) = [f(0), \lim_{+\infty} f[= \mathbb{R}_+$. Cette fonction est donc une bijection de $\mathbb{R}_+$ sur $\mathbb{R}_+$ et sa réciproque est continue, strictement croissante : c'est la fonction « racine n ème ». Si n est impair, on peut se placer sur $\mathbb{R}$ tout entier.

Exemple. La fonction $f : x \mapsto \sin x$ est continue sur $[-\frac{\pi}{2}, \frac{\pi}{2}]$, strictement croissante. Donc $f([-\frac{\pi}{2}, \frac{\pi}{2}]) = [f(-\frac{\pi}{2}), f(\frac{\pi}{2})] = [-1, 1]$. f est donc une bijection de $[-\frac{\pi}{2}, \frac{\pi}{2}]$ sur $[-1, 1]$ et sa réciproque est continue, strictement croissante : c'est la fonction « arc sinus ».

Remarque 12. La fonction arc sinus n'est donc PAS la réciproque de la fonction sinus (qui, soit dit en passant, n'est pas une bijection, et n'a donc pas de réciproque). Nous en reparlerons dans le chapitre sur les fonctions usuelles.

2.6 Fonctions lipschitziennes

Définition 11. Soit $f : I \longrightarrow \mathbb{R}$. Soit $k \in \mathbb{R}_+$. On dit que f est k -lipschitzienne sur I lorsque

$$\forall x, y \in I, |f(y) - f(x)| \leq k|y - x|$$

Exemple. La fonction $x \mapsto \sin x$ est 1-lipschitzienne sur $\mathbb{R}$. La fonction $x \mapsto |x|$ est 1-lipschitzienne sur $\mathbb{R}$. La fonction $x \mapsto \sqrt{x}$, la fonction $x \mapsto x^2$, ne sont pas lipschitziennes sur $\mathbb{R}_+$.

Remarque 13. Si la fonction f est k -lipschitzienne, elle est aussi k' -lipschitzienne pour tout $k' \geq k$. En fait, l'ensemble des réels k tels que f soit k -lipschitzienne est un intervalle de la forme $[k_0, +\infty[$. Le réel k_0 est le coefficient de Lipschitz optimal pour la fonction f . Par exemple, $k_0 = 1$ pour la fonction sinus.

Proposition 32. *Toute fonction lipschitzienne est continue. La réciproque est fausse.*

Démonstration. Soit f k -lipschitzienne sur I ($k > 0$). Soit $\varepsilon > 0$. Soit $\delta = \frac{\varepsilon}{k}$. Alors, pour tous $x, y \in I$, $|x - y| \leq \delta$ implique $|f(x) - f(y)| \leq k|x - y| \leq k\delta = \varepsilon$. En revanche, la fonction $x \mapsto x^2$ est continue sur $\mathbb{R}$ mais pas lipschitzienne. $\square$

3 Continuité uniforme

3.1 Définitions

Définition 12. Soit $f : I \rightarrow \mathbb{R}$. On dit que f est uniformément continue sur I lorsque

$$\forall \varepsilon > 0, \exists \alpha > 0, \forall x, y \in I, |x - y| \leq \alpha \Rightarrow |f(x) - f(y)| \leq \varepsilon$$

Proposition 33. *Si f est uniformément continue sur I , alors f est continue sur I . La réciproque est fausse.*

Démonstration. Supposons f uniformément continue sur I . Soit $a \in I$. Soit $\varepsilon > 0$. Il existe $\alpha > 0$ tel que pour tous $x, y \in I$, $|x - y| \leq \alpha$ implique $|f(x) - f(y)| \leq \varepsilon$. C'est en particulier vrai pour $y = a$, d'où la continuité de f en a . Pour la réciproque, voir les deux exemples ci-dessous. $\square$

Exemple. Soit $f : \mathbb{R}_+ \rightarrow \mathbb{R}$ définie par $x \mapsto x^2$. Nous allons montrer que f n'est pas uniformément continue sur $\mathbb{R}_+$, c'est à dire que

$$\exists \varepsilon > 0, \forall \alpha > 0, \exists x, y \in I, |x - y| \leq \alpha \text{ et } |f(x) - f(y)| > \varepsilon$$

En d'autres termes, on peut trouver des x et des y aussi proches que l'on désire mais dont les images sont éloignées de plus de ε , ce ε restant à trouver. Pour $n \geq 1$, soient $x_n = n$ et $y_n = n + \frac{1}{n}$. Alors

$$|x_n - y_n| = \frac{1}{n}$$

qui est aussi petit qu'on le désire. Mais

$$|x_n^2 - y_n^2| = 2 + \frac{1}{n^2} > 2$$

Donc, $\varepsilon = 2$ fait l'affaire.

Exemple. Soit $f : \mathbb{R}_+^* \rightarrow \mathbb{R}$ définie par $x \mapsto \frac{1}{x}$. Nous allons montrer que f n'est pas uniformément continue sur $\mathbb{R}_+^*$. Pour $n \geq 1$, soient $x_n = \frac{1}{n}$ et $y_n = \frac{2}{n}$. Alors

$$|x_n - y_n| = \frac{1}{n}$$

qui est aussi petit qu'on le désire. Mais

$$\left| \frac{1}{x_n} - \frac{1}{y_n} \right| = \frac{n}{2} > 123$$

pour $n > 246$. Donc $\varepsilon = 123$ fait l'affaire.

Vici un exemple important de toute une classe de fonctions uniformément continues.

Proposition 34. *Toute fonction lipschitzienne est uniformément continue. La réciproque est fausse.*

Démonstration. Soit f k -lipschitzienne sur I ($k > 0$). Soit $\varepsilon > 0$. Soit $\delta = \frac{\varepsilon}{k}$. Alors, pour tous $x, y \in I$, $|x - y| \leq \delta$ implique

$$|f(x) - f(y)| \leq k|x - y| \leq k\delta = \varepsilon$$

Pour la réciproque prendre par exemple $x \mapsto x^2$ sur $\mathbb{R}$. Pour un contre exemple de fonction uniformément continue sur un segment mais pas lipschitzienne, voir un peu plus bas. $\square$

Proposition 35. THÉORÈME DE HEINE

Toute fonction continue sur un segment est uniformément continue.

Remarque 14. La fonction $x \mapsto \sqrt{x}$ est donc uniformément continue sur $[0, 1]$, alors qu'elle n'y est pas lipschitzienne.

Démonstration. Supposons f continue sur le segment $[a, b]$ mais pas uniformément continue. Il existe donc $\varepsilon > 0$ tel que, pour tout $\delta > 0$, il existe $x, y \in [a, b]$ tels que $|x - y| \leq \delta$ et $|f(x) - f(y)| > \varepsilon$.

Pour tout entier $n \geq 1$, prenons $\delta = \frac{1}{n}$. On a donc deux suites (x_n) et (y_n) d'éléments de $[a, b]$ telles que $|x_n - y_n| \leq \frac{1}{n}$ et $|f(x_n) - f(y_n)| > \varepsilon$. La suite (x_n) est bornée. d'après le théorème de Bolzano-Weierstrass, on peut en extraire une suite $(x_{\varphi(n)})$ convergeant vers c (évidemment, $c \in [a, b]$).

Comme $|x_{\varphi(n)} - y_{\varphi(n)}| \leq \frac{1}{\varphi(n)} \leq \frac{1}{n}$, la suite $(y_{\varphi(n)})$ converge aussi vers c . Mais f est continue en c donc $f(x_{\varphi(n)}) \rightarrow f(c)$ et $f(y_{\varphi(n)}) \rightarrow f(c)$. Par passage à la limite dans la

minoration par ε , on obtient $0 \geq \varepsilon$, contradiction. La fonction f est donc uniformément continue. $\square$

Remarque 15. La notion de continuité uniforme est une notion difficile à se représenter. Être continu ne suffit pas, sauf si l'on est sur un segment. Être lipschitzien est trop fort ... On peut démontrer que la fonction $f : I \rightarrow \mathbb{R}$ est uniformément continue sur I si et seulement si

$$\forall \varepsilon > 0, \exists k \in \mathbb{R}_+, \forall x, y \in I, |f(x) - f(y)| \leq k|x - y| + \varepsilon$$

c'est à dire qu'une fonction est uniformément continue si et seulement si, pour tout $\varepsilon > 0$, elle est k -lipschitzienne « à ε près », le k dépendant du ε . À méditer ...

4 Fonctions à valeurs complexes

Tous ce qui précède passe sans problème aux fonctions à valeurs complexes, à condition évidemment de ne pas faire intervenir de comparaisons entre des valeurs prises par la fonction. Passent donc aux fonctions à valeurs dans $\mathbb{C}$: définition de limite, opérations sur les limites, fonctions bornées, limites à gauche et à droite, relations de comparaison, continuité uniforme. Ne passent pas aux fonctions à valeurs dans $\mathbb{C}$: passage à la limite dans les inégalités, encadrement, image d'un intervalle, image d'un segment, et tout ce qui concerne les fonctions monotones. Bien entendu une fonction tend vers $\ell \in \mathbb{C}$ si et seulement si sa partie réelle tend vers $\operatorname{Re} \ell$ et sa partie imaginaire vers $\operatorname{Im} \ell$.

Chapitre 12

Dérivation

1 Notion de dérivée

Dans toute cette section, I désigne un intervalle non trivial de $\mathbb{R}$, c'est à dire contenant au moins deux points.

1.1 Dérivée d'ordre 1

Définition 1. Soit $f : I \rightarrow \mathbb{R}$. Soit $a \in I$. On dit que f est dérivable en a lorsque la quantité

$$\frac{f(x) - f(a)}{x - a}$$

a une limite réelle lorsque x tend vers a , $x \neq a$.

On appelle dérivée de f en a , et on note $f'(a)$ (ou $\frac{df}{dx}(a)$) la valeur de cette limite.

Proposition 1. Soit $f : I \rightarrow \mathbb{R}$. Soit $a \in I$. f est dérivable en a si et seulement si il existe un réel λ tel que, pour tout $x \in I$,

$$(*) \quad f(x) = f(a) + \lambda(x - a) + (x - a)\varepsilon(x)$$

où $\varepsilon : I \rightarrow \mathbb{R}$ est telle que $\varepsilon(x) \rightarrow 0$ lorsque $x \rightarrow a$. Lorsque f est dérivable en a , on a $\lambda = f'(a)$.

Démonstration. Supposons f dérivable en a . Posons, pour tout $x \in I \setminus \{a\}$,

$$\varepsilon(x) = \frac{f(x) - f(a)}{x - a} - f'(a)$$

Posons également $\varepsilon(a) = 0$. La fonction ε tend vers 0 en a , et on a pour tout $x \in I$,

$$f(x) - f(a) = (x - a)f'(a) + (x - a)\varepsilon(x)$$

On a donc $(*)$ avec $\lambda = f'(a)$. Inversement, supposons $(*)$. Pour tout $x \in I$ différent de a , on a

$$\frac{f(x) - f(a)}{x - a} = \lambda + \varepsilon(x)$$

Ainsi, lorsque x tend vers a ,

$$\frac{f(x) - f(a)}{x - a} \rightarrow \lambda$$

La fonction f est donc dérivable en a , et $f'(a) = \lambda$. $\square$

Proposition 2. Soit $f : I \rightarrow \mathbb{R}$. Soit $a \in I$. Si f est dérivable en a , alors f est continue en a . La réciproque est fausse.

Démonstration. La relation ci-dessus nous dit qu'en particulier $f(x) = f(a) + \varepsilon'(x)$ où $\varepsilon'(x) \rightarrow 0$ lorsque $x \rightarrow a$. Donc $f(x) \rightarrow f(a)$ lorsque $x \rightarrow a$. La réciproque est fausse, comme le montre l'exemple de fonction valeur absolue, continue en 0 mais pas dérivable en 0. $\square$

Définition 2. Soit $f : I \rightarrow \mathbb{R}$. On dit que f est dérivable *sur* I lorsque f est dérivable en tout point de I . On dispose alors de la fonction $f' : I \rightarrow \mathbb{R}$.

1.2 Tangente à la courbe

Soit $f : I \rightarrow \mathbb{R}$. Soit $a \in I$. Supposons f dérivable en a . Rappelons que l'on a alors pour tout $x \in I$ voisin de a ,

$$f(x) = f(a) + (x - a)f'(a) + (x - a)\varepsilon(x)$$

où $\varepsilon(x) \rightarrow 0$ lorsque $x \rightarrow a$. Soit $\varphi : I \rightarrow \mathbb{R}$ définie pour tout $x \in I$ par

$$\varphi(x) = f(a) + (x - a)f'(a)$$

On a donc pour tout x voisin de a

$$f(x) = \varphi(x) + (x - a)\varepsilon(x)$$

La fonction φ est une fonction affine et, d'une certaine façon, f est « proche » de φ au voisinage de a . La courbe de f et la droite d'équation $y = \varphi(x)$ sont elles aussi, d'une certaine façon, « proches ». Cette droite est particulièrement importante.

Définition 3. La tangente à la courbe de f au point d'abscisse a est la droite d'équation

$$(T) \quad y = f(a) + (x - a)f'(a)$$

Notons $A = (a, f(a))$. Pour tout x voisin de a différent de a , soit $B = (x, f(x))$. La droite D_x passant par A et B est appelée une *sécante* à la courbe de f . Une équation cartésienne de cette sécante est

$$(D_x) \quad y = f(a) + \Delta(x - a)$$

où

$$\Delta = \frac{f(x) - f(a)}{x - a}$$

Lorsque x tend vers a , la droite D_x pivote autour du point A , et tend (notion qu'il conviendrait de préciser) vers la tangente T . Ainsi, *la tangente est la limite des sécantes*.

1.3 Dérivées d'ordre supérieur

Définition 4. Soit $f : I \rightarrow \mathbb{R}$.

On définit par récurrence sur k , pour $k \in \mathbb{N}$, la dérivabilité à l'ordre k de f sur I , ainsi que la notation $f^{(k)}$.

- La fonction f est 0 fois dérivable sur I . On pose $f^{(0)} = f$.
- Pour tout entier $k \geq 1$, f est k fois dérivable sur I si et seulement si
 - ★ f est $k - 1$ fois dérivable sur I .
 - ★ $f^{(k-1)}$ est dérivable sur I .

On pose alors

$$f^{(k)} = \left(f^{(k-1)} \right)'$$

On dit que f est indéfiniment dérivable sur I lorsque f est k fois dérivable sur I , ceci pour tout $k \geq 1$.

Notation.

On utilise aussi les notations $\frac{d^k f}{dx^k}$ pour la dérivée d'ordre k de la fonction f .

Pour tout $k \in \mathbb{N} \cup \{+\infty\}$ on note $\mathcal{D}^k(I)$ l'ensemble des fonctions k fois dérivables sur I .

On note enfin $\mathcal{D}^\infty(I)$ l'ensemble des fonctions indéfiniment dérivables sur I . Remarquons que

$$\mathcal{D}^0(I) = \mathbb{R}^I$$

et

$$\mathcal{D}^\infty(I) = \bigcap_{k \geq 1} \mathcal{D}^k(I)$$

Remarque 1. On peut également définir par récurrence sur k la notion de fonction k fois dérivable en un point a : pour tout $k \geq 1$, la fonction f est k fois dérivable en a lorsque

- f est $k - 1$ fois dérivable au voisinage de a .
- $f^{(k-1)}$ est dérivable en a .

On pose alors $f^{(k)}(a) = \left(f^{(k-1)} \right)'(a)$.

1.4 Dérivées à droite et à gauche

Définition 5. Soit $f : I \rightarrow \mathbb{R}$. Soit $a \in I$.

- Si a n'est pas la borne inférieure de I , la dérivée à gauche de f en a , notée $f'_g(a)$, est la dérivée en a de la restriction de f à $I \cap]-\infty, a]$

- Si a n'est pas la borne supérieure de I , la dérivée à droite de f en a , notée $f'_d(a)$, est la dérivée en a de la restriction de f à $I \cap [a, +\infty[$.

On a le résultat suivant.

Proposition 3. Soit $f : I \rightarrow \mathbb{R}$. Soit a un point qui n'est pas une borne de I . f est dérivable en a si et seulement si f est dérivable à gauche et à droite en a , et que les dérivées sont égales. On a alors

$$f'(a) = f'_g(a) = f'_d(a)$$

Démonstration. Les taux d'accroissement ont une limite en a si et seulement si ils ont une limite à gauche et une limite à droite, et ces deux limites sont égales. $\square$

1.5 Classes de fonctions

Définition 6. Soit $f : I \rightarrow \mathbb{R}$.

On dit que f est de classe $\mathcal{C}^0$ sur I lorsque f est continue sur I .

Soit $k \in \mathbb{N}^*$. On dit que f est de classe $\mathcal{C}^k$ sur I lorsque f est k fois dérivable sur I , et que $f^{(k)}$ est continue sur I .

On dit enfin que f est de classe $\mathcal{C}^\infty$ sur I lorsque f est de classe $\mathcal{C}^k$ sur I pour tout entier k .

Notation. Pour tout $k \in \mathbb{N} \cup \{+\infty\}$, on note $\mathcal{C}^k(I)$ l'ensemble des fonctions de classe $\mathcal{C}^k$ sur I . Remarquons que

$$\mathcal{C}^\infty(I) = \bigcap_{k \geq 0} \mathcal{C}^k(I)$$

Proposition 4. On a les inclusions suivantes :

$$\mathbb{R}^I = \mathcal{D}^0(I) \supset \mathcal{C}^0(I) \supset \mathcal{D}^1(I) \supset \mathcal{C}^1(I) \supset \mathcal{D}^2(I) \supset \mathcal{C}^2(I) \supset \dots \supset \mathcal{C}^\infty(I) = \mathcal{D}^\infty(I)$$

Démonstration. Soit $k \in \mathbb{N}$.

L'inclusion $\mathcal{C}^k(I) \subset \mathcal{D}^k(I)$ est évidente.

Supposons que $f \in \mathcal{D}^{k+1}(I)$. La fonction f est donc k fois dérivable sur I , et $f^{(k)}$ est dérivable, donc continue, sur I . De là, $f \in \mathcal{C}^k(I)$. Ainsi,

$$\mathcal{D}^{k+1}(I) \subset \mathcal{C}^k(I)$$

L'égalité $\mathcal{D}^\infty(I) = \mathcal{C}^\infty(I)$ est laissée en exercice. $\square$

Remarque 2. Ces inclusions sont en réalité toutes strictes. Nous l'admettrons dans le cas général. Montrons cependant que $\mathcal{C}^1(\mathbb{R}) \neq \mathcal{D}^1(\mathbb{R})$. Pour cela, considérons $f : \mathbb{R} \rightarrow \mathbb{R}$ définie par $f(0) = 0$ et, pour tout réel $x \neq 0$,

$$f(x) = x^2 \sin \frac{1}{x}$$

Par les théorèmes sur les opérations sur les fonctions dérivables que nous allons bientôt voir, f est dérivable en tout réel non nul. De plus, pour tout $x \neq 0$,

$$\left| \frac{f(x) - f(0)}{x - 0} \right| = \left| x \sin \frac{1}{x} \right| \leq |x|$$

Ainsi, les taux d'accroissements de f en 0 tendent vers 0 lorsque x tend vers 0. On en déduit que f est dérivable en 0, et que $f'(0) = 0$. En conclusion, $f \in \mathcal{D}^1(\mathbb{R})$.

Remarquons maintenant, toujours grâce aux futurs théorèmes sur les opérations sur les dérivées, que pour tout $x \neq 0$,

$$f'(x) = 2x \sin \frac{1}{x} - \cos \frac{1}{x}$$

Cette quantité n'a pas de limite lorsque x tend vers 0. Ainsi, f' n'est pas continue en 0, et donc $f \notin \mathcal{C}^1(\mathbb{R})$.

Remarque 3. Pour tout $k \geq 1$, si $f \in \mathcal{C}^k(I)$, alors $f' \in \mathcal{C}^{k-1}(I)$. De même pour les ensembles $\mathcal{D}^k(I)$.

1.6 Opérations sur les dérivées

Dérivées d'ordre 1

Proposition 5. Soient $f, g : I \rightarrow \mathbb{R}$. Soit $a \in I$. Soit $\lambda \in \mathbb{R}$. On suppose f et g dérivables en a . Alors

- $f + g$ est dérivable en a et

$$(f + g)'(a) = f'(a) + g'(a)$$

- λf est dérivable en a et

$$(\lambda f)'(a) = \lambda f'(a)$$

- fg est dérivable en a et

$$(fg)'(a) = f'(a)g(a) + f(a)g'(a)$$

- Si $g(a) \neq 0$, alors g ne s'annule pas au voisinage de a , $\frac{f}{g}$ est dérivable en a et

$$\left(\frac{f}{g}\right)'(a) = \frac{f'(a)g(a) - f(a)g'(a)}{g(a)^2}$$

Démonstration. Faisons la preuve pour le produit. On a pour tout $x \neq a$,

$$\frac{(fg)(x) - (fg)(a)}{x - a} = f(x)\frac{g(x) - g(a)}{x - a} + g(a)\frac{f(x) - f(a)}{x - a}$$

On passe ensuite à la limite en remarquant que, comme f est dérivable en a , f y est aussi continue. $\square$

Proposition 6. Soit $g : I \rightarrow \mathbb{R}$. Soit $a \in I$. On suppose g dérivable en a et $g(a) \neq 0$. Alors

- g ne s'annule pas au voisinage de a .
- $\frac{1}{g}$ est dérivable en a et

$$\left(\frac{1}{g}\right)'(a) = \frac{-g'(a)}{g(a)^2}$$

Démonstration. $g(a) \neq 0$ et g est continue en a , donc g est non nulle au voisinage de a . On a pour tout $x \neq a$,

$$\frac{\frac{1}{g}(x) - \frac{1}{g}(a)}{x - a} = \frac{-(g(x) - g(a))}{(x - a)g(x)g(a)}$$

On passe ensuite à la limite. $\square$

Corollaire 7. Soient $f, g : I \rightarrow \mathbb{R}$. Soit $a \in I$. On suppose f et g dérivables en a et $g(a) \neq 0$. Alors

- g ne s'annule pas au voisinage de a .
- $\frac{f}{g}$ est dérivable en a et

$$\left(\frac{f}{g}\right)'(a) = \frac{f'(a)g(a) - f(a)g'(a)}{g(a)^2}$$

Démonstration. Il suffit de remarquer que

$$\frac{f}{g} = f \times \frac{1}{g}$$

□

Proposition 8. Soit J un intervalle de $\mathbb{R}$. Soit $f : I \longrightarrow J$. Soit $g : J \longrightarrow \mathbb{R}$. Soit $a \in I$. Si f est dérivable en a et g est dérivable en $f(a)$, alors $g \circ f$ est dérivable en a , et on a :

$$(g \circ f)'(a) = g'(f(a)) \times f'(a)$$

Démonstration.

Notons $b = f(a)$. on a pour tout y voisin de b ,

$$g(y) = g(b) + (y - b)g'(b) + (y - b)\varepsilon(y)$$

où $\varepsilon(y) \longrightarrow 0$ lorsque y tend vers b . Lorsque x tend vers a , $f(x)$ tend vers b donc $\varepsilon'(x) = \varepsilon(f(x))$ tend vers 0. Prenons $y = f(x)$. Il vient

$$\begin{aligned} g(f(x)) &= g(f(a)) + (f(x) - f(a))g'(b) + (f(x) - f(a))\varepsilon'(x) \\ &= g(f(a))((x - a)f'(a) + (x - a)\varepsilon''(x))g'(f(a)) \\ &\quad + ((x - a)f'(a) + (x - a)\varepsilon''(x))\varepsilon'(x) \end{aligned}$$

où $\varepsilon''(x) \longrightarrow 0$ lorsque x tend vers a . On regroupe les termes, et en particulier tout ce qui tend vers 0 :

$$g(f(x)) = g(f(a)) + (x - a)f'(a)g'(f(a)) + (x - a)\varepsilon'''(x)$$

où $\varepsilon'''(x)$ tend vers 0 lorsque x tend vers a . La fonction $g \circ f$ est donc dérivable en a , et sa dérivée est $f'(a)g'(f(a))$. □

Proposition 9. Soit $f : I \longrightarrow J$ une bijection continue (et donc strictement monotone). Soit $a \in I$. Soit $b = f(a)$. On suppose que f est dérivable en a . Si $f'(a) \neq 0$, alors f^{-1} est dérivable en b et

$$(f^{-1})'(b) = \frac{1}{f'(a)}$$

Si $f'(a) = 0$, alors f^{-1} n'est pas dérivable en b . Plus précisément, la courbe représentative de f^{-1} possède au point b une tangente verticale.

Démonstration. Supposons par exemple f strictement croissante. Posons $g = f^{-1}$. Pour tout $y \in J \setminus \{b\}$, posons

$$\Delta(y) = \frac{g(y) - g(b)}{y - b}$$

Posons $y = f(x)$, de sorte que $x = g(y)$. Il vient

$$\Delta(y) = \frac{x - a}{f(x) - f(a)}$$

Lorsque y tend vers b , $x = g(y)$ tend vers $g(b) = a$, car g est continue en b .

Si $f'(a) \neq 0$, $\Delta(y)$ tend vers $\frac{1}{f'(a)}$.

Si $f'(a) = 0$, comme f est strictement croissante, $\Delta(y) > 0$. Ainsi, $\Delta(y)$ tend vers $+\infty$ lorsque y tend vers b . $\square$

Remarque 4. La formule qui nous donne la dérivée s'écrit aussi

$$(f^{-1})'(b) = \frac{1}{f'(f^{-1}(b))}$$

C'est sous cette forme qu'il convient de la retenir.

Exemple. Soit $f : [-\frac{\pi}{2}, \frac{\pi}{2}] \rightarrow \mathbb{R}$ définie par $f(x) = \sin x$. Nous avons déjà vu que f est une bijection continue strictement croissante. Sa réciproque est la fonction arc sinus, notée $\arcsin$. f est dérivable sur $[-\frac{\pi}{2}, \frac{\pi}{2}]$. Sa dérivée est la fonction $\cos$, qui s'annule uniquement en $\pm\frac{\pi}{2}$. Donc, $\arcsin$ est dérivable en tout point de $[-1, 1]$ différent de $\sin(\pm\frac{\pi}{2})$, c'est à dire sur $] -1, 1[$ et la courbe représentative de $\arcsin$ admet des tangentes verticales en -1 et 1 . On a pour tout $x \in] -1, 1[$,

$$\arcsin' x = \frac{1}{f'(\arcsin x)} = \frac{1}{\cos(\arcsin x)} = \frac{1}{\sqrt{1 - \sin^2(\arcsin x)}}$$

Donc

$$\arcsin' x = \frac{1}{\sqrt{1 - x^2}}$$

Dérivées d'ordre supérieur

Proposition 10. Soit $k \in \mathbb{N}$. Soient $f, g \in \mathcal{D}^k(I)$. Soit $\lambda \in \mathbb{R}$. Alors

- $f + g \in \mathcal{D}^k(I)$ et

$$(f + g)^{(k)} = f^{(k)} + g^{(k)}$$

- $\lambda f \in \mathcal{D}^k(I)$ et

$$(\lambda f)^{(k)} = \lambda f^{(k)}$$

- $fg \in \mathcal{D}^k(I)$ et on a la formule de Leibniz :

$$(fg)^{(k)} = \sum_{i=0}^k \binom{k}{i} f^{(i)} g^{(k-i)}$$

Démonstration. La somme et le produit par un réel se montrent par une récurrence facile sur k . Le produit est, quant à lui, une récurrence moins triviale.

Pour $k = 0$, c'est évident. Soit $k \in \mathbb{N}$. Supposons la propriété vraie pour k . Soient $f, g \in \mathcal{D}^{k+1}(I)$. Elles sont donc a fortiori k fois dérivables sur I . D'après l'hypothèse de récurrence, fg est k fois dérivable sur I et on a

$$(fg)^{(k)} = \sum_{i=0}^k \binom{k}{i} f^{(i)} g^{(k-i)}$$

Toutes les fonctions mises en jeu dans le membre de droite sont dérivables sur I . $(fg)^{(k)}$ l'est donc aussi et on peut dériver l'égalité.

$$(fg)^{(k+1)} = \sum_{i=0}^k \binom{k}{i} \left(f^{(i+1)} g^{(k-i)} + f^{(i)} g^{(k+1-i)} \right)$$

Un calcul analogue à celui qui a été fait pour prouver la formule du binôme mène à la formule de Leibniz. $\square$

Remarque 5. On peut, dans la proposition précédente, remplacer l'hypothèse de dérivabilité sur I par celle de dérivabilité en un point $a \in I$.

Corollaire 11. Soit $k \in \mathbb{N} \cup \{+\infty\}$. Soient $f, g \in \mathcal{C}^k(I)$. Soit $\lambda \in \mathbb{R}$. Alors

- $f + g \in \mathcal{C}^k(I)$.
- $\lambda f \in \mathcal{C}^k(I)$.
- $fg \in \mathcal{C}^k(I)$.

De même en remplaçant $\mathcal{C}^k$ par $\mathcal{D}^k$.

Démonstration. L'essentiel a déjà été prouvé. On laisse au lecteur en exercice le soin de compléter la démonstration. $\square$

Remarque 6. Soit $k \in \mathbb{N} \cup \{+\infty\}$. Les ensembles $\mathcal{C}^k(I)$ et $\mathcal{D}^k(I)$ sont des algèbres. L'application

$$D : f \mapsto f'$$

de dérivation, qui, pour $k \geq 1$, envoie $\mathcal{C}^k(I)$ sur $\mathcal{C}^{k-1}(I)$ (en posant $+\infty - 1 = +\infty$), est une application linéaire mais n'est pas un morphisme d'anneaux.

Proposition 12. Soit $k \in \mathbb{N} \cup \{+\infty\}$. Soient $f : I \longrightarrow J$ et $g : J \longrightarrow \mathbb{R}$ deux fonctions de classe $\mathcal{C}^k$. Alors, $g \circ f$ est de classe $\mathcal{C}^k$.

Démonstration. Si c'est vrai pour tout entier k , c'est clairement vrai pour $k = \infty$. Faisons donc une récurrence sur k . Pour $k = 0$, c'est connu : la composée de deux fonctions continues est continue. Soit $k \in \mathbb{N}$. Supposons que ce soit vrai pour k . Soit f, g de classe $\mathcal{C}^{k+1}$. Elles sont donc a fortiori dérivables donc $g \circ f$ est dérivable sur I . Regardons l'égalité

$$(g \circ f)' = f' \times g' \circ f$$

Les fonctions f' , f et g' sont $\mathcal{C}^k$, donc (hypothèse de récurrence et Leibniz) $(g \circ f)'$ est $\mathcal{C}^k$. Et donc, $g \circ f \in \mathcal{C}^{k+1}(I)$. $\square$

Remarque 7. Pas de formule donnée pour la dérivée k ème d'une composée. Cette formule existe, c'est la *formule de Faa di Bruno*. La voici :

$$(g \circ f)^{(n)}(x) = \sum_{m_1+2m_2+\dots+nm_n=n} \alpha_{m_1,\dots,m_n} g^{(m_1+\dots+m_n)}(f(x)) \prod_{j=1}^n (f^{(j)}(x))^{m_j}$$

où

$$\alpha_{m_1,\dots,m_n} = \frac{n!}{m_1!1!^{m_1} m_2!2!^{m_2} \dots m_n!n!^{m_n}}$$

Proposition 13. Soit $k \geq 0$. Soit $g \in \mathcal{C}^k(I)$ une fonction qui ne s'annule pas. Alors $\frac{1}{g} \in \mathcal{C}^k(I)$.

Démonstration. Laissez en exercice. S'inspirer de la démonstration précédente. $\square$

Proposition 14. Soit $k \geq 1$. Soit $f : I \longrightarrow J \in \mathcal{C}^k(I)$ une bijection dont la dérivée ne s'annule pas. Alors $f^{-1} \in \mathcal{C}^k(J)$.

Démonstration. S'inspirer des preuves précédentes, en écrivant que f^{-1} est dérivable sur J et que $(f^{-1})' = \frac{1}{f' \circ f^{-1}}$. $\square$

2 Accroissements finis

2.1 Extrema locaux

Proposition 15. Soit I un intervalle de $\mathbb{R}$. Soit $f : I \rightarrow \mathbb{R}$. Soit a un point de I qui n'est pas une extrémité de I . Si f admet un extremum local en a , alors $f'(a) = 0$. La réciproque est fausse.

Démonstration. Supposons par exemple que f admet un maximum local en a . Pour tout $x < a$ voisin de a , on a

$$\frac{f(x) - f(a)}{x - a} \geq 0$$

car numérateur et dénominateur sont négatifs. On passe à la limite : $f'(a) = f'_g(a) \geq 0$. En faisant de même à droite de a , on obtient $f'(a) = f'_d(a) \leq 0$. $\square$

2.2 Théorème de Rolle

Théorème 16. THÉORÈME DE ROLLE

Soient a et b deux réels tels que $a < b$. Soit $f : [a, b] \rightarrow \mathbb{R}$. On suppose :

- f continue sur $[a, b]$.
- f dérivable sur $]a, b[$.
- $f(a) = f(b)$.

Il existe alors un réel $c \in]a, b[$ tel que $f'(c) = 0$.

Démonstration. Si f est constante, c'est évident car f' s'annule en tout point.

Supposons maintenant f non constante. Par exemple, f prend une valeur strictement inférieure à $f(a)$ (et donc à $f(b)$). La fonction f est continue sur le segment $[a, b]$. Elle y possède donc un minimum global (et donc local) en un point c qui, selon les hypothèses, n'est ni a ni b . On a donc $f'(c) = 0$. $\square$

2.3 Accroissements finis

Théorème 17. ÉGALITÉ DES ACCROISSEMENTS FINIS.

Soient $a < b$ deux réels. Soit $f : [a, b] \rightarrow \mathbb{R}$. On suppose :

- f continue sur $[a, b]$.
- f dérivable sur $]a, b[$.

Il existe alors $c \in]a, b[$ tel que

$$f(b) - f(a) = (b - a)f'(c)$$

Démonstration. On considère la fonction $g : [a, b] \longrightarrow \mathbb{R}$ définie par

$$g(x) = f(x) - f(a) - A(x - a)$$

où

$$A = \frac{f(b) - f(a)}{b - a}$$

La fonction g est continue sur $[a, b]$, dérivable sur $]a, b[$, et $g(a) = g(b) = 0$. D'après le théorème de Rolle, il existe $c \in [a, b]$ tel que

$$g'(c) = f'(c) - A = 0$$

□

Corollaire 18. INÉGALITÉ DES ACCROISSEMENTS FINIS

Soient $a < b$ deux réels. Soit $f : [a, b] \longrightarrow \mathbb{R}$. On suppose :

- f continue sur $[a, b]$.
- f dérivable sur $]a, b[$.
- Il existe deux réels m et M tels que $m \leq f' \leq M$ sur $]a, b[$.

Alors,

$$m(b - a) \leq f(b) - f(a) \leq M(b - a)$$

Démonstration. Majorer et minorer dans l'égalité des accroissements finis. □

Notation. Étant donné un intervalle I de $\mathbb{R}$, on note $\overset{o}{I}$ l'intervalle I privé de ses bornes (l'intérieur de I).

Corollaire 19. INÉGALITÉ DES ACCROISSEMENTS FINIS (BIS)

Soit $f : I \longrightarrow \mathbb{R}$. On suppose :

- f continue sur I .
- f dérivable sur $\overset{o}{I}$.
- Il existe un réel M tel que $|f'| \leq M$ sur $\overset{o}{I}$.

Alors,

$$\forall x, y \in I, |f(y) - f(x)| \leq M|y - x|$$

En d'autres termes, une fonction dérivable et dont la dérivée est bornée est lipschitzienne.

Démonstration. Soient $x, y \in I$. Supposons par exemple $x < y$. La fonction f vérifie toutes les hypothèses du théorème des accroissements finis sur le segment $[x, y]$. Il existe

donc $c \in]x, y[\subset \overset{o}{I}$ tel que

$$f(y) - f(x) = f'(c)(y - x)$$

Donc,

$$|f(y) - f(x)| \leq M(y - x)$$

Si $x > y$, on applique ce qui précède en intervertissant x et y . Si $x = y$, l'inégalité à prouver est évidente. $\square$

2.4 Fonctions monotones

Théorème 20. Soit $f : I \rightarrow \mathbb{R}$, continue sur I et dérivable sur $\overset{o}{I}$. Alors,

- f est croissante sur I si et seulement si $f' \geq 0$ sur $\overset{o}{I}$.
- f est décroissante sur I si et seulement si $f' \leq 0$ sur $\overset{o}{I}$.
- f est constante sur I si et seulement si $f' = 0$ sur $\overset{o}{I}$.

Démonstration. Supposons f croissante sur I . Soit $a \in \overset{o}{I}$. Pour tout $x \neq a$, on a

$$\frac{f(x) - f(a)}{x - a} \geq 0$$

On passe à la limite : $f'(a) \geq 0$. Inversement, supposons $f' \geq 0$ sur $\overset{o}{I}$. Soient $x, y \in I$ tels que $x < y$. On applique le théorème des accroissements finis sur le segment $[x, y]$: il existe $c \in [x, y] \subset \overset{o}{I}$ tel que

$$f(y) - f(x) = (y - x)f'(c)$$

Le membre de droite est positif, donc $f(x) \leq f(y)$ et f est effectivement croissante.

Démonstration bien sûr identique pour f décroissante. Enfin, f est constante si et seulement si elle est à la fois croissante et décroissante. $\square$

Proposition 21. Soit $f : I \rightarrow \mathbb{R}$, continue sur I , dérivable sur $\overset{o}{I}$, et croissante. Soit

$$\mathcal{E} = \{x \in \overset{o}{I}, f'(x) = 0\}$$

Alors f est strictement croissante sur I si et seulement si $\mathcal{E}$ ne contient aucun intervalle non trivial.

Démonstration. Une fonction croissante f n'est pas strictement croissante si et seulement si elle est constante sur un intervalle non trivial. $\square$

Remarque 8. Dans les cas les plus fréquents, l'ensemble $\mathcal{E}$ est un ensemble fini, ou alors $\mathcal{E}$ ne contient que des points « isolés ». Par exemple, $\mathcal{E} = \mathbb{Z}$.

2.5 Passage à la limite dans une dérivée

Soient $a < b$. Soit $f : [a, b] \rightarrow \mathbb{R}$ une fonction continue sur $[a, b]$ et dérivable sur $]a, b[$. On se demande si f est dérivable en a . Une méthode consiste à étudier des taux d'accroissement. Il est cependant logique de se demander ce qui se passe pour $f'(x)$ lorsque x tend vers a .

Proposition 22.

- Si $f'(x)$ a une limite réelle ℓ lorsque x tend vers a , $x \neq a$, alors f est dérivable en a et $f'(a) = \ell$.
- Si $f'(x)$ tend vers l'infini lorsque x tend vers a , $x \neq a$, alors f n'est pas dérivable en a : plus précisément, la courbe représentative de f admet une tangente verticale en a .
- Si f' n'a pas de limite en a , alors on ne peut rien dire.

Démonstration. Pour tout $x \in I$ tel que $x > a$, appliquons le théorème des accroissements finis à f sur le segment $[a, x]$: il existe $a < c_x < x$ tel que

$$\frac{f(x) - f(a)}{x - a} = f'(c_x)$$

Quand x tend vers a , c_x aussi. Si l'on suppose que $f'(c_x)$ tend vers ℓ ou ∞ , le taux d'accroissement fait de même par le théorème de composition des limites.

Pour le dernier point, reprenons un exemple déjà vu plus haut. Considérons la fonction $f : \mathbb{R} \rightarrow \mathbb{R}$ définie par $f(0) = 0$ et, pour tout $x \neq 0$,

$$f(x) = x^2 \sin \frac{1}{x}$$

La fonction f est continue sur $\mathbb{R}$ et, par les théorèmes sur les opérations sur les dérivées, de classe $\mathcal{C}^\infty$ sur $\mathbb{R}_+^*$ et $\mathbb{R}_-^*$. De plus,

$$\frac{f(x) - f(0)}{x - 0} = x \sin \frac{1}{x}$$

tend vers 0 lorsque x tend vers 0. Donc, f est dérivable en 0 et $f'(0)$. Enfin, on a pour tout x non nul

$$f'(x) = 2x \sin \frac{1}{x} - \cos \frac{1}{x}$$

Le premier terme tend vers 0 lorsque x tend vers 0, mais le second terme n'a pas de limite en 0. La fonction f' n'a donc pas de limite en 0. $\square$

Remarque 9. Comment interpréter le contre-exemple ci-dessus ? Eh bien, quand une fonction est définie en un point mais qu'elle n'a pas de limite en ce point, cela signifie tout simplement... qu'elle n'y est pas continue. Bref, la fonction f' ci-dessus n'est pas continue en 0. Dit autrement, f est dérivable, mais n'est pas de classe $\mathcal{C}^1$.

2.6 Passage à la limite pour les fonctions de classe $\mathcal{C}^k$

Proposition 23. Soit $k \in \mathbb{N}$. Soient $a, b \in \mathbb{R}$ tels que $a < b$. Soit $f :]a, b] \rightarrow \mathbb{R}$ une fonction de classe $\mathcal{C}^k$. On suppose que pour tout $i \in \llbracket 0, k \rrbracket$, $f^{(i)}(x)$ a une limite finie lorsque x tend vers a . Alors f est prolongeable en une fonction de classe $\mathcal{C}^k$ sur $[a, b]$.

Démonstration. La fonction f a une limite finie en a , elle est donc prolongeable en une fonction continue sur $[a, b]$ que nous continuons à appeler f . Nous allons montrer par récurrence sur i que pour tout $i \in \llbracket 0, k \rrbracket$ f est de classe $\mathcal{C}^i$ sur $[a, b]$. Pour $i = 0$, c'est clair puisque f est continue. Soit $i < k$. Supposons la propriété vraie pour i . La fonction $f^{(i)}$ est continue sur $[a, b]$, dérivable sur $]a, b]$, et, par hypothèse, sa dérivée a une limite finie en a . Grâce au théorème de passage à la limite dans une dérivée, on conclut que $f^{(i)}$ est de classe $\mathcal{C}^1$ sur $[a, b]$, donc que f est de classe $\mathcal{C}^{i+1}$ sur $[a, b]$. $\square$

Corollaire 24. Soient $a, b \in \mathbb{R}$ tels que $a < b$. Soit $f :]a, b] \rightarrow \mathbb{R}$ une fonction de classe $\mathcal{C}^\infty$. On suppose que pour tout $i \in \mathbb{N}$, $f^{(i)}(x)$ a une limite finie lorsque x tend vers a . Alors f est prolongeable en une fonction de classe $\mathcal{C}^\infty$ sur $[a, b]$.

Démonstration. La fonction f est prolongeable par continuité à $[a, b]$ et d'après le théorème précédent son prolongement est de classe $\mathcal{C}^k$ pour tout entier naturel k . $\square$

3 Fonctions à valeurs complexes

3.1 Dérivation

Une dérivée est une limite. La notion de dérivée s'étend donc aux fonctions à valeurs complexes. Les opérations sur les dérivées, la formule de Leibniz, s'étendent sans difficulté. Signalons quand même que pour dériver $g \circ f$, la fonction f doit être à valeurs réelles. g , en revanche, peut être à valeurs complexes.

Signalons le résultat suivant. Sa preuve est laissée en exercice.

Proposition 25. Soit $f : I \rightarrow \mathbb{C}$. La fonction f est dérivable au point $a \in I$ si et seulement si $\operatorname{Re} f$ et $\operatorname{Im} f$ le sont. On a alors

$$f'(a) = (\operatorname{Re} f)'(a) + i (\operatorname{Im} f)'(a)$$

On en déduit facilement la propriété suivante :

Proposition 26. Soit $\alpha \in \mathbb{C}$. La fonction $f : x \mapsto e^{\alpha x}$ est de classe $\mathcal{C}^\infty$ sur $\mathbb{R}$ et on a pour tout réel x ,

$$f'(x) = \alpha e^{\alpha x}$$

3.2 Rolle et accroissements finis

Attention ! Le théorème de Rolle et le théorème des accroissements finis ne sont pas valables pour des fonctions à valeurs complexes. Par exemple, si l'on prend

$$f(x) = x(1-x) + ix^2(1-x)$$

alors f est de classe $\mathcal{C}^\infty$ sur $[0, 1]$, $f(0) = f(1) = 0$, et pourtant

$$f'(x) = (1-2x) + ix(2-3x)$$

ne s'annule pas.

En revanche, l'inégalité des accroissements finis reste vraie :

Théorème 27. Soient $a, b \in \mathbb{R}$ tels que $a < b$. Soit $f : [a, b] \rightarrow \mathbb{C}$ continue sur $[a, b]$, dérivable sur $]a, b[$. On suppose qu'il existe un réel M tel que

$$\forall x \in]a, b[, |f'(x)| \leq M$$

Alors, pour tous $x, y \in [a, b]$,

$$|f(y) - f(x)| \leq M|y - x|$$

Démonstration. Contentons-nous d'une démonstration lorsque $f \in \mathcal{C}^1([a, b])$ (le cas général est plus délicat). Prenons par exemple $x < y$. On utilise un résultat d'intégration démontré un peu plus bas :

$$|f(y) - f(x)| = \left| \int_x^y f'(t) dt \right| \leq \int_x^y |f'(t)| dt \leq \max_{[x, y]} |f'| \int_x^y dt \leq M(y - x)$$

d'où le résultat. $\square$

3.3 Intégration

Cette section est un prématurée et nécessite des connaissances en algèbre linéaire et en intégration. On y démontre un résultat qui est utilisé ci-dessus.

Soit $f : [a, b] \rightarrow \mathbb{C}$ une fonction continue par morceaux. On appelle intégrale de f sur $[a, b]$ le nombre complexe

$$\int_a^b f = \int_a^b \operatorname{Re} f + i \int_a^b \operatorname{Im} f$$

La linéarité de l'intégrale reste vérifiée. La formule de Chasles également. On peut aussi intégrer par parties ou faire des changements de variable.

En revanche, on ne peut plus parler de croissance ou de positivité de l'intégrale, puisqu'on ne peut pas comparer deux fonctions à valeurs complexes. On a cependant le résultat suivant :

Proposition 28. Soient $a, b \in \mathbb{R}$ tels que $a < b$. Soit $f : [a, b] \rightarrow \mathbb{C}$ continue par morceaux. Alors

$$\left| \int_a^b f(x) dx \right| \leq \int_a^b |f(x)| dx$$

Démonstration. Les barres verticales sont bien entendu des modules de nombres complexes. Si $\int_a^b f = 0$ c'est évident. Supposons donc $\int_a^b f \neq 0$.

Pour tout réel θ , on a

$$\operatorname{Re} \left(e^{-i\theta} \int_a^b f \right) = \int_a^b \operatorname{Re} (e^{-i\theta} f) \leq \int_a^b |e^{-i\theta} f| = \int_a^b |f|$$

Prenons pour θ un argument de $\int_a^b f$, qui est, rappelons-le, non nulle. On a donc, par définition de l'argument d'un nombre complexe,

$$\int_a^b f = \left| \int_a^b f \right| e^{i\theta}$$

où la première égalité résulte de la linéarité de l'intégrale. On a alors

$$\operatorname{Re} \left(e^{-i\theta} \int_a^b f \right) = \operatorname{Re} \left(\left| \int_a^b f \right| \right) = \left| \int_a^b f \right|$$

d'où l'inégalité voulue. $\square$

Chapitre 13

Fonctions Usuelles

1 Logarithmes et exponentielles

1.1 Logarithme népérien

Proposition 1. Soit $\varphi : \mathbb{R}_+^* \rightarrow \mathbb{R}$ définie par $\varphi(x) = \frac{1}{x}$. Il existe une unique fonction dérivable $F : \mathbb{R}_+^* \rightarrow \mathbb{R}$ telle que $F' = \varphi$ et $F(1) = 0$.

Démonstration. Nous admettons pour l'instant l'existence d'une fonction dérivable $F_0 : \mathbb{R}_+^* \rightarrow \mathbb{R}$ telle que $F_0' = \varphi$. Ce résultat sera démontré dans le chapitre d'intégration.

Soit $F \in \mathcal{D}^1(\mathbb{R}_+^*)$ telle que $F' = \varphi$ et $F(1) = 0$. On a alors $F' = F_0'$, ou encore $(F - F_0)' = 0$. De là, $F - F_0$ est constante sur $\mathbb{R}_+^*$. Il existe ainsi un réel c tel que

$$F = F_0 + c$$

De plus, $F(1) = 0 = F_0(1) + c$, donc $c = -F_0(1)$. On a donc l'unicité de F . In versement, soit $F = F_0 - F_0(1)$. La fonction F est dérivable sur $\mathbb{R}_+^*$, $F' = F_0' = \varphi$, et $F(1) = F_0(1) - F_0(1) = 0$, d'où l'existence. $\square$

Définition 1. La fonction F est appelée le logarithme népérien. On la note $\ln$.

Proposition 2. La fonction $\ln$ est de classe $\mathcal{C}^\infty$ sur $\mathbb{R}_+^*$.

Démonstration. En effet, $\ln' = \varphi \in \mathcal{C}^\infty(\mathbb{R}_+^*)$. $\square$

Proposition 3. Pour tous $x, y \in \mathbb{R}_+^*$,

$$\ln(xy) = \ln x + \ln y$$

Démonstration. Soit $y > 0$. La fonction $f_y : x \mapsto \ln(xy)$ est dérivable sur $\mathbb{R}_+^*$ et sa dérivée vaut, pour tout $x > 0$, $f_y'(x) = \frac{1}{x}$. La fonction $x \mapsto f_y(x) - \ln x$ est donc constante. Il existe ainsi $c_y \in \mathbb{R}$ tel que, pour tout $x > 0$,

$$f_y(x) = \ln x + c_y$$

Il reste à remarquer que

$$f_y(1) = \ln y = \ln 1 + c_y = c_y$$

Ainsi, $c_y = \ln y$. $\square$

Remarque 1. On en déduit pour tous $x, y > 0$ et tout $n \in \mathbb{Z}$,

- $\ln \frac{1}{x} = -\ln x$
- $\ln \frac{x}{y} = \ln x - \ln y$
- $\ln x^n = n \ln x$

Proposition 4. La fonction logarithme est une bijection strictement croissante de $]0, +\infty[$ sur $\mathbb{R}$, vérifiant

$$\lim_{x \rightarrow +\infty} \ln x = +\infty \text{ et } \lim_{x \rightarrow 0} \ln x = -\infty$$

Démonstration. La dérivée de $\ln$ est strictement positive, d'où la croissance stricte. La fonction $\ln$ a donc une limite α en 0 et une limite β en $+\infty$ et, par le théorème sur les fonctions continues strictement monotones, $\ln$ est une bijection de $]0, +\infty[$ sur $]\alpha, \beta[$.

Comme $2 > 1$, $\ln 2 > \ln 1 = 0$. Donc,

$$\ln(2^n) = n \ln 2 \longrightarrow +\infty$$

lorsque $n \longrightarrow +\infty$. Donc, la fonction logarithme est croissante, pas majorée, et tend ainsi vers $\beta = +\infty$ en $+\infty$.

Enfin $\ln x = -\ln \frac{1}{x} \longrightarrow \alpha = -\infty$ lorsque $x \rightarrow 0$. $\square$

Définition 2. On appelle e l'unique réel strictement positif tel que $\ln e = 1$.

Remarque 2. Le logarithme népérien est ainsi un isomorphisme du groupe $(\mathbb{R}_+^*, \times)$ sur le groupe $(\mathbb{R}, +)$.

Proposition 5. On a

$$\frac{\ln x}{x} \xrightarrow{x \rightarrow +\infty} 0 \quad \text{et} \quad x \ln x \xrightarrow{x \rightarrow 0} 0$$

Démonstration. La fonction $f : x \mapsto \frac{\ln x}{x}$ est dérivable sur $\mathbb{R}_+^*$. On a pour tout $x > 0$,

$$f'(x) = \frac{1 - \ln x}{x^2}$$

Cette quantité est négative dès que $x \geq e$. Ainsi, f décroît sur $[e, +\infty[$. f a donc une limite ℓ en $+\infty$. De plus, $f \geq 0$ au voisinage de $+\infty$, donc $\ell \in \mathbb{R}_+$. Remarquons que pour tout $x > 0$,

$$f(x^2) = \frac{2}{x} f(x)$$

Par composition des limites, $f(x^2) \rightarrow \ell$ en $+\infty$. Par opérations sur les limites, le membre de droite tend vers $0 \times \ell = 0$. Ainsi, $\ell = 0$.

Remarquons enfin, en posant $t = \frac{1}{x}$, que

$$t \ln t = -\frac{\ln x}{x}$$

Quand $t \rightarrow 0$, $x \rightarrow +\infty$. Par composition des limites, $t \ln t$ tend vers 0 lorsque t tend vers 0. $\square$

Remarque 3. On déduit de ce qui précède la *courbe* de la fonction $\ln$.

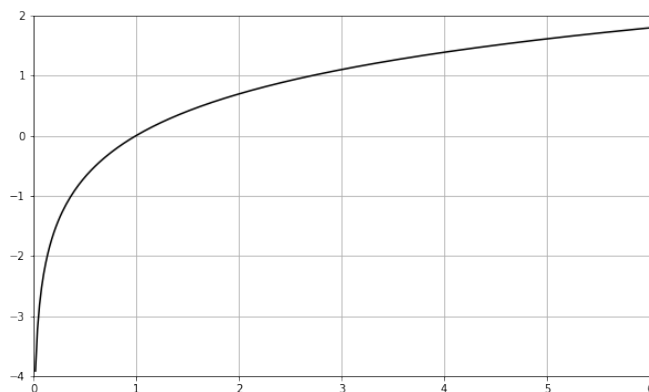


FIGURE 13.1 – Logarithme népérien

1.2 Exponentielle

Définition 3. La fonction exponentielle, notée $\exp$, est la fonction réciproque de la fonction logarithme népérien. C'est une bijection strictement croissante de $\mathbb{R}$ sur $\mathbb{R}_+^*$ vérifiant

$$\lim_{x \rightarrow +\infty} \exp x = +\infty \text{ et } \lim_{x \rightarrow -\infty} \exp x = 0$$

Elle est de classe $\mathcal{C}^\infty$ sur $\mathbb{R}$ et égale à sa dérivée.

Démonstration. On a $\exp = \ln^{-1}$. D'où la monotonie stricte, la continuité, les limites. De plus, $\ln$ est $\mathcal{C}^\infty$ et $\ln' > 0$ sur $\mathbb{R}_+^*$ donc $\exp \in \mathcal{C}^\infty(\mathbb{R})$. Enfin, pour tout x réel,

$$\exp'(x) = \frac{1}{\ln'(\exp x)} = \exp x$$

$\square$

Remarque 4. En tant que réciproque d'un isomorphisme, l'exponentielle est un isomorphisme du groupe $(\mathbb{R}, +)$ sur le groupe $(\mathbb{R}_+^*, \times)$.

Proposition 6. On a pour $x, y \in \mathbb{R}$ et $n \in \mathbb{Z}$:

- $\exp 0 = 1$
- $\exp(x + y) = \exp x \exp y$
- $\exp(x - y) = \frac{\exp x}{\exp y}$
- $\exp(nx) = (\exp x)^n$

Démonstration. L'exponentielle est un morphisme de groupes. $\square$

Remarque 5. Rappelons-nous que e est défini par $\ln e = 1$. On en déduit que $e = \exp 1$. De là, pour tout $n \in \mathbb{Z}$,

$$e^n = (\exp 1)^n = \exp n$$

Plus généralement, soit $x \in \mathbb{Q}$. Posons $x = \frac{p}{q}$ où $p \in \mathbb{Z}$ et $q \in \mathbb{N}^*$. On a $qx = p$, d'où

$$(\exp x)^q = \exp qx = \exp p = e^p$$

Lorsque nous aurons défini les puissances rationnelles, nous pourrons alors écrire

$$\exp x = (e^p)^{\frac{1}{q}} = e^{\frac{p}{q}} = e^x$$

Plus généralement, nous nous autoriserons à noter, pour tout $x \in \mathbb{R}$, $e^x = \exp x$.

1.3 Logarithmes et exponentielles en base quelconque

Définition 4. Soit $a > 0$, $a \neq 1$. On appelle logarithme de base a l'application $\log_a : \mathbb{R}_+^* \rightarrow \mathbb{R}$ définie pour tout $x > 0$ par

$$\log_a x = \frac{\ln x}{\ln a}$$

Exemple.

- Pour $a = 10$, on a le logarithme décimal, noté simplement $\log$.
- Pour $a = e$, on retrouve le logarithme népérien.
- Pour $a = 2$, on a le logarithme binaire, souvent noté $\lg$.

Les propriétés du logarithme de base a sont tout à fait similaires à celles du logarithme népérien : les fonctions $\log_a$ et $\ln$ ne diffèrent que de la constante multiplicative $\frac{1}{\ln a}$.

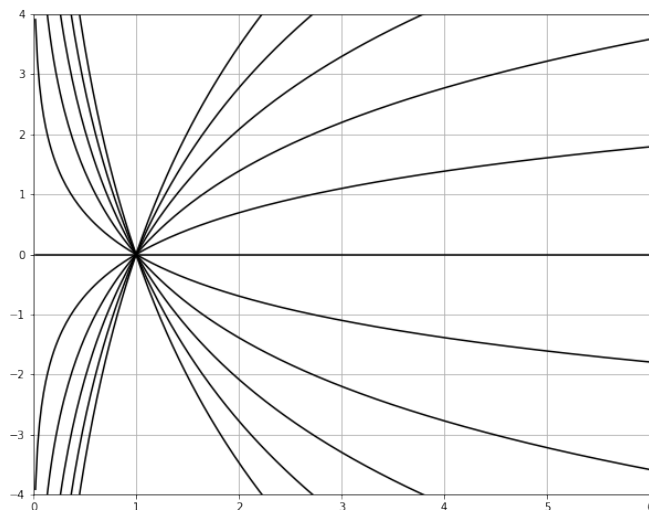


FIGURE 13.2 – Logarithmes

Définition 5. Soit $a > 0, a \neq 1$. La fonction logarithme de base a est une bijection de $\mathbb{R}_+^*$ sur $\mathbb{R}$. Sa réciproque est donc une bijection de $\mathbb{R}$ sur $\mathbb{R}_+^*$. On l'appelle exponentielle de base a , et on la note (provisoirement) $\exp_a$.

Les propriétés de l'exponentielle de base a sont identiques à celles de l'exponentielle.

Proposition 7. Soit $a > 0, a \neq 1$. On a pour tout réel x

$$\exp_a x = \exp(x \ln a)$$

Démonstration. On a $\log_a(\exp(x \ln a)) = \frac{\ln(\exp(x \ln a))}{\ln a} = x$.

Donc $\exp_a(x) = \exp(x \ln a)$. $\square$

Remarque 6. La fonction $\exp_a$, réciproque d'un isomorphisme, est elle aussi un isomorphisme. On a donc pour tout entier relatif n ,

$$\exp_a n = a^n$$

On étend la notation à tout réel x en posant

$$a^x = \exp(x \ln a)$$

Dorénavant, nous n'utiliserons plus la notation « $\exp_a$ ».

Remarque 7. À retenir, pour tout réel $a > 0$ (y compris $a = 1$) et tout réel x ,

$$a^x = e^{x \ln a}$$

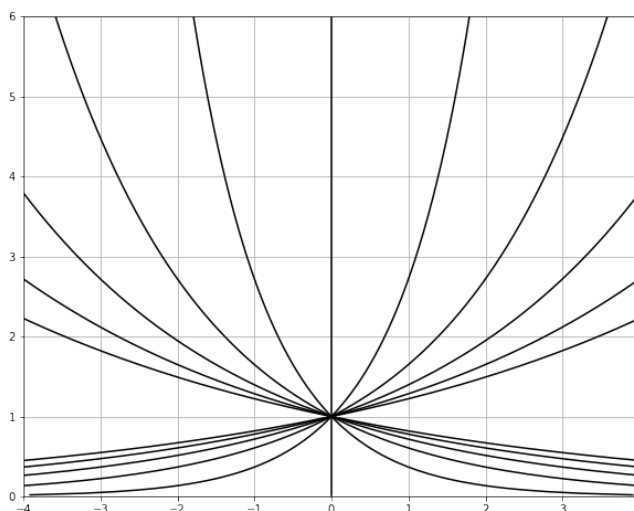


FIGURE 13.3 – Exponentielles

2 Puissances

2.1 Définition

Définition 6. Soit $a \in \mathbb{R}$. On appelle « fonction puissance de base a » la fonction $\varphi_a : \mathbb{R}_+^* \rightarrow \mathbb{R}$ définie pour tout $x \in \mathbb{R}_+^*$ par

$$\varphi_a(x) = x^a$$

Remarque 8. L'ensemble de définition des fonctions puissance est $\mathbb{R}_+^*$. Clairement, pour certaines valeurs de a , cet ensemble de définition peut être étendu :

- Si a est un entier naturel, $x^a = x \times x \times \dots \times x$ (a fois) et la fonction $x \mapsto x^a$ peut être définie sur $\mathbb{R}$.
- Si a est un entier négatif, cette même fonction peut être définie sur $\mathbb{R}^*$.

- Nous verrons un peu plus bas que si a est l'inverse d'un entier impair, la fonction peut aussi être définie sur $\mathbb{R}$.

Bref, pour un a quelconque, les fonctions puissances sont définies sur $\mathbb{R}_+^*$, mais leur ensemble de définition peut être plus gros pour certaines valeurs de a .

Proposition 8. Pour a, b réels et $x, y > 0$ on a :

$$\begin{aligned} x^a y^a &= (xy)^a & x^a x^b &= x^{a+b} & (x^a)^b &= x^{ab} \\ 1^a &= 1 & x^0 &= 1 & \ln x^a &= a \ln x \end{aligned}$$

Démonstration.

$$(xy)^a = \exp(a \ln(xy)) = \exp(a(\ln x + \ln y))$$

qui se développe en

$$\exp(a \ln x) \exp(a \ln y) = x^a y^a$$

Même démonstration pour toutes les égalités. $\square$

2.2 variations

Soit $a \in \mathbb{R}$. On a pour tout $x > 0$,

$$\varphi'_a(x) = ax^{a-1}$$

On en déduit les variations de φ_a , les limites en 0 et à l'infini, et la courbe représentative.

Remarque 9. Si $a > 0$, la fonction $x \mapsto x^a$ tend vers 0 lorsque $x \rightarrow 0$. Elle est donc prolongeable par continuité en 0 en posant $0^a = 0$. Attention, cependant : on n'a évidemment pas $0^a = \exp(a \ln 0)$! Quand ce prolongement est-il dérivable en 0 ? Formons des taux d'accroissement :

$$\frac{x^a - 0}{x - 0} = x^{a-1}$$

a une limite en 0 si et seulement si $a \geq 1$. Pour $a = 1$, la dérivée en 0 vaut 1. Pour $a > 1$, cette dérivée est nulle.

Remarque 10. On rappelle que pour tout réel x , on a $x^0 = 1$. L'élévation à la puissance 0 ne pose AUCUN problème. Ce qui est problématique, c'est le calcul de limites de puissances dont l'exposant TEND vers 0 (et pas dont l'exposant VAUT zéro).

2.3 Racines

Soit $n \in \mathbb{N}^*$. La fonction $\mathbb{R}_+ \rightarrow \mathbb{R}_+$ définie par $x \mapsto x^n$ est une bijection de $\mathbb{R}_+$ sur $\mathbb{R}_+$. Sa réciproque est appelée fonction racine n ème. On note $\sqrt[n]{x}$ l'image du réel $x \geq 0$ par cette application. Si n est un entier naturel impair la fonction $x \mapsto x^n$ est une bijection de $\mathbb{R}$ sur $\mathbb{R}$, ce qui permet de définir une fonction racine n ème sur $\mathbb{R}$.

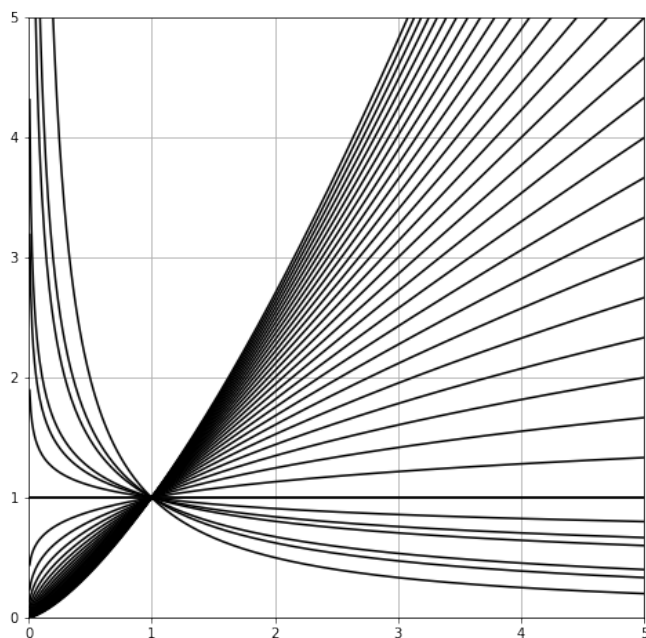


FIGURE 13.4 – Puissances

- Pour $n \in \mathbb{N}^*$ et $x > 0$, on a $(\sqrt[n]{x})^n = x$, donc $\sqrt[n]{x} = x^{\frac{1}{n}}$. Ainsi, pour $x > 0$,

$$\frac{d}{dx} \sqrt[n]{x} = \frac{d}{dx} x^{\frac{1}{n}} = \frac{1}{n} x^{\frac{1}{n}-1} = \frac{1}{n \sqrt[n]{x^{n-1}}}$$

- Lorsque n est impair et $x < 0$ la formule de dérivation ci-dessus reste encore valable.

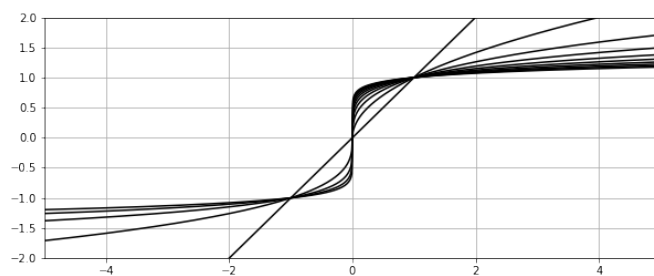


FIGURE 13.5 – Racines

2.4 Comparaison des logarithmes, puissances et exponentielles

Proposition 9. Soient a et b sont deux réels strictement positifs. On a :

$$\frac{(\ln x)^b}{x^a} \xrightarrow{x \rightarrow +\infty} 0 \quad \text{et} \quad x^a (\ln x)^b \xrightarrow{x \rightarrow 0} 0$$

Démonstration. Nous avons déjà vu que $\frac{\ln x}{x} \rightarrow 0$ lorsque x tend vers l'infini. De là,

$$\frac{(\ln x)^b}{x^a} = \left(\frac{b}{a}\right)^b \left(\frac{\ln(x^{a/b})}{x^{a/b}}\right)^b$$

En remplaçant x par $\frac{1}{x}$ on a le résultat lorsque x tend vers 0. $\square$

Proposition 10. Soient $a > 1$ et $b > 0$. On a

$$\frac{x^b}{a^x} \xrightarrow{x \rightarrow +\infty} 0 \quad \text{et} \quad |x|^b \exp(ax) \xrightarrow{x \rightarrow -\infty} 0$$

Démonstration. Pour la première limite, poser $x = \ln t$. La deuxième limite est une conséquence de la première : poser $x = -t$. $\square$

Remarque 11. Si $a, b < 0$, un passage à l'inverse donne les limites demandées. Si a et b sont de signes contraires, il n'y a aucune indétermination.

3 Fonctions circulaires

3.1 Fonctions cosinus, sinus et tangente

Nous avons déjà parlé de ces fonctions dans le chapitre sur les nombres complexes. Contenons nous d'un petit exercice de révision.

Proposition 11. Soit $x \in \mathbb{R} \setminus \{\pi + 2k\pi, k \in \mathbb{Z}\}$, Soit $t = \tan \frac{x}{2}$. On a

$$\begin{aligned}\cos x &= \frac{1 - t^2}{1 + t^2} \\ \sin x &= \frac{2t}{1 + t^2} \\ \tan x &= \frac{2t}{1 - t^2}\end{aligned}$$

Démonstration. Remplaçons t par $\frac{\sin \frac{x}{2}}{\cos \frac{x}{2}}$. Il vient, en mettant tout au même dénominateur,

$$\begin{aligned}\frac{1 - t^2}{1 + t^2} &= \frac{\cos^2 \frac{x}{2} - \sin^2 \frac{x}{2}}{\cos^2 \frac{x}{2} + \sin^2 \frac{x}{2}} \\ &= \cos x\end{aligned}$$

On procède de même pour $\sin x$. On fait le quotient pour $\tan x$. $\square$

3.2 Fonctions Arc sinus et Arc cosinus

Proposition 12. La fonction $f : x \mapsto \sin x$ est une bijection continue et strictement croissante de $[-\frac{\pi}{2}, \frac{\pi}{2}]$ sur $[-1, 1]$

Définition 7. La réciproque de f est appelée le fonction Arc sinus. On la note $\arcsin$.

Remarque 12. $\arcsin$ n'est PAS la réciproque de $\sin$, mais de l'application que nous avons appelée f . Pour $x \in [-1, 1]$, $\arcsin x$ est l'unique élément de $[-\frac{\pi}{2}, \frac{\pi}{2}]$ dont le sinus vaut x . Par exemple, $\arcsin 0 = 0$ puisque $\sin 0 = 0$. Autre exemple, $\arcsin 1 = \frac{\pi}{2}$ puisque $\sin \frac{\pi}{2} = 1$.

Proposition 13. La fonction $\arcsin$ est une bijection continue strictement croissante et continue de $[-1, 1]$ sur $[-\frac{\pi}{2}, \frac{\pi}{2}]$. Elle est impaire, puisque réciproque d'une fonction impaire.

Remarque 13. On a pour tout $x \in [-1, 1]$, $\sin(\arcsin x) = x$. En revanche, la relation $\arcsin(\sin x) = x$ n'est valable que pour $x \in [-\frac{\pi}{2}, \frac{\pi}{2}]$. À titre d'exercice, tracer le graphe de la fonction $f : \mathbb{R} \rightarrow \mathbb{R}$.

$$x \mapsto \arcsin(\sin x)$$

Proposition 14. *La fonction $g : x \mapsto \cos x$ est une bijection continue et strictement décroissante de $[0, \pi]$ sur $[-1, 1]$*

Définition 8. La réciproque de g est appelée le fonction Arc cosinus. On la note $\arccos$.

Remarque 14. $\arccos$ n'est PAS la réciproque de $\cos$, mais de l'application que nous avons appelée g . Pour $x \in [-1, 1]$, $\arccos x$ est l'unique élément de $[0, \pi]$ dont le cosinus vaut x . Par exemple, $\arccos 0 = \frac{\pi}{2}$ puisque $\cos \frac{\pi}{2} = 0$. Autre exemple, $\arccos 1 = 0$ puisque $\cos 0 = 1$.

Proposition 15. *La fonction $\arccos$ est une bijection continue strictement décroissante de $[-1, 1]$ sur $[0, \pi]$. Elle n'est ni paire, ni impaire.*

Remarque 15. On a pour tout $x \in [-1, 1]$, $\cos(\arccos x) = x$. En revanche, la relation $\arccos(\cos x) = x$ n'est valable que pour $x \in [0, \pi]$. À titre d'exercice, tracer le graphe de la fonction $f : \mathbb{R} \rightarrow \mathbb{R}$:

$$x \mapsto \arccos(\cos x)$$

Exercice.

1. Simplifier $\cos(\arcsin x)$ où $x \in [-1, 1]$.
2. Faire de même avec $\sin(\arccos x)$.
3. Montrer la relation $\arcsin x + \arccos x = \frac{\pi}{2}$ pour $x \in [-1, 1]$. On pose pour cela $y = \frac{\pi}{2} - \arcsin x$ et on prouve que $y \in [0, \pi]$ et $\cos y = x$.

Proposition 16. *Les fonctions Arc sinus et Arc cosinus sont de classe $\mathcal{C}^\infty$ sur $] -1, 1[$ et*

$$\forall x \in] -1, 1[, \arcsin' x = \frac{1}{\sqrt{1-x^2}}$$

$$\forall x \in] -1, 1[, \arccos' x = -\frac{1}{\sqrt{1-x^2}}$$

Remarque 16. Ces fonctions ne sont pas dérivables en ± 1 , puisque les dérivées de leurs réciproques au point correspondant sont nulles. On a en ces points une tangente verticale à la courbe.

Représentations graphiques

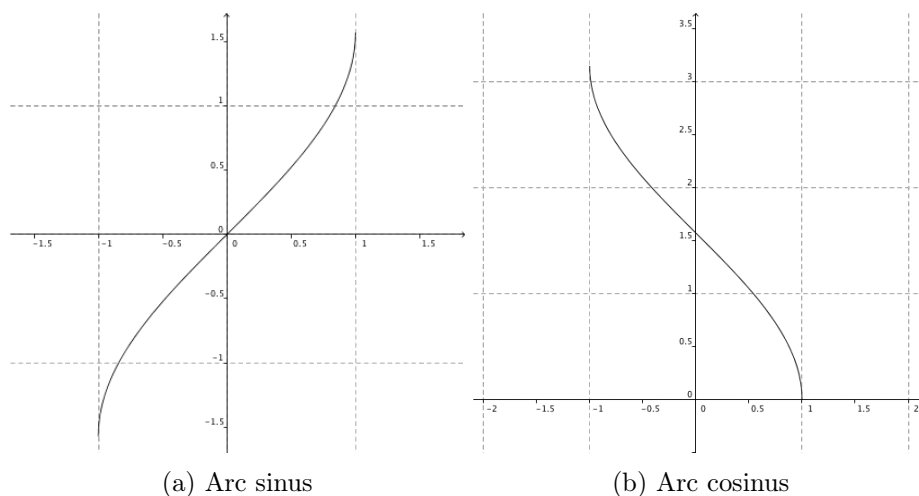


FIGURE 13.6 – Fonctions circulaires inverses

3.3 Fonction Arc tangente

Proposition 17. *La fonction $h : x \mapsto \tan x$ est une bijection continue strictement croissante de $] -\frac{\pi}{2}, \frac{\pi}{2}[$ sur $\mathbb{R}$*

Définition 9. La réciproque de h est appelée le fonction Arc tangente. On la note $\arctan$.

Remarque 17. $\arctan$ n'est PAS la réciproque de $\tan$, mais de l'application que nous avons appelée h . Pour $x \in \mathbb{R}$, $\arctan x$ est l'unique élément de $] -\frac{\pi}{2}, \frac{\pi}{2}[$ dont la tangente vaut x . Par exemple, $\arctan 0 = 0$ puisque $\tan 0 = 0$. Autre exemple, $\arctan 1 = \frac{\pi}{4}$ puisque $\tan \frac{\pi}{4} = 1$.

Proposition 18. *La fonction $\arctan$ est une bijection continue strictement croissante de $\mathbb{R}$ sur $] -\frac{\pi}{2}, \frac{\pi}{2}[$. Elle est impaire, puisque réciproque d'une fonction impaire.*

Remarque 18. On a pour tout $x \in \mathbb{R}$, $\tan(\arctan x) = x$. En revanche, la relation $\arctan(\tan x) = x$ n'est valable que pour $x \in] -\frac{\pi}{2}, \frac{\pi}{2}[$. À titre d'exercice, tracer le graphe de la fonction $f : \mathbb{R} \rightarrow \mathbb{R}$:

$$x \mapsto \arctan(\tan x)$$

Proposition 19. *La fonction Arc tangente est de classe $\mathcal{C}^\infty$ sur $\mathbb{R}$, et :*

$$\forall x \in \mathbb{R}, \arctan' x = \frac{1}{x^2 + 1}$$

Remarque 19. Pour $x \in \mathbb{R}^*$, posons $f(x) = \arctan x + \arctan \frac{1}{x}$. La fonction f est dérivable sur les intervalles $\mathbb{R}_+^*$ et $\mathbb{R}_-^*$, et sa dérivée y est nulle (le vérifier). Elle y est donc constante. De plus, $f(1) = \frac{\pi}{2}$ et $f(-1) = -\frac{\pi}{2}$. On en déduit que $\forall x > 0, \arctan x + \arctan \frac{1}{x} = \frac{\pi}{2}$ et $\forall x < 0, \arctan x + \arctan \frac{1}{x} = -\frac{\pi}{2}$.

Représentation graphique

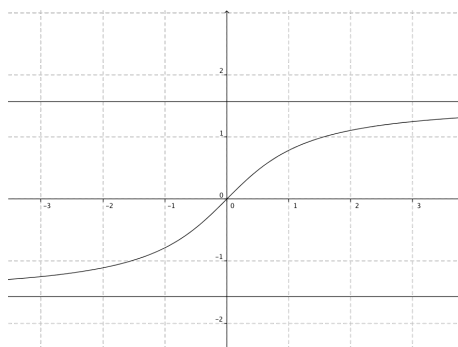


FIGURE 13.7 – Arc tangente

4 Fonctions hyperboliques

4.1 Fonctions cosinus et sinus hyperbolique

Définition 10. On définit les fonctions sinus et cosinus hyperbolique pour tout x réel par :

$$\sinh x = \frac{e^x - e^{-x}}{2} \text{ et } \cosh x = \frac{e^x + e^{-x}}{2}$$

Proposition 20. *La fonction $\sinh$ est impaire, la fonction $\cosh$ est paire. Ces deux fonctions sont de classe $\mathcal{C}^\infty$ sur $\mathbb{R}$ et on a pour tout $x \in \mathbb{R}$:*

$$\sinh' x = \cosh x \text{ et } \cosh' x = \sinh x$$

4.2 Trigonométrie hyperbolique

Proposition 21. *Les formules ci-dessous sont vraies pour tous réels x, a, b :*

$$\begin{aligned}
 e^x &= \sinh x + \cosh x \\
 \cosh^2 x - \sinh^2 x &= 1 \\
 \sinh(a+b) &= \sinh a \cosh b + \cosh a \sinh b \\
 \sinh(a-b) &= \sinh a \cosh b - \cosh a \sinh b \\
 \cosh(a+b) &= \cosh a \cosh b + \sinh a \sinh b \\
 \cosh(a-b) &= \cosh a \cosh b - \sinh a \sinh b \\
 \sinh 2x &= 2 \sinh x \cosh x \\
 \cosh 2x &= \cosh^2 x + \sinh^2 x = 2 \cosh^2 x - 1 = 1 + 2 \sinh^2 x
 \end{aligned}$$

Remarque 20. La formule $\cosh^2 t - \sinh^2 t = 1$ permet de paramétrer l'hyperbole $\frac{x^2}{a^2} - \frac{y^2}{b^2} = 1$ par $x = \pm a \cosh t, y = b \sinh t$.

4.3 Fonction tangente hyperbolique

Définition 11. La fonction tangente hyperbolique, notée $\tanh$, est définie sur $\mathbb{R}$ par

$$\tanh x = \frac{\sinh x}{\cosh x} = \frac{e^{2x} - 1}{e^{2x} + 1} = \frac{e^x - e^{-x}}{e^x + e^{-x}}$$

Proposition 22. *La fonction $\tanh$ est impaire, de classe $\mathcal{C}^\infty$ sur $\mathbb{R}$, et :*

$$\forall x \in \mathbb{R}, \tanh' x = \frac{1}{\cosh^2 x} = 1 - \tanh^2 x$$

Représentation graphique

4.4 Fonctions hyperboliques réciproques

Définition 12. La fonction « sinus hyperbolique » est continue et strictement croissante sur $\mathbb{R}$. elle définit une bijection de $\mathbb{R}$ sur $\mathbb{R}$ dont la réciproque est appelée *Argument sinus hyperbolique* et notée argsh .

La fonction argsh est ainsi une bijection continue, strictement croissante, et impaire, de $\mathbb{R}$ sur $\mathbb{R}$.

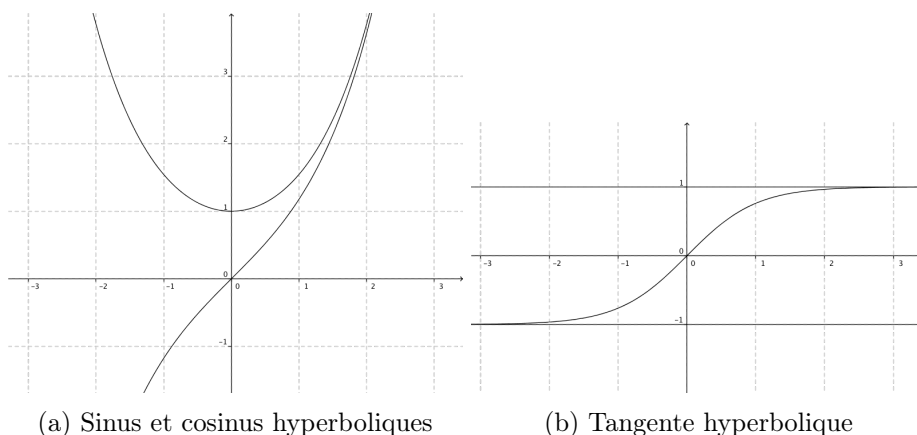


FIGURE 13.8 – Fonctions hyperboliques

Définition 13. La fonction cosinus hyperbolique est continue et strictement croissante sur $\mathbb{R}_+$. elle définit une bijection de $\mathbb{R}_+$ sur $[1, +\infty[$ dont la réciproque est appelée *Argument cosinus hyperbolique* et notée argch .

La fonction argch est ainsi une bijection continue, strictement croissante, de $[1, +\infty[$ sur $\mathbb{R}_+$.

Proposition 23.

- La fonction argsh est dérivable sur $\mathbb{R}$, et

$$\forall x \in \mathbb{R}, \operatorname{argsh}'x = \frac{1}{\sqrt{1+x^2}}$$

- La fonction argch est dérivable sur $[1, +\infty[$, et

$$\forall x > 1, \operatorname{argch}'x = \frac{1}{\sqrt{x^2-1}}$$

Démonstration. Faisons-le pour argsh . On a pour tout réel x ,

$$\operatorname{argsh}'x = \frac{1}{\sinh'(\operatorname{argsh} x)} = \frac{1}{\cosh(\operatorname{argsh} x)} = \frac{1}{\sqrt{1 + \sinh^2(\operatorname{argsh} x)}} = \frac{1}{\sqrt{1 + x^2}}$$

□

Définition 14. La fonction tangente hyperbolique est continue et strictement croissante sur $\mathbb{R}$. elle définit une bijection de $\mathbb{R}$ sur $] -1, 1[$ dont la réciproque est appelée *Argument tangente hyperbolique* et notée argth .

La fonction argth est ainsi une bijection continue, strictement croissante, de $] -1, 1[$ sur $\mathbb{R}$.

Proposition 24. La fonction argth est dérivable sur $\mathbb{R}$, et

$$\forall x \in \mathbb{R}, \operatorname{argth}' x = \frac{1}{1 - x^2}$$

Représentations graphiques

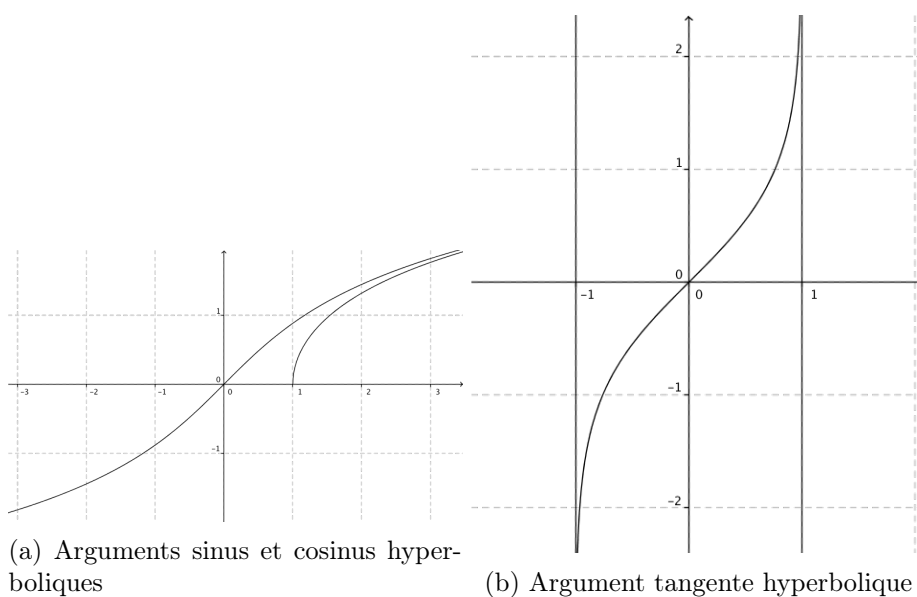


FIGURE 13.9 – Fonctions hyperboliques inverses

Expression à l'aide logarithmes

Proposition 25. On a

- $\forall x \in \mathbb{R}, \operatorname{argsh} x = \ln(x + \sqrt{x^2 + 1})$.
- $\forall x \in [1, +\infty[, \operatorname{argch} x = \ln(x + \sqrt{x^2 - 1})$.

$$\bullet \forall x \in]-1, 1[, \operatorname{argth} x = \frac{1}{2} \ln \left(\frac{1+x}{1-x} \right).$$

Démonstration. Faisons-le pour l'argument sinus hyperbolique. Soient $x, t \in \mathbb{R}$. On a $\sinh t = x$ si et seulement si $e^t - e^{-t} = 2x$ ou encore $T^2 - 2xT - 1 = 0$, où l'on a posé $T = e^t$. $T > 0$, et l'unique racine strictement positive de l'équation en T est $x + \sqrt{x^2 + 1}$. On en tire $t = \operatorname{argsh} x = \ln(x + \sqrt{x^2 + 1})$. $\square$

Chapitre 14

Fonctions convexes

1 Intervalles de $\mathbb{R}$

Cette section contient des rappels et quelques compléments sur les intervalles de $\mathbb{R}$.

1.1 Intervalles et parties convexes

Définition 1. Soit $A \subset \mathbb{R}$. On dit que A est convexe lorsque pour tous $x, y \in A$,

$$x \leq y \implies [x, y] \subset A$$

Proposition 1. Les intervalles de $\mathbb{R}$ sont les parties convexes de $\mathbb{R}$.

Démonstration. Cette propriété a déjà été démontrée dans le cours sur les réels. Redonnons-en la preuve.

- Soit I un intervalle. Il y a différents cas selon le type d'intervalle. Traitons par exemple le cas où $I =]-\infty, b]$, où $b \in \mathbb{R}$. Soient $x, y \in I$. Supposons $x \leq y$. Soit $z \in [x, y]$. On a

$$-\infty < x \leq z \leq y \leq b$$

et donc $z \in I$. Ainsi, I est convexe.

- Passons au sens intéressant. Soit I une partie convexe de $\mathbb{R}$. Ici aussi il y a différents cas. Traitons par exemple le cas où I est non vide, majoré et non minoré.

I étant non vide et majoré, il possède une borne supérieure b . Ainsi, $I \subset]-\infty, b]$.

Montrons à présent que $] -\infty, b[\subset I$. Soit $z < b$. Comme I n'est pas minoré, z ne minore pas I . Ainsi, il existe $x \in I$ tel que $x < z$. De même, b étant le plus petit majorant de I , z ne majore pas I . Il existe donc $y \in I$ tels que $z < y$. Ainsi, $x, y \in I$ et $z \in [x, y]$. Par la convexité de I , $z \in I$. Nous avons donc

$$]-\infty, b[\subset I \subset]-\infty, b]$$

Il n'y a que deux possibilités : $I =]-\infty, b[$ ou $I =]-\infty, b]$. Dans les deux cas, I est un intervalle.

□

1.2 Paramétrage d'un segment

Proposition 2. Soient $x, y \in \mathbb{R}$ tels que $x \leq y$. Soit $z \in \mathbb{R}$. On a $z \in [x, y]$ si et seulement si il existe $\lambda \in [0, 1]$ tel que

$$z = (1 - \lambda)x + \lambda y$$

Ainsi,

$$[x, y] = \{(1 - \lambda)x + \lambda y, \lambda \in [0, 1]\}$$

Démonstration. La propriété est évidente si $x = y$. Supposons donc $x < y$. Soit $z \in \mathbb{R}$. On a

$$\begin{aligned} z \in [x, y] &\iff x \leq z \leq y \\ &\iff 0 \leq z - x \leq y - x \\ &\iff 0 \leq \frac{z - x}{y - x} \leq 1 \\ &\iff \exists \lambda \in [0, 1], \frac{z - x}{y - x} = \lambda \\ &\iff \exists \lambda \in [0, 1], z = (1 - \lambda)x + \lambda y \end{aligned}$$

□

Corollaire 3. Soit $I \subset \mathbb{R}$. I est un intervalle de $\mathbb{R}$ si et seulement si pour tous $x, y \in I$, pour tout $\lambda \in [0, 1]$,

$$(1 - \lambda)x + \lambda y \in I$$

Démonstration. Immédiate. Ce n'est qu'une reformulation de l'appartenance à un segment. Remarquons qu'il est inutile ici de supposer que $x \leq y$. Cette hypothèse n'a besoin d'être faite que lorsqu'on a besoin d'écrire « $[x, y]$ ». □

Remarque. Cette propriété peut être reformulée de la façon suivante : pour tous $x, y \in I$, pour tous $\lambda_1, \lambda_2 \geq 0$ tels que $\lambda_1 + \lambda_2 = 1$, on a $\lambda_1 x + \lambda_2 y \in I$. L'inconvénient de cette formulation est l'apparition de deux réels λ_1, λ_2 au lieu d'un unique réel λ . Son avantage est qu'elle rend les choses plus symétriques et que maintenant nous pouvons nous demander s'il est possible de généraliser la propriété à plus de deux points. La réponse est oui.

1.3 Barycentres

Proposition 4. Soit I un intervalle de $\mathbb{R}$. Pour tout $n \geq 1$, pour tous $x_1, \dots, x_n \in I$, pour tous $\lambda_1, \dots, \lambda_n \geq 0$ tels que $\sum_{i=1}^n \lambda_i = 1$, on a

$$\sum_{i=1}^n \lambda_i x_i \in I$$

Démonstration. On fait une récurrence sur n .

- Le cas $n = 1$ est évident.
- Le cas $n = 2$ résulte du corollaire ci-dessus, en posant $\lambda = \lambda_2$, puisqu'alors $1 - \lambda = \lambda_1$.
- Soit $n \geq 1$. Supposons la propriété vraie pour l'entier n . Donnons nous $x_1, \dots, x_{n+1} \in I$. Soient $\lambda_1, \dots, \lambda_{n+1} \geq 0$ tels que $\sum_{i=1}^{n+1} \lambda_i = 1$. Deux cas se présentent.

- Cas 1 : $\lambda_{n+1} = 1$. On a alors $\lambda_1 = \dots = \lambda_n = 0$ et la propriété à montrer est évidente.
- Cas 2 : $\lambda_{n+1} < 1$. Soit $x = \sum_{i=1}^{n+1} \lambda_i x_i$. Posons également $\lambda = \lambda_{n+1}$. On peut alors écrire

$$x = (1 - \lambda) \sum_{i=1}^n \lambda'_i x_i + \lambda x_{n+1}$$

où

$$\lambda'_i = \frac{\lambda_i}{1 - \lambda}$$

Remarquons que les λ'_i sont positifs ou nuls, et que

$$\sum_{i=1}^n \lambda'_i = \frac{1}{1 - \lambda} \sum_{i=1}^n \lambda_i = \frac{1 - \lambda}{1 - \lambda} = 1$$

Par l'hypothèse de récurrence, $x' = \sum_{i=1}^n \lambda'_i x_i \in I$. Pour terminer,

$$x = (1 - \lambda)x' + \lambda x_{n+1} \in I$$

par le cas $n = 2$.

□

2 Notion de fonction convexe

2.1 C'est quoi ?

Définition 2. Soit I un intervalle de $\mathbb{R}$. Soit $f : I \rightarrow \mathbb{R}$. On dit que f est convexe sur I lorsque pour tous $x, y \in I$ et tout $\lambda \in [0, 1]$,

$$f((1 - \lambda)x + \lambda y) \leq (1 - \lambda)f(x) + \lambda f(y)$$

Lorsque l'inégalité est renversée, on dit que f est concave sur I .

Remarque. Par le paragraphe précédent, cette définition a bien un sens. Lorsque $x, y \in I$ et $\lambda \in [0, 1]$, le réel $(1 - \lambda)x + \lambda y$ est encore dans I , et on peut donc calculer $f((1 - \lambda)x + \lambda y)$.

Remarque. Dans la suite nous parlerons essentiellement de fonctions convexes. Tous les résultats que nous donneront s'adaptent évidemment à des fonctions concaves, la plupart du temps en renversant des inégalités ou en changeant des signes. Il suffit en général de remarquer que f est convexe si et seulement si $-f$ est concave pour conclure.

Exemples.

- Les fonctions affines $x \mapsto ax + b$ sont à la fois convexes et concaves sur $\mathbb{R}$.
- Soit $f : x \mapsto |x|$. Soient $x, y \in \mathbb{R}$. Soit $\lambda \in [0, 1]$. On a

$$\begin{aligned} f((1 - \lambda)x + \lambda y) &= |(1 - \lambda)x + \lambda y| \\ &\leq (1 - \lambda)|x| + \lambda|y| \quad \text{par l'inégalité triangulaire} \end{aligned}$$

Remarquons que nous avons utilisé $\lambda \geq 0$ et $1 - \lambda \geq 0$. Ainsi, la fonction valeur absolue est convexe sur $\mathbb{R}$.

- Soit $f : x \mapsto x^2$. Soient $x, y \in \mathbb{R}$. Soit $\lambda \in [0, 1]$. On a

$$\begin{aligned} (1 - \lambda)f(x) + \lambda f(y) - (f((1 - \lambda)x + \lambda y)) &= (1 - \lambda)x^2 + \lambda y^2 - ((1 - \lambda)x + \lambda y)^2 \\ &= \lambda(1 - \lambda)(x - y)^2 \\ &\geq 0 \end{aligned}$$

La fonction f est donc convexe sur $\mathbb{R}$.

2.2 Interprétation graphique

Définition 3. Soient $A, B \in \mathbb{R}^2$. Le segment $[A, B]$ est l'ensemble des points $M \in \mathbb{R}^2$ vérifiant

$$\exists \lambda \in [0, 1], \overrightarrow{AM} = \lambda \overrightarrow{AB}$$

Soit $f : I \rightarrow \mathbb{R}$. Soient $x, y \in I$. Soient $A = (x, f(x))$ et $B = (y, f(y))$. Soit $M = (u, v) \in \mathbb{R}^2$. Le point M est sur le segment $[A, B]$ si et seulement si il existe $\lambda \in [0, 1]$ tel que

$$\overrightarrow{OM} = \overrightarrow{OA} + \lambda \overrightarrow{AB}$$

c'est à dire, en termes de coordonnées,

$$\begin{cases} u &= (1 - \lambda)x + \lambda y \\ v &= (1 - \lambda)f(x) + \lambda f(y) \end{cases}$$

La fonction f est convexe sur I si et seulement si $f(u) \leq v$, ceci pour tous x, y, λ . Dit autrement, si et seulement si le point M est au-dessus du point $(u, f(u))$. Prenant deux points quelconques A et B de la courbe de f et traçant le segment $[A, B]$, la portion de la courbe comprise (l'arc) entre A et B se trouve sous le segment $[A, B]$ (la corde). Le dessin ci-dessous illustre la propriété, toutes les cordes se trouvent dans l'épigraphé de la fonction f , c'est à dire dans la partie du plan au-dessus de la courbe.

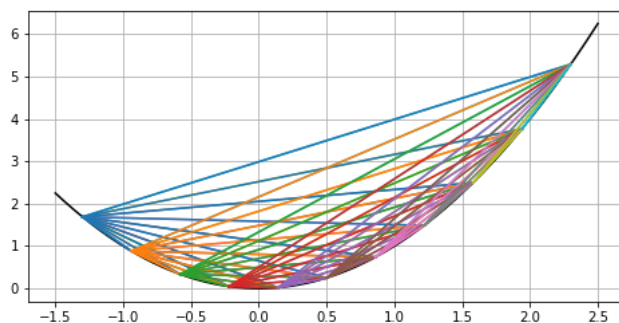


FIGURE 14.1 – Arcs et cordes

2.3 Généralisation de l'inégalité de convexité

Proposition 5. Soit $f : I \rightarrow \mathbb{R}$. La fonction f est convexe sur I si et seulement si pour tout entier $n \geq 1$, pour tous $x_1, \dots, x_n \in I$, pour tous $\lambda_1, \dots, \lambda_n \geq 0$ tels que

$$\sum_{i=1}^n \lambda_i = 1,$$

$$f\left(\sum_{i=1}^n \lambda_i x_i\right) \leq \sum_{i=1}^n \lambda_i f(x_i)$$

Démonstration. Remarquons tout d'abord que si cette inégalité très générale est vraie, le cas $n = 2$ nous dit que f est convexe.

Inversement, supposons f convexe sur I . On reprend les idées de la démonstration que nous avons faite dans la section sur les intervalles. On procède par récurrence sur n .

- Le cas $n = 1$ est évident.
- Le cas $n = 2$ est la définition de la convexité.
- Soit $n \geq 1$. Supposons la propriété vraie pour l'entier n . Donnons nous $x_1, \dots, x_{n+1} \in$

I . Soient $\lambda_1, \dots, \lambda_{n+1} \geq 0$ tels que $\sum_{i=1}^{n+1} \lambda_i = 1$. Deux cas se présentent.

Cas 1. $\lambda_{n+1} = 1$. On a alors $\lambda_0 = \dots = \lambda_n = 0$ et la propriété à montrer est évidente.

Cas 2. $\lambda_{n+1} < 1$. Soit $x = \sum_{i=1}^{n+1} \lambda_i x_i$. Posons également $\lambda = \lambda_{n+1}$. On peut alors écrire

$$f(x) = f((1 - \lambda)x' + \lambda x_{n+1})$$

où, en posant $\lambda'_i = \frac{\lambda_i}{1 - \lambda}$,

$$x' = \sum_{i=1}^n \lambda'_i x_i \in I$$

Par la convexité de f ,

$$f(x) \leq (1 - \lambda)f(x') + \lambda f(x_{n+1})$$

Par ailleurs, les λ'_i , $1 \leq i \leq n$, sont positifs ou nuls et de somme 1. Par l'hypothèse de récurrence,

$$f(x') = f\left(\sum_{i=1}^n \lambda'_i x_i\right) \leq \sum_{i=1}^n \lambda'_i f(x_i)$$

En reportant dans l'inégalité précédente et en se souvenant que $(1 - \lambda)\lambda'_i = \lambda_i$, on obtient

$$f\left(\sum_{i=1}^{n+1} \lambda_i x_i\right) \leq \sum_{i=1}^{n+1} \lambda_i f(x_i)$$

□

2.4 Croissance des pentes

Notation. Soit $f : I \rightarrow \mathbb{R}$. Pour $x, y \in I$ tels que $x \neq y$, notons

$$\Delta f(x, y) = \frac{f(y) - f(x)}{y - x}$$

Si $A = (x, f(x))$ et $B = (y, f(y))$, $\Delta(x, y)$ est la *pente* du segment $[A, B]$.

Proposition 6. CROISSANCE DES PENTES

Soit $f : I \rightarrow \mathbb{R}$. Les quatre propriétés ci-dessous sont équivalentes.

1. Pour tous $x, y, z \in I$ vérifiant $x < z < y$, $\Delta f(x, z) \leq \Delta f(z, y)$.
2. Pour tous $x, y, z \in I$ vérifiant $x < z < y$, $\Delta f(x, z) \leq \Delta f(x, y)$.
3. Pour tous $x, y, z \in I$ vérifiant $x < z < y$, $\Delta f(x, y) \leq \Delta f(z, y)$.
4. f est convexe sur I .

Il est facile de se souvenir de cette propriété. Elle nous dit que les pentes des cordes de la courbe d'une fonction convexe obéissent à des inégalités. Tout est résumé le dessin ci-dessous, où est dessiné un triangle dont les trois côtés sont rouge, vert et bleu. Si f est convexe, les pentes des côtés de *tous* les triangles vérifient

$$(1) \text{ rouge} \leq \text{bleu} \quad (2) \text{ rouge} \leq \text{vert} \quad (3) \text{ vert} \leq \text{bleu}$$

Inversement, si l'une de ces inégalités est valable pour *tous* les triangles, alors la fonction est convexe.

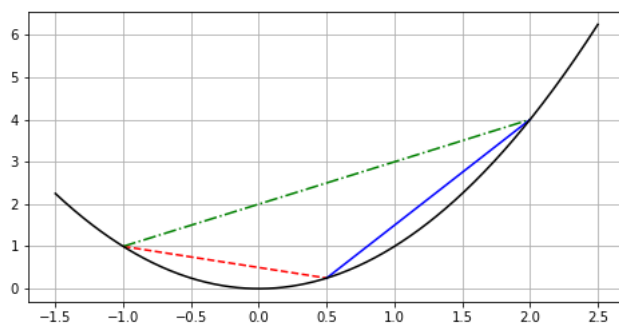


FIGURE 14.2 – inégalités fondamentales

Démonstration. Soient $x, y, z \in I$ vérifiant $x < z < y$. Posons $z = (1 - \lambda)x + \lambda y$ où $0 < \lambda < 1$. On a

$$\Delta f(x, z) = \frac{f(z) - f(x)}{\lambda(y - x)}$$

$$\Delta f(z, y) = \frac{f(y) - f(z)}{(1 - \lambda)(y - x)}$$

$$\Delta f(x, y) = \frac{f(y) - f(x)}{y - x}$$

Notons $\alpha = (1 - \lambda)f(x) + \lambda f(y) - f(z)$. On a

$$\begin{aligned} (1) \quad \Delta f(z, y) - \Delta f(x, z) &= \frac{f(y) - f(z)}{(1 - \lambda)(y - x)} - \frac{f(z) - f(x)}{\lambda(y - x)} = \frac{\alpha}{\lambda(1 - \lambda)(y - x)} \\ (2) \quad \Delta f(x, y) - \Delta f(x, z) &= \frac{f(y) - f(x)}{y - x} - \frac{f(z) - f(x)}{\lambda(y - x)} = \frac{\alpha}{\lambda(y - x)} \\ (3) \quad \Delta f(z, y) - \Delta f(x, y) &= \frac{f(y) - f(z)}{(1 - \lambda)(y - x)} - \frac{f(y) - f(x)}{y - x} = \frac{\alpha}{(1 - \lambda)(y - x)} \end{aligned}$$

La quantité (1) est positive pour tous $x < z < y$ si et seulement si $\alpha \geq 0$ pour tous $x < z < y$, ce qui équivaut à la convexité de f . Il en est de même pour (2) et (3). $\square$

Remarque. Lorsque $f : I \rightarrow \mathbb{R}$ est une fonction convexe, on a donc pour tous $x, y, z \in I$ vérifiant $x < z < y$,

$$\Delta f(x, z) \leq \Delta f(x, y) \leq \Delta f(y, z)$$

Corollaire 7. Soit $f : I \rightarrow \mathbb{R}$ une fonction convexe. Soient $x, y \in I$ tels que $x < y$. Soient $A = (x, f(x))$ et $B = (y, f(y))$. Soient enfin $z \in I$ et $C = (z, f(z))$.

- Si $x \leq z \leq y$, alors le point C est au-dessous de la droite (AB) .
- Si $z < x$ ou $y < z$, alors le point C est au-dessus de la droite (AB) .

Démonstration. Remarquons tout d'abord qu'une équation cartésienne de la droite $D = (AB)$ est

$$(D) \quad Y = f(x) + (X - x)\Delta f(x, y)$$

Le point C' de la droite D qui a la même abscisse que C est donc $C' = (z, w)$ où

$$w = f(x) + (z - x)\Delta f(x, y)$$

- Nous avons déjà montré le premier point plus haut.
- Supposons par exemple $y < z$ (le cas $z < x$ se montre de la même manière). Par la croissance des pentes (rouge $\leq$ vert), on a

$$\Delta f(x, y) \leq \Delta f(x, z)$$

ou encore

$$f(z) \geq f(x) + (z - x)\Delta f(x, y)$$

c'est à dire $f(z) \geq w$. Le point C est donc au-dessus du point C' .

$\square$

3 Fonctions convexes dérivables

3.1 Convexité et dérivées

Proposition 8. Soit $f : I \rightarrow \mathbb{R}$ une fonction convexe. Alors

- f est dérivable à gauche et à droite en tout point de $\overset{\circ}{I}$.
- Pour tout $x \in \overset{\circ}{I}$, $f'_g(x) \leq f'_d(x)$.
- f'_g et f'_d sont croissantes sur $\overset{\circ}{I}$.

Démonstration.

- Soit $a \in \overset{\circ}{I}$. Soient $x, x' \in I$ vérifiant $a < x < x'$. Par l'inégalité fondamentale numéro 2, on a

$$\frac{f(x) - f(a)}{x - a} \leq \frac{f(x') - f(a)}{x' - a}$$

La fonction $\varphi : x \mapsto \frac{f(x) - f(a)}{x - a}$, définie sur $J = I \cap]a, +\infty[$, est donc croissante sur J . Donnons nous par ailleurs $t < a$, $t \in I$. Par l'inégalité fondamentale numéro 1, on a pour tout $x \in J$,

$$\frac{f(a) - f(t)}{a - t} \leq \varphi(x)$$

Ainsi, φ est minorée sur J . La fonction φ , croissante et minorée, admet donc une limite réelle en a . En conclusion, f est dérivable à droite en a . La dérivabilité à gauche se montre de la même façon.

- Soit $a \in \overset{\circ}{I}$. Soient $x, x' \in I$ vérifiant $x' < a < x$. Par l'inégalité fondamentale numéro 1, on a

$$\frac{f(a) - f(x')}{a - x'} \leq \frac{f(x) - f(a)}{x - a}$$

Faisant tendre x et x' vers a , on obtient $f'_g(a) \leq f'_d(a)$.

- Soient $a, b \in \overset{\circ}{I}$ vérifiant $a < b$. Soient $x, x' \in I$ tels que $a < x < b < x'$. En appliquant deux fois l'inégalité fondamentale numéro 1, on obtient

$$\Delta f(a, x) \leq \Delta f(x, b) \leq \Delta f(b, x')$$

d'où

$$\Delta f(a, x) \leq \Delta f(b, x')$$

Faisant tendre x vers a et x' vers b , on obtient $f'_d(a) \leq f'_d(b)$. Ainsi, la fonction f'_d est croissante sur $\overset{\circ}{I}$. Il en est de même pour f'_g .

□

Corollaire 9. Soit $f : I \rightarrow \mathbb{R}$ une fonction convexe. Alors f est continue sur $\overset{\circ}{I}$.

Démonstration. f est dérivable à gauche sur $\overset{\circ}{I}$, elle est donc continue à gauche en tout point de $\overset{\circ}{I}$. Il en est de même à droite. □

Remarque. Une fonction convexe peut fort bien être discontinue aux bornes de son intervalle de définition. Il suffit de prendre pour exemple la fonction $f : [-1, 1] \rightarrow \mathbb{R}$ définie par $f(x) = x^2$ si $-1 < x < 1$, et $f(\pm 1) = 2$.

Proposition 10. Soit $f : I \rightarrow \mathbb{R}$ une fonction continue sur I et dérivable sur $\overset{\circ}{I}$. Alors f est convexe sur I si et seulement si f' est croissante sur $\overset{\circ}{I}$.

Démonstration.

($\Rightarrow$) Si f est convexe, sa dérivée est croissante, nous l'avons déjà vu.
 ($\Leftarrow$) Supposons f' croissante sur $\overset{\circ}{I}$. Soient $x < z < y$ trois points de I . La fonction f est continue sur $[x, z]$ et dérivable sur $]x, z[$. Par le théorème des accroissements finis, il existe donc $a \in]x, z[$ tel que $\Delta f(x, z) = f'(a)$. De même, il existe $b \in]z, y[$ tel que $\Delta f(z, y) = f'(b)$. On a $x < a < z < b < y$, et donc $a < b$. Par la croissance de f' , on a donc $f'(a) \leq f'(b)$. Ainsi, $\Delta f(x, z) \leq \Delta f(z, y)$. Nous venons donc de prouver la première inégalité fondamentale pour tous $x, y, z \in I$ vérifiant $x < z < y$. Mais ceci équivaut à la convexité de f .

□

Proposition 11. Soit $f : I \rightarrow \mathbb{R}$ une fonction continue sur I et deux fois dérivable sur $\overset{\circ}{I}$. Alors f est convexe sur I si et seulement si $f'' \geq 0$ sur $\overset{\circ}{I}$.

Démonstration. Immédiat par la proposition précédente. En effet, f' croît sur $\overset{\circ}{I}$ si et seulement si sa dérivée y est positive. □

3.2 Exemples

Nous disposons maintenant de critères simples pour montrer qu'une fonction est convexe ou concave. Si elle est deux fois dérivable, il suffit par exemple de regarder le signe de sa dérivée seconde. Chaque fois que l'on possède une fonction convexe, l'inégalité de convexité et sa généralisation sont vérifiées. Nous allons voir que cela nous permet de montrer des inégalités qui sont loin d'être évidentes.

Exemple 1. L'exponentielle est convexe sur $\mathbb{R}$. En effet, pour tout réel x , $\exp''(x) = e^x \geq 0$. On en déduit que pour tous réels x, y , pour tout $t \in [0, 1]$,

$$e^{(1-t)x+ty} \leq (1-t)e^x + te^y$$

Exemple 2. La fonction sin est concave sur $[0, \pi]$ car sa dérivée seconde, $-\sin$, y est négative. Il en résulte que

$$\forall a, b \in [0, \pi], \sin \frac{a+b}{2} \leq \frac{1}{2} (\sin a + \sin b)$$

Exemple 3. La fonction $\ln$ est concave sur $\mathbb{R}_+^*$. En effet, elle y est deux fois dérivable, et pour tout $x > 0$, $\ln''(x) = -\frac{1}{x^2} \leq 0$. On en déduit que pour tous réels $x_1, \dots, x_n > 0$,

$$\ln \left(\frac{1}{n} \sum_{k=1}^n x_k \right) \geq \frac{1}{n} \sum_{k=1}^n \ln x_k$$

Passons à l'exponentielle dans cette inégalité. Il vient

$$\frac{1}{n} \sum_{k=1}^n x_k \geq \sqrt[n]{\prod_{k=1}^n x_k}$$

Exemple 4. Soit p un réel, $p \geq 1$. La fonction $x \mapsto x^p$ est convexe sur $\mathbb{R}_+$, car elle est continue sur $\mathbb{R}_+$, deux fois dérivable sur $\mathbb{R}_+^*$, et sa dérivée seconde en tout $x > 0$ est $p(p-1)x^{p-2} \geq 0$. On a donc pour tout $n \geq 1$ et tous $x_1, \dots, x_n \geq 0$,

$$\left(\frac{1}{n} \sum_{k=1}^n x_k \right)^p \leq \frac{1}{n} \sum_{k=1}^n x_k^p$$

Si, en revanche, $p < 1$, l'inégalité est renversée puisque alors la fonction est concave (si $p \leq 0$ il convient alors de se placer sur $\mathbb{R}_+^*$).

3.3 Position courbe $\longleftrightarrow$ tangente

Proposition 12. Soit $f : I \rightarrow \mathbb{R}$ dérivable. f est convexe sur I si et seulement si la courbe de f est, en tout point, au-dessus de sa tangente.

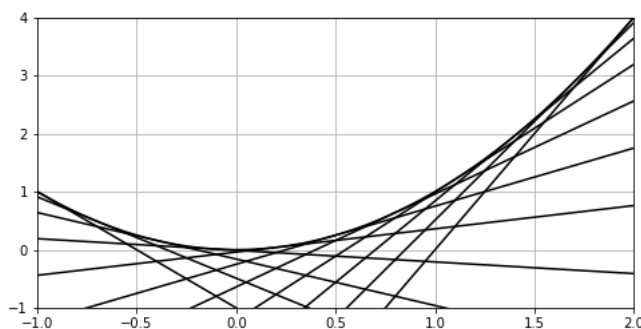


FIGURE 14.3 – courbe et tangentes

Démonstration. Soit $a \in I$. Commençons par traduire la condition « la courbe est au-dessus de sa tangente au point d'abscisse a », que nous noterons $\mathcal{P}(a)$. La tangente à la courbe au point d'abscisse a est la droite d'équation

$$(\mathcal{T}_a) \quad y = f(a) + (x - a)f'(a)$$

La propriété $\mathcal{P}(a)$ est donc

$$\forall x \in I, f(x) \geq f(a) + (x - a)f'(a)$$

($\Rightarrow$) Supposons f convexe. Soit $a \in I$. Supposons que a n'est pas la borne supérieure de I et regardons ce qui se passe à droite de a (ce qui se passe à gauche se montre de la même manière). Soit $x > a$. Pour tout t vérifiant $a < t < x$, on a, par les inégalités fondamentales,

$$\Delta f(a, t) \leq \Delta f(a, x)$$

Lorsque t tend vers a , $\Delta f(a, t)$ tend vers $f'(a)$. Par un passage à la limite, on obtient donc

$$f'(a) \leq \Delta f(a, x) = \frac{f(x) - f(a)}{x - a}$$

Multipliant l'inégalité par $x - a > 0$ et ajoutant $f(a)$, on obtient

$$f(x) \geq f(a) + (x - a)f'(a)$$

qui est l'inégalité cherchée.

($\Leftarrow$) Supposons que $\mathcal{P}(a)$ est vérifiée pour tout $a \in I$. Soient $x, y, z \in I$ vérifiant $x < z < y$. La courbe est au-dessus de sa tangente au point d'abscisse z , donc

$$f(y) \geq f(z) + (y - z)f'(z)$$

On en déduit

$$f'(z) \leq \Delta f(z, y)$$

De même,

$$f(x) \geq f(z) + (x - z)f'(z)$$

et donc

$$f'(z) \geq \Delta f(x, z)$$

Ainsi, $\Delta f(x, z) \leq \Delta f(z, y)$. Ceci étant vrai pour tous $x < z < y$, la fonction f est convexe sur I (inégalités fondamentales).

□

3.4 Inflexions

La notion générale d'inflexion est difficile à manipuler.

Définition 4. Soit $\mathcal{C}$ une courbe possédant une tangente en un point A . On dit que la courbe $\mathcal{C}$ possède une inflexion en A lorsque, au voisinage de A , la courbe traverse sa tangente en A .

Il existe cependant une condition *suffisante* d'inflexion, qui s'avère être réalisée dans beaucoup de situations.

Proposition 13. Soit $f : I \rightarrow \mathbb{R}$ une fonction deux fois dérivable sur I . Soit $a \in \overset{\circ}{I}$. On suppose que f'' s'annule en a en changeant de signe. Alors la courbe de f possède une inflexion au point d'abscisse a .

Démonstration. Supposons par exemple que pour tout x voisin de a , $x \leq a$, $f''(x) \leq 0$ et pour tout x voisin de a , $x \geq a$, $f''(x) \geq 0$. La fonction f est donc concave à gauche de a et convexe à droite de a . La courbe est donc, à gauche de a , sous sa tangente au point d'abscisse a , et, à droite de a , sur sa tangente en ce même point. Il y a donc bien inflexion. $\square$

Remarque. Cette condition d'inflexion est suffisante mais pas nécessaire. Considérons la fonction $f : \mathbb{R} \rightarrow \mathbb{R}$ définie par

$$f(x) = x^5 \left(1 + \sin \frac{1}{x} \right) \text{ si } x \neq 0$$

et $f(0) = 1$. Comme on peut le vérifier, la fonction f est deux fois dérivable sur $\mathbb{R}$. On a $f''(0) = 0$ et donc la tangente à la courbe au point d'abscisse 0 est l'axe Ox . Par ailleurs, f est positive sur $\mathbb{R}_+$ et négative sur $\mathbb{R}_-$, la courbe traverse donc sa tangente en O . En revanche, la dérivée seconde de f change une infinité de fois de signe au voisinage de 0, comme on peut s'en convaincre sur le dessin ci-dessous.

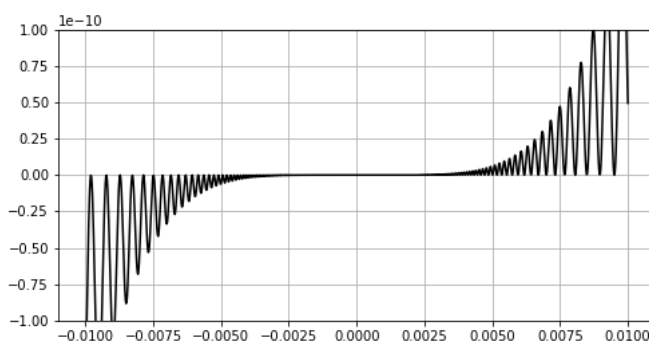


FIGURE 14.4 – une inflexion pathologique

4 Annexe - p -normes

4.1 Introduction

Dans cette section on se donne un réel $p \geq 1$ et un entier $n \geq 1$.

Pour tout $x = (x_1, \dots, x_n) \in \mathbb{R}^n$, notons

$$\|x\|_p = \left(\sum_{i=1}^n |x_i|^p \right)^{\frac{1}{p}}$$

Remarquons que si $p = 2$ on a là la norme euclidienne usuelle du vecteur x .

Commençons par quelques propriétés faciles.

Proposition 14. *Pour tous $x \in \mathbb{R}^n$ et $\lambda \in \mathbb{R}$,*

- $\|x\|_p \geq 0$.
- $\|x\|_p = 0 \implies x = 0$.
- $\|\lambda x\|_p = |\lambda| \|x\|_p$.

Démonstration. Sans difficulté. $\square$

Ainsi, $\|\cdot\|_p$ possède des propriétés propres aux normes. Il reste évidemment à prouver l'inégalité triangulaire. C'est ce que nous allons faire dans le reste de la section. Cette inégalité s'appelle *l'inégalité de Minkowski*.

Proposition 15. *Pour tous $x, y \in \mathbb{R}^n$,*

$$\|x + y\|_p \leq \|x\|_p + \|y\|_p$$

4.2 Une fonction convexe

Soit $\varphi : [0, 1] \rightarrow \mathbb{R}$ définie pour tout $x \in [0, 1]$ par

$$\varphi(x) = \left(1 - x^{\frac{1}{p}}\right)^p$$

La fonction φ est continue sur $[0, 1]$, dérivable sur $]0, 1[$, et pour tout $x \in]0, 1[$

$$\varphi'(x) = -x^{\frac{1}{p}-1} \left(1 - x^{\frac{1}{p}}\right)^{p-1}$$

Comme $\frac{1}{p} - 1 \leq 0$, la fonction $x \mapsto x^{\frac{1}{p}-1}$ est décroissante. Comme $p \geq 1$, la fonction $x \mapsto \left(1 - x^{\frac{1}{p}}\right)^{p-1}$ est croissante. Ces deux fonctions sont de plus positives. Il en résulte que φ' croît sur $]0, 1[$. Ainsi, φ est convexe sur $[0, 1]$.

4.3 Inégalité de convexité

Proposition 16. Soient $a_1, \dots, a_n, b_1, \dots, b_n$ $2n$ réels positifs. On a

$$\left(\sum_{i=1}^n (a_i + b_i)^p \right)^{\frac{1}{p}} \leq \left(\sum_{i=1}^n a_i^p \right)^{\frac{1}{p}} + \left(\sum_{i=1}^n b_i^p \right)^{\frac{1}{p}}$$

Démonstration. Quitte à changer la valeur de n , on peut supposer que pour chaque i , l'un des deux réels a_i et b_i est non nul (en effet, si les deux sont nuls, les termes correspondants dans les sommes mises en jeu n'apparaissent pas).

Posons

$$\begin{aligned} S &= \sum_{i=1}^n (a_i + b_i)^p \\ S_1 &= \sum_{i=1}^n a_i^p \\ S_2 &= \sum_{i=1}^n b_i^p \end{aligned}$$

Pour tout $i \in \llbracket 1, n \rrbracket$, posons

$$X_i = \frac{a_i^p}{(a_i + b_i)^p}, \quad Y_i = \frac{b_i^p}{(a_i + b_i)^p}, \quad t_i = \frac{1}{S}(a_i + b_i)^p$$

Les X_i appartiennent à $[0, 1]$, et les t_i sont positifs et de somme 1. On peut donc appliquer l'inégalité de convexité généralisée à φ avec les X_i et les t_i :

$$\varphi \left(\sum_{i=1}^n t_i X_i \right) \leq \sum_{i=1}^n t_i \varphi(X_i)$$

Calculons séparément les différentes quantités apparaissant dans cette inégalité. Tout d'abord,

$$\sum_{i=1}^n t_i X_i = \frac{1}{S} \sum_{i=1}^n a_i^p = \frac{S_1}{S}$$

Puis

$$\begin{aligned} \varphi \left(\sum_{i=1}^n t_i X_i \right) &= \varphi \left(\frac{S_1}{S} \right) \\ &= \left(1 - \frac{S_1^{\frac{1}{p}}}{S^{\frac{1}{p}}} \right)^p \\ &= \frac{1}{S} \left(S^{\frac{1}{p}} - S_1^{\frac{1}{p}} \right)^p \end{aligned}$$

Poursuivons :

$$\varphi(X_i) = \frac{b_i^p}{(a_i + b_i)^p} = Y_i$$

Et enfin (calculs identiques à ce qui précède) :

$$\sum_{i=1}^n t_i \varphi(X_i) = \frac{S_2}{S}$$

Réunissant le tout, on obtient par l'inégalité de convexité

$$\frac{1}{S} \left(S^{\frac{1}{p}} - S_1^{\frac{1}{p}} \right)^p \leq \frac{S_2}{S}$$

Multipliant par S puis élevant à la puissance $\frac{1}{p}$, on obtient

$$S^{\frac{1}{p}} \leq S_1^{\frac{1}{p}} + S_2^{\frac{1}{p}}$$

qui est l'inégalité voulue. $\square$

4.4 Preuve de l'inégalité de Minkowski

Nous en avons quasiment terminé. Soient $x = (x_1, \dots, x_n)$ et $y = (y_1, \dots, y_n)$ deux éléments de $\mathbb{R}^n$. Il suffit de poser, pour $i \in \llbracket 1, n \rrbracket$, $a_i = |x_i|$ et $b_i = |y_i|$, puis d'appliquer le résultat précédent, pour obtenir que

$$\|x + y\|_p \leq \|x\|_p + \|y\|_p$$

4.5 L'inégalité de Hölder

Rappelons que le produit scalaire des vecteurs x et y de $\mathbb{R}^n$ est

$$\langle x, y \rangle = \sum_{i=1}^n x_i y_i$$

Ce produit scalaire vérifie l'inégalité de Schwarz :

$$|\langle x, y \rangle| \leq \|x\|_2 \|y\|_2$$

Cette inégalité possède-t-elle une analogue avec des p -normes ? La réponse est oui, il s'agit de l'inégalité de Hölder.

Soit p un nombre réel, $p > 1$. Soit q défini par l'égalité

$$\frac{1}{p} + \frac{1}{q} = 1$$

Remarquons que le réel q vérifie également $q > 1$.

Commençons par un lemme.

Lemme 17. Soient a et b deux réels positifs. On a

$$ab \leq \frac{a^p}{p} + \frac{b^q}{q}$$

Démonstration. Si a ou b est nul cette inégalité est évidente. Supposons donc a et b strictement positifs. La fonction $\ln$ est concave sur $\mathbb{R}_+^*$. Comme $\frac{1}{p} + \frac{1}{q} = 1$, l'inégalité de concavité donne donc

$$\ln \left(\frac{1}{p} a^p + \frac{1}{q} b^q \right) \geq \frac{1}{p} \ln a^p + \frac{1}{q} \ln b^q = \ln(ab)$$

d'où le résultat en passant à l'exponentielle. $\square$

Proposition 18. [Hölder] Pour tous $x, y \in \mathbb{R}^n$,

$$| \langle x, y \rangle | \leq \|x\|_p \|y\|_q$$

Remarque. Dans le cas où $p = 2$, on a aussi $q = 2$ et cette inégalité n'est autre que l'inégalité de Schwarz. Ainsi, l'inégalité de Hölder est une généralisation de l'inégalité de Schwarz.

Démonstration. Si x ou y est nul, l'inégalité est évidente. Supposons donc x et y non nuls. Ensuite, quitte à diviser x par $\|x\|_p$ et y par $\|y\|_q$, on peut supposer que $\|x\|_p = \|y\|_q = 1$.

Pour tout $i \in \llbracket 1, n \rrbracket$, le lemme nous dit que

$$|x_i| |y_i| \leq \frac{1}{p} |x_i|^p + \frac{1}{q} |y_i|^q$$

Sommons le tout. Nous obtenons

$$\sum_{i=1}^n |x_i y_i| \leq \frac{1}{p} \|x\|_p^p + \frac{1}{q} \|y\|_q^q = \frac{1}{p} + \frac{1}{q} = 1 = \|x\|_p \|y\|_q$$

Comme le membre de gauche de cette inégalité est clairement supérieur à $| \langle x, y \rangle | = | \sum_{i=1}^n x_i y_i |$, on obtient bien l'inégalité de Hölder. $\square$

4.6 Le cas $p = 1$

Lorsque $p = 1$, le réel q n'existe pas puisqu'on devrait avoir $\frac{1}{q} = 0$. Ou alors il faudrait se lancer dans l'hypothèse hasardeuse que $q = +\infty$. Alors lançons nous. Le problème est de

définir ce que pourrait être $\|x\|_\infty$. Nous discuterons dans le paragraphe suivant du choix que nous allons maintenant faire.

Définition 5. Pour tout $x = (x_1, \dots, x_n) \in \mathbb{R}^n$, on pose

$$\|x\|_\infty = \max_{1 \leq i \leq n} |x_i|$$

Proposition 19. Pour tous $x, y \in \mathbb{R}^n$,

$$| \langle x, y \rangle | \leq \|x\|_1 \|y\|_\infty$$

Démonstration. On a

$$\begin{aligned} | \langle x, y \rangle | &= \left| \sum_{i=1}^n x_i y_i \right| \\ &\leq \sum_{i=1}^n |x_i| |y_i| \end{aligned}$$

Pour tout $i \in \llbracket 1, n \rrbracket$, $|y_i| \leq \|y\|_\infty$. La somme ci-dessus est donc majorée par

$$\|y\|_\infty \sum_{i=1}^n |x_i| = \|x\|_1 \|y\|_\infty$$

□

4.7 Une justification

Soit $x \in \mathbb{R}^n$. Que vaut $\lim_{p \rightarrow \infty} \|x\|_p$? Il s'avère que non seulement cette limite existe, mais qu'elle est précisément égale à $\dots \|x\|_\infty$.

Proposition 20. Pour tout $x \in \mathbb{R}^n$,

$$\|x\|_p \xrightarrow{p \rightarrow \infty} \|x\|_\infty$$

Démonstration. Si $x = 0$, c'est évident. Supposons donc $x \neq 0$. Soit $i_0 \in \llbracket 1, n \rrbracket$ tel que $|x_{i_0}| = \|x\|_\infty$. Comme $x \neq 0$, on a $|x_{i_0}| > 0$. Pour tout réel $p \geq 1$,

$$\|x\|_p^p = \sum_{i=1}^n |x_i|^p = |x_{i_0}|^p \sum_{i=1}^n \mu_i^p$$

où $0 \leq \mu_i = \left| \frac{x_i}{x_{i_0}} \right| \leq 1$. On a donc l'encadrement

$$|x_{i_0}|^p \leq \|x\|_p^p \leq n|x_{i_0}|^p$$

d'où

$$|x_{i_0}| \leq \|x\|_p \leq n^{\frac{1}{p}} |x_{i_0}|$$

Lorsque p tend vers l'infini, $n^{\frac{1}{p}}$ tend vers 1. Ainsi, par le théorème d'encadrement des limites, $\|x\|_p$ tend vers $|x_{i_0}| = \|x\|_\infty$.

□

Chapitre 15

Primitives et Équations différentielles

1 Calculs de primitives

1.1 Notion de primitive

Dans tout ce qui suit, $\mathbb{K}$ désigne $\mathbb{R}$ ou $\mathbb{C}$. Lorsque ce n'est pas précisé, I désigne un intervalle de $\mathbb{R}$.

Définition 1. Soit $f : I \rightarrow \mathbb{K}$ une fonction définie sur un intervalle I . Soit $F : I \rightarrow \mathbb{K}$. On dit que F est une primitive de f sur I lorsque F est dérivable sur I , et $F' = f$.

Exemple. Soit $\alpha \in \mathbb{C} \setminus \{-1\}$. Une primitive de $x \mapsto x^\alpha$ sur $\mathbb{R}_+^*$ (ou plus selon les valeurs de α) est

$$x \mapsto \frac{x^{\alpha+1}}{\alpha+1}$$

Une primitive sur $\mathbb{R}_+^*$ de $x \mapsto \frac{1}{x}$ est $x \mapsto \ln x$.

Une primitive sur $\mathbb{R}_-^*$ de $x \mapsto \frac{1}{x}$ est $x \mapsto \ln(-x)$.

Nous admettons le théorème suivant. Il sera démontré dans le chapitre d'intégration.

Proposition 1. Soit $f : I \rightarrow \mathbb{K}$ une fonction continue. La fonction f admet des primitives sur I . Si F_0 est une primitive de f sur I , alors les primitives de f sur I sont les fonctions $F = F_0 + c$ où $c \in \mathbb{K}$.

Démonstration. Nous pouvons tout de même, avec ce que nous savons, démontrer « l'unicité ». Soit $F : I \rightarrow \mathbb{K}$ une fonction dérivable. Alors F est une primitive de f sur I si et seulement si

$$F' = f = F_0'$$

c'est à dire

$$(F - F_0)' = 0$$

ou encore, si et seulement si $F - F_0$ est constante sur I . $\square$

Notation. Soit $f : I \rightarrow \mathbb{K}$ continue. Soit $F : I \rightarrow \mathbb{K}$ une primitive de f sur I . Pour tous $a, b \in I$, on note

$$\int_a^b f(x) dx = F(b) - F(a)$$

ou $\int_a^b f(t) dt$ (ou la lettre que l'on voudra), ou enfin plus simplement $\int_a^b f$. Cette quantité ne dépend pas de la primitive F choisie.

Nous noterons également de façon abusive

$$\int^x f(t)dt = F(x)$$

où F est une primitive quelconque de f sur l'intervalle I .

Remarque 1. Nous verrons plus tard que tout ceci est plus qu'une notation, mais il faudra d'abord avoir étudié la notion d'intégrale.

Notons également la *formule de Chasles*, qui est évidente.

Proposition 2. Soit $f : I \rightarrow \mathbb{K}$ une fonction continue. Soient $a, b, c \in I$. On a

$$\int_a^c f(x) dx = \int_a^b f(x) dx + \int_b^c f(x) dx$$

Démonstration. En effet, en appelant F une primitive de f sur I , on a

$$F(c) - F(a) = F(b) - F(a) + F(c) - F(b)$$

□

1.2 Primitives usuelles

Le tableau ci-dessous résume les primitives à connaître. On trouve dans la première colonne la fonction à primitiver. Dans la deuxième colonne sont indiqués le ou les intervalles (les plus grands possibles) où la fonction possède des primitives. La dernière colonne contient l'expression d'une primitive de la fonction, valable sur tous les intervalles cités dans la colonne précédente. Pour avoir toutes les primitives, il convient bien entendu de rajouter des constantes.

Fonction f	Intervalle(s) I	Primitive $\int^x f(t)dt$ sur I
e^x	$\mathbb{R}$	e^x
$\operatorname{ch} x$	$\mathbb{R}$	$\operatorname{sh} x$
$\operatorname{sh} x$	$\mathbb{R}$	$\operatorname{ch} x$
$\cos x$	$\mathbb{R}$	$\sin x$
$\sin x$	$\mathbb{R}$	$-\cos x$
$x^\alpha (\alpha \in \mathbb{R} \setminus \{-1\})$	$\mathbb{R}_+^*$ (au moins)	$\frac{x^{\alpha+1}}{\alpha+1}$
$\frac{1}{x}$	$\mathbb{R}_+^*$ ou $\mathbb{R}_-^*$	$\ln x $
$\ln x $	$\mathbb{R}_+^*$ ou $\mathbb{R}_-^*$	$x \ln x - x$
$\frac{1}{\operatorname{ch}^2 x} = 1 - \operatorname{th}^2 x$	$\mathbb{R}$	$\operatorname{th} x$
$\frac{1}{\operatorname{sh}^2 x} = \coth^2 x - 1$	$\mathbb{R}_+^*$ ou $\mathbb{R}_-^*$	$-\coth x$
$\frac{1}{\cos^2 x} = 1 + \tan^2 x$	$] -\frac{\pi}{2} + k\pi, \frac{\pi}{2} + k\pi[, k \in \mathbb{Z}$	$\tan x$
$\frac{1}{\sin^2 x} = 1 + \cotan^2 x$	$]k\pi, (k+1)\pi[, k \in \mathbb{Z}$	$-\cotan x$
$\tan x$	$] -\frac{\pi}{2} + k\pi, \frac{\pi}{2} + k\pi[, k \in \mathbb{Z}$	$-\ln \cos x $
$\cotan x$	$]k\pi, (k+1)\pi[, k \in \mathbb{Z}$	$\ln \sin x $
$\tanh x$	$\mathbb{R}$	$\ln(\cosh x)$
$\coth x$	$\mathbb{R}_+^*$ ou $\mathbb{R}_-^*$	$\ln \sinh x $
$\frac{1}{\operatorname{ch} x}$	$\mathbb{R}$	$2 \arctan e^x$
$\frac{1}{\operatorname{sh} x}$	$\mathbb{R}_+^*$ ou $\mathbb{R}_-^*$	$\ln \left \operatorname{th} \frac{x}{2} \right $
$\frac{1}{\cos x}$	$] -\frac{\pi}{2} + k\pi, \frac{\pi}{2} + k\pi[, k \in \mathbb{Z}$	$\ln \left \tan \left(\frac{x}{2} + \frac{\pi}{4} \right) \right $
$\frac{1}{\sin x}$	$]k\pi, (k+1)\pi[, k \in \mathbb{Z}$	$\ln \left \tan \frac{x}{2} \right $
$\frac{1}{x^2+a^2} (a \in \mathbb{R}_+^*)$	$\mathbb{R}$	$\frac{1}{a} \arctan \frac{x}{a}$
$\frac{1}{a^2-x^2} (a \in \mathbb{R}_+^*)$	$\mathbb{R} \setminus \{-a, a\}$	$\frac{1}{2a} \ln \left \frac{x+a}{x-a} \right $
$\frac{1}{\sqrt{a^2-x^2}} (a \in \mathbb{R}_+^*)$	$] -a, a[$	$\arcsin \frac{x}{a}$
$\frac{1}{\sqrt{x^2+h}} (h \in \mathbb{R}^*)$	tout intervalle où $x^2+h \neq 0$	$\ln \left x + \sqrt{x^2+h} \right $

Exercice. Soient $a, b \in \mathbb{R}$. Quelles sont les primitives sur $\mathbb{R}$ des fonctions $x \mapsto e^{ax} \cos bx$ et $x \mapsto e^{ax} \sin bx$?

Posons $\alpha = a + ib$. Considérons la fonction $f : \mathbb{R} \rightarrow \mathbb{C}$ définie pour tout réel x par $f(x) = e^{\alpha x}$. Une primitive de f sur $\mathbb{R}$ est la fonction $F : \mathbb{R} \rightarrow \mathbb{C}$ définie par

$$F(x) = \frac{1}{\alpha} e^{\alpha x}$$

Il suffit, pour répondre à la question posée, de prendre la partie réelle et la partie imaginaire de F . Pour tout $x \in \mathbb{R}$, on a

$$\begin{aligned}
 F(x) &= \frac{1}{a+ib} e^{ax} (\cos bx + i \sin bx) \\
 &= \frac{a-ib}{a^2+b^2} e^{ax} (\cos bx + i \sin bx) \\
 &= \frac{1}{a^2+b^2} e^{ax} ((a \cos bx + b \sin bx) + i(a \sin bx - b \cos bx))
 \end{aligned}$$

Ainsi,

$$\int^x e^{at} \cos bt \, dt = \frac{e^{ax}}{a^2 + b^2} (a \cos bx + b \sin bx)$$

et

$$\int^x e^{at} \sin bt \, dt = \frac{e^{ax}}{a^2 + b^2} (a \sin bx - b \cos bx)$$

1.3 Intégration par parties

Proposition 3. Soient $f, g \in \mathcal{C}^1(I)$. Soient $a, b \in I$. On a

$$\int_a^b f'g = f(b)g(b) - f(a)g(a) - \int_a^b fg'$$

ce que l'on note plus simplement

$$\int_a^b f'g = [fg]_a^b - \int_a^b fg'$$

Démonstration. La fonction fg est une primitive sur I de la fonction $(fg)'$, continue sur I . On a donc

$$\int_a^b (fg)' = [fg]_a^b$$

Mais

$$\int_a^b (fg)' = \int_a^b (f'g + fg') = \int_a^b f'g + \int_a^b fg'$$

□

Exemple. Pour tout $n \in \mathbb{N}$, notons

$$I_n(x) = \int_0^x t^n e^{-t} \, dt$$

Soit $n \geq 1$. On intègre par parties en posant $f(t) = t^n$ et $g'(t) = e^{-t}$. Il vient

$$I_n(x) = -[t^n e^{-t}]_0^x + nI_{n-1}(x) = -e^{-x} + nI_{n-1}(x)$$

Exercice. Calculer $I_n(x)$ sous forme d'une somme puis trouver la limite de $I_n(x)$ lorsque x tend vers $+\infty$.

Exercice. Déterminer les primitives sur $[-1, 1]$ de la fonction arc sinus.

Exercice. Calculer, pour tout $n \in \mathbb{N}$ et tout $x \geq 1$,

$$\int_1^x t^n \ln t \, dt$$

1.4 Changement de variable

Proposition 4. Soient I et J deux intervalles de $\mathbb{R}$. Soit $f : I \rightarrow \mathbb{R}$ (ou $\mathbb{C}$) une fonction continue sur I à valeurs réelles ou complexes. Soit $\varphi : J \rightarrow I$ une surjection de classe $\mathcal{C}^1$. Soient $a, b \in I$. Soient $\alpha, \beta \in J$ tels que $\varphi(\alpha) = a$ et $\varphi(\beta) = b$. On a alors

$$\int_a^b f(x) dx = \int_\alpha^\beta f(\varphi(t))\varphi'(t) dt$$

Démonstration. Soit F une primitive de f sur l'intervalle I . Une telle primitive existe car f est continue. On a

$$\int_\alpha^\beta f(\varphi(t))\varphi'(t) dt = \int_\alpha^\beta F'(\varphi(t))\varphi'(t) dt = \int_\alpha^\beta (F \circ \varphi)'(t) dt = F \circ \varphi|_\alpha^\beta$$

Mais cette dernière quantité est aussi

$$[F]_{\varphi(\alpha)}^{\varphi(\beta)} = [F]_a^b = \int_a^b f(x) dx$$

□

Pratiquement, on procède comme suit. On désire calculer $\int_a^b f(x) dx$.

- On pose $x = \varphi(t)$.
- On calcule $dx = \varphi'(t) dt$.
- On remarque que si $t = \alpha$ (resp β) alors $\varphi(t) = a$ (resp b).
- On recolle tous les morceaux.

Exercice. Calculer $\int_1^{e^{\pi/2}} \sin \ln x dx$.

Exercice. Calculer $\int_0^4 e^{\sqrt{x}} dx$.

Exercice. Soit $n \in \mathbb{N}$. Montrer que $\int_0^{\frac{\pi}{2}} \sin^n x dx = \int_0^{\frac{\pi}{2}} \cos^n x dx$.

Remarque 2. Il nous arrivera de faire des changements de variable « à l'envers », c'est à dire d'avoir une intégrale $\int_\alpha^\beta g(t) dt$ et d'avoir envie de poser $x = \varphi(t)$. Cela est-il possible ? Oui, si on arrive aussi à « voir » que $g(t) = f(\varphi(t))\varphi'(t)$ où f est une fonction judicieuse. Dans ce cas,

$$\begin{aligned} \int_\alpha^\beta g(t) dt &= \int_\alpha^\beta f(\varphi(t))\varphi'(t) dt \\ &= \int_{\varphi(\alpha)}^{\varphi(\beta)} f(x) dx \end{aligned}$$

Prenons un exemple. Calculons $\int_0^{\frac{\pi}{4}} \frac{dt}{\cos t}$. Le changement de variable « inversé » $\varphi(t) = \sin t$ est particulièrement adapté, en remarquant que

$$\frac{1}{\cos t} = \frac{\cos t}{\cos^2 t} = \frac{\cos t}{1 - \sin^2 t}$$

et que le numérateur $\cos t$ est la dérivée de $\sin t$. Ainsi,

$$\begin{aligned} \int_0^{\frac{\pi}{4}} \frac{dt}{\cos t} &= \int_0^{\frac{\pi}{4}} \frac{\cos t dt}{1 - \sin^2 t} \\ &= \int_0^{\frac{\sqrt{2}}{2}} \frac{dx}{1 - x^2} \end{aligned}$$

Nous avons là une primitive usuelle.

$$\int_0^{\frac{\pi}{4}} \frac{dt}{\cos t} = \frac{1}{2} \ln \left| \frac{1+x}{1-x} \right|_0^{\frac{\sqrt{2}}{2}} = \frac{1}{2} \ln \frac{2+\sqrt{2}}{2-\sqrt{2}}$$

Remarquons que nous aurions aussi pu remarquer que $\frac{1}{\cos x}$ est dans la liste des primitives usuelles.

$$\int_0^{\frac{\pi}{4}} \frac{dt}{\cos t} = \ln \left| \tan \left(\frac{x}{2} + \frac{\pi}{4} \right) \right|_0^{\frac{\pi}{4}} = \ln \left| \tan \left(\frac{3\pi}{8} \right) \right|$$

1.5 Primitives des fractions rationnelles

Nous aborderons de façon plus complète ce problème lorsque nous aurons traité le chapitre sur les fractions rationnelles. Nous y verrons que toute fraction rationnelle, c'est à dire toute fonction du type $\frac{P(x)}{Q(x)}$ où P et Q sont des fonctions polynômes, s'écrit comme une somme de fractions particulières, appelées éléments simples. Nous calculerons alors des primitives de tous les éléments simples. Pour l'instant, concentrons nous sur les cas faciles.

- Primitives de $x \mapsto \frac{1}{x-a}$, où $a \in \mathbb{R}$. C'est immédiat, $\ln |x-a|$ est une telle primitive, sur $] -\infty, a[$ et sur $]a, +\infty[$.
- Primitives de $x \mapsto \frac{px+q}{x^2+bx+c}$ où $p, q, b, c \in \mathbb{R}$. Ici, c'est plus délicat. Commençons par remarquer que

$$x^2 + bx + c = \left(x + \frac{b}{2} \right)^2 + c'$$

Le changement de variable $x = t - \frac{b}{2}$ permet alors de se ramener au calcul de primitives de

$$t \mapsto \frac{pt+q}{t^2+c}$$

où p, q, c ne sont plus les mêmes que précédemment (pour simplifier l'écriture). Il suffit évidemment pour cela de savoir intégrer $\frac{t}{t^2+c}$ et $\frac{1}{t^2+c}$. La première primitive est évidente, car on reconnaît dans la fonction à intégrer une dérivée :

$$\int^t \frac{u \, du}{u^2 + c} = \frac{1}{2} \ln |t^2 + c|$$

sur des intervalles qui dépendent, en particulier, du signe de c . Concernant la seconde primitive, se ramener au tableau des primitives usuelles. La discussion se fait sur la nullité et le signe de c . Pour simplifier, prenons par exemple $c = 0, 1, -1$.

- Sur $] -\infty, -1[$, $] -1, 1[$ et $] 1, +\infty[$,

$$\int^t \frac{du}{u^2 - 1} = -\frac{1}{2} \ln \left| \frac{1+t}{1-t} \right|$$

- Sur $] -\infty, 0[$ et $] 0, +\infty[$,

$$\int^t \frac{du}{u^2} = -\frac{1}{t}$$

- Sur $\mathbb{R}$,

$$\int^t \frac{du}{u^2 + 1} = \arctan t$$

2 Notion d'équation différentielle

2.1 C'est quoi ?

Une équation différentielle est une équation où l'inconnue est une fonction, définie sur un intervalle I de $\mathbb{R}$, à valeurs dans $\mathbb{R}$ ou $\mathbb{C}$, et régulière (dérivable un certain nombre de fois). L'équation différentielle nous donne l'existence d'une relation entre la fonction inconnue, la variable (x), et les dérivées de la fonction jusqu'à un certain ordre (ce que l'on appelle l'ordre de l'équation différentielle).

Une équation différentielle est donc quelque chose du genre :

$$(\mathcal{E}) \quad F(x, y, y', y'', \dots, y^{(n)}) = 0$$

Résoudre $(\mathcal{E})$, c'est trouver tous les intervalles I et toutes les fonctions $f : I \rightarrow \mathbb{R}$ ou $\mathbb{C}$, n fois dérivables sur I , telles que

$$\forall x \in I, F(x, f(x), f'(x), f''(x), \dots, f^{(n)}(x)) = 0$$

On ne sait pas en général résoudre les équations différentielles. Dans certains cas, il existe des théorèmes permettant d'affirmer l'existence de solutions, voire leur unicité lorsque l'on impose des *conditions initiales* ou des *conditions aux limites*. Nous nous intéresserons dans les sections suivantes uniquement à des équations d'un type très particulier, dites *linéaires*.

2.2 Exemple

Considérons l'équation

$$(\mathcal{E}) \quad y' = y$$

Soit $f : \mathbb{R} \rightarrow \mathbb{R}$ dérivable. Posons, pour tout réel x ,

$$g(x) = e^{-x}f(x)$$

g est dérivable sur $\mathbb{R}$, et pour tout réel x , $f(x) = e^x g(x)$. De là,

$$f'(x) = e^x g(x) + e^x g'(x)$$

Ainsi, f est solution de $\mathcal{E}$ sur $\mathbb{R}$ si et seulement si pour tout réel x on a

$$e^x g(x) + e^x g'(x) = e^x g(x)$$

c'est à dire si et seulement si $g' = 0$ c'est à dire g est constante. Les solutions de $(\mathcal{E})$ sur $\mathbb{R}$ sont donc les multiples de la fonction exponentielle.

2.3 Exemple

Considérons l'équation

$$(\mathcal{E}) \quad y' = 2\sqrt{y}$$

Cherchons les solutions de $(\mathcal{E})$ sur $\mathbb{R}$.

Soit $f : \mathbb{R} \rightarrow \mathbb{R}$. La fonction f est solution de $(\mathcal{E})$ si et seulement si $f \in \mathcal{D}^1(\mathbb{R})$, $f \geq 0$, et

$$\forall x \in \mathbb{R}, f'(x) = 2\sqrt{f(x)}$$

Soit f une solution de $(\mathcal{E})$.

Remarquons que $f' \geq 0$ et donc que f est croissante.

Si f s'annule en un réel a , alors, puisque f est positive et croissante, f est nulle sur $]-\infty, a]$. Et si f est non nulle en a , alors $f > 0$ sur $[a, +\infty[$, toujours par la croissance.

Soit $E = \{x \in \mathbb{R}, f(x) = 0\}$. Les remarques précédentes montrent que les seules possibilités pour E sont (1) $\emptyset$, les intervalles (2) $]-\infty, a[$ ou (3) $]-\infty, a]$ avec a réel, et (4) $\mathbb{R}$.

Le cas 2 est à exclure, car si $f(x) = 0$ pour tout $x < a$, alors, par continuité de f , $f(a) = 0$. Le cas 4 est évident : la fonction nulle est solution de $(\mathcal{E})$. Regardons donc les cas 1 et 3.

sur $I = \mathbb{R}$ ou $]a, +\infty]$, f ne s'annule pas, et $f' = 2\sqrt{f}$, ou encore $(\sqrt{f})' = 1$. Il existe donc un réel b tel que pour tout $x \in I$, $f(x) = (x - b)^2$. Mais cette fonction n'est croissante et non nulle que sur $]b, +\infty[$. On voit donc que le cas 1 est aussi à exclure, et que $b \leq a$. Mieux, f est nulle à gauche de a , et dérivable en a , donc $f'(a) = 0$. Mais $f'_d(a) = 2(a - b)$, donc $b = a$. Finalement, $f(x) = 0$ si $x \leq a$ et $f(x) = (x - a)^2$ si $x \geq a$.

Pour résumer notre analyse, si f est solution de $(\mathcal{E})$ sur $\mathbb{R}$, on est dans l'un des deux cas suivants :

- f est nulle sur $\mathbb{R}$, ou alors
- Il existe $a \in \mathbb{R}$ tel que f est nulle sur $] -\infty, a]$ et, pour tout $x \geq a$, $f(x) = (x - a)^2$.

Inversement, on vérifie que ces fonctions sont bien solutions de $\mathcal{E}$.

Remarque 3. La résolution de notre second exemple, pourtant très simple, nous a donné plus de mal que celle du premier exemple. Pour quoi ? Parce que l'équation différentielle $y' = y$ fait partie de la vaste famille des « équations différentielles linéaires ». Il existe pour cette famille d'équations des méthodes de résolution spécifiques. Leurs ensembles de solutions possèdent des structures algébriques intéressantes, etc. Nous allons dans ce qui suit nous intéresser aux équations différentielles linéaires d'ordre 1, et à certaines équations différentielles linéaires d'ordre 2.

3 Équations linéaires du premier ordre

Dans ce qui suit, $\mathbb{K} = \mathbb{R}$ ou $\mathbb{C}$.

3.1 Équation homogène

Proposition 5. Soit $a : I \rightarrow \mathbb{K}$ une fonction continue. Les solutions sur I de l'équation différentielle

$$(\mathcal{H}) \quad y' + a(x)y = 0$$

sont les fonctions

$$\varphi : x \mapsto k \exp(-A(x))$$

où A est une primitive sur I de a et $k \in \mathbb{K}$.

Démonstration. Soit $\varphi_0 : x \mapsto \exp(-A(x))$. La fonction φ_0 est dérivable sur I , et on a pour tout $x \in I$,

$$\varphi_0'(x) = -A'(x)\varphi_0(x) = -a(x)\varphi_0(x)$$

La fonction φ_0 est donc solution de $(\mathcal{H})$, et elle ne s'annule pas. Soit maintenant une fonction φ dérivable sur I . On écrit $\varphi = g\varphi_0$, où g est dérivable sur I . Alors φ est solution de $(\mathcal{H})$ si et seulement si

$$\varphi' + a\varphi = 0$$

c'est à dire

$$g'\varphi_0 + g\varphi_0' + a\varphi_0 = 0$$

Mais φ_0 étant solution de $(\mathcal{H})$, ceci équivaut à $g'\varphi_0 = 0$, ou encore $g' = 0$ puisque φ_0 ne s'annule pas. En d'autres termes, si et seulement si g est constante. $\square$

Corollaire 6. *En reprenant les notations de la preuve précédente, l'ensemble $\mathcal{S}_{\mathcal{H}}$ des solutions de $(\mathcal{H})$ sur I est la droite vectorielle engendrée par φ_0 .*

$$\mathcal{S}_{\mathcal{H}} = \{k\varphi_0, k \in \mathbb{K}\}$$

Exercice. Résoudre $(x+1)y' - xy + 1 = 0$ sur $] -\infty, -1[$, $] -1, +\infty[$, puis sur $\mathbb{R}$.

3.2 Équation avec second membre

Proposition 7. *Soient $a, b : I \rightarrow \mathbb{K}$ deux fonctions continues. Les solutions sur I de l'équation différentielle*

$$(\mathcal{E}) \quad y' + a(x)y = b(x)$$

sont les fonctions

$$\varphi : x \mapsto k \exp(-A(x)) + g(x) \exp(-A(x))$$

où A est une primitive sur I de la fonction a , $k \in \mathbb{K}$ et g est une primitive sur I de $x \mapsto b(x) \exp A(x)$

Démonstration. Supposons trouvée une solution φ_1 de $(\mathcal{E})$. Alors, pour toute fonction φ , φ est solution de $(\mathcal{E})$ si et seulement si $\varphi - \varphi_1$ est solution de $(\mathcal{H})$. Il suffit donc de trouver *une* solution de $(\mathcal{E})$ pour avoir *toutes* les solutions de $(\mathcal{H})$.

Pour trouver une solution de $(\mathcal{E})$ nous allons utiliser une technique qui s'appelle la *variation de la constante*. Cherchons *une* solution de $(\mathcal{E})$ de la forme

$$\varphi_1(x) = g(x)\varphi_0(x)$$

où, rappelons le,

$$\varphi_0 : x \mapsto \exp(-A(x))$$

Un calcul analogue à celui fait pour l'équation homogène conduit à

$$g'(x)\varphi_0(x) = b(x)$$

Il suffit donc de prendre pour g une primitive de la fonction

$$x \mapsto b(x) \exp A(x)$$

□

Corollaire 8. *L'ensemble $\mathcal{S}_{\mathcal{E}}$ des solutions de $(\mathcal{E})$ sur I est une droite affine dont la direction est la droite vectorielle $\mathcal{S}_{\mathcal{H}}$*

Exercice. Résoudre l'équation différentielle $(x-1)y' + (x-2)y = x(x-1)^2$ sur $] -\infty, 1[$, $]1, +\infty[$, puis sur $\mathbb{R}$.

3.3 Un exemple : deux équations fonctionnelles

Proposition 9. *Les applications $f : \mathbb{R} \rightarrow \mathbb{C}$ dérivables vérifiant pour tous réels x, y ,*

$$(\mathcal{E}) \quad f(x+y) = f(x)f(y)$$

sont la fonction nulle et les applications $x \mapsto e^{\lambda x}$ où $\lambda \in \mathbb{C}$.

Démonstration. Soit f une solution de $(\mathcal{E})$. Remarquons tout d'abord que

$$f(0+0) = f(0)f(0)$$

ou encore

$$f(0) = f(0)^2$$

Ainsi, $f(0) = 0$ ou $f(0) = 1$. Si $f(0) = 0$ alors, pour tout réel x ,

$$f(x) = f(x+0) = f(x)f(0) = 0$$

et f est donc la fonction nulle. Supposons maintenant que f n'est pas la fonction nulle, et donc que $f(0) = 1$. On a pour tous réels x, y (en dérivant « par rapport à y »),

$$f'(x+y) = f(x)f'(y)$$

En particulier, en prenant $y = 0$, on a

$$f'(x) = f'(0)f(x)$$

ou encore, en posant $\lambda = f'(0)$,

$$f' = \lambda f$$

On en déduit qu'il existe $k \in \mathbb{C}$ tel que, pour tout réel x ,

$$f(x) = ke^{\lambda x}$$

Comme $f(0) = 1$, on a $k = 1$.

Inversement, la fonction nulle et les fonctions $x \mapsto e^{\lambda x}$, où $\lambda \in \mathbb{C}$, sont bien solutions de $(\mathcal{E})$. $\square$

Proposition 10. Les applications $f : \mathbb{R}_+^* \rightarrow \mathbb{R}$ dérivables vérifiant pour tous $x, y \in \mathbb{R}_+^*$,

$$f(xy) = f(x) + f(y)$$

sont les applications $x \mapsto \lambda \ln x$ où $\lambda \in \mathbb{R}$.

Démonstration. Nous pourrions utiliser la même méthode que précédemment mais cherchons plutôt à nous *ramener* à ce qui précède.

Soit f une telle application. Posons pour tout réel x ,

$$g(x) = \exp(f(e^x))$$

La fonction g est dérivable sur $\mathbb{R}$ et on a pour tous réels x, y ,

$$\begin{aligned} g(x+y) &= \exp(f(e^{x+y})) \\ &= \exp(f(e^x e^y)) \\ &= \exp(f(e^x) + f(e^y)) \\ &= \exp(f(e^x)) \exp(f(e^y)) \\ &= g(x)g(y) \end{aligned}$$

On en déduit que g est solution de $(\mathcal{E})$. Ainsi, pour tout réel x , $g(x) = e^{\lambda x}$ avec $\lambda \in \mathbb{C}$ indépendant de x . La fonction g étant à valeurs réelles, $\lambda = g'(0)$ est aussi réel. On a donc, pour tout réel x ,

$$\exp(f(e^x)) = \exp(\lambda x)$$

d'où

$$f(e^x) = \lambda x$$

Soit maintenant $t \in \mathbb{R}_+^*$. Il existe un (unique) réel x tel que $t = e^x$, à savoir $x = \ln t$. On a alors

$$f(t) = \lambda \ln t$$

Inversement, une telle application convient. $\square$

3.4 Le problème de Cauchy

Soient $a, b : I \rightarrow \mathbb{R}$ deux fonctions continues sur un intervalle I . Soit l'équation différentielle

$$(\mathcal{E}) \quad y' + a(x)y = b(x)$$

On s'intéresse ici à la résolution de l'équation $(\mathcal{E})$ avec conditions initiales.

Proposition 11. Soit $x_0 \in I$ et $y_0 \in \mathbb{K}$. Il existe une unique solution φ de l'équation $(\mathcal{E})$ telle que $\varphi(x_0) = y_0$.

Démonstration. Soit φ une solution générale de $(\mathcal{E})$. Nous avons vu que

$$\varphi = \varphi_1 + k\varphi_0$$

où $k \in \mathbb{K}$ et la fonction φ_0 ne s'annule pas sur I . On a alors $\varphi(x_0) = y_0$ si et seulement si

$$\varphi_1(x_0) + k\varphi_0(x_0) = y_0$$

c'est à dire si et seulement si

$$k = \frac{y_0 - \varphi_1(x_0)}{\varphi_0(x_0)}$$

On en déduit l'existence et l'unicité de la fonction φ vérifiant la condition initiale. $\square$

4 Résolution approchée

La plupart des équations différentielles ne peuvent pas être résolues de façon explicite. Il existe néanmoins des méthodes numériques permettant d'en obtenir une solution approchée. Nous allons dans cette section parler d'une méthode permettant d'obtenir une valeur approchée d'une solution φ de l'équation différentielle

$$(\mathcal{E}) \quad y' = F(x, y)$$

avec condition initiale : la *méthode d'Euler*. Nous resterons à un niveau informel. Nous ne justifierons pas les résultats.

On se fixe un réel $\delta > 0$ supposé « petit ». On pose $x_n = x_0 + n\delta$, et $y_n = \varphi(x_n)$ pour tout entier n tel que $x_n \in I$.

On part de $\varphi(x_0) = y_0$. Le réel δ étant supposé « petit », on a

$$\frac{\varphi(x_0 + \delta) - \varphi(x_0)}{\delta} \approx \varphi'(x_0) = F(x_0, \varphi(x_0))$$

ou encore

$$\varphi(x_0 + \delta) \approx \varphi(x_0) + \delta F(x_0, y_0)$$

Ainsi,

$$y_1 \approx y_0 + \delta F(x_0, y_0)$$

Posons donc

$$\hat{y}_1 = y_0 + \delta F(x_0, y_0)$$

L'erreur commise est liée à l'identification entre le taux d'accroissement avec la dérivée. Rien n'empêche de poursuivre avec l'approximation obtenue et d'itérer le procédé. On pose donc $\hat{y}_0 = y_0$ et, pour tout n tel que $x_n \in I$,

$$\hat{y}_{n+1} = \hat{y}_n + \delta F(x_n, \hat{y}_n)$$

On peut alors montrer que sous certaines conditions, $\hat{y}_n$ est une approximation « intéressante » de y_n . La théorie permet en particulier d'estimer l'erreur commise en approchant y_n par $\hat{y}_n$.

Prenons un exemple. Considérons l'équation

$$(\mathcal{E})y' = y$$

avec $x_0 = 0$ et $y_0 = 1$, c'est à dire, avec les notations ci-dessus, $F(x, y) = y$. La solution φ est bien sûr $\varphi(x) = e^x$. La relation de récurrence devient

$$y_{n+1} = y_n + \delta y_n = (1 + \delta)y_n$$

c'est à dire que

$$y_n = (1 + \delta)^n$$

Particularisons encore, en posant $\delta = \frac{1}{N}$. On a alors

$$x_N = 1, \quad y_N = \left(1 + \frac{1}{N}\right)^N$$

La « vraie » valeur de y_N est bien sûr $e^1 = e$. On peut montrer (et on le fera) que

$$\left(1 + \frac{1}{N}\right)^N \xrightarrow{N \rightarrow \infty} e$$

lorsque N tend vers l'infini : si l'on diminue la taille du pas δ , on obtient a priori une meilleure approximation des solutions.

5 Équations du second ordre

Contrairement à l'ordre 1, on ne s'intéresse qu'à des équations à *coefficients constants*.

5.1 Équation homogène

Soient a, b, c trois nombres complexes, $a \neq 0$.

On considère dans ce paragraphe l'équation différentielle

$$(\mathcal{H}) \quad ay'' + by' + cy = 0$$

Nous noterons $\mathcal{C}$ l'équation caractéristique de $(\mathcal{H})$:

$$(\mathcal{C}) \quad aX^2 + bX + c = 0$$

et $\Delta = b^2 - 4ac$ le discriminant de cette équation.

Proposition 12.

- Si $\Delta \neq 0$, les solutions de $(\mathcal{H})$ sont les fonctions

$$x \longmapsto \lambda_1 \exp \alpha_1 x + \lambda_2 \exp \alpha_2 x$$

où λ_1 et λ_2 sont des nombres complexes, et α_1 et α_2 sont les deux racines distinctes de $(\mathcal{C})$.

- Si $\Delta = 0$, les solutions de $(\mathcal{H})$ sont les fonctions

$$x \longmapsto (\lambda_1 x + \lambda_2) \exp \alpha x$$

où λ_1 et λ_2 sont des nombres complexes, et α est l'unique racine de $(\mathcal{C})$.

Démonstration. Soit $\mu \in \mathbb{C}$. $f : x \longmapsto e^{\mu x}$. f est solution de $(\mathcal{H})$ si et seulement si μ est une racine de l'équation caractéristique. Soit maintenant $\varphi : \mathbb{R} \rightarrow \mathbb{C}$ deux fois dérivable sur $\mathbb{R}$. on pose

$$\varphi(x) = g(x)e^{\alpha x}$$

où α est une racine de $(\mathcal{C})$. La fonction φ est solution de $\mathcal{H}$ si et seulement si

$$a\varphi'' + b\varphi' + c\varphi = 0$$

ou encore

$$ag'' + (2a\alpha + b)g' + (a\alpha^2 + b\alpha + c)g = 0$$

Mais α est racine de $(\mathcal{C})$, donc g convient si et seulement si

$$ah' + (2a\alpha + b)h = 0$$

où $h = g'$. Ceci est une équation linéaire d'ordre 1 en h . Il s'agit maintenant de distinguer deux cas.

Première possibilité, $2a\alpha + b \neq 0$, c'est à dire que le discriminant de l'équation caractéristique est non nul. On obtient

$$h(x) = k \exp(\lambda x)$$

où $k \in \mathbb{C}$ et

$$\lambda = -\frac{2a\alpha + b}{a}$$

De là,

$$g(x) = K \exp(\lambda x) + K'$$

avec $K, K' \in \mathbb{C}$, et

$$\varphi(x) = K \exp((\lambda + \alpha)x) + K' \exp(\alpha x)$$

Mais $\lambda + \alpha = \beta$, où β est l'autre racine de l'équation caractéristique. Ainsi, les solutions de $\mathcal{H}$ sont les fonctions de la forme

$$\varphi(x) = K \exp(\alpha x) + K' \exp(\beta x)$$

Deuxième cas, $2a\alpha + b = 0$. On obtient

$$g(x) = Kx + K'$$

d'où

$$\varphi(x) = (Kx + K') \exp(\alpha x)$$

□

Corollaire 13. *L'ensemble $\mathcal{S}_{\mathcal{H}}$ des solutions de $(\mathcal{H})$ est un plan vectoriel, engendré par*

- *Si $\Delta \neq 0$: $x \mapsto e^{\alpha_1 x}$ et $x \mapsto e^{\alpha_2 x}$.*
- *Si $\Delta = 0$: $x \mapsto e^{\alpha x}$ et $x \mapsto x e^{\alpha x}$.*

Dit autrement, dans chacun des deux points ci-dessus, en appelant φ_1 et φ_2 les fonctions citées, l'ensemble des solutions de $(\mathcal{H})$ est

$$\mathcal{S}_{\mathcal{H}} = \{\lambda_1 \varphi_1 + \lambda_2 \varphi_2, \lambda_1, \lambda_2 \in \mathbb{C}\}$$

Exercice. Résoudre les équations différentielles $y'' + y = 0$, $y'' - y = 0$ et $y'' - 2y' + y = 0$.

5.2 Le cas réel

On suppose ici que $a, b, c \in \mathbb{R}$ et on cherche les solutions réelles de l'équation $(\mathcal{H})$. Il y a maintenant trois cas, selon que le discriminant de l'équation caractéristique est > 0 , nul ou < 0 . Les deux premiers cas donnent les mêmes solutions que dans le cas complexes (prendre évidemment K et K' réels).

Traisons le cas où $\Delta < 0$. Les racines de l'équation caractéristique sont $\lambda \pm i\omega$ où $\lambda \in \mathbb{R}$, $\omega \in \mathbb{R}^*$. Soit

$$\varphi(x) = K \exp((\lambda + i\omega)x) + K' \exp((\lambda - i\omega)x)$$

Cette fonction est à valeurs réelles si et seulement si, pour tout réel x ,

$$\varphi(x) = \overline{\varphi(x)}$$

c'est à dire

$$K \exp((\lambda + i\omega)x) + K' \exp((\lambda - i\omega)x) = \overline{K} \exp((\lambda - i\omega)x) + \overline{K'} \exp((\lambda + i\omega)x)$$

ou encore

$$K \exp(i\omega x) + K' \exp(-i\omega x) = \overline{K} \exp(-i\omega x) + \overline{K'} \exp(i\omega x)$$

La famille $(x \mapsto \exp(-i\omega x), x \mapsto \exp(i\omega x))$ est libre dans le $\mathbb{C}$ -espace vectoriel $\mathbb{C}^{\mathbb{R}}$, donc, une telle fonction convient si et seulement si $K' = \overline{K}$. Les solutions réelles de $(\mathcal{H})$ sont les fonctions

$$\varphi(x) = e^{\lambda x}((p + iq)e^{i\omega x} + (p - iq)e^{-i\omega x}) = e^{\lambda x}(2p \cos \omega x - 2q \sin \omega x)$$

ou encore

$$\varphi(x) = e^{\lambda x}(A \cos \omega x + B \sin \omega x)$$

avec $A, B \in \mathbb{R}$.

5.3 Équation avec second membre

On s'intéresse ici à l'équation

$$(\mathcal{E}) \quad ay'' + by' + cy = g(x)$$

où $a, b, c \in \mathbb{C}$, $a \neq 0$ et g est une fonction continue sur $\mathbb{R}$.

Proposition 14. *On suppose connue une solution φ_0 de $(\mathcal{E})$. Alors une fonction φ est solution de $(\mathcal{E})$ si et seulement si $\varphi - \varphi_0$ est solution de $(\mathcal{H})$.*

Corollaire 15. *L'ensemble $\mathcal{S}_{\mathcal{E}}$ des solutions de $(\mathcal{E})$, s'il est non vide, est un plan affine dont la direction est $\mathcal{S}_{\mathcal{H}}$.*

Dit autrement, si φ_0 est une solution de $(\mathcal{E})$, alors

$$\mathcal{S}_{\mathcal{E}} = \varphi_0 + \mathcal{S}_{\mathcal{H}}$$

Il reste évidemment à trouver *UNE* solution de $\mathcal{E}$. Nous allons nous concentrer sur un cas particulier : on considère un second membre de la forme

$$g(x) = Ae^{\alpha x}$$

où $A, \alpha \in \mathbb{C}^*$.

Proposition 16. *Il existe une solution de $(\mathcal{E})$ de la forme $Q(x)e^{\alpha x}$, où Q est un polynôme. Plus précisément :*

- Si α n'est pas racine de l'équation caractéristique, alors $\deg Q = 0$.

- Si α est racine de l'équation caractéristique, et que l'équation caractéristique a deux racines distinctes, alors $\deg Q = 1$.
- Si α est racine de l'équation caractéristique, et que l'équation caractéristique a une racine double, alors $\deg Q = 2$.

Démonstration. Soit $\varphi(x) = Q(x)e^{\alpha x}$. La fonction φ est solution de $(\mathcal{E})$ si et seulement si

$$aQ'' + (2a\alpha + b)Q' + (a\alpha^2 + b\alpha + c)Q = A$$

Si α n'est pas racine de l'équation caractéristique, le degré du premier membre est celui de Q . Ainsi, si une solution existe, on doit avoir $d^o Q = 0$.

Si α est racine de l'équation caractéristique, l'équation devient

$$aQ'' + (2a\alpha + b)Q' = A$$

Encore deux cas. Si l'équation caractéristique a deux racines distinctes, alors $2a\alpha + b \neq 0$ et on doit donc avoir $d^o Q = 1$. Sinon, l'équation devient

$$aQ'' = A$$

et Q est alors évident : on prend une primitive d'une primitive de $\frac{A}{a}X^2$. $\square$

Remarque 4. Soient $(\mathcal{E}_1)$ et $(\mathcal{E}_2)$ les équations différentielles

$$(\mathcal{E}_1) \quad ay'' + by' + c_y = g_1(x)$$

$$(\mathcal{E}_2) \quad ay'' + by' + c_y = g_2(x)$$

Si f_1 est solution de $(\mathcal{E}_1)$ et f_2 est solution de $(\mathcal{E}_2)$, alors $f_1 + f_2$ est solution de :

$$(\mathcal{E}) \quad ay'' + by' + c_y = g_1(x) + g_2(x)$$

C'est ce que l'on appelle le *principe de superposition*.

On peut donc résoudre toute une classe d'équations différentielles, puisque, par exemple, un sinus ou un cosinus sont des combinaisons d'exponentielles.

Exercice. Résoudre les équations différentielles $y'' + y = e^x$ et $y'' - y = e^x$.

Exercice. Résoudre l'équation différentielle $y'' + y = \cos x$.

5.4 Problème de Cauchy

Proposition 17. *Soit l'équation différentielle*

$$(\mathcal{E}) \quad ay'' + by' + cy = g(x)$$

où g est continue sur l'intervalle I . On suppose que l'équation $(\mathcal{E})$ a une solution. Soit $x_0 \in I$. Soient $y_0, y_1 \in \mathbb{C}$. Il existe une unique solution φ de $(\mathcal{E})$ telle que $\varphi(x_0) = y_0$ et $\varphi'(x_0) = y_1$.

Démonstration. On a pour tout réel x ,

$$\varphi(x) = \lambda_1 \varphi_1(x) + \lambda_2 \varphi_2(x) + \varphi_0(x)$$

où (φ_1, φ_2) est une base de l'espace des solutions de $(\mathcal{H})$ (tout dépend du discriminant de l'équation caractéristique), et φ_0 est une solution particulière de $(\mathcal{E})$. L'écriture des conditions initiales fournit le système

$$\begin{cases} \lambda_1 \varphi_1(x_0) + \lambda_2 \varphi_2(x_0) + \varphi_0(x_0) &= y_0 \\ \lambda_1 \varphi_1'(x_0) + \lambda_2 \varphi_2'(x_0) + \varphi_0'(x_0) &= y_1 \end{cases}$$

Ce système de deux équations à deux inconnues a pour déterminant

$$D = (\varphi_1 \varphi_2' - \varphi_2 \varphi_1')(x_0)$$

On trouve, si $\Delta \neq 0$,

$$D = (\alpha_2 - \alpha_1) e^{(\alpha_1 + \alpha_2)x_0} \neq 0$$

et si $\Delta = 0$,

$$D = e^{2\alpha x_0} \neq 0$$

Ce système a donc une unique solution. $\square$

Chapitre 16

Arithmétique

1 Divisibilité dans $\mathbb{Z}$

On suppose connues les propriétés suivantes : $(\mathbb{Z}, +, \times)$ est un anneau. De plus, dans cet anneau, lorsqu'un produit est nul l'un des facteurs est nul (anneau intègre). Cet anneau est totalement ordonné par la relation $\leq$, qui est compatible avec l'addition et la multiplication. On dispose également sur $\mathbb{Z}$ du concept de valeur absolue.

1.1 Diviseurs, multiples

Définition 1. Soient $a, b \in \mathbb{Z}$. On dit que a divise b ou que b est un multiple de a , et on note $a \mid b$, lorsqu'il existe un $c \in \mathbb{Z}$ tel que $b = ac$:

$$a \mid b \iff \exists c \in \mathbb{Z}, b = ac$$

Remarque 1.

Pour tout entier a , $0 \mid a$ si et seulement si $a = 0$.

Pour tout entier a , $1 \mid a$ et $a \mid 0$.

Si a est non nul, l'entier c tel que $b = ac$ est unique car $\mathbb{Z}$ est intègre. On l'appelle le quotient de b par a et on le note $\frac{b}{a}$.

Notation. Soit $a \in \mathbb{Z}$. On note

$$a\mathbb{Z} = \{na, n \in \mathbb{Z}\}$$

l'ensemble des multiples de a .

Remarque 2. On a $a \mid b$ si et seulement si $b\mathbb{Z} \subset a\mathbb{Z}$.

Proposition 1.

- Pour tout $a \in \mathbb{Z}$, $a \mid a$.
- Pour tous $a, b, c \in \mathbb{Z}$, si $a \mid b$ et $b \mid c$ alors $a \mid c$.
- Pour tous entiers $a, b \in \mathbb{Z}$, $a \mid b$ et $b \mid a$ si et seulement si $b = \pm a$.

La relation « divise » n'est donc pas tout à fait une relation d'ordre. C'est ce que l'on appelle un pré-ordre.

Démonstration.

- Soit $a \in \mathbb{Z}$. On a $a = 1a$ donc $a \mid a$.

- Soient $a, b, c \in \mathbb{Z}$. Supposons que $a \mid b$ et $b \mid c$. Il existe alors $p, q \in \mathbb{Z}$ tels que $b = pa$ et $c = qb$. On en déduit que $c = pqa$ donc $a \mid c$.
- Soient $a, b \in \mathbb{Z}$. Supposons que $a \mid b$ et $b \mid a$. Il existe $p, q \in \mathbb{Z}$ tels que $b = pa$ et $a = qb$. On en déduit $b = pqb$. Si $b \neq 0$, il vient $pq = 1$ donc $p = q = 1$ (et $a = b$) ou $p = q = -1$ (et $a = -b$). Si $b = 0$, alors $a = q0 = 0$ et on a toujours $b = \pm a$. Enfin, si $b = \pm a$, on a évidemment $a \mid b$ et $b \mid a$.

□

Remarque 3. La dernière propriété nous dit que

$$a\mathbb{Z} = b\mathbb{Z} \iff b = \pm a$$

1.2 Division euclidienne

Proposition 2. Soient $a, b \in \mathbb{N}$, $b > 0$. Il existe un entier naturel n tel que $nb > a$.

Démonstration. On a $b(a+1) = ab + b \geq a + b > a$. □

Corollaire 3. Soient a et b deux entiers naturels, $b > 0$. Il existe un unique entier naturel n tel que

$$nb \leq a < (n+1)b$$

Démonstration. Soit

$$E = \{m \in \mathbb{N}, mb > a\}$$

L'ensemble E est une partie non vide de $\mathbb{N}$, il admet donc un plus petit élément. Notons le m_0 . On a donc $(m_0 - 1)b \leq a < m_0b$. On a $m_0 \geq 1$ puisque $m_0b > a \geq 0$. Posons $n = m_0 - 1 \in \mathbb{N}$. On a $nb \leq a < (n+1)b$ d'où l'existence.

Supposons maintenant que n et n' conviennent. On a alors $nb \leq a < (n' + 1)b$ donc $n < n' + 1$ et ainsi $n \leq n'$. De même, $n' \leq n$ donc $n = n'$ d'où l'unicité. □

Corollaire 4. PROPRIÉTÉ D'ARCHIMÈDE

Soient a et b deux entiers relatifs tels que $b > 0$. Il existe un unique entier relatif n tel que

$$nb \leq a < (n+1)b$$

Démonstration. Si $a \geq 0$, c'est le corollaire précédent. Sinon, on applique le dit corollaire à $-a$. Il existe $m \in \mathbb{N}$ tel que

$$mb \leq -a < (m+1)b$$

Si $mb = -a$, alors $-mb = a < (-m+1)b$ et $n = -m$ convient. Sinon, $(-m-1)b < a < -mb$ et $n = -m - 1$ convient.

L'unicité se prouve comme dans le corollaire précédent. $\square$

Proposition 5. DIVISION EUCLIDIENNE

Soient a et b deux entiers relatifs, $b \neq 0$. Il existe un unique couple (q, r) d'entiers relatifs tel que

$$\begin{cases} a = bq + r \\ 0 \leq r < |b| \end{cases}$$

L'entier q est appelé le quotient de la division euclidienne de a par b . L'entier r est le reste de cette même division.

Démonstration. Prenons d'abord $b > 0$. Il existe $q \in \mathbb{Z}$ tel que $qb \leq a < (q+1)b$. Posons $r = a - bq$. On a alors $0 \leq r < b$ d'où l'existence.

Si $b < 0$, il existe $q, r \in \mathbb{Z}$ tels que $a = (-b)q + r = b(-q) + r$ avec $0 \leq r < -b = |b|$. L'unicité se prouve comme dans les corollaires ci-dessus. $\square$

Proposition 6. Les sous-groupes de $\mathbb{Z}$ sont les parties de $\mathbb{Z}$ de la forme $n\mathbb{Z}$, $n \in \mathbb{N}$.

Démonstration. Il est clair que, pour tout $n \in \mathbb{N}$, $n\mathbb{Z}$ est un sous-groupe de $\mathbb{Z}$.

Inversement, soit G un sous-groupe de $\mathbb{Z}$, différent de $\{0\}$. G possède alors un plus petit élément strictement positif, a . Comme $a \in G$, on a $a\mathbb{Z} \subset G$. Réciproquement, si $x \in G$, on fait la division euclidienne de x par a : $x = qa + r$ avec $0 \leq r < a$. Mais x et qa étant dans G , $r = x - qa$ l'est aussi. Comme a est le *plus petit* élément strictement positif de G , on a nécessairement $r = 0$. $\square$

2 PGCD, PPCM

2.1 Somme de deux sous-groupes

Définition 2. Soient H, H' deux sous-groupes du groupe abélien $(G, +)$. On appelle somme de H et H' le sous-ensemble de G défini par $H + H' = \{x + x', x \in H, x' \in H'\}$.

Proposition 7. La somme de deux sous-groupes de G est un sous-groupe de G .

Démonstration.

- Notons 0 le neutre de G . On a $0 \in H$ et $0 \in H'$ donc

$$0 = 0 + 0 \in H + H'$$

- Soient $u, v \in H + H'$. On a $u = x + x'$, $v = y + y'$ où $x, y \in H$ et $x', y' \in H'$. De là,

$$u - v = (x - y) + (x' - y') \in H + H'$$

□

2.2 PGCD

Proposition 8. Soient $a, b \in \mathbb{Z}$. Il existe $\delta \in \mathbb{Z}$ tel que $a\mathbb{Z} + b\mathbb{Z} = \delta\mathbb{Z}$. L'entier δ est unique au signe près.

Démonstration. $a\mathbb{Z} + b\mathbb{Z}$ est un sous-groupe de $\mathbb{Z}$, d'où l'existence de δ . De plus, δ' convient si et seulement si $\delta\mathbb{Z} = \delta'\mathbb{Z}$ ou encore $\delta' = \pm\delta$. □

Définition 3. On appelle plus grands communs diviseurs (en abrégé : pgcd) de a et b les entiers δ définis ci-dessus.

Remarque 4. Deux entiers relatifs admettent deux pgcd (sauf 0 et 0). On abuse en disant LE pgcd de a et b .

Notation. On note $a \wedge b$ le (un) pgcd de a et b . D'autres notations existent, comme (a, b) ou plus évidemment $\text{pgcd}(a, b)$.

Proposition 9. Soient $a, b \in \mathbb{Z}$. Il existe un entier δ , unique au signe près, tel que

$$\forall d \in \mathbb{Z}, d \mid a \text{ et } d \mid b \iff d \mid \delta$$

Démonstration. Soit $\delta = a \wedge b$. Montrons que δ convient. Remarquons tout d'abord que $a = a1 + b0 \in \delta\mathbb{Z}$ donc $\delta \mid a$. De même, $\delta \mid b$.

Soit $d \in \mathbb{Z}$. Supposons que $d \mid \delta$. Alors, par transitivité de la relation « divise », d divise a et b .

Inversement, supposons que $d \mid a$ et $d \mid b$. On a $\delta = 1\delta \in \delta\mathbb{Z}$. Donc $\delta \in a\mathbb{Z} + b\mathbb{Z}$. Il existe ainsi $u, v \in \mathbb{Z}$ tels que $\delta = ua + vb$. Comme a et b sont multiples de d , il en est de même de δ .

Soit maintenant δ' un autre entier qui convient. On a $\delta' \mid \delta$, donc $\delta' \mid a$ et $\delta' \mid b$, donc $\delta' \mid \delta$. De même, $\delta \mid \delta'$ donc $\delta' = \pm\delta$. Les entiers qui conviennent sont donc exactement les pgcd de a et b . □

2.3 Théorème de Bézout

Proposition 10. *Soient a et b et d trois entiers relatifs. Alors*

$$a \wedge b \mid d \iff \exists u, v \in \mathbb{Z}, ua + vb = d$$

Démonstration. Ce théorème est en fait une évidence. En effet, soit $\delta = a \wedge b$. $d \in \delta\mathbb{Z}$ si et seulement si $d \in a\mathbb{Z} + b\mathbb{Z}$. C'est exactement ce qu'il fallait montrer. $\square$

Un cas particulier très important est le théorème de Bézout.

Définition 4. Deux entiers sont dits premiers entre eux lorsque leur pgcd est égal à 1 (ou -1, évidemment).

Proposition 11. THÉORÈME DE BÉZOUT

Soient a et b deux entiers relatifs. Alors, a et b sont premiers entre-eux si et seulement si il existe un couple (u, v) d'entiers relatifs tel que $ua + vb = 1$.

Démonstration. Soit $\delta = a \wedge b$. On a l'existence d'un tel couple (u, v) si et seulement si δ divise 1, c'est à dire si et seulement si $\delta = 1$. $\square$

2.4 Théorème de Gauss

Proposition 12. THÉORÈME DE GAUSS

Soient a, b, c trois entiers relatifs. Alors

$$\left. \begin{array}{l} a \mid bc \\ a \wedge b = 1 \end{array} \right\} \implies a \mid c$$

Démonstration. Les hypothèses montrent l'existence de trois entiers u, v, d tels que $bc = da$ et $ua + vb = 1$. De là

$$uac + vbc = c$$

d'où

$$a(uc + vd) = uac + vad = c$$

Ainsi $a \mid c$. $\square$

2.5 Algorithme d'Euclide

Soient a et b deux entiers. Prenons $a, b \in \mathbb{N}$ pour simplifier. On définit par récurrence sur n une suite $(r_n)_{n \geq 0}$ d'entiers naturels comme suit. On pose

$$r_0 = a, r_1 = b$$

Puis, pour tout entier $n \geq 1$,

- Si $r_n \neq 0$, on note r_{n+1} le reste de la division euclidienne de r_{n-1} par r_n .
- Sinon, on pose $r_{n+1} = 0$.

Remarquons que, par une récurrence évidente, s'il existe un entier $n \geq 1$ tel que $r_n = 0$, alors la suite est nulle à partir du rang n . En fait, il existe un tel entier.

Proposition 13. *Il existe un entier $n \geq 1$ tel que $r_n = 0$.*

Démonstration. Raisonnons par l'absurde et supposons que pour tout n , r_n soit non nul. On a alors pour tout $n \geq 1$, $r_{n+1} < r_n$, puisque r_{n+1} est le reste d'une division euclidienne par r_n . On en déduit facilement par récurrence que pour tout $n \geq 1$,

$$r_n \leq r_1 - n + 1$$

Mais alors,

$$r_{b+1} \leq r_1 - (b + 1) + 1 = 0$$

ce qui est impossible car $r_{b+1} > 0$. $\square$

Proposition 14. *Soit n_0 le plus petit entier $n \geq 0$ tel que $r_{n+1} = 0$. Alors, $a \wedge b = r_{n_0}$.*

Démonstration. Soit $d \in \mathbb{N}$. On voit facilement que pour tout entier $1 \leq n \leq n_0$, $d \mid r_{n-1}$ et $d \mid r_n$ si et seulement si $d \mid r_n$ et $d \mid r_{n+1}$. De là, $d \mid a$ et $d \mid b$ si et seulement si $d \mid r_{n_0}$ et $d \mid r_{n_0+1} = 0$, c'est à dire si et seulement si $d \mid r_{n_0}$. $\square$

Nous disposons donc d'un algorithme permettant de calculer le pgcd de deux entiers. Cet algorithme est-il efficace ? Nous allons voir que oui.

2.6 Complexité de l'algorithme d'Euclide

Soient a et b deux entiers naturels, $a > b$. Nous allons estimer le nombre n_0 de divisions nécessaires à l'obtention de leur pgcd par l'algorithme d'Euclide.

Tout d'abord, quelques notations : $r_0 = a$, $r_1 = b$, et pour $1 \leq k \leq n_0$,

$$r_{k-1} = q_k r_k + r_{k+1}$$

où $r_{n_0} = a \wedge b$ et $r_{n_0+1} = 0$.

Les quotients successifs sont tous au moins égaux à 1. On a donc pour $1 \leq k \leq n_0$,

$$r_{k-1} \geq r_k + r_{k+1}$$

Notons $(F_n)_{n \geq 0}$ la *suite de Fibonacci*, définie par $F_0 = 0$, $F_1 = 1$, et, pour tout $n \in \mathbb{N}$,

$$F_{n+2} = F_{n+1} + F_n$$

On a $r_{n_0+1} = 0 = F_0$ et $r_{n_0} = a \wedge b \geq 1 = F_1$. Donc

$$r_{n_0-2} \geq 1 + 0 = F_2$$

Plus généralement, par une récurrence facile, on a pour $0 \leq k \leq n_0 + 1$,

$$r_{n_0+1-k} \geq F_k$$

Pour $k = n_0$, on obtient

$$r_1 = b \geq F_{n_0}$$

On peut montrer que pour tout $n \in \mathbb{N}$,

$$F_n = \frac{\varphi^n - \bar{\varphi}^n}{\varphi - \bar{\varphi}}$$

où φ et $\bar{\varphi}$ sont les racines de l'équation $x^2 = x + 1$. Prenons pour φ la racine positive de cette équation. On a alors $|\bar{\varphi}| < 1$ et $\varphi - \bar{\varphi} = \sqrt{5}$ donc

$$F_n \geq \frac{\varphi^n - 1}{\sqrt{5}}$$

Ainsi,

$$\frac{\varphi^{n_0} - 1}{\sqrt{5}} \leq b$$

d'où

$$n_0 \leq \frac{\ln(b\sqrt{5} + 1)}{\ln \varphi}$$

Le majorant trouvé est en gros proportionnel au nombre de chiffres de l'entier b , le coefficient de proportionnalité étant à peu près égal à 4,8. L'algorithme d'Euclide est donc très efficace.

2.7 Coefficients de Bézout

Reprenons les notations ci-dessus. Soient $a, b \in \mathbb{Z}$. Soit $\delta = a \wedge b$. L'algorithme d'Euclide permet de trouver un couple (u, v) tel que $ua + vb = 1$. On procède comme suit.

On pose $u_0 = 1$, $v_0 = 0$, $u_1 = 0$, $v_1 = 1$. On a

$$u_0a + v_0b = a = r_0$$

et

$$u_1a + v_1b = b = r_1$$

Soit q_1 tel que $r_0 = q_1r_1 + r_2$. En combinant les deux égalités précédentes, on a

$$u_2a + v_2b = r_2$$

où

$$u_2 = u_0 - q_1u_1 \text{ et } v_2 = v_0 - q_1v_1$$

Soit $0 \leq k \leq n_0 - 2$. Supposons trouvés $u_k, v_k, u_{k+1}, v_{k+1}$ tels que

$$u_ka + v_kb = r_k \text{ et } u_{k+1}a + v_{k+1}b = r_{k+1}$$

Soit q_{k+1} le quotient de la division de r_k par r_{k+1} . En posant

$$u_{k+2} = u_k - q_{k+1}u_{k+1} \text{ et } v_{k+2} = v_k - q_{k+1}v_{k+1}$$

on obtient

$$u_{k+2}a + v_{k+2}b = r_{k+2}$$

On a donc montré comment calculer, pour tout entier $0 \leq k \leq n_0$, deux entiers u_k et v_k tels que $u_ka + v_kb = r_k$. Mais pour $k = n_0$, on a $r_k = \delta$. D'où la construction effective de deux entiers u et v tels que $ua + vb = \delta$.

Un exemple illuminera tout cela.

Exemple. Prenons $a = 33$ et $b = 21$. On présente le tout dans un tableau. Dans la deuxième colonne, les u_k . Dans la troisième colonne, les v_k . Dans la quatrième colonne, les r_k . Dans la dernière colonne, les q_k .

k	$33 \times$	$+21 \times$	$=$	
0	1	0	33	—
1	0	1	21	1
2	1	−1	12	1
3	−1	2	9	1
4	2	−3	3	3
5	−7	11	0	—

On en déduit $33 \wedge 21 = 3$, $u = 2$, $v = -3$.

2.8 Propriétés utiles

Proposition 15. Soient $n \geq 1$ et $a, b_1, \dots, b_n \in \mathbb{Z}$. On a

$$a \wedge \prod_{k=1}^n b_k = 1 \iff \forall k \in \llbracket 1, n \rrbracket, a \wedge b_k = 1$$

Démonstration. Pour $n = 1$ il n'y a rien à montrer.

Montrons la propriété pour $n = 2$. Supposons $a \wedge b_1 b_2 = 1$. D'après le théorème de Bézout, il existe deux entiers u et v tels que $ua + vb_1 b_2 = 1$. En lisant cette égalité de deux façons, on constate, toujours d'après Bézout, que a est premier avec b_1 et b_2 .

Supposons maintenant que a est premier avec b_1 et b_2 . Il existe des entiers u, v, u', v' tels que $ua + vb_1 = 1$ et $u'a + v'b_2 = 1$. En multipliant ces égalités on obtient $Ua + Vb_1 b_2 = 1$ avec $U = uu' + uv'b_2 + u'vb_1$ et $V = vv'$. a est donc premier avec $b_1 b_2$, toujours et encore grâce à Bézout.

La propriété pour n quelconque est une récurrence sur n . Elle est laissée en exercice. $\square$

Proposition 16. Soient $n \geq 1$ et $a_1, \dots, a_n, b \in \mathbb{Z}$. On suppose les a_k premiers entre-eux deux à deux. Alors,

$$\prod_{k=1}^n a_k \mid b \iff \forall k \in \llbracket 1, n \rrbracket, a_k \mid b$$

Démonstration. Ici encore, le cas général est laissé en exercice. Traitons le cas où $n = 2$.

Remarquons tout d'abord que le sens de gauche à droite est évident.

Pour la réciproque, supposons que a_1 et a_2 sont premiers entre-eux et que a_1 et a_2 divisent b . Comme $a_1 \mid b$, il existe $c \in \mathbb{Z}$ tel que $b = a_1 c$. De là, $a_2 \mid a_1 c$. Comme $a_2 \wedge a_1 = 1$, par le théorème de Gauss, $a_2 \mid c$. Il existe $c' \in \mathbb{Z}$ tel que $c = a_2 c'$, d'où $b = a_1 a_2 c'$: $a_1 a_2$ divise b . $\square$

Proposition 17. Soient $a, b, k \in \mathbb{Z}$. On a

$$(ka) \wedge (kb) = k(a \wedge b)$$

Démonstration. Posons $\delta = a \wedge b$ et $\Delta = (ka) \wedge (kb)$. Comme δ divise a , on a facilement que $k\delta$ divise ka . De même, $k\delta$ divise kb et donc $k\delta$ divise Δ .

Inversement, il existe $u, v \in \mathbb{Z}$ tels que $ua + vb = \delta$. En multipliant cette égalité par k , il vient

$$u(ka) + v(kb) = k\delta$$

Ainsi, Δ divise $k\delta$.

Les entiers $k\delta$ et Δ se divisent l'un l'autre, ils sont donc égaux (au signe près, bien entendu). $\square$

Proposition 18. Soient $a, b, \delta \in \mathbb{Z}$. Alors $\delta = a \wedge b$ si et seulement si il existe deux entiers a_1, b_1 tels que :

$$\begin{cases} a &= \delta a_1 \\ b &= \delta b_1 \\ a_1 \wedge b_1 &= 1 \end{cases}$$

Démonstration. Si a et b sont nuls, le résultat est évident. Supposons donc dans ce qui suit que a ou b est non nul.

Supposons $\delta = a \wedge b$ (on a donc $\delta \neq 0$). On peut alors écrire $a = \delta a_1$ et $b = \delta b_1$, où $a_1, b_1 \in \mathbb{Z}$. De plus, il existe u et v tels que $ua + vb = \delta$. En simplifiant par δ , on en tire $ua_1 + vb_1 = 1$, donc $a_1 \wedge b_1 = 1$ par le théorème de Bézout.

Inversement, supposons $a = \delta a_1, b = \delta b_1$, avec $a_1 \wedge b_1 = 1$. Appelons Δ le pgcd de a et b . Par Bézout appliqué à a_1 et b_1 et une multiplication par δ , on obtient deux entiers u et v tels que $ua + vb = \delta$. On en déduit que $\Delta \mid \delta$. Mais δ est clairement un diviseur commun de a et b , donc $\delta \mid \Delta$. Ainsi, $\delta = \Delta$ (au signe près, comme toujours). $\square$

Appliquons ce qui précède pour montrer un résultat bien connu sur les rationnels.

Proposition 19. Soit $x \in \mathbb{Q}$. Il existe un unique $p \in \mathbb{Z}$ et un unique $q \in \mathbb{N}^*$ tels que $p \wedge q = 1$ et $x = \frac{p}{q}$.

Une telle écriture s'appelle l'écriture du rationnel x sous forme irréductible.

Démonstration. Commençons par l'existence. On a $x = \frac{a}{b}$ où $a \in \mathbb{Z}$ et $b \in \mathbb{Z}^*$. Soit $\delta = a \wedge b$. D'après la propriété précédente, on a alors $a = \delta p$ et $b = \delta q$ avec $p \wedge q = 1$. De là, $x = \frac{p}{q}$. Quitte à remplacer p et q par $-p$ et $-q$, on peut supposer $q > 0$.

Passons à l'unicité. Supposons

$$x = \frac{p}{q} = \frac{p'}{q'}$$

où $p, p', q, q' \in \mathbb{Z}$, $q, q' > 0$, $p \wedge q = 1$ et $p' \wedge q' = 1$. On a $pq' = p'q$. Par le théorème de Gauss, $q \mid q'$ et aussi $q' \mid q$. Mais $q, q' \in \mathbb{N}$, donc $q = q'$. On reporte et on obtient $p = p'$. $\square$

2.9 PPCM

Proposition 20. Soient a et b deux entiers relatifs. Il existe un entier μ , unique au signe près, vérifiant

$$\forall m \in \mathbb{Z}, a \mid m \text{ et } b \mid m \iff \mu \mid m$$

Définition 5. « L' » entier μ ci-dessus est appelé plus petit commun multiple (en abrégé : ppcm) de a et b . On abuse en parlant *du* ppcm de a et b . On note $\mu = a \vee b$ ou encore $\text{ppcm}(a, b)$.

Démonstration. Écrivons $a = \delta a_1$, $b = \delta b_1$, avec $a_1 \wedge b_1 = 1$. Soit $m \in \mathbb{Z}$. Supposons que $a \mid m$ et $b \mid m$. Comme δ divise a et b , on a aussi par transitivité $\delta \mid m$. Écrivons donc $m = \delta m_1$ où $m_1 \in \mathbb{Z}$. On a alors facilement $a_1 \mid m_1$ et $b_1 \mid m_1$, donc, par Gauss, $a_1 b_1 \mid m_1$ et donc $\delta a_1 b_1 \mid m$. Ainsi, tout multiple commun de a et b est multiple de l'entier

$$\mu = \delta a_1 b_1$$

Inversement, l'entier μ ci dessus est bien un multiple commun de a et b .

Enfin, si μ_1 et μ_2 sont des ppcm de a et b , chacun divise l'autre, et ils sont donc égaux au signe près. $\square$

Proposition 21. Soient $a, b \in \mathbb{Z}$. Soient $\delta = a \wedge b$ et $\mu = a \vee b$. Alors

$$\delta \mu = ab$$

Démonstration. $\delta \mu = \delta(\delta a_1 b_1) = (\delta a_1)(\delta b_1) = ab$. $\square$

Remarque 5. Dire que les multiples de μ sont les multiples communs de a et b revient à dire que

$$a\mathbb{Z} \cap b\mathbb{Z} = \mu\mathbb{Z}$$

Ceci serait une façon d'introduire le ppcm analogue à celle utilisée pour le pgcd, en remarquant que l'intersection de deux sous-groupes d'un groupe est encore un sous-groupe.

2.10 Résolution d'une équation diophantienne simple

Une *équation diophantienne* est une équation dont les inconnues sont des entiers. Ce type d'équation se résout par des méthodes très différentes de celles utilisées pour les équations à inconnues réelles ou complexes. Nous allons ici nous intéresser aux équations diophantienne de degré 1 à deux inconnues.

Soient a, b, c trois entiers relatifs. On considère l'équation d'inconnues $x, y \in \mathbb{Z}$

$$(E) \quad ax + by = c$$

Si $a = b = 0$, la résolution de (E) est évidente. Supposons dorénavant que $(a, b) \neq (0, 0)$, et donc $\delta \neq 0$.

Soit $\delta = a \wedge b$. On voit que si (E) admet une solution, alors $\delta \mid c$. Première conclusion :

Si $\delta \nmid c$, alors l'équation (E) n'a pas de solution.

Supposons maintenant que $\delta \mid c$. Écrivons $c = \delta c_1$, où $c_1 \in \mathbb{Z}$. On écrit $a = \delta a_1, b = \delta b_1$, où $a_1, b_1 \in \mathbb{Z}$ et on est ramenés à l'équation

$$(E_1) \quad x, y \in \mathbb{Z}, a_1 x + b_1 y = c_1$$

avec $a_1 \wedge b_1 = 1$. Cette équation a les mêmes solutions que l'équation (E) .

On sait trouver une solution à cette équation : il suffit d'appliquer l'algorithme d'Euclide. Celui-ci fournit un couple (u, v) tel que $ua_1 + vb_1 = 1$. Le couple $(x_1, y_1) = (uc_1, vc_1)$ est alors solution de (E_1) .

Soit maintenant un couple $(x, y) \in \mathbb{Z}^2$. Alors, (x, y) est solution de (E_1) (ou (E)) si et seulement si $a_1 x + b_1 y = a_1 x_1 + b_1 y_1$ ou encore

$$a_1(x - x_1) = b_1(y_1 - y)$$

Le théorème de Gauss permet alors de conclure que (x, y) est solution si et seulement si il existe un entier relatif k tel que

$$\begin{cases} x &= x_1 + kb_1 \\ y &= y_1 - ka_1 \end{cases}$$

2.11 PGCD d'un nombre fini d'entiers

Soient $a_1, \dots, a_n$ n entiers relatifs. L'ensemble

$$a_1\mathbb{Z} + \dots + a_n\mathbb{Z} = \{u_1 a_1 + \dots + u_n a_n, u_1, \dots, u_n \in \mathbb{Z}\}$$

est un sous-groupe de $\mathbb{Z}$. Il existe donc $\delta \in \mathbb{Z}$, unique au signe près, tel que

$$a_1\mathbb{Z} + \dots + a_n\mathbb{Z} = \delta\mathbb{Z}$$

L'entier δ est appelé le pgcd de $a_1, \dots, a_n$. On le note

$$\delta = \wedge_{k=1}^n a_k$$

Remarquons que

$$\sum_{k=1}^n a_k \mathbb{Z} = \sum_{k=1}^{n-1} a_k \mathbb{Z} + a_n \mathbb{Z}$$

On a donc

$$\wedge_{k=1}^n a_k \mathbb{Z} = (\wedge_{k=1}^{n-1} a_k) \mathbb{Z} + a_n \mathbb{Z} = ((\wedge_{k=1}^{n-1} a_k) \wedge a_n) \mathbb{Z}$$

On en déduit qu'au signe près,

$$\wedge_{k=1}^n a_k = (\wedge_{k=1}^{n-1} a_k) \wedge a_n$$

On peut donc calculer le pgcd de n entiers de proche en proche, en calculant des pgcd de deux entiers. On peut en fait montrer que le pgcd est associatif, et les parenthèses peuvent être placées comme bon nous semble. Nous ne prouverons pas cette propriété.

Définition 6. Soient $a_1, \dots, a_n \in \mathbb{Z}$.

On dit que les a_k sont premiers entre-eux deux à deux lorsque $\forall i \neq j, a_i \wedge a_j = 1$.

On dit que les a_k sont premiers entre-eux dans leur ensemble lorsque $\wedge_{k=1}^n a_k = 1$.

Proposition 22. *Si des entiers sont premiers entre-eux deux à deux, ils sont premiers entre-eux dans leur ensemble. La réciproque est fausse.*

Démonstration. Le sens direct est évident. Pour le contre-exemple, prendre par exemple les nombres 6, 10 et 15. $\square$

Proposition 23. Soient $a_1, \dots, a_n \in \mathbb{Z}$. Soit $\delta = \wedge_{k=1}^n a_k$. Soit $d \in \mathbb{Z}$. Alors,

$$\delta \mid d \iff \exists u_1, \dots, u_n \in \mathbb{Z}, \sum_{k=1}^n u_k a_k = d$$

Démonstration. Il suffit de reformuler : $\delta \mid d$ si et seulement si $d \in \delta\mathbb{Z}$. Or, $\delta\mathbb{Z} = a_1\mathbb{Z} + \dots + a_n\mathbb{Z}$. $\square$

Corollaire 24. THÉORÈME DE BÉZOUT GÉNÉRALISÉ

Soient $a_1, \dots, a_n \in \mathbb{Z}$. Les a_k sont premiers entre-eux dans leur ensemble si et seulement si il existe $u_1, \dots, u_n \in \mathbb{Z}$ tels que

$$\sum_{k=1}^n u_k a_k = 1$$

Démonstration. Il suffit de remarquer que $\delta \mid 1$ si et seulement si $\delta = 1$. $\square$

3 Nombres premiers

3.1 Définition

Définition 7. Soit $p \in \mathbb{Z}$. On dit que p est *premier* lorsque $p \neq \pm 1$, et que p est divisible uniquement par ± 1 et $\pm p$. Sinon, on dit que p est *composé*.

Remarque 6. Dorénavant, nous nous occuperons uniquement de nombres premiers *positifs*. Lorsque nous dirons « le nombre premier p » il conviendra de comprendre « le nombre premier positif p ».

3.2 Propriétés

Proposition 25. Soient $a, b \in \mathbb{Z}$. On suppose p premier. Alors,

$$a \wedge p = 1 \iff p \nmid a$$

Démonstration. Soit $\delta = a \wedge p$. Alors δ divise p , donc $\delta = 1$ ou $\delta = p$.

Supposons $\delta \neq 1$. Alors $\delta = p$, donc $p \mid a$.

Inversement, supposons que $p \mid a$. Alors, comme $p \mid p$, p divise δ , donc $\delta \neq 1$. $\square$

Proposition 26. LEMME D'EUCLIDE

Soit $p \in \mathbb{Z}$, $p \neq 0, 1$. Alors, p est premier si et seulement si

$$(*) \quad \forall a, b \in \mathbb{Z}, p \mid ab \implies p \mid a \text{ ou } p \mid b$$

Démonstration. Supposons p premier. Soient $a, b \in \mathbb{Z}$. Supposons que $p \mid ab$. Supposons aussi que $p \nmid a$. Par la proposition précédente, $p \wedge a = 1$. Donc, par le théorème de Gauss, $p \mid b$.

Inversement, supposons $(*)$. Soit d un diviseur de p . Écrivons $p = dc$, où $c \in \mathbb{Z}$. On a donc $p \mid dc$ d'où, par $(*)$, $p \mid d$ ou $p \mid c$. Mais c et d divisent p , donc

- $d = \pm p$, ou
- $c = \pm p$ et donc $d = \pm 1$.

$\square$

Proposition 27. *Deux nombres premiers (positifs) distincts sont premiers entre-eux.*

Démonstration. Au signe près, les seuls diviseurs de p sont 1 et p . Ceux de q sont 1 et q . Donc le seul diviseur commun de p et q est effectivement 1. $\square$

Proposition 28. *Tout entier $n \geq 2$ possède au moins un diviseur premier.*

Démonstration. On le prouve par récurrence forte sur n . Pour $n = 2$, c'est clair, vu que 2 est premier. Supposons maintenant que tout entier k , $2 \leq k < n$, possède un diviseur premier.

Si n est premier, alors n a bien un diviseur premier : lui-même.

Si $n = ab$ est composé, alors par exemple $2 \leq a < n$ a un diviseur premier p . Mais $p \mid a$ et $a \mid n$ donc $p \mid n$. $\square$

Proposition 29. *Il existe une infinité de nombres premiers.*

Démonstration. Supposons le contraire. Notons

$$\mathcal{P} = \{p_1, \dots, p_n\}$$

l'ensemble des nombres premiers. Soit

$$N = \prod_{k=1}^n p_k + 1$$

L'entier N est clairement supérieur à 2 (il est même très grand!). Donc il possède un diviseur premier. Mais ce diviseur ne peut être aucun des éléments de $\mathcal{P}$. Contradiction. $\square$

3.3 Décomposition en produit de facteurs premiers

Proposition 30. *Tout entier naturel $n \geq 2$ s'écrit comme un produit de nombres premiers.*

Démonstration. On procède par récurrence forte sur n . C'est clair si $n = 2$. Soit maintenant un entier n tel que tout entier $m < n$ s'écrit comme produit de nombres premiers. Il y a deux possibilités.

- Si n est premier, c'est un produit de 1 nombre premier.

- Sinon, $n = ab$ avec $2 \leq a, b < n$ est composé. Mais a et b sont produits de nombres premiers d'après l'hypothèse de récurrence. Donc, n aussi.

□

Proposition 31. *La décomposition en produit de nombres premiers est unique à l'ordre près des facteurs.*

On s'affranchit du problème de « à l'ordre près » en normalisant l'écriture :

Notons $\mathcal{P}$ l'ensemble des nombres premiers positifs. Le théorème d'existence de la décomposition nous dit que tout entier $n \geq 1$ s'écrit

$$n = \prod_{p \in \mathcal{P}} p^{\alpha_p}$$

où les $\alpha_p \in \mathbb{N}$ sont tous nuls sauf un nombre fini. On peut maintenant reformuler les deux propositions précédentes.

Proposition 32. *Tout entier $n \geq 1$ s'écrit de façon unique*

$$n = \prod_{p \in \mathcal{P}} p^{\alpha_p}$$

où les $\alpha_p \in \mathbb{N}$ sont tous nuls sauf un nombre fini.

Démonstration. Supposons

$$n = \prod_{p \in \mathcal{P}} p^{\alpha_p} = \prod_{p \in \mathcal{P}} p^{\beta_p}$$

où les α_p, β_p sont tous nuls sauf un nombre fini. Soit $q \in \mathcal{P}$. On isole le facteur correspondant :

$$q^{\alpha_q} A = q^{\beta_q} B$$

où A et B sont des produits de nombres premiers différents de q , et donc sont premiers avec q . En appliquant deux fois le théorème de Gauss, on obtient facilement que $q^{\alpha_q} = q^{\beta_q}$, d'où $\alpha_q = \beta_q$. □

3.4 Valuation p-adique

Définition 8. Soit $n \in \mathbb{N}^*$. Soit $p \in \mathcal{P}$. La valuation p -adique de n est le plus grand entier naturel k tel que $p^k \mid n$. On la note $\nu_p(n)$.

Exemple. De $72 = 2^3 3^2$, on déduit

$$\nu_2(72) = 3, \nu_3(72) = 2, \nu_7(72) = 0$$

Remarque 7. En d'autres termes, la valuation p -adique de n est l'exposant de p dans la décomposition de n en produit de facteurs premiers. Ainsi,

$$n = \prod_{p \in \mathcal{P}} p^{\nu_p(n)}$$

Proposition 33. Soient $a, b \in \mathbb{N}^*$. Soit $p \in \mathcal{P}$. On a

$$\nu_p(ab) = \nu_p(a) + \nu_p(b)$$

et

$$\nu_p(a + b) \geq \min(\nu_p(a), \nu_p(b))$$

Si $\nu_p(a) \neq \nu_p(b)$, l'inégalité ci-dessus est une égalité.

Démonstration. Laissée en exercice. $\square$

3.5 Application au pgcd et au ppcm

Proposition 34. Soient $a = \prod_{p \in \mathcal{P}} p^{\alpha_p}$ et $b = \prod_{p \in \mathcal{P}} p^{\beta_p}$ deux entiers naturels non nuls. Alors

$$a \wedge b = \prod_{p \in \mathcal{P}} p^{\min(\alpha_p, \beta_p)} \quad \text{et} \quad a \vee b = \prod_{p \in \mathcal{P}} p^{\max(\alpha_p, \beta_p)}$$

Démonstration. Soit $d = \prod_{p \in \mathcal{P}} p^{\gamma_p}$ un entier non nul. Alors, d divise a si et seulement si pour tout $p \in \mathcal{P}$, on a $\gamma_p \leq \alpha_p$: il suffit dans le sens non trivial d'utiliser le théorème de Gauss. Il en est de même pour b , donc d divise a et b si et seulement si pour tout $p \in \mathcal{P}$, on a $\gamma_p \leq \min(\alpha_p, \beta_p)$. Autrement dit, d divise a et b si et seulement si d divise $\prod_{p \in \mathcal{P}} p^{\min(\alpha_p, \beta_p)}$. On procède de même pour le ppcm. $\square$

Remarque 8. On peut encore écrire ce résultat avec des valuations :

$$\nu_p(a \wedge b) = \min(\nu_p(a), \nu_p(b))$$

et

$$\nu_p(a \vee b) = \max(\nu_p(a), \nu_p(b))$$

4 Congruences

4.1 Rappels

Nous avons déjà parlé de congruence dans le chapitre sur les relations. Rappelons que si $n \in \mathbb{N}$ et $a, b \in \mathbb{Z}$, on $a \equiv b[n]$ lorsque $b - a \in n\mathbb{Z}$. La relation de congruence modulo n est une relation d'équivalence sur $\mathbb{Z}$. Si $n \neq 0$, cette relation possède exactement n classes, qui sont $\overline{0}, \overline{1}, \dots, \overline{n-1}$. L'ensemble des classes est ainsi

$$\mathbb{Z}/n\mathbb{Z} = \{\overline{0}, \dots, \overline{n-1}\}$$

4.2 Opérations sur les congruences

Proposition 35. *Soit $n \in \mathbb{N}$. La relation de congruence modulo n est compatible avec l'addition et la multiplication. Pour tous $a, b, a', b' \in \mathbb{Z}$ tels que $a \equiv b[n]$ et $a' \equiv b'[n]$,*

$$a + a' \equiv b + b'[n]$$

et

$$aa' \equiv bb'[n]$$

Démonstration. Il existe $k, k' \in \mathbb{Z}$ tels que $b - a = kn$ et $b' - a' = k'n$. De là,

$$(b + b') - (a + a') = (b - a) + (b' - a') = (k + k')n \in n\mathbb{Z}$$

et

$$bb' - aa' = b(b' - a') + (b - a)a' = (k'b + ka')n \in n\mathbb{Z}$$

□

4.3 Le petit théorème de Fermat

Lemme 36. *Soit p un nombre premier. On a, pour tout entier $1 \leq k \leq p-1$,*

$$p \mid \binom{p}{k}$$

Démonstration. On a

$$p! = k!(p-k)! \binom{p}{k}$$

donc p divise $k!(p-k)!\binom{p}{k}$. Mais p est premier avec $1, 2, \dots, k$ donc avec $k!$. De même, p est premier avec $1, \dots, p-k$, donc avec $(p-k)!$. Donc p est premier avec $k!(p-k)!$. Par le théorème de Gauss, p divise $\binom{p}{k}$. $\square$

Proposition 37. *Soit p un nombre premier. On a pour tout $a \in \mathbb{Z}$*

$$a^p \equiv a[p]$$

Démonstration. Montrons d'abord la propriété pour a entier naturel, par récurrence sur a . C'est clair pour $a = 0$. Supposons la propriété vraie pour a . On a

$$(a+1)^p = \sum_{k=0}^p \binom{p}{k} a^k$$

D'après le lemme, cette somme est congrue à

$$\binom{p}{0}a^0 + \binom{p}{p}a^p = a^p + 1$$

modulo p . D'après l'hypothèse de récurrence, elle est bien congrue à $a + 1$.

Supposons maintenant $a < 0$. $-a \in \mathbb{N}$, donc $(-a)^p \equiv -a[p]$. Si $p \geq 3$, alors p est impair (car premier) donc $(-a)^p = -a^p$ d'où le résultat. Si $p = 2$, $(-a)^p = a^p$, mais $-a \equiv a[p]$, la propriété est encore vérifiée. $\square$

Corollaire 38. *Soit p un nombre premier. On a, pour tout $a \in \mathbb{Z}$ non multiple de p ,*

$$a^{p-1} \equiv 1[p]$$

Démonstration. Le nombre premier p divise $a^p - a = a(a^{p-1} - 1)$. Mais p ne divise pas a donc p est premier avec a . Par le théorème de Gauss, p divise $a^{p-1} - 1$. $\square$

5 Annexe : les anneaux $\mathbb{Z}/n\mathbb{Z}$

Nous allons dans cette section introduire les anneaux $\mathbb{Z}/n\mathbb{Z}$. Ce sont des anneaux finis qui possèdent de très nombreuses propriétés. Nous nous contenterons d'une courte introduction.

Dans toute cette section, n désigne un entier naturel non nul. Pour tout $a \in \mathbb{Z}$, on note $\bar{a}$ la classe de a modulo n . Rappelons la notation

$$\mathbb{Z}/n\mathbb{Z} = \{\bar{0}, \dots, \overline{n-1}\}$$

L'ensemble $\mathbb{Z}/n\mathbb{Z}$ est l'ensemble des *entiers modulo n* . C'est un ensemble de cardinal n .

5.1 Une addition dans $\mathbb{Z}/n\mathbb{Z}$

Pour tous $a, b \in \mathbb{Z}$, posons

$$\bar{a} + \bar{b} = \overline{a + b}$$

Il s'agit dans un premier temps de montrer que ceci a un sens. Précisément, il faut prouver que $\overline{a + b}$ ne dépend *que* de $\bar{a}$ et $\bar{b}$, et pas de a et b eux-mêmes. Pour cela, soient $a, a', b, b' \in \mathbb{Z}$. Supposons que $\bar{a} = \bar{a}'$ et $\bar{b} = \bar{b}'$. Cela signifie que $a \equiv a'[n]$ et $b \equiv b'[n]$. On a alors

$$a + b \equiv a' + b'[n]$$

c'est à dire

$$\overline{a + b} = \overline{a' + b'}$$

Nous avons donc bien défini une opération $+$ dans $\mathbb{Z}/n\mathbb{Z}$.

Proposition 39. $(\mathbb{Z}/n\mathbb{Z}, +)$ est un groupe abélien.

Démonstration. Ce sont de simples vérifications. Le neutre pour l'addition est $\bar{0}$, l'opposé de $\bar{a}$ est $\overline{-a}$. Si $a, b, c \in \mathbb{Z}$, on a

$$\begin{aligned} (\bar{a} + \bar{b}) + \bar{c} &= \overline{a + b} + \bar{c} \\ &= \overline{(a + b) + c} \\ &= \overline{a + (b + c)} \\ &= \bar{a} + \overline{b + c} \\ &= \bar{a} + (\bar{b} + \bar{c}) \end{aligned}$$

d'où l'associativité. La commutativité est laissée en exercice. $\square$

Exemple. Voici la table d'addition dans $\mathbb{Z}/6\mathbb{Z}$. Pour simplifier, on n'écrit pas les barres sur les classes.

+	0	1	2	3	4	5
0	0	1	2	3	4	5
1	1	2	3	4	5	0
2	2	3	4	5	0	1
3	3	4	5	0	1	2
4	4	5	0	1	2	3
5	5	0	1	2	3	4

5.2 Une multiplication dans $\mathbb{Z}/n\mathbb{Z}$

Pour tous $a, b \in \mathbb{Z}$, posons

$$\bar{a} \times \bar{b} = \overline{a \times b}$$

Il s'agit dans un premier temps de montrer que ceci a un sens. Précisément, il faut prouver que $\overline{a \times b}$ ne dépend *que* de $\bar{a}$ et $\bar{b}$, et pas de a et b eux-mêmes. Pour cela, soient $a, a', b, b' \in \mathbb{Z}$. Supposons que $\bar{a} = \bar{a'}$ et $\bar{b} = \bar{b'}$. Cela signifie que $a \equiv a'[n]$ et $b \equiv b'[n]$. On a alors

$$a \times b \equiv a' \times b'[n]$$

c'est à dire

$$\overline{a \times b} = \overline{a' \times b'}$$

Nous avons donc bien défini une opération $\times$ dans $\mathbb{Z}/n\mathbb{Z}$.

Proposition 40. $(\mathbb{Z}/n\mathbb{Z}, +, \times)$ est un anneau commutatif.

Démonstration. Ce sont de simples vérifications. Le neutre pour la multiplication est $\bar{1}$. Le reste des propriétés est laissé en exercice. $\square$

Exemple. Voici la table de multiplication dans $\mathbb{Z}/6\mathbb{Z}$. Pour simplifier, on n'écrit pas les barres sur les classes.

$\times$	0	1	2	3	4	5
0	0	0	0	0	0	0
1	0	1	2	3	4	5
2	0	2	4	0	2	4
3	0	3	0	3	0	3
4	0	4	2	0	4	2
5	0	5	4	3	2	1

Exemple. Voici maintenant la table de multiplication dans $\mathbb{Z}/7\mathbb{Z}$.

$\times$	0	1	2	3	4	5	6
0	0	0	0	0	0	0	0
1	0	1	2	3	4	5	6
2	0	2	4	6	1	3	5
3	0	3	6	2	5	1	4
4	0	4	1	5	2	6	3
5	0	5	3	1	6	4	2
6	0	6	5	4	3	2	1

Exercice. Dresser les tables de multiplication des anneaux $\mathbb{Z}/5\mathbb{Z}$ et $\mathbb{Z}/8\mathbb{Z}$.

5.3 $\mathbb{Z}/n\mathbb{Z}$ est-il un corps ?

Lorsqu'on compare les deux exemples ci-dessus, on voit une différence : $\mathbb{Z}/6\mathbb{Z}$ n'est pas un corps (en fait cet anneau possède des diviseurs de zéro), alors que $\mathbb{Z}/7\mathbb{Z}$ est un corps.

Proposition 41.

- Si n est premier, alors $\mathbb{Z}/n\mathbb{Z}$ est un corps.
- Si n est composé, alors $\mathbb{Z}/n\mathbb{Z}$ n'est pas intègre (et n'est donc pas un corps).

Démonstration.

- Supposons n premier. On a donc $n \neq 1$, et donc $\bar{0} \neq \bar{1}$. Il s'agit de montrer que tout élément non nul de $\mathbb{Z}/n\mathbb{Z}$ est inversible pour la multiplication. Soit donc $x \in \mathbb{Z}/n\mathbb{Z}$, $x \neq \bar{0}$. On a donc $x = \bar{a}$, où $a \in \mathbb{Z}$ et a n'est pas multiple de n . Comme n est premier, a est premier avec n et, par le théorème de Bézout, il existe $u, v \in \mathbb{Z}$ tels que

$$ua + vn = 1$$

En passant aux classes, il vient

$$\bar{u} \times \bar{a} = \bar{1}$$

Ainsi, $\bar{a}$ est inversible et son inverse est $\bar{u}$. Remarquons que l'algorithme d'Euclide étendu nous permet de calculer effectivement l'inverse de $\bar{a}$.

- Supposons n composé. On a donc $n = ab$, où $1 \leq a, b < n$. En passant aux classes,

$$\bar{a} \times \bar{b} = \bar{n} = \bar{0}$$

Cependant $\bar{a} \neq \bar{0}$ et $\bar{b} \neq \bar{0}$. L'anneau $\mathbb{Z}/n\mathbb{Z}$ possède donc des diviseurs de zéro.

□

Exercice. Quels sont les éléments inversibles de $\mathbb{Z}/6\mathbb{Z}$? Plus généralement, montrer que si $n \geq 1$ et $a \in \mathbb{Z}$, alors $\bar{a}$ est inversible dans $\mathbb{Z}/n\mathbb{Z}$ si et seulement si $a \wedge n = 1$.

Chapitre 17

Calcul matriciel

1 Notion de matrice

1.1 Introduction

Définition 1. Soit E un ensemble. Soient p, q deux entiers non nuls. On appelle matrice à p lignes et q colonnes à coefficients dans E toute application

$$A : \llbracket 1, p \rrbracket \times \llbracket 1, q \rrbracket \rightarrow E$$

Une matrice est donc une suite finie à deux indices. On représente la matrice

$$A = (A_{ij})_{1 \leq i \leq p, 1 \leq j \leq q}$$

par un tableau à p lignes et q colonnes. Ce tableau est noté indifféremment avec des $\llbracket$ ou des $()$ mais pas avec des $||$ (notation qui sera réservée aux déterminants).

Notation. $\mathcal{M}_{pq}(E)$ désigne l'ensemble des matrices à p lignes et q colonnes à coefficients dans E . On parle également de matrices de taille $p \times q$.

Lorsque $p = q$, on note plus simplement $\mathcal{M}_p(E)$ l'ensemble des *matrices carrées* de taille $p \times p$ (ou de taille p) à coefficients dans E .

Remarque 1.

- Une *matrice colonne* est une matrice qui a une seule colonne, c'est à dire un élément de $\mathcal{M}_{p1}(\mathbb{K})$.
- Une *matrice ligne* est une matrice qui a une seule ligne, c'est à dire un élément de $\mathcal{M}_{1q}(\mathbb{K})$.

Remarque 2. On généralise, et on appelle encore matrice à p lignes et q colonnes toute application $M : I \times J \rightarrow E$ lorsque I et J sont deux ensembles finis non vides de cardinaux respectifs p et q .

Dorénavant, nous supposons que $E = \mathbb{K}$ est un corps.

Nous n'explicitons pas à chaque fois les conditions sur les tailles des matrices. Par exemple, dans l'expression « $\mathcal{M}_{pq}(\mathbb{K})$ », p et q sont systématiquement des entiers naturels non nuls.

1.2 Addition, produit par un scalaire

On définit les deux opérations suivantes.

Définition 2. Soient $A, B \in \mathcal{M}_{pq}(\mathbb{K})$. Soit $\lambda \in \mathbb{K}$.

- La somme de A et B est la matrice $A + B \in \mathcal{M}_{pq}(\mathbb{K})$ définie, pour $i \in \llbracket 1, p \rrbracket$ et $j \in \llbracket 1, q \rrbracket$, par

$$(A + B)_{ij} = A_{ij} + B_{ij}$$

- Le produit de A par le scalaire λ est la matrice $\lambda A \in \mathcal{M}_{pq}(\mathbb{K})$ définie, pour $i \in \llbracket 1, p \rrbracket$ et $j \in \llbracket 1, q \rrbracket$, par

$$(\lambda A)_{ij} = \lambda A_{ij}$$

Proposition 1. $(\mathcal{M}_{pq}(\mathbb{K}), +, \cdot)$ est un $\mathbb{K}$ -espace vectoriel.

Démonstration. C'est une simple vérification. Signalons que

- Le neutre pour l'addition est la matrice nulle, qui a tous ses coefficients égaux à 0.
- L'opposé d'une matrice est la matrice ayant tous ses coefficients opposés à la matrice de départ.

□

Notation. Étant donnés deux « objets » i et j , on note $\delta_{ij} = 0$ si $i \neq j$ et $\delta_{ij} = 1$ si $i = j$. Le symbole δ_{ij} est appelé le *symbole de Kronecker*.

Définition 3. Soient $i \in \llbracket 1, p \rrbracket$ et $j \in \llbracket 1, q \rrbracket$. On appelle *matrice élémentaire* d'indices i, j la matrice $E_{ij} \in \mathcal{M}_{pq}(\mathbb{K})$ définie pour tout $k \in \llbracket 1, p \rrbracket$ et $\ell \in \llbracket 1, q \rrbracket$ par

$$(E_{ij})_{k\ell} = \delta_{ik}\delta_{j\ell}$$

Remarquons que pour chaque taille $p \times q$ de matrices, il existe pq matrices élémentaires. La taille des matrices E_{ij} n'est pas indiquée dans la façon dont on les note, mais il n'y a en général aucune ambiguïté.

Proposition 2. Toute matrice $A \in \mathcal{M}_{pq}(\mathbb{K})$ s'écrit de façon unique comme combinaison linéaire des matrices élémentaires. Plus précisément,

$$A = \sum_{i=1}^p \sum_{j=1}^q A_{ij} E_{ij}$$

Démonstration. Soit $(\lambda_{ij})_{1 \leq i \leq p, 1 \leq j \leq q}$ une famille de scalaires. On a pour tous k, ℓ

$$\begin{aligned} \left(\sum_{i=1}^p \sum_{j=1}^q \lambda_{ij} E_{ij} \right)_{k\ell} &= \sum_{i=1}^p \sum_{j=1}^q \lambda_{ij} (E_{ij})_{k\ell} \\ &= \sum_{i=1}^p \sum_{j=1}^q \lambda_{ij} \delta_{ik} \delta_{j\ell} \\ &= \lambda_{k\ell} \end{aligned}$$

En effet, tous les termes de la somme double sont nuls, sauf éventuellement si $i = k$ et $j = \ell$. On en déduit existence et unicité. $\square$

1.3 Produit matriciel

Définition 4. Soient $A \in \mathcal{M}_{pq}(\mathbb{K})$ et $B \in \mathcal{M}_{qr}(\mathbb{K})$. Le produit de A et B est la matrice $AB \in \mathcal{M}_{pr}(\mathbb{K})$ définie, pour $i \in \llbracket 1, p \rrbracket$ et $j \in \llbracket 1, r \rrbracket$, par

$$(AB)_{ij} = \sum_{k=1}^q A_{ik} B_{kj}$$

Exemple. Calculer le produit de $A = \begin{pmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \end{pmatrix}$ par $B = \begin{pmatrix} 5 & 1 & 2 \\ 0 & 1 & 1 \\ -1 & 1 & 0 \end{pmatrix}$.

Proposition 3. Soient $i \in \llbracket 1, p \rrbracket$, $j, k \in \llbracket 1, q \rrbracket$ et $\ell \in \llbracket 1, r \rrbracket$. On a

$$E_{ij} E_{kl} = \delta_{jk} E_{il}$$

Démonstration. Soient $1 \leq m \leq p$ et $1 \leq n \leq r$. On a

$$\begin{aligned} (E_{ij} E_{kl})_{mn} &= \sum_{s=1}^q (E_{ij})_{ms} (E_{kl})_{sn} \\ &= \sum_{s=1}^q \delta_{im} \delta_{js} \delta_{ks} \delta_{ln} \\ &= \delta_{im} \delta_{ln} \sum_{s=1}^q \delta_{js} \delta_{ks} \\ &= \delta_{im} \delta_{ln} \delta_{jk} \\ &= (\delta_{jk} E_{il})_{mn} \end{aligned}$$

Ceci est vrai pour tous m, n , donc

$$E_{ij}E_{kl} = \delta_{jk}E_{il}$$

□

Proposition 4. Soient $A \in \mathcal{M}_{p,q}(\mathbb{K})$, $B \in \mathcal{M}_{q,r}(\mathbb{K})$ et $C \in \mathcal{M}_{r,s}(\mathbb{K})$. Alors

$$(AB)C = A(BC)$$

Démonstration. Soient $1 \leq i \leq p$ et $1 \leq j \leq s$. On a

$$\begin{aligned} (A(BC))_{ij} &= \sum_{k=1}^q A_{ik}(BC)_{kj} \\ &= \sum_{k=1}^q \sum_{\ell=1}^r A_{ik}B_{k\ell}C_{\ell j} \end{aligned}$$

et de même pour $((AB)C)_{ij}$. Ainsi, $A(BC) = (AB)C$. □

Proposition 5. Soient $A \in \mathcal{M}_{pq}(\mathbb{K})$ et $B, C \in \mathcal{M}_{qr}(\mathbb{K})$. On a

$$A(B + C) = AB + AC$$

Démonstration. Exercice. □

Remarque 3. La multiplication des matrices n'est pas commutative. Par exemple, soient $A = \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix}$ et $B = \begin{pmatrix} 0 & 0 \\ 1 & 0 \end{pmatrix}$. On a $AB = \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix}$ alors que $BA = \begin{pmatrix} 0 & 0 \\ 0 & 1 \end{pmatrix}$.

Exercice. Généraliser, en fabriquant deux matrices $n \times n$ A et B telles que $AB \neq BA$.

1.4 Transposée

Définition 5. Soit $A \in \mathcal{M}_{p,q}(\mathbb{K})$. On appelle transposée de A la matrice $B \in \mathcal{M}_{q,p}(\mathbb{K})$ définie pour $1 \leq i \leq q$ et $1 \leq j \leq p$ par $B_{ij} = A_{ji}$.

Notation. On note A^T la transposée de la matrice A .

Proposition 6. Avec des notations évidentes, on a

- Si $A, B \in \mathcal{M}_{pq}(\mathbb{K})$, $(A + B)^T = A^T + B^T$
- Si $A \in \mathcal{M}_{pq}(\mathbb{K})$ et $\lambda \in \mathbb{K}$, $(\lambda A)^T = \lambda A^T$
- Si $A \in \mathcal{M}_{pq}(\mathbb{K})$, $A^{TT} = A$
- Si $A \in \mathcal{M}_{pq}(\mathbb{K})$ et $B \in \mathcal{M}_{qr}(\mathbb{K})$, $(AB)^T = B^T A^T$
- Si $A \in \mathcal{M}_{pq}(\mathbb{K})$ est inversible, A^T est inversible et $(A^T)^{-1} = (A^{-1})^T$.

Démonstration. Par exemple,

$$((AB)^T)_{ij} = (AB)_{ji} = \sum_k A_{jk} B_{ki} = \sum_k (B^T)_{ik} (A^T)_{kj} = (B^T A^T)_{ij}$$

□

Exercice. Démontrer les autres points.

Exercice. Calculer XX^T et $X^T X$ lorsque X est une matrice colonne.

1.5 Matrices symétriques, antisymétriques

Définition 6. Une matrice carrée A est dite

- symétrique lorsque $A^T = A$
- antisymétrique lorsque $A^T = -A$.

On note

$$\mathcal{S}_n(\mathbb{K}) = \{A \in \mathcal{M}_n(\mathbb{K}), A^T = A\}$$

et

$$\mathcal{A}_n(\mathbb{K}) = \{A \in \mathcal{M}_n(\mathbb{K}), A^T = -A\}$$

Proposition 7. Soit $\mathbb{K}$ un corps dans lequel $0 \neq 2$.

Toute matrice de $\mathcal{M}_n(\mathbb{K})$ s'écrit de façon unique comme somme d'une matrice symétrique et d'une matrice antisymétrique.

Démonstration. Soit $A \in \mathcal{M}_n(\mathbb{K})$. Supposons que

$$(1) \quad A = B + C$$

où $B \in \mathcal{S}_n(\mathbb{K})$ et $C \in \mathcal{A}_n(\mathbb{K})$. On a alors

$$(2) \quad A^T = (B + C)^T = B^T + C^T = B - C$$

En additionnant et en soustrayant (1) et (2), on obtient

$$B = \frac{1}{2}(A + A^T) \text{ et } C = \frac{1}{2}(A - A^T)$$

Inversement, on vérifie facilement que B est symétrique, C est antisymétrique, et $B + C = A$. $\square$

Exemple. Soit $A = \begin{pmatrix} 1 & 4 & 7 \\ 2 & 5 & 8 \\ 3 & 6 & 9 \end{pmatrix}$. On a $A = B + C$ avec $B = \begin{pmatrix} 1 & 3 & 5 \\ 3 & 5 & 7 \\ 5 & 7 & 9 \end{pmatrix}$ et $C = \begin{pmatrix} 0 & 1 & 2 \\ -1 & 0 & 1 \\ -2 & -1 & 0 \end{pmatrix}$

2 Systèmes linéaires

2.1 Produit d'une matrice par une matrice élémentaire

Soit $A \in \mathcal{M}_{pq}(\mathbb{K})$. Soient $i, j \in \llbracket 1, p \rrbracket$. Que vaut la matrice $E_{ij}A$? Soient $k \in \llbracket 1, p \rrbracket$ et $\ell \in \llbracket 1, q \rrbracket$. On a

$$\begin{aligned} (E_{ij}A)_{k\ell} &= \sum_{m=1}^p (E_{ij})_{km} A_{m\ell} \\ &= \sum_{m=1}^p \delta_{ik} \delta_{jm} A_{m\ell} \\ &= \delta_{ik} \sum_{m=1}^p \delta_{jm} A_{m\ell} \\ &= \delta_{ik} A_{j\ell} \end{aligned}$$

Ainsi, si $k \neq i$, on trouve à la ligne k de la matrice $E_{ij}A$ une ligne de zéros. En revanche, à la ligne i de la matrice $E_{ij}A$, on trouve la ligne j de la matrice A .

Le produit à gauche de A par E_{ij} a donc eu pour effet de *déplacer* la ligne j de A vers la ligne i .

Exercice. Soit $A \in \mathcal{M}_{pq}(\mathbb{K})$. Soient $i, j \in \llbracket 1, q \rrbracket$. Que vaut la matrice AE_{ij} ? On constatera que l'on obtient cette fois-ci un déplacement de la colonne i de A vers la colonne j .

2.2 Opérations sur les lignes et les colonnes des matrices

Il existe trois types d'opérations importantes sur les lignes et les colonnes des matrices. Nous allons dans ce qui suit discuter surtout de lignes, en faisant simplement des remarques pour ce qui concerne les colonnes.

Multiplier une ligne par un scalaire

Soit $A \in \mathcal{M}_{pq}(\mathbb{K})$. Soit $\lambda \in \mathbb{K}$. Soit $i \in \llbracket 1, p \rrbracket$. L'opération consistant à multiplier la i ème ligne de A par λ sera notée

$$L_i \longleftarrow \lambda L_i$$

Cette opération peut être faite par un produit matriciel. Pour cela, il suffit de multiplier A à gauche par $I + (\lambda - 1)E_{ii}$.

Proposition 8. *Si $\lambda \neq 0$, la matrice $I + (\lambda - 1)E_{ii}$ est inversible. Son inverse est $I + (\frac{1}{\lambda} - 1)E_{ii}$.*

Démonstration. Soient $B = I + (\lambda - 1)E_{ii}$ et $B' = I + (\frac{1}{\lambda} - 1)E_{ii}$. Pour toute matrice A , le produit $B'BA$ a pour effet

- De multiplier la ligne i de A par λ , puis
- De multiplier la ligne i de A par $\frac{1}{\lambda}$.

Ainsi, $B'BA = A$. En particulier, en prenant $A = I$, on obtient $B'B = I$. De même, $BB' = I$.

□

Échanger deux lignes

Soit $A \in \mathcal{M}_{pq}(\mathbb{K})$. Soient $i, j \in \llbracket 1, p \rrbracket$. L'opération consistant à échanger la i ème ligne et la j ème ligne de A sera notée

$$L_i \longleftrightarrow L_j$$

Cette opération peut être faite par un produit matriciel. Pour cela, il suffit de multiplier A à gauche par $I + E_{ij} + E_{ji} - E_{ii} - E_{jj}$.

Proposition 9. *La matrice $I + E_{ij} + E_{ji} - E_{ii} - E_{jj}$ est inversible. Elle est son propre inverse.*

Démonstration. Soit $B = I + E_{ij} + E_{ji} - E_{ii} - E_{jj}$. Pour toute matrice A , le produit B^2A a pour effet

- D'échanger les lignes i et j de A , puis
- D'échanger les lignes i et j de A .

Ainsi, $B^2A = A$. En particulier, en prenant $A = I$, on obtient $B^2 = I$. □

Ajouter à une ligne une combinaison linéaire des autres lignes

Soit $A \in \mathcal{M}_{pq}(\mathbb{K})$. Soit $\lambda \in \mathbb{K}$. Soient $i, j \in \llbracket 1, p \rrbracket$. L'opération consistant à ajouter à la i ème ligne λ fois la j ème ligne de A sera notée

$$L_i \longleftarrow L_i + \lambda L_j$$

Cette opération peut être faite par un produit matriciel. Pour cela, il suffit de multiplier A à gauche par $I + \lambda E_{ji}$.

Proposition 10. *Si $i \neq j$, la matrice $I + \lambda E_{ji}$ est inversible. Son inverse est $I - \lambda E_{ji}$.*

Démonstration. Soient $B = I + \lambda E_{ji}$ et $B' = I - \lambda E_{ji}$. Pour toute matrice A , le produit $B'BA$ a pour effet

- D'ajouter à la ligne i de A λ fois la ligne j , puis
- De soustraire à la ligne i de A λ fois la ligne j .

Comme $i \neq j$, la première opération n'a pas modifié la ligne j . Ainsi, $B'BA = A$. En particulier, en prenant $A = I$, on obtient $B'B = I$. De même, $BB' = I$. $\square$

2.3 Systèmes linéaires

Un système linéaire de p équations à q inconnues est un système d'équations de la forme

$$(S) \quad \begin{cases} a_{11}x_1 + a_{12}x_2 + \dots + a_{1q}x_q &= b_1 \\ a_{21}x_1 + a_{22}x_2 + \dots + a_{2q}x_q &= b_2 \\ \dots & \\ a_{p1}x_1 + a_{p2}x_2 + \dots + a_{pq}x_q &= b_p \end{cases}$$

où les a_{ij} et les b_i sont donnés, et les x_j sont les inconnues. Notons

$$A = (a_{ij})_{1 \leq i \leq p, 1 \leq j \leq q} \in \mathcal{M}_{pq}(\mathbb{K})$$

$$X = (x_1 \ x_2 \ \dots \ x_q)^T \in \mathcal{M}_{q1}(\mathbb{K})$$

$$B = (b_1 \ b_2 \ \dots \ b_p)^T \in \mathcal{M}_{p1}(\mathbb{K})$$

Le système peut alors être écrit de façon beaucoup plus compacte

$$(S) \quad AX = B$$

La matrice A est la *matrice du système*. La matrice B est le *second membre*. La matrice X est la *matrice des inconnues*.

Une *solution* de (S) est donc une matrice $X \in \mathcal{M}_{q1}(\mathbb{K})$ telle que $AX = B$. *Résoudre* (S) c'est trouver l'ensemble des solutions de (S) , que nous noterons Sol_S . Au système (S) nous pouvons systématiquement associer le *système homogène*

$$(H) \quad AX = 0$$

que les mauvaises langues appellent aussi le système *sans second membre* (ce qui est absurde car 0, ce n'est pas rien).

Le système (S) est dit *compatible* lorsqu'il possède au moins une solution. Nous avons le résultat facile suivant.

Proposition 11. Soit $(S) AX = B$ un système compatible de p équations à q inconnues. Soit X_0 une solution de (S) . Soit $X \in \mathcal{M}_{pq}(\mathbb{K})$. Alors, $X \in \text{Sol}_S$ si et seulement si

$$X = X_0 + Y$$

où $Y \in \text{Sol}_H$.

Démonstration. On a

$$AX = B \iff AX = AX_0 \iff A(X - X_0) = 0 \iff X - X_0 \in \text{Sol}_H$$

□

Autre résultat important,

Proposition 12. Le système $(S) AX = B$ est compatible si et seulement si B est combinaison linéaire des colonnes de A .

Démonstration. Notons $C_1, \dots, C_q$ les colonnes de A . La proposition est démontrée en remarquant simplement que le système peut aussi s'écrire

$$(S) x_1 C_1 + \dots + x_q C_q = B$$

□

2.4 L'algorithme du pivot de Gauss

Définition 7. Soient $(S)AX = B$ et $(S')A'X = B$ deux systèmes de p équations à q inconnues. On dit que (S) et (S') sont *équivalents* lorsque $\text{Sol}_S = \text{Sol}_{S'}$.

Proposition 13. Les opérations élémentaires sur les équations d'un système linéaire transforment ce système en un système équivalent.

Démonstration. Soit $(S) AX = B$ un système de p équations à q inconnues. Rappelons les trois opérations élémentaires.

- $L_i \leftarrow \lambda L_i$, où $\lambda \neq 0$.
- $L_i \longleftrightarrow L_j$.
- $L_i \leftarrow L_i + \lambda L_j$, où $j \neq i$

Nous avons vu que ces trois opérations sont réalisées en multipliant à gauche par une matrice M judicieuse. On obtient donc le système

$$(S') \quad MAX = MB$$

Mais M est inversible, donc

$$MAX = MB \iff AX = B$$

Ainsi, $\text{Sol}_S = \text{Sol}_{S'}$. $\square$

Considérons le système $(S) \quad AX = B$ de p équations à q inconnues.

L'algorithme du pivot de Gauss permet d'éliminer une inconnue de (S) de toutes les équations sauf une. Supposons que $A_{kj} \neq 0$.

On effectue les opérations suivantes sur les lignes de A ET de B :

- $L_k \leftarrow \frac{1}{A_{kj}} L_k$. Ceci a pour effet de ramener le coefficient A_{kj} à 1.
- Pour $i \neq k$, $L_i \leftarrow L_i - A_{ij} L_k$. Ceci a pour effet d'annuler tous les coefficients de la colonne j de A (sauf A_{kj} , qui reste égal à 1).

On obtient ainsi un nouveau système

$$(S') \quad A'X = B'$$

qui a les mêmes solutions que le système (S) , mais pour lequel x_j n'apparaît dans aucune des équations, sauf dans l'équation k .

En réitérant cette opération, on résout le système. Regardons deux exemples.

2.5 Un exemple

Considérons le système de 3 équations à 4 inconnues.

$$(S) \quad \begin{cases} x & & & + 4t & = & 1 \\ 2x & + & 3y & - & z & & = & 3 \\ x & + & 3y & - & z & - & 4t & = & 2 \end{cases}$$

- Étape 1. Le pivot est $A_{11} = 1$. On effectue les opérations $L_2 \leftarrow L_2 - 2L_1$ et $L_3 \leftarrow L_3 - L_1$. On obtient un nouveau système équivalent

$$\begin{cases} x & & & + 4t & = & 1 \\ & 3y & - & z & - & 8t & = & 1 \\ & 3y & - & z & - & 8t & = & 1 \end{cases}$$

- Étape 2. Le pivot est le $A_{22} = 3$. Inutile, dans la pratique, de diviser cette équation par 3 ! On effectue l'opération $L_3 \leftarrow L_3 - L_2$. On obtient un nouveau système équivalent

$$\begin{cases} x & & + 4t & = & 1 \\ & 3y & - z & - 8t & = & 1 \\ & & & 0 & = & 0 \end{cases}$$

La résolution est terminée. On obtient toutes les solutions du système en choisissant arbitrairement les inconnues sur lesquelles on n'a pas « pivoté ». Les inconnues sur lesquelles on a pivoté s'expriment en fonction de celles-ci. Ici, on choisit arbitrairement z et t , et on obtient

$$\begin{cases} x & = & 1 - 4t \\ y & = & \frac{1}{3}(1 + z + 8t) \end{cases}$$

Dit autrement, l'ensemble des solutions du système est

$$\text{Sol}_S = \left\{ \left(1 - 4t, \frac{1}{3}(1 + z + 8t), z, t \right), z, t \in \mathbb{R} \right\}$$

Remarque 4. En mettant 1 au lieu de 2 dans le second membre de la troisième équation de (S) , on aurait obtenu pour dernière équation à la fin de l'étape 2, $0 = -2$. Le système (S) aurait donc été incompatible.

2.6 Un autre exemple

Considérons le système

$$(S) \begin{cases} x + 2y + 3z & = & 6 \\ 4x + 5y + 6z & = & 15 \\ 7x + 8y + 9z & = & 24 \end{cases}$$

- Étape 1 : le pivot est $A_{11} = 1$. On effectue les opérations

$$L_2 \leftarrow L_2 - 4L_1$$

$$L_3 \leftarrow L_3 - 7L_1$$

On obtient un système équivalent à (S)

$$(S_1) \begin{cases} x + 2y + 3z & = & 6 \\ -3y - 6z & = & -9 \\ -6y - 12z & = & -18 \end{cases}$$

- Étape 2 : le pivot est $A_{22} = -3$. On effectue les opérations

$$L_2 \leftarrow -\frac{1}{3}L_2$$

$$L_1 \leftarrow L_1 - 2L_2$$

$$L_3 \longleftarrow L_3 - (-6)L_2$$

On obtient un nouveau système équivalent à (S)

$$(S_2) \begin{cases} x - z &= 0 \\ y + 2z &= 3 \\ 0 &= 0 \end{cases}$$

La résolution est terminée. On prend z arbitrairement, on a x et y en fonction de z .

$$\text{Sol}_S = \{z, 3 - 2z, z\}, z \in \mathbb{R}\}$$

3 L'anneau des matrices carrées

3.1 Structure d'anneau

Proposition 14. Soit n un entier non nul. $\mathcal{M}_n(\mathbb{K})$ est un anneau, non commutatif si $n \geq 2$.

Démonstration. Nous avons déjà presque tout vérifié. Le neutre pour la multiplication des matrices est la *matrice identité d'ordre n*

$$I_n = \begin{pmatrix} 1 & 0 & \dots & 0 \\ 0 & 1 & \dots & 0 \\ & & \ddots & \\ 0 & 0 & \dots & 1 \end{pmatrix}$$

Gardons la non commutativité pour un peu plus tard, nous exhiberons même des matrices *inversibles* qui ne commutent pas. $\square$

Tout ce que nous avons vu sur les anneaux s'applique donc aux matrices carrées. En particulier, si $A, B \in \mathcal{M}_p(\mathbb{K})$ vérifient $AB = BA$, et $n \in \mathbb{N}$, nous avons

- La formule du binôme de Newton :

$$(A + B)^n = \sum_{k=0}^n \binom{n}{k} A^k B^{n-k}$$

- L'identité remarquable

$$A^{n+1} - B^{n+1} = (A - B) \sum_{k=0}^n A^k B^{n-k}$$

Un cas particulier important de la seconde égalité est le cas où l'une des deux matrices est la matrice identité. Pour tout $n \geq 1$,

$$I - A^n = (I - A) \sum_{k=0}^{n-1} A^k$$

3.2 Matrices inversibles

On note $GL_n(\mathbb{K})$ l'ensemble des matrices carrées inversibles de $\mathcal{M}_n(\mathbb{K})$.

Proposition 15. $GL_n(\mathbb{K})$ est un groupe.

Démonstration. $GL_n(\mathbb{K})$ est l'ensemble des éléments inversibles de l'anneau $\mathcal{M}_n(\mathbb{K})$. C'est donc un groupe. $\square$

Le cas $n = 2$ est facile à traiter.

Proposition 16. Soit $A = \begin{pmatrix} a & c \\ b & d \end{pmatrix} \in \mathcal{M}_2(\mathbb{K})$. La matrice A est inversible si et seulement si $ad - bc \neq 0$. On a alors

$$A^{-1} = \frac{1}{ad - bc} \begin{pmatrix} d & -c \\ -b & a \end{pmatrix}$$

La quantité $ad - bc$ est appelée le déterminant de A .

Démonstration. On a

$$\begin{pmatrix} a & c \\ b & d \end{pmatrix} \begin{pmatrix} d & -c \\ -b & a \end{pmatrix} = (ad - bc)I_2$$

Si $ad - bc = 0$, la matrice A est un diviseur de 0 dans l'anneau $\mathcal{M}_2(\mathbb{K})$. Elle n'est donc pas inversible. Si, au contraire, $ad - bc \neq 0$, on peut diviser l'égalité par $ad - bc$. On en déduit l'inversibilité de A et la valeur de A^{-1} . $\square$

Proposition 17. Si $n \geq 2$, $GL_n(\mathbb{K})$ est un groupe non commutatif.

Démonstration. Commençons par donner un exemple lorsque $n = 2$. Soient

$$A = \begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix} \text{ et } B = \begin{pmatrix} 1 & 0 \\ 1 & 1 \end{pmatrix}$$

On a

$$AB = \begin{pmatrix} 2 & 1 \\ 1 & 1 \end{pmatrix} \text{ et } BA = \begin{pmatrix} 1 & 1 \\ 1 & 2 \end{pmatrix}$$

Ainsi, $AB \neq BA$. De plus, A et B ont un déterminant non nul, donc elles sont inversibles.

Pour $n \geq 2$ quelconque, on peut bâtir un exemple de matrices qui ne commutent pas en prenant

$$A = \begin{pmatrix} 1 & 1 & 0 & \dots & 0 \\ 0 & 1 & 0 & \dots & 0 \\ 0 & 0 & 1 & \dots & 0 \\ \vdots & & & \ddots & \\ 0 & 0 & 0 & \dots & 1 \end{pmatrix} \quad \text{et} \quad B = \begin{pmatrix} 1 & 0 & 0 & \dots & 0 \\ 1 & 1 & 0 & \dots & 0 \\ 0 & 0 & 1 & \dots & 0 \\ \vdots & & & \ddots & \\ 0 & 0 & 0 & \dots & 1 \end{pmatrix}$$

□

Proposition 18. Soit $A \in \mathcal{M}_n(\mathbb{K})$. A est inversible si et seulement si elle est inversible à gauche, si et seulement si elle est inversible à droite.

Démonstration. Nous admettons provisoirement ce résultat. Lorsque nous aurons, au second semestre, fait le lien entre les matrices et les applications linéaires, il sera facile de le démontrer grâce à un théorème qui s'appelle le *théorème du rang*. □

Remarque 5. Grâce à ce même théorème du rang, on peut montrer qu'une matrice qui n'est pas carrée ne peut pas être inversible. Elle peut en revanche être inversible à gauche ou à droite.

3.3 Pratique de l'inversion des matrices

Proposition 19. Soit $A \in \mathcal{M}_n(\mathbb{K})$. Pour tout $Y \in \mathcal{M}_{n1}(\mathbb{K})$, notons

$$(S_Y) \quad AX = Y$$

le système d'inconnue $X \in \mathcal{M}_{n1}(\mathbb{K})$.

La matrice A est inversible, si et seulement si pour tout $Y \in \mathcal{M}_{n1}(\mathbb{K})$, (S_Y) a une unique solution. L'unique solution de (S_Y) est alors

$$X = A^{-1}Y$$

Démonstration. Supposons A inversible. On a pour tous $X, Y \in \mathcal{M}_{n1}(\mathbb{K})$,

$$AX = Y \iff A^{-1}(AX) = A^{-1}Y \iff X = A^{-1}Y$$

Inversement, supposons que pour tout $Y \in \mathcal{M}_{n1}(\mathbb{K})$, le système (S) a une unique solution. Pour $i = 1, \dots, n$, soit $Y_j = (0 \ 0 \dots 0 \ 1 \ 0 \dots 0)^T$ où le 1 apparaît en position j . Soit X_j l'unique solution du système S_{Y_j} . Soit B la matrice dont les colonnes sont les X_j . La matrice AB est alors la matrice dont les colonnes sont les Y_j , c'est à dire I_n . Ainsi, A est inversible à droite, donc inversible. □

Une technique pratique d'inversion de matrice est donc la suivante : on résout le système d'équations $AX = Y$ avec un second membre Y « quelconque ». Si, pour un certain Y , le système n'a pas de solution (ou bien en a plusieurs), alors A n'est pas inversible. Sinon, on obtient la solution unique $X = A^{-1}Y$ et on en déduit A^{-1} .

3.4 Un exemple

Inversons la matrice $A = \begin{pmatrix} 1 & 2 & 1 \\ -1 & 3 & 1 \\ 0 & 1 & 0 \end{pmatrix}$. Donnons nous $Y = (a \ b \ c)^T \in \mathcal{M}_{31}(\mathbb{R})$. Soit $X = (x \ y \ z)^T \in \mathcal{M}_{31}(\mathbb{R})$. Résolvons le système $AX = Y$ par le pivot de Gauss. Le système s'écrit

$$(S) \quad \begin{cases} x + 2y + z &= a \\ -x + 3y + z &= b \\ y &= c \end{cases}$$

On en déduit facilement que (S) a pour unique solution

$$\begin{cases} x &= \frac{1}{2}a - \frac{1}{2}b + \frac{1}{2}c \\ y &= c \\ z &= \frac{1}{2}a + \frac{1}{2}b - \frac{5}{2}c \end{cases}$$

Ceci est vrai pour tout Y . La matrice A est donc inversible, et de plus

$$A^{-1} = \begin{pmatrix} \frac{1}{2} & -\frac{1}{2} & \frac{1}{2} \\ 0 & 0 & 1 \\ \frac{1}{2} & \frac{1}{2} & -\frac{5}{2} \end{pmatrix}$$

4 matrices triangulaires

4.1 C'est quoi ?

Définition 8. Soit $A \in \mathcal{M}_n(\mathbb{K})$. On dit que la matrice A est triangulaire supérieure lorsque $\forall i > j, A_{ij} = 0$. De même, on dit que A est triangulaire inférieure lorsque $\forall i < j, A_{ij} = 0$. Enfin, on dit que A est diagonale lorsque $\forall i \neq j, A_{ij} = 0$.

Dans la suite, on note $\mathcal{T}_n(\mathbb{K})$ l'ensemble des matrices triangulaires supérieures à coefficients dans le corps $\mathbb{K}$. Tous les résultats énoncés sont également vrais pour des matrices triangulaires inférieures.

Remarque 6. Soit $A \in \mathcal{M}_n(\mathbb{K})$. La *diagonale* de A est l'ensemble des coefficients A_{ii} , $1 \leq i \leq n$. Ainsi,

- A est triangulaire supérieure lorsque les coefficients de A strictement au-dessous de sa diagonale sont nuls.

- A est triangulaire inférieure lorsque les coefficients de A strictement au-dessus de sa diagonale sont nuls.
- A est diagonale lorsque les coefficients de A en dehors de sa diagonale sont nuls.

4.2 Structure d'anneau

Proposition 20. $(\mathcal{T}_n(\mathbb{K}), +, \times)$ est un anneau, non commutatif dès que $n \geq 2$.

Démonstration. Montrons que $\mathcal{T}_n(\mathbb{K})$ est un sous-anneau de $\mathcal{M}_n(\mathbb{K})$.

La matrice identité I_n , neutre pour la multiplication des matrices, est évidemment triangulaire.

La stabilité par soustraction a déjà été montrée.

Le seul point non évident est la stabilité par produit. Soient $A, B \in \mathcal{T}_n(\mathbb{K})$. On a

$$(AB)_{i,j} = \sum_{k=1}^n A_{ik}B_{kj}$$

Remarquons que

$$A_{ik}B_{kj} \neq 0 \implies A_{ik} \neq 0 \wedge B_{kj} \neq 0 \implies i \leq k \leq j$$

Ainsi, si $i > j$, $(AB)_{ij} = 0$. Bilan : le produit de deux matrices triangulaires supérieures est encore une matrice triangulaire supérieure.

Pour la non-commutativité, considérons tout d'abord le cas $n = 2$. Soient $A = \begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix}$ et $B = \begin{pmatrix} 1 & 1 \\ 0 & 2 \end{pmatrix}$. On a $AB = \begin{pmatrix} 1 & 3 \\ 0 & 2 \end{pmatrix}$ et $BA = \begin{pmatrix} 1 & 2 \\ 0 & 2 \end{pmatrix}$. Ainsi, $AB \neq BA$.

On laisse au lecteur le soin de trouver pour $n \geq 2$ quelconque deux matrices triangulaires supérieures qui ne commutent pas. $\square$

4.3 Le cas des matrices diagonales

Regardons tout d'abord un cas facile, celui des matrices diagonales. Notons $A = \text{Diag}(\alpha_1, \dots, \alpha_n)$ une matrice diagonale dont les coefficients de la diagonale sont les scalaires α_i . Notons $\mathcal{D}_n(\mathbb{K})$ l'ensemble des matrices diagonales de $\mathcal{M}_n(\mathbb{K})$.

Proposition 21. $\mathcal{D}_n(\mathbb{K})$ est un anneau commutatif. Une matrice diagonale est inversible si et seulement si ses coefficients diagonaux sont non nuls.

Démonstration. On vérifie facilement les résultats suivants :

- $\text{Diag}(\alpha_1, \dots, \alpha_n) - \text{Diag}(\beta_1, \dots, \beta_n) = \text{Diag}(\alpha_1 - \beta_1, \dots, \alpha_n - \beta_n)$.
- $\text{Diag}(\alpha_1, \dots, \alpha_n)\text{Diag}(\beta_1, \dots, \beta_n) = \text{Diag}(\alpha_1\beta_1, \dots, \alpha_n\beta_n)$.
- $I_n = \text{Diag}(1, \dots, 1)$.

Ainsi, $\mathcal{D}_n(\mathbb{K})$ est un sous-anneau de $\mathcal{M}_n(\mathbb{K})$. La commutativité est claire par la formule donnant le produit de deux matrices diagonales. Enfin, $\text{Diag}(\alpha_1, \dots, \alpha_n)\text{Diag}(\beta_1, \dots, \beta_n) = I_n$ si et seulement si pour tout $i \in \llbracket 1, n \rrbracket$, $\alpha_i\beta_i = 1$. La matrice $\text{Diag}(\alpha_1, \dots, \alpha_n)$ est donc inversible si et seulement si les α_i sont non nuls. Son inverse est alors $\text{Diag}(\frac{1}{\alpha_1}, \dots, \frac{1}{\alpha_n})$. $\square$

Remarque 7. Soit $A = \text{Diag}(\alpha_1, \dots, \alpha_n)$. On a pour tout entier naturel p ,

$$A^p = \text{Diag}(\alpha_1^p, \dots, \alpha_n^p)$$

Si A est inversible, ceci est encore vrai lorsque p est entier négatif.

4.4 Matrices triangulaires inversibles

Proposition 22. Soit $A \in \mathcal{T}_n(\mathbb{K})$. La matrice A est inversible si et seulement si ses coefficients diagonaux sont non nuls. On a alors $A^{-1} \in \mathcal{T}_n(\mathbb{K})$.

Démonstration. Soit $A \in \mathcal{T}_n(\mathbb{K})$. Soient $X, Y \in \mathcal{M}_{n1}(\mathbb{K})$. On a $AX = Y$ si et seulement si pour tout $i \in \llbracket 1, n \rrbracket$,

$$\sum_{j=1}^n A_{ij}x_j = y_i$$

Si $i > j$, $A_{ij} = 0$. L'égalité ci-dessus est donc équivalente à

$$\sum_{j=i}^n A_{ij}x_j = y_i$$

ou encore

$$(E_i) \quad A_{ii}x_i = y_i - \sum_{j=i+1}^n A_{ij}x_j$$

Supposons tous les A_{ii} non nuls. Le système ci-dessus admet une unique solution, donnée par

$$x_n = \frac{1}{A_{nn}}y_n$$

et pour tout $i \in \llbracket 1, n-1 \rrbracket$

$$x_i = \frac{1}{A_{ii}} \left(y_i - \sum_{j=i+1}^n A_{ij}x_j \right)$$

Le système $AX = Y$ a donc une unique solution, ceci pour tout Y , ce qui prouve que A est inversible. Nous avons même « calculé » l'inverse de A . On constate que cet inverse est une matrice triangulaire supérieure car pour tout $i \in \llbracket 1, n \rrbracket$, x_i est combinaison linéaire de $y_i, \dots, y_n$ (récurrence facile).

Supposons maintenant que l'un des A_{ii} est nul. Prenons un tel i maximal. Pour tout Y , les équations $E_n, E_{n-1}, \dots, E_{i+1}$ fournissent une unique solution $x_n, x_{n-1}, \dots, x_{i+1}$. Si l'on choisit y_i tel que $y_i \neq \sum_{j=i+1}^n A_{ij}x_j$, l'équation i n'a alors pas de solution. Il existe donc Y tel que l'équation $AX = Y$ n'ait pas de solution. La matrice A n'est donc pas inversible.

□

Exemple. Soit

$$A = \begin{pmatrix} 1 & 2 & 3 & 4 \\ 0 & 1 & 5 & 6 \\ 0 & 0 & 1 & 7 \\ 0 & 0 & 0 & 1 \end{pmatrix}$$

Il s'agit de résoudre le système

$$\begin{cases} a + 2b + 3c + 4d = x \\ b + 5c + 6d = y \\ c + 7d = z \\ d = t \end{cases}$$

Il se résout en

$$\begin{cases} a = x - 2(y - 5z + 29t) - 3(z - 7t) - 4t = x - 2y + 7z - 41t \\ b = y - 5(z - 7t) - 6t = y - 5z + 29t \\ c = z - 7t \\ d = t \end{cases}$$

Ainsi,

$$A^{-1} = \begin{pmatrix} 1 & -2 & 7 & -41 \\ 0 & 1 & -5 & 29 \\ 0 & 0 & 1 & -7 \\ 0 & 0 & 0 & 1 \end{pmatrix}$$

Chapitre 18

Polynômes

1 L'algèbre des polynômes

1.1 C'est quoi ?

Dans tout le chapitre, $\mathbb{K}$ désigne un corps.

Proposition 1. *Il existe un couple $(\mathbb{A}, X)$ vérifiant*

- $\mathbb{A}$ est une $\mathbb{K}$ algèbre commutative contenant X .
- Tout élément P de $\mathbb{A}$ s'écrit de façon unique

$$P = \sum_{k=0}^{\infty} a_k X^k$$

où les $a_k \in \mathbb{K}$ sont tous nuls sauf un nombre fini (nous dirons presque tous nuls).

Si $(\mathbb{A}, X)$ et $(\mathbb{B}, Y)$ sont deux tels couples, il existe un unique isomorphisme d'algèbres $f : \mathbb{A} \rightarrow \mathbb{B}$ tel que $f(X) = Y$.

Démonstration. Admise. Remarquons que toutes les propriétés des espaces vectoriels et des anneaux commutatifs s'appliquent à $\mathbb{A}$. On a ainsi par exemple les identités remarquables, la formule du binôme de Newton, etc. $\square$

Remarque 1.

- X est appelé *l'indéterminée*.
- L'algèbre $\mathbb{A}$ sera dorénavant notée $\mathbb{K}[X]$.
- Les éléments de $\mathbb{K}[X]$ sont appelés les *polynômes* à coefficients dans $\mathbb{K}$.
- Les a_k sont les *coefficients* du polynôme P .

1.2 Degré

Définition 1. Soit $P = \sum_{k=0}^{\infty} a_k X^k \in \mathbb{K}[X] \setminus \{0\}$. On appelle degré de P l'entier

$$d^o P = \max\{k \in \mathbb{N}, a_k \neq 0\}$$

On pose également $d^o 0 = -\infty$.

On convient dans la suite que pour tout $n \in \mathbb{N}$, $-\infty < n$ et $-\infty + n = -\infty$. On convient également que $-\infty - \infty = -\infty$.

Remarque 2. Soit $d \in \mathbb{N}$.

- Si $a_d \neq 0$ alors $d^o P \geq d$.
- Si $\forall k > d, a_k = 0$, alors $d^o P \leq d$.

1.3 Somme et produit par un scalaire

Soient $P = \sum_{k=0}^{\infty} a_k X^k$ et $Q = \sum_{k=0}^{\infty} b_k X^k$ deux polynômes. Soit $\lambda \in \mathbb{K}$. On a

$$P + Q = \sum_{k=0}^{\infty} (a_k + b_k) X^k$$

$$\lambda P = \sum_{k=0}^{\infty} (\lambda a_k) X^k$$

Proposition 2.

- Pour tous $P, Q \in \mathbb{K}[X]$, $d^o(P + Q) \leq \max(d^o P, d^o Q)$.
- Si $d^o P \neq d^o Q$ on a égalité ci-dessus. Sinon, on ne peut rien dire a priori.
- Pour tout $P \in \mathbb{K}[X]$, pour tout $\lambda \in \mathbb{K}^*$, $d^o(\lambda P) = d^o P$.
- Pour tout $P \in \mathbb{K}[X]$, $d^o(0P) = -\infty$.

Démonstration. Faisons la preuve pour la somme. Le produit par un scalaire est laissé au lecteur. Si P ou Q est nul c'est évident. Supposons donc P et Q non nuls. Posons $P = \sum_{k=0}^{\infty} a_k X^k$ et $Q = \sum_{k=0}^{\infty} b_k X^k$. Soient m et n les degrés respectifs de P et Q . Par exemple, $m \leq n$. Soit $k > n$. On a $a_k = b_k = 0$ donc $a_k + b_k = 0$. Ainsi, $d^o(P + Q) \leq n$.

Si, de plus $m < n$, alors $a_n = 0$ donc $a_n + b_n = b_n \neq 0$ donc $d^o(P + Q) = n = \max(d^o P, d^o Q)$.

Si $d^o P = d^o Q$ on ne peut rien dire de plus que l'inégalité. Prenons par exemple

$$P = X^3 + X^2 - 2X + 5 \quad \text{et} \quad Q = -X^3 + X + 7$$

On a alors

$$P + Q = X^2 - X + 12$$

qui est de degré $2 < 3$. Si, en revanche, $Q = X^3 + X + 7$, alors $P + Q = 2X^3 - X + 12$ est de degré 3. $\square$

Vocabulaire :

Soit $P = \sum_{k=0}^{\infty} a_k X^k \in \mathbb{K}[X]$.

- Pour tout $k \in \mathbb{N}$, a_k est le *coefficient de degré k* de P .
- Pour tout $k \in \mathbb{N}$, $a_k X^k$ est le *terme de degré k* de P .
- Si $P \neq 0$ est de degré d , a_d est le *coefficient dominant* de P et $a_d X^d$ est le *terme de plus haut degré*. Si $a_d = 1$ on dit que P est *unitaire*. Nous noterons $cd(P)$ le coefficient dominant du polynôme P . Remarquons que le polynôme nul n'a pas de coefficient dominant.

- Un polynôme P est dit *constant* lorsque $d^o P \leq 0$.

Notation : Pour tout $n \in \mathbb{N}$, on note $\mathbb{K}_n[X]$ l'ensemble des polynômes de degré inférieur ou égal à n .

En anticipant un peu sur la théorie des espaces vectoriels, nous avons le résultat suivant.

Proposition 3. *Pour tout $n \in \mathbb{N}$, $\mathbb{K}_n[X]$ est un sous-espace vectoriel de $\mathbb{K}[X]$ de dimension $n + 1$.*

Démonstration. En effet, c'est l'espace engendré par les polynômes $1, X, \dots, X^n$. $\square$

1.4 Produit

Soient $P = \sum_{k=0}^{\infty} a_k X^k$ et $Q = \sum_{k=0}^{\infty} b_k X^k$. Que valent les coefficients du produit PQ ? Calculons-les. On a

$$PQ = \left(\sum_{i=0}^{\infty} a_i X^i \right) \left(\sum_{j=0}^{\infty} b_j X^j \right) = \sum_{(i,j) \in \mathbb{N}^2} a_i b_j X^{i+j}$$

Pour tout $(i, j) \in \mathbb{N}^2$, posons $k = i + j$. On a $i \geq 0$ et $j \geq 0$ si et seulement si $k \geq 0$ et $0 \leq i \leq k$. Ainsi en prenant pour indice dans la somme ci-dessus le couple (i, k) ,

$$\sum_{i,j \in \mathbb{N}} a_i b_j X^{i+j} = \sum_{k=0}^{\infty} c_k X^k$$

où

$$c_k = \sum_{i=0}^k a_i b_{k-i}$$

Proposition 4. *Soient $P, Q \in \mathbb{K}[X]$. On a $d^o(PQ) = d^o P + d^o Q$.*

Démonstration. C'est évident si P ou Q est nul. Supposons donc nos deux polynômes non nuls. Soient m et n leurs degrés respectifs. Notons comme d'habitude $P = \sum_{k=0}^{\infty} a_k X^k$ et $Q = \sum_{k=0}^{\infty} b_k X^k$. Notons aussi $PQ = \sum_{k=0}^{\infty} c_k X^k$.

- Soit $k > m + n$. On a

$$c_k = \sum_{i=0}^k a_i b_{k-i} = \sum_{i=0}^k a_i b_{k-i} + \sum_{i=m+1}^k a_i b_{k-i}$$

Dans la seconde somme, $a_i = 0$ puisque $i > m = d^o P$. Dans la première somme, $k - i \geq k - m > n$, donc $b_{k-i} = 0$. Ainsi, $c_k = 0$ et donc $d^o(PQ) \leq m + n$.

- Le même calcul montre que $c_{m+n} = a_m b_n \neq 0$. Ainsi, $d^o(PQ) \geq m + n$.

□

Remarque 3. Si P et Q sont deux polynômes non nuls, alors $PQ \neq 0$ et

$$cd(PQ) = cd(P)cd(Q)$$

où $cd(P)$ désigne le coefficient dominant de P .

Corollaire 5. *L'anneau $\mathbb{K}[X]$ est un anneau intègre.*

Démonstration. Soient $P, Q \in \mathbb{K}[X]$. On a $PQ = 0$ si et seulement si $d^o(PQ) = -\infty$, ou encore $d^o P + d^o Q = -\infty$. Les degrés ne peuvent donc pas être deux entiers. Donc $d^o P = -\infty$ ou $d^o Q = -\infty$, c'est à dire $P = 0$ ou $Q = 0$. □

Proposition 6. *Soit $P \in \mathbb{K}[X] \setminus \{0\}$. Le polynôme P est inversible si et seulement si il est constant et non nul, c'est à dire si et seulement si $d^o P = 0$.*

Démonstration. Soit $\lambda \in \mathbb{K}^*$. Soit $P = \lambda$. Alors $PQ = 1$, où $Q = \frac{1}{\lambda}$. Donc P est inversible. Supposons maintenant P inversible. Il existe $Q \in \mathbb{K}[X]$ non nul tel que $PQ = 1$. En passant aux degrés, on obtient $d^o P = d^o Q = d^o 1 = 0$. Mais les degrés de P et Q sont des entiers naturels, donc $d^o P = d^o Q = 0$. □

1.5 Composée

Définition 2. Soit $P = \sum_{k=0}^{\infty} a_k X^k \in \mathbb{K}[X]$. Soit $Q \in \mathbb{K}[X]$. On pose

$$P \circ Q = \sum_{k=0}^{\infty} a_k Q^k$$

Proposition 7. *Avec les notations ci-dessus, si Q n'est pas constant alors*

$$d^o(P \circ Q) = d^o P \times d^o Q$$

Démonstration. L'hypothèse sur Q dit que $d^o Q \geq 1$. C'est évident si $P = 0$. Si $P \neq 0$, posons $m = d^o P$. On a

$$P \circ Q = a_m Q^m + \sum_{k=0}^{m-1} a_k Q^k$$

Pour tout $k \leq m-1$, on a $d^o(a_k Q^k) \leq d^o Q^k = k d^o Q < m d^o Q$. Donc le degré de la somme est strictement inférieur à $m d^o Q$. De plus, $a_m \neq 0$ donc $d^o(a_m Q^m) = d^o Q^m = m d^o Q$. En appliquant une nouvelle fois le théorème sur le degré d'une somme, il vient $d^o(P \circ Q) = m d^o Q = d^o P \times d^o Q$. $\square$

Remarque 4. Si Q est constant, c'est à dire si $d^o Q \leq 0$, alors pour tout $k \in \mathbb{N}$,

$$d^o Q^k = k d^o Q \leq 0$$

Ainsi, $P \circ Q$ est constant lui aussi. Son degré est donc $-\infty$ ou 0. Pour certaines valeurs de Q , le degré peut être $-\infty$, même si P et Q ne sont pas nuls. Par exemple,

$$(X^2 - 3X + 2) \circ 2 = 0$$

Nous en reparlerons longuement plus loin lorsque nous parlerons des *racines* d'un polynôme.

Proposition 8. *La composition des polynômes est associative.*

Démonstration. Dans un premier temps, on montre que pour tous polynômes P, Q, R , on a $(P+Q) \circ R = P \circ R + Q \circ R$. Il suffit pour cela de revenir à la définition de la composée.

On montre ensuite que pour tous polynômes P, Q, R , on a $(PQ) \circ R = (P \circ R)(Q \circ R)$. En utilisant le résultat pour la somme il suffit de vérifier que $X^{i+j} \circ R = (X^i \circ R)(X^j \circ R)$, ce qui est évident.

Soient enfin trois polynômes P, Q, R . Posons $P = \sum_{k=0}^{\infty} a_k X^k$. On a

$$(P \circ Q) \circ R = \left(\sum_{k=0}^{\infty} a_k Q^k \right) \circ R = \sum_{k=0}^{\infty} a_k (Q^k \circ R)$$

en utilisant la « semi-distributivité » de $\circ$ par rapport à l'addition. On remarque ensuite que $Q^k \circ R = (Q \circ R)^k$ en utilisant la « semi-distributivité » de $\circ$ par rapport à la multiplication. Ainsi,

$$(P \circ Q) \circ R = \sum_{k=0}^{\infty} a_k (Q \circ R)^k = P \circ (Q \circ R)$$

$\square$

Proposition 9.

X est l'élément neutre pour la composition.

La composition des polynômes n'est pas commutative.

Démonstration. Concernant la non-commutativité, constater que

$$X^2 \circ (X + 1) = X^2 + 2X + 1$$

alors que

$$(X + 1) \circ X^2 = X^2 + 1$$

Si $2 \neq 0$ on obtient donc deux polynômes différents. Si jamais $2 = 0$, alors $3 \neq 0$ (pourquoi?). Considérer X^3 au lieu de X^2 . La neutralité de X est quant à elle évidente. $\square$

2 L'anneau des polynômes

2.1 Multiples, diviseurs, d'un polynôme

Définition 3. Soient A et B deux polynômes. On dit que A divise B , et on note $A \mid B$, lorsqu'il existe $C \in \mathbb{K}[X]$ tel que $B = AC$.

Proposition 10. La relation « divise » est réflexive et transitive. En revanche, ce n'est pas une relation d'ordre. Plus précisément, soient $A, B \in \mathbb{K}[X]$. On a

$$\begin{cases} A \mid B \\ B \mid A \end{cases} \iff \exists \lambda \in \mathbb{K}^*, B = \lambda A$$

On dit alors que les polynômes A et B sont **associés**.

Démonstration. Dans le sens direct, on voit que A et B sont nuls tous les deux, ou non nuls tous les deux. Dans le second cas, on a $d^0 A \leq d^0 B$, puisque $B = AC$. De même, $d^0 B \leq d^0 A$. On a donc $B = AC$, avec $d^0 C = 0$. $\square$

2.2 Division euclidienne

Proposition 11. Soient $A, B \in \mathbb{K}[X], B \neq 0$. Il existe un unique couple (Q, R) de polynômes tel que

$$\begin{cases} A = BQ + R \\ d^0 R < d^0 B \end{cases}$$

Q et R sont appelés le quotient et le reste de la division euclidienne de A par B .

Démonstration. Montrons d'abord l'existence. Si $d^0 A < d^0 B$, $Q = 0$ et $R = A$ conviennent. On fait ensuite une récurrence forte sur $n = d^0 A \geq d^0 B$.

Soit $A \in \mathbb{K}[X]$ de degré $n \geq d^o B$. Supposons l'existence du quotient et du reste de la division par B de tout polynôme de degré $< n$. Soit m le degré de B . Considérons

$$A' = A - \frac{A_n}{B_m} X^{n-m} B$$

Le degré de A' est clairement inférieur ou égal à n , celui de A . Mieux, $A'_n = A_n - \frac{A_n}{B_m} B_m = 0$. Donc, $d^o A' < n$. On peut donc appliquer l'hypothèse de récurrence. Il existe Q', R tels que $A' = BQ' + R$ et $d^o R < d^o B$. D'où

$$A = A' + \frac{A_n}{B_m} X^{n-m} B = (Q' + \frac{A_n}{B_m} X^{n-m}) B + R = BQ + R$$

Pour l'unicité, supposons $A = BQ + R = BQ' + R'$ où R et R' sont de degré strictement plus petit que le degré de B . On a $B(Q - Q') = R' - R$ donc $d^o B + d^o(Q - Q') = d^o(R' - R) < d^o B$. Ainsi, $d^o(Q - Q') < 0$. Le seul polynôme de degré strictement négatif est 0. Donc $Q = Q'$ puis, en reportant, $R = R'$. $\square$

3 Fonctions polynômes

3.1 C'est quoi ?

Définition 4. Soit $f : \mathbb{K} \rightarrow \mathbb{K}$. On dit que f est une fonction polynôme lorsqu'il existe une suite presque nulle $(a_0, a_1, \dots)$ d'éléments de $\mathbb{K}$ telle que

$$\forall x \in \mathbb{K}, f(x) = \sum_{k=0}^{\infty} a_k x^k$$

On note $\mathbb{K}[x]$ (j'aurais préféré $\mathbb{K}[id]$) l'ensemble des fonctions polynômes.

Remarque 5. Un polynôme est un objet abstrait. On ne peut pas donner une valeur à X . En revanche, si f est une fonction polynôme, on peut évidemment calculer $f(a)$ pour tout $a \in \mathbb{K}$.

Proposition 12. $\mathbb{K}[x]$ est une $\mathbb{K}$ -algèbre commutative.

Démonstration. C'est une sous-algèbre de $\mathbb{K}^{\mathbb{K}}$, l'algèbre des fonctions de $\mathbb{K}$ vers $\mathbb{K}$. Il suffit de vérifier que la somme et le produit de deux fonctions polynômes sont encore des fonctions polynômes, et de même pour le produit d'une fonction polynôme par un élément de $\mathbb{K}$. $\square$

3.2 Polynômes et fonctions polynômes

Soit $\varphi : \mathbb{K}[X] \rightarrow \mathbb{K}[id]$ l'application qui à $P = \sum_{k=0}^{\infty} P_k X^k$ associe la fonction $\widehat{P}$ définie par $\forall x \in \mathbb{K}, \widehat{P}(x) = \sum_{k=0}^{\infty} P_k x^k$.

Proposition 13. *La fonction φ est un morphisme d'algèbres surjectif.*

Démonstration. La preuve consiste en quelques vérifications que nous ne ferons pas : les lecteurs qui auront scrupuleusement vérifié que $\mathbb{K}[x]$ est une algèbre auront d'ailleurs déjà fait tout le travail (et seuls ceux-ci pourront comprendre cette phrase). Le problème reste l'injectivité de φ : peut-on « identifier » les polynômes et les fonctions polynômes ? Nous allons voir que c'est le cas lorsque le corps $\mathbb{K}$ est infini (c'est donc toujours le cas lorsque $\mathbb{K}$ est un sous-corps de $\mathbb{C}$). $\square$

3.3 Racines d'un polynôme

Définition 5. Soit $P \in \mathbb{K}[X]$. Soit $a \in \mathbb{K}$. On dit que a est racine de P lorsque $\widehat{P}(a) = 0$.

Proposition 14. *Soit $P \in \mathbb{K}[X]$. Soit $a \in \mathbb{K}$. Alors a est racine de P si et seulement si $X - a$ divise P .*

Démonstration. Supposons que a est racine de P . On a donc

$$P = (X - a)Q$$

où $Q \in \mathbb{K}[X]$. De là,

$$\widehat{P}(a) = (a - a)\widehat{Q}(a) = 0$$

Inversement, supposons que $\widehat{P}(a) = 0$. On écrit

$$P = (X - a)Q + R$$

avec $d^0 R < d^0(X - a) = 1$. $R = \lambda \in \mathbb{K}$ est donc constant. En passant aux fonctions polynômes, on en tire

$$0 = (a - a)\widehat{Q}(a) + \lambda$$

d'où $\lambda = 0$ et $P = (X - a)Q$. $\square$

Proposition 15. *Soient $a_1, \dots, a_n$ des éléments distincts de $\mathbb{K}$. Ce sont des racines de P si et seulement si $\prod_{k=1}^n (X - a_k)$ divise P .*

Démonstration. Dans un sens, c'est évident. Dans l'autre sens, c'est une récurrence sur n . Faisons-le pour $n = 2$. Soient $a \neq b$ deux racines de P . On a $P = (X - a)Q$, puisque a est racine de P . Mais alors, $\widehat{P}(b) = 0 = (b - a)\widehat{Q}(b)$. Comme $a \neq b$, b est racine de Q , et $Q = (X - b)R$. Donc,

$$P = (X - a)(X - b)R$$

□

Proposition 16. Soit P un polynôme non nul, de degré $n \in \mathbb{N}$. Le polynôme P possède au plus n racines.

Démonstration. Si $a_1, \dots, a_k$ sont des racines distinctes de P , alors,

$$P = Q \prod_{j=1}^k (X - a_j)$$

Donc, $d^o P = d^o Q + k$ d'où $d^o P \geq k$. □

Théorème 17. La fonction $\varphi : \mathbb{K}[X] \rightarrow \mathbb{K}[id]$ définie plus haut est un isomorphisme si et seulement si le corps $\mathbb{K}$ est infini.

Démonstration. Supposons le corps $\mathbb{K}$ infini. Soit $P \in \ker \varphi$. La fonction polynôme $\widehat{P}$ est donc la fonction nulle. P a ainsi une infinité de racines et ne peut être que le polynôme nul.

Inversement, supposons $\mathbb{K}$ fini. Soit

$$P = \prod_{a \in \mathbb{K}} (X - a)$$

Le polynôme P est un polynôme de degré $\text{card } K$, et pourtant $\widehat{P} = 0$. □

Remarque 6. Dans tout corps $\mathbb{K}$ infini comme $\mathbb{Q}, \mathbb{R}, \mathbb{C}$, on peut donc identifier les polynômes et les fonctions polynômes, c'est à dire que l'on n'écrira plus les chapeaux. On écrira aussi parfois $P(X)$ pour un polynôme P . Bref, les fonctions polynômes se prennent pour des polynômes, et les polynômes se prennent pour des fonctions.

3.4 Racines multiples

Définition 6. Soit P un polynôme, et a un élément de $\mathbb{K}$. Soit $k \in \mathbb{N}$. On dit que a est racine de P de multiplicité k lorsque $(X - a)^k$ divise P , mais $(X - a)^{k+1}$ ne divise pas P .

On parle ainsi de racine simple, double, etc. Une racine de multiplicité 0 est une non-racine. Nous reviendrons un peu plus loin sur les propriétés des racines multiples.

4 Dérivation

4.1 Dérivée d'un polynôme

Définition 7. Soit $P = \sum_{k=0}^{\infty} P_k X^k$. On appelle dérivée de P le polynôme

$$P' = \sum_{k=1}^{\infty} k P_k X^{k-1}$$

Les formules classiques de dérivation d'une somme, d'un produit, d'une composée, se démontrent sans problème de façon purement formelle. On a bien sûr la notion de dérivée d'ordre supérieur, la formule de Leibniz, etc. Il convient de remarquer que si $\mathbb{K}$ est un sous-corps de $\mathbb{C}$ et $P \in \mathbb{K}[X]$ est non constant, on a $d^o P' = d^o P - 1$. Ainsi, en dérivant $d^o P + 1$ fois le polynôme P , on tombe sur le polynôme nul. On remarque enfin que $\widehat{P'} = (\widehat{P})'$ lorsque la notion de dérivée d'une fonction polynôme a un sens, c'est à dire, pour nous, lorsque $\mathbb{K} = \mathbb{R}$.

4.2 Formule de Taylor

Dans ce paragraphe, on suppose que $\mathbb{K}$ est un sous-corps de $\mathbb{C}$. Dans la pratique, $\mathbb{K}$ est égal à $\mathbb{Q}$, $\mathbb{R}$ ou $\mathbb{C}$. En fait, ce dont nous allons parler est valable dans tout corps $\mathbb{K}$ où, pour tout entier $n \geq 1$, $n \neq 0$.

Proposition 18. FORMULE DE TAYLOR Soit $P \in \mathbb{K}[X]$. Soit $a \in \mathbb{K}$. On a

$$P(X) = \sum_{k=0}^{\infty} \frac{(X-a)^k}{k!} P^{(k)}(a)$$

Démonstration. On procède par récurrence sur $n = d^o P$. C'est clair si P est constant. Soit $n \geq 1$. Supposons la propriété vraie pour les polynômes de degré $n-1$. Soit P de degré n . Le polynôme P' étant de degré $n-1$, on a par l'hypothèse de récurrence

$$P'(X) = \sum_{k=0}^{n-1} \frac{(X-a)^k}{k!} P^{(k+1)}(a)$$

d'où le résultat en intégrant. $\square$

Corollaire 19. Soit $P \in \mathbb{K}[X]$. Soit $a \in \mathbb{K}$ et $k \in \mathbb{N}$. Alors, a est racine d'ordre k du polynôme P si et seulement si

$$P(a) = P'(a) = \dots = P^{(k-1)}(a) = 0 \text{ et } P^{(k)}(a) \neq 0$$

Démonstration. Si P et ses dérivées jusqu'à l'ordre $k - 1$ sont nulles en a , la formule de Taylor s'écrit

$$P(X) = \sum_{j=k}^n \frac{(X-a)^j}{j!} P^{(j)}(a)$$

On peut donc factoriser $(X-a)^k$ dans le polynôme. Une fois factorisé, le polynôme restant ne s'annule pas en a .

Inversement, on fait une récurrence sur k . Si $k \leq 0$, c'est évident. Soit $k \geq 1$. Supposons $P = (X-a)^k Q$, où $Q(a) \neq 0$. Alors,

$$P' = (X-a)^{k-1} R$$

où $R(a) \neq 0$. Par l'hypothèse de récurrence, on a donc

$$P'(a) = \dots = P'^{(k-2)}(a) = 0 \text{ et } P'^{(k-1)}(a) \neq 0$$

Enfin, on a bien entendu $P(a) = 0$. $\square$

5 Factorisation des polynômes

5.1 Polynômes irréductibles

Définition 8. Soit $P \in \mathbb{K}[X]$. On dit que P est irréductible sur $\mathbb{K}$ lorsque

- P n'est pas constant.
- Pour tous polynômes Q et R ,

$$P = QR \implies Q \text{ est constant ou } R \text{ est constant}$$

Exemple.

- Les polynômes de degré 1
- $X^2 + X + 1$ sur $\mathbb{R}$, mais pas sur $\mathbb{C}$.
- $X^3 - 2$ sur $\mathbb{Q}$.
- $X^2 + X + 1$ sur $\mathbb{Z}/2\mathbb{Z}$.

Proposition 20. *Tout polynôme non constant s'écrit de façon unique, à l'ordre près, comme un produit de polynômes irréductibles.*

Démonstration. Montrons tout d'abord l'existence de la décomposition. On fait une récurrence forte sur $n = d^o P$. Si $n = 1$ c'est clair puisqu'un polynôme de degré 1 est irréductible. Soit $n \geq 1$. Supposons l'existence de la décomposition pour tout polynôme de degré strictement inférieur à n . Soit P un polynôme de degré n . Deux cas se présentent. Première possibilité, P est irréductible. Dans ce cas on a fini. Deuxième possibilité, $P = QR$ avec Q et R non constants. Mais alors Q et R sont de degré strictement inférieur à n . D'après l'hypothèse de récurrence, ils se décomposent en produit de polynômes irréductibles, donc P aussi. L'unicité de la décomposition repose sur le théorème de Gauss que nous n'avons pas encore. Nous l'admettons pour l'instant. $\square$

Théorème 21. THÉORÈME DE D'ALEMBERT-GAUSS *Soit $P \in \mathbb{C}[X]$ non constant. Alors, P a au moins une racine.*

Démonstration. La démonstration de ce théorème est hors programme. Il en existe plusieurs preuve, je donne ici les grandes lignes de l'une d'entre elles.

Soit $P \in \mathbb{C}[X]$ un polynôme de degré $n \geq 1$. Supposons que P ne s'annule pas sur $\mathbb{C}$. Posons $P = \sum_{k=0}^n a_k X^k$, où les $a_k \in \mathbb{C}$ et $a_n \neq 0$.

Pour tout $z \in \mathbb{C}$,

$$|P(z)| \geq |a_n||z|^n - \left| \sum_{k=0}^{n-1} a_k z^k \right| = |z|^n \left(|a_n| - \left| \sum_{k=0}^{n-1} a_k z^{k-n} \right| \right)$$

La parenthèse tend vers $|a_n|$ lorsque $|z|$ tend vers l'infini. Ainsi, il existe un réel M tel que pour tout $z \in \mathbb{C}$,

$$|z| \geq M \implies |P(z)| \geq |P(0)|$$

La fonction $|P|$ est une fonction continue du *disque fermé* D de centre O et de rayon M vers $\mathbb{R}$. Ce disque étant fermé et borné, $|P|$ possède un minimum sur D . Autrement dit, il existe $z_0 \in D$ tel que pour tout $z \in D$, $|P(z)| \geq |P(z_0)|$. Remarquons que $0 \in D$, et donc, pour tout $z \notin D$, $|P(z)| \geq |P(0)| \geq |P(z_0)|$. Finalement,

$$\forall z \in \mathbb{C}, |P(z)| \geq |P(z_0)|$$

La fonction $|P|$ possède donc un minimum global sur $\mathbb{C}$ en z_0 .

Maintenant, grâce à la formule de Taylor appliquée au polynôme P au point z_0 , on a

$$P(z) = P(z_0) + \alpha(z - z_0)^k + (z - z_0)^k \varepsilon(z)$$

où $\alpha \in \mathbb{C}^*$, $k \geq 1$ et $\varepsilon(z) \rightarrow 0$ lorsque z tend vers z_0 . Divisons par $P(z_0)$, qui est non nul car P est supposé sans racine. On obtient

$$\frac{P(z)}{P(z_0)} = 1 + \beta(z - z_0)^k + (z - z_0)^k \varepsilon(z)$$

où $\beta \in \mathbb{C}^*$. Posons

$$z = z_0 + \rho e^{i\theta}$$

où $\rho \in \mathbb{R}_+$ et $\theta \in [0, 2\pi]$. On a donc

$$(1) \quad \frac{P(z)}{P(z_0)} = 1 + \beta \rho^k e^{ik\theta} + \rho^k e^{ik\theta} \varepsilon(\rho)$$

où $\varepsilon(\rho)$ tend vers 0 lorsque ρ tend vers 0. En passant au conjugué, on a aussi

$$(2) \quad \overline{\frac{P(z)}{P(z_0)}} = 1 + \bar{\beta} \rho^k e^{-ik\theta} + \rho^k e^{-ik\theta} \overline{\varepsilon(\rho)}$$

Multiplions (1) et (2). Il vient

$$\begin{aligned} \left| \frac{P(z)}{P(z_0)} \right|^2 &= \left(1 + \beta \rho^k e^{ik\theta} + \rho^k e^{ik\theta} \varepsilon(\rho) \right) \left(1 + \bar{\beta} \rho^k e^{-ik\theta} + \rho^k e^{-ik\theta} \overline{\varepsilon(\rho)} \right) \\ &= 1 + 2\rho^k \operatorname{Re}(\beta e^{-ik\theta}) + \rho^k \hat{\varepsilon}(\rho) \end{aligned}$$

où $\hat{\varepsilon}(\rho)$ tend vers 0 lorsque ρ tend vers 0. Ainsi, pour $\rho \neq 0$,

$$\frac{1}{\rho^k} \left(\left| \frac{P(z)}{P(z_0)} \right|^2 - 1 \right) = 2\operatorname{Re}(\beta e^{-ik\theta}) + \hat{\varepsilon}(\rho)$$

Choisissons maintenant θ tel que $\operatorname{Re}(\beta e^{-ik\theta}) < 0$. Pour ρ suffisamment petit, on a alors

$$\frac{1}{\rho^k} \left(\left| \frac{P(z)}{P(z_0)} \right|^2 - 1 \right) < 0$$

et donc

$$|P(z)| < |P(z_0)|$$

Contradiction. $\square$

Proposition 22.

- Les polynômes irréductibles sur $\mathbb{C}$ sont les polynômes de degré 1.
- Les polynômes irréductibles sur $\mathbb{R}$ sont les polynômes de degré 1 et les polynômes de degré 2 à discriminant strictement négatif.

Démonstration.

- Commençons par $\mathbb{C}$. Soit $P \in \mathbb{C}[X]$. Si P est de degré 1, il est irréductible. Supposons P de degré supérieur ou égal à 2. Le théorème de d'Alembert affirme que P a une racine a . On a donc $P = (X - a)Q$ où Q n'est pas constant. Donc P n'est pas irréductible.
- Passons à $\mathbb{R}$. Soit $P \in \mathbb{R}[X]$.

Si P est de degré 1, il est irréductible.

Supposons P de degré 2. Si $P = QR$ avec Q et R non constants alors Q et R sont de degré 1, donc P a une racine. Donc, si son discriminant est strictement négatif, P est irréductible. Inversement, si le discriminant de P est positif ou nul, P a une racine et n'est donc pas irréductible.

Supposons pour terminer que $d^o P \geq 3$. Si P a une racine réelle a , alors $X - a$ divise P et P n'est pas irréductible. Sinon, P a une racine non réelle ζ . Mais P étant à coefficients réels, on voit facilement que $\bar{\zeta} \neq \zeta$ est aussi racine de P . On peut donc écrire

$$P = (X - \zeta)(X - \bar{\zeta})Q$$

où Q est a priori dans $\mathbb{C}[X]$. Remarquons que

$$A = (X - \zeta)(X - \bar{\zeta}) = X^2 - 2(\operatorname{Re} \zeta)X + |\zeta|^2 \in \mathbb{R}[X]$$

Il nous reste à prouver que l'on a en fait $Q \in \mathbb{R}[X]$. On a d'une part $P = AQ$ dans $\mathbb{C}[X]$, et d'autre part (division euclidienne dans $\mathbb{R}[X]$) $P = AQ' + R$ avec Q' et R dans $\mathbb{R}[X]$. Mais cette dernière égalité est aussi une égalité dans $\mathbb{C}[X]$, et donc une division euclidienne dans $\mathbb{C}[X]$. Par unicité du quotient et du reste de la division euclidienne, on en déduit $Q = Q'$ et $R = 0$, d'où $P = AQ$ avec $A, Q \in \mathbb{R}[X]$ et $d^o A = 2$ donc Q non constant. P n'est donc pas irréductible.

□

5.2 Polynômes scindés

Définition 9. Soit $P \in \mathbb{K}[X]$ non nul. On dit que P est scindé (sur $\mathbb{K}$) lorsque P est un produit de polynômes de degré 1.

Exemple. Tout polynôme non nul de $\mathbb{C}[X]$ est scindé. C'est le théorème de d'Alembert.

Proposition 23. Soit $P \in \mathbb{K}[X]$ un polynôme non nul de racines $\alpha_1, \dots, \alpha_q$ d'ordres de multiplicité respectifs $k_1, \dots, k_q$. Alors $k_1 + \dots + k_q \leq d^o P$ et on a égalité si et seulement si P est scindé sur $\mathbb{K}$.

Démonstration. On remarque que si $P = QR$ et si a est racine d'ordre k de P , mais pas de R , alors a est racine d'ordre k de Q . En effet a est clairement racine de Q . Donc, $Q = (X - a)Q_1$. On termine par récurrence sur k .

Ainsi, $\prod_{i=1}^q (X - \alpha_i)^{k_i}$ divise P , donc on a bien l'inégalité demandée. La CNS d'égalité est évidente. $\square$

5.3 Relations coefficients-racines

Soient $a, b, c \in \mathbb{K}$, $a \neq 0$. Soit $P = aX^2 + bX + c$. Supposons P scindé. On a donc

$$P = a(X - x_1)(X - x_2)$$

où $x_1, x_2 \in \mathbb{K}$. En développant on trouve les relations bien connues

$$\begin{cases} x_1 + x_2 &= -\frac{b}{a} \\ x_1 x_2 &= \frac{c}{a} \end{cases}$$

Prenons maintenant un polynôme scindé de degré 3,

$$P = aX^3 + bX^2 + cX + d = a(X - x_1)(X - x_2)(X - x_3)$$

Un développement donne

$$\begin{cases} x_1 + x_2 + x_3 &= -\frac{b}{a} \\ x_1 x_2 + x_2 x_3 + x_1 x_3 &= \frac{c}{a} \\ x_1 x_2 x_3 &= -\frac{d}{a} \end{cases}$$

On commence à voir apparaître quelque chose.

Définition 10. Soit $n \geq 1$, et k un entier entre 1 et n . On appelle k -ième polynôme symétrique élémentaire le « polynôme à n variables »

$$\sigma_k(X_1, \dots, X_n) = \sum_{1 \leq i_1 < \dots < i_k \leq n} X_{i_1} X_{i_2} \dots X_{i_k}$$

En particulier $\sigma_1 = X_1 + X_2 + \dots + X_n$ et $\sigma_n = X_1 X_2 \dots X_n$.

Proposition 24. Soit $n \geq 1$. Soit

$$P = \sum_{k=0}^n a_k X^k$$

un polynôme de degré n scindé, de racines $\alpha_1, \dots, \alpha_n$ (éventuellement égales s'il y a des racines multiples). En clair,

$$P = a_n \prod_{i=1}^n (X - \alpha_i)$$

On a alors pour tout $k \in \llbracket 1, n \rrbracket$

$$\sigma_k(\alpha_1, \dots, \alpha_n) = (-1)^k \frac{a_{n-k}}{a_n}$$

Démonstration. On donne juste l'idée de la preuve, qui se fait par récurrence sur n . Pour $n = 1, 2, 3$ nous avons déjà vu que c'était vrai. Supposons que l'on sait faire pour un polynôme de degré $n \geq 1$, et prenons

$$P = \lambda \prod_{k=1}^{n+1} (X - \alpha_k)$$

où $\lambda \neq 0$ est le coefficient dominant de P et les α_k sont les racines de P . On a alors, par l'hypothèse de récurrence,

$$P = \lambda \left(X^n + \sum_{k=1}^n (-1)^{n-k} \sigma_{n-k}(\alpha_1, \dots, \alpha_n) X^k \right) (X - \alpha_{n+1})$$

On réordonne ensuite suivant les puissances de X pour constater que l'on obtient bien les polynômes symétriques élémentaires en $\alpha_1, \dots, \alpha_{n+1}$. $\square$

Exercice. Trouver tous les réels x, y, z vérifiant

$$\begin{cases} x + y + z &= 2 \\ \frac{1}{x} + \frac{1}{y} + \frac{1}{z} &= \frac{5}{6} \\ xyz &= -6 \end{cases}$$

6 Pgcd, Ppcm

6.1 L'algorithme d'Euclide

Proposition 25. Soient A et B deux polynômes à coefficients dans $\mathbb{K}$. Il existe un polynôme Δ vérifiant

$$\forall D \in \mathbb{K}[X], \begin{cases} D \mid A \\ D \mid B \end{cases} \iff D \mid \Delta$$

Le polynôme Δ est unique à constante multiplicative non nulle près.

Définition 11. Δ est appelé un pgcd de A et B . On parlera abusivement du pgcd de A et B , et on notera $\Delta = A \wedge B$ ou $\Delta = \text{pgcd}(A, B)$.

Démonstration. L'unicité est claire : si Δ_1 et Δ_2 sont deux pgcd de A et B , alors chacun divise l'autre, et ils sont donc associés.

Si $B = 0$, l'existence est assurée, puisque $\Delta = A$ convient. Sinon, un polynôme D divise A et B si et seulement si il divise B et R , où R est le reste de la division euclidienne de A par B . On crée ainsi une suite de restes successifs de divisions euclidiennes : c'est l'algorithme d'Euclide. On arrive forcément à un reste nul, car la suite des degrés des restes serait sinon une suite strictement décroissante d'entiers naturels. Le dernier reste non nul est le polynôme Δ recherché. $\square$

Exemple.

- Pour tout polynôme A , on a $A \wedge 1 = 1$, $A \wedge 0 = A$, $A \wedge A = A$.
- On a $(X^2 - 3X + 2) \wedge (X^3 - 2X^2 + X - 2) = X - 2$.

6.2 Théorème de Bézout

Définition 12. Soient $A, B \in \mathbb{K}[X]$. On dit que A et B sont premiers entre eux lorsque $A \wedge B = 1$.

Exemple. Deux polynômes irréductibles non associés sont premiers entre-eux. En effet, soient P et Q irréductibles, non associés. Soit D un polynôme. Si D divise P et Q alors $D = 1$ ou P et $D = 1$ ou Q (àcmnnp). Donc $D = 1$ àcmnnp.

Théorème 26. Soient $A, B \in \mathbb{K}[X]$. Alors, A et B sont premiers entre eux si et seulement si il existe $U, V \in \mathbb{K}[X]$ tels que $AU + BV = 1$.

Démonstration. Supposons l'existence de U et V . Alors, Si D divise A et B , il divise $AU + BV$, donc 1. Donc, les seuls diviseurs communs de A et B sont les polynômes constants.

Inversement, supposons que $A \wedge B = 1$. On a $1 \times A + 0 \times B = A$, et $0 \times A + 1 \times B = B$. Soient (R_n) et (Q_n) les suites des restes et des quotients des divisions de l'algorithme d'Euclide : $R_0 = A$, $R_1 = B$, et $R_n = R_{n+1}Q_{n+2} + R_{n+2}$, avec $R_{n_0+1} = 0$, et $R_{n_0} = 1$. Si l'on a $U_n A + V_n B = R_n$, et de même au rang $n + 1$, il est alors facile de combiner ces deux égalités (la première moins Q_{n+2} fois la seconde) pour obtenir la même égalité au rang $n + 2$. On termine donc par récurrence. $\square$

Remarque 7. Si $A \wedge B = \Delta$, le raisonnement précédent montre qu'il existe U et V tels que $UA + VB = \Delta$. La réciproque est fausse lorsque Δ n'est pas constant.

6.3 Théorème de Gauss

Théorème 27. Soient $A, B, C \in \mathbb{K}[X]$. On suppose que $A \mid BC$ et que $A \wedge B = 1$. Alors, $A \mid C$.

Démonstration. Il existe $U, V, Q \in \mathbb{K}[X]$ tels que $UA + VB = 1$ et $BC = QA$. Donc,

$$C = UAC + VBC = UAC + VQA = (UC + VQ)A$$

□

6.4 Propriétés utiles

Les trois propriétés ci-dessous sont analogues aux propriétés vues en arithmétique sur $\mathbb{Z}$. Elles se démontrent de la même façon, grâce aux théorèmes de Bézout et Gauss.

Proposition 28. Soient $A, B, \Delta \in \mathbb{K}[X]$. Alors $\Delta = A \wedge B$ si et seulement si il existe $A_1, B_1 \in \mathbb{K}[X]$ tels que

$$\begin{cases} A = \Delta A_1 \\ B = \Delta B_1 \\ A_1 \wedge B_1 = 1 \end{cases}$$

Proposition 29. Soit $n \geq 1$. Soient $A, B_1, \dots, B_n \in \mathbb{K}[X]$. Alors,

$$A \wedge \prod_{k=1}^n B_k = 1 \iff \forall k \in \llbracket 1, n \rrbracket, A \wedge B_k = 1$$

Exemple. Soient $a, b \in \mathbb{K}$, $a \neq b$. Soient $\alpha, \beta \in \mathbb{N}$. Alors $(X - a)^\alpha \wedge (X - b)^\beta = 1$.

Proposition 30. Soit $n \geq 1$. Soient $B, A_1, \dots, A_n \in \mathbb{K}[X]$. On suppose $A_1, \dots, A_n$ sont premiers entre eux deux à deux. Alors

$$\prod_{k=1}^n A_k \mid B \iff \forall k \in \llbracket 1, n \rrbracket, A_k \mid B$$

6.5 Décomposition en produit de polynômes irréductibles (bis)

Rappelons-nous que nous avons jusqu'à présent laissé de côté la preuve de l'unicité de la décomposition en produit de polynômes irréductibles. Nous en savons maintenant assez pour montrer ce résultat.

Notation. Notons $\mathcal{I}$ l'ensemble des polynômes irréductibles et unitaires sur le corps $\mathbb{K}$.

Proposition 31. *Tout polynôme $Q \in \mathbb{K}[X]$ non nul s'écrit de façon unique*

$$Q = c \prod_{P \in \mathcal{I}} P^{\alpha_P}$$

où $c \in \mathbb{K}^$ et les α_P sont des entiers naturels tous nuls sauf un nombre fini.*

Démonstration. L'existence a déjà été montrée, elle se prouve par récurrence forte sur le degré de Q . Montrons l'unicité.

Soit Q un polynôme non nul. Supposons

$$Q = c \prod_{P \in \mathcal{I}} P^{\alpha_P} = c' \prod_{P \in \mathcal{I}} P^{\beta_P}$$

Le coefficient dominant de Q est $c = c'$. On a donc

$$\prod_{P \in \mathcal{I}} P^{\alpha_P} = \prod_{P \in \mathcal{I}} P^{\beta_P}$$

Soit $P_0 \in \mathcal{I}$. On a

$$P_0^{\alpha_{P_0}} \mid \left(P_0^{\beta_{P_0}} \prod_{P \in \mathcal{I}, P \neq P_0} P^{\beta_P} \right)$$

Mais $P_0^{\alpha_{P_0}}$ est premier avec $\prod_{P \in \mathcal{I}, P \neq P_0} P^{\beta_P}$. Donc, par le théorème de Gauss,

$$P_0^{\alpha_{P_0}} \mid P_0^{\beta_{P_0}}$$

On en déduit que

$$\alpha_{P_0} \leq \beta_{P_0}$$

De la même façon, $\alpha_{P_0} \geq \beta_{P_0}$ et donc $\alpha_{P_0} = \beta_{P_0}$. Il y a bien unicité de l'écriture. $\square$

6.6 Ppcm

Proposition 32. Soient $A, B \in \mathbb{K}[X]$. Il existe un polynôme μ , unique à constante multiplicative non nulle près, tel que

$$\forall M \in \mathbb{K}[X], \left\{ \begin{array}{l} A \mid M \\ B \mid M \end{array} \right\} \iff \mu \mid M$$

Démonstration. Voir le cours sur les entiers relatifs, la preuve est identique. On écrit $A = \Delta A_1$, $B = \Delta B_1$, où $\Delta = A \wedge B$ et $A_1 \wedge B_1 = 1$. On pose

$$\mu = \Delta A_1 B_1$$

et on montre ensuite que μ est, à constante multiplicative non nulle près, l'unique solution du problème. $\square$

Définition 13. μ est appelé un ppcm de A et B . On parlera abusivement du ppcm de A et B , et on notera $\mu = A \vee B$ ou $\mu = \text{ppcm}(A, B)$.

Proposition 33. Pour tous $A, B \in \mathbb{K}[K]$, $(A \wedge B)(A \vee B) = AB$.

Démonstration. Cette égalité est bien sûr vraie à constante multiplicative non nulle près. Avec les notations ci-dessus,

$$(A \wedge B)(A \vee B) = \Delta(\Delta A_1 B_1) = AB$$

$\square$

6.7 Calculs pratiques

Utilisation de polynômes irréductibles

Soient $A = \prod_{P \in \mathcal{I}} P^{\alpha_P}$ et $B = \prod_{P \in \mathcal{I}} P^{\beta_P}$ où les α_P, β_P sont des entiers naturels presque tous nuls. Alors,

$$A \wedge B = \prod_{P \in \mathcal{I}} P^{\min(\alpha_P, \beta_P)}$$

et

$$A \vee B = \prod_{P \in \mathcal{I}} P^{\max(\alpha_P, \beta_P)}$$

Se reporter au cours d'arithmétique dans $\mathbb{Z}$ pour une démonstration. Cette méthode de calcul n'est évidemment utile que pour des polynômes que l'on sait factoriser.

Calcul des coefficients de Bézout

Se reporter au cours d'arithmétique dans $\mathbb{Z}$. La méthode est identique. Elle permet de résoudre complètement l'équation $AU + BV = C$, d'inconnues $U, V \in \mathbb{K}[X]$.

7 Interpolation de Lagrange

7.1 Qu'est-ce que l'interpolation ?

Soit $f : \mathbb{K} \rightarrow \mathbb{K}$. Soit $n \in \mathbb{N}$. Soient $x_0 < x_1 < \dots < x_n$ $n+1$ éléments distincts de $\mathbb{K}$. Sans chercher la généralité, *interpoler* f en les x_k , c'est trouver un polynôme P tel que pour tout $k \in \llbracket 0, n \rrbracket$, on ait $P(x_k) = f(x_k)$. Il est clair que nous parlons ici d'interpolation *polynomiale*, on pourrait chercher à interpoler f par d'autres types de fonctions. Par exemple, si $\mathbb{K} = \mathbb{R}$, par des combinaisons de fonctions trigonométriques. Ou alors, imposer également que certaines dérivées de P soient égales aux dérivées de f aux points correspondants.

En réalité, on peut se débarrasser de la fonction f dans la formulation de la question : étant donnés $y_0, \dots, y_n \in \mathbb{K}$, existe-t-il un polynôme $P \in \mathbb{K}[X]$ tel que pour tout $k \in \llbracket 0, n \rrbracket$, on ait $P(x_k) = y_k$?

Pour les toutes petites valeurs de n , la réponse est évidente.

Lorsque $n = 0$, la réponse est clairement oui, on peut prendre par exemple un polynôme constant.

Lorsque $n = 1$, c'est encore oui en prenant pour P une fonction affine : par deux points distincts du « plan » il passe une unique droite. Ceci peut se généraliser.

7.2 Existence et unicité

Proposition 34. *Soit $n \in \mathbb{N}$. Soient x_0, x_1, x_n $n+1$ éléments distincts de $\mathbb{K}$. Soient $y_0, y_1, \dots, y_n$ $n+1$ éléments de $\mathbb{K}$. Il existe un unique polynôme $P \in \mathbb{K}_n[X]$ tel que*

$$\forall k \in \llbracket 0, n \rrbracket, P(x_k) = y_k$$

Démonstration. Commençons par l'unicité. Supposons que deux tels polynômes P et Q conviennent. Le polynôme $P - Q$ s'annule en $x_0, \dots, x_n$, et a donc au moins $n+1$ racines distinctes. Mais ce polynôme est de degré inférieur ou égal à n . C'est le polynôme nul. Ainsi, $P = Q$.

Passons à l'existence. Pour $i \in \llbracket 0, n \rrbracket$, soit

$$L_i = \frac{\prod_{k \neq i} (X - x_k)}{\prod_{k \neq i} (x_i - x_k)}$$

On vérifie facilement que ces polynômes sont de degré n , et que pour tous $i, j \in \llbracket 0, n \rrbracket$, on a

$$L_i(x_j) = \delta_{ij}$$

c'est à dire 0 si $i \neq j$ et 1 si $i = j$. Ces polynômes sont les *polynômes de Lagrange élémentaires*. Soit maintenant

$$P = \sum_{i=0}^n y_i L_i$$

Le polynôme P est de degré inférieur ou égal à n (degré d'une somme). De plus, pour tout $j \in \llbracket 0, n \rrbracket$,

$$P(x_j) = \sum_{i=0}^n y_i L_i(x_j) = \sum_{i=0}^n y_i \delta_{ij} = y_j$$

Le polynôme P convient donc. $\square$

Que se passe-t-il si l'on supprime la condition de degré ? On perd l'unicité. Plus précisément :

Proposition 35. *Avec les notations ci-dessus, soit $Q \in \mathbb{R}[X]$. On a $Q(x_j) = y_j$ pour tout $j \in \llbracket 0, n \rrbracket$ si et seulement si il existe $A \in \mathbb{R}[X]$ tel que*

$$Q = P + A \prod_{k=0}^n (X - x_k)$$

Démonstration. Un tel polynôme vérifie clairement les conditions. Inversement, soit Q un polynôme prenant les valeurs y_j aux points x_j . Le polynôme $Q - P$ s'annule en $x_0, \dots, x_n$. On en déduit que

$$\prod_{k=0}^n (X - x_k) \mid Q - P$$

$\square$

Chapitre 19

Fractions rationnelles

1 Le corps des fractions rationnelles

1.1 Notion de fraction rationnelle

Soit $\mathbb{K}$ un corps. L'anneau $\mathbb{K}[X]$ des polynômes à coefficients dans $\mathbb{K}$ est un anneau intègre. Un résultat très général affirme que $\mathbb{K}[X]$ possède un *corps des fractions*, unique à isomorphisme près, noté $\mathbb{K}(X)$: le corps des fractions rationnelles à une indéterminée.

Formellement, une fraction rationnelle est le quotient $F = \frac{P}{Q}$ de deux polynômes, où le polynôme Q est non nul. Si $P, Q, R, S \in \mathbb{K}[X]$ et $Q, S \neq 0$, on a

$$\begin{aligned} \frac{P}{Q} = \frac{R}{S} &\iff PS = QR \\ \frac{P}{Q} + \frac{R}{S} &= \frac{PS + QR}{QS} \\ \frac{P}{Q} \frac{R}{S} &= \frac{PR}{QS} \end{aligned}$$

Proposition 1. FORME IRREDUCTIBLE

Soit $F \in \mathbb{K}(X)$. F s'écrit de façon unique

$$F = \frac{P}{Q}$$

où $P, Q \in \mathbb{K}[X]$, Q est unitaire, et $P \wedge Q = 1$.

La fraction F est alors dite *sous forme irréductible*.

Démonstration.

- Pour l'unicité, supposons $F = \frac{P}{Q} = \frac{R}{S}$, où $P \wedge Q = 1$, $R \wedge S = 1$, et Q, S sont unitaires. On a alors $PS = QR$.

Q divise PS et Q est premier avec P . Par le théorème de Gauss, Q divise S . De même S divise Q . Ainsi, il existe $\lambda \in \mathbb{K}^*$ tel que $S = \lambda Q$. Comme Q et S sont unitaires, $\lambda = 1$ et ainsi $Q = S$. En reportant, on obtient $P = R$.

- Soient $P, Q \in \mathbb{K}[X]$ tels que $F = \frac{P}{Q}$, où $Q \neq 0$. Soit $\Delta = P \wedge Q$. On peut alors écrire $P = \Delta P_1$, $Q = \Delta Q_1$, où $P_1 \wedge Q_1 = 1$. De là, $F = \frac{P_1}{Q_1}$. Quitte à diviser P_1 et Q_1 par le coefficient dominant de Q_1 , on peut supposer que Q_1 est unitaire, d'où l'existence.

□

1.2 Degré d'une fraction rationnelle

Définition 1. Soit $F = \frac{P}{Q} \in \mathbb{K}(X)$. On appelle degré de F la quantité

$$d^o F = d^o P - d^o Q$$

Le degré d'une fraction est soit un entier relatif, soit $-\infty$ lorsque la fraction est nulle. Attention à quelques pièges : une fraction de degré 0 n'est pas forcément constante, par exemple.

Il convient de vérifier que cette définition a un sens, c'est à dire qu'elle dépend uniquement de la fraction F et pas des polynômes P et Q choisis pour la représenter.

Proposition 2. Cette quantité ne dépend pas des polynômes P et Q choisis pour représenter F .

Démonstration. Avec des notations évidentes, supposons

$$F = \frac{P}{Q} = \frac{R}{S}$$

On a alors $PS = QR$. De là,

$$d^o(PS) = d^o P + d^o S = d^o(QR) = d^o Q + d^o R$$

d'où

$$d^o P - d^o Q = d^o R - d^o S$$

□

Proposition 3. Soient $F, G \in \mathbb{K}(X)$. On a

$$d^o(F + G) \leq \max(d^o F, d^o G)$$

Si les fractions ont des degrés différents, c'est une égalité. Sinon, on ne peut rien dire.

Démonstration. On pose $F = \frac{P}{Q}$ et $G = \frac{R}{S}$. On a

$$F + G = \frac{PS + QR}{QS}$$

d'où

$$d^o(F + G) = d^o(PS + QR) - d^o(QS) \leq \max(d^o(PS), d^o(QR)) - d^o(QS)$$

Cette quantité vaut aussi

$$\max(d^o(PS) - d^o(QS), d^o(QR) - d^o(QS)) = \max(d^oP - d^oQ, d^oR - d^oS)$$

Si $d^oF \neq d^oG$, alors $d^oP - d^oQ \neq d^oR - d^oS$, donc $d^o(PS) \neq d^o(QR)$ et on a bien égalité.

□

Proposition 4. Soient $F, G \in \mathbb{K}(X)$. On a $d^o(FG) = d^oF + d^oG$.

Démonstration. Avec les mêmes notations que ci-dessus,

$$d^o(FG) = d^o(PR) - d^o(QS) = d^oP + d^oR - d^oQ - d^oS = d^oF + d^oG$$

□

1.3 Racines, pôles

Définition 2. Soit $F = \frac{P}{Q}$ une fraction écrite sous forme irréductible.

- On appelle racine de F toute racine de P .
- On appelle pôle de F toute racine de Q .

Une fraction a un nombre fini de pôles. Une fraction non nulle a un nombre fini de racines.

1.4 Fonctions rationnelles

Une fonction rationnelle est une fonction $f : D \subset K \rightarrow \mathbb{K}$ vérifiant la propriété suivante :

Il existe $P, Q \in \mathbb{K}[X]$, $Q \neq 0$, tels que pour tout $x \in D$, $Q(x) \neq 0$ et

$$f(x) = \frac{P(x)}{Q(x)}$$

On ne refait pas ce qui a été fait pour les polynômes. Noter simplement qu'une fonction rationnelle est définie sur $\mathbb{K}$ sauf peut-être en un nombre fini de points, les pôles de la fraction rationnelle correspondante. On définit également la dérivée d'une fraction rationnelle.

2 Éléments simples

2.1 C'est quoi ?

Définition 3. On appelle élément simple sur le corps $\mathbb{K}$ toute fraction rationnelle du type

$$F = \frac{P}{Q^n}$$

où $n \in \mathbb{N}^*$, Q est irréductible et $d^\circ P < d^\circ Q$.

2.2 Exemples essentiels

- Lorsque Q est de degré 1, on a ce que l'on appelle un élément simple de première espèce d'ordre n . On a alors

$$F = \frac{\alpha}{(X - a)^n}$$

avec $\alpha \in \mathbb{K}^*$, $a \in \mathbb{K}$ et $n \geq 1$. Ce type d'élément simple est le seul lorsque $\mathbb{K} = \mathbb{C}$.

- Si $\mathbb{K} = \mathbb{R}$, on trouve également des éléments simples de deuxième espèce : les fractions

$$F = \frac{\alpha X + \beta}{(X^2 + pX + q)^n}$$

avec $\alpha, \beta, p, q \in \mathbb{R}$, $(\alpha, \beta) \neq (0, 0)$, $n \geq 1$, et $p^2 - 4q < 0$.

2.3 Partie entière d'une fraction rationnelle

Proposition 5. Soit $F \in \mathbb{K}(X)$. Il existe un unique polynôme $E \in \mathbb{K}[X]$, et une unique fraction rationnelle G de degré strictement négatif, tels que $F = E + G$.

Le polynôme E est la *partie entière* de la fraction F .

Démonstration. On écrit $F = \frac{A}{B}$. On fait la division euclidienne de A par B :

$$A = BE + R$$

avec $d^\circ R < d^\circ B$. D'où

$$F = E + \frac{R}{B}$$

On a montré l'existence, en posant $G = \frac{R}{B}$.

Pour l'unicité, supposons, avec des notations évidentes,

$$F = E_1 + G_1 = E_2 + G_2$$

Alors $G_1 - G_2 = E_2 - E_1$ est un polynôme de degré strictement négatif. C'est donc le polynôme nul, d'où l'unicité. $\square$

2.4 Parties polaires d'une fraction

Proposition 6. Soit $n \geq 1$. Soient $A, B_1, \dots, B_n \in \mathbb{K}[X]$, où les B_k sont non nuls et premiers entre eux deux à deux. Soit

$$F = \frac{A}{B_1 \cdots B_n}$$

On suppose

$$d^o A < \sum_{i=1}^n d^o B_i$$

F s'écrit alors de façon unique

$$F = \sum_{i=1}^n \frac{A_i}{B_i}$$

avec pour tout i , $d^o A_i < d^o B_i$.

Démonstration. On fait la démonstration dans le cas $n = 2$. On étend ensuite par récurrence sur n .

- Existence : Soit $F = \frac{A}{B_1 B_2}$. D'après Bézout, il existe deux polynômes U et V tels que $UB_1 + VB_2 = 1$. On peut donc écrire

$$F = \frac{A(UB_1 + VB_2)}{B_1 B_2} = \frac{VA}{B_1} + \frac{UA}{B_2}$$

On peut ensuite isoler les parties entières de ces fractions, et écrire

$$F = E_1 + E_2 + \frac{A_1}{B_1} + \frac{A_2}{B_2}$$

Mais, forcément, $E_1 + E_2 = 0$ puisque $d^o F < 0$.

- Unicité : Supposons, avec des notations claires, que

$$F = \frac{A_1}{B_1} + \frac{A_2}{B_2} = \frac{A'_1}{B_1} + \frac{A'_2}{B_2}$$

On a alors

$$\frac{A_1 - A'_1}{B_1} = \frac{A'_2 - A_2}{B_2}$$

d'où

$$(A_1 - A'_1)B_2 = (A'_2 - A_2)B_1$$

B_1 divise $(A_1 - A'_1)B_2$. Or, $B_1 \wedge B_2 = 1$. Par le théorème de Gauss, B_1 divise $A_1 - A'_1$. Comme $d^o(A_1 - A'_1) < d^o B_1$, on a nécessairement $A_1 - A'_1 = 0$. De même, $A_2 = A'_2$.

□

Remarque 1. Le résultat précédent permet de « séparer » les pôles d'une fraction.

2.5 Décomposition en éléments simples sur $\mathbb{C}$

Proposition 7. Soit $F \in \mathbb{C}(X)$ une fraction de pôles $a_1, \dots, a_n$ d'ordres de multiplicités respectifs $k_1, \dots, k_n$. Alors

$$F = E + \sum_{j=1}^n \sum_{i=1}^{k_j} \frac{\alpha_{i,j}}{(X - a_j)^i}$$

où les $\alpha_{i,j} \in \mathbb{C}$ et $E \in \mathbb{C}[X]$. La décomposition est unique.

Démonstration. On admet l'unicité (qui a presque été déjà faite, d'ailleurs). Le polynôme E est la partie entière de F . Le j ème terme de la somme est la partie polaire relative au pôle a_j . Elle se décompose en somme d'éléments simples en écrivant la formule de Taylor (ou en faisant des divisions euclidiennes par $X - a_j$). $\square$

2.6 Pôles simples

Soit F une fraction possédant a pour pôle simple. Le théorème précédent nous dit que

$$F = \frac{\alpha}{X - a} + G$$

où a n'est pas un pôle de G . Le scalaire α est appelé le *résidu* de F au point a . Il existe essentiellement 2 méthodes pratiques de calcul de α

- Supposons ici que l'on sait écrire F sous la forme

$$F = \frac{P}{(X - a)Q}$$

les polynômes P et Q n'ayant pas a pour racine. On a alors

$$(X - a)F = \frac{P}{Q} = \alpha + (X - a)G$$

On a donc, en évaluant en a , $\alpha = \frac{P(a)}{Q(a)}$.

Exemple. Soit $F = \frac{X^2+1}{X^3-1}$. Calculons le résidu de F en 1. On a

$$(X - 1)F = \frac{X^2 + 1}{X^2 + X + 1}$$

et cette quantité vaut $2/3$ en 1. Donc,

$$F = \frac{2}{3(X - 1)} + G$$

où G a pour pôles j et $\bar{j}$. Le lecteur calculera de même les résidus en j et $\bar{j}$.

• Il arrive que l'on sache que a est pôle simple de F , sans pouvoir factoriser facilement son dénominateur. On suppose ici que

$$F = \frac{P}{Q}$$

où a est racine simple de Q , et pas racine de P . Par la formule de Taylor pour Q , on obtient :

$$Q(X) = (X - a)Q'(a) + (X - a)^2Q_1$$

où $Q_1 \in \mathbb{K}[X]$. De là,

$$F = \frac{P}{(X - a)(Q'(a) + (X - a)Q_1)}$$

On a donc réussi à factoriser $X - a$ au dénominateur. En appliquant la première méthode, on obtient ainsi que le résidu de F en a est

$$\frac{P(a)}{Q'(a)}$$

Exemple. Soit

$$F = \frac{1}{X^n - 1}$$

Les pôles de F sont les racines n èmes de 1. On a donc (rappelons que $\mathcal{U}_n$ est l'ensemble des racines n^{es} de l'unité)

$$F = \sum_{\omega \in \mathcal{U}_n} \frac{\alpha_\omega}{X - \omega}$$

avec

$$\alpha_\omega = \frac{1}{n\omega^{n-1}} = \frac{\omega}{n}$$

2.7 Pôles multiples

On traite un exemple dans lequel il y a des pôles doubles ...

Soit

$$F = \frac{X^5}{(X - 1)^2(X + 2)^2}$$

Alors

$$F = E + \frac{aX + b}{(X - 1)^2} + \frac{cX + d}{(X + 2)^2}$$

La théorie nous dit que E s'obtient par une division euclidienne, assez pénible à effectuer ici, car il faut développer le dénominateur. On trouve $E = X - 2$. Pour obtenir a et b , on multiplie F par $(X - 1)^2$. On a

$$aX + b + (X - 1)^2G = \frac{X^5}{(X + 2)^2}$$

où la fraction G n'a pas 1 pour pôle. En faisant $X = 1$, on obtient

$$a + b = \frac{1}{9}$$

En dérivant, puis en faisant $X = 1$, on obtient $a = \frac{13}{27}$, d'où $b = -\frac{10}{27}$. On laisse le lecteur terminer la décomposition. On trouve

$$F = X - 2 + \frac{1}{9(X-1)^2} + \frac{13}{27(X-1)} - \frac{32}{9(X+2)^2} + \frac{176}{27(X+2)}$$

2.8 Décomposition en éléments simples sur $\mathbb{R}$

Proposition 8. Soit $F = \frac{P}{Q} \in \mathbb{R}(X)$ une fraction écrite sous forme irréductible. On suppose que Q se factorise en

$$Q = c \prod_{i=1}^m (X - a_i)^{\alpha_i} \prod_{i=1}^n (X^2 + p_i X + q_i)^{\beta_i}$$

où $c \in \mathbb{R}^*$, les α_i, β_i sont des entiers non nuls, les $a_i \in \mathbb{R}$ sont distincts deux à deux, les $(p_i, q_i) \in \mathbb{R}$ sont distincts deux à deux et vérifient $p_i^2 - 4q_i < 0$. La fraction F s'écrit alors de façon unique

$$F = E + \sum_{i=1}^m \sum_{j=1}^{\alpha_i} \frac{\lambda_{ij}}{(X - a_i)^j} + \sum_{i=1}^n \sum_{j=1}^{\beta_i} \frac{\mu_{ij} X + \nu_{ij}}{(X^2 + p_i X + q_i)^j}$$

où E est un polynôme, et les $\lambda_{ij}, \mu_{ij}, \nu_{ij}$ sont des réels.

Démonstration. Démonstration admise. $\square$

On a souvent à décomposer des fractions à coefficients réels en éléments simples. Lorsque les pôles non réels sont des pôles simples, on peut décomposer sur $\mathbb{C}$ et regrouper les parties correspondant à des pôles conjugués. On peut aussi, écrire la décomposition a priori et donner à X des valeurs particulières.

Exemple. Soit

$$F = \frac{1}{X^4 + X^2 + 1}$$

On a

$$F = \frac{1}{(X^2 - X + 1)(X^2 + X + 1)} = \frac{aX + b}{X^2 - X + 1} + \frac{cX + d}{X^2 + X + 1}$$

Tout d'abord, F est paire. Le changement de X en $-X$ donne la même DES. On en déduit que $a = -c$ et $b = d$. Maintenant, multiplions par $X^2 + X + 1$. Il vient

$$\frac{1}{X^2 - X + 1} = cX + d + (X^2 + X + 1)G$$

où G n'a pas j pour pôle. On fait $X = j$:

$$\frac{1}{j^2 - j + 1} = cj + d$$

On utilise les relations $j^2 + j + 1 = 0$ et $j^3 = 1$ pour obtenir que

$$\frac{1}{j^2 - j + 1} = \frac{j + 1}{2}$$

La famille $(1, j)$ étant libre dans le $\mathbb{R}$ -espace vectoriel $\mathbb{C}$, on en déduit $c = d = \frac{1}{2}$. Ainsi

$$\frac{1}{X^4 + X^2 + 1} = \frac{1}{2} \frac{-X + 1}{X^2 - X + 1} + \frac{1}{2} \frac{X + 1}{X^2 + X + 1}$$

3 Primitives des fractions rationnelles

Toute fraction rationnelle se décompose en une somme contenant un polynôme et des éléments simples. Ainsi, pour obtenir une primitive d'une F.R., il suffit de savoir calculer une primitive d'un élément simple. Le lecteur est évidemment supposé savoir calculer une primitive d'un polynôme...

3.1 Primitives d'un élément simple de première espèce

On se donne $F = \frac{1}{(X-a)^n}$, avec $a \in \mathbb{R}$ ou $\mathbb{C}$, et $n \geq 1$. On désire calculer $\int^x F(t) dt$. La variable x reste évidemment réelle.

Lorsque $n \geq 2$, une primitive de F est

$$\int^x F(t) dt = \frac{-1}{(n-1)(x-a)^{n-1}}$$

que a soit réel ou complexe.

Lorsque $n = 1$ et a est réel, une primitive de F sur $] -\infty, a[$ et sur $]a, +\infty[$ est

$$\int^x F(t) dt = \ln |x - a|$$

Lorsque $n = 1$ et $a = p + iq$ avec p, q réels, $q \neq 0$, on a

$$F(x) = \frac{1}{(x-p) - iq} = \frac{(x-p) + iq}{(x-p)^2 + q^2}$$

Une primitive de F sur $\mathbb{R}$ est donc

$$\int^x F(t) dt = \frac{1}{2} \ln((x-p)^2 + q^2) + i \arctan \frac{x-p}{q}$$

Exemple. On a

$$\frac{1}{X-i} = \frac{X+i}{X^2+1}$$

Ainsi, pour tout réel x ,

$$\begin{aligned} \int^x \frac{dt}{t-i} &= \int^x \frac{t+i}{t^2+1} dt \\ &= \int^x \frac{t}{t^2+1} dt + i \int^x \frac{dt}{t^2+1} \\ &= \frac{1}{2} \ln(x^2+1) + i \arctan x \end{aligned}$$

3.2 Primitives d'un élément simple de deuxième espèce d'ordre 1

On suppose ici les fractions à coefficients réels. On se donne

$$F = \frac{\alpha X + \beta}{X^2 + pX + q}$$

avec un dénominateur irréductible. Une simple translation ($t = x + \frac{p}{2}$) permet de se ramener au calcul d'une primitive de

$$G(t) = \frac{ut + v}{t^2 + a^2}$$

avec a réel non nul, ce qui s'intègre en

$$\frac{u}{2} \ln(t^2 + a^2) + \frac{v}{a} \arctan \frac{t}{a}$$

3.3 Primitives d'un élément simple de deuxième espèce d'ordre au moins 2

Une translation de la variable, comme au paragraphe ci-dessus, permet de se ramener au problème suivant : calculer

$$I_n = \int^x \frac{t}{(t^2 + a^2)^n} dt$$

et

$$J_n = \int^x \frac{1}{(t^2 + a^2)^n} dt$$

où $n \geq 2$.

Le calcul de I_n est immédiat :

$$I_n = \frac{-1}{2(n-1)(x^2 + a^2)^{n-1}}$$

Pour calculer J_n , on intègre par parties en posant $u' = 1$ et $v = \frac{1}{(x^2+a^2)^n}$. Il vient

$$J_n = \frac{x}{(x^2 + a^2)^n} + 2n \int^x \frac{t^2}{(t^2 + a^2)^{n+1}} dt$$

L'intégrale vaut $J_n - a^2 J_{n+1}$. Connaissant J_1 , on peut ainsi calculer de proche en proche les intégrales J_n .

Exercice. Calculer

$$I = \int_0^1 \frac{dx}{(x^2 + 1)^2}$$

On évalue $\int_0^1 \frac{dx}{x^2+1}$ en intégrant par parties. Posons $u' = 1$ et $v = \frac{1}{x^2+1}$. Il vient

$$\begin{aligned} \int_0^1 \frac{dx}{x^2 + 1} &= \left[\frac{x}{x^2 + 1} \right]_0^1 + 2 \int_0^1 \frac{x^2}{(x^2 + 1)^2} dx \\ &= \frac{1}{2} + 2 \int_0^1 \frac{(x^2 + 1) - 1}{(x^2 + 1)^2} dx \\ &= \frac{1}{2} + 2 \int_0^1 \frac{1}{x^2 + 1} dx - 2 \int_0^1 \frac{1}{(x^2 + 1)^2} dx \\ &= \frac{1}{2} + 2 \int_0^1 \frac{1}{x^2 + 1} dx - 2I \end{aligned}$$

De là,

$$2I = \int_0^1 \frac{1}{x^2 + 1} dx - \frac{1}{2} = \frac{\pi}{4} + \frac{1}{2}$$

et finalement

$$I = \frac{\pi}{8} + \frac{1}{4}$$

Chapitre 20

Analyse asymptotique

1 Comparaison des fonctions au voisinage d'un point

Dans cette section, nous allons explorer la notion de fonction négligeable ou dominée par une autre fonction, ainsi que la notion de fonctions équivalentes. Nous allons retrouver des notions que nous avons déjà étudiées pour les suites. La différence essentielle est que pour les fonctions tout se passe au voisinage d'un point $a \in \overline{\mathbb{R}}$ à préciser, alors que pour les suites c'est toujours lorsque l'indice n tend vers l'infini.

1.1 Négligeabilité, domination

Définition 1. Soit $a \in \overline{\mathbb{R}}$. Soient f et g deux fonctions définies au voisinage de a (sauf peut-être en a). On dit que

- f est dominée par g en a lorsque $f = gh$, où la fonction h est bornée au voisinage de a . On note alors $f = O_a(g)$.
- f est négligeable devant g en a lorsque $f = \varepsilon g$, où ε est une fonction qui tend vers 0 en a . On note alors $f = o_a(g)$ ou $f \prec_a g$.

Remarque 1. Lorsque la fonction g ne s'annule pas au voisinage de a , les définitions ci-dessus équivalent respectivement à $\left(\frac{f}{g}\right)$ est bornée au voisinage de a , et à $\left(\frac{f}{g}\right)$ tend vers 0 en a .

Remarque 2. Si le point a est clair et sans ambiguïté, on notera plus légèrement $o(g)$ et $O(g)$.

Exemple.

- On a $3x^2 - 1 = O_{+\infty}(x^2)$. En effet,

$$\frac{3x^2 - 1}{x^2} = 3 - \frac{1}{x^2} \longrightarrow 3$$

lorsque x tend vers $+\infty$. Ce quotient, ayant une limite finie, est borné au voisinage de $+\infty$.

- Étant donnés deux réels $\alpha < \beta$, on a $x^\alpha = o_{+\infty}(x^\beta)$ mais $x^\beta = o_0(x^\alpha)$. Il suffit pour le voir de considérer le quotient $\frac{x^\beta}{x^\alpha} = x^{\beta-\alpha}$.

Proposition 1. Soit f une fonction définie au voisinage de a . Alors

$$o_a(f) + o_a(f) = o_a(f)$$

et

$$o_a(f) \times o_a(f) = o_a(f^2)$$

Si f est bornée au voisinage de a , alors

$$o_a(f) \times o_a(f) = o_a(f)$$

Démonstration. On a

$$o_a(f) + o_a(f) = (\varepsilon + \varepsilon')f = \varepsilon''f$$

où toutes les fonctions ε tendent vers 0 en a . Preuves analogues dans tous les cas. $\square$

Remarque 3. On remarque, dans la proposition précédente, que $1 + 1$ ça ne fait pas 2 et que les égalités avec des petits o ou des grands O sont à considérer avec précaution. Il faut effectivement faire attention lorsqu'on manipule des o , car $o_a(f)$ représente UNE fonction négligeable devant f en a . Donc, par exemple lorsqu'on écrit $o_a(f) + o_a(f)$ on additionne deux fonctions négligeables devant f , mais ces deux fonctions sont a priori différentes. En cas de doute, il n'y a qu'à revenir à la définition : $o_a(f) = \varepsilon f$ où la fonction ε tend vers 0 au point a considéré.

1.2 Fonctions équivalentes

Définition 2. Soient deux fonctions f et g définies au voisinage de $a \in \overline{\mathbb{R}}$, sauf peut-être en a . On dit que ces fonctions sont équivalentes en a , et on note $f \sim_a g$, lorsque $f = \lambda g$, où $\lambda(x) \rightarrow 1$ lorsque x tend vers a .

Remarque 4. Celà revient à dire que $\frac{f}{g}$ tend vers 1 en a , dans le cas où la fonction g ne s'annule pas au voisinage de a .

Remarque 5. Lorsque le point a est clair, on note $f \sim g$ plutôt que $f \sim_a g$, histoire d'alléger.

Proposition 2. *L'équivalence des fonctions en un point est réflexive, symétrique et transitive : c'est une relation d'équivalence.*

Proposition 3. *Si deux fonctions sont équivalentes en un point, et si l'une des deux tend vers une limite, alors l'autre tend aussi vers cette même limite.*

Démonstration. Supposons $f \sim_a g$ et $g \rightarrow_a \ell$. On a $f = \lambda g$ avec $\lambda \rightarrow_a 1$, d'où le résultat. $\square$

Proposition 4. *Si deux fonctions sont équivalentes en un point, elles ont même signe au voisinage de ce point.*

Démonstration. Supposons $f \sim_a g$. On a $f = \lambda g$ avec $\lambda \rightarrow_a 1$. En appliquant la définition de limite, $\lambda(x) > 0$ pour tout x voisin de a , d'où le résultat. $\square$

Proposition 5. *Soient f et g deux fonctions définies au voisinage de a . On a*

$$f \sim_a g \iff f = g + o_a(g)$$

Démonstration. Tout simplement parce qu'une fonction λ tend vers 1 si et seulement si on peut écrire $\lambda = 1 + \varepsilon$ où $\varepsilon \rightarrow 0$. $\square$

Exemple. Cet exemple est essentiel. Soit P un polynôme. En $\pm\infty$ P est équivalent à son terme de plus haut degré. En 0, il est équivalent à son terme de plus bas degré. Par exemple,

$$3x^4 + 7x^2 - 2x - 1 \sim_{+\infty} 3x^4$$

et

$$5x^3 - 2x^2 - 3x \sim_0 -3x$$

Proposition 6. *L'équivalence des fonctions en un point est compatible avec la multiplication, la division, et l'élevation à une puissance fixée. En revanche, elle n'est pas compatible avec l'addition.*

Démonstration. De façon plus explicite, soit $a \in \overline{\mathbb{R}}$. Soient f, g, f', g' des fonctions définies au voisinage de a , sauf peut-être en a . Supposons $f \sim_a g$ et $f' \sim_a g'$. Il s'agit de montrer que

- $ff' \sim_a gg'$
- $\frac{f}{f'} \sim_a \frac{g}{g'}$
- Pour tout réel α , $f^\alpha \sim_a g^\alpha$.

Le produit et le quotient de deux fonctions qui tendent vers 1 tend encore vers 1. Si la fonction λ tend vers 1 et si $\alpha \in \mathbb{R}$, alors λ^α tend encore vers 1. L'addition pose en revanche problème. Par exemple, on a $x^2 \sim_0 x^2 + x^3$ et $-x^2 + x^4 \sim_0 -x^2$. Mais on n'a PAS $x^4 \sim_0 x^3$. $\square$

1.3 Équivalents classiques

Proposition 7. Soit $\alpha \in \mathbb{R}$. On a les équivalents suivants en 0 :

$e^x - 1$	$\sim_0$	x
$\ln(1+x)$	$\sim_0$	x
$(1+x)^\alpha - 1$	$\sim_0$	αx
$\sin x$	$\sim_0$	x
$1 - \cos x$	$\sim_0$	$\frac{x^2}{2}$
$\tan x$	$\sim_0$	x
$\sinh x$	$\sim_0$	x
$\cosh x - 1$	$\sim_0$	$\frac{x^2}{2}$
$\tanh x$	$\sim_0$	x
$\arcsin x$	$\sim_0$	x
$\arctan x$	$\sim_0$	x
$\operatorname{argsh} x$	$\sim_0$	x
$\operatorname{argth} x$	$\sim_0$	x

Démonstration. La plupart de ces formules se montrent de la même façon, en montrant qu'un taux d'accroissement tend vers une certaine dérivée. Par exemple,

$$\frac{e^x - 1}{x} = \frac{e^x - e^0}{x - 0}$$

tend vers la dérivée de l'exponentielle en 0, c'est à dire 1. Seules deux formules, celles du cosinus et celle du cosinus hyperbolique, se traitent à part. On a par exemple, pour le cosinus,

$$\frac{1 - \cos x}{x^2} = \frac{2 \sin^2 \frac{x}{2}}{x^2} \sim \frac{2 \frac{x^2}{4}}{x^2} = \frac{1}{2} \longrightarrow \frac{1}{2}$$

Donc,

$$1 - \cos x \sim_0 \frac{x^2}{2}$$

□

1.4 Fonctions de référence

Proposition 8. Soient $a > 1$, $\alpha > 0$, $\beta > 0$. On a en $+\infty$:

- $(\ln x)^\beta = o_{+\infty}(x^\alpha)$
- $x^\alpha = o_{+\infty}(a^x)$

Démonstration. Voir le chapitre sur les fonctions usuelles. $\square$

Remarque 6. Ce résultat permet de lever des indéterminations pour certaines limites en $\pm\infty$ et en 0. Par exemple, $\frac{x^3}{e^x} \rightarrow 0$ en $+\infty$ puisque $x^3 = o(e^x)$.

Autre exemple, considérons $f(x) = x^2 \ln^7 x$, définie au voisinage de 0. Posons $t = \frac{1}{x}$. On a

$$f(x) = -\frac{\ln^7 t}{t^2}$$

Lorsque $x \rightarrow 0$, $t \rightarrow +\infty$ donc $f(x) \rightarrow 0$ lorsque $x \rightarrow 0$.

Dernier exemple, soit $f(x) = x^3 e^x$. Posons $t = -x$. On a

$$f(x) = -\frac{t^3}{e^t}$$

Lorsque $x \rightarrow -\infty$, on a $t \rightarrow +\infty$, donc $f(x)$ tend vers 0.

1.5 Un exemple de calcul d'équivalent

Les techniques de calcul sont les mêmes que celles vues dans le chapitre de comparaison de suites :

- Se ramener à une variable qui tend vers 0
- Chercher à faire apparaître des équivalents usuels
- Se méfier des sommes

Prenons un exemple déjà vu pour les suites, en le transposant à des fonctions.

Soit a un réel strictement positif, $a \neq 1$. Cherchons un équivalent de $f(x) = a^x - x^a$ lorsque x tend vers a .

- Se ramener à 0. Posons $x = a + h$ de sorte que h tend vers 0 lorsque x tend vers a .
- Faire apparaître des équivalents usuels. On a

$$f(x) = a^{a+h} - (a+h)^a = a^a \left(a^h - \left(1 + \frac{h}{a} \right)^a \right)$$

Tout d'abord, $a^h = e^{h \ln a}$ donc

$$a^h - 1 = e^{h \ln a} - 1 \sim h \ln a$$

Ensuite,

$$\left(1 + \frac{h}{a} \right)^a - 1 \sim a \frac{h}{a} = h$$

- Se méfier des sommes. Il faut additionner les deux équivalents précédents, ce qui est interdit. On écrit donc ces équivalents comme des égalités avec des petits o .

$$a^h - 1 = h \ln a + o(h \ln a) = h \ln a + o(h)$$

et

$$\left(1 + \frac{h}{a}\right)^a - 1 = h + o(h)$$

Regroupant le tout, il vient

$$f(x) = a^a(h \ln a - h + o(h)) = a^a(\ln a - 1)h + o(h)$$

Finalement, si $a \neq e$,

$$f(x) \sim_a a^a(\ln a - 1)(x - a)$$

Pourquoi si $a \neq e$, parce que pour $a = e$ on obtient

$$f(x) = o(h) = o_a(x - a)$$

Dans ce cas nous avons pas trouvé d'équivalent de f en a . Il nous faudra attendre la section sur les développements limités pour en obtenir un. Patience!

2 Développements limités

2.1 Formule de Taylor-Young

Proposition 9. Soit $n \geq 0$. Soit I un intervalle non trivial de $\mathbb{R}$. Soit $a \in I$. Soit $f : I \rightarrow \mathbb{R}$ n fois dérivable en a . Alors

$$f(x) = \sum_{k=0}^n \frac{(x-a)^k}{k!} f^{(k)}(a) + o_a((x-a)^n)$$

Remarque 7. Pour $n = 0$, à condition d'interpréter « 0 fois dérivable » comme « continue » ce théorème nous dit que si f est continue en a , alors $f(x) = f(a) + o(1)$, c'est à dire que $f(x)$ tend vers $f(a)$ lorsque x tend vers a . Normal, c'est la définition de la continuité.

Pour $n = 1$, le théorème affirme que si f est dérivable en a alors

$$f(x) = f(a) + (x-a)f'(a) + o_a(x-a)$$

Nous connaissons déjà ce résultat, vu dans le cours de dérivation.

Démonstration. Supposons dans ce qui suit $n \geq 2$. Posons

$$g(x) = f(x) - \sum_{k=0}^n \frac{(x-a)^k}{k!} f^{(k)}(a)$$

On constate (exercice!) que g , définie au voisinage de a , est n fois dérivable en a , et que

$$g(a) = g'(a) = \dots = g^{(n)}(a) = 0$$

Cela signifie en particulier que g est $n - 1$ fois dérivable dans un voisinage de a . On utilise le théorème des accroissements finis pour écrire

$$g(x) = g(x) - g(a) = (x - a)g'(c_1)$$

avec $a < c_1 < x$ (ou $x < c_1 < a$). De même

$$g'(c_1) = g'(c_1) - g'(a) = (c_1 - a)g''(c_2)$$

avec $a < c_2 < c_1$. On continue jusqu'à

$$g^{(n-2)}(c_{n-2}) = (c_{n-2} - a)g^{(n-1)}(c_{n-1})$$

On a donc

$$g(x) = (x - a)(c_1 - a) \cdots (c_{n-2} - a)g^{(n-1)}(c_{n-1})$$

avec $a < c_{n-1} < \cdots < c_1 < x$. Ainsi,

$$\left| \frac{g(x)}{(x - a)^n} \right| \leq \left| \frac{g^{(n-1)}(c_{n-1})}{x - a} \right|$$

Or,

$$\left| \frac{g^{(n-1)}(c_{n-1})}{x - a} \right| \leq \left| \frac{g^{(n-1)}(c_{n-1}) - g^{(n-1)}(a)}{c_{n-1} - a} \right|$$

et cette quantité tend vers $g^{(n)}(a) = 0$ lorsque x tend vers a . $\square$

2.2 Notion de DL

Définition 3. Soit f une fonction définie au voisinage du réel a . Soit $n \in \mathbb{N}$. On dit que f admet un développement limité à l'ordre n en a lorsqu'il existe $a_0, \dots, a_n \in \mathbb{R}$ tels que

$$f(x) = \sum_{k=0}^n a_k (x - a)^k + o_a((x - a)^n)$$

Remarque 8. Lorsque le point a est clair, nous noterons simplement o au lieu de o_a . De plus, nous écrirons $o(x - a)^n$ au lieu de $o((x - a)^n)$, histoire de ne pas alourdir les notations.

Remarque 9. On peut encore écrire

$$f(x) = P_n(x - a) + o(x - a)^n$$

avec $P_n \in \mathbb{R}_n[X]$. La quantité $P_n(x - a)$ est appelée la *partie régulière* du DL, et $o(x - a)^n$ est le *reste*.

Proposition 10. *Si f est n fois dérivable en $a \in \mathbb{R}$, elle admet un DL à l'ordre n en a .*

Démonstration. C'est exactement ce que dit le théorème de Taylor-Young. Ce théorème donne également les coefficients du DL. $\square$

Proposition 11. *Soit f une fonction définie au voisinage du réel a .*

- *f admet un DL à l'ordre 0 en a si et seulement si f est continue en a .*
- *f admet un DL à l'ordre 1 en a si et seulement si f est dérivable en a .*
- *On ne peut rien dire d'autre.*

Démonstration.

- f a un DL à l'ordre 0 en a si et seulement si il existe $\alpha \in \mathbb{R}$ tel que $f(x) = \alpha + o(1)$. Ceci revient à dire que f a une limite réelle α en a . C'est la définition de la continuité en a .
- f a un DL à l'ordre 1 en a si et seulement si il existe $\alpha, \beta \in \mathbb{R}$ tels que

$$f(x) = \alpha + \beta(x - a) + o(x - a)$$

Nécessairement, $\alpha = f(a)$ (faire tendre x vers a). On a vu dans le cours de dérivation que l'égalité

$$f(x) = f(a) + \beta(x - a) + o(x - a)$$

équivalait à la dérivabilité de f en a , et que $f'(a) = \beta$.

- Soit $f : \mathbb{R} \rightarrow \mathbb{R}$ définie par $f(0) = 0$ et, pour tout $x \neq 0$,

$$f(x) = x^3 \sin \frac{1}{x}$$

On vérifie point par point les résultats suivants :

- ★ f est de classe $\mathcal{C}^\infty$ sur $\mathbb{R}_+^*$ et sur $\mathbb{R}_-^*$.
- ★ Pour tout $x \neq 0$, $f'(x) = 3x^2 \sin \frac{1}{x} - x \cos \frac{1}{x}$.
- ★ f est dérivable en 0 et $f'(0) = 0$.

Pour tout $x \neq 0$, on a donc

$$\frac{f'(x) - f'(0)}{x - 0} = 3x \sin \frac{1}{x} - \cos \frac{1}{x}$$

Il est clair que cette quantité n'a pas de limite en 0. La fonction f n'est donc pas deux fois dérivable en 0. Cependant,

$$f(x) = x^2 \times \left(x \sin \frac{1}{x}\right) = o(x^2)$$

puisque $x \sin \frac{1}{x}$ tend vers 0 lorsque x tend vers 0. Ainsi, f a un développement limité à l'ordre 2 en 0.

□

Proposition 12. *Soit f une fonction définie au voisinage du réel a . Soit $n \in \mathbb{N}$. Si f admet un DL à l'ordre n au point a , ce DL est unique.*

Démonstration. Supposons

$$\sum_{k=0}^n a_k(x-a)^k + o(x-a)^n = \sum_{k=0}^n b_k(x-a)^k + o(x-a)^n$$

On en déduit

$$\sum_{k=0}^n c_k(x-a)^k = o(x-a)^n$$

où l'on a posé $c_k = a_k - b_k$. Supposons l'un au moins des c_k non nul. Soit j le plus petit indice tel que $c_j \neq 0$. Alors,

$$\sum_{k=0}^n c_k(x-a)^k \sim_a c_j(x-a)^j$$

et cette quantité n'est pas négligeable devant $(x-a)^n$. On a donc une contradiction. Ainsi, tous les c_k sont nuls, c'est à dire

$$\forall k \in \llbracket 0, n \rrbracket, a_k = b_k$$

□

Exemple. Si une fonction est définie au voisinage de 0, et paire, tout D.L. de cette fonction en 0 a ses coefficients de degré impair nuls. Remarque analogue pour une fonction impaire.

2.3 Utilité des développements limités

Les développements limités permettent de connaître le comportement d'une fonction au voisinage d'un point : limite éventuelle de cette fonction, signe de la fonction, extrema,...

Par exemple, nous savons que si une fonction est dérivable en un point et y possède un extremum, alors sa dérivée s'annule. Mais nous savons aussi que la réciproque est fautive. Si l'on possède un développement limité à l'ordre 2 de la fonction au point considéré, on peut souvent trancher (mais pas toujours, hélas).

Proposition 13. *Soit f une fonction définie au voisinage du réel a . On suppose que f admet en a un DL à l'ordre 2 de la forme*

$$f(x) = f(a) + \lambda(x - a)^2 + o(x - a)^2$$

- Si $\lambda > 0$, f admet un minimum local en a .
- Si $\lambda < 0$, f admet un maximum local en a .
- Si $\lambda = 0$ on ne peut rien dire.

Démonstration. Supposons $\lambda \neq 0$. On a

$$f(x) - f(a) \sim \lambda(x - a)^2$$

du signe de λ au voisinage de a . $\square$

Remarque 10. Si $\lambda = 0$, une solution peut consister à faire un DL de f en a à un ordre plus élevé. lequel ? Suffisamment élevé pour qu'un coefficient non nul apparaisse dans le DL.

2.4 Développements classiques

Voici ci-dessous les développements limités **EN 0** pour les fonctions usuelles, à un ordre quelconque.

Proposition 14. *Pour tout entier naturel n on a, lorsque x tend vers 0,*

$$\begin{aligned}
 e^x &= \sum_{k=0}^n \frac{x^k}{k!} + o(x^n) \\
 \ln(1+x) &= \sum_{k=1}^n (-1)^{k+1} \frac{x^k}{k} + o(x^n) \\
 \arctan x &= \sum_{k=0}^n (-1)^k \frac{x^{2k+1}}{2k+1} + o(x^{2n+2}) \\
 \operatorname{argth} x &= \sum_{k=0}^n \frac{x^{2k+1}}{2k+1} + o(x^{2n+2}) \\
 \sin x &= \sum_{k=0}^n (-1)^k \frac{x^{2k+1}}{(2k+1)!} + o(x^{2n+2}) \\
 \cos x &= \sum_{k=0}^n (-1)^k \frac{x^{2k}}{(2k)!} + o(x^{2n+1}) \\
 \sinh x &= \sum_{k=0}^n \frac{x^{2k+1}}{(2k+1)!} + o(x^{2n+2}) \\
 \cosh x &= \sum_{k=0}^n \frac{x^{2k}}{(2k)!} + o(x^{2n+1}) \\
 (1+x)^\alpha &= \sum_{k=0}^n \binom{\alpha}{k} x^k + o(x^n) \quad (\alpha \in \mathbb{R}) \\
 \tan x &= x + \frac{1}{3}x^3 + \frac{2}{15}x^5 + o(x^5)
 \end{aligned}$$

Remarque 11. Pour le lecteur qui se demanderait ce que vaut $\binom{\frac{1}{2}}{3}$, on pose, pour tout $\alpha \in \mathbb{R}$,

$$\binom{\alpha}{k} = \frac{\alpha(\alpha-1)\dots(\alpha-k+1)}{k!}$$

si $k \neq 0$, et $\binom{\alpha}{0} = 1$. Ainsi, par exemple,

$$\binom{\frac{1}{2}}{3} = \frac{\frac{1}{2} \frac{-1}{2} \frac{-3}{2}}{3!} = \frac{1}{16}$$

Démonstration. Toutes les fonctions ci-dessus sont de classe $\mathcal{C}^\infty$ au voisinage de 0. Par le théorème de Taylor-Young elles admettent donc des développements limités à tous les ordres en 0. Mises à part les fonctions $\arctan$ et argth , pour lesquelles la preuve attendra un

peu, la démonstration consiste à calculer la dérivée d'ordre k de la fonction puis à l'évaluer en 0. Traitons juste le cas important de $f(x) = (1+x)^\alpha$.

On a, pour tout x voisin de 0, $f'(x) = \alpha(1+x)^{\alpha-1}$, $f''(x) = \alpha(\alpha-1)(1+x)^{\alpha-2}$, et par une récurrence facile, pour tout $k \geq 0$,

$$f^{(k)}(x) = \alpha(\alpha-1) \dots (\alpha-k+1)(1+x)^{\alpha-k}$$

Ainsi,

$$f^{(k)}(0) = \alpha(\alpha-1) \dots (\alpha-k+1)$$

On a donc

$$\frac{f^{(k)}(0)}{k!} = \frac{\alpha(\alpha-1) \dots (\alpha-k+1)}{k!} = \binom{\alpha}{k}$$

Il ne reste plus qu'à appliquer Taylor-Young pour obtenir le D.L. de f à l'ordre n en 0. $\square$

Exemple. On a $\binom{\frac{1}{2}}{0} = 1$, $\binom{\frac{1}{2}}{1} = \frac{1}{2}$ et $\binom{\frac{1}{2}}{2} = \frac{\frac{1}{2}(\frac{1}{2}-1)}{2} = -\frac{1}{8}$. Ainsi, on a en 0

$$\sqrt{1+x} = 1 + \frac{1}{2}x - \frac{1}{8}x^2 + o(x^2)$$

Exercice. Déterminer le D.L. à l'ordre 3 en 0 de $\frac{1}{\sqrt{1-x^2}}$.

Remarque 12. Un autre cas important est le cas $\alpha = -1$. On a pour tout $k \geq 1$,

$$\binom{-1}{k} = \frac{(-1)(-2) \dots (-k)}{k!} = (-1)^k$$

Ainsi, un DL de $\frac{1}{1+x}$ en 0 à l'ordre n est

$$\frac{1}{1+x} = \sum_{k=0}^n (-1)^k x^k + o(x^n)$$

En remplaçant x par $-x$, un DL à l'ordre n en 0 de $\frac{1}{1-x}$ est

$$\frac{1}{1-x} = \sum_{k=0}^n x^k + o(x^n)$$

En remplaçant x par $-x^2$, un DL à l'ordre $2n$ en 0 de $\frac{1}{1-x^2}$ est

$$\frac{1}{1-x^2} = \sum_{k=0}^n x^{2k} + o(x^{2n})$$

et même $o(x^{2n+1})$ puisque la fonction $x \mapsto \frac{1}{1-x^2}$ est paire.

Exemple. Depuis la fin de la section précédente, une question nous obsède : quel est un équivalent de $f(x) = e^x - x^e$ lorsque x tend vers e ? Nous pouvons maintenant y répondre

en faisant un DL à l'ordre 2 de f en e . Posons $x = e + h$, de sorte que $h \rightarrow 0$ lorsque $x \rightarrow e$.

$$\begin{aligned}
 f(x) &= e^{e+h} - (e+h)^e \\
 &= e^e \left(e^h - \left(1 + \frac{h}{e} \right)^e \right) \\
 &= e^e \left(1 + h + \frac{1}{2}h^2 + o(h^2) - \left(1 + \binom{e}{1} \frac{h}{e} + \binom{e}{2} \frac{h^2}{e^2} + o(h^2) \right) \right) \\
 &= e^e \left(1 + h + \frac{1}{2}h^2 + o(h^2) - \left(1 + e \frac{h}{e} + \frac{e(e-1)}{2} \frac{h^2}{e^2} + o(h^2) \right) \right) \\
 &= \frac{1}{2} e^{e-1} h^2 + o(h^2) \\
 &\sim \frac{1}{2} e^{e-1} h^2
 \end{aligned}$$

Ainsi,

$$e^x - x^e \sim_e \frac{1}{2} e^{e-1} (x - e)^2$$

Remarque 13. Comment estimer *numériquement* la validité d'un tel résultat ? Si nous posons $g(x) = \frac{1}{2} e^{e-1} (x - e)^2$, alors $f(x) - g(x) = o(x - e)^2$: en prenant un réel x « proche » de e , $f(x) - g(x)$ devrait être « très proche » de 0. Ici, en prenant $x = e + 10^{-1}$ et une calculatrice, on obtient $f(x) - g(x) \simeq 2.2 \times 10^{-3}$, ce qui ne prouve rien mais qui est assez rassurant. En poussant le calcul un peu plus loin et en faisant un DL à l'ordre 3, on montre qu'en fait

$$f(x) - g(x) \sim \frac{1}{6} e^{e-2} (3e - 2) (x - e)^3$$

Il s'avère que cette quantité, évaluée en $e + 10^{-1}$, vaut environ 2.1×10^{-3} .

3 Opérations sur les développements limités

Nous allons voir dans cette section que les DLs peuvent s'ajouter, se soustraire, se multiplier, se diviser, se composer, et même s'intégrer... mais pas se dériver ! Les exemples donnés pour chaque opération sont évidemment essentiels pour bien comprendre comment opérer en pratique sur les DLs.

Notation. Si $P = \sum_{k=0}^{\infty} a_k X^k \in \mathbb{R}[X]$, et $n \in \mathbb{N}$, on note

$$P|_n = \sum_{k=0}^n a_k X^k$$

le polynôme P que l'on a tronqué au degré n (suppression des termes de degré $> n$).

N.B : le point fondamental dans les calculs qui suivent est la manipulation des petits o . Nous utiliserons sans modération les deux résultats suivants :

- Soient $p \in \mathbb{N}$, $\lambda \in \mathbb{R}^*$ et $a \in \mathbb{R}$. Alors

$$o_a(\lambda(x-a)^p) = o_a((x-a)^p)$$

- Si $0 \leq p \leq q$ sont deux entiers et $a \in \mathbb{R}$, alors le plus petit gagne :

$$o_a(x-a)^p + o_a(x-a)^q = o_a(x-a)^p$$

3.1 Somme de DLs

Soient f et g deux fonctions définies au voisinage du point a , et admettant en a des DLs d'ordres respectifs p et q . On a

$$f(x) = P(x-a) + o(x-a)^p$$

et

$$g(x) = Q(x-a) + o(x-a)^q$$

où $P \in \mathbb{R}_p[X]$ et $Q \in \mathbb{R}_q[X]$. On en tire

$$f(x) + g(x) = P(x-a) + Q(x-a) + o(x-a)^p + o(x-a)^q$$

ou encore

$$(f+g)(x) = (P+Q)|_r(x) + o(x-a)^r$$

où $r = \min(p, q)$. Bref, additionner des DLs ne pose aucune difficulté.

Exemple. Un DL à l'ordre 4 en 0 de $\sin x$ est

$$\sin x = x - \frac{1}{6}x^3 + o(x^4)$$

Un DL à l'ordre 4 en 0 de $\cos x$ est

$$\cos x = 1 - \frac{1}{2}x^2 + \frac{1}{24}x^4 + o(x^4)$$

Donc, un DL à l'ordre 4 en 0 de $\sin x + \cos x$ est

$$\sin x + \cos x = 1 + x - \frac{1}{2}x^2 - \frac{1}{6}x^3 + \frac{1}{24}x^4 + o(x^4)$$

3.2 Produit de DLs

Soient f et g deux fonctions définies au voisinage du point a , et admettant en a des DLs d'ordres respectifs p et q . On a

$$f(x) = P(x-a) + o(x-a)^p$$

et

$$g(x) = Q(x-a) + o(x-a)^q$$

où $P \in \mathbb{R}_p[X]$ et $Q \in \mathbb{R}_q[X]$. Posons plus précisément

$$P(x-a) = (x-a)^\mu P_1(x-a)$$

où $P_1(0) \neq 0$, et de même

$$Q(x-a) = (x-a)^\nu Q_1(x-a)$$

où $Q_1(0) \neq 0$. Il vient alors

$$\begin{aligned} f(x)g(x) &= ((x-a)^\mu P_1(x-a) + o(x-a)^p)((x-a)^\nu Q_1(x-a) + o(x-a)^q) \\ &= P(x-a)Q(x-a) + o(x-a)^r \\ &= P(x-a)Q(x-a)|_r + o(x-a)^r \end{aligned}$$

où $r = \min(p+\nu, q+\mu)$. Le lecteur est invité à bien comprendre ces quelques lignes, surtout la seconde : tout l'intérêt des DLs est de n'écrire que ce qui est nécessaire !

Exemple. Déterminons un DL à l'ordre 3 en 0 de $e^x \sin x$. Avec les notations ci-dessus, on a $\mu = 0$ et $\nu = 1$. Il suffit donc de prendre $p = 2$ et $q = 3$.

$$e^x \sin x = (1 + x + \frac{1}{2}x^2 + o(x^2))(x - \frac{1}{6}x^3 + o(x^3)) = x + x^2 + \frac{1}{3}x^3 + o(x^3)$$

Lorsque les calculs deviennent un peu compliqués, ce n'est pas une mauvaise idée de calculer le produit en stockant les coefficients dans un tableau.

Calculons par exemple le DL à l'ordre 4 en 0 de $e^x \ln(1+x)$. Un DL à l'ordre 3 de e^x suffit car le DL de $\ln(1+x)$ « commence » par x . On a

$$\begin{array}{rclclcl} e^x & = & 1 & +x & +\frac{1}{2}x^2 & +\frac{1}{6}x^3 & +o(x^3) \\ \ln(1+x) & = & & x & -\frac{1}{2}x^2 & +\frac{1}{3}x^3 & -\frac{1}{4}x^4 & +o(x^4) \end{array}$$

On effectue le produit en rangeant chaque coefficient dans la bonne case. Les termes de degré strictement supérieur à 4 ne sont pas pris en compte, ils sont « absorbés » par le $o(x^4)$.

x^0	0	= 0
x^1	1	= 1
x^2	$-\frac{1}{2} + 1$	= $\frac{1}{2}$
x^3	$\frac{1}{3} - \frac{1}{2} + \frac{1}{2}$	= $\frac{1}{3}$
x^4	$-\frac{1}{4} + \frac{1}{3} - \frac{1}{2} + \frac{1}{6}$	= 0

Ainsi,

$$e^x \ln(1+x) = x + \frac{1}{2}x^2 + \frac{1}{3}x^3 + o(x^4)$$

3.3 Inverse d'un DL

Soit f ayant un développement limité à l'ordre n au point a . On suppose que $f(a) \neq 0$, et on cherche un DL en a de $\frac{1}{f(x)}$. Pour simplifier, on suppose $f(a) = 1$ (si ce n'est pas le cas on peut facilement s'y ramener en factorisant $f(a)!$), et on écrit

$$f(x) = 1 + P(x - a) + o(x - a)^p$$

où $P(0)=0$. On a

$$\frac{1}{f(x)} - \frac{1}{1 + P(x - a)} = \frac{o(x - a)^p}{f(x)(1 + P(x - a))} = o(x - a)^p$$

ou encore

$$\frac{1}{f(x)} = \frac{1}{1 + P(x - a)} + o(x - a)^p$$

Utilisons le DL usuel de $(1 + x)^{-1} = \frac{1}{1+x}$. On a

$$\frac{1}{1 + x} = 1 - x + x^2 - \dots + (-1)^k x^k + o_0(x^k)$$

et donc, puisque $P(x - a)$ tend vers 0 lorsque x tend vers a :

$$\frac{1}{1 + P(x - a)} = 1 - P(x - a) + P(x - a)^2 - \dots + (-1)^k P(x - a)^k + o(P(x - a)^k)$$

Écrivons $P(x - a) = (x - a)^\mu P_1(x - a)$ où $P_1(0) \neq 0$. Alors $o(P(x - a)^k) = o(x - a)^{\mu k}$. Donc,

$$\frac{1}{1 + P(x - a)} = Q(x - a) + o(P(x - a)^p)$$

où Q est un polynôme de degré inférieur ou égal à p , à condition de choisir k tel que $\mu k \geq p$. N'explicitons pas plus le calcul général et prenons un exemple.

Exemple. Déterminons un D.L. à l'ordre 4 en 0 de $\frac{1}{\cos x}$. On a, par les DLs usuels,

$$\frac{1}{\cos x} = \frac{1}{1 - u}$$

avec

$$u = \frac{1}{2}x^2 - \frac{1}{24}x^4 + o(x^4)$$

Remarquons que u est « de l'ordre de » x^2 . Donc $o(u^2) = o(x^4)$. De là,

$$\begin{aligned} \frac{1}{\cos x} &= 1 + u + u^2 + o(u^2) \\ &= 1 + u + u^2 + o(x^4) \\ &= 1 + \frac{1}{2}x^2 + \frac{5}{24}x^4 + o(x^4) \end{aligned}$$

Déduisons-en un D.L. à l'ordre 5 en 0 de $\tan x$:

$$\begin{aligned}\tan x &= \sin x \frac{1}{\cos x} \\ &= \left(x - \frac{1}{6}x^3 + \frac{1}{120}x^5 + o(x^5)\right)\left(1 + \frac{1}{2}x^2 + \frac{5}{24}x^4 + o(x^4)\right) \\ &= x + \frac{1}{3}x^3 + \frac{2}{15}x^5 + o(x^5)\end{aligned}$$

3.4 Composition de DLs

Nous traitons juste deux exemples.

Exemple. Soit $f(x) = \ln \cos x$. Faisons un DL de f en 0 à l'ordre 4. On a

$$f(x) = \ln(1 + u)$$

où

$$u = -\frac{1}{2}x^2 + \frac{1}{24}x^4 + o(x^4)$$

tend vers 0 lorsque x tend vers 0. Donc,

$$f(x) = u - \frac{1}{2}u^2 + o(u^2)$$

et on s'arrête à l'ordre 2 parce que $o(u^2) = o(x^4)$. On a donc à additionner et multiplier des DLs, ce que l'on sait déjà faire. Le lecteur vérifiera que

$$f(x) = \frac{1}{2}x^2 - \frac{1}{12}x^4 + o(x^4)$$

Exemple. Faisons un DL à l'ordre 3 en 0 de $f(x) = e^{e^x}$. On a

$$f(x) = \exp\left(1 + x + \frac{1}{2}x^2 + \frac{1}{6}x^3 + o(x^3)\right) = e \exp u$$

où

$$u = x + \frac{1}{2}x^2 + \frac{1}{6}x^3 + o(x^3)$$

tend vers 0 lorsque x tend vers 0. Donc,

$$f(x) = e\left(1 + u + \frac{1}{2}u^2 + \frac{1}{6}u^3 + o(u^3)\right)$$

Le lecteur est prié ici de calculer des DLs de u^2 et u^3 , puis de vérifier que, en fin de compte,

$$f(x) = e\left(1 + x + x^2 + \frac{5}{6}x^3 + o(x^3)\right)$$

3.5 Primitivation de DLs

Proposition 15. Soit f une fonction dérivable au voisinage du point a . On suppose que f' admet en a un D.L. à l'ordre n :

$$f'(x) = \sum_{k=0}^n a_k (x-a)^k + o(x-a)^n$$

Alors, f admet en a un D.L. à l'ordre $n+1$ qui est

$$f(x) = f(a) + \sum_{k=0}^n \frac{a_k}{k+1} (x-a)^{k+1} + o(x-a)^{n+1}$$

Remarque 14. Ce théorème est bien un théorème d'intégration de DL : si g a un DL à l'ordre n , alors toute primitive de g aura un DL à l'ordre $n+1$ en ce même point, obtenu en intégrant « formellement » le DL de g . En revanche, si g a un DL à l'ordre n en a , si g est dérivable au voisinage de a , on ne peut absolument pas affirmer que g' aura un DL en a .

Démonstration. Posons

$$\varphi(x) = f(x) - f(a) - \sum_{k=0}^n \frac{a_k}{k+1} (x-a)^{k+1}$$

La fonction φ est dérivable au voisinage de a , et on a

$$\varphi'(x) = o(x-a)^n = (x-a)^n \varepsilon(x)$$

où $\varepsilon(x)$ tend vers 0 lorsque x tend vers a . Utilisons le théorème des accroissements finis :

$$\varphi(x) = \varphi(x) - \varphi(a) = (x-a)\varphi'(c)$$

où c est situé entre a et x . On a donc

$$\varphi(x) = (x-a)(c-a)^n \varepsilon(c) = (x-a)^{n+1} \varepsilon(c) = o(x-a)^{n+1}$$

car comme c est situé entre a et x , $\varepsilon(c)$ tend vers 0 lorsque x tend vers a (théorème de composition des limites). $\square$

Exemple. Soit $f : x \mapsto \arctan x$. La fonction f est de classe $\mathcal{C}^\infty$ sur $\mathbb{R}$, elle admet donc des développements limités de tous ordres en 0. On a

$$f'(x) = \frac{1}{1+x^2} = \sum_{k=0}^n (-1)^k x^{2k} + o(x^{2n+1})$$

Le théorème d'intégration des DL nous dit que $\arctan x$ a un DL à l'ordre $2n + 2$ en 0, obtenu en intégrant formellement le DL ci-dessus. Ainsi, comme $\arctan 0 = 0$,

$$\arctan x = \sum_{k=0}^n (-1)^k \frac{x^{2k+1}}{2k+1} + o(x^{2n+2})$$

Exercice. Déterminer un DL à l'ordre $2n + 2$ en 0 de $\operatorname{argth} x$.

Exemple. Déterminons un DL à l'ordre 7 en 0 de $\arcsin x$.

La fonction $\arcsin$ est dérivable au voisinage de 0 et

$$\arcsin' x = \frac{1}{\sqrt{1-x^2}} = (1-x^2)^{-\frac{1}{2}}$$

On reconnaît un DL usuel :

$$\arcsin' x = 1 - \frac{1}{2}(-x^2) + \frac{3}{8}(-x^2)^2 - \frac{5}{16}(-x^2)^3 + o(x^6)$$

où l'on a utilisé $(-\frac{1}{2}) = \frac{3}{8}$ et $(-\frac{1}{2})^2 = -\frac{5}{16}$. Comme $\arcsin 0 = 0$, on en déduit, par le théorème d'intégration des DLs,

$$\arcsin x = x + \frac{1}{6}x^3 + \frac{3}{40}x^5 + \frac{5}{112}x^7 + o(x^7)$$

Exemple. Soit $f(x) = \arccos(1-x^2)$, définie sur $[-1, 1]$. La fonction f est continue sur $[-1, 1]$ et dérivable sur $[-1, 1]$, sauf peut-être en 0. Vérifions que f est en fait dérivable en 0. Pour tout $x \in [-1, 1] \setminus \{0\}$, on a

$$f'(x) = \frac{2}{\sqrt{2-x^2}}$$

Donc, $f'(x)$ tend vers $\sqrt{2}$ lorsque x tend vers 0. On en déduit par un corollaire du théorème des accroissements finis que, comme f' a une limite finie en 0, f est dérivable en 0 et $f'(0) = \sqrt{2}$, la limite de f' en 0.

Un DL en 0 de f' à l'ordre 4 est (DLs usuels)

$$f'(x) = \sqrt{2} + \frac{x^2}{2\sqrt{2}} + \frac{3x^4}{16\sqrt{2}} + o(x^4)$$

On en déduit grâce au théorème d'intégration des DL, que

$$f(x) = \sqrt{2}x + \frac{x^3}{6\sqrt{2}} + \frac{3x^5}{80\sqrt{2}} + o(x^5)$$

Posant $x = \sqrt{1-t}$, on en déduit ce que l'on appelle un développement asymptotique de la fonction $\arccos$ en 1 :

$$\arccos t = \sqrt{2}(1-t)^{\frac{1}{2}} + \frac{(1-t)^{\frac{3}{2}}}{6\sqrt{2}} + \frac{3(1-t)^{\frac{5}{2}}}{80\sqrt{2}} + o((1-t)^{\frac{5}{2}})$$

4 Développement asymptotiques

4.1 Notion de DA

Si l'on reprend l'exemple de la fonction arccos du paragraphe précédent, on voit que l'on a effectué au voisinage de 1 quelque chose qui ressemble à un développement limité, sauf que ce n'en est pas un. D'ailleurs, aucun développement limité de arccos au voisinage de 1 n'existe (sauf un DL à l'ordre 0), puisque cette fonction n'est pas dérivable en 1. Si l'on regarde attentivement ce que l'on a fait, on a écrit $\arccos x$ sous forme d'une somme de fonctions, chaque fonction étant négligeable en 1 devant la précédente.

De façon générale, soit f une fonction définie au voisinage de $a \in \overline{\mathbb{R}}$. On effectue un développement asymptotique de f en a lorsqu'on écrit

$$f(x) = \varphi_0(x) + \varphi_1(x) + \dots + \varphi_n(x) + o_a(\varphi_n(x))$$

où pour tout $k \in \llbracket 1, n \rrbracket$, $\varphi_k(x) = o_a(\varphi_{k-1}(x))$.

On voit qu'un développement limité de f en $a \in \mathbb{R}$ est un cas particulier de DA, obtenu en prenant $\varphi_k(x) = (x - a)^k$.

4.2 Quelques exemples

Exemple. On prend $f(x) = x^x$ et $a = 0$. On a

$$f(x) = \exp(x \ln x) = 1 + x \ln x + \frac{1}{2}x^2 \ln^2 x + o(x^2 \ln^2 x)$$

Exemple. On prend $f(x) = \sqrt{\sinh x}$ et $a = +\infty$. On a

$$f(x) = \sqrt{\frac{1}{2}(e^x - e^{-x})}$$

On essaie évidemment de se ramener à ce que l'on sait faire, c'est à dire à un développement limité. Pour cela, on factorise e^x sous la racine :

$$f(x) = \frac{1}{\sqrt{2}} e^{\frac{x}{2}} \sqrt{1 - e^{-2x}}$$

Comme e^{-2x} tend vers 0 lorsque x tend vers l'infini, on peut faire un « DL » de la racine carrée :

$$\begin{aligned} f(x) &= \frac{1}{\sqrt{2}} e^{\frac{x}{2}} \left(1 - \frac{1}{2}e^{-2x} - \frac{1}{8}e^{-4x} + o(e^{-4x}) \right) \\ &= \frac{1}{\sqrt{2}} \left(e^{\frac{x}{2}} - \frac{1}{2}e^{-\frac{3}{2}x} - \frac{1}{8}e^{-\frac{7}{2}x} + o(e^{-\frac{7}{2}x}) \right) \end{aligned}$$

Exemple. On prend $f(x) = \arctan x$ et $a = +\infty$. On a pour $x > 0$

$$f(x) = \frac{\pi}{2} - \arctan \frac{1}{x} = \frac{\pi}{2} - \frac{1}{x} + \frac{1}{3x^3} + o\left(\frac{1}{x^3}\right)$$

Exemple. On prend $f(x) = \cotan x$ et $a = 0$. On a

$$f(x) = \frac{\cos x}{\sin x} = \frac{1}{x} \frac{\cos x}{g(x)}$$

où

$$g(x) = \frac{\sin x}{x} = 1 - \frac{1}{6}x^2 + \frac{1}{120}x^4 + o(x^4)$$

On fait ensuite un DL (à l'ordre 5) de $\frac{1}{g(x)}$ que l'on multiplie par un DL de $\cos x$:

$$\frac{\cos x}{g(x)} = 1 - \frac{1}{3}x^2 - \frac{1}{45}x^4 + o(x^5)$$

d'où

$$\cotan x = \frac{1}{x} - \frac{1}{3}x - \frac{1}{45}x^3 + o(x^4)$$

Le lecteur est invité à compléter le détail des calculs.

4.3 La formule de Stirling

Il existe des équivalents pas évidents du tout à démontrer. C'est le cas de la formule de Stirling.

Le lecteur angoissé par la démonstration de 4 pages qui va suivre est autorisé à ne pas lire les détails. Cela dit, on y raconte des choses très instructives.

Proposition 16. On a $n! \sim \left(\frac{n}{e}\right)^n \sqrt{2\pi n}$.

La démonstration de la formule de Stirling est longue et nécessite un certain nombre de lemmes.

Posons pour tout $n \geq 1$,

$$u_n = \sum_{k=1}^n \ln k - n \ln n + n - \frac{1}{2} \ln n$$

Lemme 17. On a $u_{n+1} - u_n \sim -\frac{1}{12n^2}$.

Démonstration.

$$\begin{aligned} u_{n+1} - u_n &= \ln(n+1) - (n+1)\ln(n+1) + (n+1) - \frac{1}{2}\ln(n+1) + n\ln n - n + \frac{1}{2}\ln n \\ &= 1 - (n + \frac{1}{2})\ln(1 + \frac{1}{n}) \end{aligned}$$

Comme $\frac{1}{n}$ tend vers 0, on peut utiliser un DL de $\ln(1+x)$ en 0 :

$$\begin{aligned} u_{n+1} - u_n &= 1 - (n + \frac{1}{2})(\frac{1}{n} - \frac{1}{2}\frac{1}{n^2} + \frac{1}{3}\frac{1}{n^3} + o(\frac{1}{n^3})) \\ &= -\frac{1}{12n^2} + o(\frac{1}{n^2}) \\ &\sim -\frac{1}{12n^2} \end{aligned}$$

□

On voit donc que la différence $u_{n+1} - u_n$ entre deux termes successifs de la suite (u_n) tend vers 0. En fait, cette différence tend assez vite vers 0 pour que la suite (u_n) converge.

Lemme 18. *La suite (u_n) est convergente.*

Démonstration. Pour tout n assez grand, on a

$$\frac{1}{2} \leq \frac{u_{n+1} - u_n}{-\frac{1}{12n^2}} \leq \frac{3}{2}$$

ou encore

$$-\frac{1}{24n^2} \geq u_{n+1} - u_n \geq -\frac{1}{8n^2}$$

On déduit de la première inégalité que la suite (u_n) décroît à partir d'un certain rang N . Il reste à prouver qu'elle est minorée. C'est l'autre inégalité qui permettra de conclure. Soit $n \geq N$. On a

$$u_n = \sum_{k=N}^{n-1} (u_{k+1} - u_k) + u_N \geq -\frac{1}{8} \sum_{k=N}^{n-1} \frac{1}{k^2} + u_N$$

Il nous faut un lemme pour montrer le lemme ... □

Lemme 19. *Soit $v_n = \sum_{k=1}^n \frac{1}{k^2}$. La suite (v_n) est majorée.*

Démonstration. La suite (v_n) est croissante. Pour montrer qu'elle est majorée, il suffit de montrer que la suite extraite (v_{2^n}) est majorée. Or,

$$v_{2^{n+1}} - v_{2^n} = \sum_{k=2^n}^{2^{n+1}-1} \frac{1}{k^2} \leq \frac{1}{2^{2n}} \times 2^n$$

où 2^n est le nombre de termes de la somme et $\frac{1}{2^{2n}}$ est le plus grand terme, obtenu pour $k = 2^n$. De là,

$$v_{2^{n+1}} - v_{2^n} \leq \frac{1}{2^n}$$

Pour terminer,

$$\begin{aligned} v_{2^n} &= (v_2 - v_1) + (v_4 - v_2) + (v_8 - v_4) + \dots + (v_{2^n} - v_{2^{n-1}}) + v_1 \\ &\leq \frac{1}{2^0} + \frac{1}{2^1} + \dots + \frac{1}{2^{n-1}} + 1 \\ &= 3 - \frac{1}{2^{n-1}} \leq 3 \end{aligned}$$

□

Démonstration. Retour à u_n ... On a pour tout $n \geq N$, $u_n \geq u_N - \frac{3}{8}$. La suite (u_n) est décroissante à partir d'un certain rang, minorée, donc converge vers un réel ℓ . □

Lemme 20. *Il existe un réel $C > 0$ tel que*

$$n! \sim \left(\frac{n}{e}\right)^n C \sqrt{n}$$

Démonstration. Il s'avère que

$$u_n = \ln \frac{n!}{\left(\frac{n}{e}\right)^n \sqrt{n}}$$

Cette quantité tend vers ℓ , donc $\exp(u_n)$ tend vers $C = e^\ell$, d'où l'équivalent cherché. □

Il reste maintenant à calculer C , et même cela n'est pas simple. À cet effet, introduisons les intégrales de Wallis. Pour tout $n \in \mathbb{N}$, soit

$$W_n = \int_0^{\frac{\pi}{2}} \sin^n t \, dt$$

Lemme 21. *La suite (W_n) est positive, décroissante.*

Démonstration. La positivité est évidente : on intègre une fonction positive. Pour tout entier n , on a

$$W_{n+1} - W_n = \int_0^{\frac{\pi}{2}} \sin^n t (\sin t - 1) \, dt$$

Cette fois on intègre une fonction négative, d'où la décroissance. □

Lemme 22. On a pour tout entier $n \geq 2$,

$$W_n = \frac{n-1}{n} W_{n-2}$$

Démonstration. On intègre par parties W_n , en posant $u' \sin t$ et $v = \sin^{n-1} t$. Le calcul est laissé au lecteur (s'il reste un lecteur). $\square$

Lemme 23. On a pour tout entier n ,

$$W_{2n} = \frac{(2n)!}{2^{2n} n!^2} \frac{\pi}{2}$$

et

$$W_{2n+1} = \frac{2^{2n} n!^2}{(2n+1)!}$$

Démonstration. Récurrence sur n , utiliser le lemme précédent. $\square$

Lemme 24. $nW_n W_{n+1} \longrightarrow \frac{\pi}{2}$ lorsque n tend vers l'infini.

Démonstration. Posons $u_n = nW_n W_{n+1}$. On a

$$u_{2n} = \frac{2n}{2n+1} \frac{\pi}{2}$$

qui tend bien vers $\frac{\pi}{2}$ lorsque n tend vers l'infini. Et

$$u_{2n+1} = (2n+1)W_{2n+1}W_{2n+2} = \frac{2n+1}{2n+2} \frac{\pi}{2}$$

qui tend aussi vers $\frac{\pi}{2}$. On conclut avec la « réciproque » du théorème sur les suites extraites. $\square$

Lemme 25. $W_n \sim W_{n+1}$.

Démonstration. On a

$$W_{n+2} = \frac{n+1}{n+2} W_n \leq W_{n+1} \leq W_n$$

d'où

$$\frac{n+1}{n+2} \leq \frac{W_{n+1}}{W_n} \leq 1$$

On conclut avec le théorème d'encadrement des limites. $\square$

Lemme 26. $W_n \sim \sqrt{\frac{\pi}{2n}}$.

Démonstration. En, effet, $nW_nW_{n+1} \sim nW_n^2 \rightarrow \frac{\pi}{2}$. $\square$

On y est !

Lemme 27. $C = \sqrt{2\pi}$

Démonstration. On a

$$W_{2n} = \frac{(2n)!}{2^{2n}n!^2} \frac{\pi}{2} \sim \frac{\left(\frac{2n}{e}\right)^{2n} C\sqrt{2n}}{2^{2n} \left(\left(\frac{n}{e}\right)^n C\sqrt{n}\right)^2} \frac{\pi}{2} = \frac{1}{C} \frac{\pi}{\sqrt{2n}}$$

Par ailleurs

$$W_{2n} \sim \sqrt{\frac{\pi}{4n}}$$

D'où

$$C \sim \sqrt{2\pi}$$

Mais C est constant. Donc

$$C = \sqrt{2\pi}$$

$\square$

Moralité : il n'est pas toujours facile de trouver un équivalent...

Chapitre 21

Étude des Fonctions

Voici un plan d'étude des fonctions d'une variable réelle. Certaines parties de cette étude peuvent être abordées dans un ordre différent, ou être parfois omises. Tout ce qui suit n'est bien entendu valable que pour des fonctions raisonnablement régulières : on exclut les cas **pathologiques** et les courbes **monstrueuses**.

1 Ensemble de définition

L'étude de la fonction f commence forcément par la recherche de son ensemble de définition $\mathcal{D}_f$. Cette recherche est souvent, mais pas toujours, immédiate.

2 Réduction de l'ensemble d'étude

La fonction est-elle paire ? Impaire ? Périodique ? L'expression de f présente-t-elle des **symétries** ? À partir des symétries *algébriques* on peut souvent déduire des symétries *géométriques* de la courbe représentative, et donc une réduction de l'ensemble d'étude. Plus on réduit cet ensemble, et moins on a de travail par la suite.

3 Régularité de la fonction

Les **théorèmes du cours** permettent d'affirmer que la fonction est continue, dérivable, $\mathcal{C}^1, \dots, \mathcal{C}^\infty$ sur certains intervalles.

Tout ce qui précède a permis de mettre en évidence un certain nombre de points que nous appellerons **remarquables**. Ce sont

- $+\infty$ ou $-\infty$ si l'ensemble d'étude de f est non borné.
 - Les bornes des intervalles trouvés lors de la recherche de $\mathcal{D}_f$, ou lors de l'étude de la régularité de f . En fait, tout point ayant été nommé est remarquable. En général, en un tel point, la fonction n'est pas définie, ou alors elle est définie mais les théorèmes généraux du cours ne permettent pas d'affirmer si elle y est continue, dérivable, ... D'autres points remarquables risquent d'apparaître dans les paragraphes qui suivent.
-

4 Variations

Le **signe** de f' permet de déterminer des intervalles sur lesquels f est monotone (on rappelle que f est supposée être raisonnablement gentille). Les points **d'annulation** de f' sont également des points remarquables (tangente horizontale).

On résume ce qui précède dans le **tableau de variations** de f , qui doit faire apparaître :

- Le signe de f' .
- Les points remarquables.
- Les points que nous trouvons rigolos ou attrayants (zéros de f , de f'' , ...).

Ce tableau est **incomplet** : il y manque les limites éventuelles de f en certains points remarquables.

5 Points remarquables réels

5.1 Point réel où la fonction n'est pas définie

La fonction a-t-elle une limite infinie en ce point (éventuellement à droite ou à gauche) ? On a alors une **asymptote verticale** .

A-t-on une limite finie ? On prolonge alors par continuité, et on passe au numéro suivant.

Pas de limite ? On abandonne tout espoir d'étude et on dessine comme on peut (cf $\sin \frac{1}{x}$).

5.2 Points où la fonction est définie

On recherche (en revenant éventuellement à la définition) si la fonction est continue au point, si elle y est dérivable, si on a une tangente verticale à la courbe ... Lorsqu'il y a une tangente oblique on recherche la position de la courbe par rapport à cette tangente au voisinage du point considéré.

6 Étude à l'infini

on suppose ici à titre d'exemple que la fonction est définie au voisinage de $+\infty$. C'est pareil en $-\infty$.

6.1 Limite à l'infini

Limite finie

Si f a une limite finie à l'infini, on a alors une **asymptote horizontale** . Le tableau de variations permet de préciser la position de la courbe par rapport à cette asymptote.

Limite infinie

On passe au paragraphe intitulé « Direction asymptotique ».

Pas de limite

On abandonne l'étude et trace comme on le sent (exemple : $\sin(x^2)$).

6.2 Direction asymptotique

On suppose ici que $f(x)$ tend vers l'infini (plus ou moins) lorsque $x \rightarrow +\infty$. On considère $\frac{f(x)}{x}$ et on regarde ce que fait cette quantité lorsque $x \rightarrow \infty$. A-t-on

- $f(x) \sim ax$ avec $a \neq 0$? Il y a **direction asymptotique** dans la direction de la droite $Y = aX$. On passe au paragraphe suivant.
- $f(x) = o(x)$? Il y a **direction asymptotique** dans la direction de la droite Ox . On parle également de **branche parabolique** dans la direction de l'axe Ox .
- $x = o(f(x))$? Il y a **direction asymptotique** dans la direction de la droite Oy . On parle également de **branche parabolique** dans la direction de l'axe Oy .
- Rien de tout cela? On abandonne l'étude à l'infini et on trace au feeling.

En résumé, c'est l'étude d'une limite éventuelle de $\frac{f(x)}{x}$ à l'infini qui permet de savoir si la courbe admet une direction asymptotique.

6.3 Asymptote, branche parabolique

On suppose ici que $f(x) \rightarrow \pm\infty$ lorsque $x \rightarrow +\infty$, et que $f(x) \sim ax$ où $a \neq 0$. On a donc $f(x) = ax + o(x)$ et il est alors **naturel** de considérer la quantité $f(x) - ax$. Cette quantité admet-elle en $+\infty$

- Une limite finie b ? On a alors $f(x) = ax + b + o(1)$ et on dit que la droite $Y = aX + b$ est **asymptote** à la courbe au voisinage de $+\infty$. On passe ensuite au paragraphe suivant.
- Une limite infinie? On a une branche parabolique dans la direction de la droite $Y = aX$.
- Rien du tout? On se contente de la direction asymptotique, et on arrête l'étude.

6.4 Position courbe/asymptote

On suppose ici que la courbe admet l'asymptote $Y = aX + b$ ($a \neq 0$) au voisinage de $+\infty$. Il faut alors étudier la position de la courbe par rapport à cette asymptote, en considérant la quantité $f(x) - ax - b$. Le **signe** de cette quantité détermine la position. On peut chercher ce signe « algébriquement » lorsque les expressions mises en jeu sont simples, ou sinon au voisinage de l'infini par des **calculs d'équivalents** ou des **développements asymptotiques**.

7 Concavité

Cette étude est souvent omise. L'étude du signe de f'' permet de préciser les **inflexions** de la courbe, et les intervalles où la fonction est **convexe** ou **concave**.

Lorsque $f'' > 0$ sur un intervalle, cela signifie que f' croît sur cet intervalle. Les pentes augmentent, la courbe « tourne » vers la gauche. On dit dans ce cas que la fonction f est **convexe**.

Si, au contraire, $f'' < 0$, la fonction est **concave**.

Lorsque f'' s'annule en changeant de signe, la fonction passe de convexe à concave (ou le contraire). On tournait vers la gauche et on se met à tourner vers la droite. La courbe présente une **inflexion**.

8 Tracé

C'est le but ultime de l'étude, et il ne doit pas être négligé. Tout ce qui a été étudié doit se voir sur le dessin : les points remarquables et leurs éventuelles tangentes, les asymptotes éventuelles, ...et des axes avec une indication d'échelle. En revanche, en mathématiques un tableau de valeurs est totalement inutile en général. Ce qui est important est la forme de la courbe, beaucoup plus que son exactitude.

9 Un exemple complet

Nous allons étudier la fonction $f : x \mapsto \sqrt[3]{x^3 - x^2}$.

9.1 Ensemble de définition, symétries, régularité

La fonction f est continue sur $\mathbb{R}$. Elle ne présente aucune symétrie apparente. Elle est de classe $\mathcal{C}^\infty$ sur $] -\infty, 0[$, $]0, 1[$ et $]1, +\infty[$. Nous voyons donc se dégager 4 points remarquables : $-\infty, 0, 1$ et $+\infty$.

9.2 Variations

Pour tout $x \neq 0, 1$, on a

$$f'(x) = \frac{x(3x - 2)}{3\sqrt[3]{(x^3 - x^2)^2}}$$

qui est du signe de $x(3x - 2)$. La fonction f est donc

- Strictement croissante sur $] -\infty, 0]$.
- Strictement décroissante sur $[0, \frac{2}{3}]$.
- Strictement croissante sur $[\frac{2}{3}, +\infty[$.

9.3 Étude en 0

La fonction f est-elle dérivable en 0 ? Considérons des taux d'accroissement. Pour tout x non nul voisin de 0, on a

$$\frac{f(x) - f(0)}{x - 0} = \frac{f(x)}{x} = \sqrt[3]{1 - \frac{1}{x}}$$

Lorsque x tend vers 0, $x > 0$, cette quantité tend vers $-\infty$. À gauche de 0, cette même quantité tend vers $+\infty$. La courbe de f a donc une tangente verticale en 0. Cette tangente verticale s'accompagne d'un changement de sens de variation : nous avons ce que l'on appelle un **point de rebroussement**.

9.4 Étude en 1

La fonction f est-elle dérivable en 1 ? Considérons des taux d'accroissement. Pour tout x différent de 1 voisin de 1, on a

$$\frac{f(x) - f(1)}{x - 1} = \frac{\sqrt[3]{x^2(x-1)}}{x - 1} = \frac{\sqrt[3]{x^2}}{\sqrt[3]{(x-1)^2}}$$

Lorsque x tend vers 1, cette quantité tend vers $+\infty$. La courbe de f a donc une tangente verticale en 1.

9.5 Étude en $\pm\infty$

On a pour tout $x \neq 0$

$$\frac{f(x)}{x} = \sqrt[3]{1 - \frac{1}{x}}$$

Cette quantité tend vers 1 lorsque x tend vers $\pm\infty$. On a donc à l'infini une direction asymptotique, qui est celle de la droite d'équation $Y = X$.

Considérons donc $f(x) - x$:

$$f(x) - x = \sqrt[3]{x^3 - x^2} - x = x \left(\sqrt[3]{1 - \frac{1}{x}} - 1 \right)$$

Comme $\frac{1}{x}$ tend vers 0 en $\pm\infty$, les équivalents usuels nous donnent

$$\sqrt[3]{1 - \frac{1}{x}} - 1 \sim_{\pm\infty} -\frac{1}{3x}$$

Ainsi,

$$f(x) - x \sim -\frac{1}{3}$$

Comme deux fonctions équivalentes ont même limite, il en résulte que $f(x) - x \rightarrow -\frac{1}{3}$ lorsque x tend vers $\pm\infty$. La courbe admet donc à l'infini une asymptote qui est la droite d'équation

$$Y = X - \frac{1}{3}$$

Il s'agit maintenant de déterminer la position de la courbe par rapport à l'asymptote. Pour cela il faut déterminer le signe de $f(x) - x + \frac{1}{3}$ au voisinage de l'infini. Il s'agit pour cela

de peaufiner notre calcul d'équivalent, en faisant un développement asymptotique. On a, à l'infini,

$$\begin{aligned}\sqrt[3]{1 - \frac{1}{x}} &= \left(1 - \frac{1}{x}\right)^{\frac{1}{3}} \\ &= 1 - \frac{1}{3} \frac{1}{x} + \left(\frac{\frac{1}{3}}{2}\right) \frac{1}{x^2} + o\left(\frac{1}{x^2}\right) \\ &= 1 - \frac{1}{3} \frac{1}{x} - \frac{1}{9} \frac{1}{x^2} + o\left(\frac{1}{x^2}\right)\end{aligned}$$

Multipliant le tout par x , il vient

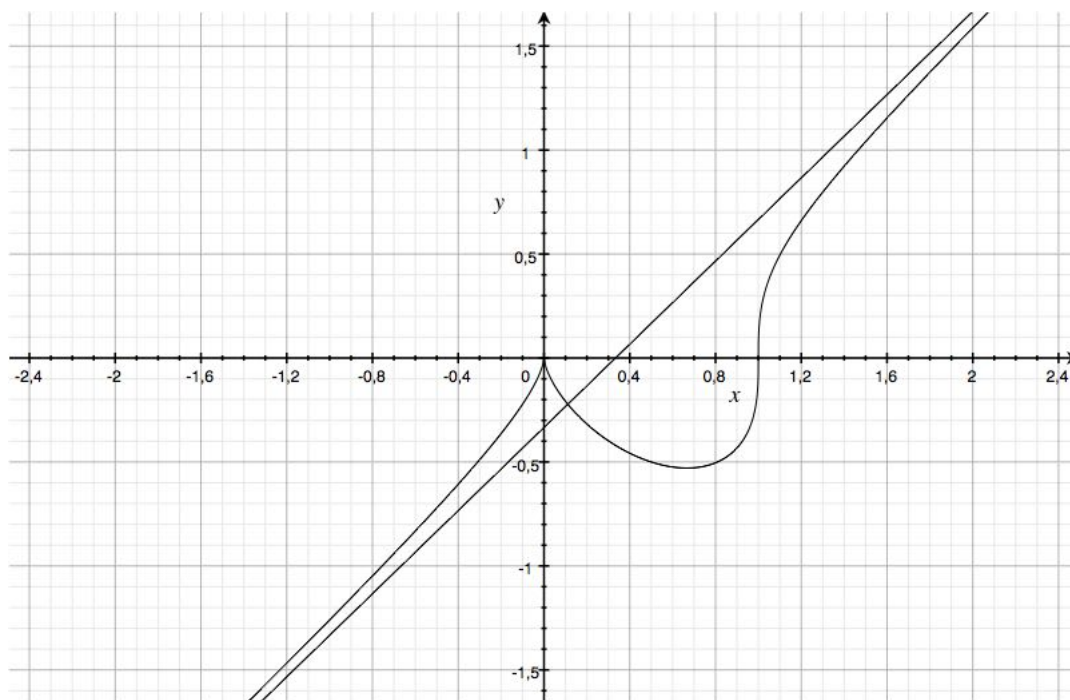
$$f(x) = x - \frac{1}{3} - \frac{1}{9} \frac{1}{x} + o\left(\frac{1}{x}\right)$$

Ainsi,

$$f(x) - x + \frac{1}{3} \sim -\frac{1}{9} \frac{1}{x}$$

Au voisinage de $-\infty$, la courbe est sur l'asymptote. Au voisinage de $+\infty$, la courbe est sous l'asymptote.

9.6 Courbe



9.7 Concavité

On trouve avec un peu de courage que pour $x \neq 0, 1$,

$$f''(x) = -\frac{2}{9(x^2(x-1))^{\frac{2}{3}}(x-1)}$$

Cette quantité est positive si $x < 1$ et négative si $x > 1$. Ainsi,

La courbe tourne vers la gauche sur $] -\infty, 1[$.

La courbe tourne vers la droite sur $]1, +\infty[$.

On a une inflexion en 1.

10 Un autre exemple

Nous allons étudier la fonction $f : x \mapsto \sin x^{\tan x}$.

10.1 Ensemble de définition, symétries, régularité

On a très précisément $f(x) = e^{\tan x \ln \sin x}$. La fonction f est donc définie en tout réel x tel que $\tan x$ est défini et $\sin x > 0$. Ainsi,

$$\mathcal{D}_f = \left(\bigcup_{k \in \mathbb{Z}}]2k\pi, (2k+1)\pi[\right) \setminus \left\{ 2k\pi + \frac{\pi}{2}, k \in \mathbb{Z} \right\}$$

Cela fait très peur, mais f est de période 2π . Nous allons donc l'étudier sur

$$\mathcal{D} =]0, \pi[\setminus \left\{ \frac{\pi}{2} \right\}$$

La fonction f est de classe $\mathcal{C}^\infty$ sur $]0, \frac{\pi}{2}[$ et $] \frac{\pi}{2}, \pi[$.

Nous voyons se dégager au final 3 points remarquables : $0, \frac{\pi}{2}$ et π .

10.2 Variations

Pour tout $x \in \mathcal{D}$,

$$f'(x) = f(x) \left(\frac{1}{\cos^2 x} \ln \sin x + \tan x \frac{\cos x}{\sin x} \right) = \frac{f(x)}{\cos^2 x} (\ln \sin x + \cos^2 x)$$

Le signe de f' n'est pas évident. Posons $\varphi(x) = \ln \sin x + \cos^2 x$. La fonction φ est dérivable sur $]0, \pi[$ et, pour tout $x \in]0, \pi[$,

$$\varphi'(x) = \frac{\cos x}{\sin x} - 2 \sin x \cos x = \frac{\cos x(1 - 2 \sin^2 x)}{\sin x}$$

On en déduit les variations de φ .

x	0	$\frac{\pi}{4}$	$\frac{\pi}{2}$	$\frac{3\pi}{4}$	π
$\varphi'(x)$	+ 0 - 0 + 0 -				
$\varphi(x)$	$-\infty$ > 0 0 > 0 $-\infty$				

Par le TVI et la monotonie stricte, φ s'annule en un réel $\alpha \in]0, \frac{\pi}{4}[$, en $\frac{\pi}{2}$ et en $\beta = \pi - \alpha$. Comme le signe de f' est celui de φ , on en déduit les variations de f .

x	0	α	$\frac{\pi}{2}$	β	π	
$f'(x)$		−	0	+	0	−
$f(x)$		?		?		?

Avis aux lecteurs : il faudrait une double barre à la verticale de $\frac{\pi}{2}$ mais, assez paresseusement, je n'ai pas envie de chercher comment faire.

10.3 Étude en 0

Pour tout $x > 0$ voisin de 0, on a

$$\tan x \ln \sin x \sim x \ln x$$

Or, $x \ln x$ tend vers 0 lorsque x tend vers 0. On en déduit que $\tan x \ln \sin x$ tend aussi vers 0 et, donc, $f(x)$ tend vers $e^0 = 1$ lorsque x tend vers 0.

On peut donc prolonger f par continuité en 0 en posant $f(0) = 1$. Le prolongement est-il dérivable ? On a

$$\frac{f(x) - f(0)}{x - 0} = \frac{1}{x} \left(e^{\tan x \ln \sin x} - 1 \right) \sim \frac{1}{x} \tan x \ln \sin x \sim \frac{1}{x} x \ln x = \ln x$$

Les taux d'accroissement tendent donc vers $-\infty$ lorsque x tend vers 0 : la courbe de f possède au point $(0, 1)$ une tangente verticale.

10.4 Étude en π

L'étude en π est tout à fait analogue. Il suffit de poser $x = \pi - h$ où h tend vers 0 lorsque x tend vers π . On obtient que $f(x)$ tend vers 1 lorsque x tend vers π et que la courbe a une tangente verticale au point $(\pi, 1)$.

10.5 Étude en $\frac{\pi}{2}$

Posons $x = \frac{\pi}{2} + h$, de sorte que h tend vers 0 lorsque x tend vers $\frac{\pi}{2}$. On a

$$f(x) = \exp(\tan x \ln \sin x) = \exp\left(-\frac{1}{\tan h} \ln \cos h\right)$$

Nous allons faire un DL à l'ordre 1 de l'expression entre parenthèses. Tout d'abord

$$\ln \cos h = \ln\left(1 - \frac{1}{2}h^2 + o(h^2)\right) = -\frac{1}{2}h^2 + o(h^2)$$

Ensuite,

$$\tan h = h + o(h^2) = h(1 + o(h))$$

et donc

$$\frac{1}{\tan h} = \frac{1}{h} \frac{1}{1 + o(h)} = \frac{1}{h}(1 + o(h))$$

Regroupons le tout. Il vient

$$\begin{aligned} -\frac{1}{\tan h} \ln \cos h &= \frac{1}{h}(1 + o(h)) \left(\frac{1}{2}h^2 + o(h^2)\right) \\ &= (1 + o(h)) \left(\frac{1}{2}h + o(h)\right) \\ &= \frac{1}{2}h + o(h) \end{aligned}$$

Pour finir,

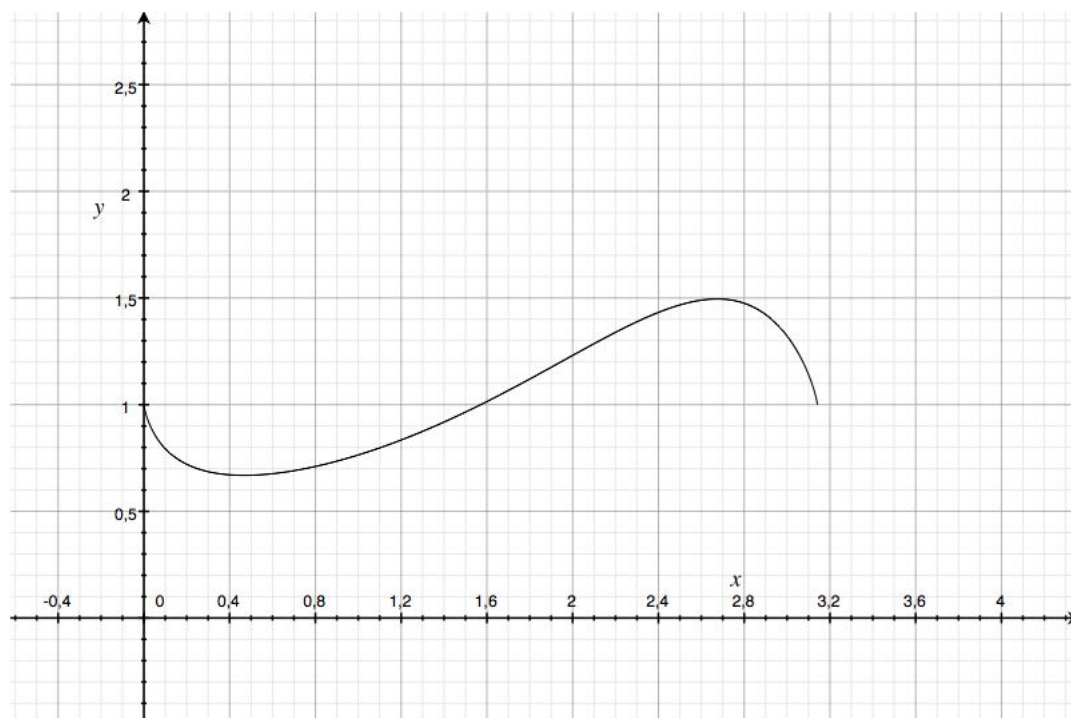
$$f(x) = \exp\left(\frac{1}{2}h + o(h)\right) = 1 + \frac{1}{2}h + o(h) = 1 + \frac{1}{2}\left(x - \frac{\pi}{2}\right) + o\left(x - \frac{\pi}{2}\right)$$

On en déduit que $f(x)$ tend vers 1 lorsque x tend vers $\frac{\pi}{2}$, puis que le prolongement par continuité de f est dérivable en $\frac{\pi}{2}$, de dérivée $\frac{1}{2}$.

Un développement limité à l'ordre 2 aurait permis de connaître la position de la courbe par rapport à sa tangente au voisinage de $\frac{\pi}{2}$. Nous ne le ferons pas.

10.6 Courbe

Ne connaissant pas précisément α , et donc $\beta = \pi - \alpha$, on trace la courbe un peu comme on le peut. Ce qui est important, c'est de faire clairement apparaître les variations de f et l'allure de f au voisinage des points remarquables. La courbe ci-dessous est un exemple de ce qu'il ne faut pas faire. Obtenue automatiquement, la courbe ne présente pas de tangente clairement verticale pour $x = 0$ et $x = \pi$. Un tracé à la main serait préférable.



10.7 Concavité

Nous avons déjà eu du mal à étudier les variations de f . L'étude de la concavité serait sans doute truffée de calculs inextricables, nous nous en dispenserons.

Chapitre 22

Espaces vectoriels

1 Généralités

1.1 Notion d'espace vectoriel

Définition 1. Soit $\mathbb{K}$ un corps. Soit E un ensemble. On suppose E muni

- d'une loi de composition interne $+: E \times E \longrightarrow E$.
- d'une loi de composition externe $.: \mathbb{K} \times E \longrightarrow E$.

On dit que $(E, +, .)$ est un $\mathbb{K}$ -espace vectoriel lorsque

- $(E, +)$ est un groupe commutatif.
- $\forall x \in E, 1.x = x$.
- $\forall \lambda \in \mathbb{K}, \forall x, y \in E, \lambda.(x + y) = \lambda.x + \lambda.y$.
- $\forall \lambda, \mu \in \mathbb{K}, \forall x \in E, (\lambda + \mu).x = \lambda.x + \mu.x$.
- $\forall \lambda, \mu \in \mathbb{K}, \forall x \in E, (\lambda\mu).x = \lambda.(\mu.x)$.

Les éléments de E sont appelés des *vecteurs*. Les éléments de $\mathbb{K}$ sont appelés des *scalaires*.

1.2 Propriétés immédiates

Soit E un $\mathbb{K}$ -espace vectoriel.

Proposition 1. $\forall \lambda \in \mathbb{K}, \forall x \in E, \lambda x = 0 \iff \lambda = 0 \text{ ou } x = 0$.

Démonstration. Soit $\lambda \in \mathbb{K}$. On a $\lambda 0 = \lambda(0 + 0) = \lambda 0 + \lambda 0$ d'où, en soustrayant $\lambda 0$: $\lambda 0 = 0$. De même, pour tout $x \in E, 0x = 0$. Inversement, soient $\lambda \in \mathbb{K}$ et $x \in E$ tels que $\lambda x = 0$. Supposons $\lambda \neq 0$. On multiplie par $\frac{1}{\lambda}$:

$$0 = \frac{1}{\lambda}(\lambda x) = \left(\frac{1}{\lambda}\lambda\right)x = 1x = x$$

□

Proposition 2. $\forall x \in E, (-1)x = -x$.

Démonstration. $1x + (-1)x = (1 + (-1))x = 0x = 0$. □

1.3 Exemples

- $\mathbb{K}^n$

On munit $\mathbb{K}^n$ d'une structure de $\mathbb{K}$ -espace vectoriel avec les opérations ci-dessous.

$$\begin{aligned} + : \mathbb{K}^n &\times \mathbb{K}^n &\longrightarrow \mathbb{K}^n \\ (x_1, \dots, x_n) &, (y_1, \dots, y_n) &\mapsto (x_1 + y_1, \dots, x_n + y_n) \end{aligned}$$

et

$$\begin{aligned} \cdot : \mathbb{K} &\times \mathbb{K}^n &\longrightarrow \mathbb{K}^n \\ \lambda &, (x_1, \dots, x_n) &\mapsto (\lambda x_1, \dots, \lambda x_n) \end{aligned}$$

Remarque 1. $\mathbb{C}$ peut être vu à la fois comme le $\mathbb{C}$ -espace vectoriel $\mathbb{C}^1$ et le $\mathbb{R}$ -espace vectoriel $\mathbb{R}^2$. Ces deux structures sont différentes. Lorsque le corps de base n'est pas précisé, c'est le contexte qui permet de savoir.

- F^X

On se donne un ensemble X et un $\mathbb{K}$ -espace vectoriel F . L'ensemble des fonctions de X dans F , muni des opérations

- d'addition des fonctions : $(f + g)(x) = f(x) + g(x)$
- de multiplication d'une fonction par un scalaire : $(\lambda.f)(x) = \lambda.(f(x))$.

est lui-même un $\mathbb{K}$ -espace vectoriel. Comme cas particuliers, les fonctions de $\mathbb{R}$ vers $\mathbb{R}$, les suites à valeurs complexes, etc.

- $\mathbb{K}[X]$

- $\mathcal{M}_{pq}(\mathbb{K})$

- **Produit d'espaces vectoriels**

Soient E et F deux $\mathbb{K}$ -espaces vectoriels. L'ensemble $E \times F$ est muni de façon évidente d'une structure d'espace vectoriel. On peut évidemment faire un produit de 3, ..., n espaces vectoriels. On remarque que l'espace $\mathbb{K}^n$ est en fait le produit $\mathbb{K} \times \dots \times \mathbb{K}$ de n copies du $\mathbb{K}$ -espace vectoriel $\mathbb{K}$.

1.4 Combinaisons linéaires

Définition 2. Soit E un $\mathbb{K}$ -espace vectoriel, $x_1, \dots, x_n$ n vecteurs de E , et $\lambda_1, \dots, \lambda_n$ n scalaires. On appelle combinaison linéaire des x_i affectés des coefficients λ_i le vecteur

$$u = \sum_{i=1}^n \lambda_i . x_i$$

Exemple. Dans $\mathbb{R}^2$, soient $e_1 = (1, 1)$ et $e_2 = (0, 1)$. Tout vecteur de $\mathbb{R}^2$ s'écrit de façon unique comme combinaison linéaire de e_1 et e_2 . On dit que (e_1, e_2) est une base du plan.

Exemple. Dans $\mathbb{R}^3$, soient $e_1 = (1, 2, 1)$, $e_2 = (2, -1, 0)$ et $e_3 = (3, 1, 1)$. Soit $u = (x, y, z) \in \mathbb{R}^3$. Le vecteur u est combinaison linéaire des e_i si et seulement si il existe 3 scalaires $\lambda_1, \lambda_2, \lambda_3$ tels que $u = \lambda_1 e_1 + \lambda_2 e_2 + \lambda_3 e_3$. Il est équivalent de dire que le système ci-dessous admet au moins une solution :

$$(\mathcal{S}) \begin{cases} x = \lambda_1 + 2\lambda_2 + 3\lambda_3 \\ y = 2\lambda_1 - \lambda_2 + \lambda_3 \\ z = \lambda_1 + \lambda_3 \end{cases}$$

On voit qu'une condition nécessaire d'existence d'une solution pour $(\mathcal{S})$ est que $x + 2y = 5z$. Inversement, si cette égalité est vérifiée, on peut trouver une solution avec $\lambda_2 = 0$.

On donne maintenant une définition plus générale de la notion de combinaison linéaire.

Définition 3. Soit I un ensemble quelconque, $(x_i)_{i \in I}$ une famille de vecteurs de E , et $(\lambda_i)_{i \in I}$ une famille de scalaires telle que tous les λ_i , sauf un nombre fini, soient nuls. On utilisera la notation $\sum_{i \in I} \lambda_i x_i$ pour représenter $\lambda_{i_1} x_{i_1} + \dots + \lambda_{i_p} x_{i_p}$, où $i_1, \dots, i_p$ sont des éléments de I en dehors desquels les λ sont nuls. On convient également qu'une combinaison linéaire vide est le vecteur nul.

Remarque 2. Une famille de scalaires tous nuls sauf un nombre fini d'entre-eux est dite *presque nulle*.

Exemple. Dans $E = \mathbb{C}^{\mathbb{R}}$, on considère la famille $\mathcal{F} = (f_n)_{n \in \mathbb{Z}}$ définie par $\forall n \in \mathbb{Z}, \forall x \in \mathbb{R}, f_n(x) = e^{inx}$. Alors, la fonction $\cos^3$ est combinaison des éléments de la famille $\mathcal{F}$. En effet, pour tout réel x , on a

$$\cos^3 x = \frac{3}{4} \cos 3x + \frac{1}{4} \cos x$$

donc

$$\cos^3 = \frac{3}{4} f_3 + \frac{1}{4} f_1$$

2 Applications linéaires

2.1 C'est quoi ?

Définition 4. Soient E et F deux $\mathbb{K}$ -espaces vectoriels. Soit $f : E \longrightarrow F$. On dit que f est une application linéaire lorsque

- $\forall x, y \in E, f(x + y) = f(x) + f(y)$.
- $\forall \lambda \in \mathbb{K}, \forall x \in E, f(\lambda x) = \lambda f(x)$.

Bref, les applications linéaires sont les morphismes d'espaces vectoriels.

Remarque 3. Une application linéaire f est avant tout un morphisme de groupes additifs. Elle vérifie donc $f(0) = 0$.

Proposition 3. Soit $f : E \longrightarrow F$.

f est linéaire $\iff f$ conserve les combinaisons linéaires.

Démonstration. Supposons f linéaire. Montrons par récurrence sur $n \in \mathbb{N}$ que f conserve les combinaisons linéaires de n vecteurs. Pour $n = 0$, cela vient du fait que $f(0) = 0$. Pour $n = 1$, c'est parce que $f(\lambda x) = \lambda f(x)$. Pour $n = 2$,

$$f(\lambda x + \mu y) = f(\lambda x) + f(\mu y) = \lambda f(x) + \mu f(y)$$

Pour passer de n à $n + 1$, il suffit de remarquer que combiner $n + 1$ vecteurs revient à combiner les n premiers, puis à combiner le résultat avec le dernier. Donc, si on sait faire pour n et pour 2, on sait faire pour $n + 1$. Inversement, si f conserve les combinaisons linéaires, elle est évidemment linéaire, puisque l'addition et le produit par un scalaire sont des cas particuliers de combinaisons linéaires. $\square$

Remarque 4. En clair, la linéarité de f permet d'écrire des égalités du type

$$f\left(\sum_{i \in I} \lambda_i x_i\right) = \sum_{i \in I} \lambda_i f(x_i)$$

2.2 Quelques exemples

Cinq exemples choisis parmi les milliers que nous rencontrerons au fil du cours ...

- Les homothéties : $f : E \longrightarrow E; x \mapsto \alpha x$. Lorsque $E = \mathbb{K}$, ce sont les seules applications linéaires.
- $f : \mathbb{K}[X] \longrightarrow \mathbb{K}$ définie par $f(P) = P(1)$.
- $f(u) = \lim_{+\infty} u$, définie sur l'ensemble des suites réelles convergentes, à valeurs dans $\mathbb{R}$.
- $f : \mathcal{D}^1(\mathbb{R}) \longrightarrow \mathbb{R}^{\mathbb{R}}$ définie par $f(\varphi) = \varphi'$.
- $f : \mathcal{C}^0(\mathbb{R}) \longrightarrow \mathcal{C}^1(\mathbb{R})$ définie par $f(\varphi)(x) = \int_0^x \varphi(t) dt$.

2.3 Endomorphismes du plan

On se propose de rechercher toutes les applications linéaires $\mathbb{R}^2 \longrightarrow \mathbb{R}^2$ et de déterminer lesquelles sont bijectives. Soient $e_1 = (1, 0)$ et $e_2 = (0, 1)$ les vecteurs de la *base canonique*

de $\mathbb{R}^2$. Soit $u = (x, y) = xe_1 + ye_2$ un vecteur quelconque. On se donne f linéaire. Posons $u' = (x', y') = f(u)$. Alors, $f(u) = xf(e_1) + yf(e_2)$, d'où

$$(\mathcal{S}) \begin{cases} x' &= ax + cy \\ y' &= bx + dy \end{cases}$$

où $f(e_1) = (a, b)$ et $f(e_2) = (c, d)$. Il est clair qu'inversement l'application ci-dessus est linéaire. Les endomorphismes de $\mathbb{R}^2$ sont donc caractérisés par la donnée de 4 réels a, b, c, d .

Cherchons maintenant lesquels sont injectifs. Si f est injectif, l'équation $f(u) = f(0)$ doit avoir une unique solution : $(0, 0)$. On s'intéresse donc au système

$$\begin{cases} ax + cy &= 0 \\ bx + dy &= 0 \end{cases}$$

À quelle condition ce système admet-il pour unique solution le couple $(0, 0)$? On voit que si un couple (x, y) est solution, alors (combinaison judicieuse) $(ad - bc)x = 0$ et $(ad - bc)y = 0$. Donc, si $ad - bc \neq 0$, on aura $x = y = 0$. Inversement, si $ad - bc = 0$, alors une petite discussion sur la nullité ou non du couple (a, c) permet de trouver une solution non nulle au système.

Conclusion : f est injective si et seulement si $ad - bc \neq 0$.

Cherchons maintenant une condition de surjectivité. On constate, en résolvant le système $(\mathcal{S})$, que si $ad - bc \neq 0$, alors f est surjective. Mieux, dans ce cas le système admet une solution unique : f est donc bijective. La réciproque de f est d'ailleurs linéaire, et est donnée par $(x', y') = f^{-1}(x, y)$ où

$$\begin{cases} x' &= \frac{1}{ad - bc} (dx - cy) \\ y' &= \frac{1}{ad - bc} (-bx + ay) \end{cases}$$

Si $ad - bc = 0$, alors, dans le cas où le couple (a, b) n'est pas le couple nul, on constate que si $(\mathcal{S})$ possède une solution, alors $bx' - ay' = 0$. L'application f n'est donc pas surjective. Si $(a, b) = (0, 0)$, on trouve de même des éléments de $\mathbb{R}^2$ qui n'ont pas d'antécédent par f .

Conclusion : f est surjective si et seulement si $ad - bc \neq 0$.

2.4 Vocabulaire, notations

Soient E et F deux espaces vectoriels.

On note $\mathcal{L}(E, F)$ l'ensemble des applications linéaires de E dans F .

Lorsque $E = F$, on note plus simplement $\mathcal{L}(E)$.

Lorsque $F = \mathbb{K}$ on note $E^* = \mathcal{L}(E, \mathbb{K})$ et cet ensemble est appelé le *dual* de E .

Définition 5. Soit $f \in \mathcal{L}(E, F)$. on dit que

- f est un endomorphisme de E lorsque $E = F$.
- f est un isomorphisme lorsque f est bijective.
- f est un automorphisme de E lorsque f est un endomorphisme bijectif.
- f est une forme linéaire sur E lorsque $F = \mathbb{K}$ (le corps de base).

2.5 Opérations sur les applications linéaires

On donne ici quelques théorèmes relatifs aux opérations sur les applications linéaires.

Proposition 4. Soient E et F deux $\mathbb{K}$ -espaces vectoriels. L'ensemble $\mathcal{L}(E, F)$, muni de l'addition des fonctions et de la multiplication par un scalaire, est un $\mathbb{K}$ -espace vectoriel.

Démonstration. Vérifications immédiates. $\square$

En clair, la somme de deux applications linéaires est linéaire, le produit d'une application linéaire par un scalaire est linéaire, et ces opérations vérifient tous les axiomes de la structure d'espace vectoriel. Noter que $\mathcal{L}(E, F) \subset F^E$ est ce que l'on appellera bientôt un sous-espace vectoriel de F^E .

Proposition 5. Soient E, F, G trois $\mathbb{K}$ -espaces vectoriels. Soit $f \in \mathcal{L}(E, F)$ et $g \in \mathcal{L}(F, G)$. Alors

- $g \circ f \in \mathcal{L}(E, G)$.
- Si f est un isomorphisme, alors $f^{-1} \in \mathcal{L}(F, E)$.

Démonstration. Soient $x, y \in E$ et $\lambda, \mu \in \mathbb{K}$. On a

$$g \circ f(\lambda x + \mu y) = g(f(\lambda x + \mu y)) = g(\lambda f(x) + \mu f(y)) = \lambda g \circ f(x) + \mu g \circ f(y)$$

Supposons $f \in \mathcal{L}(E, F)$ bijectif et posons $g = f^{-1}$. Soient $x, y \in F$. Posons $x = f(X)$ et $y = f(Y)$. On a

$$g(\lambda x + \mu y) = g(\lambda f(X) + \mu f(Y)) = g(f(\lambda X + \mu Y)) = \lambda X + \mu Y = \lambda g(X) + \mu g(Y)$$

$\square$

Proposition 6. Soit E un $\mathbb{K}$ -espace vectoriel. L'ensemble $\mathcal{L}(E)$, muni des lois $+$, $\cdot$ et $\circ$ est une $\mathbb{K}$ -algèbre (non commutative en général).

Démonstration. Nous savons déjà que $\mathcal{L}(E)$ est un $\mathbb{K}$ espace vectoriel. Il reste à vérifier les propriétés d'anneau et une petite propriété dont nous n'avons pas encore parlé. Le neutre pour la composition est id_E , qui est clairement linéaire. La composition est associative, ceci est une propriété universelle. Reste à voir la distributivité par rapport à l'addition. Soient donc $f, g, h \in \mathcal{L}(E)$ et $x \in E$. On a

$$(g + h) \circ f(x) = (g + h)(f(x)) = g(f(x)) + h(f(x)) = (g \circ f + h \circ f)(x)$$

D'où

$$(g + h) \circ f = g \circ f + h \circ f$$

Remarquons que ceci est vrai sans utiliser la linéarité. En revanche

$$f \circ (g + h)(x) = f(g(x) + h(x)) = (f \circ g + f \circ h)(x)$$

donc

$$f \circ (g + h) = f \circ g + f \circ h$$

Et là, on a eu besoin de la linéarité de f .

La dernière propriété à vérifier relie le produit par un scalaire à la composition. Il s'agit de voir que

$$\forall f, g \in \mathcal{L}(E), \forall \lambda \in \mathbb{K}, \lambda(g \circ f) = (\lambda g) \circ f = g \circ (\lambda f)$$

C'est vrai parce que g est linéaire. Laissons la non commutativité de côté pour l'instant. Nous verrons que dès que E est « suffisamment gros » (de dimension strictement supérieure à 1), $\mathcal{L}(E)$ n'est pas commutatif. $\square$

Proposition 7. *Soit E un $\mathbb{K}$ -espace vectoriel. L'ensemble des automorphismes de E est un groupe pour la composition, non commutatif en général.*

Démonstration. C'est l'ensemble des éléments inversibles de l'anneau $\mathcal{L}(E)$. Pour la non commutativité, même remarque que ci-dessus. $\square$

Définition 6. Le groupe des automorphismes de l'espace vectoriel E est appelé le *groupe linéaire* de E , et est noté $GL(E)$.

3 Sous-espaces vectoriels

3.1 C'est quoi ?

Définition 7. Soit E un $\mathbb{K}$ -espace vectoriel. Soit $F \subset E$. On dit que F est un sous-espace vectoriel de E lorsque, F est stable pour les lois $+$ et $\cdot$, et que, muni de ces deux lois, F est lui-même un espace vectoriel.

3.2 Caractérisation

Proposition 8. Soit E un $\mathbb{K}$ -espace vectoriel. Soit $F \subset E$. F est un sous-espace vectoriel de E si et seulement si

- F est non vide.
- $\forall x, y \in F, x + y \in F$.
- $\forall \lambda \in \mathbb{K}, \forall x \in F, \lambda x \in F$.

Démonstration. Exercice. $\square$

Remarque 5.

- Cette proposition nous dit que pour prouver qu'un ensemble est un espace vectoriel, il suffit tout d'abord de remarquer qu'il est inclus dans un « gros » espace vectoriel, puis de vérifier trois petites propriétés. Comparer avec la définition d'espace vectoriel, qui nécessite de prouver 10 propriétés. Moralité, on utilisera ce théorème à chaque fois que ce sera possible.
- Les deux dernières conditions peuvent être remplacées par

$$\forall \lambda \in \mathbb{K}, \forall x, y \in E, \lambda x + y \in E$$

- Une partie non vide de E est un s.e.v. de E si et seulement si elle est stable par combinaisons linéaires.

3.3 Exemples

- Pour tout espace vectoriel E , $\{0\}$ et E sont des sev de E .
- Les droites vectorielles sont des s.e.v. de $\mathbb{R}^2$. Les droites et les plans vectoriels sont des sev de $\mathbb{R}^3$.
- L'ensemble des suites réelles convergentes est un sev de l'espace $\mathbb{R}^{\mathbb{N}}$ des suites réelles.
- L'ensemble des fonctions $f : \mathbb{R} \rightarrow \mathbb{R}$ de classe $\mathcal{C}^3$ telles que $f(0) = 0$ est un sev de $\mathbb{R}^{\mathbb{R}}$.
- Pour tout $n \geq 0$, l'ensemble $\mathbb{R}_n[X]$ des polynômes de degré inférieur ou égal à n est un sev de $\mathbb{R}[X]$.
- L'ensemble $\mathcal{T}_n(\mathbb{K})$ des matrices triangulaires supérieures $n \times n$ est un sev de $\mathcal{M}_n(\mathbb{K})$.

3.4 Image directe et réciproque par une application linéaire

Proposition 9. Soient E, F deux $\mathbb{K}$ -espaces vectoriels, et $f \in \mathcal{L}(E, F)$. Soient E' un s.e.v de E , et F' un s.e.v de F . Alors, $f(E')$ est un s.e.v de F , et $f^{-1}(F')$ est un s.e.v de E .

Démonstration. Prouvons le résultat pour l'image directe. L'image réciproque est laissée en exercice. Tout d'abord, $0 \in E'$, donc $f(0) = 0 \in f(E') : f(E')$ est non vide. Soient $u, v \in f(E')$ et $\lambda \in \mathbb{K}$. Il existe $x, y \in E'$ tels que $u = f(x)$ et $v = f(y)$. On a alors

$$\lambda u + v = \lambda f(x) + f(y) = f(\lambda x + y) \in f(E')$$

□

Deux cas fondamentaux vont suivre.

3.5 Noyau et image d'une application linéaire

Définition 8. Soit $f \in \mathcal{L}(E, F)$ une application linéaire. On appelle noyau de f le sous-espace vectoriel de E

$$\ker f = f^{-1}(\{0\})$$

et image de f le sous-espace vectoriel de F

$$\operatorname{Im} f = f(E)$$

Exemple.

- $f : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ définie par $f(x, y) = (x + 3y, 2x + 6y)$. Soit $u = (x, y) \in \mathbb{R}^2$. Alors $u \in \ker f$ si et seulement si $x + 3y = 0$. $\ker f$ est la droite de $\mathbb{R}^2$ engendrée par $(-3, 1)$. On a $u \in \operatorname{Im} f$ si et seulement si il existe $a, b \in \mathbb{R}$ tels que $x = a + 3b$ et $y = 2a + 6b$. Il est facile de voir que ceci équivaut à $y = 2x$. Donc, l'image de f est la droite engendrée par $(1, 2)$.
- Soit $\varphi : \mathcal{C}^1(\mathbb{R}) \rightarrow \mathbb{R}^{\mathbb{R}}$ définie par $\varphi(f) = f'$. Alors $\ker \varphi$ est l'ensemble des fonctions constantes et le cours d'intégration nous montrera que $\operatorname{Im} \varphi = \mathcal{C}^0(\mathbb{R})$.

Proposition 10. Soit $f \in \mathcal{L}(E, F)$.

- f est injective si et seulement si $\ker f = \{0\}$.
- f est surjective si et seulement si $\operatorname{Im} f = F$.

Démonstration. Voir le chapitre sur les groupes. En effet, une application linéaire est avant tout un morphisme de groupes. □

4 Opérations sur les s.e.v

4.1 Intersection

Proposition 11. *Soit E un espace vectoriel. Toute intersection de sous-espaces vectoriels de E est un sous-espace vectoriel de E .*

Démonstration. Soit $(E_i)_{i \in I}$ une famille de sev de E . Soit $F = \bigcap_{i \in I} E_i$.

0 est dans tous les E_i , donc dans F . Ainsi, F est non vide.

Soient $x, y \in F$ et $\lambda \in \mathbb{K}$. x et y sont dans tous les E_i , donc $\lambda x + y$ aussi. $\square$

4.2 s.e.v engendré par une partie

Proposition 12. *Soit A une partie d'un espace vectoriel E . L'intersection de tous les s.e.v de E contenant A est encore un s.e.v de E contenant A , et c'est même le plus petit au sens de l'inclusion. On l'appelle le s.e.v de E engendré par A et on le note $\text{Vect}(A)$ ou $\langle A \rangle$.*

Démonstration. Toute intersection de sev de E est encore un sev de E . De plus, toute intersection de parties de E contenant A contient aussi A . $\square$

Exemple.

- $\langle \emptyset \rangle = \{0\}$.
- Si F est un sev de E , alors $\langle F \rangle = F$.
- Soit $u = (1, 2) \in \mathbb{R}^2$: $\langle \{u\} \rangle$ (noté plus simplement $\langle u \rangle$) est la droite vectorielle engendrée par u . Ce sera évident dans deux secondes.

Proposition 13. *Soit A une partie d'un espace vectoriel E . Le s.e.v de E engendré par A est l'ensemble des combinaisons linéaires d'éléments de A .*

Démonstration. Notons F l'ensemble des combinaisons linéaires d'éléments de A . Il est clair que F est un s.e.v de E contenant A . De plus tout s.e.v de E contenant A doit contenir F , car un s.e.v est stable par combinaison linéaire. Donc, $F = \langle A \rangle$. $\square$

Un cas important est celui où la partie A est finie. Par exemple, si $u \in E$, alors

$$\langle u \rangle = \{\lambda u, \lambda \in \mathbb{K}\}$$

Si u et v sont deux vecteurs de E , alors

$$\langle u, v \rangle = \{\lambda u + \mu v, \lambda, \mu \in \mathbb{K}\}$$

Remarquer l'abus de notation qui consiste à ne pas écrire les accolades.

Exercice. On se place dans $\mathbb{R}^3$, et on pose $e_1 = (1, 2, 3)$, $e_2 = (4, 5, 6)$ et $e_3 = (7, 8, 9)$. Déterminer $\langle e_1, e_2, e_3 \rangle$.

4.3 Quelques propriétés

Proposition 14. Soient A et B deux parties d'un espace vectoriel E . On a

- $A \subset \langle A \rangle$ avec égalité si et seulement si A est un sev de E .
- $\langle \langle A \rangle \rangle = \langle A \rangle$.
- $A \subset B \implies \langle A \rangle \subset \langle B \rangle$.
- $\langle A \cap B \rangle \subset \langle A \rangle \cap \langle B \rangle$.

Démonstration.

Premier point : $\langle A \rangle$ est le plus petit ...contenant A , donc il contient A . Si on a égalité, alors $\langle A \rangle$ est le plus petit sev de E ..., c'est donc un sev de E . Inversement, si A est un sev de E , il est clair que le plus petit sev de E qui contient A est A .

Second point : puisque $\langle A \rangle$ est un sev de E , on a $\langle \langle A \rangle \rangle = \langle A \rangle$.

Troisième point : supposons $A \subset B$. Alors $A \subset \langle B \rangle$. Mais $\langle A \rangle$ est le plus petit sev de E contenant A , et $\langle B \rangle$ en est un, donc $\langle A \rangle \subset \langle B \rangle$.

Dernier point : $A \cap B \subset A$, donc $\langle A \cap B \rangle \subset \langle A \rangle$. De même, $\langle A \cap B \rangle \subset \langle B \rangle$ d'où l'inclusion demandée. $\square$

Remarque 6. Soient $A = \{(1, 2)\}$ et $B = \{(2, 4)\}$ deux parties de $\mathbb{R}^2$. On a

$$\langle A \rangle = \langle B \rangle = \langle (1, 2) \rangle$$

donc

$$\langle A \rangle \cap \langle B \rangle = \langle (1, 2) \rangle$$

En revanche,

$$\langle A \cap B \rangle = \{0\}$$

L'inclusion réciproque dans le dernier point de la proposition ci-dessus est donc fausse en général.

5 Sous-espaces supplémentaires

5.1 Somme de deux s.e.v

Soient F et G deux s.e.v d'un espace vectoriel E . On appelle somme de F et G la partie de E

$$F + G = \{x + y, x \in F, y \in G\}$$

Proposition 15. $F + G$ est un sev de E . Précisément,

$$F + G = \langle F \cup G \rangle$$

Démonstration. $F + G$ est non vide et stable par combinaison linéaire. C'est donc un sev de E .

Pour tout $x \in F$, on a $x = x + 0 \in F + G$ donc $F \subset F + G$. De même pour G , d'où $F \cup G \subset F + G$.

Soit H un sev de E contenant $F \cup G$. Pour tout $x \in F$ et $y \in G$, $x, y \in H$ donc $x + y \in H$. Ainsi, $F + G \subset H$. $F + G$ est donc le plus petit sev de E contenant $F \cup G$, c'est à dire que $F + G = \langle F \cup G \rangle$. $\square$

Remarque 7. On généralise à la somme de n sev $E_1, \dots, E_n$ d'un espace vectoriel E .

5.2 Sommes directes - s.e.v supplémentaires

Définition 9. On dit que F et G sont en somme directe lorsque tout élément z de $F + G$ s'écrit de façon unique sous la forme $z = x + y$, avec $x \in F$ et $y \in G$.

Notation. Lorsque F et G sont en somme directe, on note $F \oplus G$ la somme de F et G .

C'est l'unicité qui importe, vu que l'existence est assurée par la définition même de la somme de deux sev.

Proposition 16. Deux sev de E , F et G , sont en somme directe, si et seulement si

$$F \cap G = \{0\}$$

Démonstration. Supposons F et G en somme directe. Soit $x \in F \cap G$. Alors, $x = x + 0 = 0 + x$ et on a deux décompositions de x sur $F + G$. Comme la somme est directe, la décomposition est unique, et ainsi $x = 0$.

Inversement, supposons que $F \cap G = \{0\}$. Soit $x = x_1 + x_2 = y_1 + y_2$ un élément de $F + G$, donné avec deux décompositions sur la somme. On a alors $x_1 - y_1 = y_2 - x_2$ et ces deux quantités sont dans $F \cap G$. Elles sont donc nulles. $\square$

Définition 10. Deux sev de E , F et G sont dit supplémentaires lorsque

$$F \oplus G = E$$

En d'autres termes, F et G sont supplémentaires lorsque tout vecteur de E s'écrit de façon unique comme somme d'un vecteur de F et d'un vecteur de G .

Exemple.

- Deux droites vectorielles distinctes du plan sont supplémentaires.
- Les ensembles $\mathcal{P}$ des fonctions paires et $\mathcal{I}$ des fonctions impaires de $\mathbb{R}$ vers $\mathbb{R}$ sont des sev supplémentaires de $\mathbb{R}^{\mathbb{R}}$.
- Les ensembles $\mathcal{S}_n(\mathbb{K})$ et $\mathcal{A}_n(\mathbb{K})$ des matrices symétriques et des matrices antisymétriques $n \times n$ sont des sev supplémentaires de $\mathcal{M}_n(\mathbb{K})$.

5.3 Projecteurs

Définition 11. Soit E un espace vectoriel. On appelle projecteur de E tout endomorphisme p de E vérifiant

$$p \circ p = p$$

Proposition 17. Soit E un espace vectoriel. Soit $p \in \mathcal{L}(E)$. L'endomorphisme p est un projecteur si et seulement si il existe deux sous-espaces vectoriels de E supplémentaires, F et G tels que

- $\forall x \in F, p(x) = x$.
- $\forall x \in G, p(x) = 0$.

Les deux sev F et G sont uniquement déterminés par la donnée du projecteur p . Précisément,

$$F = \ker(p - id) = Im\ p \text{ et } G = \ker p$$

Démonstration. Supposons que p est un projecteur de E . Soient $F = \ker(p - id)$ et $G = \ker p$. Alors, F et G sont des sev de E puisque ce sont des noyaux d'applications linéaires. De plus, $F \cap G = \{0\}$ puisque un vecteur x de leur intersection doit vérifier $p(x) = x = 0$. Enfin, pour tout x de E , on a

$$x = (x - p(x)) + p(x)$$

et $p(x - p(x)) = 0$, donc $x - p(x) \in G$, et $p(p(x)) = p(x)$, donc $p(x) \in F$, d'où $F \oplus G = E$.

La réciproque est immédiate.

Si F_1 et G_1 vérifient les mêmes propriétés que F et G , alors pour tout x de F_1 , on a $p(x) = x$, donc $x \in F = \ker(p - id)$. De même, $G_1 \subset G$. Soit $x \in F$. Alors, $x = x_1 + x_2$, avec $x_1 \in F_1$ et $x_2 \in G_1$. On en tire $p(x) = p(x_1) + p(x_2)$ ou encore $x = x_1$. Donc, $x \in F_1$. De même, $G_1 = G$. $\square$

Définition 12. Le projecteur p est appelé projecteur sur F parallèlement à G .

5.4 Symétries

Définition 13. Soit E un espace vectoriel. On appelle symétrie de E tout endomorphisme f de E vérifiant

$$f \circ f = id$$

Proposition 18. Soit E un espace vectoriel sur un corps $\mathbb{K}$ dans lequel $2 \neq 0$. Soit $f \in \mathcal{L}(E)$. L'endomorphisme f est une symétrie si et seulement si il existe deux sous-espaces vectoriels de E supplémentaires, F et G tels que

- $\forall x \in F, f(x) = x$.
- $\forall x \in G, f(x) = -x$.

Les deux sev F et G sont uniquement déterminés par la donnée dde la symétrie f .

Démonstration. Supposons que f est une symétrie de E . Soient $F = \ker(f - id)$ et $G = \ker(f + id)$. Alors, F et G sont des sev de E puisque ce sont des noyaux d'applications linéaires. De plus, $F \cap G = \{0\}$ puisque un vecteur x de leur intersection doit vérifier $f(x) = x = -x$. Enfin, pour tout x de E , on a

$$x = \frac{x + f(x)}{2} + \frac{x - f(x)}{2}$$

Remarquons que

$$f\left(\frac{x + f(x)}{2}\right) = \frac{x + f(x)}{2}$$

donc $\frac{x+f(x)}{2} \in F$. De même,

$$f\left(\frac{x-f(x)}{2}\right) = \frac{x-f(x)}{2}$$

donc $\frac{x-f(x)}{2} \in G$, d'où $F \oplus G = E$.

La réciproque est immédiate.

Si F_1 et G_1 vérifient les mêmes propriétés que F et G , alors pour tout x de F_1 , on a $f(x) = x$, donc $x \in F = \text{Ker}(f - \text{id})$. De même, $G_1 \subset G$. Soit $x \in F$. Alors, $x = x_1 + x_2$, avec $x_1 \in F_1$ et $x_2 \in G_1$. On en tire $f(x) = f(x_1) + f(x_2)$ ou encore $x = x_1 - x_2$. Donc, en additionnant les deux égalités donnant x , $x = x_1 \in F_1$. De même, $G_1 = G$. $\square$

Définition 14. On dit que f est la symétrie par rapport à F , parallèlement à G .

6 Familles remarquables de vecteurs

6.1 Famille libre, famille génératrice, base

Définition 15. Soit $\mathcal{F} = (e_i)_{i \in I}$ une famille de vecteurs d'un espace vectoriel E . On dit que la famille $\mathcal{F}$ est

- Une famille libre lorsque tout vecteur de E s'écrit d'une façon unique comme combinaison linéaire des vecteurs de $\mathcal{F}$.
- Une famille génératrice de E lorsque tout vecteur de E s'écrit d'au moins une façon comme combinaison linéaire des vecteurs de $\mathcal{F}$.
- Une base de E lorsque tout vecteur de E s'écrit d'exactement une façon comme combinaison linéaire des vecteurs de $\mathcal{F}$.

Exemple.

- La famille vide est libre.
- Une famille (u) de 1 vecteur est libre si et seulement si ce vecteur est non nul.
- Dans $\mathbb{K}^n$, la famille $(e_1, \dots, e_n)$ où e_i est le n -uplet dont toutes les coordonnées sont nulles, sauf la i ème qui vaut 1, est une base. On l'appelle la base canonique de $\mathbb{K}^n$.
- La famille $(X^n)_{n \geq 0}$ est une base de $\mathbb{K}[X]$. On l'appelle la *base canonique* de $\mathbb{K}[X]$.
- La famille $(E_{ij})_{1 \leq i \leq p, 1 \leq j \leq q}$ des matrices élémentaires de taille $p \times q$ est une base de $\mathcal{M}_{pq}(\mathbb{K})$. On l'appelle la *base canonique* de $\mathcal{M}_{pq}(\mathbb{K})$.

6.2 Vocabulaire

Les vecteurs d'une famille libre sont dits *indépendants*.

Une famille qui n'est pas libre est dite *liée*.

Si $\mathcal{B} = (e_i)_{i \in I}$ est une base de E , et $x = \sum_{i \in I} x_i e_i$ est un vecteur de E , les scalaires x_i , tous nuls sauf un nombre fini d'entre eux, sont appelés les *coordonnées* du vecteur x dans la base $\mathcal{B}$.

6.3 Propriétés faciles

Proposition 19. *Une famille $\mathcal{F} = (e_i)_{i \in I}$ de vecteurs de E est libre si et seulement si pour toute famille presque nulle de scalaires $(\lambda_i)_{i \in I}$, on a*

$$\sum_{i \in I} \lambda_i e_i = 0 \implies \forall i \in I, \lambda_i = 0$$

Démonstration. Supposons la famille $\mathcal{F}$ libre. Le vecteur nul s'écrit d'au plus une façon comme combinaison des e_i . Or, $0 = \sum_{i \in I} 0 e_i$. D'où l'implication. Inversement, supposons $\mathcal{F}$ liée. Il existe donc un vecteur x de E s'écrivant de deux façons différentes comme combinaison des vecteurs de $\mathcal{F}$. On a donc

$$x = \sum_{i \in I} \lambda_i e_i = \sum_{i \in I} \mu_i e_i$$

et il existe au moins un $i \in I$ tel que $\lambda_i \neq \mu_i$. En posant pour tout $i \in I$, $\nu_i = \lambda_i - \mu_i$, il vient $0 = \sum_{i \in I} \nu_i e_i$ avec les ν_i pas tous nuls, ce qui est la négation de l'implication. $\square$

Exercice. Montrer que les vecteurs $e_1 = (1, 1, 2)$, $e_2 = (1, 1, 1)$ et $e_3 = (2, 1, 0)$ forment une base $\mathcal{B}$ de $\mathbb{R}^3$, et calculer, pour tout vecteur $u = (x, y, z) \in \mathbb{R}^3$, ses coordonnées dans la base $\mathcal{B}$.

Proposition 20. *Une famille $\mathcal{F} = (e_i)_{i \in I}$ de vecteurs de E est liée si et seulement si l'un des vecteurs de la famille est combinaison linéaire des autres.*

Démonstration. Supposons la famille liée. On a alors une combinaison nulle $\sum_{i \in I} \lambda_i e_i = 0$ avec au moins un des λ_i non nul. Soit $i_0 \in I$ tel que $\lambda_{i_0} \neq 0$. Alors, on peut exprimer e_{i_0} comme combinaison des autres vecteurs :

$$e_{i_0} = \sum_{i \neq i_0} \left(-\frac{\lambda_i}{\lambda_{i_0}} \right) e_i$$

Inversement, supposons qu'il existe $i_0 \in I$ tel que

$$e_{i_0} = \sum_{i \neq i_0} \mu_i e_i$$

où les $\mu_i \in \mathbb{K}$ sont presque tous nuls. Alors, on a une combinaison nulle $\sum_{i \in I} \lambda_i e_i = 0$, en posant $\lambda_i = \mu_i$ si $i \neq i_0$, et $\lambda_{i_0} = -1$. Combinaison nulle, donc, mais au moins un des coefficients est non nul, puisqu'il vaut -1 . $\square$

Histoire d'insister, $\mathcal{F}$ est liée si et seulement si

$$\exists i_0 \in I, \exists (\lambda_i)_{i \neq i_0}, e_{i_0} = \sum_{i \neq i_0} \lambda_i e_i$$

Il s'agit bien de « $\exists i_0$ ». L'un des vecteurs est combinaison des autres, mais on ne peut pas a priori dire lequel. Prenons un exemple. Dans $E = \mathbb{R}^3$, soient $e_1, (1, 2, 3)$, $e_2 = (2, 4, 6)$ et $e_3 = (1, 2, 4)$. La famille (e_1, e_2, e_3) est liée car

$$e_2 = 2e_1 + 0e_3$$

Remarquons cependant qu'il est impossible d'exprimer e_3 comme combinaison linéaire de e_1 et e_2 (exercice).

Les preuves des trois propriétés ci-dessous sont laissées en exercice.

Proposition 21. *Toute « sous-famille » d'une famille libre est libre.*

Proposition 22. *Toute « sur-famille » d'une famille liée est liée.*

Proposition 23. *Toute « sur-famille » d'une famille génératrice est génératrice.*

Remarque 8. Les problèmes intéressants se posent donc lorsqu'on considère des sur-familles de familles libres ou des sous-familles de familles génératrices, c'est à dire des « grosses » familles libres ou des « petites » familles génératrices.

6.4 Familles remarquables et applications linéaires

Proposition 24. *Soient E et F deux $\mathbb{K}$ -espaces vectoriels. Soit $\mathcal{F}$ une famille de vecteurs de E . Soit $f \in \mathcal{L}(E, F)$. On a*

- $f(<\mathcal{F}>) = <f(\mathcal{F})>$.
- Si f est injective, et $\mathcal{F}$ est libre, alors $f(\mathcal{F})$ est libre.
- Si f est surjective, et $\mathcal{F}$ est génératrice de E , alors $f(\mathcal{F})$ est génératrice de F .

Démonstration. Posons $\mathcal{F} = (e_i)_{i \in I}$.

Premier point : Soit $y \in E$. Alors $y \in f(< \mathcal{F} >)$ si et seulement si $\exists (\lambda_i)_{i \in I}, y = f(\sum_i \lambda_i e_i)$. Comme

$$f\left(\sum_i \lambda_i e_i\right) = \sum_i \lambda_i f(e_i)$$

ceci équivaut à $y \in < f(\mathcal{F}) >$.

Deuxième point : Supposons $\sum_i \lambda_i f(e_i) = 0$. Ceci équivaut à $f(\sum_i \lambda_i e_i) = 0$, c'est à dire $\sum_i \lambda_i e_i \in \ker f$. Mais f est injective, donc $\sum_i \lambda_i e_i = 0$. Donc les λ_i sont nuls puisque les e_i sont libres.

Troisième point : $< f(\mathcal{F}) > = f(< \mathcal{F} >) = f(< E >) = F$. $\square$

Dans le cas des bases, on a des résultats plus précis.

Proposition 25. Soient E et F deux $\mathbb{K}$ -espaces vectoriels. Soit $\mathcal{B}$ une base de l'espace vectoriel E . Soit $f \in \mathcal{L}(E, F)$. Alors

- f est injective, si et seulement si $f(\mathcal{B})$ est libre.
- f est surjective, si et seulement si $f(\mathcal{B})$ est génératrice de F .
- f est bijective, si et seulement si $f(\mathcal{B})$ est une base de F .

Démonstration. Posons $\mathcal{B} = (e_i)_{i \in I}$. Le troisième point est évidemment une conséquence des deux premiers, et les implications des deux premiers points nous sont déjà connues. Supposons $f(\mathcal{B})$ libre. Soit $x = \sum_i \lambda_i e_i \in \ker f$. On a donc $0 = f(x) = \sum_i \lambda_i f(e_i)$. Les $f(e_i)$ sont libres, donc les λ_i sont nuls. Donc $x = 0$ et f est bien injective. $f(\mathcal{B})$ génératrice $\Rightarrow f$ surjective est laissé en exercice. $\square$

Proposition 26. Soient E et F deux $\mathbb{K}$ -espaces vectoriels. Soit $\mathcal{B} = (e_i)_{i \in I}$ une base de E . Soit $\mathcal{F}' = (e'_i)_{i \in I}$ une famille de vecteurs de F indexée par le même ensemble I que la base $\mathcal{B}$. Il existe alors une unique application linéaire $f : E \rightarrow F$ telle que

$$\forall i \in I, f(e_i) = e'_i$$

Démonstration. Supposons qu'une telle application linéaire f existe. Soit $x \in E$. On écrit $x = \sum_{j \in I} x_j e_j$. On a alors $f(x) = \sum_{j \in I} x_j e'_j$; d'où l'existence d'au plus une telle application. Inversement, l'application que nous venons de trouver est bien linéaire (exercice) et envoie les e_j sur les e'_j . D'où l'existence. $\square$

Remarque 9. Ce théorème est d'une **importance capitale**. Une application linéaire est caractérisée par la donnée de l'image des vecteurs d'une base de l'espace de départ. Grâce à ce théorème, on peut se donner de façon simple des applications linéaires.

Ainsi, par exemple, soit (e_1, e_2, e_3) la base canonique de R^3 . Soit f l'application linéaire de R^3 dans $\mathbb{R}^2$ définie par $f(e_1) = (1, 2)$, $f(e_2) = (3, 4)$ et $f(e_3) = (5, 6)$. On obtient facilement, en cas de besoin,

$$f(x, y, z) = (x + 3y + 5z, 2x + 4y + 6z)$$

6.5 Familles libres maximales, génératrices minimales

Proposition 27. *Soit $\mathcal{F}$ une famille de vecteurs de l'espace vectoriel E . Il y a équivalence entre*

- $\mathcal{F}$ est libre et toute sur-famille de $\mathcal{F}$ est liée.
- $\mathcal{F}$ est génératrice et toute sous-famille de $\mathcal{F}$ est non-génératrice.
- $\mathcal{F}$ est une base de E .

Ainsi, une famille de vecteurs de E est une base si et seulement si elle est libre maximale ou génératrice minimale.

Démonstration. Posons $\mathcal{F} = (e_i)_{i \in I}$. Supposons $\mathcal{F}$ libre maximale. Soit $e \in E$. La famille $\mathcal{F} \cup (e)$ est donc liée. Il existe une combinaison nulle $0 = \sum_i \lambda_i e_i + \lambda e$ dans laquelle tous les coefficients ne sont pas nuls. En fait, $\lambda \neq 0$ puisque la famille $\mathcal{F}$ est libre. On peut donc écrire e comme combinaison des vecteurs de $\mathcal{F}$: $\mathcal{F}$ est génératrice et c'est donc une base.

Inversement, supposons que $\mathcal{F}$ est une base. Soit $e \in E$. Soit $\mathcal{G} = \mathcal{F} \cup (e)$. On peut écrire e comme combinaison des e_i puisque $\mathcal{F}$ est une base. Donc $\mathcal{G}$ est liée. $\mathcal{F}$ est bien libre maximale.

L'équivalence entre base et génératrice minimale est laissée en exercice. $\square$

6.6 bases et sev supplémentaires

Proposition 28. *Soit E un espace vectoriel, et $F_1, \dots, F_n$ des sev de E en somme directe tels que $\bigoplus_{k=1}^n F_k = E$. Pour $k = 1, \dots, n$, soit $\mathcal{B}_k$ une base de F_k . Alors, $\mathcal{B}_1 \cup \dots \cup \mathcal{B}_n$ est une base de E .*

Démonstration. Exercice. Il suffit de le faire pour $n = 2$, puis d'effectuer une récurrence sur n . $\square$

Chapitre 23

Dimension finie

1 Dimension

1.1 Espaces de dimension finie

Définition 1. Soit E un $\mathbb{K}$ -espace vectoriel. On dit que E est de dimension finie lorsque E possède au moins une famille génératrice finie.

Exemple.

- $\mathbb{K}^n$, dont la base canonique a n éléments.
- $\mathcal{M}_{pq}(\mathbb{K})$, dont la base canonique a pq éléments.
- $\mathbb{K}_n[X]$, l'espace des polynômes de degré inférieur ou égal à n , est engendré par les polynômes X^k , $k = 0, \dots, n$.
- En revanche, l'espace $\mathbb{R}[X]$ des polynômes à coefficients dans $\mathbb{R}$, ne l'est pas. En effet, les polynômes d'une famille génératrice finie ont leurs degrés bornés et leurs combinaisons linéaires aussi. Une famille finie ne peut donc pas engendrer l'espace tout entier.

1.2 Cardinal des familles libres

Lemme 1. Soit $n \geq 0$. Soit $\mathcal{F} = (e_i)_{1 \leq i \leq n}$ une famille de n vecteurs de l'espace vectoriel E . Soit $\mathcal{F}' = (e'_i)_{1 \leq i \leq n+1}$ une famille de $n+1$ vecteurs combinaisons linéaires des vecteurs de $\mathcal{F}$. Alors, $\mathcal{F}'$ est liée.

Démonstration. On fait une récurrence sur n .

- C'est clair lorsque $n = 0$. La famille $\mathcal{F}$ est vide et la famille $\mathcal{F}' = (0)$ est liée.
- Soit $n \in \mathbb{N}$. Supposons la propriété vraie pour n . On prend $\mathcal{F} = (e_1, \dots, e_{n+1})$, et $\mathcal{F}' = (e'_1, \dots, e'_{n+2})$ avec

$$e'_j = \sum_{i=1}^{n+1} a_{ij} e_i$$

pour $j = 1, \dots, n+2$. Si tous les $a_{n+1,j}$ sont nuls, alors les vecteurs de $\mathcal{F}'$ sont combinaisons de n vecteurs, donc sont liés d'après l'hypothèse de récurrence. Sinon, quitte à renuméroter, on peut supposer que $a_{n+1,n+2} \neq 0$.

On pose pour $1 \leq j \leq n+1$,

$$e''_j = e'_j - \lambda_j e'_{n+2}$$

avec λ_j choisi de sorte que les e''_j soient combinaisons de $e_1, \dots, e_n$. Alors, les e''_j sont liés, et on en tire aisément que les e'_j aussi.

□

Corollaire 2. *Soit E un espace vectoriel ayant une famille génératrice de n éléments. Alors, toutes les familles libres de E sont finies, et ont au plus n éléments.*

Démonstration. Soit $\mathcal{G}$ une famille génératrice de E de cardinal n . Toute famille de cardinal $n + 1$ est une famille dont les vecteurs sont des combinaisons des vecteurs de $\mathcal{G}$, et est donc liée. Toute famille de cardinal au moins égal à $n + 1$ étant une sur-famille d'une famille liée est donc aussi liée. Conclusion, les familles libres sont finies, de cardinal inférieur ou égal à n . □

Corollaire 3. *Toutes les bases d'un espace vectoriel de dimension finie sont finies, de même cardinal.*

Démonstration. Soient $\mathcal{B}$ et $\mathcal{B}'$ deux bases d'un espace de dimension finie. La famille $\mathcal{B}$ est libre et la famille $\mathcal{B}'$ est génératrice, donc $\mathcal{B}$ est finie et $\text{card } \mathcal{B} \leq \text{card } \mathcal{B}'$. De même pour l'autre inégalité en inversant les rôles de $\mathcal{B}$ et $\mathcal{B}'$. □

Il ne reste plus qu'à prouver que tout espace de dimension finie possède des bases ...

1.3 Le théorème de la base incomplète

Proposition 4. *Soit E un espace vectoriel de dimension finie. Soient $\mathcal{F}$ une famille libre de E , et $\mathcal{G}$ une famille génératrice de E . Alors on peut fabriquer une base de E en rajoutant à la famille $\mathcal{F}$ certains des vecteurs de la famille $\mathcal{G}$.*

Démonstration. Rappelons que si E a une famille génératrice finie de n vecteurs, alors toutes les familles libres de E sont finies, de cardinal inférieur ou égal à n .

Quitte à rajouter à $\mathcal{G}$ les vecteurs de $\mathcal{F}$, on peut supposer que « $\mathcal{F} \subset \mathcal{G}$ ». On pose $\mathcal{F} = (e_i)_{i \in I}$ et $\mathcal{G} = (e_i)_{i \in J}$, avec I, J finis et $I \subset J$. Parmi toutes les familles libres $(e_i)_{i \in K}$ avec $I \subset K \subset J$, il y en a une dont le cardinal est maximal (car les cardinaux de telles familles sont majorés par n). On montre facilement que cette famille est une base de E . □

Corollaire 5. THÉORÈME DE LA BASE INCOMPLÈTE

Soit $\mathcal{F}$ une famille libre de l'espace vectoriel de dimension finie E . Alors il existe une sur-famille de $\mathcal{F}$ qui est une base de E .

Démonstration. Il suffit d'appliquer la proposition précédente à $\mathcal{F}$ et à $\mathcal{G} = (x)_{x \in E}$ qui est évidemment génératrice. □

Corollaire 6. THÉORÈME DE LA BASE EXTRAITE

Soit $\mathcal{G}$ une famille génératrice de l'espace vectoriel de dimension finie E . Alors il existe une sous-famille de $\mathcal{G}$ qui est une base de E .

Démonstration. Il suffit d'appliquer la proposition à la famille vide $\mathcal{F} = ()$, qui est libre, et à la famille génératrice $\mathcal{G}$. $\square$

Corollaire 7. *Tout espace vectoriel de dimension finie possède des bases.*

Démonstration. On applique par exemple le théorème de la base incomplète à la famille vide $\mathcal{F} = ()$, qui est libre. Ou, si l'on préfère, on applique le théorème de la base extraite à $\mathcal{G} = (x)_{x \in E}$, qui est génératrice. $\square$

1.4 Bilan

Dans un espace vectoriel de dimension finie, il y a des bases, et elles sont toutes finies, et de même cardinal.

Définition 2. Le cardinal commun des bases d'un espace vectoriel de dimension finie E est appelé la dimension de E , et noté $\dim E$, (ou $\dim_{\mathbb{K}} E$ en cas de confusion possible).

Exemple.

- L'espace nul est de dimension 0. Il a pour seule et unique base la famille vide.
- Une *droite* est un espace vectoriel de dimension 1, un *plan* un espace de dimension 2.
- $\mathbb{K}^n$ est de dimension n puisque sa base canonique a n vecteurs. Et donc toutes ses bases aussi.
- $\mathbb{R}_n[X]$, l'espace des polynômes de degré inférieur ou égal à n , est de dimension $n+1$, puisque sa base canonique $(X^k)_{0 \leq k \leq n}$ est de cardinal $n+1$.
- $\mathcal{M}_{pq}(\mathbb{K})$ est de dimension pq , puisque sa base canonique, constituée des matrices élémentaires, est de cardinal pq .
- $\mathcal{M}_n(\mathbb{K})$, l'espace des matrices carrées $n \times n$ à coefficients dans $\mathbb{K}$, est donc de dimension n^2 .

Remarque 1. Si l'on voit $\mathbb{C}$ comme un $\mathbb{R}$ -espace vectoriel, une base de $\mathbb{C}$ est $(1, i)$. Ainsi,

$$\dim_{\mathbb{R}} \mathbb{C} = 2$$

En revanche, si l'on voit $\mathbb{C}$ comme un $\mathbb{C}$ -espace vectoriel, alors une base de $\mathbb{C}$ est (1) . Ainsi,

$$\dim_{\mathbb{C}} \mathbb{C} = 1$$

Il peut être bizarre de parler de la *droite complexe*, mais c'est ainsi.

Ci-dessous, quelques faits à toujours avoir à l'esprit.

Proposition 8. *Soit E un espace vectoriel de dimension n .*

- *Les familles libres de E sont finies, de cardinal $\leq n$.*
- *Toute famille de vecteurs de E de cardinal $> n$ est liée.*
- *Les familles génératrices de E sont infinies, ou finies de cardinal $\geq n$.*
- *Toute famille de vecteurs de E de cardinal $< n$ est non génératrice.*
- *Une famille de vecteurs de E est une base, si et seulement si elle est libre, et de cardinal n .*
- *Une famille de vecteurs de E est une base, si et seulement si elle est génératrice, et de cardinal n .*

Démonstration.

- Déjà montré.
- C'est la contraposée du premier point.
- Par le théorème de la base extraite, de toute famille génératrice de E on peut extraire une base. Or, les bases sont de cardinal n .
- C'est la contraposée du troisième point.
- Dans un sens c'est évident. Inversement, une famille libre de cardinal n est, par le premier point, libre maximale. C'est donc une base.
- Dans un sens c'est évident. Inversement, une famille génératrice de cardinal n est, par le troisième point, génératrice minimale. C'est donc une base.

□

2 Sev d'un espace de dimension finie

2.1 Dimension d'un sev

Proposition 9. *Soit E un espace de dimension finie n . Soit F un sev de E . Alors*

- *F est de dimension finie p .*
- *$p \leq n$.*
- *$p = n$ si et seulement si $F = E$.*

Démonstration. Soit $\mathcal{F}$ une famille libre de F . C'est aussi une famille libre de E , donc $\mathcal{F}$ possède au plus n éléments. On en déduit les deux premiers points. Si $p = n$, alors F possède une base $\mathcal{B}$ de cardinal n . Mais $\mathcal{B}$ est alors une famille libre de E , de cardinal $n = \dim E$. On en déduit que $\mathcal{B}$ est aussi une base de E . D'où $\langle \mathcal{B} \rangle = E = F$. □

Définition 3. Soit E un espace de dimension finie n . On appelle droite de E , tout sev de E de dimension 1, plan de E tout sev de dimension 2, et hyperplan de E tout sev de dimension $n - 1$.

2.2 Sev supplémentaires

Proposition 10. Soit E un espace de dimension finie n . Soit F un sev de E , de dimension p . Alors, F possède un supplémentaire, et les supplémentaires de F sont tous de dimension $n - p$.

Démonstration. La réunion d'une base de F et d'une base d'un supplémentaire de F est une base de E . D'où la dimension des supplémentaires. Pour l'existence, c'est le théorème de la dimension. $\square$

3 Dimensions classiques

3.1 Image d'un sev par une application linéaire

Proposition 11. Soient E et F deux espaces vectoriels. Soit E' un sev de dimension finie de E . Alors, $f(E')$ est aussi de dimension finie. Précisément :

- $\dim f(E') \leq \dim E'$.
- Si f est injective, on a égalité.

Démonstration. Soit $\mathcal{B}$ une base de E' . Alors, $f(\mathcal{B})$ engendre $f(E')$. Si, de plus, f est injective, $f(\mathcal{B})$ est libre, et est donc une base de $f(E')$. $\square$

3.2 Espaces isomorphes

Proposition 12. Deux espaces vectoriels de dimension finie sur un même corps $\mathbb{K}$ sont isomorphes si et seulement si ils ont la même dimension.

Démonstration. Le sens direct résulte du théorème précédent. Dans l'autre sens, il n'y a qu'à envoyer une base de l'un sur une base de l'autre pour fabriquer un isomorphisme. $\square$

Corollaire 13. Soit E un $\mathbb{K}$ -espace de dimension n . Alors, E est isomorphe à $\mathbb{K}^n$.

Ce résultat affirme que tous les espaces de même dimension sur un corps $\mathbb{K}$ sont « identiques » du point de vue de la structure d'espace vectoriel. Cependant, certains espaces possèdent aussi des propriétés non vectorielles. Par exemple, lorsqu'on travaille sur les polynômes, on dispose du produit de polynômes, de la notion de degré, ..., qui ne sont pas des propriétés vectorielles.

3.3 Somme de sous-espaces vectoriels

Proposition 14. *Soit E un $\mathbb{K}$ -espace de dimension finie. Soient F, G deux sev de E . On a alors*

$$\dim F + G \leq \dim F + \dim G$$

On a égalité si et seulement si F et G sont en somme directe.

Démonstration. Soient $\mathcal{B}'$ une base de F et $\mathcal{B}''$ une base de G . La réunion

$$\mathcal{B} = \mathcal{B}' \cup \mathcal{B}''$$

est une famille génératrice de $F + G$. D'où l'inégalité demandée.

Supposons maintenant que F et G sont en somme directe. Alors $\mathcal{B}$ est une base de $F \oplus G$, donc

$$\dim(F \oplus G) = \text{card } \mathcal{B} = \text{card } \mathcal{B}' + \text{card } \mathcal{B}'' = \dim F + \dim G$$

Inversement, supposons que F et G ne sont pas en somme directe. Il existe un vecteur de $F + G$ qui s'écrit de deux façons différentes comme somme d'un vecteur de F et d'un vecteur de G . On en déduit, en écrivant ces deux vecteurs dans les bases $\mathcal{B}'$ et $\mathcal{B}''$, qu'il existe une combinaison linéaire non triviale des vecteurs de $\mathcal{B}$. La famille $\mathcal{B}$ est génératrice mais pas libre, donc

$$\dim(F + G) < \text{card } \mathcal{B} = \dim F + \dim G$$

□

3.4 Produit d'espaces vectoriels

Proposition 15. *Soient E et F deux $\mathbb{K}$ -espaces de dimension finie. Alors $E \times F$ est aussi de dimension finie, et*

$$\dim E \times F = \dim E + \dim F$$

Remarque 2. Généralisation immédiate à un produit de n espaces vectoriels par récurrence sur n .

Démonstration. Soient $(e_i)_{1 \leq i \leq n}$ une base de E et $(f_j)_{1 \leq j \leq p}$ une base de F . La famille constituée par les couples $(e_i, 0)$, $1 \leq i \leq n$ et $(0, f_j)$, $1 \leq j \leq p$ est une base de $E \times F$. $\square$

Remarque 3. Soient F, G deux sev d'un espace E . Supposons que F et G sont en somme directe. Alors, $F \oplus G$ est isomorphe à $F \times G$ par l'application

$$(x, y) \mapsto x + y$$

Ces deux espaces ont donc bien la même dimension, à savoir $\dim F + \dim G$.

3.5 Dimension des espaces d'applications linéaires

Proposition 16. Soient E et F deux $\mathbb{K}$ -espaces vectoriels de dimensions respectives q et p . Alors, $\mathcal{L}(E, F)$ est de dimension finie, et

$$\dim \mathcal{L}(E, F) = p \times q$$

Démonstration. Soient $\mathcal{B} = (e_j)_{1 \leq j \leq q}$ une base de E , et $\mathcal{B}' = (e'_i)_{1 \leq i \leq p}$ une base de F . On définit, pour $1 \leq i \leq p$ et $1 \leq j \leq q$, l'application $g_{ij} \in \mathcal{L}(E, F)$, par

$$\forall k \in \llbracket 1, q \rrbracket, g_{ij}(e_k) = \delta_{jk} e'_i$$

Soit $f \in \mathcal{L}(E, F)$. Soit $x = \sum_{j=1}^q x_j e_j \in E$. On a

$$f(x) = \sum_{j=1}^q x_j f(e_j)$$

Mais $f(e_j) \in F$, donc est combinaison linéaire des e'_i . Posons

$$f(e_j) = \sum_{i=1}^p a_{ij} e'_i$$

Il vient alors

$$f(x) = \sum_{i=1}^p \sum_{j=1}^q a_{ij} x_j e'_i$$

Par ailleurs,

$$g_{ij}(x) = \sum_{k=1}^q x_k g_{ij}(e_k) = \sum_{k=1}^q x_k \delta_j^k e'_i = x_j e'_i$$

Ainsi,

$$f(x) = \sum_{i=1}^p \sum_{j=1}^q a_{ij} g_{ij}(x)$$

Ceci est vrai pour tout $x \in E$, donc

$$f = \sum_{i=1}^p \sum_{j=1}^q a_{ij} g_{ij}$$

La famille des g_{ij} est donc génératrice de $\mathcal{L}(E, F)$.

Supposons maintenant que

$$\sum_{i=1}^p \sum_{j=1}^q a_{ij} g_{ij} = 0$$

On a donc pour tout $k \in \llbracket 1, q \rrbracket$,

$$0 = \sum_{i=1}^p \sum_{j=1}^q a_{ij} g_{ij}(e_k) = \sum_{i=1}^p \sum_{j=1}^q a_{ij} \delta_j^k e'_i = \sum_{i=1}^p a_{ik} e'_i$$

Les e'_i étant libres, il s'ensuit qu'on a pour tout $i \in \llbracket 1, p \rrbracket$ et tout $k \in \llbracket 1, q \rrbracket$ $a_{ik} = 0$. La famille des g_{ij} est donc libre, c'est ainsi une base de $\mathcal{L}(E, F)$. $\square$

Un certain nombre de cas particuliers sont à noter. Si E est de dimension n , alors

- $\dim \mathcal{L}(E) = n^2$.
- $\dim E^* = \dim E = n$.

Exercice. On prend $E = F$ et $n = p$. Avec les notations du théorème, calculer $g_{i,j} \circ g_{k,l}$ pour $1 \leq i, j, k, l \leq n$.

4 Notion de rang

4.1 Rang d'une famille de vecteurs

Définition 4. Soit $\mathcal{F}$ une famille de vecteurs dans un espace vectoriel E . Si $\langle \mathcal{F} \rangle$ est de dimension finie, on appelle rang de $\mathcal{F}$ l'entier

$$\text{rg } \mathcal{F} = \dim \langle \mathcal{F} \rangle$$

Proposition 17. Soit $\mathcal{F}$ une famille finie de vecteurs dans un espace vectoriel E . On a

- $\text{rg } \mathcal{F} \leq \text{card } \mathcal{F}$.
- $\text{rg } \mathcal{F} = \text{card } \mathcal{F}$ si et seulement si $\mathcal{F}$ est libre.

Démonstration.

- $\mathcal{F}$ engendre $\langle \mathcal{F} \rangle$, d'où l'inégalité.
- On a égalité si et seulement si $\mathcal{F}$ est une base de $\langle \mathcal{F} \rangle$, c'est à dire si et seulement si $\mathcal{F}$ est libre (puisqu'elle est par le premier point toujours génératrice).

□

4.2 Rang d'une application linéaire

Définition 5. Soit $f \in \mathcal{L}(E, F)$. On appelle rang de f la quantité

$$\text{rg } f = \dim(\text{Im } f)$$

Bien entendu, on a des liens entre rang d'une famille et rang d'une application linéaire. Par exemple, si $\mathcal{B}$ est une base de E , alors

$$\text{rg } f = \text{rg } f(\mathcal{B})$$

4.3 Le théorème du rang

Proposition 18. Soient E et F deux espaces de dimension finie sur un même corps. Soit $f \in \mathcal{L}(E, F)$. Alors

- Pour tout supplémentaire G de $\ker f$, l'application $\hat{f} : G \longrightarrow \text{Im } f$ définie, pour tout $x \in G$, par $\hat{f}(x) = f(x)$ est un isomorphisme de G sur $\text{Im } f$.
- $\dim \ker f + \dim \text{Im } f = \dim E$.

Démonstration. $\hat{f}$ est évidemment linéaire. Soit $x \in \ker \hat{f}$. Alors, $f(x) = 0$ et $x \in G$, c'est à dire $x \in \ker f \cap G$. Donc $x = 0$ puisque ces sev sont en somme directe.

Soit maintenant $y \in \text{Im } f$. Il existe $x \in E$ tel que $y = f(x)$. Écrivons $x = x' + x''$ où $x' \in \ker f$ et $x'' \in G$. On a

$$y = f(x) = f(x') + f(x'') = f(x'') = \hat{f}(x'')$$

$\hat{f}$ est donc surjective, c'est un isomorphisme. Conséquence,

$$\dim \text{Im } f = \dim G = \dim E - \dim \ker f$$

□

4.4 Quelques conséquences

Proposition 19. Soit $f : E \rightarrow F$ une application linéaire entre deux espaces **de même dimension, finie**. Alors,

$$f \text{ est injective} \iff f \text{ est surjective} \iff f \text{ est bijective.}$$

Démonstration. C'est immédiat avec le théorème du rang. $\square$

Proposition 20. Soient F et G deux sev d'un espace E . Alors

$$\dim(F + G) = \dim F + \dim G - \dim F \cap G$$

Démonstration. Considérons $f : F \times G \rightarrow F + G$ définie pour $x \in F$ et $y \in G$ par

$$f(x, y) = x + y$$

L'application f est surjective, et son noyau est

$$H = \{(x, -x), x \in F \cap G\}$$

H est clairement un sev de $F \times G$ isomorphe à $F \cap G$ par l'application $x \mapsto (x, -x)$. Il ne reste qu'à appliquer le théorème du rang. $\square$

Exercice. Refaire la démonstration en considérant une base de $F \cap G$ et en la complétant avec le théorème de la base incomplète.

Exercice. Soient $f \in \mathcal{L}(E, F)$ et $g \in \mathcal{L}(F, G)$. Montrer

- $rg(f) \leq \min(\dim E, \dim F)$.
- $rg(g \circ f) \leq \min(rg(f), rg(g))$
- Si g est injective, $rg(g \circ f) = rg(f)$.
- Si f est surjective, $rg(g \circ f) = rg(g)$.

4.5 Calcul pratique d'un rang

Des méthodes du type « pivot de Gauss » permettent de calculer un rang. Nous allons le voir sur un exemple. Supposons donnés

- un espace vectoriel E de dimension p .
- une base $\mathcal{B} = (e_i)_{1 \leq i \leq p}$ de E .
- une famille $\mathcal{F} = (u_j)_{1 \leq j \leq n}$ de n vecteurs de E dont on connaît les coordonnées dans la base $\mathcal{B}$:

$$u_j = \sum_{i=1}^p a_{ij} e_i$$

Proposition 21. *On ne change pas l'espace engendré par $\mathcal{F}$ lorsque :*

- *On échange deux vecteurs de $\mathcal{F}$.*
- *On multiplie un vecteur de $\mathcal{F}$ par un scalaire non nul.*
- *On ajoute à l'un des vecteurs de $\mathcal{F}$ une combinaison linéaire des autres vecteurs.*

Démonstration. Laissée en exercice. $\square$

Ce résultat permet pratiquement de calculer le rang d'une famille de vecteurs. On va prendre deux exemples.

Exemple. Dans $\mathbb{R}^3$, $u_1 = (1, 2, 5)$, $u_2 = (2, 1, 4)$ et $u_3 = (1, -1, -1)$. On écrit les vecteurs en colonnes

$$\begin{array}{c|c|c} u_1 & u_2 & u_3 \\ \hline 1 & 2 & 1 \\ 2 & 1 & -1 \\ 5 & 4 & -1 \end{array} \xrightarrow{(1)} \begin{array}{c|c|c} 1 & 0 & 0 \\ 2 & -3 & -3 \\ 5 & -6 & -6 \end{array} \xrightarrow{(2)} \begin{array}{c|c|c} 1 & 0 & 0 \\ 2 & -3 & 0 \\ 5 & -6 & 0 \end{array} \xrightarrow{(3)} \begin{array}{c|c|c} v_1 & v_2 & v_3 \\ \hline 1 & 0 & 0 \\ 2 & 1 & 0 \\ 5 & 2 & 0 \end{array}$$

avec

1. $C_2 \leftarrow C_2 - 2C_1$, $C_3 \leftarrow C_3 - C_1$
2. $C_3 \leftarrow C_3 - C_2$
3. $C_2 \leftarrow -\frac{1}{3}C_2$

On a donc une famille de rang 2.

Exemple. Dans $\mathbb{R}^5$, $u_1 = (2, 3, -3, 4, 2)$, $u_2 = (3, 6, -2, 5, 9)$, $u_3 = (7, 18, -2, 7, 7)$ et $u_4 = (2, 4, -2, 3, 1)$.

On met le tout dans un tableau.

$$\begin{array}{c|c|c|c} u_1 & u_2 & u_3 & u_4 \\ \hline 2 & 3 & 7 & 2 \\ 3 & 6 & 18 & 4 \\ -3 & -2 & -2 & -2 \\ 4 & 5 & 7 & 3 \\ 2 & 9 & 7 & 1 \end{array}$$

Étape 1 : $C_2 \leftarrow 2C_2 - 3C_1$, $C_3 \leftarrow 2C_3 - 7C_1$, $C_4 \leftarrow C_4 - C_1$.

2	0	0	0
3	3	15	1
-3	5	17	1
4	-2	-14	-1
2	12	0	-1

Étape 2 : $C_3 \leftarrow C_3 - 5C_2$, $C_4 \leftarrow 3C_4 - C_2$.

2	0	0	0
3	3	0	0
-3	5	-8	-2
4	-2	-4	-1
2	12	-60	-15

Étape 3 : $C_4 \leftarrow 4C_4 - C_3$, $C_3 \leftarrow -\frac{1}{4}C_3$.

v_1	v_2	v_3	v_4
2	0	0	0
3	3	0	0
-3	5	2	0
4	-2	1	0
2	12	15	0

Le rang de la famille est donc 3.

Exercice. Déterminer noyau et image de $f \in \mathcal{L}(\mathbb{R}^3)$ définie par $f(e_1) = e_1 + 2e_2 + e_3$, $f(e_2) = -2e_1 + 3e_2 + 3e_3$, $f(e_3) = 4e_1 + e_2 - e_3$.

5 Dual d'un espace vectoriel

5.1 Rappels

Soit E un $\mathbb{K}$ -espace vectoriel. Rappelons que le dual de E est l'espace $E^* = \mathcal{L}(E, \mathbb{K})$ des formes linéaires sur E . Lorsque E est de dimension finie, E et E^* ont la même dimension (et sont donc isomorphes). Nous allons dans cette section étudier quelques aspects de ce que l'on appelle la dualité.

5.2 Formes coordonnées, base duale

Définition 6. Soit E un $\mathbb{K}$ -espace vectoriel de dimension finie n . Soit $\mathcal{B} = (e_1, \dots, e_n)$ une base de E . On appelle *formes linéaires coordonnées* (relatives à la base $\mathcal{B}$) les formes linéaires $e_i^*, i = 1, \dots, n$ définies par

$$\forall i, j \in \{1, \dots, n\}, e_i^*(e_j) = \delta_{ij}$$

Remarque 4. Soit $x = \sum_{j=1}^n x_j e_j$ un vecteur de E . On a

$$e_i^*(x) = \sum_{j=1}^n x_j e_i^*(e_j) = \sum_{j=1}^n x_j \delta_{ij} = x_i$$

Ainsi, la i -ème forme coordonnée renvoie tout simplement la i -ème coordonnée de x dans la base $\mathcal{B}$. On a donc eu raison de l'appeler une « forme coordonnée ».

Notation. On note $\mathcal{B}^*$ la famille des formes coordonnées dans la base $\mathcal{B}$.

Proposition 22. La famille $\mathcal{B}^*$ est une base de E^* . On l'appelle la base duale de la base $\mathcal{B}$.

Démonstration. Comme cette famille est de cardinal n , il suffit de montrer qu'elle est libre. Supposons donc

$$\sum_{i=1}^n \lambda_i e_i^* = 0$$

où les λ_i sont des scalaires. On applique au vecteur e_j :

$$\sum_{i=1}^n \lambda_i e_i^*(e_j) = \lambda_j = 0$$

La famille $\mathcal{B}^*$ est bien libre. $\square$

Exemple. Dans $E = \mathbb{R}^2$, on pose $e_1 = (1, 2)$ et $e_2 = (3, 4)$. Soit $u = (x, y) \in E$. On vérifie aisément que

$$u = \left(-2x + \frac{3}{2}y\right) e_1 + \left(x - \frac{1}{2}y\right) e_2$$

Ainsi,

$$e_1^*(x, y) = -2x + \frac{3}{2}y \text{ et } e_2^*(x, y) = x - \frac{1}{2}y$$

Exemple. Soient $x_0, x_1, \dots, x_n$ $n+1$ scalaires distincts. Soit $\mathcal{B} = (L_j)_{0 \leq j \leq n}$ la famille des polynômes élémentaires de Lagrange associés aux x_k (voir le chapitre sur les polynômes).

La famille $\mathcal{B}$ est une famille de polynômes de $\mathbb{K}_n[X]$ de cardinal $n + 1$. Montrons qu'elle est libre.

Supposons que

$$\sum_{j=0}^n \lambda_j L_j = 0$$

En évaluant en x_j , on obtient $\lambda_j = 0$. La famille $\mathcal{B}$ est donc une base de $\mathbb{K}_n[X]$. Quelle est sa base duale ? La réponse est claire. On doit avoir

$$\forall i, j, e_i^*(L_j) = 0$$

et on sait que $\forall i, j, L_j(x_i) = 0$. Par linéarité, on a pour tout $P \in \mathbb{K}[X]$,

$$e_i^*(P) = P(x_i)$$

5.3 Hyperplans

Définition 7. Soit E un $\mathbb{K}$ -espace vectoriel. Soit H un sev de E . On dit que H est un hyperplan de E lorsque H possède un supplémentaire qui est une droite.

Remarque 5. En dimension finie n on retrouve la notion usuelle : un hyperplan de E est un sev de E de dimension $n - 1$.

Proposition 23. Soit E un $\mathbb{K}$ -espace vectoriel. Soit H un sev de E . Soit $a \in E \setminus H$. H est un hyperplan de E si et seulement si

$$H \oplus \langle a \rangle = E$$

Démonstration. Le sens réciproque est évident, par définition d'un hyperplan.

Inversement, soit H un hyperplan de E . Il existe $b \in E \setminus H$ tel que $H \oplus \langle b \rangle = E$. Soit $a \in E \setminus H$. On a

$$a = h + \lambda b$$

où $h \in H$ et $\lambda \in \mathbb{K}$. Le scalaire λ ne peut pas être nul car $a \notin H$. On en déduit que

$$b = \frac{1}{\lambda}(a - h)$$

Soit maintenant $x \in E$. on a

$$x = h' + \mu b = h'' + \nu a$$

donc $E = H + \langle a \rangle$. Enfin, $H \cap \langle a \rangle$ est clairement réduit à 0, donc

$$E = H \oplus \langle a \rangle$$

□

5.4 Noyau d'une forme linéaire non nulle

Proposition 24. Soit E un $\mathbb{K}$ -espace vectoriel. Soit $\varphi \in E^* \setminus \{0\}$. Le noyau de φ est un hyperplan de E .

Démonstration. Soit $H = \ker \varphi$. Soit $a \in E$ tel que $\varphi(a) \neq 0$. Il en existe puisque φ n'est pas nulle. Soit $x \in E$. On a

$$x = h + \frac{\varphi(x)}{\varphi(a)}a$$

où

$$h = x - \frac{\varphi(x)}{\varphi(a)}a$$

On a facilement $\varphi(h) = 0$, donc $x \in H + \langle a \rangle$. Donc $E = H + \langle a \rangle$. La somme est clairement directe. $\square$

Proposition 25. Soit E un $\mathbb{K}$ -espace vectoriel. Soit H un hyperplan de E . Il existe une forme linéaire φ non nulle telle que $H = \ker \varphi$. De plus, étant données deux formes linéaires non nulles sur E , φ et ψ , on a $\ker \varphi = \ker \psi$ si et seulement si il existe $\mu \in \mathbb{K}^*$ tel que $\psi = \mu\varphi$.

Démonstration. Soit a un vecteur qui n'est pas dans H , de sorte que $H \oplus \langle a \rangle = E$. La forme linéaire définie par $\forall x \in H, \varphi(x) = 0$ et $\varphi(a) = 1$ a bien pour noyau H .

Deux formes linéaires non nulles ont clairement même noyau.

Soient enfin deux formes linéaires non nulles φ et ψ ayant même noyau H . Soit $a \notin H$. Soit $\mu = \frac{\psi(a)}{\varphi(a)}$. On a $\psi(a) = \mu\varphi(a)$ donc, par linéarité, ψ et $\mu\varphi$ coïncident sur $\langle a \rangle$. Elles coïncident aussi sur H , qui est leur noyau. Donc, par linéarité, elles sont égales sur E . $\square$

5.5 Équations d'un hyperplan

Soit E un $\mathbb{K}$ -espace vectoriel de dimension finie n . Soit H un hyperplan de E . L'hyperplan H est le noyau d'une forme linéaire non nulle φ . Il est plus exactement le noyau commun des formes linéaires non nulles multiples de φ . Donnons nous une base $\mathcal{B} = (e_1, \dots, e_n)$ de E . Posons $\varphi(e_i) = a_i$ pour $i = 1, \dots, n$. Le n -uplet $(a_1, \dots, a_n)$ est un n -uplet non nul de scalaires. Pour tout $x = \sum_{i=1}^n x_i e_i \in E$, on a par linéarité $\varphi(x) = \sum_{i=1}^n a_i x_i$. On a alors $x \in H$ si et seulement si $\varphi(x) = 0$, c'est à dire $\sum_{i=1}^n a_i x_i = 0$.

Définition 8. La « formule »

$$(H) \quad \sum_{i=1}^n a_i x_i = 0$$

est appelée une équation cartésienne de l'hyperplan H dans la base $\mathcal{B}$.

Proposition 26.

- *Tout hyperplan possède une équation cartésienne.*
- *Toute équation cartésienne est l'équation d'un hyperplan.*
- *Deux équations cartésiennes sont des équations d'un même hyperplan si et seulement si leurs coefficients sont proportionnels.*

Démonstration. Ce n'est qu'une reformulation de résultats déjà démontrés : un hyperplan est le noyau d'une forme linéaire non nulle. $\square$

Exemple.

- Soit $E = \mathbb{R}^2$, muni de la base canonique. Les hyperplans de E sont les droites vectorielles. Une équation d'une droite vectorielle (dans la base canonique) est donc

$$(D) \quad ax + by = 0$$

où le couple (a, b) est différent du couple $(0, 0)$. Deux équations représentent une même droite si et seulement si leurs coefficients sont proportionnels. remarquons qu'un vecteur directeur de D est n'importe quel vecteur non nul de D , par exemple le vecteur $(-b, a)$.

- Soit $E = \mathbb{R}^3$, muni de la base canonique. Les hyperplans de E sont les plans vectoriels. Une équation d'un plan vectoriel (dans la base canonique) est donc

$$(P) \quad ax + by + cz = 0$$

où le triplet (a, b, c) est différent du triplet $(0, 0, 0)$. Deux équations représentent un même plan si et seulement si leurs coefficients sont proportionnels.

- etc.

5.6 Intersections d'hyperplans

Soit E un $\mathbb{K}$ -espace de dimension finie n . Soit $m \in \mathbb{N}^*, m \leq n$. Toute intersection de m hyperplans de E est évidemment un sev de E . Peut-on affirmer quelque chose sur sa dimension ? Peut-on énoncer une réciproque ?

Proposition 27.

- *Toute intersection de m hyperplans de E est un sev de E de dimension supérieure ou égale à $n - m$.*
- *Tout sev de E de dimension $n - m$ est l'intersection de m hyperplans de E .*

Démonstration.

- Pour $i = 1, \dots, m$, soit $H_i = \ker \varphi_i$ où φ_i est une forme linéaire non nulle sur E . Soit

$$F = H_1 \cap \dots \cap H_m$$

Considérons l'application linéaire $f : E \rightarrow \mathbb{K}^m$ définie par

$$f(x) = (\varphi_1(x), \dots, \varphi_m(x))$$

On a précisément $F = \ker f$. Appliquons le théorème du rang :

$$\dim E = n = \dim F + r$$

où $r = \operatorname{rg} f$. Mais $r \leq \dim \mathbb{K}^m = m$. Donc, $\dim F = n - r \geq n - m$.

- Soit F un sev de E de dimension $n - m$. $\mathcal{B}_0 = (e_1, \dots, e_{n-m})$ une base de F . On complète en $\mathcal{B} = (e_1, \dots, e_n)$ base de E . Soit $x = \sum_{i=1}^n x_i e_i$ un vecteur de E . Alors

$$\begin{aligned} x \in F &\iff x_{n-m+1} = \dots = x_n = 0 \\ &\iff e_{n-m+1}^*(x) = \dots = e_n^*(x) \\ &\iff x \in \ker e_{n-m+1}^* \cap \dots \cap \ker e_n^* \end{aligned}$$

Ainsi,

$$F = \ker e_{n-m+1}^* \cap \dots \cap \ker e_n^*$$

F est donc l'intersection de $n - (n - m + 1) + 1 = m$ hyperplans de E .

□

Chapitre 24

Matrices

1 Vecteurs, applications linéaires et matrices

1.1 Matrice d'un vecteur, d'une famille de vecteurs

Soit E un $\mathbb{K}$ -espace vectoriel de dimension p et $\mathcal{B} = (e_1, \dots, e_p)$ une base de E . Soit $x \in E$. Le vecteur x s'écrit de façon unique $x = \sum_{k=1}^p x_k e_k$ où les $x_k \in \mathbb{K}$ sont les coordonnées de x dans la base $\mathcal{B}$.

Définition 1. On appelle matrice de x dans la base $\mathcal{B}$ la matrice colonne

$$\text{mat}_{\mathcal{B}}(x) = (x_1 \quad \dots \quad x_p)^T$$

Proposition 1. Soit E un $\mathbb{K}$ -espace vectoriel de dimension p muni d'une base $\mathcal{B}$. L'application

$$\text{mat}_{\mathcal{B}} : E \longrightarrow \mathcal{M}_{p1}(\mathbb{K})$$

qui associe à tout x de E sa matrice dans la base $\mathcal{B}$ est un isomorphisme d'espaces vectoriels.

Démonstration. Il est immédiat de vérifier que pour tous vecteurs $x, y \in E$ et tout $\lambda \in \mathbb{K}$, on a

$$\text{mat}_{\mathcal{B}}(x + y) = \text{mat}_{\mathcal{B}}(x) + \text{mat}_{\mathcal{B}}(y)$$

et

$$\text{mat}_{\mathcal{B}}(\lambda x) = \lambda \text{mat}_{\mathcal{B}}(x)$$

De plus, toute matrice de $\mathcal{M}_{p1}(\mathbb{K})$ est la matrice dans la base $\mathcal{B}$ d'un unique vecteur de E . $\square$

Pour toutes les opérations sur les vecteurs, on peut donc travailler indifféremment sur les vecteurs ou leurs matrices, une fois que l'on s'est fixé une base de E .

Définition 2. Plus généralement, étant donnée une famille $\mathcal{F} = (x_1, \dots, x_q)$ de q vecteurs de E , on définit $\text{mat}_{\mathcal{B}}(\mathcal{F})$ comme la matrice à p lignes et q colonnes dont la j ème colonne ($j \in \llbracket 1, q \rrbracket$) est formée des coordonnées du vecteur x_j dans la base $\mathcal{B}$.

1.2 Matrice d'une application linéaire

Définition 3. Soient E et F deux $\mathbb{K}$ -espaces vectoriels de dimensions respectives q et p et munis de bases $\mathcal{B} = (e_1, \dots, e_q)$ et $\mathcal{B}' = (e'_1, \dots, e'_p)$. Soit $f \in \mathcal{L}(E, F)$. On appelle matrice de f dans les bases $\mathcal{B}$ et $\mathcal{B}'$ la matrice de $\mathcal{M}_{pq}(\mathbb{K})$

$$\text{mat}_{\mathcal{B}, \mathcal{B}'}(f) = \text{mat}_{\mathcal{B}'}(f(\mathcal{B}))$$

C'est donc la matrice à p lignes et q colonnes dont la j ème colonne est formée des coordonnées du j ème vecteur de $f(\mathcal{B})$ dans la base $\mathcal{B}'$.

Notation. Lorsque $F = E$ et $\mathcal{B}' = \mathcal{B}$ on note de façon plus légère $\text{mat}_{\mathcal{B}}(f)$.

Exemple. Soit $f = \text{id}_E$ où $\dim E = n$. Alors,

$$\text{mat}_{\mathcal{B}}(f) = I_n$$

où I_n est la *matrice identité d'ordre n* . Tous les coefficients diagonaux de I_n sont égaux à 1 et les autres coefficients sont nuls.

Exemple. Soit $f : \mathbb{R}^3 \rightarrow \mathbb{R}^3$ définie par

$$f(x, y, z) = (x + 4y + 7z, 2x + 5y + 8z, 3x + 6y + 9z)$$

La matrice de f dans la base canonique de $\mathbb{R}^3$ est $\begin{pmatrix} 1 & 4 & 7 \\ 2 & 5 & 8 \\ 3 & 6 & 9 \end{pmatrix}$.

Proposition 2. Soient E et F deux $\mathbb{K}$ -espaces vectoriels de dimensions respectives q et p , munis de bases respectives $\mathcal{B}$ et $\mathcal{B}'$. L'application

$$\text{mat}_{\mathcal{B}, \mathcal{B}'} : \mathcal{L}(E, F) \longrightarrow \mathcal{M}_{pq}(\mathbb{K})$$

qui à chaque application associe sa matrice dans les bases $\mathcal{B}$ et $\mathcal{B}'$ est un isomorphisme d'espaces vectoriels.

Démonstration. La linéarité est une simple vérification.

L'isomorphisme résulte du fait qu'une application linéaire est caractérisée par la donnée des images des vecteurs d'une base. $\square$

Remarque 1. A retenir :

$$\text{mat}_{\mathcal{B}, \mathcal{B}'}(f + g) = \text{mat}_{\mathcal{B}, \mathcal{B}'}(f) + \text{mat}_{\mathcal{B}, \mathcal{B}'}(g)$$

et

$$\text{mat}_{\mathcal{B},\mathcal{B}'}(\lambda f) = \lambda \text{mat}_{\mathcal{B},\mathcal{B}'}(f)$$

Proposition 3. *On se donne trois espaces vectoriels E, F, G , de dimensions respectives r, q et p , munis de bases $\mathcal{B}, \mathcal{B}'$ et $\mathcal{B}''$. Soient $f \in \mathcal{L}(E, F)$ et $g \in \mathcal{L}(F, G)$. Alors*

$$\text{mat}_{\mathcal{B},\mathcal{B}''}(g \circ f) = \text{mat}_{\mathcal{B}',\mathcal{B}''}(g) \times \text{mat}_{\mathcal{B},\mathcal{B}'}(f)$$

Démonstration. C'est une conséquence directe de la définition du produit matriciel. $\square$

Proposition 4. *Soit E un $\mathbb{K}$ -espace vectoriel de dimension finie n . Soit $\mathcal{B}$ une base de E . L'application $\text{mat}_{\mathcal{B}} : \mathcal{L}(E) \rightarrow \mathcal{M}_n(\mathbb{K})$ définie par*

$$A \mapsto \text{mat}_{\mathcal{B}} A$$

est un isomorphisme d'algèbres.

Démonstration. Nous savons déjà que c'est un isomorphisme d'espaces vectoriels. De plus, $\text{mat}_{\mathcal{B}} \text{id}_E = I_n$, qui est l'élément neutre pour la multiplication matricielle, et la proposition précédente nous dit que l'application $\text{mat}_{\mathcal{B}}$ transforme les composées d'applications linéaires en produits de matrices. $\square$

Proposition 5. *Soient E et F deux $\mathbb{K}$ -espaces vectoriel de même dimension finie n . Soient $\mathcal{B}$ une base de E et $\mathcal{B}'$ une base de F . Soit $f \in \mathcal{L}(E, F)$. L'application f est un isomorphisme si et seulement si $\text{mat}_{\mathcal{B},\mathcal{B}'} f \in \text{GL}_n(\mathbb{K})$. On a alors*

$$(\text{mat}_{\mathcal{B},\mathcal{B}'} f)^{-1} = \text{mat}_{\mathcal{B}',\mathcal{B}}(f^{-1})$$

Démonstration. Posons $A = \text{mat}_{\mathcal{B},\mathcal{B}'} f$. Supposons que f est un isomorphisme. Soit $g = f^{-1}$. De $g \circ f = \text{id}_E$, on déduit

$$\text{mat}_{\mathcal{B}',\mathcal{B}} g \times A = I_n$$

De même,

$$A \times \text{mat}_{\mathcal{B}',\mathcal{B}} g = I_n$$

Ainsi, A est inversible et

$$A^{-1} = \text{mat}_{\mathcal{B}',\mathcal{B}} g$$

Inversement, supposons A inversible. Posons $B = A^{-1}$. Soit $g \in \mathcal{L}(F, E)$ telle que $\text{mat}_{\mathcal{B}',\mathcal{B}} g = B$. De $BA = I_n$, on déduit $g \circ f = \text{id}_E$. De même, $f \circ g = \text{id}_F$. Ainsi, f est un isomorphisme et $f^{-1} = g$. $\square$

Proposition 6. Soit E un $\mathbb{K}$ -espace vectoriel de dimension finie n . Soit $\mathcal{B}$ une base de E . L'application $\text{mat}_{\mathcal{B}} : \text{GL}(E) \longrightarrow \text{GL}_n(\mathbb{K})$ définie par

$$A \longmapsto \text{mat}_{\mathcal{B}} A$$

est un isomorphisme de groupes.

Démonstration. C'est une conséquence immédiate des propositions précédentes. $\square$

Proposition 7. Proposition Soit $A \in \mathcal{M}_n(\mathbb{K})$. Les propriétés suivantes sont équivalentes.

- A est inversible.
- A est inversible à gauche.
- A est inversible à droite.

Démonstration. Soit E un $\mathbb{K}$ -espace vectoriel de dimension n . Soit $\mathcal{B}$ une base de E . Soit $f \in \mathcal{L}(E)$ tel que $\text{mat}_{\mathcal{B}} f = A$. Supposons que A est inversible à gauche. Facilement, f est inversible à gauche, c'est à dire injectif. Par un corollaire du théorème du rang, f est bijectif, donc A est inversible. De même, si A est inversible à droite alors A est inversible. $\square$

Remarque 2. Soit $A \in \mathcal{M}_{p,q}(\mathbb{K})$ où $p \neq q$. On a $A = \text{mat}_{\mathcal{B},\mathcal{B}'} f$ où $f \in \mathcal{L}(\mathbb{K}^q, \mathbb{K}^p)$ et $\mathcal{B}, \mathcal{B}'$ sont des bases de $\mathbb{K}^q$ et $\mathbb{K}^p$. Comme $p \neq q$, f n'est pas un isomorphisme donc A n'est pas inversible. En revanche, une matrice non carrée peut fort bien être inversible à gauche, ou inversible à droite (mais pas les deux). Par exemple, soit

$$A = \begin{pmatrix} 1 & 2 & 1 \\ 1 & 0 & 0 \end{pmatrix}$$

Facilement, $f \in \mathcal{L}(\mathbb{R}^3, \mathbb{R}^2)$ est de rang 2. Ainsi, f est surjectif, donc inversible à droite. Un inverse à droite de A (ce n'est pas le seul) est

$$B = \begin{pmatrix} 0 & 1 \\ 0 & 0 \\ 1 & -1 \end{pmatrix}$$

Exercice. Trouver tous les inverses à droite de la matrice A ci-dessus.

1.3 Image d'un vecteur par une application linéaire

Proposition 8. Soit $f \in \mathcal{L}(E, F)$, les deux espaces E et F étant munis de bases respectives $\mathcal{B}$ et $\mathcal{B}'$. Soit $x \in E$. Alors

$$\text{mat}_{\mathcal{B}'}(f(x)) = \text{mat}_{\mathcal{B}, \mathcal{B}'}(f) \times \text{mat}_{\mathcal{B}}(x)$$

Démonstration. Posons $\mathcal{B} = (e_1, \dots, e_q)$, $\mathcal{B}' = (e'_1, \dots, e'_p)$, $A = \text{mat}_{\mathcal{B}, \mathcal{B}'} f$, $X = \text{mat}_{\mathcal{B}} x$ et $X' = \text{mat}_{\mathcal{B}'}(f(x))$. On a, avec des notations évidentes,

$$f(x) = f\left(\sum_{j=1}^q x_j e_j\right) = \sum_{i=1}^p \left(\sum_{j=1}^q A_{ij} x_j\right) e'_i$$

Ainsi, pour tout $i \in \llbracket 1, p \rrbracket$,

$$X'_{i1} = \sum_{j=1}^q A_{ij} x_j = \sum_{j=1}^q A_{ij} X_{j1}$$

Effectivement, $X' = AX$. $\square$

1.4 Application linéaire canoniquement associée à une matrice

Soit $A \in \mathcal{M}_{pq}(\mathbb{K})$. Il existe alors une unique application linéaire $f \in \mathcal{L}(\mathbb{K}^q, \mathbb{K}^p)$ telle que $A = \text{mat}_{\mathcal{B}, \mathcal{B}'} f$, où $\mathcal{B}$ et $\mathcal{B}'$ sont les bases canoniques respectives de $\mathbb{K}^q$ et $\mathbb{K}^p$.

Cette application f , que nous noterons f_A , est appelée *l'application linéaire canoniquement associée à la matrice A* . Remarquons que l'on a

- Pour toutes matrices $A, B \in \mathcal{M}_{pq}(\mathbb{K})$,

$$f_{A+B} = f_A + f_B$$

- Pour toute matrice $A \in \mathcal{M}_{pq}(\mathbb{K})$ et tout $\lambda \in \mathbb{K}$,

$$f_{\lambda A} = \lambda f_A$$

- Pour toutes matrices $A \in \mathcal{M}_{pq}(\mathbb{K})$ et $B \in \mathcal{M}_{qr}(\mathbb{K})$,

$$f_{AB} = f_A \circ f_B$$

Nous éviterons de le faire par la suite, mais il est bon de noter que le vocabulaire relatif aux applications linéaires peut être utilisé dans l'univers des matrices. En identifiant $\mathcal{M}_{p1}(\mathbb{K})$ et $\mathbb{K}^p$, on peut donc parler :

- Du noyau de A .

$$\ker A = \ker f_A = \{X \in \mathcal{M}_{q1}(\mathbb{K}), AX = 0\} \subset \mathbb{K}^q$$

- De l'image de A .

$$\text{Im } A = \text{Im } f_A = \{AX, X \in \mathcal{M}_{q1}(\mathbb{K})\} \subset \mathbb{K}^p$$

L'image de A est l'espace engendré par les colonnes de A .

- Du rang de A .

$$\text{rg } A = \text{rg } f_A$$

C'est le rang de la famille des colonnes de A .

2 Changements de base

2.1 Matrices de passage

Définition 4. Soit E un $\mathbb{K}$ -espace vectoriel de dimension finie. Soient $\mathcal{B}$ et $\mathcal{B}'$ deux bases de E . On appelle matrice de passage de $\mathcal{B}$ à $\mathcal{B}'$ la matrice

$$P_{\mathcal{B}}^{\mathcal{B}'} = \text{mat}_{\mathcal{B}}(\mathcal{B}') = \text{mat}_{\mathcal{B}', \mathcal{B}}(id_E)$$

Proposition 9. Soit E un $\mathbb{K}$ -espace vectoriel de dimension finie. Soient $\mathcal{B}$ et $\mathcal{B}'$ deux bases de E . La matrice $P_{\mathcal{B}}^{\mathcal{B}'}$ est inversible et son inverse est $P_{\mathcal{B}'}^{\mathcal{B}}$.

Démonstration. On a

$$I_n = \text{mat}_{\mathcal{B}, \mathcal{B}} id = \text{mat}_{\mathcal{B}, \mathcal{B}}(id \circ id) = \text{mat}_{\mathcal{B}', \mathcal{B}} id \times \text{mat}_{\mathcal{B}, \mathcal{B}'} id = P_{\mathcal{B}}^{\mathcal{B}'} P_{\mathcal{B}'}^{\mathcal{B}}$$

De même, $P_{\mathcal{B}'}^{\mathcal{B}} P_{\mathcal{B}}^{\mathcal{B}'} = I$. $\square$

2.2 Changement de base pour un vecteur

Proposition 10. Soit E un $\mathbb{K}$ -espace vectoriel muni de deux bases $\mathcal{B}$ et $\mathcal{B}'$. Soit $x \in E$. Posons $P = P_{\mathcal{B}}^{\mathcal{B}'}$, $X = \text{mat}_{\mathcal{B}}(x)$ et $X' = \text{mat}_{\mathcal{B}'}(x)$. Alors

$$X = PX'$$

Démonstration. On écrit que $x = id_E(x)$, où id_E part de E muni de la base $\mathcal{B}'$ et arrive dans E muni de la base $\mathcal{B}$. $\square$

2.3 Changement de bases pour une application linéaire

Proposition 11. Soient E et F deux $\mathbb{K}$ -espaces vectoriels. Soient $\mathcal{B}_1$ et $\mathcal{B}'_1$ deux bases de E , et $\mathcal{B}_2$ et $\mathcal{B}'_2$ deux bases de F . Soit $f \in \mathcal{L}(E, F)$. Notons $M = \text{mat}_{\mathcal{B}_1, \mathcal{B}_2}(f)$ et $M' = \text{mat}_{\mathcal{B}'_1, \mathcal{B}'_2}(f)$. Notons également $P_1 = P_{\mathcal{B}_1}^{\mathcal{B}'_1}$ et $P_2 = P_{\mathcal{B}_2}^{\mathcal{B}'_2}$. Alors

$$M' = P_2^{-1} M P_1$$

Démonstration. On écrit $f = \text{id}_F \circ f \circ \text{id}_E$ avec $\text{id}_E : E, \mathcal{B}'_1 \rightarrow E, \mathcal{B}_1$, $f : E, \mathcal{B}_1 \rightarrow F, \mathcal{B}_2$ et $\text{id}_F : F, \mathcal{B}_2 \rightarrow F, \mathcal{B}'_2$. $\square$

Remarque 3. Un cas particulier très important est celui où $E = F$, $\mathcal{B}_1 = \mathcal{B}_2 = \mathcal{B}$, et $\mathcal{B}'_1 = \mathcal{B}'_2 = \mathcal{B}'$. On a alors la formule plus simple

$$M' = P^{-1} M P$$

où $P = P_{\mathcal{B}}^{\mathcal{B}'}$.

3 Rang

3.1 Rang d'une matrice

Rappelons la définition du rang d'une matrice.

Définition 5. Soit $A \in \mathcal{M}_{\mathbb{K}}(p, q)$. On appelle rang de A le rang de la famille des colonnes de A , identifiées à des vecteurs de $\mathbb{K}^p$.

Exemple.

$$\text{rg} \begin{pmatrix} 1 & 2 & 3 \\ 0 & 2 & 4 \end{pmatrix} = 2.$$

$$\text{rg} \begin{pmatrix} 1 & 4 & 7 \\ 2 & 5 & 8 \\ 3 & 6 & 9 \end{pmatrix} = 2.$$

$$\text{rg} \begin{pmatrix} 1 & 4 & 7 \\ 2 & 5 & 8 \\ 3 & 6 & 10 \end{pmatrix} = 3.$$

$$\text{rg} I_n = n, \text{rg} 0 = 0.$$

Remarque 4.

- Si $f \in \mathcal{L}(E, F)$, le rang de f est égal au rang de sa matrice dans n'importe quel couple de bases.
- Une matrice $A \in \mathcal{M}_n(\mathbb{K})$ est inversible si et seulement si $rg A = n$.

3.2 Matrice canonique d'une application linéaire

Proposition 12. Soit $f \in \mathcal{L}(E, F)$ où E et F sont deux espaces vectoriels de dimensions q et p . On suppose que $rg f = r$. Il existe alors une base $\mathcal{B}$ de E et une base $\mathcal{B}'$ de F telles que

$$\text{mat}_{\mathcal{B}, \mathcal{B}'}(f) = \begin{pmatrix} I_r & 0_{r, q-r} \\ 0_{p-r, r} & 0_{p-r, q-r} \end{pmatrix}$$

Notation. On note dorénavant $J_{p,q,r}$ la matrice ci-dessus. A noter que son rang est égal à r . A noter également, la forme de $J_{p,q,r}$ lorsque r est égal à p et/ou q (quand f est une surjection et/ou une injection).

Démonstration. $\ker f$ est de dimension $q - r$. Soit G un supplémentaire de $\ker f$. Soit $\mathcal{B}_1 = (e_1, \dots, e_r)$ une base de G , et $\mathcal{B}_2 = (e_{r+1}, \dots, e_q)$ une base de $\ker f$. La réunion $\mathcal{B}$ de ces deux familles est une base de E . Soit $\mathcal{B}'_1 = (e'_1, \dots, e'_r)$ l'image par f de la famille $\mathcal{B}_1$. D'après le théorème du rang, cette famille est libre. On la complète en une base $\mathcal{B}' = (e'_1, \dots, e'_p)$ de F . La matrice de f dans les bases $\mathcal{B}$ et $\mathcal{B}'$ a la forme cherchée. $\square$

3.3 Équivalence des matrices

Définition 6. Soient $A, B \in \mathcal{M}_{pq}(\mathbb{K})$ deux matrices de même taille. On dit que A est équivalente à B lorsqu'il existe $Q \in GL_q(\mathbb{K})$ et $P \in GL_p(\mathbb{K})$ telles que

$$B = P^{-1}AQ$$

Proposition 13. L'équivalence des matrices est une relation d'équivalence.

Démonstration. Exercice. $\square$

Proposition 14. Deux matrices sont équivalentes si et seulement si elles représentent une même application linéaire dans deux couples de bases.

Démonstration. Supposons que $B = P^{-1}AQ$. On interprète A comme $\text{mat}_{\mathcal{B}, \mathcal{B}'}(f)$ où $f \in \mathcal{L}(\mathbb{K}^q, \mathbb{K}^p)$ (par exemple). On interprète ensuite P comme $P_{\mathcal{B}'}^{\mathcal{B}'_1}$ et Q comme $P_{\mathcal{B}}^{\mathcal{B}_1}$. Il en

résulte que $B = \text{mat}_{\mathcal{B}_1, \mathcal{B}'_1}(f)$. La réciproque provient des formules de changement de base. $\square$

Proposition 15. *Deux matrices sont équivalentes si et seulement si elles ont le même rang.*

Démonstration. Supposons A et B équivalentes. Elles sont donc les matrices d'une même application linéaire f dans deux couples de bases différentes. Ainsi,

$$\text{rg} B = \text{rg} f = \text{rg} A$$

Inversement, si A et B ont pour rang r , alors elles sont toutes deux équivalentes à $J_{p,q,r}$, donc elles sont équivalentes. $\square$

3.4 Rang et transposition

Proposition 16. *Soit $A \in \mathcal{M}_{p,q}(\mathbb{K})$. Alors, $\text{rg} A = \text{rg} A^T$.*

Remarque 5. Ceci implique en particulier que pour chercher le rang d'une famille de vecteurs ou d'une matrice, on peut tout aussi bien opérer sur les lignes que sur les colonnes. En revanche, bien noter que des opérations sur les lignes détruisent toute possibilité de trouver une image ou un noyau.

Démonstration. Soit $r = \text{rg} A$. La matrice A est équivalente à $J_{p,q,r}$. On voit alors facilement que A^T équivaut à $J_{q,p,r}$ et est donc de même rang que A . $\square$

3.5 Rang et matrices extraites

Définition 7. Soit $M \in \mathcal{M}_{p,q}(\mathbb{K})$. On appelle matrice extraite de M toute matrice M' qui est la restriction de M à $I \times J$, I et J étant des sous-ensembles non vides respectivement de $[1, p]$ et $[1, q]$.

En pratique, extraire une matrice de M c'est rayer des lignes et des colonnes de M .

Exemple. $M = \begin{pmatrix} a & c \\ b & d \end{pmatrix}$ possède 9 matrices extraites (les écrire).

Exercice. Soit $M \in \mathcal{M}_{p,q}(\mathbb{K})$. Montrer qu'on peut extraire $(2^p - 1)(2^q - 1)$ matrices de M .

Proposition 17. *Soit $A \in \mathcal{M}_{p,q}(\mathbb{K})$ une matrice de rang r . Alors :*

- Toutes les matrices extraites de A sont de rang inférieur ou égal à r .
- Il existe une matrice carrée de taille $r \times r$, extraite de A , qui est inversible.

Démonstration. Soit $r = \text{rg} A$. C'est le rang de la famille des colonnes de A . Si l'on extrait $r + 1$ colonnes ou plus de A , elles seront donc liées. Le fait de rayer des lignes conservera la dépendance linéaire des colonnes. Il reste à extraire de A une matrice $r \times r$ inversible. On sait tout d'abord qu'il existe r colonnes de A formant une famille libre. On les extrait. La matrice $p \times r$ ainsi extraite est toujours de rang r . Il y a donc r lignes de cette matrice qui forment une famille libre. On les extrait, et on obtient ce que l'on cherche. $\square$

4 Trace

4.1 Trace d'une matrice carrée

Définition 8. Soit $A \in \mathcal{M}_n(\mathbb{K})$. On appelle trace de A le scalaire

$$\text{Tr} A = \sum_{k=1}^n A_{kk}$$

Proposition 18. Soient A et B deux matrices carrées de même taille, et λ un scalaire. Alors

- $\text{Tr}(A + B) = \text{Tr} A + \text{Tr} B$
- $\text{Tr}(\lambda A) = \lambda \text{Tr} A$
- $\text{Tr}(A^T) = \text{Tr} A$
- $\text{Tr}(AB) = \text{Tr}(BA)$

Démonstration. Montrons la dernière égalité. On a

$$\text{Tr}(AB) = \sum_i (AB)_{ii} = \sum_{i,k} A_{ik} B_{ki}$$

De même pour $\text{Tr}(BA)$. $\square$

4.2 Trace d'un endomorphisme

Proposition 19. *Soit f un endomorphisme d'un espace vectoriel E . La quantité $\text{Tr}(\text{mat}_{\mathcal{B}}(f))$ ne dépend pas de la base $\mathcal{B}$ choisie. On l'appelle la trace de f , et on la note $\text{Tr}f$.*

Démonstration. Si A et A' sont les matrices de f dans deux bases de E , on a alors $A' = P^{-1}AP$ où P est une matrice de passage. Ainsi,

$$\text{Tr}A' = \text{Tr}(P^{-1}AP) = \text{Tr}(APP^{-1}) = \text{Tr}A$$

□

Exemple. La trace de id_E est égale à la dimension de E .

Exemple. Soit p le projecteur sur F (de dimension k) parallèlement à G , où F et G sont deux sev supplémentaires de E . Soit $\mathcal{B}$ une base de E formée de la réunion d'une base de F et d'une base de G . La matrice de p dans la base $\mathcal{B}$ est particulièrement simple. Elle est diagonale, les k premiers éléments de la diagonale valent 1, les autres valent 0. Ainsi, la trace de p est égale à $k = \dim F$, c'est à dire au rang de p .

Exercice. Quelle est la trace d'une symétrie ?

5 Similitude

Définition 9. Deux matrices carrées A et B sont semblables lorsqu'il existe une matrice P inversible telle que $B = P^{-1}AP$.

Remarque 6. La relation de similitude est aussi une relation d'équivalence, et deux matrices semblables sont a fortiori équivalentes et ont donc le même rang.

Proposition 20. *Deux matrices sont semblables si et seulement si elles représentent un même endomorphisme dans deux bases.*

Démonstration. Identique à ce qui a été fait pour les matrices équivalentes. La seule différence, c'est qu'ici $P = Q$. □

Proposition 21. *Deux matrices semblables ont même trace. La réciproque est fausse.*

Démonstration. Soient A, B deux matrices carrées semblables. Il existe une matrice inversible P telle que $B = P^{-1}AP$. On a alors

$$\operatorname{tr} B = \operatorname{tr} P^{-1}AP = \operatorname{tr} P P^{-1}A = \operatorname{tr} A$$

□

Il y a en fait bien plus à raconter, mais nous n'en dirons pas plus pour l'instant. Finissons ce chapitre par un exercice.

Exercice. Soit $A = \begin{pmatrix} 0 & 1 \\ 1 & 1 \end{pmatrix}$. Soit f l'endomorphisme de $\mathbb{R}^2$ dont la matrice dans la base canonique $\mathcal{B}$ est A .

1. Soient φ et $\hat{\varphi}$ les racines de l'équation $X^2 - X - 1 = 0$. Montrons qu'il existe deux vecteurs non nuls e_1 et e_2 de $\mathbb{R}^2$ tels que $f(e_1) = \varphi e_1$ et $f(e_2) = \hat{\varphi} e_2$.
2. Montrer que $\mathcal{B}' = (e'_1, e'_2)$ est une base de $\mathbb{R}^2$.
3. En déduire que $A = PDP^{-1}$ où P est une matrice inversible (à calculer explicitement) et $D = \begin{pmatrix} \varphi & 0 \\ 0 & \hat{\varphi} \end{pmatrix}$.
4. Calculer A^n pour tout entier naturel n .
5. On pose $F_0 = 0$, $F_1 = 1$, et, pour tout entier n , $F_{n+2} = F_{n+1} + F_n$. On pose également $Y_n = (F_n \ F_{n+1})^T$.
 - (a) Montrer que pour tout $n \geq 0$, $Y_{n+1} = AY_n$.
 - (b) En déduire que $Y_n = A^n Y_0$.
 - (c) Déduire de ce qui précède une expression de F_n en fonction de n .

Chapitre 25

Groupes Symétriques

1 Compléments sur les groupes finis

Ce paragraphe éclaire ce que nous ferons dans les sections suivantes. Tous les résultats sont énoncés pour des groupes « multiplicatifs ». Il convient bien entendu d'adapter dans le cas de groupes dont l'opération est notée additivement.

1.1 Groupes cycliques

Définition 1. Soit $(G, \times)$ un groupe. On dit que G est cyclique lorsque G est fini et qu'il existe $a \in G$ tel que

$$G = \{a^k, k \in \mathbb{Z}\}$$

Notation. On note $G = \langle a \rangle$. Le cardinal de G est aussi appelé son *ordre*. On le note $|G|$.

Exemple. $(\mathcal{U}_n, \times)$, $(\mathbb{Z}/n\mathbb{Z}, +)$ pour tout $n \geq 1$.

Proposition 1. Soit $G = \langle a \rangle$ un groupe cyclique d'ordre n . On a

- $G = \{e, a, a^2, \dots, a^{n-1}\}$.
- $\forall i, j \in \mathbb{Z}, a^i = a^j \Leftrightarrow j - i \in n\mathbb{Z}$.
- $\forall k \in \mathbb{Z}, a^k = e \Leftrightarrow k \in n\mathbb{Z}$.

Démonstration. Notons $E = \{m \in \mathbb{Z}, a^m = e\}$. E est clairement un sous-groupe de $\mathbb{Z}$. Il existe donc $m \in \mathbb{N}$ tel que $E = m\mathbb{Z}$.

Comme le groupe G est fini, on peut trouver $i, j \in \mathbb{Z}, i < j$ tels que $a^i = a^j$. On a alors $a^{j-i} = e$: il existe $m \in \mathbb{N}^*$ tel que $a^m = e$. Ainsi, $E \neq \{0\}$ et $m > 0$.

Soient maintenant $i, j \in \mathbb{Z}$. On a $a^i = a^j$ si et seulement si $a^{j-i} = e$, c'est à dire $j - i \in E$. Enfin, une division euclidienne par m montre facilement que tout élément de G est de la forme a^k avec $k \in \llbracket 0, m-1 \rrbracket$. D'où $G = \{e, a, \dots, a^{m-1}\}$.

Il reste à vérifier que toutes ces puissances sont distinctes deux à deux pour en déduire que $|G| = m$, donc $m = n$. $\square$

1.2 Ordre d'un élément, ordre d'un sous-groupe

Définition 2. Soit G un groupe. Soit $a \in G$. Si le sous-groupe de G $\langle a \rangle$ est fini, on appelle ordre de a l'entier

$$|a| = |\langle a \rangle|$$

Proposition 2.

- $|a|$ est le plus petit entier naturel k tel que $a^k = e$.
- $\forall i, j \in \mathbb{Z}, a^i = a^j \Leftrightarrow j - i \in n\mathbb{Z}$.
- $\forall k \in \mathbb{Z}, a^k = e \Leftrightarrow k \in n\mathbb{Z}$.

Démonstration. Voir la démonstration faite pour les groupes cycliques : $\langle a \rangle$ est un groupe cyclique, engendré par a . $\square$

Exemple. Ordre des éléments de $\mathcal{U}_6$, de $\mathbb{Z}/12\mathbb{Z}$.

Exercice. Soit G un groupe. Soient $a, b \in G$ tels que $ab = ba$. Alors $|ab|$ divise $|a| \vee |b|$. La relation peut être stricte.

1.3 Le théorème de Lagrange

Proposition 3. Soit G un groupe fini. Soit H un sous-groupe de G . Alors, $|H|$ divise $|G|$.

Démonstration. Soit $\sim$ la relation binaire sur G définie par

$$\forall x, y \in G, x \sim y \Leftrightarrow xy^{-1} \in H$$

On montre facilement que l'on a là une relation d'équivalence. Notons, pour tout $x \in G$, C_x la classe de x . On remarque que $H = C_e$. Soit $C = C_g$ une classe modulo $\sim$. On a pour tout $x \in H, xg \in C$, puisque $(xg)g^{-1} = x \in H$. On définit ainsi une application $\varphi : H \rightarrow C$ en posant $\varphi(x) = xg$. On voit facilement que φ est une bijection : ainsi, H et C ont le même cardinal. Toutes les classes modulo $\sim$ sont de même cardinal. Notons $C_1, C_2, \dots, C_p$ les classes distinctes modulo $\sim$. On a

$$G = \bigcup_{k=1}^p C_k$$

et les classes sont disjointes. Donc

$$|G| = \sum_{k=1}^p |C_k| = p|H|$$

l'ordre de H divise l'ordre de G et le quotient de ces ordres est le nombre de classes modulo $\sim$. $\square$

Exercice. Montrer qu'un groupe d'ordre premier est cyclique.

Exemple. Sous-groupes de $\mathfrak{D}_4$, le groupe des isométries du plan laissant invariant l'ensemble des sommets d'un carré.

2 Notion de permutation

2.1 Permutations d'un ensemble

Définition 3. Soit E un ensemble non vide. On appelle permutation de E toute bijection de E dans E . On note $\mathfrak{S}(E)$ l'ensemble des permutations de E .

Proposition 4. Soit E un ensemble. L'ensemble $\mathfrak{S}(E)$ est un groupe pour la composition des applications. Ce groupe est non abélien dès que E possède au moins 3 éléments.

Démonstration. La composée de deux bijections, la réciproque d'une bijection, sont des bijections. La composition des applications est associative. L'identité de E est le neutre pour la composition.

En ce qui concerne la non-commutativité, soient a, b, c trois éléments distincts de E . Soient $f, g \in \mathfrak{S}(E)$ définies comme suit : $f(a) = b, f(b) = c, f(c) = a$ et pour tout $x \neq a, b, c$, $f(x) = x$. $g(a) = b, g(b) = a$ et pour tout $x \neq a, b$, $g(x) = x$. On a alors $g \circ f(a) = g(b) = b$ et $f \circ g(a) = f(b) = c$. Donc, $f \circ g \neq g \circ f$. $\square$

Définition 4. Le groupe $\mathfrak{S}(E)$ est appelé le *groupe symétrique* de E . Ses éléments sont appelés les *permutations* de E .

Proposition 5. Soient E et F deux ensembles. On suppose qu'il existe une bijection de E sur F . Alors les groupes $\mathfrak{S}(E)$ et $\mathfrak{S}(F)$ sont isomorphes.

Démonstration. Soit $\varphi : E \longrightarrow F$ une bijection. Soit $F : \mathfrak{S}(E) \longrightarrow \mathfrak{S}(F)$ définie, pour toute $f \in \mathfrak{S}(E)$, par

$$F(f) = \varphi \circ f \circ \varphi^{-1}$$

On vérifie aisément que F est une bijection. De plus, pour toutes $f, g \in \mathfrak{S}(E)$,

$$F(g \circ f) = \varphi \circ (g \circ f) \circ \varphi^{-1} = (\varphi \circ g \circ \varphi^{-1}) \circ (\varphi \circ f \circ \varphi^{-1}) = F(g) \circ F(f)$$

Ainsi, les groupes $\mathfrak{S}(E)$ et $\mathfrak{S}(F)$ sont isomorphes. $\square$

Proposition 6. Soit E un ensemble fini de cardinal $n \in \mathbb{N}$. Le groupe $(\mathfrak{S}(E), \circ)$ est un groupe fini, de cardinal $n!$.

Démonstration. Nous démontrerons de façon rigoureuse ce résultat dans le cours de dénombrement. Contentons nous pour l'instant de l'argument suivant. Notons $E = \{x_1, \dots, x_n\}$. Une application $f \in \mathfrak{S}(E)$ est caractérisée par

- La valeur de $f(x_1)$: n possibilités.
- La valeur de $f(x_2)$: $n - 1$ possibilités, car f est injective, donc $f(x_2) \neq f(x_1)$.
- La valeur de $f(x_3)$: $n - 2$ possibilités, car f est injective, donc $f(x_3) \neq f(x_1)$ et $f(x_3) \neq f(x_2)$.
- ...
- La valeur de $f(x_n)$: 1 possibilité, car f est injective, donc $f(x_n) \neq f(x_1), \dots, f(x_n) \neq f(x_{n-1})$.

Le nombre total de bijections de E sur E est donc

$$n \times (n - 1) \times (n - 2) \dots \times 1 = n!$$

□

On ne s'occupera dans ce qui suit que des permutations de l'ensemble $\llbracket 1, n \rrbracket$.

Notation.

- On note $\mathfrak{S}_n = \mathfrak{S}(\llbracket 1, n \rrbracket)$.
- Une permutation $\sigma \in \mathfrak{S}_n$ est notée par le tableau de ses valeurs :

$$\sigma = \begin{pmatrix} 1 & 2 & \dots & n \\ \sigma(1) & \sigma(2) & \dots & \sigma(n) \end{pmatrix}$$

- L'usage est de noter *multiplicativement* la composition des permutations. Attention, donc, aux confusions.

2.2 Exemples

On a $\mathfrak{S}_1 = \{id\}$, $\mathfrak{S}_2 = \{id, \tau\}$ où $\tau(1) = 2, \tau(2) = 1$. Regardons d'un peu plus près $\mathfrak{S}_3 = \{id, \rho, \rho^2, \tau_{12}, \tau_{13}, \tau_{23}\}$ où $\rho = \begin{pmatrix} 1 & 2 & 3 \\ 2 & 3 & 1 \end{pmatrix}$ et τ_{ij} échange i et j et laisse le troisième élément invariant. Voici la table de multiplication dans $\mathfrak{S}_3$:

	id	ρ	ρ^2	τ_{12}	τ_{23}	τ_{13}
id	id	ρ	ρ^2	τ_{12}	τ_{23}	τ_{13}
ρ	ρ	ρ^2	id	τ_{13}	τ_{12}	τ_{23}
ρ^2	ρ^2	id	ρ	τ_{23}	τ_{13}	τ_{12}
τ_{12}	τ_{12}	τ_{23}	τ_{13}	id	ρ	ρ^2
τ_{23}	τ_{23}	τ_{13}	τ_{12}	ρ^2	id	ρ
τ_{13}	τ_{13}	τ_{12}	τ_{23}	ρ	ρ^2	id

Exercice. Trouver les sous-groupes de $\mathfrak{S}_3$.

2.3 Ordre d'une permutation

Soit $\sigma \in \mathfrak{S}_n$. Considérons l'ensemble

$$E = \{\sigma^k, k \in \mathbb{Z}\}$$

L'ensemble E est inclus dans $\mathfrak{S}_n$, c'est donc un ensemble fini. Cependant, il existe une infinité d'entiers relatifs. Il existe donc $i, j \in \mathbb{Z}$, $i < j$, tels que $\sigma^i = \sigma^j$. Soit $k = j - i \geq 1$. On a alors

$$\sigma^k = id$$

Définition 5. Soit $\sigma \in \mathfrak{S}_n$. L'ordre de σ est le plus petit entier $k \geq 1$ tel que $\sigma^k = id$.

Exercice. Quel est l'ordre des éléments de $\mathfrak{S}_3$?

Proposition 7. Soit $\sigma \in \mathfrak{S}_n$ une permutation d'ordre ω . Pour tout $k \in \mathbb{Z}$, on a $\sigma^k = id$ si et seulement si $\omega \mid k$. Dit autrement,

$$\{k \in \mathbb{Z}, \sigma^k = id\} = \omega\mathbb{Z}$$

Démonstration. Soit $k = c\omega$ un multiple de ω . On a

$$\sigma^k = (\sigma^\omega)^c = id^c = id$$

Ainsi,

$$\{k \in \mathbb{Z}, \sigma^k = id\} \supset \omega\mathbb{Z}$$

Inversement, soit $k \in \mathbb{Z}$ tel que $\sigma^k = id$. Par une division euclidienne, on écrit $k = q\omega + r$ où $q, r \in \mathbb{Z}$ et $0 \leq r < \omega$. De là,

$$id = \sigma^k = (\sigma^\omega)^q \sigma^r = \sigma^r$$

Mais ω est le plus petit entier ≥ 1 tel que $\sigma^\omega = id$. Ainsi, on n'a pas $r \geq 1$, donc $r = 0$. On en déduit que $k = q\omega$, d'où

$$\{k \in \mathbb{Z}, \sigma^k = id\} \subset \omega\mathbb{Z}$$

□

3 Orbites

3.1 Notion d'orbite

Définition 6. Soit $a \in \llbracket 1, n \rrbracket$. Soit $\sigma \in \mathfrak{S}_n$. On appelle orbite de a suivant σ l'ensemble

$$\mathcal{O}_\sigma(a) = \{\sigma^k(a), k \in \mathbb{Z}\}$$

Proposition 8. Les orbites suivant une permutation $\sigma \in \mathfrak{S}_n$ forment une **partition** de $\llbracket 1, n \rrbracket$: les orbites sont non vides, disjointes, et leur réunion est $\llbracket 1, n \rrbracket$.

Démonstration. Soit $a \in \llbracket 1, n \rrbracket$. On a $a = \sigma^0(a)$, donc $a \in \mathcal{O}_\sigma(a)$. Ainsi, les orbites sont non vides et leur réunion est $\llbracket 1, n \rrbracket$.

Prenons maintenant deux orbites $\mathcal{O} = \mathcal{O}_\sigma(a)$ et $\mathcal{O}' = \mathcal{O}_\sigma(a')$. Supposons que $\mathcal{O} \cap \mathcal{O}' \neq \emptyset$. Soit $b \in \mathcal{O} \cap \mathcal{O}'$. Il existe donc $k, k' \in \mathbb{Z}$ tels que $b = \sigma^k(a) = \sigma^{k'}(a')$. Donc $a' = \sigma^{k-k'}(a)$. On en déduit facilement que $\mathcal{O}' \subset \mathcal{O}$. De même, $\mathcal{O} \subset \mathcal{O}'$ donc $\mathcal{O}' = \mathcal{O}$. $\square$

3.2 Étude des orbites

Proposition 9. Soit $\sigma \in \mathfrak{S}_n$. Soit $a \in \llbracket 1, n \rrbracket$. Soit p le cardinal de l'orbite de a . Alors

- Pour tout entier relatif k , on a $\sigma^k(a) = a$ si et seulement si $k \in p\mathbb{Z}$.
- $\mathcal{O}_\sigma(a) = \{\sigma^k(a), k \in \llbracket 0, p-1 \rrbracket\}$.

Démonstration. Soit $E = \{k \in \mathbb{Z}, \sigma^k(a) = a\}$. On vérifie facilement que E est un sous-groupe de $\mathbb{Z}$. Il existe donc $m \in \mathbb{N}$ tel que $E = m\mathbb{Z}$.

L'entier m est non nul. En effet, l'orbite de a est un ensemble fini, il existe donc deux entiers distincts i et j tels que $\sigma^i(a) = \sigma^j(a)$. Donc, $0 \neq i - j \in E$.

Soit maintenant $k \in \mathbb{Z}$. On a, par une division euclidienne, $k = mq + r$, $0 \leq r < m$. De là

$$\sigma^k(a) = \sigma^r(\sigma^{mq}(a)) = \sigma^r(a)$$

Donc, $\mathcal{O}_\sigma(a) = \{a, \sigma(a), \dots, \sigma^{m-1}(a)\}$.

De plus, si $i, j \in \llbracket 0, m-1 \rrbracket$, $i \neq j$ et $\sigma^i(a) = \sigma^j(a)$, alors $i - j \in E = m\mathbb{Z}$. Mais $|i - j| \leq m-1$. D'où $i = j$. Donc, $\mathcal{O}_\sigma(a)$ a exactement m éléments, donc $m = p$. $\square$

3.3 Cycles, Transpositions

Définition 7.

- Un *cycle* est une permutation ayant exactement une orbite non réduite à un point.
- L'orbite en question est appelée le *support* du cycle.
- La *longueur* du cycle est le cardinal de son support.
- Une *transposition* est un cycle de longueur 2.

Notation. Si σ est un cycle de longueur ℓ , et a appartient à l'unique orbite de σ non réduite à un point, on note

$$\sigma = (a \ \sigma(a) \ \sigma^2(a) \ \dots \ \sigma^{\ell-1}(a))$$

Remarque 1. Cette notation est ambiguë, puisque l'on ne sait plus à quel $\mathfrak{S}_n$ le cycle appartient. En général, ceci n'a pas d'importance. Mieux, c'est plutôt un avantage puisqu'un cycle dont le support est inclus dans $\llbracket 1, n \rrbracket$ peut être vu comme un élément de $\mathfrak{S}_n, \mathfrak{S}_{n+1}$, etc.

Proposition 10. *L'ordre d'un cycle est égal à sa longueur.*

Démonstration. Soit $\sigma = (a \ \sigma(a) \ \dots \ \sigma^{p-1}(a))$ un p -cycle. On a pour tout $k < p$, $\sigma^k(a) \neq a$, donc $\sigma^k \neq id$. On vérifie ensuite facilement que $\sigma^p = id$. Ainsi, l'ordre de σ est p . $\square$

3.4 Cycles de supports disjoints

Proposition 11. *Deux cycles de supports disjoints commutent.*

Démonstration. Soient σ, σ' deux cycles de supports disjoints. Soit $a \in \llbracket 1, n \rrbracket$. Il y a trois possibilités.

- Si a est dans le support de σ , alors il n'est pas dans le support de σ' et $\sigma(a)$ non plus. De là,

$$\sigma\sigma'(a) = \sigma(a) = \sigma'\sigma(a)$$

- Même raisonnement si a est dans le support de σ' .
- Enfin, si a n'est dans aucun des deux supports, alors

$$\sigma(a) = \sigma'(a) = \sigma\sigma'(a) = \sigma'\sigma(a) = a$$

$\square$

Proposition 12. *Soient σ et σ' deux cycles de supports disjoints, de longueurs p et p' . Alors l'ordre de $\sigma\sigma'$ est $p \vee p'$.*

Démonstration. Posons $\gamma = \sigma\sigma'$. On a pour tout $k \in \mathbb{N}$, $(\sigma\sigma')^k = \sigma^k \sigma'^k$ puisque σ et σ' commutent.

Supposons $(\sigma\sigma')^k = id$. Soit a un élément du support de σ . a n'appartient donc pas au support de σ' , donc $\sigma'(a) = a$ et, a fortiori, $\sigma'^k(a) = a$. Ainsi,

$$a = (\sigma\sigma')^k(a) = \sigma^k(\sigma'^k(a)) = \sigma^k(a)$$

donc k est multiple de p . De même, en prenant a dans le support de σ' , on obtient que k est un multiple de p' . Ainsi, $p \vee p'$ divise k .

Inversement, supposons que $p \vee p'$ divise k . Alors p et p' divisent k , donc $\sigma^k = \sigma'^k = id$. Ainsi, $(\sigma\sigma')^k = id$.

En conclusion, $(\sigma\sigma')^k = id$ si et seulement si $p \vee p'$ divise k . L'ordre de $\sigma\sigma'$ est effectivement $p \vee p'$. $\square$

Proposition 13. *Toute permutation s'écrit de façon unique, à l'ordre près des facteurs, comme un produit de cycles de supports disjoints.*

Démonstration. Nous allons donner uniquement les grandes lignes de la preuve de l'existence. Soit $\sigma \in \mathfrak{S}_n$. Supposons que σ possède q orbites $\mathcal{O}_1, \dots, \mathcal{O}_q$ de cardinaux $p_1, \dots, p_q$. Pour $i = 1, \dots, q$, soit

$$\gamma_i = (a \ \sigma(a) \dots \sigma^{p_i-1}(a))$$

où $a \in \mathcal{O}_i$. Le support de γ_i est précisément $\mathcal{O}_i$. Les orbites étant disjointes, les γ_i sont des cycles de supports disjoints.

Soit maintenant $b \in \llbracket 1, n \rrbracket$. b appartient à une, et une seule des orbites suivant σ , mettons l'orbite $\mathcal{O}_i$. On a donc pour tout $j \neq i$, $\gamma_j(b) = b$, et $\gamma_i(b) = \sigma(b)$. Il en résulte facilement que

$$(\gamma_1 \dots \gamma_q)(b) = \sigma(b)$$

Ainsi,

$$\sigma = \gamma_1 \dots \gamma_q$$

$\square$

La preuve ci-dessus nous dit comment décomposer de façon effective une permutation σ en produit de cycles de supports disjoints.

- On choisit un point a non fixe de σ .

- On construit le cycle $\gamma_1 = (a \ \sigma(a) \dots \sigma^{p-1}(a))$ où p est le cardinal de $\mathcal{O}_\sigma(a)$.
- On recommence avec un point non fixe b qui n'est pas dans le support de γ_1 , s'il existe un tel point. On obtient un cycle γ_2 .
- etc.

Exemple. Soit

$$\sigma = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 & 6 & 7 & 8 & 9 & 10 & 11 \\ 3 & 1 & 7 & 8 & 5 & 6 & 2 & 10 & 11 & 4 & 9 \end{pmatrix}$$

On prend par exemple $a = 1$ et on obtient le cycle $(1 \ 3 \ 7 \ 2)$.

On recommence avec par exemple $b = 4$ et on obtient le cycle $(4 \ 8 \ 10)$.

On recommence avec par exemple $c = 9$ et on obtient le cycle $(9 \ 11)$.

Tous les points non fixes ont été considérés. Ainsi,

$$\sigma = (1 \ 3 \ 7 \ 2)(4 \ 8 \ 10)(9 \ 11)$$

Exercice. Décrire le groupe symétrique $\mathfrak{S}_4$.

Nous avons maintenant une technique pour énumérer de façon logique les éléments d'un groupe symétrique. On laisse les détails au lecteur. Voici les éléments de $\mathfrak{S}_4$:

- id .
- Les transpositions. Il y en a 6.
- Les 3-cycles. Il y en a 8.
- Les 4 cycles. Il y en a 6.
- Les produits de deux transpositions de supports disjoints. Il y en a 3.

Remarquons que $1 + 6 + 8 + 6 + 3 = 4!$, ce qui est heureux.

4 Signature

4.1 Décomposition en produit de transpositions

Proposition 14. Soit $n \in \mathbb{N}^*$. Toute permutation de $\mathfrak{S}_n$ est un produit d'au plus $n - 1$ transpositions.

Démonstration. On raisonne par récurrence sur n , où $\sigma \in \mathfrak{S}_n$.

- Si $n = 1$, c'est trivial. $\sigma = id$ est un produit de 0 transposition.
- Supposons la propriété vraie pour les permutations de $\mathfrak{S}_n$. Soit $\sigma \in \mathfrak{S}_{n+1}$.

Si $\sigma(n+1) = n+1$, il n'y a rien à faire : σ , identifiée à un élément de $\mathfrak{S}_n$, est un produit d'au plus $n - 1$ transpositions, donc d'au plus n transpositions.

Sinon, notons $\tau = ((n+1)\sigma(n+1))$. Soit $\sigma' = \tau\sigma$. Alors $\sigma'(n+1) = n+1$ donc $\sigma' \in \mathfrak{S}_n$ est un produit d'au plus $n-1$ transpositions par l'hypothèse de récurrence. Donc, $\sigma = \tau\sigma'$ est un produit d'au plus n transpositions.

□

Exemple. Décomposition d'un cycle. Soit $s = (a_1 a_2 \dots a_p)$ un cycle. On a alors une décomposition de s en produit de transpositions (ce n'est pas la seule) qui est

$$s = (a_1 a_2)(a_2 a_3) \dots (a_{p-1} a_p)$$

Remarque 2. Il n'y a pas unicité de la décomposition. Par exemple,

$$(1\ 2)(2\ 3) = (1\ 2\ 3) = (2\ 3\ 1) = (2\ 3)(3\ 1)$$

4.2 Signature d'une permutation

Définition 8. Soit $\sigma \in \mathfrak{S}_n$. On appelle signature de σ le nombre

$$\varepsilon(\sigma) = (-1)^{n-m}$$

où m est le nombre d'orbites suivant σ .

Exemple. $\varepsilon(id) = 1$. Pour toute transposition τ , $\varepsilon(\tau) = -1$. Plus généralement, si γ est un cycle de longueur $\ell \geq 2$, on a

$$\varepsilon(\gamma) = (-1)^{\ell-1}$$

Proposition 15. Soit $\sigma \in \mathfrak{S}_n$. Soit τ une transposition. On a

$$\varepsilon(\sigma\tau) = -\varepsilon(\sigma)$$

Démonstration. Notons $\tau = (i\ j)$ et $\sigma' = \sigma\tau$. Comment les orbites suivant σ se comportent-elles suivant σ' ? Soit γ une orbite suivant σ .

Tout d'abord, si ni i ni j n'appartiennent à γ , alors pour tout $x \in \gamma$ on a $\sigma'(x) = \sigma(x)$. γ est donc aussi une orbite suivant σ' . Nous avons réglé le cas de presque toutes les orbites, il n'en restent plus qu'une ou deux à considérer.

- Première possibilité, i et j sont dans une même orbite suivant σ . Il ne reste donc à considérer que l'orbite γ qui contient i et j . On a

$$\gamma = \mathcal{O}_\sigma(i) = \{i, \sigma(i), \dots, j = \sigma^k(i), \dots, \sigma^{p-1}(i)\}$$

où p est le cardinal de l'orbite γ . On a $\sigma'(i) = \sigma(j) = \sigma^{k+1}(i)$. Puis

$$\sigma'(\sigma^{k+1}(i)) = \sigma^{k+2}(i), \dots, \sigma'(\sigma^{p-1}(i)) = i$$

Nous avons une orbite suivant σ' , qui est

$$\gamma'_1 = \{i, \sigma(j), \dots, \sigma^{p-1}(i)\}$$

Cela dit, tous les points de σ n'ont pas été passés en revue. Le même raisonnement nous donne une seconde orbite suivant σ' , qui est

$$\gamma'_2 = \{j, \sigma(i), \dots, \sigma^{q-1}(j)\}$$

Ainsi, γ s'est transformée en *deux* orbites suivant σ' .

- Deuxième possibilité, i et j sont dans deux orbites distinctes suivant σ ,

$$\gamma_1 = \mathcal{O}_\sigma(i) = \{i, \sigma(i), \dots, \sigma^{p-1}(i)\}$$

et

$$\gamma_2 = \mathcal{O}_\sigma(j) = \{j, \sigma(j), \dots, \sigma^{q-1}(j)\}$$

Cette fois-ci, γ_1 et γ_2 se réunissent en une unique orbite suivant σ' , qui est

$$\gamma = \{i, \sigma(j), \dots, \sigma^{q-1}(j), \sigma(i), \dots, \sigma^{p-1}(i)\}$$

Ainsi, γ_1 et γ_2 se sont transformées en une orbite suivant σ' .

Bilan, le nombre d'orbites suivant σ' est égal au nombre d'orbites suivant $\sigma \pm 1$. On a donc effectivement $\varepsilon(\sigma') = -\varepsilon(\sigma)$. $\square$

Proposition 16. Soient $\tau_1, \dots, \tau_p$ p transpositions de $\mathfrak{S}_n$. On a

$$\varepsilon(\tau_1 \dots \tau_p) = (-1)^p$$

Démonstration. On fait une récurrence sur p . C'est évident par la proposition précédente. $\square$

Remarque 3. Nous avons vu qu'une permutation pouvait s'écrire de différentes façons comme produit de transpositions. Il y a en revanche *unicité de la parité* du nombre de transpositions qui interviennent dans le produit.

Proposition 17. La fonction $\varepsilon : \mathfrak{S}_n \longrightarrow \{-1, 1\}$ est un morphisme de groupes :

$$\forall \sigma, \sigma' \in \mathfrak{S}_n, \varepsilon(\sigma\sigma') = \varepsilon(\sigma)\varepsilon(\sigma')$$

Démonstration. Si σ est le produit de p transpositions et σ' est le produit de q transpositions, alors $\sigma\sigma'$ est le produit de $p + q$ transpositions. On a donc

$$\varepsilon(\sigma\sigma') = (-1)^{p+q} = (-1)^p(-1)^q = \varepsilon(\sigma)\varepsilon(\sigma')$$

□

Exercice. Trouver la signature de

$$\sigma = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 & 6 & 7 & 8 & 9 & 10 & 11 & 12 \\ 7 & 11 & 4 & 3 & 1 & 6 & 12 & 5 & 10 & 2 & 9 & 8 \end{pmatrix}$$

On peut par exemple décomposer σ en produit de cycles de supports disjoints.

$$\sigma = (1\ 7\ 12\ 8\ 5)(2\ 11\ 9\ 10)(3\ 4)$$

De là,

$$\varepsilon(\sigma) = \varepsilon(1\ 7\ 12\ 8\ 5)\varepsilon(2\ 11\ 9\ 10)\varepsilon(3\ 4) = 1 \times (-1) \times (-1) = 1$$

Exercice. Trouver la signature de tous les éléments de $\mathfrak{S}_3$, $\mathfrak{S}_4$ et $\mathfrak{S}_5$.

4.3 Groupe alterné

Définition 9. On appelle groupe alterné d'ordre n l'ensemble

$$\mathfrak{A}_n = \ker \varepsilon$$

Une permutation est dite paire lorsqu'elle est dans $\mathfrak{A}_n$, impaire sinon.

Exemple. Un p -cycle est pair si et seulement si p est impair.

Proposition 18. L'ensemble $\mathfrak{A}_n$ est un sous-groupe de $\mathfrak{S}_n$. Pour $n \geq 2$, on a

$$\text{card } \mathfrak{A}_n = \frac{n!}{2}$$

Démonstration. $\mathfrak{A}_n$ est un sous-groupe de $\mathfrak{S}_n$ en tant que noyau d'un morphisme de groupes. De plus, si $n \geq 2$, il existe dans $\mathfrak{S}_n$ au moins une permutation impaire (une transposition, par exemple). Notons-la τ . L'application $\sigma \mapsto \tau\sigma$ est alors une bijection de $\mathfrak{A}_n$ dans $\mathfrak{S}_n \setminus \mathfrak{A}_n$: il y a autant de permutations paires que de permutations impaires. D'où le résultat. □

Exercice. Décrire $\mathfrak{A}_3$, $\mathfrak{A}_4$ et $\mathfrak{A}_5$.

4.4 Complément

Quels sont *tous* les morphismes de groupes de $(\mathfrak{S}_n, \circ)$ vers $(\mathbb{C}^*, \times)$?

Nous avons bien entendu la fonction $\mathbb{1}$ constante égale à 1, et la signature ε . Y en a-t-il d'autres ? La réponse est non. Commençons par un lemme.

Lemme 19. Soit $n \geq 2$. Soient τ, τ' deux transpositions de $\mathfrak{S}_n$. Il existe $\sigma \in \mathfrak{S}_n$ telle que

$$\tau' = \sigma \tau \sigma^{-1}$$

Démonstration. Posons $\tau = (a \ b)$ et $\tau' = (c \ d)$, où $a \neq b$ et $c \neq d$. Soit $\sigma \in \mathfrak{S}_n$ telle que $\sigma(a) = c$ et $\sigma(b) = d$. Soit $x \in \llbracket 1, n \rrbracket$.

- Cas 1, $x \neq c, d$. On a alors $\sigma^{-1}(x) \neq a, b$, donc $\tau(\sigma^{-1}(x)) = \sigma^{-1}(x)$, d'où

$$\sigma \tau \sigma^{-1}(x) = \sigma \sigma^{-1}(x) = x = \tau'(x)$$

- Cas 2, $x = c$. On a alors

$$\sigma \tau \sigma^{-1}(x) = \sigma \tau(a) \sigma(b) = d = \tau'(x)$$

- Cas 3, $x = d$. Ce cas est identique au précédent.

τ' et $\sigma \tau \sigma^{-1}$ coïncident en tout point, ces deux applications sont donc égales. $\square$

On dit que les transpositions τ et τ' sont *conjuguées*.

Exercice. Montrer, plus généralement, que deux cycles de même longueur sont conjugués.

Proposition 20. Soit $n \geq 2$. Les seuls morphismes de groupes de $(\mathfrak{S}_n, \circ)$ vers $(\mathbb{C}^*, \times)$ sont $\mathbb{1}$ et ε .

Démonstration. Soit f un tel morphisme. Soient τ, τ' deux transpositions. Soit $\sigma \in \mathfrak{S}_n$ telle que $\tau' = \sigma \tau \sigma^{-1}$. On a alors

$$f(\tau') = f(\sigma \tau \sigma^{-1}) = f(\sigma) f(\tau) f(\sigma)^{-1} = f(\tau)$$

Ainsi, f prend la même valeur sur toutes les transpositions.

Quelle est cette valeur ? On a $(1 \ 2)^2 = id$, donc

$$f((1 \ 2))^2 = f(id) = 1$$

Ainsi, $f((1 \ 2)) = \pm 1$. Deux cas se présentent.

- Cas 1, $f((1\ 2)) = -1$. On a alors pour toute transposition τ , $f(\tau) = -1$. Soit $\sigma \in \mathfrak{S}_n$. Écrivons $\sigma = \tau_1 \dots \tau_p$ où les τ_i sont des transpositions. On a alors

$$f(\sigma) = f(\tau_1) \dots f(\tau_p) = (-1)^p = \varepsilon(\sigma)$$

Ainsi,

$$f = \varepsilon$$

- Cas 2, $f((1\ 2)) = 1$. Par le même raisonnement, on obtient

$$f = \mathbb{1}$$

□

5 Annexe : Inversions

Définition 10. Pour $\sigma \in \mathfrak{S}_n$, on appelle inversion de σ tout couple (i, j) tel que $i < j$ et $\sigma(i) > \sigma(j)$. On note $I(\sigma)$ le nombre d'inversions de σ . On note également $e(\sigma) = (-1)^{I(\sigma)}$.

Exemple. $I(id) = 0$. Soit $\tau = (pq)$ une transposition. Prenons par exemple $p < q$. Les couples (i, j) tels que $i < j$ et $\tau(i) < \tau(j)$ ont soit leur premier élément égal à p , soit leur second élément égal à q . Précisément, les couples concernés sont $(p, p+1), (p, p+2), \dots, (p, q)$, et aussi $(p+1, q), (p+2, q), \dots, (q-1, q)$. Donc $I(\tau) = 2(q-p) - 1$.

Notation. On note $\mathcal{P}_2(n)$ l'ensemble des parties de $\{1, \dots, n\}$ à deux éléments. On a ainsi $\text{card } \mathcal{P}_2(n) = \frac{n(n-1)}{2}$.

Lemme 21. Soit $\sigma \in \mathfrak{S}_n$. L'application, notée encore σ , $\mathcal{P}_2(n) \longrightarrow \mathcal{P}_2(n)$ définie par $\sigma(\{i, j\}) = \{\sigma(i), \sigma(j)\}$ est une bijection.

Démonstration. Il suffit de prouver l'injectivité, qui est évidente. □

Proposition 22. Soit $\sigma \in \mathfrak{S}_n$. On a

$$e(\sigma) = \prod_{\{i, j\} \in \mathcal{P}_2(n)} \frac{\sigma(i) - \sigma(j)}{i - j}$$

Démonstration. Soit

$$\varphi(\sigma) = \prod_{\{i,j\} \in \mathcal{P}_2(n)} \frac{\sigma(i) - \sigma(j)}{i - j}$$

On constate tout d'abord que

$$\varphi(\sigma) = \prod_{i < j} \frac{\sigma(i) - \sigma(j)}{i - j}$$

Donc, le signe de $\varphi(\sigma)$ est le même que celui de $e(\sigma)$.

Considérons ensuite

$$P = \left| \prod_{\{i,j\} \in \mathcal{P}_2(n)} (\sigma(i) - \sigma(j)) \right|$$

Posons $\{i', j'\} = \sigma(\{i, j\})$. Alors,

$$P = \left| \prod_{\{i',j'\} \in \mathcal{P}_2(n)} (i' - j') \right|$$

d'après le lemme démontré plus haut. Ainsi, $|\varphi(\sigma)| = 1$.

Les quantités $e(\sigma)$ et $\varphi(\sigma)$ ont même signe et même valeur absolue, elles sont donc égales.

□

Proposition 23. *La fonction $e : \mathfrak{S}_n \longrightarrow \{-1, 1\}$ est un morphisme de groupes.*

Démonstration. Soient σ et τ deux permutations de $\mathfrak{S}_n$. On a

$$e(\sigma\tau) = \prod_{\{i,j\} \in \mathcal{P}_2(n)} \frac{\sigma\tau(i) - \sigma\tau(j)}{i - j} = \prod_{\{i,j\} \in \mathcal{P}_2(n)} \frac{\sigma(\tau(i)) - \sigma(\tau(j))}{\tau(i) - \tau(j)} \prod_{\{i,j\} \in \mathcal{P}_2(n)} \frac{\tau(i) - \tau(j)}{i - j}$$

Le second produit est $\varepsilon(\tau)$. Quant au premier produit, il suffit de poser $\{i', j'\} = \tau(\{i, j\})$, pour voir qu'il est égal à $\varepsilon(\sigma)$. Ainsi,

$$e(\sigma\tau) = e(\sigma)e(\tau)$$

□

Proposition 24. $e = \varepsilon$.

Démonstration. Rappelons-nous que pour toute transposition $\tau = (pq)$ (où $p < q$) on a $I(\tau) = 2(q - p) - 1$. On a donc $e(\tau) = -1 = \varepsilon(\tau)$. Ainsi, e et ε sont deux morphismes de groupes qui coïncident sur les transpositions. Comme toute permutation est produit de transpositions on en déduit aisément que $e = \varepsilon$. □

Chapitre 26

Déterminants

1 Applications multilinéaires

1.1 Définitions

Définition 1. Soient $E_1, \dots, E_n, F$ $n + 1$ $\mathbb{K}$ -espaces vectoriels. Soit

$$f : E_1 \times \dots \times E_n \longrightarrow F$$

une application. On dit que f est n -linéaire lorsque f est « linéaire en chacune de ses variables » : $\forall i \in \llbracket 1, n \rrbracket, \forall x_1, \dots, x_{i-1}, x_{i+1}, \dots, x_n \in E_1 \times \dots \times E_{i-1} \times E_{i+1} \times \dots \times E_n$, l'application $f_i : E_i \longrightarrow F$ définie par

$$f_i(x) = f(x_1, \dots, x, \dots, x_n)$$

est linéaire.

Si $F = \mathbb{K}$, on dit que f est une forme n -linéaire.

Exemple.

- La multiplication des réels, le produit matriciel, la multiplication des polynômes, sont bilinéaires.
- Le produit scalaire sur $\mathbb{R}^n$ est une forme bilinéaire.
- Le produit vectoriel sur $\mathbb{R}^3$ est une application bilinéaire.
- Le déterminant sur $\mathbb{R}^2$ est une forme bilinéaire (cf chapitre de début d'année).
- Le déterminant sur $\mathbb{R}^3$ est une forme trilinéaire.

Notation. On note $\mathcal{L}(E_1, \dots, E_n; F)$ l'ensemble des applications n -linéaires de $E_1 \times \dots \times E_n$ vers F . Lorsque $E_1 = \dots = E_n$, on note plus simplement $\mathcal{L}_n(E, F)$.

Exemple. Déterminons toutes les formes bilinéaires $\mathbb{R}^2 \times \mathbb{R}^2 \longrightarrow \mathbb{R}$. Soit (e_1, e_2) une base du plan. Soit f une forme bilinéaire sur $\mathbb{R}^2$. Soient $u = xe_1 + ye_2$ et $v = x'e_1 + y'e_2$. On a

$$f(u, v) = a_{11}xx' + a_{12}xy' + a_{21}x'y + a_{22}yy'$$

où $a_{ij} = f(e_i, e_j)$.

Inversement, toute application de cette forme est clairement bilinéaire. Creusons un peu. Pour $i, j \in \llbracket 1, 2 \rrbracket$, soit f_{ij} définie par

$$f_{ij}(x_1e_1 + x_2e_2, x'_1e_1 + x'_2e_2) = x_ix'_j$$

Toute forme bilinéaire sur $\mathbb{R}^2$ s'écrit de façon unique $f = \sum_{i,j} a_{ij}f_{ij}$ où les $a_{ij} \in \mathbb{R}$. En résumé,

$$\mathcal{L}_2(\mathbb{R}^2, \mathbb{R}) = \langle f_{ij}, 1 \leq i \leq 2, 1 \leq j \leq 2 \rangle$$

est un $\mathbb{R}$ -espace vectoriel de dimension 4.

1.2 Structure des ensembles d'applications n-linéaires

Proposition 1. $\mathcal{L}(E_1, \dots, E_n; F)$ est un $\mathbb{K}$ -espace vectoriel de dimension

$$\left(\prod_{k=1}^n \dim E_k \right) \times \dim F$$

lorsque tous les espaces mis en jeu sont de dimension finie.

Démonstration. Nous avons vu un cas particulier de ce théorème dans l'exemple ci-dessus. Le cas général se démontre de la même façon, en introduisant des bases des E_k et de F . Nous ne donnerons pas les détails. $\square$

1.3 Formes n-linéaires alternées

Définition 2. Soit $f \in \mathcal{L}_n(E, F)$. On dit que f est

- symétrique lorsque pour tous $x_1, \dots, x_n \in E$, pour tous $i, j \in \llbracket 1, n \rrbracket$,

$$f(x_1, \dots, x_j, \dots, x_i, \dots, x_n) = f(x_1, \dots, x_i, \dots, x_j, \dots, x_n)$$

- antisymétrique lorsque pour tous $x_1, \dots, x_n \in E$, pour tous $i, j \in \llbracket 1, n \rrbracket$, $i \neq j$,

$$f(x_1, \dots, x_j, \dots, x_i, \dots, x_n) = -f(x_1, \dots, x_i, \dots, x_j, \dots, x_n)$$

- alternée lorsque pour tous $x_1, \dots, x_n \in E$, pour tous $i, j \in \llbracket 1, n \rrbracket$, $i \neq j$,

$$x_i = x_j \implies f(x_1, \dots, x_n) = 0$$

Exemple. Recherchons toutes les formes antisymétriques sur $\mathbb{R}^2$. Nous avons déjà la forme générale d'une forme bilinéaire. En reprenant les notations ci-dessus, supposons que f est antisymétrique. Alors,

$$a_{11} = f(e_1, e_1) = -f(e_1, e_1) = -a_{1,1}$$

Donc $a_{11} = 0$. De même, $a_{22} = 0$. Et

$$a_{12} = f(e_1, e_2) = -f(e_2, e_1) = -a_{21}$$

Finalement, $f(u, v) = k(xy' - x'y)$ où $k \in \mathbb{R}$. Inversement, une telle application est bien antisymétrique.

Exercice. Déterminer toutes les formes bilinéaires alternées sur $\mathbb{R}^2$. *Indication : lire le théorème suivant.*

Proposition 2. Soit $\mathbb{K}$ un corps dans lequel $2 \neq 0$. Soit E un $\mathbb{K}$ -espace vectoriel. Soit $f \in \mathcal{L}_n(E, \mathbb{K})$. Alors, f est antisymétrique si et seulement si f est alternée.

Démonstration. On le fait pour $n = 2$, histoire de simplifier. Soient $x, y \in E$.

Si f est antisymétrique, alors, pour tout $x \in E$, $f(x, x) = -f(x, x)$, d'où $2f(x, x) = 0$. Comme 2 est non nul, f est donc alternée.

Si f est alternée, alors, pour tous $x, y \in E$,

$$0 = f(x + y, x + y) = f(x, y) + f(y, x)$$

donc f est antisymétrique. $\square$

Notation. On note $\Lambda_n(E)$ l'ensemble des formes n -linéaires alternées sur E .

Proposition 3. Soit f une forme n -linéaire sur E . Alors, $f \in \Lambda_n(E)$ si et seulement si pour tout $\sigma \in \mathfrak{S}_n$,

$$\forall x_1, \dots, x_n \in E, f(x_{\sigma(1)}, \dots, x_{\sigma(n)}) = \varepsilon(\sigma) f(x_1, \dots, x_n)$$

Démonstration. Dans un sens, c'est évident : prendre pour σ une transposition.

Inversement, toute permutation est un produit de transpositions. On raisonne par récurrence sur le nombre de transpositions du produit. $\square$

2 Déterminant d'une famille de vecteurs

2.1 Notations

Dans ce qui va suivre :

- E est un $\mathbb{K}$ -espace vectoriel de dimension $n \geq 1$.
- $\mathcal{B} = (e_1, \dots, e_n)$ est une base de E .
- $x_1, \dots, x_n$ sont n vecteurs de E .
- $x_j = \sum_{i=1}^n x_{ij} e_i$ est l'écriture de x_j dans la base $\mathcal{B}$.
- f est une forme n -linéaire alternée sur E .

On cherche à exprimer $f(x_1, \dots, x_n)$ à l'aide des coordonnées des x_i .

2.2 Calculs

On utilise tout d'abord la n -linéarité de f :

$$\begin{aligned} f(x_1, \dots, x_n) &= f\left(\sum_{i_1=1}^n x_{i_1 1} e_{i_1}, \dots, \sum_{i_n=1}^n x_{i_n n} e_{i_n}\right) \\ &= \sum_{i_1=1}^n \dots \sum_{i_n=1}^n x_{i_1 1} \dots x_{i_n n} f(e_{i_1}, \dots, e_{i_n}) \\ &= \sum_{(i_1, \dots, i_n) \in \mathcal{A}} x_{i_1 1} \dots x_{i_n n} f(e_{i_1}, \dots, e_{i_n}) \end{aligned}$$

où

$$A = \{(i_1, \dots, i_n) \in \llbracket 1, n \rrbracket^n, \forall j \neq k, i_j \neq i_k\}$$

En effet, $f(e_{i_1}, \dots, e_{i_n}) = 0$ dès que deux des e_{i_j} sont égaux (alternance de f).

Soit $\psi : \mathfrak{S}_n \longrightarrow \mathcal{A}$ définie par

$$\psi(\sigma) = (\sigma(1), \dots, \sigma(n))$$

L'application ψ est une bijection. On peut grâce à elle effectuer un changement d'indice dans la somme qui nous (pré)occupe, et écrire :

$$\begin{aligned} f(x_1, \dots, x_n) &= \sum_{\sigma \in \mathfrak{S}_n} x_{\sigma(1)1} \dots x_{\sigma(n)n} f(e_{\sigma(1)1}, \dots, e_{\sigma(n)n}) \\ &= f(e_1, \dots, e_n) \times \sum_{\sigma \in \mathfrak{S}_n} \varepsilon(\sigma) x_{\sigma(1)1} \dots x_{\sigma(n)n} \end{aligned}$$

On a donc $f(x_1, \dots, x_n) = C\varphi(x_1, \dots, x_n)$ où $\varphi : E^n \longrightarrow \mathbb{K}$ est définie par

$$\varphi(x_1, \dots, x_n) = \sum_{\sigma \in \mathfrak{S}_n} \varepsilon(\sigma) x_{\sigma(1)1} \dots x_{\sigma(n)n}$$

et

$$C = f(e_1, \dots, e_n)$$

ne dépend pas des vecteurs x_j .

Conclusion : toute forme n -linéaire alternée sur E est un multiple de φ .

Il reste à prouver deux choses :

- φ n'est pas la fonction nulle.
- φ est n -linéaire alternée.
- Les coordonnées du vecteur e_j dans la base $\mathcal{B}$ sont les δ_{ij} , $1 \leq i \leq n$:

$$e_j = \sum_{i_j=1}^n \delta_{i_j j} e_i$$

Donc,

$$\varphi(e_1, \dots, e_n) = \sum_{\sigma \in \mathfrak{S}_n} \varepsilon(\sigma) \delta_{\sigma(1)1} \dots \delta_{\sigma(n)n}$$

Or, tous les produits dans la somme ci-dessus sont nuls, sauf lorsque $\sigma = id$. Donc, $\varphi(e_1, \dots, e_n) = 1$. Ainsi, φ n'est pas la fonction nulle.

• Soit $\gamma \in \mathfrak{S}_n$. Posons $y_j = x_{\gamma(j)}$. On a

$$\begin{aligned} \varphi(x_{\gamma(1)}, \dots, x_{\gamma(n)}) &= \sum_{\sigma \in \mathfrak{S}_n} \varepsilon(\sigma) y_{\sigma(1)1} \dots y_{\sigma(n)n} \\ &= \sum_{\sigma \in \mathfrak{S}_n} \varepsilon(\sigma) x_{\sigma(1)\gamma(1)} \dots x_{\sigma(n)\gamma(n)} \\ &= \sum_{\sigma \in \mathfrak{S}_n} \varepsilon(\sigma) \prod_{k=1}^n x_{\sigma(k)\gamma(k)} \end{aligned}$$

On pose dans le produit $k = \gamma(k')$ (γ est une bijection). Il vient

$$\varphi(x_{\gamma(1)}, \dots, x_{\gamma(n)}) = \sum_{\sigma \in \mathfrak{S}_n} \varepsilon(\sigma) \prod_{k=1}^n x_{\sigma\gamma(k)k}$$

Maintenant, on pose dans la somme $\sigma = \sigma'\gamma$, et on obtient

$$\varphi(x_{\gamma(1)}, \dots, x_{\gamma(n)}) = \sum_{\sigma \in \mathfrak{S}_n} \varepsilon(\sigma\gamma) \prod_{k=1}^n x_{\sigma(k)k} = \varepsilon(\gamma) \varphi(x_1, \dots, x_n)$$

puisque $\varepsilon(\sigma\gamma) = \varepsilon(\gamma)\varepsilon(\sigma)$. Ainsi, φ est antisymétrique, donc alternée.

Remarque 1. Pour être parfaitement exacts, notre preuve ne fonctionne que dans un corps $\mathbb{K}$ où $2 \neq 0$.

2.3 Bilan

Proposition 4. Les formes n -linéaires alternées sur E sont les multiples de l'application $\varphi : E^n \longrightarrow \mathbb{K}$ définie par

$$\varphi(x_1, \dots, x_n) = \sum_{\sigma \in \mathfrak{S}_n} \varepsilon(\sigma) x_{\sigma(1)1} \dots x_{\sigma(n)n}$$

De plus, φ est l'unique application n -linéaire alternée sur E qui prend la valeur 1 en $(e_1, \dots, e_n)$.

Corollaire 5. $\Lambda_n(E)$ est un $\mathbb{K}$ -espace vectoriel de dimension 1.

Démonstration. En effet, $\Lambda_n(E) = \langle \varphi \rangle$ et $\varphi \neq 0$. $\square$

2.4 Déterminant d'une famille de vecteurs

Définition 3. Avec les notations qui précèdent, l'application φ est appelée « déterminant dans la base $\mathcal{B}$ », et est notée $\det_{\mathcal{B}}$:

$$\det_{\mathcal{B}}(x_1, \dots, x_n) = \sum_{\sigma \in \mathfrak{S}_n} \varepsilon(\sigma) x_{\sigma(1)1} \dots x_{\sigma(n)n}$$

Remarque 2. L'application « déterminant dans la base $\mathcal{B}$ » est donc l'unique application n -linéaire alternée f sur E telle que $f(\mathcal{B}) = 1$.

Remarque 3. Lorsque σ décrit $\mathfrak{S}_n$, σ^{-1} aussi. Le changement d'indice $\sigma = \sigma'^{-1}$ dans l'expression de φ montre (calcul omis) que l'on a également :

$$\det_{\mathcal{B}}(x_1, \dots, x_n) = \sum_{\sigma \in \mathfrak{S}_n} \varepsilon(\sigma) x_{1\sigma(1)} \dots x_{n\sigma(n)}$$

2.5 Déterminant d'une matrice carrée

Définition 4. Soit $A \in \mathcal{M}_n(\mathbb{K})$. On appelle déterminant de A le déterminant de la famille des colonnes de A , identifiées à des vecteurs de $\mathbb{K}^n$, dans la base canonique de $\mathbb{K}^n$.

Notation. Si $A = (a_{ij})_{1 \leq i, j \leq n}$, on notera le déterminant de A :

$$\det A = \begin{vmatrix} a_{11} & \cdots & a_{1n} \\ \vdots & & \vdots \\ a_{n1} & \cdots & a_{nn} \end{vmatrix}$$

Une formule pour le déterminant de A est

$$\det A = \sum_{\sigma \in \mathfrak{S}_n} \varepsilon(\sigma) a_{\sigma(1)1} \dots a_{\sigma(n)n}$$

Proposition 6. Soit $A \in \mathcal{M}_n(\mathbb{K})$. On a $\det {}^t A = \det A$.

Démonstration. C'est une conséquence directe de « l'autre » écriture de φ dans le paragraphe précédent. $\square$

On a donc aussi

$$\det A = \sum_{\sigma \in \mathfrak{S}_n} \varepsilon(\sigma) a_{1\sigma(1)} \cdots a_{n\sigma(n)}$$

2.6 Exemples

Pour $n = 2$, on a

$$\begin{vmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{vmatrix} = a_{11}a_{22} - a_{12}a_{21}$$

Pour $n = 3$, on a

$$\begin{vmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{vmatrix} = a_{11}a_{22}a_{33} + a_{21}a_{32}a_{13} + a_{31}a_{12}a_{23} \\ - a_{21}a_{12}a_{33} - a_{11}a_{32}a_{23} - a_{31}a_{22}a_{13}$$

De façon générale, pour n quelconque, un déterminant est une expression possédant $n!$ termes. La moitié de ces termes est affectée d'un signe plus (permutations paires), l'autre moitié d'un signe moins (permutations impaires).

2.7 Déterminants, bases, matrices inversibles

Proposition 7. Soient E un $\mathbb{K}$ -espace vectoriel de dimension n , $\mathcal{B}$ une base de E , et $\mathcal{F}$ une famille de n vecteurs de E . Alors $\mathcal{F}$ est une base de E si et seulement si

$$\det_{\mathcal{B}} \mathcal{F} \neq 0$$

Proposition 8. Soit $A \in \mathcal{M}_n(\mathbb{K})$. Alors, A est inversible si et seulement si

$$\det A \neq 0$$

Démonstration. Ces deux propositions sont évidemment deux points de vues d'un même théorème.

($\Rightarrow$) Supposons que $\mathcal{F}$ est une base de E . On dispose alors des *deux* applications n -linéaires alternées non nulles, $\det_{\mathcal{B}}$ et aussi $\det_{\mathcal{F}}$. Ces deux applications sont proportionnelles : il existe donc un réel λ (non nul) tel que $\det_{\mathcal{B}} = \lambda \det_{\mathcal{F}}$. En particulier

$$\det_{\mathcal{B}}(\mathcal{F}) = \lambda \det_{\mathcal{F}}(\mathcal{F}) = \lambda \neq 0$$

($\Leftarrow$) Supposon que $\mathcal{F} = (x_1, \dots, x_n)$ est liée. Alors, l'un des vecteurs est combinaison linéaire des autres :

$$\exists i \in \llbracket 1, n \rrbracket, x_i = \sum_{j \neq i} \lambda_j x_j$$

Mais alors,

$$\begin{aligned} \det_{\mathcal{B}}(\mathcal{F}) &= \det_{\mathcal{B}}(x_1, \dots, \sum_{j \neq i} \lambda_j x_j, \dots, x_n) \\ &= \sum_{j \neq i} \lambda_j \det_{\mathcal{B}}(x_1, \dots, x_j, \dots, x_n) \end{aligned}$$

où x_j se trouve à la position i . Comme x_j se trouve aussi à la position j , par l'alternance, tous les termes de la somme sont nuls. $\square$

3 Calcul des déterminants

3.1 Déterminant d'une matrice triangulaire

Proposition 9. Soit $A = (a_{ij})_{1 \leq i, j \leq n}$ une matrice triangulaire. Alors

$$\det A = \prod_{k=1}^n a_{kk}$$

Démonstration. Prenons par exemple A triangulaire supérieure. Pour tout $i > j$, $a_{ij} = 0$. On a

$$\det A = \sum_{\sigma \in \mathfrak{S}_n} \varepsilon(\sigma) T_{\sigma}$$

où $T_{\sigma} = a_{\sigma(1)1} \dots a_{\sigma(n)n}$. Mais si $T_{\sigma} \neq 0$ alors, par la triangularité de A ,

$$\sigma(1) \leq 1, \dots, \sigma(n) \leq n$$

D'où $\sigma = id$ par une petite récurrence. Il n'y a donc au plus qu'un terme non nul dans la somme. $\square$

Exemple. Soit $A = \begin{pmatrix} 2 & 2 & 3 & 4 \\ 0 & 1 & 6 & 5 \\ 0 & 0 & 4 & 3 \\ 0 & 0 & 0 & 1 \end{pmatrix}$. On a

$$\det A = 2 \times 1 \times 4 \times 1 = 8$$

3.2 Opérations sur les lignes et les colonnes

Proposition 10.

- Échanger deux colonnes d'un déterminant change le signe de celui-ci.
- Multiplier une colonne d'un déterminant par un scalaire λ revient à multiplier ce déterminant par λ .
- Ajouter à une colonne d'un déterminant une combinaison linéaire des autres colonnes ne change pas la valeur du déterminant.

Démonstration. Ce sont des conséquences directes de l'antisymétrie, de la multilinéarité, et de l'alternance. $\square$

Remarque 4. Bien entendu, ces résultats sont également valables pour les opérations sur les lignes. Il suffit de transposer.

Exemple. Soit $x \in \mathbb{C}$. Calculons le déterminant de taille n

$$D_n = \begin{vmatrix} x & 1 & \cdots & 1 \\ 1 & x & \ddots & 1 \\ \vdots & \ddots & \ddots & 1 \\ 1 & \cdots & 1 & x \end{vmatrix}$$

On effectue d'abord $C_1 \leftarrow C_1 + C_2 + \dots + C_n$. Il vient

$$D_n = (x + n - 1) \begin{vmatrix} 1 & 1 & \cdots & 1 \\ 1 & x & \ddots & 1 \\ \vdots & \ddots & \ddots & 1 \\ 1 & \cdots & 1 & x \end{vmatrix}$$

On effectue ensuite $L_i \leftarrow L_i - L_1$, pour $i = 2, 3, \dots, n$. Il vient

$$D_n = (x + n - 1) \begin{vmatrix} 1 & \cdots & \cdots & 1 \\ 0 & x - 1 & \ddots & \vdots \\ \vdots & \ddots & \ddots & 0 \\ 0 & \cdots & 0 & x - 1 \end{vmatrix} = (x + n - 1)(x - 1)^{n-1}$$

3.3 Développement suivant une ligne ou une colonne

Définition 5. Soit $A = (a_{ij}) \in \mathcal{M}_n(\mathbb{K})$. On appelle :

- Mineur de A d'indices i, j le déterminant de la matrice obtenue en rayant la ligne i et la colonne j de A .
- Cofacteur de A d'indices i, j le mineur de A d'indices i, j multiplié par $(-1)^{i+j}$.
- Comatrice de A la matrice des cofacteurs de A .

Notation. Nous noterons

- $\text{Min}(A, i, j)$ le mineur de A , ligne i , colonne j .
- $\text{Cof}(A, i, j)$ le cofacteur de A , ligne i , colonne j .
- $\text{Com } A$ la comatrice de A .

Exemple.

$$\text{Com} \begin{pmatrix} a & c \\ b & d \end{pmatrix} = \begin{pmatrix} d & -b \\ -c & a \end{pmatrix}$$

$$\text{Com} \begin{pmatrix} 1 & 4 & 7 \\ 2 & 5 & 8 \\ 3 & 6 & 9 \end{pmatrix} = \begin{pmatrix} -3 & 6 & -3 \\ 6 & -12 & 6 \\ -3 & 6 & -3 \end{pmatrix}$$

On va maintenant établir une formule permettant de calculer le déterminant d'une matrice en fonction de certains de ses cofacteurs.

Soit $A = (a_{ij})_{1 \leq i, j \leq n} \in \mathcal{M}_n(\mathbb{K})$. Appelons $A_1, \dots, A_n$ les colonnes de A , identifiées à des vecteurs de $\mathbb{K}^n$, et soit $\mathcal{B} = (e_1, \dots, e_n)$ la base canonique de $\mathbb{K}^n$. Soit j un entier fixé entre 1 et n . On a

$$\begin{aligned} \det A &= \det_{\mathcal{B}}(A_1, \dots, A_n) \\ &= \det_{\mathcal{B}}(A_1, \dots, \sum_{i=1}^n a_{ij} e_i, \dots, A_n) \\ &= \sum_{i=1}^n a_{ij} \det_{\mathcal{B}}(A_1, \dots, e_i, \dots, A_n) \end{aligned}$$

où e_i est en j ème position.

Posons $D_{ij} = \det_{\mathcal{B}}(A_1, \dots, e_i, \dots, A_n)$:

$$D_{ij} = \begin{vmatrix} a_{11} & \cdots & 0 & \cdots & a_{1n} \\ \vdots & & \vdots & & \vdots \\ \vdots & & 0 & & \vdots \\ \vdots & & 1 & & \vdots \\ \vdots & & 0 & & \vdots \\ \vdots & & \vdots & & \vdots \\ a_{n1} & \cdots & 0 & \cdots & a_{nn} \end{vmatrix}$$

le 1 étant à la ligne i , colonne j . On effectue les opérations

$$L_i \longleftrightarrow L_{i+1}, L_{i+1} \longleftrightarrow L_{i+2}, \dots, L_{n-1} \longleftrightarrow L_n$$

puis

$$C_j \longleftrightarrow C_{j+1}, C_{j+1} \longleftrightarrow C_{j+2}, \dots, C_{n-1} \longleftrightarrow C_n$$

ce qui donne

$$D_{ij} = (-1)^{i+j} \begin{vmatrix} a_{11} & \cdots & a_{1n} & 0 \\ \vdots & & \vdots & \vdots \\ a_{n1} & \cdots & a_{nn} & 0 \\ \cdots & \cdots & \cdots & 1 \end{vmatrix}$$

Attention, il n'y a « plus » de ligne i et de colonne j dans le déterminant ci-dessus. Le calcul de ce déterminant par la formule de la définition fait apparaître une somme sur les $\sigma \in \mathfrak{S}_n$ tels que $\sigma(n) = n$. On change d'indice en remarquant que l'ensemble de ces permutations est en bijection avec $\mathfrak{S}_{n-1}$. On trouve alors que

$$D_{ij} = \text{Cof}(A, i, j)$$

On a ainsi prouvé la

Proposition 11. Soit $A = (a_{ij}) \in \mathcal{M}_n(\mathbb{K})$. Soit $j \in [1, n]$. On a

$$\det A = \sum_{i=1}^n a_{ij} \text{Cof}(A, i, j)$$

Définition 6. La formule précédente est appelée développement de $\det A$ par rapport à la colonne j .

Tout ce qui vient d'être vu peut être bien sûr transposé :

Proposition 12. Soit $A = (a_{ij}) \in \mathcal{M}_n(\mathbb{K})$. Soit $i \in [1, n]$. On a

$$\det A = \sum_{j=1}^n a_{ij} \text{Cof}(A, i, j)$$

Définition 7. La formule précédente est appelée développement de $\det A$ par rapport à la ligne i .

Exemple. Développons

$$D = \begin{vmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \\ 7 & 8 & 9 \end{vmatrix}$$

par rapport à la première colonne. On a

$$\begin{aligned} D &= 1 \times \begin{vmatrix} 5 & 6 \\ 8 & 9 \end{vmatrix} - 4 \times \begin{vmatrix} 2 & 3 \\ 8 & 9 \end{vmatrix} + 7 \times \begin{vmatrix} 2 & 3 \\ 5 & 6 \end{vmatrix} \\ &= 1 \times (-3) - 4 \times (-6) + 7 \times (-3) \\ &= 0 \end{aligned}$$

Exemple. Calculons

$$D_n = \begin{vmatrix} 1+x^2 & x & 0 & \cdots & \cdots & 0 \\ x & 1+x^2 & x & \ddots & & \vdots \\ 0 & \ddots & \ddots & \ddots & \ddots & \vdots \\ \vdots & \ddots & & & & \end{vmatrix}$$

On développe par rapport à la première colonne :

$$D_n = (1+x^2)D_{n-1} - x\Delta_{n-1}$$

où Δ_{n-1} est le mineur « de x ». En développant Δ_{n-1} par rapport à la première colonne, on voit que $\Delta_{n-1} = xD_{n-2}$. D'où

$$D_n = (1+x^2)D_{n-1} - x^2D_{n-2}$$

Il ne reste qu'à résoudre cette récurrence, qui est une récurrence linéaire à deux termes. La fin du calcul ne pose pas de problème, elle est laissée au lecteur.

3.4 Un exemple plus compliqué

Soit $n \in \mathbb{N}^*$. Soient $a_0, \dots, a_{n-1}$ n réels ou complexes. Soit $V_n(a_0, \dots, a_{n-1})$ (que nous noterons juste V_n s'il n'y a pas d'ambiguïté) le déterminant $n \times n$ dans lequel le coefficient ligne i , colonne j est a_i^j . Attention, une fois n'est pas coutume, on numérote les lignes et les colonnes à partir de 0 et pas de 1. Ce déterminant s'appelle le déterminant de *Vandermonde* et on se propose de le calculer. Plus explicitement, on a

$$V_n = \begin{vmatrix} x_0^0 & x_0^1 & \dots & x_0^{n-1} \\ x_1^0 & x_1^1 & \dots & x_1^{n-1} \\ \dots & \dots & \dots & \dots \\ x_{n-1}^0 & x_{n-1}^1 & \dots & x_{n-1}^{n-1} \end{vmatrix}$$

On effectue les opérations suivantes : $C_n \leftarrow C_k - x_{n-1}C_{k-1}$ pour $k = n, n-1, \dots, 2$. On en déduit

$$V_n = \begin{vmatrix} 1 & x_0 - x_{n-1} & x_0(x_0 - x_{n-1}) & \dots & x_0^{n-2}(x_0 - x_{n-1}) \\ 1 & x_1 - x_{n-1} & x_1(x_1 - x_{n-1}) & \dots & x_1^{n-2}(x_1 - x_{n-1}) \\ \dots & \dots & \dots & \dots & \dots \\ 1 & x_{n-2} - x_{n-1} & x_{n-2}(x_{n-2} - x_{n-1}) & \dots & x_{n-2}^{n-2}(x_{n-2} - x_{n-1}) \\ 1 & 0 & 0 & \dots & 0 \end{vmatrix}$$

Développons par rapport à la dernière ligne. Nous obtenons

$$V_n = (-1)^{n+1} \begin{vmatrix} x_0 - x_{n-1} & x_0(x_0 - x_{n-1}) & \dots & x_0^{n-2}(x_0 - x_{n-1}) \\ x_1 - x_{n-1} & x_1(x_1 - x_{n-1}) & \dots & x_1^{n-2}(x_1 - x_{n-1}) \\ \dots & \dots & \dots & \dots \\ x_{n-2} - x_{n-1} & x_{n-2}(x_{n-2} - x_{n-1}) & \dots & x_{n-2}^{n-2}(x_{n-2} - x_{n-1}) \end{vmatrix}$$

Sur chaque ligne i , il y a une factorisation possible par $x_i - x_{n-1}$. On obtient

$$V_n(x_0, \dots, x_{n-1}) = V_{n-1}(x_0, \dots, x_{n-2}) \prod_{i=0}^{n-2} (x_{n-1} - x_i)$$

Par une récurrence simple, on en tire

$$V_n(x_0, \dots, x_{n-1}) = \prod_{0 \leq i < j \leq n-1} (x_j - x_i)$$

On a donc obtenu une factorisation complète du déterminant de Vandermonde, avec en prime une condition nécessaire et suffisante de nullité :

Proposition 13. $V_n = 0$ si et seulement si deux des x_k sont égaux.

Remarque 5. Dans un sens, c'était évident, puisque si deux des x_k sont égaux on a deux lignes égales dans le déterminant. L'autre implication n'était pas triviale.

Remarque 6. On rencontre des déterminants de Vandermonde dans un grand nombre de situations. Prenons l'exemple de l'interpolation de Lagrange. Soient $x_0, \dots, x_{n-1} \in \mathbb{C}$ distincts deux à deux, soient $y_0, \dots, y_{n-1} \in \mathbb{C}$. Soit

$$P = \sum_{k=0}^{n-1} a_k x^k$$

un polynôme de $\mathbb{C}[X]$ de degré inférieur ou égal à $n-1$. On a pour tout $j \in \llbracket 0, n-1 \rrbracket$, $P(x_j) = y_j$, si et seulement si $\forall j \in \llbracket 0, n-1 \rrbracket$,

$$\sum_{k=0}^{n-1} a_k x_j^k = y_j$$

Interprétons matriciellement ce système dont les inconnues sont les a_k . On peut encore l'écrire $AX = Y$ où

$$X = \begin{pmatrix} a_0 \\ a_1 \\ \vdots \\ a_{n-1} \end{pmatrix}, Y = \begin{pmatrix} y_0 \\ y_1 \\ \vdots \\ y_{n-1} \end{pmatrix}, A = (x_j^k)_{0 \leq j, k \leq n-1}$$

$\det A$ est un déterminant de Vandermonde (transposé), *donc* A est inversible et $AX = Y$ si et seulement si $X = A^{-1}Y$. Le système a une unique solution, il existe un unique polynôme répondant à la question. Évidemment, cette méthode ne donne pas le polynôme, en tout cas pas sans résoudre le système, ce que nous ne ferons pas.

4 Déterminant d'un endomorphisme

4.1 Construction

Proposition 14. Soit E un $\mathbb{K}$ -espace vectoriel de dimension n . Soit $f \in \mathcal{L}(E)$. Il existe un unique $\lambda \in \mathbb{K}$ tel que pour tout $\varphi \in \Lambda_n(E)$, pour tous $x_1, \dots, x_n \in E$,

$$\varphi(f(x_1), \dots, f(x_n)) = \lambda \varphi(x_1, \dots, x_n)$$

Définition 8. Ce scalaire λ est appelé le déterminant de f . Il est noté $\det f$.

Démonstration. À toute application $\varphi \in \Lambda_n(E) \setminus \{0\}$, associons l'application $\bar{\varphi}$ définie pour $x_1, \dots, x_n \in E$ par

$$\bar{\varphi}(x_1, \dots, x_n) = \varphi(f(x_1), \dots, f(x_n))$$

$\bar{\varphi}$ est encore une forme n -linéaire alternée sur E . De plus, si $\varphi = 0$ alors $\bar{\varphi} = 0$. Donc, puisque $\Lambda_n(E)$ est une droite, il existe $\alpha \in \mathbb{K}$ tel que

$$\bar{\varphi} = \alpha\varphi$$

Il s'agit maintenant de prouver que α ne dépend pas de φ . Pour cela, soient $\varphi, \psi \in \Lambda_n(E)$. Prenons φ, ψ non nulles. Il existe $\alpha, \beta \in \mathbb{K}$ tels que

$$\bar{\varphi} = \alpha\varphi \text{ et } \bar{\psi} = \beta\psi$$

Mais comme $\varphi \neq 0$, il existe aussi $\gamma \in \mathbb{K}^*$ tel que

$$\psi = \gamma\varphi$$

d'où, facilement,

$$\bar{\psi} = \gamma\bar{\varphi}$$

De là,

$$\beta\psi = \beta\gamma\varphi = \gamma\alpha\varphi$$

Comme γ et φ sont non nuls, il vient $\alpha = \beta$. $\square$

Exercice. Pourquoi ce scalaire λ est-il unique ? Indication : supposer qu'il y en a eux.

Remarque 7. $\det f$ est caractérisé par $\forall \varphi \in \Lambda_n(E), \forall x_1, \dots, x_n \in E$,

$$\varphi(f(x_1), \dots, f(x_n)) = (\det f)\varphi(x_1, \dots, x_n)$$

4.2 Calcul pratique

Proposition 15. Soit $f \in \mathcal{L}(E)$. Soit $\mathcal{B}$ une base de E . Soit A la matrice de f dans la base $\mathcal{B}$. Alors

$$\det f = \det A$$

Démonstration. $\varphi = \det_{\mathcal{B}}$ est une forme n -linéaire alternée. On en déduit

$$\bar{\varphi} = (\det f)\varphi$$

On applique aux vecteurs de la base :

$$\bar{\varphi}(\mathcal{B}) = (\det f)\varphi(\mathcal{B})$$

ou encore

$$\varphi(f(\mathcal{B})) = \det f$$

puisque $\varphi(\mathcal{B}) = 1$. Finalement :

$$\det f = \det_{\mathcal{B}}(f(\mathcal{B}))$$

ou encore :

$$\det f = \det(\text{mat}_{\mathcal{B}}(f))$$

□

Exemple. Calculons le déterminant d'une symétrie. Soit f est une symétrie de E . On a

$$\det f = (-1)^{\dim G}$$

où G est l'ensemble des vecteurs changés en leur opposé par f . Il suffit pour cela d'écrire la matrice de f dans une base judicieuse (laquelle?).

4.3 Interprétation géométrique

On suppose ici que $\mathbb{K} = \mathbb{R}$. Nous ne justifierons pas les résultats qui vont suivre. Soit $\mathcal{B}$ la base canonique de $\mathbb{R}^n$. Soit $\mathcal{B}' = (u_1, \dots, u_n)$ une base de $\mathbb{R}^n$.

Si $n = 2$, $\det_{\mathcal{B}} \mathcal{B}'$ est l'aire du parallélogramme défini par les deux vecteurs de $\mathcal{B}'$.

Si $n = 3$, $\det_{\mathcal{B}} \mathcal{B}'$ est le volume du parallélépipède défini par les trois vecteurs de $\mathcal{B}'$.

Nous admettrons que ce résultat se généralise à une dimension quelconque. Soit alors $f \in \mathcal{L}(\mathbb{R}^n)$. Les calculs du paragraphe précédent montrent que :

$$\det f = \frac{\det_{\mathcal{B}}(f(u_1), \dots, f(u_n))}{\det_{\mathcal{B}}(u_1, \dots, u_n)}$$

En d'autres termes, f transforme les volumes dans un rapport constant, et le coefficient de proportionnalité est précisément le déterminant de f .

4.4 Propriétés de morphisme du déterminant

Proposition 16. Soient f et g deux endomorphismes d'un $\mathbb{K}$ -espace vectoriel de dimension finie E . On a

$$\det(g \circ f) = \det g \times \det f$$

Proposition 17. Soient $A, B \in \mathcal{M}_n(\mathbb{K})$. On a

$$\det(AB) = \det A \times \det B$$

Démonstration. Ces deux théorèmes sont évidemment deux interprétations d'une même chose. Soit $\mathcal{B}$ une base de E . Soit φ l'application $\det_{\mathcal{B}}$. On a

$$\varphi(g \circ f(\mathcal{B})) = \det(g \circ f)\varphi(\mathcal{B}) = \det(g \circ f)$$

Mais

$$\varphi(g \circ f(\mathcal{B})) = \varphi(g(f(\mathcal{B}))) = (\det g)\varphi(f(\mathcal{B})) = (\det g)(\det f)\varphi(\mathcal{B}) = \det g \det f$$

d'où le résultat. $\square$

Proposition 18. *Soit $f \in \mathcal{L}(E)$. Alors,*

$$f \in GL(E) \iff \det f \neq 0$$

Démonstration. Soit $\mathcal{B}$ une base de E . Soit $A = \text{mat}_{\mathcal{B}} f$. Alors

$$f \in GL(E) \iff A \in \mathcal{M}_n(\mathbb{K}) \iff \det A \neq 0$$

$\square$

Proposition 19. *Soit E un $\mathbb{K}$ -espace vectoriel de dimension $n \geq 1$. L'application $\det : GL(E) \rightarrow \mathbb{K}^*$ est un morphisme surjectif du groupe $(GL(E), \circ)$ sur le groupe $(\mathbb{K}^*, \times)$.*

Proposition 20. *Soit $n \geq 1$. L'application $\det : GL_n(\mathbb{K}) \rightarrow \mathbb{K}^*$ est un morphisme surjectif du groupe $(GL_n(\mathbb{K}), \times)$ sur le groupe $(\mathbb{K}^*, \times)$.*

Démonstration. Ces deux résultats sont deux versions d'un même théorème. Il reste seulement à prouver la surjectivité. Faisons-le pour les matrices. Soit $\alpha \in \mathbb{K}^*$. Soit $A \in \mathcal{M}_n(\mathbb{K})$ la matrice diagonale dont tous les éléments diagonaux sont égaux à 1, sauf celui ligne 1, colonne 1, qui vaut α . Alors, $\det A = \alpha$. $\square$

Remarque 8. Le cas $n = 0$ est un peu à part. Un $\mathbb{K}$ -espace vectoriel de dimension 0 est $E = \{0\}$. Il possède une unique base qui est la famille vide $\mathcal{B} = ()$. Le seul endomorphisme de E est $f = 0$, et on a $\det f = 1$ car

$$\det f = \det_{\mathcal{B}} f(\mathcal{B}) = \det_{\mathcal{B}} \mathcal{B} = 1$$

Du point de vue matriciel, le déterminant de la « matrice vide » est 1.

Proposition 21. *Deux matrices semblables ont le même déterminant.*

Démonstration. Soient $A, B \in \mathcal{M}_n(\mathbb{K})$ deux matrices semblables. Il existe donc $P \in GL_n(\mathbb{K})$ telle que

$$B = P^{-1}AP$$

De là,

$$\det B = \det(P^{-1}AP) = (\det P)^{-1} \det A \det P = \det A$$

□

4.5 Orientation d'un espace vectoriel réel

Soit E un $\mathbb{R}$ -espace vectoriel de dimension finie. Notons $\mathcal{E}$ l'ensemble des bases de E . On définit une relation binaire $\sim$ sur E comme suit :

Définition 9. Étant données deux bases $\mathcal{B}$ et $\mathcal{B}'$ de E , $\mathcal{B} \sim \mathcal{B}'$ si et seulement $\det f > 0$, où f est l'unique automorphisme de E qui envoie $\mathcal{B}$ sur $\mathcal{B}'$.

Proposition 22. La relation $\sim$ est une relation d'équivalence sur E .

Démonstration. L'identité envoie toute base $\mathcal{B}$ sur elle-même, et $\det id = 1 > 0$. Donc, $\mathcal{B} \sim \mathcal{B}$.

Si $f \in GL(E)$ envoie $\mathcal{B}$ sur $\mathcal{B}'$, et $\det f > 0$, alors f^{-1} envoie $\mathcal{B}'$ sur $\mathcal{B}$ et

$$\det(f^{-1}) = \frac{1}{\det f} > 0$$

Donc, $\sim$ est symétrique.

La transitivité est laissée en exercice, il suffit de constater que la composée de deux automorphismes de déterminant strictement positif en est encore un. □

Proposition 23. Lorsque $\dim E > 0$, la relation $\sim$ possède exactement deux classes d'équivalence.

Démonstration. Soit $f_0 \in GL(E)$ telle que $\det f_0 < 0$.

Soit $\mathcal{B}_1$ une base de E . Notons $\mathcal{C}_1$ la classe de $\mathcal{B}_1$ modulo $\sim$. Soit $\mathcal{B}_2 = f_0(\mathcal{B}_1)$. Soit $\mathcal{C}_2$ la classe de $\mathcal{B}_2$. Comme $\det f_0 < 0$, on a $\mathcal{B}_1 \not\sim \mathcal{B}_2$, donc $\mathcal{C}_1 \neq \mathcal{C}_2$. On a au moins deux classes d'équivalence, celle de $\mathcal{B}_1$ et celle de $\mathcal{B}_2$.

Soit maintenant $\mathcal{B}$ une base quelconque de E . Supposons que $\mathcal{B} \notin \mathcal{C}_1$, c'est à dire $\mathcal{B} \not\sim \mathcal{B}_1$. Soit f l'unique automorphisme de E qui envoie $\mathcal{B}_1$ sur $\mathcal{B}$. On a donc $\det f < 0$. Mais $f \circ f_0^{-1}$ envoie $\mathcal{B}_2$ sur $\mathcal{B}$ et

$$\det(f \circ f_0^{-1}) = \det f \det(f_0^{-1}) > 0$$

en tant que produit de deux nombres négatifs. Donc $\mathcal{B} \in \mathcal{C}_2$, on a au plus deux classes. □

Définition 10. Ces deux classes sont appelées les *orientations* de E . *Orienter* E , c'est choisir l'une des deux classes.

Les bases de la classe choisie sont dites *directes*. Les bases de l'autre classe sont dites *indirectes*.

Remarque 9. Une fois l'espace orienté, on passe d'une base directe à une base directe par un automorphisme de déterminant > 0 , d'une base directe à une base indirecte par un automorphisme de déterminant < 0 , et d'une base indirecte à une base indirecte par un automorphisme de déterminant > 0 .

Remarque 10. Le choix de l'orientation est a priori arbitraire. Cependant, en dimensions 2 et 3, il y a des préférences pour ce choix. Sens trigonométrique dans le plan, règle des trois doigts dans l'espace.

5 Inversion des matrices

5.1 Définitions, Notations

Rappelons la définition ci-dessous.

Définition 11. Soit A une matrice carrée. On appelle comatrice de A la matrice dont les coefficients sont les cofacteurs de A .

Notation. On note $\text{Com } A$ la comatrice de A et $\tilde{A}$ la transposée de la comatrice de A . $\tilde{A}$ est aussi appelée la *transcomatrice* de A .

Exercice. Calculer la transposée de la comatrice de $\begin{pmatrix} 1 & 4 & 7 \\ 2 & 5 & 8 \\ 3 & 6 & 9 \end{pmatrix}$.

5.2 Produit d'une matrice et de sa transcomatrice

Proposition 24. Soit $A \in \mathcal{M}_n(\mathbb{K})$. On a :

$$A\tilde{A} = \tilde{A}A = (\det A)I_n$$

Démonstration. On calcule $A\tilde{A}$:

- Soit $i \in \llbracket 1, n \rrbracket$: on a

$$\begin{aligned}(A\tilde{A})_{ii} &= \sum_{k=1}^n A_{ik}(\tilde{A})_{ki} \\ &= \sum_{k=1}^n A_{ik} \operatorname{Cof}(A, i, k) \\ &= \det A\end{aligned}$$

- Soient $i, j \in \llbracket 1, n \rrbracket, i \neq j$:

$$(A\tilde{A})_{ij} = \sum_{k=1}^n A_{ik} \operatorname{Cof}(A, j, k)$$

Soit B la matrice fabriquée à partir de A comme suit : les lignes de B sont identiques aux lignes de A , sauf la ligne j de B que l'on prend égale à la ligne i de A . B a donc deux lignes identiques, d'où $\det B = 0$. De plus, $A_{ik} = B_{ik} = B_{jk}$, et $\operatorname{Cof}(A, j, k) = \operatorname{Cof}(B, j, k)$, car la ligne j de B , la seule qui diffère des lignes de A , n'intervient pas dans ce calcul de cofacteur. Donc,

$$(A\tilde{A})_{ij} = \sum_{k=1}^n B_{jk} \operatorname{Cof}(B, j, k) = \det B = 0$$

□

5.3 Application au calcul de l'inverse

Proposition 25. Soit $A \in \mathcal{M}_n(\mathbb{K})$, $\det A \neq 0$. Alors :

$$A^{-1} = \frac{1}{\det A} \tilde{A}$$

Démonstration. C'est une conséquence immédiate du théorème précédent. □

Exemple. Soit $A = \begin{pmatrix} a & c \\ b & d \end{pmatrix}$ une matrice 2×2 . On a $A^{-1} = \frac{1}{\det A} \begin{pmatrix} d & -c \\ -b & a \end{pmatrix}$.

Exercice. Soit $A = \begin{pmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \\ 7 & 8 & 8 \end{pmatrix}$. Calculer $\det A$, puis $\tilde{A}$, et enfin l'inverse de A .

Chapitre 27

Systemes lineaires

1 Sous-espaces affines d'un espace vectoriel

1.1 Espaces affines - Points et vecteurs

Voici la définition générale d'un espace affine.

Définition 1. Soit E un $\mathbb{K}$ -espace vectoriel. Soit $\mathcal{E}$ un ensemble dont les éléments sont appelés des *points*. On suppose donnée une application $\oplus : \mathcal{E} \times E \longrightarrow \mathcal{E}$.

On dit que $(\mathcal{E}, E, +)$ est un $\mathbb{K}$ -espace affine lorsque

- Pour tout $a \in \mathcal{E}$ l'application $x \longmapsto a \oplus x$ est une bijection de E sur $\mathcal{E}$.
- Pour tout $a \in \mathcal{E}$, pour tous $x, y \in E$, $(a \oplus x) \oplus y = a \oplus (x + y)$.

Nous ne chercherons pas dans ce qui suit la généralité. Contentons-nous de remarquer que si E est un $\mathbb{K}$ -espace vectoriel, alors E lui-même est muni de façon évidente d'une structure d'espace affine en posant pour tous $a, x \in E$, $a \oplus x = a + x$.

Ainsi, selon les circonstances, les éléments de E sont vus comme des points ou comme des vecteurs. Le *translaté* d'un « point » $a \in E$ par le « vecteur » $u \in E$ est le « point » $a + u \in E$, la somme classique de a et de u .

Définition 2. Soient $a, b \in E$. Le vecteur d'origine a et d'extrémité b est le vecteur

$$\overrightarrow{ab} = b - a$$

Concrètement, au niveau où nous nous plaçons, on a simplement un jeu de notations :

$$b = a + u \iff \overrightarrow{ab} = u \iff b - a = u$$

On a par exemple la *relation de Chasles*. Si $a, b, c \in E$,

$$c - a = (b - a) + (c - b)$$

ou encore

$$\overrightarrow{ac} = \overrightarrow{ab} + \overrightarrow{bc}$$

1.2 Translations - sous-espaces affines

Définition 3. Soit E un $\mathbb{K}$ -espace vectoriel. Soit $u \in E$. On appelle *translation* de vecteur u l'application $\tau_u : E \rightarrow E$ définie pour tout $x \in E$ par

$$\tau_u(x) = x + u$$

Proposition 1. L'ensemble $\mathcal{T}(E)$ des translations de E , muni de la composition des applications, est un groupe isomorphe au groupe additif de E .

Démonstration. Pour tous vecteurs u, v de E , pour tout $x \in E$, on a

$$\tau_{u+v}(x) = x + (u + v) = (x + v) + u = (\tau_u \circ \tau_v)(x)$$

Ainsi,

$$\tau_{u+v} = \tau_u \circ \tau_v$$

On vérifie alors sans peine que l'application $\tau : E \rightarrow \mathcal{T}(E)$ définie par $u \mapsto \tau_u$ est un isomorphisme de groupes. $\square$

Notation. Pour tout $a \in E$ et tout sev F de E , on note

$$a + F = \{a + u, u \in F\}$$

Définition 4. Soit $\mathcal{F}$ une partie de E . On dit que $\mathcal{F}$ est un sous-espace affine de E lorsqu'il existe un point $a \in E$ et un sous-espace vectoriel F de E tels que

$$\mathcal{F} = a + F$$

Proposition 2. Avec les notations ci-dessus, le sous-espace vectoriel F est caractérisé de façon unique. Quant à a , on peut prendre n'importe quel point de $\mathcal{F}$.

Démonstration. Supposons $\mathcal{F} = a + F = b + G$ où $a, b \in E$ et F, G sont des sev de E . En remarquant que $0 \in F$, on obtient que $a \in \mathcal{F}$. Il en est de même évidemment pour b .

Toujours avec $0 \in F$, on voit que $a \in b + G$, donc que $\overrightarrow{ba} \in G$ (et aussi $\overrightarrow{ab} = -\overrightarrow{ba}$, et de même pour F).

Soit $x \in F$. On a $a + x \in b + G$, c'est à dire il existe $y \in G$ tels que $a + x = b + y$. D'où $x = (b - a) + y \in G$. Donc, $F \subset G$. L'inclusion réciproque se montre de la même façon, donc $F = G$. $\square$

Définition 5. On dit que $\mathcal{F} = a + F$ est le sous-espace affine de E passant par a et dirigé par F .

Remarque 1. La dimension d'un sous-espace affine est par définition celle de sa direction. On parle ainsi de droite affine, de plan affine, etc.

1.3 Parallélisme

Définition 6. Soient $\mathcal{F}$ et $\mathcal{G}$ deux sous-espaces affines de E de directions respectives F et G . On dit que $\mathcal{F}$ est parallèle à $\mathcal{G}$ lorsque $F = G$.

Remarque 2. La relation de parallélisme est clairement une relation d'équivalence sur l'ensemble des sous-espaces affines de E . De plus, deux sous-espaces affines de E parallèles ont la même dimension.

Cela peut paraître parfois gênant : on ne peut pas, par exemple, dire qu'une droite est parallèle à un plan.

Il est possible de changer notre définition en décrétant que $\mathcal{F}$ est parallèle à $\mathcal{G}$ lorsque $F \subset G$ ou $G \subset F$. Dans ce cas on peut parler de droite parallèle à un plan, mais ...on perd la transitivité du parallélisme.

1.4 Intersection de sous-espaces affines

Proposition 3. Soit $(\mathcal{F}_i)_{i \in I}$ une famille de sous-espaces affines de E . Alors, si l'intersection $\mathcal{F} = \bigcap_{i \in I} \mathcal{F}_i$ est non vide, c'est un sous-espace affine de E .

Démonstration. Soit $a \in \mathcal{F}$. Soit $F = \bigcap_{i \in I} F_i$. Alors, on montre facilement que $\mathcal{F} = a + F$. $\square$

1.5 Barycentres-Parties convexes

Proposition 4. Soient $a_1, \dots, a_n$ n points de l'espace vectoriel E . Soient $\lambda_1, \dots, \lambda_n \in \mathbb{K}$ tels que $\sum_{k=1}^n \lambda_k \neq 0$. Il existe un unique point $\omega \in E$ tel que

$$\sum_{k=1}^n \lambda_k \overrightarrow{\omega a_k} = 0$$

Définition 7. Le point ω est appelé le barycentre des a_k affectés des coefficients λ_k .

Démonstration. Soit $\omega \in E$. On a

$$\sum_{k=1}^n \lambda_k \overrightarrow{\omega a_k} = \sum_{k=1}^n \lambda_k (\overrightarrow{Oa_k} - \overrightarrow{O\omega})$$

où O est un point quelconque de E . Notons $s = \sum_{k=1}^n \lambda_k$. Le scalaire s est non nul. Le point ω convient si et seulement si

$$s\overrightarrow{O\omega} = \sum_{k=1}^n \lambda_k \overrightarrow{Oa_k}$$

d'où l'existence et l'unicité de ω . $\square$

Proposition 5. Soit $\mathcal{F}$ une partie non vide du $\mathbb{K}$ -espace vectoriel E . Alors $\mathcal{F}$ est un sous-espace affine de E si et seulement si tout barycentre d'éléments de $\mathcal{F}$ est encore un élément de $\mathcal{F}$.

Démonstration. Supposons que $\mathcal{F} = a + F$ est un sous-espace affine de E , passant par a , dirigé par F . Soit $n \geq 1$. Soient $a_i = a + x_i$, $1 \leq i \leq n$ n points de $\mathcal{F}$. Soient $\lambda_1, \dots, \lambda_n$ n scalaires dont la somme s est non nulle. Soit ω le barycentre des a_i affectés des coefficients λ_i . On a

$$\omega = \frac{1}{s} \sum_i \lambda_i (a + x_i) = \frac{1}{s} (sa + \sum_i \lambda_i x_i) = a + z$$

où $z = \frac{1}{s} \sum_i \lambda_i x_i \in F$. Donc, $\omega \in \mathcal{F}$.

Supposons inversement que $\mathcal{F}$ est stable par barycentres. Soit $a \in \mathcal{F}$. Soit

$$F = \{p - a, p \in \mathcal{F}\}$$

On a donc $\mathcal{F} = a + F$, et il s'agit de prouver que F est un sev de E .

- Tout d'abord, $a - a = 0 \in F$, donc F est non vide.
- Soient $x = p - a, y = q - a \in F$. Soit $z = x + y$. On a donc $z = (p + q - a) - a$. Mais $p + q - a$ est le barycentre de p, q, a affectés des coefficients 1, 1, -1. Donc $p + q - a \in \mathcal{F}$ et $z \in F$.
- Enfin, soit $x = p - a \in F$. Soit $\lambda \in \mathbb{K}$. On a $\lambda x = \lambda(p - a) = \lambda p - (\lambda - 1)a - a = q - a$ où q n'est autre que le barycentre de p et a affectés des coefficients λ et $1 - \lambda$. Donc, $q \in \mathcal{F}$ et $\lambda x \in F$.

Ainsi, F est bien un sev de E . $\square$

Définition 8. Soit E un $\mathbb{K}$ -espace vectoriel. Soient $x, y \in E$. Le segment x, y est l'ensemble

$$[x, y] = \{(1-t)x + ty, t \in [0, 1]\}$$

Remarque 3. Soit $z \in E$. Alors, $z \in [x, y]$ si et seulement si il existe $t \in [0, 1]$ tel que

$$z = (1-t)x + ty$$

Ceci s'écrit encore $z - x = t(y - x)$ ou encore

$$\overrightarrow{xz} = t\overrightarrow{xy}$$

où $t \in [0, 1]$.

Définition 9. Soit A une partie du $\mathbb{K}$ -ev E . A est dite *convexe* lorsque pour tous $x, y \in A$ on a $[x, y] \subset A$.

Proposition 6. Soit A une partie du $\mathbb{K}$ -espace vectoriel E . Alors A est convexe si et seulement si tout barycentre d'éléments de A à coefficients positifs est encore un élément de A .

Démonstration. Supposons A stable par barycentres à coefficients positifs. Soient $x, y \in A$. Le segment $[x, y]$ est justement constitué des barycentres de x et de y à coefficients positifs. Donc $[x, y] \subset A$ et A est convexe.

Supposons inversement que A est convexe. Montrons par récurrence sur n que pour tout $n \geq 1$, A est stable par barycentre de n points à coefficients positifs.

Pour $n = 1$, il n'y a rien à montrer.

Pour $n = 2$, c'est la définition de la convexité.

Soit $n \geq 1$. Supposons la propriété vraie pour n . Soient $a_1, \dots, a_{n+1}$ $n+1$ points de A . Soient $\lambda_1, \dots, \lambda_{n+1}$ $n+1$ réels positifs non tous nuls. Soit s leur somme. Soit p le barycentre des a_k affectés des coefficients λ_k . Si $\lambda_1 = \dots = \lambda_n = 0$, on a $p = a_{n+1} \in \mathcal{F}$. Sinon, on a

$$p = \frac{1}{s} \sum_{k=1}^{n+1} \lambda_k a_k = \left(1 - \frac{\lambda_{n+1}}{s}\right) b + \frac{\lambda_{n+1}}{s} a_{n+1}$$

où

$$b = \frac{1}{\sum_{k=1}^n \lambda_k} \sum_{k=1}^n \lambda_k a_k$$

Par l'hypothèse de récurrence, on a $b \in \mathcal{F}$. Or p est un barycentre de a et b , donc $p \in \mathcal{F}$.

□

1.6 hyperplans affines

Soit E un $\mathbb{K}$ -espace vectoriel de dimension n . Soit $\mathcal{B} = (e_1, \dots, e_n)$ une base de E . Soit $F : \mathbb{K}^n \rightarrow \mathbb{K}$. On peut alors considérer l'ensemble

$$A = \{x \in E, F(x_1, \dots, x_n) = 0\}$$

où $x = \sum_{k=1}^n x_k e_k$. Pour écrire les choses plus simplement, on note juste

$$(A) \quad F(x_1, \dots, x_n) = 0$$

et on dit que l'on a là une *équation cartésienne* de A dans la base $\mathcal{B}$.

Définition 10. Soit E un $\mathbb{K}$ -espace vectoriel de dimension finie n . On appelle hyperplan affine de E tout sous-espace affine de E de dimension $n - 1$.

Proposition 7. Une base de E étant choisie, tout hyperplan affine de E admet une équation cartésienne du type

$$\sum_{k=1}^n a_k x_k = b$$

où les a_k ne sont pas tous nuls. Deux équations représentent le même hyperplan si et seulement si leurs coefficients (second membre y compris) sont proportionnels. Enfin, toute équation de ce type est celle d'un hyperplan affine.

Démonstration. Remarquons tout d'abord qu'une forme linéaire non nulle sur E est surjective. En effet, son image est un sev non nul de $\mathbb{K}$. Comme $\mathbb{K}$ est de dimension 1, ce sev est donc $\mathbb{K}$.

Maintenant, toute équation du type $\sum_{k=1}^n a_k x_k = b$ s'écrit aussi $f(x) = b$ où f est une forme linéaire non nulle sur E , ou encore $f(x) = f(a)$ puisque f est surjective, ou encore $f(x - a) = 0$ c'est à dire $x - a \in H$ où $H = \ker f$ est un hyperplan vectoriel de E . Le reste de la preuve consiste à écrire ce que l'on sait sur les équations des hyperplans vectoriels. Le lecteur est amené à relire la section « dualité » du cours d'algèbre linéaire. $\square$

Exemple. Plaçons nous dans $\mathbb{R}^3$. L'équation

$$(\mathcal{P}) \quad 3x - 2y + 5z = 6$$

est l'équation d'un plan affine $\mathcal{P}$. La direction de $\mathcal{P}$ est le plan vectoriel d'équation

$$(P) \quad 3x - 2y + 5z = 0$$

On a donc $\mathcal{P} = a + P$, où a est n'importe quel point de $\mathcal{P}$, par exemple $a = (1, 1, 1)$.

2 Rappels sur les systèmes linéaires

2.1 Position du problème

Un système linéaire de p équations à q inconnues est un système d'équations du type

$$(\mathcal{S}) \sum_{j=1}^q a_{ij}x_j = b_i, \quad 1 \leq i \leq p$$

Un tel système peut plus simplement être écrit

$$(\mathcal{S}) \quad AX = B$$

où $A = (a_{ij}) \in \mathcal{M}_{p,q}(\mathbb{K})$, $X = (x_j) \in \mathcal{M}_{q,1}(\mathbb{K})$ et $B = (b_i) \in \mathcal{M}_{p,1}(\mathbb{K})$.

Définition 11. Avec les notations ci-dessus :

- La matrice A est la matrice du système.
- La matrice B est le second membre.
- La matrice X est la matrice des inconnues.
- Le rang du système est le rang de A .
- Le système est dit compatible lorsqu'il a au moins une solution, et incompatible sinon.
- Le système est dit homogène lorsque $B = 0$: tout système homogène est compatible, puisqu'il admet au moins la solution triviale $X = 0$.

Remarque 4. Soit $(\mathcal{S}) \quad AX = B$ un système de p équations à q inconnues, de rang r . Alors

$$r \leq \min(p, q)$$

Proposition 8. *Un système de une équation à q inconnues, ayant un premier membre non trivial, est toujours compatible. L'ensemble de ses solutions est un hyperplan affine de $\mathbb{K}^q$*

Démonstration. Déjà faite, à la fin de la section précédente. $\square$

2.2 Interprétations

Il y a plusieurs façons d'interpréter l'ensemble des solutions d'un système. Chacune a ses avantages et ses inconvénients. En voici quelques unes.

- Résoudre p équations à q inconnues : vue naïve du problème, aucune interprétation, on résout et c'est tout.
- Résoudre l'équation matricielle $AX = B$ à une inconnue X : vue matricielle, finalement on a 1 équation à 1 inconnue.
- Résoudre l'équation $f(x) = b$ où $f \in \mathcal{L}(\mathbb{K}^q, \mathbb{K}^p)$: et là, tout l'arsenal du cours d'algèbre linéaire est à notre disposition.
- Intersection d'hyperplans : vue géométrique.

Selon les besoins de telle ou telle démonstration, nous utiliserons l'interprétation qui nous convient le mieux.

2.3 Structure des solutions - Systèmes homogènes

Proposition 9. *Soit*

$$(\mathcal{H}) \quad AX = 0$$

un système de p équations à q inconnues, de rang r . L'ensemble $Sol_{\mathcal{H}}$ des solutions de $(\mathcal{H})$ est un sous-espace vectoriel de $\mathbb{K}^q$ de dimension $q - r$.

Démonstration. Soit $f \in \mathcal{L}(\mathbb{K}^q, \mathbb{K}^p)$, de matrice A dans les bases canoniques de $\mathbb{K}^q$ et $\mathbb{K}^p$. L'ensemble des solutions de $(\mathcal{H})$ est

$$Sol_{\mathcal{H}} = \ker f$$

D'où le résultat par le théorème du rang. $\square$

Remarque 5. Si $p < q$ (moins d'équations que d'inconnues), alors $r \leq p < q$ donc $q - r > 0$: il y a des solutions non triviales. En revanche, si $q \leq p$, on ne peut rien dire sans un calcul précis du rang r .

2.4 Structure des solutions - Cas général

Proposition 10. *Soit*

$$(\mathcal{E}) \quad AX = B$$

un système de p équations à q inconnues de rang r . Soit

$$(\mathcal{H}) \quad AX = 0$$

le système homogène associé.

SI le système $(\mathcal{E})$ est compatible, alors l'ensemble $Sol_{\mathcal{E}}$ des solutions de $(\mathcal{E})$ est un sous-espace affine de $\mathbb{K}^q$ dont la direction est $Sol_{\mathcal{H}}$.

Démonstration. Soit x_0 une solution de $(\mathcal{E})$. Soit $x \in \mathbb{K}^q$. Alors, x est solution de $(\mathcal{E})$ si et seulement si $x - x_0$ est solution de $(\mathcal{H})$. $\square$

Exercice. L'intérêt de cet exercice n'est pas de résoudre des systèmes mais de regarder la *structure* de leurs ensembles de solutions. Examiner les systèmes très simples suivants : trouver leur matrice, leur rang, puis résoudre et regarder la structure de l'ensemble des solutions.

1. $\begin{cases} x + y = 5 \\ x - y = 3 \end{cases}$
2. $\begin{cases} x + y + z = 5 \\ x - y - z = 3 \end{cases}$
3. $\begin{cases} x + y = 1 \\ x + 2y = 3 \\ 3x - y = \lambda \end{cases}$
4. $\begin{cases} x & & & + 4t = 1 \\ 2x + 3y - z & & & = 3 \\ x + 3y - z - 4t = -2 \end{cases}$

3 Systèmes de Cramer

3.1 Définitions

Définition 12. Un système $AX = B$ est dit de Cramer lorsque

- Il y a autant d'équations que d'inconnues.
- La matrice A du système est inversible.

En résumé, avec les notations précédentes, lorsque $p = q = r$.

Proposition 11. Soit $A \in \mathcal{M}_n(\mathbb{K})$. Il y a équivalence entre

- A est inversible.
- $\forall B \in \mathcal{M}_{n,1}(\mathbb{K})$, le système $AX = B$ a une unique solution.
- $\forall B \in \mathcal{M}_{n,1}(\mathbb{K})$, le système $AX = B$ est compatible.
- $\forall B \in \mathcal{M}_{n,1}(\mathbb{K})$, le système $AX = B$ a au plus une solution.
- La seule solution du système $AX = 0$ est $X = 0$.

Démonstration. On interprète en termes de vecteurs et d'applications linéaires. Soit E un $\mathbb{K}$ -espace vectoriel de dimension n , muni d'une base $\mathcal{B}$. Soit $f \in \mathcal{L}(E)$ telle que $\text{mat}_{\mathcal{B}} f = A$. Nous avons à montrer qu'il y a équivalence entre

- f est bijective.

- f est bijective (eh oui, encore!).
- f est surjective.
- f est injective.
- $\ker f = \{0\}$.

Mais tout cela, nous le savons déjà. $\square$

3.2 Formules de Cramer

Proposition 12. Soit $A \in GL_n(\mathbb{K})$. Soit (S) $AX = B$ un système de Cramer. Soit $(x_1, \dots, x_n)$ l'unique solution de (S) . Pour $1 \leq j \leq n$, soit A_j la matrice obtenue en remplaçant la j ème colonne de A par B . On a alors

$$\forall j \in \llbracket 1, n \rrbracket, x_j = \frac{\det A_j}{\det A}$$

Démonstration. Soient $C_1, \dots, C_n$ les colonnes de A , vues comme vecteurs de $\mathbb{K}^n$. Soit $\mathcal{B}$ la base canonique de $\mathbb{K}^n$. Le système s'écrit alors

$$x_1 C_1 + \dots + x_n C_n = B$$

Or,

$$\begin{aligned} \det A_j &= \det_{\mathcal{B}}(C_1, \dots, B, \dots, C_n) \\ &= \det_{\mathcal{B}}(C_1, \dots, \sum_{i=1}^n x_i C_i, \dots, C_n) \\ &= x_j \det_{\mathcal{B}}(C_1, \dots, C_n) = x_j \det A \end{aligned}$$

tous les autres termes étant nuls à cause de la propriété d'alternance. $\square$

Exemple. Soit le système $\begin{cases} ax + cy = \alpha \\ bx + dy = \beta \end{cases}$ où $ad - bc \neq 0$. L'unique solution de ce système est

$$x = \frac{\alpha d - \beta c}{ad - bc}, \quad y = \frac{-\alpha b + \beta a}{ad - bc}$$

Exercice. Résoudre le système $\begin{cases} x + 2y - z = 1 \\ 3x - y - z = 0 \\ 2x + 2y + z = 0 \end{cases}$

La matrice du système est

$$A = \begin{pmatrix} 1 & 2 & -1 \\ 3 & -1 & -1 \\ 2 & 2 & 1 \end{pmatrix}$$

Le déterminant de A vaut $-17 \neq 0$. On a donc bien un système de Cramer. La matrice A_1 est

$$A_1 = \begin{pmatrix} 1 & 2 & -1 \\ 0 & -1 & -1 \\ 0 & 2 & 1 \end{pmatrix}$$

Le déterminant de A_1 est 1. On a donc

$$x = \frac{\det A_1}{\det A} = \frac{-1}{17}$$

On laisse au lecteur le calcul de y et z .

Remarque 6. Les formules de Cramer n'ont d'intérêt pratique que pour les petites valeurs de n (2 ou 3). Elles sont en revanche d'un intérêt théorique certain, puisqu'elles donnent la *forme* des solutions du système.

4 L'algorithme du Pivot de Gauss

Nous avons déjà parlé de l'algorithme du pivot de Gauss. Nous allons ici formaliser de façon plus détaillée cet algorithme.

Soit le système (S) $AX = B$ de p équations à q inconnues, de rang r . On se propose de transformer ce système en un système équivalent, de même rang, et de matrice plus « simple ». La discussion va porter sur le rang de A .

4.1 Première étape

Si $r = 0$, alors $A = 0$ et on n'a rien à faire :

- Si $B \neq 0$, le système n'a aucune solution : $Sol_S = \emptyset$.
- Si $B = 0$, $Sol_S = \mathbb{K}^q$.

Si $r \neq 0$, il existe i, j tels que $a_{ij} \neq 0$. Quitte à échanger deux équations, et peut-être également deux inconnues, on peut supposer $i = j = 1$. On effectue les opérations suivantes :

- $L_1 \leftarrow \frac{1}{a_{11}} L_1$
- Pour $i = 2, \dots, p$, $L_i \leftarrow L_i - a_{i1} L_1$

On obtient ainsi un nouveau système équivalent au premier, et de même rang,

$$(S_1) \quad A_1 X = B_1$$

où

$$A_1 = \begin{pmatrix} 1 & a_{12}^1 & \dots & a_{1q}^1 \\ 0 & a_{22}^1 & \dots & a_{2q}^1 \\ \vdots & \vdots & & \vdots \\ 0 & a_{p2}^1 & \dots & a_{pq}^1 \end{pmatrix}$$

4.2 Étape générale

Supposons qu'au bout de k étapes, on soit parvenu à un système équivalent à (S) et de même rang,

$$(S_k) \quad A_k X = B_k$$

où

$$A_k = \begin{pmatrix} I_k & \cdots \\ O_{p-k,k} & \cdots \end{pmatrix}$$

Si $r = k$, alors on a

$$A_k = \begin{pmatrix} I_k & \cdots \\ O_{p-k,k} & O_{p-k,q-k} \end{pmatrix}$$

En effet, les $k = r$ premières colonnes de A_k sont indépendantes. Les autres colonnes sont donc combinaisons linéaires des k premières. En conséquence, on en déduit les solutions du système :

- Si $b_{r+1}^r = \dots = b_p^r = 0$, alors les solutions du système sont obtenues en choisissant arbitrairement les inconnues $x_{r+1}, \dots, x_q$. Les autres inconnues, $x_1, \dots, x_r$ sont alors uniquement déterminées.
- sinon, il n'y a pas de solution : le système est incompatible.

Si $r > k$, alors il existe $i, j > k$ tels que $a_{ij} \neq 0$ (sinon la matrice A_k serait de rang k). Quitte à échanger deux équations, voire deux inconnues, on peut supposer que $i = j = k + 1$. Le nombre $a_{(k+1)(k+1)} \neq 0$ est appelé le *pivot*. On effectue alors les opérations suivantes :

- $L_{k+1} \leftarrow \frac{1}{a_{(k+1)(k+1)}} L_{k+1}$
- $L_i \leftarrow L_i - a_{i(k+1)} L_{k+1}$ pour $i \neq k + 1$.

On obtient ainsi un nouveau système équivalent au premier, et de même rang,

$$(S_{k+1}) \quad A_{k+1} X = B_{k+1}$$

où

$$A_{k+1} = \begin{pmatrix} I_{k+1} & \cdots \\ O_{p-k-1,k+1} & \cdots \end{pmatrix}$$

4.3 Bon, et alors ?

Partant d'un système (S) de rang r , on fabrique des systèmes $(S_1), (S_2), \dots, (S_r)$ équivalents à (S) et tels que l'on connaisse les solutions de (S_r) . La méthode du pivot permet donc de résoudre le système (S) en exactement r étapes, où r est le rang du système.

Remarquons que le rang de (S) n'a pas à être connu a priori. La méthode du pivot termine au bout d'un certain nombre d'étapes. À la fin du calcul on en *déduit* le rang de (S) .

4.4 Un exemple

Considérons le système

$$(S) \begin{cases} x + 2y + 3z &= 6 \\ 4x + 5y + 6z &= 15 \\ 7x + 8y + 9z &= 24 \end{cases}$$

- Étape 1 : le pivot est $a_{11} = 1$. On effectue les opérations

$$L_1 \leftarrow \frac{1}{1}L_1$$

$$L_2 \leftarrow L_2 - 4L_1$$

$$L_3 \leftarrow L_3 - 7L_1$$

On obtient un système équivalent à (S)

$$(S_1) \begin{cases} x + 2y + 3z &= 6 \\ -3y - 6z &= -9 \\ -6y - 12z &= -18 \end{cases}$$

- Étape 2 : le pivot est $a_{22} = -3$. On effectue les opérations

$$L_2 \leftarrow \frac{1}{-3}L_2$$

$$L_1 \leftarrow L_1 - 2L_2$$

$$L_3 \leftarrow L_3 - (-6)L_2$$

On obtient un système équivalent à (S)

$$(S_2) \begin{cases} x - z &= 0 \\ y + 2z &= 3 \\ 0 &= 0 \end{cases}$$

La résolution est terminée. On prend z arbitrairement, on a x et y en fonction de z . Un triplet (x, y, z) est solution si et seulement si $x = z$ et $y = 3 - 2z$, ce qui s'écrit encore

$$(x, y, z) = (z, 3 - 2z, z) = (0, 3, 0) + z(1, -2, 1)$$

Si nous notons $a = (0, 3, 0)$ et $D = \langle (1, -2, 1) \rangle$, nous avons pour conclure

$$\text{Sol}_S = a + D$$

L'ensemble des solutions de (S) est la droite affine passant par a et dirigée par D .

5 Compléments

5.1 Calcul du déterminant d'une matrice carrée

Soit $A \in \mathcal{M}_n(\mathbb{K})$. L'algorithme du pivot permet de calculer le déterminant de A . Pour cela, on simule la résolution du système $AX = B$ avec un second membre B virtuel (inutile de l'écrire). À chaque étape de l'algorithme, on obtient un nouveau système dont on connaît le déterminant en fonction du déterminant du système précédent : on a multiplié une ligne par un scalaire, ce qui a multiplié le déterminant par ce même scalaire. On a peut-être aussi effectué des échanges de lignes, ce qui a changé le signe du déterminant. Si l'algorithme termine en strictement moins de n étapes, alors $\det A = 0$ puisque $\text{rg } A < n$. Sinon, le dernier système est $IX = B$, de déterminant 1. D'où la valeur de $\det A$.

5.2 Inversion d'une matrice carrée

La méthode du pivot peut s'appliquer à un second membre et à des inconnues « matricielles » : ceci revient à effectuer une résolution simultanée de systèmes ayant tous la même matrice mais des seconds membres différents. Prenons un exemple. Nous allons inverser en

trois étapes $A = \begin{pmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \\ 7 & 8 & 8 \end{pmatrix}$. On résout simultanément les trois systèmes $AX = \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}$,

$AX = \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix}$ et $AX = \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix}$, ce qui revient à résoudre les système d'inconnue matricielle

$X \in \mathcal{M}_3(\mathbb{R})$, $AX = I_3$. En d'autres termes, $X = A^{-1}$. Présentons tout cela dans des tableaux. On part du tableau

$$\begin{array}{ccc|ccc} 1 & 2 & 3 & 1 & 0 & 0 \\ 4 & 5 & 6 & 0 & 1 & 0 \\ 7 & 8 & 8 & 0 & 0 & 1 \end{array}$$

et on utilise l'algorithme du pivot.

Première étape : $L_2 \leftarrow L_2 - 4L_1$ et $L_3 \leftarrow L_3 - 7L_1$. On obtient

$$\begin{array}{ccc|ccc} 1 & 2 & 3 & 1 & 0 & 0 \\ 0 & -3 & -6 & -4 & 1 & 0 \\ 0 & -6 & -13 & -7 & 0 & 1 \end{array}$$

Deuxième étape : $L_2 \leftarrow -\frac{1}{3}L_2$ puis $L_1 \leftarrow L_1 - 2L_2$ et $L_3 \leftarrow L_3 + 6L_2$. On obtient :

$$\begin{array}{ccc|ccc} 1 & 0 & -1 & -\frac{5}{3} & \frac{2}{3} & 0 \\ 0 & 1 & 2 & \frac{4}{3} & -\frac{1}{3} & 0 \\ 0 & 0 & -1 & 1 & -2 & 1 \end{array}$$

Troisième et dernière étape, $L_3 \leftarrow -L_3$ puis $L_1 \leftarrow L_1 + L_3$ et $L_2 \leftarrow L_2 - 2L_3$. On obtient :

$$\begin{array}{c|c|c||c|c|c} 1 & 0 & 0 & -\frac{8}{3} & \frac{8}{3} & -1 \\ \hline 0 & 1 & 0 & \frac{10}{3} & -\frac{13}{3} & 2 \\ \hline 0 & 0 & 1 & -1 & 2 & -1 \end{array}$$

On lit l'inverse de la matrice A dans les trois dernières colonnes du tableau :

$$A^{-1} = \begin{pmatrix} -\frac{8}{3} & \frac{8}{3} & -1 \\ \frac{10}{3} & -\frac{13}{3} & 2 \\ -1 & 2 & -1 \end{pmatrix}$$

5.3 Évaluation de la complexité de la méthode du pivot

Soit à résoudre un système de p équations, q inconnues, rang r . Cette résolution se fait en r étapes. À chaque étape, on fait les opérations suivantes :

- Échange possible de 2 lignes : $q + 1$ échanges.
- Échange possible de 2 colonnes : p échanges.
- Normalisation d'une ligne : $p + 1$ divisions.
- Ajout à $p - 1$ lignes d'un multiple de la p ième : $(p - 1)((q + 1) + (q + 1))$ additions ou multiplications.

On fait donc au total de l'ordre de $2rpq$ opérations. Pour une système de Cramer $n \times n$, on a donc $2n^3$ opérations. Ce nombre est également le nombre d'opérations nécessaire pour calculer un déterminant avec la méthode du pivot. Pour inverser une matrice, il faut deux fois plus d'opérations, et on a donc de l'ordre de $4n^3$ opérations. Aussi non-intuitif que cela puisse paraître,

Inverser une matrice par l'algorithme du pivot a un coût du même ordre de grandeur que multiplier naïvement deux matrices.

Faisons pour clore ce chapitre une petite application numérique. La formule générale pour calculer un déterminant de taille n avec une somme sur les permutations de $\mathfrak{S}_n$ nécessite de l'ordre de $(n + 1)!$ opérations. Pour inverser une matrice A avec la formule

$$A^{-1} = \frac{1}{\det A} \tilde{A}$$

il faut calculer 1 déterminant de taille n et n^2 déterminants de taille $n - 1$, ce qui donne de l'ordre de $n^2 n!$ opérations. Notons $t_n = 4n^3$ et $t'_n = n^2 n!$.

Supposons donnée une machine effectuant 10^{10} opérations par seconde. Évaluons le temps t_n nécessaire à cette machine pour résoudre un système $n \times n$ de Cramer avec la méthode du pivot, et le temps t'_n avec la méthode naïve :

n	t_n	t'_n
3	$1.8 \cdot 10^{-8} s$	$5.4 \cdot 10^{-9} s$
5	$5 \cdot 10^{-8} s$	$3 \cdot 10^{-7} s$
10	$4 \cdot 10^{-7} s$	$3 \cdot 10^{-2} s$
15	$1.3 \cdot 10^{-6} s$	8 heures
20	$3.2 \cdot 10^{-6} s$	3085 ans
100	$4 \cdot 10^{-4} s$	$2.9 \cdot 10^{144}$ ans

Remarque 7. Prenons un ordinateur d'une puissance de 100 Watts. Pour inverser une matrice de taille 20 avec les formules naïves, il faudra une énergie de 10^{11} Joules. Sachant qu'une centrale nucléaire fournit une puissance de 10^9 Watts, cela représente la quantité d'énergie fournie par cette centrale pendant 5 minutes. C'est à dire la consommation d'environ un demi-kilo d'uranium enrichi...

Chapitre 28

Intégration

L'intégrale d'une fonction f entre deux réels a et b est un réel qui mesure l'aire sous la courbe de f . Mais qu'est-ce qu'une aire ?

Le mathématicien prend les choses ... à l'envers. Il définit l'aire comme l'intégrale de la fonction. Mais alors, me direz-vous, qu'est-ce que l'intégrale ? Le but de ce chapitre est de répondre à cette question.

Il existe différentes théories de l'intégration : intégrale de Riemann, intégrale de Lebesgue, intégrale de Stieljes, et bien d'autres. La théorie présentée ici est celle de l'intégrale de Riemann.

- Nous allons commencer par définir l'intégrale de fonctions très simples, les fonctions en escalier, qui sont tout simplement des fonctions constantes par morceaux.
- Nous montrerons ensuite que, d'une certaine façon, les fonctions en escalier sont *denses* dans un ensemble qui contient, entre-autres, les fonctions continues par morceaux.
- Cette propriété nous permettra d'étendre la définition de l'intégrale aux fonctions continues par morceaux.

1 Intégration des fonctions en escalier

1.1 Subdivisions d'un segment

Définition 1. Soient $a, b \in \mathbb{R}, a < b$. On appelle subdivision du segment $[a, b]$ toute suite finie $\sigma = (x_0, x_1, \dots, x_n)$ ($n \geq 1$) de points de $[a, b]$ vérifiant

$$a = x_0 < x_1 < \dots < x_n = b$$

Notation. On note $\mathcal{S}([a, b])$ l'ensemble des subdivisions du segment $[a, b]$.

Définition 2. Soit $\sigma = (x_i)_{0 \leq i \leq n} \in \mathcal{S}([a, b])$.

- On appelle pas de σ le réel

$$\mu(\sigma) = \max_{0 \leq i < n} (x_{i+1} - x_i)$$

- On appelle support de σ l'ensemble

$$\text{Supp}(\sigma) = \{x_i, 0 \leq i \leq n\}$$

Exemple. Soient $a < b$, et $n \geq 1$. Pour $0 \leq k \leq n$ posons

$$x_k = a + k \frac{b-a}{n}$$

On obtient une *subdivision à pas constant* dont le pas est

$$\mu_n = \frac{b-a}{n}$$

Définition 3. Soit $\sigma, \sigma' \in \mathcal{S}([a, b])$. On dit que σ est moins fine que σ' , ou σ' est plus fine que σ , et on note $\sigma \leq \sigma'$ lorsque

$$\text{Supp}(\sigma) \subset \text{Supp}(\sigma')$$

Proposition 1. Soient $\sigma', \sigma'' \in \mathcal{S}([a, b])$. Il existe une subdivision de $[a, b]$ à la fois plus fine que σ' et que σ'' .

Démonstration. C'est par exemple la subdivision dont le support est la réunion des supports de σ' et de σ'' . $\square$

1.2 Notion de fonction en escalier

Définition 4. Soit $\varphi : [a, b] \rightarrow \mathbb{R}$. On dit que φ est en escalier sur $[a, b]$ lorsqu'il existe une subdivision $\sigma = (x_i)_{0 \leq i \leq n}$ de $[a, b]$ telle que φ soit constante sur les intervalles $]x_i, x_{i+1}[$.

Définition 5. Une telle subdivision est dite adaptée à φ .

Remarque 1. Il est clair que si σ est une subdivision adaptée à φ , alors toute subdivision plus fine que σ est encore adaptée.

Exemple. Les fonctions constantes sont en escalier.

Proposition 2. L'ensemble $\mathcal{E}([a, b])$ des fonctions en escalier sur $[a, b]$ est une $\mathbb{R}$ -algèbre.

Démonstration. Montrons que c'est une sous-algèbre de l'algèbre $\mathbb{R}^{[a, b]}$.

- La fonction constante égale à 1 est en escalier.
- Soient φ, ψ en escalier sur $[a, b]$. Soit σ une subdivision adaptée à φ et à ψ . Il suffit pour en créer une de prendre une subdivision adaptée à φ , une subdivision adaptée à ψ , et de créer la subdivision dont le support est la réunion des supports de ces deux subdivisions. Il est alors clair que cette nouvelle subdivision est adaptée à $\varphi + \psi$, à $\varphi\psi$ et à $\lambda\varphi$ pour tout λ réel.

□

1.3 Intégrale d'une fonction en escalier

Proposition 3. Soit $\varphi \in \mathcal{E}([a, b])$. Soit $\sigma = (x_i)_{0 \leq i \leq n}$ une subdivision adaptée à φ . Pour $0 \leq i \leq n-1$, soit y_i la valeur (constante) de φ sur $]x_i, x_{i+1}[$. La quantité

$$\sum_{i=0}^{n-1} y_i (x_{i+1} - x_i)$$

ne dépend pas de la subdivision σ .

Démonstration. Sans entrer dans les détails de la preuve, cela signifie que si l'on coupe un rectangle en plusieurs morceaux, l'aire totale est la somme des aires des morceaux. □

Définition 6. La quantité introduite dans le théorème précédent est appelée intégrale de φ sur $[a, b]$, et notée $\int_a^b \varphi$, ou $\int_{[a,b]} \varphi$, ou encore $\int_a^b \varphi(x) dx$.

Exercice. Dessinez une fonction en escalier. Puis hachurez la partie du plan entre sa « courbe » et l'axe Ox .

Remarque 2. Si vous avez fait l'exercice, bravo, vous avez devant les yeux $\int_a^b \varphi$! C'est l'aire de ce que vous avez hachuré.

1.4 Propriétés

Proposition 4. Soient $a, b \in \mathbb{R}$, $a < b$. On a pour toutes fonctions $\varphi, \psi \in \mathcal{E}([a, b])$ et tout réel λ :

- Linéarité 1 : $\int_a^b (\varphi + \psi) = \int_a^b \varphi + \int_a^b \psi$.
- Linéarité 2 : $\int_a^b \lambda \varphi = \lambda \int_a^b \varphi$.
- Croissance : $\varphi \leq \psi \implies \int_a^b \varphi \leq \int_a^b \psi$.
- $\left| \int_a^b \varphi \right| \leq \int_a^b |\varphi|$

Démonstration.

- Pour la linéarité 1, on prend une subdivision adaptée à la fois à φ et à ψ .
- La linéarité 2 et la croissance sont évidentes.
- Pour la dernière inégalité, remarquons d'abord que si $\varphi \in \mathcal{E}([a, b])$, il en va de même pour $|\varphi|$. Maintenant, $\varphi \leq |\varphi|$, donc par la croissance, on a

$$\int_a^b \varphi \leq \int_a^b |\varphi|$$

De même, $-\varphi \leq |\varphi|$, donc par croissance et linéarité,

$$-\int_a^b \varphi \leq \int_a^b |\varphi|$$

D'où l'inégalité voulue, puisque $|\int_a^b \varphi|$ est l'un des deux nombres $\int_a^b \varphi$ et $-\int_a^b \varphi$.

□

Proposition 5. FORMULE DE CHASLES

Soient $a, b, c \in \mathbb{R}$, $a < b < c$. Soit $\varphi : [a, c] \rightarrow \mathbb{R}$. Alors

- $\varphi \in \mathcal{E}([a, c])$ si et seulement si $\varphi \in \mathcal{E}([a, b])$ et $\varphi \in \mathcal{E}([b, c])$.
- Si $\varphi \in \mathcal{E}([a, c])$ alors

$$\int_a^c \varphi = \int_a^b \varphi + \int_b^c \varphi$$

Démonstration. On n'entre pas dans les détails. Disons simplement qu'un réunissant une subdivision adaptée sur $[a, b]$ et une subdivision adaptée sur $[b, c]$, on fabrique une subdivision adaptée sur $[a, c]$. Inversement, si on a une subdivision adaptée sur $[a, c]$, on obtient des subdivisions adaptées sur $[a, b]$ et $[b, c]$ en la coupant en deux et en rajoutant éventuellement b aux deux subdivisions obtenues. □

Remarque 3. On généralise la notation intégrale en posant pour $a > b$

$$\int_a^b \varphi = - \int_b^a \varphi$$

et pour tout a

$$\int_a^a \varphi = 0$$

La formule de Chasles est alors vérifiée dans tous les sens (si j'ose dire). La linéarité fonctionne encore sans problème. En revanche, **attention aux inégalités** ! Elles ont toutes tendance à se renverser.

Moralité : quand on écrit des inégalités sur des intégrales du genre $\int_a^b \varphi$ on contrôle toujours que $a \leq b$.

2 Construction de l'intégrale

2.1 Fonctions continues par morceaux

Définition 7. Soit $f : [a, b] \rightarrow \mathbb{R}$. On dit que f est continue par morceaux sur $[a, b]$ lorsqu'il existe une subdivision $\sigma = (x_i)_{0 \leq i \leq n}$ de $[a, b]$ telle que

- f est continue sur les intervalles ouverts $]x_i, x_{i+1}[$.
- f admet une limite à gauche et à droite réelles en tous les x_i , $1 \leq i \leq n-1$, une limite à droite réelle en a et une limite à gauche réelle en b .

Une telle subdivision est, comme pour les fonctions en escalier, dite *adaptée* à f . Toute subdivision plus fine qu'une subdivision adaptée est encore adaptée.

Remarque 4. Toute fonction continue, toute fonction en escalier, est évidemment continue par morceaux. Nous avons donc les inclusions

$$\mathcal{C}^0([a, b]) \subset \mathcal{C}_m([a, b]) \quad \text{et} \quad \mathcal{E}([a, b]) \subset \mathcal{C}_m([a, b])$$

Définition 8. On étend cette notion à des fonctions définies sur un intervalle quelconque. Soit $f : I \rightarrow \mathbb{R}$. On dit que f est continue par morceaux sur I lorsque sa restriction à tout segment inclus dans I est continue par morceaux.

Exemple. La fonction partie entière, la fonction $f : \mathbb{R} \rightarrow \mathbb{R}$ définie par $f(x) = x - \lfloor x \rfloor$, sont continues par morceaux sur $\mathbb{R}$.

Proposition 6. L'ensemble $\mathcal{C}_m([a, b])$ des fonctions continues par morceaux sur $[a, b]$ est une $\mathbb{R}$ -algèbre et $\mathcal{E}([a, b])$ en est une sous-algèbre.

Démonstration. C'est la même démonstration que celle que nous avons faite pour les fonctions en escalier. $\square$

Proposition 7. Soient $a, b \in \mathbb{R}$, $a < b$. L'ensemble $\mathcal{E}([a, b])$ est dense dans $\mathcal{C}_m([a, b])$ au sens suivant : pour toute fonction $f \in \mathcal{C}_m([a, b])$, pour tout réel $\varepsilon > 0$, il existe deux fonctions $\varphi, \psi \in \mathcal{E}([a, b])$ telles que

$$\varphi \leq f \leq \psi \quad \text{et} \quad \psi - \varphi \leq \varepsilon$$

Démonstration. Pour alléger la preuve, on la fait seulement pour f continue sur $[a, b]$.

Soit $\varepsilon > 0$. D'après le théorème de Heine, f est uniformément continue sur $[a, b]$. Il existe donc $\alpha > 0$ tel que

$$\forall x, y \in [a, b], |x - y| \leq \alpha \implies |f(x) - f(y)| \leq \varepsilon$$

Soit $\sigma = (x_0 = a, \dots, x_n = b)$ une subdivision de $[a, b]$ de pas $\mu(\sigma) \leq \alpha$. Soit φ en escalier définie comme suit.

Pour tout $k \in \llbracket 0, n-1 \rrbracket$, pour tout $x \in [x_k, x_{k+1}[$,

$$\varphi(x) = \min_{[x_k, x_{k+1}]} f$$

Il reste à définir $\varphi(b)$. Posons $\varphi(b) = f(b)$.

On définit de même ψ en mettant $\max$ à la place de $\min$. On a clairement

$$\varphi \leq f \leq \psi$$

Par ailleurs, soit $x \in [a, b[$. Soit k tel que $x \in [x_k, x_{k+1}[$. On a

$$\psi(x) - \varphi(x) = \max_{[x_k, x_{k+1}]} f - \min_{[x_k, x_{k+1}]} f = f(v) - f(u)$$

où $u, v \in [x_k, x_{k+1}]$ (théorème des valeurs intermédiaires). Mais alors,

$$|v - u| \leq \mu(\sigma) \leq \alpha$$

donc

$$|f(v) - f(u)| \leq \varepsilon$$

Ainsi, pour tout x différent de b , on a $\psi(x) - \varphi(x) \leq \varepsilon$. Pour $x = b$ aussi. Donc, $\psi - \varphi \leq \varepsilon$.

□

2.2 Intégrale d'une fonction continue par morceaux

Notation. Soit $f : [a, b] \longrightarrow \mathbb{R}$ bornée. On note

$$E_-(f) = \left\{ \int_{[a,b]} \varphi, \varphi \in \mathcal{E}([a, b]), \varphi \leq f \right\}$$

et

$$E_+(f) = \left\{ \int_{[a,b]} \psi, \psi \in \mathcal{E}([a, b]), f \leq \psi \right\}$$

Remarque 5. $E_-(f)$ est donc une partie de $\mathbb{R}$. C'est l'ensemble de toutes les intégrales de toutes les fonctions en escalier inférieures ou égales à f . De même pour $E_+(f)$.

Remarque 6. Toute fonction continue par morceaux sur $[a, b]$ est bornée sur $[a, b]$. En effet, elle est bornée sur chacun des morceaux d'une subdivision adaptée.

Proposition 8. *Les ensembles $E_-(f)$ et $E_+(f)$ possèdent respectivement une borne supérieure et une borne inférieure, et on a*

$$\inf E_+(f) \geq \sup E_-(f)$$

Démonstration. f est bornée sur $[a, b]$, c'est à dire minorée et majorée par des fonctions constantes. Comme une fonction constante est en escalier, les ensembles E_- et E_+ sont donc non vides.

De plus, soient $x \in E_-$ et $y \in E_+$. On a

$$x = \int_a^b \varphi \quad \text{et} \quad y = \int_a^b \psi$$

où $\varphi, \psi \in \mathcal{E}([a, b])$ et $\varphi \leq f \leq \psi$. Ainsi, par la croissance de l'intégrale des fonctions en escalier, $x \leq y$. On en déduit que $\sup E_-$ et $\inf E_+$ existent, et

$$\sup E_- \leq \inf E_+$$

□

Notation. Pour toute fonction $f : [a, b] \rightarrow \mathbb{R}$ bornée, nous noterons

$$\int_{*a}^b f = \sup E_-(f)$$

et

$$\int_a^{*b} f = \inf E_+(f)$$

Exercice. Les réels $\sup E_-$ et $\inf E_+$ peuvent être distincts lorsque la fonction f est « compliquée ». Prenons par exemple la fonction $f : [0, 1] \rightarrow \mathbb{R}$ définie par $f(x) = 1$ si $x \in \mathbb{Q}$ et $f(x) = 0$ sinon. Montrer que

$$E_-(f) = \mathbb{R}_- \quad \text{et} \quad E_+(f) = [1, +\infty[$$

En déduire que

$$\int_{*0}^1 f = 0 \quad \text{et} \quad \int_0^{*1} f = 1$$

Proposition 9. *Soit $f \in \mathcal{C}_m([a, b])$. On a $\sup E_-(f) = \inf E_+(f)$.*

Démonstration. Soit $\varepsilon > 0$. Il existe $\varphi, \psi \in \mathcal{E}([a, b])$ telles que

$$\varphi \leq f \leq \psi \quad \text{et} \quad \psi - \varphi \leq \frac{\varepsilon}{b-a}$$

De là,

$$\int_a^b \psi - \int_a^b \varphi \leq \int_a^b \frac{\varepsilon}{b-a} = \varepsilon$$

On en déduit que

$$\inf E_+(f) - \sup E_-(f) \leq \int_a^b \psi - \int_a^b \varphi \leq \varepsilon$$

Ceci étant vrai pour tout $\varepsilon > 0$, on en déduit que $\inf E_+(f) \leq \sup E_-(f)$. L'autre inégalité étant vraie pour toute fonction f , on a l'égalité recherchée. $\square$

Définition 9. Soit $f : [a, b] \rightarrow \mathbb{R}$ une fonction bornée. On dit que f est *intégrable* sur $[a, b]$ lorsque $\int_{*a}^b f = \int_a^{*b} f$. On appelle alors *intégrale* de f sur $[a, b]$ le réel

$$\int_a^b f = \int_{*a}^b f = \int_a^{*b} f$$

Notation. On note aussi, comme pour les fonctions en escalier, $\int_{[a,b]} f$, ou encore $\int_a^b f(x) dx$.

Exercice. Soit $f : [a, b] \rightarrow \mathbb{R}$ une fonction en escalier. On dispose alors a priori pour f de deux notions d'intégrale. Montrer qu'il n'y a pas d'ambiguïté.

Précisément, prouver que $E_-(f)$ a un plus grand élément, qui est $\int_a^b(f)$ au sens de l'intégrale des fonctions en escalier. Ainsi, l'intégrale de f au sens du I et l'intégrale de f au sens du II sont égales.

Exemple. Calculons $\int_0^1 x^2 dx$.

Soit $f : [0, 1] \rightarrow \mathbb{R}$ définie par $f(x) = x^2$. Considérons pour tout n la subdivision régulière $(x_k)_{0 \leq k \leq n}$ où

$$x_k = \frac{k}{n}$$

Soit φ en escalier définie par $\varphi(x) = f(x_k)$ sur l'intervalle $[x_k, x_{k+1}[$ (et $\varphi(1) = 1$). On a $\varphi \leq f$. Calculons son intégrale :

$$\int_0^1 \varphi = \frac{1}{n} \sum_{k=0}^{n-1} \left(\frac{k}{n}\right)^2 = \frac{(n-1)(2n-1)}{6n^2}$$

Soit de même ψ en escalier définie par $\psi(x) = f(x_{k+1})$ sur l'intervalle $[x_k, x_{k+1}[$ (et $\psi(1) = 1$). On a $\psi \geq f$. Calculons son intégrale :

$$\int_0^1 \psi = \frac{1}{n} \sum_{k=0}^{n-1} \left(\frac{k+1}{n} \right)^2 = \frac{(n+1)(2n+1)}{6n^2}$$

On a donc pour tout $n \geq 1$,

$$\frac{(n-1)(2n-1)}{6n^2} \leq \int_0^1 f \leq \frac{(n+1)(2n+1)}{6n^2}$$

Le minorant et le majorant tendent vers $\frac{1}{3}$ lorsque n tend vers l'infini. On a donc

$$\int_0^1 x^2 dx = \frac{1}{3}$$

Exercice. Calculer de la même manière $\int_0^1 e^x dx$.

3 Propriétés de l'intégrale

3.1 Linéarité

La propriété qui suit est fondamentale.

Proposition 10. Soient $f, g \in \mathcal{C}_m([a, b])$ et $\lambda \in \mathbb{R}$. On a

- $\int_a^b (f + g) = \int_a^b f + \int_a^b g$
- $\int_a^b \lambda f = \lambda \int_a^b f$

Démonstration. Soit $\varepsilon > 0$. Il existe $\varphi_1, \psi_1, \varphi_2, \psi_2$ en escalier sur $[a, b]$ telles que $\varphi_1 \leq f \leq \psi_1$, $\varphi_2 \leq g \leq \psi_2$, $\psi_1 - \varphi_1 \leq \frac{\varepsilon}{2(b-a)}$ et $\psi_2 - \varphi_2 \leq \frac{\varepsilon}{2(b-a)}$.

Par définition même de l'intégrale de f et de g , on a

$$(1) \quad \int_a^b \varphi_1 \leq \int_a^b f \leq \int_a^b \psi_1$$

$$(2) \quad \int_a^b \varphi_2 \leq \int_a^b g \leq \int_a^b \psi_2$$

Remarquons que $\varphi_1 + \varphi_2$ est en escalier et minore $f + g$. De même, $\psi_1 + \psi_2$ est en escalier et majore $f + g$. Par définition de l'intégrale de $f + g$ et la linéarité de l'intégrale pour les fonctions en escalier,

$$(3) \quad \int_a^b \varphi_1 + \int_a^b \varphi_2 \leq \int_a^b (f + g) \leq \int_a^b \psi_1 + \int_a^b \psi_2$$

On ajoute à (3) les inégalités opposées de (1) et (2). On obtient

$$\int_a^b (\varphi_1 - \psi_1) + \int_a^b (\varphi_2 - \psi_2) \leq \int_a^b (f + g) - \int_a^b f - \int_a^b g \leq \int_a^b (\psi_1 - \varphi_1) + \int_a^b (\psi_2 - \varphi_2)$$

d'où

$$\left| \int_a^b (f + g) - \int_a^b f - \int_a^b g \right| \leq \varepsilon$$

Ceci est vrai pour tout $\varepsilon > 0$, donc l'intégrale de la somme est la somme des intégrales. Pour le produit par un réel, c'est bien plus facile, nous nous dispenserons ici de donner la preuve. $\square$

3.2 Positivité et croissance

La propriété de positivité de l'intégrale, ainsi que ses corollaires (croissance, inégalité triangulaire) sont fondamentales.

Proposition 11. POSITIVITÉ DE L'INTÉGRALE

Soient $a, b \in \mathbb{R}$ tels que $a < b$. Soit $f \in \mathcal{C}_m([a, b])$. On a

$$f \geq 0 \implies \int_a^b f \geq 0$$

Démonstration. La fonction nulle est en escalier et minore f , donc

$$0 = \int_a^b 0 \leq \int_a^b f$$

$\square$

Proposition 12. CROISSANCE DE L'INTÉGRALE

Soient $a, b \in \mathbb{R}$ tels que $a < b$. Soient $f, g \in \mathcal{C}_m([a, b])$. On a

$$f \leq g \implies \int_a^b f \leq \int_a^b g$$

Démonstration. On a $g - f \geq 0$, donc

$$\int_a^b (g - f) \geq 0$$

On termine en utilisant la linéarité de l'intégrale. $\square$

Proposition 13. INÉGALITÉ TRIANGULAIRE

Soient $a, b \in \mathbb{R}$ tels que $a < b$. Soit $f \in \mathcal{C}_m([a, b])$. Alors $|f| \in \mathcal{C}_m([a, b])$ et

$$\left| \int_a^b f \right| \leq \int_a^b |f|$$

Démonstration. Une subdivision adaptée à f est aussi adaptée à $|f|$. Pour montrer l'inégalité, il suffit de remarquer que $f \leq |f|$ et $-f \leq |f|$, puis d'utiliser la croissance de l'intégrale et sa linéarité. $\square$

3.3 Formule de Chasles

La propriété qui suit est fondamentale.

Proposition 14. FORMULE DE CHASLES

Soient $a, b, c \in \mathbb{R}$ tels que $a < b < c$. Soit $f : [a, c] \rightarrow \mathbb{R}$. Alors

- $f \in \mathcal{C}_m([a, c])$ si et seulement si $f \in \mathcal{C}_m([a, b])$ et $f \in \mathcal{C}_m([b, c])$.
- Si $f \in \mathcal{C}_m([a, c])$ alors

$$\int_a^c f = \int_a^b f + \int_b^c f$$

Démonstration. Nous l'admettons. L'idée est de travailler avec des subdivisions adaptées auxquelles on rajoute éventuellement le point b . $\square$

Remarque 7. Comme pour les fonctions en escalier, on généralise la notation intégrale en posant, si $a > b$,

$$\int_a^b f = - \int_b^a f$$

et aussi

$$\int_a^a f = 0$$

La formule de Chasles est alors vérifiée dans tous les sens.

Attention cependant : les inégalités concernant les intégrales (comme la positivité ou la croissance) ne sont valables que lorsque $a \leq b$.

3.4 Nullité de l'intégrale

Proposition 15. Soit $f : [a, b] \rightarrow \mathbb{R}$ continue et positive, telle que

$$\int_a^b f = 0$$

Alors

$$\forall x \in [a, b], f(x) = 0$$

Démonstration. Supposons que f n'est pas identiquement nulle sur $[a, b]$. Il existe $c \in [a, b]$ tel que $f(c) > 0$.

Supposons $c \neq a, b$ (si $c = a$ ou b , la preuve est analogue). Comme f est continue en c , il existe $\delta > 0$, tel que

$$\forall x \in [c - \delta, c + \delta], f(x) \geq \frac{1}{2}f(c)$$

Prenons δ suffisamment petit pour que $a \leq c - \delta$ et $c + \delta \leq b$. On a alors

$$\int_a^b f = \int_a^{c-\delta} f + \int_{c-\delta}^{c+\delta} f + \int_{c+\delta}^b f \geq 0 + \int_{c-\delta}^{c+\delta} f + 0 \geq \delta f(c) > 0$$

□

3.5 Inégalité de Schwarz

Soient $a, b \in \mathbb{R}$, $a < b$. Pour $f, g \in \mathcal{C}_m([a, b])$, notons

$$\langle f, g \rangle = \int_a^b fg$$

Proposition 16. L'application $(f, g) \mapsto \langle f, g \rangle$ vérifie les propriétés suivantes :

- C'est une forme bilinéaire symétrique sur $\mathcal{C}_m([a, b])$.
- $\forall f \in \mathcal{C}_m([a, b]), \langle f, f \rangle \geq 0$.
- $\forall f \in \mathcal{C}^0([a, b]), \langle f, f \rangle = 0 \implies f = 0$.

Démonstration.

- La symétrie est évidente car la multiplication des réels est commutative. La bilinéarité résulte de la linéarité de l'intégrale.
- Soit $f \in \mathcal{C}_m([a, b])$. La fonction f^2 est positive donc, par positivité de l'intégrale,

$$\langle f, f \rangle = \int_a^b f^2 \geq 0$$

- Soit $f \in \mathcal{C}^0([a, b])$. Supposons $\langle f, f \rangle = 0$. La fonction f^2 est continue et positive sur $[a, b]$ et son intégrale est nulle. Par le théorème de nullité de l'intégrale, on a donc $f = 0$.

□

Remarque 8. Nous reviendrons sur ce sujet dans le chapitre sur les espaces préhilbertiens réels. Ce que nous raconte le théorème, c'est que la fonction $(f, g) \mapsto \langle f, g \rangle$ est un *produit scalaire* sur $\mathcal{C}^0([a, b])$.

Proposition 17. INÉGALITÉ DE SCHWARZ

Soient $a, b \in \mathbb{R}$ vérifiant $a < b$. Soient $f, g \in \mathcal{C}_m([a, b])$. Alors

$$\left| \int_a^b fg \right| \leq \sqrt{\int_a^b f^2} \sqrt{\int_a^b g^2}$$

Si f et g sont continues sur $[a, b]$, l'inégalité précédente est une égalité si et seulement si f et g sont proportionnelles.

Démonstration. Considérons la fonction $\varphi : \mathbb{R} \rightarrow \mathbb{R}$ définie par

$$\varphi(t) = \langle tf + g, tf + g \rangle$$

Remarquons tout d'abord que, par positivité de l'intégrale, on a pour tout réel t , $\varphi(t) \geq 0$. De plus, par bilinéarité,

$$\varphi(t) = \langle f, f \rangle t^2 + 2 \langle f, g \rangle t + \langle g, g \rangle$$

La fonction φ est donc une fonction polynôme de degré inférieur ou égal à 2.

Commençons par le cas où $\langle f, f \rangle = 0$. La fonction φ est alors une fonction affine de signe constant. C'est donc une fonction constante, ce qui entraîne $\langle f, g \rangle = 0$. L'inégalité demandée est alors évidente, elle devient $0 \leq 0$. C'est même une égalité.

Supposons maintenant $\langle f, f \rangle \neq 0$. La fonction φ est une fonction polynôme de degré 2 et de signe constant. Son discriminant est donc négatif ou nul. Que vaut le discriminant en question ? Eh bien c'est

$$\Delta = 4 \langle f, g \rangle^2 - 4 \langle f, f \rangle \langle g, g \rangle$$

On a donc

$$\langle f, g \rangle^2 \leq \langle f, f \rangle \langle g, g \rangle$$

d'où l'inégalité de Schwarz en passant à la racine carrée.

Passons maintenant au cas d'égalité. Supposons f et g continues sur $[a, b]$. Reprenons les deux cas ci-dessus.

- Le cas $\langle f, f \rangle = 0$. Dans ce cas, nous avons vu que $\langle f, g \rangle = 0$. On a donc égalité. De plus, comme f est continue, $\langle f, f \rangle = 0$ entraîne $f = 0$. Ainsi, $f = 0g$: f et g sont proportionnelles. L'équivalence est donc vérifiée.
- Le cas $\langle f, f \rangle \neq 0$ (et donc $f \neq 0$).

★ Supposons f et g proportionnelles : il existe $t \in \mathbb{R}$ tel que $g = tf$. On a

$$\left| \int_a^b fg \right| = \left| \int_a^b tf^2 \right| = |t| \int_a^b f^2$$

Par ailleurs,

$$\int_a^b f^2 \int_a^b g^2 = t^2 \left(\int_a^b f^2 \right)^2$$

En prenant la racine carrée on voit qu'on a égalité dans l'inégalité de Schwarz.

- ★ Supposons maintenant l'égalité dans l'inégalité de Schwarz. Rappelons-nous la démonstration de l'inégalité : l'inégalité de Schwarz nous dit qu'un discriminant est négatif ou nul. Maintenant, ce discriminant est nul. Qu'est-ce que cela veut dire ? Eh bien que la fonction φ a une racine (double). Il existe donc $t \in \mathbb{R}$ tel que $\varphi(t) = 0$, c'est à dire

$$\langle tf + g, tf + g \rangle = 0$$

La fonction $tf + g$ étant continue sur $[a, b]$, cela implique que $tf + g = 0$, c'est à dire

$$g = -tf$$

Les fonctions f et g sont effectivement proportionnelles.

□

Exercice. Voyons sur un exemple comment utiliser ce théorème. Soient $a, b \in \mathbb{R}$, $a < b$. Posons

$$E = \{f \in C^0([a, b]), f > 0\}$$

Soit $\varphi : E \rightarrow \mathbb{R}$ définie, pour toute $f \in E$, par

$$\varphi(f) = \int_a^b f \int_a^b \frac{1}{f}$$

Montrons que φ possède un minimum sur l'ensemble E et calculons ce minimum.

Par l'inégalité de Schwarz, on a pour tout $f \in E$,

$$\varphi(f) \geq \left(\int_a^b f \frac{1}{f} \right)^2 = (b - a)^2$$

La fonction φ est donc minorée par $(b - a)^2$. Par ailleurs, on a égalité ci-dessus si et seulement si il existe un réel λ tel que

$$f = \lambda \frac{1}{f}$$

c'est à dire si et seulement si

$$f^2 = \lambda$$

bref, si et seulement si f est constante. Conclusion :

φ admet un minimum sur E qui est $(b-a)^2$. Ce minimum est atteint pour les fonctions constantes et elles seulement.

4 Approximations de l'intégrale

4.1 Sommes de Riemann - Méthode des rectangles

Définition 10. Soient $a, b \in \mathbb{R}$, $a < b$. On appelle *subdivision pointée* de $[a, b]$ tout couple $S = (\sigma, \tau)$ où $\sigma = (x_i)_{0 \leq i \leq n}$ est une subdivision « normale » de $[a, b]$ et $\tau = (\xi_i)_{0 \leq i \leq n-1}$ vérifie

$$\forall i \in \llbracket 0, n-1 \rrbracket, \xi_i \in [x_i, x_{i+1}]$$

Une subdivision pointée est donc une subdivision pour laquelle, dans chacun des morceaux de la subdivision, on se donne un point particulier.

Notation. On note $\dot{\mathcal{S}}([a, b])$ l'ensemble des subdivisions pointées de $[a, b]$.

Définition 11. Soient $a, b \in \mathbb{R}$ tels que $a < b$. Soit $f : [a, b] \rightarrow \mathbb{R}$. Soit S une subdivision pointée de $[a, b]$. On appelle *somme de Riemann* associée à f et S la quantité

$$\mathcal{R}_S(f) = \sum_{i=0}^{n-1} (x_{i+1} - x_i) f(\xi_i)$$

Exercice. Faire un petit dessin. $\mathcal{R}_S(f)$ est l'aire d'une réunion de rectangles : dessiner la courbe de f , une subdivision de quelques points, hachurer l'aire correspondant à $\mathcal{R}_S(f)$.

Proposition 18. Soient $a, b \in \mathbb{R}$ tels que $a < b$. Soit $f : [a, b] \rightarrow \mathbb{R}$ une fonction continue. Alors

$$\forall \varepsilon > 0, \exists \delta > 0, \forall S \in \dot{\mathcal{S}}([a, b]), \mu(S) \leq \delta \implies \left| \int_a^b f - \mathcal{R}_S(f) \right| \leq \varepsilon$$

Démonstration. Soit $\varepsilon > 0$. f est uniformément continue sur $[a, b]$, donc il existe $\delta > 0$ tel que

$$\forall x, y \in [a, b], |x - y| \leq \delta \implies |f(x) - f(y)| \leq \frac{\varepsilon}{b - a}$$

Soit S une subdivision pointée de $[a, b]$ (mêmes notations que dans la définition) de pas inférieur ou égal à δ . On a

$$\begin{aligned} \left| \int_a^b f - \mathcal{R}_S(f) \right| &= \left| \sum_{i=0}^{n-1} \left(\int_{x_i}^{x_{i+1}} f(x) dx - \int_{x_i}^{x_{i+1}} f(\xi_i) dx \right) \right| \\ &\leq \sum_{i=0}^{n-1} \int_{x_i}^{x_{i+1}} |f(x) - f(\xi_i)| dx \\ &\leq \sum_{i=0}^{n-1} \int_{x_i}^{x_{i+1}} \frac{\varepsilon}{b - a} dx \\ &= \int_a^b \frac{\varepsilon}{b - a} dx \\ &= \varepsilon \end{aligned}$$

□

La proposition ci-dessus nous dit les sommes de Riemann associées à la fonction continue f convergent vers $\int_a^b f$ lorsque le pas des subdivisions tend vers 0. Le problème, c'est qu'on ne maîtrise pas l'erreur commise en approchant l'intégrale par une somme de Riemann. Tout ce que l'on sait, c'est que la somme de Riemann est *proche* de l'intégrale lorsque le pas de la subdivision est *assez petit*. Tout cela est évidemment assez vague.

Si l'on a des renseignements supplémentaires sur f , on peut dire mieux. Par exemple, lorsque f est lipschitzienne ...

Proposition 19. Soient $a, b \in \mathbb{R}$ tels que $a < b$. Soit $k \geq 0$. Soit $f : [a, b] \rightarrow \mathbb{R}$ k -lipschitzienne. On a alors pour toute subdivision pointée S de $[a, b]$,

$$\left| \int_a^b f - \mathcal{R}_S(f) \right| \leq k(b - a)\mu(S)$$

Démonstration.

$$\begin{aligned}
 \left| \int_a^b f - \mathcal{R}_S(f) \right| &= \left| \sum_{i=0}^{n-1} \left(\int_{x_i}^{x_{i+1}} f(x) dx - \int_{x_i}^{x_{i+1}} f(\xi_i) dx \right) \right| \\
 &\leq \sum_{i=0}^{n-1} \int_{x_i}^{x_{i+1}} |f(x) - f(\xi_i)| dx \\
 &\leq k \sum_{i=0}^{n-1} \int_{x_i}^{x_{i+1}} |x - \xi_i| dx \\
 &\leq k \mu(S) \sum_{i=0}^{n-1} \int_{x_i}^{x_{i+1}} dx \\
 &= k(b-a)\mu(S)
 \end{aligned}$$

□

Exemple. Si S est une subdivision régulière de $[a, b]$ en n morceaux, alors

$$\left| \int_a^b f - \mathcal{R}_S(f) \right| \leq k \frac{(b-a)^2}{n}$$

Exemple. Appliquons la méthode des rectangles à $f(x) = \exp x$ sur $[0, 1]$. Soit $n \geq 1$. Avec les notations ci-dessus, prenons une subdivision régulière en posant $x_k = \frac{k}{n}$, $0 \leq k \leq n$. Posons également $\xi_k = x_k$, $0 \leq k \leq n-1$. On a

$$\begin{aligned}
 \mathcal{R}_S(f) &= \sum_{k=0}^{n-1} (x_{k+1} - x_k) f(\xi_k) \\
 &= \sum_{k=0}^{n-1} \frac{1}{n} e^{k/n} \\
 &= \frac{1}{n} \sum_{k=0}^{n-1} \left(e^{1/n} \right)^k
 \end{aligned}$$

On reconnaît la somme des termes d'une suite géométrique :

$$\begin{aligned}
 \mathcal{R}_S(f) &= \frac{1}{n} \frac{(e^{1/n})^n - 1}{e^{1/n} - 1} \\
 &= \frac{1}{n} \frac{e - 1}{e^{1/n} - 1}
 \end{aligned}$$

Lorsque n tend vers l'infini, $e^{1/n} - 1 \sim \frac{1}{n}$. Ainsi,

$$\mathcal{R}_S(f) \sim \frac{1}{n} \frac{e - 1}{\frac{1}{n}} = e - 1$$

On en déduit que $R_S(f)$ tend vers $e - 1$ lorsque n tend vers l'infini. Ainsi,

$$\int_0^1 e^x dx = e - 1$$

Remarque 9. Bien évidemment, calculer des intégrales par cette méthode n'est pas de tout repos. Il existe cependant des cas où l'on ne sait pas trouver de primitives d'une fonction, mais où la méthode des rectangles permet de calculer quand même l'intégrale de cette fonction. Voir feuille d'exercices.

Remarque 10. La méthode des rectangles n'est pas une bonne méthode pour approcher une intégrale. Le majorant de l'erreur en $\frac{1}{n}$ nous dit que si l'on veut, par exemple, obtenir une intégrale avec 9 chiffres exacts, il faudra additionner 1 milliard de termes. C'est évidemment peu rassurant.

Dans le paragraphe suivant, nous allons examiner une méthode plus efficace, où l'erreur est majorée par $\frac{1}{n^2}$.

4.2 Méthode des trapèzes

Soient $a, b \in \mathbb{R}$ tels que $a < b$. Soit $f : [a, b] \rightarrow \mathbb{R}$. On suppose dans ce qui suit que f est de classe $\mathcal{C}^2$. On note $M_2(f)$ le maximum de $|f''|$ sur le segment $[a, b]$.

Dans un premier temps, considérons la fonction affine $\varphi : [a, b] \rightarrow \mathbb{R}$ égale à f aux points a et b .

Proposition 20. On a pour tout $t \in [a, b]$,

$$|f(t) - \varphi(t)| \leq \frac{(t-a)(b-t)}{2} M_2(f)$$

Démonstration. Soit

$$g(x) = f(x) - \varphi(x) - \frac{M_2}{2}(x-a)(b-x)$$

On a $g \in \mathcal{C}^2[a, b]$, $g(a) = g(b) = 0$, et $g'' = f'' + M_2 \geq 0$. La fonction g' est donc croissante. De plus, par Rolle, il existe $c \in]a, b[$, tel que $g'(c) = 0$. Donc, $g' \leq 0$ sur $[a, c]$ et $g' \geq 0$ sur $[c, b]$. On en tire les variations de g sur $[a, b]$, et on constate (exercice, faire le tableau de variations de g) que $g \leq 0$ sur $[a, b]$. Donc,

$$f(x) - \varphi(x) \leq \frac{M_2}{2}(x-a)(b-x)$$

En faisant de même avec

$$h(x) = f(x) - \varphi(x) + \frac{M_2}{2}(x-a)(b-x)$$

on trouve l'autre inégalité. $\square$

Proposition 21. On a

$$\left| \int_a^b f - \frac{1}{2}(b-a)(f(a) + f(b)) \right| \leq M_2 \frac{(b-a)^3}{12}$$

Démonstration. On a

$$\left| \int_a^b f - \int_a^b \varphi \right| \leq \int_a^b |f - \varphi| \leq \frac{M_2}{2} \int_a^b (t-a)(b-t) dt = M_2 \frac{(b-a)^3}{12}$$

Par ailleurs, $\int_a^b \varphi$ est l'aire d'un trapèze (faire un dessin), et donc

$$\int_a^b \varphi = \frac{1}{2}(b-a)(f(a) + f(b))$$

$\square$

L'idée suivante est de couper le segment $[a, b]$ en n morceaux, et d'appliquer ce qui précède sur chacun des morceaux. On obtient la proposition ci-dessous.

Proposition 22. Soit $n \geq 1$. Pour $i \in \llbracket 0, n \rrbracket$, soit $x_i = a + i \frac{b-a}{n}$. Alors

$$\left| \int_a^b f - \mathcal{T}_n(f) \right| \leq \frac{(b-a)^3}{12n^2} M_2(f)$$

où

$$\mathcal{T}_n(f) = \frac{b-a}{2n} \sum_{i=0}^{n-1} (f(x_i) + f(x_{i+1}))$$

Démonstration. Soit φ la fonction égale à f aux points x_i et affine sur chacun des morceaux de la subdivision régulière. On a

$$\begin{aligned} \left| \int_a^b f - \int_a^b \varphi \right| &\leq \sum_{i=0}^{n-1} \left| \int_{x_i}^{x_{i+1}} f - \int_{x_i}^{x_{i+1}} \varphi \right| \\ &\leq \sum_{i=0}^{n-1} M_2 \frac{(b-a)^3}{12n^3} \\ &= M_2 \frac{(b-a)^3}{12n^2} \end{aligned}$$

Par ailleurs,

$$\int_a^b \varphi = \sum_{i=0}^{n-1} \int_{x_i}^{x_{i+1}} \varphi = \mathcal{T}_n(f)$$

□

Remarque 11. On a

$$\mathcal{T}_n(f) = \frac{b-a}{n} \sum_{i=0}^{n-1} f(x_i) + \frac{(b-a)(f(b) - f(a))}{2n}$$

On constate donc que la somme associée à la méthode des trapèzes est égale à une somme de Riemann à laquelle on rajoute le terme correcteur

$$\frac{(b-a)(f(b) - f(a))}{2n}$$

Exemple. Prenons $a = 0, b = 1$ et $f(x) = x^2$. Il vient

$$\mathcal{T}_n(f) = \frac{1}{n} \sum_{i=0}^{n-1} \left(\frac{i}{n}\right)^2 + \frac{1}{2n} = \frac{(n-1)n(2n-1)}{6n^3} + \frac{1}{2n} = \frac{1}{3} + \frac{1}{6n^2}$$

Or, $\int_0^1 f = \frac{1}{3}$. Ainsi, l'erreur commise est $\frac{1}{6n^2}$. Par ailleurs, l'erreur théorique est

$$M_2 \frac{(b-a)^3}{12n^2} = \frac{1}{6n^2}$$

Ainsi, notre calcul de l'erreur était optimal. Le majorant que nous avons trouvé est atteint pour notre exemple.

5 Primitives

Nous allons enfin voir comment la notion d'intégrale permet de prouver que toute fonction continue sur un intervalle y admet des primitives.

5.1 Notion de primitive

Définition 12. Soit I un intervalle de $\mathbb{R}$. Soient f et F deux fonctions définies sur I . On dit que F est une primitive de f sur I lorsque $F \in \mathcal{D}^1(I)$ et que $F' = f$ sur I .

5.2 Existence et unicité des primitives

Proposition 23. UNICITÉ DES PRIMITIVES

Soit I un intervalle de $\mathbb{R}$. Soit $f : I \longrightarrow \mathbb{R}$. On suppose que f admet une primitive F sur I . Alors, les primitives de f sont les fonctions $F + c$ où $c \in \mathbb{R}$.

Démonstration. Soit $G \in \mathcal{D}^1(I)$. G est une primitive de f sur I si et seulement si $G' = f$, c'est à dire $G' = F'$, ou encore $(G - F)' = 0$. C'est à dire si et seulement si $G - F$ est constante sur I . $\square$

Nous y voilà ...

Proposition 24. EXISTENCE DES PRIMITIVES

Soit I un intervalle de $\mathbb{R}$. Soit $f : I \longrightarrow \mathbb{R}$ une fonction continue. Soit $a \in I$. L'application $F : I \longrightarrow \mathbb{R}$ définie pour tout $x \in I$ par

$$F(x) = \int_a^x f(t) dt$$

est une primitive de f sur I .

Démonstration.

- La fonction F est bien définie sur I puisque pour tout $x \in I$, f est continue sur $[a, x]$ (ou $[x, a]$ si $x \leq a$) : l'intégrale existe donc.
- Soit $x_0 \in I$. Soit $x \in I$. Alors

$$F(x) - F(x_0) - (x - x_0)f(x_0) = \int_{x_0}^x f(t) dt - (x - x_0)f(x_0) = \int_{x_0}^x (f(t) - f(x_0)) dt$$

Soit $\varepsilon > 0$. Puisque f est continue en x_0 , il existe $\alpha > 0$ tel que, pour tout $t \in I$,

$$|t - x_0| \leq \alpha \implies |f(t) - f(x_0)| \leq \varepsilon$$

On a alors pour tout $x > x_0$ tel que $|x - x_0| \leq \alpha$,

$$|F(x) - F(x_0) - (x - x_0)f(x_0)| \leq \int_{x_0}^x |f(t) - f(x_0)| dt \leq |x - x_0|\varepsilon$$

D'où

$$\left| \frac{F(x) - F(x_0)}{x - x_0} - f(x_0) \right| \leq \varepsilon$$

Le même résultat est obtenu pour $x < x_0$. Tout ceci n'est autre que la définition de limite. Nous venons de prouver que

$$\frac{F(x) - F(x_0)}{x - x_0} \xrightarrow{x \rightarrow x_0} f(x_0)$$

Ainsi, F est dérivable en x_0 et $F'(x_0) = f(x_0)$.

□

5.3 Application

Proposition 25. THÉORÈME FONDAMENTAL DU CALCUL INTÉGRAL

Soit $f : [a, b] \longrightarrow \mathbb{R}$ une fonction continue. Soit F une primitive de f sur $[a, b]$. Alors

$$\int_a^b f(x) dx = F(b) - F(a)$$

Démonstration. Posons, pour tout $x \in [a, b]$,

$$G(x) = \int_a^x f(t) dt$$

F et G étant des primitives de f sur $[a, b]$, il existe un réel c tel que pour tout $x \in [a, b]$, on ait

$$F(x) = G(x) + c$$

Mais alors

$$F(b) - F(a) = G(b) - G(a) = G(b) = \int_a^b f(x) dx$$

□

Le reste de la section est déjà connu (voir le chapitre sur les primitives). Redonnons les résultats.

5.4 Intégration par parties

Proposition 26. Soient $u, v : [a, b] \longrightarrow \mathbb{R}$ deux fonctions de classe $\mathcal{C}^1$. On a

$$\int_a^b u'v = [uv]_a^b - \int_a^b uv'$$

où $[\varphi]_a^b = \varphi(b) - \varphi(a)$.

Démonstration. On a $(uv)' = u'v + uv'$. Et on intègre. Comme $(uv)'$ est continue, on peut appliquer le théorème fondamental du calcul intégral dans le membre de gauche. Bien évidemment, une primitive de $(uv)'$ est uv . □

Exemple. Cherchons une primitive de la fonction arcsin sur $[-1, 1]$. Pour cela, considérons

$$F(x) = \int_0^x \arcsin t dt$$

pour $-1 \leq x \leq 1$. On voudrait poser $u' = 1$ et $v = \arcsin x$, malheureusement v n'est pas dérivable en -1 et 1 . Supposons donc d'abord $-1 < x < 1$ pour rendre notre IPP valide. Il vient

$$F(x) = [t \arcsin t]_0^x - \int_0^x \frac{t}{\sqrt{1-t^2}} dt = x \arcsin x + \sqrt{1-x^2} + 1$$

Une primitive de la fonction $\arcsin$ sur $] -1, 1[$ est donc la fonction

$$F : x \mapsto x \arcsin x + \sqrt{1-x^2}$$

Maintenant, la fonction F ci-dessus est parfaitement définie, et même continue, sur $[-1, 1]$. Nous savons que F est une primitive de $\arcsin$ sur $] -1, 1[$, et nous savons que F est continue sur $[-1, 1]$. Pourquoi F est-elle une primitive de $\arcsin$ sur $[-1, 1]$?

La fonction $\arcsin$ est continue sur $[-1, 1]$, elle y admet donc une primitive. Choisissons une telle primitive G de sorte que $G(0) = F(0) = -1$. F et G sont deux primitives de $\arcsin$ sur $] -1, 1[$, elles diffèrent donc sur $] -1, 1[$ d'une constante, nulle par notre choix de $G(0)$. Bref, $F = G$ sur $] -1, 1[$.

Pour terminer, le lecteur est invité à faire l'exercice ci-dessous.

Exercice. Soient F et G deux fonctions continues sur $[-1, 1]$. On suppose que $F = G$ sur $] -1, 1[$. Montrer qu'on a en fait $F = G$ sur $[-1, 1]$. La continuité est évidemment la clé, faire tendre x vers 1 (et -1).

5.5 Changement de variable

Proposition 27. Soient $\alpha, \beta \in \mathbb{R}$. Soit $\varphi \in \mathcal{C}^1([\alpha, \beta])$. Soit f une fonction continue sur un intervalle I contenant $\varphi([\alpha, \beta])$. On a

$$\int_{\alpha}^{\beta} f(\varphi(t))\varphi'(t) dt = \int_{\varphi(\alpha)}^{\varphi(\beta)} f(x) dx$$

Démonstration. Soit F une primitive de f sur I . On a pour tout $t \in [\alpha, \beta]$,

$$f(\varphi(t))\varphi'(t) = F'(\varphi(t))\varphi'(t) = (F \circ \varphi)'(t)$$

Ainsi,

$$\begin{aligned} \int_{\alpha}^{\beta} f(\varphi(t))\varphi'(t) dt &= \int_{\alpha}^{\beta} (F \circ \varphi)'(t) dt \\ &= F \circ \varphi|_{\alpha}^{\beta} \\ &= F(\varphi(\beta)) - F(\varphi(\alpha)) \\ &= \int_{\varphi(\alpha)}^{\varphi(\beta)} f(x) dx \end{aligned}$$

□

Remarque 12. L'énoncé du théorème de changement de variable est compliqué, mais son utilisation est en réalité simple. On désire calculer $\int_a^b f(x) dx$. Pour cela, on décide de poser $x = \varphi(t)$. On cherche alors α et β tels que $a = \varphi(\alpha)$ et $b = \varphi(\beta)$ (quand x vaut a , t vaut α). La question est ensuite : que devient dx ? Le théorème nous dit qu'il devient $\varphi'(t) dt$ (facile de s'en souvenir : $dx = \frac{dx}{dt} dt$). Ainsi

$$\int_{\alpha}^{\beta} f(\varphi(t))\varphi'(t) dt = \int_a^b f(x) dx$$

Exercice. Soit $X \in [-1, 1]$. Calculer

$$\mathcal{I} = \int_0^X \arcsin x dx$$

en posant $x = \sin t$. Trouver tout d'abord

$$\mathcal{I} = \int_0^{\arcsin X} t \cos t dt$$

Puis faire une intégration par parties et finir le calcul.

6 Formules de Taylor

6.1 Introduction

Nous avons déjà rencontré deux formules de Taylor : la formule de Taylor Young, qui permet de calculer des limites, des équivalents, des développements limités. Et la formule de Taylor pour les polynômes, qui est une formule « exacte », permettant d'exprimer un polynôme à l'aide des puissances de $X - a$ et des dérivées de ce polynôme en a .

Nous allons dans ce paragraphe donner deux autres formules de Taylor, toujours bâties sur le même modèle : f est une fonction possédant une certaine régularité, a, b (ou a, x , ou autre chose) sont deux réels, et la formule est toujours

$$f(b) = T_n(b - a) + R_n$$

où

$$T_n = \sum_{k=0}^n \frac{X^k}{k!} f^{(k)}(a)$$

est un polynôme de degré inférieur ou égal à n , et R_n est un « reste ». C'est ce reste qui va changer selon le théorème.

6.2 Formule de Taylor avec reste intégral

Proposition 28. Soit $n \geq 0$. Soient $a, b \in \mathbb{R}$. Soit $f : [a, b] \rightarrow \mathbb{R}$ (ou $[b, a]$, si $a > b$) de classe $\mathcal{C}^{n+1}$. Alors

$$f(b) = \sum_{k=0}^n \frac{(b-a)^k}{k!} f^{(k)}(a) + R_n$$

où

$$R_n = \int_a^b \frac{(b-t)^n}{n!} f^{(n+1)}(t) dt$$

Démonstration. On fait une récurrence sur n . Pour $n = 0$, la formule s'écrit

$$f(b) = f(a) + \int_a^b f'(t) dt$$

qui n'est autre que le théorème fondamental du calcul intégral. Supposons maintenant le théorème vérifié pour un entier n . Soit $f \in \mathcal{C}^{n+2}([a, b])$. On a alors

$$f(b) = \sum_{k=0}^n \frac{(b-a)^k}{k!} f^{(k)}(a) + R_n$$

d'après l'hypothèse de récurrence. Intégrons R_n par parties en posant

$$u' = \frac{(b-t)^n}{n!} \quad \text{et} \quad v = f^{(n+1)}(t)$$

Il vient

$$R_n = \frac{(b-a)^{n+1}}{(n+1)!} + \int_a^b \frac{(b-t)^{n+1}}{(n+1)!} f^{(n+2)}(t) dt$$

d'où le résultat pour l'entier $n+1$. $\square$

6.3 Inégalité de Taylor-Lagrange

Nous allons majorer le R_n de la formule de Taylor avec reste intégral pour obtenir une inégalité : l'inégalité de Taylor-Lagrange.

Proposition 29. Soit $n \geq 0$. Soient $a, b \in \mathbb{R}, a < b$. Soit $f : [a, b] \rightarrow \mathbb{R}$ de classe $\mathcal{C}^{n+1}$. Alors

$$f(b) = \sum_{k=0}^n \frac{(b-a)^k}{k!} f^{(k)}(a) + R_n$$

où

$$|R_n| \leq \frac{|b-a|^{n+1}}{(n+1)!} \max_{[a,b]} |f^{(n+1)}|$$

*Le résultat reste valable si $a > b$ en remplaçant le segment $[a, b]$ par le segment $[b, a]$.
Et il est trivial lorsque $a = b$.*

Démonstration. Il suffit de majorer l'intégrale $|R_n|$ de la formule de Taylor avec reste intégral. $\square$

Remarque 13. Contrairement à la formule de Taylor-Young qui ne raconte des choses que localement, au voisinage d'un certain point, l'inégalité de Taylor-Lagrange est une inégalité globale. Elle sert donc à prouver...des inégalités globales. Prenons un exemple. On prend $f(x) = e^x$, $a = 0$, $b = x$. Et n quelconque. Il vient donc

$$e^x = \sum_{k=0}^n \frac{x^k}{k!} + R_n$$

où

$$|R_n| \leq \frac{|x|^{n+1}}{(n+1)!} \max_S |\exp|$$

avec $S = [0, x]$ ou $[x, 0]$ selon que x est positif ou négatif. Ce max vaut en fait 1 si $x \leq 0$ et e^x si $x > 0$. Mais il ne dépend pas de n . Lorsque n tend vers l'infini, le majorant tend vers 0. On vient ainsi de prouver que

$$\sum_{k=0}^n \frac{x^k}{k!} \xrightarrow{n \rightarrow \infty} e^x$$

Lorsque nous aurons vu le chapitre sur les séries, nous écrirons cela

$$e^x = \sum_{k=0}^{\infty} \frac{x^k}{k!}$$

Exercice. Soit $x \in \mathbb{R}$. Montrer que

$$\sum_{k=0}^n (-1)^k \frac{x^{2k}}{(2k)!} \xrightarrow{n \rightarrow \infty} \cos x$$

Énoncer et montrer un résultat analogue pour $\sin x$.

7 Fonctions à valeurs complexes

Cette section contient une très brève extension de la notion d'intégrale aux fonctions toujours définies sur un segment $[a, b]$, mais à valeurs dans $\mathbb{C}$.

Soit $f : [a, b] \rightarrow \mathbb{C}$ une fonction continue par morceaux (i.e. sa partie réelle et sa partie imaginaire sont continues par morceaux). On appelle intégrale de f sur $[a, b]$ le nombre complexe

$$\int_a^b f(x) dx = \int_a^b \operatorname{Re}(f(x)) dx + i \int_a^b \operatorname{Im}(f(x)) dx$$

La linéarité de l'intégrale reste vérifiée. La formule de Chasles également. On peut évidemment intégrer par parties ou faire des changements de variable.

En revanche, on ne peut plus parler de croissance ou de positivité de l'intégrale, puisqu'on ne peut pas comparer deux nombres complexes. On a cependant le résultat suivant :

Proposition 30. Soient $a, b \in \mathbb{R}$, $a < b$. Soit $f : [a, b] \rightarrow \mathbb{C}$ continue par morceaux. Alors

$$\left| \int_a^b f(x) dx \right| \leq \int_a^b |f(x)| dx$$

Démonstration. Les barres verticales sont bien entendu des modules de nombres complexes. Si $\int_a^b f = 0$ c'est évident. Supposons donc $\int_a^b f \neq 0$.

Pour tout réel θ , on a

$$\operatorname{Re} \left(e^{-i\theta} \int_a^b f \right) = \int_a^b \operatorname{Re} \left(e^{-i\theta} f \right) \leq \int_a^b |e^{-i\theta} f| = \int_a^b |f|$$

Prenons pour θ un argument de $\int_a^b f$, qui est, rappelons-le, non nulle. On a donc, par définition de l'argument d'un nombre complexe,

$$\int_a^b f = \left| \int_a^b f \right| e^{i\theta}$$

où la première égalité résulte de la linéarité de l'intégrale. On a alors

$$\operatorname{Re} \left(e^{-i\theta} \int_a^b f \right) = \operatorname{Re} \left(\left| \int_a^b f \right| \right) = \left| \int_a^b f \right|$$

d'où l'inégalité voulue. $\square$

Chapitre 29

Dénombrements

1 Notion d'ensemble fini

1.1 Cardinal

Lemme 1. Soient E et F deux ensembles non vides. Soient $a \in E$ et $b \in F$. Il existe une bijection de E vers F si et seulement si il existe une bijection de $E' = E \setminus \{a\}$ vers $F' = F \setminus \{b\}$. Le résultat reste vrai si l'on remplace bijection par injection ou surjection.

Démonstration. Supposons qu'il existe une bijection $f : E' \rightarrow F'$. Soit $g : E \rightarrow F$ définie par $g(x) = f(x)$ si $x \in E'$ et $g(a) = b$. Alors g est une bijection de E sur F .

Supposons maintenant qu'il existe une bijection $f : E \rightarrow F$. Soit $\varphi : F \rightarrow F$ définie par $\varphi(x) = x$ si $x \neq f(a)$ et $x \neq b$, $\varphi(f(a)) = b$ et $\varphi(b) = f(a)$. La fonction φ échange simplement b et $f(a)$. Elle est bijective, et d'ailleurs égale à sa réciproque. Considérons maintenant la fonction $f' = \varphi \circ f$. f' est une bijection de E vers F (composée de bijections) et, de plus $f'(a) = b$. La fonction $g : E' \rightarrow F'$ définie par $g(x) = f'(x)$ est alors une bijection de E' sur F' .

La démonstration est identique pour les injections et les surjections. $\square$

Notation. Pour tout entier non nul, on note $\bar{n} = \llbracket 0, n-1 \rrbracket$ l'ensemble des entiers entre 0 et $n-1$. On note également $\bar{0} = \emptyset$.

Notation. Soient A et B deux ensembles. On écrit $A \simeq B$ lorsqu'il existe une bijection de A sur B .

Proposition 2. Soient n et p deux entiers. On a

$$\bar{n} \simeq \bar{p} \iff n = p$$

Démonstration. Dans un sens c'est trivial. En effet, si $n = p$, alors $\bar{n} = \bar{p}$ et il existe bien une bijection entre $\bar{n}$ et $\bar{p}$: l'identité.

Inversement, on effectue une récurrence sur p . Pour $p = 0$ c'est clair, car s'il existe une application d'un ensemble vers $\emptyset$, alors cet ensemble est vide.

Soit maintenant $p \in \mathbb{N}$. Supposons que $\forall n \in \mathbb{N}$, s'il existe une bijection de $\bar{n}$ sur $\bar{p}$ alors $n = p$. Soit $n \in \mathbb{N}$. Supposons qu'il y a une bijection de $\bar{n}$ sur $\bar{p} + \bar{1}$. On a forcément $n \geq 1$, car $p + 1 \neq 0$. D'après le lemme, il existe une bijection de $\overline{n-1}$ sur $\bar{p}$. Par l'hypothèse de récurrence, on a $n-1 = p$, c'est à dire $n = p + 1$. $\square$

Lemme 3.

- Soient n et p deux entiers. Il existe une injection de $\bar{n}$ sur $\bar{p}$ si et seulement si $n \leq p$.
- Soient n et p deux entiers. Il existe une surjection de $\bar{n}$ sur $\bar{p}$ si et seulement si $n \geq p$.

Démonstration. Pour les injections, on peut par exemple faire une démonstration analogue à celle du résultat sur les bijections.

Pour les surjections, on peut remarquer qu'il existe une surjection de $\bar{n}$ sur $\bar{p}$ si et seulement si il existe une injection de $\bar{p}$ vers $\bar{n}$. On est ainsi ramenés au cas des injections. $\square$

Définition 1. Soit E un ensemble. On dit que E est *fini* lorsqu'il existe un entier n tel que $E \simeq \bar{n}$. D'après la proposition précédente, cet entier est unique. On l'appelle le cardinal de E .

Notation. Il existe un certain nombre de notations pour le cardinal d'un ensemble fini E . Les plus fréquentes sont $\text{Card } E$, $\#E$ ou $|E|$. Dans ce chapitre nous utiliserons la notation $|E|$.

1.2 Propriétés

Proposition 4. Soient E et F deux ensembles finis. Alors

$$E \simeq F \iff |E| = |F|$$

Démonstration. Soit $n = |E|$. Soit $p = |F|$. Soient $f : E \longrightarrow \bar{n}$ une bijection et $g : F \longrightarrow \bar{p}$ une bijection. Si $n = p$, alors $g^{-1} \circ f$ est une bijection de E vers F . Inversement, si h est une bijection de E vers F , alors, $g \circ h \circ f^{-1}$ est une bijection de $\bar{n}$ vers $\bar{p}$, donc $n = p$. $\square$

Proposition 5. Soient E et F deux ensembles finis. Alors il existe une injection de E vers F si et seulement si $|E| \leq |F|$. Et il existe une surjection de E vers F si et seulement si $|E| \geq |F|$.

Démonstration. Démonstration analogue à celle du théorème sur les bijections. $\square$

Proposition 6. *Soit E un ensemble fini. Soit $E' \subset E$. Alors*

- E' est fini.
- $|E'| \leq |E|$.
- $|E'| = |E| \iff E' = E$.

Démonstration. On fait une récurrence sur $n = |E|$.

Pour $n = 0$, c'est clair, puisque $E' = \emptyset$.

Soit $n \in \mathbb{N}$. Supposons la propriété vérifiée pour n . Soit E un ensemble fini de cardinal $n + 1$. Soit $E' \subset E$.

Si $E' = E$, il n'y a rien à faire : E' est fini et $|E'| = |E|$.

Sinon, soit $a \in E \setminus E'$. On a $E' \subset E \setminus \{a\}$. On a une bijection de E sur $\overline{n+1}$ donc, d'après le lemme au début du paragraphe, on a une bijection entre $E \setminus \{a\}$ et $\bar{n}$. Donc $E \setminus \{a\}$ est fini, de cardinal n . Par l'hypothèse de récurrence, E' est fini, de cardinal inférieur ou égal à n (et donc différent de $n + 1$). $\square$

Proposition 7. *Soit $f : E \longrightarrow F$ une application d'un ensemble fini E vers un ensemble fini F de même cardinal. Alors,*

$$f \text{ est bijective} \iff f \text{ est injective} \iff f \text{ est surjective.}$$

Démonstration. Supposons f injective. f est alors une bijection de E sur $f(E) \subset F$. Donc, $|E| = |f(E)| = |F|$. Donc $f(E) = F$ et f est surjective.

Supposons maintenant que f est surjective. f possède un inverse à gauche g qui est injectif. D'après ce que nous venons de dire sur les injections, g est bijectif et sa réciproque, f , l'est aussi. $\square$

2 Opérations sur les ensembles finis

2.1 Union disjointe

Proposition 8. *Soit $n \geq 1$. Soient $E_1, \dots, E_n$ n ensembles finis disjoints deux à deux. Leur réunion est alors finie et*

$$\left| \bigcup_{k=1}^n E_k \right| = \sum_{k=1}^n |E_k|$$

Démonstration. On le montre pour $n = 2$. Une récurrence facile sur n permet de conclure. Soient donc A et B deux ensembles finis disjoints, de cardinaux respectifs n et p . Soient $f : A \rightarrow \bar{n}$ et $g : B \rightarrow \bar{p}$ deux bijections. Considérons l'application

$$h : A \cup B \rightarrow \overline{n+p}$$

définie par $h(x) = f(x)$ si $x \in A$ et $h(x) = g(x) + n$ si $x \in B$. On montre alors facilement que h est une bijection. $\square$

2.2 Différence, union

Proposition 9. Soient A et B deux ensembles finis, $B \subset A$. Alors $A \setminus B$ est fini et

$$|A \setminus B| = |A| - |B|$$

Démonstration. On a $A = B \cup (A \setminus B)$ et la réunion est disjointe. $\square$

Proposition 10. Soient E et F deux ensembles finis. Alors, $E \cup F$ et $E \cap F$ sont finis, et

$$|E \cup F| = |E| + |F| - |E \cap F|$$

Démonstration. $E \cup F = (E \setminus (E \cap F)) \cup F$ et cette réunion est disjointe. On a donc

$$|E \cup F| = |E \setminus (E \cap F)| + |F|$$

Puis on utilise que $E \cap F \subset E$ et, donc, que

$$|E \setminus (E \cap F)| = |E| - |E \cap F|$$

$\square$

Cette formule se généralise à la réunion de n ensembles finis. La formule que l'on obtient alors s'appelle la formule du crible.

Exemple. Écrivons la formule pour trois ensembles finis E, F, G . On a

$$\begin{aligned} |E \cup F \cup G| &= |(E \cup F) \cup G| \\ &= |E \cup F| + |G| - |(E \cup F) \cap G| \\ &= |E| + |F| - |E \cap F| + |G| - |(E \cap G) \cup (F \cap G)| \\ &= |E| + |F| - |E \cap F| + |G| - (|E \cap G| + |F \cap G| - |(E \cap G) \cap (F \cap G)|) \\ &= |E| + |F| + |G| - |E \cap F| - |E \cap G| - |F \cap G| + |E \cap F \cap G| \end{aligned}$$

Pour 4 ensembles finis E, F, G, H on obtient

$$\begin{aligned} |E \cup F \cup G \cup H| &= |E| + |F| + |G| + |H| \\ &\quad - |E \cap F| - |E \cap G| - |E \cap H| - |F \cap G| - |F \cap H| - |G \cap H| \\ &\quad + |E \cap F \cap G| + |E \cap F \cap H| + |E \cap G \cap H| + |F \cap G \cap H| \\ &\quad - |E \cap F \cap G \cap H| \end{aligned}$$

La formule générale pour n ensembles finis est énoncée et montrée en annexe du chapitre.

2.3 Produit cartésien

Proposition 11. *Soient E et F deux ensembles finis. Alors $E \times F$ est fini, et*

$$|E \times F| = |E| \times |F|$$

Démonstration. On a

$$E \times F = \bigcup_{y \in F} E \times \{y\}$$

Les ensembles $E \times \{y\}$ sont disjoints. On a donc

$$|E \times F| = \sum_{y \in F} |E \times \{y\}|$$

Soit $y \in F$. La fonction $F : E \rightarrow E \times \{y\}$ définie par $f(x) = (x, y)$ est clairement bijective. Ainsi, $|E \times \{y\}| = |E|$. Ainsi,

$$|E \times F| = \sum_{y \in F} |E| = |E| \sum_{y \in F} 1 = |E| \times |F|$$

□

Remarque 1. Généralisation évidente par récurrence sur p à un produit de p ensembles finis : si $E_1, \dots, E_p$ sont p ensembles finis, alors $E_1 \times \dots \times E_p$ est fini et

$$|E_1 \times \dots \times E_p| = |E_1| \times \dots \times |E_p|$$

Un cas particulier important est le suivant.

Proposition 12. *Soit F un ensemble fini. Soit $p \in \mathbb{N}$. L'ensemble F^p des p -uplets d'éléments de F est fini, et*

$$|F^p| = |F|^p$$

Démonstration. En effet,

$$F^p = F \times F \times \dots \times F$$

p fois, et donc

$$|F^p| = |F| \times |F| \times \dots \times |F| = |F|^p$$

□

3 Applications entre ensembles finis

3.1 Dénombrements d'applications

Proposition 13. Soient E et F deux ensembles finis. L'ensemble F^E des applications de E dans F est fini, et

$$|F^E| = |F|^{|E|}$$

Démonstration. On fait une récurrence sur $n = |E|$.

Pour $n = 0$, on a $E = \emptyset$ et il y a une seule application de E vers $F : (\emptyset, F, \emptyset)$. Donc, le cardinal de F^E est $1 = |F|^0$.

Soit maintenant $n \in \mathbb{N}$. Supposons la propriété vérifiée pour n . Soit E de cardinal $n + 1$. Soit $a \in E$. Posons $E' = E \setminus \{a\}$. Alors $E = E' \cup \{a\}$ et $|E'| = n$. Considérons l'application

$$\varphi : F^{E'} \times F \longrightarrow F^E$$

définie comme suit. À tout couple (f, b) on associe l'application g définie par $g(x) = f(x)$ si $x \neq a$, et $g(a) = b$. On démontre facilement que l'application φ est bijective. Donc

$$|F^E| = |F^{E'} \times F| = |F^{E'}| \times |F| = |F|^n \times |F| = |F|^{n+1}$$

□

3.2 Dénombrements d'injections, de bijections

Notation. Étant donnés deux ensembles A et B , on note $\mathcal{I}(A, B)$ l'ensemble des injections de A vers B .

Notation. Soient $n, p \in \mathbb{N}$. On note $n^{\underline{p}} = n(n-1) \dots (n-p+1)$.

Proposition 14. Soient E et F deux ensembles finis. L'ensemble $\mathcal{I}(E, F)$ des injections de E dans F est fini, et

$$|\mathcal{I}(E, F)| = |F|^{\underline{|E|}}$$

Démonstration. On montre le résultat par une récurrence sur $p = |E|$.

Pour $p = 0$, c'est clair. Il y a une seule application de E vers F et elle est injective.

Soit maintenant $p \in \mathbb{N}$. Supposons la propriété vraie pour tout ensemble E fini de cardinal p . Soit E un ensemble de cardinal $p + 1$. $E = E' \cup \{a\}$ où E' est de cardinal p et $a \in E$. Soit

$$\mathcal{Z} = \bigcup_{b \in F} \mathcal{I}(E \setminus \{a\}, F \setminus \{b\}) \times \{b\}$$

Considérons l'application $\varphi : \mathcal{Z} \rightarrow \mathcal{I}(E, F)$ définie comme suit : au couple (f, b) on associe l'application g définie par $g(x) = f(x)$ si $x \neq a$ et $g(a) = b$ (on vérifie que g est bien injective). Alors, l'application φ est bijective, et sa réciproque associe à l'injection $g : E \rightarrow F$ le couple (f, b) où $b = g(a)$ et f est la restriction de g à $E \setminus \{a\}$, corestreinte à $F \setminus \{b\}$ (oui, je sais, plus personne ne suit). On a donc égalité des cardinaux des ensembles de départ et d'arrivée de φ :

$$\begin{aligned} |\mathcal{I}(E, F)| &= \sum_{b \in F} |\mathcal{I}(E \setminus \{a\}, F \setminus \{b\}) \times \{b\}| \\ &= \sum_{b \in F} |\mathcal{I}(E \setminus \{a\}, F \setminus \{b\})| \\ &= \sum_{b \in F} (n-1)^p \\ &= n(n-1)^p \\ &= \underline{n^{p+1}} \end{aligned}$$

□

Remarque 2. Soient E et F deux ensembles de cardinaux respectifs p et n . Supposons $p > n$. On a $n^p = 0$, ce qui est heureux puisque dans ce cas $\mathcal{I}(E, F) = \emptyset$.

Remarque 3. Le lecteur aura sûrement remarqué la complexité de la démonstration précédente. Nous sommes en présence d'un problème (déterminer le nombre d'injections) où une idée simple conduit à une preuve compliquée. Voici l'idée simple en question :

Une injection f de E (de cardinal p) vers F (de cardinal $n \geq p$) se fabrique comme suit. Appelons $x_1, \dots, x_p$ les éléments de E :

- On choisit la valeur de $f(x_1)$: n possibilités. Puis,
- On choisit la valeur de $f(x_2)$: $n - 1$ possibilités, car $f(x_2) \neq f(x_1)$ puisque f est injective. Puis,
- On choisit la valeur de $f(x_3)$: $n - 2$ possibilités, car deux valeurs dans l'ensemble d'arrivées sont déjà « prises »..
- ...
- On choisit enfin la valeur de $f(x_p)$: $n - p + 1$ possibilités.

On a au total $n(n-1) \dots (n-p+1) = n^p$ façons de construire une telle injection f .

Cette façon de procéder n'est certes pas totalement rigoureuse, mais nous ne nous privons pas de l'utiliser pour gagner du temps. En gardant toujours à l'esprit que *si nous le voulions* nous pourrions toujours fabriquer des bijections pour déterminer des cardinaux. Voir l'appendice pour une discussion plus approfondie à ce sujet.

Corollaire 15. *Soit E un ensemble fini de cardinal n . L'ensemble $\mathfrak{S}(E)$ des bijections de E dans E est fini, et*

$$|\mathfrak{S}(E)| = n!$$

Démonstration. Une application de E vers E est bijective si et seulement si elle est injective. Il suffit d'appliquer le théorème précédent et de remarquer que

$$n^n = n!$$

□

Corollaire 16. *Soit F un ensemble fini de cardinal n . Soit $p \leq n$. L'ensemble des p -uplets d'éléments distincts de F est fini, de cardinal n^p .*

Démonstration. Notons F^p l'ensemble des p -uplets d'éléments distincts de F . Rappelons que $\mathcal{I}(\llbracket 1, p \rrbracket, F)$ est l'ensemble des injections de $\llbracket 1, p \rrbracket$ vers F .

La fonction $\varphi : \mathcal{I}(\llbracket 1, p \rrbracket, F) \longrightarrow F^p$ définie par

$$\varphi(f) = (f(1), \dots, f(p))$$

est une bijection. Sa réciproque est la fonction $\psi : F^p \longrightarrow \mathcal{I}(\llbracket 1, p \rrbracket, F)$ définie par

$$\psi(x_1, \dots, x_p) = f$$

où f est définie par $\forall i \in \llbracket 1, p \rrbracket, f(i) = x_i$. Ainsi,

$$|F^p| = |\mathcal{I}(\llbracket 1, p \rrbracket, F)| = |F|^p$$

□

4 Parties d'un ensemble fini

4.1 Nombre total de parties

Proposition 17. *Soit E un ensemble fini. L'ensemble $\mathcal{P}(E)$ des parties de E est fini, et*

$$|\mathcal{P}(E)| = 2^{|E|}$$

Démonstration. On procède par récurrence sur $n = |E|$.

Pour $n = 0$ c'est clair puisque $\mathcal{P}(\emptyset) = \{\emptyset\}$.

Soit $n \in \mathbb{N}$. Supposons la propriété vraie pour tout ensemble de cardinal n . Soit $E = E' \cup \{a\}$ un ensemble de cardinal $n + 1$. Nous avons $\mathcal{P}(E) = \mathcal{P}'(E) \cup \mathcal{P}''(E)$ où

- $\mathcal{P}'(E)$ est l'ensemble des parties de E qui ne contiennent pas a .
- $\mathcal{P}''(E)$ est l'ensemble des parties de E qui contiennent a .

Une partie de E ne peut pas à la fois contenir a et ne pas le contenir. Ces deux ensembles sont donc disjoints et ainsi

$$|\mathcal{P}(E)| = |\mathcal{P}'(E)| + |\mathcal{P}''(E)|$$

Nous allons maintenant calculer les deux cardinaux du membre de droite de cette égalité.

- Pour $\mathcal{P}'(E)$ c'est facile. En effet, les parties de E qui ne contiennent pas a sont les parties de E' . Ainsi, $\mathcal{P}'(E) = \mathcal{P}(E')$. Par l'hypothèse de récurrence,

$$|\mathcal{P}'(E)| = 2^n$$

- Un élément de $\mathcal{P}''(E)$ est une partie A de E qui contient a . Une telle partie A est de la forme $A' \cup \{a\}$ où $A' \subset E'$. On trouve donc facilement une bijection entre $\mathcal{P}''(E)$ et $\mathcal{P}(E')$. C'est la fonction $\varphi : \mathcal{P}''(E) \rightarrow \mathcal{P}(E')$ définie par $\varphi(A) = A \setminus \{a\}$. Ainsi,

$$|\mathcal{P}''(E)| = |\mathcal{P}(E')| = 2^n$$

Finalement,

$$|\mathcal{P}(E)| = |\mathcal{P}'(E)| + |\mathcal{P}''(E)| = 2^n + 2^n = 2^{n+1}$$

□

Remarque 4. Il existe une autre démonstration tout aussi instructive de ce résultat. Soit E un ensemble fini. Rappelons que la *fonction indicatrice* d'une partie A de E est la fonction $\mathbb{1}_A : E \rightarrow \{0, 1\}$ définie par $\mathbb{1}_A(x) = 1$ si $x \in A$ et 0 sinon.

Considérons la fonction $\mathbb{1} : \mathcal{P}(E) \rightarrow \{0, 1\}^E$ définie pour tout $A \subset E$ par $\mathbb{1}(A) = \mathbb{1}_A$. On montre facilement que cette fonction est une bijection. Sa réciproque est la fonction $\mathbb{1}^{-1} : \{0, 1\}^E \rightarrow \mathcal{P}(E)$ définie, pour toute fonction $f : E \rightarrow \{0, 1\}$, par

$$\mathbb{1}^{-1}(f) = f^{-1}(\{1\})$$

Ainsi,

$$|\mathcal{P}(E)| = |\{0, 1\}^E| = |\{0, 1\}|^{|E|} = 2^{|E|}$$

Nous allons maintenant introduire une notation déjà utilisée pour les coefficients binomiaux au début de l'année. Évidemment, nous allons prouver que ces nombres notés comme des coefficients binomiaux sont bien les coefficients binomiaux.

4.2 Nombre de parties de cardinal donné

Notation. Soit E un ensemble fini de cardinal n . Soit $k \in \mathbb{N}$. On note $\mathcal{P}_k(E)$ l'ensemble des parties de E de cardinal k . On note $\binom{n}{k}$ le cardinal de $\mathcal{P}_k(E)$, c'est à dire le nombre de parties de E de cardinal k .

Remarque 5.

- Pour que notre notation ait un sens, il conviendrait de vérifier que le cardinal de $\mathcal{P}_k(E)$ ne dépend que du cardinal de E et pas de E . En d'autres termes, si E et E' sont deux ensembles finis de même cardinal alors $\mathcal{P}_k(E)$ et $\mathcal{P}_k(E')$ ont le même cardinal. La vérification est laissée au lecteur.
- On a bien évidemment $\binom{n}{k} = 0$ dès que $k > n$ puisque, dans ce cas, $\mathcal{P}_k(E) = \emptyset$. Le problème est de calculer cette quantité lorsque $0 \leq k \leq n$.

Proposition 18. Soit $n \in \mathbb{N}$. On a les propriétés suivantes :

- Pour tout $k \in \llbracket 0, n \rrbracket$, $\binom{n}{k} = \binom{n}{n-k}$
- $\sum_{k=0}^n \binom{n}{k} = 2^n$
- Pour tout $k \in \mathbb{N}$, $\binom{n+1}{k+1} = \binom{n}{k} + \binom{n}{k+1}$ (triangle de Pascal)

Démonstration.

- Soit E un ensemble de cardinal n . Pour toute partie A de E de cardinal k , $E \setminus A$ est de cardinal $n - k$. On vérifie alors que l'application $A \mapsto E \setminus A$ est une bijection de $\mathcal{P}_k(E)$ sur $\mathcal{P}_{n-k}(E)$ d'où la première égalité.
- Soit E un ensemble de cardinal n . La seconde égalité résulte de

$$\mathcal{P}(E) = \bigcup_{k=0}^n \mathcal{P}_k(E)$$

puis de la formule donnant le cardinal d'une union disjointe.

- Pour la troisième égalité, soit E un ensemble de cardinal $n+1$. posons $E = E' \cup \{a\}$ où E' est de cardinal n . Les parties de E de cardinal $k+1$ sont de deux sortes. Il y a tout d'abord celles qui ne contiennent pas a , et qui sont les éléments de $\mathcal{P}_{k+1}(E')$. Il y en a $\binom{n}{k+1}$. Et il y a celles qui contiennent a , et qui sont de la forme $A \cup \{a\}$ où $A \in \mathcal{P}_k(E')$. Il y en a $\binom{n}{k}$. D'où l'égalité cherchée.

□

Proposition 19. Soient $n, k \in \mathbb{N}$ tels que $0 \leq k \leq n$. Alors :

$$\binom{n}{k} = \frac{n!}{k!(n-k)!}$$

Démonstration. On montre par récurrence sur n que pour tout entier $n \in \mathbb{N}$,

$$\forall k \in \llbracket 0, n \rrbracket, \binom{n}{k} = \frac{n!}{k!(n-k)!}$$

Pour $n = 0$ c'est évident.

Soit $n \in \mathbb{N}$. Supposons la propriété vraie pour n . Soit $1 \leq k \leq n$. On a

$$\binom{n+1}{k} = \binom{n}{k} + \binom{n}{k-1}$$

Comme $1 \leq k \leq n$ on peut utiliser l'hypothèse de récurrence sur les deux termes du membre de droite :

$$\begin{aligned} \binom{n+1}{k} &= \frac{n!}{k!(n-k)!} + \frac{n!}{(k-1)!(n-k+1)!} \\ &= \frac{n!}{k!(n+1-k)!}((n+1-k) + k) \\ &= \frac{(n+1)!}{k!(n+1-k)!} \end{aligned}$$

Il reste à vérifier la propriété pour $k = 0$ et $k = n + 1$. Soit E un ensemble de cardinal $n + 1$. $\mathcal{P}_0(E) = \{\emptyset\}$ est un singleton, donc

$$\binom{n+1}{0} = 1 = \frac{(n+1)!}{0!(n+1-0)!}$$

De même, $\mathcal{P}_{n+1}(E) = \{E\}$ est un singleton, donc

$$\binom{n+1}{n+1} = 1 = \frac{(n+1)!}{(n+1)!(n+1-(n+1))!}$$

□

Proposition 20. Soit $\mathbb{A}$ un anneau. Soient $a, b \in \mathbb{A}$ tels que $ab = ba$. Soit $n \in \mathbb{N}$. Alors

$$(a+b)^n = \sum_{k=0}^n \binom{n}{k} a^k b^{n-k}$$

Démonstration. Il s'agit de la formule du binôme, déjà prouvée par récurrence dans un chapitre antérieur. Nous allons en faire ici une « démonstration combinatoire », sans entrer dans les micro-détails. On a

$$(a+b)^n = (a+b)(a+b) \dots (a+b)$$

Le développement de ce produit de n sommes de deux termes donne 2^n termes. Chacun de ces termes est de la forme

$$a^k b^{n-k}$$

où k est le nombre de fois que l'on a choisi a dans un facteur. On a alors évidemment choisi b dans les $n - k$ autres facteurs. Ainsi,

$$(a + b)^n = \sum_{k=0}^n C(n, k) a^k b^{n-k}$$

où $C(n, k)$ est le nombre de façons de choisir a dans k facteurs parmi n . Mais ce choix de k facteurs n'est rien d'autre que le choix d'une partie de cardinal k dans l'ensemble des n facteurs. On a donc

$$C(n, k) = \binom{n}{k}$$

d'où la formule cherchée. $\square$

5 Annexe : la formule du crible

On donne dans ce paragraphe une égalité qui permet d'écrire le cardinal d'une réunion comme une somme de cardinaux d'intersections : la *formule du crible*, aussi connue sous le nom de *formule de Poincaré*. La démonstration donnée ici utilise les relations coefficients-racines pour les polynômes à l'aide des polynômes symétriques élémentaires.

Proposition 21. Soit E un ensemble fini. Soient $A_1, \dots, A_n$ n parties de E . On a

$$\left| \bigcup_{i=1}^n A_i \right| = \sum_{k=1}^n (-1)^{k+1} \sum_{1 \leq i_1 < \dots < i_k \leq n} |A_{i_1} \cap \dots \cap A_{i_k}|$$

Démonstration. Soit E un ensemble fini. Soit A une partie de E . On rappelle que la fonction indicatrice de A est la fonction $\mathbb{1}_A : E \rightarrow \mathbb{R}$ définie par $\mathbb{1}_A(x) = 0$ si $x \notin A$ et $\mathbb{1}_A(x) = 1$ si $x \in A$. Les propriétés ci-dessous sont faciles à démontrer :

- $\forall A \subset E, \mathbb{1}_{E \setminus A} = 1 - \mathbb{1}_A$.
- $\forall A, B \subset E, \mathbb{1}_{A \cap B} = \mathbb{1}_A \mathbb{1}_B$.
- $\forall A \subset E, |A| = \sum_{x \in E} \mathbb{1}_A(x)$.

Donnons-nous n parties de E , $A_1, A_2, \dots, A_n$. Les lois de Morgan nous permettent d'écrire

$$\bigcup_{i=1}^n A_i = E \setminus \bigcap_{i=1}^n (E \setminus A_i)$$

Un passage à la fonction caractéristique donne

$$\begin{aligned}
 \mathbb{1}_{\cup_{i=1}^n A_i} &= 1 - \mathbb{1}_{\cap_{i=1}^n E \setminus A_i} \\
 &= 1 - \prod_{i=1}^n \mathbb{1}_{E \setminus A_i} \\
 \mathbb{1}_{\cup_{i=1}^n A_i} &= 1 - \prod_{i=1}^n (1 - \mathbb{1}_{A_i})
 \end{aligned} \tag{29.1}$$

Soient maintenant $x_1, \dots, x_n$ n éléments d'un anneau commutatif $\mathbb{A}$. Le polynôme

$$P = \prod_{i=1}^n (X - x_i) \in \mathbb{A}[X]$$

se développe en

$$\begin{aligned}
 P &= X^n - \sigma_1 X^{n-1} + \dots + (-1)^n \sigma_n \\
 &= X^n + \sum_{k=1}^n (-1)^k \sigma_k X^{n-k}
 \end{aligned}$$

où

$$\sigma_k = \sum_{1 \leq i_1 < \dots < i_k \leq n} x_{i_1} \dots x_{i_k}$$

est le k ième polynôme symétrique élémentaire en $x_1, \dots, x_n$. En évaluant le polynôme P en $X = 1$, on obtient

$$\prod_{i=1}^n (1 - x_i) = 1 + \sum_{k=1}^n (-1)^k \sigma_k \tag{29.2}$$

Revenons à nos ensembles. Appliquons la formule (2) à $x_i = \mathbb{1}_{A_i}$, $i = 1, \dots, n$. Ceci est une opération parfaitement licite, puisque les fonctions $\mathbb{1}_{A_i}$ sont des éléments de l'anneau commutatif $\mathbb{R}^E$. Nous obtenons

$$\mathbb{1}_{\cup_{i=1}^n A_i} = 1 - 1 + \sum_{k=1}^n (-1)^k \sigma_k = \sum_{k=1}^n (-1)^{k+1} \sigma_k \tag{29.3}$$

où

$$\sigma_k = \sum_{1 \leq i_1 < \dots < i_k \leq n} \mathbb{1}_{A_{i_1}} \dots \mathbb{1}_{A_{i_k}} = \sum_{1 \leq i_1 < \dots < i_k \leq n} \mathbb{1}_{A_{i_1} \cap \dots \cap A_{i_k}}$$

On reporte cette valeur dans l'équation (3), ce qui nous donne l'égalité entre fonctions indicatrices

$$\mathbb{1}_{\cup_{i=1}^n A_i} = \sum_{k=1}^n (-1)^{k+1} \sum_{1 \leq i_1 < \dots < i_k \leq n} \mathbb{1}_{A_{i_1} \cap \dots \cap A_{i_k}} \tag{29.4}$$

Il ne nous reste plus qu'à déduire de l'égalité (4) une égalité entre cardinaux. Pour cela, on somme cette égalité sur tous les éléments x de E

$$\begin{aligned}
 \left| \bigcup_{i=1}^n A_i \right| &= \sum_{x \in E} \sum_{k=1}^n (-1)^{k+1} \sum_{1 \leq i_1 < \dots < i_k \leq n} \mathbb{1}_{A_{i_1} \cap \dots \cap A_{i_k}}(x) \\
 &= \sum_{k=1}^n (-1)^{k+1} \sum_{1 \leq i_1 < \dots < i_k \leq n} \sum_{x \in E} \mathbb{1}_{A_{i_1} \cap \dots \cap A_{i_k}}(x) \\
 &= \sum_{k=1}^n (-1)^{k+1} \sum_{1 \leq i_1 < \dots < i_k \leq n} |A_{i_1} \cap \dots \cap A_{i_k}| \tag{29.5}
 \end{aligned}$$

□

Chapitre 30

Probabilités

1 Espaces probabilisables

1.1 Univers

L'objet de la théorie des probabilités est de modéliser puis d'étudier certains types d'expériences dont le résultat est « aléatoire ». Pour prendre des exemples issus du monde des jeux, citons : lancer une pièce, lancer deux dés, tirer une carte, etc. À chaque répétition d'une telle expérience, *l'issue* de l'expérience change d'une façon qui semble imprévisible. Bien qu'imprévisible, cette issue est quand même restreinte à appartenir à un certain ensemble de valeurs.

Pour étudier une telle expérience, on commence par *modéliser* les issues de l'expérience par des objets mathématiques. L'ensemble des issues de l'expérience est appelé *l'univers*.

Définition 1. Un *univers* est un ensemble non vide.

Exemple. On lance une pièce. Un univers qui semble adapté est l'ensemble $\Omega = \{P, F\}$.

Exemple. On jette deux dés à 6 faces. On peut prendre pour univers l'ensemble $\Omega = \llbracket 1, 6 \rrbracket^2$ des couples d'entiers entre 1 et 6.

Exemple. On jette deux fois de suite un dé à 6 faces. Bien que l'expérience soit différente de celle de l'exemple précédent, le même univers la modélisera aussi.

Exemple. On tire 5 cartes dans un jeu de 52 cartes. On peut prendre pour univers l'ensemble $\mathcal{P}_5(\llbracket 1, 52 \rrbracket)$ des parties à 5 éléments de l'ensemble $\llbracket 1, 52 \rrbracket$.

Exemple. On place 8 reines sur un échiquier. On peut modéliser cette expérience par l'univers

$$\Omega = (\llbracket 1, 8 \rrbracket^2)^\mathbb{S}$$

des octuplets de couples distincts d'entiers entre 1 et 8.

1.2 Événements, tribus

Lorsqu'on effectue une expérience, on s'intéresse la plupart du temps à la réalisation ou pas de certaines propriétés de l'issue de l'expérience. Par exemple, on lance deux pièces et on veut savoir si l'une des pièces est tombée sur pile. Ou on lance deux dés, et on veut savoir si la somme des dés est un nombre premier. Ou encore on tire 5 cartes et on veut savoir si on a obtenu deux paires. Plusieurs issues de l'expérience ont la propriété voulue. On s'intéresse ainsi à une partie de l'univers, ce que l'on appelle un *événement*.

Une première définition possible serait la suivante.

Définition 2. Un événement est une partie de l'univers Ω .

En d'autres termes, les événements sont les éléments de $\mathcal{P}(\Omega)$. Il s'avère que dans certains cas, en particuliers lorsque l'univers est très gros, cette définition pose problème. On est ainsi amenés à ne considérer comme des événements que certaines parties de Ω .

Pour tout $A \subset \Omega$, notons $\overline{A} = \Omega \setminus A$ le complémentaire de A .

Définition 3. Une *tribu* sur l'univers Ω est un ensemble $\mathfrak{T}$ de parties de Ω vérifiant

- NON VIDE. $\mathfrak{T} \neq \emptyset$.
- STABILITÉ PAR COMPLÉMENTAIRE. Pour tout $A \in \mathfrak{T}$, $\overline{A} \in \mathfrak{T}$.
- STABILITÉ PAR UNION DÉNOMBRABLE. Pour toute suite $(A_n)_{n \in \mathbb{N}}$ d'éléments de $\mathfrak{T}$, $\bigcup_{n \in \mathbb{N}} A_n \in \mathfrak{T}$.

Les tribus vérifient certaines propriétés qu'il est bon de connaître.

Proposition 1. Une tribu est stable par union finie, intersection finie et intersection dénombrable.

Démonstration. Soit $\mathfrak{T}$ une tribu sur un univers Ω . Soient $A, B \in \mathfrak{T}$. La suite $(A, B, B, B, \dots)$ est une suite d'éléments de $\mathfrak{T}$, donc

$$A \cup B = A \cup B \cup B \cup \dots \in \mathfrak{T}$$

$\mathfrak{T}$ est stable par réunion de deux ensembles. Par récurrence sur n , pour tout $n \geq 1$, $\mathfrak{T}$ est stable par réunion de n ensembles.

Soit $(A_n)_{n \in \mathbb{N}}$ une suite d'éléments de $\mathfrak{T}$. On a

$$\bigcap_{n \in \mathbb{N}} A_n = \overline{\bigcup_{n \in \mathbb{N}} \overline{A_n}}$$

d'où la stabilité par intersection dénombrable. La stabilité par intersection finie résulte de la stabilité par passage au complémentaire et par réunion finie. $\square$

Proposition 2. Soit $\mathfrak{T}$ une tribu sur Ω . Alors, $\emptyset \in \mathfrak{T}$ et $\Omega \in \mathfrak{T}$.

Démonstration. Soit $A \in \mathfrak{T}$ (un tel A existe car une tribu n'est pas vide). On a aussi $\overline{A} \in \mathfrak{T}$. De là, $\Omega = A \cup \overline{A} \in \mathfrak{T}$, puis $\emptyset = \overline{\Omega} \in \mathfrak{T}$. $\square$

Exemple.

- La plus grande tribu sur Ω est $\mathcal{P}(\Omega)$.
- La plus petite tribu sur Ω est $\{\emptyset, \Omega\}$.

Proposition 3. TRIBUS SUR UN UNIVERS FINI OU DÉNOMBRABLE

Soit Ω un univers fini ou dénombrable. Soit $\mathfrak{T}$ une tribu sur Ω qui contient les singletons. Alors, $\mathfrak{T} = \mathcal{P}(\Omega)$.

Démonstration. Soit $A \subset \Omega$. On a

$$A = \bigcup_{\omega \in \Omega} \{\omega\}$$

Pour tout $\omega \in \Omega$, $\{\omega\} \in \mathfrak{T}$. De plus, cette réunion est finie ou dénombrable et donc $A \in \mathfrak{T}$. $\square$

Remarque 1. Lorsqu'on fait des probabilités, il est souhaitable que les singletons soient des événements. Nous les appellerons d'ailleurs des événements élémentaires. Ainsi, par la proposition précédente, si Ω est fini ou dénombrable, la seule tribu dont on puisse munir Ω pour en faire un espace probabilisé est $\mathcal{P}(\Omega)$. Notre définition 1 de la notion d'événement est dans ce cas correcte.

Si Ω est infini mais pas dénombrable, il faut être plus précis.

Définition 4. Un *espace probabilisable* est un couple $(\Omega, \mathfrak{T})$ où Ω est un univers et $\mathfrak{T}$ est une tribu sur Ω qui contient les singletons. Un *événement* est un élément de $\mathfrak{T}$.

1.3 Une tribu sur $\Omega = [0, 1]$

Alors, beaucoup de bruit pour rien ? Non, car si l'univers Ω est infini non dénombrable, la situation est très différente. Dans ce cas, il n'est pas souhaitable de munir Ω de la tribu $\mathcal{P}(\Omega)$. Les raisons pour lesquelles cela n'est pas souhaitable dépassent le cadre de ce cours.

Donnons un exemple important de tribu sur l'univers $\Omega = [0, 1]$. L'intersection $\mathcal{B}$ des tribus sur Ω qui contiennent tous les intervalles inclus dans Ω est encore une tribu sur Ω qui contient les intervalles. Les éléments de la tribu $\mathcal{B}$ sont les *ensembles boréliens*.

Comme exemples de boréliens, nous avons évidemment les intervalles, et donc aussi les singletons (qui sont des intervalles). Et donc aussi tous les ensembles dénombrables, comme par exemple $A = \mathbb{Q} \cap [0, 1]$. Et aussi les complémentaires des ensembles dénombrables. Et encore beaucoup d'autres ensembles. En fait, il n'est pas du tout facile de donner un exemple explicite d'une partie de Ω qui ne soit pas un borélien. Et pourtant, il existe beaucoup d'ensembles non boréliens : on peut montrer qu'il existe une bijection entre $\mathcal{B}$ et Ω alors que $\mathcal{P}(\Omega)$ est beaucoup plus « gros » que Ω . Ainsi, *la plupart des parties de Ω ne sont pas des boréliens*.

Exercice. Soit Ω un univers. Soit E un ensemble de parties de Ω . Montrer que l'intersection $\mathfrak{T}$ des tribus sur Ω qui contiennent E est encore une tribu sur Ω qui contient E . Montrer que $\mathfrak{T}$ est la plus petite tribu sur Ω qui contient E .

1.4 Le vocabulaire des probabilités

Les probabilistes utilisent leur propre vocabulaire, un ensemble de mots qui sont des synonymes de mots que nous utilisons couramment dans les autres branches des mathématiques. Voici un petit dictionnaire, qui permet de passer du vocabulaire de la théorie des ensembles à celui des probabilités, et vice versa. On suppose donné un espace probabilisable $(\Omega, \mathfrak{T})$. Dans le tableau, A et B sont des événements et $\omega \in \Omega$.

	Théorie des ensembles	Probabilités
A	Élément de $\mathfrak{T}$	Événement
$\{\omega\}$	Singleton	Événement élémentaire
ω	Élément de Ω	Issue d'une expérience
$\Omega \setminus A$	Complémentaire de A	Contraire de A
$A \cup B$	Réunion de A et B	A ou B
$A \cap B$	Intersection de A et B	A et B
$\emptyset$	Ensemble vide	Événement impossible
Ω		Univers, événement certain
$A \cap B = \emptyset$	Ensembles disjoints	Événements incompatibles

Dans la suite, hormis les mots « univers » et « événement », j'utiliserai le vocabulaire ensembliste.

Encore un mot probabiliste.

Définition 5. SYSTÈME COMPLET D'ÉVÉNEMENTS

Soit $(\Omega, \mathfrak{T})$ un espace probabilisable. Un *système complet d'événements* est un ensemble $\mathcal{S}$ d'événements fini ou dénombrable vérifiant :

- Les éléments de $\mathcal{S}$ sont deux à deux disjoints.
- $\bigcup_{A \in \mathcal{S}} A = \Omega$.

2 Espaces probabilisés

2.1 Notion de probabilité

Définition 6. Soit $(\Omega, \mathfrak{T})$ un espace probabilisable. Une *probabilité* sur $(\Omega, \mathfrak{T})$ est une fonction $\mathbb{P} : \mathfrak{T} \longrightarrow [0, 1]$ vérifiant les deux propriétés suivantes.

- MASSE TOTALE 1. $\mathbb{P}(\Omega) = 1$.
- ADDITIVITÉ DÉNOMBRABLE. Pour toute suite $(A_n)_{n \in \mathbb{N}}$ d'éléments de $\mathfrak{T}$ deux à deux disjoints,

$$\mathbb{P}\left(\bigcup_{n \in \mathbb{N}} A_n\right) = \sum_{n=0}^{\infty} \mathbb{P}(A_n)$$

La propriété d'additivité dénombrable signifie précisément que la suite

$$\left(\sum_{n=0}^N \mathbb{P}(A_n)\right)_{N \in \mathbb{N}}$$

est convergente, et que sa limite est

$$\mathbb{P}\left(\bigcup_{n \in \mathbb{N}} A_n\right)$$

Remarquons qu'une telle suite est croissante. Elle a donc une limite, réelle ou égale à $+\infty$. La définition nous dit que si les A_n sont deux à deux disjoints, alors la limite est réelle, et elle est égale à la probabilité de la réunion des A_n .

Définition 7. Un *espace probabilisé* est un triplet $(\Omega, \mathfrak{T}, \mathbb{P})$ où $(\Omega, \mathfrak{T})$ est un espace probabilisable et $\mathbb{P}$ est une probabilité sur $(\Omega, \mathfrak{T})$.

Proposition 4. PROBABILITÉ DU VIDE

Soit $(\Omega, \mathfrak{T}, \mathbb{P})$ un espace probabilisé. On a $\mathbb{P}(\emptyset) = 0$.

Démonstration. La famille $(\emptyset)_{n \in \mathbb{N}}$ est une suite d'événements disjoints deux à deux et sa réunion est $\emptyset$. On a donc

$$\mathbb{P}(\emptyset) = \sum_{n=0}^{\infty} \mathbb{P}(\emptyset)$$

Ceci signifie que lorsque N tend vers l'infini,

$$\sum_{n=0}^N \mathbb{P}(\emptyset) = (N+1)\mathbb{P}(\emptyset) \longrightarrow \mathbb{P}(\emptyset)$$

Clairement, $\mathbb{P}(\emptyset) \neq 0$ amène une contradiction. $\square$

Proposition 5. ADDITIVITÉ FINIE

Soit $(\Omega, \mathfrak{T}, \mathbb{P})$ un espace probabilisé. Pour tout $n \geq 1$ et tous événements $A_1, \dots, A_n$ disjoints deux à deux,

$$\mathbb{P}\left(\bigcup_{k=1}^n A_k\right) = \sum_{k=1}^n \mathbb{P}(A_k)$$

Démonstration. Considérer la suite $(A_1, \dots, A_n, \emptyset, \emptyset, \dots)$ et appliquer la propriété d'additivité dénombrable. $\square$

2.2 Le cas des univers finis

Si l'univers Ω est fini, les tribus sur Ω ne contiennent qu'un nombre fini d'ensembles. Il est probablement inutile de demander à une probabilité sur Ω de vérifier l'additivité dénombrable. Il suffit en fait d'imposer l'additivité finie, et même, en fait, l'additivité pour deux ensembles disjoints.

Proposition 6. Soit $(\Omega, \mathcal{P}(\Omega))$ un espace probabilisable fini. Soit $\mathbb{P} : \mathcal{P}(\Omega) \rightarrow [0, 1]$. La fonction $\mathbb{P}$ est une probabilité si et seulement si

- $\mathbb{P}(\Omega) = 1$.
- Pour tous événements A et B disjoints,

$$\mathbb{P}(A \cup B) = \mathbb{P}(A) + \mathbb{P}(B)$$

Démonstration. Exercice. $\square$

2.3 Un gros espace probabilisé

Nous avons parlé plus haut de l'univers $\Omega = [0, 1]$ et de la tribu $\mathcal{B}$ des boréliens. On peut montrer (et ce n'est pas facile) qu'il existe une probabilité $\mathbb{P}$ sur $(\Omega, \mathcal{B})$ telle que pour tout intervalle I inclus dans $[0, 1]$, $\mathbb{P}(I)$ soit la longueur de I . La probabilité $\mathbb{P}$ est la *mesure de Lebesgue* sur $[0, 1]$.

Considérons l'ensemble $A = \mathbb{Q} \cap \Omega$. L'ensemble A est un ensemble dénombrable, donc $A \in \mathcal{B}$. Que vaut $\mathbb{P}(A)$? On a

$$A = \bigcup_{x \in A} \{x\}$$

La réunion ci-dessus est une union disjointe donc, par la propriété d'additivité dénombrable,

$$\mathbb{P}(A) = \sum_{x \in A} \mathbb{P}(\{x\}) = \sum_{x \in A} 0 = 0$$

Nous avons donc ici un exemple d'événement non vide, et même infini, et pourtant de probabilité nulle. Ne confondons donc pas impossibilité et probabilité nulle !

Moralité : ne disons plus « il n'y a aucune chance pour que cela arrive », parce que cela peut quand même arriver.

2.4 Fonction de masse

Nous nous plaçons dans ce paragraphe sur des univers **finis**. Les résultats que nous allons montrer peuvent s'étendre à des univers dénombrables, mais nous ne le ferons pas ici. Ce que nous allons appeler *fonction de masse* s'appelle parfois autrement chez certains auteurs, par exemple *fonction de distribution*.

Définition 8. Soit $(\Omega, \mathcal{P}(\Omega), \mathbb{P})$ un espace probabilisé fini. La *fonction de masse* de la probabilité $\mathbb{P}$ est la fonction $\mu_{\mathbb{P}} : \Omega \longrightarrow [0, 1]$ définie pour tout $\omega \in \Omega$ par

$$\mu_{\mathbb{P}}(\omega) = \mathbb{P}(\{\omega\})$$

Proposition 7. Soit $(\Omega, \mathcal{P}(\Omega), \mathbb{P})$ un espace probabilisé fini. On a

$$\sum_{\omega \in \Omega} \mu_{\mathbb{P}}(\omega) = 1$$

Démonstration. En effet,

$$\Omega = \bigcup_{\omega \in \Omega} \{\omega\}$$

et cette union est disjointe. On a donc

$$1 = \mathbb{P}(\Omega) = \sum_{\omega \in \Omega} \mathbb{P}(\{\omega\}) = \sum_{\omega \in \Omega} \mu_{\mathbb{P}}(\omega)$$

□

Ceci suggère une définition.

Définition 9. FONCTION DE MASSE

Soit Ω un univers fini. Une fonction $f : \Omega \longrightarrow [0, 1]$ est une *fonction de masse* si

$$\sum_{\omega \in \Omega} f(\omega) = 1$$

Proposition 8. *Proposition Soit $(\Omega, \mathcal{P}(\Omega))$ un espace probabilisable fini. Soit f une fonction de masse définie sur Ω . Alors, il existe une unique probabilité $\mathbb{P}$ sur $(\Omega, \mathcal{P}(\Omega))$ telle que $\mu_{\mathbb{P}} = f$.*

Démonstration. Soit $\mathbb{P}$ une probabilité sur $(\Omega, \mathcal{P}(\Omega))$. Supposons que $\mu_{\mathbb{P}} = f$. On a alors pour tout $A \in \mathcal{P}(\Omega)$

$$\mathbb{P}(A) = \mathbb{P}\left(\bigcup_{\omega \in A} \{\omega\}\right)$$

L'union étant disjointe,

$$\mathbb{P}(A) = \sum_{\omega \in A} \mathbb{P}(\{\omega\}) = \sum_{\omega \in A} f(\omega)$$

On en déduit l'unicité. Inversement on vérifie que la fonction $\mathbb{P}$ ci-dessus est une probabilité sur $(\Omega, \mathcal{P}(\Omega))$ et $\mu_{\mathbb{P}} = f$, d'où l'existence. $\square$

Exemple. Soit Ω un univers fini de cardinal $n \in \mathbb{N}^*$. Un exemple essentiel de fonction de masse sur Ω est la fonction $f : \Omega \rightarrow [0, 1]$ définie par

$$f(\omega) = \frac{1}{n}$$

L'unique probabilité $\mathbb{P}$ sur $(\Omega, \mathcal{P}(\Omega))$ telle que $f_{\mathbb{P}} = f$ est la *probabilité uniforme* sur Ω . On a dans ce cas pour tout $A \subset \Omega$,

$$\mathbb{P}(A) = \frac{1}{n} \sum_{\omega \in A} 1 = \frac{|A|}{|\Omega|}$$

Remarque 2. Les énoncés de probabilités suggèrent parfois l'utilisation de probabilités uniformes : on lance un *dé parfait*, on jette une *pièce équilibrée*, on tire 5 cartes *au hasard* dans un jeu de cartes, etc. Les mots en italiques indiquent l'uniformité de la probabilité sous-jacente. Attention, toutefois, à ne pas tomber dans l'excès de croire que toutes les probabilités sont uniformes !

2.5 Premières propriétés

Dans ce paragraphe, $(\Omega, \mathfrak{T}, \mathbb{P})$ est un espace probabilisé.

Proposition 9. *Pour tous $A, B \in \mathfrak{T}$ tels que $A \subset B$,*

- $\mathbb{P}(A) \leq \mathbb{P}(B)$.
- $\mathbb{P}(B \setminus A) = \mathbb{P}(B) - \mathbb{P}(A)$.

Démonstration. Soient $A, B \in \mathfrak{T}$ tels que $A \subset B$. On a $B = A \cup (B \setminus A)$ et cette union est disjointe, donc

$$\mathbb{P}(B) = \mathbb{P}(A) + \mathbb{P}(B \setminus A)$$

d'où le résultat. $\square$

Corollaire 10. Pour tout $A \in \mathfrak{T}$, $\mathbb{P}(\overline{A}) = 1 - \mathbb{P}(A)$.

Démonstration. Prendre $B = \Omega$ dans la proposition précédente. $\square$

Proposition 11. Pour tous $A, B \in \mathfrak{T}$,

$$\mathbb{P}(A \cup B) = \mathbb{P}(A) + \mathbb{P}(B) - \mathbb{P}(A \cap B)$$

Démonstration. Écrivons

$$A \cup B = (A \setminus (A \cap B)) \cup B$$

Cette union est disjointe, donc

$$\mathbb{P}(A \cup B) = \mathbb{P}(A \setminus (A \cap B)) + \mathbb{P}(B) = \mathbb{P}(A) - \mathbb{P}(A \cap B) + \mathbb{P}(B)$$

$\square$

Remarque 3. Soient $A, B, C \in \mathfrak{T}$. On a

$$\begin{aligned} \mathbb{P}(A \cup B \cup C) &= \mathbb{P}(A \cup B) + \mathbb{P}(C) - \mathbb{P}((A \cup B) \cap C) \\ &= \mathbb{P}(A) + \mathbb{P}(B) - \mathbb{P}(A \cap B) + \mathbb{P}(C) - \mathbb{P}((A \cap C) \cup (B \cap C)) \end{aligned}$$

et une dernière application de la formule pour la réunion de deux ensembles fournit

$$\mathbb{P}(A \cup B \cup C) = \mathbb{P}(A) + \mathbb{P}(B) + \mathbb{P}(C) - \mathbb{P}(A \cap B) - \mathbb{P}(A \cap C) - \mathbb{P}(B \cap C) + \mathbb{P}(A \cap B \cap C)$$

Cette formule se généralise à des réunions de n ensembles : c'est la *formule de Poincaré*, appelée aussi la *formule du crible*.

Proposition 12. Soit $n \in \mathbb{N}$. Soient $A_1, \dots, A_n \in \mathfrak{T}$. Alors,

$$\mathbb{P}\left(\bigcup_{i=1}^n A_i\right) \leq \sum_{i=1}^n \mathbb{P}(A_i)$$

Démonstration. Faisons une récurrence sur n . C'est clair pour $n = 0$. Soit $n \in \mathbb{N}$. Supposons la propriété vraie pour n . Soient $A_1, \dots, A_{n+1} \in \mathfrak{T}$. Posons

$$\begin{aligned}
 B &= \bigcup_{i=1}^n A_i \\
 \mathbb{P}\left(\bigcup_{i=1}^{n+1} A_i\right) &= \mathbb{P}(B \cup A_{n+1}) \\
 &= \mathbb{P}(B) + \mathbb{P}(A_{n+1}) - \mathbb{P}(B \cap A_{n+1}) \\
 &\leq \mathbb{P}(B) + \mathbb{P}(A_{n+1}) \\
 &\leq \sum_{i=1}^n \mathbb{P}(A_i) + \mathbb{P}(A_{n+1}) \\
 &= \sum_{i=1}^{n+1} \mathbb{P}(A_i)
 \end{aligned}$$

□

3 Urnes

De nombreux problèmes de probabilités font intervenir des *tirages de boules dans des urnes* : boules numérotées, boules colorées, plusieurs urnes, rajout de boules, etc. L'intérêt de la formulation en termes d'urnes est la plus ou moins grande universalité de ce genre de problème. Par exemple, on peut modéliser le lancer d'une pièce parfaite avec une urne contenant une boule F et une boule P . Si on veut lancer une pièce parfaite ayant une probabilité $1/3$ de tomber sur pile, il suffit de mettre une boule F et deux boules P dans une urne, etc.

Nous allons considérer dans ce qui suit deux problèmes très classiques :

- Les successions de tirages *avec remise* dans une urne.
- Les successions de tirages *sans remise* dans une urne.

Le premier problème nous conduira à une loi de probabilités appelée la *loi binomiale*. Le second problème mettra en évidence la *loi hypergéométrique*.

3.1 Tirages avec remise dans une urne : la loi binomiale

Considérons l'expérience suivante.

Données. $N \in \mathbb{N}^*$, $n \in \mathbb{N}$ et $r \in \llbracket 0, N \rrbracket$. On pose $v = N - r$.

Expérience. Une urne contient N boules numérotées de 1 à N . Les boules numérotées de 1 à r sont rouges, les autres sont vertes. On effectue n fois l'opération suivante : on prend

une boule dans l'urne, on regarde sa couleur, on remet la boule dans l'urne, et on mélange bien. Soit $k \in \llbracket 0, n \rrbracket$. Quelle est la probabilité d'avoir tiré k boules rouges ?

Modèle. On choisit pour univers $\Omega = \llbracket 1, N \rrbracket^n$, l'ensemble des n -uplets d'éléments de $\llbracket 1, N \rrbracket$. Le cardinal de Ω est N^n . Chaque élément de Ω représente la liste des boules tirées, dans l'ordre chronologique. Une même boule peut y apparaître plusieurs fois car à chaque tirage on remet la boule dans l'urne. On munit Ω de la tribu $\mathcal{P}(\Omega)$ et de la probabilité uniforme $\mathbb{P}$.

Pour tout $k \in \llbracket 0, n \rrbracket$ soit A_k l'événement « on a tiré k boules rouges ». Pour construire un n -uplet élément de A_k ,

- On choisit les k positions du n -uplet où il y a une boule rouge : il y a $\binom{n}{k}$ possibilités.
- On remplit ces k positions avec une boule rouge arbitraire : r^k possibilités.
- On remplit les $n - k$ positions restantes avec une boule verte arbitraire : v^{n-k} possibilités.

On a donc

$$|A_k| = \binom{n}{k} r^k v^{n-k}$$

En posant $p = r/N$, la probabilité de A_k est

$$\mathbb{P}(A_k) = \binom{n}{k} \frac{r^k v^{n-k}}{N^n} = \binom{n}{k} \left(\frac{r}{N}\right)^k \left(\frac{v}{N}\right)^{n-k} = \binom{n}{k} p^k (1-p)^{n-k}$$

Nous reparlerons de ces probabilités plus loin, lorsque nous aborderons la loi binomiale $\mathcal{B}(n, p)$.

3.2 Tirages sans remise dans une urne : la loi hypergéométrique

Données. $N \geq 1$, $n \in \llbracket 0, N \rrbracket$ et $r \in \llbracket 0, N \rrbracket$. On pose $v = N - r$.

Expérience. Une urne contient N boules numérotées de 1 à N . Les boules numérotées de 1 à r sont rouges, les autres sont vertes. On effectue n fois l'opération suivante : on prend une boule dans l'urne, on regarde sa couleur, ON NE REMET PAS la boule dans l'urne (et on ne mélange pas, ça ne servirait à rien). Soit $k \in \llbracket 0, n \rrbracket$. Quelle est la probabilité d'avoir tiré k boules rouges ?

Modèle. On choisit pour univers $\Omega = \mathcal{P}_n(\llbracket 1, N \rrbracket)$, l'ensemble des parties à n éléments de $\llbracket 1, N \rrbracket$. Chaque élément de Ω représente l'ensemble des n boules tirées. Le cardinal de Ω est $\binom{N}{n}$.

On munit Ω de la tribu $\mathcal{P}(\Omega)$ et de la probabilité uniforme $\mathbb{P}$.

Pour tout $k \in \llbracket 0, n \rrbracket$ soit A_k l'événement « on a tiré k boules rouges ». Pour construire un élément de A_k ,

- On choisit l'ensemble des k boules rouges qui ont été tirées : il y a $\binom{r}{k}$ possibilités.

- On choisit l'ensemble des $n - k$ boules vertes qui ont été tirées : il y a $\binom{v}{n-k}$ possibilités.

On a donc

$$|A_k| = \binom{r}{k} \binom{v}{n-k}$$

On en déduit

$$\mathbb{P}(A_k) = \frac{\binom{r}{k} \binom{v}{n-k}}{\binom{N}{n}}$$

Remarque 4. Ces probabilités sont elles aussi liées à une loi importante, la loi hypergéométrique $\mathcal{H}(N, n, p)$. Nous ne l'étudierons pas de façon aussi approfondie que la loi binomiale. Ci-dessous, $\mathbb{P}(A_k)$ en fonction de k . On a pris $r = 60$, $v = 40$, $N = 100$ et $n = 30$.

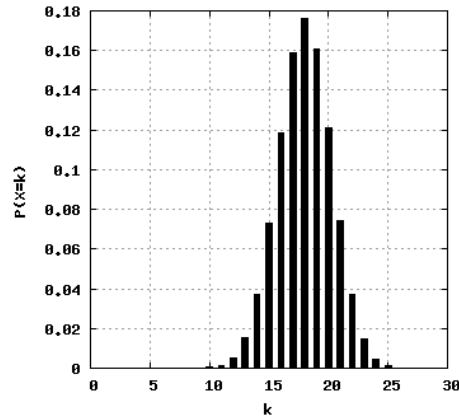


FIGURE 30.1 – Loi hypergéométrique : $r = 60, v = 40, n = 30$

Remarque 5. On peut très bien faire les calculs en choisissant comme univers $\Omega' = \llbracket 1, N \rrbracket^n$, l'ensemble des n -uplets d'éléments distincts de $\llbracket 1, N \rrbracket$, muni de la probabilité uniforme. On a

$$|\Omega'| = N^n$$

Soit B_k l'événement « on a tiré k boules rouges. Pour construire un élément de B_k ,

- On choisit les k indices du n -uplet qui correspondent aux boules rouges : $\binom{n}{k}$ possibilités.
- Pour ces k positions, on choisit k boules rouges distinctes : $r^{\underline{k}}$ possibilités.
- Pour les $n - k$ positions restantes, on choisit $n - k$ boules vertes distinctes : $v^{\underline{n-k}}$ possibilités.

On a donc

$$|B_k| = \binom{n}{k} r^{\underline{k}} v^{\underline{n-k}}$$

d'où

$$\mathbb{P}(B_k) = \binom{n}{k} \frac{r^k v^{n-k}}{N^n} = \frac{\binom{r}{k} \binom{v}{n-k}}{\binom{N}{n}}$$

On retrouve **évidemment** le même résultat. Comment ça, évidemment ? Euh, mais on a deux univers différents Ω et Ω' ! Et on cherche la probabilité de deux événements différents A_k et B_k !!! Miracle ...

4 Probabilités conditionnelles

4.1 Conditionnement

Définition 10. Soit $(\Omega, \mathfrak{T}, \mathbb{P})$ un espace probabilisé. Soient $A, B \in \mathfrak{T}$. On suppose que $\mathbb{P}(B) > 0$. La *probabilité conditionnelle de A sachant B* est

$$\mathbb{P}(A|B) = \frac{\mathbb{P}(A \cap B)}{\mathbb{P}(B)}$$

Exemple. Une urne contient r boules rouges et v boules vertes. Soit $N = r + v$. On tire sans remise deux boules de l'urne. Quelle est la probabilité que la seconde boule soit verte, sachant que la première boule est verte ?

Informellement, une fois qu'on a tiré une boule verte, il reste $N - 1$ boules dans l'urne dont $v - 1$ sont vertes. Si cette histoire de probabilités conditionnelles fonctionne, nous devrions trouver $(v - 1)/(N - 1)$ comme valeur.

Prenons $\Omega = \{(x, y) \in \llbracket 1, N \rrbracket^2, x \neq y\}$ muni de la probabilité uniforme. Pour $i = 1, 2$ soit V_i l'événement « la i ème boule tirée est verte ». On a

$$|V_1| = v(N - 1)$$

$$|V_1 \cap V_2| = v(v - 1)$$

d'où

$$P(V_2|V_1) = \frac{v(v - 1)}{v(N - 1)} = \frac{v - 1}{N - 1}$$

Les probabilités conditionnelles, ça a l'air de marcher.

Proposition 13. Soit $(\Omega, \mathfrak{T}, \mathbb{P})$ un espace probabilisé. Soit $B \in \mathfrak{T}$ tel que $\mathbb{P}(B) > 0$. L'application $\mathbb{P}_B : \mathfrak{T} \rightarrow [0, 1]$ définie par $\mathbb{P}_B(A) = \mathbb{P}(A|B)$ est une probabilité sur $(\Omega, \mathfrak{T})$.

Démonstration. Il y a deux propriétés à vérifier. Tout d'abord,

$$\mathbb{P}_B(\Omega) = \frac{\mathbb{P}(\Omega \cap B)}{\mathbb{P}(B)} = \frac{\mathbb{P}(B)}{\mathbb{P}(B)} = 1$$

Soit $(A_n)_{n \in \mathbb{N}}$ une suite d'événements disjoints deux à deux. On a

$$\left(\bigcup_{n \in \mathbb{N}} A_n \right) \cap B = \bigcup_{n \in \mathbb{N}} (A_n \cap B)$$

Les événements $A_n \cap B$ sont disjoints deux à deux, donc

$$\mathbb{P} \left(\bigcup_{n \in \mathbb{N}} (A_n \cap B) \right) = \sum_{n=0}^{\infty} \mathbb{P}(A_n \cap B)$$

d'où le résultat en divisant par $\mathbb{P}(B)$. $\square$

Remarque 6. Cela va sans dire mais mieux vaut le dire : « $A|B$ », cela n'existe pas. La notation $\mathbb{P}(A|B)$ n'est qu'une notation pratique pour la probabilité de A ...sauf qu'il ne s'agit pas de la probabilité $\mathbb{P}$ mais de la probabilité $\mathbb{P}_B$.

4.2 Probabilités composées, probabilités totales

Proposition 14. PROBABILITÉS COMPOSÉES

Soit $(\Omega, \mathfrak{F}, \mathbb{P})$ un espace probabilisé. Soit $n \in \mathbb{N}^*$. Soient $A_1, \dots, A_n$ n événements tels que

$$\mathbb{P}(A_1 \cap \dots \cap A_n) > 0$$

Alors,

$$\mathbb{P}(A_1 \cap \dots \cap A_n) = \mathbb{P}(A_1) \mathbb{P}(A_2|A_1) \mathbb{P}(A_3|A_1 \cap A_2) \dots \mathbb{P}(A_n|A_1 \cap \dots \cap A_{n-1})$$

Démonstration. On fait une récurrence sur n . Pour $n = 1, 2$ c'est évident. Soit $n \in \mathbb{N}^*$. Supposons la propriété vraie pour n et donnons nous $n + 1$ événements $A_1, \dots, A_{n+1}$ tels que $\mathbb{P}(A_1 \cap \dots \cap A_{n+1}) > 0$. On a

$$\mathbb{P}(A_1 \cap \dots \cap A_{n+1}) = \mathbb{P}(A_{n+1}|A_1 \cap \dots \cap A_n) \mathbb{P}(A_1 \cap \dots \cap A_n)$$

On termine grâce à l'hypothèse de récurrence. $\square$

Exemple. On dispose des objets suivants

- Une pièce parfaite.
- Une urne numéro 1 qui contient 5 boules rouges et 5 boules blanches.

- Une urne numéro 2 qui contient r boules rouges et $10 - r$ boules blanches, où $r \in \llbracket 0, 10 \rrbracket$.
- Un dé rouge parfait.
- Un dé blanc ayant une probabilité $p \in [0, 1]$ de faire 6 et une probabilité égale pour les autres faces.

On fait l'expérience suivante. On lance la pièce. Si elle tombe sur pile, on tire une boule dans l'urne 1. Sinon, on tire une boule dans l'urne 2. On lance le dé qui est de la même couleur que la boule tirée.

Quelle est la probabilité d'avoir fait face, d'avoir tiré une boule blanche, et d'avoir fait 4 avec le dé ?

On pourrait évidemment construire un espace probabilisé fini $(\Omega, \mathcal{P}(\Omega), \mathbb{P})$ qui modélise cette expérience mais il ne serait vraiment pas joli. Nous allons plutôt opérer de la façon suivante :

- Nous disposons d'un espace probabilisé fini $(\Omega, \mathcal{P}(\Omega), \mathbb{P})$ qui modélise cette expérience.
- Les probabilités conditionnelles, ça marche.

Notons

- A_1 l'événement « la pièce est tombée sur face ».
- A_2 l'événement « on a tiré une boule blanche ».
- A_3 l'événement « on a fait 4 avec le dé ».

Par le théorème des probabilités composées,

$$\begin{aligned} \mathbb{P}(A_1 \cap A_2 \cap A_3) &= \mathbb{P}(A_1)\mathbb{P}(A_2|A_1)\mathbb{P}(A_3|A_1 \cap A_2) \\ &= \frac{1}{2} \left(1 - \frac{r}{10}\right) \frac{1}{5} (1-p) \\ &= \frac{1}{10} \left(1 - \frac{r}{10}\right) (1-p) \end{aligned}$$

Proposition 15. PROBABILITÉS TOTALES

Soit $(\Omega, \mathfrak{T}, \mathbb{P})$ un espace probabilisé. Soit $n \in \mathbb{N}^*$. Soit $\{A_1, \dots, A_n\}$ un système complet d'événements de probabilités non nulles. Soit $B \in \mathfrak{T}$. Alors :

$$\mathbb{P}(B) = \sum_{i=1}^n \mathbb{P}(B|A_i)\mathbb{P}(A_i)$$

Démonstration. On a

$$B = B \cap \Omega = B \cap \left(\bigcup_{i=1}^n A_i \right) = \bigcup_{i=1}^n (B \cap A_i)$$

Les A_i sont disjoints deux à deux, donc les $B \cap A_i$ aussi. On a donc

$$\mathbb{P}(B) = \sum_{i=1}^n \mathbb{P}(B \cap A_i) = \sum_{i=1}^n \mathbb{P}(B|A_i)\mathbb{P}(A_i)$$

□

Remarque 7. Pour simplifier le théorème des probabilités totales, on fait fréquemment la convention suivante. Si A et B sont deux événements et $\mathbb{P}(A) = 0$, alors $\mathbb{P}(B|A)$ n'a pas de sens. On convient cependant de poser $\mathbb{P}(B|A)\mathbb{P}(A) = 0$. Via cette convention, on n'a plus besoin de supposer, dans le théorème des probabilités totales, que les A_i ont des probabilités non nulles.

Proposition 16. PROBABILITÉS TOTALES (BIS)

Soit $(\Omega, \mathfrak{T}, \mathbb{P})$ un espace probabilisé. Soit $n \in \mathbb{N}^$. Soit $\{A_1, \dots, A_n\}$ un système complet d'événements. Soit $B \in \mathfrak{T}$. Alors :*

$$\mathbb{P}(B) = \sum_{i=1}^n \mathbb{P}(B|A_i)\mathbb{P}(A_i)$$

Démonstration. Reprenons la preuve ci-dessus. Soit $I = \{i \in \llbracket 1, n \rrbracket, \mathbb{P}(A_i) = 0\}$. Pour tout $i \in I$, $B \cap A_i \subset A_i$ et donc $\mathbb{P}(B \cap A_i) = 0$. Ainsi,

$$\begin{aligned} \mathbb{P}(B) &= \sum_{i=1}^n \mathbb{P}(B \cap A_i) \\ &= \sum_{i \notin I} \mathbb{P}(B \cap A_i) \\ &= \sum_{i \notin I} \mathbb{P}(B|A_i)\mathbb{P}(A_i) \\ &= \sum_{i=1}^n \mathbb{P}(B|A_i)\mathbb{P}(A_i) \end{aligned}$$

où la dernière ligne résulte de notre convention (on a ajouté des zéros à la somme). □

Exercice. Un système quasi-complet d'événements est un ensemble $\{A_1, \dots, A_n\}$ d'événements disjoints tel que

$$\mathbb{P}\left(\bigcup_{k=1}^n A_k\right) = 1$$

Montrer que le théorème des probabilités totales reste vérifié pour un système quasi-complet d'événements de probabilités non nulles. *Indication : ajouter à ce système l'événement $A_{n+1} = \Omega \setminus (A_1 \cup \dots \cup A_n)$. Que vaut $\mathbb{P}(A_{n+1})$? Reprendre ensuite la preuve de la proposition ci-dessus.*

Exemple. Une urne contient $r \geq 1$ boules rouges et $v \geq 1$ boules vertes. Soit $N = r + v$. On tire deux boules de l'urne sans remise. Quelle est la probabilité que la deuxième boule tirée soit verte ?

Prenons $\Omega = \{(x, y) \in \llbracket 1, N \rrbracket^2, x \neq y\}$ muni de la tribu $\mathcal{P}(\Omega)$ et de la probabilité uniforme $\mathbb{P}$. Pour $i = 1, 2$ soit V_i (resp : R_i) l'événement « la i ème boule tirée est verte » (resp. : rouge). L'ensemble $\{V_1, R_1\}$ est un système complet d'événements de probabilités non nulles (v/N et r/N). Par le théorème des probabilités totales,

$$\begin{aligned} \mathbb{P}(V_2) &= \mathbb{P}(V_2|V_1)\mathbb{P}(V_1) + \mathbb{P}(V_2|R_1)\mathbb{P}(R_1) \\ &= \frac{v-1}{N-1} \frac{v}{N} + \frac{v}{N-1} \frac{r}{N} \\ &= \frac{v}{N} \end{aligned}$$

On remarque que la probabilité que la seconde boule soit verte est égale à la probabilité que la première boule soit verte, ce qui n'était pas si évident que cela. Et si on tirait 3 boules ? 8 boules ?

Explication. L'application $(x, y) \mapsto (y, x)$ est une bijection de V_1 sur V_2 . Ces deux événements ont donc même cardinal, et comme nous avons choisi la probabilité uniforme ils ont même probabilité. L'avantage de ce raisonnement est qu'il s'étend sans difficulté au tirage de n boules sans remise.

4.3 Formule de Bayes

Proposition 17. FORMULE DE BAYES

Soit $(\Omega, \mathfrak{F}, \mathbb{P})$ un espace probabilisé. Soient A et B deux événements tels que $P(A) > 0$ et $P(B) > 0$. On a

$$\mathbb{P}(A|B)\mathbb{P}(B) = \mathbb{P}(B|A)\mathbb{P}(A)$$

La formule de Bayes nous dit que si on connaît $\mathbb{P}(A)$, $\mathbb{P}(B)$ et, par exemple, $\mathbb{P}(B|A)$ alors on connaît aussi $\mathbb{P}(A|B)$.

Démonstration. Immédiate, vu que

$$\mathbb{P}(A \cap B) = \mathbb{P}(A|B)\mathbb{P}(B) = \mathbb{P}(B|A)\mathbb{P}(A)$$

□

Remarque 8. La formule de Bayes est généralement donnée sous une forme beaucoup plus compliquée, faisant intervenir des systèmes complets d'événements et le théorème des probabilités totales. Nous nous en dispenserons ici parce que nous pensons qu'il est plus simple de retenir « notre » formule.

4.4 Événements indépendants

Définition 11. Soit $(\Omega, \mathfrak{T}, \mathbb{P})$ un espace probabilisé. Soient $A, B \in \mathfrak{T}$. les événements A et B sont *indépendants* si

$$\mathbb{P}(A \cap B) = \mathbb{P}(A)\mathbb{P}(B)$$

Remarque 9. Si $\mathbb{P}(B) > 0$, A et B sont indépendants si et seulement si

$$\mathbb{P}(A|B) = \mathbb{P}(A)$$

En d'autres termes, lors d'une expérience aléatoire, savoir que B s'est réalisé n'augmente pas nos connaissances concernant la réalisation de A .

Exercice. Soient A et B deux événements indépendants d'un espace probabilisé $(\Omega, \mathfrak{T}, \mathbb{P})$. Montrer que les couples $(A, \overline{B})$, $(\overline{A}, B)$ et $(\overline{A}, \overline{B})$ sont encore des couples d'événements indépendants.

Généralisons à un ensemble fini quelconque d'événements.

Définition 12. Soit $(\Omega, \mathfrak{T}, \mathbb{P})$ un espace probabilisé. Une famille finie $(A_i)_{i \in I}$ d'événements est *indépendante* si pour tout $J \subset I$,

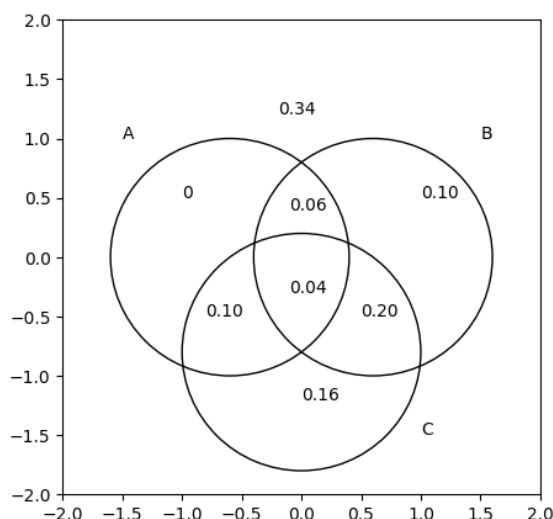
$$\mathbb{P}\left(\bigcap_{i \in J} A_i\right) = \prod_{i \in J} \mathbb{P}(A_i)$$

Remarque 10. Clairement, toute sous-famille d'une famille d'événements indépendants est encore une famille d'événements indépendants.

Pourquoi la définition de l'indépendance de plus de trois événements est-elle si compliquée ? Prenons deux exemples.

Exemple. Des événements peuvent être indépendants deux à deux sans que la famille le soit. Prenons par exemple $\Omega = \{1, 2, 3, 4\}$ muni de la tribu $\mathcal{P}(\Omega)$ et de la probabilité uniforme. Soient $A = \{1, 2\}$, $B = \{1, 3\}$ et $C = \{1, 4\}$. On a $\mathbb{P}(A \cap B) = \frac{1}{4} = \mathbb{P}(A)\mathbb{P}(B)$, et de même pour $A \cap C$ et $B \cap C$. Mais $\mathbb{P}(A \cap B \cap C) = \frac{1}{4}$ alors que $\mathbb{P}(A)\mathbb{P}(B)\mathbb{P}(C) = \frac{1}{8}$.

Exemple. Des événements peuvent être « globalement indépendants » (notion inexistante!) sans être indépendants. Soit $(\Omega, \mathfrak{T}, \mathbb{P})$ un espace probabilisé. Trois événements A, B, C définissent un système complet de 8 événements donné dans la figure ci-dessous. On suppose l'espace probabilisé assez « riche » pour que ces 8 événements aient les probabilités indiquées dans la figure.



On a

$$\mathbb{P}(A) = 0.2, \mathbb{P}(B) = 0.4, \mathbb{P}(C) = 0.5$$

De là,

$$\mathbb{P}(A)\mathbb{P}(B)\mathbb{P}(C) = 0.2 \times 0.4 \times 0.5 = 0.04 = \mathbb{P}(A \cap B \cap C)$$

Cependant,

$$\mathbb{P}(A \cap B) = 0.1 \neq \mathbb{P}(A)\mathbb{P}(B) = 0.08$$

Les événements A, B, C ne sont donc pas indépendants.

En fait, on a aussi

$$\mathbb{P}(A \cap C) \neq \mathbb{P}(A)\mathbb{P}(C)$$

$$\mathbb{P}(B \cap C) \neq \mathbb{P}(B)\mathbb{P}(C)$$

Les événements A, B, C sont donc indépendants « globalement » et pourtant ils ne sont pas indépendants deux à deux.

Exercice. Soient A, B et C trois événements indépendants d'un espace probabilisé $(\Omega, \mathfrak{T}, \mathbb{P})$. Montrer que les couples $(A \cap B, C)$ et $(A \cup B, C)$ sont encore des couples d'événements indépendants.

4.5 Utilité de l'indépendance et des probabilités conditionnelles

À quoi servent l'indépendance et les probabilités conditionnelles? Regardons un premier exemple qui nous montre comment *construire* un espace probabilisé lorsqu'on fait des hypothèses d'indépendance.

Exemple. Une pièce a une « probabilité » $p \in [0, 1]$ de faire pile. Soit $n \geq 1$. On lance indépendamment n fois la pièce. Pour tout $i \in \llbracket 0, n \rrbracket$, quelle est la probabilité de faire exactement i piles?

On prend $\Omega = \llbracket 0, 1 \rrbracket^n$, muni de la tribu $\mathcal{P}(\Omega)$ (où 1 symbolise pile et 0 symbolise face). Que choisit-on comme probabilité? L'idée est de choisir la probabilité « produit », définie par sa fonction de masse. Posons $\varphi(1) = p$ et $\varphi(0) = 1 - p$. Pour tout $(x_1, \dots, x_n) \in \Omega$, posons

$$\mu_{\mathbb{P}}((x_1, \dots, x_n)) = \mathbb{P}(\{(x_1, \dots, x_n)\}) = \prod_{k=1}^n \varphi(x_k)$$

Soit A l'événement qui nous intéresse. Pour $i \in \llbracket 0, n \rrbracket$, soit A_i l'événement « exactement i piles ». Pour tout $i \in \llbracket 1, n \rrbracket$,

$$A_i = \bigcup_{x \in A_i} \{x\}$$

et la réunion est disjointe. Or, pour tout $x \in A_i$,

$$\mathbb{P}(\{x\}) = p^i (1 - p)^{n-i}$$

On a donc

$$\mathbb{P}(A_i) = p^i (1 - p)^{n-i} |A_i| = \binom{n}{i} p^i (1 - p)^{n-i}$$

La loi binomiale n'est évidemment pas bien loin ...

Regardons maintenant comment les probabilités conditionnelles nous permettent de calculer des probabilités *sans construire d'espace probabilisé*. Un grand nombre de problèmes de probabilités fournissent des données pour lesquelles on connaît des probabilités conditionnelles. Dans ce cas, les problèmes en question supposent l'existence d'un espace probabilisé $(\Omega, \mathcal{F}, \mathbb{P})$ et d'événements A, B , etc. pour lesquels certaines probabilités conditionnelles ont les valeurs données dans l'énoncé. En appliquant les théorèmes du cours (probabilités composées, probabilités totales), on en déduit d'autres probabilités conditionnelles, voire des probabilités « absolues ».

Exemple. Soit $p \in]0, 1[$. On dispose d'une suite $(O_n)_{n \in \mathbb{N}}$ d'ordinateurs. L'ordinateur O_0 envoie un bit à l'ordinateur O_1 . Celui-ci le réexpédie à O_2 , etc. Le souci est que le canal de transmission n'est pas fiable. Pour tout $n \in \mathbb{N}$, le bit expédié par O_n à O_{n+1} a une probabilité p d'être inversé (0 devient 1 ou 1 devient 0).

Le problème est le suivant. Pour tout $n \in \mathbb{N}$, quelle est la probabilité que le bit reçu par O_n soit le bit envoyé par O_0 ?

On *admet* l'existence d'un espace probabilisé $(\Omega, \mathfrak{F}, \mathbb{P})$ dans lequel ce problème peut être modélisé. Pour tout $n \in \mathbb{N}$, soit A_n l'événement « le bit reçu par O_n est égal au bit envoyé par O_0 » (on convient de prendre $A_0 = \Omega$). Que nous dit l'énoncé ? Il nous dit que pour tout $n \geq 1$,

$$\mathbb{P}(A_n | A_{n-1}) = p$$

Notons, pour tout $n \in \mathbb{N}$, $a_n = \mathbb{P}(A_n)$. On a donc $a_0 = 1$. Soit $n \geq 1$. L'ensemble $\{A_{n-1}, \bar{A}_{n-1}\}$ est un système complet d'événements de probabilités non nulles (?). On a donc, par le théorème des probabilités totales, pour tout $n \geq 1$,

$$\mathbb{P}(A_n) = \mathbb{P}(A_n | A_{n-1})\mathbb{P}(A_{n-1}) + \mathbb{P}(A_n | \bar{A}_{n-1})\mathbb{P}(\bar{A}_{n-1})$$

ou encore

$$a_n = pa_{n-1} + (1-p)(1-a_{n-1}) = (2p-1)a_{n-1} + 1-p$$

La suite $(a_n)_{n \in \mathbb{N}}$ est donc une suite arithmético-géométrique de raison $2p-1$. Finies les probas, il ne reste qu'à résoudre la récurrence.

Soit $x \in \mathbb{R}$. On a

$$x = (2p-1)x + 1-p$$

si et seulement si (car $p \neq 1$)

$$x = \frac{1}{2}$$

De là, pour tout $n \in \mathbb{N}$,

$$a_n = \frac{1}{2} + \left(a_0 - \frac{1}{2}\right)(2p-1)^n = \frac{1}{2} + \frac{1}{2}(2p-1)^n$$

Remarquons que comme $p \in]0, 1[$, $-1 < 2p-1 < 1$. Ainsi,

$$a_n \xrightarrow[n \rightarrow \infty]{} \frac{1}{2}$$

Bref, une chance sur 2 ...

Chapitre 31

Variables aléatoires

1 Notion de variable aléatoire

1.1 C'est quoi ?

Définition 1. Soit $(\Omega, \mathfrak{T}, \mathbb{P})$ un espace probabilisé. Soit E un ensemble. Une *variable aléatoire discrète* définie sur Ω à valeurs dans E est une application $X : \Omega \longrightarrow E$ vérifiant les propriétés suivantes.

- L'ensemble $X(\Omega)$ est fini ou dénombrable.
- Pour tout $x \in E$, $X^{-1}(\{x\}) \in \mathfrak{T}$.

Il existe également une notion de variable aléatoire non discrète, pour laquelle $X(\Omega)$ est un ensemble infini non dénombrable. La théorie des variables aléatoires non discrètes présente des différences avec celle des variables aléatoires discrètes. Dans toute la suite, quand nous parlerons de « variable aléatoire », il sera sous-entendu qu'il s'agit d'une variable aléatoire discrète.

Lorsque $E = \mathbb{R}$ on parle de variable aléatoire réelle.

Les conditions de la définition sont assez techniques. Dans la pratique, les univers Ω que nous considérons à notre niveau sont finis ou dénombrables et on choisit comme tribu $\mathcal{P}(\Omega)$. Les deux conditions sont alors automatiquement vérifiées.

Remarque 1. Une « variable » « aléatoire » n'est pas une variable mais une fonction. Et cette fonction n'a rien d'aléatoire.

Proposition 1. *Proposition* Soit $(\Omega, \mathfrak{T}, \mathbb{P})$ un espace probabilisé. Soit E un ensemble. Soit $X : \Omega \longrightarrow E$ une variable aléatoire. Pour tout $A \subset E$, $X^{-1}(A) \in \mathfrak{T}$.

Démonstration. Soit $A \subset E$. Posons $A = A' \cup A''$ où $A' = A \cap X(\Omega)$ et $A'' = A \setminus A'$. On a

$$X^{-1}(A) = X^{-1}(A') \cup X^{-1}(A'') = X^{-1}(A')$$

Comme $X(\Omega)$ est fini ou dénombrable, A' l'est aussi. On a

$$X^{-1}(A') = X^{-1} \left(\bigcup_{x \in A'} \{x\} \right) = \bigcup_{x \in A'} X^{-1}(\{x\})$$

L'ensemble $X^{-1}(A')$, union dénombrable d'éléments de $\mathfrak{T}$, est donc aussi un élément de $\mathfrak{T}$. $\square$

Notation. Le monde des variables aléatoires est plein de notations spécifiques. Elles permettent d'écrire les choses de façon plus concise.

- Soit $A \subset E$. $\{X \in A\} = X^{-1}(A)$.
- Si $x \in E$, $\{X = x\} = \{X \in \{x\}\}$.
- Si X est une variable aléatoire réelle et $x \in \mathbb{R}$, $\{X \leq x\} = \{X \in]-\infty, x]\}$. On utilise évidemment des notations analogues pour d'autres types d'intervalles.

Avec nos nouvelles notations, le second point de la définition d'une variable aléatoire devient

$$\text{Pour tout } x \in E, \{X = x\} \in \mathfrak{T}.$$

La conclusion du théorème ci-dessus s'écrit

$$\text{Pour tout } A \subset E, \{X \in A\} \in \mathfrak{T}.$$

1.2 Loi d'une variable aléatoire

Définition 2. Soit $(\Omega, \mathfrak{T}, \mathbb{P})$ un espace probablisé. Soit $X : \Omega \longrightarrow E$ une variable aléatoire. La *loi* de X est la fonction $\mathbb{P}_X : \mathcal{P}(E) \longrightarrow [0, 1]$ définie pour tout $A \subset E$ par

$$\mathbb{P}_X(A) = \mathbb{P}(\{X \in A\})$$

Nous noterons plus simplement $\mathbb{P}(\{X \in A\})$ par $\mathbb{P}(X \in A)$.

Proposition 2. Avec les notations ci-dessus, $\mathbb{P}_X$ est une probabilité sur les espaces probablisables $(E, \mathcal{P}(E))$ et $(X(\Omega), \mathcal{P}(X(\Omega)))$.

Démonstration. Faisons la preuve pour $(E, \mathcal{P}(E))$. L'autre preuve est identique.

- Tout d'abord, $\{X \in E\} = \Omega$ et donc

$$\mathbb{P}_X(E) = \mathbb{P}(\Omega) = 1$$

- Soit $(A_n)_{n \in \mathbb{N}}$ une famille de parties de E disjointes deux à deux. On a (exercice)

$$\{X \in \bigcup_{n \in \mathbb{N}} A_n\} = \bigcup_{n \in \mathbb{N}} \{X \in A_n\}$$

et cette union est disjointe. De là,

$$\begin{aligned} \mathbb{P}_X \left(\bigcup_{n \in \mathbb{N}} A_n \right) &= \mathbb{P} \left(\bigcup_{n \in \mathbb{N}} \{X \in A_n\} \right) \\ &= \sum_{n=0}^{\infty} \mathbb{P}(X \in A_n) \\ &= \sum_{n=0}^{\infty} \mathbb{P}_X(A_n) \end{aligned}$$

□

Remarque 2. Le triplet $(E, \mathcal{P}(E), \mathbb{P}_X)$ est ainsi un espace probabilisé.

Notation. Les notations introduites précédemment se répercutent bien entendu sur la notation des probabilités. Comme nous l'avons déjà dit, on note $\mathbb{P}(X \in A)$ au lieu de $\mathbb{P}_X(A)$. Mais aussi, $\mathbb{P}(\{X = x\})$ est noté $\mathbb{P}(X = x)$. Si X est une variable aléatoire réelle, $\mathbb{P}(X \leq x)$ veut dire $\mathbb{P}(\{X \leq x\})$, etc.

Remarque 3. Soit $(\Omega, \mathfrak{T}, \mathbb{P})$ un espace probabilisé. Soit $X : \Omega \longrightarrow E$ une variable aléatoire discrète. Pour tout $A \subset E$, écrivons

$$A = (X(\Omega) \setminus A) \cup \bigcup_{x \in X(\Omega) \cap A} \{x\}$$

La réunion ci-dessus est une union finie ou dénombrable disjointe, donc

$$\begin{aligned} \mathbb{P}_X(A) &= \mathbb{P}_X(X(\Omega) \setminus A) + \sum_{x \in X(\Omega) \cap A} \mathbb{P}_X(\{x\}) \\ &= 0 + \sum_{x \in X(\Omega) \cap A} \mathbb{P}(X = x) \end{aligned}$$

Ainsi, la loi de X est entièrement déterminée par les réels $\mathbb{P}(X = x)$ pour $x \in X(\Omega)$.

Exemple. Prenons l'exemple du lancer de deux dés parfaits et de la variable aléatoire « somme des dés ». On choisit l'univers $\Omega = \llbracket 1, 6 \rrbracket^2$, muni de la tribu $\mathcal{P}(\Omega)$ et de la probabilité uniforme $\mathbb{P}$. La variable aléatoire qui nous intéresse est $X : \Omega \longrightarrow \mathbb{R}$ définie par $X(x, y) = x + y$. On a $X(\Omega) = \llbracket 2, 12 \rrbracket$. La loi de X est caractérisée par le tableau ci-dessous :

x	2	3	4	5	6	7	8	9	10	11	12
$\mathbb{P}(X = x)$	$\frac{1}{36}$	$\frac{2}{36}$	$\frac{3}{36}$	$\frac{4}{36}$	$\frac{5}{36}$	$\frac{6}{36}$	$\frac{5}{36}$	$\frac{4}{36}$	$\frac{3}{36}$	$\frac{2}{36}$	$\frac{1}{36}$

1.3 Image d'une variable aléatoire par une fonction

Soit $(\Omega, \mathfrak{T}, \mathbb{P})$ un espace probabilisé. Soit $X : \Omega \longrightarrow E$ une variable aléatoire. Soit $f : E \longrightarrow F$ une application. La fonction $f \circ X : \Omega \longrightarrow F$ est elle-aussi une variable aléatoire, que l'on note $f(X)$.

Pour tout $A \subset F$,

$$\mathbb{P}(f(X) \in A) = \mathbb{P}(f \circ X \in A) = \mathbb{P}(X \in f^{-1}(A))$$

Ainsi, déterminer la loi de $f(X)$ nécessite de calculer des images réciproques par f . Si la fonction f est « compliquée », ce genre de calcul peut être délicat.

Exemple. Soit $X : \Omega \rightarrow \llbracket -2, 2 \rrbracket$. On suppose que $\mathbb{P}_X$ est la probabilité uniforme sur $\llbracket -2, 2 \rrbracket$. Quelle est la loi de X^2 ?

Posons $Y = X^2$. Tout d'abord, $Y(\Omega) = \{0, 1, 4\}$. Ensuite,

- $\mathbb{P}(Y = 0) = \mathbb{P}(X = 0) = \frac{1}{5}$.
- $\mathbb{P}(Y = 1) = \mathbb{P}(\{X = 1\} \cup \{X = -1\}) = \mathbb{P}(X = 1) + \mathbb{P}(X = -1) = \frac{2}{5}$.
- $\mathbb{P}(Y = 4) = \mathbb{P}(\{X = 2\} \cup \{X = -2\}) = \mathbb{P}(X = 2) + \mathbb{P}(X = -2) = \frac{2}{5}$.

2 Lois usuelles

2.1 Loi uniforme

Définition 3. Soit $(\Omega, \mathfrak{T}, \mathbb{P})$ un espace probabilisé. Soit $X : \Omega \longrightarrow E$ une variable aléatoire, où E est un ensemble fini de cardinal $n \in \mathbb{N}^*$. X suit la *loi uniforme sur E* lorsque pour tout $x \in E$,

$$\mathbb{P}(X = x) = \frac{1}{n}$$

On note $X \sim \mathcal{U}(E)$.

Remarque 4. Bien entendu,

$$\forall A \subset E, \mathbb{P}(X \in A) = \frac{|A|}{n}$$

Exemple. Lors du lancer d'un dé parfait, la variable aléatoire X égale à la valeur marquée par le dé suit une loi uniforme.

2.2 Loi de Bernoulli

Définition 4. Soit $(\Omega, \mathfrak{T}, \mathbb{P})$ un espace probabilisé. Soit $X : \Omega \longrightarrow \mathbb{R}$ une variable aléatoire réelle. X est une *variable de Bernoulli* si $X(\Omega) \subset \{0, 1\}$.

Le *paramètre* de X est le réel $p = \mathbb{P}(X = 1)$.

On note $X \sim \mathcal{B}(p)$.

Remarque 5. On a $\{X = 0\} \cup \{X = 1\} = \Omega$ et cette union est disjointe, donc

$$\mathbb{P}(X = 0) + \mathbb{P}(X = 1) = \mathbb{P}(\Omega) = 1$$

Ainsi,

$$\mathbb{P}(X = 0) = 1 - p$$

Remarque 6. La loi de Bernoulli intervient dans tout problème où le résultat peut prendre deux valeurs, comme le jeu de pile ou face, ou encore le succès ou l'échec d'une expérience,

la réponse oui ou non à une question posée, garçon ou fille, boule rouge ou verte, gaucher ou droitier, etc.

Le réel p est la probabilité de *succès* de l'expérience. Le réel $q = 1 - p$ est la probabilité d'*échec*.

Exemple. Une urne contient r boules rouges et v boules vertes. Soit $p = \frac{r}{N}$ où $N = r + v$ est le nombre total de boules présentes dans l'urne. On tire une boule au hasard. La variable aléatoire X qui prend la valeur 1 lorsque la boule est rouge et 0 lorsque la boule est verte suit une loi de Bernoulli de paramètre p .

2.3 Loi binomiale

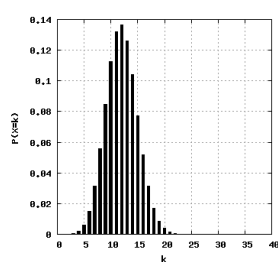
Définition 5. Soit $(\Omega, \mathfrak{T}, \mathbb{P})$ un espace probabilisé. Soit $X : \Omega \rightarrow \mathbb{R}$ une variable aléatoire. Soit $p \in [0, 1]$. X suit une loi binomiale de paramètres n et p si $X(\Omega) = \llbracket 0, n \rrbracket$ et pour tout $k \in \llbracket 0, n \rrbracket$,

$$\mathbb{P}(X = k) = \binom{n}{k} p^k (1 - p)^{n-k}$$

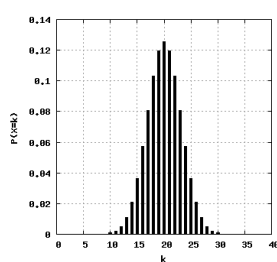
On note $X \sim \mathcal{B}(n, p)$.

Remarque 7. Il convient de vérifier qu'on a bien là une loi de probabilité. Le « poids total » est-il égal à 1 ? Oui, car

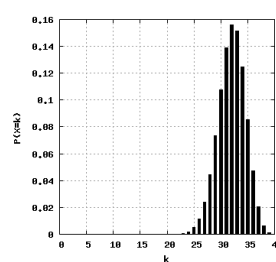
$$\sum_{k=0}^n \binom{n}{k} p^k (1 - p)^{n-k} = (p + (1 - p))^n = 1$$



(a) Loi $\mathcal{B}(40, 0.3)$



(b) Loi $\mathcal{B}(40, 0.5)$



(c) Loi $\mathcal{B}(40, 0.8)$

FIGURE 31.1 – La loi binomiale

Remarque 8. Remarquons que $X \sim \mathcal{B}(1, p) \iff X \sim \mathcal{B}(p)$. Nous verrons également plus loin que si une variable aléatoire réelle X est la somme de n variables aléatoires indépendantes suivant une loi de Bernoulli $\mathcal{B}(p)$, alors $X \sim \mathcal{B}(n, p)$. Mais pour cela il nous faudra d'abord définir l'indépendance des variables aléatoires ...

3 Couples de variables aléatoires

3.1 Loi conjointe, lois marginales

Soit $(\Omega, \mathfrak{T}, \mathbb{P})$ un espace probabilisé. Soient $X : \Omega \longrightarrow E$ et $Y : \Omega \longrightarrow F$ deux variables aléatoires. On dispose alors de la variable aléatoire $(X, Y) : \Omega \longrightarrow E \times F$ définie par

$$(X, Y)(\omega) = (X(\omega), Y(\omega))$$

Quelques remarques et quelques notations ...

Soient $a \in E$ et $b \in F$. On a

$$\{X = a\} \cap \{Y = b\} = \{(X, Y) = (a, b)\}$$

Nous noterons aussi cet événement $\{X = a, Y = b\}$. La virgule signifie « et » et nous rappelle qu'il s'agit en réalité de l'intersection de deux événements. Plus généralement, si $A \subset E$ et $B \subset F$,

$$\{X \in A, Y \in B\} = \{X \in A\} \cap \{Y \in B\} = \{(X, Y) \in A \times B\}$$

Ces notations s'étendent aux probabilités. Par exemple,

$$\mathbb{P}(X = a, Y = b) = \mathbb{P}(\{X = a, Y = b\})$$

Définition 6. Soient X et Y deux variables aléatoires définies sur un même espace probabilisé.

- La *loi conjointe* de X et Y est la loi du couple (X, Y) .
- Les *lois marginales* du couple (X, Y) sont les lois de X et de Y .

3.2 Un exemple

Soient X et Y deux variables aléatoires à valeurs dans $\{0, 1, 2\}$. La loi du couple (X, Y) est donnée dans le tableau ci-dessous. Par exemple,

$$\mathbb{P}(X = 2, Y = 1) = 0.2$$

On a également indiqué dans ce tableau la somme des lignes et la la somme des colonnes. Bien entendu, la somme de toutes les probabilités est 1.

$\mathbb{P}$	$\{X = 0\}$	$\{X = 1\}$	$\{X = 2\}$	
$\{Y = 0\}$	0.1	0.3	0	0.4
$\{Y = 1\}$	0	0.2	0.2	0.4
$\{Y = 2\}$	0.1	0	0.1	0.2
	0.2	0.5	0.3	

La dernière colonne du tableau est la loi de Y , et la dernière ligne est la loi de X . De façon générale, si l'on connaît la loi conjointe de X et Y on peut retrouver les lois de X et Y . En effet, soit $x \in X(\Omega)$. On a

$$\{X = x\} = \bigcup_{y \in Y(\Omega)} \{X = x, Y = y\}$$

et les événements de cette réunion sont disjoints deux à deux. Donc

$$\mathbb{P}(X = x) = \sum_{y \in Y(\Omega)} \mathbb{P}(X = x, Y = y)$$

La loi $\mathbb{P}_X$ est donc déterminée de façon unique à partir de la loi $\mathbb{P}_{(X,Y)}$. De même, pour tout $y \in Y(\Omega)$,

$$\mathbb{P}(Y = y) = \sum_{x \in X(\Omega)} \mathbb{P}(X = x, Y = y)$$

Remarque 9. Les lois $\mathbb{P}_X$ et $\mathbb{P}_Y$ sont appelées les *lois marginales* parce qu'elles sont dans les marges du tableau.

Remarque 10. Connaissant les lois de X et de Y , il n'est pas possible, sauf informations supplémentaires (par exemple l'indépendance, voir un peu plus loin), de trouver la loi conjointe de X et Y .

3.3 Un autre exemple

On munit $\mathfrak{S}_5$ de la tribu $\mathcal{P}(\mathfrak{S}_5)$ et de la probabilité uniforme $\mathbb{P}$. Soient $X : \mathfrak{S}_5 \rightarrow \mathbb{R}$ et $Y : \mathfrak{S}_5 \rightarrow \mathbb{R}$ définies par $X(\sigma) =$ l'ordre de σ et $Y(\sigma) =$ la signature de σ . Déterminons la loi conjointe de X et Y .

Voici, regroupés dans un tableau, les ordres et les signatures des diverses permutations de $\mathfrak{S}_5$, regroupées par « types ». On indique également sur chaque ligne combien il y a de permutations de chaque type, classées par leur décomposition en produit de cycles de supports disjoints.

σ	$X(\sigma)$	$Y(\sigma)$	nb
id	1	1	1
$(i\ j)$	2	-1	10
$(i\ j\ k)$	3	1	20
$(i\ j\ k\ l)$	4	-1	30
$(i\ j\ k\ l\ m)$	5	1	24
$(i\ j)(k\ l)$	2	1	15
$(i\ j)(k\ l\ m)$	6	-1	20

On en déduit la loi du couple (X, Y) , résumée dans un tableau, ainsi, évidemment, que les

lois marginales.

$\mathbb{P}$	$\{X = 1\}$	$\{X = 2\}$	$\{X = 3\}$	$\{X = 4\}$	$\{X = 5\}$	$\{X = 6\}$	
$\{Y = -1\}$	0	$\frac{10}{120}$	0	$\frac{30}{120}$	0	$\frac{20}{120}$	$\frac{1}{2}$
$\{Y = 1\}$	$\frac{1}{120}$	$\frac{15}{120}$	$\frac{20}{120}$	0	$\frac{24}{120}$	0	$\frac{1}{2}$
	$\frac{1}{120}$	$\frac{25}{120}$	$\frac{20}{120}$	$\frac{30}{120}$	$\frac{24}{120}$	$\frac{20}{120}$	

3.4 Loi conditionnelle

Définition 7. Soit $(\Omega, \mathfrak{F}, \mathbb{P})$ un espace probabilisé. Soient $X : \Omega \longrightarrow E$ et $Y : \Omega \longrightarrow F$ deux variables aléatoires. Soit $x \in E$ tel que $\mathbb{P}(X = x) > 0$. La loi conditionnelle de Y sachant que $X = x$ est la fonction

$$\mathbb{P}_{Y|X=x} : \mathcal{P}(F) \longrightarrow [0, 1]$$

définie pour tout $B \in \mathcal{P}(F)$ par

$$\mathbb{P}_{Y|X=x}(B) = \mathbb{P}(Y \in B | X = x)$$

Exemple. Reprenons l'exemple vu plus haut. Rappelons la loi conjointe de X et Y .

$\mathbb{P}$	$\{X = 0\}$	$\{X = 1\}$	$\{X = 2\}$	
$\{Y = 0\}$	0.1	0.3	0	0.4
$\{Y = 1\}$	0	0.2	0.2	0.4
$\{Y = 2\}$	0.1	0	0.1	0.2
	0.2	0.5	0.3	

Pour $x = 0, 1, 2$, $\mathbb{P}(X = x) \neq 0$. On peut donc considérer $\mathbb{P}_{Y|X=x}$. Il suffit de diviser chacune des colonnes du tableau ci-dessus par $\mathbb{P}(X = x)$.

$\mathbb{P}_{Y X=x}$	$\{X = 0\}$	$\{X = 1\}$	$\{X = 2\}$
$\{Y = 0\}$	0.5	0.6	0
$\{Y = 1\}$	0	0.4	0.666...
$\{Y = 2\}$	0.5	0	0.333...

Remarque 11. On vérifie facilement que $\mathbb{P}_{Y|X=x}$ est une probabilité sur $(F, \mathcal{P}(F))$. En conséquences, les théorèmes vérifiés par les lois des variables aléatoires restent valables pour les lois conditionnelles. Par exemple, nous définirons bientôt *l'espérance* d'une variable aléatoire réelle. Si Y est une variable aléatoire réelle on peut alors sans difficulté définir l'espérance conditionnelle de Y sachant $X = x$. L'espérance conditionnelle vérifie toutes les propriétés de l'espérance « classique » comme par exemple, nous le verrons, la positivité ou la linéarité.

Proposition 3. Soient $(\Omega, \mathfrak{F}, \mathbb{P})$ un espace probabilisé. Soient $X : \Omega \longrightarrow E$ et $Y : \Omega \longrightarrow F$ deux variables aléatoires. La loi conjointe de X et Y est caractérisée par la loi de X et, pour tout $x \in E$ tel que $\mathbb{P}(X = x) > 0$, de la loi conditionnelle de Y sachant $X = x$.

Démonstration. Supposons connue la loi conjointe de X et Y . On en déduit, comme on l'a déjà vu, la loi de X . De plus, pour tout $x \in E$ tel que $\mathbb{P}(X = x) > 0$ et tout $B \subset F$,

$$\mathbb{P}_{Y|X=x}(B) = \mathbb{P}(Y \in B | X = x) = \frac{\mathbb{P}(Y \in B, X = x)}{\mathbb{P}(X = x)}$$

Inversement, supposons connus $\mathbb{P}_X$ et $\mathbb{P}_{Y|X=x}$. Soient $x \in E$ et $y \in F$.

Si $\mathbb{P}(X = x) > 0$, alors

$$\mathbb{P}(X = x, Y = y) = \mathbb{P}(Y = y | X = x) \mathbb{P}(X = x)$$

Si $\mathbb{P}(X = x) = 0$ alors, comme $\{X = x, Y = y\} \subset \{X = x\}$,

$$\mathbb{P}(X = x, Y = y) = 0$$

Nous avons donc déterminé la loi conjointe de X et Y à partir de $\mathbb{P}_X$ et de $\mathbb{P}_{Y|X=x}$ pour tout $x \in E$ tel que $\mathbb{P}(X = x) > 0$. $\square$

3.5 Généralisation à n variables aléatoires

Ce que nous avons fait pour des couples de variables aléatoires peut être généralisé à des triplet, des quadruplets, etc. de variables aléatoires. Par exemple, pour un triplet (X, Y, Z) de variables aléatoires, nous avons une loi conjointe qui est la loi de (X, Y, Z) . À partir de cette loi, on peut déterminer la loi de X , la loi de Y , la loi de Z , mais aussi la loi du couple (X, Y) , du couple (Y, Z) et du couple (X, Z) . Par exemple, pour tout $x \in X(\Omega)$,

$$\mathbb{P}(X = x) = \sum_{y \in Y(\Omega)} \sum_{z \in Z(\Omega)} \mathbb{P}(X = x, Y = y, Z = z)$$

ou encore, pour tout $x \in X(\Omega)$ et tout $z \in Z(\Omega)$,

$$\mathbb{P}(X = x, Z = z) = \sum_{y \in Y(\Omega)} \mathbb{P}(X = x, Y = y, Z = z)$$

etc. De même, on peut définir un certain nombre de lois conditionnelles, comme par exemple $\mathbb{P}(X = x, Y = y | Z = z)$ ou encore $\mathbb{P}(X = x | Y = y, Z = z)$, etc.

4 Indépendance

4.1 Variables aléatoires indépendantes

Définition 8. Soit $(\Omega, \mathfrak{T}, \mathbb{P})$ un espace probabilisé. Soient $X : \Omega \longrightarrow E$ et $Y : \Omega \longrightarrow F$ deux variables aléatoires. X et Y sont *indépendantes* si pour tout $A \subset X(\Omega)$ et tout $B \subset Y(\Omega)$,

$$\mathbb{P}(X \in A, Y \in B) = \mathbb{P}(X \in A)\mathbb{P}(Y \in B)$$

Dit autrement, les événements $\{X \in A\}$ et $\{Y \in B\}$ sont indépendants.

Proposition 4. Soit $(\Omega, \mathfrak{T}, \mathbb{P})$ un espace probabilisé. Soient $X : \Omega \longrightarrow E$ et $Y : \Omega \longrightarrow F$ deux variables aléatoires. Les variables aléatoires X et Y sont indépendantes si et seulement si pour tout $x \in X(\Omega)$ et tout $y \in Y(\Omega)$ on a

$$\mathbb{P}(X = x, Y = y) = \mathbb{P}(X = x)\mathbb{P}(Y = y)$$

Démonstration. Dans un sens c'est évident. Il suffit de prendre des singletons pour A et B . Inversement, supposons que pour tout $x \in X(\Omega)$ et tout $y \in Y(\Omega)$ on a

$$\mathbb{P}(X = x, Y = y) = \mathbb{P}(X = x)\mathbb{P}(Y = y)$$

Soient $A \subset X(\Omega)$ et $B \subset Y(\Omega)$. On a

$$\{X \in A, Y \in B\} = \bigcup_{x \in A, y \in B} \{X = x, Y = y\}$$

et cette union est disjointe. De là,

$$\begin{aligned} \mathbb{P}(X \in A, Y \in B) &= \sum_{x \in A, y \in B} \mathbb{P}(X = x, Y = y) \\ &= \sum_{x \in A, y \in B} \mathbb{P}(X = x)\mathbb{P}(Y = y) \\ &= \sum_{x \in A} \mathbb{P}(X = x) \sum_{y \in B} \mathbb{P}(Y = y) \\ &= \mathbb{P}(X \in A)\mathbb{P}(Y \in B) \end{aligned}$$

□

Généralisons maintenant la notion d'indépendance à plus de deux variables aléatoires.

Définition 9. Soit $(\Omega, \mathfrak{F}, \mathbb{P})$ un espace probabilisé. Soit $n \in \mathbb{N}^*$. Pour tout $i \in \llbracket 1, n \rrbracket$, soit $X_i : \Omega \longrightarrow E_i$ une variable aléatoire. Les variables aléatoires $X_1, \dots, X_n$ sont *indépendantes* si pour tous ensembles $A_k \subset X_k(\Omega)$ ($k \in \llbracket 1, n \rrbracket$),

$$\mathbb{P}(X_1 \in A_1, \dots, X_n \in A_n) = \mathbb{P}(X_1 \in A_1) \dots \mathbb{P}(X_n \in A_n)$$

Proposition 5. Soit $(\Omega, \mathfrak{F}, \mathbb{P})$ un espace probabilisé. Soit $n \in \mathbb{N}^*$. Pour tout $i \in \llbracket 1, n \rrbracket$, soit $X_i : \Omega \longrightarrow E_i$ une variable aléatoire. Les variables aléatoires $X_1, \dots, X_n$ sont *indépendantes* si et seulement si pour tous $x_1 \in X_1(\Omega), \dots, x_n \in X_n(\Omega)$,

$$\mathbb{P}(X_1 = x_1, \dots, X_n = x_n) = \mathbb{P}(X_1 = x_1) \dots \mathbb{P}(X_n = x_n)$$

Démonstration. La démonstration est analogue à celle vue pour deux variables aléatoires. $\square$

Remarque 12. La définition ci-dessus *ne dit pas* que les événements $\{X_k \in A_k\}$ sont indépendants. Il s'avère qu'ils le sont. C'est l'objet de la proposition suivante.

Proposition 6. Soit $(\Omega, \mathfrak{F}, \mathbb{P})$ un espace probabilisé. Soit $n \in \mathbb{N}^*$. Pour tout $i \in \llbracket 1, n \rrbracket$, soit $X_i : \Omega \longrightarrow E_i$ une variable aléatoire. Les variables aléatoires $X_1, \dots, X_n$ sont *indépendantes* si et seulement si pour tous $A_k \subset X_k(\Omega)$ ($k \in \llbracket 1, n \rrbracket$), les événements $\{X_k \in A_k\}$ sont *indépendants*.

Démonstration. De droite à gauche c'est immédiat puisque l'indépendance entraîne « l'indépendance globale ». Inversement, supposons que $X_1, \dots, X_n$ sont indépendantes. Nous devons montrer que pour *tout* ensemble $I \subset \llbracket 1, n \rrbracket$,

$$\mathbb{P}(\{X_k \in A_k, k \in I\}) = \prod_{k \in I} \mathbb{P}(X_k \in A_k)$$

Posons pour tout $k \notin I$, $A_k = X_k(\Omega)$. On a alors

$$\{X_k \in A_k, k \in I\} = \{X_1 \in A_1, \dots, X_n \in A_n\}$$

De là, par la définition de l'indépendance de n variables aléatoires,

$$\begin{aligned} \mathbb{P}(X_k \in A_k, k \in I) &= \mathbb{P}(X_1 \in A_1, \dots, X_n \in A_n) \\ &= \prod_{k=1}^n \mathbb{P}(X_k \in A_k) \\ &= \prod_{k \in I} \mathbb{P}(X_k \in A_k) \end{aligned}$$

parce que si $k \notin I$,

$$\mathbb{P}(X_k \in A_k) = \mathbb{P}(X_k \in X_k(\Omega)) = \mathbb{P}(\Omega) = 1$$

□

Corollaire 7. *Toute sous-famille d'une famille de variables aléatoires indépendantes est encore une famille de variables aléatoires indépendantes.*

4.2 Le lemme des coalitions

Nous allons voir dans ce paragraphe qu'en « combinant » des variables aléatoires indépendantes on ne peut fabriquer que des variables aléatoires indépendantes. Commençons par deux cas particuliers.

Proposition 8. *Soit $(\Omega, \mathfrak{F}, \mathbb{P})$ un espace probabilisé. Soient $X : \Omega \longrightarrow E$ et $Y : \Omega \longrightarrow F$ deux variables aléatoires indépendantes. Soient $f : E \longrightarrow E'$ et $g : F \longrightarrow F'$. Alors, les variables aléatoires $f(X)$ et $g(Y)$ sont indépendantes.*

Démonstration. Soient $x \in E'$ et $y \in F'$. On a

$$\begin{aligned} \mathbb{P}(f(X) = x, g(Y) = y) &= \mathbb{P}(X \in f^{-1}(\{x\}), Y \in g^{-1}(\{y\})) \\ &= \mathbb{P}(X \in f^{-1}(\{x\}))\mathbb{P}(Y \in g^{-1}(\{y\})) \\ &= \mathbb{P}(f(X) = x)\mathbb{P}(g(Y) = y) \end{aligned}$$

□

Proposition 9. *Soit $(\Omega, \mathfrak{F}, \mathbb{P})$ un espace probabilisé. Soit $n \geq 1$. Pour tout $i \in \llbracket 1, n \rrbracket$, soit $X_i : \Omega \longrightarrow E_i$ une variable aléatoire. On suppose que $X_1, \dots, X_n$ sont indépendantes. Soit $m \in \llbracket 1, n-1 \rrbracket$. Alors, les variables aléatoires $(X_1, \dots, X_m)$ et $(X_{m+1}, \dots, X_n)$ sont indépendantes.*

Démonstration. Pour faire plus simple, faisons la preuve pour $n = 3$ et $m = 1$. Le cas général est identique. Soient X, Y, Z trois variables aléatoires indépendantes définies sur Ω à valeurs dans E, F, G . Montrons que X et (Y, Z) sont aussi indépendantes. Soient $x \in E$, $y \in F$ et $z \in G$. On a

$$\begin{aligned} \mathbb{P}(X = x, (Y, Z) = (y, z)) &= \mathbb{P}(X = x, Y = y, Z = z) \\ &= \mathbb{P}(X = x)\mathbb{P}(Y = y)\mathbb{P}(Z = z) \end{aligned}$$

où la dernière ligne résulte de l'indépendance de Y et Z . Il reste à remarquer que

$$\begin{aligned} \mathbb{P}(Y = y)\mathbb{P}(Z = z) &= \mathbb{P}(Y = y, Z = z) \\ &= \mathbb{P}((Y, Z) = (y, z)) \end{aligned}$$

□

En combinant ces deux cas particuliers, on obtient le lemme des coalitions.

Proposition 10. LEMME DES COALITIONS

Soit $(\Omega, \mathfrak{F}, \mathbb{P})$ un espace probabilisé. Soit $n \geq 1$. Pour tout $i \in \llbracket 1, n \rrbracket$, soit $X_i : \Omega \longrightarrow E_i$ une variable aléatoire. On suppose que $X_1, \dots, X_n$ sont indépendantes. Soit $m \in \llbracket 1, n-1 \rrbracket$. Soient $f : E_1 \times \dots \times E_m \longrightarrow E'$ et $g : E_{m+1} \times \dots \times E_n \longrightarrow F'$ deux fonctions. Alors, les variables aléatoires $f(X_1, \dots, X_m)$ et $g(X_{m+1}, \dots, X_n)$ sont indépendantes.

Démonstration. Tout d'abord, $X = (X_1, \dots, X_m)$ et $Y = (X_{m+1}, \dots, X_n)$ sont indépendantes. Puis, $f(X)$ et $g(Y)$ sont indépendantes. □

Remarque 13. On peut généraliser le lemme des coalitions en découpant le n -uplet $(X_1, \dots, X_n)$ en plus de deux coalitions.

Remarque 14. Voici quelques exemples :

- Si X, Y, Z, T sont des variables aléatoires réelles indépendantes, alors $X + Y$ et $Z - T$ sont indépendantes.
- Si X, Y, Z, T, U sont des variables aléatoires réelles indépendantes, alors $\sqrt{X^2 + Y^2}$ et $e^{U(Z-T)}$ sont indépendantes.
- etc.

4.3 Retour à la loi binomiale

Nous avons maintenant suffisamment de résultats à notre disposition pour montrer un résultat important.

Proposition 11. Soit $(\Omega, \mathfrak{F}, \mathbb{P})$ un espace probabilisé. Soit $n \in \mathbb{N}$. Soit $p \in [0, 1]$. Soient $X_1, \dots, X_n : \Omega \longrightarrow \{0, 1\}$ n variables aléatoires indépendantes suivant la loi de Bernoulli de paramètre p . Alors, la variable aléatoire

$$S = X_1 + \dots + X_n$$

suit la loi binomiale $\mathcal{B}(n, p)$.

Démonstration. On fait une récurrence sur n . Pour $n = 0$, c'est évident : la fonction nulle suit une loi binomiale de paramètres 0 et p . Soit $n \in \mathbb{N}$. Supposons la propriété vraie pour n . Soient $X_1, \dots, X_{n+1} : \Omega \longrightarrow \{0, 1\}$ $n+1$ variables aléatoires indépendantes suivant la loi de Bernoulli de paramètre p . Posons

$$S_n = X_1 + \dots + X_n$$

$$S_{n+1} = X_1 + \dots + X_{n+1}$$

Remarquons que S_n ne prend que des valeurs entre 0 et $n + 1$.

Soit tout d'abord $k \in \llbracket 1, n + 1 \rrbracket$. Écrivons

$$\{S_{n+1} = k\} = \{X_{n+1} = 0, S_n = k\} \cup \{X_{n+1} = 1, S_n = k - 1\}$$

Les deux événements de la réunion sont disjoints, donc

$$\mathbb{P}(S_{n+1} = k) = \mathbb{P}(X_{n+1} = 0, S_n = k) + \mathbb{P}(X_{n+1} = 1, S_n = k - 1)$$

Par le lemme des coalitions, S_n et X_{n+1} sont indépendantes. On a donc

$$\mathbb{P}(S_{n+1} = k) = \mathbb{P}(X_{n+1} = 0)\mathbb{P}(S_n = k) + \mathbb{P}(X_{n+1} = 1)\mathbb{P}(S_n = k - 1)$$

On applique l'hypothèse de récurrence.

$$\begin{aligned} \mathbb{P}(S_{n+1} = k) &= (1 - p) \binom{n}{k} p^k (1 - p)^{n-k} + p \binom{n}{k-1} p^{k-1} (1 - p)^{n-k+1} \\ &= \left[\binom{n}{k} + \binom{n}{k-1} \right] p^k (1 - p)^{n+1-k} \\ &= \binom{n+1}{k} p^k (1 - p)^{n+1-k} \end{aligned}$$

qui est bien la probabilité voulue.

Il reste à regarder ce qui se passe pour $k = 0$. On a

$$\{S_{n+1} = 0\} = \{X_1 = 0, \dots, X_{n+1} = 0\}$$

d'où, par l'indépendance des X_i ,

$$\mathbb{P}(S_{n+1} = 0) = \mathbb{P}(X_1 = 0) \dots \mathbb{P}(X_{n+1} = 0) = (1 - p)^{n+1} = \binom{n+1}{0} p^0 (1 - p)^{n+1}$$

□

Remarque 15. La proposition que nous venons de montrer n'est pas une équivalence. Si n variables aléatoires indépendantes suivent la loi de Bernoulli $\mathcal{B}(p)$, ALORS leur somme suit la loi binomiale $\mathcal{B}(n, p)$.

En revanche, la réciproque est fautive. Il existe des variables aléatoires qui suivent une loi binomiale et qui ne sont pas la somme de variables de Bernoulli. Cela dit, dans la plupart des situations, on veut montrer des propriétés des variables aléatoires qui ne dépendent que de leur loi, et pas vraiment de la variable aléatoire elle-même. Lorsqu'on se trouve dans cette situation, on peut, au lieu d'étudier la variable qui suit une loi binomiale, étudier une autre variable qui, elle, est la somme de variables de Bernoulli indépendantes et de même loi.

Chapitre 32

Espérance, variance et covariance

1 Espérance

Dans ce chapitre, toutes les variables aléatoires que nous considérons

- sont des variables aléatoires réelles.
- prennent un *nombre fini* de valeurs.

Nous ne préciserons pas à chaque fois le second point. Si $X : \Omega \longrightarrow \mathbb{R}$, nous supposons implicitement que $X(\Omega)$ est un ensemble fini.

1.1 Notion d'espérance

Définition 1. Soit X une variable aléatoire réelle définie sur un espace probabilisé $(\Omega, \mathfrak{T}, \mathbb{P})$. L'*espérance* de X est le réel

$$\mathbb{E}(X) = \sum_{x \in X(\Omega)} x \mathbb{P}(X = x)$$

Exemple. On lance deux dés parfaits. Prenons pour univers $\Omega = \llbracket 1, 6 \rrbracket^2$, muni de la tribu $\mathcal{P}(\Omega)$ et de la probabilité uniforme $\mathbb{P}$. Soit $X : \Omega \longrightarrow \llbracket 2, 12 \rrbracket$ la somme des dés. Le tableau ci-dessous donne le cardinal de $\{X = k\}$ pour $k \in \llbracket 2, 12 \rrbracket$.

k	2	3	4	5	6	7	8	9	10	11	12
$ \{X = k\} $	1	2	3	4	5	6	5	4	3	2	1

On a

$$\mathbb{E}(X) = 2 \times \frac{1}{36} + 3 \times \frac{2}{36} + \dots + 11 \times \frac{2}{36} + 12 \times \frac{1}{36} = \frac{252}{36} = 7$$

Clairement, si nous devions lancer 10 dés et nous demander quelle est l'espérance de la somme des dés, le travail à fournir serait peu gratifiant. Nous verrons bientôt un résultat (linéarité de l'espérance) qui nous permettra de donner la réponse à cette question (35, en fait) immédiatement, *sans calculer la loi de X* .

1.2 Espérance des lois usuelles

Variable aléatoire constante. Soit X une variable aléatoire réelle constante, de valeur m . On a évidemment $\mathbb{E}(X) = m$.

Si X est une variable aléatoire réelle, le réel $\mathbb{E}(X)$ peut être considéré comme une variable aléatoire constante. Donc, un peu abusivement,

$$\mathbb{E}(\mathbb{E}(X)) = \mathbb{E}(X)$$

Loi uniforme. soit X une variable aléatoire réelle. Soit $E \subset \mathbb{R}$ l'ensemble des valeurs prises par X . Supposons que $X \sim \mathcal{U}(E)$. On a

$$\mathbb{E}(X) = \sum_{x \in E} x \mathbb{P}(X = x) = \frac{1}{n} \sum_{x \in E} x$$

c'est à dire la moyenne arithmétique des valeurs prises par X .

Loi de Bernoulli. soit $p \in [0, 1]$. soit $X \sim \mathcal{B}(p)$. On a

$$\mathbb{E}(X) = \mathbb{P}(X = 0) \times 0 + \mathbb{P}(X = 1) \times 1 = (1 - p) \times 0 + p \times 1 = p$$

Loi binomiale. Soient $n \in \mathbb{N}$ et $p \in [0, 1]$. soit $X \sim \mathcal{B}(n, p)$. On a

$$\mathbb{E}(X) = \sum_{k=0}^n k \binom{n}{k} p^k (1-p)^{n-k}$$

Si $n = 0$, $\mathbb{E}(X) = 0 = np$. Supposons $n \geq 1$. Nous allons calculer cette somme, mais nous verrons très bientôt que l'espérance de la loi binomiale peut être obtenue beaucoup plus simplement. Commençons par remarquer que pour tout $k \in \llbracket 1, n \rrbracket$,

$$\binom{n}{k} = \frac{n}{k} \binom{n-1}{k-1}$$

On a

$$\begin{aligned} \sum_{k=0}^n k \binom{n}{k} p^k (1-p)^{n-k} &= n \sum_{k=0}^n \frac{k}{n} \binom{n}{k} p^k (1-p)^{n-k} \\ &= n \sum_{k=1}^n \frac{k}{n} \binom{n}{k} p^k (1-p)^{n-k} \\ &= n \sum_{k=1}^n \binom{n-1}{k-1} p^k (1-p)^{n-k} \\ &= np \sum_{k=1}^n \binom{n-1}{k-1} p^{k-1} (1-p)^{n-k} \end{aligned}$$

En faisant le changement d'indice $k' = k - 1$, on obtient

$$\mathbb{E}(X) = np \sum_{k=0}^{n-1} \binom{n-1}{k} p^k (1-p)^{n-1-k} = np(1-p+p)^{n-1} = np$$

1.3 Linéarité de l'espérance

Lemme 1. Soit $(\Omega, \mathfrak{F}, \mathbb{P})$ un espace probabilisé. Soit $n \in \mathbb{N}^*$. Soient $A_1, \dots, A_n$ n événements de Ω deux à deux disjoints. Soient $x_1, \dots, x_n \in \mathbb{R}$. Soit

$$X = \sum_{k=1}^n x_k \mathbb{1}_{A_k}$$

Alors,

$$\mathbb{E}(X) = \sum_{k=1}^n x_k \mathbb{P}(A_k)$$

Démonstration. Comme les A_k sont disjoints, pour tout $\omega \in \Omega$, un au plus des réels $\mathbb{1}_{A_k}(\omega)$ est non nul. L'ensemble des valeurs prises par X est donc

$$X(\Omega) \subset \{x_1, \dots, x_n\}$$

Décomposons la somme $\sum_{k=1}^n x_k \mathbb{P}(A_k)$ en regroupant les x_k égaux :

$$\begin{aligned} \sum_{k=1}^n x_k \mathbb{P}(A_k) &= \sum_{x \in X(\Omega)} \sum_{k \in \llbracket 1, n \rrbracket, x_k = x} x_k \mathbb{P}(A_k) \\ &= \sum_{x \in X(\Omega)} x \sum_{k \in \llbracket 1, n \rrbracket, x_k = x} \mathbb{P}(A_k) \end{aligned}$$

Soit $x \in X(\Omega)$. On a

$$\{X = x\} = \bigcup_{k \in \llbracket 1, n \rrbracket, x_k = x} A_k$$

et donc (union disjointe)

$$\mathbb{P}(X = x) = \mathbb{P}\left(\bigcup_{k \in \llbracket 1, n \rrbracket, x_k = x} A_k\right) = \sum_{k \in \llbracket 1, n \rrbracket, x_k = x} \mathbb{P}(A_k)$$

□

Proposition 2. LINÉARITÉ DE L'ESPÉRANCE

Soient X, Y deux variables aléatoires réelles définies sur un même espace probabilisé $(\Omega, \mathfrak{F}, \mathbb{P})$, et $\alpha, \beta \in \mathbb{R}$. On a

$$\mathbb{E}(\alpha X + \beta Y) = \alpha \mathbb{E}(X) + \beta \mathbb{E}(Y)$$

Démonstration. On commence par écrire

$$\alpha X + \beta Y = \sum_{x \in X(\Omega)} \sum_{y \in Y(\Omega)} (\alpha x + \beta y) \mathbb{1}_{\{X=x, Y=y\}}$$

Pour montrer cette égalité, il suffit de prendre $\omega \in \Omega$ et de constater que les deux fonctions de l'égalité ont même image en ω . Les fonctions indicatrices qui apparaissent ci-dessus sont des fonctions indicatrices d'ensembles disjoints deux à deux. On a donc, par le lemme,

$$\begin{aligned} \mathbb{E}(\alpha X + \beta Y) &= \sum_{x \in X(\Omega)} \sum_{y \in Y(\Omega)} (\alpha x + \beta y) \mathbb{P}(X = x, Y = y) \\ &= \alpha \sum_{x \in X(\Omega)} x \sum_{y \in Y(\Omega)} \mathbb{P}(X = x, Y = y) + \beta \sum_{y \in Y(\Omega)} y \sum_{x \in X(\Omega)} \mathbb{P}(X = x, Y = y) \end{aligned}$$

Remarquons maintenant que

$$\begin{aligned} \sum_{y \in Y(\Omega)} \mathbb{P}(X = x, Y = y) &= \sum_{y \in Y(\Omega)} \mathbb{P}(\{X = x\} \cap \{Y = y\}) \\ &= \mathbb{P}\left(\bigcup_{y \in Y(\Omega)} \{X = x\} \cap \{Y = y\}\right) \\ &= \mathbb{P}\left(\{X = x\} \cap \bigcup_{y \in Y(\Omega)} \{Y = y\}\right) \\ &= \mathbb{P}(\{X = x\} \cap \Omega) \\ &= \mathbb{P}(X = x) \end{aligned}$$

De même,

$$\sum_{x \in X(\Omega)} \mathbb{P}(X = x, Y = y) = \mathbb{P}(Y = y)$$

Ainsi,

$$\mathbb{E}(\alpha X + \beta Y) = \alpha \sum_{x \in X(\Omega)} x \mathbb{P}(X = x) + \beta \sum_{y \in Y(\Omega)} y \mathbb{P}(Y = y) = \alpha \mathbb{E}(X) + \beta \mathbb{E}(Y)$$

□

Exemple. Revenons à l'espérance d'une variable aléatoire X suivant une loi binomiale $\mathcal{B}(n, p)$ où $n \in \mathbb{N}$ et $p \in [0, 1]$. L'espérance de X ne dépend que de la loi de X , et pas de X elle-même. Ainsi, toutes les variables aléatoires suivant la loi $\mathcal{B}(n, p)$ ont la même espérance. On peut donc prendre

$$X = X_1 + \dots + X_n$$

où les X_k sont n variables aléatoires (indépendantes, mais on n'en a pas besoin) suivant la loi de Bernoulli $\mathcal{B}(p)$. Par linéarité de l'espérance,

$$\mathbb{E}(X) = \sum_{k=1}^n \mathbb{E}(X_k) = \sum_{k=1}^n p = np$$

C'est quand même plus facile que le calcul fait plus haut.

Exemple. Soit $n \in \mathbb{N}$. On lance n dés parfaits. On modélise l'expérience par l'univers $\Omega = \llbracket 1, 6 \rrbracket^n$, muni de la tribu $\mathcal{P}(\Omega)$ et de la probabilité uniforme $\mathbb{P}$. Pour tout $i \in \llbracket 1, n \rrbracket$, soit $X_i : \Omega \longrightarrow \llbracket 1, 6 \rrbracket$ définie par

$$X_i(x_1, \dots, x_n) = x_i$$

La variable aléatoire X_i « est la valeur du dé numéro i ». On vérifie facilement que pour tout $i \in \llbracket 1, n \rrbracket$, $X_i \sim \mathcal{U}(\llbracket 1, 6 \rrbracket)$, et donc

$$\mathbb{E}(X_i) = \frac{7}{2}$$

Soit $X = \sum_{i=1}^n X_i$. La variable aléatoire X « est la somme des dés ». Par la linéarité de l'espérance

$$\mathbb{E}(X) = \sum_{i=1}^n \mathbb{E}(X_i) = \frac{7}{2}n$$

Exercice. Montrer que les variables aléatoires $X_1, \dots, X_n$ sont indépendantes.

Exemple. Espérance de la loi hypergéométrique. Soit S le nombre de boules rouges tirées d'une urne lors d'un tirage de n boules dans une urne sans remise. Alors,

$$S = X_1 + \dots + X_n$$

où $X_k = 1$ si la k ième boule tirée est rouge, et $X_k = 0$ sinon. Les X_k sont des variables de Bernoulli de paramètre $p = \mathbb{P}(X_k = 1)$. Nous avons montré au chapitre précédent que $p = \frac{r}{N}$, où N est le nombre de boules et r est le nombre de boules rouges. Par la linéarité de l'espérance,

$$\mathbb{E}(S) = \sum_{k=1}^n \mathbb{E}(X_k) = np$$

Remarque 1. Nous venons de voir des exemples très instructifs. Pour calculer l'espérance d'une variable aléatoire réelle X , on n'a pas toujours besoin de connaître sa loi. Si X s'exprime en fonction d'autres variables aléatoires dont on connaît l'espérance, on peut alors espérer obtenir $\mathbb{E}(X)$. Voyons un autre exemple.

Exemple. On munit le groupe symétrique $\Omega = \mathfrak{S}_n$ de la tribu $\mathcal{P}(\Omega)$ et de la probabilité uniforme $\mathbb{P}$. Soit $X : \mathfrak{S}_n \longrightarrow \mathbb{R}$ définie par $X(\sigma) =$ le nombre de points fixes de σ . Calculons $\mathbb{E}(X)$.

Pour tout $i \in \llbracket 1, n \rrbracket$, soit $X_i : \Omega \longrightarrow \{0, 1\}$ la variable de Bernoulli définie par $X_i(\sigma) = 1$ si $\sigma(i) = i$, et 0 sinon. $\{X_i = 1\}$ est l'ensemble des permutations $\sigma \in \mathfrak{S}_n$ telles que $\sigma(i) = i$ qui est clairement en bijection avec l'ensemble des permutations de l'ensemble $\llbracket 1, n \rrbracket \setminus \{i\}$. On a donc

$$|\{X_i = 1\}| = (n - 1)!$$

De là,

$$\mathbb{P}(X_i = 1) = \frac{|\{X_i = 1\}|}{|\mathfrak{S}_n|} = \frac{(n - 1)!}{n!} = \frac{1}{n}$$

On a donc

$$\mathbb{E}(X_i) = \mathbb{P}(X_i = 1) = \frac{1}{n}$$

Enfin,

$$X = \sum_{k=1}^n X_k$$

donc, par linéarité de l'espérance,

$$\mathbb{E}(X) = \sum_{k=1}^n \mathbb{E}(X_k) = 1$$

Une permutation de $\mathfrak{S}_n$ possède donc « en moyenne » 1 point fixe.

1.4 Positivité, croissance

Proposition 3. POSITIVITÉ DE L'ESPÉRANCE

Soit X une variable aléatoire réelle à valeurs positives, définie sur un espace probabilisé $(\Omega, \mathfrak{T}, P)$. On a

- $\mathbb{E}(X) \geq 0$.
- $\mathbb{E}(X) = 0$ si et seulement si $\mathbb{P}(X = 0) = 1$.

Démonstration. Remarquons tout d'abord que

$$\mathbb{E}(X) = \sum_{x \in X(\Omega)} x \mathbb{P}(X = x) = \sum_{x \in X(\Omega), x \neq 0} x \mathbb{P}(X = x)$$

Cette somme est clairement positive puisque $X(\Omega) \subset \mathbb{R}_+$. De plus elle est nulle si et seulement si

$$\forall x \in X(\Omega) \setminus \{0\}, \mathbb{P}(X = x) = 0$$

ce qui équivaut à

$$\sum_{x \in X(\Omega), x \neq 0} \mathbb{P}(X = x) = 0$$

Cette dernière somme n'est autre que $\mathbb{P}(X \neq 0)$. Finalement,

$$\mathbb{E}(X) = 0 \iff \mathbb{P}(X \neq 0) = 0 \iff \mathbb{P}(X = 0) = 1$$

□

Définition 2. Soit $(\Omega, \mathfrak{T}, \mathbb{P})$ un espace probabilisé.

Un événement A de Ω est *presque sûr* lorsque $\mathbb{P}(A) = 1$.

Si $P(\omega)$ est une propriété des éléments de Ω , la propriété P est *presque sûrement vérifiée* sur Ω lorsque l'événement $\{\omega \in \Omega, P(\omega)\}$ est presque sûr.

Notation. Nous noterons p.s. pour « presque sûrement ».

Par exemple, dans la proposition précédente, $\mathbb{E}(X) = 0$ si et seulement si $X = 0$ p.s.

Proposition 4. CROISSANCE DE L'ESPÉRANCE

Soient X, Y deux variables aléatoires réelles définies sur un même espace probabilisé $(\Omega, \mathfrak{T}, \mathbb{P})$. On suppose que $X \leq Y$. Alors,

- $\mathbb{E}(X) \leq \mathbb{E}(Y)$.
- $\mathbb{E}(X) = \mathbb{E}(Y)$ si et seulement si $X = Y$ p.s.

Démonstration. On a $Y - X \geq 0$ donc $\mathbb{E}(Y - X) \geq 0$. Nous sommes donc ramenés à la proposition précédente. □

1.5 Formule de transfert

Proposition 5. FORMULE DE TRANSFERT

Soit $(\Omega, \mathfrak{T}, \mathbb{P})$ un espace probabilisé. Soit $X : \Omega \rightarrow E$ une variable aléatoire. Soit $f : E \rightarrow \mathbb{R}$. On a

$$\mathbb{E}(f(X)) = \sum_{x \in X(\Omega)} f(x) \mathbb{P}(X = x)$$

Démonstration. Écrivons

$$X = \sum_{x \in X(\Omega)} x \mathbb{1}_{\{X=x\}}$$

Soit

$$Y = \sum_{x \in X(\Omega)} f(x) \mathbb{1}_{\{X=x\}}$$

Soit $\omega \in \Omega$. On a

$$\begin{aligned} Y(\omega) &= \sum_{x \in X(\Omega)} f(x) \mathbb{1}_{\{X=x\}}(\omega) \\ &= \sum_{x \in X(\Omega)} f(x) \delta_{X(\omega), x} \\ &= f(X(\omega)) \\ &= f(X)(\omega) \end{aligned}$$

Ainsi, $Y = f(X)$. On en déduit, par linéarité de l'espérance, que

$$E(f(X)) = \sum_{x \in X(\Omega)} f(x) E(\mathbb{1}_{\{X=x\}}) = \sum_{x \in X(\Omega)} f(x) \mathbb{P}(X = x)$$

□

Remarque 2. Si X est une variable aléatoire réelle et $f = id_{\mathbb{R}}$, on obtient

$$\mathbb{E}(X) = \sum_{x \in X(\Omega)} x \mathbb{P}(X = x)$$

ce qui est heureux.

Exemple. La variable aléatoire réelle X suit une loi uniforme sur $\{-1, 0, 1, 2\}$. On a $\mathbb{E}(X) = \frac{1}{2}$ et, par la formule de transfert,

$$\mathbb{E}(X^2) = \frac{1}{4}(-1)^2 + \frac{1}{4}0^2 + \frac{1}{4}1^2 + \frac{1}{4}2^2 = \frac{3}{2}$$

1.6 Espérance d'un produit de variables aléatoires indépendantes

Proposition 6. ESPÉRANCE D'UN PRODUIT

Soit $(\Omega, \mathfrak{F}, \mathbb{P})$ un espace probabilisé. Soient X et Y deux variables aléatoires réelles définies sur Ω . On suppose que X et Y sont indépendantes. Alors,

$$\mathbb{E}(XY) = \mathbb{E}(X)\mathbb{E}(Y)$$

Démonstration. Écrivons

$$XY = \sum_{x \in X(\Omega)} \sum_{y \in Y(\Omega)} xy \mathbb{1}_{\{X=x, Y=y\}}$$

Par la linéarité de l'espérance,

$$\begin{aligned} \mathbb{E}(XY) &= \sum_{x \in X(\Omega)} \sum_{y \in Y(\Omega)} xy \mathbb{E}(\mathbb{1}_{\{X=x, Y=y\}}) \\ &= \sum_{x \in X(\Omega)} \sum_{y \in Y(\Omega)} xy \mathbb{P}(X = x, Y = y) \end{aligned}$$

Comme X et Y sont indépendantes, $\mathbb{P}(X = x, Y = y) = \mathbb{P}(X = x)\mathbb{P}(Y = y)$. De là,

$$\begin{aligned}\mathbb{E}(XY) &= \sum_{x \in X(\Omega)} \sum_{y \in Y(\Omega)} xy \mathbb{P}(X = x) \mathbb{P}(Y = y) \\ &= \sum_{x \in X(\Omega)} x \mathbb{P}(X = x) \sum_{y \in Y(\Omega)} y \mathbb{P}(Y = y) \\ &= \mathbb{E}(X) \mathbb{E}(Y)\end{aligned}$$

□

Remarque 3. L'indépendance de X et Y est inutile pour la somme de deux variables aléatoires, mais **indispensable** pour leur produit. Prenons par exemple X de loi de Bernoulli $\mathcal{B}(p)$ où $p \in]0, 1[$ et $Y = X$. On a $XY = X^2 = X$, donc $\mathbb{E}(XY) = p$. Mais $\mathbb{E}(X)\mathbb{E}(Y) = \mathbb{E}(X)^2 = p^2$, donc $\mathbb{E}(X)^2 \neq \mathbb{E}(X^2)$.

Proposition 7. Soit $n \geq 1$. Soient $X_1, \dots, X_n$ n variables aléatoires réelles indépendantes. On a

$$\mathbb{E} \left(\prod_{k=1}^n X_k \right) = \prod_{k=1}^n \mathbb{E}(X_k)$$

Démonstration. On fait une récurrence sur n . C'est clair pour $n = 1$. Soit $n \geq 1$. Supposons la propriété vraie pour n . Soient $X_1, \dots, X_{n+1}$ $n+1$ variables aléatoires réelles indépendantes. Par le lemme des coalitions, les variables aléatoires $\prod_{k=1}^n X_k$ et X_{n+1} sont indépendantes, donc

$$\mathbb{E} \left(\prod_{k=1}^{n+1} X_k \right) = \mathbb{E} \left(\prod_{k=1}^n X_k \right) \mathbb{E}(X_{n+1})$$

On termine grâce à l'hypothèse de récurrence. □

1.7 Inégalité de Schwarz

Proposition 8. INÉGALITÉ DE SCHWARZ

Soit $(\Omega, \mathfrak{F}, \mathbb{P})$ un espace probabilisé. Soient $X, Y : \Omega \rightarrow \mathbb{R}$ deux variables aléatoires réelles. On a

$$\mathbb{E}(XY)^2 \leq \mathbb{E}(X^2) \mathbb{E}(Y^2)$$

Il y a égalité si et seulement si il existe $a, b \in \mathbb{R}$ pas tous les deux nuls tels que

$$aX + bY = 0$$

Démonstration. Soit $f : \mathbb{R} \longrightarrow \mathbb{R}$ définie pour tout $t \in \mathbb{R}$ par

$$f(t) = \mathbb{E}((X + tY)^2)$$

Par la linéarité de l'espérance, pour tout $t \in \mathbb{R}$,

$$f(t) = \mathbb{E}(X^2 + 2tXY + t^2Y^2) = \mathbb{E}(Y^2)t^2 + 2\mathbb{E}(XY)t + \mathbb{E}(X^2)$$

Deux cas se présentent.

- Cas 1, $\mathbb{E}(Y^2) = 0$. Dans ce cas, par le théorème de nullité de l'espérance, $Y = 0$ p.s. De là, $XY = 0$ p.S. et donc $\mathbb{E}(XY) = 0$. On a donc l'égalité.
- Cas 2, $\mathbb{E}(Y^2) = 0$. La fonction f est une fonction polynomiale de degré 2. Par la positivité de l'espérance, $f \geq 0$. On en déduit que le discriminant Δ de f est nul :

$$4(\mathbb{E}(XY)^2 - 4\mathbb{E}(X^2)\mathbb{E}(Y^2)) \geq 0$$

d'où l'inégalité voulue. De plus, on a l'égalité si et seulement si $\Delta = 0$, ou encore si et seulement si il existe $t_0 \in \mathbb{R}$ tel que

$$f(t_0) = \mathbb{E}(X + t_0Y)^2 = 0$$

Par le théorème de nullité de l'espérance, ceci équivaut à

$$X + t_0Y = 0 \text{ p.s.}$$

□

2 Variance

2.1 Variance et écart-type

Soit X une variable aléatoire réelle. Quelle est la fonction constante qui « approche le mieux » X ? Une première idée consiste à déterminer le réel m qui minimise $|X - m|$, ou plus exactement, $\mathbb{E}(|X - m|)$. Malheureusement, les mauvaises propriétés de la valeur absolue rendent ce problème, intéressant au demeurant, difficile à résoudre. La bonne approche est de considérer $\mathbb{E}((X - m)^2)$. Soit $\varphi : \mathbb{R} \longrightarrow \mathbb{R}$ définie pour tout $m \in \mathbb{R}$ par

$$\varphi(m) = \mathbb{E}((X - m)^2)$$

Par la linéarité de l'espérance,

$$\varphi(m) = \mathbb{E}(X^2 - 2mX + m^2) = m^2 - 2\mathbb{E}(X)m + \mathbb{E}(X^2)$$

De là,

$$\varphi'(m) = 2(m - \mathbb{E}(X))$$

La fonction φ admet un minimum en $m = \mathbb{E}(X)$ et ce minimum est égal à

$$\varphi(\mathbb{E}(X)) = \mathbb{E}((X - \mathbb{E}(X))^2) = \mathbb{E}(X)^2 - \mathbb{E}(X)^2$$

Le réel $\varphi(\mathbb{E}(X))$ est d'une importance suffisante pour lui donner un nom.

Définition 3. Soit X une variable aléatoire réelle.

- La *variance* de X est le réel

$$\mathbb{V}(X) = \mathbb{E}((X - \mathbb{E}(X))^2) = \mathbb{E}(X)^2 - \mathbb{E}(X)^2$$

- L'*écart-type* de X est le réel

$$\sigma(X) = \sqrt{\mathbb{V}(X)}$$

Remarque 4. L'intérêt de l'écart-type est d'avoir la même dimension que X . Si X est exprimé en mètres, par exemple, alors $\mathbb{V}(X)$ est exprimée en mètres carrés alors que $\sigma(X)$ est exprimé en mètres, comme X . L'écart-type est « la racine carrée de la moyenne du carré de l'écart de X à sa moyenne ». D'une certaine façon, c'est un indicateur de la « dispersion » de la fonction de masse de la loi de X .

Proposition 9. Soit X une variable aléatoire réelle. Soient $a, b \in \mathbb{R}$. On a

$$\mathbb{V}(aX + b) = a^2 \mathbb{V}(X)$$

et

$$\sigma(aX + b) = |a| \sigma(X)$$

Démonstration. Soit $Y = aX + b$. En appelant avec un léger abus b la variable aléatoire constante égale à b ,

$$\mathbb{E}(Y) = a\mathbb{E}(X) + \mathbb{E}(b) = a\mathbb{E}(X) + b$$

Donc

$$\mathbb{E}((Y - \mathbb{E}(Y))^2) = \mathbb{E}(a^2(X - \mathbb{E}(X))^2) = a^2 \mathbb{E}((X - \mathbb{E}(X))^2) = a^2 \mathbb{V}(X)$$

Pour l'écart-type, il suffit de passer à la racine carrée. $\square$

Définition 4. Soit X une variable aléatoire réelle.

- X est *centrée* lorsque $\mathbb{E}(X) = 0$.
- X est *réduite* lorsque $\mathbb{V}(X) = 1$.

Proposition 10. Soit X une variable aléatoire réelle. On suppose $\sigma(X) > 0$. Soit

$$Y = \frac{X - \mathbb{E}(X)}{\sigma(X)}$$

La variable aléatoire réelle Y est centrée et réduite.

Démonstration. Par la linéarité de l'espérance,

$$\mathbb{E}(Y) = \frac{1}{\sigma(X)}(\mathbb{E}(X) - \mathbb{E}(\mathbb{E}(X))) = \frac{1}{\sigma(X)}(\mathbb{E}(X) - \mathbb{E}(X)) = 0$$

et donc Y est centrée. On a aussi

$$\sigma(Y) = \sigma\left(\frac{1}{\sigma(X)}X - \frac{\mathbb{E}(X)}{\sigma(X)}\right) = \frac{1}{\sigma(X)}\sigma(X) = 1$$

et donc Y est réduite. $\square$

Proposition 11. NULLITÉ DE LA VARIANCE

Soit X une variable aléatoire réelle. On a $\mathbb{V}(X) = 0$ si et seulement si X est constante p.s.

Démonstration. Soit $Y = (X - \mathbb{E}(X))^2$. On a $Y \geq 0$. Par le théorème de nullité de l'espérance,

$$\mathbb{V}(X) = 0 \iff Y = 0 \text{ p.s.} \iff X = \mathbb{E}(X) \text{ p.s.}$$

$\square$

2.2 Variance des variables aléatoires usuelles

Quelles sont les variances des variables aléatoires qui suivent nos lois favorites (uniforme, Bernoulli, binomiale) ?

Proposition 12. Soit $n \geq 1$. Soit $a \in \mathbb{R}$. Soit $n \in \mathbb{N}^*$. Soit X une variable aléatoire de loi uniforme sur l'ensemble $\{a, a+1, \dots, a+n-1\}$. On a

$$\mathbb{V}(X) = \frac{n^2 - 1}{12}$$

Démonstration. Soit $Y = X - a$. La variable aléatoire Y suit une loi uniforme sur $\llbracket 0, n-1 \rrbracket$ et $\mathbb{V}(Y) = \mathbb{V}(X)$. Calculons donc plutôt $\mathbb{V}(Y)$.

Tout d'abord,

$$\mathbb{E}(Y) = \frac{1}{n} \sum_{k=0}^{n-1} k = \frac{1}{2}(n-1)$$

et donc

$$\mathbb{E}(Y)^2 = \frac{1}{4}(n^2 - 2n + 1)$$

Ensuite, par la formule de transfert,

$$\begin{aligned} \mathbb{E}(Y^2) &= \frac{1}{n} \sum_{k=0}^{n-1} k^2 \\ &= \frac{1}{6n}(n-1)n(2n-1) \\ &= \frac{1}{6}(2n^2 - 3n + 1) \end{aligned}$$

De là,

$$\begin{aligned} \mathbb{V}(Y) &= \mathbb{E}(Y^2) - \mathbb{E}(Y)^2 \\ &= \frac{1}{6}(2n^2 - 3n + 1) - \frac{1}{4}(n^2 - 2n + 1) \\ &= \frac{n^2 - 1}{12} \end{aligned}$$

□

Proposition 13. Soit $p \in [0, 1]$. Soit X une variable de Bernoulli de paramètre p . On a

$$\mathbb{V}(X) = p(1-p)$$

Démonstration. C'est immédiat. □

Exemple. On lance une pièce équilibrée. Soit X la variable aléatoire qui vaut 1 lorsque la pièce tombe sur pile et 0 lorsque la pièce tombe sur face. On a

$$\mathbb{E}(X) = \frac{1}{2} \text{ et } \sigma(X) = \frac{1}{2}$$

Proposition 14. Soit $p \in [0, 1]$. Soit $n \in \mathbb{N}$. Soit X une variable aléatoire suivant une loi binomiale $\mathcal{B}(n, p)$. On a

$$\mathbb{V}(X) = np(1-p)$$

Démonstration. Le calcul peut être fait de façon « directe », comme nous l'avons fait pour l'espérance. Mais allons droit au but : admettons pour l'instant que si $X_1, \dots, X_n$ sont des variables aléatoires réelles indépendantes, alors

$$\mathbb{V}(X_1 + \dots + X_n) = \mathbb{V}(X_1) + \dots + \mathbb{V}(X_n)$$

Nous montrerons ce résultat à la fin du chapitre. La variance d'une variable aléatoire ne dépendant que de la loi de celle-ci, on peut supposer que X est la somme de n variables aléatoires indépendantes $X_1, \dots, X_n$ suivant une loi de Bernoulli de même paramètre p . On a alors

$$\mathbb{V}(X) = \mathbb{V}\left(\sum_{k=1}^n X_k\right) = \sum_{k=1}^n \mathbb{V}(X_k) = np(1-p)$$

□

2.3 L'inégalité de Bienaymé-Tchebychev

Lemme 15. INÉGALITÉ DE MARKOV

Soit $(\Omega, \mathfrak{F}, \mathbb{P})$ un espace probabilisé. Soit $X : \Omega \longrightarrow \mathbb{R}$ une variable aléatoire réelle. Soit $t > 0$. On suppose $X \geq 0$. Alors,

$$\mathbb{P}(X \geq t) \leq \frac{\mathbb{E}(X)}{t}$$

Démonstration.

$$\begin{aligned} t\mathbb{P}(X \geq t) &= t \sum_{x \in X(\Omega), x \geq t} \mathbb{P}(X = x) \\ &= \sum_{x \in X(\Omega), x \geq t} t\mathbb{P}(X = x) \\ &\leq \sum_{x \in X(\Omega), x \geq t} x\mathbb{P}(X = x) \\ &\leq \sum_{x \in X(\Omega)} x\mathbb{P}(X = x) \\ &= \mathbb{E}(X) \end{aligned}$$

□

Proposition 16. INÉGALITÉ DE BIENAYMÉ-TCHEBYCHEV

Soit X une variable aléatoire réelle. Soit $\lambda > 0$. On a

$$\mathbb{P}(|X - \mathbb{E}(X)| \geq \lambda) \leq \frac{\mathbb{V}(X)}{\lambda^2}$$

Démonstration. Soit $Y = (X - \mathbb{E}(X))^2$. Remarquons que

$$\{Y \geq \lambda^2\} = \{|X - \mathbb{E}(X)| \geq \lambda\}$$

Appliquons l'inégalité de Markov à la variable aléatoire Y , avec $t = \lambda^2$. Il vient

$$\mathbb{P}(Y \geq \lambda^2) \leq \frac{\mathbb{E}(Y)}{\lambda^2}$$

c'est à dire

$$\mathbb{P}(|X - m| \geq \lambda) \leq \frac{\mathbb{V}(X)}{\lambda^2}$$

□

Remarque 5. Supposons $\sigma(X) > 0$. Si on remplace λ par $\lambda\sigma(X)$ dans l'inégalité, on obtient une autre forme de l'inégalité de Bienaymé-Tchebychev :

$$\mathbb{P}(|X - \mathbb{E}(X)| \geq \lambda\sigma(X)) \leq \frac{1}{\lambda^2}$$

Exemple. Soit $n \geq 1$. On lance n fois une pièce parfaite. Soit X le nombre de piles obtenus. Chaque lancer de pièce est une expérience aléatoire de type succès-échec, avec une probabilité de succès $p = \frac{1}{2}$. On a donc $X \sim \mathcal{B}(n, \frac{1}{2})$, d'espérance $n/2$ et de variance $n/4$. L'inégalité de Bienaymé-Tchebychev appliquée à X s'écrit, pour tout $\lambda > 0$,

$$\mathbb{P}\left(\left|X - \frac{n}{2}\right| \geq \lambda\right) \leq \frac{n}{4\lambda^2}$$

prenons $\lambda = \frac{n}{4}$.

$$\mathbb{P}\left(\left|X - \frac{n}{2}\right| \geq \frac{n}{4}\right) \leq \frac{4}{n}$$

Le complémentaire de l'événement $\{|X - \frac{n}{2}| \geq n/4\}$ est l'événement

$$\left\{\left|X - \frac{n}{2}\right| < n/4\right\} = \left\{\frac{n}{4} < X < \frac{3n}{4}\right\}$$

On a donc, par passage au complémentaire,

$$\mathbb{P}\left(\frac{n}{4} < X < \frac{3n}{4}\right) \geq 1 - \frac{4}{n}$$

Pour que la probabilité d'obtenir un nombre de piles entre $n/4$ et $3n/4$ au moins égale à $1/2$, il suffit de choisir n tel que $1 - \frac{4}{n} \geq 1/2$, c'est à dire $n \geq 8$. Si on veut que cette

même probabilité soit supérieure à $9/10$, il suffira de prendre n tel que $1 - \frac{4}{n} \geq 9/10$, c'est à dire $n \geq 40$.

Remarque 6. On connaît la loi de X , donc cette même probabilité s'écrit explicitement sous la forme d'une somme. Une petite simulation numérique montre qu'il suffit de prendre $n = 12$ pour que la probabilité cherchée soit supérieure à $9/10$. L'inégalité de Bienaymé-Tchebychev n'est pas optimale, mais elle a l'immense avantage d'être vraie pour **toutes** les variables aléatoires réelles, quelle que soit leur loi.

2.4 La loi faible des grands nombres

Proposition 17. LOI FAIBLE DES GRANDS NOMBRES

Soit $(\Omega, \mathcal{F}, \mathbb{P})$ un espace probabilisé. Soit $(X_n)_{n \geq 1}$ une suite de variables aléatoires réelles définies sur Ω . On suppose que les X_n sont indépendantes et de même loi. Soit μ l'espérance commune des X_n . Alors,

$$\mathbb{P} \left(\left| \frac{X_1 + \dots + X_n}{n} - \mu \right| < \varepsilon \right) \xrightarrow{n \rightarrow \infty} 1$$

Démonstration. Pour tout $n \geq 1$, soit

$$Y_n = \frac{1}{n} \sum_{k=1}^n X_k$$

On a, par la linéarité de l'espérance,

$$\mathbb{E}(Y_n) = \frac{1}{n} \sum_{k=1}^n \mathbb{E}(X_k) = \mu$$

Notons σ l'écart-type commun des X_k . En admettant (nous le montrerons plus loin) que la variance d'une somme de variables aléatoires indépendantes est la somme des variances,

$$\mathbb{V}(Y_n) = \frac{1}{n^2} \sum_{k=1}^n \mathbb{V}(X_k) = \frac{\sigma^2}{n}$$

Appliquons maintenant l'inégalité de Bienaymé-Tchebychev à Y_n . Pour tout $\varepsilon > 0$,

$$\mathbb{P}(|Y_n - \mu| \geq \varepsilon) \leq \frac{\sigma^2}{n\varepsilon^2}$$

Lorsque n tend vers l'infini, le majorant tend vers 0, d'où le résultat en passant au complémentaire. $\square$

Remarque 7. Le lecteur curieux se demandera évidemment ce qu'est la *loi forte des grands nombres* : on peut montrer (mais c'est plus difficile) que, sous les mêmes hypothèses que ci-dessus,

$$\frac{X_1 + \dots + X_n}{n} \xrightarrow[n \rightarrow \infty]{} \mu \text{ p.s.}$$

Nous n'en dirons pas plus.

3 Covariance

3.1 Notion de covariance

Définition 5. Soient X et Y deux variables aléatoires réelles. La *covariance* de X et Y est le réel

$$\text{Cov}(X, Y) = \mathbb{E}((X - \mathbb{E}(X))(Y - \mathbb{E}(Y)))$$

Remarquons que $\mathbb{V}(X) = \text{Cov}(X, X)$.

Proposition 18. Soient X et Y deux variables aléatoires réelles. On a

$$\text{Cov}(X, Y) = \mathbb{E}(XY) - \mathbb{E}(X)\mathbb{E}(Y)$$

Démonstration. On développe et on utilise la linéarité de l'espérance. $\square$

Corollaire 19. Si X et Y sont deux variables aléatoires réelles indépendantes, alors $\text{Cov}(X, Y) = 0$.

Démonstration. Si deux variables aléatoires réelles sont indépendantes, l'espérance du produit est le produit des espérances. $\square$

Remarque 8. Soit X une variable aléatoire suivant une loi uniforme sur $\{-1, 0, 1\}$. Soit $Y = X^2$. Les variables aléatoires X et Y ne sont pas indépendantes. en effet,

$$\mathbb{P}(X = -1, Y = 0) = 0$$

alors que

$$\mathbb{P}(X = -1)\mathbb{P}(Y = 0) = \frac{1}{9}$$

On a $XY = X^3 = X$ donc $\mathbb{E}(XY) = \mathbb{E}(X) = 0$. On a aussi $\mathbb{E}(X)\mathbb{E}(Y) = 0$ donc $\text{Cov}(X, Y) = 0$. La réciproque de la proposition ci-dessus est donc fausse.

Ceci suggère une définition.

Définition 6. Les variables aléatoires réelles X et Y sont *décorrélées* lorsque $\text{Cov}(X, Y) = 0$.

Ainsi, deux variables indépendantes sont décorréliées, mais la réciproque est fausse.

Exemple. On lance trois dés parfaits. Soit X la somme des deux premiers dés et Y la somme des deux derniers dés. Quelle est la covariance de X et Y ? Il nous faudrait pour l'instant calculer la loi du couple (X, Y) . Faisable, mais pénible. Attendons un peu ...

3.2 Propriétés de la covariance

Soit $(\Omega, \mathfrak{T}, \mathbb{P})$ un espace probabilisé. Notons $\mathcal{V}$ l'ensemble des variables aléatoires réelles définies sur Ω prenant un nombre fini de valeurs.

Proposition 20. Proposition $\mathcal{V}$ est une $\mathbb{R}$ -algèbre.

Démonstration. Montrons que $\mathcal{V}$ est une sous-algèbre de $\mathbb{R}^\Omega$. Soient $X, Y \in \mathcal{V}$. On a

$$(X + Y)(\Omega) \subset X(\Omega) + Y(\Omega)$$

et donc $(X + Y)(\Omega)$ est un ensemble fini. Soit $z \in (X + Y)(\Omega)$. On a

$$\{X + Y = z\} = \bigcup_{x \in X(\Omega)} (\{X = x\} \cap \{Y = z - x\})$$

Les ensembles $\{X = x\}$ et $\{Y = z - x\}$ appartiennent à $\mathfrak{T}$ donc, par stabilité d'une tribu par intersection et union finie, $\{X + Y = z\} \in \mathfrak{T}$. Ainsi, $X + Y \in \mathcal{V}$.

On démontre de même la stabilité de $\mathcal{V}$ pour la multiplication et le produit par un réel. Enfin, la fonction constante égale à 1 est clairement dans $\mathcal{V}$. $\square$

Proposition 21. La fonction $\text{Cov} : \mathcal{V} \times \mathcal{V} \longrightarrow \mathbb{R}$ est une forme bilinéaire symétrique.

Démonstration. Laissée en exercice. Cette propriété signifie :

- $\forall X, Y \in \mathcal{V}, \text{Cov}(X, Y) = \text{Cov}(Y, X)$.
- $\forall X, Y, Z \in \mathcal{V}, \text{Cov}(X + Y, Z) = \text{Cov}(X, Z) + \text{Cov}(Y, Z)$.
- $\forall X, Z \in \mathcal{V}, \forall \alpha \in \mathbb{R}, \text{Cov}(\alpha X, Z) = \alpha \text{Cov}(X, Z)$.
- Et de même pour la linéarité en la deuxième variable.

$\square$

Exemple. On reprend l'exemple avorté du lancer de 3 dés. Soient X et Y la somme des deux premiers dés et la somme des deux derniers dés. La bilinéarité de la covariance permet de calculer très simplement $\text{Cov}(X, Y)$. En effet, soient A, B, C les valeurs des 3 dés. Les variables aléatoires A, B, C sont indépendantes et on a $X = A + B$ et $Y = B + C$. D'où

$$\begin{aligned}\text{Cov}(X, Y) &= \text{Cov}(A + B, B + C) \\ &= \text{Cov}(A, B) + \text{Cov}(A, C) + \text{Cov}(B, B) + \text{Cov}(B, C) \\ &= \text{Cov}(A, B) + \text{Cov}(A, C) + \text{Cov}(B, C) + \mathbb{V}(B)\end{aligned}$$

Les trois covariances sont nulles puisque les variables aléatoires sont indépendantes. De plus, B suit une loi uniforme sur $\llbracket 1, 6 \rrbracket$, donc

$$\text{Cov}(X, Y) = \mathbb{V}(B) = \frac{6^2 - 1}{12} = \frac{35}{12}$$

Proposition 22. INÉGALITÉ DE SCHWARZ

Soient $X, Y \in \mathcal{V}$. On a

$$|\text{Cov}(X, Y)| \leq \sigma(X)\sigma(Y)$$

Il y a égalité si et seulement si il existe $a, b, c \in \mathbb{R}$ tels que $(a, b) \neq (0, 0)$ vérifiant

$$aX + bY + c = 0 \text{ p.s.}$$

Démonstration. Appliquons l'inégalité de Schwarz aux variables aléatoires $X - \mathbb{E}(X)$ et $Y - \mathbb{E}(Y)$ pour obtenir

$$|\mathbb{E}((X - \mathbb{E}(X))(Y - \mathbb{E}(Y)))| \leq \sqrt{\mathbb{E}((X - \mathbb{E}(X))^2)}\sqrt{\mathbb{E}((Y - \mathbb{E}(Y))^2)}$$

C'est exactement l'inégalité voulue. De plus, on a égalité si et seulement si il existe $a, b \in \mathbb{R}$ pas nuls tous les deux tels que

$$a(X - \mathbb{E}(X)) + b(Y - \mathbb{E}(Y)) = 0 \text{ p.s.}$$

qui s'écrit encore

$$aX + bY + c = 0 \text{ p.s.}$$

où $c = -a\mathbb{E}(X) - b\mathbb{E}(Y)$. Il reste à remarquer que si $aX + bY + c = 0$ p.s. alors, par linéarité de l'espérance, $a\mathbb{E}(X) + b\mathbb{E}(Y) + c = 0$ et donc $c = -a\mathbb{E}(X) - b\mathbb{E}(Y)$. $\square$

Remarque 9. Soient X et Y deux variables aléatoires réelles. Supposons $\sigma(X) > 0$ et $\sigma(Y) > 0$. Le réel

$$\rho(X, Y) = \frac{\text{Cov}(X, Y)}{\sigma(X)\sigma(Y)}$$

est le *coefficient de corrélation* de X et Y . Si nous réinterprétons ce qui précède :

- $-1 \leq \rho(X, Y) \leq 1$.
- $\rho(X, Y) = 0$ si et seulement si X et Y sont décorrélées.
- $\rho(X, Y) = \pm 1$ si et seulement si X et Y sont presque sûrement *affinement liées*.

3.3 Variance d'une somme de variables aléatoires

Proposition 23. VARIANCE D'UNE SOMME

Soient $X, Y \in \mathcal{V}$. On a

$$\mathbb{V}(X + Y) = \mathbb{V}(X) + \mathbb{V}(Y) + 2\text{Cov}(X, Y)$$

Proposition 24. Soit $n \geq 1$. Soient $X_1, \dots, X_n \in \mathcal{V}$. Alors,

$$\mathbb{V}(X_1 + \dots + X_n) = \mathbb{V}(X_1) + \dots + \mathbb{V}(X_n) + 2 \sum_{1 \leq i < j \leq n} \text{Cov}(X_i, X_j)$$

Démonstration. On fait la démonstration pour deux variables aléatoires. Pour n variables, c'est pareil. Par bilinéarité et symétrie.

$$\begin{aligned} \mathbb{V}(X + Y) &= \text{Cov}(X + Y, X + Y) \\ &= \text{Cov}(X, X) + 2\text{Cov}(X, Y) + \text{Cov}(Y, Y) \\ &= \mathbb{V}(X) + 2\text{Cov}(X, Y) + \mathbb{V}(Y) \end{aligned}$$

□

Proposition 25. Soit $n \geq 1$. Soient $X_1, \dots, X_n \in \mathcal{V}$. Si les variables aléatoires $X_1, \dots, X_n$ sont indépendantes, alors

$$\mathbb{V}\left(\sum_{k=1}^n X_k\right) = \sum_{k=1}^n \mathbb{V}(X_k)$$

Démonstration. Pour tous $i, j \in \llbracket 1, n \rrbracket$ tels que $i \neq j$, X_i et X_j sont indépendantes, donc $\text{Cov}(X_i, X_j) = 0$. D'où le résultat. □

Exemple. On lance n dés parfaits. Soit X la somme des valeurs des dés. Pour tout $k \in \llbracket 1, n \rrbracket$, soit X_k la variable aléatoire « valeur du dé numéro k ». On a

$$X = \sum_{k=1}^n X_k$$

Les X_k sont indépendantes, donc

$$\mathbb{V}(X) = \sum_{k=1}^n \mathbb{V}(X_k) = \frac{35n}{12}$$

3.4 La loi hypergéométrique

Pour terminer ce chapitre, nous allons calculer la variance de la loi hypergéométrique. Remettons d'abord toutes les notations en place. Une urne contient $N \geq 2$ boules, dont r sont rouges et $v = N - r$ sont vertes. Soit $1 \leq n \leq N$. On effectue n tirages sans remise dans l'urne. Choisissons pour univers l'ensemble $\Omega = \llbracket 1, N \rrbracket^n$ des n -uplets d'entiers distincts de l'ensemble $\llbracket 1, N \rrbracket$, muni de la tribu $\mathcal{P}(\Omega)$ et de la probabilité uniforme $\mathbb{P}$.

Pour tout $k \in \llbracket 1, n \rrbracket$, soit $X_k : \Omega \rightarrow \{0, 1\}$ la variable de Bernoulli définie par $X_k(x_1, \dots, x_n) = 1$ si x_k est rouge, et 0 sinon. Notons enfin

$$S = \sum_{k=1}^n X_k$$

Nous avons déjà vu que les variables aléatoires X_k suivent une loi de Bernoulli de paramètre $p = \mathbb{E}(X_k) = r/N$. De là,

$$\mathbb{E}(S) = \sum_{k=1}^n \mathbb{E}(X_k) = \frac{nr}{N}$$

Calculons maintenant la variance de S . Nous avons

$$\mathbb{V}(S) = \sum_{k=1}^n \mathbb{V}(X_k) + 2 \sum_{1 \leq i < j \leq n} \text{Cov}(X_i, X_j)$$

Les variances des X_k nous sont connues. Pour tout $k \in \llbracket 1, n \rrbracket$,

$$\mathbb{E}(X_k) = p(1 - p) = \frac{r(N - r)}{N^2}$$

Il reste à calculer les covariances.

Commençons par remarquer le fait suivant : Soient $1 \leq i < j \leq n$. La fonction

$$(x_1, x_2, \dots, x_n) \mapsto (x_i, x_j, \dots, x_n)$$

qui échange x_1 et x_i , ainsi que x_2 et x_j , est une bijection de $\{X_1 = 1, X_2 = 1\}$ sur $\{X_i = 1, X_j = 1\}$. Comme la probabilité $\mathbb{P}$ est la probabilité uniforme, ces deux événements ont la même probabilité. La variable aléatoire $X_i X_j$ ne prend que les valeurs 0 et 1, c'est donc une variable de Bernoulli. De plus,

$$\{X_i = 1, X_j = 1\} = \{X_i X_j = 1\}$$

et donc

$$\mathbb{P}(X_i = 1, X_j = 1) = \mathbb{P}(X_i X_j = 1) = \mathbb{E}(X_i X_j)$$

Ainsi,

$$\mathbb{E}(X_i X_j) = \mathbb{E}(X_1 X_2)$$

$$\text{Cov}(X_i, X_j) = \text{Cov}(X_1, X_2) = \mathbb{E}(X_1 X_2) - \mathbb{E}(X_1)\mathbb{E}(X_2)$$

Il reste à déterminer $\mathbb{E}(X_1 X_2)$.

$$\mathbb{E}(X_1 X_2) = \mathbb{P}(X_1 = 1, X_2 = 1) = \mathbb{P}(X_2 = 1 | X_1 = 1) \mathbb{P}(X_1 = 1) = p \frac{r-1}{N-1}$$

et donc

$$\text{Cov}(X_1, X_2) = p \frac{r-1}{N-1} - p^2 = p \left(\frac{r-1}{N-1} - \frac{r}{N} \right) = -p \frac{N-r}{N(N-1)} = -\frac{p(1-p)}{(N-1)}$$

Remarquons enfin que le nombre de couples (i, j) tels que $1 \leq i < j \leq n$ est $\frac{1}{2}n(n-1)$. De là,

$$\begin{aligned} \mathbb{V}(S) &= n\mathbb{V}(X_1) + 2\frac{1}{2}n(n-1)\text{Cov}(X_1, X_2) \\ &= np(1-p) - n(n-1)\frac{p(1-p)}{(N-1)} \\ &= \lambda np(1-p) \end{aligned}$$

où

$$\lambda = \frac{N-n}{N-1} \leq 1$$

Nous constatons que cette variance est inférieure à celle que nous aurions obtenue pour un tirage avec remise. Nous voyons également que si n est petit devant N , $\lambda \simeq 1$ et la variance de S s'approche alors de la variance $np(1-p)$ de la loi binomiale du tirage sans remise. Enfin, si $n = N$, on a $\mathbb{V}(S) = 0$ ce qui est normal puisque dans ce cas S est constante égale à r (si on retire toutes les boules de l'urne, le nombre de boules rouges tirées est constant égal à r).

Chapitre 33

Espaces préhilbertiens

1 Produits scalaires

1.1 Notion de produit scalaire

Définition 1. Soit E un $\mathbb{R}$ -espace vectoriel. On appelle produit scalaire sur E toute application $\varphi : E \times E \longrightarrow \mathbb{R}$ vérifiant :

- φ est bilinéaire.
- φ est symétrique : $\forall x, y \in E, \varphi(x, y) = \varphi(y, x)$.
- φ est « positive » : $\forall x \in E, \varphi(x, x) \geq 0$.
- φ est « définie » : $\forall x \in E, \varphi(x, x) = 0 \Rightarrow x = 0$.

Notation. Les produits scalaires sont notés de beaucoup de façons : $x.y$ ou $(x|y)$ ou $\langle x, y \rangle$, etc. Dans la suite, nous utiliserons la notation $\langle \bullet, \bullet \rangle$.

Exemple.

- Prenons $E = \mathbb{R}^n$. Le *produit scalaire canonique* sur E est

$$\langle x, y \rangle = \sum_{k=1}^n x_k y_k$$

- Prenons $E = \mathcal{C}^0([a, b])$, où $a < b$:

$$\langle f, g \rangle = \int_{-1}^1 f(x)g(x) dx$$

- Prenons $E = \mathbb{R}[X]$:

$$\langle P, Q \rangle = \int_{-1}^1 P(x)Q(x) dx$$

- Prenons $E = \mathcal{M}_n(\mathbb{R})$:

$$\langle A, B \rangle = \text{tr}(A^T B)$$

Le lecteur est invité à vérifier que l'on a bien là des produits scalaires.

Exercice. Trouver tous les produits scalaires sur $\mathbb{R}^n$ pour $n = 1$ et $n = 2$.

Pour $n = 1$, c'est immédiat. Soit E un $\mathbb{R}$ -espace vectoriel de dimension 1. Soit e un vecteur non nul de E . Soit $\langle, \rangle$ un produit scalaire sur E . On a pour tous $x, y \in \mathbb{R}$,

$$\langle xe, ye \rangle = xy \langle e, e \rangle = \lambda xy$$

où $\lambda = \langle e, e \rangle > 0$. Inversement, pour tout $\lambda > 0$, la fonction $(xe, ye) \mapsto \lambda xy$ est clairement un produit scalaire sur E .

Prenons maintenant un $\mathbb{R}$ espace vectoriel E de dimension 2. Soit $\mathcal{B} = (e_1, e_2)$ une base de E . Soit $\langle, \rangle$ un produit scalaire sur E . Soient $u, v \in E$. Écrivons $u = xe_1 + ye_2$ et $v = x'e_1 + y'e_2$. Il vient, par bilinéarité et symétrie,

$$\begin{aligned}\langle u, v \rangle &= xx' \langle e_1, e_1 \rangle + (xy' + x'y) \langle e_1, e_2 \rangle + yy' \langle e_2, e_2 \rangle \\ &= axx' + b(xy' + x'y) + cyy'\end{aligned}$$

où $a = \langle e_1, e_1 \rangle$, $b = \langle e_1, e_2 \rangle$ et $c = \langle e_2, e_2 \rangle$. En particulier,

$$\langle u, u \rangle = ax^2 + 2bxy + cy^2$$

Par la « positivité » et la « définition », $a > 0$ et $c > 0$. Toujours par ces deux propriétés, on a pour tous $x, y \in \mathbb{R}$,

$$ax^2 + 2bxy + cy^2 \geq 0$$

avec nullité si et seulement si $x = y = 0$. Cette contrainte est évidemment très forte. Prenant $y = 1$, en particulier on obtient pour tout $x \in \mathbb{R}$,

$$ax^2 + 2bx + c > 0$$

On a donc un trinôme du second degré qui ne s'annule pas. Cela impose $b^2 - ac < 0$.

On laisse au lecteur de vérifier la réciproque. En résumé, les produits scalaires sur E sont les fonctions de la forme

$$(u, v) \mapsto axx' + b(xy' + x'y) + cyy'$$

où

$$a > 0, c > 0, b^2 - ac < 0$$

Définition 2. On appelle espace préhilbertien réel tout espace vectoriel réel muni d'un produit scalaire. On appelle espace euclidien tout espace préhilbertien réel de dimension finie.

1.2 Inégalité de Cauchy-Schwarz

Proposition 1. INÉGALITÉ DE SCHWARZ

Soit E un espace préhilbertien réel. Alors :

- Pour tous $x, y \in E$,

$$|\langle x, y \rangle| \leq \sqrt{\langle x, x \rangle \langle y, y \rangle}$$

- Il y a égalité si et seulement si x et y sont liés.

Démonstration.

- Si y est nul, on a égalité, ce qui prouve la proposition dans ce cas particulier.
- Supposons donc $y \neq 0$. Soit λ un réel. Alors,

$$0 \leq \langle x + \lambda y, x + \lambda y \rangle = \langle x, x \rangle + 2\lambda \langle x, y \rangle + \lambda^2 \langle y, y \rangle$$

On a là un trinôme du second degré qui garde un signe constant : son discriminant est donc négatif ou nul : c'est l'inégalité de Schwarz. Si x et y sont colinéaires, l'égalité est claire. Inversement, si il y a égalité, le trinôme ci-dessus a un discriminant nul. Le trinôme a ainsi une racine réelle α . On a alors $\langle x + \alpha y, x + \alpha y \rangle = 0$, d'où $x + \alpha y = 0$. Les vecteurs x et y sont liés.

□

1.3 Norme euclidienne

Définition 3. Soit E un espace vectoriel réel (ou complexe). Une norme sur E est une application $\Phi : E \rightarrow \mathbb{R}$ vérifiant :

- $\forall x \in E, \Phi(x) \geq 0$.
- $\forall x \in E, \Phi(x) = 0 \Rightarrow x = 0$.
- $\forall x \in E, \forall \lambda \in \mathbb{K}, \Phi(\lambda x) = |\lambda| \Phi(x)$ (homogénéité).
- $\forall x, y \in E, \Phi(x + y) \leq \Phi(x) + \Phi(y)$ (inégalité de Minkowski).

Proposition 2. Soit E un espace préhilbertien. L'application $E \rightarrow \mathbb{R}$ qui à tout vecteur $x \in E$ associe $\|x\| = \sqrt{\langle x, x \rangle}$ est une norme sur E , appelée norme euclidienne (associée au produit scalaire en question).

Démonstration. L'inégalité triangulaire est la seule des propriétés qui soit non triviale :

$$\|x + y\|^2 = \langle x + y, x + y \rangle = \|x\|^2 + \|y\|^2 + 2 \langle x, y \rangle$$

On termine grâce à l'inégalité de Schwarz. □

Remarque 1. L'inégalité de Schwarz s'écrit dorénavant :

$$|\langle x, y \rangle| \leq \|x\| \cdot \|y\|$$

On constate donc que si x et y sont non nuls, on a

$$-1 \leq \frac{\langle x, y \rangle}{\|x\| \cdot \|y\|} \leq 1$$

Définition 4. Soient x et y deux vecteurs non nuls de l'espace préhilbertien E . On appelle écart angulaire entre les vecteurs x et y le réel $\theta \in [0, \pi]$ défini par :

$$\theta = \arccos \frac{\langle x, y \rangle}{\|x\| \cdot \|y\|}$$

Exercice. Montrer les propriétés suivantes :

- $\|x + y\| = \|x\| + \|y\|$ si et seulement si x et y sont colinéaires et de même sens.
- $|\|x\| - \|y\|| \leq \|x - y\|$

Indications : Pour le premier point, reprendre la démonstration de l'inégalité triangulaire. Puis utiliser le cas d'égalité dans l'inégalité de Schwarz. Pour le second point, écrire $x = (x - y) + y$ puis utiliser l'inégalité triangulaire : vous aurez fait la moitié du travail.

Définition 5. On appelle vecteur unitaire tout vecteur de norme 1.

Proposition 3. Soit E un espace préhilbertien. Tout vecteur $u \neq 0$ est colinéaire à exactement deux vecteurs unitaires : $\frac{u}{\|u\|}$ et $-\frac{u}{\|u\|}$.

Démonstration. Soit $\lambda \in \mathbb{R}$. On a

$$\|\lambda u\| = |\lambda| \|u\|$$

Ainsi, $\|\lambda u\| = 1$ si et seulement si

$$\lambda = \pm \frac{1}{\|u\|}$$

□

1.4 Distance euclidienne

Définition 6. Soit E un ensemble non vide. On appelle distance sur E toute application $d : E \times E \rightarrow \mathbb{R}$ vérifiant :

- $\forall x, y \in E, d(x, y) \geq 0$.
- $\forall x, y \in E, d(x, y) = 0 \Leftrightarrow x = y$.
- $\forall x, y \in E, d(x, y) = d(y, x)$.
- $\forall x, y, z \in E, d(x, z) \leq d(x, y) + d(y, z)$.

Proposition 4. Soit E un espace préhilbertien. L'application $d : E \times E \rightarrow \mathbb{R}$ définie par $d(x, y) = \|y - x\|$ est une distance sur E , appelée distance euclidienne (associée au produit scalaire de E).

Démonstration. Exercice. Cela résulte facilement des propriétés de la norme. $\square$

1.5 Relations utiles

Proposition 5. On a dans tout espace préhilbertien les égalités suivantes :

- $\|x + y\|^2 = \|x\|^2 + \|y\|^2 + 2 \langle x, y \rangle$
- $\|x - y\|^2 = \|x\|^2 + \|y\|^2 - 2 \langle x, y \rangle$
- $\langle x, y \rangle = \frac{1}{4}(\|x + y\|^2 - \|x - y\|^2)$ (identité de polarisation).
- $\langle x, y \rangle = \frac{1}{2}(\|x + y\|^2 - \|x\|^2 - \|y\|^2)$ (identité de polarisation).
- $\langle x, y \rangle = \frac{1}{2}(-\|x - y\|^2 + \|x\|^2 + \|y\|^2)$ (identité de polarisation).
- $\|x + y\|^2 + \|x - y\|^2 = 2(\|x\|^2 + \|y\|^2)$ (identité du parallélogramme).

Démonstration.

- $\|x + y\|^2 = \langle x + y, x + y \rangle = \|x\|^2 + \|y\|^2 + 2 \langle x, y \rangle$ par bilinéarité du produit scalaire.
- Preuve identique au premier point.
- Soustraire les égalités du premier point.
- Réarranger l'égalité du premier point.
- Réarranger l'égalité du deuxième point.
- Additionner les égalités du premier point.

$\square$

Exercice. On définit sur $\mathbb{R}^2$ $\|(x, y)\| = |x| + |y|$. Vérifier que l'on a là une norme, mais que cette norme n'est pas euclidienne.

La propriété de norme est facile à vérifier. Supposons maintenant que l'on a là une norme euclidienne, associée à un produit scalaire $\langle \bullet, \bullet \rangle$ sur $\mathbb{R}^2$.

Prenons $u = (1, 1)$, $v = (1, -2)$ et $w = (-1, 1)$. On a

$$\|u\| = 2, \|v\| = 3, \|w\| = 2$$

On a aussi

$$\|u + v\| = 3, \|u + w\| = 2, \|v + w\| = 1, \|u + v + w\| = 1$$

De là, par l'identité de polarisation,

$$\langle u, w \rangle = \frac{1}{2}(\|u + w\|^2 - \|u\|^2 - \|w\|^2) = \frac{1}{2}(4 - 4 - 4) = -2$$

De même,

$$\langle v, w \rangle = \frac{1}{2}(1 - 9 - 4) = -6$$

et donc, par bilinéarité du produit scalaire,

$$\langle u + v, w \rangle = \langle u, w \rangle + \langle v, w \rangle = -8$$

Mais par l'identité de polarisation,

$$\langle u + v, w \rangle = \frac{1}{2}(\|u + v + w\|^2 - \|u + v\|^2 - \|w\|^2) = \frac{1}{2}(1 - 9 - 4) = -6$$

On a donc une contradiction. Ainsi, cette norme ne peut pas être une norme euclidienne.

2 Orthogonalité

2.1 Vecteurs orthogonaux

Définition 7. VECTEURS ORTHOGONAUX

Deux vecteurs x et y de l'espace préhilbertien E sont dits orthogonaux, et on note $x \perp y$, lorsque $\langle x, y \rangle = 0$.

Proposition 6. Soit E un espace préhilbertien.

- Pour tous $x, y \in E$, $x \perp y \iff y \perp x$.
- Pour tout $x \in E$, $x \perp x \iff x = 0$.
- Le seul vecteur orthogonal à tous les vecteurs de E est le vecteur nul.

Démonstration.

- Le produit scalaire est symétrique, donc l'orthogonalité des vecteurs est une relation symétrique.
- $x \perp x \iff \langle x, x \rangle = 0 \iff x = 0$.
- si x est orthogonal à tous les vecteurs de E il est orthogonal à lui-même. Donc il est nul.

□

Définition 8. FAMILLES ORTHOGONALES

Une famille $\mathcal{F} = (e_i)_{i \in I}$ de vecteurs de E est dite orthogonale lorsque :

$$\forall i, j \in I, i \neq j \implies \langle e_i, e_j \rangle = 0$$

Si, de plus, on a $\forall i \in I, \langle e_i, e_i \rangle = 1$, on dit que la famille est orthonormée (ou orthonormale).

Proposition 7. THÉORÈME DE PYTHAGORE

Soit $(e_i)_{1 \leq i \leq n}$ une famille finie de vecteurs. Si cette famille est orthogonale, alors

$$\left\| \sum_{k=1}^n e_k \right\|^2 = \sum_{k=1}^n \|e_k\|^2$$

La réciproque est vraie dans le cas où $n = 2$.

Démonstration.

$$\begin{aligned} \left\| \sum_{k=1}^n e_k \right\|^2 &= \left\langle \sum_{k=1}^n e_k, \sum_{k=1}^n e_k \right\rangle \\ &= \sum_{k=1}^n \|e_k\|^2 + 2 \sum_{i < j} \langle e_i, e_j \rangle \end{aligned}$$

par bilinéarité du produit scalaire.

Si la famille est orthogonale, les produits scalaires $\langle e_i, e_j \rangle$ sont nuls pour $i \neq j$. On a donc l'égalité de Pythagore.

Pour $n = 2$, l'égalité ci-dessus devient

$$\|e_1 + e_2\|^2 = \|e_1\|^2 + \|e_2\|^2 + 2 \langle e_1, e_2 \rangle$$

On a l'égalité de Pythagore si et seulement si $\langle e_1, e_2 \rangle = 0$. $\square$

Proposition 8. Toute famille orthogonale de vecteurs non nuls est libre.

Démonstration. Soit $(e_i)_{i \in I}$ une famille orthogonale de vecteurs non nuls. Soit $(\lambda_i)_{i \in I}$ une famille presque nulle de réels. Supposons $\sum_{i \in I} \lambda_i e_i = 0$.

Soit $j \in I$. On a, en faisant le produit scalaire par e_j ,

$$\left\langle \sum_{i \in I} \lambda_i e_i, e_j \right\rangle = \lambda_j \langle e_j, e_j \rangle = \langle 0, e_j \rangle = 0$$

D'où $\lambda_j = 0$ puisque e_j est non nul et, donc, $\langle e_j, e_j \rangle \neq 0$. $\square$

2.2 Sous-ensembles orthogonaux

Définition 9. Soient A et B deux parties de E . On dit que A est orthogonale à B , et on note $A \perp B$, lorsque

$$\forall (x, y) \in A \times B, x \perp y$$

On parle ainsi par exemple de deux droites orthogonales du plan ou de l'espace, ou encore d'une droite et d'un plan orthogonaux dans $\mathbb{R}^3$. En revanche, deux plans de $\mathbb{R}^3$ ne sont jamais orthogonaux.

Exercice. Pourquoi? Et est-ce que deux plans peuvent être orthogonaux en dimension 4?

2.3 Orthogonal d'une partie

Définition 10. Soit $A \subset E$. On appelle orthogonal de A l'ensemble

$$A^\perp = \{y \in E, \forall x \in A, \langle x, y \rangle = 0\}$$

Proposition 9.

1. $E^\perp = \{0\}$ et $\{0\}^\perp = E$.
2. A et $A^\perp$ sont orthogonaux.
3. $A^\perp$ est un sous-espace vectoriel de E .
4. $A \subset B \Rightarrow B^\perp \subset A^\perp$
5. $A^\perp = \langle A \rangle^\perp$
6. $A \subset A^{\perp\perp}$

Remarque 2. Nous verrons bientôt que la dernière inclusion est en fait une *égalité* en dimension finie.

Démonstration.

1. 0 est le seul vecteur de E orthogonal à tous les vecteurs de E , donc $E^\perp = \{0\}$. Tous les vecteurs de E sont orthogonaux à 0 donc $\{0\}^\perp = E$.
2. Soit A une partie de E . Soient $x \in A$ et $y \in A^\perp$. On a $y \perp x$ donc A et $A^\perp$ sont bien orthogonaux.
3. $A^\perp$ est non vide, car il contient 0. Soient $x, y \in A^\perp$ et $z \in A$. Soit $\lambda \in \mathbb{R}$. On a $\langle x, z \rangle = \langle y, z \rangle = 0$, donc $\langle \lambda x + y, z \rangle = \lambda \langle x, z \rangle + \langle y, z \rangle = 0$.

4. Supposons $A \subset B$. Un vecteur orthogonal à tous les vecteurs de B est a fortiori orthogonal à tous les vecteurs de A . Donc $B^\perp \subset A^\perp$.
5. On a $A \subset \langle A \rangle$, donc $\langle A \rangle^\perp \subset A^\perp$. Inversement, Soit $y \in A^\perp$. Il est orthogonal à tous les vecteurs de A , donc, par bilinéarité du produit scalaire, à toutes les combinaisons linéaires de vecteurs de A . Donc $y \in \langle A \rangle^\perp$.
6. Soit $x \in A$. On a pour tout $y \in A^\perp$, $x \perp y$. Donc, $x \in A^{\perp\perp}$.

□

2.4 Algorithme de Gram-Schmidt

Dans cette partie, E est un espace euclidien (de dimension n sauf mention contraire).

Proposition 10. Soit $\mathcal{B} = (e_1, \dots, e_n)$ une base de E . Il existe une base orthogonale $\mathcal{B}' = (u_1, \dots, u_n)$ de E telle que

$$\forall i \in \llbracket 1, n \rrbracket, \text{Vect}(e_1, \dots, e_i) = \text{Vect}(u_1, \dots, u_i)$$

Démonstration. Posons $u_1 = e_1$, et, pour $k = 1, \dots, n-1$,

$$u_{k+1} = e_{k+1} - \sum_{i=1}^k \frac{\langle e_{k+1}, u_i \rangle}{\langle u_i, u_i \rangle} u_i$$

Montrons par récurrence sur k que pour tout $k \in \llbracket 1, n \rrbracket$,

- $u_1, \dots, u_k$ sont non nuls.
- $u_1, \dots, u_k$ sont orthogonaux deux à deux.
- $\text{Vect}(e_1, \dots, e_k) = \text{Vect}(u_1, \dots, u_k)$

Pour $k = 1$ c'est immédiat. Soit donc $1 \leq k \leq n-1$. Supposons la propriété vraie pour k .

- Par l'hypothèse de récurrence, $u_1, \dots, u_k$ sont combinaisons de $e_1, \dots, e_k$. Comme la famille $\mathcal{B}$ est libre, e_{k+1} n'est donc pas combinaison de $u_1, \dots, u_k$. Ainsi, $u_{k+1} \neq 0$. Et par l'hypothèse de récurrence, $u_1, \dots, u_k$ ne sont pas nuls non plus.
- Soit $1 \leq j \leq k$. Montrons que u_{k+1} est orthogonal à u_j . Pour cela effectuons un produit scalaire :

$$\begin{aligned} \langle u_{k+1}, u_j \rangle &= \left\langle e_{k+1} - \sum_{i=1}^k \frac{\langle e_{k+1}, u_i \rangle}{\langle u_i, u_i \rangle} u_i, u_j \right\rangle \\ &= \langle e_{k+1}, u_j \rangle - \sum_{i=1}^k \frac{\langle e_{k+1}, u_i \rangle}{\langle u_i, u_i \rangle} \langle u_i, u_j \rangle \end{aligned}$$

Dans la somme ci-dessus, tous les termes sont nuls (orthogonalité par l'hypothèse de récurrence), sauf lorsque $i = j$. Il reste donc

$$\langle u_{k+1}, u_j \rangle = \langle e_{k+1}, u_j \rangle - \frac{\langle e_{k+1}, u_j \rangle}{\langle u_j, u_j \rangle} \langle u_j, u_j \rangle = 0$$

- Par l'hypothèse de récurrence, $u_1, \dots, u_k$ sont combinaisons linéaires de $e_1, \dots, e_k$. De là, l'expression de u_{k+1} nous dit que u_{k+1} , combinaison linéaire de $u_1, \dots, u_k, e_{k+1}$, est combinaison linéaire de $e_1, \dots, e_k, e_{k+1}$. Ainsi, $u_1, \dots, u_{k+1}$ sont combinaisons linéaires de $e_1, \dots, e_{k+1}$. On a donc

$$\text{Vect}(u_1, \dots, u_{k+1}) \subset \text{Vect}(e_1, \dots, e_{k+1})$$

L'inclusion réciproque se montre de façon analogue.

En prenant $k = n$, on obtient donc que $\mathcal{B}'$ est famille orthogonale de vecteurs non nuls de E . C'est donc une famille libre. Comme son cardinal est la dimension de E , $\mathcal{B}'$ est bien une base orthogonale de E .

cette famille répond aux conditions du théorème. $\square$

Remarque 3. Nous verrons bientôt que $e_{k+1} - u_{k+1}$ n'est autre que la projection orthogonale de e_{k+1} sur $\langle u_1, \dots, u_k \rangle$.

Corollaire 11. *Tout espace euclidien possède des bases orthogonales (et même ortho-normales).*

Remarque 4. Toute famille orthogonale de vecteurs non nuls de E peut être complétée en une base orthogonale de E : il suffit pour cela d'appliquer :

- le théorème de la base incomplète, puis
- l'algorithme de Schmidt (à partir d'un certain rang).

Exercice. Orthogonaliser la famille des trois vecteurs de l'espace, $(1, 2, 3)$, $(4, 5, 6)$ et $(7, 8, 8)$.

On reprend les notations du théorème de Gram-Schmidt. Posons $e_1 = (1, 2, 3)$, $e_2 = (4, 5, 6)$ et $e_3 = (7, 8, 8)$. La famille $\mathcal{B} = (e_1, e_2, e_3)$ est bien une base de E (calculer par exemple son déterminant pour s'en convaincre). On a $u_1 = e_1 = (1, 2, 3)$, puis

$$u_2 = e_2 - \frac{\langle e_2, u_1 \rangle}{\langle u_1, u_1 \rangle} u_1 = (4, 5, 6) - \frac{16}{7} (1, 2, 3) = \frac{3}{7} (4, 1, -2)$$

Inutile de s'embarasser de la fraction, prenons plutôt $u_2 = (4, 1, -2)$. Enfin,

$$\begin{aligned} u_3 &= e_3 - \frac{\langle e_3, u_1 \rangle}{\langle u_1, u_1 \rangle} u_1 - \frac{\langle e_3, u_2 \rangle}{\langle u_2, u_2 \rangle} u_2 \\ &= (7, 8, 8) - \frac{47}{14} (1, 2, 3) - \frac{20}{21} (4, 1, -2) \end{aligned}$$

On laisse au lecteur le soin de finir le calcul.

2.5 Conséquences

Proposition 12. Soit E un espace euclidien. Soit F un sous-espace vectoriel de E . Alors

- $F \oplus F^\perp = E$
- $F^{\perp\perp} = F$

Démonstration. Soit $(e_1, \dots, e_p)$ une b.o.n. de F . On la complète en une b.o.n. $(e_1, \dots, e_n)$ de E . Soit $G = \langle e_{p+1}, \dots, e_n \rangle$. On montre facilement que $G = F^\perp$. $\square$

Remarque 5. Parmi tous les supplémentaires de F , il y en a donc un qui se distingue des autres : c'est $F^\perp$, LE supplémentaire orthogonal de F .

Proposition 13. Soit E un espace euclidien. Pour tout vecteur $a \in E$, soit $\varphi_a : E \rightarrow \mathbb{R}$ définie par $\varphi_a(x) = \langle a, x \rangle$. L'application $\varphi : a \mapsto \varphi_a$ est un isomorphisme de E sur son dual E^* .

Démonstration. L'application φ est clairement linéaire, et aboutit bien dans E^* . De plus, si $a \in \ker \varphi$, alors, a est orthogonal à tous les vecteurs de E , donc $a = 0$: l'application φ est injective. En raison de l'égalité des dimensions de E et E^* , φ est donc un isomorphisme d'espaces vectoriels. $\square$

Remarque 6. Ce résultat a priori abstrait est intéressant lorsqu'on le relie à ce que l'on sait sur les hyperplans et les formes linéaires. Rappelons que les hyperplans de E sont les noyaux des formes linéaires non nulles sur E . Or, le noyau de la forme linéaire φ_a est l'orthogonal de la droite $\text{Vect}(a)$.

Encore plus concrètement, plaçons nous dans $\mathbb{R}^3$ et considérons le plan d'équation

$$(P) \quad x - 2y + 3z = 0$$

Soit $a = (1, -2, 3)$. Le vecteur u est dans P si et seulement si $\varphi_a(u) = \langle a, u \rangle = 0$. Bref, le théorème ci-dessus est une façon très abstraite d'exprimer que P est l'orthogonal de la droite $D = \text{Vect}(a)$.

2.6 Calculs en bases orthonormées

Proposition 14. Soit E un espace euclidien. Soit $\mathcal{B} = (e_1, \dots, e_n)$ une base orthonormée de E . Soient $x, y \in E$ de coordonnées x_i, y_i dans la base $\mathcal{B}$. On a alors

- $x_i = \langle x, e_i \rangle$
- $\langle x, y \rangle = \sum x_i y_i$
- $\|x\|^2 = \sum x_i^2$

En définitive, tous les calculs sont identiques à ceux que l'on ferait dans $\mathbb{R}^n$ muni du produit scalaire canonique.

Démonstration. Soit n est la dimension de E . La base $\mathcal{B}$ étant orthonormée, on a

$$\forall i, j \in \llbracket 1, n \rrbracket \quad \langle e_i, e_j \rangle = \delta_{ij}$$

où δ_{ij} , le symbole de Kronecker, vaut 0 si $i \neq j$ et 1 si $i = j$.

- Soit $i \in \llbracket 1, n \rrbracket$. On a

$$\langle x, e_i \rangle = \left\langle \sum_{j=1}^n x_j e_j, e_i \right\rangle = \sum_{j=1}^n x_j \langle e_j, e_i \rangle = \sum_{j=1}^n x_j \delta_{ji} = x_i$$

- On a

$$\langle x, y \rangle = \left\langle \sum_{i=1}^n x_i e_i, \sum_{j=1}^n y_j e_j \right\rangle = \sum_{i,j=1}^n x_i y_j \langle e_i, e_j \rangle = \sum_{i,j=1}^n x_i y_j \delta_{ij} = \sum_{i=1}^n x_i y_i$$

- L'égalité $\|x\|^2 = \sum x_i^2$ est une conséquence du point précédent (prendre $y = x$).

□

Remarque 7. En base orthogonale, ces formules se compliquent un peu parce que les normes $\|e_i\|$ ne sont plus égales à 1. On a alors (exercice)

- $x_i = \frac{\langle x, e_i \rangle}{\|e_i\|^2}$
- $\langle x, y \rangle = \sum \|e_i\|^2 x_i y_i$
- $\|x\|^2 = \sum \|e_i\|^2 x_i^2$

Exercice. Que deviennent ces formules dans une base quelconque ?

3 Projecteurs orthogonaux - Symétries orthogonales

3.1 Définitions

Définition 11. Soit F un sous-espace vectoriel de l'espace euclidien E . Le projecteur orthogonal sur F est le projecteur sur F parallèlement à $F^\perp$. De même, la symétrie orthogonale par rapport à F est la symétrie par rapport à F parallèlement à $F^\perp$.

Remarque 8. Supposons donnée une base orthogonale de F , $\mathcal{F} = (u_1, \dots, u_k)$. Soit x un vecteur de E et $p_F(x)$ son projeté orthogonal sur F . Alors,

$$p_F(x) = \sum_{i=1}^k \frac{\langle x, u_i \rangle}{\langle u_i, u_i \rangle} u_i$$

On reconnaît là les expressions qui intervenaient dans l'algorithme de Schmidt.

Bien entendu, si la base est orthonormée, tout ceci se simplifie et devient :

$$p_F(x) = \sum_{i=1}^k \langle x, u_i \rangle u_i$$

3.2 Cas particuliers

Soit $D = \langle e \rangle$ où e est un vecteur non nul. Alors

$$p_D(x) = \frac{\langle x, e \rangle}{\|e\|^2} e$$

Soit $H = e^\perp$ un hyperplan de E . Alors

$$p_H(x) = x - \frac{\langle x, e \rangle}{\|e\|^2} e$$

puisque $p_D + p_H = id_E$.

Il est donc immédiat de calculer un projeté orthogonal sur une droite ou sur un hyperplan.

3.3 Distance d'un vecteur à un sous-espace vectoriel

Proposition 15. *Soit F un sous-espace vectoriel de l'espace euclidien E . Soit $x \in E$. Soit $p(x)$ le projeté orthogonal de x sur F . On a alors*

- $\forall y \in F, d(x, y) \geq d(x, p(x)).$
- $\forall y \in F, d(x, y) = d(x, p(x)) \Rightarrow y = p(x).$

En clair, $p(x)$, et lui seul, réalise le minimum de distance entre x et les vecteurs de F .

Démonstration. Soit $y \in F$. On a

$$y - x = (y - p(x)) + (p(x) - x)$$

Comme $x - p(x) \in F^\perp$ et $y - p(x) \in F$, ces deux vecteurs sont orthogonaux. Ainsi, par Pythagore,

$$\|y - x\|^2 = \|y - p(x)\|^2 + \|p(x) - x\|^2$$

ou encore $d(x, y)^2 = d(x, p(x))^2 + d(p(x), y)^2$. Comme un carré est positif, on a donc $d(x, y)^2 \geq d(x, p(x))^2$. Et comme seul le carré de 0 est nul, on a égalité si et seulement si $d(p(x), y) = 0$, c'est à dire $y = p(x)$. $\square$

Définition 12. On appelle distance de x à F la quantité $d(x, F) = \|x - p(x)\|$.

Exemple. En reprenant les expressions du projeté sur un hyperplan et une droite vus un peu plus haut, on obtient que

- Si $H = e^\perp$ est un hyperplan de E , alors

$$d(x, H) = \frac{\langle x, e \rangle}{\|e\|}$$

- Si $D = \langle e \rangle$ est une droite vectorielle, alors

$$d(x, D)^2 = \left\| x - \frac{\langle x, e \rangle}{\|e\|^2} e \right\|^2$$

Développons cette expression. On a

$$\begin{aligned} d(x, D)^2 &= \left\langle x - \frac{\langle x, e \rangle}{\|e\|^2} e, x - \frac{\langle x, e \rangle}{\|e\|^2} e \right\rangle \\ &= \|x\|^2 - 2 \left\langle x, \frac{\langle x, e \rangle}{\|e\|^2} e \right\rangle + \left\| \frac{\langle x, e \rangle}{\|e\|^2} e \right\|^2 \\ &= \|x\|^2 - \frac{\langle x, e \rangle^2}{\|e\|^2} \\ &= \frac{1}{\|e\|^2} (\|x\|^2 \|e\|^2 - \langle x, e \rangle^2) \end{aligned}$$

Ainsi,

$$d(x, D) = \frac{\sqrt{\|x\|^2 \|e\|^2 - \langle x, e \rangle^2}}{\|e\|}$$

Chapitre 34

Endomorphismes orthogonaux

1 Généralités

Dans tout le chapitre, $(E, <, >)$ désigne un espace euclidien.

1.1 Notion d'endomorphisme orthogonal

Définition 1. Soit f un endomorphisme de E . On dit que f est un endomorphisme orthogonal lorsque f conserve les produits scalaires, c'est à dire :

$$\forall x, y \in E, < f(x), f(y) > = < x, y >$$

Remarque 1. En fait, la linéarité résulte de la conservation des produits scalaires. En effet, soit $f : E \rightarrow E$ conservant les produits scalaires. Soit $\mathcal{B} = (e_1, \dots, e_n)$ une base orthonormée de E . On a pour tous i, j ,

$$< f(e_i), f(e_j) > = < e_i, e_j > = \delta_{ij}$$

Donc la famille $\mathcal{B}' = f(\mathcal{B})$ est une base orthonormée de E .

Soient maintenant $x, y \in E$ et $i \in \llbracket 1, n \rrbracket$. Soit

$$X = < f(x + y) - f(x) - f(y), f(e_i) >$$

On a

$$X = < f(x + y), f(e_i) > - < f(x), f(e_i) > - < f(y), f(e_i) >$$

par bilinéarité du produit scalaire. Comme f conserve les produits scalaires, on a donc

$$X = < x + y, e_i > - < x, e_i > - < y, e_i > = < x + y - x - y, e_i > = 0$$

Le vecteur $f(x + y) - f(x) - f(y)$ est donc orthogonal à tous les vecteurs d'une base de E , et donc, par linéarité, à tous les vecteurs de E . Ainsi, $f(x + y) = f(x) + f(y)$. On démontre de façon analogue que pour tout vecteur x de E et tout réel λ , on a $f(\lambda x) = \lambda f(x)$.

Exemple. Les symétries orthogonales sont un exemple important. Soit F un sous-espace vectoriel de E . Soit s la symétrie par rapport à F , parallèlement à $F^\perp$. Alors, s est un endomorphisme orthogonal de E . En effet, soient $x = x' + x''$ et $y = y' + y''$ deux vecteurs de E décomposés suivant $F \oplus F^\perp$. On a $f(x) = x' - x''$ et $f(y) = y' - y''$. D'où

$$< f(x), f(y) > = < x', y' > + < x'', y'' >$$

car les produits scalaires des « primes » par les « secondes » sont nuls. Mais cette quantité est aussi $< x, y >$.

Définition 2. Une symétrie orthogonale par rapport à un hyperplan est appelée une réflexion. Une symétrie orthogonale par rapport à un sous-espace vectoriel de dimension $n - 2$ (où n est la dimension de E) est appelée un demi-tour.

Remarque 2. Le déterminant d'une réflexion vaut -1 . Celui d'un demi-tour est égal à 1 .

Notation. On note $\mathcal{O}(E)$ l'ensemble des endomorphismes orthogonaux de E .

1.2 Propriétés élémentaires

Proposition 1. L'ensemble $\mathcal{O}(E)$, muni de la loi $\circ$, est un groupe, sous-groupe de $GL(E)$.

Démonstration. On a vu un peu plus haut qu'une base orthonormée était transformée par un endomorphisme orthogonal en une base orthonormée. Un endomorphisme orthogonal est donc un automorphisme. De plus, id_E est clairement dans $\mathcal{O}(E)$. Enfin, $\mathcal{O}(E)$ est stable par composition et réciproque. $\square$

Proposition 2. Soit $f \in \mathcal{L}(E)$. Alors, $f \in \mathcal{O}(E)$ si et seulement si f conserve la norme euclidienne :

$$\forall x \in E, \|f(x)\| = \|x\|$$

Démonstration. Si f conserve les produits scalaires, elle conserve évidemment la norme euclidienne. Inversement, supposons que f conserve la norme. On a pour tous vecteurs $x, y \in E$

$$\langle f(x), f(y) \rangle = \frac{1}{4}(\|f(x) + f(y)\|^2 - \|f(x) - f(y)\|^2)$$

La linéarité de f et la conservation de la norme permettent de conclure. $\square$

Proposition 3. Soit $f \in \mathcal{O}(E)$. On a pour tous $x, y \in E$,

$$d(f(x), f(y)) = d(x, y)$$

Les endomorphismes orthogonaux sont des isométries (vectorielles).

Démonstration. Une distance est la norme d'une différence, et les endomorphismes orthogonaux conservent la norme. $\square$

Remarque 3. La réciproque est fautive : par exemple, les translations conservent les distances mais ne sont pas linéaires. On peut montrer en revanche que si f conserve les distances et $f(0) = 0$, alors $f \in \mathcal{O}(E)$. Nous laissons ce résultat en exercice.

Proposition 4. On rappelle que l'écart angulaire entre deux vecteurs non nuls x et y est l'unique réel $\theta \in [0, \pi]$ tel que

$$\cos \theta = \frac{\langle x, y \rangle}{\|x\| \|y\|}$$

Les endomorphismes orthogonaux conservent les écarts angulaires.

Démonstration. Un endomorphisme orthogonal conserve le produit scalaire et la norme. $\square$

1.3 Endomorphismes orthogonaux et bases orthonormées

Proposition 5. Soit $\mathcal{B}$ une base orthonormée de E . Soit $f \in \mathcal{L}(E)$. Alors, $f \in \mathcal{O}(E)$ si et seulement si $f(\mathcal{B})$ est une base orthonormée de E .

Démonstration. Si $f \in \mathcal{O}(E)$, alors f conserve les produits scalaires, donc les vecteurs orthogonaux. f conserve également la norme, donc les vecteurs unitaires. Donc f conserve les bases orthonormées.

Inversement, supposons que $f(\mathcal{B}) = \mathcal{B}'$ où $\mathcal{B}$ et $\mathcal{B}'$ sont deux bases orthonormées de E . Avec des notations évidentes, on a alors pour tous indices i et j ,

$$\langle f(e_i), f(e_j) \rangle = \delta_{ij} = \langle e_i, e_j \rangle$$

Mais cela entraîne que pour tous vecteurs x et y de E ,

$$\begin{aligned} \langle f(x), f(y) \rangle &= \sum_{i,j} x_i y_j \langle f(e_i), f(e_j) \rangle \\ &= \sum_{i,j} x_i y_j \langle e_i, e_j \rangle \\ &= \langle x, y \rangle \end{aligned}$$

L'application f conserve bien les produits scalaires. $\square$

1.4 Matrices orthogonales

Proposition 6. Soit $f \in \mathcal{L}(E)$. Soit $\mathcal{B}$ une base orthonormée de E . Soit $M = \text{mat}_{\mathcal{B}}f$. Les assertions suivantes sont équivalentes :

- ${}^tMM = I$
- $M {}^tM = I$
- M est inversible et $M^{-1} = {}^tM$
- Les colonnes de M , vues comme vecteurs de $\mathbb{R}^n$, forment une base orthonormée de $\mathbb{R}^n$.
- $f \in \mathcal{O}(E)$

Démonstration. L'équivalence des trois premiers points est un résultat connu : une matrice carrée est inversible à gauche si et seulement si elle est inversible à droite, et inverse à gauche et à droite sont alors uniques et égaux (à l'inverse de la matrice).

Le premier et le quatrième point sont clairement équivalents : à la ligne i , colonne j de tMM , on trouve justement le produit scalaire des colonnes i et j de la matrice M .

Enfin, le quatrième et le cinquième point sont équivalents, puisque f est un endomorphisme orthogonal si et seulement si $f(\mathcal{B})$ est une base orthonormée de E , c'est-à-dire si et seulement si les colonnes de M sont unitaires et orthogonales deux à deux. $\square$

Proposition 7. On appelle matrice orthogonale toute matrice $M \in \mathcal{M}_n(\mathbb{R})$ telle que ${}^tMM = I$ ou, ce qui est équivalent, $M {}^tM = I$, ou encore M est inversible et $M^{-1} = {}^tM$.

Notation. On note $\mathcal{O}_n(\mathbb{R})$ l'ensemble des matrices orthogonales $n \times n$.

Exemple. Fabriquons une matrice orthogonale 3×3 . Cela revient à fabriquer une base orthonormée de $\mathbb{R}^3$ pour le produit scalaire canonique. Par exemple, la matrice

$$M = \frac{1}{4} \begin{pmatrix} 3 & 1 & \sqrt{6} \\ 1 & 3 & -\sqrt{6} \\ -\sqrt{6} & \sqrt{6} & 2 \end{pmatrix}$$

convient car ses colonnes sont de norme 1, et orthogonales deux à deux (le vérifier!).

Proposition 8. L'ensemble $\mathcal{O}_n(\mathbb{R})$ est un groupe pour la multiplication matricielle, sous-groupe de $GL_n(\mathbb{R})$. Ce groupe est isomorphe à $\mathcal{O}(E)$ lorsque $n = \dim E$.

Démonstration. Soit $\mathcal{B}$ une base orthonormée de E . L'application $\Phi : \mathcal{O}(E) \rightarrow \mathcal{O}_n(\mathbb{R})$ définie par $\Phi(f) = \text{mat}_{\mathcal{B}}f$ est un isomorphisme de groupes. $\square$

Proposition 9. Soit $M \in \mathcal{O}_n(\mathbb{R})$. Alors, $\det M = \pm 1$.

Démonstration. C'est une conséquence immédiate de ${}^tMM = I$. $\square$

Proposition 10. Les ensembles $\mathcal{SO}(E)$ et $\mathcal{SO}_n(\mathbb{R})$ formés respectivement des endomorphismes orthogonaux de déterminant 1 et des matrices orthogonales de déterminant 1 sont des groupes. Ces deux groupes sont isomorphes. On les appelle respectivement le groupe spécial-orthogonal de E et le groupe spécial-orthogonal d'ordre n .

Démonstration. Ce sont des groupes, car ce sont les noyaux du morphisme « déterminant ». Ils sont isomorphes par le même isomorphisme que celui vu un peu avant. $\square$

Définition 3. On appelle rotations les éléments de $\mathcal{SO}(E)$.

Remarque 4. On ne voit pas pourquoi. C'est normal. Le reste du chapitre consiste à essayer de s'en faire une idée en dimensions 2 et 3.

1.5 Produit mixte

E désigne un espace euclidien orienté de dimension n .

Proposition 11. Soient $\mathcal{B}, \mathcal{B}'$ deux bases de E . Soit $P = P_{\mathcal{B}}^{\mathcal{B}'}$ la matrice de passage de $\mathcal{B}$ à $\mathcal{B}'$. Soient $x_1, \dots, x_n$ n vecteurs de E . Soient $U = \text{mat}_{\mathcal{B}}(x_1, \dots, x_n)$ et $U' = \text{mat}_{\mathcal{B}'}(x_1, \dots, x_n)$. On a

$$U = PU'$$

Démonstration. Soient X_1, X'_1 les matrices respectives de x_1 dans les bases $\mathcal{B}$ et $\mathcal{B}'$. On a $X_1 = PX'_1$. Ainsi, la première colonne de la matrice U est la première colonne de la matrice PU' . Il en est de même pour les $n - 1$ autres colonnes. $\square$

Corollaire 12. Avec les notations ci-dessus, on a

$$\det_{\mathcal{B}}(x_1, \dots, x_n) = \det P \det_{\mathcal{B}'}(x_1, \dots, x_n)$$

Corollaire 13. *Le déterminant $\det_{\mathcal{B}}(x_1, \dots, x_n)$ est le même dans toutes les bases orthonormées directes $\mathcal{B}$.*

Démonstration. La matrice de passage d'une b.o.n.d à une autre b.o.n.d est une matrice de rotation. Son déterminant est donc égal à 1. $\square$

Définition 4. On appelle produit mixte des vecteurs $x_1, \dots, x_n$, et on note $[x_1, \dots, x_n]$ le déterminant des vecteurs $x_1, \dots, x_n$ dans n'importe quelle b.o.n.d de E . Ainsi,

$$[x_1, \dots, x_n] = \det_{\mathcal{B}}(x_1, \dots, x_n)$$

où $\mathcal{B}$ est n'importe quelle b.o.n.d. de E .

Exemple. Soit $(e_1, \dots, e_n)$ une b.o.n.d de E . On a $[e_1, \dots, e_n] = 1$. En effet, c'est le déterminant de la base dans n'importe quelle b.o.n.d, par exemple elle-même. En revanche, le produit mixte des vecteurs d'une b.o.n.i vaut -1 .

2 Le groupe orthogonal du plan

On se donne dans cette section $f \in \mathcal{O}(E)$, où E est un plan euclidien orienté. On suppose fixée une base orthonormée $\mathcal{B}$ de E . On appelle M la matrice de f dans cette base. La matrice M est donc orthogonale. Posons $M = \begin{pmatrix} a & c \\ b & d \end{pmatrix}$.

2.1 Recherche des matrices orthogonales

On écrit ${}^tMM = I$. Cela fournit le système

$$\begin{cases} a^2 + b^2 &= 1 \\ ac + bd &= 0 \\ c^2 + d^2 &= 1 \end{cases}$$

Les deux premières égalités nous disent qu'il existe des réels θ et φ tels que $a = \cos \theta$, $b = \sin \theta$, $c = \cos \varphi$, $d = \sin \varphi$. La seconde égalité montre qu'alors $\cos(\varphi - \theta) = 0$. Deux cas se présentent :

- Si $\varphi \equiv \theta + \pi/2 \pmod{2\pi}$, alors

$$M = \begin{pmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{pmatrix}$$

On voit qu'alors $\det M = 1$, c'est à dire que $M \in \text{SO}_2(\mathbb{R})$ et $f \in \text{SO}(E)$.

- Si $\varphi \equiv \theta - \pi/2 \pmod{2\pi}$, alors

$$M = \begin{pmatrix} \cos \theta & \sin \theta \\ \sin \theta & -\cos \theta \end{pmatrix}$$

On voit qu'alors $\det M = -1$.

2.2 Étude des endomorphismes orthogonaux de déterminant -1

On a ici ${}^tM = M$. Comme ${}^tMM = I$, il vient donc $M^2 = I$, d'où $f^2 = id$. f est donc une symétrie. On vérifie sans peine que f est la symétrie orthogonale (la réflexion) par rapport à la droite dirigée par le vecteur de coordonnées $(\cos \frac{\theta}{2}, \sin \frac{\theta}{2})$ dans la base $\mathcal{B}$. Il suffit de vérifier que ce vecteur est invariant par f est que le vecteur $(-\sin \frac{\theta}{2}, \cos \frac{\theta}{2})$, orthogonal au premier, est changé par f en son opposé.

2.3 Étude des endomorphismes orthogonaux de déterminant 1

Ici, f est une rotation. Notons $M(\theta) = \begin{pmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{pmatrix}$. On a donc

$$\mathcal{SO}_2(\mathbb{R}) = \{M(\theta), \theta \in \mathbb{R}\}$$

Proposition 14. *L'application $\varphi : \mathbb{R} \rightarrow \mathcal{SO}_2(\mathbb{R})$ définie par $\varphi(\theta) = M(\theta)$ est un morphisme de groupes surjectif.*

Démonstration. Un simple calcul montre que pour tous $\theta, \theta' \in \mathbb{R}$,

$$M(\theta)M(\theta') = M(\theta + \theta')$$

D'où la propriété de morphisme. Cette application est surjective par définition même de $\mathcal{SO}_2(\mathbb{R})$. Remarquons enfin que $M(\theta) = I$ si et seulement si $\theta \in 2\pi\mathbb{Z}$. Notre morphisme a donc l'ensemble $2\pi\mathbb{Z}$ pour noyau, et n'est donc pas un isomorphisme. $\square$

Si l'on veut fabriquer un isomorphisme entre $\mathcal{SO}_2(\mathbb{R})$ et un groupe connu, il nous faut quelque chose qui représenterait les réels « modulo 2π ». Cela existe : c'est par exemple le groupe $\mathcal{U}$ des nombres complexes de module 1.

Proposition 15. *L'application $\varphi : \mathcal{U} \rightarrow \mathcal{SO}_2(\mathbb{R})$ définie par $\varphi(e^{i\theta}) = M(\theta)$ est un isomorphisme de groupes.*

Démonstration. Laissée en exercice. $\square$

Proposition 16. *Soit f une rotation de E . La matrice de f ne dépend pas de la base orthonormée directe choisie. Elle est de la forme $M(\theta)$, où $\theta \in \mathbb{R}$ est défini de façon unique modulo 2π .*

Démonstration. Soient $\mathcal{B}$ et $\mathcal{B}'$ deux bases de E . La matrice de f dans la base $\mathcal{B}$ est $M = M(\theta)$, et celle de f dans la base $\mathcal{B}'$ est $M' = M(\theta')$. La matrice de passage de $\mathcal{B}$ à $\mathcal{B}'$ est une matrice de rotation, puisque, ces deux bases ont la même orientation. Elle est donc de la forme $P = M(\varphi)$ où $\varphi \in \mathbb{R}$. Les formules de changement de base donnent alors

$$\begin{aligned} M' &= P^{-1}MP \\ &= M(-\varphi)M(\theta)M(\varphi) \\ &= M(-\varphi + \theta + \varphi) \\ &= M \end{aligned}$$

□

Définition 5. Le réel θ est appelé l'angle de la rotation f , et caractérise complètement cette rotation.

Exemple. Quelle est la matrice de la rotation de E d'angle $\frac{\pi}{3}$? C'est

$$\begin{pmatrix} \frac{1}{2} & -\frac{\sqrt{3}}{2} \\ \frac{\sqrt{3}}{2} & \frac{1}{2} \end{pmatrix}$$

2.4 Angle orienté de deux vecteurs non nuls

Proposition 17. *Soient x et y deux vecteurs non nuls de E . Il existe une unique rotation f telle que*

$$f\left(\frac{x}{\|x\|}\right) = \frac{y}{\|y\|}$$

L'angle de cette rotation (défini modulo 2π) est appelé angle (orienté) entre les vecteurs x et y .

Démonstration. On peut supposer x et y unitaires pour faire la démonstration. Fixons une base de E . Dans cette base, la matrice de x est $X = \begin{pmatrix} \cos \alpha \\ \sin \alpha \end{pmatrix}$ et celle de y est $Y = \begin{pmatrix} \cos \beta \\ \sin \beta \end{pmatrix}$. Soit f une rotation, de matrice $M(\theta)$ dans la même base. On a $f(x) = y$ si et seulement si $M(\theta)X = Y$ ce qui équivaut (calcul immédiat) à $\theta \equiv \beta - \alpha \pmod{2\pi}$. D'où l'existence et l'unicité de f . □

Proposition 18. Soient x et y deux vecteurs non nuls du plan. Soit θ l'angle orienté entre x et y . On a

$$\begin{cases} \langle x, y \rangle &= \|x\| \|y\| \cos \theta \\ [x, y] &= \|x\| \|y\| \sin \theta \end{cases}$$

Démonstration. Immédiat, en reprenant les matrices de x et y ci-dessus. $\square$

Exemple. Dans $\mathbb{R}^2$ euclidien, orienté, soient $u = (1, 2)$ et $v = (3, 4)$. Soit θ l'angle orienté entre u et v . On a

$$\|u\| = \sqrt{5}, \|v\| = 5, \langle u, v \rangle = 11, [u, v] = -2$$

On en déduit par la proposition précédente que

$$\begin{cases} 5\sqrt{5} \cos \theta &= 11 \\ 5\sqrt{5} \sin \theta &= -2 \end{cases}$$

De la première égalité on obtient que

$$\theta \equiv \pm \arccos \frac{11}{5\sqrt{5}} \pmod{2\pi}$$

La deuxième inégalité nous dit que $\sin \theta < 0$, donc

$$\theta \equiv -\arccos \frac{11}{5\sqrt{5}} \pmod{2\pi}$$

2.5 Générateurs du groupe orthogonal

Proposition 19. Les réflexions engendrent $\mathcal{O}(E)$.

Démonstration. Le produit de deux matrices de réflexion est évidemment une matrice de rotation : en effet, deux matrices de déterminant -1 se multiplient en une matrice de déterminant 1. Mieux, en effectuant effectivement ce calcul, on voit que l'on peut obtenir toutes les matrices de rotation. On peut même choisir l'une des deux matrices de réflexion de façon complètement arbitraire.

Par exemple, on a pour tout réel θ

$$\begin{pmatrix} \cos \theta & \sin \theta \\ \sin \theta & -\cos \theta \end{pmatrix} \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} = \begin{pmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{pmatrix}$$

$\square$

3 Produit vectoriel

Dans toute cette section, E désigne un espace euclidien orienté de dimension 3.

3.1 Notion de produit vectoriel

Proposition 20. Soient $u, v \in E$. Il existe un unique vecteur $a \in E$ tel que

$$\forall w \in E, [u, v, w] = \langle a, w \rangle$$

Démonstration. L'application $w \mapsto [u, v, w]$ est une forme linéaire sur E . Et on sait que l'application $a \mapsto \langle a, \cdot \rangle$ est un isomorphisme de E sur son dual E^* . $\square$

Définition 6. Le vecteur a est appelé produit vectoriel de u et v et est noté $u \wedge v$. Il est donc caractérisé par

$$\forall w \in E, [u, v, w] = \langle u \wedge v, w \rangle$$

Proposition 21. L'application $(u, v) \mapsto u \wedge v$ est bilinéaire, antisymétrique, alternée.

Démonstration. Pour tous vecteurs u, u', v, w , on a

$$\begin{aligned} \langle (u + u') \wedge v, w \rangle &= [u + u', v, w] \\ &= [u, v, w] + [u', v, w] \\ &= \langle u \wedge v, w \rangle + \langle u' \wedge v, w \rangle \\ &= \langle u \wedge v + u' \wedge v, w \rangle \end{aligned}$$

D'où

$$(u + u') \wedge v = u \wedge v + u' \wedge v$$

De même pour toutes les autres propriétés de la bilinéarité.

Pour l'antisymétrie, $[v, u, w] = -[u, v, w]$ donc $\langle v \wedge u, w \rangle = \langle -u \wedge v, w \rangle$ d'où $v \wedge u = -u \wedge v$. Pour l'alternance, prendre $v = u$ dans l'antisymétrie. $\square$

3.2 Propriétés essentielles

Proposition 22. Soient $u, v \in E$. On a $u \wedge v = 0$ si et seulement si u et v sont colinéaires.

Démonstration. $u \wedge v = 0$ si et seulement si $u \wedge v$ est orthogonal à tous les vecteurs w de E c'est à dire $\forall w \in E, [u, v, w] = 0$ ou encore $\forall w \in E, (u, v, w)$ est liée. Si (u, v) est liée, c'est bien le cas. Si (u, v) est libre, le théorème de la base incomplète fournit un vecteur w tel que (u, v, w) soit une base de E , et donc $[u, v, w] \neq 0$. $\square$

Proposition 23. Soient $u, v \in E$. Le vecteur $u \wedge v$ est orthogonal à u et à v .

Démonstration. $\langle u \wedge v, u \rangle = [u, v, u] = 0$ et de même pour v . $\square$

Proposition 24. Soit (e_1, e_2, e_3) une b.o.n.d. de E . On a

$$e_1 \wedge e_2 = e_3, e_2 \wedge e_3 = e_1, e_3 \wedge e_1 = e_2$$

Démonstration. Le vecteur $e_1 \wedge e_2$ est orthogonal à e_1 et e_2 . Il appartient donc à l'orthogonal de $\text{Vect}(e_1, e_2) = \text{Vect}(e_3)$. Il existe ainsi un réel λ tel que $e_1 \wedge e_2 = \lambda e_3$. De là,

$$[e_1, e_2, e_3] = 1 = \langle e_1 \wedge e_2, e_3 \rangle = \lambda \|e_3\|^2 = \lambda$$

Les autres relations se montrent de façon analogue. $\square$

Remarque 5. Par antisymétrie, $e_2 \wedge e_1 = -e_3$ et relations analogues. Si la base est orthonormée indirecte, changer tous les signes.

Remarque 6. Soient $u, v \in E$ deux vecteurs non colinéaires. Soit $a = u \wedge v$. Le vecteur a est non nul et orthogonal à u et v . De plus,

$$[u, v, a] = \langle u \wedge v, a \rangle = \langle a, a \rangle > 0$$

Ainsi, la famille (u, v, a) est une base (pas orthogonale en général!) directe puisque son déterminant est strictement positif. Pour caractériser complètement a , il nous faudrait sa norme. C'est ce que nous allons faire.

Proposition 25. Soient u et v deux vecteurs non nuls. Soit θ l'écart angulaire entre u et v . On a

$$\|u \wedge v\| = \|u\| \|v\| \sin \theta$$

Démonstration. Supposons tout d'abord u et v unitaires. Si u et v sont liés, l'égalité est évidente. Supposons-les donc libres. Soit

$$u_1 = v - \langle v, u \rangle u = v - \cos \theta u$$

(Gram-Schmidt). Les vecteurs u et u_1 sont orthogonaux et non nuls. On a

$$\|u_1\|^2 = \|v\|^2 + \cos^2 \theta \|u\|^2 - 2 \cos \theta \langle u, v \rangle = 1 - \cos^2 \theta = \sin^2 \theta$$

Donc

$$\|u_1\| = \sin \theta$$

Posons

$$u_1 = \sin \theta u'$$

de sorte que u' est unitaire et orthogonal à u . On a donc

$$v = \cos \theta u + \sin \theta u'$$

Soit enfin u'' tel que (u, u', u'') soit une base orthonormée directe. On a alors

$$u \wedge v = u \wedge (\cos \theta u + \sin \theta u') = \cos \theta u \wedge u + \sin \theta u \wedge u' = \sin \theta u''$$

et on en tire

$$\|u \wedge v\| = \sin \theta$$

Lorsque u et v ne sont pas forcément unitaires, il suffit de poser $u = \|u\|u_1$, $v = \|v\|v_1$. Les vecteurs u_1 et v_1 sont unitaires, et de même écart angulaire que u et v . On a alors

$$\|u \wedge v\| = \|u\| \|v\| \|u_1 \wedge v_1\| = \|u\| \|v\| \sin \theta$$

□

Corollaire 26. Pour tous vecteurs $u, v \in E$, on a

$$\langle u, v \rangle^2 + \|u \wedge v\|^2 = \|u\|^2 \|v\|^2$$

Démonstration. Si u ou v est nul, c'est évident. Sinon, il suffit d'écrire

$$\langle u, v \rangle = \|u\| \|v\| \cos \theta$$

où θ est l'écart angulaire entre u et v . □

3.3 Calculs en base orthonormée directe

Soit $\mathcal{B} = (e_1, e_2, e_3)$ une base orthonormée directe de E . Soient u, v deux vecteurs de E , de coordonnées u_i et v_i dans la base $\mathcal{B}$. Calculons les coordonnées de $u \wedge v$. On a

$$\begin{aligned} u \wedge v &= \langle u_1 e_1 + u_2 e_2 + u_3 e_3, v_1 e_1 + v_2 e_2 + v_3 e_3 \rangle \\ &= (u_2 v_3 - u_3 v_2) e_1 + (u_3 v_1 - u_1 v_3) e_2 + (u_1 v_2 - u_2 v_1) e_3 \end{aligned}$$

La formule n'est pas très jolie, un bon moyen de s'en souvenir est de la voir comme un « déterminant » 3×3 dont la troisième colonne contient des vecteurs. Plus précisément, on a la formule formelle

$$u \wedge v = \begin{vmatrix} u_1 & v_1 & e_1 \\ u_2 & v_2 & e_2 \\ u_3 & v_3 & e_3 \end{vmatrix}$$

Exemple. Dans $\mathbb{R}^3$ euclidien canonique orienté, le produit vectoriel de $(1, 2, 3)$ et $(4, 5, 6)$ est

$$\begin{vmatrix} 1 & 4 & e_1 \\ 2 & 5 & e_2 \\ 3 & 6 & e_3 \end{vmatrix} = (-3, 6, -3)$$

3.4 Annexe : double produit vectoriel

Proposition 27. Soient $u, v, w \in E$. On a

$$u \wedge (v \wedge w) = \langle u, w \rangle v - \langle u, v \rangle w$$

Démonstration. Écrire les vecteurs u, v, w dans une b.o.n.d et développer les deux membres de l'égalité. $\square$

Remarque 7. Le produit vectoriel n'est donc pas associatif.

Exercice. Soient $u, v, w \in E$. Montrer que

$$[u \wedge v, v \wedge w, w \wedge u] = [u, v, w]^2$$

On a

$$[u \wedge v, v \wedge w, w \wedge u] = \langle u \wedge v, (v \wedge w) \wedge (w \wedge u) \rangle$$

Par la formule du double produit vectoriel,

$$(v \wedge w) \wedge (w \wedge u) = \langle v \wedge w, u \rangle w - \langle v \wedge w, w \rangle u$$

Le produit scalaire $\langle v \wedge w, w \rangle$ est nul puisque $v \wedge w$ est orthogonal à w . Donc

$$\begin{aligned}(v \wedge w) \wedge (w \wedge u) &= \langle v \wedge w, u \rangle w \\ &= [v, w, u]w \\ &= [u, v, w]w\end{aligned}$$

De là,

$$\begin{aligned}[u \wedge v, v \wedge w, w \wedge u] &= \langle u \wedge v, [u, v, w]w \rangle \\ &= [u, v, w] \langle u \wedge v, w \rangle \\ &= [u, v, w]^2\end{aligned}$$

4 Le groupe orthogonal de l'espace

4.1 Sous espaces stables

Soit E un espace euclidien de dimension quelconque. Soit $f \in \mathcal{O}(E)$.

Valeurs propres de f

Soit $x \in E \setminus \{0\}$. Supposons qu'il existe un réel λ tel que $f(x) = \lambda x$. Alors, $\lambda = \pm 1$. En effet, $\|f(x)\| = \|\lambda\| \|x\|$, mais aussi $\|f(x)\| = \|x\|$ par conservation de la norme. Comme x est non nul, il vient $|\lambda| = 1$.

Droites stables par f

Supposons qu'il existe une droite D stable par f . Soit e un vecteur directeur de D . Comme $f(e) \in D$, il existe donc un réel λ tel que $f(e) = \lambda e$. Donc, $f(e) = \pm e$. Par linéarité, on a donc $f(x) = x$ pour tout $x \in D$, ou $f(x) = -x$ pour tout $x \in D$.

Ce qui vient d'être dit ne s'applique bien sûr qu'aux droites stables. Par exemple, dans un plan stable, il peut arriver qu'aucun vecteur (non nul) ne soit changé par f en un multiple de lui-même (exemple : une rotation du plan, tout simplement).

Orthogonal d'un sous-espace stable

Proposition 28. *Soit F un sous-espace de E stable par f . Alors,*

- *f induit sur F un endomorphisme orthogonal de F .*
- *$F^\perp$ est stable par f .*

Démonstration. Le premier point est évident.

Soit maintenant $x \in F^\perp$. Soit $y \in F$. On a

$$\langle f(x), f(y) \rangle = \langle x, y \rangle = 0$$

Donc, $f(x)$ appartient à l'orthogonal de $f(F)$. Mais F est stable par f , donc $f(F) \subset F$. De plus f est un isomorphisme, donc f conserve les dimensions, et donc $f(F) = F$. Finalement, $f(x) \in F^\perp$. Donc, $F^\perp$ est lui aussi stable par f . $\square$

bilan

L'idée pour étudier un endomorphisme orthogonal de E est donc la suivante : on cherche les vecteurs invariants de f (on pourrait tout aussi bien chercher les vecteurs changés en leur opposé). L'ensemble E_1 de ces vecteurs invariants est un sous-espace stable par f . Si ce sous-espace est non trivial, alors :

- f est l'identité sur E_1 .
- f est un endomorphisme orthogonal de $E_1^\perp$. Or, $\dim E_1^\perp < \dim E$. On s'est donc ramenés à l'étude des endomorphismes orthogonaux d'un espace de dimension strictement plus petite.

Dans ce qui suit, on se place dans un espace euclidien orienté de dimension 3.

4.2 Cas 1 : Les invariants forment tout l'espace

Ce cas est trivial : $f = id_E$.

4.3 Cas 2 : Les invariants forment un plan

Soit P le plan en question. La droite $D = P^\perp$ est alors stable par f . Les vecteurs de D sont donc changés en leur opposé par f (hormis le vecteur nul, ils ne peuvent pas être invariants, puisque les invariants sont tous dans P , et que $D \cap P = \{0\}$). On constate donc que f est la réflexion par rapport à P .

4.4 Cas 3 : Les invariants forment une droite

Soit D la droite en question. Le sous-espace $P = D^\perp$ est un plan stable par f , et f est un endomorphisme orthogonal de P . Ce ne peut être une réflexion de P , puisque tous les vecteurs invariants par f sont sur D . C'est donc une rotation de P (différente de l'identité).

Un petit problème se pose alors : il convient d'orienter P , mais il n'y a aucune façon « logique » d'effectuer cette orientation de manière compatible avec l'orientation de l'espace E . On procède comme suit :

- On oriente la droite D en choisissant sur D un vecteur unitaire e_3 (deux choix arbitraires sont donc possibles).
- Une base orthonormée (e_1, e_2) de P est déclarée directe lorsque la base (e_1, e_2, e_3) est une base orthonormée directe de E .

La restriction de f à P est alors la rotation d'angle $\theta \not\equiv 0 \pmod{2\pi}$.

On remarque que la matrice de f dans une base orthonormée directe de dernier vecteur e_3 est alors

$$M = \begin{pmatrix} \cos \theta & -\sin \theta & 0 \\ \sin \theta & \cos \theta & 0 \\ 0 & 0 & 1 \end{pmatrix}$$

On voit que $\det f = 1$: f est donc une rotation.

Définition 7. f est la rotation d'axe D orienté par e_3 , et d'angle θ . L'axe orienté et l'angle caractérisent complètement f .

Remarque 8. Si l'on change l'orientation de D , on change θ en $-\theta$.

4.5 Cas 4 : Le seul invariant est 0

Là, on est un peu embêtés. Cependant, on peut s'en tirer par une petite acrobatie. Soit λ un réel. Il existe un vecteur x non nul tel que $f(x) = \lambda x$ si et seulement si $\det(f - \lambda \text{id}_E) = 0$. Or l'application $\lambda \mapsto \det(f - \lambda \text{id}_E)$ est une fonction polynomiale de degré 3. Or, tout polynôme de degré 3 a une racine réelle (théorème des valeurs intermédiaires). Il existe donc un vecteur non nul x et un réel λ tel que $f(x) = \lambda x$. Ce λ vaut -1 ou 1 , mais ce n'est pas 1 . C'est donc -1 .

Ainsi, l'ensemble G des vecteurs changés en leur opposé est un sous-espace vectoriel non trivial de E . Il peut être de dimension 1, 2 ou 3. Considérons chacun des trois cas :

- $G = E$: dans ce cas, $f = -\text{id}_E$. Ce n'est pas une rotation.
- G est un plan. C'est impossible, car alors $G^\perp$ serait une droite stable, nécessairement composée de vecteurs invariants. Mais f n'a pas d'invariants à part le vecteur nul.
- $G = D$ est une droite. Soit P l'orthogonal de D . Soit σ la réflexion par rapport à P .

On constate alors que dans ce cas $f \circ \sigma$ admet G pour droite d'invariants. Donc, $f \circ \sigma$ est une rotation ρ d'axe D . De là, $f = \rho \circ \sigma$. Et pour finir, $\det f = \det \rho \det \sigma = -1$. Donc f n'est pas une rotation.

Remarque 9. On vérifie facilement dans le dernier cas que $\rho \circ \sigma = \sigma \circ \rho$. On peut également montrer l'unicité des applications σ et ρ .

4.6 Conclusion

Proposition 29. Les endomorphismes orthogonaux de E de dimension 3 sont :

- Les rotations qui, mise à part l'identité, sont caractérisées par un axe et un angle.
- Les réflexions, caractérisées par un plan.
- Les composées d'une réflexion et d'une rotation différente de l'identité d'axe orthogonal au plan de la réflexion, caractérisées par un axe orienté et un angle.

Exercice. Vérifier que $-id_E$ entre bien la dernière catégorie.

4.7 Générateurs de $\mathcal{O}(E)$ et de $\mathcal{SO}(E)$

Nous montrons dans cette section que les réflexions engendrent $\mathcal{O}(E)$, c'est à dire que tout endomorphisme orthogonal est une composée de réflexions. Ainsi, les réflexions sont des endomorphismes orthogonaux élémentaires, en ce sens que l'on peut avec celles-ci fabriquer tous les autres. Nous donnons une démonstration purement matricielle. Elle a l'avantage d'être rapide, mais de cacher le sens géométrique de ce théorème. Il est conseillé de faire l'exercice qui suit la proposition.

Proposition 30. *Les réflexions engendrent $\mathcal{O}(E)$.*

Démonstration. Nous allons démontrer que toute rotation est la composée de deux réflexions. Comme les endomorphismes orthogonaux de E sont les rotations et les composées d'une réflexion et d'une rotation, cela prouvera le théorème. Soit f une rotation différente de l'identité, d'angle θ . Dans une base $\mathcal{B}$ orthonormée directe judicieuse, la matrice de f est

$$M = \begin{pmatrix} \cos \theta & -\sin \theta & 0 \\ \sin \theta & \cos \theta & 0 \\ 0 & 0 & 1 \end{pmatrix}$$

Soit

$$A = \begin{pmatrix} \cos \theta & \sin \theta & 0 \\ \sin \theta & -\cos \theta & 0 \\ 0 & 0 & 1 \end{pmatrix}$$

A est une matrice orthogonale et $A^2 = I$. C'est donc une matrice de symétrie orthogonale. De plus, $\det A = -1$ et $A \neq -I$, c'est donc la matrice d'une réflexion. De même,

$$B = \begin{pmatrix} 1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & 1 \end{pmatrix}$$

est une matrice de réflexion. On a $AB = M$, ce qui prouve notre assertion. $\square$

Exercice. Déterminer les caractéristiques des réflexions qui apparaissent dans la démonstration de la proposition ci-dessus.

Nous énonçons et démontrons un théorème analogue au précédent pour $\mathcal{SO}(E)$, mais cette fois-ci avec les demi-tours.

Proposition 31. *Les demi-tours engendrent $\mathcal{SO}(E)$.*

Démonstration. Nous allons démontrer que toute rotation est la composée de deux demi-tours. Soit f une rotation différente de l'identité, d'angle θ . Dans une base $\mathcal{B}$ orthonormée directe judicieuse, la matrice de f est

$$M = \begin{pmatrix} \cos \theta & -\sin \theta & 0 \\ \sin \theta & \cos \theta & 0 \\ 0 & 0 & 1 \end{pmatrix}$$

Soit

$$A = \begin{pmatrix} \cos \theta & \sin \theta & 0 \\ \sin \theta & -\cos \theta & 0 \\ 0 & 0 & -1 \end{pmatrix}$$

A est une matrice orthogonale et $A^2 = I$, donc A est une matrice de symétrie orthogonale. De plus, $\det A = 1$ et $A \neq I$, c'est donc la matrice d'un demi-tour. De même,

$$B = \begin{pmatrix} 1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & -1 \end{pmatrix}$$

est une matrice de demi-tour. On a $AB = M$, ce qui prouve notre assertion. $\square$

Exercice. Déterminer les caractéristiques des demi-tours qui apparaissent dans la démonstration de la proposition ci-dessus.

5 Compléments

5.1 Écart angulaire entre un vecteur et son image par une rotation

Soit f une rotation de l'espace d'angle θ . Dans une base $\mathcal{B} = (e_1, e_2, e_3)$ judicieusement choisie, la matrice de f est

$$M = \begin{pmatrix} \cos \theta & -\sin \theta & 0 \\ \sin \theta & \cos \theta & 0 \\ 0 & 0 & 1 \end{pmatrix}$$

Soit $x = ae_1 + be_2 + ce_3$ un vecteur *unitaire* de E . On a

$$f(x) = (a \cos \theta - b \sin \theta)e_1 + (a \sin \theta + b \cos \theta)e_2 + ce_3$$

Ainsi,

$$\begin{aligned} \langle x, f(x) \rangle &= a(a \cos \theta - b \sin \theta) + b(a \sin \theta + b \cos \theta) + c^2 \\ &= (a^2 + b^2) \cos \theta + c^2 \\ &= \cos \theta + c^2(1 - \cos \theta) \end{aligned}$$

Les vecteurs x et $f(x)$ étant unitaires, cette quantité est égale à $\cos \varphi$ où φ est l'écart angulaire entre x et $f(x)$: $\cos \varphi = \cos \theta + c^2(1 - \cos \theta)$. On constate que $\cos \varphi \geq \cos \theta$, c'est à dire $\varphi \leq \theta$, avec égalité si et seulement si $c = 0$, c'est à dire si et seulement si le vecteur x est orthogonal à l'axe de la rotation.

5.2 La formule d'Euler-Rodrigues

On reprend

$$\begin{aligned} f(x) &= (a \cos \theta - b \sin \theta)e_1 + (a \sin \theta + b \cos \theta)e_2 + ce_3 \\ &= \cos \theta(ae_1 + be_2) + \sin \theta(-be_1 + ae_2) + ce_3 \end{aligned}$$

Par ailleurs, $ae_1 + be_2 = x - ce_3$, $-be_1 + ae_2 = e_3 \wedge x$, et $c = \langle x, e_3 \rangle$. Donc,

$$f(x) = \cos \theta(x - \langle e_3, x \rangle e_3) + \sin \theta e_3 \wedge x + \langle e_3, x \rangle e_3$$

Appelons ω le vecteur e_3 (vecteur unitaire orientant l'axe). On obtient

Proposition 32. Soit f la rotation d'axe orienté par le vecteur unitaire ω , d'angle θ . On a pour tout $x \in E$,

$$f(x) = (\cos \theta)x + (\sin \theta)\omega \wedge x + (1 - \cos \theta)\langle \omega, x \rangle \omega$$

Cette formule s'appelle la formule d'Euler-Rodrigues.

Exemple. Calculons la matrice dans la base canonique de la rotation d'angle $\frac{\pi}{2}$ d'axe orienté par le vecteur $(1, 1, 1)$.

Dans la formule précédente, on prend donc $\theta = \frac{\pi}{2}$ et $\omega = \frac{1}{\sqrt{3}}(1, 1, 1)$. Ainsi, pour tout $u = (x, y, z) \in \mathbb{R}^3$,

$$\begin{aligned} f(u) &= \omega \wedge u + \langle \omega, u \rangle \omega \\ &= \frac{1}{\sqrt{3}}(z - y, x - z, y - x) + \frac{1}{3}(x + y + z)(1, 1, 1) \end{aligned}$$

On en déduit la matrice de f en mettant en colonnes les vecteurs $f(1, 0, 0)$, $f(0, 1, 0)$ et $f(0, 0, 1)$:

$$A = \begin{pmatrix} \frac{1}{3} & \frac{1}{3} - \frac{1}{\sqrt{3}} & \frac{1}{3} - \frac{1}{\sqrt{3}} \\ \frac{1}{3} + \frac{1}{\sqrt{3}} & \frac{1}{3} & \frac{1}{3} + \frac{1}{\sqrt{3}} \\ \frac{1}{3} - \frac{1}{\sqrt{3}} & \frac{1}{3} + \frac{1}{\sqrt{3}} & \frac{1}{3} \end{pmatrix}$$

5.3 Étude pratique d'un endomorphisme orthogonal

Soit E euclidien orienté de dimension 3. Soit $f \in \mathcal{L}(E)$. Soit $A = \text{mat}_{\mathcal{B}} f$ la matrice de f dans la base $\mathcal{B}$. Dans un grand nombre d'exercices, on demande d'étudier f . Comment faire? L'étude se fait en plusieurs étapes. Le plan qui suit est systématique (cela dit on peut parfois aller plus vite, mais bref ...). On suppose dans ce qui suit que $A \neq \pm I_3$, car sinon l'étude est évidemment immédiate.

1. f est-il un endomorphisme orthogonal? Pour cela il suffit de vérifier que $A \in \mathcal{O}_3(\mathbb{R})$. C'est facile, il suffit de constater que les colonnes de A sont de norme 1 et orthogonales. Un simple coup d'oeil suffit. Supposons dans la suite que c'est le cas.
2. De quel type est l'endomorphisme f ? Réflexion? Rotation? Composée de ces deux types? On recherche les invariants par f . Pour cela, on résout l'équation $AX = X$, où $X \in \mathcal{M}_{31}(\mathbb{R})$. Il y a trois possibilités.
 - (a) Les invariants sont un plan P . Alors f est la réflexion de plan P .
 - (b) Le seul invariant est le vecteur nul 0. Alors, pour aller au plus vite, $-f$ est une rotation différente de l'identité, et on est ramenés au dernier cas.
 - (c) Les invariants sont une droite D . C'est le cas intéressant : f est une rotation d'axe D . On passe à l'étape 3.
3. Dans cette étape, on suppose que f est une rotation d'axe D connu, d'angle θ à déterminer.
 - (a) On oriente D en choisissant un vecteur non nul e sur la droite D . On note dans ce qui suit $\omega = \frac{e}{\|e\|}$.
 - (b) Comment trouver $\cos \theta$? Assez curieusement, c'est immédiat. En effet, les matrices de f dans toutes les bases de E ont la même trace. Or, on sait qu'il existe une base de E dans laquelle la matrice de f est

$$\begin{pmatrix} \cos \theta & -\sin \theta & 0 \\ \sin \theta & \cos \theta & 0 \\ 0 & 0 & 1 \end{pmatrix}$$

La trace de cette matrice est $1 + 2 \cos \theta$. Ainsi,

$$1 + 2 \cos \theta = \text{Tr } A$$

- (c) Enfin, comment trouver $\sin \theta$? Connaissant déjà le cosinus, il nous suffit d'obtenir le signe du sinus. Pour cela, reprenons la formule d'Euler-Rodrigues, vue au paragraphe précédent. Soit u un vecteur de E . Calculons un produit mixte :

$$\begin{aligned} [u, f(u), e] &= [u, (\cos \theta)u + (\sin \theta)\omega \wedge u + (1 - \cos \theta) < \omega, u > \omega, e] \\ &= (\cos \theta)[u, u, e] + (\sin \theta)[u, \omega \wedge u, e] + (1 - \cos \theta) < \omega, u > [u, \omega, e] \end{aligned}$$

par trilinearité du produit mixte. Dans la somme ci-dessus, le premier et le dernier terme sont nuls, car le produit mixte est alterné. Ainsi,

$$\begin{aligned}[u, f(u), e] &= (\sin \theta)[u, \omega \wedge u, e] \\ &= (\sin \theta)\|e\|[u, \omega \wedge u, \omega] \\ &= (\sin \theta)\|e\|[\omega, u, \omega \wedge u]\end{aligned}$$

Si u n'est pas colinéaire à ω , le produit mixte ci-dessus est strictement positif. On en déduit la recette suivante.

Proposition 33. *Le signe de $\sin \theta$ est le signe de $[u, f(u), e]$ où u est n'importe quel vecteur non colinéaire à e .*

5.4 Un exemple

Soit f l'endomorphisme de $\mathbb{R}^3$ dont la matrice dans la base canonique est

$$A = \frac{1}{9} \begin{pmatrix} 8 & 1 & -4 \\ -4 & 4 & -7 \\ 1 & 8 & 4 \end{pmatrix}$$

Suivons le plan ...

1. Les colonnes de A sont unitaires et orthogonales (immédiat), donc $A \in \mathcal{O}_3(\mathbb{R})$ et $f \in \mathcal{O}(\mathbb{R}^3)$.
2. Cherchons les invariants par f en résolvant le système $AX = X$. Ce système, après multiplication par 9 et simplifications mineures, s'écrit

$$\begin{cases} -x & +y & -4z & = & 0 \\ 4x & +5y & +7z & = & 0 \\ x & +8y & -5z & = & 0 \end{cases}$$

Sa résolution ne pose aucune difficulté. On obtient que les invariants par f sont la droite $D = \langle e \rangle$ où $e = (3, -1, -1)$. f est donc une rotation d'axe D .

- (a) Décidons d'orienter D par le vecteur e . Il reste à déterminer l'angle θ de la rotation f .
- (b) On a $1 + 2 \cos \theta = \text{Tr } A = \frac{16}{9}$. Ainsi, $\cos \theta = \frac{7}{18}$ et donc

$$\theta \equiv \pm \arccos \frac{7}{18} [2\pi]$$

- (c) Déterminons le signe de $\sin \theta$. Pour cela, prenons un vecteur non colinéaire à e : $u = (1, 0, 0)$ fera l'affaire ! On lit $f(u) = \frac{1}{9}(8, -4, 1)$ dans la première colonne de

A. Le signe de $\sin \theta$ est celui du produit mixte

$$[u, f(u), e] = \frac{1}{9} \begin{vmatrix} 1 & 8 & 3 \\ 0 & -4 & -1 \\ 0 & 1 & -1 \end{vmatrix} = \frac{5}{9} > 0$$

Ainsi, $\sin \theta > 0$.

Pour conclure, f est la rotation d'axe D dirigé par $(3, -1, -1)$ et d'angle $\arccos \frac{7}{18}$ modulo 2π .

Remarque. Il aurait été tout à fait possible d'orienter l'axe D par le vecteur $(-3, 1, 1)$. Dans ce cas, l'angle de la rotation aurait été $-\arccos \frac{7}{18}$.

Chapitre 35

Séries

1 Généralités

1.1 Notion de série

Définition 1. Soit $(u_n)_{n \geq 0}$ une suite de nombres complexes.

- On appelle *série de terme général* u_n la suite $(S_n)_{n \geq 0}$ définie pour tout entier n par

$$S_n = \sum_{k=0}^n u_k$$

- On notera $\sum u_n$ (ou $\sum u_k$, ou ce que l'on veut comme indice) la série de terme général u_n .
- Pour $n \in \mathbb{N}$, les sommes S_n sont les sommes partielles de la série $\sum u_n$.
- Pour $n \in \mathbb{N}$, u_n est le n ième terme de la série $\sum u_n$.

Remarque 1. Comme pour les suites, les séries peuvent démarrer à un rang autre que 0. Pour les questions de convergence que nous allons bientôt aborder, les premiers termes de la série ne jouent aucun rôle. En revanche, pour le calcul de ce que nous appellerons la somme de la série, chaque terme importe.

L'inconvénient de la notation $\sum u_n$ est que l'on ne sait pas quel est le premier terme de la série. Il faudra parfois être plus précis.

Remarque 2. Il convient de faire très attention à la notation $\sum u_n$: ce n'est qu'une notation abstraite, avec laquelle il ne s'agit pas de faire des calculs. La soi-disant « somme » $\sum u_n$ est une fonction $\mathbb{N} \rightarrow \mathbb{R}$ ou $\mathbb{C}$.

$$\sum u_n = (S_n)_{n \geq 0}$$

Exemple. Les deux séries ci-dessous joueront un rôle important dans l'étude de la convergence d'une série.

- Soit $q \in \mathbb{C}$. On appelle série géométrique de raison q la série

$$\sum q^n$$

- Soit $\alpha \in \mathbb{C}$. On appelle série de Riemann de paramètre α la série

$$\sum \frac{1}{n^\alpha}$$

Remarque 3. Au cas où vous vous demanderiez ce que vaut 2^{3+i} , c'est

$$2^{3+i} = 2^3 2^i = 8e^{i \ln 2} = 8 \cos(\ln 2) + 8i \sin(\ln 2)$$

Définition 2. Soit $(u_n)_{n \geq 0}$ une suite de nombres complexes. On dit que la série $\sum u_n$ converge lorsque la suite de ses sommes partielles est convergente. Sinon, on dit que la série diverge. En cas de convergence, on appelle somme de la série, et on note

$$\sum_{n=0}^{\infty} u_n$$

la limite en question.

Ainsi,

$$\sum_{n=0}^{\infty} u_n = \lim_{N \rightarrow +\infty} \sum_{n=0}^N u_n$$

Remarque 4. Cette définition est en réalité un pléonasme puisque une série EST la suite de ses sommes partielles.

Remarque 5. On adapte évidemment la notation si la série démarre à un rang autre que 0.

Exercice. Soit $(u_n)_{n \geq 0}$ une suite complexe. Soit $N \in \mathbb{N}$. Montrer que la série $\sum u_n$ (avec $n \geq 0$) est convergente si et seulement si la série $\sum u_n$ (avec $n \geq N + 1$) est convergente et que, lorsqu'il y a convergence,

$$\sum_{n=0}^{\infty} u_n = \sum_{n=0}^N u_n + \sum_{n=N+1}^{\infty} u_n$$

En d'autres termes, le fait de converger ou de diverger ne dépend pas des premiers termes de la série.

Exemple. Soit $q \in \mathbb{C}$. Si $q = 1$,

$$\sum_{k=0}^n q^k = n + 1$$

qui tend vers $+\infty$ lorsque n tend vers l'infini. Si $q \neq 1$, on a

$$\sum_{k=0}^n q^k = \frac{1 - q^{n+1}}{1 - q}$$

Cette quantité converge si et seulement si $|q| < 1$, et sa limite est alors $\frac{1}{1-q}$.

Proposition 1. Soit $q \in \mathbb{C}$. La série $\sum q^n$ converge si et seulement si $|q| < 1$. On a alors

$$\sum_{n=0}^{\infty} q^n = \frac{1}{1 - q}$$

Remarque 6. Pour la convergence des séries de Riemann, il va falloir attendre un peu.

Remarque 7. Il est temps de faire un petit bilan sur le vocabulaire. Soit $\sum u_n$ une série. Soit $n \in \mathbb{N}$.

- $\sum u_k$ c'est la **série**. Ce n'est pas un nombre mais une suite.
- u_n c'est le n ième **terme** de la série.
- $S_n = \sum_{k=0}^n u_k$ c'est la n ième **somme partielle** de la série.
- $S = \sum_{k=0}^{\infty} u_k$ c'est la **somme** de la série. Et ce n'est PAS une somme mais une limite de sommes.

Un autre exemple important est celui de la fonction exponentielle.

Proposition 2. Soit $z \in \mathbb{C}$. La série $\sum \frac{z^n}{n!}$ est convergente. Sa somme est

$$\sum_{n=0}^{\infty} \frac{z^n}{n!} = e^z$$

Démonstration. Soit $f : [0, 1] \rightarrow \mathbb{C}$ définie par $f(t) = e^{tz}$. La fonction f est de classe $\mathcal{C}^\infty$ sur $[0, 1]$. On peut donc lui appliquer l'inégalité de Taylor-Lagrange. Remarquons que pour tout $t \in [0, 1]$ et tout entier n , on a

$$f^{(n)}(t) = z^n e^{tz}$$

et en particulier, $f^{(n)}(0) = z^n$. Ainsi,

$$f(1) = e^z = \sum_{k=0}^n \frac{z^k}{k!} + R_n$$

où

$$|R_n| \leq \frac{1}{(n+1)!} \max_{t \in [0,1]} |z^{n+1} e^{tz}|$$

Soit $x = \operatorname{Re} z$. Remarquons que pour tout $t \in [0, 1]$,

$$|z^{n+1} e^{tz}| = |z|^{n+1} e^{tx} \leq |z|^{n+1} \max(1, e^x)$$

On a donc

$$|R_n| \leq \frac{|z|^{n+1}}{(n+1)!} \max(1, e^x)$$

Comme $|z|^{n+1} = o((n+1)!)$, on en déduit que $R_n \rightarrow 0$ lorsque n tend vers l'infini. Ainsi,

$$\sum_{k=0}^n \frac{z^k}{k!} \xrightarrow{n \rightarrow \infty} e^z$$

□

Exercice. Soit $x \in \mathbb{R}$. Montrer que les séries $\sum (-1)^n \frac{x^{2n}}{(2n)!}$ et $\sum (-1)^n \frac{x^{2n+1}}{(2n+1)!}$ sont convergentes, et que

$$\sum_{n=0}^{\infty} (-1)^n \frac{x^{2n}}{(2n)!} = \cos x$$

$$\sum_{n=0}^{\infty} (-1)^n \frac{x^{2n+1}}{(2n+1)!} = \sin x$$

1.2 linéarité de la somme

Proposition 3. Soient $\sum u_n$ et $\sum v_n$ deux séries convergentes. Soient $\lambda, \mu \in \mathbb{R}$ ou $\mathbb{C}$. La série $\sum (\lambda u_n + \mu v_n)$ est alors convergente et on a

$$\sum_{n=0}^{\infty} (\lambda u_n + \mu v_n) = \lambda \sum_{n=0}^{\infty} u_n + \mu \sum_{n=0}^{\infty} v_n$$

Démonstration. Soit $n \in \mathbb{N}$. On a

$$\sum_{k=0}^n (\lambda u_k + \mu v_k) = \lambda \sum_{k=0}^n u_k + \mu \sum_{k=0}^n v_k$$

Les sommes du membre de droite convergent respectivement vers $\sum_{k=0}^{\infty} u_k$ et $\sum_{k=0}^{\infty} v_k$. Par les opérations sur les limites, la somme du membre de gauche tend vers

$$\lambda \sum_{k=0}^{\infty} u_k + \mu \sum_{k=0}^{\infty} v_k$$

□

1.3 Restes

Définition 3. Soit $\sum u_n$ une série convergente, de somme S . Les *restes* de la série sont les nombres

$$R_n = S - S_n$$

où $S_n = \sum_{k=0}^n u_k$ est la n ième somme partielle de la série.

Remarque 8. La suite $(R_n)_{n \geq 0}$ des restes tend évidemment vers 0 lorsque n tend vers l'infini. Chaque reste R_n est en fait lui-même la somme de la série $\sum u_k$ (où $k \geq n+1$) :

$$R_n = \sum_{k=n+1}^{\infty} u_k$$

Exemple. Reprenons les séries géométriques de raison $q \in \mathbb{C}$, vérifiant $|q| < 1$. On a

$$S_n = \frac{1 - q^{n+1}}{1 - q} \text{ et } S = \frac{1}{1 - q}$$

Donc

$$R_n = S - S_n = \sum_{k=n+1}^{\infty} q^k = \frac{q^{n+1}}{1 - q}$$

1.4 Condition nécessaire de convergence

Proposition 4. Soit $\sum u_n$ une série.

- Si cette série converge, alors son terme général u_n tend vers 0.
- La réciproque est fautive : il existe des (tas de) séries divergentes dont le terme général tend vers 0.

Démonstration. Supposons que la série converge. La suite de terme général $S_n = \sum_{k=0}^n u_k$ converge donc vers une limite S qui est la somme de notre série. On a pour tout $n \geq 1$,

$$u_n = \sum_{k=0}^n u_k - \sum_{k=0}^{n-1} u_k = S_n - S_{n-1}$$

Cette quantité tend vers $S - S = 0$.

Pour un contre-exemple de la réciproque, posons par exemple pour tout $n \geq 1$,

$$u_n = \sqrt{n} - \sqrt{n-1}$$

On a

$$u_n = \frac{(\sqrt{n} - \sqrt{n-1})(\sqrt{n} + \sqrt{n-1})}{\sqrt{n} + \sqrt{n-1}} = \frac{1}{\sqrt{n} + \sqrt{n-1}}$$

qui tend vers 0 lorsque n tend vers l'infini. En revanche,

$$S_n = \sum_{k=1}^n (\sqrt{k} - \sqrt{k-1}) = \sqrt{n}$$

tend vers l'infini lorsque n tend vers l'infini. Nous avons donc là un exemple d'une série divergente dont le terme général tend pourtant vers 0. $\square$

Remarque 9. Ce théorème sert à démontrer qu'une série *diverge*. Lorsque le terme général d'une série ne tend pas vers 0, on dit que la série **diverge grossièrement**. À retenir absolument : à chaque fois qu'on s'attaque à l'étude d'une série, on regarde la limite éventuelle du terme général. Si cette limite n'existe pas, ou n'est pas nulle, c'est RÉGLÉ : la série diverge. Sinon, eh bien ON NE SAIT PAS. Il faut travailler davantage.

1.5 Lien entre suites et séries

S'intéresser à la convergence d'une série, c'est par définition même s'intéresser à la convergence d'une suite : la suite de ses sommes partielles. La fin de la démonstration précédente suggère que pour s'intéresser à la convergence d'une suite, on peut si l'on veut s'intéresser à la convergence d'une série.

Proposition 5. Soit $(u_n)_{n \geq 0}$ une suite réelle ou complexe. La suite $(u_n)_{n \geq 0}$ est convergente si et seulement si la série $\sum (u_{n+1} - u_n)$ est convergente.

Démonstration. C'est immédiat, puisque la somme partielle de cette série est

$$S_n = \sum_{k=0}^n (u_{k+1} - u_k) = u_{n+1} - u_0$$

par télescopage. On a même mieux : en cas de convergence,

$$\sum_{n=0}^{\infty} (u_{n+1} - u_n) = \ell - u_0$$

où ℓ est la limite de u_n . $\square$

Remarque 10. Pourquoi se compliquer la vie ? Pourquoi étudier une série pour montrer la convergence d'une suite ? Tout simplement parce que bientôt nous aurons des théorèmes qui nous diront que sous certaines conditions une série converge. En utilisant ces théorèmes sur des séries judicieuses nous pourrons donc aussi montrer que des suites convergent.

2 Séries à termes positifs

Les résultats qui suivent s'appliquent *uniquement* à des séries à termes réels positifs. Ils sont essentiels dans l'étude de la convergence des séries.

2.1 Majorer (ou pas) les sommes partielles

Proposition 6. Soit $\sum u_n$ une série à termes réels. On suppose que pour tout n assez grand, $u_n \geq 0$. La série converge si et seulement si ses sommes partielles sont majorées.

Démonstration. Il existe un entier N tel que pour tout $n \geq N$, $u_n \geq 0$. On a donc pour tout $n \geq N - 1$, $S_{n+1} - S_n = u_{n+1} \geq 0$. La suite $(S_n)_{n \geq 0}$ des sommes partielles de la série est donc croissante à partir d'un certain rang. Si la suite $(S_n)_{n \geq 0}$ est majorée alors elle converge. Si elle n'est pas majorée alors elle tend vers l'infini et, donc, elle diverge. $\square$

2.2 Majorer (ou minorer) le terme général

Proposition 7. Soient $\sum u_n$ et $\sum v_n$ deux séries à termes réels. On suppose

- Pour tout entier n assez grand on a

$$0 \leq u_n \leq v_n$$

- La série $\sum v_n$ converge.

Alors la série $\sum u_n$ converge aussi.

Démonstration. Pour simplifier on fait la démonstration en supposant $0 \leq u_n \leq v_n$ pour tout entier n . La preuve générale est à peine plus compliquée. Notons

$$S = \sum_{k=0}^{\infty} v_k$$

Pour tout entier $n \in \mathbb{N}$, on a

$$\sum_{k=0}^n u_k \leq \sum_{k=0}^n v_k \leq \sum_{k=0}^{\infty} v_k = S$$

Les sommes partielles de la série $\sum u_n$ sont donc majorées, ainsi la série $\sum u_n$ converge. $\square$

Corollaire 8. Soient $\sum u_n$ et $\sum v_n$ deux séries à termes réels. On suppose

- Pour tout entier n assez grand on a

$$0 \leq u_n \leq v_n$$

- La série $\sum u_n$ diverge.

Alors la série $\sum v_n$ diverge aussi.

Démonstration. C'est la contraposée de la proposition précédente. $\square$

Remarque 11. Nous avons maintenant une quasi-règle de base pour les séries À TERMES RÉELS POSITIFS :

- POUR MONTRER QU'UNE SÉRIE À TERMES POSITIFS CONVERGE ON MAJORE LE TERME GÉNÉRAL DE CETTE SÉRIE PAR LE TERME GÉNÉRAL D'UNE SÉRIE À TERMES POSITIFS CONVERGENTE.
- POUR MONTRER QU'UNE SÉRIE À TERMES POSITIFS DIVERGE ON MINORE LE TERME GÉNÉRAL DE CETTE SÉRIE PAR LE TERME GÉNÉRAL D'UNE SÉRIE À TERMES POSITIFS DIVERGENTE.
- ON MAJORE OU ON MINORE LE **terme général** ET PAS LES SOMMES PARTIELLES !

Mais ça tourne en rond cette histoire ? Non, l'idée est de posséder un « catalogue » de séries convergentes « de base », et de tenter des majorations par les termes généraux de ces séries. Nous avons déjà les séries géométriques dans notre catalogue. Bientôt les séries de Riemann. ...

2.3 Termes généraux équivalents

Proposition 9. Soient $\sum u_n$ et $\sum v_n$ deux séries à termes réels. On suppose

- Pour tout entier n assez grand, $v_n \geq 0$.
- $u_n \sim v_n$ lorsque n tend vers l'infini.

Alors, la série $\sum u_n$ converge si et seulement si la série $\sum v_n$ converge.

Démonstration. Tout d'abord, deux suites équivalentes ont même signe pour n assez grand. Ainsi, pour tout n assez grand on a u_n et v_n positifs.

Supposons que $\sum v_n$ converge. Comme $u_n \sim v_n$, on a pour tout n assez grand $0 \leq u_n \leq \frac{3}{2}v_n$. La série $\sum (\frac{3}{2}v_n)$ est à termes positifs et convergente, donc $\sum u_n$ converge.

Supposons inversement que $\sum u_n$ converge. Comme $v_n \sim u_n$, on a pour tout n assez grand $0 \leq v_n \leq \frac{3}{2}u_n$. Donc de même que deux lignes plus haut $\sum v_n$ converge. $\square$

Exemple. Considérons la série de terme général

$$u_n = \frac{n^2 + 1}{3n^2 + 2} e^{-n}$$

On a pour tout $n \in \mathbb{N}$, $u_n \geq 0$. De plus, $u_n \sim \frac{1}{3} \left(\frac{1}{e}\right)^n$, terme général d'une série géométrique de raison $\frac{1}{e}$. Donc la série $\sum u_n$ converge.

2.4 Équivalence des restes

Proposition 10. Soient $\sum u_n$ et $\sum v_n$ deux séries à termes positifs pour n assez grand. On suppose que

- Les deux séries convergent.
- $u_n \sim v_n$ lorsque n tend vers l'infini.

Alors les restes des deux séries sont équivalents : lorsque n tend vers $+\infty$,

$$\sum_{k=n+1}^{\infty} u_k \sim \sum_{k=n+1}^{\infty} v_k$$

Démonstration. Soit $\varepsilon > 0$. prenons $\varepsilon < 1$. Il existe un entier N tel que pour tout $n \geq N$ on ait $u_n, v_n \geq 0$ et

$$u_n(1 - \varepsilon) \leq v_n \leq u_n(1 + \varepsilon)$$

Soient $n, p \geq N$ tels que $n < p$. On a alors

$$(1 - \varepsilon) \sum_{k=n+1}^p u_k \leq \sum_{k=n+1}^p v_k \leq (1 + \varepsilon) \sum_{k=n+1}^p u_k$$

Faisons tendre p vers $+\infty$ ci-dessus. On obtient

$$(1 - \varepsilon) \sum_{k=n+1}^{\infty} u_k \leq \sum_{k=n+1}^{\infty} v_k \leq (1 + \varepsilon) \sum_{k=n+1}^{\infty} u_k$$

c'est à dire

$$(1 - \varepsilon)R_n \leq R'_n \leq (1 + \varepsilon)R_n$$

où R_n et R'_n sont les restes des séries $\sum_{n \geq 0} u_n$ et $\sum_{n \geq 0} v_n$. En résumé,

$$\forall \varepsilon > 0, \exists N \in \mathbb{N}, \forall n \geq N, (1 - \varepsilon)R_n \leq R'_n \leq (1 + \varepsilon)R_n$$

Par la définition de limite, cela signifie que $R_n \sim R'_n$. $\square$

3 Comparaison entre séries et intégrales

3.1 Séries vs intégrales

Soit $f : \mathbb{R}_+ \rightarrow \mathbb{R}_+$ une fonction décroissante. Considérons la série $\sum f(n)$. On a pour tout $k \geq 0$, pour tout $t \in [k, k+1]$,

$$f(k+1) \leq f(t) \leq f(k)$$

Intégrons entre k et $k+1$. Il vient

$$f(k+1) \leq \int_k^{k+1} f(t) dt \leq f(k)$$

Soient maintenant $0 \leq n < p$. Sommons l'inégalité ci-dessus pour $n \leq k \leq p-1$. On obtient

$$\sum_{k=n}^{p-1} f(k+1) \leq \sum_{k=n}^{p-1} \int_k^{k+1} f(t) dt \leq \sum_{k=n}^{p-1} f(k)$$

c'est à dire

$$\sum_{k=n+1}^p f(k) \leq \int_n^p f(t) dt \leq \sum_{k=n}^{p-1} f(k)$$

Ceci peut encore s'écrire

$$\int_n^{p+1} f(t) dt \leq_{(1)} \sum_{k=n}^p f(k) \leq_{(2)} \int_{n-1}^p f(t) dt$$

où l'inégalité (1) est valable pour $0 \leq n \leq p$ et l'inégalité (2) est valable pour $1 \leq n \leq p$.

Si l'on connaît une primitive de f , on a alors un encadrement des sommes partielles de la série qui permet très souvent de conclure à sa convergence ou à sa divergence. Nous pouvons enfin régler le cas des séries de Riemann pour un exposant réel.

Proposition 11. Soit $\alpha \in \mathbb{R}$.

- Si $\alpha \leq 0$, la série $\sum \frac{1}{n^\alpha}$ est grossièrement divergente.
- Si $0 < \alpha \leq 1$, la série $\sum \frac{1}{n^\alpha}$ diverge, bien que son terme général tende vers 0.
- Si $\alpha > 1$, la série $\sum \frac{1}{n^\alpha}$ converge.

Démonstration.

- Pour $\alpha \leq 0$, le terme général ne tend pas vers 0, d'où la divergence grossière.
- Supposons $0 < \alpha < 1$. On prend dans l'inégalité (1) ci-dessus $f(x) = \frac{1}{x^\alpha}$, $n \leftarrow 1$ et $p \leftarrow n$. Il vient donc

$$\begin{aligned} \sum_{k=1}^n \frac{1}{k^\alpha} &\geq \int_1^{n+1} \frac{dt}{t^\alpha} \\ &= \left[\frac{t^{1-\alpha}}{1-\alpha} \right]_1^{n+1} \\ &= \frac{(n+1)^{1-\alpha} - 1}{1-\alpha} \end{aligned}$$

Comme $\alpha < 1$, le minorant tend vers $+\infty$ lorsque n tend vers l'infini. Ainsi, par le théorème de comparaison des limites,

$$\sum_{k=1}^n \frac{1}{k^\alpha} \longrightarrow +\infty$$

lorsque n tend vers l'infini. La série diverge.

- Supposons maintenant $\alpha = 1$. On prend dans l'inégalité (1) ci-dessus $f(x) = \frac{1}{x}$, $n \leftarrow 1$ et $p \leftarrow n$. Il vient

$$\begin{aligned} \sum_{k=1}^n \frac{1}{k} &\geq \int_1^{n+1} \frac{dt}{t} \\ &= [\ln t]_1^{n+1} \\ &= \ln(n+1) \end{aligned}$$

Le minorant tend vers $+\infty$ lorsque n tend vers l'infini. Ainsi,

$$\sum_{k=1}^n \frac{1}{k} \longrightarrow +\infty$$

lorsque n tend vers l'infini. La série diverge.

- Supposons enfin $\alpha > 1$. On prend dans l'inégalité (2) ci-dessus $f(x) = \frac{1}{x^\alpha}$, $n \leftarrow 2$ et $p \leftarrow n$. Il vient donc

$$\begin{aligned} \sum_{k=1}^n \frac{1}{k^\alpha} &= 1 + \sum_{k=2}^n \frac{1}{k^\alpha} \\ &\leq 1 + \int_1^n \frac{dt}{t^\alpha} \\ &= 1 - \left[\frac{t^{1-\alpha}}{\alpha-1} \right]_1^n \\ &= 1 + \frac{1}{\alpha-1} - \frac{n^{1-\alpha}}{\alpha-1} \\ &\leq 1 - \frac{1}{\alpha-1} = \frac{\alpha}{\alpha-1} \end{aligned}$$

La série est à termes positifs et ses sommes partielles sont majorées. Elle est donc convergente.

□

Exercice. La comparaison séries-intégrales permet également d'évaluer des restes de séries convergentes. Considérons par exemple la série $\sum \frac{1}{n^2}$. Cette série converge, donc

$$R_n = \sum_{k=n+1}^{\infty} \frac{1}{k^2}$$

tend vers 0 lorsque n tend vers l'infini. Cherchons un équivalent de R_n . Pour cela, écrivons l'inégalité :

$$\int_{n+1}^{p+1} \frac{dt}{t^2} \leq_{(1)} \sum_{k=n+1}^p \frac{1}{k^2} \leq_{(2)} \int_n^p \frac{dt}{t^2}$$

ou encore

$$\frac{1}{n+1} - \frac{1}{p+1} \leq \sum_{k=n+1}^p \frac{1}{k^2} \leq \frac{1}{n} - \frac{1}{p}$$

Faisons maintenant tendre p vers l'infini :

$$\frac{1}{n+1} \leq \sum_{k=n+1}^{\infty} \frac{1}{k^2} \leq \frac{1}{n}$$

Ainsi

$$R_n \sim \frac{1}{n}$$

lorsque n tend vers l'infini.

3.2 Un exemple complet : la série harmonique

On considère ici la *série harmonique*

$$\sum \frac{1}{n}$$

Nous savons que cette série est divergente, mais nous allons estimer plus précisément ses sommes partielles S_n . Tout d'abord, en appliquant le paragraphe précédent, on obtient

$$\int_1^{n+1} \frac{dt}{t} \leq \sum_{k=1}^n \frac{1}{k} \leq 1 + \int_1^n \frac{dt}{t}$$

ou encore

$$\ln(n+1) \leq S_n \leq 1 + \ln n$$

On en déduit aisément que $S_n \sim \ln n$ lorsque n tend vers l'infini.

Considérons maintenant $u_n = S_n - \ln n$. Nous savons que la suite u et la série de terme général $v_n = u_{n+1} - u_n$ sont de même nature. Or,

$$\begin{aligned} v_n &= \sum_{k=1}^{n+1} \frac{1}{k} - \ln(n+1) - \sum_{k=1}^n \frac{1}{k} + \ln n \\ &= \frac{1}{n+1} - \ln \left(1 + \frac{1}{n} \right) \\ &= \frac{1}{n} \left(1 - \frac{1}{n} + o \left(\frac{1}{n} \right) \right) - \frac{1}{n} + \frac{1}{2} \frac{1}{n^2} + o \left(\frac{1}{n^2} \right) \\ &\sim -\frac{1}{2n^2} \end{aligned}$$

La série $\sum v_n$ est donc convergente, puisque son terme général est équivalent, au facteur $-\frac{1}{2}$ près, au terme général d'une série de Riemann convergente. Ainsi, $(u_n)_{n \geq 1}$ converge également. Sa limite, γ , est appelée la constante d'Euler. Remarquons que nous avons

$$\sum_{k=n+1}^{\infty} v_k \sim -\frac{1}{2n}$$

lorsque n tend vers l'infini. Cela résulte de l'exemple du paragraphe précédent et du théorème de l'équivalence des restes.

Estimons maintenant

$$w_n = u_n - \gamma = \sum_{k=1}^n \frac{1}{k} - \ln n - \gamma$$

Dans le calcul ci-dessus, nous avons

$$\sum_{k=1}^n v_k = u_{n+1} - u_1 = u_{n+1} - 1$$

d'où

$$\gamma = 1 + \sum_{k=1}^{\infty} v_k$$

Ainsi,

$$w_n = u_n - \gamma = 1 + \sum_{k=1}^{n-1} v_k - \left(1 + \sum_{k=1}^{\infty} v_k\right) = - \sum_{k=n}^{\infty} v_k$$

La remarque faite un peu plus haut nous montre que

$$w_n \sim \frac{1}{2n}$$

En conclusion :

$$\sum_{k=1}^n \frac{1}{k} = \ln n + \gamma + \frac{1}{2n} + o\left(\frac{1}{n}\right)$$

On pourrait continuer l'expérience en considérant

$$z_n = \sum_{k=1}^n \frac{1}{k} - \ln n - \gamma - \frac{1}{2n}$$

En appliquant à nouveau le théorème suite-série. On obtiendrait que $z_n \sim -\frac{1}{12n^2}$, et donc

$$\sum_{k=1}^n \frac{1}{k} = \ln n + \gamma + \frac{1}{2n} - \frac{1}{12n^2} + o\left(\frac{1}{n^2}\right)$$

Application, si l'on veut calculer une valeur approchée de la constante d'Euler on a intérêt à considérer par exemple la quantité

$$\sum_{k=1}^n \frac{1}{k} - \ln n - \frac{1}{2n} + \frac{1}{12n^2}$$

qui va converger vers γ aussi vite qu'un $o\left(\frac{1}{n^2}\right)$.

Considérons la fonction Python ci-dessous :

```
def approx_gamma(n):
    s = 0
    for k in range(1, n+1):
        s = s + 1 / k
    return s - log(n) - 1 / (2 * n) + 1 / (12 * n ** 2)
```

`approx_gamma(1000)` renvoie instantanément la valeur 0.577215664902 : tous les chiffres sont exacts. Si l'on s'était contentés de `return s - log(n)`, il aurait fallu additionner 10^{12} termes, c'est à dire attendre quelques semaines.

Remarque 12. Pourquoi tous les chiffres exacts ? Eh bien, croyez moi sur parole, le terme suivant dans le développement asymptotique de $\sum_{k=1}^n \frac{1}{k}$ est $\frac{1}{120n^4}$. Notre fonction `approx_gamma` renvoie donc la valeur de γ avec une erreur de l'ordre de $\frac{1}{120n^4}$, et pour $n = 1000$, cette quantité est inférieure à 10^{-14} . Attention, ceci n'est pas une preuve ! Nous n'avons pas un encadrement de l'erreur commise mais seulement un équivalent.

4 Convergence absolue

4.1 Séries absolument convergentes

Définition 4. Soit $\sum u_n$ une série à termes dans $\mathbb{R}$ ou $\mathbb{C}$. On dit que la série $\sum u_n$ converge absolument lorsque la série $\sum |u_n|$ converge.

Exemple. Soit $\alpha \in \mathbb{R}$. La série $\sum \frac{(-1)^n}{n^\alpha}$ est absolument convergente si et seulement si $\alpha > 1$.

Proposition 12. Soit $\alpha \in \mathbb{C}$. La série $\sum \frac{1}{n^\alpha}$ est absolument convergente si et seulement si $\operatorname{Re}(\alpha) > 1$.

Démonstration. Posons $\alpha = p + iq$, où $p, q \in \mathbb{R}$. On a

$$|n^\alpha| = |n^p e^{iq \ln n}| = n^p = n^{\operatorname{Re}(\alpha)}$$

D'où le résultat par le théorème sur les séries de Riemann à paramètre réel. $\square$

4.2 Convergence absolue et convergence

Proposition 13. Il existe des séries qui convergent, mais ne convergent pas absolument.

Démonstration. Considérons la série

$$\sum \frac{(-1)^n}{2n+1}$$

Cette série ne converge pas absolument, puisque $\frac{1}{2n+1} \sim \frac{1}{2} \frac{1}{n}$. Or la série $\sum \frac{1}{n}$ est une série de Riemann divergente.

En revanche, nous avons vu dans l'un des exercices du chapitre d'intégration que cette série converge, avec de plus

$$\sum_{n=0}^{\infty} \frac{(-1)^n}{2n+1} = \frac{\pi}{4}$$

□

Proposition 14. *Toute série absolument convergente est convergente.*

Remarque 13. En résumé :

- Absolument convergente $\Rightarrow$ convergente.
- Convergente $\not\Rightarrow$ absolument convergente.

Une série convergente mais pas absolument convergente est dite *semi-convergente*. Ce sont les séries semi-convergentes qui sont le plus difficiles à étudier.

Avant de démontrer ce théorème, introduisons une notation.

Notation. Pour tout réel x , on pose $x^+ = x$ si $x \geq 0$ et $x^+ = 0$ sinon. De même, on pose $x^- = -x$ si $x \leq 0$ et $x^- = 0$ sinon.

Proposition 15. *Pour tout réel x , les réels x^+ et x^- sont positifs,*

$$x = x^+ - x^- \text{ et } |x| = x^+ + x^-$$

Démonstration. C'est évident, il n'y a qu'à considérer deux cas, selon le signe de x . □

- Prouvons d'abord le théorème pour une série à termes réels.

Démonstration. Soit $(u_n)_{n \geq 0}$ une suite réelle. Supposons que la série $\sum |u_n|$ est convergente. On a, pour tout entier n ,

$$0 \leq u_n^+ \leq |u_n| \text{ et } 0 \leq u_n^- \leq |u_n|$$

Les séries à termes positifs $\sum u_n^+$ et $\sum u_n^-$ ont donc leur terme général majoré par le terme général d'une série convergente. On en déduit que ces deux séries convergent. Par linéarité, la série

$$\sum u_n = \sum (u_n^+ - u_n^-)$$

est également convergente. $\square$

- Prouvons maintenant le théorème pour une série à termes complexes.

Démonstration. Soit $(u_n)_{n \geq 0}$ une suite complexe. Supposons que la série $\sum |u_n|$ est convergente. On a pour tout entier n ,

$$0 \leq |\operatorname{Re}(u_n)| \leq |u_n|$$

La série à termes réels $\sum \operatorname{Re}(u_n)$ est ainsi absolument convergente, donc convergente. De même pour la série $\sum \operatorname{Im}(u_n)$ et, par linéarité, pour la série $\sum u_n$. $\square$

Remarque 14. Soit $n \in \mathbb{N}$. On a

$$\left| \sum_{k=0}^n u_k \right| \leq \sum_{k=0}^n |u_k|$$

par l'inégalité triangulaire pour les « sommes finies ». On passe à la limite dans l'inégalité en faisant tendre n vers l'infini (les deux séries mises en jeu convergent), et on obtient

$$\left| \sum_{k=0}^{\infty} u_k \right| \leq \sum_{k=0}^{\infty} |u_k|$$

Corollaire 16. Soit $(u_n)_{n \geq 0}$ une suite complexe. Soit $(v_n)_{n \geq 0}$ une suite de réels positifs. On suppose que $u_n = O(v_n)$ et que la série $\sum v_n$ est convergente. Alors, la série $\sum u_n$ est absolument convergente, donc convergente.

Démonstration. C'est évident, par définition même de la relation de domination. Il existe un réel $K > 0$ tel que pour tout n assez grand, on ait $|u_n| \leq K v_n$. $\square$

4.3 Séries alternées

Que faire lorsqu'une série n'est pas absolument convergente? Dans ce cas, il n'est pas possible de ramener l'étude de celle-ci à l'étude d'une série à termes positifs. Il existe cependant une classe de séries pour lesquelles on peut conclure à la convergence. Il s'agit des *séries alternées*.

Définition 5. Une *série alternée* est une série de la forme $\sum (-1)^n u_n$ où la suite $(u_n)_{n \geq 0}$ est une suite réelle **décroissante** qui **tend vers 0**.

Proposition 17. Soit $(u_n)_{n \geq 0}$ est une suite réelle décroissante qui tend vers 0. Notons S_n la n ième somme partielle de la série et R_n le n ième reste. Alors,

- La série $\sum (-1)^n u_n$ est convergente.
- Pour tout $n \geq 0$, R_n est du signe de $(-1)^{n+1}$ et

$$|R_n| \leq u_{n+1}$$

Démonstration. Pour tout entier naturel n , notons $S'_n = S_{2n}$ et $S''_n = S_{2n+1}$. On a

$$\begin{aligned} S'_{n+1} - S'_n &= S_{2n+2} - S_{2n} \\ &= \sum_{k=0}^{2n+2} (-1)^k u_k - \sum_{k=0}^{2n+2} (-1)^k u_k \\ &= u_{2n+2} - u_{2n+1} \\ &\leq 0 \end{aligned}$$

Ainsi, la suite $(S'_n)_{n \geq 0}$ est décroissante. Par un calcul analogue, la suite $(S''_n)_{n \geq 0}$ est croissante. Enfin, pour tout $n \geq 0$,

$$\begin{aligned} S'_n - S''_n &= \sum_{k=0}^{2n} (-1)^k u_k - \sum_{k=0}^{2n+1} (-1)^k u_k \\ &= u_{2n+1} \\ &\xrightarrow{n \rightarrow \infty} 0 \end{aligned}$$

Les suites $(S'_n)_{n \geq 0}$ et $(S''_n)_{n \geq 0}$ sont donc adjacentes. Elles convergent vers une même limite S . Par la « réciproque » du théorème des suites extraites, la suite $(S_n)_{n \geq 0}$ converge aussi vers S .

Venons-en au signe et à la majoration de $|R_n|$. Par un résultat classique sur les suites adjacentes, on a pour tout entier $n \geq 0$,

$$S_{2n+1} \leq S \leq S_{2n}$$

Soustrayons S_{2n} pour obtenir

$$-u_{2n+1} = S_{2n+1} - S_{2n} \leq R_{2n} \leq 0$$

Ainsi,

$$|R_{2n}| \leq u_{2n+1}$$

De même,

$$S_{2n+1} \leq S \leq S_{2n+2}$$

et donc, en soustrayant S_{2n+1} ,

$$0 \leq R_{2n+1} \leq S_{2n+2} - S_{2n+1} = u_{2n+2}$$

On a donc bien, pour tout entier naturel n ,

$$|R_n| \leq u_{n+1}$$

□

5 Appendice - Représentations p-adiques des réels

Dans toute cette section, p désigne un entier supérieur ou égal à 2. Nous allons montrer que « presque » tous les réels $x \in [0, 1[$ s'écrivent de façon unique comme la somme d'une série :

$$x = \sum_{k=1}^{\infty} \frac{a_k}{p^k}$$

où les a_k sont des entiers entre 0 et $p - 1$.

5.1 Représentation en base p des entiers naturels

Ce paragraphe n'a rien à voir avec les séries, mais avant de parler de la représentation p-adique des réels il est logique, en guise d'échauffement, de commencer par s'intéresser aux entiers.

Proposition 18. *Soit $n \in \mathbb{N}$. L'entier n s'écrit de façon unique*

$$\sum_{k=0}^{\infty} a_k p^k$$

où les $a_k \in \llbracket 0, p - 1 \rrbracket$ sont tous nuls sauf un nombre fini.

Démonstration.

- Pour l'existence, on fait une récurrence forte sur n .

Pour $n = 0$ l'existence est immédiate : on prend tous les a_k nuls.

Soit $n \geq 1$. Supposons que les entiers entre 0 et $n - 1$ s'écrivent sous la forme voulue. Effectuons la division euclidienne de n par p :

$$n = mp + a_0$$

où $m \in \mathbb{N}$ et $a_0 \in \llbracket 0, p - 1 \rrbracket$. On a $0 \leq m < n$, et on peut donc grâce à l'hypothèse de récurrence écrire

$$m = \sum_{k=0}^{\infty} b_k p^k$$

où les b_k sont des entiers entre 0 et $p-1$ presque tous nuls. On reporte et on obtient

$$n = \sum_{k=0}^{\infty} b_k p^{k+1} + a_0$$

qui est bien de la forme recherchée, en posant pour tout $k \geq 1$, $a_k = b_{k-1}$.

- Prouvons maintenant l'unicité en raisonnant par l'absurde. Soit $n \in \mathbb{N}$. Supposons que l'entier n admet deux écritures distinctes

$$n = \sum_{k=0}^{\infty} a_k p^k = \sum_{k=0}^{\infty} b_k p^k$$

L'ensemble des entiers k tels que $a_k \neq b_k$ est une partie de $\mathbb{N}$ non vide par hypothèse, et majorée puisque les a_k et les b_k sont nuls (donc égaux) pour k assez grand. Elle possède donc un plus grand élément que nous appellerons j . On a pour fixer les idées $a_j > b_j$. Soustrayons, il vient

$$(a_j - b_j)p^j = \sum_{k=0}^{j-1} (b_k - a_k)p^k$$

Le membre de gauche est supérieur ou égal à p^j . Majorons le membre de droite :

$$\sum_{k=0}^{j-1} (b_k - a_k)p^k \leq (p-1) \sum_{k=0}^{j-1} p^k = (p-1) \frac{p^j - 1}{p-1} = p^j - 1$$

Le membre de gauche et celui de droite ne peuvent pas être égaux, contradiction.

□

Notation. Pour $n = \sum_{k=0}^N a_k p^k$ avec les conditions requises sur les a_k , on écrit $n = \overline{a_N \dots a_0}_p$, ou plus simplement $n = a_N \dots a_0$ lorsque aucune confusion n'est à craindre. Les a_k sont appelés les chiffres de l'écriture de n en base p .

Exemple. Les deux valeurs de p les plus utilisées sont sans aucun doute $p = 10$ (écriture décimale) et $p = 2$ (écriture binaire). Prenons $n = 1789$. On a

$$1789 = 1 \times 10^3 + 7 \times 10^2 + 8 \times 10^1 + 9 \times 10^0$$

donc l'écriture décimale de 1789 est ...1789.

On utilise aussi souvent en informatique la base 16 (hexadécimale) dont les chiffres sont notés $0, 1, \dots, 9, A, B, C, D, E, F$.

Exercice. Quelle est l'écriture binaire de 1789 ? Regrouper les chiffres de l'écriture binaire 4 par 4 et en déduire l'écriture hexadécimale de 1789.

5.2 Représentations en base p d'un réel

Tout réel $x \geq 0$ s'écrit sous la forme $n + y$ où $n = \lfloor x \rfloor \in \mathbb{N}$ et $y \in [0, 1[$. Nous allons donc nous concentrer sur les réels de l'intervalle $[0, 1[$.

Soit $x \in [0, 1[$. Posons $x_0 = x$ et $a_1 = \lfloor px_0 \rfloor$. On a

$$a_1 \leq px < a_1 + 1$$

donc

$$0 \leq x - \frac{a_1}{p} < \frac{1}{p}$$

De plus

$$0 \leq a_1 \leq px < p$$

donc, comme a_1 est entier,

$$0 \leq a_1 \leq p - 1$$

Ceci suggère de recommencer, en posant $x_1 = x - \frac{a_1}{p}$.

On définit par récurrence sur n une suite $(x_n)_{n \geq 0}$ et une suite $(a_n)_{n \geq 0}$ comme suit : $a_0 = 0$, $x_0 = x$ et, pour tout $n \in \mathbb{N}$,

$$a_{n+1} = \lfloor p^{n+1} x_n \rfloor \quad \text{et} \quad x_{n+1} = x_n - \frac{a_{n+1}}{p^{n+1}}$$

Proposition 19. *On a pour tout entier $n \geq 0$:*

- $0 \leq a_n \leq p - 1$.
- $x_n = x - \sum_{k=0}^n \frac{a_k}{p^k}$.
- $0 \leq x_n < \frac{1}{p^n}$.

Démonstration. Par récurrence sur n . Pour $n = 0$, c'est évident. Supposons donc la propriété vérifiée pour un entier n . On a

$$a_{n+1} \leq p^{n+1} x_n < a_{n+1} + 1$$

Comme $x_n < \frac{1}{p^n}$, on a $a_{n+1} < p$ donc $a_{n+1} \leq p - 1$. Comme $x_n \geq 0$, a_{n+1} l'est aussi (c'est sa partie entière). Enfin,

$$\frac{a_{n+1}}{p^{n+1}} \leq x_n < \frac{a_{n+1}}{p^{n+1}} + \frac{1}{p^{n+1}}$$

d'où, après une petite soustraction,

$$0 \leq x_{n+1} < \frac{1}{p^{n+1}}$$

□

Corollaire 20. On a $x = \sum_{k=0}^{\infty} \frac{a_k}{p^k}$.

Démonstration. Évident : d'après la proposition précédente, x_n tend vers 0 lorsque n tend vers l'infini. $\square$

Définition 6. La suite $(a_n)_{n \geq 1}$ est appelé un développement p -adique du réel x . On note $x = 0.a_1a_2a_3 \dots$

Exemple. Un développement décimal du réel $\frac{1}{7}$ est $0.142857142857 \dots$. Montrons-le. Soit $y = 0.\overline{142857}$. La barre indique que l'on répète à l'infini cette séquence de 6 chiffres. On a donc

$$10^6 y = 142857.\overline{142857} = 142857 + y$$

De là, $999999y = 142857$ et donc

$$y = \frac{142857}{999999} = \frac{1}{7}$$

Exercice. Exhiber un candidat pour développement décimal de $\frac{1}{13}$ et **montrer** que vous avez bien choisi.

Remarque 15. Si x est un réel positif quelconque, un développement p -adique de x est $x = b_n \dots b_0.a_1a_2a_3 \dots$ où $b_n \dots b_0$ est un développement p -adique de la partie entière de x . Si $x < 0$, on écrit $x = -y$ où $y = -x$.

5.3 Unicité du développement

Eh bien il n'y a pas unicité! Considérons par exemple le réel $x = 0.49999 \dots$. On a

$$10x - 4 = 0.9999 \dots$$

d'où

$$10(10x - 4) = 9 + 10x - 4$$

c'est-à-dire $90x = 45$, d'où $x = \frac{1}{2}$. Mais on a aussi $\frac{1}{2} = 0.50000 \dots$. Ainsi, $\frac{1}{2}$ admet deux développements décimaux distincts.

Nous allons maintenant regarder quels sont précisément les réels admettant plusieurs développements p -adiques, combien ils en ont, et à quoi ressemblent ces développements.

Soit donc $x \in [0, 1]$. On suppose

$$x = \sum_{k=1}^{\infty} \frac{a_k}{p^k} = \sum_{k=1}^{\infty} \frac{b_k}{p^k}$$

où les a_k et les b_k sont des entiers entre 0 et $p-1$, et au moins un des a_k est différent du b_k correspondant. Prenons un tel k minimal, notons le k_0 et supposons par exemple $a_{k_0} > b_{k_0}$. On a donc

$$\frac{a_{k_0} - b_{k_0}}{p^{k_0}} = \sum_{k=k_0+1}^{\infty} \frac{b_k - a_k}{p^k}$$

La somme du membre de droite est inférieure ou égale à

$$\sum_{k=k_0+1}^{\infty} \frac{p-1}{p^k} = (p-1) \sum_{k=k_0+1}^{\infty} \frac{1}{p^k} = \frac{1}{p^{k_0}}$$

tout le monde ayant reconnu ci-dessus un reste de série géométrique. Et elle est égale à $\frac{1}{p^{k_0}}$ si et seulement si pour tout $k \geq k_0$, on a $b_k = p-1$ et $a_k = 0$ (le lecteur est invité à justifier cette affirmation pas si évidente).

Le membre de gauche est quant à lui supérieur ou égal à $\frac{1}{p^{k_0}}$, et est égal à cette valeur si et seulement si $a_{k_0} = b_{k_0} + 1$. Mais ces deux membres sont égaux. On en déduit donc :

- Pour tout $k \geq k_0$, on a $b_k = p-1$ et $a_k = 0$
- $a_{k_0} = b_{k_0} + 1$.

On voit donc que si x possède deux développements p -adiques, l'un de ces deux développements est nul à partir d'un certain rang. Ceci suggère de définir l'ensemble

$$E = \left\{ \frac{n}{p^k}, n \in \mathbb{Z}, k \in \mathbb{N} \right\}$$

Muni de l'addition et de la multiplication, cet ensemble est un sous-anneau de $\mathbb{R}$, l'ensemble des nombres p -adiques (ont dit les nombres décimaux lorsque $p = 10$).

On en déduit (presque) la proposition :

Proposition 21. *Soit $x \in [0, 1[$. Le nombre x possède deux développements p -adiques distincts si et seulement si $x \in E$. Dans le cas où $x \in E$, x possède exactement deux développements p -adiques. L'un est nul à partir d'un certain rang, de la forme*

$$x = \sum_{k=1}^N \frac{a_k}{p^k}$$

où les a_k sont des entiers entre 0 et $p-1$, et $a_N \neq 0$. L'autre est

$$x = \sum_{k=1}^{N-1} \frac{a_k}{p^k} + \frac{a_N - 1}{p^N} + \sum_{k=N+1}^{\infty} \frac{p-1}{p^k}$$

Démonstration. Il reste une vérification à faire, à savoir que les deux développements ci-dessus sont bien égaux à x . Elle est laissée au lecteur. $\square$

Exemple. On prend $p = 10$. Soit $x = \frac{4675}{10000}$. Les deux développements décimaux de x sont $x = 0.4675000\dots$ et $x = 0.4674999\dots$

Exemple. Soit $x = \frac{11}{13}$. Le réel x n'étant pas un nombre décimal, il possède un unique développement. L'algorithme de division appris au CM2 nous donne $x = 0.\overline{846153}$. Une vérification serait de poser $y = 0.\overline{846153}$ et de constater que $1000000y = 846153 + y$, donc $999999y = 846153$ d'où $y = \frac{846153}{999999} = \frac{11}{13} = x$.

5.4 Développements périodiques

NB : Ce paragraphe est traité à titre de complément. Il peut être sauté en première lecture.

Quels sont précisément les réels admettant un développement p -adique périodique à partir d'un certain rang ?

Périodique implique rationnel

Commençons par le sens facile. Soit

$$x = n + 0.b_1 \dots b_N \overline{a_1 \dots a_k}$$

où $n \in \mathbb{Z}$, $N \in \mathbb{N}$, $k \in \mathbb{N}^*$ et les $a_i, b_i \in \llbracket 0, p-1 \rrbracket$. En clair, x est un réel dont le développement p -adique est périodique à partir du rang $N+1$, de période k .

Proposition 22. On a $x \in \mathbb{Q}$.

Démonstration. Commençons par faire un peu de ménage. Soit $y = 0.\overline{a_1 \dots a_k}$. On a

$$y = 10^N(x - n) - \overline{b_1 \dots b_N}$$

Pour montrer ce que l'on veut, il suffit de montrer que $y \in \mathbb{Q}$. Or

$$10^k y = a_1 \dots a_k \cdot \overline{a_1 \dots a_k} = a_1 \dots a_k + 0.\overline{a_1 \dots a_k} = a + y$$

où $a = a_1 \dots a_k$. Ainsi,

$$y = \frac{a}{10^k - 1} \in \mathbb{Q}$$

$\square$

C'était le sens facile. Passons au sens *intéressant*. Soit $x \in \mathbb{Q}_+$. Si x est entier il a évidemment un développement p -adique périodique (de période 1). Supposons donc x non entier et posons $x = \frac{a}{b}$ où $a, b \in \mathbb{N}$, $b \geq 2$.

Rationnel implique périodique - Cas 1

Supposons dans un premier temps que $b \wedge p = 1$. Par exemple, si $p = 10$, nous supposons que b n'est ni pair ni multiple de 5. Il existe des démonstrations « élémentaires » de ce qui va suivre mais la théorie des anneaux et la théorie des groupes finis vont nous fournir une preuve très élégante. Rappelons que $\mathbb{Z}/b\mathbb{Z}$ est un anneau commutatif. Ses éléments inversibles forment un groupe G pour la multiplication. Notons g le cardinal de ce groupe. Comme $b \wedge p = 1$, la classe de p appartient à G . Le théorème de Lagrange nous dit alors que $\bar{p}^g = \bar{1}$ dans $\mathbb{Z}/b\mathbb{Z}$. En termes de congruences, cela signifie que

$$p^g \equiv 1 \pmod{b}$$

ou encore qu'il existe $m \in \mathbb{N}$ (forcément positif) tel que

$$p^g = 1 + mb$$

Multiplions tout cela par a pour obtenir

$$p^g a = a + m'b$$

où $m' = ma \in \mathbb{N}$. Puis divisons par b :

$$p^g x = x + m'$$

Notons $x = 0.a_1 a_2 \dots$ le développement p -adique de x . Nous avons donc

$$a_1 a_2 \dots a_g + 0.a_{g+1} a_{g+2} \dots = m' + 0.a_1 a_2 \dots$$

d'où, en identifiant les parties entières,

$$m' = a_1 a_2 \dots a_g$$

et

$$0.a_{g+1} a_{g+2} \dots = 0.a_1 a_2 \dots$$

Peut-on faire des identifications dans l'égalité ci-dessus ? Oui, car b est premier avec p , donc le réel $p^g x - m'$ qui apparaît dans le membre de gauche de l'égalité ci-dessus a un unique développement p -adique. Ainsi, $a_{g+1} = a_1$, $a_{g+2} = a_2$, ... Le développement p -adique de x est périodique de période g .

Exemple. Prenons $x = \frac{1}{7}$. L'anneau $\mathbb{Z}/7\mathbb{Z}$ est un corps car 7 est premier. On a donc $g = 6$. De là, $10^6 \equiv 1 \pmod{7}$. Effectivement, une division euclidienne nous donne

$$10^6 = 1 + 142857 \times 7$$

En divisant par 7, on obtient ainsi

$$10^6 x = x + 142857$$

d'où

$$x = 0.\overline{142857}$$

Rationnel implique périodique - Cas 2

Il nous reste à traiter les rationnels dont le dénominateur n'est pas nécessairement premier avec p . Pour ne pas surcharger les notations, je vais prendre dans ce qui suit $p = 10$, mais les raisonnements s'adaptent sans problème à p quelconque.

Soit $x = \frac{a}{b} \in \mathbb{Q}_+$ où $a, b \in \mathbb{N}$, $b \geq 2$. Écrivons

$$b = 2^\alpha 5^\beta b'$$

où $\alpha, \beta \in \mathbb{N}$ et b' est premier avec p . Prenons pour fixer les idées $\alpha \leq \beta$. On a

$$10^\beta x = 10^\beta \frac{a}{2^\alpha 5^\beta b'} \frac{2^{\beta-\alpha} a}{b'} = \frac{a'}{b'}$$

On est ainsi ramenés au cas du dénominateur premier avec p : $10^\beta x$ a un développement p -adique périodique, ce qui veut dire que x a un développement périodique à partir du rang $\beta + 1$.

Exemple. Prenons $x = \frac{1}{35}$. Avec les notations ci-dessus, on a $\alpha = 0$ et $\beta = 1$. On a $10x = \frac{2}{7} = 0.\overline{285714}$. Ainsi, $x = 0.0\overline{285714}$ a un développement périodique à partir du rang 2, et de période 6.

Conclusion

Nous avons donc prouvé le théorème suivant.

Proposition 23. *Soit $x \in \mathbb{R}_+$. Le développement p -adique de x est périodique à partir d'un certain rang si et seulement si x est rationnel.*

Chapitre 36

Familles Sommables

1 Familles sommables de réels positifs

1.1 Introduction

Nous avons déjà étudié la droite réelle achevée dans le chapitre sur les réels. Voici quelques rappels.

On étend l'addition des réels à l'ensemble $[0, +\infty]$ en posant, pour tout $x \in [0, +\infty]$, $x + \infty = \infty$.

On étend également l'ordre usuel sur $\mathbb{R}_+$ à $[0, +\infty]$ en posant pour tout $x \in \mathbb{R}_+$, $x < +\infty$. Le plus grand élément de $[0, +\infty]$ est donc $+\infty$.

Pour toute partie E de $[0, +\infty]$, la borne supérieure de E est, comme d'habitude, le plus petit majorant de E . On montre facilement que

- Si $E = \emptyset$, $\sup E = 0$.
- Si E est non vide et majorée par un réel positif, $\sup E \in \mathbb{R}_+$ est la borne supérieure usuelle de E .
- Si E est non vide et n'est pas majorée par un réel positif, alors $\sup E = +\infty$.

1.2 Notion de famille sommable

Soit $(u_i)_{i \in I}$ une famille de réels positifs indexée par un ensemble I quelconque. Nous allons donner un sens à l'expression $\sum_{i \in I} u_i$. C'est évidemment facile si l'ensemble I est fini (prendre la somme usuelle), et cela nous suggère la définition suivante dans le cas général.

Si $u = (u_i)_{i \in I}$ est une famille de réels positifs, posons

$$\mathcal{S}(u) = \left\{ \sum_{i \in F} u_i, F \text{ fini}, F \subset I \right\}$$

Définition 1. Soit $u = (u_i)_{i \in I}$ une famille de réels positifs. La *somme* de la famille u est

$$\sum_{i \in I} u_i = \sup \mathcal{S}(u)$$

La famille u est dite *sommable* lorsque sa somme est finie.

Remarque 1. La somme d'une famille de réels positifs est donc un réel positif, ou alors $+\infty$.

Exemple. Soit $x \in [0, 1[$. Pour tout $n \in \mathbb{Z}$, posons $u_n = x^{|n|}$. Montrons que la famille

$(u_n)_{n \in \mathbb{Z}}$ est sommable. Pour cela, soit $F \subset \mathbb{Z}$ fini. On a

$$\sum_{n \in F} u_n = S_1 + S_2$$

où

$$S_1 = \sum_{n \in F, n < 0} x^{|n|} \text{ et } S_2 = \sum_{n \in F, n \geq 0} x^{|n|}$$

Majorons S_1 et S_2 . Soit $M \in \mathbb{N}^*$ tel que pour tout $n \in F$, $|n| \leq M$. On a alors

$$S_1 \leq \sum_{n=-M}^{-1} x^{-n} = \sum_{n=1}^M x^n = \frac{x - x^{M+1}}{1 - x} \leq \frac{x}{1 - x}$$

De même,

$$S_2 \leq \sum_{n=0}^M x^n = \frac{1 - x^{M+1}}{1 - x} \leq \frac{1}{1 - x}$$

d'où

$$\sum_{n \in F} x^{|n|} \leq \frac{1 + x}{1 - x}$$

Ainsi,

$$\sup \mathcal{S}(u) \leq \frac{1 + x}{1 - x}$$

donc la famille $(u_n)_{n \in \mathbb{Z}}$ est sommable.

Montrons maintenant que la somme de cette famille est $\frac{1+x}{1-x}$. Pour cela, donnons-nous $\varepsilon > 0$. Il existe un entier N tel que

$$\sum_{n=-N}^{-1} x^{-n} \geq \frac{x}{1 - x} - \frac{1}{2}\varepsilon$$

et

$$\sum_{n=0}^N x^n \geq \frac{1}{1 - x} - \frac{1}{2}\varepsilon$$

De là, en posant $F = \llbracket -N, N \rrbracket$,

$$\sum_{n \in F} u_n \geq \frac{1 + x}{1 - x} - \varepsilon$$

Il en résulte que

$$\sup \mathcal{S}(u) \geq \frac{1 + x}{1 - x} - \varepsilon$$

Ceci étant vrai pour tout $\varepsilon > 0$, on a donc

$$\sup \mathcal{S}(u) \geq \frac{1 + x}{1 - x}$$

En conclusion,

$$\sup \mathcal{S}(u) = \sum_{n \in \mathbb{Z}} u_n = \frac{1+x}{1-x}$$

Remarquons l'inégalité entre la somme d'une famille de réels positifs et la somme d'une sous-famille :

Proposition 1. *Soit $u = (u_i)_{i \in I}$ une famille de réels positifs. Soit $I' \subset I$. Soit $u' = (u_i)_{i \in I'}$. Alors,*

$$\sum_{i \in I'} u_i \leq \sum_{i \in I} u_i$$

En particulier, si la famille u est sommable, alors la famille u' est aussi sommable.

Démonstration. Soit F fini inclus dans I' . Alors, $F \subset I$ et donc

$$\sum_{i \in F} u_i \in \mathcal{S}(u)$$

Ainsi,

$$\mathcal{S}(u') \subset \mathcal{S}(u)$$

d'où

$$\sum_{i \in I'} u_i = \sup \mathcal{S}(u') \leq \sup \mathcal{S}(u) = \sum_{i \in I} u_i$$

□

1.3 Deux cas particuliers importants

L'ensemble I des indices des familles que nous considérons est a priori quelconque. Il y a deux cas particuliers importants à noter.

Proposition 2. *Soit $u = (u_i)_{i \in I}$ une famille de réels positifs indexée par un ensemble fini I . Alors, la famille $(u_i)_{i \in I}$ est sommable et sa somme est la somme usuelle des u_i .*

Démonstration. Posons $I = \{i_1, \dots, i_n\}$. Pour toute partie F finie de I , on a

$$\sum_{i \in F} u_i \leq S = \sum_{k=1}^n u_{i_k}$$

L'ensemble $\mathcal{S}(u)$ est majoré par S . De plus, en prenant $F = I$, on constate que $S \in \mathcal{S}(u)$. Donc $\mathcal{S}(u)$ possède un plus grand élément, qui est S . □

Proposition 3. Soit $(u_n)_{n \in \mathbb{N}}$ une suite de réels positifs. On a

$$\sum_{n \in \mathbb{N}} u_n = \sum_{n=0}^{\infty} u_n$$

où la somme de droite vaut par convention $+\infty$ lorsque la série diverge. En particulier, la famille $(u_n)_{n \in \mathbb{N}}$ est sommable si et seulement si la série $\sum u_n$ est convergente.

Démonstration. Posons

$$S = \sum_{n=0}^{\infty} u_n \in [0, +\infty]$$

Soit F une partie finie de $\mathbb{N}$. L'ensemble F possède un plus grand élément n_0 . Comme les u_i sont positifs, on a

$$\sum_{i \in F} u_i \leq \sum_{i=0}^{n_0} u_i \leq S$$

Ainsi, par passage à la borne supérieure,

$$(1) \quad \sum_{i \in I} u_i \leq S$$

Soit $\alpha < S$. Par définition de la somme d'une série, il existe un entier N tel que, pour tout $n \geq N$,

$$\sum_{i=0}^n u_i > \alpha$$

En prenant $F = \llbracket 0, N \rrbracket$, on a donc

$$\sum_{i \in F} u_i > \alpha$$

De là,

$$\sum_{i \in I} u_i \geq \sum_{i \in F} u_i > \alpha$$

Ceci étant vrai pour tout $\alpha < S$, il vient donc

$$(2) \quad \sum_{i \in I} u_i \geq S$$

En combinant (1) et (2), on obtient

$$\sum_{i \in I} u_i = S = \sum_{n=0}^{\infty} u_n$$

□

1.4 Invariance par permutation

Proposition 4. Soit $(u_i)_{i \in I}$ une famille de réels positifs. Soit $\sigma : I \rightarrow I$ une bijection. Pour tout $i \in I$, soit $v_i = u_{\sigma(i)}$. Alors

$$\sum_{i \in I} v_i = \sum_{i \in I} u_i$$

En particulier, la famille $(v_i)_{i \in I}$ est sommable si et seulement si la famille $(u_i)_{i \in I}$ est sommable.

Démonstration. Notons $S = \sum_{i \in I} u_i$ et $S' = \sum_{i \in I} v_i$.

Soit F fini inclus dans I . Notons $G = \sigma(F)$. Remarquons que G est fini, et inclus dans I .

On a

$$\sum_{i \in F} v_i = \sum_{i \in F} u_{\sigma(i)} = \sum_{k \in G} u_k \leq S$$

De là, par passage à la borne supérieure,

$$S' \leq S$$

Soit maintenant $\alpha < S$. Il existe $F \subset I$ tel que

$$\sum_{i \in F} u_i > \alpha$$

Soit $G = \sigma^{-1}(F)$. On a alors

$$\sum_{k \in G} v_k = \sum_{i \in F} u_i > \alpha$$

Ainsi,

$$S' = \sum_{k \in I} v_k \geq \sum_{k \in G} v_k > \alpha$$

En résumé, pour tout $\alpha < S$,

$$\alpha < S' \leq S$$

d'où

$$S' = S$$

□

Corollaire 5. Soit $(u_n)_{n \geq 0}$ une suite de réels positifs. Soit $\sigma : \mathbb{N} \rightarrow \mathbb{N}$ une bijection. Pour tout $n \in \mathbb{N}$, soit $v_n = u_{\sigma(n)}$. Alors

$$\sum_{n=0}^{\infty} v_n = \sum_{n=0}^{\infty} u_n$$

En particulier, la série $\sum v_n$ converge si et seulement si la série $\sum u_n$ converge.

1.5 Linéarité

Proposition 6. *Soient $(u_i)_{i \in I}$ et $(v_i)_{i \in I}$ deux familles de réels positifs. Alors*

$$\sum_{i \in I} (u_i + v_i) = \sum_{i \in I} u_i + \sum_{i \in I} v_i$$

Démonstration. Posons

$$S' = \sum_{i \in I} u_i, \quad S'' = \sum_{i \in I} v_i \quad \text{et} \quad S = \sum_{i \in I} (u_i + v_i)$$

Il s'agit de prouver que $S = S' + S''$.

Soit F fini, $F \subset I$. On a

$$\sum_{i \in F} (u_i + v_i) = \sum_{i \in F} u_i + \sum_{i \in F} v_i \leq S' + S''$$

Par passage à la borne supérieure,

$$(1) \quad S \leq S' + S''$$

Soit maintenant $\gamma < S' + S''$. On laisse le soin au lecteur de montrer que $\gamma = \alpha + \beta$, où $\alpha < S'$ et $\beta < S''$. Il existe F' fini, $F' \subset I$, tel que

$$\sum_{i \in F'} u_i > \alpha$$

De même, il existe F'' fini, $F'' \subset I$, tel que

$$\sum_{i \in F''} v_i > \beta$$

Posons $F = F' \cup F''$. On a alors

$$\sum_{i \in F} u_i \geq \sum_{i \in F'} u_i > \alpha$$

et

$$\sum_{i \in F} v_i \geq \sum_{i \in F''} v_i > \beta$$

De là,

$$\sum_{i \in I} (u_i + v_i) \geq \sum_{i \in F} (u_i + v_i) = \sum_{i \in F} u_i + \sum_{i \in F} v_i > \alpha + \beta = \gamma$$

Ceci étant vrai pour tout $\gamma < S' + S''$, il en résulte

$$(2) \quad S \geq S' + S''$$

De (1) et (2), on déduit le résultat. $\square$

Proposition 7. Soit $(u_i)_{i \in I}$ une famille de réels positifs. Soit $\lambda \in \mathbb{R}_+$. Alors

$$\sum_{i \in I} (\lambda u_i) = \lambda \sum_{i \in I} u_i$$

Démonstration. Exercice. $\square$

1.6 Sommation par paquets

Proposition 8. Soit $(u_i)_{i \in I}$ une famille de réels positifs. On suppose que

$$I = \bigcup_{j \in J} I_j$$

où les ensembles I_j sont disjoints deux à deux. Alors

$$\sum_{i \in I} u_i = \sum_{j \in J} \left(\sum_{i \in I_j} u_i \right)$$

On convient, dans l'égalité ci-dessus, que si l'un des termes d'une « somme » est égal à $+\infty$, alors la somme est égale à $+\infty$.

Démonstration. Nous admettrons ce résultat. $\square$

En reprenant les mêmes notations, on en déduit le résultat suivant.

Proposition 9. La famille $(u_i)_{i \in I}$ est sommable si et seulement si

- Pour tout $j \in J$, la famille $(u_i)_{i \in I_j}$ est sommable, de somme S_j .
- La famille $(S_j)_{j \in J}$ est sommable.

Exemple. Reprenons l'exemple vu au début du chapitre. Soit $x \in [0, 1[$. Considérons la famille $(x^{|n|})_{n \in \mathbb{Z}}$. Décomposons $\mathbb{Z}$ en $\mathbb{Z} = \mathbb{Z}_-^* \cup \mathbb{N}$. La famille $(x^{|n|})_{n \in \mathbb{N}}$ est sommable, de somme

$$\sum_{n \in \mathbb{N}} x^n = \sum_{n=0}^{\infty} x^n = \frac{1}{1-x}$$

De même, la famille $(x^{|n|})_{n \in \mathbb{Z}_-^*}$ est sommable, de somme

$$\sum_{n \in \mathbb{Z}_-^*} x^{|n|} = \sum_{n \in \mathbb{N}^*} x^n = \sum_{n=1}^{\infty} x^n = \frac{x}{1-x}$$

Par le théorème de sommation par paquets, la famille $(x^{|n|})_{n \in \mathbb{Z}}$ est sommable, de somme

$$\sum_{n \in \mathbb{Z}} x^{|n|} = \sum_{n \in \mathbb{Z}_-^*} x^{|n|} + \sum_{n \in \mathbb{N}} x^{|n|} = \frac{x}{1-x} + \frac{1}{1-x}$$

Un cas particulier de la proposition précédente est celui où $I = A \times B$ est un produit cartésien. En effet,

$$A \times B = \bigcup_{a \in A} \{a\} \times B$$

et les ensembles $\{a\} \times B$ sont disjoints deux à deux. Ainsi,

$$\sum_{(a,b) \in A \times B} u_{ab} = \sum_{a \in A} \left(\sum_{(a,b) \in \{a\} \times B} u_{ab} \right) = \sum_{a \in A} \left(\sum_{b \in B} u_{ab} \right)$$

On peut bien entendu renverser dans le calcul que nous venons de faire le rôle de A et B . On a ainsi le résultat ci-dessous, appelé le *théorème de Fubini*.

Proposition 10. THÉORÈME DE FUBINI

Soit $(u_{ab})_{a \in A, b \in B}$ une famille de réels positifs. Alors

$$\sum_{a \in A, b \in B} u_{ab} = \sum_{a \in A} \left(\sum_{b \in B} u_{ab} \right) = \sum_{b \in B} \left(\sum_{a \in A} u_{ab} \right)$$

Par la proposition 9, on a de plus le résultat suivant.

Proposition 11. La famille $(u_{ab})_{a \in A, b \in B}$ est sommable si et seulement si

- Pour tout $a \in A$, la famille $(u_{ab})_{b \in B}$ est sommable, de somme S_a .

- La famille $(S_a)_{a \in A}$ est sommable.

Et un résultat analogue en échangeant les rôles de A et B .

Exemple. Soit $x \in [0, 1[$. Pour $(p, q) \in \mathbb{N}^2$, posons $u_{pq} = x^{p+q}$. Par le théorème de Fubini, on a

$$\begin{aligned}
 \sum_{(p,q) \in \mathbb{N}^2} x^{p+q} &= \sum_{p \in \mathbb{N}} \left(\sum_{q \in \mathbb{N}} x^p x^q \right) \quad (\text{Fubini}) \\
 &= \sum_{p \in \mathbb{N}} \left(x^p \sum_{q \in \mathbb{N}} x^q \right) \quad (\text{linéarité}) \\
 &= \sum_{p \in \mathbb{N}} \left(x^p \frac{1}{1-x} \right) \quad (\text{série géométrique}) \\
 &= \frac{1}{1-x} \sum_{p \in \mathbb{N}} x^p \quad (\text{linéarité}) \\
 &= \frac{1}{(1-x)^2} \quad (\text{série géométrique})
 \end{aligned}$$

2 Familles sommables de nombres complexes

Dans la partie I, manipuler des sommes de réels positifs ne posait pas de difficultés particulière : nos « sommes » étaient finies, ou égales à $+\infty$.

Nous allons maintenant considérer des familles de nombres complexes. Nous allons voir que les résultats de la section I s'étendent pour la plupart, mais il va falloir être plus prudents.

2.1 Sommabilité

Définition 2. Soit $(u_i)_{i \in I}$ une famille de nombres complexes indexée par un ensemble I . On dit que la famille est *sommable* lorsque la famille de réels positifs $(|u_i|)_{i \in I}$ est sommable, c'est à dire lorsque

$$\sum_{i \in I} |u_i| < +\infty$$

Notation. On note $\ell^1(I)$ l'ensemble des familles sommables de nombres complexes (ou réels) indexées par I . Pour être plus précis, on peut utiliser la notation $\ell^1(I, \mathbb{R})$ ou $\ell^1(I, \mathbb{C})$.

Les deux résultats ci-dessous permettent de prouver qu'une famille est sommable.

Proposition 12. *Toute sous-famille d'une famille sommable est sommable.*

Démonstration. Même démonstration que pour les familles de réels positifs. On ne peut plus, évidemment, écrire des inégalités reliant la somme de la sous-famille et la somme de la famille. $\square$

Proposition 13. *Soient $(u_i)_{i \in I}$ une famille de nombres complexes et $(v_i)_{i \in I}$ une famille de réels positifs. On suppose que la famille $(v_i)_{i \in I}$ est sommable et que pour tout $i \in I$, $|u_i| \leq v_i$. Alors, la famille $(u_i)_{i \in I}$ est sommable.*

Démonstration. Notons $S = \sum_{i \in I} v_i$. Soit J une partie finie de I . Alors

$$\sum_{i \in J} |u_i| \leq \sum_{i \in J} v_i \leq S$$

La famille $(|u_i|)_{i \in I}$ est donc sommable. $\square$

Exemple. Soit $x \in]-1, 1[$. La famille $(x^{|n|})_{n \in \mathbb{Z}}$ est sommable. En effet, pour tout $n \in \mathbb{Z}$,

$$|x^{|n|}| = |x|^{|n|}$$

Or, $|x| \in [0, 1[$ et nous avons vu au début du chapitre que la famille $(|x|^{|n|})_{n \in \mathbb{Z}}$ est une famille sommable de réels positifs.

Il s'agit maintenant de donner un sens à la *somme* d'une famille sommable de nombres réels ou de nombres complexes.

2.2 Somme d'une famille sommable

Première étape, ramenons nous à ce que nous savons faire, à savoir des familles de réels positifs.

Pour tout réel x , notons $x^+ = \max(x, 0)$ et $x^- = -\min(x, 0)$. Remarquons que x^+ et x^- sont des réels positifs et que, de plus

$$x = x^+ - x^- \text{ et } |x| = x^+ + x^-$$

Proposition 14. *Soit $(u_i)_{i \in I}$ une famille de nombres réels. La famille $(u_i)_{i \in I}$ est sommable si et seulement si les familles $(u_i^+)_{i \in I}$ et $(u_i^-)_{i \in I}$ sont sommables.*

Démonstration. Supposons la famille $(u_i)_{i \in I}$ sommable, c'est à dire la famille $(|u_i|)_{i \in I}$. On a pour tout $i \in I$,

$$0 \leq u_i^+ \leq |u_i|$$

donc la famille $(u_i^+)_{i \in I}$ de réels positifs, dont chacun des termes est majoré par le terme d'une famille sommable, est elle aussi sommable. Il en est de même pour $(u_i^-)_{i \in I}$.

Supposons inversement les familles $(u_i^+)_{i \in I}$ et $(u_i^-)_{i \in I}$ sommables. La famille $(u_i^+ + u_i^-)_{i \in I}$ est ainsi sommable. Mais cette famille n'est autre que $(|u_i|)_{i \in I}$. $\square$

La définition de la somme d'une famille sommable de réels s'impose clairement.

Définition 3. Soit $(u_i)_{i \in I}$ une famille sommable de réels. La *somme* de la famille est

$$\sum_{i \in I} u_i = \sum_{i \in I} u_i^+ - \sum_{i \in I} u_i^-$$

Proposition 15. Soit $(u_i)_{i \in I}$ une famille de nombres complexes. La famille $(u_i)_{i \in I}$ est sommable si et seulement si les familles $(\operatorname{Re} u_i)_{i \in I}$ et $(\operatorname{Im} u_i)_{i \in I}$ sont sommables.

Démonstration. On a pour tout $i \in I$,

$$|\operatorname{Re} u_i| \leq |u_i|$$

$$|\operatorname{Im} u_i| \leq |u_i|$$

$$|u_i| \leq |\operatorname{Re} u_i| + |\operatorname{Im} u_i|$$

Grâce à ces inégalités, l'équivalence est facile à montrer. $\square$

La définition de la somme d'une famille sommable de nombres complexes s'impose.

Définition 4. Soit $(u_i)_{i \in I}$ une famille sommable de nombres complexes. La *somme* de la famille est

$$\sum_{i \in I} u_i = \sum_{i \in I} \operatorname{Re} u_i + i \sum_{i \in I} \operatorname{Im} u_i$$

Remarque 2. Le lecteur pardonnera au rédacteur de la définition l'emploi du i comme nombre complexe, mais aussi comme indice.

Ces deux définitions, parfaitement correctes, nous laissent un peu dubitatifs. Ne pourrait-on pas donner une définition de la somme d'une famille sommable de réels ou de complexes

qui n'oblige pas à décomposer les réels ou les complexes en petits morceaux ? C'est tout à fait possible.

Proposition 16. *Soit $(u_i)_{i \in I}$ une famille sommable de nombres complexes. Il existe un unique nombre complexe S tel que pour tout $\varepsilon > 0$, il existe F fini inclus dans I tel que pour tout G fini vérifiant $F \subset G \subset I$,*

$$\left| \sum_{i \in G} u_i - S \right| \leq \varepsilon$$

Le nombre complexe S est précisément la somme de la famille :

$$S = \sum_{i \in I} u_i$$

Démonstration. Supposons qu'il existe deux tels nombres complexes S et S' . Soit $\varepsilon > 0$. Soient F, F' associés à S, S' et ε . Posons $G = F \cup F'$. On a alors

$$|S - S'| \leq \left| \sum_{i \in G} u_i - S \right| + \left| \sum_{i \in G} u_i - S' \right| \leq 2\varepsilon$$

Ceci étant vrai pour tout $\varepsilon > 0$, il en résulte que $S = S'$.

On fait la démonstration de l'existence pour une famille de nombres réels. La preuve pour les familles de nombres complexes est analogue (on se ramène à des familles réelles).

Soit $S = \sum_{i \in I} u_i$. La famille de réels positifs $(u_i^+)_{i \in I}$ est sommable, de somme S^+ , et la famille de réels positifs $(u_i^-)_{i \in I}$ est sommable, de somme S^- .

Soit $\varepsilon > 0$. Il existe $F_1 \subset I$ fini tel que pour tout G fini, $F_1 \subset G \subset I$,

$$\left| \sum_{i \in G} u_i^+ - S^+ \right| \leq \frac{1}{2}\varepsilon$$

De même, il existe $F_2 \subset I$ fini tel que pour tout G fini, $F_2 \subset G \subset I$,

$$\left| \sum_{i \in G} u_i^- - S^- \right| \leq \frac{1}{2}\varepsilon$$

Posons $F = F_1 \cup F_2$. Pour tout G fini, $F \subset G \subset I$, on a

$$\begin{aligned} \left| \sum_{i \in G} u_i - S \right| &= \left| \left(\sum_{i \in G} u_i^+ - S^+ \right) - \left(\sum_{i \in G} u_i^- - S^- \right) \right| \\ &\leq \left| \sum_{i \in G} u_i^+ - S^+ \right| + \left| \sum_{i \in G} u_i^- - S^- \right| \\ &\leq \varepsilon \end{aligned}$$

□

Proposition 17. LINÉARITÉ DE LA SOMME

Soient $(u_i)_{i \in I}$ et $(v_i)_{i \in I}$ deux familles sommables de nombres complexes. Soient $\lambda, \mu \in \mathbb{C}$. Alors, la famille $(\lambda u_i + \mu v_i)_{i \in I}$ est sommable, et

$$\sum_{i \in I} (\lambda u_i + \mu v_i) = \lambda \sum_{i \in I} u_i + \mu \sum_{i \in I} v_i$$

Démonstration. Remarquons tout d'abord que pour tout $i \in I$,

$$|\lambda u_i + \mu v_i| \leq |\lambda| |u_i| + |\mu| |v_i|$$

Les familles $(|u_i|)_{i \in I}$ et $(|v_i|)_{i \in I}$ étant sommables, les résultats sur les familles sommables de réels positifs nous montrent que la famille $(|\lambda u_i + \mu v_i|)_{i \in I}$ l'est aussi. Il reste à montrer l'égalité.

Notons $S' = \sum_{i \in I} u_i$, $S'' = \sum_{i \in I} v_i$ et $S = \sum_{i \in I} \lambda u_i + \mu v_i$. Soit $\varepsilon > 0$. Il existe F fini inclus dans I tel que, pour tout $F \subset G \subset I$, G fini, on ait

$$\left| \sum_{i \in G} u_i - S' \right| \leq \frac{\varepsilon}{2(|\lambda| + 1)}$$

et

$$\left| \sum_{i \in G} v_i - S'' \right| \leq \frac{\varepsilon}{2(|\mu| + 1)}$$

De là,

$$\begin{aligned} \left| \sum_{i \in G} (\lambda u_i + \mu v_i) - (\lambda S' + \mu S'') \right| &= \left| \lambda \left(\sum_{i \in G} u_i - S' \right) + \mu \left(\sum_{i \in G} v_i - S'' \right) \right| \\ &\leq |\lambda| \left| \sum_{i \in G} u_i - S' \right| + |\mu| \left| \sum_{i \in G} v_i - S'' \right| \\ &\leq \frac{|\lambda| \varepsilon}{2(|\lambda| + 1)} + \frac{|\mu| \varepsilon}{2(|\mu| + 1)} \\ &\leq \varepsilon \end{aligned}$$

Par unicité de la somme d'une famille sommable, on conclut que

$$S = \lambda S' + \mu S''$$

□

Remarque 3. L'ensemble $\ell^1(I)$ est ainsi un $\mathbb{C}$ -espace vectoriel, sous-espace vectoriel de $\mathbb{C}^I$.

2.3 Séries et sommabilité

Proposition 18. SÉRIES ET SOMMABILITÉ

Soit $(u_n)_{n \in \mathbb{N}}$ une suite de nombres complexes. La famille $(u_n)_{n \in \mathbb{N}}$ est sommable si et seulement si la série $\sum u_n$ est absolument convergente. On a alors

$$\sum_{n \in \mathbb{N}} u_n = \sum_{n=0}^{\infty} u_n$$

Démonstration. La famille $(u_n)_{n \in \mathbb{N}}$ sommable si et seulement si la famille de réels positifs $(|u_n|)_{n \in \mathbb{N}}$ est sommable. Ceci équivaut, comme nous l'avons vu au I, à la convergence de la série $\sum |u_n|$, c'est à dire à la convergence absolue de la série $\sum u_n$.

Supposons qu'on a la convergence. Soient $S = \sum_{n \in \mathbb{N}} u_n$ et $S' = \sum_{n=0}^{\infty} u_n$. Montrons que $S' = S$. Donnons-nous $\varepsilon > 0$. Il existe un ensemble fini F inclus dans $\mathbb{N}$ tel que pour tout G fini, $F \subset G \subset \mathbb{N}$, on ait

$$(1) \quad \left| \sum_{i \in G} u_i - S \right| \leq \frac{1}{2} \varepsilon$$

Par ailleurs, il existe un entier N tel que pour tout $n \geq N$,

$$(2) \quad \left| \sum_{i=0}^n u_i - S' \right| \leq \frac{1}{2} \varepsilon$$

Prenons $n \geq N$ tel que $F \subset \llbracket 0, n \rrbracket$. Pour $G = \llbracket 0, n \rrbracket$, on a alors (1), et aussi (2). De là,

$$|S - S'| \leq \left| S - \sum_{i \in G} u_i \right| + \left| \sum_{i \in G} u_i - S' \right| \leq \varepsilon$$

Ceci étant vrai pour tout $\varepsilon > 0$, on en déduit que $S = S'$. □

2.4 Le théorème de Fubini

Proposition 19. SOMMATION PAR PAQUETS

Soit $(u_i)_{i \in I}$ une famille sommable de nombres complexes. On suppose que

$$I = \bigcup_{j \in J} I_j$$

où les ensembles I_j sont disjoints deux à deux. Alors,

- Pour tout $j \in J$, la famille $(u_i)_{i \in I_j}$ est sommable, de somme S_j .
- La famille $(S_j)_{j \in J}$ est sommable.
- On a

$$\sum_{i \in I} u_i = \sum_{j \in J} \sum_{i \in I_j} u_i$$

Démonstration. Nous admettrons ce résultat. $\square$

Remarque 4. Remarquons que l'on n'a plus ici l'équivalence que l'on avait pour les familles de réels positifs. Les paquets peuvent être sommables, ainsi que la famille des sommes des paquets, sans que la famille complète le soit. Prenons par exemple $I = \mathbb{N}$ et, pour tout $j \in \mathbb{N}$, $I_j = \{2j, 2j+1\}$. Posons, pour tout $i \in \mathbb{N}$, $u_i = (-1)^i$. Pour tout $j \in \mathbb{N}$, la famille $(u_i)_{i \in I_j}$ est évidemment sommable (famille finie), et sa somme est

$$S_j = (-1)^{2j} + (-1)^{2j+1} = 0$$

La famille $(S_j)_{j \in \mathbb{N}}$ est la famille nulle, elle est donc sommable. En revanche, la famille $(u_i)_{i \in \mathbb{N}}$ n'est pas sommable, puisque la série $\sum (-1)^i$ diverge grossièrement.

Comme dans le cas des familles de réels positifs, le cas où $I = A \times B$ fournit le *théorème de Fubini*.

Proposition 20. THÉORÈME DE FUBINI

Soit $(u_{ab})_{a \in A, b \in B}$ une famille sommable de nombres complexes. Alors,

- Pour tout $a \in A$, la famille $(u_{ab})_{b \in B}$ est sommable, de somme S_a .
- La famille $(S_a)_{a \in A}$ est sommable.
- On a

$$\sum_{a \in A, b \in B} u_{ab} = \sum_{a \in A} \left(\sum_{b \in B} u_{ab} \right)$$

Démonstration. Identique à celle faite pour les familles sommables de réels positifs. $\square$

Proposition 21. Soient $(a_i)_{i \in I}$ et $(b_{i'})_{i' \in I'}$ deux familles sommables de nombres complexes. Alors, la famille $(a_i b_{i'})_{(i,i') \in I \times I'}$ est sommable, et

$$\sum_{(i,i') \in I \times I'} a_i b_{i'} = \sum_{i \in I} a_i \sum_{i' \in I'} b_{i'}$$

Démonstration. On applique le théorème de Fubini à la famille $(u_{ii'} = a_i b_{i'})_{(i,i') \in I \times I'}$. On a alors

$$\begin{aligned} \sum_{i \in I, i' \in I'} u_{ii'} &= \sum_{i \in I} \left(\sum_{i' \in I'} a_i b_{i'} \right) \\ &= \sum_{i \in I} \left(a_i \sum_{i' \in I'} b_{i'} \right) \\ &= \sum_{i \in I} a_i \sum_{i' \in I'} b_{i'} \end{aligned}$$

□

Remarque 5. Ce résultat s'étend par récurrence à un nombre fini quelconque de familles sommables.

Exemple. Reprenons un exemple déjà vu plus haut. Soit $x \in]-1, 1[$. Les familles $(x^p)_{p \in \mathbb{N}}$ et $(x^q)_{q \in \mathbb{N}}$ sont sommables (série géométrique), donc la famille $(x^{p+q})_{(p,q) \in \mathbb{N}^2}$ l'est aussi, et

$$\sum_{(p,q) \in \mathbb{N}^2} x^{p+q} = \sum_{p \in \mathbb{N}} x^p \sum_{q \in \mathbb{N}} x^q = \frac{1}{(1-x)^2}$$

2.5 Un exemple moins évident

Soit $x \in]-1, 1[$. Considérons la famille $(x^{mn})_{(m,n) \in (\mathbb{N}^*)^2}$.

Montrons tout d'abord que cette famille est sommable. Le théorème de Fubini pour les familles de réels positifs nous dit que

$$\sum_{m,n \in \mathbb{N}^*} |x^{mn}| = \sum_{m \in \mathbb{N}^*} \sum_{n \in \mathbb{N}^*} (|x|^m)^n = \sum_{m \in \mathbb{N}^*} \frac{|x|^m}{1 - |x|^m}$$

Lorsque m tend vers l'infini,

$$\frac{|x|^m}{1 - |x|^m} \sim |x|^m$$

Or, la série $\sum |x|^m$ est une série géométrique convergente, donc

$$\sum_{m,n \in \mathbb{N}^*} |x^{mn}| < +\infty$$

Maintenant que nous savons que notre famille est sommable nous pouvons lui appliquer le théorème de sommation par paquets et ses corollaires.

L'ensemble des indices étant un produit cartésien, notre première idée est d'appliquer le théorème de Fubini.

On a pour tout $m \in \mathbb{N}^*$,

$$\sum_{n \in \mathbb{N}^*} x^{mn} = \sum_{n \in \mathbb{N}^*} (x^m)^n = \frac{x^m}{1 - x^m}$$

Ainsi, par le théorème de Fubini,

$$\sum_{(m,n) \in (\mathbb{N}^*)^2} x^{mn} = \sum_{m=1}^{\infty} \frac{x^m}{1 - x^m}$$

Regroupons maintenant les termes de façon différente. Écrivons

$$(\mathbb{N}^*)^2 = \bigcup_{k \in \mathbb{N}^*} I_k$$

où

$$I_k = \{(m, n) \in (\mathbb{N}^*)^2, mn = k\}$$

Les I_k étant des ensembles disjoints, on peut appliquer le théorème de sommation par paquets.

$$\sum_{(m,n) \in (\mathbb{N}^*)^2} x^{mn} = \sum_{k \in \mathbb{N}^*} \sum_{(m,n) \in I_k} x^{mn} = \sum_{k \in \mathbb{N}^*} |I_k| x^k$$

Remarquons que $|I_k| = d(k)$, où $d(k)$ est le nombre de diviseurs de k . Nous venons donc de montrer que pour tout $x \in]-1, 1[$, les séries $\sum d(n)x^n$ et $\sum \frac{x^n}{1-x^n}$ sont convergentes, et

$$\sum_{n=1}^{\infty} d(n)x^n = \sum_{n=1}^{\infty} \frac{x^n}{1-x^n}$$

3 Produit de Cauchy

3.1 Produit de Cauchy de deux séries

Définition 5. Soient $\sum u_n$ et $\sum v_n$ deux séries de nombres complexes. Le *produit de Cauchy* de ces deux séries est la série $\sum w_n$, où, pour tout $n \in \mathbb{N}$,

$$w_n = \sum_{k=0}^n u_k v_{n-k}$$

Proposition 22. *Le produit de Cauchy $\sum w_n$ de deux séries absolument convergentes $\sum u_n$ et $\sum v_n$ est une série absolument convergente. De plus,*

$$\sum_{n=0}^{\infty} w_n = \sum_{n=0}^{\infty} u_n \sum_{n=0}^{\infty} v_n$$

Démonstration. Pour $i, i' \in \mathbb{N}$, posons $z_{ii'} = u_i v_{i'}$. La famille $(u_i)_{i \in \mathbb{N}}$ et la famille $(v_{i'})_{i' \in \mathbb{N}}$ sont sommables. Par la proposition 16, la famille $(z_{ii'})_{(i,i') \in \mathbb{N} \times \mathbb{N}}$ est sommable, et

$$\sum_{i \in \mathbb{N}, i' \in \mathbb{N}} z_{ii'} = \sum_{i \in \mathbb{N}} u_i \sum_{i' \in \mathbb{N}} v_{i'}$$

Pour tout $n \in \mathbb{N}$, posons

$$I_n = \{(i, i') \in \mathbb{N} \times \mathbb{N}, i + i' = n\}$$

Les I_n sont disjoints deux à deux et leur réunion est $\mathbb{N} \times \mathbb{N}$. Par le théorème de sommation par paquets,

$$\sum_{i \in \mathbb{N}, i' \in \mathbb{N}} z_{ii'} = \sum_{n=0}^{\infty} \sum_{(i,i') \in I_n} z_{ii'}$$

Il reste à remarquer que

$$\sum_{(i,i') \in I_n} z_{ii'} = \sum_{i+i'=n} u_i v_{i'} = \sum_{i=0}^n u_i v_{n-i} = w_n$$

□

3.2 Un exemple

Soient $a, b \in \mathbb{C}$ tels que $|a| < 1$ et $|b| < 1$. Supposons $a \neq b$. On a pour tout $n \in \mathbb{N}$,

$$\frac{a^{n+1} - b^{n+1}}{a - b} = \sum_{k=0}^n a^k b^{n-k}$$

On reconnaît là le n ième terme du produit de Cauchy des séries $\sum a^n$ et $\sum b^n$, qui sont toutes les deux absolument convergentes. Ainsi, la série $\sum \frac{a^{n+1} - b^{n+1}}{a - b}$ est convergente, et

$$\sum_{n=0}^{\infty} \frac{a^{n+1} - b^{n+1}}{a - b} = \sum_{n=0}^{\infty} a^n \sum_{n=0}^{\infty} b^n = \frac{1}{(1-a)(1-b)}$$

Exercice. Nous aurions pu étudier la série de l'exemple ci-dessus dans le chapitre sur les séries, en utilisant uniquement des résultats élémentaires. Comment ?

3.3 L'exponentielle

En application de la proposition précédente, prenons l'exemple de l'exponentielle complexe. Soient $x, x' \in \mathbb{C}$. Les séries $\sum \frac{x^n}{n!}$ et $\sum \frac{y^n}{n!}$ sont absolument convergentes, de sommes

$$\sum_{n=0}^{\infty} \frac{x^n}{n!} = e^x \text{ et } \sum_{n=0}^{\infty} \frac{y^n}{n!} = e^y$$

Quel est leur produit de Cauchy ? C'est la série de terme général

$$w_n = \sum_{k=0}^n \frac{x^k}{k!} \frac{y^{n-k}}{(n-k)!} = \frac{1}{n!} \sum_{k=0}^n \binom{n}{k} x^k y^{n-k} = \frac{(x+y)^n}{n!}$$

Ainsi, par la proposition précédente,

$$e^{x+y} = e^x e^y$$

On retrouve donc la propriété de morphisme de l'exponentielle complexe.

4 Complément - Le cardinal des familles sommables

Le résultat que nous énonçons dans ce paragraphe n'est pas à connaître, il est cependant très instructif. Un ensemble *dénombrable* est un ensemble en bijection avec l'ensemble $\mathbb{N}$. On peut prouver, par exemple, que $\mathbb{Z}$ et $\mathbb{Q}$ sont dénombrables, mais que $\mathbb{R}$ ne l'est pas.

Est-il possible de considérer des familles sommables $(u_i)_{i \in I}$ de nombres complexes où, par exemple, $I = \mathbb{R}$? Oui, mais alors la plupart des u_i sont obligatoirement nuls. Précisément, nous allons montrer que seuls un nombre dénombrable des u_i sont différents de 0. Commençons par un lemme que nous admettrons.

Lemme 23. Soit $(E_n)_{n \in \mathbb{N}}$ une suite d'ensembles finis. Alors, $\bigcup_{n \in \mathbb{N}} E_n$ est un ensemble fini ou dénombrable.

Démonstration. Admise. $\square$

Venons-en au résultat qui nous intéresse.

Proposition 24. *Soit $(u_i)_{i \in I}$ une famille sommable de nombres complexes. Soit*

$$E = \{i \in I, u_i \neq 0\}$$

Alors, l'ensemble E est fini ou dénombrable.

Démonstration. Notons $S = \sum_{i \in I} |u_i|$. Pour tout entier $n \geq 0$, soit

$$E_n = \left\{ i \in I, |u_i| \geq \frac{1}{n+1} \right\}$$

Pour toute partie finie F de E_n ,

$$\frac{1}{n+1}|F| \leq \sum_{i \in F} |u_i| \leq \sum_{i \in E_n} |u_i| \leq \sum_{i \in I} |u_i| = S$$

Ainsi, $|F| \leq (n+1)S$. L'ensemble E_n ne contient aucun sous-ensemble fini de cardinal strictement supérieur à $(n+1)S$, donc E_n est lui aussi un ensemble fini. Remarquons maintenant que

$$E = \bigcup_{n \in \mathbb{N}} E_n$$

L'ensemble E est une réunion dénombrable d'ensembles finis, c'est donc un ensemble fini ou dénombrable. $\square$

Chapitre 37

Fonctions de deux variables

Dans tout ce chapitre, on se place dans le plan $\mathbb{R}^2$ muni du produit scalaire canonique. Lorsqu'on parle de norme, il s'agit de la norme euclidienne associée. La plupart des résultats présentés ici restent valables si l'on se place dans $\mathbb{R}^p$ (où p est un entier ≥ 2) et avec une norme quelconque.

1 Un peu de topologie

1.1 Disques

Définition 1. Soit $a \in \mathbb{R}^2$. Soit $r \in \mathbb{R}_+$. On appelle

- Disque ouvert de centre a et de rayon r l'ensemble

$$D(a, r) = \{x \in E, \|x - a\| < r\}$$

- Disque fermé de centre a et de rayon r l'ensemble

$$\overline{D}(a, r) = \{x \in E, \|x - a\| \leq r\}$$

- Cercle de centre a et de rayon r l'ensemble

$$C(a, r) = \{x \in E, \|x - a\| = r\}$$

Le cas où $r = 0$ est un peu à part. Les disques ouverts de rayon 0 sont vides, les disques fermés de rayon 0 sont les singletons. Dans la suite, nous supposons implicitement que le rayon des disques est non nul.

Remarque 1. En dimension 3 ou plus, on peut définir les mêmes notions. Mais on dit alors boule et sphère au lieu de disque et cercle.

Proposition 1. *Un disque est convexe.*

Démonstration. Soit $D = D(a, r)$ un disque (ouvert, pour fixer les idées). Soient $x, y \in D$ et $\lambda \in [0, 1]$. Soit $z = (1 - \lambda)x + \lambda y$. Alors

$$\|z - a\| = \|(1 - \lambda)(x - a) + \lambda(y - a)\|$$

On applique l'inégalité triangulaire :

$$\|z - a\| < (1 - \lambda)r + \lambda r = r$$

Ainsi, $z \in D$. $\square$

1.2 Ensembles ouverts

Définition 2. Soit $A \subset \mathbb{R}^2$. On dit que A est ouvert lorsque A est une réunion de disques ouverts.

Exemple. L'ensemble vide, $\mathbb{R}^2$, les disques ouverts, sont des ensembles ouverts.

Nous donnons maintenant une caractérisation plus manipulable des ouverts.

Proposition 2. Soit $A \subset \mathbb{R}^2$. L'ensemble A est ouvert si et seulement si pour tout point $a \in A$, il existe un réel $\varepsilon > 0$ tel que $D(a, \varepsilon) \subset A$.

Démonstration. Supposons A ouvert. Soit $a \in A$. Comme A est réunion de disques ouverts, on a $a \in D(b, r)$ pour un certain b et un certain r . On vérifie alors que

$$D(a, r - \|a - b\|) \subset D(b, r)$$

En posant $\varepsilon = r - \|a - b\|$, on a $D(a, \varepsilon) \subset A$.

Inversement, supposons que pour tout $a \in A$, on ait l'existence d'un $\varepsilon(a) > 0$ tel que $D(a, \varepsilon(a)) \subset A$. Alors,

$$A = \bigcup_{a \in A} D(a, \varepsilon(a))$$

donc A est ouvert. $\square$

Exercice. Montrons que le complémentaire d'un disque fermé est ouvert. Soient $c \in \mathbb{R}^2$ et $r \geq 0$. Soit $D = \overline{D}(c, r)$ le disque fermé de centre c et de rayon r . Posons

$$A = \mathbb{R}^2 \setminus D$$

Pour tout $x \in A$, on a $\|x - c\| > r$. Soit

$$\varepsilon_x = \|x - c\| - r > 0$$

Soit D_x le disque ouvert de centre x et de rayon ε_x . Pour tout $y \in D_x$, on a

$$\|x - c\| = \|(x - y) + (y - c)\| \leq \|x - y\| + \|y - c\| < \varepsilon_x + \|y - c\|$$

Ainsi,

$$\|y - c\| > \|x - c\| - \varepsilon_x = r$$

et donc $y \in A$. On en déduit que $D_x \subset A$. Par la proposition précédente, A est ouvert.

Exercice. Montrons qu'un disque fermé n'est jamais ouvert.

Soit $D = \overline{D}(c, r)$ le disque fermé de rayon $c \in \mathbb{R}^2$ et de rayon $r \geq 0$. Soit $a = c + (r, 0)$. On a $a \in D$ puisque $\|a - c\| = r$.

Donnons-nous maintenant $\varepsilon > 0$ quelconque. Considérons le disque ouvert $D' = D(a, \varepsilon)$ de centre a et de rayon ε . Soit $x = a + (\frac{1}{2}\varepsilon, 0)$. On a

$$\|x - a\| = \frac{1}{2}\varepsilon < \varepsilon$$

et donc $x \in D'$. Cependant,

$$\|x - c\| = r + \frac{1}{2}\varepsilon > r$$

et donc $x \notin D$. Ainsi, $D' \not\subset D$. Ceci étant vrai pour tout $\varepsilon > 0$, D n'est pas ouvert.

Exercice. Toute réunion d'ouverts est un ouvert, et toute intersection finie d'ouverts est un ouvert. En revanche, une intersection infinie d'ouverts peut fort bien ne pas être ouverte.

2 Continuité

2.1 Notion de continuité

Définition 3. Soit Ω un ouvert de $\mathbb{R}^2$. Soit $f : \Omega \longrightarrow \mathbb{R}$. Soit $a \in \Omega$. On dit que f est continue en a lorsque

$$\forall \varepsilon > 0, \exists \alpha > 0, \forall x \in \Omega, \|x - a\| \leq \alpha \implies |f(x) - f(a)| \leq \varepsilon$$

On dit que f est continue sur Ω lorsque f est continue en tout point $a \in \Omega$.

2.2 Exemples

Proposition 3. La norme euclidienne est continue sur $\mathbb{R}^2$.

Démonstration. Soit $a \in \mathbb{R}^2$. Pour tout $x \in \mathbb{R}^2$, on a

$$|||x|| - ||a||| \leq \|x - a\|$$

Donnons-nous un réel $\varepsilon > 0$. Soit $\alpha = \varepsilon$. Alors

$$\|x - a\| \leq \alpha \implies |||x|| - ||a||| \leq \varepsilon$$

d'où la continuité de la norme en a . $\square$

Proposition 4. Les applications $(x, y) \longmapsto x$ et $(x, y) \longmapsto y$ sont continues sur $\mathbb{R}^2$.

Démonstration. Exercice. $\square$

Proposition 5. *L'application $+: \mathbb{R}^2 \rightarrow \mathbb{R}$ définie par $+(x, y) = x + y$ est continue. On a des résultats analogues pour les autres opérations.*

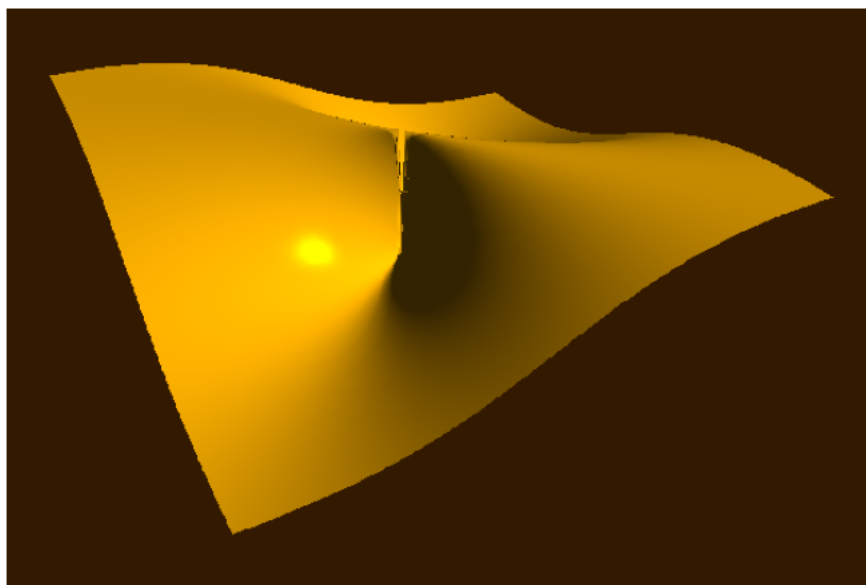
Démonstration. Exercice. $\square$

Exercice. Quel est l'ensemble de définition de la fonction « division » ? Où est-elle continue ?

Exemple. Soit $c \in \mathbb{R}$. Soit $f : \mathbb{R}^2 \rightarrow \mathbb{R}$ définie par $f(0, 0) = c$ et, pour tout $(x, y) \in \mathbb{R}^2$,

$$f(x, y) = \frac{xy}{x^2 + y^2}$$

Géométriquement, l'ensemble des points $(x, y, z) \in \mathbb{R}^3$ tels que $z = f(x, y)$ est une surface. En voici une représentation graphique.



On voit que pour $x = y = 0$ la surface présente un aspect « pincé ». Effectivement, montrons que f n'est pas continue en $(0, 0)$. Pour cela, posons pour tout $(x, y) \neq (0, 0)$, $x = r \cos \theta$ et $y = r \sin \theta$, de sorte que $r = \|(x, y)\|$. On a

$$f(x, y) = \frac{r^2 \cos \theta \sin \theta}{r^2} = \frac{1}{2} \sin(2\theta)$$

Ainsi,

$$f(x, y) - f(0, 0) = \frac{1}{2} \sin(2\theta) - c$$

Choisissons θ tel que $\sin \theta \neq 2c$. Soit $\varepsilon = \frac{1}{2} |\frac{1}{2} \sin(2\theta) - c|$. Soit $\alpha > 0$. Pour tout $r \leq \alpha$, le point $(x, y) = (r \cos \theta, r \sin \theta)$ vérifie $\|(x, y)\| \leq \alpha$. Pourtant, $|f(x, y) - f(0, 0)| > \varepsilon$. Ainsi, f n'est pas continue en $(0, 0)$.

2.3 Opérations sur les fonctions continues

Citons sans preuve les deux résultats suivants. Leur démonstration est analogue à celle faite pour les fonctions d'une variable réelle.

Proposition 6. Soit Ω un ouvert de $\mathbb{R}^2$. Soient $f, g : \Omega \rightarrow \mathbb{R}$. Soit $\lambda \in \mathbb{R}$. Soit $a \in \Omega$. On suppose que f et g sont continues en a . Alors

- $f + g$ est continue en a .
- λf est continue en a .
- fg est continue en a .
- Si $g(a) \neq 0$, il existe une disque ouvert $D \subset \Omega$ centré en a tel que g ne s'annule pas sur D , et $\frac{f}{g}$ est continue en a .

Proposition 7. Soient Ω un ouvert de $\mathbb{R}^2$ et I un intervalle de $\mathbb{R}$. Soient $f : \Omega \rightarrow I$ et $g : I \rightarrow \mathbb{R}$. Soit $a \in I$. On suppose que f est continue en a et g est continue en $f(a)$. Alors $g \circ f$ est continue en a .

Remarque 2. Rappelons l'exercice vu un peu plus haut : on y a montré que les applications $(x, y) \mapsto x$ et $(x, y) \mapsto y$ sont continues sur $\mathbb{R}^2$. Les théorèmes sur les opérations sur les fonctions continues prouvent alors que les formes linéaires sont continues sur $\mathbb{R}^2$. De même, les fonctions polynômes à plusieurs variables sont continues sur $\mathbb{R}^2$. Les fractions rationnelles à plusieurs variables sont continues là où elles sont définies.

Exemple. Soit $c \in \mathbb{R}$. Soit $f : \mathbb{R}^2 \rightarrow \mathbb{R}$ définie par $f(0, 0) = c$ et, pour tout $(x, y) \in \mathbb{R}^2$,

$$f(x, y) = \frac{xy}{x^2 + y^2}$$

Nous avons déjà vu que f n'est pas continue en $(0, 0)$. En revanche, par les théorèmes sur les fonctions continues, f est continue sur $\mathbb{R}^2 \setminus \{(0, 0)\}$.

2.4 Calculs pratiques

Comment prouver qu'une fonction de plusieurs variables est continue ? Les théorèmes sur les opérations sur les fonctions continues permettent de régler le problème, sauf peut-être

en quelques points (nous avons déjà vu l'exemple de $\frac{xy}{x^2+y^2}$). En ces points, un passage en « coordonnées polaires » permet souvent de régler la question, du moins pour les fonctions auxquelles nous nous intéressons.

Exemple. Prenons

$$f(x, y) = \frac{x^3 - y^3}{x^2 + y^2} \text{ si } (x, y) \neq (0, 0)$$

et $f(0, 0) = 0$. Par les opérations sur les fonctions continues, f est continue sur $\mathbb{R}^2 \setminus \{(0, 0)\}$. Pour l'étude en 0, posons $x = r \cos \theta$ et $y = r \sin \theta$, où $r \geq 0$ et $\theta \in \mathbb{R}$. On a donc $\|(x, y)\| = r$. Il vient alors

$$f(x, y) = \frac{r^3(\cos^3 \theta - \sin^3 \theta)}{r^2} = r(\cos^3 \theta - \sin^3 \theta)$$

Donc,

$$|f(x, y)| \leq 2r = 2\|(x, y)\|$$

Ainsi, pour tout $\varepsilon > 0$, $|f(x, y)| \leq \varepsilon$ dès que $\|(x, y)\| \leq \frac{\varepsilon}{2}$. La fonction f est continue en $(0, 0)$.

3 Dérivées partielles

3.1 Dérivée selon un vecteur

Soit Ω un ouvert de $\mathbb{R}^2$. Soit $f : \Omega \rightarrow \mathbb{R}$. Soit $a \in \Omega$. Soit $u \in \mathbb{R}^2$.

Proposition 8. Pour tout réel t assez petit, on a $a + tu \in \Omega$.

Démonstration. L'ensemble Ω est ouvert. Il existe donc $\varepsilon > 0$ tel que $D(a, \varepsilon) \subset \Omega$. On a donc $a + tu \in \Omega$ dès que $|t|\|u\| < \varepsilon$. $\square$

Définition 4. On dit que f est dérivable en a selon le vecteur u lorsque la quantité

$$\frac{f(a + tu) - f(a)}{t}$$

admet une limite lorsque le réel $t \neq 0$ tend vers 0. On appelle alors dérivée de f en a selon le vecteur u la limite en question.

Notation. La dérivée de f en a selon u est notée $df(a)(u)$.

Exemple. Soit $f(x, y) = x^2 - y^2$. Soient $a = (x_0, y_0)$ et $u = (h, k)$. On vérifie aisément que pour tout $u \in \mathbb{R}^2$, $df(a)(u)$ existe et vaut $2x_0h - 2y_0k$. Remarquons que

$$df(a) : \mathbb{R}^2 \rightarrow \mathbb{R}$$

est ici une forme linéaire sur $\mathbb{R}^2$. Nous verrons au cours de ce chapitre que sous certaines conditions de régularité sur f , ce sera toujours le cas.

3.2 Dérivées partielles

On note (e_1, e_2) la base canonique du plan $\mathbb{R}^2$.

Définition 5. On appelle dérivées partielles de f en a les quantités $df(a)(e_1)$ et $df(a)(e_2)$, lorsqu'elles existent.

Notation. On note ces dérivées partielles $\frac{\partial f}{\partial x}(a)$ et $\frac{\partial f}{\partial y}(a)$, en sous-entendant que x et y sont les noms des « variables ».

Lorsque les dérivées partielles de f existent en tout point de l'ouvert Ω , on dispose donc des deux fonctions $\frac{\partial f}{\partial x} : \Omega \longrightarrow \mathbb{R}$ et $\frac{\partial f}{\partial y} : \Omega \longrightarrow \mathbb{R}$.

Définition 6. On dit que f est de classe $\mathcal{C}^1$ sur Ω lorsque $\frac{\partial f}{\partial x}$ et $\frac{\partial f}{\partial y}$ existent et sont continues sur Ω . On note $\mathcal{C}^1(\Omega)$ l'ensemble des fonctions de classe $\mathcal{C}^1$ sur l'ouvert Ω .

3.3 Calculs pratiques

Posons $a = (a_1, a_2)$. La dérivée partielle $\frac{\partial f}{\partial x}(a)$, lorsqu'elle existe, est

$$\frac{\partial f}{\partial x}(a) = \lim_{t \rightarrow 0} \frac{f(a_1 + t, a_2) - f(a_1, a_2)}{t}$$

De même,

$$\frac{\partial f}{\partial y}(a) = \lim_{t \rightarrow 0} \frac{f(a_1, a_2 + t) - f(a_1, a_2)}{t}$$

En réalité, dans la plupart des cas il n'est pas utile de considérer un taux d'accroissement. Par exemple, pour dériver « par rapport à x », tout se passe comme si on fixait la deuxième variable et que l'on dérivait la fonction d'une variable restante.

Exemple. Soit $f : \mathbb{R}^2 \longrightarrow \mathbb{R}$ définie par $f(0, 0) = 0$ et, pour tout $(x, y) \neq (0, 0)$,

$$f(x, y) = \frac{xy}{x^2 + y^2}$$

Soit $a = (x, y) \neq (0, 0)$. Alors, $\frac{\partial f}{\partial x}(a)$ existe et

$$\frac{\partial f}{\partial x}(a) = \frac{y(y^2 - x^2)}{(x^2 + y^2)^2}$$

Bref, on fixe y et on dérive par rapport à x , comme s'il n'y avait qu'une variable.

Reprenons maintenant ce qui se passe en $(0, 0)$. Pour voir si f a des dérivées partielles en $(0, 0)$ on ne peut évidemment plus dériver « bêtement ». On revient à la définition avec des taux d'accroissement. Pour tout $t \neq 0$,

$$\frac{f(0+t, 0) - f(0, 0)}{t} = \frac{f(t, 0) - f(0, 0)}{t} = 0$$

On en déduit $\frac{\partial f}{\partial x}(0, 0)$ existe et vaut 0. Un calcul analogue montre que la dérivée partielle par rapport à y en $(0, 0)$ existe et vaut aussi 0.

Ce résultat est remarquable si l'on se souvient que la fonction f n'est pas continue en 0. Ainsi, pour les fonctions de plusieurs variables, la « dérivabilité » n'entraîne pas la continuité.

Exercice. Montrer que la fonction f de l'exemple ci-dessus possède en $(0, 0)$ des dérivées suivant tout vecteur.

3.4 Opérations sur les dérivées

Une dérivée partielle est en fait une dérivée « normale ». On a donc des formules pour les dérivées partielles d'une somme, d'un produit, d'un quotient, ... analogues à celles déjà vues pour les fonctions d'une seule variable. Par souci de complétude, voici le résultat.

Proposition 9. Soit Ω un ouvert de $\mathbb{R}^2$. Soient $f, g : \Omega \rightarrow \mathbb{R}$. Soit $\lambda \in \mathbb{R}$. Soient $a \in \Omega$ et $u \in \mathbb{R}^2$. On suppose que f et g sont dérivables en a selon le vecteur u . Alors

- $f + g$ est dérivable en a selon le vecteur u , et

$$d(f + g)(a)(u) = df(a)(u) + dg(a)(u)$$

- λf est dérivable en a selon le vecteur u , et

$$d(\lambda f)(a)(u) = \lambda df(a)(u)$$

- fg est dérivable en a selon le vecteur u , et

$$d(fg)(a)(u) = df(a)(u)g(a) + f(a)dg(a)(u)$$

- Si g est non nulle dans un disque centré en a , alors $\frac{f}{g}$ est dérivable en a selon le vecteur u , et

$$d\left(\frac{f}{g}\right)(a)(u) = \frac{df(a)(u)g(a) - f(a)dg(a)(u)}{g(a)^2}$$

3.5 Développement limité à l'ordre 1

Proposition 10. DÉVELOPPEMENT LIMITÉ À L'ORDRE 1

Soit $f : \Omega \longrightarrow \mathbb{R}$ une fonction de classe $\mathcal{C}^1$. Soit $a = (x_0, y_0) \in \Omega$. Alors,

- Pour tout $u = (h, k) \in \mathbb{R}^2$, $df(a)(u)$ existe et

$$df(a)(u) = h \frac{\partial f}{\partial x}(a) + k \frac{\partial f}{\partial y}(a)$$

- Lorsque u tend vers $(0, 0)$, on a

$$f(a + u) = f(a) + df(a)(u) + o(\|u\|)$$

Démonstration. Théorème admis.

Explicitons tout de même le sens de $o(\|u\|)$. Soit $\varphi : D \longrightarrow \mathbb{R}$ une fonction définie sur un disque ouvert D centré en 0. On dit que $\varphi(u) = o(\|u\|)$ lorsque pour tout $\varepsilon > 0$, il existe $\alpha > 0$, tel que pour tout $u \in D$,

$$\|u\| \leq \alpha \implies |\varphi(u)| \leq \varepsilon \|u\|$$

où $\varepsilon(u)$ tend vers 0 lorsque u tend vers 0.

□

Remarquons également que ce résultat est extrêmement fort : l'existence ET la continuité des dérivées partielles, c'est à dire des dérivées selon les deux vecteurs de la base canonique, entraîne l'existence des dérivées selon **tous** les vecteurs. De plus, l'application $df(a) : \mathbb{R}^2 \longrightarrow \mathbb{R}$ est une **forme linéaire** sur $\mathbb{R}^2$.

Remarque 3. Nous avons vu plus haut que l'ensemble des points $(x, y, z) \in \mathbb{R}^3$ tels que $z = f(x, y)$ est une *surface*. Soit $a = (x_0, y_0)$. Soit $b = (x, y) = (x_0 + h, y_0 + k)$ un point de Ω voisin de a . On a par la proposition précédente

$$f(b) = f(a) + \alpha(x - x_0) + \beta(y - y_0) + o(\|(h, k)\|)$$

où $\alpha = \frac{\partial f}{\partial x}(a)$ et $\beta = \frac{\partial f}{\partial y}(a)$. Le sous-ensemble $\mathcal{P}$ de $\mathbb{R}^3$ d'équation cartésienne

$$(\mathcal{P}) \quad z = f(a) + \alpha(x - x_0) + \beta(y - y_0)$$

est un plan affine. C'est le *plan tangent* à la surface représentative de f au point $(x_0, y_0, f(x_0, y_0))$. Au voisinage de ce point, la surface représentative de f est, d'une certaine façon, « collée » au plan $\mathcal{P}$. Nous n'entrerons pas plus avant dans les détails.

Corollaire 11. Une fonction de classe $\mathcal{C}^1$ est continue.

Démonstration. Reprenons les notations ci-dessus. Lorsque u tend vers 0, ses composantes h et k également, donc $f(a+u)$ tend vers $f(a)$. C'est précisément la définition de la continuité de f en a . $\square$

3.6 Gradient

Définition 7. Soit $f \in \mathcal{C}^1(\Omega)$. Soit $a \in \Omega$. Le gradient de f en a est le vecteur de $\mathbb{R}^2$

$$\nabla f(a) = \left(\frac{\partial f}{\partial x}(a), \frac{\partial f}{\partial y}(a) \right)$$

Le développement limité à l'ordre 1 de f en a s'écrit maintenant

$$f(a+u) = f(a) + \langle \nabla f(a), u \rangle + o(\|u\|)$$

Dit autrement,

$$df(a)(u) = \langle \nabla f(a), u \rangle$$

Prenons $u \neq 0$ et supposons que $\nabla f(a) \neq 0$. On notons θ l'écart angulaire entre u et $\nabla f(a)$. On a alors

$$f(a+u) - f(a) = \|\nabla f(a)\| \|u\| \cos \theta + o(\|u\|)$$

En « négligeant » le $o(\|u\|)$, on constate ainsi que, pour des vecteurs u de norme fixée, $|f(a+u) - f(a)|$ est maximal lorsque $\theta = 0$ ou $\theta = \pi$, c'est à dire lorsque u et $\nabla f(a)$ sont colinéaires. Le gradient de f en a est ainsi la direction dans laquelle f croît le plus vite au voisinage de a .

3.7 Extrema locaux

Définition 8. Soit $f : \Omega \rightarrow \mathbb{R}$. Soit $a \in \Omega$. On dit que f admet un maximum local en a lorsqu'il existe $\varepsilon > 0$ tel que $\forall x \in D(a, \varepsilon), f(x) \leq f(a)$. On définit de même un minimum local, et un extremum local.

Proposition 12. Si f est de classe $\mathcal{C}^1$ sur Ω et admet un extremum local en $a \in \Omega$, alors toutes ses dérivées suivant tout vecteur (et en particulier ses dérivées partielles) s'annulent en a .

Démonstration. Soit $u \in \mathbb{R}^2$. Supposons que f a un minimum local en a . On a alors pour tout t assez petit $f(a + tu) \geq f(a)$. Donc, pour tout $t > 0$ assez petit, le taux d'accroissement $\frac{f(a+tu)-f(a)}{t}$ est positif. En passant à la limite, on en déduit

$$df(a)(u) \geq 0$$

En recommençant pour $t < 0$, on obtient l'inégalité inverse. $\square$

Exemple. La fonction définie sur $\mathbb{R}^2$ par $f(x, y) = x^2 + y^2$ a toutes ses dérivées partielles qui s'annulent en $(0, 0)$. On vérifie qu'inversement, ce point donne lieu à un minimum de f .

Exemple. En revanche, la fonction définie sur $\mathbb{R}^2$ par $f(x, y) = x^2 - y^2$ a toutes ses dérivées partielles qui s'annulent en $(0, 0)$, mais ce point ne donne pas lieu à un extremum local de f . En effet, $f(0, 0) = 0$. Pour tout $h > 0$,

$$f(h, 0) = h^2 > 0$$

donc f n'a pas un maximum local en $(0, 0)$. De même,

$$f(0, h) = -h^2 < 0$$

donc f n'a pas un minimum local en $(0, 0)$.

4 Composition

4.1 La règle de la chaîne

On se donne un intervalle I de $\mathbb{R}$ et un ouvert Ω de $\mathbb{R}^2$. Soit $\varphi : I \longrightarrow \Omega$ une fonction de classe $\mathcal{C}^1$. De façon précise, pour tout $t \in I$,

$$\varphi(t) = (x(t), y(t))$$

où $x, y : I \longrightarrow \mathbb{R}$ sont deux fonctions de classe $\mathcal{C}^1$ sur I .

Soit $f : \Omega \longrightarrow \mathbb{R}$ une fonction de classe $\mathcal{C}^1$. La fonction $g = f \circ \varphi : I \longrightarrow \mathbb{R}$ est donc une fonction réelle d'une variable réelle.

Proposition 13. RÈGLE DE LA CHAÎNE

La fonction g est de classe $\mathcal{C}^1$ sur I , et on a pour tout réel $t \in I$:

$$g'(t) = \frac{\partial f}{\partial x}(\varphi(t))x'(t) + \frac{\partial f}{\partial y}(\varphi(t))y'(t)$$

Démonstration. Soit $h \neq 0$ un réel assez petit pour que $t + h \in I$. On a

$$g(t + h) - g(t) = f(T + H) - f(T)$$

où $T = \varphi(t)$ et $T + H = \varphi(t + h)$. On fait un DL de f à l'ordre 1 en T :

$$f(T + H) - f(T) = df(T)(H) + o(\|H\|)$$

Remarquons maintenant que

$$H = \varphi(t + h) - \varphi(t) = h\varphi'(t) + o(h) = (hx'(t) + o(h), hy'(t) + o(h))$$

Donc,

$$df(T)(H) = (hx'(t) + o(h)) \frac{\partial f}{\partial x}(\varphi(t)) + (hy'(t) + o(h)) \frac{\partial f}{\partial y}(\varphi(t))$$

De plus, $o(\|H\|) = o(h)$, donc

$$\frac{g(t + h) - g(t)}{h} = (x'(t) + o(1)) \frac{\partial f}{\partial x}(T) + (y'(t) + o(1)) \frac{\partial f}{\partial y}(T) + o(1)$$

tend vers

$$x'(t) \frac{\partial f}{\partial x}(T) + y'(t) \frac{\partial f}{\partial y}(T)$$

lorsque h tend vers 0. $\square$

Comment retenir cette formule ? Nous avons déjà mis en place les « bonnes notations ». Remarquons que x et y désignent deux choses

- Le nom des variables relatives à la fonction f .
- Le nom des fonctions qui définissent φ .

Il n'y a pas de confusion possible. Lorsque x et y apparaissent au *numérateur* d'une dérivée, ce sont des fonctions. Lorsqu'elles apparaissent au *dénominateur*, ce sont des variables. De façon abrégée, la règle de la chaîne s'écrit

$$\frac{dg}{dt} = \frac{\partial f}{\partial x} \frac{dx}{dt} + \frac{\partial f}{\partial y} \frac{dy}{dt}$$

Bien évidemment, il faut préciser en quels points les dérivées sont calculées. C'est facile :

- Si c'est une dérivée par rapport à t , c'est au point t .
- Si c'est une dérivée par rapport à x ou y , c'est au point « (x, y) », qu'il faut comprendre comme étant $(x(t), y(t))$.

Exemple. Un point se déplace au cours d'un intervalle de temps I . Son abscisse est son ordonnée, x et y , sont deux fonctions du temps et définissent une fonction $\varphi : I \rightarrow \mathbb{R}^2$. Pour tout $t \in I$, $\varphi(t) = (x(t), y(t))$.

Comment trouver les points de la trajectoire qui sont le plus près (ou le plus loin) de l'origine O ? Il suffit de considérer la fonction $f : \mathbb{R}^2 \rightarrow \mathbb{R}$ définie par

$$f(x, y) = x^2 + y^2$$

On cherche alors les extrema de la fonction $g : t \mapsto f(\varphi(t))$.

En supposant toutes les fonctions de classe $\mathcal{C}^1$, on a pour tout $t \in I$,

$$g'(t) = \frac{\partial f}{\partial x} \frac{dx}{dt} + \frac{\partial f}{\partial y} \frac{dy}{dt} = 2x \frac{dx}{dt} + 2y \frac{dy}{dt}$$

Pour prendre un exemple concret, le point décrit une demi-hyperbole. La fonction φ , définie sur $\mathbb{R}$, est donnée par

$$\varphi(t) = (x(t), y(t)) = (\operatorname{ch} t, \operatorname{sh} t)$$

On a donc

$$g'(t) = 2x \operatorname{sh} t + 2y \operatorname{ch} t = 4 \operatorname{ch} t \operatorname{sh} t$$

g' s'annule donc en $t = 0$. On a alors $\varphi(0) = (1, 0)$ et $f(0, 0) = 1$. Si la demi-hyperbole possède un point à distance minimale, il s'agit du point $(1, 0)$.

Il s'avère que ce point réalise effectivement le minimum des distances des points de la demi-hyperbole à O . En effet, pour tout $t \in \mathbb{R}$,

$$\operatorname{ch}^2 t + \operatorname{sh}^2 t = \operatorname{ch} 2t \geq 1$$

avec égalité si et seulement si $t = 0$.

4.2 Généralisation

Soient Ω et Ω' deux ouverts de $\mathbb{R}^2$. Soient $\varphi : \Omega' \rightarrow \Omega$ et $f : \Omega \rightarrow \mathbb{R}$. Soit

$$g = f \circ \varphi : \Omega' \rightarrow \mathbb{R}$$

Pour tout $(x, y) \in \Omega'$, on a

$$\varphi(x, y) = (u(x, y), v(x, y))$$

où $u, v : \Omega \rightarrow \mathbb{R}$. Supposons les fonctions f, u, v de classe $\mathcal{C}^1$ sur leurs domaines respectifs.

Proposition 14. *La fonction g est de classe $\mathcal{C}^1$ sur Ω' . Pour tout $a \in \Omega'$, on a*

$$\frac{\partial g}{\partial x}(a) = \frac{\partial f}{\partial u}(\varphi(a)) \frac{\partial u}{\partial x}(a) + \frac{\partial f}{\partial v}(\varphi(a)) \frac{\partial v}{\partial x}(a)$$

et

$$\frac{\partial g}{\partial y}(a) = \frac{\partial f}{\partial u}(\varphi(a)) \frac{\partial u}{\partial y}(a) + \frac{\partial f}{\partial v}(\varphi(a)) \frac{\partial v}{\partial y}(a)$$

Remarquer ici encore comment la règle de la chaîne s'applique. Dit rapidement, on a

$$\frac{\partial g}{\partial x} = \frac{\partial f}{\partial u} \frac{\partial u}{\partial x} + \frac{\partial f}{\partial v} \frac{\partial v}{\partial x}$$

et

$$\frac{\partial g}{\partial y} = \frac{\partial f}{\partial u} \frac{\partial u}{\partial y} + \frac{\partial f}{\partial v} \frac{\partial v}{\partial y}$$

4.3 Un exemple

Nous allons déterminer toutes les fonctions $f : \mathbb{R}^2 \longrightarrow \mathbb{R}$ de classe $\mathcal{C}^1$ vérifiant l'équation aux dérivées partielles

$$(\mathcal{E}) \quad \frac{\partial f}{\partial x} + \frac{\partial f}{\partial y} = 0$$

Considérons pour cela l'application $\varphi : \mathbb{R}^2 \longrightarrow \mathbb{R}^2$ définie par

$$\varphi(x, y) = (u(x, y), v(x, y)) = (x + y, x - y)$$

La fonction φ est clairement de classe $\mathcal{C}^1$ et bijective. Sa réciproque est aussi $\mathcal{C}^1$.

Soit f une fonction de classe $\mathcal{C}^1$ sur $\mathbb{R}^2$. Posons $g = f \circ \varphi^{-1}$, de sorte que $f = g \circ \varphi$. Il vient, pour tout $a = (x, y) \in \mathbb{R}^2$,

$$\begin{aligned} \frac{\partial f}{\partial x}(a) &= \frac{\partial g}{\partial u}(\varphi(a)) \frac{\partial u}{\partial x}(a) + \frac{\partial g}{\partial v}(\varphi(a)) \frac{\partial v}{\partial x}(a) \\ &= \frac{\partial g}{\partial u}(\varphi(a)) + \frac{\partial g}{\partial v}(\varphi(a)) \end{aligned}$$

De même,

$$\begin{aligned} \frac{\partial f}{\partial y}(a) &= \frac{\partial g}{\partial u}(\varphi(a)) \frac{\partial u}{\partial y}(a) + \frac{\partial g}{\partial v}(\varphi(a)) \frac{\partial v}{\partial y}(a) \\ &= \frac{\partial g}{\partial u}(\varphi(a)) - \frac{\partial g}{\partial v}(\varphi(a)) \end{aligned}$$

Ainsi,

$$\frac{\partial f}{\partial x}(a) + \frac{\partial f}{\partial y}(a) = 2 \frac{\partial g}{\partial u}(\varphi(a))$$

La fonction f est solution de l'équation $(\mathcal{E})$ si et seulement si pour tout $a \in \mathbb{R}^2$,

$$\frac{\partial g}{\partial u}(\varphi(a)) = 0$$

La fonction h étant bijective, cela équivaut à dire que pour tout $b \in \mathbb{R}^2$,

$$\frac{\partial g}{\partial u}(b) = 0$$

On peut montrer que cela équivaut à l'existence d'une fonction $G : \mathbb{R} \longrightarrow \mathbb{R}$ de classe $\mathcal{C}^1$ telle que pour tout $(u, v) \in \mathbb{R}^2$, $g(u, v) = G(v)$, ou encore, pour tout $(x, y) \in \mathbb{R}^2$,

$$f(x, y) = g(u(x, y), v(x, y)) = g(x + y, x - y) = G(x - y)$$

Pour résumer, les fonctions solutions de l'équation aux dérivées partielles sont les « fonctions de $x - y$ ».

Index

Index

- $\mathbb{C}$, 58
- $\mathbb{N}$, 32
- $\mathbb{R}$, 85
- $\mathbb{Z}$, 250

- adhérence, 156
- affiche, 59
- algorithme d'Euclide
 - dans $\mathbb{Z}$, 255
 - pour les polynômes, 309
- algorithme de Gram-Schmidt, 590
- algorithme du pivot de Gauss, 468
- angle d'une rotation, 605
- angle orienté, 605
- anneau, 140
- antécédent, 45
- application, 44
 - bilinéaire, 436
 - identité, 46
 - linéaire, 370
 - multilinéaire, 436
 - multilinéaire alternée, 437
 - multilinéaire antisymétrique, 437
 - multilinéaire symétrique, 437
 - réciproque, 47
- arc
 - cosinus, 201
 - sinus, 201
 - tangente, 203
- Archimède, propriété d', 86
- argument, 66
 - cosinus hyperbolique, 205
 - sinus hyperbolique, 205
 - tangente hyperbolique, 206
- associativité, 130
- asymptote, 358
- automorphisme, 137, 372

- barycentre, 212, 460
- base, 382
- base duale, 400
- bijection, 45
- borne supérieure, inférieure, 83
- branche parabolique, 358

- cardinal, 505
 - d'ensembles d'applications, 509
 - d'ensembles de parties, 511
 - d'ensembles en bijection, 505
 - d'un produit cartésien, 508
 - d'un sous-ensemble, 505
 - d'une différence, 507
 - d'une famille sommable, 666
 - d'une réunion, 507
 - d'une union disjointe, 506
- changement de variable, 234, 498
- Chasles, formule de, 486
- classe d'équivalence, 54
- coefficient binomial, 37
- cofacteur, 445
- comatrice, 454
- combinaison linéaire, 369
- commutativité, 130
- complémentaire, 27
- composée, 46
- congruence

- dans $\mathbb{R}$, 53
 - dans $\mathbb{Z}$, 52
- congruences, 267
- conjonction, 11
- conjugué, 60
- convexe, 90, 462
- corps, 140
- corps ordonné, 80
- cosinus, 62, 200
- cosinus hyperbolique, 204
- couple, 27
 - de variables aléatoires, 547
- covariance, 574
- croissance des pentes, 215
- cycle, 425
- degré
 - d'une fraction rationnelle, 319
- demi-tour, 598
- densité de $\mathbb{Q}$ et $\mathbb{R} \setminus \mathbb{Q}$ dans $\mathbb{R}$, 88
- dérivée
 - partielle, 676
 - selon un vecteur, 675
- différence de deux ensembles, 27
- dimension
 - d'espaces vectoriels isomorphes, 392
 - d'un espace vectoriel, 390
 - d'un produit d'espaces vectoriels, 393
 - d'un sous-espace vectoriel, 391
 - d'une somme, 393
 - des espaces d'applications linéaires, 394
- direction asymptotique, 358
- disjonction, 13
- disque, 670
- distance, 585
 - d'un vecteur à un sous-espace, 594
 - euclidienne, 586
- diviseur, 250
- diviseur d'un polynôme, 299
- division euclidienne
 - dans $\mathbb{Z}$, 252
 - des polynômes, 299
- double produit vectoriel, 610
- droite réelle achevée, 85
- dual d'un espace vectoriel, 399
- dérivée, 174
 - d'ordre supérieur, 176
 - d'un polynôme, 303
 - à droite, à gauche, 176
- déterminant
 - d'un endomorphisme, 449
 - d'une famille de vecteurs, 441
 - d'une matrice carrée, 441
 - d'une matrice triangulaire, 443
 - de Vandermonde, 448
 - développement, 445
 - opérations lignes-colonnes, 444
- développement asymptotique, 349
- développement limité, 336
- écart angulaire, 584
- écart-type, 568
- élément
 - absorbant, 131
 - inversible, 131
 - inversible, anneau, 141
 - neutre, 131
 - régulier, 132
 - simple, 320
- endomorphisme, 137, 372
 - orthogonal, 598
- ensemble, 22
 - ouvert, 671
 - vide, 24
- ensemble fini, 505
- ensembles égaux, 23
- équation du second degré, 72
- équations d'un hyperplan, 402
- équivalence, 11
- équivalence logique, 16
- espace affine, 458
- espace préhilbertien, 583
- espace probabilisé, 524
- espace vectoriel, 368
 - de dimension finie, 388
- espérance, 558
 - croissance, 564
 - d'un produit, 565

- linéarité, 559
- loi binomiale, 558
- loi de Bernoulli, 558
- positivité, 563
- variable aléatoire constante, 558
- et, 11
- Euler, formule d', 65
- événement, 520
 - certain, 523
 - impossible, 523
 - presque sûr, 564
 - élémentaire, 523
- événements
 - incompatibles, 523
 - indépendants, 537
- exponentielle, 194
- exponentielle complexe, 68
- extremum local, 183
- factorielle, 37
- famille, 28, 44
- famille génératrice, 382
- famille libre, 382
- famille liée, 382
- famille sommable
 - de nombres complexes, 656
 - de réels positifs, 648
- fonction
 - bornée, minorée, majorée, 149
 - caractéristique, 49
 - continue, 158
 - continue de deux variables, 672
 - continue par morceaux, 480
 - continue strictement monotone, 167
 - convexe, 213
 - de classe $\mathcal{C}^k, \mathcal{D}^k$, 177
 - de masse, 526
 - dominée par une autre, 330
 - en escalier, 477
 - indicatrice, 49
 - lipschitzienne, 168
 - monotone, 150
 - négligeable devant une autre, 330
 - paire, impaire, 151
 - polynôme, 300
 - puissance, 197
 - périodique, 152
 - rationnelle, 320
 - uniformément continue, 169
- fonctions équivalentes, 331
- forme linéaire, 372
- formes coordonnées, 400
- formule
 - d'Euler, 65
 - de Bayes, 536
 - de Moivre, 65
 - de Stirling, 350
 - de Taylor avec reste intégral, 500
 - de Taylor pour les polynômes, 303
 - de Taylor-Young, 335
 - de transfert, 564
 - du binôme, 38, 143
- formule de Chasles, 486
- formules de Cramer, 467
- fraction rationnelle, 318
- gradient, 679
- groupe, 133
 - abélien, 133
 - alterné, 431
 - cyclique, 420
 - orthogonal, 599
 - spécial-orthogonal, 602
 - symétrique, 422
- homothétie, 75
- hyperplan, 401
- hyperplan affine, 463
- identité, 46
- identité du parallélogramme, 586
- identités de polarisation, 586
- identités remarquables, 142
- image
 - d'un morphisme de groupes, 138
 - d'un élément, 44
 - d'une application linéaire, 376
 - directe, 49
 - directe d'un sous-espace vectoriel, 376

- réciproque, 49
 - réciproque d'un sous-espace vectoriel, 376
- image, nombre complexe, 59
- implication, 14
- inclusion, 23
- indépendants, vecteurs, 382
- inégalité
 - de Bienaymé-Tchebychev, 571
 - de Markov, 571
 - de Schwarz pour les intégrales, 487
 - de Schwarz pour les v.a., 566
 - de Schwarz pour les vecteurs, 583
 - de Taylor-Lagrange, 500
 - des accroissements finis, 185
- inflexion, 221
- injection, 45
- intégrale
 - fonction en escalier, 478
- intersection, 26
 - d'hyperplans, 403
 - de sous-espaces affines, 460
- intervalle, 90
 - image par une fonction continue, 166
- intégrale
 - fonction continue par morceaux, 481
- intégration par parties, 233, 497
- intérieur d'un intervalle, 185
- inverse, 131
- inversion, 433
- isomorphisme, 137, 372
- isométrie, 599
- issue d'une expérience, 523
- lemme
 - des coalitions, 554
- limite
 - d'une fonction, 157
 - d'une fonction monotone, 163
 - directionnelle, 162
 - infinie d'une suite, 97
 - réelle d'une suite, 94
- logarithme népérien, 192
- loi
 - binomiale, 529, 546
 - conditionnelle, 549
 - conjointe, 547
 - d'une variable aléatoire, 543
 - de Bernoulli, 545
 - hypergéométrique, 530
 - uniforme, 545
- loi de composition interne, 130
- lois de Morgan, 27
- lois marginales, 547
- majorant, 51
- matrice, 274
 - antisymétrique, 278
 - canonique d'une application linéaire, 413
 - carrée, 274
 - colonne, 274
 - d'un système, 464
 - d'un vecteur, 406
 - d'une application linéaire, 407
 - d'une famille de vecteurs, 406
 - de l'image d'un vecteur, 409
 - de passage, 411
 - élémentaire, 275
 - extraite, 414
 - invertible, 286
 - ligne, 274
 - orthogonale, 600
 - symétrique, 278
 - triangulaire, 288
- matrices
 - équivalentes, 413
 - semblables, 416
- maximum, 51
 - global, 150
 - local, 150
- mineur, 445
- minimum, 51
- minorant, 51
- module, 61
- Moivre, formule de, 65
- morphisme
 - d'anneaux, 142
 - de groupes, 136
- multiple, 250

- multiple d'un polynôme, 299
- multiples
 - dans un anneau, 141
 - dans un groupe, 135
- méthode des rectangles, 490
- méthode des trapèzes, 493
- négation, 13
- nombre complexe, 58
 - de module 1, 64
- nombre decimal, 87
- nombre premier, 263
- non, 13
- norme, 584
 - euclidienne, 584
- noyau
 - d'un morphisme de groupes, 138
 - d'une application linéaire, 376
- opération, 130
- orbite d'une permutation, 424
- ordre d'un élément, d'un sous-groupe, 420
- orientation, 453
- orthogonal
 - d'une partie, 589
- orthogonalité
 - d'une famille de vecteurs, 587
 - de deux ensembles, 589
 - de deux vecteurs, 587
- ou, 13
- paradoxe de Russell, 22
- paramétrage d'un segment, 211
- partie dense, 88
- partie entière
 - d'un réel, 87
 - d'une fraction rationnelle, 321
- partie imaginaire, 58
- partie réelle, 58
- parties d'un ensemble, 24
- parties polaires, 322
- partition, 54
- pas d'une subdivision, 476
- passage à la limite dans une inégalité
 - fonctions, 161
 - suites, 101
- période, 152
- permutation, 422
- plus grand commun diviseur
 - de deux entiers, 253
- plus grand élément, 51
- plus petit commun multiple
 - de deux entiers, 259
 - de deux polynômes, 312
- plus petit élément, 51
- point adhérent, 156
- pôle d'une fraction rationnelle, 320
- polynôme
 - d'interpolation de Lagrange, 314
 - irréductible, 304
 - scindé, 307
 - symétrique élémentaire, 308
- primitive, 230, 495
- probabilité, 523
 - conditionnelle, 532
 - uniforme, 527
- probabilités
 - composées, 533
 - totales, 534
- produit
 - d'une matrice par un scalaire, 274
 - de Cauchy, 664
 - de deux matrices, 276
 - mixte, 603
- produit cartésien, 28
- produit scalaire, 582
- produit vectoriel, 607
- produits
 - dans un anneau, 140
- projecteur, 380
 - orthogonal, 593
- prolongement, 45
- prolongement par continuité, 164
- proposition, 10
- propriété
 - presque sûre, 564
- propriété
 - vraie au voisinage d'un point, 156
- propriété d'Archimède, 251

- puissances
 - dans un anneau, 141
 - dans un groupe, 133
- quantificateur existentiel, 17
- quantificateur universel, 17
- racine
 - carrée d'un nombre complexe, 70
 - d'un polynôme, 301
 - d'une fraction rationnelle, 320
 - de l'unité, 73
 - multiple, 302
 - nième d'un nombre complexe, 74
- rang
 - d'un système, 464
 - d'une application linéaire, 396
 - d'une famille de vecteurs, 395
 - d'une matrice, 412
- récurrence
 - définir une suite, 41
 - faible, 33
 - forte, 39
 - à deux termes, 39
- réflexion, 598
- règle de la chaîne, 680
- relation
 - binaire, 50
 - d'ordre, 50
 - d'ordre total, 51
 - d'équivalence, 52
- restes d'une série convergente, 625
- restriction, 45
- réunion, 25
- rotation, 76, 602
- réciproque, 47
- segment, 90, 461
 - image par une fonction continue, 166
- série, 622
 - absolument convergente, 635
 - alternée, 637
 - convergente, divergente, 622
 - exponentielle, 624
 - géométrique, 622
 - harmonique, 633
- signature d'une permutation, 429
- similitude, 76
- sinus, 62, 200
- sinus hyperbolique, 204
- somme
 - de deux matrices, 274
 - de deux sous-groupes, 252
 - de Riemann, 490
 - de sous-espaces vectoriels, 379
- somme directe, 379
- sommes
 - dans un anneau, 140
- sous-anneau
 - caractérisation, 141
 - définition, 141
- sous-espace affine, 459
- sous-espace vectoriel
 - caractérisation, 375
 - définition, 374
 - engendré par une partie, 377
- sous-espaces affines parallèles, 460
- sous-espaces vectoriels supplémentaires, 379
- sous-groupe
 - caractérisation, 136
 - de $\mathbb{Z}$, 252
 - définition, 136
- subdivision, 476
 - adaptée, 477, 480
 - plus fine qu'une autre, 477
 - pointée, 490
 - régulière, à pas constant, 477
- suite
 - arithmétique, 112
 - convergente, divergente, 96
 - géométrique, 112
 - monotone, 104
 - récurrente linéaire à 2 termes, 113
- suite extraite, 102
- suites adjacentes, 105
- sup de deux fonctions, 149
- support d'un cycle, 425
- surjection, 45
- symétrie

- orthogonale, 593
- symétrie, 381
- système
 - compatible, 464
 - de Cramer, 466
 - homogène, 464
 - linéaire, 281
- système complet d'événements, 523
- série
 - de Riemann, 622
- tangente, 62, 200
- tangente hyperbolique, 205
- théorème
 - de Fubini, 655, 662
 - de Pythagore, 588
 - de sommation par paquets, 661
 - des accroissements finis, 184
- théorème
 - de Bolzano-Weierstraß, 108
 - de Bézout dans $\mathbb{Z}$, 254
 - de Bézout pour les polynômes, 310
 - de caractérisation séquentielle des limites, 161
 - de composition des limites, 160
 - de d'Alembert-Gauss, 305
 - de Fermat (petit), 267
 - de Gauss dans $\mathbb{Z}$, 254
 - de Gauss pour les polynômes, 311
 - de Heine, 170
 - de Lagrange, 421
 - de nullité de l'intégrale, 487
 - de Rolle, 184
 - des segments emboîtés, 107
 - des valeurs intermédiaires, 165
 - du rang, 396
 - fondamental du calcul intégral, 497
- théorème d'encadrement des limites
 - fonctions, 161
 - suites, 101
- tirages
 - avec remise, 529
 - sans remise, 530
- trace
 - d'un endomorphisme, 415
 - d'une matrice carrée, 415
- translation, 75, 458
- transposition, 425
- transposée d'une matrice, 277
- trigonométrique, forme, 67
- univers, 520
- valeur absolue, 82
- variable aléatoire
 - centrée, 568
 - réduite, 568
- variable aléatoire, 542
- variables aléatoires
 - décorrélées, 575
- variables aléatoires indépendantes, 551
- variance, 568
 - d'une somme, 577
- loi binomiale, 570
- loi de Bernoulli, 570
- loi hypergéométrique, 578
- loi uniforme, 569
- nullité, 569
- vecteur
 - unitaire, 585
- voisinage, 156