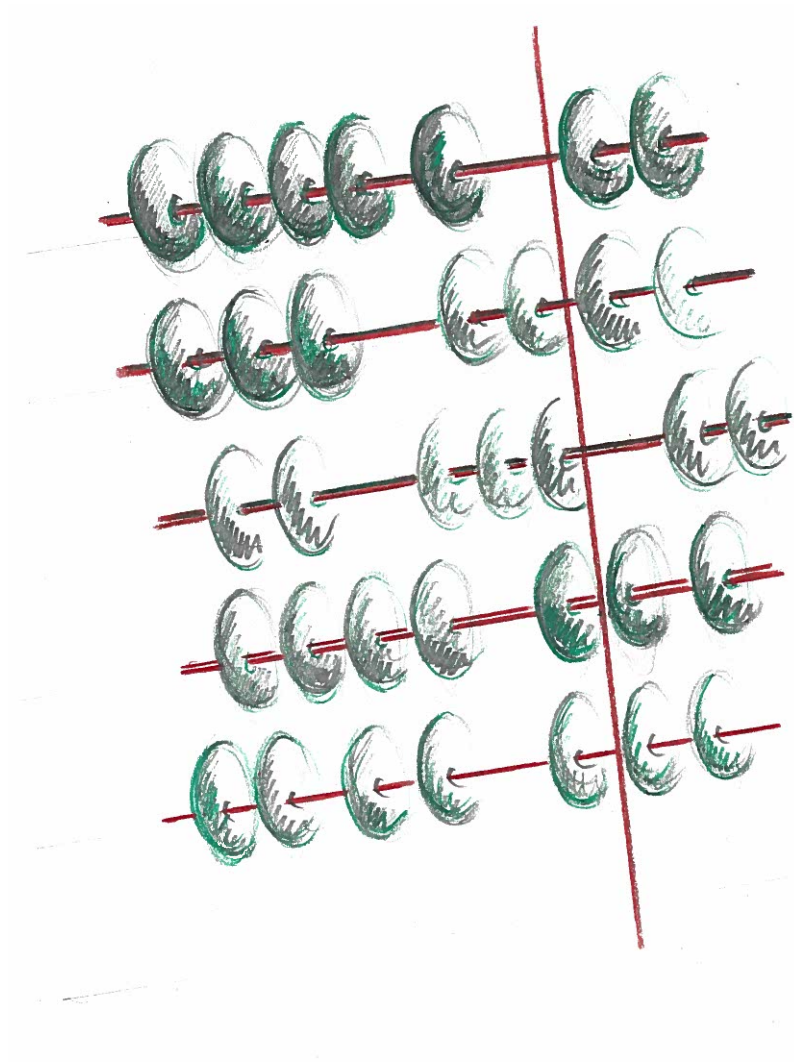


Algèbre linéaire avancée 1 pour physiciens

Marc Troyanov



Version du 1er février 2023).
(© marc.troyanov, EPFL)

Table des matières

1	Rudiments de logique formelle	5
1.1	De quoi parle-t-on en logique ?	5
1.2	La notion de proposition	6
1.3	Connecteurs logiques	6
1.4	Formules logiques	8
1.5	Tautologies et contradictions	10
1.6	Formules équivalentes	10
1.7	Prédicats et quantificateurs	12
2	Notions sur les ensembles	14
2.1	Qu'est-ce qu'un ensemble ?	14
2.2	Opérations sur les ensembles	16
2.3	Produit cartésien de deux ensembles	17
2.4	Relations	17
2.5	Applications	18
2.6	Images directe et inverse d'un sous-ensemble	20
2.7	Ensembles finis et cardinal	21
2.8	Ensembles infinis dénombrables et non dénombrables	22
2.9	Quelques mots sur les fondements de la théorie des ensembles	25
2.9.1	Le paradoxe de Russel	25
2.9.2	L'axiome de fondation	26
2.9.3	L'axiome de la réunion	26
2.9.4	L'axiome du choix	27
2.9.5	L'hypothèse du continu	27
3	Groupes, Anneaux et Corps	28
3.1	La notion de groupe	28
3.1.1	Sous-groupes et homomorphismes	29
3.2	Les permutations et le groupe symétrique	31
3.2.1	Signature d'une permutation.	33
3.3	Anneaux et corps	35
3.3.1	L'anneau $\mathbb{Z}/m\mathbb{Z}$	37
3.3.2	Carrés et formule quadratique.	40
3.4	Le corps des nombres complexes	41
3.4.1	Définition et propriétés de bases	41

3.4.2	Racine carrée d'un nombre complexe	43
3.4.3	Interprétation géométrique des nombres complexes	45
3.4.4	L'exponentielle d'un nombre complexe.	47
3.4.5	Logarithme naturel d'un nombre complexe.	49
4	Espaces vectoriels	51
4.1	Définitions et premiers exemples.	51
4.2	Sous-espaces vectoriels et applications linéaires	53
4.3	Familles libres, liées, génératrices et bases.	56
4.3.1	Définitions	56
4.3.2	Bases et dimension	59
4.3.3	Le théorème d'échange de Grassmann	62
4.3.4	Sommes directes et sous-espaces supplémentaires	63
4.3.5	Résumé des notions fondamentales sur l'indépendance linéaire	66
5	Applications linéaires et matrices	67
5.1	Rappels sur les applications linéaires	67
5.2	Opérations sur les applications linéaires	67
5.3	Le théorème du rang et ses conséquences	69
5.4	La matrice d'une application linéaire	72
5.4.1	Exemples de matrices d'applications linéaires	74
5.5	L'espace vectoriel des matrices	75
5.6	Produit matriciel	77
5.7	Matrices carrées : diagonale, transposée et matrices symétriques	79
5.8	Matrices inversibles.	80
5.9	Des applications linéaires vers les matrices.	81
5.10	Forme matricielle spéciale d'une application linéaire	84
5.11	Matrice de changement de bases	84
5.12	Des matrices vers les applications linéaires.	86
6	Systèmes linéaires	89
6.1	Sous-espaces affines d'un espace vectoriel	89
6.2	Systèmes d'équations linéaires	90
6.3	La méthode de Gauss-Jordan	91
6.4	Matrices élémentaires	94
6.5	Systèmes matriciels et inversion d'une matrice par la méthode de Gauss-Jordan	95
7	Déterminants	97
7.1	Déterminants des 2×2 matrices.	97
7.2	Déterminants des 3×3 matrices.	98
7.3	Définition générale du déterminant et premières propriétés	99
7.4	Théorème fondamental	103
7.5	Cofacteurs et formule de Laplace	104
7.6	Calcul de déterminants par l'algorithme de Gauss-Jordan	108

8	Vecteurs propres et valeurs propres	110
8.1	Introduction	110
8.2	L'algèbre des endomorphismes	110
8.3	Sous-espaces invariants	111
8.4	Valeurs propres et vecteurs propres	112
8.5	Le polynôme caractéristique	114
8.6	Endomorphismes et matrices diagonalisables	116
8.7	Multiplicités des valeurs propres et diagonalisabilité	119
8.8	Puissances d'une matrice diagonalisable	121
8.9	Application aux récurrences linéaires	123
A	Notions sur le polynômes	127
B	Un petit guide pour la diagonalisation	131
B.1	Aspects théoriques de la diagonalisation	131
B.2	Aspects pratiques de la diagonalisation	133
C	Sur les puissances d'une matrice	135
C.1	Matrices infra-périodiques	135
C.2	Puissances d'une matrice diagonalisable	136
C.3	Polynôme annulateur et puissances d'une matrice	137
C.4	Le cas d'une somme commutative de deux matrices	137
D	La notion de groupe quotient	139

Quelques livres de références :

- Robert J. Valenza, Linear Algebra - An Introduction to Abstract Mathematics, Springer, 1993.
- J. Grifone, Algèbre linéaire, Cepadues-Editions, 1990.
- R. Cairoli, Algèbre linéaire, Presses Polytechniques Universitaires Romandes, 2e édition 1999.
- K. Hoffman, R. Kunze, Linear Algebra, Prentice-Hall, second edition, 1971.

De nombreux exercices se trouvent dans les livres

- S. Lipschutz, Algèbre linéaire, série Schaum, Mc Graw-Hill.
- R. Dalang, A. Chabouni, Algèbre linéaire, PPUR.

Chapitre 1

Rudiments de logique formelle

1.1 De quoi parle-t-on en logique ?

La logique est la théorie des raisonnements valides ; elle ne peut pas décider si une affirmation est vraie ou fausse, mais seulement si un raisonnement est valide ou non. La forme la plus élémentaire du raisonnement logique est le *syllogisme* qui permet d'appliquer une hypothèse générale à un cas particulier pour en déduire une conclusion. L'exemple standard (peut-être depuis Aristote) est le suivant :

Tous les hommes sont mortels, or Socrate est un homme, donc Socrate est mortel.

La structure formelle de ce raisonnement s'appelle le *modus ponens*, on peut l'écrire sous une forme symbolique

$$[(p \Rightarrow q) \text{ et } p] \Rightarrow q,$$

que l'on peut lire «*si (p implique q) est vrai, et si p est vrai, alors q est vrai*».

Un raisonnement peut avoir l'apparence d'un raisonnement valide sans être correct. C'est parfois inoffensif :

Tous les hommes sont mortels, or Socrate est mortel donc Socrate est un homme,

ou au contraire avoir des conséquences potentiellement dramatiques (par exemple conduire à de graves erreurs judiciaires).

- *L'assassin roulait dans une voiture verte, or Socrate possède une voiture verte, donc le coupable est Socrate.*

Un tel raisonnement, qui semble valide sans être conclusif s'appelle un *sophisme*. Il serait illusoire de prétendre décrire toutes les formes possibles de sophismes. En voici deux exemples :

- *L'homo sapiens est apparu il y a plus de 200'000 ans, or Socrate est un homo sapiens, donc Socrate est apparu il y a plus de 200'000 ans.*
- *Paris est la capitale de la France, or Paris est un mot de 5 lettres, donc la capitale de la France est un mot de 5 lettres.*

Voici encore un raisonnement erroné, bien que sa conclusion nous semble être une évidence :

- *La salle de cours ne possède que deux portes d'entrée, Socrate est venu au cours en passant par la première porte, donc Socrate n'est pas passé par la seconde porte.*

Problème. Expliquer pourquoi ces quatre derniers raisonnements ne sont pas conclusifs.

1.2 La notion de proposition

L'unité fondamentale en logique est la *proposition*. Une proposition logique est une affirmation qu'on peut énoncer dans une langue naturelle (français, anglais) ou formelle (symbolisme mathématique, langage informatique) et à laquelle on peut attribuer une *valeur de vérité*.

On convient que les seules valeurs de vérité sont «vrai» et «faux» et qu'une proposition ne peut être à la fois vraie et fausse. Des exemples simples de propositions sont :

p_1 : Socrate est un homme.

p_2 : $65 + 4 = 69$.

p_3 : $65 \times 4 = 100$.

p_4 : La molécule d'eau contient deux atomes d'oxygène.

p_5 : L'électron a été découvert par Joseph Thompson.

Noter que p_3 et p_4 sont fausses, alors que les autres propositions sont vraies.

Nous admettons que dans certains cas nous ignorons si une proposition donnée est vraie ou fausse, mais si un énoncé n'est par essence ni vrai ni faux, alors cet énoncé n'est pas considéré comme une proposition logique. Par exemple une question du type «fera-t-il beau demain?» n'est pas une proposition (c'est une question), de même une injonction du type «tu dois étudier maintenant!» n'est pas une proposition. L'énoncé « x est un nombre pair » n'est *a priori* pas non plus une proposition, car sa valeur de vérité ne peut-être décidée sans information supplémentaire sur x .

Une *antinomie* (ou *paradoxe*) est un énoncé que l'on ne peut supposer vrai (car on en conclurait qu'il est faux) et qu'on ne peut supposer faux (car on conclurait qu'il est vrai). L'exemple fondamental d'antinomie est le *paradoxe d'Epiménide* (aussi appelé *paradoxe du menteur*), dont la forme la plus courte est :

Cet énoncé est faux !

Une variante est "*toutes les règles ont des exceptions*".

1.3 Connecteurs logiques

À la base de la logique, il y a donc la proposition, qui peut être vraie ou fausse. A partir d'une ou plusieurs propositions, on peut en formuler des nouvelles au moyen des *connecteurs logiques*. Il y a cinq connecteurs fondamentaux : la *négation*, la *conjonction*, la *disjonction*, l'*implication* et la *double implication*.

La *négation* de la proposition p se note $\neg p$ et se lit «non- p ». La proposition $\neg p$ est vraie si p est fausse et fausse si p est vraie, cela peut s'exprimer dans le tableau suivant, qu'on appelle *table de vérité*.

p	$\neg p$
V	F
F	V

La *conjonction* des propositions p et q se note $p \wedge q$ et se lit « p et q ». La proposition $p \wedge q$ est vraie si et seulement si p et q sont vraies, la table de vérité correspondante se présente donc ainsi :

p	q	$p \wedge q$
V	V	V
V	F	F
F	V	F
F	F	F

La *disjonction* des propositions p et q se note $p \vee q$ et se lit « p ou q ». La proposition $p \vee q$ est vraie si et seulement si p est vraie ou si q est vraie, la table de vérité correspondante est :

p	q	$p \vee q$
V	V	V
V	F	V
F	V	V
F	F	F

L'*implication logique* de p vers q se note $p \Rightarrow q$ et se lit « p implique q », ou «si p , alors q ». La proposition $p \Rightarrow q$ est vraie si et seulement si lorsque p est vraie alors q est également vraie, la table de vérité est :

p	q	$p \Rightarrow q$
V	V	V
V	F	F
F	V	V
F	F	V

Remarquons qu'en logique, on considère que toute proposition est conséquence d'une proposition fausse. La phrase «*si Chicago est la capitale de l'Italie, alors les hommes vivent en paix*» est considérée comme logiquement vraie (pour que cette proposition soit fausse, il faudrait à la fois que Chicago devienne la capitale de l'Italie et que, hélas, les hommes perpétuent leurs conflits). En mathématiques ce principe signifie que l'on peut déduire toute affirmation d'une proposition fausse, et donc, qu'un raisonnement rigoureux ne peut être conclusif s'il est basé sur une ou plusieurs hypothèses fausses. Voyons un exemple : la proposition «*si $4 + 4 = 5$, alors 39 est un nombre pair*» est du type $p \Rightarrow q$ avec l'hypothèse fausse $p : 4 + 4 = 5$. Cette proposition est donc vraie et voici une preuve :

Si $4 + 4 = 5$, alors $39 = 34 + 5 = 34 + (4 + 4) = 42 = 2 \times 21$ est divisible par 2, c'est donc un nombre pair.

La *double implication* de p vers q se note $p \Leftrightarrow q$ et se lit « p implique q et réciproquement», ou « p si et seulement si q ». La proposition $p \Leftrightarrow q$ est vraie si et seulement si p et q ont la même valeur de vérité (p et q sont en même temps vraies ou fausses).

p	q	$p \Leftrightarrow q$
V	V	V
V	F	F
F	V	F
F	F	V

Analogie arithmétique.

Il est intéressant de comparer les connecteurs logiques avec des opérations arithmétiques. Supposons que x et y désignent des nombres ne prenant que les valeurs 0 et 1, et définissons les opérations suivantes :

$$x' = 1 - x, \quad x \vee y = \max\{x, y\}, \quad x \wedge y = \min\{x, y\} = x \cdot y.$$

Les tables pour ces opérations sont

x	x'	x	y	$x \wedge y$	x	y	$x \vee y$
1	0	1	1	1	1	1	1
1	0	1	0	0	1	0	1
0	1	0	1	0	0	1	1
0	1	0	0	0	0	0	0

On voit qu'elles correspondent aux tables de vérité pour $\neg p$, $p \wedge q$ et $p \vee q$ (où 1 correspond à «vrai» et 0 correspond à «faux»).

1.4 Formules logiques

On appelle *formule logique* une expression construite à partir d'une famille finie de symboles $p, q, r, \dots$ qu'on appelle *variables propositionnelles*, en utilisant les connecteurs logiques ainsi que des parenthèses. On dit qu'une formule est *bien formée* si les règles de cohérence suivantes sont respectées :

- i.) Toute variable propositionnelle est une formule bien formée.
- ii.) Si F et G sont des formules bien formées, alors $\neg F$, $(F \wedge G)$, $(F \vee G)$, $(F \Rightarrow G)$ et $(F \Leftrightarrow G)$ sont aussi des formules bien formées.

Observons que toute parenthèse ouverte doit être refermée, les parenthèses peuvent être emboîtées. Si par exemple p, q, r, s sont des propositions, alors

$$[(p \vee (\neg q)) \Rightarrow (r \wedge ((\neg s) \vee q))]$$

est une formule bien formée, alors que

$$(p \vee (r \wedge) (\neg s \vee q)) \quad \text{et} \quad (p \Rightarrow q \vee r)$$

ne sont pas bien formées.

Remarque. Le rôle des parenthèses est d'éviter qu'une formule logique puisse avoir plusieurs interprétations. Lorsqu'il n'y a pas d'ambiguïté possible, on peut s'autoriser à supprimer certaines parenthèses.

Par exemple

$(p \Rightarrow q)$ peut s'écrire $p \Rightarrow q$.

$p \vee (q \vee r)$ peut s'écrire $p \vee q \vee r$.

$p \wedge (\neg q)$ peut s'écrire $p \wedge \neg q$.

En revanche on ne peut pas écrire $p \vee q \wedge r$ pour $p \vee (q \wedge r)$.

On peut faire une analogie utile avec les expressions algébriques. Dans cette analogie, une expression algébrique du type $2 \cdot x - 3 \cdot (x + y) + 2^z$ est l'analogue d'une formule logique. Les variables numériques $x, y, z, \dots$ sont les analogues des variables propositionnelles $p, q, r, \dots$, les opérateurs $+, -, \cdot, \dots$ sont les analogues des connecteurs logiques $\neg, \vee, \wedge, \Rightarrow$ et les règles de bon usage des parenthèses en algèbre sont analogues aux règles de cohérence pour les formules logiques bien formées.

Si l'on fixe les valeurs de vérité de chaque variable propositionnelle d'une formule logique, alors cette formule représente une proposition. On peut en déterminer la valeur de vérité, à partir des valeurs de vérité des variables propositionnelles qui la composent au moyen des tables de vérité, cela peut se faire en analysant la formule en sous-formules plus simples et en décrivant les tables de vérité de chaque sous-formule. Par exemple

$$\neg(\neg p \vee q)$$

est une formule bien formée, les sous formules sont $p, q, \neg p$ et $(\neg p \vee q)$. La table de vérité se construit ainsi :

p	q	$\neg p$	$(\neg p \vee q)$	$\neg(\neg p \vee q)$
V	V	F	V	F
V	F	F	F	V
F	V	V	V	F
F	F	V	V	F

Si la formule contient trois variables propositionnelles, la table de vérité correspondante contient 8 lignes. Par exemple la table de vérité de $((p \vee \neg q) \Rightarrow r)$ se présente ainsi :

p	q	r	$\neg q$	$(p \vee \neg q)$	$((p \vee \neg q) \Rightarrow r)$
V	V	V	F	V	V
V	V	F	F	V	F
V	F	V	V	V	V
V	F	F	V	V	F
F	V	V	F	F	V
F	V	F	F	F	V
F	F	V	V	V	V
F	F	F	V	V	F

Plus généralement, si une formule contient n variables propositionnelles, la table de vérité aura 2^n lignes.

Une table de vérité est donc une méthode (un « calcul ») qui permet de décider dans quels cas une formule propositionnelle est vraie en fonction des valeurs de vérité des propositions qui la composent.

Remarque. Dans la suite de ce chapitre, le mot *formule* signifiera *formule logique bien formée*.

1.5 Tautologies et contradictions

Une *tautologie* est une formule qui est vraie quelque soit la valeur de vérité prise par les variables propositionnelles qui la constituent. Une tautologie est donc vraie “dans tous les cas”. L’exemple le plus simple de tautologie est sans doute $(p \vee \neg p)$, un autre exemple est $(p \Rightarrow p)$, ce que l’on peut constater sur les tables de vérité respectives :

p	$\neg p$	$(p \vee \neg p)$	p	p	$(p \Rightarrow p)$
V	F	V	V	V	V
F	V	V	F	F	V

Un exemple de tautologie un peu moins banal est la formule $((p \Rightarrow q) \vee \neg q)$. Il s’agit d’une tautologie car le seul cas où $(p \Rightarrow q)$ est faux est lorsque p est vrai et q est faux, et donc $\neg q$ est vrai. De la même manière, $((p \Rightarrow q) \vee p)$ est aussi une tautologie. Cela se vérifie avec les tables de vérité :

p	q	$\neg q$	$(p \Rightarrow q)$	$((p \Rightarrow q) \vee \neg q)$	$((p \Rightarrow q) \vee p)$
V	V	F	V	V	V
V	F	V	F	V	V
F	V	F	V	V	V
F	F	V	V	V	V

Le modus ponens est la formule $((p \Rightarrow q) \wedge p) \Rightarrow q$, la table de vérité correspondante est

p	q	$(p \Rightarrow q)$	$(p \Rightarrow q) \wedge p$	$((p \Rightarrow q) \wedge p) \Rightarrow q$
V	V	V	V	V
V	F	F	F	V
F	V	V	F	V
F	F	V	F	V

c’est une tautologie puisque la dernière colonne ne contient que des V. La nature tautologique du modus ponens confirme qu’il s’agit bien d’un raisonnement valide.

Une *contradiction* est une formule qui est fausse quelque soit la valeur de vérité prise par les variables propositionnelles qui la constituent. Une contradiction est donc fausse “dans tous les cas”. L’exemple le plus simple de contradiction est sans doute $(p \wedge \neg p)$. Remarquons qu’une formule F est une contradiction si et seulement si $\neg F$ est une tautologie.

1.6 Formules équivalentes

On dit que deux formules F et G sont *équivalentes* si elles ont les mêmes variables propositionnelles et si elles prennent les mêmes valeurs de vérité en fonction de celles des variables propositionnelles. Elles sont donc équivalentes si et seulement si elles ont les mêmes tables de vérité.

Par exemple $p \Rightarrow q$, $\neg p \vee q$ et $\neg q \Rightarrow \neg p$ sont des formules équivalentes, ce que l’on constate sur les tables de vérité suivantes :

p	q	$\neg p$	$\neg q$	$p \Rightarrow q$	$\neg p \vee q$	$\neg q \Rightarrow \neg p$
V	V	F	F	V	V	V
V	F	F	V	F	F	F
F	V	V	F	V	V	V
F	F	V	V	V	V	V

Pour illustrer ces équivalences, considérons l'implication «*si tu dépasses la vitesse, alors tu risques une amende*», il est clair que cette phrase peut se reformuler «*ou bien tu ne dépasses pas la vitesse, ou bien tu risques une amende*».

Il est moins intuitif, mais néanmoins correct d'admettre que cette phrase peut se reformuler «*si tu ne risques pas d'amende, alors tu ne dépasses pas la vitesse*». Cet exemple montre que le langage naturel n'a pas exactement la même structure que la logique symbolique. Dans la langue française, l'implication $\Rightarrow$ est souvent perçue comme une causalité, le dépassement de vitesse est la cause du risque d'amende, mais la logique ne considère que les valeurs de vérité. Si on souhaite une interprétation intuitive de $\neg q \Rightarrow \neg p$, on peut l'énoncer : «*si tu n'as pas risqué d'amende, c'est que tu n'as pas dépassé la vitesse*», ou encore «*pour ne pas risquer d'amende, ne dépasses pas la vitesse!*». Dans le premier cas on introduit une notion de temporalité et dans le second cas il y a injonction. Le symbolisme logique est beaucoup plus pauvre que les langues naturelles, mais c'est précisément cette pauvreté qui permet de vérifier la validité d'un raisonnement en se débarrassant des ambiguïtés et des marges d'interprétations.

Définition. On dit que $(\neg q \Rightarrow \neg p)$ est la *contraposée* de $(p \Rightarrow q)$. Une implication est donc toujours logiquement équivalente à sa contraposée.

D'autres équivalences importantes sont les *Lois de De Morgan*. Elles disent que $\neg(p \vee q)$ est équivalent à $(\neg p \wedge \neg q)$, et que $\neg(p \wedge q)$ est équivalent à $(\neg p \vee \neg q)$, nous laissons au lecteur le soin de vérifier ces équivalences.

L'équivalence de deux formules F et G se note souvent $F \equiv G$, ainsi les lois de Morgan s'écrivent symboliquement :

$$\neg(p \vee q) \equiv (\neg p \wedge \neg q) \quad \text{et} \quad \neg(p \wedge q) \equiv (\neg p \vee \neg q).$$

Quelques équivalences de base à retenir sont :

Commutativité	$p \vee q \equiv q \vee p$	$p \wedge q \equiv q \wedge p$
Associativité	$(p \vee q) \vee r \equiv p \vee (q \vee r)$	$(p \wedge q) \wedge r \equiv p \wedge (q \wedge r)$
Distributivité	$p \vee (q \wedge r) \equiv (p \vee q) \wedge (p \vee r)$	$p \wedge (q \vee r) \equiv (p \wedge q) \vee (p \wedge r)$
Lois de De Morgan	$\neg(p \vee q) \equiv (\neg p) \wedge (\neg q)$	$\neg(p \wedge q) \equiv (\neg p) \vee (\neg q)$

Remarque. Toute équivalence logique $F \equiv G$ peut se reformuler en affirmant que la double implication $F \Leftrightarrow G$ est une tautologie. D'un point de vue "pratique", les symboles $\equiv$ et $\Leftrightarrow$ sont donc interchangeables. La différence est toutefois importante au point de vue de la théorie : lorsque on écrit $F \equiv G$, on signifie qu'il y a deux formules propositionnelles et que ces deux formules sont logiquement équivalentes (elles ont les mêmes tables de vérité), mais lorsqu'on écrit $F \Leftrightarrow G$, on signifie qu'il n'y a qu'une formule et que c'est une tautologie (sa table de vérité ne contient que des V).

1.7 Prédicats et quantificateurs

On a vu plus haut qu'un énoncé du type « x est un nombre pair» n'est a priori pas une proposition. La raison est que, n'ayant pas d'information sur l'objet « x », nous ne pouvons pas donner de valeur de vérité à cet énoncé. Nous pouvons toutefois considérer que notre énoncé est «une proposition dépendant d'une variable», cette idée intuitive nous conduit à la notion de *prédicat*.

Définition. Un *prédicat* est une propriété d'une ou plusieurs variables libres telle que si l'on fixe la valeur de la ou des variables libres, on obtient une proposition. En particulier, une unique valeur de vérité lui est alors attribuée.

Par exemple

- $p_1(x) : x$ est un nombre pair,
- $p_2(x) : x^2 = -1$,
- $p_3(x, y) : x < y$,
- $p_4(x, y) : 2x + 3y^2 - 1 = 0$,

sont des prédicats de une ou deux variables.

Si $p(x)$ est un prédicat à une variable, on peut lui associer la proposition «pour tout x , $p(x)$ est vrai». On abrège cette proposition sous la forme :

$$\forall x; p(x).$$

On peut aussi lui associer la proposition «il existe x tel que $p(x)$ est vrai», ce qu'on abrège ainsi :

$$\exists x; p(x).$$

Définitions. Le symbole $\forall$ se lit «pour tout» et s'appelle le *quantificateur universel*. Le symbole $\exists$ se lit «il existe» et s'appelle le *quantificateur existentiel*.

Remarquons que ces quantificateurs sont des généralisations des connecteurs logiques $\wedge$ et $\vee$. En effet, si la variable x ne prend par exemple que trois valeurs $x = 1$, $x = 2$, ou $x = 3$, alors $(\forall x; p(x))$ est équivalent à la proposition $p(1) \wedge p(2) \wedge p(3)$. De même $(\exists x; p(x))$ est équivalent à la proposition $p(1) \vee p(2) \vee p(3)$.

La négation d'une propriété obtenue par quantification est une généralisation des lois de De Morgan :

$$\neg (\forall x; p(x)) \text{ est équivalente à } (\exists x; \neg p(x))$$

et

$$\neg (\exists x; p(x)) \text{ est équivalente à } (\forall x; \neg p(x))$$

Par exemple la négation de «pour tout nombre x , x^2 est positif» est «il existe un nombre x tel que x^2 n'est pas positif».

Remarque. On utilise parfois un troisième quantificateur, noté $\exists!$, et qu'on lit *il existe un et un seul*.... Par exemple la proposition

$$\exists! x : x \text{ est un nombre réel et } x^3 = 5$$

signifie qu'*il existe un unique nombre réel dont le cube est égal à 5*¹.

Ce quantificateur peut s'exprimer à partir des deux autres quantificateurs sous la forme suivante :

$$(\exists! x : p(x)) \quad \equiv \quad (\exists x : p(x)) \wedge (\forall y : p(y) \Rightarrow (y = x)).$$

Remarque. Nous avons défini les quantificateurs logiques pour un prédicat d'une variable. On peut les appliquer de la même manière pour un prédicat de deux variables $p(x, y)$. Dans ce cas $\forall x; p(x, y)$ et $\exists x; p(x, y)$ sont des prédicats d'une seule variable.

Plus généralement les quantificateurs $\forall, \exists$ s'appliquent à tout prédicat de n variables $p(x_1, x_2, \dots, x_n)$ et donnent des prédicats de $n - 1$ variables.

1. On note ce nombre $\sqrt[3]{5}$.

Chapitre 2

Notions sur les ensembles

2.1 Qu'est-ce qu'un ensemble ?

Lorsqu'on considère qu'une collection d'objets forme en elle-même un nouvel objet, on obtient un *ensemble*. Le concept d'ensemble joue un rôle fondamental en mathématiques. Depuis la fin du XIX^{ème} siècle la logique symbolique et la théorie des ensembles se sont développées en parallèle et les deux théories se complètent et s'éclairent l'une l'autre.

Nous ne chercherons pas à donner une définition de la notion d'ensemble ; lorsqu'on dit qu'une famille ou une collection d'objets forme un ensemble nous n'avons pas vraiment défini le concept d'ensemble puisqu'il faudrait d'abord définir les mots "famille" ou "collection". Nous ne chercherons pas non plus à définir les mots "objet" ou "élément". Nous supposons donc que la notion d'ensemble est comprise intuitivement et nous décrivons quelques règles de base pour travailler avec les ensembles.

La première règle est qu'*il existe (au moins un) ensemble*. Nous supposons donc que la théorie des ensembles décrit une certaine réalité.

La seconde règle dit que *l'appartenance ou la non appartenance d'un objet est une proposition au sens de la logique* et obéit donc à la logique propositionnelle. On note

$$a \in A$$

pour dire que l'objet a appartient à l'ensemble A (on dit aussi que a est un *élément* de A). Ainsi ($a \in A$) est une proposition, il s'agit donc d'une affirmation qui est vraie ou fausse (et qui ne peut pas être vraie et fausse en même temps). On note $a \notin A$ pour dire que a n'est pas élément de A (cette notation est donc une abréviation de $\neg(a \in A)$).

La troisième règle est que *deux ensembles sont égaux si et seulement s'ils ont les mêmes éléments*, ce qu'on écrit formellement par

$$A = B \quad \text{si et seulement si} \quad \forall x : (x \in A \Leftrightarrow x \in B).$$

Cette règle s'appelle le *principe ou axiome d'extensionnalité*.

La façon la plus simple (et la plus naïve) de construire un ensemble est d'écrire la liste de ses éléments. On peut par exemple définir l'ensemble

$$A := \{1, 2, 3, 4, 5\}.$$

Le principe d'extensionnalité nous dit que l'on peut modifier l'ordre des éléments, ou en répéter certains, sans modifier l'ensemble. On a par exemple les égalités ensemblistes suivantes :

$$A = \{1, 2, 3, 4, 5\} = \{2, 5, 3, 1, 4\} = \{1, 2, 2, 3, 3, 4, 3, 5, \}.$$

Nous voudrions aussi pouvoir décrire un ensemble par une règle (un prédicat) qui précise les conditions pour qu'un élément en fasse partie. Par exemple on peut déclarer que l'ensemble A est l'ensemble des nombres compris entre 1 et 5, mais cette description n'est complète qu'à la condition qu'on précise qu'il s'agit de nombres entiers. La règle suivante nous dit d'une façon générale qu'à partir de la donnée d'un prédicat et d'un ensemble, on obtient un nouvel ensemble : Quatrième règle : si E est un ensemble et $p(x)$ est un prédicat, alors *on peut former l'ensemble des éléments de E qui vérifient la propriété p* . On note cet ensemble

$$\{x \in E \mid p(x)\}.$$

Cette règle s'appelle l'*axiome de compréhension* ou l'*axiome de séparation* (il permet de séparer dans l'ensemble E ceux qui vérifient la propriété x). Par exemple l'ensemble A défini plus haut peut s'écrire $A = \{x \in \mathbb{N} \mid 1 \leq x \leq 5\}$.

Ces quatre règles ne suffisent pas à établir une construction rigoureuse de la théorie des ensembles, mais elles nous donnent une bonne base pour utiliser le concept d'ensemble dans notre pratique mathématique. Voyons un premier résultat :

Théorème 2.1.1. *Il existe un ensemble qui ne contient aucun élément, et cet ensemble est unique.*

On note cet ensemble $\emptyset$ et on l'appelle l'*ensemble vide*.

Preuve. La première règle nous dit qu'il existe ou moins un ensemble, notons le E . La deuxième règle nous dit que $p(x) : x \neq x$ est un prédicat (i.e. pour tout ensemble x , la proposition $x \neq x$ vraie ou fausse). En appliquant la quatrième règle à cet ensemble et au prédicat $p(x) : x \neq x$, on obtient l'ensemble

$$\{x \in E \mid x \neq x\}$$

qui ne peut clairement contenir aucun élément (puisque $x = x$ est toujours vrai). Cela démontre l'existence d'un ensemble vide.

L'unicité découle immédiatement du principe d'extensionnalité.

□

En mathématiques, un rôle important est porté par les ensembles de nombres qui sont familiers au lecteur. Un ensemble particulièrement important est l'ensemble des entiers naturels¹

$$\mathbb{N} = \{0, 1, 2, 3, 4, 5, \dots\}.$$

Nous avons aussi les *intervalles entiers*

$$\llbracket n, m \rrbracket = \{n, n+1, \dots, m\}, \quad \llbracket n \rrbracket = \{1, 2, \dots, n\}$$

1. Notons que $\mathbb{N}$ n'est pas un ensemble fini et on ne peut pas le décrire ou l'explicitier en listant ses éléments (c'est la signification des points $\dots$). L'existence de $\mathbb{N}$ en tant qu'ensemble est donc une *hypothèse* et non une réalité qu'on pourrait prouver. Dans l'histoire de la philosophie, cette hypothèse renvoie à la notion d'*infini actuel*.

(où $1 \leq n \leq m$ sont des entiers) ainsi que l'ensemble des entiers relatifs $\mathbb{Z}$, l'ensemble des nombres rationnels $\mathbb{Q}$ et celui des nombres réels $\mathbb{R}$. Un rôle important sera aussi joué par l'ensemble $\mathbb{C}$ des *nombres complexes*.

Définitions. Un ensemble est un *singleton* s'il contient un seul élément. Deux ensembles sont *disjoints* si leur intersection est nulle (ils n'ont donc aucun élément en commun).

2.2 Opérations sur les ensembles

Voyons quelques définitions.

- i.) Si A et B sont des ensembles, on dit que B est une *partie* de A , ou un *sous-ensemble* de A , et on note $B \subset A$ (ou $B \subseteq A$) si tout élément de B est aussi un élément de A . On peut donc écrire cette définition ainsi :

$$B \subset A \quad \Leftrightarrow \quad [\forall x : x \in B \Rightarrow x \in A].$$

Attention : Il ne faut jamais confondre $A \in B$ et $A \subset B$!

- ii.) L'*intersection* des ensembles A et B est l'ensemble des éléments de A qui sont aussi éléments de B , on note cet ensemble $A \cap B$, ainsi

$$A \cap B = \{x \in A \mid x \in B\}.$$

Remarquons que cet ensemble est bien défini d'après la quatrième règle. On a clairement

$$x \in A \cap B \quad \Leftrightarrow \quad [(x \in A) \wedge (x \in B)].$$

- iii.) La *réunion* des ensembles A et B est l'ensemble des éléments qui appartiennent à A ou qui appartiennent à B , on note cet ensemble $A \cup B$. La définition formelle s'écrit

$$x \in A \cup B \quad \Leftrightarrow \quad [(x \in A) \vee (x \in B)].$$

- iv.) Le *complémentaire* de l'ensemble B dans l'ensemble A est l'ensemble des éléments de A qui n'appartiennent pas à B , on note cet ensemble $A \setminus B$:

$$A \setminus B = \{x \in A \mid x \notin B\}.$$

- v.) Si E est un ensemble, on note $\mathcal{P}(E)$ l'ensemble des parties de E :

$$A \in \mathcal{P}(E) \quad \Leftrightarrow \quad A \subset E.$$

Les propriétés de bases des opérations de réunion, intersection et complémentaire sont les suivantes :

Commutativité	$A \cup B = B \cup A$	$A \cap B = B \cap A$
Associativité	$(A \cup B) \cup C = A \cup (B \cup C)$	$(A \cap B) \cap C = A \cap (B \cap C)$
Distributivité	$A \cup (B \cap C) = (A \cup B) \cap (A \cup C)$	$A \cap (B \cup C) = (A \cap B) \cup (A \cap C)$
Lois de De Morgan	$E \setminus (A \cup B) = (E \setminus A) \cap (E \setminus B)$	$E \setminus (A \cap B) = (E \setminus A) \cup (E \setminus B)$

2.3 Produit cartésien de deux ensembles

Si A et B sont deux ensembles non vides, on forme un nouvel ensemble noté $A \times B$ dont les éléments sont les *couples* (a, b) avec $a \in A$ et $b \in B$. La notion de couple se distingue de celle d'ensemble à deux éléments par la règle suivante : $(a, b) = (a', b')$ si et seulement si $a = a'$ et $b = b'$ (en particulier $(a, b) \neq (b, a)$ sauf si $a = b$). L'ensemble $A \times B$ s'appelle le *produit cartésien* de A et B :

$$A \times B = \{(a, b) \mid a \in A \text{ et } b \in B\}.$$

Lorsque $A = B$, on note $A^2 = A \times A$.

Exemple : Si $A = \{+1, -1\}$ et $B = \{2, 4, 6\}$, alors

$$A \times B = \{(+1, 2), (+1, 4), (+1, 6), (-1, 2), (-1, 4), (-1, 6)\}$$

et

$$A^2 = \{(+1, +1), (+1, -1), (-1, +1), (-1, -1)\}.$$

Lorsque A ou B est l'ensemble vide, alors on convient que le produit cartésien est l'ensemble vide :

$$\emptyset \times B = A \times \emptyset = \emptyset.$$

On définit de même le produit cartésien de trois ensembles : si A, B et C sont des ensembles, alors $A \times B \times C$ est l'ensemble des *triples* (a, b, c) avec $a \in A$, $b \in B$ et $c \in C$; toujours avec la convention que $(a, b, c) = (a', b', c')$ si et seulement si $a = a'$, $b = b'$ et $c = c'$. Lorsque $A = B = C$, on note ce produit A^3 . La construction se généralise à un nombre quelconque d'ensembles.

2.4 Relations

Définition Une *relation* $\mathcal{R}$ sur un ensemble A est un sous-ensemble du produit cartésien $A \times A$:

$$\mathcal{R} \subset A \times A.$$

Lorsque une relation $\mathcal{R}$ sur A est donnée, il est habituel de noter $a\mathcal{R}b$ pour $(a, b) \in \mathcal{R}$ et de dire que a et b sont en *relation par rapport à* $\mathcal{R}$.

Dans la pratique, le symbole $\mathcal{R}$ est souvent remplacé par un symbole plus adéquat tel que

$$\sim, \simeq, \approx, \perp, <, \leq, \prec, \text{ etc.}$$

Considérons par exemple l'ensemble $A = \{1, 2, 3, 4, 5\}$, alors l'ensemble de couples suivant décrit une relation sur A :

$$\mathcal{R} = \{(1, 2), (1, 3), (1, 4), (1, 5), (2, 3), (2, 4), (2, 5), (3, 4), (3, 5), (4, 5)\} \subset A \times A.$$

On constate que l'ensemble $\mathcal{R}$ peut aussi être décrit par

$$\mathcal{R} = \{(x, y) \in A \times A \mid x < y\}.$$

En d'autres termes $x\mathcal{R}y$ si et seulement si $x < y$: la relation $\mathcal{R}$ n'est autre que la relation d'inégalité stricte sur A .

Il existe plusieurs types de relations importantes en mathématiques, en particulier les *relations d'ordre* et les *relations d'équivalence*.

Définition. Une relation $\preccurlyeq$ sur un ensemble A est une *relation d'ordre* si elle vérifie les conditions suivantes :

- i.) $x \preccurlyeq x$ pour tout $x \in A$.
- ii.) Si $x \preccurlyeq y$ et $y \preccurlyeq z$, alors $x \preccurlyeq z$ (transitivité).
- iii.) Si $x \preccurlyeq y$ et $y \preccurlyeq x$, alors $x = y$.

Exemples 1. L'ordre naturel sur les entiers $\mathbb{N}$ se définit par

$$x \leq y \Leftrightarrow \exists z \in \mathbb{N} \text{ tel que } y = x + z.$$

On vérifie facilement qu'il s'agit d'une relation d'ordre au sens de la définition ci-dessus.

2. La relation d'inclusion $\subset$ est une relation d'ordre sur l'ensemble $\mathcal{P}(X)$ des parties d'un ensemble X .

Définition. Une relation $\sim$ sur un ensemble X est une *relation d'équivalence* si elle vérifie les conditions suivantes :

- i.) $x \sim x$ pour tout $x \in X$ (réflexivité).
- ii.) Si $x \sim y$ alors $y \sim x$ (symétrie).
- iii.) Si $x \sim y$ et $y \sim z$, alors $x \sim z$ (transitivité).

D'autres symboles souvent utilisés pour les relations d'équivalences sont $\simeq, \approx, \cong, \equiv$ etc.

Exemples 3. En géométrie, la notion de *parallélisme* définit une relation d'équivalence sur l'ensemble $\mathcal{L}$ des droites dans le plan euclidien. On note $L \parallel L'$ cette relation.

4. Un exemple important de relation d'équivalence est donné par la construction suivante. Fixons un entier naturel $m \in \mathbb{N}$, on dit alors que deux entiers $x, y \in \mathbb{Z}$ sont *congrus modulo m* (ou simplement qu'ils sont *égaux modulo m*) si leur différence est un multiple de m : On note cette relation $x \equiv y \pmod{m}$ (ou simplement $x = y \pmod{m}$). Ainsi

$$x \equiv y \pmod{m} \iff (x - y) \in m \cdot \mathbb{Z}.$$

2.5 Applications

Une *application* f entre deux ensembles non vides X et Y est une règle qui à chaque élément x de X associe un et un seul élément y de Y . On note alors $y = f(x)$ ou $x \mapsto y$ et on dit que y est l'*image* de x par f . On dit aussi que x est une *préimage* de y (ou que c'est un *antécédent* de y). Notons qu'il peut exister zéro, une ou plusieurs préimage(s) pour un élément donné $y \in Y$.

La donnée des ensembles X et Y fait partie de la définition de l'application f , et on note souvent $f : X \rightarrow Y$. On dit que X est le *domaine* (ou la *source* ou l'*ensemble de départ*) de f et que Y est le *codomaine* (ou le *but* ou l'*ensemble d'arrivée*) de f . Une *fonction* est une application dont le but est une partie de l'ensemble des nombres réels $\mathbb{R}$ (ou des nombres complexes $\mathbb{C}$).

Voyons quelques exemples :

- (a) L'application *identité* de X est l'application $\text{Id}_X : X \rightarrow X$ telle que $\text{Id}_X(x) = x$ pour tout $x \in X$. C'est donc l'application "qui ne fait rien", l'image d'un élément x est l'élément x lui-même.

- (b) Si $f : X \rightarrow Y$ est une application et $A \subset X$ est un sous-ensemble, alors on définit une nouvelle application de A vers Y et qu'on note $f|_A : A \rightarrow Y$ par

$$f|_A(x) = f(x) \quad \text{pour tout } x \in A.$$

cette application s'appelle la *restriction* de f à l'ensemble A .

- (c) Si b est un élément de Y , on peut définir l'application $C_b : X \rightarrow Y$ telle que $C_b(x) = b$ pour tout $x \in X$. C'est l'*application constante* de valeur b .
- (d) Une suite de nombres réels $x_1, x_2, x_3, \dots$ peut être vue comme une fonction $f : \mathbb{N} \rightarrow \mathbb{R}$ en posant $f(k) = x_k$ pour tout $k \in \mathbb{N}$.
- (e) Si A est un sous-ensemble de X , on peut définir une fonction $\chi_A : X \rightarrow \{0, 1\}$ par la condition

$$\chi_A(x) = 1 \text{ si } x \in A \quad \text{et} \quad \chi_A(x) = 0 \text{ si } x \notin A,$$

cette fonction s'appelle la *fonction caractéristique* de l'ensemble A .

Définitions. Une application $f : X \rightarrow Y$ est *surjective* si tout élément de Y a au moins une préimage, elle est *injective* si tout élément de Y a au plus une préimage et elle est *bijective* si tout élément de Y a une et une seule préimage (c'est-à-dire si elle est à la fois injective et surjective). En notations logiques, on a

$$f : X \rightarrow Y \text{ est } \textit{surjective} \iff \forall y \in Y, \exists x \in X : f(x) = y.$$

et

$$f : X \rightarrow Y \text{ est } \textit{injective} \iff \forall x_1, x_2 \in X : [f(x_1) = f(x_2) \Rightarrow x_1 = x_2].$$

Définition. Si $f : X \rightarrow Y$ et $g : Y \rightarrow Z$ sont deux applications, alors on définit une nouvelle application de X vers Z , que l'on note $g \circ f : X \rightarrow Z$ par la condition

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

pour tout $x \in X$. Cette application s'appelle la *composition* de f et g .

Proposition 2.5.1. *La composition est une opération associative : si $f : X \rightarrow Y$, $g : Y \rightarrow Z$ et $h : Z \rightarrow W$ sont trois applications, alors*

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

Preuve. Pour tout $x \in X$ on a

$$((h \circ g) \circ f)(x) = (h \circ g)(f(x)) = h(g(f(x))) = h((g \circ f)(x)) = h \circ (g \circ f)(x).$$

Puisque cette égalité est vraie pour tout $x \in X$, on a $(h \circ g) \circ f = h \circ (g \circ f)$. □

Proposition 2.5.2. *Une application $f : X \rightarrow Y$ est bijective si et seulement s'il existe une application $g : Y \rightarrow X$ telle que $g \circ f = \text{Id}_X$ et $f \circ g = \text{Id}_Y$. Cette application est unique.*

Preuve. Supposons que f est bijective, alors pour tout $y \in Y$ il existe un et un seul élément $x \in X$ tel que $f(x) = y$. Notons cet élément $g(y)$, alors on a $f \circ g(y) = f(g(y)) = y$ par définition de g . Par conséquent $f \circ g = \text{Id}_Y$.

Pour vérifier que $g \circ f = \text{Id}_X$, on se donne un élément quelconque $x \in X$, et notons $z = g(f(x))$. Alors par définition de g , z est l'unique élément de X tel que $f(z) = f(x)$, donc $z = x$. Ainsi $g(f(x)) = x$ pour tout $x \in X$, ce qui montre que $g \circ f = \text{Id}_X$.

Inversément, supposons qu'il existe $g : Y \rightarrow X$ telle que $g \circ f = \text{Id}_X$ et $f \circ g = \text{Id}_Y$ et prouvons que f est bijective. L'application f est surjective car pour tout $y \in Y$ il existe au moins un élément $x \in X$ tel que $f(x) = y$, il suffit en effet de choisir $x = g(y)$.

L'application f est aussi injective, car si $f(x_1) = f(x_2)$, alors on a

$$x_1 = g(f(x_1)) = g(f(x_2)) = x_2.$$

Il nous faut encore démontrer l'unicité de l'application g . Supposons que l'application $h : Y \rightarrow X$ vérifie aussi $h \circ f = \text{Id}_X$ et $f \circ h = \text{Id}_Y$, alors on a en utilisant l'associativité de la composition :

$$h = h \circ \text{Id}_Y = h \circ (f \circ g) = (h \circ f) \circ g = \text{Id}_X \circ g = g.$$

□

Définition. L'application $g : Y \rightarrow X$ de la proposition précédente s'appelle l'*application inverse* ou l'*application réciproque* de f . On la note f^{-1} ou ${}^r f$.

Une application qui est injective ou surjective (mais non nécessairement bijective) admet un inverse partiel (appelé inverse à droite ou inverse à gauche).

Proposition 2.5.3. a) Une application $f : X \rightarrow Y$ est surjective si et seulement si il existe une injection $g : Y \rightarrow X$ vérifiant $f \circ g(y) = y$ pour tout $y \in Y$ (on dit que g est un inverse à droite de f .)

b) Une application $f : X \rightarrow Y$ est injective si et seulement si il existe une surjection $h : Y \rightarrow X$ vérifiant $h \circ f(x) = x$ pour tout $x \in X$ (on dit que h est un inverse à gauche de f).

Nous laissons la preuve en exercice. Soulignons que lorsque l'application est bijective, la proposition précédente implique que l'inverse à droite et l'inverse à gauche coïncident.

2.6 Images directe et inverse d'un sous-ensemble

Si une application $f : X \rightarrow Y$ est donnée, on peut définir deux nouvelles applications

$$f_* : \mathcal{P}(X) \rightarrow \mathcal{P}(Y) \quad \text{et} \quad f^* : \mathcal{P}(Y) \rightarrow \mathcal{P}(X)$$

de la manière suivante : si $A \subset X$, alors $f_*(A) \in \mathcal{P}(Y)$ est l'ensemble des images par f de tous les éléments de A :

$$f_*(A) := \{f(x) \mid x \in A\} = \{y \in Y \mid \exists x \in A : f(x) = y\}.$$

Si $B \in \mathcal{P}(Y)$, alors $f^*(B) \in \mathcal{P}(X)$ est l'ensemble de toutes les préimages des éléments de B :

$$f^*(B) := \{x \in X \mid f(x) \in B\}.$$

On dit que $f_*(A) \subset Y$ est l'*image directe* de $A \subset X$ par f et on la note habituellement simplement $f(A)$, et on dit que $f^*(B) \subset X$ est l'*image inverse* de $B \subset Y$ par f et on la note habituellement $f^{-1}(B)$.

Remarquer que $f : X \rightarrow Y$ est surjective si et seulement si $f^{-1}(\{y\}) \neq \emptyset$ pour tout élément $y \in Y$ et que f est injective si et seulement si $f^{-1}(\{y\})$ contient au plus un élément pour tout $y \in Y$. Finalement f est bijective si et seulement si $f^{-1}(\{y\})$ contient exactement un élément pour tout $y \in Y$.

2.7 Ensembles finis et cardinal

Un ensemble A est *fini* s'il est vide ou s'il existe un entier $n \in \mathbb{N}$ tel que A est en bijection avec

$$[n] = \{1, 2, \dots, n\} = \{k \in \mathbb{N} \mid 1 \leq k \leq n\}.$$

Sinon on dit que A est *infini*. On note alors $n = \text{Card}(A)$, c'est le *cardinal* de A (= le nombre d'éléments), si A est l'ensemble vide, on convient que $\text{Card}(\emptyset) = 0$.

Le cardinal des ensembles finis possède les propriétés suivantes :

Proposition 2.7.1. *Soient A et B deux ensembles finis, alors*

- (a) $\text{Card}(A) = \text{Card}(B)$ si et seulement s'il existe une bijection entre A et B .
- (b) $\text{Card}(A \cup B) = \text{Card}(A) + \text{Card}(B) - \text{Card}(A \cap B)$.
- (c) S'il existe une application $f : A \rightarrow B$ injective, alors $\text{Card}(A) \leq \text{Card}(B)$.
- (d) S'il existe une application $f : A \rightarrow B$ surjective, alors $\text{Card}(A) \geq \text{Card}(B)$.
- (e) $\text{Card}(A \times B) = \text{Card}(A) \cdot \text{Card}(B)$.
- (f) $\text{Card}(\mathcal{P}(A)) = 2^{\text{Card}(A)}$.
- (g) Si F est l'ensemble des applications de A vers B , alors $\text{Card}(F) = \text{Card}(B)^{\text{Card}(A)}$.

Ces propriétés sont plus au moins évidentes (le cardinal d'un ensemble s'obtient en comptant le nombre d'éléments), mais c'est tout de même un bon exercice d'écrire les preuves à partir de la définition du cardinal (qui est basée sur la notion de bijection).

Remarque. L'ensemble des applications de A vers B se note souvent B^A . Avec cette notation, on a donc

$$\text{Card}(B^A) = \text{Card}(B)^{\text{Card}(A)}.$$

Notons aussi que la propriété (f) ci-dessus peut-être vue comme un cas particulier de la propriété (g). En effet, il existe une bijection entre les ensembles $\mathcal{P}(A)$ et $\{0, 1\}^A$. Cette bijection est donnée par la fonction caractéristique (à tout $A \in \mathcal{P}(A)$ on associe la fonction $\chi_A \in \{0, 1\}^A$).

Une conséquence utile de la propriété (b) est que *Si A et B sont disjoints, alors*

$$\text{Card}(A \cup B) = \text{Card}(A) + \text{Card}(B).$$

Plus généralement, si $A_1, A_2, \dots, A_r$ sont des ensembles deux-à-deux disjoints², alors

$$\text{Card}(A_1 \cup A_2 \cdots \cup A_r) = \text{Card}(A_1) + \text{Card}(A_2) + \cdots + \text{Card}(A_r),$$

2. On dit que les ensembles A_i sont *deux-à-deux disjoints* si $A_i \cap A_j = \emptyset$ dès que $i \neq j$.

ce qu'on préfère écrire sous la forme

$$\text{Card}\left(\bigcup_{i=1}^r A_i\right) = \sum_{i=1}^r \text{Card}(A_i).$$

Un cas particulier de cette formule se rencontre lorsqu'on a une application $f : A \rightarrow B$ entre deux ensembles finis. L'ensemble A est alors réunion disjointe des préimages $f^{-1}(\{b\})$ pour les différents éléments de B :

$$A = \bigcup_{b \in B} f^{-1}(\{b\}) \quad \text{et} \quad f^{-1}(\{b\}) \cap f^{-1}(\{c\}) = \emptyset \quad \text{si} \quad b \neq c.$$

Par conséquent

$$\text{Card}(A) = \sum_{b \in B} \text{Card}(f^{-1}(\{b\})).$$

Théorème 2.7.2. *Soit $f : A \rightarrow B$ une application entre deux ensembles finis de même cardinal.*

- a) *Si f est injective, alors f est bijective.*
- b) *Si f est surjective, alors f est bijective.*

Nous laissons la preuve en exercice ; elle peut se déduire de la formule précédente.

2.8 Ensembles infinis dénombrables et non dénombrables

Définition. Un ensemble A est dit *dénombrable* s'il existe une bijection $f : \mathbb{N} \rightarrow A$. Intuitivement, un ensemble est dénombrable si on peut faire la liste (infinie) de ses éléments.

Exemple. L'ensemble des entiers relatifs $\mathbb{Z}$ est dénombrable. Un exemple de bijection $f : \mathbb{N} \rightarrow \mathbb{Z}$ est donné par

$$f(k) = \frac{k}{2} \quad \text{si } k \text{ est pair,} \quad f(k) = -\frac{k+1}{2} \quad \text{si } k \text{ est impair.}$$

Notons que cette bijection revient à lister les éléments de $\mathbb{Z}$ de la façon suivante :

$$\mathbb{Z} = \{0, -1, +1, -2, +2, -3, +3, \dots\}.$$

Proposition 2.8.1. A) *Tout sous-ensemble de $\mathbb{N}$ est ou bien fini ou bien dénombrable.*

B) *S'il existe une surjection $f : \mathbb{N} \rightarrow B$, alors B est ou bien fini ou bien dénombrable.*

Preuve. (A) Soit $A \subset \mathbb{N}$ un sous ensemble. Si A est fini il n'y a rien à montrer. Sinon on peut construire une application $f : \mathbb{N} \rightarrow A$ de la façon suivante : $f(1)$ est le plus petit élément de A , puis $f(2)$ est le plus petit élément de $A \setminus \{f(1)\}$ et, de façon inductive,

$$f(k) := \min(A \setminus \{f(1), f(2), \dots, f(k-1)\}).$$

Il est clair par construction que f est injective. Cette fonction est aussi surjective car $f(k) \geq k$ pour tout $k \in \mathbb{N}$, donc si $m \in A$, alors $m \in \{f(1), f(2), \dots, f(m)\}$.

(B) Pour tout $b \in B$, l'ensemble des préimages $f^{-1}(\{b\}) = \{k \in \mathbb{N} \mid f(k) = b\}$ est non vide car f est supposée surjective. Notons $g(b)$ le plus petit élément de $f^{-1}(\{b\})$. Cela définit une application

$g : B \rightarrow \mathbb{N}$, et cette application est injective car si $b_1 \neq b_2$, alors $f^{-1}(\{b_1\})$ et $f^{-1}(\{b_2\})$ sont disjoints, en particulier leurs plus petits éléments sont différents et donc $g(b_1) \neq g(b_2)$.

On a montré que B est en bijection avec un sous-ensemble $g(B) \subset \mathbb{N}$ et par le point (A) on en déduit que B est fini ou dénombrable. □

Corollaire 2.8.2. *L'ensemble $\mathbb{N} \times \mathbb{N}$ est dénombrable.*

Preuve. L'application $\varphi : \mathbb{N} \times \mathbb{N} \rightarrow \mathbb{N}$ définie par $\varphi(m, n) = 2^m \cdot 3^n$ est injective. □

Corollaire 2.8.3. *i) Si $A_1, A_2, \dots, A_n$ sont des ensembles dénombrables, alors $A_1 \times A_2 \times \dots \times A_n$ est dénombrable.*

ii) L'ensemble des nombres rationnels $\mathbb{Q}$ est dénombrable.

Preuve. L'affirmation (i) se démontre par récurrence à partir de la proposition précédente. L'affirmation (ii) est une conséquence du fait que l'application

$$\begin{aligned} \mathbb{Z} \times \mathbb{N}^* &\rightarrow \mathbb{Q} \\ (m, n) &\mapsto \frac{m}{n} \end{aligned}$$

est surjective (où $\mathbb{N}^* = \mathbb{N} \setminus \{0\}$). □

Théorème 2.8.4. *Soit A un ensemble, alors les conditions suivantes sont équivalentes :*

a) A est infini.

b) Il existe une injection de $\mathbb{N}$ dans A .

c) Il existe une bijection de A sur un sous-ensemble strict $B \subset A$.

Preuve. $(a) \Rightarrow (b)$. Soit A un ensemble infini. Choisissons un élément $a_1 \in A$, puis un élément $a_2 \in A \setminus \{a_1\}$, puis $a_3 \in A \setminus \{a_1, a_2\}$, et en général on choisit

$$a_k \in A \setminus \{a_1, a_2, \dots, a_{k-1}\}$$

Le processus ne s'arrête jamais car A est un ensemble infini. En posant $f(k) = a_k$ on obtient ainsi une injection $f : \mathbb{N} \rightarrow A$.

$(b) \Rightarrow (c)$. Maintenant on suppose qu'il existe une injection $f : \mathbb{N} \rightarrow A$, notons $a_k = f(k)$ et $S = f(\mathbb{N}) = \{a_0, a_1, a_2, \dots\} \subset A$. On peut alors définir une application $g : A \rightarrow A \setminus \{a_0\}$ de la façon suivante :

$$g(x) = \begin{cases} x & \text{si } x \notin S \\ a_{k+1} & \text{si } x = a_k \in S, \end{cases}$$

et cette application est clairement bijective (il est facile de construire son inverse $g^{-1} : A \setminus \{a_0\} \rightarrow A$).

$(c) \Rightarrow (a)$. La contraposée $\neg(a) \Rightarrow \neg(c)$ dit que si A est un ensemble fini, alors il n'existe aucune bijection de A sur un sous-ensemble strict, ce qui est évident. □

Avec le même type d'arguments, on montre facilement le résultat suivant :

Proposition 2.8.5. *Si A est un ensemble infini et $F \subset A$ est un ensemble fini, alors il existe une bijection de A vers $A \setminus F$.*

Par exemple il existe une bijection entre les intervalles $[0, 1]$ et $]0, 1[$.

On pourrait imaginer que tout ensemble infini est dénombrable. Ça n'est pas le cas :

Théorème 2.8.6 (Cantor). *L'intervalle $[0, 1]$ est non dénombrable.*

Preuve Donnons-nous une application $f : \mathbb{N} \rightarrow [0, 1]$ et prouvons qu'elle ne peut pas être surjective. Notons $a_k = f(k)$, c'est un nombre réel compris entre 0 et 1 et on peut écrire son développement décimal³ :

$$a_k = 0, \alpha_{k,1}\alpha_{k,2}\alpha_{k,3}\alpha_{k,4}\dots = \sum_{j=1}^{\infty} \frac{\alpha_{k,j}}{10^j}$$

où $\alpha_{k,j} \in \{0, 1, 2, 3, 4, 5, 6, 7, 8, 9\}$ sont les décimales de a_k . Pour prouver que f n'est pas surjective, il suffit d'exhiber un nombre $c \in [0, 1]$ tel que $\forall k : c \neq a_k$. L'idée de Cantor est de choisir le nombre

$$c := 0, \gamma_1\gamma_2\gamma_3\gamma_4\dots$$

dont les décimales γ_k sont données par

$$\gamma_j = \begin{cases} 4 & \text{si } \alpha_{j,j} = 5 \\ 5 & \text{si } \alpha_{j,j} \neq 5. \end{cases}$$

Le nombre c n'appartient pas à l'image de f , car si l'on suppose que $c = f(k)$ pour un certain entier k , alors γ_k doit être la k -ième décimale de a_k , c'est à dire $\gamma_k = \alpha_{k,k}$, ce qui contredit la définition du nombre c . □

Remarque Cet argument s'appelle l'*argument diagonal de Cantor*. La définition exacte du nombre c n'est pas essentielle, ce qui est importe est que $\gamma_k \neq \alpha_{k,k}$ pour tout $k \in \mathbb{N}$.

Corollaire 2.8.7. *Si un ensemble A contient un sous ensemble en bijection avec $[0, 1]$, alors il n'est pas dénombrable. En particulier $\mathbb{R}$, $\mathbb{R}^2$ et tout intervalle non vide $I \subset \mathbb{R}$ sont des ensembles non dénombrables.*

Nous concluons cette partie avec le résultat surprenant suivant :

Théorème 2.8.8. *Il existe une application injective de $[0, 1] \times [0, 1]$ vers $[0, 1]$.*

Preuve. On définit une application $f : [0, 1] \times [0, 1] \rightarrow [0, 1]$ de la façon suivante : Si $(x, y) \in [0, 1] \times [0, 1]$ on considère les développements décimaux (propres)

$$x = \sum_{i=1}^{\infty} \frac{a_i}{10^i}, \quad y = \sum_{i=1}^{\infty} \frac{b_i}{10^i},$$

3. pour éviter l'ambiguïté liée aux développements décimaux, on n'utilise dans ce type de raisonnement que les développements *propres*, c'est à dire les développements infinis sans suite illimitée de zéros successifs. Par exemple 0.3 s'écrit $0.2999\dots = 0.2\bar{9}$ (la seule exception à cette règle est $x = 0$).

et on pose $z = f(x, y)$ où z est le nombre donné par le développement décimal

$$z = \sum_{i=1}^{\infty} \frac{c_i}{10^i}$$

avec

$$c_i = \begin{cases} a_{(i+1)/2} & \text{si } i \text{ est impair,} \\ b_{i/2} & \text{si } i \text{ est pair.} \end{cases}$$

Par exemple si

$$(x, y) = (0.3808687\dots, 0.1398139\dots) \quad \text{alors} \quad f(x, y) = 0.31830988618379\dots$$

Cette application est injective car si $f(x, y) = f(x', y')$ alors nécessairement x et x' ont le même développement décimal (donc $x = x'$), de même y et y' ont le même développement décimal (donc $y = y'$), et par conséquent $(x, y) = (x', y')$. □

Théorème 2.8.9 (Théorème de Cantor-Bernstein-Schroeder). *Soient X et Y deux ensembles. Supposons qu'il existe des applications injectives $f : X \rightarrow Y$ et $g : Y \rightarrow X$ alors il existe une application bijective $h : X \rightarrow Y$.*

La preuve de ce théorème est délicate, nous ne la donnons pas ici.

Corollaire 2.8.10. *Il existe une bijection de $[0, 1] \times [0, 1]$ vers $[0, 1]$.*

Preuve. Le théorème précédent dit qu'il existe une application injective de $[0, 1] \times [0, 1]$ vers $[0, 1]$. Réciproquement il est facile de construire une application injective de $[0, 1]$ vers $[0, 1] \times [0, 1]$. Le théorème de Bernstein-Schroeder permet de conclure. □

Corollaire 2.8.11. *Il existe une bijection de $\mathbb{R}$ vers $\mathbb{R}^2$.*

Preuve. Cela découle immédiatement du résultat précédent puisque les ensembles $\mathbb{R}$ et $[0, 1]$ sont en bijection. □

2.9 Quelques mots sur les fondements de la théorie des ensembles

2.9.1 Le paradoxe de Russel

Nous n'avons pas donné de définition rigoureuse de ce qu'est un ensemble, mais nous nous sommes donnés quatre règles naturelles pour travailler avec les ensembles. Avec ces quatre règles, nous pouvons démontrer le théorème suivant :

Théorème 2.9.1. *Il n'existe pas d'ensemble dont les éléments sont tous les ensembles.*

Preuve La preuve procède par contradiction. On suppose qu'il existe un ensemble $\mathcal{E}$ tel que

$$A \in \mathcal{E} \Leftrightarrow A \text{ est un ensemble,}$$

alors notre deuxième règle dit que si $A \in \mathcal{E}$ et $B \in \mathcal{E}$, alors $A \notin B$ est une proposition (qui est donc ou bien vraie ou bien fausse). C'est également le cas si $A = B$. Ainsi pour tout $A \in \mathcal{E}$, la

proposition $A \notin A$ est ou bien vraie, ou bien fausse, et elle ne peut pas être à la fois vraie et fausse.

En utilisant la quatrième règle, on peut alors construire l'ensemble $\mathcal{R}$ de *tous les ensembles qui n'appartiennent pas à eux même*

$$\mathcal{R} := \{A \in \mathcal{E} \mid A \notin A\}.$$

Pour cet ensemble également, la proposition $\mathcal{R} \in \mathcal{R}$ doit être vraie ou fausse. Mais par définition de cet ensemble : $\mathcal{R} \in \mathcal{R}$ signifie $\mathcal{R} \notin \mathcal{R}$ et inversement $\mathcal{R} \notin \mathcal{R}$ implique $\mathcal{R} \in \mathcal{R}$.

Cette contradiction signifie qu'on ne peut pas admettre que $\mathcal{R}$ est un ensemble, et que donc l'hypothèse qu'il existe un ensemble $\mathcal{E}$ de tous les ensembles doit être rejetée. □

Remarque. La nature paradoxale de l' "ensemble" $\mathcal{R}$ a été découverte par Bertrand Russel en 1901. Le paradoxe de Russel nous enseigne que la théorie des ensembles n'est pas aussi simple qu'on pourrait imaginer à l'occasion d'un premier contact avec le sujet. Pour développer rigoureusement la théorie des ensembles, on se fonde habituellement sur une liste de 8 ou 9 axiomes (il y a quelques variantes) qui ont été proposés par Ernst Zermelo et Abraham Fraenkel à partir de 1908. Nous ne donnons pas ici la liste des axiomes de Zermelo-Fraenkel, ni les développements associés, mais nous discutons brièvement dans les paragraphes qui suivent quelques points importants de la théorie.

2.9.2 L'axiome de fondation

L'*axiome de fondation* dit que tout ensemble non vide $\mathcal{A}$, dont les éléments sont des ensembles, contient un élément qui est disjoint de $\mathcal{A}$:

$$\forall \mathcal{A} : \mathcal{A} \neq \emptyset \Rightarrow \exists C [(C \in \mathcal{A}) \wedge (C \cap \mathcal{A} = \emptyset)].$$

Avec cet axiome, on peut démontrer le théorème suivant :

Théorème 2.9.2. *Aucun ensemble n'est élément de lui-même.*

Preuve. Supposons qu'il existe un ensemble X tel que $X \in X$, alors le singleton $\{X\}$ ne vérifie pas l'axiome de fondation. On en conclut qu'un tel ensemble n'existe pas. □

Exercice. Montrer qu'il n'existe pas de paire d'ensembles X et Y tels que $X \in Y$ et $Y \in X$.

2.9.3 L'axiome de la réunion

L'axiome de la réunion dit que pour toute collection d'ensembles $\mathcal{A}$ (donc un ensemble dont les éléments sont eux-mêmes des ensembles), il existe un nouvel ensemble U qui contient tous les éléments des ensembles qui appartiennent à $\mathcal{A}$ et seulement ceux-ci. Formellement :

$$x \in U \iff \exists A \in \mathcal{A} \text{ et } x \in A.$$

On note souvent cet ensemble par $\cup \mathcal{A}$ et on l'appelle la *réunion* de $\mathcal{A}$.

2.9.4 L'axiome du choix

Etant donné une collection non vide d'ensembles non vides $\mathcal{A}$ (donc un ensemble qui contient au moins un élément et dont chaque élément est lui-même un ensemble non vide), on souhaite parfois “choisir” simultanément un élément particulier de chaque ensemble de cette collection. L'axiome du choix justifie cette construction, il s'énonce formellement ainsi :

$$\exists f : \mathcal{A} \rightarrow \cup \mathcal{A} \quad \text{tel que} \quad \forall A \in \mathcal{A} : f(A) \in A.$$

On appelle parfois cette fonction la *fonction de choix*.

Considérons un exemple simple. Supposons que $\mathcal{A} = \mathcal{P}(\mathbb{N}) \setminus \{\emptyset\}$ est l'ensemble des parties non vides de $\mathbb{N}$. Pour tout $A \in \mathcal{A}$ on note $f(A)$ le plus petit élément de A :

$$f(A) := \min A \in \mathbb{N}.$$

Alors f est un exemple de fonction de choix pour $\mathcal{A}$. L'axiome du choix nous dit qu'une telle fonction existe toujours, indépendamment de l'existence ou non d'une construction explicite.

Remarque. Cet axiome a été controversé durant les premières décennies de la théorie des ensembles. Il semblait en effet douteux que de faire une infinité (peut-être non dénombrable) de choix eût un sens. En raison de sa puissance et de son rôle structurant, cet axiome est aujourd'hui largement accepté en mathématiques.

2.9.5 L'hypothèse du continu

L'*hypothèse du continu* énonce que tout sous-ensemble infini de $\mathbb{R}$ est ou bien dénombrable (en bijection avec $\mathbb{N}$) ou bien en bijection avec $\mathbb{R}$. Georg Cantor pensait que cette hypothèse est vraie et qu'elle pouvait être démontrée (auquel cas elle deviendrait un théorème de la théorie des ensembles). Ce problème a occupé Cantor et d'autres chercheurs en logiques et en théorie des ensembles durant de nombreuses années, mais les travaux des logiciens Kurt Gödel (en 1938) puis ceux de Paul Cohen (en 1963) ont montré que l'hypothèse du continu est *indécidable* à partir des axiomes de Zermelo-Fraenkel. Cela signifie qu'on peut ajouter aux axiomes de ZF ou bien l'hypothèse du continu ou bien sa négation sans briser la cohérence de la théorie des ensembles, et, par conséquent, ni l'hypothèse du continu, ni sa négation ne peuvent être démontrée à partir des axiomes de ZF.

Chapitre 3

Groupes, Anneaux et Corps

3.1 La notion de groupe

Définition. Une *loi de composition interne* sur un ensemble E est une application

$$\mu : E \times E \rightarrow E.$$

Il est commode de noter $x * y = \mu(x, y)$, où $*$ est un méta-symbole représentant une opération concrète et qui peut être remplacé suivant le contexte par un symbole tel que $+$, $\cdot$, $\circ$, $\times$, $\vee$, $\wedge$, ...

Définition. Un *groupe* est un ensemble G muni d'une loi de composition interne $*$ qui est associative, pour laquelle il existe un élément neutre et telle que chaque élément de G admet un inverse.

(G1) $x * (y * z) = (x * y) * z, \forall x, y, z \in G$ (associativité).

(G2) Il existe $e \in G$ tel que $e * x = x * e = x, \forall x, y, z \in G$ (élément neutre).

(G3) Pour tout $x \in G$, il existe un élément $x' \in G$ tel que $x' * x = x * x' = e$.

L'élément $x' \in G$ s'appelle l'*inverse* de x . Suivant les cas il est souvent noté x^{-1} ou $-x$.

Le groupe $(G, *)$ est *abélien* ou *commutatif* s'il vérifie la propriété supplémentaire suivante :

(G4) $x * y = y * x$ pour tous $x, y \in G$.

Exemples.

- (1) $(\mathbb{R}, +)$ et $(\mathbb{R}^*, \cdot)$ sont des groupes.
- (2) Le plus petit groupe ne contient qu'un seul élément, qui est l'élément neutre : $G = \{e\}$ (on l'appelle le groupe trivial) (remarquons que l'ensemble vide ne peut pas être un groupe).
- (3) Il n'y a essentiellement qu'une façon de construire un groupe à deux éléments : $G = \{e, a\}$ avec e comme élément neutre. La loi de groupe est donnée dans la table

$*$	e	a
e	e	a
a	a	e

deux réalisations concrètes de ce groupe sont $\{+1, -1\}$ avec la multiplication usuelle et $\{0, 1\}$ avec l'addition modulo 2 :

$*$	$+1$	-1
$+1$	$+1$	-1
-1	-1	$+1$

$*$	0	1
0	0	1
1	1	0

- (4) Il n'y a essentiellement qu'une façon de construire un groupe à trois éléments : $G = \{e, a, b\}$ avec e comme élément neutre. La loi de groupe est donnée dans la table

$*$	e	a	b
e	e	a	b
a	a	b	e
b	b	e	a

- (5) La soustraction $x - y$ ne définit pas une loi de groupe sur $\mathbb{Z}$ (car la soustraction n'est pas associative).

Proposition 3.1.1. *L'élément neutre dans un groupe est unique. De plus chaque élément d'un groupe possède un unique inverse.*

Preuve. Supposons que e et e sont deux éléments neutres du groupe $(G, *)$. La définition de la notion d'élément neutre entraîne alors que

$$e = e * e = e.$$

Supposons maintenant que x' et x'' sont deux inverses de l'élément $x \in G$. La définition de la notion d'inverse et l'associativité de $*$ entraînent que

$$x' = e * x' = (x'' * x) * x' = x'' * (x * x') = x'' * e = x''.$$

□

Proposition 3.1.2 (Règles de simplification). *Soit $(G, *)$ un groupe, alors pour tous $x, y, u, v \in G$, on a :*

$$(i) \quad u * x = u * y \Rightarrow x = y;$$

$$(ii) \quad x * v = y * v \Rightarrow x = y.$$

Nous laissons la preuve en exercice.

3.1.1 Sous-groupes et homomorphismes

Définition. Un sous-ensemble H d'un groupe $(G, *)$ est un *sous-groupe* si c'est lui-même un groupe pour la loi $*$.

Lemme 3.1.3. *Soit H un sous-groupe de $(G, *)$, alors l'élément neutre de G est aussi l'élément neutre de H et l'inverse dans H d'un élément $x \in H$ coïncide avec l'inverse de x dans G .*

Preuve. Notons e_G l'élément neutre dans le groupe G et e_H l'élément neutre dans H , nous avons alors

$$e_H * e_G = e_H = e_H * e_H,$$

donc $e_G = e_H$ par la proposition 3.1.2. Pour prouver la seconde affirmation, on note x'_H l'inverse dans H d'un élément $x \in H$ et x'_G son inverse dans G . Alors

$$x'_H * x = e_H = e_G = x'_G * x,$$

ce qui implique que $x'_H = x'_G$.

□

Proposition 3.1.4. *Soit $(G, *)$ un groupe. Un sous-ensemble $H \subset G$ est un sous-groupe si et seulement si les conditions suivantes sont satisfaites :*

$$(a) \quad H \neq \emptyset.$$

$$(b) \quad x, y \in H \Rightarrow x * y \in H.$$

(c) $x \in H \Rightarrow x' \in H$ (où on note x' l'inverse de x dans G).

Preuve. Supposons que $H \subset G$ est un sous-groupe. Alors les conditions (i) et (ii) sont vérifiées par définition de la notion de groupe. La condition (iii) découle du lemme précédent car l'inverse de x dans H coïncide avec l'inverse de x dans G .

Supposons dans l'autre sens que $H \subset G$ est un sous-ensemble vérifiant les conditions (i), (ii) et (iii). La seconde hypothèse nous dit alors que la restriction de $*$ au sous-ensemble H définit une loi de composition interne $H \times H \xrightarrow{*} H$. La troisième hypothèse nous dit que chaque élément de H admet un inverse dans H .

La première hypothèse nous dit qu'il existe au moins un élément $x \in H$ (car $H \neq \emptyset$). Avec les deux autres hypothèses on a alors $x' \in H$ et donc

$$e = x * x' \in H.$$

Ainsi l'élément neutre appartient à H qui est donc bien un sous groupe de G . □

Exemples de sous-groupes.

- (i.) $\mathbb{Z}$ et $\mathbb{Q}$ sont des sous-groupes du groupe additif $(\mathbb{R}, +)$.
- (ii.) Si m est entier, l'ensemble de ses multiples, $\mathbb{Z} \cdot m = \{0, m, \pm 2m, \pm 3m, \dots\}$ est un sous-groupe de $(\mathbb{Z}, +)$.
- (iii.) Pour tout groupe $(G, *)$ l'ensemble $\{e\} \subset G$ est un sous-groupe.
- (iv.) $\{+1, -1\}$ est un sous-groupe de $(\mathbb{R}^*, *)$.
- (v.) Si $a \in \mathbb{R}^*$, alors l'ensemble des puissances de a , $\{a^k \mid k \in \mathbb{Z}\}$ est un sous-groupe de $(\mathbb{R}^*, *)$.
- (vi.) $\mathbb{N}$ et $\{0, 1\}$ ne sont pas des sous-groupes de $(\mathbb{Z}, +)$.

Tous ces groupes sont abéliens.

Définition. Si $(H, \times)$ et $(G, *)$ sont deux groupes, alors une application $f : H \rightarrow G$ est un *homomorphisme de groupes* (parfois on dit aussi un *morphisme*) si pour tous $x, y \in H$ on a

$$f(x \times y) = f(x) * f(y).$$

Un *isomorphisme de groupes* est un homomorphisme bijectif entre deux groupes.

Exemples.

- 1.) Si H est un sous groupe de $(G, *)$, alors l'application $i : H \rightarrow G$ définie par $i(x) = x$ pour tout $x \in H$ est un homomorphisme de groupes. C'est un isomorphisme si et seulement si $H = G$.
- 2.) Si m est un nombre entier, alors l'application $f : \mathbb{Z} \rightarrow \mathbb{Z}$, définie par $f(k) = k \cdot m$ est un homomorphisme du groupe additif $(\mathbb{Z}, +)$ vers lui-même. C'est un isomorphisme si et seulement si $m = \pm 1$.
- 3.) Si a est un nombre réel non nul, alors l'application $h : \mathbb{Z} \rightarrow \mathbb{R}^*$, définie par $h(k) = a^k$ est un homomorphisme du groupe additif $(\mathbb{Z}, +)$ vers le groupe multiplicatif $(\mathbb{R}^*, \cdot)$ car

$$h(k + n) = a^{k+n} = a^k \cdot a^n = h(k) \cdot h(n)$$

pour tous $k, n \in \mathbb{Z}$.

- 4.) L'application $\varphi : \mathbb{Z} \rightarrow \{+1, -1\}$ définie par

$$\varphi(k) = \begin{cases} +1 & \text{si } k \text{ est pair,} \\ -1 & \text{si } k \text{ est impair} \end{cases}$$

est un homomorphisme du groupe additif des entiers vers le groupe multiplicatif à deux éléments.

- 5.) Le logarithme naturel $\ln : \mathbb{R}_+ \rightarrow \mathbb{R}$ est un isomorphisme du groupe $(\mathbb{R}_+, \cdot)$ vers le groupe $(\mathbb{R}, +)$ car le logarithme est une bijection de l'ensemble des nombres réels strictement positifs vers l'ensemble de tous les nombre réels et

$$\ln(x \cdot y) = \ln(x) + \ln(y)$$

pour tous $x, y \in \mathbb{R}_+$.

Proposition 3.1.5. Soit $f : H \rightarrow G$ un homomorphisme de groupes. Notons e_G , respectivement e_H , les éléments neutres de ces deux groupes. Alors

- i.) $f(e_H) = e_G$.
- ii.) $f(x^{-1}) = (f(x))^{-1}$ pour tout $x \in H$.
- iii.) $\text{Im}(f) = f(H) = \{f(x) \mid x \in H\} \subset G$ est un sous-groupe de G (appelé l'image de f).
- iv.) $\ker(f) = \{x \in H \mid f(x) = e_G\} \subset H$ est un sous-groupe de H (appelé le noyau de f)

Les deux résultats suivants seront démontrés en exercices.

Théorème 3.1.6. Un homomorphisme de groupes $f : H \rightarrow G$ est injectif si et seulement si $\ker(f) = \{e_H\}$.

Proposition 3.1.7. Si $f : H \rightarrow G$ est un isomorphisme de groupes, alors l'application inverse $f^{-1} : G \rightarrow H$ est aussi un isomorphisme. .

3.2 Les permutations et le groupe symétrique

On appelle *permutation* d'un ensemble non vide X toute bijection $f : X \rightarrow X$, on note $\text{Perm}(X)$ l'ensemble des permutations de l'ensemble X .

Proposition 3.2.1. Pour tout ensemble non vide X , $\text{Perm}(X)$ est un groupe si on le munit de la composition $\circ$. L'élément neutre de ce groupe est l'application identité $\text{id}_X : X \rightarrow X$.

Preuve. L'associativité $(f \circ g) \circ h = f \circ (g \circ h)$ a déjà été démontrée à la proposition 2.5.1 et l'identité est clairement un élément neutre. Comme tout $f \in \text{Perm}(X)$ est bijectif, il existe un inverse à gauche et un inverse à droite, il ne reste qu'à vérifier que l'inverse à gauche coïncide avec l'inverse à droite. En effet, si $f \circ g = \text{Id}_X$ et $h \circ f = \text{Id}_X$ alors

$$h = h \circ \text{Id}_X = h \circ (f \circ g) = (h \circ f) \circ g = \text{Id}_X \circ g = g.$$

□

Définition. Le groupe *symétrique* sur n symboles est le groupe des permutations de l'ensemble

$$[n] = \{1, 2, 3, \dots, n\}.$$

On le note $\mathcal{S}_n = \text{Perm}([n])$.

Proposition 3.2.2. L'ordre du groupe $\mathcal{S}_n$ (son cardinal) est $n! = 1 \cdot 2 \cdot 3 \dots (n-1) \cdot n$.

La preuve est facile, par récurrence.

Un exemple d'élément de $\mathcal{S}_4$ est la bijection σ telle que $\sigma(1) = 3$, $\sigma(2) = 4$, $\sigma(3) = 1$ et $\sigma(4) = 2$, que l'on peut noter

$$\sigma = \begin{pmatrix} 1 & 2 & 3 & 4 \\ \downarrow & \downarrow & \downarrow & \downarrow \\ 3 & 4 & 1 & 2 \end{pmatrix}$$

ou, de manière plus simple,

$$\sigma = \begin{pmatrix} 1 & 2 & 3 & 4 \\ 3 & 4 & 1 & 2 \end{pmatrix}$$

(un tel tableau s'appelle une *matrice* à deux lignes et quatre colonnes). Si on veut composer deux éléments σ et τ de $\mathcal{S}_n$, il suffit de juxtaposer les deux tableaux, par exemple

$$\text{si } \sigma = \begin{pmatrix} 1 & 2 & 3 & 4 \\ 3 & 4 & 1 & 2 \end{pmatrix} \text{ et } \tau = \begin{pmatrix} 1 & 2 & 3 & 4 \\ 2 & 1 & 3 & 4 \end{pmatrix}$$

alors

$$\sigma \circ \tau = \begin{pmatrix} 1 & 2 & 3 & 4 \\ 2 & 1 & 3 & 4 \\ 4 & 3 & 1 & 2 \end{pmatrix} \text{ c'est à dire } \sigma \circ \tau = \begin{pmatrix} 1 & 2 & 3 & 4 \\ 4 & 3 & 1 & 2 \end{pmatrix}.$$

(se souvenir qu'on obtient la permutation $\sigma \circ \tau$ en appliquant d'abord τ , puis σ .)

Voyons quelques définitions importantes :

Définitions 1. On dit que $\sigma \in \mathcal{S}_n$ est un *k-cycle* (ou une *permutation cyclique d'ordre k*) s'il existe un sous-ensemble $C = \{c_1, c_2, c_3, \dots, c_k\} \subset \{1, 2, 3, \dots, n\}$ de cardinal k tel que $\sigma(j) = j$ si $j \notin C$ et $\sigma(c_1) = c_2, \sigma(c_2) = c_3, \dots, \sigma(c_{k-1}) = c_k$ et $\sigma(c_k) = c_1$.

On peut noter un tel *k-cycle* par $\sigma : c_1 \mapsto c_2 \mapsto c_3 \mapsto \dots \mapsto c_k \mapsto c_1$, ou, plus simplement

$$\sigma = (c_1 c_2 c_3 \dots c_k).$$

2. Deux cycles $\sigma = (c_1 c_2 \dots c_k)$ et $\rho = (d_1 d_2 \dots d_m)$ sont *disjoints* si

$$\{c_1, c_2, \dots, c_k\} \cap \{d_1, d_2, \dots, d_m\} = \emptyset.$$

Définition. Une *transposition* est un 2-cycle. Par exemple pour $1 \leq r < s \leq n$ la transposition $\tau = (rs)$ envoie r sur s et s sur r . Tous les autres éléments de $\{1, 2, 3, \dots, n\}$ sont invariants.

Remarque. On observe que toute transposition τ est sa propre inverse : $\tau \circ \tau = \text{id}$.

Lemme 3.2.3. Si deux cycles sont disjoints, alors ils commutent.

La preuve de ce lemme n'est pas difficile, nous la laissons en exercice.

Théorème 3.2.4. Toute permutation $\rho \in \mathcal{S}_n$ (telle que $\rho \neq \text{Id}$) peut s'écrire comme une composition de cycles deux à deux disjoints. Cette écriture est unique à l'ordre des cycles près.

Par exemple la permutation

$$\rho = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 \\ 3 & 5 & 1 & 2 & 4 \end{pmatrix}$$

peut s'écrire

$$\rho = (13)(254).$$

Preuve. Soit $\rho \in \mathcal{S}_n$. On appelle *support* de ρ l'ensemble des éléments de $\{1, 2, \dots, n\}$ qui sont déplacés par ρ .

$$\text{supp}(\rho) = \{j \mid \rho(j) \neq j\} \subset \{1, 2, \dots, n\}.$$

Notons $m = \text{Card}(\text{supp}(\rho))$, c'est le nombre d'éléments qui sont déplacés par ρ . Observons que si $m = 0$, alors $\rho = \text{id}$ et que $m = 1$ est impossible (car ρ est une bijection). De plus, si $m = 2$, alors ρ est une transposition, c'est donc un 2-cycle.

On continue la preuve par récurrence sur m . Si $m \geq 2$, il existe $c_1 \in \{1, 2, \dots, n\}$ tel que $\rho(c_1) = c_2 \neq c_1$. Définissons ensuite $c_3 = \rho(c_2), \dots, c_j = \rho(c_{j-1})$; comme le nombre de possibilités est fini, il existe un plus petit entier k tel que $\rho(c_k) = c_1$.

Remarquons que, en raison de l'injectivité de ρ , les éléments $\{c_1, \dots, c_k\} \subset \{1, \dots, n\}$ sont deux à deux distincts. On définit alors le k -cycle $\sigma = (c_1, c_2, \dots, c_k)$ et on considère la permutation $\varphi \in \mathcal{S}_n$ définie par $\varphi = \sigma^{-1} \circ \rho$. Par construction, le cardinal du support de φ est égal à $m - k < m$. Par hypothèse de récurrence, φ est ou bien l'identité ou bien un produit de cycles de supports deux à deux disjoints, c'est donc aussi le cas de $\rho = \sigma \circ \varphi$ car les supports de σ et φ sont disjoints. $\square$

Lemme 3.2.5. *Tout k -cycle est un produit de $k - 1$ transpositions.*

Preuve. On vérifie facilement que

$$(c_1 c_2 c_3 \dots c_k) = (c_1 c_k)(c_1 c_{(k-1)})(c_1 c_{(k-2)}) \dots (c_1, c_3)(c_1, c_2).$$

$\square$

Corollaire 3.2.6. *Toute permutation $\sigma \in \mathcal{S}_n$ (telle que $\sigma \neq \text{Id}$) peut s'écrire comme une composition de transpositions (en général cette écriture n'est pas unique).*

Par exemple la permutation $\rho = (1\ 3)(2\ 5\ 4)$ de l'exemple précédent peut s'écrire $\rho = (1\ 3)(2\ 4)(2\ 5)$.

Preuve. Le corollaire est une conséquence immédiate du théorème et du lemme précédents. $\square$

3.2.1 Signature d'une permutation.

Soit $n \in \mathbb{N}$ un entier ≥ 2 et $\phi : \llbracket n \rrbracket \rightarrow \llbracket n \rrbracket$ une application (non forcément bijective) de l'ensemble $\llbracket n \rrbracket = \{1, 2, \dots, n\}$ dans lui-même. On définit

$$\Omega(\phi) = \prod_{i < j} \left(\frac{\phi(j) - \phi(i)}{j - i} \right).$$

Ce produit contient $\binom{n}{2}$ termes, à savoir un terme pour chaque paire $\{i, j\} \subset \llbracket n \rrbracket$ telle que $i \neq j$.

Proposition 3.2.7. *Si ϕ est bijectif, alors $\Omega(\phi) = \pm 1$. Si ϕ n'est pas bijectif alors $\Omega(\phi) = 0$.*

Preuve. Comme $\{1, 2, \dots, n\}$ est un ensemble fini, on sait que toute application de cet ensemble dans lui-même qui est injective est en fait bijective. On a donc les implications suivantes :

$$\phi \text{ non bijective} \Rightarrow \phi \text{ non injective} \Rightarrow \exists i \neq j \text{ tels que } \phi(j) = \phi(i) \Rightarrow \Omega(\phi) = 0.$$

Supposons maintenant que ϕ est bijective et écrivons $\Omega(\phi)$ comme une fraction $\Omega(\phi) = \frac{N}{D}$, avec

$$N = \prod_{i < j} (\phi(j) - \phi(i)) \quad \text{et} \quad D = \prod_{i < j} (j - i).$$

En utilisant que ϕ est bijectif et que la multiplication est commutative, et en notant $i' = \phi(i)$ et $j' = \phi(j)$ on voit que

$$|N| = \prod_{i < j} |\phi(j) - \phi(i)| = \prod_{\phi(i) < \phi(j)} |\phi(j) - \phi(i)| = \prod_{i' < j'} |j' - i'| = D.$$

Nous avons donc $N = \pm D$ et par conséquent $\Omega(\phi) = \pm 1$. $\square$

Définitions. Une *paire d'inversion* d'une permutation $\sigma \in \mathcal{S}_n$ est un couple (i, j) tel que $1 \leq i < j \leq n$ et $\sigma(i) > \sigma(j)$. On note $\text{inv}(\sigma)$ le nombre de paires d'inversion de σ :

$$\text{inv}(\sigma) = \text{Card} \{ (i, j) \mid 1 \leq i < j \leq n \text{ et } \sigma(i) > \sigma(j) \}.$$

On dit que $\sigma \in \mathcal{S}_n$ est une permutation *paire* si $\text{inv}(\sigma)$ est paire et que c'est une permutation *impaire* sinon.

En examinant la définition de Ω , on voit que le résultat suivant est vrai :

Proposition 3.2.8. *Pour toute permutation $\sigma \in \mathcal{S}_n$ on a*

$$\Omega(\sigma) = (-1)^{\text{inv}(\sigma)} = \begin{cases} +1 & \text{si } \sigma \text{ est paire,} \\ -1 & \text{si } \sigma \text{ est impaire.} \end{cases}$$

Définition 3.2.9. Ce nombre s'appelle la *signature* de $\sigma \in \mathcal{S}_n$ et se note

$$\text{sgn}(\sigma) = (-1)^{\text{inv}(\sigma)} = \Omega(\sigma).$$

Théorème 3.2.10. *L'application*

$$\text{sgn} : \mathcal{S}_n \rightarrow \{+1, -1\}$$

est un homomorphisme du groupe des permutations sur le groupe multiplicatif à deux éléments.

Preuve. Pour deux permutations quelconques $\tau, \rho \in \mathcal{S}_n$, on a

$$\begin{aligned} \Omega(\tau \circ \rho) &= \prod_{i < j} \left(\frac{\tau(\rho(j)) - \tau(\rho(i))}{j - i} \right) \\ &= \prod_{i < j} \left[\left(\frac{\tau(\rho(j)) - \tau(\rho(i))}{\rho(j) - \rho(i)} \right) \cdot \left(\frac{\rho(j) - \rho(i)}{j - i} \right) \right] \\ &= \prod_{i < j} \left(\frac{\tau(\rho(j)) - \tau(\rho(i))}{\rho(j) - \rho(i)} \right) \cdot \prod_{i < j} \left(\frac{\rho(j) - \rho(i)}{j - i} \right) \\ &= \prod_{i' < j'} \left(\frac{\tau(j') - \tau(i')}{j' - i'} \right) \cdot \prod_{i < j} \left(\frac{\rho(j) - \rho(i)}{j - i} \right) \\ &= \Omega(\tau) \cdot \Omega(\rho). \end{aligned}$$

En passant à la notation $\text{sgn}(\sigma) = \Omega(\sigma)$ on conclut que

$$\text{sgn}(\tau\rho) = \text{sgn}(\tau) \cdot \text{sgn}(\rho).$$

□

Proposition 3.2.11. *La signature de toute transposition est égale à -1 .*

Preuve. Soit $\tau = (rs) \in \mathcal{S}_n$ la transposition qui échange r et s . On suppose $r < s$; alors la liste complète des paires d'inversion pour τ est donnée par

$$(ij) = (rs) \quad (ij) = (rk) \quad \text{et} \quad (ij) = (ks),$$

avec $r < k < s$. Il y a ainsi $m = 1 + 2(s - r - 1)$ paires d'inversion et donc

$$\text{sgn}(\tau) = (-1)^m = (-1)^{1+2(s-r-1)} = -1.$$

□

Corollaire 3.2.12. *La signature d'un cycle d'ordre k est égale à $(-1)^{k-1}$.*

Preuve. Nous avons montré au Lemme 3.2.5 qu'un k -cycle est un produit de $(k-1)$ transpositions, or nous venons d'établir que chaque transposition est de signature -1 . On conclut avec le théorème précédent. $\square$

Pour calculer la signature d'une permutation, on peut donc la décomposer en produit de cycles et multiplier les signatures de chaque cycle. Par exemple la signature de $\pi = (13)(254)$ est

$$\text{sign}(\pi) = (-1)(-1)^2 = -1.$$

Le résultat suivant est une autre conséquence immédiate du théorème 3.2.10.

Corollaire 3.2.13. *Si on décompose une permutation en produit de transpositions de deux manières différentes :*

$$\pi = \tau_1 \circ \tau_2 \circ \dots \circ \tau_k = \rho_1 \circ \rho_2 \circ \dots \circ \rho_m$$

(avec τ_i, ρ_j des transpositions), alors les nombres de transpositions ont la même parité (i.e. $k \equiv m \pmod{2}$). De plus la signature de π est égale à $(-1)^k$.

Nous terminons ce paragraphe par un résultat qui illustre la richesse des groupes symétriques :

Théorème 3.2.14 (Cayley). *Tout groupe fini est isomorphe à un sous-groupe d'un groupe symétrique $\mathcal{S}_n$.*

Preuve. Soit $G = \{g_1, g_2, \dots, g_n\}$ un groupe multiplicatif fini d'ordre n . On peut associer à tout élément $h \in G$ une application $\lambda_h : G \rightarrow G$ définie par multiplication à gauche :

$$\lambda_h : g \mapsto h \cdot g.$$

Il est facile de vérifier que $\lambda_{h^{-1}} \circ \lambda_h = \text{id}_G$. En particulier λ_h est une bijection, donc c'est une permutation de l'ensemble G . Par conséquent il existe pour tout $h \in G$ un élément bien défini $\sigma_h \in \mathcal{S}_n$ tel que pour tout $j \in \{1, 2, \dots, n\}$ on a

$$\lambda_h(g_j) = g_{\sigma_h(j)}.$$

On vérifie que l'application $\varphi : G \rightarrow \mathcal{S}_n$ définie par $\varphi(h) = \sigma_h$ est injective et que c'est un homomorphisme de groupes, par conséquent G est isomorphe au sous-groupe $\text{Im}(\varphi) \subset \mathcal{S}_n$. $\square$

3.3 Anneaux et corps

Définition. Un *anneau* est un ensemble A muni de deux lois de composition internes, appelées *addition* et *multiplication* et notées $+$ et $\cdot$.

$$A \times A \xrightarrow{+} A \quad \text{et} \quad A \times A \xrightarrow{\cdot} A,$$

et qui vérifient les six axiomes suivants pour tous $x, y, z \in A$:

- (A1) L'addition est commutative : $x + y = y + x$.
- (A2) L'addition est associative : $(x + y) + z = x + (y + z)$.
- (A3) Il existe un élément neutre pour l'addition : $\exists 0 \in A$ tel que $x + 0 = x$.
- (A4) Tout élément de l'anneau possède un élément opposé :

$$\forall x \in A, \quad \exists (-x) \in A \quad \text{tel que} \quad x + (-x) = 0.$$

- (A5) La multiplication est associative : $x \cdot (y \cdot z) = (x \cdot y) \cdot z$.
 (A6) La multiplication est distributive par rapport à l'addition :

$$x \cdot (y + z) = x \cdot y + x \cdot z \quad \text{et} \quad (x + y) \cdot z = x \cdot z + y \cdot z.$$

Observons que les axiomes (A1)–(A4) disent que $(A, +)$ est un groupe abélien.

L'anneau $(A, +, \cdot)$ est dit *unitaire* s'il existe un élément neutre pour la multiplication (différent de 0) :

$$(U7) \quad \exists 1 \in A \text{ tel que } 1 \neq 0 \text{ et } 1 \cdot x = x \cdot 1 = x \text{ pour tout } x \in A.$$

L'anneau est dit *commutatif* si la multiplication est commutative, i.e. $x \cdot y = y \cdot x$ pour tous $x, y \in A$ (rappelons que l'addition est supposée commutative dans tout anneau).

Définition. Un *corps* est un anneau $(K, +, \cdot)$ commutatif unitaire tel que tout élément non nul possède un inverse pour la multiplication. On a donc les 7 axiomes précédents et aussi

$$(K8) \quad x \cdot y = y \cdot x \text{ pour tous } x, y \in K.$$

$$(K9) \quad \forall x \in K^* = K \setminus \{0\}, \exists x^{-1} \in K \text{ tel que } x \cdot x^{-1} = 1.$$

On peut aussi définir un corps en disant que c'est un anneau unitaire K tel que $K^* = K \setminus \{0\}$ est un groupe abélien pour la multiplication.

Ces axiomes représentent des propriétés familières de l'addition et de la multiplication. Il y a naturellement d'autres propriétés simples de ces opérations, mais nous n'avons pas besoin de les supposer axiomatiquement car nous pouvons les démontrer. En particulier on a :

Proposition 3.3.1. Dans tout anneau $(A, +, \cdot)$ on a pour tous $x, y \in A$

$$(a) \quad x \cdot 0 = 0 \cdot x = 0.$$

$$(b) \quad x \cdot (-y) = (-x) \cdot y = -(x \cdot y).$$

$$(c) \quad \text{Si l'anneau est unitaire, on a } (-1) \cdot x = x \cdot (-1) = -x \text{ et } (-1) \cdot (-1) = 1.$$

Preuve. Nous devons prouver ces propriétés à partir des seuls axiomes.

a) On a $0 = 0 + 0$ puisque 0 est neutre pour l'addition. En utilisant maintenant la distributivité on peut écrire

$$x \cdot 0 = x \cdot (0 + 0) = x \cdot 0 + x \cdot 0,$$

par la règle de simplification on a donc $x \cdot 0 = 0$. La preuve que $0 \cdot x = 0$ est semblable.

b) En utilisant à nouveau la distributivité et la propriété (a), on a

$$0 = x \cdot 0 = x \cdot (y + (-y)) = x \cdot y + x \cdot (-y),$$

donc $x \cdot (-y)$ est l'élément opposé de $x \cdot y$, c'est à dire $-(x \cdot y) = x \cdot (-y)$.

On vérifie que $-(x \cdot y) = (-x) \cdot y$ de la même manière.

c) Si l'anneau est unitaire, on a

$$(-1) \cdot x + x = (-1) \cdot x + 1 \cdot x = ((-1) + 1) \cdot x = 0 \cdot x = 0,$$

donc $(-1) \cdot x = -x$. On prouve par un argument semblable que $x \cdot (-1) = -x$. Finalement

$$(-1) \cdot (-1) = -(-1) = 1.$$

□

Notations. Dans un anneau on note habituellement x^2 pour $x \cdot x$ et plus généralement x^k pour $x \cdot x \cdot x \dots x$ (k fois). On note aussi $x - y$ pour $x + (-y)$. Dans un corps K on note aussi

$$x^{-1} = \frac{1}{x} \quad \text{et} \quad x^{-1} \cdot y = \frac{y}{x}$$

pour tous $x, y \in K$ avec $x \neq 0$.

Exemples.

- 1.) Le plus petit corps est l'ensemble à deux éléments $\{0, 1\}$ muni de la multiplication habituelle et de l'addition modulo 2.
- 2.) $\mathbb{Z}$ est un anneau unitaire pour l'addition et la multiplication usuelle des nombres entiers.
- 3.) Pour tout $m \in \mathbb{N}$, l'ensemble de ses multiples $m\mathbb{Z}$ est un anneau. Cet anneau n'est pas unitaire si $m \neq 1$.
- 4.) L'ensemble des nombres rationnels $\mathbb{Q}$ est un corps. Rappelons qu'un élément de $\mathbb{Q}$ est par définition un quotient de deux nombres entiers et que l'addition et la multiplication se définissent par

$$\frac{a}{b} \cdot \frac{c}{d} := \frac{a \cdot c}{b \cdot d} \quad \text{et} \quad \frac{a}{b} + \frac{c}{d} = \frac{a \cdot d + c \cdot b}{b \cdot d}$$

(pour $a, b, c, d \in \mathbb{Z}$, $b, d \neq 0$).

- 5.) L'ensemble $\mathbb{R}$ des nombres réels est un corps. Ce qui distingue les réels des rationnels est l'axiome de complétude qui peut se formuler en disant que *toute suite croissante et bornée d'éléments de $\mathbb{R}$ possède une limite*. La description du corps des réels se fait dans le cours d'analyse.

Théorème 3.3.2. Soit $(A, +, \cdot)$ un anneau commutatif unitaire fini. Supposons que cet anneau n'a pas de diviseurs de zéro, c'est à dire $x \cdot y = 0 \Rightarrow x = 0$ ou $y = 0$. Alors A est un corps.

Preuve. Nous devons montrer que tout élément non nul $a \in A$ admet un inverse. Considérons l'application $\lambda : A \rightarrow A$ définie par $\lambda(x) = a \cdot x$. Cette application est injective car

$$\lambda(x) = \lambda(y) \Rightarrow a \cdot x = a \cdot y \Rightarrow a \cdot (y - x) = 0.$$

Or $a \neq 0$, donc $y - x = 0$ car l'anneau A n'a pas de diviseur de zéro. Puisque l'ensemble A est fini et que $\lambda : A \rightarrow A$ est injective, alors cette application est aussi surjective. Il existe donc $b \in A$ tel que $a \cdot b = 1$. L'élément $b \in A$ est donc l'inverse de $a \in A$.

□

3.3.1 L'anneau $\mathbb{Z}/m\mathbb{Z}$.

On se propose de construire une structure naturelle d'anneau sur l'ensemble $\{0, 1, 2, \dots, m-1\}$. Cette construction est basée sur la définition suivante :

Définition. Fixons un entier naturel $m \in \mathbb{N}$, on dit alors que deux entiers $x, y \in \mathbb{Z}$ sont *congrus modulo m* (ou simplement qu'ils sont *égaux modulo m*) si leur différence est un multiple de m : On note cette relation $x \equiv y \pmod{m}$ (ou simplement $x = y \pmod{m}$). Ainsi

$$x \equiv y \pmod{m} \Leftrightarrow (x - y) \in m \cdot \mathbb{Z}.$$

Les propriétés élémentaires de cette relation sont (pour tous $x, y, z \in \mathbb{Z}$) :

- i.) $x \equiv x \pmod{m}$.
- ii.) Si $x \equiv y \pmod{m}$, alors $y \equiv x \pmod{m}$.
- iii.) Si $x \equiv y \pmod{m}$ et $y \equiv z \pmod{m}$, alors $x \equiv z \pmod{m}$.
- iv.) Tout entier $x \in \mathbb{Z}$ est congru modulo m à un unique élément $r \in \{0, 1, 2, \dots, m-1\}$.

Le preuve des trois premières propriétés est immédiate, ces propriétés disent que la congruence modulo m est une relation d'équivalence sur $\mathbb{Z}$. La dernière propriété est une conséquence du lemme suivant :

Lemme 3.3.3. Pour tout entier positif m et tout entier relatif x il existe un unique entier $q \in \mathbb{Z}$ et un unique entier $0 \leq r \leq m-1$ tels que

$$x = qm + r.$$

On dit que cette relation est la *division entière*, ou *division euclidienne* de x par m . On dit aussi que q est le *quotient* et r le *reste* de cette division.

Preuve. La preuve se divise en plusieurs cas.

Cas 1. Si $0 \leq x < m$, alors $x = r$ et $q = 0$ vérifient les conditions du lemme.

Cas 2. Supposons $x \geq m$ et définissons l'ensemble $S = \{p \in \mathbb{N} \mid (x - pm) \geq 0\}$. Cet ensemble est borné car tout $p \in S$ vérifie $p \leq x + 1$. Notons q le plus grand élément de S et $r = x - qm$. Alors $r \geq 0$ par définition de l'ensemble S car $q \in S$ et $r < m$ car sinon on aurait

$$x - (q + 1)m = r - m \geq 0,$$

qui est impossible car $(q + 1) \notin S$. On a donc trouvé q, r tels que $0 \leq r < m$ et $x = qm + r$.

Cas 3. Si $x < 0$ on considère deux sous-cas. Lorsque $x = qm$ avec q un entier négatif il n'y a rien à faire (on pose $r = 0$). Si x n'est pas un multiple de m , alors on définit q', r' par $-x = q'm + r'$ avec $0 < r' < m$. On définit ensuite $r = m - r'$ et $q = -(q' + 1)$. On a alors $0 < r < m$ et

$$qm + r = -(q' + 1)m + (m - r') = -(q'm + r') = x.$$

□

Remarque. Le preuve donne l'algorithme pour trouver la division Euclidienne d'un nombre entier par un autre. Si $x, m \in \mathbb{N}$, alors q est le plus grand entier tel que $qm \leq x$ et $r = x - qm$.

Arithmétique modulaire On définit l'addition, la différence et la multiplication modulo m sur l'ensemble $\{0, 1, 2, \dots, m - 1\}$ par les règles :

somme : $x + y \bmod(m)$ est l'unique entier de $\{0, 1, 2, \dots, m - 1\}$ congru à $x + y$ modulo m .

différence : $x - y \bmod(m)$ est l'unique entier de $\{0, 1, 2, \dots, m - 1\}$ congru à $x - y \bmod m$.

produit : $x \cdot y \bmod(m)$ est l'unique entier de $\{0, 1, 2, \dots, m - 1\}$ congru à $x \cdot y \bmod m$.

Par exemple si $m = 5$, alors

$$2 + 4 = 1 \bmod(5), \quad 2 - 4 = 3 \bmod(5), \quad 2 \cdot 4 = 3 \bmod(5).$$

On peut démontrer directement à partir des définitions que la somme et le produit modulo m définissent une structure d'anneau sur l'ensemble $\{0, 1, 2, \dots, m - 1\}$, mais il est plus efficace d'introduire d'abord quelques nouveaux concepts.

Définition. La *classe modulo m* d'un entier $x \in \mathbb{Z}$ est l'ensemble de tous les entiers congrus à x modulo m . C'est donc le sous-ensemble suivant de $\mathbb{Z}$:

$$[x]_m = \{x + k \cdot m \mid k \in \mathbb{Z}\}$$

(que l'on note aussi $[x]_m = x + m \cdot \mathbb{Z}$). Remarquons que $[x]_m = [y]_m$ si et seulement si x et y sont congrus modulo m .

Le *représentant canonique* r d'une classe $[x]_m$ est l'unique élément de cette classe compris entre 0 et $m - 1$:

$$r = [x]_m \cap \{0, 1, 2, \dots, m - 1\}.$$

C'est donc le reste de la division entière (division euclidienne) de x par m . On note $\mathbb{Z}/m\mathbb{Z}$ l'ensemble des classes modulo m , cet ensemble est de cardinal m :

$$\mathbb{Z}/m\mathbb{Z} = \{[0]_m, [1]_m, \dots, [m - 1]_m\}.$$

On peut alors définir l'addition et la multiplication des classes modulo m de la façon suivante :

$$[x]_m + [y]_m = [x + y]_m \quad \text{et} \quad [x]_m \cdot [y]_m = [x \cdot y]_m.$$

Proposition 3.3.4. *L'ensemble $\mathbb{Z}/m\mathbb{Z}$ des classes modulo m est un anneau commutatif unitaire pour l'addition et la multiplication définies ci-dessus.*

Preuve. Montrons d'abord que $\mathbb{Z}/m\mathbb{Z}$ est un groupe abélien pour l'addition modulo m . L'associativité de l'addition des classes modulo m se déduit de la propriété correspondante sur les entiers :

$$\begin{aligned}([x]_m + [y]_m) + [z]_m &= [x + y]_m + [z]_m \\ &= [(x + y) + z]_m \\ &= [x + (y + z)]_m \\ &= [x]_m + [y + z]_m \\ &= [x]_m + ([y]_m + [z]_m).\end{aligned}$$

La preuve de la commutativité de l'addition suit le même type d'argument :

$$[x]_m + [y]_m = [x + y]_m = [y + x]_m = [y]_m + [x]_m.$$

La *classe nulle modulo m* est définie par $[0]_m = m \cdot \mathbb{Z}$, c'est l'ensemble des multiples de m et c'est un élément neutre pour l'addition des classes car

$$[x]_m + [0]_m = [x + 0]_m = [x]_m.$$

La *classe opposée* de $[x]_m$ est définie par $-[x]_m := [-x]_m$ et c'est l'opposé pour l'addition car

$$[x]_m + (-[x]_m) = [x]_m + [-x]_m = [x + (-x)]_m = [0]_m.$$

On a donc démontré que $(\mathbb{Z}/m\mathbb{Z}, +)$ est un groupe abélien. Pour voir que c'est un anneau commutatif, il faut vérifier l'associativité et la commutativité de la multiplication des classes, la preuve est semblable au cas de l'addition et nous la laissons au lecteur. Il faut aussi vérifier la distributivité :

$$\begin{aligned}[x]_m \cdot ([y]_m + [z]_m) &= [x]_m \cdot [y + z]_m \\ &= [x \cdot (y + z)]_m \\ &= [x \cdot y + x \cdot z]_m \\ &= [x \cdot y]_m + [x \cdot z]_m \\ &= [x]_m \cdot [y]_m + [x]_m \cdot [z]_m\end{aligned}$$

On vérifie en outre facilement que la classe $[1]_m$ est un élément neutre pour la multiplication. L'anneau $(\mathbb{Z}/m\mathbb{Z}, +, \cdot)$ est donc unitaire. □

Notons $\pi : \mathbb{Z}/m\mathbb{Z} \rightarrow \{0, 1, \dots, m-1\}$ l'application qui envoie la classe $[x]_m \in \mathbb{Z}/m\mathbb{Z}$ sur son représentant canonique dans $\{0, 1, \dots, m-1\}$. Par construction, cette application est bijective et elle définit un isomorphisme de l'anneau $(\mathbb{Z}/m\mathbb{Z}, +, \cdot)$ vers l'ensemble $\{0, 1, \dots, m-1\}$ muni des opérations de somme et produit modulo m . Cela démontre immédiatement la proposition suivante :

Proposition 3.3.5. *L'ensemble $\{0, 1, \dots, m-1\}$ muni des opérations de somme et produit modulo m est un anneau commutatif unitaire. Cet anneau est isomorphe à $\mathbb{Z}/m\mathbb{Z}$.*

Grâce à cette proposition, on peut identifier l'ensemble $\{0, 1, \dots, m-1\}$ muni des opérations de l'arithmétique modulaire et l'anneau $\mathbb{Z}/m\mathbb{Z}$.

Proposition 3.3.6. *L'anneau $(\mathbb{Z}/p\mathbb{Z}, +, \cdot)$ est un corps si et seulement si p est un nombre premier.*

Preuve D'après le théorème 3.3.2, il suffit de montrer que $\mathbb{Z}/p\mathbb{Z}$ est un corps si et seulement s'il ne contient pas de diviseurs de zero. Cela signifie que pour tous $x, y \in \mathbb{Z}$, on a

$$[x]_p[y]_p = [0]_p \Rightarrow [x]_p = [0]_p \text{ ou } [y]_p = [0]_p$$

Mais $[0]_p = p\mathbb{Z}$ est l'ensemble des multiples de p . La conditions précédente signifie que

$$x \cdot y \in p\mathbb{Z} \Rightarrow x \in p\mathbb{Z} \text{ ou } y \in p\mathbb{Z}.$$

ou encore :

$$\text{si } (p \text{ divise } x \cdot y), \text{ alors } (p \text{ divise } x) \text{ ou } (p \text{ divise } y).$$

Cette condition caractérise précisément les nombres premiers. □

3.3.2 Carrés et formule quadratique.

On dit qu'un élément w d'un corps K est un *carré dans le corps K* s'il existe $v \in K$ tel que $w = v^2 = v \cdot v$. Voyons quelques exemples :

- i) 0 et 1 sont des carrés dans tout corps (car $0 = 0^2$ et $1 = 1^2$).
- ii) On démontre dans le cours d'analyse qu'un nombre réel w est un carré dans $\mathbb{R}$ si et seulement si $w \geq 0$.
- iii) Les carrés du corps à 3 éléments $\mathbb{F}_3 = \{0, 1, 2\}$ sont 0 et 1 (preuve : il suffit de vérifier que $0^2 = 0$, $1^2 = 1$ et $2^2 = 1$).
- iv) 2 n'est pas un carré dans $\mathbb{Q}$.

Démontrons cette dernière affirmation par l'absurde. Supposons qu'il existe $v \in \mathbb{Q}$ tel que $v^2 = 2$. Comme v est un rationnel non nul, on peut écrire $v = \frac{x}{y}$ avec $x, y \in \mathbb{Z} \setminus \{0\}$. On a donc $x^2 = 2y^2$. Si on réduit cette équation modulo 3 on obtient deux éléments (toujours notés x et y) dans $\mathbb{F}_3$ tels que $x^2 = 2y^2$. Or dans le corps $\mathbb{F}_3$ le seul carré non nul est 1, donc $x^2 = y^2 = 1$ et l'équation précédente dit que $1 = 2$, ce qui est faux dans $\mathbb{F}_3$. □

Formule quadratique. On a le résultat suivant : Supposons que K est un corps et $a, b, c \in K$ avec $2a \neq 0$. Alors l'équation quadratique

$$ax^2 + bx + c = 0$$

admet une solution dans K si et seulement si $\Delta = b^2 - 4ac$ est un carré dans le corps K .

Si c'est le cas, il existe une unique solution lorsque $\Delta = 0$, qui est donnée par $x = -\frac{b}{2a}$ et lorsque $\Delta \neq 0$ il y a deux solutions qui sont données par

$$x_1 = \frac{-b + v}{2a} \quad \text{et} \quad x_2 = \frac{-b - v}{2a},$$

avec $v \in K$ tel que $v^2 = \Delta$.

On dit que Δ est le *discriminant* du polynôme $ax^2 + bx + c$ et on écrit les formules précédentes sous la forme

$$x_{1,2} = \frac{-b \pm \sqrt{\Delta}}{2a}. \tag{3.1}$$

Preuve. Supposons que x est solution de notre équation, i.e. $ax^2 + bx + c = 0$. En multipliant cette équation par $4a$ et en utilisant les règles de calcul dans un corps, on obtient :

$$\begin{aligned} 0 &= 4a(ax^2 + bx + c) = 4a^2x^2 + 4abx + 4ac \\ &= ((2ax)^2 + 4abx + b^2) - \Delta \\ &= (2ax + b)^2 - v^2 \end{aligned}$$

car $v^2 = \Delta$. On a donc $(2ax + b)^2 = v^2$ qui est équivalent à $2ax + b = \pm v$. Inversément, supposons que x est donné par (3.1), alors en remontant les calculs précédents, on vérifie que $ax^2 + bx + c = 0$. □

Remarque. Dans le cas du corps $\mathbb{R}$ des réels, la condition “ Δ est un carré” est équivalente à $\Delta \geq 0$. Dans le cas du corps $\mathbb{C}$ des nombres complexes, cette condition est toujours satisfaite ; nous donnons la preuve dans le prochain paragraphe.

3.4 Le corps des nombres complexes

3.4.1 Définition et propriétés de bases

Sur l'ensemble $\mathbb{R}^2 = \mathbb{R} \times \mathbb{R}$ nous définissons deux lois de compositions internes.

(i) La *somme* (ou l'*addition*) de (x, y) et (u, v) est définie par

$$(x, y) + (u, v) = (x + u, y + v).$$

(ii) Le *produit* (ou la *multiplication*) de (x, y) et (u, v) est défini par

$$(x, y) \cdot (u, v) = (xu - yv, xv + yu).$$

Nous allons démontrer le résultat important suivant :

Théorème 3.4.1. *L'ensemble $\mathbb{R}^2$, lorsqu'il est muni de ces deux opérations, est un corps.*

On appelle ce corps le *corps des nombres complexes* et on le note $\mathbb{C}$.

Avant de donner la preuve du théorème, il est bon d'explorer un peu l'univers des nombres complexes. Observons d'abord que ces deux lois de compositions sont commutatives :

$$(x, y) + (u, v) = (u, v) + (x, y) \quad \text{et} \quad (x, y) \cdot (u, v) = (u, v) \cdot (x, y).$$

pour tous $(x, y), (u, v) \in \mathbb{C}$. L'élément neutre pour l'addition est $(0, 0)$ et l'élément neutre pour la multiplication est $(1, 0)$ car

$$(1, 0) \cdot (u, v) = (1 \cdot u - 0 \cdot v, 1 \cdot v + 0 \cdot u) = (u, v).$$

Il est commode de noter simplement 0 pour $(0, 0)$ et 1 pour $(1, 0)$. Nous notons également i pour l'élément $(0, 1)$. Un nombre complexe $z = (x, y) \in \mathbb{C}$ se note alors $z = x + iy$; ce qu'on justifie ainsi :

$$z = (x, y) = x \cdot (1, 0) + y \cdot (0, 1) = x \cdot 1 + y \cdot i = x + iy.$$

Nous pouvons donc redéfinir

$$\mathbb{C} = \{z = x + iy \mid x, y \in \mathbb{R}\}.$$

On dit alors que $x \in \mathbb{R}$ est la *partie réelle* du nombre complexe z et y est sa *partie imaginaire*, et on note :

$$x = \text{Ré}(z) \quad \text{et} \quad y = \text{Im}(z).$$

Les opérations de somme et produits de nombres complexes s'écrivent avec cette notation :

$$\begin{cases} (x + iy) + (u + iv) = (x + y) + i(u + v) \\ (x + iy) \cdot (u + iv) = (xu - yv) + i(xv + yu). \end{cases}$$

Notons en particulier que $i^2 = -1$, car

$$i^2 = i \cdot i = -1 \cdot 1 + i \cdot 0 = -1,$$

et on dit que i est *l'unité imaginaire* de $\mathbb{C}$. Cela nous permet de multiplier les nombres complexes comme des expressions algébriques usuelles, en remplaçant chaque occurrence de i^2 par -1 . Par exemple

$$\begin{aligned}(3 + 2i) \cdot (4 - 5i) &= 3 \cdot 4 - 3 \cdot 5i + 2i \cdot 4 - 2 \cdot 5i^2 \\ &= 12 + (8 - 15)i - 10i^2 \\ &= 22 - 7i.\end{aligned}$$

L'opération qui consiste à changer le signe de la partie imaginaire d'un nombre complexe s'appelle la *conjugaison complexe*. Si $z = x + iy$, on note $\bar{z} = x - iy$ son conjugué. La conjugaison complexe est compatible avec les opérations de somme et de produit :

$$\overline{z + w} = \bar{z} + \bar{w}, \quad \overline{z \cdot w} = \bar{z} \cdot \bar{w}$$

Vérifions la deuxième relation : si $z = x + iy$ et $w = u + iv$, alors

$$\begin{aligned}\overline{z \cdot w} &= \overline{(x + iy) \cdot (u + iv)} = \overline{(xu - yv) + i(xv + yu)} = (xu - yv) - i(xv + yu) \\ &= (x - iy) \cdot (u - iv) = \overline{(x + iy)} \cdot \overline{(u + iv)} \\ &= \bar{z} \cdot \bar{w}.\end{aligned}$$

Observons que

$$\operatorname{Ré}(z) = \frac{1}{2}(z + \bar{z}) \quad \text{et} \quad \operatorname{Im}(z) = \frac{1}{2i}(z - \bar{z}).$$

On dit que le nombre complexe z est *réel* si sa partie imaginaire est nulle et qu'il est *imaginaire* (on dit parfois *purement imaginaire*) si sa partie réelle est nulle. Par conséquent le nombre complexe z est réel si et seulement si $\bar{z} = z$ et il est imaginaire si et seulement si $\bar{z} = -z$. Il est également utile de remarquer que si $z = x + iy$, alors

$$z \cdot \bar{z} = (x + iy) \cdot (x - iy) = x^2 + y^2.$$

On appelle *module* du nombre complexe $z = x + iy$ le nombre réel $|z|$ défini par

$$|z| = \sqrt{z \cdot \bar{z}} = \sqrt{x^2 + y^2}.$$

Notons que $|z| = 0$ si et seulement si $z = 0$. Observons aussi que $|z \cdot w| = |z| \cdot |w|$, car on a

$$|z \cdot w|^2 = (z \cdot w) \cdot (\overline{z \cdot w}) = (z \cdot w) \cdot (\bar{z} \cdot \bar{w}) = (z \cdot \bar{z}) \cdot (w \cdot \bar{w}) = |z|^2 \cdot |w|^2.$$

Démontrons maintenant le théorème précédent :

Preuve du théorème. Pour prouver que $\mathbb{C}$ est un corps, nous devons vérifier 9 propriétés de l'addition et de la multiplication. Les quatre premières propriétés disent que $(\mathbb{C}, +)$ est un groupe abélien. La vérification est facile et nous la laissons au lecteur (il s'agit simplement du groupe $(\mathbb{R}^2, +)$, produit cartésien de $(\mathbb{R}, +)$ avec lui-même).

La cinquième condition dit que la multiplication est associative. En effet un calcul (un peu long) nous confirme que d'une part

$$\begin{aligned}[(x_1 + iy_1)(x_2 + iy_2)](x_3 + iy_3) &= \\ (x_1 x_2 x_3 - x_1 y_2 y_3 - y_1 x_2 y_3 - y_1 y_2 x_3) &+ i(x_1 x_2 y_3 + x_1 y_2 x_3 + y_1 x_2 x_3 - y_1 y_2 y_3),\end{aligned}$$

et d'autre part que $(x_1 + iy_1)[(x_2 + iy_2)(x_3 + iy_3)]$ est égale à la même expression. Nous encourageons le lecteur à faire ce calcul en détail.

La sixième propriété est la distributivité de la multiplication sur l'addition :

$$z_1(z_2 + z_3) = z_1 z_2 + z_1 z_3.$$

On a en effet

$$\begin{aligned}
(x_1 + iy_1)[(x_2 + iy_2) + (x_3 + iy_3)] &= (x_1 + iy_1)[(x_2 + x_3) + i(y_2 + y_3)] \\
&= (x_1 x_2 + x_1 x_3 - y_1 y_2 - y_1 y_3) + i(x_1 y_2 + x_1 y_3 + y_1 x_2 + y_1 x_3) \\
&= (x_1 x_2 - y_1 y_2) + i(x_1 y_2 + y_1 x_2) + (x_1 x_3 - y_1 y_3) + i(x_1 y_3 + y_1 x_3) \\
&= (x_1 + iy_1)(x_2 + iy_2) + (x_1 + iy_1)(x_3 + iy_3).
\end{aligned}$$

On vérifie de la même façon que $(z_1 + z_2)z_3 = z_1 z_3 + z_2 z_3$. La septième règle est l'existence d'un élément neutre pour la multiplication. Nous savons déjà que $1 = 1 + i \cdot 0$ est cet élément. La huitième règle est la commutativité de la multiplication que nous avons déjà observée.

La dernière condition dit que tout nombre complexe non nul z doit avoir un inverse pour la multiplication. Soit donc $z \in \mathbb{C} \setminus \{0\}$. On veut prouver qu'il existe $w \in \mathbb{C}$ tel que $z \cdot w = 1$; remarquons pour cela que si un tel élément existe, alors on a

$$\bar{z} = \bar{z} \cdot 1 = \bar{z} \cdot z \cdot w = |z|^2 \cdot w.$$

Par conséquent $w = \bar{z}/|z|^2$ vérifie $z \cdot w = 1$ et on note ce nombre complexe

$$z^{-1} = \frac{1}{z} = \frac{\bar{z}}{|z|^2}. \quad (3.2)$$

Nous avons ainsi vérifié que $\mathbb{C}$ vérifie toutes les conditions pour être un corps. □

Remarques 1. En particulier l'ensemble $\mathbb{C}^* = \mathbb{C} \setminus \{0\}$ est un groupe abélien pour la multiplication complexe.

2. Le quotient du nombre complexe z par $w \in \mathbb{C}^*$ est le nombre complexe noté $\frac{z}{w}$ et est défini par

$$\frac{z}{w} = w^{-1} \cdot z = z \cdot w^{-1}.$$

Remarquons que

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

Exemple.

$$\frac{3 + 2i}{4 - 5i} = \frac{(3 + 2i)(4 + 5i)}{4^2 + 5^2} = \frac{2 + 23i}{41}$$

3.4.2 Racine carrée d'un nombre complexe

On prouve dans le cours d'analyse que la fonction $u \mapsto u^2$ est une bijection de $\mathbb{R}_+$ dans lui-même. Ainsi pour tout nombre réel positif x il existe un et un seul réel positif u tel que $u^2 = x$, on le note $u = \sqrt{x}$. On sait par ailleurs que pour $x > 0$, l'équation $u^2 = x$ admet exactement deux solutions qui sont $\sqrt{x}$ et $-\sqrt{x}$.

Pour décrire la racine carrée d'un nombre complexe, on introduit les ensembles

$$\mathbb{H}_+ = \{z \in \mathbb{C} \mid \text{Im}(z) > 0\} \cup \mathbb{R}_+ \quad \text{et} \quad \mathbb{H}_- = \{z \in \mathbb{C} \mid \text{Im}(z) < 0\} \cup \mathbb{R}_-.$$

Notons que $\mathbb{H}_+$ et $\mathbb{H}_-$ représentent des demi-plans ouverts auxquels on a ajouté une demi-droite. Ces deux ensembles sont symétriques, i.e. $z \in \mathbb{H}_+$ si et seulement si $-z \in \mathbb{H}_-$. De plus $\mathbb{H}_+ \cup \mathbb{H}_- = \mathbb{C}$ et $\mathbb{H}_+ \cap \mathbb{H}_- = \{0\}$.

Proposition 3.4.2. La fonction $w \mapsto w^2$ définit une bijection de $\mathbb{H}_+$ vers $\mathbb{C}$.

Preuve. Nous devons prouver que pour tout nombre complexe $z = x + iy \in \mathbb{C}$, il existe un unique $w = u + iv \in \mathbb{H}_+$ tel que $w^2 = z$. Rappelons que $w^2 = (u + iv)^2 = (u^2 - v^2) + 2iuv$, nous avons donc les relations

$$\operatorname{Ré}(z) = (u^2 - v^2), \quad \operatorname{Im}(z) = 2uv, \quad |z| = |w|^2 = u^2 + v^2, \quad \text{et } v = \operatorname{Im}(w) \geq 0.$$

Par conséquent

$$|z| + \operatorname{Ré}(z) = 2u^2 \quad \text{et} \quad |z| - \operatorname{Ré}(z) = 2v^2,$$

et donc

$$u = \pm \sqrt{\frac{|z| + \operatorname{Ré}(z)}{2}} \quad \text{et} \quad v = \pm \sqrt{\frac{|z| - \operatorname{Ré}(z)}{2}}.$$

Par hypothèse nous avons $v = \operatorname{Im}(w) \geq 0$ car $w \in \mathbb{H}_+$. Pour déterminer le signe de u on distingue trois cas :

cas 1 Si $\operatorname{Im}(z) \neq 0$, alors u et $\operatorname{Im}(z)$ ont le même signe car $\operatorname{Im}(z) = 2uv$ et $v \geq 0$.

cas 2 Si $\operatorname{Im}(z) = 0$ et $x = \operatorname{Ré}(z) \geq 0$, alors $|z| = \operatorname{Ré}(z)$ par conséquent $v = 0$ et $u = \sqrt{x} \geq 0$.

cas 3 Si $\operatorname{Im}(z) = 0$ et $x = \operatorname{Ré}(z) < 0$, alors $|z| = -\operatorname{Ré}(z)$ par conséquent $u = 0$ et $v = \sqrt{-x} > 0$.

Noter que dans le cas 2 $z = x$ est un réel positif et $w = \sqrt{x}$ et dans le cas 3 $z = x$ est un réel négatif et $w = i\sqrt{-x}$. Dans tous les cas il existe un unique élément $w = u + iv \in \mathbb{H}_+$ tel que $w^2 = z$, et ce nombre complexe est donné par

$$w = \varepsilon \cdot \sqrt{\frac{|z| + \operatorname{Ré}(z)}{2}} + i \cdot \sqrt{\frac{|z| - \operatorname{Ré}(z)}{2}}, \quad (3.3)$$

où $\varepsilon = +1$ si $\operatorname{Im}(z) \geq 0$ et $\varepsilon = -1$ si $\operatorname{Im}(z) < 0$.

L'argument montre que s'il existe une solution de $w^2 = z$ dont la partie imaginaire est positive ou nulle, alors w est donné par (3.3). Pour conclure la démonstration, nous devons encore prouver que si w est défini par (3.3), alors $w^2 = z$. Ceci se fait par un calcul direct que nous laissons au lecteur.

□.

Nous avons immédiatement le corollaire suivant.

Corollaire 3.4.3. (a) Pour tout nombre complexe non nul z il existe exactement deux solutions de l'équation $w^2 = z$,

(b) La fonction $w \mapsto w^2$ définit aussi une bijection de $\mathbb{H}_+$ vers $\mathbb{C}$.

Définition. Pour tout nombre complexe z , l'unique élément $w \in \mathbb{H}_+$ tel que $w^2 = z$ s'appelle la *racine carrée principale* de z . On note $w = \sqrt{z}$.

Exemples 1. Si $z = i$ on a $|z| = 1$ et $\operatorname{Ré}(z) = 0$. On a aussi $\operatorname{Im}(z) = 1 > 0$, donc

$$\sqrt{i} = \sqrt{\frac{1}{2}} + i \cdot \sqrt{\frac{1}{2}} = \frac{\sqrt{2} + i\sqrt{2}}{2}.$$

2. Pour trouver $\sqrt{z}$ avec $z = 5 - 12i$, on calcule d'abord $|z| = \sqrt{5^2 + 12^2} = \sqrt{169} = 13$. On a alors

$$\sqrt{z} = \pm \sqrt{\frac{13 + 5}{2}} + i \sqrt{\frac{13 - 5}{2}} = \pm 3 + 2i.$$

Le signe dans le membre de gauche doit être négatif car $\operatorname{Im}(z) = -12 < 0$. Nous avons donc finalement $\sqrt{z} = (-3 + 2i)$.

Corollaire 3.4.4. Pour tous $a, b, c \in \mathbb{C}$ l'équation quadratique $az^2 + bz + c = 0$ possède une ou deux solutions.

Preuve. C'est une application directe de la formule quadratique. Les solutions sont données par

$$z = \frac{-b \pm \sqrt{b^2 - 4ac}}{2a}.$$

Il existe une ou deux solutions selon que le discriminant $b^2 - 4ac$ est nul ou non. □

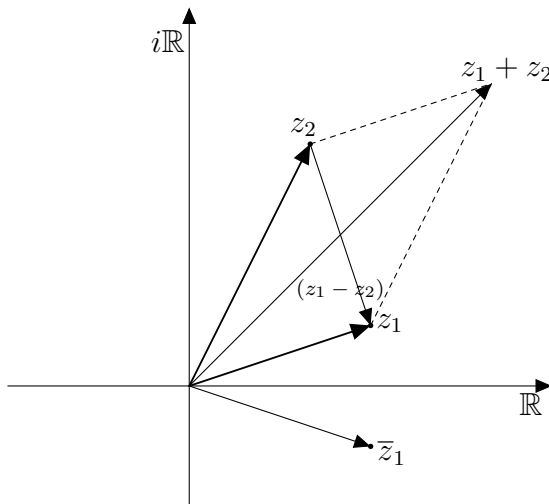
3.4.3 Interprétation géométrique des nombres complexes

Dans le plan euclidien on représente géométriquement un nombre complexe $z = x + iy \in \mathbb{C}$ par le vecteur correspondant $(x, y) \in \mathbb{R}^2$. L'axe Ox s'appelle alors l'*axe réel* et l'axe Oy est l'*axe imaginaire*. On note ces axes $\mathbb{R}$ et $i\mathbb{R}$ (vus come des droites dans le plan complexe $\mathbb{C}$).

La somme et la différence de deux nombres complexes z_1, z_2 correspondent alors aux opérations de sommes et de différence vectorielles et se visualisent sur le parallélogramme de sommets $0, z_1, z_2$ et $(z_1 + z_2)$. Le conjugué d'un nombre complexe est le symétrique de ce nombre par rapport à l'axe réel.

Le module du nombre complexe z représente la distance euclidienne entre z et 0 (par le théorème de Pythagore). Plus généralement la distance entre $z_1 = x_1 + iy_1$ et $z_2 = x_2 + iy_2$ est

$$\text{dist}(z_1, z_2) = \sqrt{(x_2 - x_1)^2 + (y_2 - y_1)^2} = |z_2 - z_1|.$$



Pour interpréter géométriquement la multiplication complexe, il est commode d'utiliser les coordonnées polaires. Rappelons que tout élément $(x, y) \in \mathbb{R}^2$ peut s'écrire sous la forme

$$(x, y) = (r \cos(\theta), r \sin(\theta)),$$

où $r = \sqrt{x^2 + y^2}$ est la distance à l'origine et $0 \leq \theta < 2\pi$ représente l'angle entre l'axe Ox et le vecteur (x, y) mesuré dans le sens trigonométrique (l'angle θ n'est bien défini que lorsque $(x, y) \neq (0, 0)$).

Appliqué aux nombres complexes, cela signifie que tout nombre complexe s'écrit

$$z = r \cdot (\cos(\theta) + i \sin(\theta)), \quad \text{avec } r = |z|.$$

On dit alors que θ est l'*argument* du nombre complexe z et on note $\theta = \arg(z)$.

Théorème 3.4.5. *Le produit des nombres complexes $z_1 = r_1 \cdot (\cos(\theta_1) + i \sin(\theta_1))$ et $z_2 = r_2 \cdot (\cos(\theta_2) + i \sin(\theta_2))$ est égal à*

$$z_1 \cdot z_2 = r_1 r_2 \cdot (\cos(\theta_1 + \theta_2) + i \sin(\theta_1 + \theta_2)). \quad (3.4)$$

Si $z_2 \neq 0$, alors le quotient de ces deux nombres complexes est

$$\frac{z_1}{z_2} = \frac{r_1}{r_2} \cdot (\cos(\theta_1 - \theta_2) + i \sin(\theta_1 - \theta_2)). \quad (3.5)$$

Nous pouvons aussi écrire ces résultats sous la forme

$$|z_1 \cdot z_2| = |z_1| \cdot |z_2|, \quad \left| \frac{z_1}{z_2} \right| = \frac{|z_1|}{|z_2|}.$$

et

$$\arg(z_1 \cdot z_2) = \arg(z_1) + \arg(z_2), \quad \arg\left(\frac{z_1}{z_2}\right) = \arg(z_1) - \arg(z_2).$$

Preuve. En appliquant la définition de la multiplication complexe nous avons :

$$\begin{aligned} z_1 \cdot z_2 &= (r_1 \cos(\theta_1) + i r_1 \sin(\theta_1)) \cdot (r_2 \cos(\theta_2) + i r_2 \sin(\theta_2)) \\ &= r_1 r_2 (\cos(\theta_1) \cos(\theta_2) - \sin(\theta_1) \sin(\theta_2)) + i r_1 r_2 (\cos(\theta_1) \sin(\theta_2) + \sin(\theta_1) \cos(\theta_2)). \end{aligned}$$

La formule (3.4) est donc une conséquence des formules d'addition d'angles pour les fonctions trigonométriques. Rappelons que ces formules s'écrivent :

$$\begin{cases} \cos(\theta_1 + \theta_2) = \cos(\theta_1) \cos(\theta_2) - \sin(\theta_1) \sin(\theta_2). \\ \sin(\theta_1 + \theta_2) = \cos(\theta_1) \sin(\theta_2) + \sin(\theta_1) \cos(\theta_2). \end{cases}$$

Pour prouver la formule du quotient (3.5), on note

$$w = \frac{z_1}{z_2} = s \cdot (\cos(\varphi) + i \sin(\varphi)).$$

On doit déterminer s et φ . En appliquant (3.4) à la relation $w \cdot z_2 = z_1$ on obtient

$$r_2 = s \cdot r_1 \quad \text{et} \quad \theta_2 = \varphi + \theta_1,$$

c'est-à-dire :

$$\left| \frac{z_1}{z_2} \right| = |w| = s = \frac{r_1}{r_2} \quad \text{et} \quad \varphi = \theta_2 - \theta_1.$$

□

Remarque. La somme de deux angles dans l'intervalle $[0, 2\pi)$ n'est pas toujours contenue dans cet intervalle. La façon la plus élégante de résoudre cette difficulté est de considérer que l'argument d'un nombre complexe est bien défini modulo 2π .

Ce théorème conduit à l'interprétation géométrique suivante : *l'effet de la multiplication du nombre complexe w par z est de multiplier sa longueur par un facteur $|z|$ et de la faire tourner dans le sens positif d'un angle égal à l'argument de z .*

En particulier *la multiplication par i correspond à une rotation d'angle $\frac{\pi}{2}$ dans le sens positif.*

Corollaire 3.4.6 (Formule de Moivre). *Pour tout nombre complexe $z = r \cdot (\cos(\theta) + i \sin(\theta))$ non nul et tout entier $n \in \mathbb{Z}$ on a*

$$z^n = r^n \cdot (\cos(n\theta) + i \sin(n\theta)). \quad (3.6)$$

De façon équivalente, on peut écrire

$$|z^n| = |z|^n, \quad \text{et} \quad \arg(z^n) = n \cdot \arg(z) \quad (\text{modulo } 2\pi).$$

Preuve. Considérons d'abord le cas $n \geq 0$. Pour $n = 0$ la formule dit que

$$1 = z^0 = r^0 \cdot (\cos(0) + i \sin(0)),$$

qui est évidemment vrai. Supposons la formule démontrée pour $n \geq 0$ et appliquons l'hypothèse de récurrence et le théorème précédent au calcul suivant :

$$\begin{aligned} z^{n+1} &= z \cdot z^n = r \cdot (\cos(\theta) + i \sin(\theta)) \cdot r^n \cdot (\cos(n\theta) + i \sin(n\theta)) \\ &= r \cdot r^n (\cos(\theta + n\theta) + i \sin(\theta + n\theta)) \\ &= r^{n+1} \cdot (\cos((n+1)\theta) + i \sin((n+1)\theta)). \end{aligned}$$

On a donc démontré (3.6) pour tout entier positif ou nul.

Si $n < 0$, alors $z^n = \frac{1}{z^{-n}}$ et la formule (3.6) se déduit du cas $n > 0$ et de (3.5). □

Une application de la formule de Moivre.

On obtient $\cos(n\theta)$ en prenant la partie réelle et $\sin(n\theta)$ en prenant la partie imaginaire de l'égalité

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

Notons pour le calcul que $i^k = \pm 1$ si k est pair et $i^k = \pm i$ si k est impair.

Exemple. La partie réelle de la formule nous donne pour $n = 4$:

$$\cos(4\theta) = \cos^4(\theta) - 6 \sin^2(\theta) \cos^2(\theta) + \sin^4(\theta),$$

et la partie imaginaire est

$$\sin(4\theta) = 4 \sin(\theta) \cos^3(\theta) - 4 \sin^3(\theta) \cos(\theta).$$

3.4.4 L'exponentielle d'un nombre complexe.

Rappelons que l'exponentielle d'un nombre réel x est définie par la série entière :

$$\exp(x) = \sum_{k=0}^{\infty} \frac{x^k}{k!} = 1 + x + \frac{x^2}{2} + \frac{x^3}{6} + \frac{x^4}{24} + \dots$$

La convergence de cette série est démontrée dans le cours d'analyse, de même que les propriétés fondamentales de cette fonction. La propriété la plus importante est que l'exponentielle est un homomorphisme du groupe additif $(\mathbb{R}, +)$ vers le groupe multiplicatif $(\mathbb{R}, \cdot)$, c'est à dire qu'on a

$$\exp(x_1 + x_2) = \exp(x_1) \exp(x_2)$$

pour tous $x_1, x_2 \in \mathbb{R}$. L'exponentielle de 1 s'appelle la *nombre d'Euler* et se note

$$e = \exp(1) = \sum_{k=0}^{\infty} \frac{1}{k!} = 1 + 1 + \frac{1}{2} + \frac{1}{6} + \frac{1}{24} + \dots = 2.718281828459 \dots$$

L'exponentielle de $x \in \mathbb{R}$ se note également $\exp(x) = e^x$.

On peut aussi démontrer que la série exponentielle converge pour tous nombre complexe z et définit un élément de $\mathbb{C}$:

$$e^z = \exp(z) = \sum_{k=0}^{\infty} \frac{z^k}{k!} \in \mathbb{C}.$$

L'application $z \mapsto e^z$ est un homomorphisme de $(\mathbb{C}, +)$ vers $(\mathbb{C}, \cdot)$:

$$e^{z_1+z_2} = e^{z_1} \cdot e^{z_2}$$

pour tous $z_1, z_2 \in \mathbb{C}$. Nous énonçons le théorème suivant sans démonstration :

Théorème 3.4.7. *L'exponentielle complexe est l'unique homomorphisme de $(\mathbb{C}, +)$ vers $(\mathbb{C}^*, \cdot)$ qui est différentiable et qui envoie 1 sur e .*

Ce théorème nous permet de démontrer le résultat fondamental suivant :

Corollaire 3.4.8. *Pour tout $z = x + iy \in \mathbb{C}$ on a*

$$e^z = e^x(\cos(y) + i \sin(y)). \quad (3.7)$$

Preuve. Notons $f(z)$ la fonction définie par $f(x + iy) = e^x(\cos(y) + i \sin(y))$. Cette fonction est différentiable car les fonctions exponentielle (réelle), sinus et cosinus sont différentiables. En utilisant (3.4), nous avons

$$\begin{aligned} f(z_1) \cdot f(z_2) &= e^{x_1}(\cos(y_1) + i \sin(y_1)) \cdot e^{x_2}(\cos(y_2) + i \sin(y_2)) \\ &= e^{x_1} e^{x_2} \cdot (\cos(y_1 + y_2) + i \sin(y_1 + y_2)) \\ &= e^{x_1+x_2} \cdot (\cos(y_1 + y_2) + i \sin(y_1 + y_2)) \\ &= f(z_1 + z_2). \end{aligned}$$

L'application f est donc un homomorphisme de $(\mathbb{C}, +)$ vers $(\mathbb{C}, \cdot)$ et comme $f(1) = f(1 + i \cdot 0) = e$, nous concluons par le théorème précédent que

$$e^z = f(z) = e^x(\cos(y) + i \sin(y)).$$

□

La formule (3.7) est due à Euler, et c'est l'une des formules fondamentales en mathématiques. Notons en particulier que pour tout $t \in \mathbb{R}$ elle dit que

$$e^{it} = \cos(t) + i \sin(t) \quad \text{et} \quad e^{-it} = \cos(t) - i \sin(t) = \overline{e^{it}}.$$

De façon équivalente

$$\cos(t) = \operatorname{Re}(e^{it}) = \frac{e^{it} + e^{-it}}{2} \quad \text{et} \quad \sin(t) = \operatorname{Im}(e^{it}) = \frac{e^{it} - e^{-it}}{2i}.$$

Soulignons aussi la magnifique identité d'Euler

$$e^{i\pi} + 1 = 0,$$

qui relie les nombres 0, 1, i , π et e .

Notons que la formule (3.7) nous permet d'écrire la forme polaire d'un nombre complexe sous la forme

$$z = re^{i\theta}, \quad (3.8)$$

où $r = |z|$ et $\theta = \arg(z)$. Nous concluons ce paragraphe par le résultat suivant qui généralise la proposition 3.4.2.

Proposition 3.4.9. *Pour tout nombre complexe non nul z et tout entier $n \geq 1$ il existe exactement n solutions de l'équation $w^n = z$.*

La preuve est une application très simple de la formule (3.8). Notons $z = re^{i\theta}$, alors les nombres complexes

$$w_k = \sqrt[n]{r} \cdot e^{i(\theta/n + 2\pi k/n)} = \sqrt[n]{r} \cdot \left[\cos\left(\frac{\theta + 2k\pi}{n}\right) + i \sin\left(\frac{\theta + 2k\pi}{n}\right) \right] \quad (3.9)$$

sont deux-à-deux distincts pour $k = 0, 1, \dots, n-1$ et vérifient tous $(w_k)^n = z$. $\square$

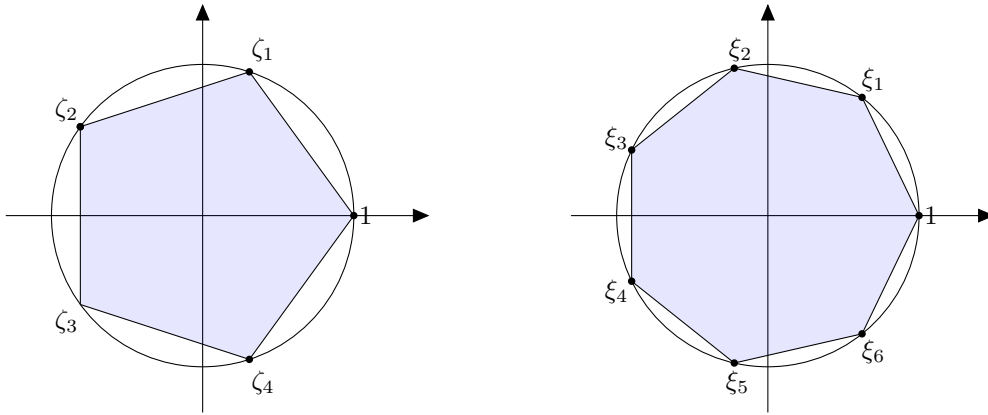
Remarque. On vérifie facilement que lorsque $n = 2$ et $k = 0$, la formule (3.9) est équivalente à (3.3). À savoir elle définit l'unique racine carrée de z telle que $\text{Im}(w) \geq 0$ (et $\text{Ré}(w) \geq 0$ si $z \in \mathbb{R}_+$).

La formule la formule (3.9) nous dit aussi que les nombres complexes

$$\zeta_k = e^{i(2\pi k/n)} = \cos\left(\frac{2k\pi}{n}\right) + i \sin\left(\frac{2k\pi}{n}\right)$$

vérifient $\zeta_k^n = 1$; on les appelle les *racines $n^{\text{ième}}$ de l'unité*. Géométriquement ils forment les sommets d'un polygone régulier sur le cercle unité.

Le pentagone ci-dessous représente les racines $5^{\text{ième}}$ de l'unité et l'heptagone indique les racines $7^{\text{ième}}$ de l'unité.



Finalement nous énonçons le théorème fondamental de l'algèbre, qui sera démontré dans le cours d'analyse complexe.

Théorème 3.4.10 (Théorème fondamental de l'algèbre). *Soit $f(z) = a_0 + a_1z + a_2z^2 + \dots + a_nz^n$ un polynôme quelconque à coefficients complexes. Si $n \geq 1$, alors ce polynôme admet (au moins une) racine réelle : il existe $z_0 \in \mathbb{C}$ tel que $f(z_0) = 0$.*

3.4.5 Logarithme naturel d'un nombre complexe.

Nous voudrions définir le logarithme naturel d'un nombre complexe comme l'inverse de l'exponentielle, toutefois un problème se pose du fait que l'exponentielle complexe n'est pas une fonction injective. En effet, pour tout $z \in \mathbb{C}$, on a par la formule (3.7)

$$\exp(z) = \exp(z') \Leftrightarrow z' = z + 2i\pi k \quad \text{pour un certain entier } k \in \mathbb{N}.$$

Ce problème se résout en prescrivant un intervalle semi-ouvert de longueur 2π pour l'argument de la variable complexe considérée. Cela nous conduit à la définition suivante :

Définition. Le *logarithme naturel* de $w \in \mathbb{C} \setminus \{0\}$ est la fonction $\ln : \mathbb{C} \setminus \{0\} \rightarrow \mathbb{C}$ définie par la formule

$$\ln(w) = \ln(r) + i\theta,$$

où $r = |w|$ est le module de w et $\theta = \arg(w)$ est l'argument de w que nous choisissons dans l'intervalle $(-\pi, \pi]$.

Le logarithme naturel est un inverse à droite de l'exponentielle :

$$\exp(\ln(w)) = w,$$

et c'est un inverse à gauche modulo $2i\pi$:

$$\ln(\exp(z)) = z \text{ (modulo } 2i\pi \text{)}.$$

Notons encore que le logarithme naturel ne définit pas un homomorphisme de groupes $(\mathbb{C} \setminus \{0\}, \cdot) \rightarrow (\mathbb{C}, +)$, mais nous avons la formule

$$\ln(w_1 \cdot w_2) = \ln(w_1) + \ln(w_2) \text{ (modulo } 2i\pi \text{)},$$

qui se déduit facilement du théorème 3.4.5.

Chapitre 4

Espaces vectoriels

4.1 Définitions et premiers exemples.

Pour motiver la définition des espaces vectoriels, nous commençons par une constructions simples sur les groupes abéliens. Soit donc $(V, +)$ est groupe abélien et $k \in \mathbb{N}$ un entier. On peut alors définir pour tout élément $v \in V$ un nouvel élément

$$k \cdot v := \underbrace{v + v + \cdots + v}_{k \text{ termes}} \in V.$$

On peut aussi définir $0 \cdot v$ comme étant l'élément neutre de V et si k est un entier négatif, on définit $k \cdot v$ comme étant l'opposé de $(-k) \cdot v$. On a ainsi défini une opération externe

$$\begin{aligned} \mathbb{Z} \times V &\longrightarrow V \\ (k, v) &\mapsto k \cdot v. \end{aligned}$$

Il n'est pas difficile de vérifier les propriétés suivantes (pour tous $v, v_1, v_2 \in V$ et $k, k_1, k_2 \in \mathbb{Z}$) :

- i.) $(k_1 + k_2) \cdot v = k_1 \cdot v + k_2 \cdot v$.
- ii.) $k \cdot (v_1 + v_2) = k \cdot v_1 + k \cdot v_2$.
- iii.) $(k_1 \cdot k_2) \cdot v = k_1 \cdot (k_2 \cdot v)$.
- iv.) $1 \cdot v = v$.

En général, il n'est pas possible de définir un élément $\frac{1}{2} \cdot v$ pour un élément $v \in V$, ni *a fortiori* $\sqrt{2} \cdot v$, $\pi \cdot v$, etc. Le concept d'espace vectoriel nous permet d'étendre l'opération $v \mapsto k \cdot v$ précédente à tout élément $k \in \mathbb{R}$, et même à tout élément d'un corps quelconque.

Définition. Soit K un corps. Un *espace vectoriel* sur le corps K (on dit aussi un K -espace vectoriel) est un groupe abélien $(V, +)$ muni d'une loi de composition externe

$$K \times V \rightarrow V,$$

notée multiplicativement $(\lambda, v) \mapsto \lambda \cdot v$, et pour laquelle on admet les quatre propriétés suivantes pour tous $v, w \in V$ et tous $\lambda, \mu \in K$.

- (i) $(\lambda + \mu) \cdot v = \lambda \cdot v + \mu \cdot v$.
- (ii) $\lambda \cdot (v + w) = \lambda \cdot v + \lambda \cdot w$.
- (iii) $(\lambda \cdot \mu) \cdot v = \lambda(\mu \cdot v)$.
- (iv) $1 \cdot v = v$ (ici 1 est l'unité dans le corps K).

Remarques. Les éléments de l'espace vectoriel V sont appelés les *vecteurs* et les éléments du corps K sont appelés les *scalaires*. Il faut se souvenir qu'il y a deux éléments nuls : l'élément neutre du groupe $(V, +)$ et l'élément neutre du groupe $(K, +)$. Si on souhaite éviter un risque de confusion, on note $0_K \in K$ le zero scalaire et $0_V \in V$ (ou **0**) le zero vectoriel.

Exemples d'espaces vectoriels.

1. Le plus petit espace vectoriel ne possède qu'un élément : $V = \{0_V\}$, on l'appelle l'*espace vectoriel nul* ou l'*espace vectoriel trivial*.
2. Tout corps K est un espace vectoriel sur lui-même.
3. Plus généralement, si K est un sous-corps d'un corps L , alors L est un K -espace vectoriel. En particulier $\mathbb{C}$ est un $\mathbb{R}$ -espace vectoriel et $\mathbb{R}$ est un $\mathbb{Q}$ -espace vectoriel.
4. L'ensemble de n -tuples d'éléments du corps, i.e. l'ensemble

$$K^n = \{(x_1, x_2, \dots, x_n) \mid x_j \in K\}$$

est un espace vectoriel si l'on définit l'addition interne et la multiplication par un scalaire terme-à-terme :

$$(x_1, x_2, \dots, x_n) + (y_1, y_2, \dots, y_n) = (x_1 + y_1, x_2 + y_2, \dots, x_n + y_n),$$

et

$$\lambda \cdot (x_1, x_2, \dots, x_n) = (\lambda \cdot x_1, \lambda \cdot x_2, \dots, \lambda \cdot x_n).$$

5. Si V_1 et V_2 sont deux espaces vectoriels, alors le produit cartésien $V_1 \times V_2$ admet aussi une structure naturelle d'espace vectoriel.
6. Si X est un ensemble non vide quelconque, on note K^X l'ensemble des fonctions f de X vers K . Alors K^X est un espace vectoriel si on définit $f + g$ et $\lambda \cdot f$ par

$$(f + g)(x) := f(x) + g(x) \quad \text{et} \quad (\lambda \cdot f)(x) = \lambda \cdot (f(x))$$

(remarquer que dans cet exemple, les "vecteurs" sont les fonctions $f : X \rightarrow K$).

7. Si $I \subset \mathbb{R}$ est un intervalle, l'ensemble des fonctions $f : I \rightarrow \mathbb{R}$ qui sont continues forme un espace vectoriel sur le corps $\mathbb{R}$ pour les opérations définies dans l'exemple précédent. On le note $C^0(I)$.
8. On appelle *monôme d'indéterminée t* à coefficient dans K une expression $a \cdot t^d$ où $a \in K$ et $d \in \mathbb{N}$ est appelé le *degré* du monôme. Un *polynôme* est une somme de monômes :

$$P(t) = a_0 + a_1 t + a_2 t^2 + \dots + a_m t^m.$$

La somme de deux polynômes s'effectue en sommant les monômes de même degré, et la multiplication par un scalaire $\lambda \in K$ s'effectue aussi sur chaque monôme. L'ensemble des polynômes à coefficients dans le corps K est ainsi un K -espace vectoriel, on le note $K[t]$.

Proposition 4.1.1. Soit V un espace vectoriel sur le corps K , on note $0_V \in V$ le zero vectoriel et $0_K \in K$ le zero scalaire. Alors, on a pour tous $\lambda, \mu \in K$ et $v \in V$:

- a) $\lambda \cdot 0_V = 0_V$.
- b) $0_K \cdot v = 0_V$.
- c) Si $\lambda \cdot v = 0_V$ alors ou bien $\lambda = 0_K$, ou bien $v = 0_V$.
- d) $(-\lambda) \cdot v = \lambda \cdot (-v) = -(\lambda \cdot v)$.
- e) $(\lambda - \mu) \cdot v = \lambda \cdot v - \mu \cdot v$.

Preuve. (a) On a pour tous $\lambda \in K$ et $v \in V$

$$\lambda \cdot 0_V = \lambda \cdot (0_V + 0_V) = \lambda \cdot 0_V + \lambda \cdot 0_V,$$

En simplifiant à gauche (i.e. en ajoutant $-(\lambda \cdot 0_V)$), on obtient $0_V = \lambda \cdot 0_V$.

(b) C'est le même type de raisonnement que pour (a) : on a

$$0_K \cdot v = (0_K + 0_K) \cdot v = 0_K \cdot v + 0_K \cdot v,$$

en simplifiant on trouve $0_K \cdot v = 0_V$.

(c) Supposons que $\lambda \cdot v = 0_V$ et que $\lambda \neq 0$. Alors il existe $\lambda^{-1} \in K$ (l'inverse multiplicatif de λ) et on a

$$v = 1 \cdot v = (\lambda^{-1} \lambda) \cdot v = \lambda^{-1} \cdot (\lambda \cdot v) = \lambda^{-1} \cdot 0_V = 0_V.$$

(d) On a

$$0_V = 0_K \cdot v = (\lambda + (-\lambda)) \cdot v = \lambda \cdot v + (-\lambda) \cdot v.$$

donc $(-\lambda) \cdot v = -(\lambda \cdot v)$. De façon similaire, on a

$$\lambda \cdot (-v) + \lambda \cdot v = \lambda \cdot (-v + v) = \lambda \cdot 0_V = 0_V$$

et on conclut comme dans le cas précédent que $\lambda \cdot (-v) = -(\lambda \cdot v)$.

(e) Finalement, on a

$$(\lambda - \mu) \cdot v = (\lambda + (-\mu)) \cdot v = \lambda \cdot v + (-\mu) \cdot v = \lambda \cdot v + (-\mu \cdot v) = \lambda \cdot v - \mu \cdot v.$$

□

4.2 Sous-espaces vectoriels et applications linéaires

Définition. Si V est un espace vectoriel sur le corps K et $W \subset V$, alors on dit que W est un *sous-espace vectoriel* (abrégé s.e.v.) si W est un espace vectoriel pour les mêmes opérations d'addition et de multiplication par un scalaire que dans l'espace vectoriel V .

Proposition 4.2.1. Une partie W d'un K -espace vectoriel V est un sous espace vectoriel si et seulement si $W \neq \emptyset$ et pour tous $\lambda \in K$ et $v, w \in W$, on a

$$\lambda \cdot v \in W \quad \text{et} \quad v + w \in W.$$

Preuve. La preuve est très simple. Montrons d'abord que pour tout $w \in W$ on a $(-w) \in W$. En effet on a d'après l'hypothèse de la proposition :

$$w \in W \Rightarrow -w = (-1)w \in W.$$

Comme d'autre part W est non vide, alors $v, w \in W \Rightarrow v + w \in W$. Le sous-ensemble $W \subset V$ est donc un sous groupe de $(V, +)$. La loi de multiplication par les scalaires $K \times W \rightarrow W$ est bien définie par hypothèse et les conditions (i) à (iv) de la définition d'espaces vectoriels sont *a fortiori* satisfaites pour W puisque V est un K -espace vectoriel. Ainsi W est bien un espace vectoriel sur le corps K .

□

Définition. Une notion importante est la suivante : on dit qu'une expression du type

$$\lambda \cdot v + \mu \cdot w$$

où $\lambda, \mu \in K$ et $v, w \in V$ est une *combinaison linéaire* de v et w .

On peut alors reformuler la proposition précédente en disant que $W \subset V$ est un sous-espace vectoriel si et seulement si W est non vide et est stable pour les combinaisons linéaires :

$$(v, w \in W \quad \text{et} \quad \lambda, \mu \in K) \implies (\lambda \cdot v + \mu \cdot w) \in W.$$

On peut généraliser cette condition à des combinaisons linéaires de plus de deux vecteurs :

$$\lambda_i \in K \quad \text{et} \quad v_i \in W \quad (i = 1, \dots, r) \implies \sum_{i=1}^r \lambda_i v_i \in W.$$

Définition. Si V et W sont deux espaces vectoriels sur le même corps K , alors on dit qu'une application $f : V \rightarrow W$ est une *application linéaire* (ou *K -linéaire* si on veut préciser le corps) si elle respecte les combinaisons linéaires, i.e. si pour tous $\lambda, \mu \in K$ et $v, w \in V$, on a

$$f(\lambda \cdot v + \mu \cdot w) = \lambda \cdot f(v) + \mu \cdot f(w).$$

A nouveau, on peut généraliser cette relation à des combinaisons linéaires de plus de deux vecteurs :

$$f\left(\sum_{i=1}^r \lambda_i \cdot v_i\right) = \sum_{i=1}^r \lambda_i \cdot f(v_i).$$

Proposition 4.2.2. $f : K^n \rightarrow K$ est une application linéaire de K^n vers K si et seulement s'il existe $a_1, a_2, \dots, a_n \in K$ tels que pour tout $x = (x_1, x_2, \dots, x_n) \in K^n$ on a

$$f(x_1, x_2, \dots, x_n) = a_1 \cdot x_1 + a_2 \cdot x_2 + \dots + a_n \cdot x_n = \sum_{i=1}^n a_i \cdot x_i.$$

Preuve. On vérifie directement qu'une telle application est linéaire. Supposons réciproquement que $f : K^n \rightarrow K$ est une application linéaire quelconque, on considère les vecteurs $e_1, e_2, \dots, e_n \in K^n$ définis par

$$e_1 = (1, 0, 0, \dots, 0), \quad e_2 = (0, 1, 0, \dots, 0), \dots, \quad e_n = (0, 0, 0, \dots, 1)$$

(i.e. toutes les coordonnées de e_i sont nulles, sauf la $i^{\text{ème}}$ coordonnée qui vaut 1). Le système de vecteurs $\{e_1, e_2, \dots, e_n\}$ s'appelle la *base canonique* de K^n .

Alors tout vecteur $x \in K^n$, s'écrit comme combinaison linéaire des e_i :

$$x = (x_1, x_2, \dots, x_n) = x_1 \cdot e_1 + x_2 \cdot e_2 + \dots + x_n \cdot e_n = \sum_{i=1}^n x_i \cdot e_i.$$

Posons $a_i := f(e_i)$, alors a_i est un élément de K et comme f est linéaire, on a

$$f(x) = f\left(\sum_{i=1}^n x_i \cdot e_i\right) = \sum_{i=1}^n x_i \cdot f(e_i) = \sum_{i=1}^n a_i \cdot x_i.$$

□

Exemples d'applications linéaires.

1. On vient de prouver que toute application linéaire de K^n vers K est de la forme $f(x_1, x_2, \dots, x_n) = a_1 \cdot x_1 + a_2 \cdot x_2 + \dots + a_n \cdot x_n$, avec $a_i \in K$ pour tout $i = 1, 2, \dots, n$.
2. Si X est un ensemble quelconque et $a \in X$, alors l'*application d'évaluation* $e_a : K^X \rightarrow K$ définie par

$$e_a(f) := f(a)$$

est linéaire.

3. Les applications $\pi_1 : V_1 \times V_2 \rightarrow V_1$ et $\pi_2 : V_1 \times V_2 \rightarrow V_2$ définies par $\pi_1(v_1, v_2) = v_1$ et $\pi_2(v_1, v_2) = v_2$ (appelées projections canoniques) sont linéaires.
4. Si $f : V \rightarrow W$ et $g : V \rightarrow W$, alors $(\lambda \cdot f + \mu \cdot g)$ est aussi une application linéaire de V vers W pour tous $\lambda, \mu \in K$.
5. Si U, V, W sont trois K -espaces vectoriels et si $g : U \rightarrow V$ et $f : V \rightarrow W$ sont des applications linéaires, alors la composition $f \circ g : U \rightarrow W$ est une application linéaire.
6. L'intégration

$$f \mapsto \int_a^b f(t)dt$$

définit une application linéaire de $C^0([a, b])$ vers $\mathbb{R}$.

Proposition 4.2.3. *Si $f : V \mapsto W$ est une application linéaire entre deux K -espaces vectoriels, alors*

- (a) $f(0_V) = 0_W$
- (b) Le noyau de f , i.e. $\text{Ker}(f) = f^{-1}(0_W) = \{x \in V \mid f(x) = 0_W\}$ est un sous-espace vectoriel de V .
- (c) L'image de f , i.e. $\text{Im}(f) = \{f(v) \mid v \in V\}$ est un sous-espace vectoriel de W .
- (d) L'application f est injective si et seulement si $\text{Ker}(f) = \{0_V\}$.

Preuve. Notons $0_K \in K$ le scalaire nul, i.e. le zéro dans le corps K .

(a) Si f est une application linéaire, alors $f(0_V) = f(0_K \cdot V) = 0_K f(0_V) = 0_W$.

(b) Notons d'abord que $\text{Ker}(f) \neq \emptyset$ puisque $0_V \in \text{Ker}(f)$. Si $v_1, v_2 \in \text{Ker}(f)$ et $\lambda, \mu \in K$, alors

$$f(\lambda v_1 + \mu v_2) = \lambda f(v_1) + \mu f(v_2) = \lambda 0_W + \mu 0_W = 0_W.$$

donc $\lambda v_1 + \mu v_2 \in \text{Ker}(f)$ ce qui prouve que $\text{Ker}(f)$ est un sous-espace vectoriel de V .

(c) Notons d'abord que $\text{Im}(f) \neq \emptyset$ puisque $0_W \in \text{Im}(f)$. Supposons que $\lambda, \mu \in K$ et $w_1, w_2 \in \text{Im}(f)$ alors par définition il existe $v_1, v_2 \in V$ tels que $f(v_1) = w_1$ et $f(v_2) = w_2$ et nous avons

$$\lambda w_1 + \mu w_2 = \lambda f(v_1) + \mu f(v_2) = f(\lambda v_1 + \mu v_2) = f(v)$$

avec $v = \lambda v_1 + \mu v_2 \in V$. Donc $\lambda w_1 + \mu w_2 \in W$.

(d) Si f est injective, alors $\text{Ker}(f) = \{0_V\}$ car on sait que $f(0_V) = 0_W$, donc

$$v \in \text{Ker}(f) \Rightarrow f(v) = 0_W = f(0_V) \Rightarrow v = 0_V.$$

Pour montrer l'affirmation réciproque, on suppose que $\text{Ker}(f) = \{0_V\}$ et on prouve que ceci implique que f est injective. En effet, on a pour tous $v_1, v_2 \in V$:

$$\begin{aligned} f(v_1) = f(v_2) &\Rightarrow f(v_1) - f(v_2) = 0_W \Rightarrow f(v_1 - v_2) = 0_W \\ &\Rightarrow (v_1 - v_2) \in \text{Ker}(f) \Rightarrow (v_1 - v_2) = 0_V \\ &\Rightarrow v_1 = v_2. \end{aligned}$$

□

Exemples de sous espaces vectoriels

1. Si V et W sont des espaces vectoriels et $f : V \rightarrow W$ est linéaire, alors $\text{Ker}(f) \subset V$ est un sous-espace vectoriel (voir plus haut).
2. En particulier,

$$\{(x_1, x_2, \dots, x_n) \in K^n \mid a_1 \cdot x_1 + a_2 \cdot x_2 + \dots + a_n \cdot x_n = 0\}$$

est un sous-espace vectoriel de K^n .

3. Si W_1 et W_2 sont deux sous-espaces vectoriels de V , alors $W_1 \cap W_2$ est aussi un sous-espace vectoriel de V .
4. Si W_1 et W_2 sont deux sous-espaces vectoriels de V , alors l'ensemble

$$\{w_1 + w_2 \mid w_1 \in W_1, w_2 \in W_2\} \subset V$$

est encore un sous-espace vectoriel de V . On le note $W_1 + W_2$ et on l'appelle la *somme vectorielle* de W_1 et W_2 .

4.3 Familles libres, liées, génératrices et bases.

4.3.1 Définitions

En algèbre linéaire, on parle souvent de *famille* de vecteurs. Ce mot est synonyme de sous-ensemble, ou partie, d'un espace vectoriel donné V . Une famille de vecteurs peut être finie ou infinie.

Définition. Si $E \subset V$ est une famille de vecteurs du K -espace vectoriel V , alors on dit qu'un vecteur $v \in V$ est *combinaison linéaire d'éléments de E* si on peut écrire

$$v = \lambda_1 v_1 + \lambda_2 v_2 + \cdots + \lambda_m v_m$$

où $\{v_1, v_2, \dots, v_m\} \subset E$ et $\{\lambda_1, \lambda_2, \dots, \lambda_m\} \subset K$. La combinaison linéaire est dite *triviale* si chaque $\lambda_i = 0$ (dans ce cas on a bien sûr $v = 0$).

Définition. Soient V un K -espace vectoriel et $E \subset V$ une famille non vide. On dit qu'un vecteur $w \in V$ est *engendré* par E si w est combinaison linéaire d'éléments de E . On note $\text{Vec}(E)$ l'ensemble des vecteurs engendrés par E , ainsi $w \in \text{Vec}(E)$ si et seulement s'il existe $\{v_1, v_2, \dots, v_m\} \subset E$ et $\{\lambda_1, \lambda_2, \dots, \lambda_m\} \subset K$ tels que

$$w = \lambda_1 v_1 + \lambda_2 v_2 + \cdots + \lambda_m v_m.$$

Si E est l'ensemble vide, alors on convient que le sous-espace vectoriel engendré est le sous-espace trivial, i.e. $\text{Vec}(\emptyset) = \{0_V\}$.

Exemple. On souhaite vérifier si les vecteurs $(1, 2, 3)$ et $(3, 0, 1)$ de $\mathbb{R}^3$ appartiennent à $W = \text{Vec}(\{(4, 5, 6), (7, 8, 9)\})$. Considérons d'abord le premier vecteur, la question est celle de l'existence (ou non) de $\lambda, \mu \in \mathbb{R}$ tels que

$$(1, 2, 3) = \lambda(4, 5, 6) + \mu(7, 8, 9) = (4\lambda + 7\mu, 5\lambda + 8\mu, 6\lambda + 9\mu).$$

Il s'agit donc de décider si on peut résoudre le système de trois équations à deux inconnues suivant :

$$\begin{aligned} 4\lambda + 7\mu &= 1 \\ 5\lambda + 8\mu &= 2 \\ 6\lambda + 9\mu &= 3 \end{aligned}$$

Ce système est simple à résoudre, il y a une unique solution qui est $\lambda = 2, \mu = -1$. Ainsi $(1, 2, 3)$ est la combinaison linéaire suivante de $(7, 8, 9)$ et $(4, 5, 6)$:

$$(1, 2, 3) = 2 \cdot (4, 5, 6) - (7, 8, 9),$$

et c'est donc un élément de W .

Pour le second vecteur, le système à résoudre est

$$\begin{aligned} 4\lambda + 7\mu &= 3 \\ 5\lambda + 8\mu &= 0 \\ 6\lambda + 9\mu &= 1 \end{aligned}$$

Ce système n'a aucune solution (on dit qu'il est incompatible). Pour le voir on peut soustraire la première équation de la seconde, puis la seconde de la troisième, cela donne

$$\begin{aligned}\lambda + \mu &= -3 \\ \lambda + \mu &= 1\end{aligned}$$

Qui n'a clairement aucune solution, par conséquent $(3, 0, 1) \notin W$.

Proposition 4.3.1. *Soit $E \subset V$ un sous-ensemble d'un espace vectoriel V . Alors $\text{Vec}(E)$ est un sous-espace vectoriel de V ; et c'est le plus petit sous-espace vectoriel qui contient E .*

Preuve. Si $E = \emptyset$ alors $\text{Vec}(E) = \{0\}$ et il n'y a rien à montrer. Supposons donc que $E \subset V$ est non vide, et notons $W = \text{Vec}(E)$, c'est donc l'ensemble de toutes les combinaisons linéaires d'éléments de E , i.e. $w \in W$ si et seulement s'il existe $v_1, v_2, \dots, v_m \in E$ et $\lambda_1, \lambda_2, \dots, \lambda_m \in K$ tels que $w = \lambda_1 v_1 + \lambda_2 v_2 + \dots + \lambda_m v_m$. On doit prouver que W est un sous-espace vectoriel de V . Il est clair que $E \subset W$, donc en particulier $W \neq \emptyset$, supposons maintenant que $w, w' \in W$ et $\alpha, \beta \in K$, alors on peut écrire

$$w = \lambda_1 v_1 + \lambda_2 v_2 + \dots + \lambda_m v_m \quad \text{et} \quad w' = \lambda'_1 v'_1 + \lambda'_2 v'_2 + \dots + \lambda'_m v'_m,$$

où tous les λ_i, λ'_j appartiennent à K et tous les v_i, v'_j appartiennent à E . On a alors

$$\alpha w + \beta w' = \alpha \lambda_1 v_1 + \alpha \lambda_2 v_2 + \dots + \alpha \lambda_m v_m + \beta \lambda'_1 v'_1 + \beta \lambda'_2 v'_2 + \dots + \beta \lambda'_m v'_m,$$

qui est encore une combinaison linéaire d'éléments de E , donc $\alpha w + \beta w' \in W$. On a donc montré que W est stable pour les combinaisons linéaires, c'est donc bien un sous-espace vectoriel.

Il est clair que $E \subset W$ (pour tout $v \in E$ on peut considérer la combinaison linéaire $1 \cdot v$), d'autre part si $U \subset V$ est un autre sous-espace vectoriel qui contient E , alors il doit contenir toutes les combinaisons linéaires d'éléments de E , il doit donc contenir W . On a donc $W \subset U \subset V$ et W est ainsi bien le plus petit sous-espace vectoriel qui contient E . □

Définition. On dit que la famille $E \subset V$ est *liée* si on peut écrire $0 \in V$ comme combinaison linéaire non triviale d'éléments de E , i.e. s'il existe $\{\lambda_1, \lambda_2, \dots, \lambda_m\} \in K$ et $\{v_1, v_2, \dots, v_m\} \in E$ tels que les λ_i ne sont pas tous nuls et

$$\lambda_1 v_1 + \lambda_2 v_2 + \dots + \lambda_m v_m = 0_V.$$

On dit aussi que les vecteurs de E sont *linéairement dépendants*.

Lemme 4.3.2. *Une famille $E \subset V$ contenant au moins deux éléments est liée si et seulement si l'un des éléments de E peut s'écrire comme combinaison linéaire des autres.*

Preuve. Si $E \subset V$ est liée, alors zero s'écrit comme combinaison linéaire non triviale d'éléments $v_i \in E$:

$$0 = \lambda_1 v_1 + \lambda_2 v_2 + \dots + \lambda_m v_m,$$

avec au moins un des $\lambda_i \neq 0$. Supposons par exemple que $\lambda_m \neq 0$, alors on a

$$v_m = \mu_1 v_1 + \mu_2 v_2 + \dots + \mu_{m-1} v_{m-1},$$

avec $\mu_i := -\frac{\lambda_i}{\lambda_m}$ pour $i = 1, 2, \dots, m-1$.

Inversément, supposons qu'un élément $v \in E$ soit combinaison linéaire d'autres éléments $v_i \in E$:

$$v = \lambda_1 v_1 + \lambda_2 v_2 + \dots + \lambda_m v_m,$$

(avec $v_i \neq v$ pour tout i), alors on a une combinaison linéaire non triviale d'éléments de E qui donne zero :

$$\lambda_1 v_1 + \lambda_2 v_2 + \cdots + \lambda_m v_m - v = 0.$$

□

Définition. La famille $E \subset V$ est *libre* si elle n'est pas liée, i.e. toute combinaison linéaire d'éléments de E qui donne le vecteur nul est triviale :

$$\lambda_1 v_1 + \lambda_2 v_2 + \cdots + \lambda_m v_m = 0_V \quad \Rightarrow \quad \lambda_1 = \lambda_2 = \cdots = \lambda_m = 0.$$

On dit aussi que les vecteurs de E sont *linéairement indépendants*.

Remarquons que, par définition, l'ensemble vide est une famille libre.

Exemples 1. Les vecteurs $(1, 0)$ et $(1, 1)$ sont linéairement indépendants dans K^2 pour tout corps K (donc la famille $\{(1, 0), (1, 1)\}$ est libre).

Preuve. Si $\lambda(1, 0) + \mu(1, 1) = (0, 0)$, alors

$$\begin{aligned} \lambda + \mu &= 0 \\ \mu &= 0 \end{aligned}$$

et donc $\mu = \lambda = 0$.

2. Si une famille $E \subset V$ contient le vecteur nul, alors elle est liée.

3. La famille $\{v, w\} \subset V$ est liée si et seulement si ces vecteurs sont colinéaires, c'est-à-dire si et seulement s'il existe $\lambda \in K$ tel que $v = \lambda w$ ou $w = \lambda v$.

4. Les fonctions sin et cos sont des vecteurs linéairement indépendants de l'espace vectoriel $C^0(\mathbb{R}, \mathbb{R})$ des fonctions continues de $\mathbb{R}$ dans $\mathbb{R}$.

Preuve. Supposons que $a \cdot \cos + b \cdot \sin$ soit la fonction identiquement nulle, i.e. $a \cos(x) + b \sin(x) = 0$ pour tout $x \in \mathbb{R}$, alors en particulier $0 = a \cdot \cos\left(\frac{\pi}{2}\right) + b \cdot \sin\left(\frac{\pi}{2}\right) = a \cdot 0 + b \cdot 1$, donc $b = 0$ et de même $0 = a \cdot \cos(0) + b \cdot \sin(0) = a \cdot 1 + b \cdot 0$, donc $a = 0$.

Proposition 4.3.3. *Tout sous-ensemble d'une famille libre d'un K -espace vectoriel V est une famille libre.*

On peut reformuler cette proposition ainsi : *Toute famille de V qui contient une famille liée est elle-même liée.*

Preuve. Soit $E \subset V$ une famille libre et $S \subset E$. S'il existe $\{\lambda_1, \lambda_2, \dots, \lambda_m\} \subset K$ et $\{v_1, v_2, \dots, v_m\} \subset S$ tels que $\lambda_1 v_1 + \lambda_2 v_2 + \cdots + \lambda_m v_m = 0_V$, alors tous les λ_i sont nuls car les v_i appartiennent à E qui est une famille libre. Par conséquent S est aussi une famille libre.

□

Définition. On dit que la famille $E \subset V$ *engendre* l'espace vectoriel V si $V = \text{Vec}(E)$, i.e. tout vecteur de V peut s'écrire comme combinaison linéaire d'éléments de E :

$$\forall w \in V, \exists \lambda_1, \lambda_2, \dots, \lambda_m \in K, \exists v_1, v_2, \dots, v_m \in E \text{ t. q. } w = \lambda_1 v_1 + \lambda_2 v_2 + \cdots + \lambda_m v_m.$$

On dit aussi que E est une *famille génératrice* pour V .

Remarquons que tout espace vectoriel contient des familles génératrices (on peut prendre l'espace vectoriel V lui même).

Il est évident à partir de la définition que toute famille de V qui contient une partie génératrice est elle-même génératrice.

4.3.2 Bases et dimension

Définitions. a.) On dit qu'un espace vectoriel V est de *génération finie* (ou de *dimension finie*) s'il existe une famille finie E de V qui engendre V . Dans le cas contraire, on dit que V est de *dimension infinie*.

b.) La famille $B \subset V$ est une *base* de l'espace vectoriel V si B est une famille libre et si B engendre V .

Exemple. La famille $E = \{(2, 1), (1, -1), (5, 1)\}$ est une famille liée dans l'espace vectoriel $\mathbb{R}^2$ car on peut trouver $\lambda, \mu, \nu \in \mathbb{R}$ tels que

$$\lambda(2, 1) + \mu(1, -1) + \nu(1, -2) = (0, 0),$$

on peut par exemple choisir $\lambda = 1, \mu = -5$ et $\nu = 3$. En revanche, on peut facilement vérifier que les trois sous-familles $\{(2, 1), (1, -1)\}$, $\{(2, 1), (5, 1)\}$ et $\{(1, -1), (5, 1)\}$ sont des bases de $\mathbb{R}^2$.

Théorème 4.3.4. *La famille B est une base de l'espace vectoriel (non nul) V si et seulement si tout vecteur $v \in V$ non nul peut s'écrire d'une unique manière comme combinaison linéaire d'éléments de B (on ne tiens pas compte des coefficients nuls).*

Preuve. Par définition toute base engendre l'espace vectoriel V , donc tout vecteur de V peut s'écrire comme combinaison linéaire d'éléments de B . Il faut montrer l'unicité de cette écriture.

Supposons que

$$x = \lambda_1 v_1 + \lambda_2 v_2 + \cdots + \lambda_m v_m = \mu_1 w_1 + \mu_2 w_2 + \cdots + \mu_n w_n$$

avec $v_i, w_j \in B$ et $\lambda_i, \mu_j \in K^*$ (rappelons que $K^* = K \setminus \{0\}$). Notons $s = \text{Card}(\{v_1, v_2, \dots, v_m\} \cap \{w_1, w_2, \dots, w_n\})$, donc s est le nombre de vecteurs communs entre les v_i et les w_j . Quitte à renuméroter les vecteurs, on peut supposer que

$$v_1 = w_1, v_2 = w_2, \dots, v_s = w_s.$$

On a alors

$$0 = x - x = (\lambda_1 - \mu_1)v_1 + \cdots + (\lambda_s - \mu_s)v_s + \lambda_{s+1}v_{s+1} + \cdots + \lambda_m v_m - \mu_{s+1}w_{s+1} - \cdots - \mu_n w_n,$$

c'est une combinaison linéaire d'éléments de B qui donne zéro. Puisque B est une famille libre, on a

$$(\lambda_1 - \mu_1) = \cdots = (\lambda_s - \mu_s) = \lambda_{s+1} = \cdots = \lambda_m = \mu_{s+1} = \cdots = \mu_n = 0.$$

Or nous avons supposé $\lambda_i \neq 0$ et $\mu_i \neq 0$. On en conclut que $m = n = s$ et que $\lambda_i = \mu_i$ pour tout i .

Pour prouver la réciproque, nous allons montrer que si B est une famille liée, alors il existe ou moins un vecteur de V qui peut s'écrire de deux manières différentes comme combinaison linéaire d'éléments de B . Nous distinguons le cas où B contient le vecteur nul et celui où tous les vecteurs de B sont non nuls.

Si $v_1 = 0 \in B$, alors il existe au moins un autre vecteur $v_2 \in B$ non nul (car on a supposé que V n'est pas l'espace vectoriel nul). Nous avons alors deux façons d'écrire le vecteur v_2 comme combinaison linéaire d'éléments de B car $v_2 = 1 \cdot v_2 = 1 \cdot v_1 + 1 \cdot v_2$.

Considérons maintenant le cas où B est une famille liée qui ne contient pas le vecteur nul. Alors on peut écrire

$$\lambda_1 v_1 + \lambda_2 v_2 + \cdots + \lambda_m v_m = 0,$$

Avec les $v_i \in B$ et les $\lambda_i \in K$ non tous nuls. Supposons par exemple $\lambda_1 \neq 0$, alors on obtient deux façons d'écrire le vecteur $\lambda_1 v_1$ comme combinaison linéaire d'éléments de B :

$$\lambda_1 v_1 = -\lambda_2 v_2 - \cdots - \lambda_m v_m$$

□

- p. 58 ; Théorème 4.3.4 : si on exclut le vecteur nul la preuve n'est pas suffisante : Pour montrer l'implication " $\leq$ ", on doit d'abord montrer que 0 n'appartient pas à B (sinon, $v_1 = 0$ est possible)

Exemples 1. Les vecteurs $e_1 = (1, 0, 0, \dots, 0)$, $e_2 = (0, 1, 0, \dots, 0)$, $\dots$, $e_n = (0, 0, \dots, 1)$ forment une base de K^n car tout $x = (x_1, x_2, \dots, x_n)$ s'écrit de façon unique comme combinaison linéaire de ces vecteurs :

$$x = \sum_{i=1}^n x_i e_i$$

On appelle cette base la *base canonique* de K^n .

2. Les monômes $\{1, t, t^2, t^3, \dots\} = \{t^k \mid k \in \mathbb{N}\}$ forment une base de l'espace vectoriel $\mathbb{R}[t]$ des polynômes à coefficients réels, car tout polynôme s'écrit de façon unique sous la forme

$$P(t) = a_0 + a_1 t + a_2 t^2 + \dots + a_n t^n,$$

avec $a_0, \dots, a_n \in \mathbb{R}$.

Définition. Si $B = \{b_1, b_2, \dots, b_n\}$ est une base du K -espace vectoriel V , alors on dit que les scalaires $x_1, x_2, \dots, x_n \in K$ sont les *composantes* du vecteur $v \in V$ par rapport à la base B si

$$v = x_1 b_1 + x_2 b_2 + \dots + x_n b_n.$$

La proposition précédente nous dit que cette notion est bien définie (à condition que la base soit finie et numérotée).

Théorème 4.3.5 (Théorème de complétion des bases). *Soit E une famille génératrice finie d'un K -espace vectoriel V (non nul) et $L \subset E$ une sous-famille libre. Alors il existe une base B de V telle que $L \subset B \subset E$.*

Preuve. Notons $L = \{b_1, b_2, \dots, b_r\}$. On dira que L est une *famille libre maximale* dans E si pour tout $v \in E \setminus L$, la famille $L \cup \{v\}$ est liée. Nous définissons maintenant B de la façon suivante : si L est une famille libre maximale dans E , on pose $B = L$. Sinon, il existe au moins un vecteur $b_{r+1} \in E \setminus L$ tel que $L \cup \{b_{r+1}\}$ est libre et on pose $L_1 = L \cup \{b_{r+1}\}$. Si L_1 est maximale, on pose $B = L_1$ et sinon on recommence. Le processus s'arrête après un nombre fini d'étapes car E est une famille finie. Il existe donc un sous-ensemble $\{b_{r+1}, \dots, b_{r+m}\} \subset E$ tel que la famille

$$B = L \cup \{b_{r+1}, \dots, b_{r+m}\} = \{b_1, b_2, \dots, b_{r+m}\} \subset E$$

est une famille libre maximale dans E .

Montrons que B est une base de V . Il suffit de montrer que B engendre l'espace vectoriel V . Prenons d'abord un vecteur $v \in E$, alors, par hypothèse, $B \cup \{v\}$ n'est pas une famille libre, il existe donc une combinaison linéaire non triviale du type

$$\lambda_1 b_1 + \lambda_2 b_2 + \dots + \lambda_n b_n + \mu v = 0.$$

On observe que μ ne peut pas être nul, car si on avait $\mu = 0$, alors on aurait une combinaison linéaire non triviale des éléments de B qui donne zéro. On obtient donc le vecteur v comme combinaison linéaire des vecteurs de B :

$$v = -\frac{1}{\mu}(\lambda_1 b_1 + \lambda_2 b_2 + \dots + \lambda_n b_n).$$

Ainsi tout vecteur de E est combinaison linéaire des vecteurs de B . Cela implique que $E \subset \text{Vec}(B)$ et donc

$$\text{Vec}(B) = \text{Vec}(E) = V.$$

On a montré que B est une partie libre de V qui engendre V . C'est donc une base. Il est par ailleurs clair par construction que $L \subset B \subset E$.

□

Corollaire 4.3.6. *Tout espace vectoriel de dimension finie contient une base, plus précisément toute famille libre peut être complétée en une base.*

Preuve. Soit V un espace vectoriel de dimension finie et $L \subset V$ une famille libre (le cas $L = \emptyset$ est possible). Par hypothèse, il existe une famille finie $S \subset V$ qui engendre V . Posons $E = L \cup S$, alors E engendre aussi V et par le théorème précédent, on sait qu'il existe une base B de V telle que $L \subset B$. $\square$

Le résultat le plus important de ce chapitre est le théorème suivant :

Théorème 4.3.7. *Si l'espace vectoriel V est engendré par une famille de m éléments, alors toute partie libre de V contient au plus m éléments.*

Preuve. On suppose qu'il existe un sous-ensemble de cardinal m qui engendre V et on doit conclure que si $L = \{b_1, \dots, b_r\} \subset V$ est une famille libre dans V alors $r \leq m$.

Notons $k \in \mathbb{N}$ le plus grand entier tel qu'il existe un sous-ensemble $E \subset V$ tel que

$$\text{Vec}(E) = V, \text{ Card}(E) = m \quad \text{et} \quad \text{Card}(L \cap E) = k.$$

Il est clair que $0 \leq k \leq \min\{r, m\}$, nous allons démontrer que $k = r$. En effet, supposons par l'absurde que $k < r$, alors il existe une famille $E \subset V$ de cardinal m qui engendre V et qui contient exactement k éléments de L . Quitte à renuméroter les éléments de L , on peut supposer que $L \cap E = \{b_1, \dots, b_k\}$ et on peut donc noter les éléments de E par

$$E = \{b_1, \dots, b_k, v_{k+1}, v_{k+2}, \dots, v_m\}.$$

Nous allons construire une nouvelle famille E' de cardinal m , qui engendre E et qui contient $k+1$ éléments de L , contredisant la définition de k .

Par hypothèse E engendre V , donc nous pouvons écrire b_{k+1} comme combinaison linéaire des éléments de E :

$$b_{k+1} = \alpha_1 b_1 + \dots + \alpha_k b_k + \alpha_{k+1} v_{k+1} + \dots + \alpha_m v_m$$

Nous affirmons qu'il existe $j > k$ tel que $\alpha_j \neq 0$. En effet, dans le cas contraire nous aurions

$$b_{k+1} = \alpha_1 b_1 + \dots + \alpha_k b_k,$$

contredisant l'hypothèse que la famille L est libre.

Quitte à renuméroter une nouvelle fois les vecteurs de E , nous pouvons supposer que $\alpha_{k+1} \neq 0$. On peut alors écrire

$$v_{k+1} = \frac{1}{\alpha_{k+1}} \left(b_{k+1} - \sum_{i=1}^k \alpha_i b_i - \sum_{j=k+2}^m \alpha_j v_j \right).$$

Nous pouvons réécrire cette combinaison linéaire de façon plus simple :

$$v_{k+1} = \sum_{i=1}^{k+1} \beta_i b_i + \sum_{j=k+2}^m \beta_j v_j. \quad (4.1)$$

On définit maintenant une nouvelle famille de vecteurs par

$$E' = (E \setminus \{v_{k+1}\}) \cup \{b_{k+1}\} = \{b_1, \dots, b_k, b_{k+1}, v_{k+2}, \dots, v_m\} \subset V$$

(on a échangé le vecteur v_{k+1} par le vecteur b_{k+1}), et nous affirmons que E' engendre l'espace vectoriel V .

En effet, soit $x \in V$ un vecteur quelconque. Puisqu'on a supposé que E engendre V , on peut trouver des scalaires λ_i tels que

$$x = \lambda_1 b_1 + \cdots + \lambda_k b_k + \lambda_{k+1} v_{k+1} + \lambda_{k+2} v_{k+2} + \cdots + \lambda_m v_m$$

En remplaçant v_{k+1} par la combinaison linéaire (4.1), on obtient x comme une combinaison linéaire du type

$$x = \mu_1 b_1 + \cdots + \mu_k b_k + \mu_{k+1} b_{k+1} + \mu_{k+2} v_{k+2} + \cdots + \mu_m v_m,$$

cela démontre que E' engendre V . Nous avons donc construit un ensemble E' de cardinal m qui engendre V et tel que

$$\text{Card}(L \cap E') = k + 1,$$

ce qui est impossible puisque k est supposé maximal avec ces propriétés. Cette contradiction prouve que $k = r$, cela signifie que $L \subset E$ et donc $\text{Card}(L) \leq \text{Card}(E)$. □

Corollaire 4.3.8. *Si V est un espace vectoriel de génération finie et si B_1 et B_2 sont deux bases de V , alors $\text{Card}(B_1) = \text{Card}(B_2)$.*

Définition. Soit V un espace vectoriel de génération finie sur un corps K . On appelle *dimension* de V le cardinal d'une base quelconque B de V et on note $\dim(V) = \text{Card}(B)$. Si on veut préciser le corps de base, on note $\dim_K(V) = \text{Card}(B)$. Si V n'est pas de génération finie, on note $\dim(V) = \infty$.

Par convention, on considère que la dimension de l'espace vectoriel trivial $\{0\}$ est nulle. Le corollaire précédent implique que la notion de dimension est bien définie (indépendamment du choix d'une base).

Voyons quelques exemples :

- K est un espace vectoriel sur lui-même. L'ensemble $\{1\} \subset V$ forme une base, donc K est un K -espace vectoriel de dimension 1.
- K^n est un K -espace vectoriel de dimension n , une base est (par exemple) donnée par la base canonique.
- Si V est un espace vectoriel de dimension finie, alors tout sous-espace vectoriel $W \subset V$ est aussi de dimension finie et $\dim(W) \leq \dim(V)$.
- La dimension d'un produit direct est donnée par $\dim(V_1 \times V_2) = \dim(V_1) + \dim(V_2)$.
- La notion de dimension d'un espace vectoriel dépend du corps de base considéré. Par exemple $\mathbb{C}$ est un espace vectoriel de dimension 1 sur le corps $\mathbb{C}$, mais de dimension 2 sur le corps $\mathbb{R}$ (une base réelle de $\mathbb{C}$ est donnée par $\{1, i\}$). Plus généralement $\mathbb{C}^n$ est un espace vectoriel de dimension n sur $\mathbb{C}$, mais de dimension $2n$ sur $\mathbb{R}$.
- L'espace des fonctions $f : \mathbb{R} \rightarrow \mathbb{R}$ et l'espace des polynômes $\mathbb{R}[t]$ sont de dimension infinie.
- Si on regarde $\mathbb{R}$ comme $\mathbb{Q}$ -espace vectoriel, alors sa dimension est infinie.

4.3.3 Le théorème d'échange de Grassmann

Le théorème d'échange de Grassmann est une généralisation du théorème 4.3.7 dans laquelle on ne suppose pas nécessairement que l'espace vectoriel V est de génération finie.

Théorème 4.3.9. *Soient V un K -espace vectoriel et $E, L \subset V$ deux familles de vecteurs. Supposons que E engendre V et que L est une famille libre finie. Alors il existe une famille $L' \subset E$ telle que*

$$\text{Card}(L') = \text{Card}(L) \quad \text{et} \quad (E \setminus L') \cup L \text{ engendre } V.$$

Ce théorème est assez subtil. Il dit que si E est une famille génératrice de V (qui peut être infinie) et que par ailleurs on a une famille de r vecteurs linéairement indépendants $\{v_1, v_2, \dots, v_r\}$, alors on peut “échanger” les v_i avec r vecteurs de E et la nouvelle famille ainsi obtenue est encore une famille génératrice.

Preuve. La preuve est une variante de la preuve du théorème 4.3.7. L’argument se fait par récurrence sur $r = \text{Card}(L)$. Si $r = 1$, alors $L = \{b_1\}$ où b_1 est un vecteur non nul de V . Ce vecteur peut s’écrire comme combinaison linéaire d’éléments de la partie génératrice E :

$$b_1 = \alpha_1 v_1 + \alpha_2 v_2 + \dots + \alpha_m v_m,$$

avec $v_i \in E$ et $\alpha_i \in K$. Au moins l’un des scalaires α_j est non nul (car sinon $b_1 = 0$), et on peut supposer (quitte à renuméroter les v_i) que $\alpha_1 \neq 0$. On a alors

$$v_1 = \frac{1}{\alpha_1}(b_1 - \alpha_2 v_2 - \dots - \alpha_m v_m).$$

Si on pose $L' = \{v_1\}$ et $E' = (E \setminus L') \cup L = (E \setminus \{v_1\}) \cup \{b_1\}$, alors on a $\text{Card}(L') = \text{Card}(L) = 1$ et

$$\text{Vec}(E') = \text{Vec}(E) = V.$$

Supposons maintenant que le théorème a été démontré pour un entier $r \in \mathbb{N}$ quelconque, et soit $L = \{b_1, \dots, b_r, b_{r+1}\}$ une famille libre de $r + 1$ vecteurs dans V . Par hypothèse de récurrence, il existe $L' = \{v_1, \dots, v_r\} \subset E$ tel que $E' = (E \setminus L') \cup L$ engendre V . Puisque E' engendre V , on peut en particulier exprimer b_{r+1} comme combinaison linéaire d’éléments de E' :

$$b_{r+1} = \beta_1 b_1 + \dots + \beta_r b_r + \gamma_1 w_1 + \dots + \gamma_s w_s,$$

avec $w_i \in E \setminus L'$ et $\beta_i, \gamma_j \in K$. L’un des scalaires γ_j est non nul car sinon b_{r+1} serait combinaison linéaire de $b_1, \dots, b_r$ et la famille L serait liée. Supposons donc que $\gamma_1 \neq 0$, on a donc

$$w_1 = \frac{1}{\gamma_1}(b_{r+1} - \beta_1 b_1 - \dots - \beta_r b_r - \gamma_2 w_2 - \dots - \gamma_s w_s).$$

On pose maintenant

$$L'' = L' \cup \{w_1\} \quad \text{et} \quad E'' = (E' \setminus \{w_1\}) \cup \{b_{r+1}\} = (E \setminus L') \cup L.$$

Alors $\text{Card}(L'') = \text{Card}(L') + 1 = r + 1 = \text{Card}(L)$ et

$$\text{Vec}(E'') = \text{Vec}(E') = \text{Vec}(E) = V.$$

□

4.3.4 Sommes directes et sous-espaces supplémentaires

Théorème 4.3.10 (Formule des dimensions). *Soient W_1 et W_2 deux sous-espaces vectoriels de dimension finie d’un K -espace vectoriel V . Alors $(W_1 + W_2)$ est de dimension finie et on a*

$$\dim(W_1 + W_2) = \dim W_1 + \dim W_2 - \dim(W_1 \cap W_2).$$

Rappelons que $(W_1 + W_2)$ est le sous-espace vectoriel de V défini par

$$(W_1 + W_2) = \{v = w_1 + w_2 \mid w_1 \in W_1, \text{ et } w_2 \in W_2\}.$$

Preuve. Soit $A = \{a_1, \dots, a_r\}$ une base de $W_1 \cap W_2$. Par le théorème de complétion de base, on peut trouver $\{b_1, \dots, b_s\} \subset W_1$ tels que $B_1 = \{a_1, \dots, a_r, b_1, \dots, b_s\}$ est une base de W_1 . De même, on peut trouver $\{c_1, \dots, c_t\} \subset W_2$ tels que $B_2 = \{a_1, \dots, a_r, c_1, \dots, c_t\}$ est une base de W_2 . Nous affirmons que

$$B = \{a_1, \dots, a_r, b_1, \dots, b_s, c_1, \dots, c_t\}$$

est une base de $W_1 + W_2$.

En effet, B engendre $W_1 + W_2$ car tout vecteurs de ce sous-espace s'écrit $x = x_1 + x_2$ avec $x_1 \in W_1$ et $x_2 \in W_2$. Donc $x_1 \in \text{Vec}(B_1) \subset \text{Vec}(B)$ et $x_2 \in \text{Vec}(B_2) \subset \text{Vec}(B)$, par conséquent $x = x_1 + x_2 \in \text{Vec}(B_1 \cup B_2) = \text{Vec}(B)$.

Prouvons que B est libre. Considérons pour cela une combinaison linéaire nulle d'éléments de B :

$$\lambda_1 a_1 + \dots \lambda_r a_r + \mu_1 b_1 + \dots + \mu_s b_s + \nu_1 c_1 + \dots + \nu_t c_t = 0. \quad (4.2)$$

Nous devons prouver que chacun des scalaires λ_j, μ_j, ν_k est nul. Posons

$$y = \lambda_1 a_1 + \dots \lambda_r a_r + \mu_1 b_1 + \dots + \mu_s b_s = -(\nu_1 c_1 + \dots + \nu_t c_t).$$

Alors $y \in \text{Vec}(B_1) \cap \text{Vec}(B_2) = W_1 \cap W_2 = \text{Vec}(\{a_1, \dots, a_r\})$. En particulier, on peut trouver des scalaires α_i tels que

$$y = -(\nu_1 c_1 + \dots + \nu_t c_t) = \alpha_1 a_1 + \dots \alpha_r a_r,$$

et donc

$$\alpha_1 a_1 + \dots \alpha_r a_r + \nu_1 c_1 + \dots + \nu_t c_t = 0.$$

Puisque les vecteurs de B_2 sont linéairement indépendants, on en conclut que $\alpha_i = \nu_k = 0$ pour tout i et tout k . En particulier $y = 0$:

$$y = \lambda_1 a_1 + \dots \lambda_r a_r + \mu_1 b_1 + \dots + \mu_s b_s = 0.$$

Puisque les vecteurs de B_1 sont linéairement indépendants, on en conclut que $\lambda_i = \mu_j = 0$ pour tout i et tout j . On a montré que tous les coefficients de la combinaison linéaire (4.2) sont nuls ; la famille B est donc libre. On conclut que

$$\begin{aligned} \dim(W_1 + W_2) &= \text{Card}(B) = r + s + t = (r + s) + (r + t) - r \\ &= \text{Card}(B_1) + \text{Card}(B_2) - \text{Card}(A) \\ &= \dim W_1 + \dim W_2 - \dim(W_1 \cap W_2). \end{aligned}$$

□

Définition 4.3.11. Soient W_1 et W_2 deux sous-espaces vectoriels de V . On dit que le sous-espace vectoriel W de V est la *somme directe* de W_1 et W_2 si

$$W = W_1 + W_2 \quad \text{et} \quad W_1 \cap W_2 = \{0\}.$$

Dans ce cas on écrit $W = W_1 \oplus W_2$.

Définition 4.3.12. On dit que $U \subset V$ est un *supplémentaire* de W si U est un sous-espace vectoriel de V et $V = W \oplus U$.

Proposition 4.3.13. Soient W_1, W_2 deux sous-espaces vectoriels d'un K -espace vectoriel V . Alors les conditions suivantes sont équivalentes :

- a) $V = W_1 \oplus W_2$
- b) $\dim(V) = \dim W_1 + \dim W_2$ et $\dim(W_1 \cap W_2) = 0$.

c) Tout vecteur $w \in V$ non nul s'écrit de façon unique comme $w = w_1 + w_2$ avec $w_1 \in W_1$ et $w_2 \in W_2$.

Preuve. (a) $\Rightarrow$ (b) : Si $V = W_1 \oplus W_2$, alors par la formule des dimensions, on a

$$\dim(V) = \dim(W_1 + W_2) = \dim(W_1) + \dim(W_2) - \dim(W_1 \cap W_2) = \dim(W_1) + \dim(W_2),$$

puisque $W_1 \cap W_2 = \{0\}$ (et donc $\dim(W_1 \cap W_2) = 0$) par hypothèse.

(b) $\Rightarrow$ (c) : Soit $\{a_1, \dots, a_p\}$ une base de W_1 et $\{b_1, \dots, b_q\}$ une base de W_2 . Montrons que la réunion $\{a_1, \dots, a_p, b_1, \dots, b_q\}$ est une famille libre. Supposons que

$$\lambda_1 a_1 + \dots + \lambda_p a_p + \mu_1 b_1 + \dots + \mu_q b_q = 0,$$

alors on a

$$y = \lambda_1 a_1 + \dots + \lambda_p a_p = -(\mu_1 b_1 + \dots + \mu_q b_q) \in W_1 \cap W_2.$$

Or on suppose que $\dim(W_1 \cap W_2) = 0$ (donc $W_1 \cap W_2 = \{0\}$), par conséquent $y = 0$. Puisque les a_i sont linéairement indépendants, on a $\lambda_i = 0$ pour tout i . De même $\mu_j = 0$ pour tout j puisque les b_j sont linéairement indépendants.

Par hypothèse on a $\dim(V) = p + q$, on obtient donc

$$\begin{aligned} \dim(W_1 + W_2) &= \dim(\text{Vec}(\{a_1, \dots, a_p, b_1, \dots, b_q\})) = p + q \\ &= \dim(W_1) + \dim(W_2) \\ &= \dim(V) \quad (\text{par hypothèse}). \end{aligned}$$

Or si $W_1 + W_2$ est un sous-espace vectoriel de V et que sa dimension est égale à celle de V , alors $W_1 + W_2 = V$. On a montré que $\{a_1, \dots, a_p, b_1, \dots, b_q\}$ est une base de V . Par conséquent tout $v \in V$ s'écrit de manière $v = \lambda_1 a_1 + \dots + \lambda_p a_p + \mu_1 b_1 + \dots + \mu_q b_q = w_1 + w_2$ avec

$$w_1 = \lambda_1 a_1 + \dots + \lambda_p a_p \in W_1 \quad \text{et} \quad w_2 = \mu_1 b_1 + \dots + \mu_q b_q \in W_2.$$

Cette écriture est unique, car si

$$v = w_1 + w_2 = w'_1 + w'_2$$

avec $w_1, w'_1 \in W_1$ et $w_2, w'_2 \in W_2$, alors

$$w_1 - w'_1 = w_2 - w'_2 \in W_1 \cap W_2 = \{0\}.$$

Donc $w_1 = w'_1$ et $w_2 = w'_2$.

(c) $\Rightarrow$ (a) : L'hypothèse entraîne que $V = W_1 + W_2$. Il faut seulement montrer que $W_1 \cap W_2 = \{0\}$.

Soit $w \in W_1 \cap W_2$, alors ce vecteur peut s'écrire $w = w_1 + 0$ (avec $w_1 \in W_1$) et aussi $w = 0 + w_2$ (avec $w_2 \in W_2$). La condition (c) entraîne alors que $w_1 = 0 = w_2$. Donc $w = 0$, par conséquent $W_1 \cap W_2 = \{0\}$. $\square$

Corollaire 4.3.14. *Tout sous-espace vectoriel d'un K -espace vectoriel de dimension finie admet un supplémentaire (ce supplémentaire n'est pas unique).*

Preuve. Soit V un K -espace vectoriel de dimension finie et soit W un sous-espace vectoriel de V . Choisissons une base $B = \{b_1, b_2, \dots, b_r\}$ de W . On sait qu'on peut compléter cette base en une base de V , notons-la

$$B' = \{b_1, \dots, b_r, c_1, \dots, c_s\}.$$

Alors $U = \text{Vec}(\{c_1, \dots, c_s\})$ vérifie $V = W \oplus U$ car $\dim V = \dim W + \dim U$ (appliquer le critère (c) de la proposition précédente). $\square$

4.3.5 Résumé des notions fondamentales sur l'indépendance linéaire

Voici un résumé de quelques propriétés importantes sur les notions vues dans ce chapitre. On considère un K -espace vectoriel V , alors :

- (a) Tout sous-ensemble d'une famille libre de V est une famille libre.
- (b) Toute partie de V qui contient une partie liée est elle-même liée.
- (c) Toute partie de V qui contient une partie génératrice est elle-même génératrice.
- (d) L'espace vectoriel V est de dimension infinie si et seulement s'il contient une famille libre infinie.
- (e) Soit $S = \{f_1, \dots, f_n\}$ une famille génératrice de V . Si $p > n$, alors toute famille $\{v_1, \dots, v_p\} \subset V$ contenant p vecteurs est liée.
De façon équivalente, si $L = \{w_1, \dots, w_q\}$ est une partie libre dans V , alors $q \leq n$
- (f) Tout espace vectoriel de dimension finie possède une base. Plus généralement, toute famille libre peut être complétée en une base. De plus, toute famille génératrice contient une base. (Ce résultat est aussi vrai pour les espaces vectoriels de dimension infinie, mais la preuve est assez élaborée et utilise l'axiome du choix).
- (g) L'espace vectoriel V est de dimension infinie si et seulement si pour tout $n \in \mathbb{N}$, il existe une partie libre L de V de cardinal n .
- (h) L'espace vectoriel V est de dimension finie si et seulement s'il existe $m \in \mathbb{N}$ tel que toute partie de V de cardinal m est liée.

Dans les propriétés qui suivent, on suppose que $\dim(V) = n < \infty$.

- (i) Toute famille contenant plus de n vecteurs est liée.
- (j) Aucune famille ayant moins de n vecteurs n'est génératrice.
- (k) Toute famille libre ayant n vecteurs est une base.
- (l) Toute famille génératrice ayant n vecteurs est une base.
- (m) Si $W \subset V$ est un sous-espace vectoriel, alors $\dim(W) \leq \dim(V)$ et on a $\dim(W) = \dim(V)$ si et seulement si $W = V$
- (n) Soit V un K -espace vectoriel de dimension finie, et soit W un sous-espace vectoriel de V . Alors
 - W est de dimension finie.
 - $\dim W \leq \dim V$.
 - Toute base de W peut être complétée en une base de V .
 - Si $\dim W = \dim V$, alors $W = V$.
- (o) (Formule des dimensions) Si W_1 et W_2 sont deux sous-espaces vectoriels de dimension finies de V alors

$$\dim(W_1 + W_2) = \dim W_1 + \dim W_2 - \dim(W_1 \cap W_2).$$

Chapitre 5

Applications linéaires et matrices

5.1 Rappels sur les applications linéaires

Soient V, W deux espaces vectoriels sur un corps K . On note $\mathcal{L}(V, W)$ ou $\text{Hom}(V, W)$ l'ensemble des applications K -linéaires de V vers W . Si on veut préciser le corps, on note $\mathcal{L}_K(V, W)$ ou $\text{Hom}_K(V, W)$.

Rappelons qu'une application $f : V \rightarrow W$ est une application K -linéaire si $f(\lambda u + v) = \lambda f(u) + f(v)$, pour tout $\lambda \in K$ et $u, v \in V$.

On rappelle les propriétés suivantes :

- L'image de l'élément neutre de V est l'élément neutre de W , i.e. $f(0_V) = 0_W$.
- f est compatible avec les combinaisons linéaires :

$$f(\lambda_1 v_1 + \cdots + \lambda_m v_m) = \lambda_1 f(v_1) + \cdots + \lambda_m f(v_m),$$

pour tout $\lambda_i \in K$ et $v_i \in V$.

- Le noyau de f est l'ensemble $\text{Ker}(f) = f^{-1}(\{0_W\}) \subset V$, et c'est un sous-espace vectoriel de V (auss appelé $\text{Null}(f)$ en anglais, pour "null space").
- L'image de f est l'ensemble $\text{Im}(f) = f(V) \subset W$, et c'est un sous-espace vectoriel de W (auss appelé $\text{Range}(f)$ en anglais, pour "range").
- f est injective si et seulement si $\text{Ker}(f) = \{0\}$.
- f est surjective si et seulement si $\text{Im}(f) = W$.

5.2 Opérations sur les applications linéaires

Soient V et W deux espaces vectoriels sur un corps K . On définit deux opérations

$$\mathcal{L}(V, W) \times \mathcal{L}(V, W) \xrightarrow{+} \mathcal{L}(V, W), \quad K \times \mathcal{L}(V, W) \xrightarrow{\cdot} \mathcal{L}(V, W)$$

de la façon suivante :

1. La somme de $f, g \in \mathcal{L}(V, W)$ est l'application $f + g : V \rightarrow W$ définie par $(f + g)(v) = f(v) + g(v)$ pour tout $v \in V$.
2. Pour $\lambda \in K$, l'application $\lambda \cdot f : V \rightarrow W$ est définie par $(\lambda \cdot f)(v) = \lambda \cdot f(v)$ pour tout $v \in V$.

On montre facilement :

Proposition 5.2.1. $\mathcal{L}(V, W)$ est un K -espace vectoriel pour les opérations d'addition et de multiplication scalaire définies ci-dessus.

On peut aussi composer des applications linéaires :

Proposition 5.2.2. Soient U, V, W trois espaces vectoriels sur le corps K , et $f : V \rightarrow W$, $g : W \rightarrow U$ deux applications K -linéaires. Alors l'application $g \circ f : V \rightarrow U$ est K -linéaire.

Preuve. On prend $\lambda \in K$ et $x, y \in V$. On a $(g \circ f)(\lambda x + y) = g(f(\lambda x + y)) = g(\lambda f(x) + f(y))$, car f est K -linéaire. Et ensuite, $g(\lambda f(x) + f(y)) = \lambda g(f(x)) + g(f(y))$ car g est K -linéaire. Mais ce dernier vecteur est précisément $\lambda(g \circ f)(x) + (g \circ f)(y)$. □

Théorème 5.2.3. Toute application linéaire définie sur un espace vectoriel de dimension finie est déterminée par l'image d'une base. Plus précisément : soient $\{b_1, b_2, \dots, b_n\}$ une base de V et $w_1, w_2, \dots, w_n$ des vecteurs (quelconques) de W . Alors il existe une unique application linéaire $f : V \rightarrow W$ telle que $f(b_i) = w_i$ pour tout $i = 1, \dots, n$.

Preuve. Tout vecteur $x \in V$ s'écrit de manière unique $x = \sum_{i=1}^n x_i b_i$ avec $x_i \in K$ (c'est l'une des propriétés des bases). On peut alors définir une application $f : V \rightarrow W$ par

$$f(x) = f\left(\sum_{i=1}^n x_i b_i\right) = \sum_{i=1}^n x_i w_i.$$

On vérifie que f est linéaire et il est clair que $f(b_i) = w_i$ pour tout $i = 1, \dots, n$. L'unicité de l'application f vient des propriétés de la linéarité. □

Définition : On dit que f a été obtenue en étendant les conditions $f(b_i) = w_i$ par linéarité.

Proposition 5.2.4. Toute application linéaire $g : K^n \rightarrow K^m$ s'écrit

$$g(x_1, \dots, x_n) = \left(\sum_{j=1}^n a_{1j} x_j, \dots, \sum_{j=1}^n a_{mj} x_j \right),$$

pour $(x_1, \dots, x_n) \in K^n$.

Preuve. Vérifions d'abord que g est bien K -linéaire. Soient $(x_1, \dots, x_n)$ et $(y_1, \dots, y_n)$ deux vecteurs de K^n , alors on a

$$\begin{aligned} g((x_1, \dots, x_n) + (y_1, \dots, y_n)) &= g(x_1 + y_1, \dots, x_n + y_n) \\ &= \left(\sum_{j=1}^n a_{1j} \cdot (x_j + y_j), \dots, \sum_{j=1}^n a_{mj} \cdot (x_j + y_j) \right) \\ &= \left(\sum_{j=1}^n a_{1j} x_j + \sum_{j=1}^n a_{1j} y_j, \dots, \sum_{j=1}^n a_{mj} x_j + \sum_{j=1}^n a_{mj} y_j \right) \\ &= \left(\sum_{j=1}^n a_{1j} x_j, \dots, \sum_{j=1}^n a_{mj} x_j \right) + \left(\sum_{j=1}^n a_{1j} y_j, \dots, \sum_{j=1}^n a_{mj} y_j \right) \\ &= g(x_1, \dots, x_n) + g(y_1, \dots, y_n), \end{aligned}$$

De même, si $\lambda \in K$, alors

$$\begin{aligned}
g(\lambda(x_1, \dots, x_n)) &= g(\lambda x_1, \dots, \lambda x_n) \\
&= \left(\sum_{j=1}^n a_{1j} \cdot (\lambda x_j), \dots, \sum_{j=1}^n a_{mj} \cdot (\lambda x_j) \right) \\
&= \left(\lambda \sum_{j=1}^n a_{1j} x_j, \dots, \lambda \sum_{j=1}^n a_{mj} x_j \right) \\
&= \lambda \left(\sum_{j=1}^n a_{1j} x_j, \dots, \sum_{j=1}^n a_{mj} x_j \right) \\
&= \lambda g(x_1, \dots, x_n),
\end{aligned}$$

ce qui prouve la linéarité de g . Pour montrer que toute application linéaire de K^n vers K^m est de ce type, on considère l'application $\pi_i : K^m \rightarrow K$ de projection sur la $i^{\text{ème}}$ coordonnée :

$$\pi_i(y_1, \dots, y_m) = y_i.$$

Cette application est clairement linéaire, et donc la composition $g_i := \pi_i \circ g : K^n \rightarrow K$ est aussi linéaire. Mais on a déjà prouvé que toute application linéaire de K^n à valeurs dans K s'écrit $f(x) = \sum_{j=1}^n \alpha_j x_j$ avec $\alpha_j \in K$. Par conséquent g_i est de la forme

$$g_i(x_1, \dots, x_n) = \sum_{j=1}^n a_{ij} x_j = a_{i1} x_1 + \dots + a_{in} x_n.$$

□

5.3 Le théorème du rang et ses conséquences

Théorème 5.3.1 (Théorème du rang). *Soient V et W deux espaces vectoriels sur un corps K et $f : V \rightarrow W$ une application K -linéaire. Si $\dim(V) < \infty$, alors*

$$\dim(V) = \dim(\text{Ker}(f)) + \dim(\text{Im}(f)).$$

Preuve. Supposons que $\dim(V) = n$. Notons $r = \dim(\text{Im}(f))$ et $s = \dim(\text{Ker}(f))$.

Choisissons une base $\{\mathbf{b}_1, \mathbf{b}_2, \dots, \mathbf{b}_s\}$ de $\text{Ker}(f)$ et complétons cette liste de vecteurs pour obtenir une base de V :

$$\{\mathbf{b}_1, \mathbf{b}_2, \dots, \mathbf{b}_s, \mathbf{b}_{s+1}, \dots, \mathbf{b}_n\}.$$

Notons $U = \text{Vec}(\{f(\mathbf{b}_{s+1}), \dots, f(\mathbf{b}_n)\}) \subset W$, c'est le sous-espace vectoriel engendré par les $(n - s)$ vecteurs :

$$\{f(\mathbf{b}_{s+1}), \dots, f(\mathbf{b}_n)\}.$$

Nous affirmons que ces $(n - s)$ vecteurs sont linéairement indépendants et que $U = \text{Im}(f)$.

En effet, il est clair que $U \subset \text{Im}(f)$. Inversement, si $y \in \text{Im}(f)$, alors il existe $x \in V$ tel que $y = f(x)$. On peut développer x dans la base $\{\mathbf{b}_i\}$:

$$x = \sum_{i=1}^n x_i \mathbf{b}_i = \sum_{i=1}^s x_i \mathbf{b}_i + \sum_{i=s+1}^n x_i \mathbf{b}_i.$$

Or $\mathbf{b}_i \in \text{Ker}(f)$ pour $1 \leq i \leq s$, donc

$$y = f(x) = f\left(\sum_{i=s+1}^n x_i \mathbf{b}_i\right) = \sum_{i=s+1}^n x_i f(\mathbf{b}_i) \in U.$$

Par conséquent $\text{Im}(f) \subset U$.

Vérifions maintenant que $f(\mathbf{b}_{s+1}), \dots, f(\mathbf{b}_n)$ sont linéairement indépendants. Supposons que

$$\sum_{i=s+1}^n \lambda_i f(\mathbf{b}_i) = 0.$$

Alors $\sum_{i=s+1}^n \lambda_i \mathbf{b}_i \in \text{Ker}(f)$, donc on peut écrire

$$\sum_{i=s+1}^n \lambda_i \mathbf{b}_i = \sum_{j=1}^s \mu_j \mathbf{b}_j.$$

Mais ceci n'est possible que si tous les λ_i et tous les μ_j sont nuls car les vecteurs $\mathbf{b}_i$ sont linéairement indépendants.

On a donc

$$\dim(\text{Im}(f)) = \dim(U) = n - s = \dim(V) - \dim(\text{Ker}(f)).$$

□

Définition. Le *rang* de l'application linéaire f est la dimension de $\text{Im}(f)$. On réécrit parfois la formule ci-dessus sous la forme

$$\text{rang}(f) = \dim(V) - \dim(\text{Ker}(f)).$$

Une première conséquence du théorème du rang est :

Corollaire 5.3.2. Soient V et W deux K -espaces vectoriels de dimensions finies et $f \in \mathcal{L}(V, W)$.

- a) Si f est injective, alors $\dim(V) \leq \dim(W)$,
- b) Si f est surjective, alors $\dim(V) \geq \dim(W)$,

Preuve (a) Supposons f injective, alors $\text{Ker}(f) = \{0\}$ et donc

$$\dim(V) = \dim(\text{Im}(f)) + \dim(\text{Ker}(f)) = \dim(\text{Im}(f)) \leq \dim(W).$$

(b) Si f est surjective, alors

$$\dim(V) = \dim(\text{Im}(f)) + \dim(\text{Ker}(f)) \geq \dim(\text{Im}(f)) = \dim(W).$$

□

Corollaire 5.3.3. Supposons $\dim(V) = \dim(W) < \infty$ et $f \in \mathcal{L}(V, W)$ alors les conditions suivantes sont équivalentes :

- a) f est surjective ;
 - b) f est injective ;
 - c) f est bijective ;
 - d) L'image par f d'une base de V est une base de W .
 - e) Il existe une application linéaire $h \in \mathcal{L}(W, V)$ telle que $h \circ f = I_V$.
 - f) Il existe une application linéaire $g \in \mathcal{L}(W, V)$ telle que $f \circ g = I_W$.
- De plus $g = h$, i.e. l'inverse à gauche est égale à l'inverse à droite, et on note $f^{-1} = g = h$.

Preuve. La preuve suit la schéma suivant : $(a) \Rightarrow (b) \Rightarrow (c) \Rightarrow (a)$, puis $(c) \Rightarrow (d) \Rightarrow (e) \Rightarrow (b)$ et $(d) \Rightarrow (f) \Rightarrow (a)$.

Montrons que (a) implique (b). Puisque f est surjective, on a par la formule du rang

$$\dim(\text{Ker}(f)) = \dim(V) - \dim(\text{Im}(f)) = \dim(V) - \dim(V) = 0,$$

donc f est injective.

Supposons maintenant que f est injective, alors par la formule du rang

$$\dim(\text{Im}(f)) = \dim(V) - \dim(\text{Ker}(f)) = \dim(V),$$

donc $\text{Im}(f) = V$. Cela montre que (b) implique (c). Il est clair que (c) implique (a) et (b), donc (a), (b) et (c) sont équivalentes.

Montrons que (c) implique (d). Soit $\{b_1, \dots, b_n\}$ une base de V . Alors les vecteurs $\{f(b_1), \dots, f(b_n)\}$ sont linéairement indépendants (sinon f ne serait pas injective) et ils engendrent W (sinon f ne serait pas surjective). Donc $\{f(b_1), \dots, f(b_n)\}$ est une base de W .

Voyons que (d) implique (e). Soit $\{b_1, \dots, b_n\}$ une base de V et notons $w_i = f(b_i)$. Notre hypothèse est que $\{w_1, \dots, w_n\}$ est une base de W . On définit $h : W \rightarrow V$ comme l'unique application linéaire telle que $h(w_i) = b_i$ (donc on étend h par linéarité), i.e.

$$h\left(\sum_{i=1}^n \lambda_i w_i\right) = \sum_{i=1}^n \lambda_i b_i,$$

et on vérifie que $h \circ f = \text{I}_V$. En effet, on a pour tout $x = \sum_{i=1}^n \mu_i b_i \in V$:

$$h \circ f(x) = h \circ f\left(\sum_{i=1}^n \mu_i b_i\right) = h\left(\sum_{i=1}^n \mu_i f(b_i)\right) = h\left(\sum_{i=1}^n \mu_i w_i\right) = \sum_{i=1}^n \mu_i b_i = x.$$

Le même argument (avec la même application $h : W \rightarrow V$) prouve (d) $\Rightarrow$ (f) car pour tout $y = \sum_{i=1}^n \lambda_i w_i \in W$ on a car

$$f \circ h(y) = f \circ h\left(\sum_{i=1}^n \lambda_i w_i\right) = f\left(\sum_{i=1}^n \lambda_i b_i\right) = \sum_{i=1}^n \lambda_i w_i = y.$$

Donc $f \circ h = \text{I}_W$.

Finalement (e) $\Rightarrow$ (b) et (f) $\Rightarrow$ (a) sont évidentes par les définitions. □

Remarque. Il est clair que si une application linéaire f est bijective, alors f admet un inverse à gauche et à droite. Le point à vérifier dans cette preuve est que cet inverse est aussi une application linéaire.

Définition 5.3.4. On dit que $f \in \mathcal{L}(V, W)$ est un *isomorphisme* si l'une des conditions précédentes est vérifiée ; un isomorphisme est donc un homomorphisme bijectif. On dit aussi que l'homomorphisme f est *invertible*.

Deux espaces vectoriels V et W sont *isomorphes* s'il existe un isomorphisme de V vers W .

On note $V \cong W$ pour dire que V et W sont isomorphes ; l'isomorphie est une relation d'équivalence, c'est à dire que les trois propriétés suivantes sont vérifiées (où V, W, U sont trois espaces vectoriels sur un même corps K) :

- i) $V \cong V$;
- ii) $V \cong W \Leftrightarrow W \cong V$
- iii) $V \cong W$ et $W \cong U \Rightarrow V \cong U$

Preuve. (i) est clair car l'identité $I_V : V \rightarrow V$ est un isomorphisme et (ii) est conséquence du fait que si $f : V \rightarrow W$ est un isomorphisme, alors $f^{-1} : W \rightarrow V$ est aussi un isomorphisme par le corollaire précédent. La condition (iii) vient du fait que si $f : V \rightarrow W$ et $g : W \rightarrow U$ sont des isomorphismes, alors $g \circ f : V \rightarrow U$ est un isomorphisme (car la composition de deux applications linéaires est linéaire et la composition de deux bijections est une bijection).

Corollaire 5.3.5. *Deux K -espaces vectoriels de dimensions finies sont isomorphes si et seulement s'ils ont la même dimension.*

Preuve. Supposons que V et W sont isomorphes, alors il existe $f \in \mathcal{L}(V, W)$ bijective et le théorème du rang 5.3.1 implique donc que $\dim(V) = \dim(W)$. Réciproquement, supposons que $\dim(V) = \dim(W) = n$, alors il existe une base $\{v_1, \dots, v_n\}$ de V et une base $\{w_1, \dots, w_n\}$ de W (toutes deux avec le même nombre n d'éléments). Il existe alors une unique application linéaire $f \in \mathcal{L}(V, W)$ telle que $f(v_i) = w_i$ pour tout $1 \leq i \leq n$ et par le point (d) du corollaire précédent, on sait que f est un isomorphisme. □

En particulier, tout K -espace vectoriel V de dimension finie n est isomorphe à K^n . De manière très concrète, à toute base $\{v_1, \dots, v_n\}$ du K -espace vectoriel V , on associe l'unique isomorphisme $f : K^n \rightarrow V$ tel que $f(e_i) = v_i$ (où $\{e_1, \dots, e_n\}$ est la base canonique de K^n), ce qui donne :

$$f(x_1, \dots, x_n) = x_1 v_1 + \dots + x_n v_n.$$

5.4 La matrice d'une application linéaire

Nous avons vu précédemment que si V et W sont deux espaces vectoriels sur un corps K , alors l'ensemble $\mathcal{L}(V, W)$ des applications linéaires de V vers W est un K -espace vectoriel pour les opérations naturelles de somme et de multiplication par des scalaires.

Théorème 5.4.1. *Si V et W sont de dimensions finies, alors*

$$\dim(\mathcal{L}(V, W)) = \dim V \cdot \dim W$$

Preuve. Supposons que $\dim(V) = n$ et $\dim(W) = m$ et fixons $B = \{v_1, \dots, v_n\}$ une base de V et $B' = \{w_1, \dots, w_m\}$ une base de W . On sait qu'une application linéaire $f : V \rightarrow W$ est entièrement déterminée par son effet sur une base donnée de V . En particulier, il existe une unique application linéaire de V vers W qui envoie v_j sur w_i et les autres v_k sur 0. Notons cette application $\mathcal{E}_{ij} \in \mathcal{L}(V, W)$. On a ainsi

$$\mathcal{E}_{ij}(v_j) = w_i, \quad \mathcal{E}_{ij}(v_k) = 0 \quad \text{si } k \neq j,$$

et donc

$$\mathcal{E}_{ij} \left(\sum_{k=1}^n x_k v_k \right) = x_j w_i.$$

Considérons maintenant une application linéaire quelconque $f \in \mathcal{L}(V, W)$, et notons $a_{1j}, a_{2j}, \dots, a_{mj}$ les composantes de $f(v_j) \in W$ dans la base $B' \subset W$, i.e.

$$f(v_j) = a_{1j}w_1 + a_{2j}w_2 + \dots + a_{mj}w_m.$$

Nous affirmons que

$$f = \sum_{i=1}^m \sum_{j=1}^n a_{ij} \mathcal{E}_{ij}. \quad (5.1)$$

En effet, si $x = \sum_{k=1}^n x_k v_k \in V$ est un vecteur quelconque de V , alors¹

$$\begin{aligned} \sum_{i=1}^m \sum_{j=1}^n a_{ij} \mathcal{E}_{ij}(x) &= \sum_{i=1}^m \sum_{j=1}^n a_{ij} \mathcal{E}_{ij} \left(\sum_{k=1}^n x_k v_k \right) = \sum_{i=1}^m \sum_{j=1}^n a_{ij} x_j w_i = \sum_{j=1}^n x_j \sum_{i=1}^m a_{ij} w_i \\ &= \sum_{j=1}^n x_j f(v_j) = f \left(\sum_{j=1}^n x_j v_j \right) = f(x). \end{aligned}$$

L'équation (5.1) signifie que tout élément de $\mathcal{L}(V, W)$ s'écrit de façon unique comme combinaison linéaire des applications linéaires $\mathcal{E}_{ij} \in \mathcal{L}(V, W)$. Il s'ensuit que $\{\mathcal{E}_{ij}\}_{i,j}$ est une base de $\mathcal{L}(V, W)$ et donc

$$\dim(\mathcal{L}(V, W)) = \text{Card}\{\mathcal{E}_{ij}\}_{i,j} = m \cdot n.$$

□

Définition. Les coefficients a_{ij} dans la preuve ci-dessus représentent donc les composantes de l'application $f \in \mathcal{L}(V, W)$ dans la base $\{\mathcal{E}_{ij}\}_{i,j}$. On les dispose sous forme d'un tableau rectangulaire à m lignes et n colonnes contenant le coefficient a_{ij} à l'intersection de la $i^{\text{ème}}$ ligne et $j^{\text{ème}}$ colonne. Ce tableau s'appelle la *matrice de f dans les bases B, B'* . On note donc

$$M_{B',B}(f) = \begin{pmatrix} a_{11} & \dots & a_{1n} \\ \vdots & & \vdots \\ a_{m1} & \dots & a_{mn} \end{pmatrix}$$

Remarque 5.4.2. (a) Il est très important de bien noter que la $j^{\text{ème}}$ colonne de la matrice est formée des composantes dans la base $B' \subset W$ de l'image $f(v_j)$ du $j^{\text{ème}}$ vecteur de la base $B \subset V$. Donc l'écriture

$$f(v_j) = \sum_{i=1}^m a_{ij} w_i = a_{1j} w_1 + a_{2j} w_2 + \dots + a_{mj} w_m.$$

nous donne les coefficients de la $j^{\text{ème}}$ colonne de la matrice $M_{B',B}(f)$ (bien que les coefficients de $f(v_j)$ dans la base B' apparaissent en ligne, ils forment la $j^{\text{ème}}$ colonne de la matrice).

(b) La matrice de f dans les bases B, B' se note parfois $M_{B'}^B(f)$ ou simplement $[f]_{B'}^B$.

Remarque importante.

Pour associer une matrice à une application linéaire $f : V \rightarrow W$, il est important non seulement de choisir des bases de V et de W , mais aussi de fixer une numérotation des éléments de ces bases, sinon les coefficients de la matrice peuvent être échangés.

Dans la suite de ce chapitre, toutes les bases sont supposées numérotées et nous considérons que renuméroter les vecteurs revient à changer de base. Par exemple $\{(1, 0), (1, 1)\}$ et $\{(1, 1), (1, 0)\}$ sont vues comme deux bases différentes de K^2 . Pour cette raison certains auteurs notent les bases comme des *listes* de vecteurs et non des *ensembles*, i.e. on note $((1, 0), (1, 1))$ au lieu de $\{(1, 0), (1, 1)\}$.

1. Observer que dans cette preuve l'indice i varie de 1 à $m = \dim(W)$ et j varie de 1 à $n = \dim(V)$.

5.4.1 Exemples de matrices d'applications linéaires

Soient V et W deux espaces vectoriels sur un corps K de dimensions n et m respectivement. On se donne une base $\{v_1, \dots, v_n\}$ de V et $\{w_1, \dots, w_m\}$ une base de W . Alors toute application linéaire $f : V \rightarrow W$ est représentée par une matrice $M(f)$ à m lignes et n colonnes (cette matrice dépend des bases choisies). Voyons quelques exemples.

- 1) Soit $f : K^2 \rightarrow K^3$ l'application linéaire

$$f(x, y) = (2x + y, x + 3y, -x),$$

et notons $\{e_1, e_2\}$ la base canonique de K^2 et $\{e'_1, e'_2, e'_3\}$ la base canonique de K^3 . Alors

$$f(e_1) = f(1, 0) = (2, 1, -1) = 2e'_1 + e'_2 - e'_3$$

et

$$f(e_2) = f(0, 1) = (1, 3, 0) = e'_1 + 3e'_2 + 0e'_3.$$

La matrice de f dans ces bases est donc

$$M(f) = \begin{pmatrix} 2 & 1 \\ 1 & 3 \\ -1 & 0 \end{pmatrix}$$

- 2) Si $f : K^n \rightarrow K$ est l'application linéaire telle que $f(x_1, \dots, x_n) = a_1x_1 + \dots + a_nx_n$, alors sa matrice dans les bases canoniques est la matrice à une ligne et n colonnes

$$M(f) = (a_1 \ a_2 \ \dots \ a_n).$$

- 3) Si $g : K \rightarrow K^m$ est l'application linéaire définie par $g(t) = (ta_1, ta_2, \dots, ta_m)$, alors sa matrice dans les bases canoniques est la matrice à m lignes et une colonne

$$M(g) = \begin{pmatrix} a_1 \\ a_2 \\ \vdots \\ a_m \end{pmatrix}$$

- 4) La matrice de l'application nulle (i.e. telle que $x \mapsto 0_W$ pour tout $x \in W$) est la matrice dont tous les coefficients sont nuls. On l'appelle la *matrice nulle* et on la note $\mathbf{0}$ ou $0_{m \times n}$:

$$0_{m \times n} = \begin{pmatrix} 0 & \dots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \dots & 0 \end{pmatrix}$$

- 5) La matrice de l'application linéaire $\mathcal{E}_{ij}$ définie par $\mathcal{E}_{ij}(v_j) = w_i$ et $\mathcal{E}_{ij}(v_k) = 0$ si $k \neq j$ est la matrice qui vaut 1 en position i, j ($i^{\text{ème}}$ ligne et $j^{\text{ème}}$ colonne) et 0 dans les autres positions. On note cette matrice E_{ij} . Par exemple

$$E_{23} = M(\mathcal{E}_{23}) = \begin{pmatrix} 0 & 0 & 0 & \dots & 0 \\ 0 & 0 & 1 & \dots & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & 0 & \dots & 0 \end{pmatrix}$$

- 6) Si $W = V$ et $w_i = v_i$ pour tout i (i.e. on a deux fois le même espace vectoriel et les mêmes bases), alors les coefficients δ_{ij} de la matrice de l'application identité $I_V : V \rightarrow V$ sont donnés² par

$$\delta_{ij} = \begin{cases} 1 & \text{si } i = j \\ 0 & \text{si } i \neq j \end{cases}$$

On note cette matrice $\mathbf{I}_n$, donc

$$M(I_V) = \mathbf{I}_n = \begin{pmatrix} 1 & 0 & \cdots & 0 \\ 0 & 1 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & 1 \end{pmatrix}$$

- 7) Si $f, g \in \mathcal{L}(V, W)$ ont pour matrices $M(f) = (a_{ij})$ et $M(g) = (b_{ij})$, et si $\lambda \in K$, alors

$$M(f + g) = (a_{ij} + b_{ij}) \quad \text{et} \quad M(\lambda f) = (\lambda a_{ij}).$$

5.5 L'espace vectoriel des matrices

Définition Une *matrice* est un tableau rectangulaire contenant des scalaires :

$$A = (a_{ij}) = \begin{pmatrix} a_{11} & \cdots & a_{1n} \\ \vdots & \ddots & \vdots \\ a_{m1} & \cdots & a_{mn} \end{pmatrix}$$

Les matrices peuvent représenter toutes sortes d'informations. Si la matrice a m lignes et n colonnes, on dit que c'est une matrice de *taille* (ou de *type*) $m \times n$. Les nombres a_{ij} sont les *coefficients* de la matrice. Ils ont un double indice : i désigne la ligne où se trouve ledit coefficient et j indique la colonne. On dit que la matrice est *carrée* si $n = m$, i.e. si elle a autant de lignes que de colonnes, sinon elle est *rectangulaire*. Une matrice de taille $1 \times n$ s'appelle un *vecteur-ligne* et une matrice de taille $m \times 1$ s'appelle un *vecteur-colonne*.

Par exemple la matrice

$$A = \begin{pmatrix} 2 & 3 & 1 & 7 \\ 0 & 1 & 4 & 1 \\ 2 & 0 & 0 & 0 \end{pmatrix}$$

est une matrice rectangulaire de taille 3×4 . Le coefficient de la deuxième ligne et troisième colonne est $a_{23} = 4$ et celui de la première ligne et quatrième colonne est $a_{14} = 7$.

On peut définir plusieurs opérations algébriques sur les matrices :

1. Multiplication d'une matrice par un scalaire

La multiplication d'une matrice A par un scalaire λ consiste simplement à multiplier chaque coefficient de la matrice par ce scalaire :

$$\lambda \cdot A = \begin{pmatrix} \lambda a_{11} & \cdots & \lambda a_{1n} \\ \vdots & \ddots & \vdots \\ \lambda a_{m1} & \cdots & \lambda a_{mn} \end{pmatrix}$$

2. on appelle δ_{ij} le *symbole de Kronecker*.

Par exemple

$$3 \cdot \begin{pmatrix} 2 & 3 & 1 & 7 \\ 0 & 1 & 4 & 1 \\ 2 & 0 & 0 & 0 \end{pmatrix} = \begin{pmatrix} 6 & 9 & 3 & 21 \\ 0 & 3 & 12 & 3 \\ 6 & 0 & 0 & 0 \end{pmatrix}$$

2. Somme de deux matrices

On peut additionner deux matrices de même taille. Cela se fait simplement en additionnant les coefficients de même position :

$$(a_{ij}) + (b_{ij}) = (a_{ij} + b_{ij})$$

Par exemple

$$\begin{pmatrix} 2 & 3 & 1 & 7 \\ 0 & 1 & 4 & 1 \\ 2 & 0 & 0 & 0 \end{pmatrix} + \begin{pmatrix} -1 & 0 & 2 & 1 \\ 3 & 2 & 1 & 1 \\ -1 & 0 & 1 & 2 \end{pmatrix} = \begin{pmatrix} 1 & 3 & 3 & 8 \\ 3 & 3 & 5 & 2 \\ 1 & 0 & 1 & 2 \end{pmatrix}.$$

La *différence* de deux matrices se définit en faisant la différence des coefficients de même position :

$$(a_{ij}) - (b_{ij}) = (a_{ij} - b_{ij}).$$

Par exemple

$$\begin{pmatrix} 2 & 3 & 1 & 7 \\ 0 & 1 & 4 & 1 \\ 2 & 0 & 0 & 0 \end{pmatrix} - \begin{pmatrix} -1 & 0 & 2 & 1 \\ 3 & 2 & 1 & 1 \\ -1 & 0 & 1 & 2 \end{pmatrix} = \begin{pmatrix} 3 & 3 & -1 & 6 \\ -3 & -1 & 3 & 0 \\ 3 & 0 & -1 & -2 \end{pmatrix}.$$

La *matrice nulle* de taille $m \times n$ est la matrice $O = O_{m \times n}$ dont tous les coefficients sont nuls, on la note

$$O_{m \times n} = \begin{pmatrix} 0 & \cdots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \cdots & 0 \end{pmatrix},$$

par exemple

$$O_{2 \times 4} = \begin{pmatrix} 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{pmatrix}, \quad O_{3 \times 1} = \begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix}, \quad O_{2 \times 2} = \begin{pmatrix} 0 & 0 \\ 0 & 0 \end{pmatrix}.$$

Si A est une $m \times n$ matrice quelconque, alors

$$A + O_{m \times n} = O_{m \times n} + A = A.$$

La *matrice opposée* d'une matrice A est la matrice obtenue en changeant le signe de chaque coefficient, elle se note $-A$:

$$-\begin{pmatrix} a_{11} & \cdots & a_{1n} \\ \vdots & \ddots & \vdots \\ a_{m1} & \cdots & a_{mn} \end{pmatrix} = \begin{pmatrix} -a_{11} & \cdots & -a_{1n} \\ \vdots & \ddots & \vdots \\ -a_{m1} & \cdots & -a_{mn} \end{pmatrix}$$

Théorème 5.5.1. *L'ensemble $M_{m \times n}(K)$ des matrices de taille $m \times n$ à coefficients dans le corps K , muni des opérations d'addition des matrices et de multiplication par les scalaires définies ci-dessus est un K -espace vectoriel.*

Il est clair que les matrices $E_{ij} \in M_{m \times n}(K)$ dont tous les coefficients sont nuls, sauf le coefficient (i, j) , qui vaut 1, forment une base de cet espace vectoriel. En particulier $\dim M_{m \times n}(K) = m \cdot n$.

5.6 Produit matriciel

On peut multiplier une matrice de taille $m \times n$ avec une matrice de taille $n \times p$, on obtient alors une matrice de taille $m \times p$.

La règle est la suivante : si $A = (a_{ij})$ et $B = (b_{jk})$, alors le produit $C = A \cdot B$ est la matrice $C = (c_{ik})$ telle que

$$c_{ik} = \sum_{j=1}^n a_{ij}b_{jk} = a_{i1}b_{1k} + a_{i2}b_{2k} + \cdots + a_{in}b_{nk}$$

Exemple :

$$\begin{pmatrix} 2 & 0 \\ 3 & 1 \\ 0 & -1 \end{pmatrix} \cdot \begin{pmatrix} 4 & 2 & 0 \\ 0 & -1 & 2 \end{pmatrix} = \begin{pmatrix} 8 & 4 & 0 \\ 12 & 5 & 2 \\ 0 & 1 & -2 \end{pmatrix}$$

Définition. On dit que deux matrices A et B sont *multipliables* (en anglais on dit *conformable*) si le nombre de colonnes de la première matrice est égal au nombre de lignes de la seconde matrice. Dans ce cas le produit $A \cdot B$ est bien défini.

Définition. La *matrice identité* de taille $n \times n$ est la matrice $\mathbf{I}_n = (\delta_{ij})$ dont les coefficients valent

$$\delta_{ij} = \begin{cases} 1 & \text{si } i = j, \\ 0 & \text{sinon.} \end{cases}$$

Par exemple

$$\mathbf{I}_2 = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} \quad \mathbf{I}_3 = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix} \quad \mathbf{I}_4 = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix}$$

Proposition 5.6.1. *Le produit matriciel possède les propriétés suivantes :*

- i) $A(B_1 + B_2) = AB_1 + AB_2$, pour tous $A \in M_{m \times n}(K)$, $B_1, B_2 \in M_{n \times p}(K)$
- ii) $(A_1 + A_2)B = A_1B + A_2B$, pour tous $A_1, A_2 \in M_{m \times n}(K)$, $B \in M_{n \times p}(K)$
- iii) $(\lambda A)(\mu B) = (\lambda\mu)(AB)$ pour tous $A \in M_{m \times n}(K)$, $B \in M_{n \times p}(K)$ et $\lambda, \mu \in K$.
- iv) $(AB)C = A(BC)$ pour tous $A \in M_{m \times n}(K)$, $B \in M_{n \times p}(K)$ et $C \in M_{p \times q}(K)$,
- v) La matrice identité est un élément neutre pour la multiplication matricielle : Si $A \in M_{m \times n}(K)$, alors

$$\mathbf{I}_m \cdot A = A \cdot \mathbf{I}_n = A.$$

Noter que le produit matriciel n'est pas commutatif. En général $AB \neq BA$.

Preuve. Nous démontrons la propriété (iv) et laissons les autres au lecteur. Notons $U = AB$ et $V = BC$, nous devons montrer que $UC = AV$. Si $A = (a_{ik})$ et $B = (b_{kj})$, alors $U = (u_{ij})$ avec

$$u_{ij} = \sum_{k=1}^n a_{ik}b_{kj},$$

donc le coefficient is de UC est

$$(UC)_{is} = \sum_{j=1}^p u_{ij}c_{js} = \sum_{j=1}^p \left(\sum_{k=1}^n a_{ik}b_{kj} \right) c_{js} = \sum_{j=1}^p \sum_{k=1}^n (a_{ik}b_{kj}c_{js}).$$

De même $V = (v_{ks})$ avec

$$v_{ks} = \sum_{j=1}^p b_{kj} c_{js},$$

donc le coefficient is de AV est

$$(AV)_{is} = \sum_{k=1}^n a_{ik} v_{ks} = \sum_{k=1}^n a_{ik} \left(\sum_{j=1}^p b_{kj} c_{js} \right) = \sum_{k=1}^n \sum_{j=1}^p (a_{ik} b_{kj} c_{js}).$$

Les matrices UC et AV ont les mêmes coefficients, elles sont donc égales et on a

$$(AB)C = UC = AV = A(BC).$$

□

Définition 5.6.2. Une *algèbre* sur le corps K est un K -espace vectoriel $\mathcal{A}$ muni d'une opération de multiplication interne $\mathcal{A} \times \mathcal{A} \rightarrow \mathcal{A}$ telle que

- a) $(\mathcal{A}, +, \cdot)$ est un anneau.
- b) $(\lambda A)(\mu B) = (\lambda\mu)(AB)$ pour tous $A, B \in \mathcal{A}$ et $\lambda, \mu \in K$.

L'algèbre est dite *unitaire* s'il existe un élément neutre pour la multiplication interne de $\mathcal{A}$; elle est dite *commutative* si $AB = BA$ pour tous $A, B \in \mathcal{A}$.

Donc $M_n(K) = M_{n \times n}(K)$ est une K -algèbre unitaire. Noter que $M_n(K)$ n'est pas commutative (en général $A \cdot B \neq B \cdot A$). Remarquons aussi que $A \cdot B = 0$ n'implique pas que $A = 0$ ou $B = 0$ (on dit que l'anneau $M_{n \times n}(K)$ admet des *diviseurs de zéro*, ou qu'il n'est pas *intègre*).

Résumons les propriétés du calcul matriciel.

Les identités du calcul matriciel	
1) $A + B = B + A$	9) $(A \cdot B) \cdot C = A \cdot (B \cdot C)$
2) $A + (B + C) = (A + B) + C$	10) $A \cdot (B + C) = A \cdot B + A \cdot C$
3) $A + O = O + A = A$	11) $(A + B) \cdot C = A \cdot C + B \cdot C$
4) $A + (-A) = O$	12) $(\lambda A) \cdot B = A \cdot (\lambda B) = \lambda(A \cdot B)$
5) $\lambda \cdot (A + B) = \lambda \cdot A + \lambda \cdot B$	13) $\mathbf{I}_m \cdot A = A \cdot \mathbf{I}_n = A$
6) $\lambda \cdot A + \mu \cdot A = (\lambda + \mu) \cdot A$	
7) $\lambda \cdot (\mu \cdot A) = (\lambda\mu) \cdot A$	
8) $1 \cdot A = A.$	

Remarques. Dans la colonne de gauche de ce tableau, toutes les matrices sont supposées de même type, dans la colonne de droite elles sont supposées multipliables, λ et μ sont des scalaires dans le corps K . Les propriétés (1) à (4) disent que $M_{m \times n}(K)$ est un groupe abélien, et les propriétés (5) à (6) disent que ce groupe abélien est un K -espace vectoriel. Lorsque $m = n$, les propriétés (9) à (11) disent que $M_n(K)$ est un anneau pour la multiplication matricielle. La propriété (12) dit que cet anneau est une K -algèbre et (13) dit que cette algèbre est unitaire. Rappelons aussi que l'anneau $M_n(K)$ n'est ni commutatif ni intègre (sauf si $n = 1$), en particulier il n'est pas un corps.

5.7 Matrices carrées : diagonale, transposée et matrices symétriques

Rappelons qu'une matrice est dite *carrée* si elle a autant de lignes que de colonnes, c'est à dire qu'elle est de taille $m \times n$ avec $m = n$.

La *diagonale principale* d'une matrice carrée A est la liste des coefficients dont les numéro de ligne et de colonne coïncident : $a_{1,1}, a_{2,2}, \dots, a_{n,n}$. Par exemple la diagonale principale de

$$\begin{pmatrix} 7 & 3 & \frac{4}{11} \\ 0 & -1 & 2 \\ 4 & 12 & 0 \end{pmatrix}$$

contient les nombres 7, -1 et 0.

Une matrice carrée est dite *diagonale* si tous les coefficients hors de la diagonale principale sont nuls. Par exemple D est diagonale et F ne l'est pas :

$$D = \begin{pmatrix} 5 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & -2 \end{pmatrix}, \quad F = \begin{pmatrix} 0 & 1 \\ 1 & 1 \end{pmatrix}$$

On note $\text{Diag}(\lambda_1, \dots, \lambda_n)$ la matrice diagonale dont le coefficient en position (i, i) est égale à λ_i :

$$\text{Diag}(\lambda_1, \dots, \lambda_n) = \begin{pmatrix} \lambda_1 & 0 & \dots & \dots & 0 \\ 0 & \lambda_2 & 0 & \dots & 0 \\ 0 & 0 & \lambda_3 & \dots & 0 \\ 0 & \vdots & \vdots & \ddots & 0 \\ 0 & 0 & \dots & 0 & \lambda_n \end{pmatrix}$$

Noter que la matrice identité est la matrice diagonale $\mathbf{I}_n = \text{Diag}(1, \dots, 1)$ (car la matrice identité de taille $n \times n$ est la matrice carrée diagonale dont les coefficients diagonaux valent tous 1).

Une matrice S est dite *scalaire* si elle est diagonale et si tous ses coefficients sont égaux. C'est donc un multiple scalaire de l'identité :

$$S = \lambda \mathbf{I}_n = \text{Diag}(\lambda, \dots, \lambda) = \begin{pmatrix} \lambda & 0 & \dots & 0 \\ 0 & \lambda & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \lambda \end{pmatrix}$$

La *transposée* d'une matrice A est la matrice obtenue en échangeant les lignes et les colonnes, on la note $A^\top$ (la transposée de A se note aussi A^t ou ${}^\top A$.)

Observons que si $A = (a_{ij})$, alors $A^\top = (a_{ji})$. Par exemple

$$\begin{pmatrix} 2 & 3 & 1 \\ 0 & 1 & 4 \\ 2 & 0 & 0 \end{pmatrix}^\top = \begin{pmatrix} 2 & 0 & 2 \\ 3 & 1 & 0 \\ 1 & 4 & 0 \end{pmatrix}$$

et

$$\begin{pmatrix} 2 \\ 7 \\ 0 \end{pmatrix}^\top = (2 \ 7 \ 0)$$

Une matrice carrée A est *symétrique* si $A = A^\top$ et elle est *antisymétrique* si $A = -A^\top$.
 Par exemple $F = \begin{pmatrix} 0 & 1 \\ 1 & 1 \end{pmatrix}$ est symétrique et

$$\begin{pmatrix} 0 & 2 & -4 \\ -2 & 0 & 5 \\ 4 & -5 & 0 \end{pmatrix}$$

est antisymétrique.

5.8 Matrices inversibles.

Définition Si A et B sont des matrices carrées de même taille $n \times n$, alors on dit que B est une matrice *inverse* de A si

$$A \cdot B = B \cdot A = \mathbf{I}_n.$$

On dit que A est *inversible* s'il existe une matrice inverse de A .

Proposition 5.8.1. *Si elle existe, l'inverse d'une matrice A est unique (et on le note $B = A^{-1}$).*

Preuve Supposons que B et C soient deux inverses de A , alors

$$B = B \cdot (A \cdot C) = (B \cdot A) \cdot C = \mathbf{I}_n \cdot C = C.$$

□

On verra plus loin que si $A \cdot B = \mathbf{I}_n$ ou $B \cdot A = \mathbf{I}_n$, alors A est inversible et $A^{-1} = B$.

Proposition 5.8.2. *L'ensemble des $n \times n$ matrices à coefficients dans le corps K qui sont inversibles forme un groupe pour la multiplication matricielle.*

On note cet ensemble $GL_n(K)$ et on l'appelle le *groupe linéaire général* d'ordre n .

Preuve. On a déjà prouvé l'associativité : si $A, B, C \in GL_n(K)$, alors $(AB)C = A(BC)$. Il est clair que $\mathbf{I}_n \in GL_n(K)$, et par définition tout élément $A \in GL_n(K)$ est inversible.

□

Proposition 5.8.3. *Le groupe des 2×2 matrices inversibles à coefficients dans le corps K est*

$$GL_2(K) = \left\{ A = \begin{pmatrix} a & b \\ c & d \end{pmatrix} \in M_2(K) \mid ad - bc \neq 0 \right\}.$$

Preuve Il faut montrer qu'une 2×2 matrice $A = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$ est inversible si et seulement si $ad - bc \neq 0$.

On part de l'identité matricielle suivante :

$$\begin{pmatrix} a & b \\ c & d \end{pmatrix} \cdot \begin{pmatrix} d & -b \\ -c & a \end{pmatrix} = \begin{pmatrix} ad - bc & 0 \\ 0 & ad - bc \end{pmatrix} = (ad - bc) \cdot \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}$$

Distinguons maintenant les cas $ad - bc \neq 0$ et $ad - bc = 0$. Si $ad - bc \neq 0$, alors l'identité ci-dessus entraîne que A est inversible et

$$A^{-1} = \frac{1}{ad - bc} \begin{pmatrix} d & -b \\ -c & a \end{pmatrix}. \quad (5.2)$$

Supposons maintenant que A est inversible, alors l'identité précédente entraîne que

$$\begin{pmatrix} d & -b \\ -c & a \end{pmatrix} = \begin{pmatrix} a & b \\ c & d \end{pmatrix}^{-1} \cdot \begin{pmatrix} ad-bc & 0 \\ 0 & ad-bc \end{pmatrix},$$

si on avait $ad-bc=0$, alors on aurait $d=-b=-c=a=0$, mais alors A serait la matrice nulle ce qui contredit qu'elle est inversible. On a bien montré que A est inversible si et seulement si $ad-bc \neq 0$. $\square$

Définition. Le *déterminant* de la 2×2 matrice $A = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$ est le nombre $\det(A) = ad-bc$. Nous avons donc montré que A est inversible si et seulement si $\det(A) \neq 0$, et dans ce cas l'inverse est donné par (5.2).

Exemples

$$\begin{pmatrix} 3 & -3 \\ 1 & 1 \end{pmatrix}^{-1} = \frac{1}{6} \begin{pmatrix} 1 & 3 \\ -1 & 3 \end{pmatrix}, \quad \begin{pmatrix} 1 & s \\ 0 & 1 \end{pmatrix}^{-1} = \begin{pmatrix} 1 & -s \\ 0 & 1 \end{pmatrix} \quad \text{et} \quad \begin{pmatrix} 0 & 3 \\ 2 & 5 \end{pmatrix}^{-1} = \begin{pmatrix} -\frac{5}{6} & \frac{1}{2} \\ \frac{1}{3} & 0 \end{pmatrix}$$

Nous étudierons plus loin l'inversion des matrices de taille 3×3 et plus.

5.9 Des applications linéaires vers les matrices.

Théorème 5.9.1. Soit $f : V \rightarrow W$ une application linéaire entre deux K -espaces vectoriels de dimensions finies. Soient $B = \{v_1, \dots, v_n\}$ et $B' = \{w_1, \dots, w_m\}$ des bases de V et W respectivement et notons $A = (a_{ij})$ la matrice de f dans ces bases. Alors pour tout élément $x = x_1v_1 + x_2v_2 + \dots + x_nv_n$ de V , on a $f(x) = y_1w_1 + y_2w_2 + \dots + y_mw_m$ avec

$$y_i = \sum_{j=1}^n a_{ij}x_j, \tag{5.3}$$

pour tout $i = 1, \dots, m$.

Remarque 5.9.2. Si on développe l'équation (5.3) pour chaque variable, on obtient le tableau

$$\begin{array}{ccccccc} y_1 & = & a_{11}x_1 & + & a_{12}x_2 & + & \dots & + & a_{1n}x_n \\ y_2 & = & a_{21}x_1 & + & a_{22}x_2 & + & \dots & + & a_{2n}x_n \\ \vdots & & & & & & & & \vdots \\ y_m & = & a_{m1}x_1 & + & a_{m2}x_2 & + & \dots & + & a_{mn}x_n. \end{array}$$

Observons que les coefficients a_{ij} se présentent de la manière dont on écrit la matrice de f dans les bases B et B' (la $i^{\text{ème}}$ ligne de ce tableau contient les coefficients de la $i^{\text{ème}}$ ligne de la matrice). C'est tout-à-fait différent du développement du vecteur $f(v_j)$ dans la base B' (comparer avec la remarque 5.4.2). Cette différence vient du fait que la coordonnée y_i ne doit en aucun cas être confondue avec le vecteur de base w_i . Nous reviendrons sur ce point au second semestre lorsque la notion d'espace dual sera introduite.

Preuve. Rappelons que les coefficients a_{ij} de la matrice de f dans les bases B et B' sont définis par

$$f(v_j) = a_{1j}w_1 + \dots + a_{mj}w_m = \sum_{i=1}^m a_{ij}w_i.$$

On a alors par linéarité de f :

$$\begin{aligned} f(x) = f\left(\sum_{j=1}^n x_j v_j\right) &= \sum_{j=1}^n x_j f(v_j) = \sum_{j=1}^n x_j \left(\sum_{i=1}^m a_{ij} w_i\right) \\ &= \sum_{j=1}^n \sum_{i=1}^m x_j a_{ij} w_i = \sum_{i=1}^m \left(\sum_{j=1}^n a_{ij} x_j\right) w_i, \end{aligned}$$

donc

$$f(x) = \sum_{i=1}^m y_i w_i \Leftrightarrow y_i = \sum_{j=1}^n a_{ij} x_j \text{ pour tout } i = 1, \dots, m$$

car $\{w_1, \dots, w_m\}$ est une base de W .

□

Une notation utile. Soit V un K -espace vectoriel de dimension finie n et $B = \{v_1, \dots, v_n\}$ une base de V . Alors tout vecteur $x \in V$ s'écrit d'une manière unique comme combinaison linéaire

$$x = x_1 v_1 + x_2 v_2 + \dots + x_n v_n.$$

On appelle *vecteur-colonne* de x associé à la base $\{v_j\}$ la $n \times 1$ matrice

$$X = M_B(x) = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix}.$$

Exemple. Si $V = K^3$ et si $x = (1, 0, -1)$, alors

$$M_{B_0}(x) = \begin{pmatrix} 1 \\ 0 \\ -1 \end{pmatrix} \quad \text{et} \quad M_{B_1}(x) = \begin{pmatrix} 1 \\ 1 \\ -1 \end{pmatrix}$$

relativement aux bases $B_0 = \{(1, 0, 0), (0, 1, 0), (0, 0, 1)\}$ et $B_1 = \{(1, 0, 0), (1, 1, 0), (1, 1, 1)\}$.

Le théorème précédent peut se reformuler de la façon suivante :

Théorème 5.9.3. Soit $f : V \rightarrow W$ une application linéaire entre deux espaces vectoriels de dimensions finies. Soient $B = \{v_1, \dots, v_n\}$ et $B' = \{w_1, \dots, w_m\}$ des bases de V et W respectivement et notons $A = (a_{ij}) = M_{B'B}(f)$ la matrice de f dans ces bases. Soient $x \in V$ et $y \in W$ deux vecteurs, et notons $X = M_B(x)$ et $Y = M_{B'}(y)$, alors

$$y = f(x) \Leftrightarrow Y = AX.$$

De manière plus concise, on peut écrire

$$M_{B'B}(f) \cdot M_B(x) = M_{B'}(f(x)).$$

Remarque. Si on identifie les éléments de K^n à des vecteurs colonnes (i.e on écrit les vecteurs de K^n comme des colonnes), et si on fait de même pour les éléments de K^m , alors on peut résumer la situation du théorème précédent par les diagrammes suivants :

$$\begin{array}{ccc} V & \xrightarrow{f} & W \\ \downarrow & & \downarrow \\ K^n & \xrightarrow{A} & K^m \end{array} \quad \begin{array}{ccc} x & \mapsto & y \quad (= f(x)) \\ \downarrow & & \downarrow \\ X & \mapsto & Y \quad (= A \cdot X) \end{array}$$

Proposition 5.9.4. Soient V et W deux K -espaces vectoriels de dimensions finies et $B = \{v_1, \dots, v_n\}$ et $B' = \{w_1, \dots, w_m\}$ des bases de V et W respectivement. Alors l'application

$$\mathcal{L}(V, W) \rightarrow M_{m \times n}(K)$$

$$f \mapsto M_{B', B}(f)$$

est un isomorphisme d'espaces vectoriels.

Preuve. Il est facile de vérifier que si $f, g \in \mathcal{L}(V, W)$ ont pour matrices $(a_{ij}) = M_{B', B}(f)$ et $(b_{ij}) = M_{B', B}(g)$, et si $\lambda \in K$, alors

$$M_{B', B}(f + g) = (a_{ij} + b_{ij}) \quad \text{et} \quad M_{B', B}(\lambda f) = (\lambda a_{ij}).$$

Cela montre que l'application $f \mapsto M_{B', B}(f)$ est linéaire. Il reste à prouver $f \mapsto M_{B, B'}(f)$ est un isomorphisme. Comme les deux espaces vectoriels $\mathcal{L}(V, W)$ et $M_{m \times n}(K)$ ont même dimension (qui vaut $\dim(V) \cdot \dim(W)$), il suffit de prouver que le noyau de $M_{B, B'}$ est nul, i.e. que la matrice d'une application linéaire est nulle si et seulement si cette application linéaire est nulle. Mais si $M_{B', B}(f) = 0$ (la matrice nulle), alors par définition $f(v_j) = 0$ pour tout $j = 1, \dots, n$, donc f est identiquement nulle. $\square$

Proposition 5.9.5. Soient U, V et W trois K -espaces vectoriels de dimensions finies et

$$B_1 = \{u_1, \dots, u_p\}, \quad B_2 = \{v_1, \dots, v_n\} \quad \text{et} \quad B_3 = \{w_1, \dots, w_m\}$$

des bases de U, V et W respectivement. Si $g \in \mathcal{L}(U, V)$ et $f \in \mathcal{L}(V, W)$, alors $f \circ g \in \mathcal{L}(U, W)$ et

$$M_{B_3 B_1}(f \circ g) = M_{B_3 B_2}(f) \cdot M_{B_2 B_1}(g)$$

Preuve. Supposons que $M(f) = A = (a_{ik})$, $M(g) = B = (b_{kj})$ et $M(f \circ g) = C = (c_{ij})$. Alors on a

$$f(v_k) = \sum_{i=1}^m a_{ik} w_i, \quad g(u_j) = \sum_{k=1}^n b_{kj} v_k$$

donc

$$f \circ g(u_j) = f\left(\sum_{k=1}^n b_{kj} v_k\right) = \sum_{k=1}^n b_{kj} f(v_k) = \sum_{k=1}^n b_{kj} \sum_{i=1}^m a_{ik} w_i = \sum_{i=1}^m c_{ij} w_i$$

avec

$$c_{ij} = \sum_{k=1}^n a_{ik} \cdot b_{kj} = a_{i1} \cdot b_{1j} + a_{i2} \cdot b_{2j} + \dots + a_{in} \cdot b_{nj}$$

Ce qui veut dire que $C = A \cdot B$. $\square$

Remarque. Cette proposition n'est en fait pas une surprise, car on a défini la multiplication matricielle de façon qu'elle corresponde à la composition des applications linéaires.

Corollaire 5.9.6. Si $f : V \rightarrow W$ est un isomorphisme entre deux espaces vectoriels de dimension finie, alors pour tout choix de bases de V et W , la matrice de f est inversible.

Nous laissons la preuve de ce corollaire en exercice.

5.10 Forme matricielle spéciale d'une application linéaire

Théorème 5.10.1. Soit $f : V \rightarrow W$ une application linéaire entre deux K -espaces vectoriels V et W , de dimensions finies n et m . Alors il existe des bases $B \subset V$ et $B' \subset W$ pour lesquelles

$$M_{B'B}(f) = \begin{pmatrix} \mathbf{I}_r & 0_{r \times (n-r)} \\ 0_{(m-r) \times r} & 0_{(m-r) \times (n-r)} \end{pmatrix}$$

où $r = \text{rang}(f)$ et $0_{p \times q}$ est la matrice nulle de taille $p \times q$.

L'idée de la preuve est de construire des sous-espaces vectoriels $V_1, V_2 \subset V$ et $W_1, W_2 \subset W$ tels que $V = V_1 \oplus V_2$ et $W = W_1 \oplus W_2$ avec $V_2 = \text{Ker}(f)$, $W_1 = \text{Im}(f)$ et la restriction de f à V_1 définit un isomorphisme $V_1 \cong W_1$.

Preuve. Rappelons que le rang de f est la dimension de $\text{Im}(f)$. Choisissons une base $w_1, \dots, w_r$ de $\text{Im}(f)$, puis complétons cette base en une base $B' = \{w_1, \dots, w_r, w_{r+1}, \dots, w_m\}$ de W . Pour tous $i = 1, \dots, r$ choisissons un élément $v_i \in V$ tel que $f(v_i) = w_i$. Il est facile de vérifier que $v_1, \dots, v_r$ sont linéairement indépendants (nous laissons la vérification en exercice).

Le noyau de f est de dimension $n - r$ (par le théorème du rang : $\dim(V) = \dim(\text{Ker}(f)) + r$) ; on peut donc choisir une base $v_{r+1}, \dots, v_n$ de $\text{Ker}(f)$. Vérifions que $B = \{v_1, \dots, v_r, v_{r+1}, \dots, v_n\}$ est une famille libre de V : soient $\lambda_1, \dots, \lambda_n \in K$ tels que $\sum_{j=1}^n \lambda_j v_j = 0$, alors

$$0 = f\left(\sum_{j=1}^n \lambda_j v_j\right) = \sum_{j=1}^n \lambda_j f(v_j) = \sum_{j=1}^r \lambda_j w_j \quad (\text{car } f(v_j) = 0 \text{ si } j > r).$$

Or les vecteurs w_j sont linéairement indépendants, donc $\lambda_1 = \lambda_2 = \dots = \lambda_r = 0$. Mais alors on a $\sum_{j=r+1}^n \lambda_j v_j = 0$ et comme $v_{r+1}, \dots, v_n$ est une base de $\text{Ker}(f)$, cela entraîne que $\lambda_{r+1} = \dots = \lambda_n = 0$, finalement tous les λ_j sont nuls et B est donc une famille libre. C'est donc une base de V puisque $\text{Card}(B) = \dim(V)$.

On a donc trouvé deux bases $B = \{v_1, \dots, v_n\} \subset V$ et $B' = \{w_1, \dots, w_m\} \subset W$ telles que

$$f(v_j) = \begin{cases} w_j & \text{si } j \leq r \\ 0 & \text{si } j > r \end{cases}$$

La matrice de f dans ces bases a donc la forme indiquée. □

5.11 Matrice de changement de bases

Définition. Soit V un K -espace vectoriel et donnons-nous deux bases $B = \{v_1, \dots, v_n\}$ et $B' = \{v'_1, \dots, v'_n\}$ de V . Il est commode d'appeler B l'ancienne base et B' la nouvelle base. Si on écrit les éléments de la nouvelle base comme combinaison linéaire des éléments de l'ancienne base, on obtient une matrice carrée de taille $n \times n$, notons-la P . Ainsi $P = (p_{ij})$ avec

$$v'_j = \sum_{i=1}^n p_{ij} v_i.$$

La matrice P s'appelle *matrice de transition* (ou matrice de passage) de la base B vers la base B' .

Il y a deux façons d'interpréter cette matrice :

- P est la matrice dans la base B de l'application linéaire $\varphi \in \mathcal{L}(V, V)$ telle que $\varphi(v_j) = v'_j$ pour tout j , c'est-à-dire $P = M_{BB}(\varphi)$.
- P est la matrice de l'application identité de la base B' vers la base B (donc de la nouvelle base vers l'ancienne, ce qui n'est pas très intuitif) :

$$P = M_{BB'}(\text{Id}_V).$$

Observons en particulier que par le corollaire 5.9.6, une matrice de transition est toujours inversible.

Proposition 5.11.1. *Si on note X le vecteur colonne dans la base B d'un vecteur $x \in V$ et X' le vecteur colonne du même vecteur x dans la base B' (i.e. $X = M_B(x)$ et $X' = M_{B'}(x)$), alors on a*

$$X = PX' \quad \text{i.e.} \quad x_i = \sum_{j=1}^n p_{ij} x'_j.$$

Preuve. C'est un cas particulier du théorème 5.9.3 :

$$PX' = M_{BB'}(\text{Id}) \cdot M_{B'}(x) = M_B(\text{Id}(x)) = M_B(x) = X.$$

On peut aussi le voir avec le calcul suivant :

$$\sum_{j=1}^n x'_j v'_j = \sum_{j=1}^n x'_j \left(\sum_{i=1}^n p_{ij} v_i \right) = \sum_{i=1}^n \left(\sum_{j=1}^n p_{ij} x'_j \right) v_i = \sum_{i=1}^n x_i v_i.$$

□

Considérons maintenant la situation suivante : soient V et W deux espaces vectoriels munis chacun de deux bases : $B = \{v_1, \dots, v_n\}$ et $B' = \{v'_1, \dots, v'_n\}$ sont deux bases de V

et $E = \{w_1, \dots, w_m\}$ et $E' = \{w'_1, \dots, w'_m\}$ sont deux bases de W .

Notons P la matrice de transition de B vers B' et Q la matrice de transition de E vers E' , i.e.

$$P = M_{BB'}(\text{Id}_V) \quad \text{et} \quad Q = M_{EE'}(\text{Id}_W).$$

Soit maintenant une application linéaire $f \in \mathcal{L}(V, W)$ et notons A sa matrice dans les bases B, E et A' sa matrice dans les bases B', E' , i.e. $A = M_{EB}(f)$ et $A' = M_{E'B'}(f)$.

Théorème 5.11.2. *Sous les conditions précédentes, les matrices P et Q sont inversibles et on a*

$$A' = Q^{-1}AP.$$

Preuve. On a

$$QA' = M_{EE'}(\text{Id}_W)M_{E'B'}(f) = M_{EB'}(\text{Id}_W \circ f) = M_{EB'}(f) = M_{EB}(f)M_{BB'}(\text{Id}_V) = AP.$$

□

Remarque : Voici deux autres preuves de ce théorème, la première est calculatoire et la seconde est matricielle.

i.) On peut aussi démontrer le théorème par un calcul direct :

$$f(v'_j) = \sum_{i=1}^m a'_{ij} w'_i = \sum_{i=1}^m \sum_{k=1}^m a'_{ij} q_{ki} w_k = \sum_{k=1}^m \left(\sum_{i=1}^m q_{ki} a'_{ij} \right) w_k$$

et

$$f\left(\sum_{s=1}^n p_{sj} v_s\right) = \sum_{s=1}^n \sum_{k=1}^m p_{sj} a_{ks} w_k = \sum_{k=1}^m \left(\sum_{s=1}^n a_{ks} p_{sj}\right) w_k,$$

or $v'_j = \sum_{s=1}^n p_{sj} v_s$, donc le calcul précédent montre que $\sum_{i=1}^m q_{ki} a'_{ij} = \sum_{s=1}^n a_{ks} p_{sj}$ pour tout k , ce qui signifie précisément que $QA' = AP$.

- ii.) Et on peut raisonner sur les vecteurs-colonnes. Soit x un vecteur quelconque de V et $y = f(x) \in W$. Notons $X = M_B(x)$, $X' = M_{B'}(x)$, $Y = M_E(y)$, et $Y' = M_{E'}(y)$. Alors $X = PX'$ et $Y = QY'$ mais on a aussi $Y = AX$ et $Y' = A'X'$, donc

$$APX' = AX = Y = QY' = QA'X'.$$

Comme cette égalité a lieu pour tout vecteur colonne X' , on en déduit que $AP = QA'$.

□

5.12 Des matrices vers les applications linéaires.

Nous avons vu comment associer une matrice à une application linéaire entre deux espaces vectoriels munis de bases. Dans ce paragraphe nous prenons le problème dans l'autre sens et nous associons à toute matrice une application linéaire.

Définition. Soit $A = (a_{ij}) \in M_{m \times n}(K)$ une matrice de taille $m \times n$ à coefficient dans un corps K . Alors on définit une application linéaire $L_A : K^n \rightarrow K^m$ en posant :

$$L_A(e_j) = \sum_{i=1}^m a_{ij} e'_i$$

où $\{e_1, \dots, e_n\}$ est la base canonique de K^n et $\{e'_1, \dots, e'_m\}$ est la base canonique de K^m . On appelle L_A l'*application linéaire associée* à la matrice A .

On démontre facilement les points suivants :

- La matrice de L_A dans les bases $\{e_j\}$ et $\{e'_i\}$ est égale à A : $M_{\{e'_i\}, \{e_j\}}(L_A) = A$.
- Si X est le vecteur colonne associé à $x \in K^n$, alors AX est le vecteur colonne associé à $L_A(x) \in K^m$.
- L'application $A \mapsto L_A$ est un isomorphisme $M_{m \times n}(K) \rightarrow \mathcal{L}(K^n, K^m)$. En particulier $L_{A_1} = L_{A_2}$ si et seulement si $A_1 = A_2$.

Nous avons aussi la proposition suivante :

Proposition 5.12.1. Si $B \in M_{n \times p}(K)$ et $A \in M_{m \times n}(K)$, alors les applications $L_B : K^p \rightarrow K^n$ et $L_A : K^n \rightarrow K^m$ peuvent être composées et on a

$$L_A \circ L_B = L_{A \cdot B} \in \mathcal{L}(K^p, K^m).$$

Ce résultat est une nouvelle illustration du fait que la multiplication matricielle a été définie de façon qu'elle corresponde à la composition des applications linéaires.

Preuve. C'est une conséquence de la proposition 5.9.5. On peut aussi le prouver directement : soit $x \in K^p$ et X son vecteur colonne, alors

$$(L_A \circ L_B)(X) = L_A(L_B(X)) = L_A(BX) = A(BX) = (AB)X = L_{A \cdot B}(X).$$

□

Convention. Lorsqu'on étudie une (ou des) matrice(s) comme application(s) linéaire(s), on se permet d'identifier les vecteurs $x \in K^n$ avec leurs vecteurs-colonnes associés $X \in M_{n \times 1}(K)$ et les matrices $A \in M_{m \times n}(K)$ avec l'application linéaire associée $L_A \in \mathcal{L}(K^n, K^m)$. Concrètement, cela signifie qu'on regarde les vecteurs de K^n comme des vecteurs-colonnes et on regarde les matrices comme des applications linéaires.

Par conséquent on note

- $\text{Ker}(A) = \{X \in K^n \mid AX = 0\}$.
- $\text{Im}(A) = \{A \cdot X \mid X \in K^n\}$.
- $\text{rang}(A) = \dim(\text{Im}(A))$.

Proposition 5.12.2. Soit $A \in M_n(K)$ une matrice carrée, alors on a

- (a) Si A admet une inverse à gauche B (i.e. $BA = \mathbf{I}_n$), alors B est aussi une inverse à droite,
- (b) En particulier A est inversible et on a $A^{-1} = B$.
- (c) De même, si $A \in M_n(K)$ admet une inverse à droite C , alors C est aussi une inverse à gauche. En particulier l'inverse à gauche est égal à l'inverse à droite.
- (d) De plus une matrice $A \in M_n(K)$ est inversible si et seulement si $\text{rang}(A) = n$.

Remarque. Attention, cette proposition est fausse pour les matrices qui ne sont pas carrées !

Preuve. (a) On suppose que $A, B \in M_n(K)$ vérifient $BA = \mathbf{I}_n$, on a donc $L_B \circ L_A = \text{id} : K^n \rightarrow K^n$. En particulier L_A est injective, or on sait que toute application linéaire injective entre deux espaces vectoriels de même dimensions est bijective.

(b) Par conséquent, pour tout $X \in K^n$ (vu comme un vecteur colonne), il existe $Z \in K^n$ tel que $L_A(Z) = X$. Ainsi :

$$ABX = AB(AZ) = A(BA)Z = AI_n Z = AZ = X.$$

Ceci est vrai pour tout $X \in K^n$, donc $AB = \mathbf{I}_n$. On prouve de même que $BA = \mathbf{I}_n$.

(c) La preuve est semblable à la preuve de (a).

(d) La dernière assertion vient du fait que $\text{rang}(A) = \dim(\text{Im}(A))$ si et seulement si l'application L_A est surjective.

□

Proposition 5.12.3. Le rang de la matrice $A \in M_{m \times n}(K)$ est le nombre maximal de colonnes de cette matrice qui sont linéairement indépendantes.

Preuve. Si on note $A_1, A_2, \dots, A_n$ les colonnes de la matrices A , alors $\text{Im}(A)$ est l'espace vectoriel engendré par ces vecteurs colonnes :

$$\text{Im}(A) = \text{Vec}(\{A_1, A_2, \dots, A_n\}) \subset K^m,$$

car A_j est l'image par A du $j^{\text{ème}}$ vecteur de base. Le théorème 4.3.5 nous dit alors que la famille $\{A_1, A_2, \dots, A_n\}$ contient une base de $\text{Im}(A)$. Cette base est formée de r vecteurs linéairement indépendants et toute famille de plus de r vecteurs dans $\text{Im}(A)$ est une famille liée.

□

Définition. On dit que deux matrices de même tailles A et A' dans $M_{m \times n}$ sont équivalentes s'il existe des matrices carrées inversibles $P \in M_n(K)$ et $Q \in M_m(K)$ telles que

$$A' = Q^{-1}AP. \tag{5.4}$$

Théorème 5.12.4. Toute matrice A dans $M_{m \times n}(K)$ est équivalente à une matrice

$$\begin{pmatrix} \mathbf{I}_r & 0 \\ 0 & 0 \end{pmatrix} \quad (5.5)$$

où r est le rang de A et les 0 représentent des matrices nulles de dimensions adéquates. En particulier deux matrices de même taille et de même rang sont toujours équivalentes.

Preuve. Le théorème 5.10.1 nous dit qu'il existe des bases $\mathcal{B}'$ de K^m et $\mathcal{E}'$ de K^n telles que la matrice de L_A dans les bases $\mathcal{B}', \mathcal{E}'$ est donnée par

$$A' = M_{\mathcal{E}', \mathcal{B}'}(L_A) = \begin{pmatrix} \mathbf{I}_r & 0 \\ 0 & 0 \end{pmatrix}.$$

Notons P la matrice de passage de la matrice canonique de K^n vers $\mathcal{B}'$ et Q la matrice de passage de la matrice canonique de K^m vers $\mathcal{E}'$, alors le théorème 5.11.2 nous dit alors que

$$Q^{-1} A P = A' = \begin{pmatrix} \mathbf{I}_r & 0 \\ 0 & 0 \end{pmatrix}.$$

□

En plus de la définition d'équivalence, on introduit trois autres types de relation entre les matrices :

Définition. i.) On dit que deux matrices de même taille A et A' dans $M_{m \times n}(K)$ sont *équivalentes à droite* s'il existe une matrice carrée inversible $P \in M_n(K)$ telle que

$$A' = AP.$$

ii.) On dit que $A, A' \in M_{m \times n}(K)$ sont *équivalentes à gauche* s'il existe une matrice carrée inversible $Q \in M_m(K)$ telles que

$$A' = Q^{-1}A.$$

iii.) Deux matrices carrées de même taille A et A' dans $M_n(K)$ sont *semblables* (ou *similaires*, ou *conjuguées*) s'il existe une matrice carrée inversible telle que

$$A' = P^{-1}AP.$$

Remarques. 1.) Chacune de ces relations vérifie les propriétés suivantes : si on note $A \simeq B$ pour dire ou bien que A et B sont équivalentes, ou bien qu'elles sont droite-équivalentes ou gauche-équivalentes, ou encore semblables, alors

- $A \simeq A$;
- $A \simeq B \Rightarrow B \simeq A$;
- $A \simeq B$ et $B \simeq C \Rightarrow A \simeq C$.

2.) L'équivalence définie par (5.4) est plus faible que les trois autres relations que nous venons de définir, i.e. équivalence-droite $\Rightarrow$ équivalence, équivalence-gauche $\Rightarrow$ équivalence, similitude $\Rightarrow$ équivalence.

3.) Nous verrons aux exercices que deux matrices carrées de même taille peuvent être équivalentes mais ne pas être semblables. En particulier il n'est pas vrai qu'une matrice carrée est toujours semblable à une matrice de type (5.5). Le problème de déterminer si deux matrices données sont semblables est un problème important. Il sera étudié aux second semestre.

Chapitre 6

Systèmes linéaires

6.1 Sous-espaces affines d'un espace vectoriel

De nombreux problèmes d'algèbre (théoriques et pratiques) se ramènent à la résolution d'un *système d'équations linéaires*. Par exemple *déterminer tous les éléments $(x, y, z) \in \mathbb{R}^3$ tels que*

$$f(x, y, z) = (2x - y, x + z) = (1, 2).$$

Pour résoudre ce problème, on écrit cette équation vectorielle comme un système de 2 équations à 3 inconnues :

$$\begin{aligned} 2x - y &= 1 \\ x + z &= 2. \end{aligned}$$

Donnons une valeur t quelconque à x , alors on obtient $x = t$, $y = 2t - 1$ et $z = 2 - t$. L'ensemble des solutions est donc donné par

$$S = \{(t, 2t - 1, 2 - t) \mid t \in \mathbb{R}\}.$$

Cet ensemble représente la droite de $\mathbb{R}^3$ passant par le point $(0, -1, 2)$ et parallèle au vecteur $(1, 2, -1)$; on peut écrire

$$S = (0, -1, 2) + \text{Vec}(\{(1, 2, -1)\}).$$

On dit que S est une *droite affine* dans l'espace vectoriel $\mathbb{R}^3$. Plus généralement, on a la définition suivante :

Définition. Soit V un espace vectoriel sur le corps K . On dit qu'un sous-ensemble $E \subset V$ est un *sous-espace affine* si ou bien $E = \emptyset$, ou bien il existe $p \in V$ et un sous-espace vectoriel $W \subset V$ tel que

$$x \in E \Leftrightarrow (x - p) \in W.$$

De manière équivalente,

$$E = p + W = \{x = p + w \mid w \in W\}.$$

On dit alors que E est le *translaté* par p du sous-espace vectoriel W .

Définition. Si $E \subset V$ est un sous-espace affine non vide, alors on dit que le sous-espace vectoriel W dont E est un translaté est l'*espace des directions* de E , on dit aussi que c'est l'*espace directeur* de E .

La *dimension* du sous-espace affine E est par définition la dimension de son espace directeur W . Si $E = \emptyset$ on convient que $\dim(\emptyset) = -1$. Un sous-espace affine de dimension 1 de l'espace vectoriel V s'appelle une *droite affine* de V ; un sous-espace affine de dimension 2 s'appelle un *plan affine*. Un *hyperplan* est un sous-espace affine de dimension $\dim(V) - 1$ (donc de codimension 1). Un sous-espace affine de dimension 0 s'appelle un *point* de V . Il ne contient qu'un seul élément et on identifie habituellement le sous-espace affine $\{p\}$ avec l'élément $p \in V$.

Proposition 6.1.1. Si $E \subset V$ est un sous-espace affine et $W \subset V$ son espace directeur, alors pour tout $q \in E$ on a $E = q + W$.

Preuve. Si $E = \emptyset$ il n'y a rien à prouver. Si $E \neq \emptyset$, alors par définition de la notion d'espace affine, il existe $p \in E$ tel que $E = p + W$. Soit q un autre élément de E , alors il existe $w_0 \in W$ tel que $q = p + w_0$ et on a

$$x \in E \Leftrightarrow (x - p) \in W \Leftrightarrow (x - (q - w_0)) = ((x - q) + w_0) \in W \Leftrightarrow (x - q) \in W,$$

donc $q + W = p + W$. □

Une conséquence importante est que si $\{w_1, w_2, \dots, w_d\}$ est une base de W et $q \in E$, alors tout élément x de E peut s'écrire sous la forme

$$x = q + \sum_{i=1}^d \lambda_i w_i$$

avec $\lambda_1, \dots, \lambda_d \in K$.

L'une des motivations pour introduire la notion de sous-espace affine est le résultat suivant :

Théorème 6.1.2. Soit $f : V_1 \rightarrow V_2$ une application linéaire entre deux espaces vectoriels sur le même corps K . Alors pour tout $b \in V_2$, l'ensemble

$$f^{-1}(b) = \{x \in V_1 | f(x) = b\}$$

est un sous-espace affine de V_1 .

Si $b \notin \text{Im}(f)$, alors ce sous-espace affine est vide, et si $b \in \text{Im}(f)$ alors l'espace directeur de $f^{-1}(b)$ est le noyau $\text{Ker}(f)$.

Preuve. Si $b \notin \text{Im}(f)$ il n'y a rien à prouver. Si $b \in \text{Im}(f)$, alors par définition il existe $p \in V_1$ tel que $f(p) = b$. On a alors

$$x \in f^{-1}(b) \Leftrightarrow f(x) = f(p) \Leftrightarrow f(x - p) = 0 \Leftrightarrow (x - p) \in \text{Ker}(f) \Leftrightarrow x \in p + \text{Ker}(f).$$

□

Conséquence. De façon explicite, ce théorème dit que si $q \in V_1$ est une solution particulière de l'équation $f(x) = b$, et si $\{w_1, w_2, \dots, w_d\}$ est une base de $\text{Ker}(f)$ (ce sont donc d solutions linéairement indépendantes de l'équation (dite homogène) $f(w) = 0$), alors la solution générale de l'équation $f(x) = b$ est de la forme $x = q + \sum_{i=1}^d \lambda_i w_i$.

6.2 Systèmes d'équations linéaires

Définition 6.2.1. Un système d'équations linéaires à m équations et n inconnues (m et n deux entiers positifs) sur le corps K est un ensemble d'équations du type

$$(\star) \quad \begin{cases} a_{11}x_1 + a_{12}x_2 + \cdots + a_{1n}x_n = b_1 \\ a_{21}x_1 + a_{22}x_2 + \cdots + a_{2n}x_n = b_2 \\ \vdots \\ a_{m1}x_1 + a_{m2}x_2 + \cdots + a_{mn}x_n = b_m \end{cases}$$

où $1 \leq i \leq m$, $1 \leq j \leq n$.

- Les coefficients du système a_{ij} et les termes inhomogènes b_i sont des scalaires (des éléments du corps de base) donnés.

- Les x_j sont les inconnues que l'on cherche à déterminer.
- Le système obtenu à partir de $(\star)$ en remplaçant les b_j par 0 s'appelle le *système homogène* associé à $(\star)$.
- Une *solution* du système est un n -tuple $(x_1, x_2, \dots, x_n) \in K^n$ qui satisfait simultanément toutes les m équations du système.

Remarque : On peut écrire le système d'équations sous la forme d'une unique équation matricielle :

$$\begin{pmatrix} a_{11} & \dots & a_{1n} \\ \vdots & \ddots & \vdots \\ a_{m1} & \dots & a_{mn} \end{pmatrix} \cdot \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} = \begin{pmatrix} b_1 \\ \vdots \\ b_m \end{pmatrix}$$

D'une manière abrégée, on peut écrire $A \cdot X = B$ où X et B sont des vecteurs-colonnes et A est une matrice rectangulaire. Le vecteur X représente les inconnues du système.

On a alors le théorème suivant :

Théorème 6.2.2. *L'ensemble*

$$\mathcal{S} = \{X \in K^n \mid AX = B\}$$

des solutions du système $(\star)$ est un sous-espace affine de K^n . Si $\mathcal{S} \neq \emptyset$, alors $\mathcal{S} = X_0 + \text{Ker}(A)$, où $X_0 \in K^n$ est une solution particulière du système.

La preuve est immédiate à partir des résultats du paragraphe précédent. □

Définition. Le système $(\star)$ est dit *incompatible* si $\mathcal{S} = \emptyset$. D'une manière équivalente, le système est incompatible si et seulement si $B \notin \text{Im}(A)$.

Dans la pratique, pour résoudre le système $(\star)$, on essaye d'abord de déterminer s'il est compatible ou incompatible. S'il est compatible, on cherche une solution particulière X_0 , puis on cherche une base $w_1, \dots, w_r$ de $\text{Ker}(A)$. Toute solution de $(\star)$ s'écrit alors de façon unique sous la forme

$$X = X_0 + \lambda_1 w_1 + \dots + \lambda_r w_r.$$

6.3 La méthode de Gauss-Jordan

Il y a plusieurs méthodes pour résoudre un système linéaire. L'une des méthodes les plus efficaces est de procéder par élimination systématique des variables. Cette méthode, que le lecteur a sans-doute déjà vue lors de ses études secondaires, semble avoir été pratiquée depuis la nuit des temps, on en trouve trace sur des tablettes cunéiformes babyloniennes (environ 2000 ans avant J.C) et un écrit chinois du premier siècle, le *Jiuzhang Suanshu*, ou les *Neuf Chapitres sur l'Art Mathématique*, en donne une description précise.

On appelle *algorithme de Gauss-Jordan* une méthode systématique pour résoudre un système linéaire par élimination.

Commençons par un exemple élémentaire, considérons le système

$$\begin{cases} 2x & - & 3y & = & 5 \\ 4x & + & y & = & -4 \end{cases}$$

Pour résoudre ce système, on peut soustraire de la seconde équation le double de la première (ce qui élimine x de la seconde équation), puis diviser chaque équation par son premier coefficient non nul (qu'on appelle le *pivot*) :

$$\begin{cases} 2x & - & 3y & = & 5 \\ & 7y & = & -14 \end{cases} \rightarrow \begin{cases} x & - & \frac{3}{2}y & = & \frac{5}{2} \\ & y & = & -2 \end{cases}$$

et on obtient la solution $y = -2$, $x = (3y+5)/2 = -1/2$.

Décrivons maintenant le cas général. Considérons le système de m équations linéaires à n inconnues sur le corps K :

$$\begin{cases} a_{11}x_1 + \cdots + a_{1n}x_n = b_1 \\ \vdots \\ a_{m1}x_1 + \cdots + a_{mn}x_n = b_m \end{cases}$$

On peut écrire ce système sous forme matricielle $AX = B$. Il est commode de considérer la matrice de taille $m \times (n+1)$ suivante :

$$\begin{pmatrix} a_{11} & \cdots & a_{1n} & b_1 \\ \vdots & & \vdots & \vdots \\ a_{m1} & \cdots & a_{mn} & b_m \end{pmatrix}$$

On appelle cette matrice la *matrice augmentée* du système. On la note aussi $(A|b)$ et on remarque qu'elle contient toute l'information nécessaire à la résolution des équations.

L'algorithme de Gauss-Jordan consiste à ramener la matrice augmentée $(A|b)$ à une forme plus simple, au moyen d'une suite de transformations dites *transformations élémentaires*. Ces transformations sont les suivantes :

Type 1 : On échange deux lignes de la matrice augmentée $(A|b)$.

Type 2 : On multiplie une ligne de $(A|b)$ par un scalaire non nul $\lambda \in K^*$.

Type 3 : On ajoute à une ligne de $(A|b)$ un multiple d'une autre ligne.

Il est facile de se convaincre que ces trois types de transformations ne changent pas les solutions du système d'équations.

Par une suite de transformations élémentaires, on peut obtenir une nouvelle matrice $(A'|b')$ qui satisfait aux trois conditions suivantes :

- (i) Toutes les lignes de $(A'|b')$ qui ne contiennent que des zéros, s'il en existe, se trouvent en bas de la matrice.
- (ii) Le premier coefficient non nul d'une ligne s'appelle le *pivot* de cette ligne. On suppose que tout coefficient en dessous d'un pivot est nul.
- (iii) Le pivot de chaque ligne non nulle est situé à droite du pivot de la ligne précédente.

Définition. Si une matrice satisfait ces trois conditions, on dit qu'elle est sous *forme échelonnée*. On dit qu'elle est sous *forme échelonnée réduite*, si elle satisfait aux deux conditions supplémentaires suivantes :

- (iv) Chaque pivot est égal à 1.
- (v) Tout coefficient en dessus d'un pivot est nul (donc le pivot est le seul coefficient non nul de sa colonne).

Reprenons l'exemple du système de deux équations à deux inconnues précédent. La méthode d'élimination, exprimée avec la matrice augmentée, se présente sous la forme suivante :

$$(A|b) = \begin{pmatrix} 2 & -3 & 5 \\ 4 & 1 & -4 \end{pmatrix} \rightarrow \begin{pmatrix} 2 & -3 & 5 \\ 0 & 7 & -14 \end{pmatrix}$$

Le système initial est donc équivalent à

$$\begin{cases} 2x - 3y = 5 \\ 7y = -14 \end{cases}$$

qui se résout facilement en $(x, y) = (-\frac{1}{2}, -2)$.

On peut aussi effectuer quelques transformations élémentaires supplémentaires et ramener la matrice augmentée à sa forme échelonnée réduite :

$$\begin{pmatrix} 2 & -3 & 5 \\ 0 & 7 & -14 \end{pmatrix} \rightarrow \begin{pmatrix} 2 & -3 & 5 \\ 0 & 1 & -2 \end{pmatrix} \rightarrow \begin{pmatrix} 2 & 0 & -1 \\ 0 & 1 & -2 \end{pmatrix} \rightarrow \begin{pmatrix} 1 & 0 & -\frac{1}{2} \\ 0 & 1 & -2 \end{pmatrix}$$

Le système initial est donc équivalent au système suivant

$$\begin{cases} x & = & -\frac{1}{2} \\ y & = & -2, \end{cases}$$

dont la solution est immédiate.

Théorème 6.3.1. *A) Toute matrice est équivalente à une unique matrice échelonnée réduite par une suite de transformations élémentaires.*

B) Le rang d'une matrice ne change pas lors des transformations élémentaires. Le rang d'une matrice échelonnée est le nombre de lignes non nulles ; c'est aussi le nombre de pivots.

C) Une matrice carrée est inversible si et seulement si sa matrice échelonnée réduite est la matrice identité.

La preuve de ce théorème consiste à formaliser précisément l'algorithme de Gauss-Jordan et à prouver par un argument de récurrence que cet algorithme s'arrête après un nombre fini d'étapes.

Exemple. On considère le système de 4 équations à 4 inconnues suivant

$$\begin{cases} x + 2z & = & 2 \\ 2x + y + 3z & = & -1 \\ 3y - 3z + 3t & = & -12 \\ 3x + y + 5z & = & 1 \end{cases}$$

En alignant les variables, on trouve la matrice augmentée :

$$\begin{cases} x & + & 2z & = & 2 \\ 2x & + & y & + & 3z & = & -1 \\ & & 3y & - & 3z & + & 3t & = & -12 \\ 3x & + & y & + & 5z & = & 1 \end{cases} \Rightarrow (A|b) = \begin{pmatrix} 1 & 0 & 2 & 0 & 2 \\ 2 & 1 & 3 & 0 & -1 \\ 0 & 3 & -3 & 3 & -12 \\ 3 & 1 & 5 & 0 & 1 \end{pmatrix}$$

La méthode de Gauss-Jordan se déroule comme suit :

$$\begin{pmatrix} 1 & 0 & 2 & 0 & 2 \\ 2 & 1 & 3 & 0 & -1 \\ 0 & 3 & -3 & 3 & -12 \\ 3 & 1 & 5 & 0 & 1 \end{pmatrix} \xrightarrow{(1)} \begin{pmatrix} 1 & 0 & 2 & 0 & 2 \\ 2 & 1 & 3 & 0 & -1 \\ 0 & 1 & -1 & 1 & -4 \\ 3 & 1 & 5 & 0 & 1 \end{pmatrix} \xrightarrow{(2)} \begin{pmatrix} 1 & 0 & 2 & 0 & 2 \\ 0 & 1 & -1 & 0 & -5 \\ 0 & 1 & -1 & 1 & -4 \\ 3 & 1 & 5 & 0 & 1 \end{pmatrix}$$

$$\xrightarrow{(3)} \begin{pmatrix} 1 & 0 & 2 & 0 & 2 \\ 0 & 1 & -1 & 0 & -5 \\ 0 & 0 & 0 & 1 & 1 \\ 0 & 1 & -1 & 0 & -5 \end{pmatrix} \xrightarrow{(4)} \begin{pmatrix} 1 & 0 & 2 & 0 & 2 \\ 0 & 1 & -1 & 0 & -5 \\ 0 & 0 & 0 & 1 & 1 \\ 0 & 0 & 0 & 0 & 0 \end{pmatrix}$$

Les transformations élémentaires appliquées sont

$$(1) L_3 \rightarrow \frac{1}{3}L_3 \quad (2) L_2 \rightarrow L_2 - 2L_1 \quad (3) L_4 \rightarrow L_4 - 3L_1 \quad (4) L_4 \rightarrow L_4 - L_2.$$

Le système initial est donc équivalent au système échelonné suivant de 3 équations à 4 inconnues :

$$\begin{cases} x + 2z = 2 \\ y - z = -5 \\ t = 1 \end{cases}$$

La solution générale est facile à trouver. Elle s'écrit

$$\begin{pmatrix} x \\ y \\ z \\ t \end{pmatrix} = \begin{pmatrix} 2 - 2z \\ z - 5 \\ z \\ 1 \end{pmatrix} = \begin{pmatrix} 2 \\ -5 \\ 0 \\ 1 \end{pmatrix} + \lambda \begin{pmatrix} -2 \\ 1 \\ 1 \\ 0 \end{pmatrix}, \quad (\lambda \in K)$$

Remarque. La méthode d'échelonnage nous donne une façon simple de savoir si un élément $b \in K^m$ appartient à l'image de la matrice A . C'est en effet le cas si et seulement si le système linéaire $AX = b$ est compatible, ce qui se traduit par la condition $\text{rang}(A|b) = \text{rang}(A)$. Il suffit donc de vérifier que les formes échelonnées de ces deux matrices ont le même nombre de lignes non nulles (ou le même nombre de pivots).

6.4 Matrices élémentaires

Les transformations élémentaires de l'algorithme de Gauss-Jordan peuvent s'appliquer à toute matrice, en particulier à la matrice identité.

Définition. On appelle *matrice élémentaire d'ordre m et de type I, II ou III* toute matrice qui s'obtient en appliquant une transformation élémentaire de type I, II ou respectivement III sur les lignes de la matrice identité $\mathbf{I}_m$. Par exemple, les matrices

$$P_{(2,3)} = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix}, \quad D_{(4)}(\lambda) = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & \lambda \end{pmatrix} \quad \text{et} \quad L_{(3,1)}(\lambda) = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ \lambda & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix}$$

sont respectivement des matrices d'ordre 4 de type I (échanger les lignes 2 et 3 pour $P_{(2,3)}$), de type II (multiplier la dernière ligne par $\lambda \neq 0$ pour $D_{(4)}(\lambda)$) et de type III (ajouter λ fois la première ligne à la troisième ligne pour $L_{(3,1)}(\lambda)$).

D'une manière générale, on définit dans $M_m(K)$ les matrices suivantes :

- La matrice élémentaire de type I est $P_{(r,s)} = (p_{ij})$ avec $p_{ij} = \delta_{ij}$ si $i \notin \{r, s\}$ et $p_{rs} = p_{sr} = 1$ (et les autres p_{ij} sont nuls). C'est la matrice obtenue à partir de $\mathbf{I}_m$ en échangeant les lignes r et s .
- La matrice élémentaire de type II est $D_{(r)}(\lambda) = (d_{ij})$ (où $\lambda \neq 0$) avec $d_{ij} = \delta_{ij}$ si $i \neq r$ et $d_{rr} = \lambda$. On peut aussi écrire

$$D_{(r)}(\lambda) = \mathbf{I}_m + (\lambda - 1)E_{rr} = \text{Diag}(1, 1, \dots, \lambda, \dots, 1),$$

c'est la matrice obtenue en multipliant la ligne r de l'identité par le scalaire λ .

- La matrice élémentaire de type III est $L_{(r,s)}(\lambda) = \mathbf{I}_m + \lambda E_{rs}$ est la matrice obtenue en ajoutant à la ligne r de la matrice identité λ fois la ligne s (avec $r \neq s$).

Les indices r et s ainsi que le scalaire λ s'appellent les *paramètres* de la matrice élémentaire.

On a alors les propriétés suivantes :

Théorème 6.4.1. (A) Toute matrice élémentaire est inversible et son inverse est une matrice élémentaire de même type.

(B) Si $A \in M_{m \times n}(K)$ est une matrice à m lignes et n colonnes, alors multiplier A à gauche par une matrice élémentaire a le même effet qu'appliquer à A la transformation de même type et de mêmes paramètres à la matrice A .

(C) Pour toute matrice $A \in M_{m \times n}(K)$, il existe une matrice inversible $Q \in M_m(K)$ telle que

i.) Q est le produit d'un nombre fini de matrices élémentaires.

ii.) $A' = Q \cdot A$ est de forme échelonnée réduite.

Preuve. (A) On vérifie facilement que

$$(P_{(r,s)})^{-1} = P_{(r,s)}, \quad (D_{(r)}(\lambda))^{-1} = D_{(r)}\left(\frac{1}{\lambda}\right) \quad \text{et} \quad (L_{(r,s)}(\lambda))^{-1} = L_{(r,s)}(-\lambda)$$

(B) Se voit en examinant chaque cas.

(C) Il s'agit d'une reformulation de l'algorithme de Gauss-Jordan.

□

Le point de vue matriciel sur l'algorithme de Gauss-Jordan nous dit que pour résoudre $AX = B$, on se ramène au système équivalent $(QA)X = QB$ où Q est un produit de matrices élémentaires et QA est échelonnée. Les deux systèmes $AX = B$ et $A'X = B'$ où $A' = QA$ et $B' = QB$ ont les mêmes solutions car Q est inversible.

6.5 Systèmes matriciels et inversion d'une matrice par la méthode de Gauss-Jordan

On a vu que la méthode de Gauss-Jordan permet de résoudre une équation matricielle $AX = B$ où X et B sont des matrices-colonnes. La méthode s'applique aussi lorsque A et B sont des matrices plus générales.

Exemple. Supposons que l'on désire résoudre le système suivant dont l'inconnue et le second membre sont des matrices de type 3×2 :

$$\begin{pmatrix} 1 & 2 & 2 \\ 2 & 1 & 3 \\ 0 & 7 & 3 \end{pmatrix} \begin{pmatrix} x_1 & y_1 \\ x_2 & y_2 \\ x_3 & y_3 \end{pmatrix} = \begin{pmatrix} 10 & 4 \\ 2 & 14 \\ 4 & 0 \end{pmatrix}$$

La matrice augmentée du système est la 3×5 matrice

$$(A|B) = \left(\begin{array}{ccc|cc} 1 & 2 & 2 & 10 & 4 \\ 2 & 1 & 3 & 2 & 14 \\ 0 & 7 & 3 & 4 & 0 \end{array} \right).$$

La méthode de Gauss-Jordan nous permet de trouver la forme échelonnée réduite de cette matrice, on obtient :

$$\left(\begin{array}{ccc|cc} 1 & 0 & 0 & 74 & -20 \\ 0 & 1 & 0 & 25 & -9 \\ 0 & 0 & 1 & -57 & 21 \end{array} \right).$$

Le système admet donc pour unique solution

$$X = \begin{pmatrix} 74 & -20 \\ 25 & -9 \\ -57 & 21 \end{pmatrix}.$$

Supposons maintenant que la matrice $A \in M_n(K)$ est une matrice carrée inversible, alors sa forme échelonnée réduite est la matrice identité $\mathbf{I}_n$. Par conséquent, si on considère le système linéaire

$$AX = \mathbf{I}_n,$$

avec $X \in M_n(K)$, alors l'algorithme de Gauss-Jordan revient à multiplier par une matrice Q , qui est produit de matrices élémentaires, pour obtenir le système équivalent

$$QAX = Q\mathbf{I}_n.$$

Mais puisque $QA = \mathbf{I}_n$, on a $Q = A^{-1}$. On a donc le résultat suivant :

Proposition 6.5.1. *Si A est une $n \times n$ matrice inversible, alors l'algorithme de Gauss-Jordan appliqué à la matrice augmentée $(A|\mathbf{I}_n)$ produit une matrice $(\mathbf{I}_n|Q)$ où $Q = A^{-1}$.*

Exemple 1. Pour inverser la matrice $\begin{pmatrix} 2 & 3 \\ 1 & 1 \end{pmatrix}$, on peut appliquer la méthode d'échelonnage à la matrice augmentée

$$\left(\begin{array}{cc|cc} 2 & 3 & 1 & 0 \\ 1 & 1 & 0 & 1 \end{array} \right)$$

on a

$$\left(\begin{array}{cccc|cccc} 2 & 3 & 1 & 0 & 1 & 0 & 0 & 0 \\ 1 & 1 & 0 & 1 & 0 & 1 & 0 & 0 \end{array} \right) \rightarrow \left(\begin{array}{cccc|cccc} 1 & 1 & 0 & 1 & 1 & 0 & 0 & 0 \\ 2 & 3 & 1 & 0 & 0 & 1 & 0 & 0 \end{array} \right) \rightarrow \left(\begin{array}{cccc|cccc} 1 & 1 & 0 & 1 & 1 & 0 & 0 & 0 \\ 0 & 1 & 1 & -2 & 0 & 1 & 0 & 0 \end{array} \right) \rightarrow \left(\begin{array}{cccc|cccc} 1 & 0 & -1 & 3 & 1 & 0 & 0 & 0 \\ 0 & 1 & 1 & -2 & 0 & 1 & 0 & 0 \end{array} \right)$$

Donc

$$\begin{pmatrix} 2 & 3 \\ 1 & 1 \end{pmatrix}^{-1} = \begin{pmatrix} -1 & 3 \\ 1 & -2 \end{pmatrix}.$$

Exemple 2. Pour inverser la matrice $A = \begin{pmatrix} 1 & 0 & 3 \\ -2 & 1 & 0 \\ 1 & 0 & 4 \end{pmatrix}$, on échelonne la matrice augmentée $(A|\mathbf{I}_3)$ de la façon suivante :

$$\begin{aligned} \left(\begin{array}{cccccc|cccc} 1 & 0 & 3 & 1 & 0 & 0 & 1 & 0 & 0 & 0 \\ -2 & 1 & 0 & 0 & 1 & 0 & 0 & 1 & 0 & 0 \\ 1 & 0 & 4 & 0 & 0 & 1 & 0 & 0 & 1 & 0 \end{array} \right) &\xrightarrow{(1)} \left(\begin{array}{cccccc|cccc} 1 & 0 & 3 & 1 & 0 & 0 & 1 & 0 & 0 & 0 \\ -2 & 1 & 0 & 0 & 1 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & -1 & 0 & 1 & 0 & 0 & 1 & 0 \end{array} \right) &\xrightarrow{(2)} \left(\begin{array}{cccccc|cccc} 1 & 0 & 0 & 4 & 0 & -3 & 1 & 0 & 0 & 0 \\ -2 & 1 & 0 & 0 & 1 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & -1 & 0 & 1 & 0 & 0 & 1 & 0 \end{array} \right) \\ &\xrightarrow{(3)} \left(\begin{array}{cccccc|cccc} 1 & 0 & 0 & 4 & 0 & -3 & 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 8 & 1 & -6 & 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & -1 & 0 & 1 & 0 & 0 & 1 & 0 \end{array} \right) \end{aligned}$$

(Les étapes sont : (1) on soustrait la première ligne de la troisième, (2) on soustrait le triple de la troisième ligne de la première, (3) on ajoute le double de la première ligne à la seconde).

On en conclut que

$$A^{-1} = \begin{pmatrix} 4 & 0 & -3 \\ 8 & 1 & -6 \\ -1 & 0 & 1 \end{pmatrix}$$

Chapitre 7

Déterminants

7.1 Déterminants des 2×2 matrices.

A chaque matrice carrée A est associé un scalaire qui s'appelle son *déterminant* et se note $\det(A)$ ou parfois $|A|$. Nous discutons d'abord le cas des matrices de taille 2×2 .

Définition Le *déterminant* d'une 2×2 matrice est défini par

$$\det \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix} = \begin{vmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{vmatrix} := a_{11}a_{22} - a_{12}a_{21}.$$

Théorème 7.1.1. *Le déterminant des 2×2 matrices vérifie les propriétés suivantes :*

- a) $\det(\lambda A) = \lambda^2 \det(A)$.
- b) *Le déterminant d'un produit de deux matrices est égal au produit des déterminants :*

$$\det(A \cdot B) = \det(A) \cdot \det(B).$$

- c) *La 2×2 matrice A est inversible si et seulement si $\det(A) \neq 0$. Dans ce cas on a*

$$\begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix}^{-1} = \frac{1}{\det(A)} \begin{pmatrix} a_{22} & -a_{12} \\ -a_{21} & a_{11} \end{pmatrix}.$$

- d) *Deux vecteur-colonnes de K^2 sont linéairement indépendants si et seulement si leurs composantes forment une matrice de déterminant non nul :*

$$\begin{pmatrix} x_1 \\ x_2 \end{pmatrix} \text{ et } \begin{pmatrix} y_1 \\ y_2 \end{pmatrix} \text{ sont linéairement indépendants } \Leftrightarrow \det \begin{pmatrix} x_1 & y_1 \\ x_2 & y_2 \end{pmatrix} \neq 0.$$

Preuve. La propriété (a) suit immédiatement de la définition. Pour (b), il suffit de calculer et simplifier :

$$\begin{aligned} \det \left(\begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix} \cdot \begin{pmatrix} b_{11} & b_{12} \\ b_{21} & b_{22} \end{pmatrix} \right) &= \det \begin{pmatrix} (a_{11}b_{11} + a_{12}b_{21}) & (a_{11}b_{12} + a_{12}b_{22}) \\ (a_{21}b_{11} + a_{22}b_{21}) & (a_{21}b_{12} + a_{22}b_{22}) \end{pmatrix} \\ &= (a_{11}b_{11} + a_{12}b_{21})(a_{21}b_{12} + a_{22}b_{22}) - (a_{11}b_{12} + a_{12}b_{22})(a_{21}b_{11} + a_{22}b_{21}) \\ &= (a_{11}a_{22} - a_{12}a_{21})(b_{11}b_{22} - b_{12}b_{21}) \\ &= \det \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix} \det \begin{pmatrix} b_{11} & b_{12} \\ b_{21} & b_{22} \end{pmatrix}. \end{aligned}$$

L'affirmation (c) a été démontrée à la Proposition 5.8.3.

Pour prouver (d) on observe que les vecteurs $\begin{pmatrix} x_1 \\ x_2 \end{pmatrix}$ et $\begin{pmatrix} y_1 \\ y_2 \end{pmatrix}$ engendrent l'image de

$$A = \begin{pmatrix} x_1 & y_1 \\ x_2 & y_2 \end{pmatrix}.$$

On a donc $\det(A) \neq 0 \Leftrightarrow A : K^2 \rightarrow K^2$ est surjective $\Leftrightarrow \begin{pmatrix} x_1 \\ x_2 \end{pmatrix}$ et $\begin{pmatrix} y_1 \\ y_2 \end{pmatrix}$ sont linéairement indépendants. □

7.2 Déterminants des 3×3 matrices.

Le déterminant d'une matrice de taille 3×3 est défini par

$$\det \begin{pmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{pmatrix} = a_{11}a_{22}a_{33} + a_{21}a_{32}a_{13} + a_{31}a_{12}a_{23} - a_{11}a_{32}a_{23} - a_{21}a_{12}a_{33} - a_{31}a_{22}a_{13}.$$

Il est parfois commode de voir le déterminant comme une fonction $\det : K^3 \times K^3 \times K^3 \rightarrow K$ des 3 vecteurs colonne de K^3 formant la matrice $A \in M_3(K)$. Il prend alors la forme

$$\det(X, Y, Z) = \det \begin{pmatrix} x_1 & y_1 & z_1 \\ x_2 & y_2 & z_2 \\ x_3 & y_3 & z_3 \end{pmatrix} = x_1y_2z_3 + x_2y_3z_1 + x_3y_1z_2 - x_1y_3z_2 - x_2y_1z_3 - x_3y_2z_1.$$

où $X = \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix}$ $Y = \begin{pmatrix} y_1 \\ y_2 \\ y_3 \end{pmatrix}$ et $Z = \begin{pmatrix} z_1 \\ z_2 \\ z_3 \end{pmatrix}$.

On note aussi $|A|$ pour le déterminant d'une matrice. Il est alors facile de voir que le déterminant d'une matrice de $M_3(K)$ s'exprime de la manière suivante comme une combinaison linéaire de trois déterminants 2×2 :

$$\begin{vmatrix} x_1 & y_1 & z_1 \\ x_2 & y_2 & z_2 \\ x_3 & y_3 & z_3 \end{vmatrix} = x_1 \cdot \begin{vmatrix} y_2 & z_2 \\ y_3 & z_3 \end{vmatrix} - x_2 \cdot \begin{vmatrix} y_1 & z_1 \\ y_3 & z_3 \end{vmatrix} + x_3 \cdot \begin{vmatrix} y_1 & z_1 \\ y_2 & z_2 \end{vmatrix}$$

Par exemple

$$\begin{vmatrix} 1 & 3 & 4 \\ 0 & 1 & 3 \\ 5 & 3 & -1 \end{vmatrix} = 1 \cdot \begin{vmatrix} 1 & 3 \\ 3 & -1 \end{vmatrix} - 0 \cdot \begin{vmatrix} 3 & 4 \\ 3 & -1 \end{vmatrix} + 5 \cdot \begin{vmatrix} 3 & 4 \\ 1 & 3 \end{vmatrix} \\ = 1 \cdot (-1 - 9) + 0 \cdot (-3 - 12) + 5 \cdot (9 - 4) \\ = 15.$$

7.3 Définition générale du déterminant et premières propriétés

Dans une matrice carrée $A = (a_{ij}) \in M_n(K)$, il y a $n!$ façons de former un monôme du type

$$a_{i_1 1} a_{i_2 2} \cdots a_{i_n n}$$

qui contient un et un seul coefficient de chaque colonne et de chaque ligne, c'est à dire que $i_1, i_2, \dots, i_n$ est une permutation de $1, 2, \dots, n$. On obtient le déterminant de la matrice A en multipliant ce monôme par la signature de la permutation associée et en sommant tous les termes obtenus.

Définition. Le *déterminant* de la matrice $A = (a_{ij}) \in M_n(K)$ est le scalaire défini par

$$\det(A) = |A| = \sum_{\sigma \in \mathcal{S}_n} \text{sgn}(\sigma) a_{\sigma(1)1} a_{\sigma(2)2} \cdots a_{\sigma(n)n}, \quad (7.1)$$

où $\mathcal{S}_n$ est le groupe symétrique des permutations de $\{1, 2, \dots, n\}$ et $\text{sgn}(\sigma) = \pm 1$ est la signature de la permutation $\sigma \in \mathcal{S}_n$.

Remarques.

1. La formule (7.1) s'appelle la *formule de Leibniz* pour le déterminant.
2. Il est parfois commode (et fréquent parmi les physiciens) d'écrire le déterminant sous la forme

$$\det(A) = \sum_{i_1=1}^n \sum_{i_2=1}^n \cdots \sum_{i_n=1}^n \varepsilon_{i_1, i_2, \dots, i_n} a_{i_1 1} a_{i_2 2} \cdots a_{i_n n}, \quad (7.2)$$

où $\varepsilon_{i_1, i_2, \dots, i_n} \in \{0, 1, -1\}$ est le *symbole de Levi-Civita* défini par

$$\varepsilon_{i_1, i_2, \dots, i_n} = \begin{cases} +1 & \text{si } i_1, i_2, \dots, i_n \text{ est une permutation paire de } 1, 2, \dots, n, \\ -1 & \text{si } i_1, i_2, \dots, i_n \text{ est une permutation impaire de } 1, 2, \dots, n, \\ 0 & \text{si } i_1, i_2, \dots, i_n \text{ n'est pas une permutation de } 1, 2, \dots, n. \end{cases}$$

3. Remarquer que la somme (7.2) possède n^n termes, mais seulement $n!$ parmi ces termes sont non nuls (par exemple si $n = 5$, la somme (7.2) contient $3125 = 5^5$ termes dont seulement $120 = 5!$ sont non nuls).
4. On peut aussi écrire (7.2) sous la forme

$$\det(A) = \sum_{\phi: \llbracket n \rrbracket \rightarrow \llbracket n \rrbracket} \Omega(\phi) a_{\phi(1)1} a_{\phi(2)2} \cdots a_{\phi(n)n},$$

où la somme est étendue à toutes les applications $\phi: \llbracket n \rrbracket \rightarrow \llbracket n \rrbracket$ et $\Omega(\phi) \in \{0, +1, -1\}$ a été défini au §3.2.1. La seule différence entre le symbole de Levi-Civita et $\Omega(\phi)$ est une différence de notations.

5. Lorsque le corps K est de caractéristique 2 (c'est-à-dire $1 + 1 = 0$), la formule de Leibniz garde son sens, mais la signature peut-être omise dans la formule de Leibniz (car $-1 = +1$). Ce cas ne sera pas considéré dans ce cours.

Un premier résultat général sur les déterminants est le suivant :

Proposition 7.3.1. *Soit A une $n \times n$ matrice dont tous les coefficients de la dernière colonne sont nuls sauf le dernier. Alors le déterminant de A est égale au produit du coefficient a_{nn} avec le déterminant de la $(n-1) \times (n-1)$ matrice A' obtenue à partir de A en supprimant la dernière ligne et la dernière colonne de A .*

En écrivant pour simplifier $n' = n - 1$ on a donc

$$\det \begin{pmatrix} a_{11} & \cdots & a_{1n'} & 0 \\ \vdots & \ddots & \vdots & \vdots \\ a_{n',1} & \cdots & a_{n',n'} & 0 \\ a_{n,1} & \cdots & a_{n,n'} & a_{n,n} \end{pmatrix} = a_{nn} \cdot \det \begin{pmatrix} a_{11} & \cdots & a_{1n'} \\ \vdots & \ddots & \vdots \\ a_{n',1} & \cdots & a_{n',n'} \end{pmatrix}$$

Preuve. En examinant la formule de Leibniz on voit que pour une telle matrice A , chaque terme de la somme (7.1) tel que $\sigma(n) \neq n$ est nul. En conservant la notation $n' = n - 1$ on a donc

$$\begin{aligned} \det(A) &= \sum_{\sigma \in \mathcal{S}_n} \text{sgn}(\sigma) a_{\sigma(1)1} a_{\sigma(2)2} \cdots a_{\sigma(n')n'} a_{\sigma(n)n} \\ &= a_{nn} \cdot \sum_{\sigma \in \mathcal{S}'} \text{sgn}(\sigma) a_{\sigma(1)1} a_{\sigma(2)2} \cdots a_{\sigma(n')n'}, \end{aligned}$$

où $\mathcal{S}'$ est l'ensemble des permutation de $\llbracket n \rrbracket$ qui fixent n , c'est à dire

$$\mathcal{S}' = \{\sigma \in \mathcal{S}_n \mid \sigma(n) = n\}.$$

Il est clair que $\mathcal{S}'$ est un sous-groupe de $\mathcal{S}_n$ et qu'il est isomorphe à $\mathcal{S}_{n-1} = \mathcal{S}_{n'}$. On conclut que

$$\det(A) = a_{nn} \cdot \det(A').$$

□

Corollaire 7.3.2. *Le déterminant d'une matrice triangulaire inférieure (i.e. $a_{ij} = 0$ si $i < j$) est égal au produit des coefficients diagonaux :*

$$\det \begin{pmatrix} a_{11} & 0 & \cdots & 0 \\ * & a_{22} & & 0 \\ \vdots & & \ddots & \vdots \\ * & \cdots & * & a_{nn} \end{pmatrix} = a_{11} a_{22} \cdots a_{nn}$$

Preuve. Par récurrence à partir de la proposition précédente. □

Proposition 7.3.3. *Le déterminant d'une matrice carrée est égal au déterminant de sa transposée*

$$\det(A^\top) = \det(A).$$

Preuve. Rappelons que la signature d'une permutation σ est égale à la signature de la permutation inverse σ^{-1} . L'opération de transposer une matrice échange les lignes et les colonnes, on a donc

$$\begin{aligned} \det(A^\top) &= \sum_{\sigma \in \mathcal{S}_n} \text{sgn}(\sigma) a_{\sigma(1)1}^\top a_{\sigma(2)2}^\top \cdots a_{\sigma(n)n}^\top \\ &= \sum_{\sigma \in \mathcal{S}_n} \text{sgn}(\sigma) a_{1\sigma(1)} a_{2\sigma(2)} \cdots a_{n\sigma(n)} \\ &= \sum_{\sigma \in \mathcal{S}_n} \text{sgn}(\sigma^{-1}) a_{\sigma^{-1}(1)1} a_{\sigma^{-1}(2)2} \cdots a_{\sigma^{-1}(n)n} \\ &= \sum_{\rho \in \mathcal{S}_n} \text{sgn}(\rho) a_{\rho(1)1} a_{\rho(2)2} \cdots a_{\rho(n)n} \quad (\text{on a posé } \rho = \sigma^{-1}) \\ &= \det(A). \end{aligned}$$

□

Une conséquence importante de cette proposition est que toute propriété du déterminant qui s'applique aux colonnes d'une matrice s'applique également aux lignes de cette matrice. En particulier le déterminant d'une matrice triangulaire supérieure (i.e. $a_{ij} = 0$ si $i > j$) est aussi égal au produit des coefficients diagonaux.

Dans la suite de ce paragraphe, il est commode d'associer à un n -tuple de vecteurs colonne

$$(A_1, A_2, \dots, A_n) \in K^n \times K^n \times \dots \times K^n$$

la matrice dont la $j^{\text{ème}}$ colonne est le vecteur A_j (on utilise donc l'isomorphisme naturel $(K^n)^n \rightarrow M_n(K)$). On peut ainsi regarder le déterminant comme une application

$$\det : \underbrace{K^n \times K^n \times \dots \times K^n}_n \rightarrow K.$$

On a alors les propriétés suivantes.

Théorème 7.3.4. (a) *Le déterminant est linéaire par rapport à chaque colonne :*

$$\begin{aligned} \det(A_1, \dots, A_{j-1}, (\lambda A'_j + \mu A''_j), A_{j+1}, \dots, A_n) &= \lambda \det(A_1, \dots, A_{j-1}, A'_j, A_{j+1}, \dots, A_n) \\ &+ \mu \det(A_1, \dots, A_{j-1}, A''_j, A_{j+1}, \dots, A_n) \end{aligned}$$

pour tous $1 \leq j \leq n$ et tous $\lambda, \mu \in K$.

(b) *Pour toute permutation $\rho \in \mathcal{S}_n$ on a*

$$\det(A_{\rho(1)}, \dots, A_{\rho(n)}) = \text{sgn}(\rho) \det(A_1, \dots, A_n).$$

(c) *Si $E_1, E_2, \dots, E_n$ est la base canonique de K^n , alors*

$$\det(E_1, \dots, E_n) = 1.$$

Preuve. (a) Si $A_j = \lambda A'_j + \mu A''_j$, alors pour chaque coefficient de la $j^{\text{ème}}$ colonne on a $a_{ij} = (\lambda a'_{ij} + \mu a''_{ij})$. Donc pour chaque monôme formant le déterminant, on a

$$a_{i_1 1} \cdots a_{i_j j} \cdots a_{i_n n} = \lambda a_{i_1 1} \cdots a'_{i_j j} \cdots a_{i_n n} + \mu a_{i_1 1} \cdots a''_{i_j j} \cdots a_{i_n n}.$$

En faisant la somme (7.1), on obtient la propriété (a).

(b) On a par définition

$$\begin{aligned} \det(A_{\rho(1)}, \dots, A_{\rho(n)}) &= \sum_{\sigma \in \mathcal{S}_n} \text{sgn}(\sigma) a_{\sigma(1)\rho(1)} a_{\sigma(2)\rho(2)} \cdots a_{\sigma(n)\rho(n)} \\ &= \sum_{\tau \in \mathcal{S}_n} \text{sgn}(\tau \rho) a_{\tau(1)1} a_{\tau(2)2} \cdots a_{\tau(n)n} \quad (\text{on a posé } \tau = \sigma \rho^{-1}) \\ &= \text{sgn}(\rho) \sum_{\tau \in \mathcal{S}_n} \text{sgn}(\tau) a_{\tau(1)1} a_{\tau(2)2} \cdots a_{\tau(n)n} \\ &= \text{sgn}(\rho) \det(A). \end{aligned}$$

Car $\text{sgn}(\tau \rho) = \text{sgn}(\tau) \text{sgn}(\rho)$ puisque $\text{sgn} : \mathcal{S} \rightarrow \{+1, -1\}$ est un homomorphisme de groupes.

c) Le déterminant $\det(E_1, \dots, E_n)$ est le déterminant de la matrice identité dont les coefficients sont les symboles de Kronecker δ_{ij} . Tous les termes de la somme (7.1) sont nuls, sauf celui correspondant à la permutation $\sigma = \text{Id}$, et le monôme correspondant est

$$\text{sgn}(\text{Id}) \cdot \delta_{11} \delta_{22} \cdots \delta_{nn} = 1.$$

□

Remarque 7.3.5. Pour simplifier, à nous supposons désormais que le corps K n'est pas de caractéristique 2. Cette condition signifie que dans ce corps. $1 + 1 \neq 0$ (ou de façon équivalente $1 \neq -1$). Cela exclut des corps tels que le corps $\mathbb{F}_2$ à deux éléments ou le corps à 4 éléments vu aux exercices, mais tous les autres corps familiers dans ce cours satisfont cette condition.

Définition. Une application $\theta : K^n \times K^n \times \dots \times K^n \rightarrow K$ vérifiant les conditions (a) et (b) de ce théorème est dite *multilinéaire* et *antisymétrique*.

Corollaire 7.3.6. Si deux colonnes sont identiques, i.e. $A_r = A_s$ avec $1 \leq r < s \leq n$, alors

$$\det(A_1, \dots, A_n) = 0.$$

Preuve. Soit $\rho = (r, s) \in \mathcal{S}_n$ la transposition qui échange r et s , alors la propriété (b) du théorème précédent implique que

$$\begin{aligned} \det(A_1, \dots, A_r, \dots, A_s, \dots, A_r, \dots, A_n) &= \operatorname{sgn}(\rho) \det(A_1, \dots, A_s, \dots, A_r, \dots, A_n) \\ &= -\det(A_1, \dots, A_s, \dots, A_r, \dots, A_n) \\ &= -\det(A_1, \dots, A_r, \dots, A_s, \dots, A_n) \quad (\text{car } A_r = A_s) \end{aligned}$$

donc ce déterminant est nul. □

Corollaire 7.3.7. a.) Si on ajoute à la $j^{\text{ème}}$ colonne de la matrice $A \in M_n(K)$ une combinaison linéaire des autres colonnes, alors le déterminant ne change pas.

b.) Si les colonnes de la matrice $A \in M_n(K)$ sont linéairement dépendantes, alors $\det(A) = 0$.

Preuve. a.) On a

$$\begin{aligned} \det(A_1, \dots, A_{r-1}, (A_r + \sum_{j \neq r} \lambda_j A_j), A_{r+1}, \dots, A_n) &= \det(A_1, \dots, A_r, \dots, A_n) \\ &\quad + \sum_{j \neq r} \lambda_j \det(A_1, \dots, A_{r-1}, A_j, A_{r+1}, \dots, A_n), \end{aligned}$$

on obtient donc l'affirmation voulue par le corollaire précédent car chaque terme de cette dernière somme est nul.

b.) Les colonnes de $A \in M_n(K)$ sont linéairement dépendantes si et seulement si l'une des colonnes est combinaison linéaire des autres colonnes. Supposons par exemple que la colonne A_r est combinaison linéaire des autres colonnes. On a alors par l'affirmation (a) :

$$\det(A_1, \dots, A_r, \dots, A_n) = \det(A_1, \dots, 0, \dots, A_n) = 0.$$

□

Exemple. On a

$$\det \begin{pmatrix} 1 & 0 & x_1 & 0 \\ 0 & 1 & x_2 & 0 \\ 0 & 0 & x_3 & 0 \\ 0 & 0 & x_4 & 1 \end{pmatrix} = \det \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & x_3 & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix} = x_3.$$

Plus généralement si $X = (x_1, x_2, \dots, x_n)^\top$, alors

$$\det(E_1, \dots, E_{j-1}, X, E_{j+1}, \dots, E_n) = x_j,$$

7.4 Théorème fondamental

Théorème 7.4.1 (Théorème fondamental de la théorie des déterminants.). Soit $\theta : K^n \times K^n \times \cdots \times K^n \rightarrow K$ une application multilinéaire et alternée, alors il existe $\gamma \in K$ tel que $\theta = \gamma \cdot \det$, i.e.

$$\theta(A_1, \dots, A_n) = \gamma \cdot \det(A_1, \dots, A_n)$$

pour tous $A_1, \dots, A_n \in K^n$. De plus, $\gamma = \theta(E_1, \dots, E_n)$.

Preuve. On a $A_1 = a_{11}E_1 + \cdots + a_{n1}E_n$, donc par multilinéarité de θ , on a

$$\theta(A_1, \dots, A_n) = \sum_{i=1}^n a_{i1} \theta(E_i, A_2, \dots, A_n).$$

On a aussi $A_2 = a_{12}E_1 + \cdots + a_{n2}E_n$, donc

$$\theta(A_1, \dots, A_n) = \sum_{i=1}^n \sum_{j=1}^n a_{i1} a_{j2} \theta(E_i, E_j, A_3, \dots, A_n).$$

en continuant avec les colonnes $A_3, \dots, A_n$, on obtient

$$\theta(A_1, \dots, A_n) = \sum_{i_1=1}^n \sum_{i_2=1}^n \cdots \sum_{i_n=1}^n a_{i_1 1} a_{i_2 2} \cdots a_{i_n n} \theta(E_{i_1}, E_{i_2}, \dots, E_{i_n}).$$

Or nous avons supposé que θ est alternée, cela entraîne que

$$\theta(E_{i_1}, \dots, E_{i_n}) = \varepsilon_{i_1 \dots i_n} \cdot \theta(E_1, \dots, E_n).$$

Posons $\gamma = \theta(E_1, \dots, E_n)$, alors nous avons

$$\begin{aligned} \theta(A_1, \dots, A_n) &= \gamma \cdot \sum_{i_1=1}^n \sum_{i_2=1}^n \cdots \sum_{i_n=1}^n \varepsilon_{i_1 \dots i_n} a_{i_1 1} a_{i_2 2} \cdots a_{i_n n} \\ &= \gamma \cdot \det(A). \end{aligned}$$

□

Corollaire 7.4.2. Si A et B appartiennent à $M_n(K)$, alors

$$\det(A \cdot B) = \det(A) \cdot \det(B).$$

Preuve. La $j^{\text{ème}}$ colonne de $A \cdot B$ est l'image par A de la $j^{\text{ème}}$ colonne de B

$$(AB)_j = A \cdot B_j = b_{1j}A_1 + b_{2j}A_2 + \cdots + b_{nj}A_n.$$

Si on note $\theta : M_n(K) \rightarrow K$ l'application définie par $\theta(B) = \det(A \cdot B)$, alors θ est une application multilinéaire et alternée. Le théorème précédent montre alors que

$$\det(AB) = \theta(B) = \gamma \cdot \det(B)$$

avec $\gamma = \theta(I_n) = \det(AI_n) = \det(A)$. Donc $\det(A \cdot B) = \det(A) \cdot \det(B)$.

□

Corollaire 7.4.3. Le déterminant d'une matrice carrée est nul si et seulement si ses colonnes sont linéairement dépendantes.

Preuve. On a déjà vu que si les colonnes de $A \in M_n(K)$ sont linéairement dépendantes, alors $\det(A) = 0$. Supposons réciproquement que les colonnes de $A \in M_n(K)$ sont linéairement indépendantes. Alors l'application linéaire $L_A : K^n \rightarrow K^n$ est injective et donc inversible, ce qui implique que la matrice A est inversible. On a alors

$$1 = \det(\mathbf{I}_n) = \det(A \cdot A^{-1}) = \det(A) \cdot \det(A^{-1}),$$

ce qui implique que $\det(A) \neq 0$. □

Observons que la preuve montre aussi que $\det(A^{-1}) = \frac{1}{\det(A)}$.

Théorème 7.4.4. (*Règle de Cramer*¹) Soient $A \in M_n(K)$ et $X, B \in K^n$. Supposons que $AX = B$, alors la $j^{\text{ème}}$ composante de X vérifie

$$\det(A) \cdot x_j = \det(A_1, \dots, A_{j-1}, B, A_{j+1}, \dots, A_n).$$

En particulier si A est inversible, alors cette formule permet de résoudre le système linéaire $AX = B$.

Preuve. Notons S_j la matrice $(E_1, \dots, E_{j-1}, X, E_{j+1}, \dots, E_n)$, alors la condition $AX = B$ est équivalente à

$$A \cdot S_j = (A_1, \dots, A_{j-1}, B, A_{j+1}, \dots, A_n).$$

Donc

$$\det(A) \cdot \det(S_j) = \det(A_1, \dots, A_{j-1}, B, A_{j+1}, \dots, A_n),$$

mais nous avons vu plus haut que $\det S_j = x_j$. □

7.5 Cofacteurs et formule de Laplace

La formule de Laplace est une formule récursive qui permet de calculer un déterminant $n \times n$ à partir de n déterminants de taille $(n-1) \times (n-1)$. Nous aurons besoin de quelques notations.

Définitions et notations. Pour une $n \times n$ matrice $A = (a_{ij}) \in M_n(K)$ et $i, j \in \{1, \dots, n\}$, on note

- $A(i|j)$ la $(n-1) \times (n-1)$ matrice obtenue en supprimant la $i^{\text{ème}}$ ligne et la $j^{\text{ème}}$ colonne de A .
- Le *cofacteur* de a_{ij} est défini par $c_{ij} = c_{ij}(A) = (-1)^{i+j} \det A(i|j)$.
- La *matrice des cofacteurs* de A est la $n \times n$ matrice $\text{Cof}(A) = (c_{ij})$.
- On note $\tilde{A}(i|j)$ la $n \times n$ matrice obtenue à partir de A en remplaçant la $j^{\text{ème}}$ colonne A_j de A par le $i^{\text{ème}}$ vecteur de la base canonique E_i :

$$\tilde{A}(i|j) = (A_1, \dots, A_{j-1}, E_i, A_{j+1}, \dots, A_n).$$

(Si $j = 1$, on a $\tilde{A}(i|1) = (E_i, A_2, \dots, A_n)$).

Proposition 7.5.1. Le déterminant de $\tilde{A}(i|j)$ est égal au cofacteur c_{ij} :

$$c_{ij}(A) = \det(\tilde{A}(i|j)).$$

1. Gabriel Cramer, mathématicien genevois 1704–1752.

Preuve. Commençons par le cas $i = j = n$. La proposition 7.3.1 nous dit que

$$\det \tilde{A}(n|n) = \begin{vmatrix} a_{1,1} & \cdots & a_{1,n-1} & 0 \\ \vdots & \ddots & \vdots & 0 \\ a_{n-1,1} & \cdots & a_{n-1,n-1} & 0 \\ a_{n,1} & \cdots & a_{n,n-1} & 1 \end{vmatrix} = \begin{vmatrix} a_{1,1} & \cdots & a_{1,n-1} \\ \vdots & \ddots & \vdots \\ a_{n-1,1} & \cdots & a_{n-1,n-1} \end{vmatrix}$$

c'est à dire

$$\det \tilde{A}(n|n) = \det A(n|n).$$

Passons au cas général, observons tout d'abord que

$$\begin{aligned} \det(\tilde{A}(i|j)) &= \det(A_1, \dots, A_{j-1}, E_i, A_{j+1}, \dots, A_n) \\ &= (-1)^{n-j} \det(A_1, \dots, A_{j-1}, A_{j+1}, \dots, A_n, E_i) \end{aligned}$$

car on a déplacé $(n-j)$ -fois le vecteur-colonne E_i pour l'amener en dernière position sans changer l'ordre des autres colonnes. En déplaçant maintenant les lignes de la matrice et non les colonnes, on obtient finalement

$$\det \tilde{A}(i|j) = (-1)^{n-i} (-1)^{n-j} \det A(i|j) = (-1)^{i+j} \det A(i|j) = c_{ij}.$$

□

Illustrons l'argument de cette preuve sur un exemple avec $n = 4$:

$$\begin{aligned} \det(\tilde{A}(2|3)) &= \begin{vmatrix} a_{11} & a_{12} & 0 & a_{14} \\ a_{21} & a_{22} & 1 & a_{24} \\ a_{31} & a_{32} & 0 & a_{34} \\ a_{41} & a_{42} & 0 & a_{44} \end{vmatrix} = - \begin{vmatrix} a_{11} & a_{12} & a_{14} & 0 \\ a_{21} & a_{22} & a_{24} & 1 \\ a_{31} & a_{32} & a_{34} & 0 \\ a_{41} & a_{42} & a_{44} & 0 \end{vmatrix} = + \begin{vmatrix} a_{11} & a_{12} & a_{14} & 0 \\ a_{31} & a_{32} & a_{34} & 0 \\ a_{21} & a_{22} & a_{24} & 1 \\ a_{41} & a_{42} & a_{44} & 0 \end{vmatrix} \\ &= - \begin{vmatrix} a_{11} & a_{12} & a_{14} & 0 \\ a_{31} & a_{32} & a_{34} & 0 \\ a_{41} & a_{42} & a_{44} & 0 \\ a_{21} & a_{22} & a_{24} & 1 \end{vmatrix} = - \begin{vmatrix} a_{11} & a_{12} & a_{14} \\ a_{31} & a_{32} & a_{34} \\ a_{41} & a_{42} & a_{434} \end{vmatrix} = (-1)^{2+3} \det(A(2|3)). \end{aligned}$$

Nous pouvons maintenant énoncer et démontrer la formule de Laplace

Théorème 7.5.2 (Formule de Laplace). *Pour toute matrice $A \in M_n(K)$ et tous $1 \leq j, k \leq n$ on a*

$$\sum_{i=1}^n a_{ik} c_{ij} = \delta_{jk} \det(A). \quad (7.3)$$

Preuve. Considérons d'abord le cas $k = j$, la colonne A_j peut s'écrire $A_j = \sum_{i=1}^n a_{ij} E_i$, donc

$$\det(A) = \sum_{i=1}^n a_{ij} \det(A_1, \dots, A_{j-1}, E_i, A_{j+1}, \dots, A_n) = \sum_{i=1}^n a_{ij} \det \tilde{A}(i|j) = \sum_{i=1}^n a_{ij} c_{ij},$$

grâce à la proposition précédente. On a donc prouvé l'équation (7.3) dans le cas $j = k$.

Considérons maintenant le cas $k \neq j$ et notons Z la matrice obtenue à partir de A en remplaçant la $j^{\text{ème}}$ colonne de A par la $k^{\text{ème}}$ colonne, i.e.

$$z_{rs} = \begin{cases} a_{rs} & \text{si } s \neq k, \\ a_{rk} & \text{si } s = j. \end{cases}$$

Alors $\det(Z) = 0$ car cette matrice a deux colonnes identiques. Par construction on a $Z(i|j) = A(i|j)$ et $z_{ij} = a_{ik}$ pour tout i , donc l'équation (7.3) appliquée à la matrice Z avec $i = k$ entraîne que

$$\sum_{i=1}^n a_{ik} c_{ij} = \sum_{i=1}^n (-1)^{i+j} a_{ik} \det A(i|j) = \sum_{i=1}^n (-1)^{i+j} z_{ij} \det Z(i|j) = \det(Z) = 0.$$

□

Remarque. La formule 7.3 peut aussi s'écrire

$$A \cdot \text{Cof}(A)^\top = \det(A) \cdot \mathbf{I}_n$$

Corollaire 7.5.3. On peut aussi calculer le déterminant d'une matrice $A = (a_{ij}) \in M_n(K)$ en développant suivant la $i^{\text{ème}}$ ligne :

$$\det(A) = \sum_{j=1}^n a_{ij} c_{ij} = \sum_{j=1}^n (-1)^{i+j} a_{ij} \det A(i|j).$$

Preuve. C'est une conséquence de la formule 7.3 et du fait que $\det(A^\top) = \det(A)$.

□

Exemple. Le déterminant de la matrice $A = \begin{pmatrix} 1 & 0 & 1 \\ 2 & 1 & 0 \\ 1 & 2 & 1 \end{pmatrix}$ peut se développer selon la première colonne, ce qui donne

$$\begin{aligned} \det(A) &= 1 \cdot \begin{vmatrix} 1 & 0 \\ 2 & 1 \end{vmatrix} - 2 \cdot \begin{vmatrix} 0 & 1 \\ 2 & 1 \end{vmatrix} + 1 \cdot \begin{vmatrix} 0 & 1 \\ 1 & 0 \end{vmatrix} \\ &= 1 \cdot (1) - 2 \cdot (-2) + 1 \cdot (-1) \\ &= 4. \end{aligned}$$

Il peut aussi se développer selon la troisième colonne, ce qui donne

$$\det(A) = 1 \cdot 3 - 0 \cdot 2 + 1 \cdot 1 = 4,$$

ou encore selon la deuxième ligne :

$$\begin{aligned} \det(A) &= -2 \cdot \begin{vmatrix} 0 & 1 \\ 2 & 1 \end{vmatrix} + 1 \cdot \begin{vmatrix} 1 & 1 \\ 1 & 1 \end{vmatrix} - 0 \cdot \begin{vmatrix} 1 & 0 \\ 1 & 2 \end{vmatrix} \\ &= -2 \cdot (-2) + 1 \cdot (0) - 0 \cdot (2) \\ &= 4. \end{aligned}$$

On simplifie le calcul d'un déterminant en choisissant de le développer selon une ligne ou une colonne qui contient un ou plusieurs zéros (si c'est possible). On simplifie encore si un ou plusieurs cofacteurs sont nuls. Dans l'exemple ci-dessus c'est le développement selon la deuxième ligne qui est le plus efficace.

Corollaire 7.5.4. Une matrice carrée A est inversible si et seulement si son déterminant est non nul. Si c'est le cas, l'inverse est donnée par la transposée de la matrice des cofacteurs divisée par le déterminant :

$$A^{-1} = \frac{1}{\det(A)} \text{Cof}(A)^\top.$$

Exemple. Reprenons la matrice

$$A = \begin{pmatrix} 1 & 0 & 1 \\ 2 & 1 & 0 \\ 1 & 2 & 1 \end{pmatrix}.$$

On a $\det(A) = 4$ et la matrice des cofacteurs $c_{ij} = (-1)^{i+j} \det(A_{ij})$ est

$$\text{Cof}(A) = \begin{pmatrix} 1 & -2 & 3 \\ 2 & 0 & -2 \\ -1 & 2 & 1 \end{pmatrix}.$$

Ainsi

$$A^{-1} = \frac{1}{\det(A)} (\text{Cof}(A))^{\top} = \frac{1}{4} \begin{pmatrix} 1 & 2 & -1 \\ -2 & 0 & 2 \\ 3 & -2 & 1 \end{pmatrix}.$$

Le lecteur est invité à vérifier qu'en effet le produit de cette matrice par A donne la matrice identité.

Résumé des propriétés du déterminant :

- 1) Le déterminant définit une application multilinéaire $\det : K^n \times \cdots \times K^n \rightarrow K$.
- 2) On a $\det(A) = \det(A^{\top})$ et $\det(\lambda \cdot A) = \lambda^n \det(A)$ (en particulier $\det(-A) = (-1)^n \det(A)$).
- 3) Le déterminant de la matrice A change de signe si on échange deux lignes ou deux colonnes de A .
- 4) Plus généralement $\det(A_{\sigma(1)}, \dots, A_{\sigma(n)}) = \text{sgn}(\sigma) \det(A_1, \dots, A_n)$ pour toute permutation σ (et on a une formule similaire sur les lignes de la matrice).
- 5) Si une matrice A' est obtenue à partir de A en ajoutant à une colonne de A une combinaison linéaire des autres colonnes de A , alors $\det(A') = \det(A)$.
- 6) Si une matrice A' est obtenue à partir de A en ajoutant à une ligne de A une combinaison linéaire des autres lignes de A , alors $\det(A') = \det(A)$.
- 7) Les vecteurs colonnes $A_1, \dots, A_n$ sont linéairement indépendants s.s. $\det(A_1, \dots, A_n) \neq 0$.
- 8) $\det(A \cdot B) = \det(A) \cdot \det(B)$.
- 9) Le déterminant peut se calculer à partir des cofacteurs en développant selon une ligne ou une colonne :

$$\det(A) = \sum_{j=1}^n a_{ij} c_{ij} = \sum_{i=1}^n a_{ij} c_{ij}.$$

où les c_{ij} sont les cofacteurs de A : $c_{ij} = (-1)^{i+j} \det A(i|j)$ (formule de Laplace).

- 10) La matrice $A \in M_n(K)$ est inversible si et seulement si $\det(A) \neq 0$. Dans ce cas on a $A^{-1} = \frac{1}{\det(A)} (\text{Cof}(A))^{\top}$.

Remarque. Il peut être efficace pour calculer un déterminant de commencer par simplifier la matrice en utilisant les propriétés (5) et (6), puis d'appliquer la formule de Laplace.

Exemple. Soit à calculer le déterminant

$$A = \det \begin{pmatrix} 4 & 2 & 0 & -2 \\ 2 & 1 & 2 & 1 \\ 2 & 3 & 1 & -1 \\ -1 & 0 & 2 & 1 \end{pmatrix}.$$

On obtient une matrice A' de même déterminant en soustrayant la quatrième ligne de la deuxième ligne. On développe ensuite selon les cofacteurs de la deuxième ligne et on obtient

$$\begin{aligned} \det(A) = \det(A') &= \begin{vmatrix} 4 & 2 & 0 & -2 \\ 3 & 1 & 0 & 0 \\ 2 & 3 & 1 & -1 \\ -1 & 0 & 2 & 1 \end{vmatrix} = -3 \cdot \begin{vmatrix} 2 & 0 & -2 \\ 3 & 1 & -1 \\ 0 & 2 & 1 \end{vmatrix} + 1 \cdot \begin{vmatrix} 4 & 0 & -2 \\ 2 & 1 & -1 \\ -1 & 2 & 1 \end{vmatrix} \\ &= -3 \cdot \left(2 \begin{vmatrix} 1 & -1 \\ 2 & 1 \end{vmatrix} - 2 \begin{vmatrix} 3 & 1 \\ 0 & 2 \end{vmatrix} \right) + 1 \cdot \left(4 \begin{vmatrix} 1 & -1 \\ 2 & 1 \end{vmatrix} - 2 \begin{vmatrix} 2 & 1 \\ -1 & 2 \end{vmatrix} \right) \\ &= -3(2 \cdot 3 - 2 \cdot 6) + (4 \cdot 3 - 2 \cdot 5) \\ &= 20. \end{aligned}$$

7.6 Calcul de déterminants par l'algorithme de Gauss-Jordan

Pour calculer un déterminant, on peut utiliser le développement par les cofacteurs. Concrètement cela signifie qu'on remplace le calcul d'un déterminant $n \times n$ par n déterminants $(n-1) \times (n-1)$. Cette méthode n'est pas utilisable pour les matrices de grande taille. On préfère alors utiliser l'algorithme de Gauss-Jordan pour calculer des déterminants.

On a vu au paragraphe 6.4 la définition des matrices élémentaires. Leurs déterminants sont très simples à calculer :

Proposition 7.6.1. *Les déterminants des matrices élémentaires sont donnés par*

$$\det(P_{(r,s)}) = -1, \quad \det(D_{(4)}(\lambda)) = \lambda \quad \text{et} \quad \det(L_{(r,s)}(\lambda)) = 1.$$

La preuve est une simple application des propriétés du déterminant. □

Rappelons que le théorème 6.4.1 nous dit que pour toute matrice $A \in M_{n \times n}(K)$ il existe $Q \in GL_n(K)$ telle que Q est le produit d'un nombre fini de matrices et $A' = Q \cdot A$ est de forme échelonnée. Cela nous conduit à la méthode suivante pour calculer le déterminant d'une matrice $A \in M_n(K)$:

1. Ramener A à une forme échelonnée A' par l'algorithme de Gauss-Jordan.
2. Comme A' est une $n \times n$ échelonnée, c'est une matrice triangulaire.
3. Noter S_k la matrice élémentaire qui correspond à la $k^{\text{ème}}$ étape de l'algorithme, alors

$$S_m S_{m-1} \cdots S_1 \cdot A = A',$$

où m est le nombre d'étapes de la réduction à la forme échelonnée. La matrice Q est donc le produit $Q = S_m S_{m-1} \cdots S_1$.

4. On a donc

$$\det(A) = \det(Q)^{-1} \det(A') = \prod_{k=1}^m \det(S_k)^{-1} \det(A')$$

et chaque $\det(S_k)$ est facile à calculer. Le déterminant $\det(A')$ est aussi facile à calculer car c'est une matrice triangulaire.

Remarquer que si A' contient une ligne nulle, alors $\text{Rang}(A) = \text{Rang}(A') < n$ et donc $\det(A) = 0$.

Exemple. Considérons la matrice

$$A = \begin{pmatrix} 4 & 12 & 4 \\ -7 & -17 & -7 \\ 3 & 9 & 4 \end{pmatrix}.$$

Cette matrice peut s'échelonner en 3 étapes :

$$A = \begin{pmatrix} 4 & 12 & 4 \\ -7 & -17 & -7 \\ 3 & 9 & 4 \end{pmatrix} \rightarrow \begin{pmatrix} 1 & 3 & 1 \\ -7 & -17 & -7 \\ 3 & 9 & 4 \end{pmatrix} \rightarrow \begin{pmatrix} 1 & 3 & 1 \\ 0 & 4 & 0 \\ 3 & 9 & 4 \end{pmatrix} \rightarrow \begin{pmatrix} 1 & 3 & 1 \\ 0 & 4 & 0 \\ 0 & 0 & 1 \end{pmatrix} = A'$$

La première étape est une opération élémentaire de type II, qui est la multiplication de la première ligne par $\frac{1}{4}$, la seconde opération de type III (on ajoute 7 fois la première ligne à la seconde) et la troisième opération est aussi de type III (on soustrait 3 fois la première ligne de la troisième). Matriciellement, cela nous donne :

$$A' = L_{(3,1)}(-3) \cdot L_{(2,1)}(7) \cdot D_{(1)}\left(\frac{1}{4}\right) \cdot A \quad \Rightarrow \quad A = D_{(1)}(4) \cdot L_{(2,1)}(-7) \cdot L_{(3,1)}(3) \cdot A'.$$

Ce qu'on peut vérifier en multipliant les matrices :

$$\begin{pmatrix} 4 & 12 & 4 \\ -7 & -17 & -7 \\ 3 & 9 & 4 \end{pmatrix} = \begin{pmatrix} 4 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix} \cdot \begin{pmatrix} 1 & 0 & 0 \\ -7 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix} \cdot \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 3 & 0 & 1 \end{pmatrix} \cdot \begin{pmatrix} 1 & 3 & 1 \\ 0 & 4 & 0 \\ 0 & 0 & 1 \end{pmatrix}.$$

Le déterminant de A' vaut 4, donc

$$\det(A) = \det(D_{(1)}(4)) \cdot \det(L_{(2,1)}(-7)) \cdot \det(L_{(3,1)}(3)) \cdot \det(A') = 4 \cdot 4 = 16.$$

Chapitre 8

Vecteurs propres et valeurs propres

8.1 Introduction

Un *endomorphisme*¹ f est une application linéaire d'un espace vectoriel V dans lui-même. Le but de ce chapitre et du suivant est d'étudier la structure d'un endomorphisme en général. L'une des idées clés est celle de *sous-espace* invariant. Si f est un endomorphisme de V , on dit que le sous-espace vectoriel $W \subset V$ est un *sous-espace invariant pour f* si l'image de tout élément de W appartient à W . L'application f définit alors, par restriction, un endomorphisme de W , et on peut alors prouver certaines propriétés de f par des arguments de récurrence puisque $\dim(W) < \dim(V)$ (en supposant V de dimension finie).

Nous verrons dans ce chapitre et le suivant, que si V est un espace vectoriel sur le corps $\mathbb{C}$ des complexes, alors on peut décomposer V comme somme directe $V = W_1 \oplus \cdots \oplus W_r$, où chaque W_i est un sous-espace invariant, de sorte que la restriction de f à chaque W_i prend une forme particulièrement simple.

8.2 L'algèbre des endomorphismes

Définitions.

- (a) On appelle *endomorphisme* d'un espace vectoriel V sur le corps K toute application linéaire $f : V \rightarrow V$. On note $\mathcal{L}(V)$ (ou $\text{End}(V)$) l'ensemble des endomorphismes de V . Ainsi $\mathcal{L}(V) = \mathcal{L}(V, V)$.
- (b) On dit que $f \in \mathcal{L}(V)$ est un *automorphisme* de V si f est un endomorphisme bijectif. On note $\text{GL}(V) \subset \mathcal{L}(V)$ l'ensemble des automorphismes de V .
- (c) On dit que deux endomorphismes f_1 et f_2 de V sont *semblables* ou *conjugués* s'il existe un automorphisme $g \in \text{GL}(V)$ tel que $f_2 = g^{-1} \circ f_1 \circ g$.

Rappelons ce qu'est une algèbre :

Définition. Soit K un corps quelconque et $(\mathcal{A}, +, *)$ un anneau. On dit que $\mathcal{A}$ est une *K -algèbre* si le groupe abélien $(\mathcal{A}, +)$ admet aussi une structure de K -espace vectoriel pour laquelle on a

$$\lambda \cdot (a * b) = (\lambda \cdot a) * b = a * (\lambda \cdot b)$$

pour tous $a, b \in \mathcal{A}$ et tout $\lambda \in K$.

1. Un endomorphisme s'appelle aussi un *opérateur*, cette terminologie est habituelle lorsque l'espace vectoriel considéré est un espace de fonctions.

Exemples.

1. Si L est un corps qui contient K comme sous-corps, alors L est une K -algèbre.
2. Les polynômes $K[t]$ forment une algèbre sur le corps K .
3. Les matrices carrées $M_n(K)$ forment une algèbre pour l'addition et la multiplication matricielle.

La proposition suivante récapitule ce que l'on sait déjà sur $\mathcal{L}(V)$:

Proposition 8.2.1. *A) L'ensemble $\mathcal{L}(V)$ des endomorphismes de V est une K -algèbre pour les opérations naturelles de sommes d'applications, de multiplication par des scalaires et de compositions. De plus $\text{GL}(V)$ est un groupe pour la composition.*

B) Si $\dim(V) = n < \infty$ et si $\mathcal{B} = \{b_1, \dots, b_n\}$ est une base de V , alors la correspondance $f \mapsto M_{\mathcal{B}}(f)$ qui à un endomorphisme $f \in \mathcal{L}(V)$ associe sa matrice $M_{\mathcal{B}}(f)$ dans la base $\mathcal{B}$ est un isomorphisme d'algèbres de $\mathcal{L}(V)$ vers $M_n(K)$.

C) L'application $M_{\mathcal{B}} : f \mapsto M_{\mathcal{B}}(f)$ définit aussi un isomorphisme du groupe $\text{GL}(V)$ vers $\text{GL}_n(K)$.

Preuve. Exercice. □

8.3 Sous-espaces invariants

Définitions.

- i.) Un sous-espace vectoriel $W \subset V$ du K -espace vectoriel V est *invariant*, ou *stable*, par rapport à l'endomorphisme $f \in \mathcal{L}(V)$ si $f(W) \subset W$.
- ii.) On dit que l'espace vectoriel V est *simple* pour f s'il ne contient aucun sous-espace vectoriel invariant par f autre que $\{0\}$ et V lui-même.
- iii.) Un sous-espace $W \subset V$ qui est f -invariant est dit *primitif* s'il est simple pour f , i.e. si les seuls sous-espaces $W' \subset W$ qui sont f -invariants sont $W' = \{0\}$ et $W' = W$.

Remarquons que si $W \subset V$ est invariant par rapport à l'endomorphisme $f \in \mathcal{L}(V)$, alors la restriction de f à W définit un endomorphisme $f|_W \in \mathcal{L}(W)$.

- Proposition 8.3.1.** *a.) Soient $f_1, f_2 \in \mathcal{L}(V)$ deux endomorphismes qui commutent, i.e. $f_1 \circ f_2 = f_2 \circ f_1$. Alors $\text{Ker}(f_1)$ et $\text{Im}(f_1)$ sont invariants par f_2 (en particulier $\text{Ker}(f)$ et $\text{Im}(f)$ sont invariants par f pour tout $f \in \mathcal{L}(V)$).*
- b.) Si $W_1 \subset V$ et $W_2 \subset V$ sont invariants par f , alors $W_1 \cap W_2 \subset V$ et $W_1 + W_2 \subset V$ sont invariants par f .*
- c.) Si $W \subset V$ est invariant par rapport à $f \in \mathcal{L}(V)$ et si $g \in \text{GL}(V)$, alors $g^{-1}W$ est invariant par rapport à $g^{-1}fg$.*

Nous laissons la preuve en exercice. □

La proposition suivante nous donne une façon naturelle de produire des sous-espaces invariants.

Proposition 8.3.2. *Soit f un endomorphisme d'un K -espace vectoriel de dimension finie V . Si w_0 est un vecteur non nul de V , alors il existe un plus petit entier $m \geq 1$ tel que $f^m(w_0)$ est combinaison linéaire de $w_0, w_1 = f(w_0), \dots, w_{m-1} = f^{m-1}(w_0)$. De plus, le sous-espace*

$$W = \text{Vec}(\{w_0, w_1, \dots, w_{m-1}\})$$

engendré par les vecteurs w_i est invariant par f .

Preuve. Exercice.

Définition. Un tel sous-espace vectoriel $W \subset V$ est dit *cyclique* pour l'endomorphisme $f \in \mathcal{L}(V)$ et on dit que ce sous-espace cyclique est *engendré* par w_0 .

8.4 Valeurs propres et vecteurs propres

Voyons quelques définitions : Soit $f \in \mathcal{L}(V)$ un endomorphisme du K -espace vectoriel V .

Définition 8.4.1. i.) Un élément $v \in V$ est un *vecteur propre* de f si $v \neq 0$ et s'il existe un scalaire $\lambda \in K$ tel que $f(v) = \lambda v$. De façon équivalente, v est vecteur propre si $\text{Vec}(\{v\}) \subset V$ est un sous-espace vectoriel de dimension 1 qui est invariant par f .

ii.) Le scalaire λ s'appelle la *valeur propre* de f associée à v .

iii.) On dit aussi que le vecteur v est *associé* à la valeur propre λ (ou que c'est un λ -*vecteur propre*).

iv.) Le *spectre*² de l'endomorphisme $f \in \mathcal{L}(V)$ est l'ensemble de ses valeurs propres. On le note $\sigma(f) \subset K$.

v.) L'*espace propre* associé à $\lambda \in \sigma(f)$ est défini par

$$E_\lambda(f) = \{v \in V \mid f(v) = \lambda v\} = \text{Ker}(\lambda \text{Id}_V - f).$$

vi.) On appelle *multiplicité géométrique*³ de la valeur propre $\lambda \in \sigma(f)$ la dimension de l'espace propre $E_\lambda(f)$, on la note $\text{multgeom}_\lambda(f)$.

Remarques.

- 1.) L'espace-propre $E_\lambda(f)$ est clairement un sous-espace vectoriel de V . C'est l'ensemble des vecteurs propres associés à la valeur propre λ auquel on ajoute le vecteur nul (qui n'est pas considéré comme un vecteur propre).
- 2.) Observons de plus que $E_\lambda(f)$ est invariant par f .
- 3.) L'endomorphisme f est injectif si et seulement si $0 \notin \sigma(f)$. Si f n'est pas injectif, alors le noyau de f est constitué de l'ensemble des vecteurs propres associés à la valeur propre 0, auquel on ajoute le vecteur nul.
- 4.) Les vecteurs propres et les valeurs propres d'une matrices $A \in M_n(K)$ sont par définition les vecteurs propres et les valeurs propres de l'endomorphisme associé $L_A \in \mathcal{L}(K^n)$.

Exemples 1. Le scalaire $0 \in K$ est une valeur propre de l'endomorphisme $f \in \mathcal{L}(V)$ si et seulement si le noyau de l'endomorphisme f est non nul. Dans ce cas tout vecteur non nul de $\text{Ker}(f)$ est un vecteur propre.

2. Si A est une matrice diagonale, $A = \text{Diag}(\alpha_1, \alpha_2, \dots, \alpha_n)$, alors les vecteurs de la base canoniques sont des vecteurs propres car $Ae_i = \alpha_i e_i$, et le spectre est donné par les coefficients diagonaux : $\sigma(A) = \{\alpha_1, \alpha_2, \dots, \alpha_n\}$.

3. On considère la matrice

$$A := \begin{pmatrix} 1 & 3 & 1 \\ 0 & 2 & 0 \\ 0 & 3 & 2 \end{pmatrix}$$

On voit clairement que $\lambda = 1$ est une valeur propre; un vecteur propre associé est $(1, 0, 0)$. Une autre valeur propre est $\lambda = 2$ et un vecteur propre associé est $(1, 0, 1)$. On peut démontrer qu'il n'y a pas d'autre valeur propre pour la matrice A , donc $\sigma(A) = \{1, 2\}$.

4. La matrice $J \in M_2(\mathbb{R})$ définie par

$$J := \begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix}$$

2. Ce que nous appelons ici le *spectre* de f s'appelle parfois le *spectre ponctuel* dans le cadre de l'analyse fonctionnelle. Dans le cas des espaces vectoriels de dimension finie, les deux terminologies sont équivalentes.

3. On verra plus loin une autre notion de multiplicité, appelée *multiplicité algébrique*, de la valeur propre, (voir définition 8.5.6).

n'admet aucune valeur propre dans $\mathbb{R}$. En effet, supposons que

$$\lambda \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix} \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} -y \\ x \end{pmatrix},$$

alors $y = -\lambda x$ et $x = \lambda y$. Donc $x = -\lambda^2 x$ et $y = -\lambda^2 y$. Donc si $x \neq 0$ ou $y \neq 0$, alors on doit avoir $\lambda^2 = -1$, ce qui est impossible dans $\mathbb{R}$.

Par contre la même matrice vue comme élément de $M_2(\mathbb{C})$ admet les valeurs propres $+i$ et $-i$, avec les vecteurs propres correspondants $(1, -i)$ et $(1, +i)$.

5. La dérivation $\frac{d}{dx}$ définit un opérateur sur l'espace vectoriel des fonctions réelles infiniment différentiables $C^\infty(\mathbb{R}, \mathbb{R})$. Tout nombre réel $\lambda \in \mathbb{R}$ est valeur propre pour $\frac{d}{dx}$, un vecteur propre associé est la fonction $x \mapsto e^{\lambda x}$ car

$$\frac{d}{dx} e^{\lambda x} = \lambda e^{\lambda x},$$

on a donc $\sigma(\frac{d}{dx}) = \mathbb{R}$.

Deux vecteurs propres associés à des valeurs propres distinctes sont linéairement indépendants, plus généralement, on a :

Proposition 8.4.2. *Soit V un K -espace vectoriel et supposons que $v_1, v_2, \dots, v_k \in V$ sont des vecteurs propres de l'endomorphisme $f \in \mathcal{L}(V)$ associés aux valeurs propres $\lambda_1, \dots, \lambda_k \in \sigma(f)$. Si $\lambda_i \neq \lambda_j$ pour $i \neq j$, alors $\{v_1, v_2, \dots, v_k\}$ est une famille libre de V .*

Preuve. La preuve se fait par récurrence sur k . Si $k = 1$, alors $\{v_1\}$ est une famille libre car un vecteur propre est non nul par définition.

On suppose la proposition démontrée pour une famille de $(k-1)$ vecteurs propres. Supposons que les scalaires $\alpha_1, \dots, \alpha_k \in K$ vérifient $\alpha_1 v_1 + \dots + \alpha_k v_k = 0$ et définissons un endomorphisme $\varphi \in \mathcal{L}(V)$ par $\varphi(v) = f(v) - \lambda_k v$, alors

$$\begin{aligned} 0 &= \varphi(\alpha_1 v_1 + \dots + \alpha_k v_k) \\ &= (\alpha_1 f(v_1) + \dots + \alpha_k f(v_k)) - \lambda_k (\alpha_1 v_1 + \dots + \alpha_k v_k) \\ &= \alpha_1 (\lambda_1 - \lambda_k) v_1 + \dots + \alpha_{k-1} (\lambda_{k-1} - \lambda_k) v_{k-1} + \alpha_k (\lambda_k - \lambda_k) v_k \\ &= \alpha_1 (\lambda_1 - \lambda_k) v_1 + \dots + \alpha_{k-1} (\lambda_{k-1} - \lambda_k) v_{k-1}. \end{aligned}$$

Mais par hypothèse de récurrence, les vecteurs $v_1, v_2, \dots, v_{k-1}$ sont linéairement indépendants, et par hypothèse $(\lambda_j - \lambda_k) \neq 0$ pour tous $1 \leq j \leq (k-1)$, on conclut que $\alpha_j = 0$ pour tous $1 \leq j \leq (k-1)$. Mais alors on a aussi $\alpha_k = 0$ et donc $v_1, v_2, \dots, v_{k-1}, v_k$ sont linéairement indépendants. □

Exemple. Si α, β, γ sont trois nombres réels distincts, alors les fonctions $e^{\alpha x}, e^{\beta x}, e^{\gamma x}$ sont linéairement indépendantes car ces fonctions sont des vecteurs propres de l'opérateur $\frac{d}{dx}$ associés à α, β et γ .

Corollaire 8.4.3. *La somme des multiplicités géométriques de toutes les valeurs propres d'un endomorphisme $f \in \mathcal{L}(V)$ n'excède pas la dimension de cet espace :*

$$\sum_{\lambda \in \sigma(f)} \text{multgeom}_\lambda(f) \leq \dim(V).$$

Preuve. Exercice. □

Définition. Si $A \in M_n(K)$ est une matrice carrée, alors on dit que λ est une *valeur propre* de A si c'est une valeur propre de l'endomorphisme associé $L_A : K^n \rightarrow K^n$, i.e. s'il existe un vecteur-colonne

non-nul $X \in K^n$ tel que $AX = \lambda X$. Toutes les autres définitions ci-dessus (espace propre, multiplicité géométrique etc.) s'étendent au cas des matrices de $M_n(K)$.

Remarque. Pour vérifier qu'un élément $\lambda \in K$ est valeur propre d'une matrice A , il faut déterminer si la matrice $(A - \lambda I_n)$ est inversible ou non. On peut alors trouver la multiplicité géométrique de la valeur propre λ en appliquant la formule du rang :

$$\text{multgeom}_\lambda(A) = \dim \text{Ker}(\lambda I_n - A) = n - \text{rang}(\lambda I_n - A).$$

Pour calculer le rang de $(A - \lambda I_n)$, on peut au besoin utiliser la méthode du pivot (Gauss-Jordan).

8.5 Le polynôme caractéristique

Les valeurs propres peuvent s'obtenir comme racines d'un certain polynôme :

Proposition 8.5.1. *Soit $A \in M_n(K)$. Alors $\det(tI_n - A)$ est un polynôme unitaire⁴ de degré n .*

Définition. Le polynôme $\det(tI_n - A)$ s'appelle le *polynôme caractéristique* de la matrice A et on le note $\chi_A(t)$.

Démonstration. Notons

$$b_{ij}(t) = \delta_{ij}t - a_{ij} \in K[t],$$

pour $1 \leq i, j \leq n$, et notons $B(t) = (b_{ij}(t))$ la matrice correspondante. C'est une matrice à coefficients dans $K[t]$, on peut alors écrire formellement

$$\det(B(t)) = \det(b_{ij}(t)) = \sum_{\sigma \in \mathfrak{S}_n} \text{sign}(\sigma) b_{1\sigma(1)} \cdots b_{n\sigma(n)}.$$

Les b_{ij} sont des polynômes de degré 0 ou 1, donc la somme précédente est une somme de polynômes de degré au plus n ; c'est donc un polynôme de degré au plus n . En fait il n'y a qu'un seul terme qui est de degré maximal, c'est le terme qui correspond à la permutation $\sigma = \text{id}$ et qui est

$$(t - a_{11})(t - a_{22}) \cdots (t - a_{nn}).$$

Ce dernier polynôme est unitaire de degré n . On constate finalement que $\chi_A(t) = \det(B(t))$ s'écrit comme une somme du monôme t^n et de termes de degrés strictement plus petits. Ceci prouve le résultat. $\square$

Remarque 8.5.2. Le terme constant du polynôme caractéristique est $(-1)^n \det A$. En effet le terme constant d'un polynôme est obtenu en remplaçant l'indéterminée par 0. Or

$$\chi_A(0) = \det(-A) = (-1)^n \det A.$$

Exemples. Dans la pratique, un polynôme caractéristique se calcule formellement comme un déterminant, par exemple :

1. Le polynôme caractéristique de la matrice $R_\theta = \begin{pmatrix} \cos(\theta) & -\sin(\theta) \\ \sin(\theta) & \cos(\theta) \end{pmatrix}$ se calcule facilement, on trouve

$$\chi_{R_\theta}(t) = t^2 - 2\cos(\theta)t + 1.$$

4. Un polynôme $p(t) \in K[t]$ est dit *unitaire* si son coefficient dominant vaut 1.

2. Plus généralement, pour une 2×2 matrice quelconque, $A = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$, on a

$$\begin{aligned}\chi_A(t) &= (t-a)(t-d) - bc \\ &= t^2 - (a+d)t + (ad-bc) \\ &= t^2 - \text{Tr}(A)t + \det(A).\end{aligned}$$

3. Si A est une matrice diagonale, $A = \text{Diag}(\alpha_1, \dots, \alpha_n)$, alors

$$\chi_A(t) = \prod_{i=1}^n (t - \alpha_i).$$

4. Si A est une matrice par blocs du type $A = \begin{pmatrix} A_1 & S \\ 0 & A_2 \end{pmatrix}$, alors

$$\chi_A(t) = \chi_{A_1}(t) \cdot \chi_{A_2}(t).$$

Proposition 8.5.3. *Deux matrices semblables ont même polynôme caractéristique.*

Démonstration. Supposons que $A, B \in M_n(K)$ sont deux matrices semblables, alors il existe une matrice inversible $P \in \text{GL}_n(K)$ telle que $B = P^{-1}AP$ et on a

$$\begin{aligned}\chi_B(t) &= \det(tI_n - B) \\ &= \det(tI_n - P^{-1}AP) \\ &= \det(P^{-1}(tI_n - A)P) \\ &= \det(P^{-1}) \det(tI_n - A) \det P \\ &= \det(tI_n - A) \\ &= \chi_A(t).\end{aligned}$$

Le passage de la deuxième ligne à la troisième provient du fait que la matrice tI_n commute avec toute matrice. On a utilisé à la 4ème ligne la multiplicativité du déterminant. $\square$

Conséquence importante : La proposition précédente permet de définir le polynôme caractéristique d'un endomorphisme. En effet, soit $f : V \rightarrow V$ un endomorphisme. On fixe une base quelconque $\mathcal{B}$ de V et on note $A = M_{\mathcal{B}}(f)$ la matrice de f dans cette base, puis on pose

$$\chi_f(t) = \chi_A(t).$$

La proposition précédente montre que le choix de la base $\mathcal{B}$ n'influe pas sur la valeur de $\chi_f(t)$ (la matrice A change mais le polynôme caractéristique ne change pas).

Théorème 8.5.4. *Soit $A \in M_n(K)$. Les racines de $\chi_A(t)$ sont exactement les valeurs propres de A .*

Démonstration. Prenons tout d'abord une valeur propre de A et montrons qu'elle est racine de $\chi_A(t)$. Soit donc $\lambda \in K$ telle qu'il existe $v \neq 0$ avec

$$Av = \lambda v.$$

On a donc $\lambda v - Av = 0$, c'est-à-dire $(\lambda I_n - A)v = 0$. Cela montre que $\text{Ker}(\lambda I_n - A) \neq \{0\}$ et donc $\det(\lambda I_n - A) = 0$.

Réciproquement, on doit prouver que toute racine λ de $\chi_A(t)$ est valeur propre de A . On a donc

$$\det(\lambda I_n - A) = 0,$$

et ceci veut bien dire que la matrice $\lambda I_n - A$ n'est pas inversible. Puisqu'elle n'est pas inversible, c'est donc que son noyau n'est pas réduit à $\{0\}$ (cette dernière implication n'est pas complètement triviale, on rappelle que pour un endomorphisme, il est équivalent d'être inversible et d'avoir un noyau nul). Il existe donc un vecteur v non nul tel que

$$(\lambda I_n - A)v = 0.$$

Ainsi $Av = \lambda v$ et λ est valeur propre. □

Corollaire 8.5.5. a.) Une matrice $A \in M_n(K)$ a au plus n valeurs propres.

b.) Un endomorphisme $f \in \mathcal{L}(V)$ a au plus n valeurs propres, où $n = \dim V$.

Preuve. Le polynôme $\chi_A(t)$ est de degré n , il contient donc au plus n racines. □

Définition 8.5.6. On appelle *multiplicité algébrique* de la valeur propre $\lambda \in \sigma(f)$, et on note $\text{multalg}_\lambda(f)$, la multiplicité de λ en tant que racine du polynôme caractéristique. C'est donc le plus grand entier m tel que

$$\chi_f(t) = (t - \lambda)^m q(t)$$

pour un certain polynôme $q \in K[t]$.

Exemple. On considère la matrice

$$A := \begin{pmatrix} 1 & 3 & 1 \\ 0 & 2 & 0 \\ 0 & 3 & 2 \end{pmatrix}.$$

Le polynôme caractéristique est $\chi_A(t) = (t - 1)(t - 2)^2$. Les valeurs propres sont 1 et 2, et on a $\text{multalg}_1(A) = 1$ et $\text{multalg}_2(A) = 2$. De plus, on a $\text{multgeom}_1(A) = 1$ car $\text{rang}(I_3 - A) = 2$ et $\text{multgeom}_2(A) = 1$ car la matrice

$$(2I_3 - A) = \begin{pmatrix} 1 & -3 & -1 \\ 0 & 0 & 0 \\ 0 & -3 & 0 \end{pmatrix}$$

est de rang 2.

8.6 Endomorphismes et matrices diagonalisables

Définitions. Un endomorphisme $f \in \mathcal{L}(V)$ d'un espace vectoriel de dimension finie est *diagonalisable* s'il existe une base de V pour laquelle la matrice de f est une matrice diagonale.

Une matrice $A \in M_n(K)$ est *diagonalisable sur le corps K* s'il existe une matrice inversible $P \in \text{GL}_n(K)$ telle que $P^{-1}AP$ est une matrice diagonale.

Proposition 8.6.1. Soit V un K -espace vectoriel de dimension n et $f \in \mathcal{L}(V)$ un endomorphisme de V . Alors f est diagonalisable si et seulement s'il existe une base de V formée de vecteurs propres de f .

Preuve. Soit $\mathcal{B} = \{v_1, v_2, \dots, v_n\}$ une base $\mathcal{B}$ de V formée de vecteurs propres. Alors $f(v_j) = \lambda_j v_j$ pour tout j et la matrice de f dans cette base est donc

$$M_{\mathcal{B}}(f) = \text{Diag}(\lambda_1, \dots, \lambda_n) = \begin{pmatrix} \lambda_1 & & \\ & \ddots & \\ & & \lambda_n \end{pmatrix}.$$

Inversement, si la matrice de f dans une base $\mathcal{B} = \{v_1, v_2, \dots, v_n\}$ est diagonale, alors chaque v_j est un vecteur propre de f . □

Corollaire 8.6.2. Si V est un K -espace vectoriel de dimension n et $f \in \mathcal{L}(V)$ un endomorphisme de V ayant n valeurs propres distinctes deux à deux, alors f est diagonalisable.

Preuve. Rappelons que des vecteurs propres associés à des valeurs propres distinctes sont linéairement indépendants (voir proposition 8.4.2). Il en découle que si f admet $n = \dim(V)$ valeurs propres deux-à-deux distinctes, alors les vecteurs propres associés forment une base de V . □

Diagonalisation d'une matrice.

Considérons une matrice $A \in M_n(K)$, et supposons qu'on a trouvé une base propre $\mathcal{B} = \{v_1, \dots, v_n\} \subset K^n$ de cette matrice, i.e. $Av_j = \lambda_j v_j$. Notons $P \in GL_n(K)$ la matrice de transition de la base canonique de K^n vers la base propre $\mathcal{B}$ (c'est donc la matrice dont la $j^{\text{ème}}$ colonne est formée des coordonnées de v_j). Notons aussi $D = \text{Diag}(\lambda_1, \dots, \lambda_n)$, alors on a :

$$AP = PD.$$

Par conséquent :

$$A = P \cdot D \cdot P^{-1} \quad \text{et} \quad D = P^{-1} \cdot A \cdot P.$$

On peut représenter la situation par le diagramme suivant :

$$\begin{array}{ccc} K^n & \xrightarrow{D} & K^n \\ P \downarrow & & \downarrow P \\ K^n & \xrightarrow{A} & K^n \end{array}$$

Définition. La matrice $D = \text{Diag}(\lambda_1, \dots, \lambda_n)$ s'appelle la *forme diagonale* de A et P s'appelle la *matrice modale*. Les formules ci-dessus s'appellent les *formules de diagonalisation* de A .

Pour diagonaliser une matrice $A \in M_n(K)$ on procède donc concrètement comme suit :

1. On calcule son spectre $\{\lambda_1, \dots, \lambda_r\}$.
2. Pour chaque valeur propre λ_i on cherche une base $\mathcal{B}_i$ de E_{λ_i} en trouvant une famille maximale de solutions *linéairement indépendantes* du système linéaire

$$AX = \lambda_i X \quad \Leftrightarrow \quad (A - \lambda_i)X = 0.$$

3. On note $\mathcal{B} = \mathcal{B}_1 \cup \dots \cup \mathcal{B}_r$ la réunion des bases partielles obtenues à l'étape précédente, on vérifie que $\text{Card}(\mathcal{B}) = n$ (ça doit être le cas si la matrice est diagonalisable). On a ainsi obtenu une base propre $\mathcal{B}$ de K^n .
4. Chaque élément de $\mathcal{B}$ est un vecteur-colonne de K^n . La matrice modale est la matrice P dont la $j^{\text{ème}}$ colonne est formée par les coordonnées du $j^{\text{ème}}$ vecteur de $\mathcal{B}$.
5. On note D la matrice diagonale dont les m_1 premiers coefficients diagonaux sont λ_1 , les m_2 coefficients diagonaux suivants sont λ_2 etc.

$$D = \text{Diag}(\underbrace{\lambda_1, \dots, \lambda_1}_{m_1}, \dots, \underbrace{\lambda_r, \dots, \lambda_r}_{m_r})$$

Cette matrice est la forme diagonale de A .

6. On vérifie que $D = P^{-1}AP$, ou si on préfère, que $AP = PD$.

Exemple 1. On veut diagonaliser la matrice

$$F = \begin{pmatrix} 0 & 1 \\ 1 & 1 \end{pmatrix}.$$

Le polynôme caractéristique de F est

$$\chi_F(t) = \det(tI_2 - F) = \det \begin{pmatrix} t & -1 \\ -1 & t-1 \end{pmatrix} = t^2 - t - 1 = (t - \lambda)(t - \mu),$$

où λ et μ sont donnés par

$$\lambda = \frac{1 + \sqrt{5}}{2}, \quad \mu = \frac{1 - \sqrt{5}}{2} = 1 - \lambda$$

(le nombre λ est le *nombre d'or*). Cherchons une base propre, i.e. deux vecteurs propres linéairement indépendants. Pour la valeur propre λ , on pose

$$\begin{pmatrix} 0 & 1 \\ 1 & 1 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} = \lambda \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} \quad \Leftrightarrow \quad \begin{cases} x_2 = \lambda x_1 \\ x_1 + x_2 = \lambda x_2 \end{cases}$$

Une solution est donnée par $x_1 = 1$ et $x_2 = \lambda$ car $1 + \lambda = \lambda^2$. Le vecteur propre cherché est donc $X = \begin{pmatrix} 1 \\ \lambda \end{pmatrix}$. De la même manière, on trouve que $Y = \begin{pmatrix} 1 \\ \mu \end{pmatrix}$ est un vecteur propre pour la valeur propre μ .

Nous obtenons donc la diagonalisation suivante de F :

$$D = P^{-1}FP$$

avec

$$D = \begin{pmatrix} \lambda & 0 \\ 0 & \mu \end{pmatrix} \quad \text{et} \quad P = \begin{pmatrix} 1 & 1 \\ \lambda & \mu \end{pmatrix}.$$

Observons que $\det(P) = \mu - \lambda = -\sqrt{5}$, donc

$$P^{-1} = -\frac{1}{\det(P)} \begin{pmatrix} \mu & -1 \\ -\lambda & 1 \end{pmatrix} = \frac{1}{\sqrt{5}} \begin{pmatrix} -\mu & 1 \\ \lambda & -1 \end{pmatrix}$$

et on vérifie que effectivement

$$P^{-1}FP = \frac{1}{\sqrt{5}} \begin{pmatrix} -\mu & 1 \\ \lambda & -1 \end{pmatrix} \begin{pmatrix} 0 & 1 \\ 1 & 1 \end{pmatrix} \begin{pmatrix} 1 & 1 \\ \lambda & \mu \end{pmatrix} = \begin{pmatrix} \lambda & 0 \\ 0 & \mu \end{pmatrix}$$

Exemple 2. Soit

$$R_\theta = \begin{pmatrix} \cos(\theta) & -\sin(\theta) \\ \sin(\theta) & \cos(\theta) \end{pmatrix}$$

la matrice de rotation d'angle θ dans le plan $\mathbb{R}^2$. Son polynôme caractéristique est

$$\chi_{R_\theta}(t) = t^2 - 2t \cos(\theta) + 1.$$

Ce polynôme est irréductible dans $\mathbb{R}[t]$ si $\cos(\theta) \neq \pm 1$, c'est-à-dire si θ n'est pas un multiple de π . En particulier $\chi_{R_\theta}(t)$ n'est donc pas scindé et R_θ n'est pas diagonalisable dans $M_2(\mathbb{R})$ si $0 < \theta < \pi$. En fait il n'y a aucun vecteur propre dans $\mathbb{R}^2$, ce qui est évident géométriquement pour une rotation.

Mais si on regarde la matrice R_θ dans $M_2(\mathbb{C})$, le polynôme caractéristique se factorise en

$$\chi_{R_\theta}(t) = (t - e^{i\theta})(t - e^{-i\theta}).$$

Il y a donc deux valeurs propres distinctes et par conséquent R_θ est diagonalisable comme matrice à coefficients complexes. On cherche une base propre en cherchant des solutions non nulles de

$$R_\theta X = e^{i\theta} X \quad \text{et} \quad R_\theta Y = e^{-i\theta} Y.$$

On trouve $X = \begin{pmatrix} 1 \\ -i \end{pmatrix}$ et $Y = \begin{pmatrix} 1 \\ i \end{pmatrix}$. La matrice modale est donc $P = \begin{pmatrix} 1 & 1 \\ -i & i \end{pmatrix}$ et on peut vérifier directement que

$$P^{-1} R_\theta P = \begin{pmatrix} e^{i\theta} & 0 \\ 0 & e^{-i\theta} \end{pmatrix}.$$

8.7 Multiplicités des valeurs propres et diagonalisabilité

Dans ce paragraphe, nous précisons les liens entre la multiplicité géométrique et la multiplicité algébrique d'une valeur propre d'un endomorphisme.

Proposition 8.7.1. *Soit f un endomorphisme d'un K espace vectoriel V de dimension finie. Pour toute valeur propre $\lambda \in \sigma(f)$ on a*

$$1 \leq \text{multgeom}_\lambda(f) \leq \text{multalg}_\lambda(f) \leq n = \dim(V).$$

Preuve. Nous devons démontrer trois inégalités. La première inégalité est évidente car par définition pour toute valeur propre λ il existe au moins un élément non nul $v \in E_\lambda(f)$, par conséquent

$$\text{multgeom}_\lambda(f) = \dim(E_\lambda(f)) \geq 1.$$

La dernière inégalité est immédiate car $\chi_f(t)$ est un polynôme de degré n , donc aucune de ses racines ne peut avoir une multiplicité dépassant n .

Notons $s = \text{multgeom}_\lambda(f)$ et démontrons que $s \leq \text{multalg}_\lambda(f)$. Donnons-nous une base $\{v_1, \dots, v_s\} \subset V$ de l'espace propre $E_\lambda(f)$ et complétons cette famille de vecteurs en une base $\mathcal{B}$ de V :

$$\mathcal{B} = \{v_1, \dots, v_s, v_{s+1}, \dots, v_n\}.$$

Observons que par construction $f(v_j) = \lambda v_j$ pour $1 \leq j \leq s$, la matrice de f dans cette base est donc du type

$$M_{\mathcal{B}}(f) = \left(\begin{array}{cccc|cc} \lambda & 0 & \cdots & 0 & * & * \\ 0 & \lambda & \cdots & 0 & * & * \\ & & \ddots & 0 & & \ddots \\ 0 & 0 & & \lambda & * & * \\ 0 & 0 & \cdots & 0 & * & * \\ & & \ddots & & & \ddots \\ 0 & 0 & \cdots & 0 & * & * \end{array} \right) = \left(\begin{array}{c|c} \lambda I_s & B \\ \hline O & C \end{array} \right)$$

et donc

$$\chi_f(t) = (t - \lambda)^s \chi_C(t).$$

or $\chi_C(t) \in K[t]$ est un polynôme de degré $n - s$. Par conséquent $\text{multalg}_\lambda(f) \geq s$. □

Rappelons que l'endomorphisme $f \in \mathcal{L}(V)$ est par définition diagonalisable s'il existe une base $\mathcal{B}$ de V pour laquelle la matrice de f est une matrice diagonale. Le théorème suivant caractérise les endomorphismes diagonalisables.

Théorème 8.7.2. *Soient V un espace vectoriel de dimension finie n sur un corps K et $f \in \mathcal{L}(V)$ un endomorphisme de V . Alors les trois conditions suivantes sont équivalentes :*

a) L'endomorphisme f est diagonalisable.

b) (i) Le polynôme caractéristique de f est scindé sur le corps K et

(ii) La multiplicité géométrique de chaque valeur propre est égale à sa multiplicité algébrique :

$$\text{multgeom}_\lambda(f) = \text{multalg}_\lambda(f) \quad \forall \lambda \in \sigma(f).$$

c) La somme des multiplicités géométriques des valeurs propres de f est égale à la dimension de l'espace vectoriel V :

$$\sum_{\lambda \in \sigma(f)} \text{multgeom}_\lambda(f) = n.$$

Preuve. (a) $\Rightarrow$ (b) Supposons que f est diagonalisable, alors il existe une base $\mathcal{B} = \{v_1, \dots, v_n\}$ de V telle que $f(v_j) = \alpha_j v_j$. Dans cette base la matrice de f est la matrice diagonale $A = \text{diag}(\alpha_1, \dots, \alpha_n)$ et donc le polynôme caractéristique

$$\chi_f(t) = \chi_A(t) = (t - \alpha_1) \cdots (t - \alpha_n) = \prod_{j=1}^n (t - \alpha_j)$$

est scindé, nous avons donc démontré la propriété (i) du point (b). Pour prouver (ii), on doit tenir compte du fait que les α_j ne sont pas forcément distincts, on peut donc écrire en général l'ensemble $\sigma(f)$ des valeurs propres de f sous la forme

$$\{\alpha_1, \dots, \alpha_n\} = \{\lambda_1, \dots, \lambda_r\},$$

où cette fois les λ_i sont deux-à-deux distincts (observer que $r \leq n$).

Notons m_j la multiplicité algébrique de la valeur propre λ_j , c'est donc le nombre d'indices i tel que $\alpha_i = \lambda_j$. Le polynôme caractéristique s'écrit alors :

$$\chi_f(t) = \prod_{j=1}^r (t - \lambda_j)^{m_j}.$$

Quitte à renuméroter les vecteurs v_i de la base propre $\mathcal{B}$, on peut supposer que

$$f(v_i) = \lambda_j v_i \quad \text{pour } s_{j-1} < i \leq s_j,$$

où on a noté $s_0 = 0$ et $s_j = m_1 + \cdots + m_j$ pour $j \geq 1$. La famille

$$\mathcal{B}_j = \{v_i \mid s_{j-1} < i \leq s_j\}$$

est donc une base de l'espace propre E_{λ_j} , par conséquent

$$\text{multgeom}_{\lambda_j}(f) = \dim(E_{\lambda_j}) = \text{Card}(\mathcal{B}_j) = s_j - s_{j-1} = m_j = \text{multalg}_{\lambda_j}(f).$$

On a prouvé que la multiplicité géométrique de chaque valeur propre est égale à sa multiplicité algébrique.

(b) $\Rightarrow$ (c) Si le polynôme caractéristique est scindé et la multiplicité géométrique de chaque valeur propre est égale à sa multiplicité algébrique, alors en particulier la somme des multiplicités géométriques est égale au degré du polynôme caractéristique, or ce degré est égal à la dimension de l'espace vectoriel V :

$$\sum_{\lambda \in \sigma(f)} \text{multgeom}_\lambda(f) = \deg(\chi_f(t)) = n = \dim(V).$$

(c) $\Rightarrow$ (a) La condition (c) signifie en particulier qu'on peut trouver n vecteurs propres linéairements indépendants, ce qui implique immédiatement que f est diagonalisable. □

8.8 Puissances d'une matrice diagonalisable

Soit $A \in M_n(K)$ une matrice carrée et $k \in \mathbb{N}$ un entier naturel. On définit les puissances de la matrice A par

$$A^0 = \mathbf{I}_n, \quad A^k = \underbrace{A \cdot A \cdots A}_k.$$

Supposons que la matrice A est diagonalisable, alors elle s'écrit

$$A = PDP^{-1},$$

où $D = \text{Diag}(\lambda_1, \dots, \lambda_n)$ est la forme diagonale de A et $P \in \text{GL}_n(K)$ est la matrice modale. On déduit de l'identité précédente que

$$A^k = PD^kP^{-1}, \quad (8.1)$$

et comme on a clairement

$$D^k = \text{Diag}(\lambda_1^k, \dots, \lambda_n^k),$$

on voit que le calcul d'une puissance quelconque de la matrice A ne pose pas de problème lorsqu'on peut diagonaliser la matrice A .

Résumons-nous : Pour calculer les puissances d'une matrice $A \in M_n(K)$, on peut essayer de la diagonaliser. On procède donc ainsi :

- (i) On calcule le polynôme caractéristique $\chi_A(t)$.
- (ii) On cherche les valeurs propres en résolvant l'équation caractéristique $\chi_A(\lambda) = 0$.
- (iii) On cherche une base de K^n formée de vecteurs propres de A .
- (iv) Nous avons donc la forme diagonale D et la matrice modale P .
- (v) Finalement : $A^k = P \cdot D^k \cdot P^{-1}$.

Exemple 1. Trouver A^{10} , où

$$A = \begin{pmatrix} 5 & 6 \\ -2 & -2 \end{pmatrix}.$$

Procédons par diagonalisation :

- i) Le polynôme caractéristique $\chi_A(t)$ est donné par

$$\chi_A(t) = \det \begin{pmatrix} t-5 & -6 \\ 2 & t+2 \end{pmatrix} = t^2 - 3t + 2 = (t-1)(t-2),$$

- ii) Les valeurs propres de A sont donc $\lambda_1 = 1$ et $\lambda_2 = 2$.
- iii) Un vecteur propre de A pour la valeur propre $\lambda_1 = 1$ s'obtient en résolvant le système $A \cdot X = X$, c'est à dire

$$\begin{pmatrix} 5 & 6 \\ -2 & -2 \end{pmatrix} \cdot \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} = \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} \Leftrightarrow \begin{cases} 5x_1 + 6x_2 = x_1 \\ -2x_1 - 2x_2 = x_2 \end{cases}$$

Une solution est donnée par $x_1 = 3$ et $x_2 = -2$; ainsi $X = \begin{pmatrix} 3 \\ -2 \end{pmatrix}$ est un vecteur propre pour

$\lambda_1 = 1$. De façon similaire, $Y = \begin{pmatrix} 2 \\ -1 \end{pmatrix}$ est un vecteur propre pour $\lambda_2 = 2$.

- iv) La forme diagonale et la matrice modale de A sont donc

$$D = \begin{pmatrix} 1 & 0 \\ 0 & 2 \end{pmatrix} \quad \text{et} \quad P = \begin{pmatrix} 3 & 2 \\ -2 & -1 \end{pmatrix}$$

Le déterminant de P vaut 1 et $P^{-1} = \begin{pmatrix} -1 & -2 \\ 2 & 3 \end{pmatrix}$. On vérifie par un simple calcul que $A = P \cdot D \cdot P^{-1}$.

v) Finalement : $A^{10} = P \cdot D^{10} \cdot P^{-1}$, c'est à dire

$$A^{10} = \begin{pmatrix} 3 & 2 \\ -2 & -1 \end{pmatrix} \begin{pmatrix} 1 & 0 \\ 0 & 2^{10} \end{pmatrix} \begin{pmatrix} -1 & -2 \\ 2 & 3 \end{pmatrix} = \begin{pmatrix} 4093 & 6138 \\ -2046 & -3068 \end{pmatrix}$$

Exemple 2. On considère la matrice

$$G = \begin{pmatrix} 1 & 1 \\ 0 & a \end{pmatrix}$$

avec $a \notin \{0, 1\}$. On vérifie par récurrence que les puissances de cette matrice sont données par

$$G^n = \begin{pmatrix} 1 & s_n \\ 0 & a^n \end{pmatrix} \quad (*)$$

où

$$s_n = 1 + a + a^2 + \dots + a^{n-1} = \sum_{k=0}^{n-1} a^k.$$

En effet, pour vérifier l'équation (*), il suffit d'observer que

$$G^{n+1} = G \cdot G^n = \begin{pmatrix} 1 & 1 \\ 0 & a \end{pmatrix} \cdot \begin{pmatrix} 1 & s_n \\ 0 & a^n \end{pmatrix} = \begin{pmatrix} 1 & s_n + a^n \\ 0 & a^{n+1} \end{pmatrix} = \begin{pmatrix} 1 & s_{n+1} \\ 0 & a^{n+1} \end{pmatrix}$$

Appliquons la méthode de diagonalisation pour calculer G^n :

i) Le polynôme caractéristique $\chi_G(t)$ est donné par

$$\chi_G(t) = \det \begin{pmatrix} t-1 & -1 \\ 0 & t-a \end{pmatrix} = (t-1)(t-a).$$

ii) Les valeurs propres sont donc $\lambda_1 = 1$ et $\lambda_2 = a$.

iii) On trouve les vecteurs propres $X = \begin{pmatrix} 1 \\ 0 \end{pmatrix}$ pour la valeur propre $\lambda_1 = 1$ et $Y = \begin{pmatrix} 1 \\ a-1 \end{pmatrix}$ pour la valeur propre $\lambda_2 = a$. Ces vecteurs propres sont linéairement indépendants car nous avons supposé $a \neq 1$.

iv) Nous avons donc la forme diagonale D et la matrice modale P suivantes

$$D = \begin{pmatrix} 1 & 0 \\ 0 & a \end{pmatrix} \quad \text{et} \quad P = \begin{pmatrix} 1 & 1 \\ 0 & a-1 \end{pmatrix}$$

La matrice inverse de P est donnée par $P^{-1} = \frac{1}{a-1} \begin{pmatrix} a-1 & -1 \\ 0 & 1 \end{pmatrix}$ et on vérifie par calcul que $G = PDP^{-1}$

v) Nous pouvons maintenant calculer :

$$\begin{aligned} G^n &= P \cdot D^n \cdot P^{-1} = \frac{1}{a-1} \begin{pmatrix} 1 & 1 \\ 0 & a-1 \end{pmatrix} \begin{pmatrix} 1 & 0 \\ 0 & a^n \end{pmatrix} \begin{pmatrix} a-1 & -1 \\ 0 & 1 \end{pmatrix} \\ &= \frac{1}{a-1} \begin{pmatrix} 1 & a^n \\ 0 & a^n(a-1) \end{pmatrix} \begin{pmatrix} a-1 & -1 \\ 0 & 1 \end{pmatrix} \\ &= \frac{1}{a-1} \begin{pmatrix} a-1 & a^n-1 \\ 0 & a^n(a-1) \end{pmatrix} \\ &= \begin{pmatrix} 1 & \frac{a^n-1}{a-1} \\ 0 & a^n \end{pmatrix}. \end{aligned}$$

Observons en particulier que

$$s_n = \sum_{k=0}^{n-1} a^k = \frac{a^n - 1}{a - 1},$$

nous avons donc retrouvé la formule de sommation d'une série géométrique au moyen du calcul matriciel.

8.9 Application aux récurrences linéaires

Définition. On dit qu'une suite $(x_k) \subset K$ est une *récurrence linéaire d'ordre m* si on a pour tout $k \in \mathbb{N}$ la relation

$$x_{k+m} = a_0 x_k + a_1 x_{k+1} + \cdots + a_{m-1} x_{k+m-1}, \quad (8.2)$$

où $a_0, \dots, a_{m-1} \in K$. La suite est alors déterminée par cette relation et les conditions initiales

$$x_0 = c_0, \quad x_1 = c_1, \dots, x_{m-1} = c_{m-1}. \quad (8.3)$$

Remarquons qu'une récurrence linéaire d'ordre 1 est simplement une suite géométrique $x_{k+1} = ax_k$, dont le terme général est donné par $x_k = a^k x_0$.

Définition. La *matrice compagne* de la récurrence (8.2) est la $m \times m$ matrice

$$A = \begin{pmatrix} 0 & 1 & 0 & \cdots & 0 \\ 0 & 0 & 1 & \cdots & 0 \\ \vdots & \vdots & \ddots & \ddots & \vdots \\ 0 & 0 & 0 & \ddots & 1 \\ a_0 & a_1 & a_2 & \cdots & a_{m-1} \end{pmatrix}.$$

Nous avons alors le résultat suivant qui résout la récurrence dans de nombreux cas.

Théorème 8.9.1. *Supposons que la matrice compagne A de la récurrence (8.2) est diagonalisable, alors la solution (x_k) s'écrit*

$$x_k = b_1 \lambda_1^k + \cdots + b_m \lambda_m^k, \quad (8.4)$$

où les λ_i sont les valeurs propres de A (pas nécessairement toutes distinctes) et les b_i sont des constantes qu'on peut déterminer à partir des conditions initiales (8.3).

Remarque. Le polynôme caractéristique de la matrice compagne est donné par

$$\chi_A(t) = t^m - a_{m-1} t^{m-1} - \cdots - a_1 t - a_0.$$

Si ce polynôme a m racines distinctes, alors A est diagonalisable.

Preuve. On considère les vecteurs X_0 et $X_k \in K^m$ définis par

$$X_0 = \begin{pmatrix} c_0 \\ c_1 \\ \vdots \\ c_{m-1} \end{pmatrix} \quad \text{et} \quad X_k = \begin{pmatrix} x_k \\ x_{k+1} \\ \vdots \\ x_{k+m-1} \end{pmatrix}.$$

La récurrence (8.2) est alors équivalente à l'équation matricielle $X_{k+1} = A \cdot X_k$ dont la solution est clairement

$$X_k = A^k \cdot X_0, \quad (8.5)$$

Il s'agit donc de calculer les puissances de la matrice A . Nous avons supposé que A est diagonalisable, on a donc $A = PDP^{-1}$ où P est une matrice inversible et $D = \text{diag}(\lambda_1, \dots, \lambda_m)$. Par conséquent

$$A^k = PD^kP^{-1}, \quad \text{avec} \quad D^k = \text{diag}(\lambda_1^k, \dots, \lambda_m^k).$$

On a donc

$$X_k = PD^kP^{-1}X_0, \tag{8.6}$$

et la première composante x_k du vecteur X_k est bien du type (8.4).

□

Exemple : la suite de Fibonacci

La suite de Fibonacci se définit récursivement par $x_0 = 0$, $x_1 = 1$ et

$$x_{k+2} = x_k + x_{k+1},$$

pour tout $k \geq 0$. Les premiers termes sont

$$x_k = 0, 1, 1, 2, 3, 5, 8, 13, 21, 34, 55, \dots$$

On peut écrire la suite de Fibonacci matriciellement $X_k = F^k \cdot X_0$ où $X_k = \begin{pmatrix} x_k \\ x_{k+1} \end{pmatrix}$ et F est la matrice compagnon :

$$F = \begin{pmatrix} 0 & 1 \\ 1 & 1 \end{pmatrix}.$$

Le polynôme caractéristique de F est $\chi_F(t) = t^2 - t - 1$, les valeurs propres sont

$$\lambda = \frac{1 + \sqrt{5}}{2} \cong 1.618 \quad \text{et} \quad \mu = \frac{1 - \sqrt{5}}{2} = 1 - \lambda \cong -0.618$$

(λ est le *nombre d'or*). Le terme général de la suite de Fibonacci peut donc s'écrire

$$x_k = b_1 \lambda^k + b_2 \mu^k.$$

où b_1 et b_2 sont déterminés par les conditions initiales :

$$0 = x_0 = b_1 + b_2, \quad 1 = x_1 = b_1 \lambda + b_2 \mu.$$

La première équation donne $b_1 = -b_2 =: b$ et la seconde dit que $b(\lambda - \mu) = 1$, c'est-à-dire

$$b = \frac{1}{\lambda - \mu} = \frac{1}{\sqrt{5}}.$$

Le $k^{\text{ème}}$ terme de la suite de Fibonacci est donc égal à

$$x_k = \frac{1}{\sqrt{5}} (\lambda^k - \mu^k) = \frac{\left(\frac{1+\sqrt{5}}{2}\right)^k - \left(\frac{1-\sqrt{5}}{2}\right)^k}{\sqrt{5}}.$$

Cette expression pour calculer le terme général de la suite de Fibonacci, s'appelle la *formule de Binet*⁵, par exemple

$$x_{100} = \frac{\left(\frac{1+\sqrt{5}}{2}\right)^{100} - \left(\frac{1-\sqrt{5}}{2}\right)^{100}}{\sqrt{5}} = 354224848179261915075$$

5. Jacques Philippe Marie Binet, mathématicien français. 1786–1856.

Terminons par quelques remarques :

1. La formule de Binet peut se vérifier directement de la manière suivante : pour $k = 0, 1, 2$, elle dit que

$$x_0 = \frac{1}{\sqrt{5}} (\lambda^0 - \mu^0) = 0,$$

$$x_1 = \frac{1}{\sqrt{5}} (\lambda - \mu) = 1,$$

et

$$x_2 = \frac{1}{\sqrt{5}} (\lambda^2 - \mu^2) = \frac{1}{\sqrt{5}} ((\lambda + 1) - (\mu + 1)) = 1.$$

Dans ce calcul, et le suivant, on utilise la propriété suivante de λ et μ :

$$\lambda^2 = \lambda + 1, \quad \mu^2 = \mu + 1.$$

On démontre maintenant le cas général de la formule de Binet par récurrence (pour $k > 2$) :

$$\begin{aligned} x_k &= \frac{1}{\sqrt{5}} (\lambda^k - \mu^k) \\ &= \frac{1}{\sqrt{5}} (\lambda^2 \lambda^{k-2} - \mu^2 \mu^{k-2}) \\ &= \frac{1}{\sqrt{5}} ((\lambda + 1) \lambda^{k-2} - (\mu + 1) \mu^{k-2}) \\ &= \frac{1}{\sqrt{5}} ((\lambda^{k-1} + \lambda^{k-2}) - (\mu^{k-1} + \mu^{k-2})) \\ &= x_{k-1} + x_{k-2}. \end{aligned}$$

2. Le nombre $\mu \cong -0.618$ est en valeur absolue plus petit que 1. Donc μ^k converge rapidement vers 0 lorsque k tend vers l'infini (on a $|\mu|^k < 0.1$ si $k \geq 5$). Pour cette raison on peut négliger le terme μ^k dans la formule de Binet et simplement calculer x_k ainsi :

$$x_k = \text{l'entier le plus proche de } \frac{1}{\sqrt{5}} \left(\frac{1 + \sqrt{5}}{2} \right)^k.$$

Par exemple $\frac{1}{\sqrt{5}} \left(\frac{1 + \sqrt{5}}{2} \right)^6 = 8.0249$, et donc $x_6 = 8$. L'ordre de grandeur de x_{1000} est

$$x_{1000} \cong \frac{1}{\sqrt{5}} \left(\frac{1 + \sqrt{5}}{2} \right)^{1000} \cong 4.36284 \cdot 10^{208}.$$

3. La suite de Fibonacci tend vers l'infini, mais la proportion entre deux termes successifs converge vers le nombre d'or :

$$\lim_{k \rightarrow \infty} \frac{x_{k+1}}{x_k} = \lambda = \frac{1 + \sqrt{5}}{2}$$

En effet, on a

$$\lim_{k \rightarrow \infty} \frac{x_{k+1}}{x_k} = \lim_{k \rightarrow \infty} \frac{\frac{1}{\sqrt{5}} (\lambda^{k+1} - \mu^{k+1})}{\frac{1}{\sqrt{5}} (\lambda^k - \mu^k)} = \lambda$$

car $\lim_{k \rightarrow \infty} \mu^k = 0$.

Réurrence générale d'ordre 2

La suite de Fibonacci est un cas particulier de récurrence linéaire d'ordre 2. Le cas général est du type

$$x_{k+2} = ax_k + bx_{k+1}.$$

Cherchons la solution de cette récurrence. Le polynôme caractéristique associé est $\chi(t) = t^2 - bt - a$ et les valeurs propres sont

$$\alpha = \frac{b + \sqrt{b^2 + 4a}}{2}, \quad \beta = \frac{b - \sqrt{b^2 + 4a}}{2}.$$

On suppose que $\alpha \neq \beta$, alors le théorème 8.9.1 nous dit que le terme général de la récurrence s'écrit

$$x_k = p\alpha^k + q\beta^k.$$

Pour trouver les coefficients p et q on résout le système linéaire

$$\begin{cases} p + q &= x_0 \\ p\alpha + q\beta &= x_1. \end{cases}$$

et on trouve que

$$p = \frac{\beta x_0 - x_1}{\beta - \alpha}, \quad q = -\frac{\alpha x_0 - x_1}{\beta - \alpha}.$$

Le terme général est alors donné par

$$x_k = \left(\frac{x_0\beta - x_1}{\beta - \alpha} \right) \alpha^k - \left(\frac{\alpha x_0 - x_1}{\beta - \alpha} \right) \beta^k,$$

que l'on peut aussi écrire sous la forme

$$x_k = \left(\frac{\alpha^k \beta - \alpha \beta^k}{\beta - \alpha} \right) x_0 + \left(\frac{\beta^k - \alpha^k}{\beta - \alpha} \right) x_1. \quad (8.7)$$

On voit clairement que la méthode ne marche pas lorsque $\alpha = \beta$. D'une manière générale les puissances d'une matrice non diagonalisables sont plus délicates à calculer, nous donnons quelques indications dans l'annexe C.

Annexe A

Notions sur le polynômes

Dans ce paragraphe on donne quelques notions sur les polynômes.

Definitions.

(1.) Un *polynôme* $f(t)$ à coefficients dans le corps K est une expression formelle

$$f(t) = a_n t^n + a_{n-1} t^{n-1} + \cdots + a_1 t + a_0,$$

où $a_i \in K$ pour $0 \leq i \leq n$. Les a_i s'appellent *les coefficients de f et t l'indéterminée*. L'ensemble de tous les polynômes à coefficients dans K est noté $K[t]$.

(2.) Le plus grand entier n tel que $a_n \neq 0$ s'appelle le *degré* de f . On le note $\deg(f)$ et a_n est le *coefficient dominant* du polynôme. Par convention on admet que si $f = 0$ (le polynôme nul) alors $\deg(f) = -\infty$.

(3.) Deux polynômes $f(t) = a_n t^n + a_{n-1} t^{n-1} + \cdots + a_1 t + a_0$ et $g(t) = b_m t^m + b_{m-1} t^{m-1} + \cdots + b_1 t + b_0$ sont égaux si et seulement s'ils ont le même degré et $a_i = b_i$ pour tout i .

(4.) Un polynôme est dit *unitaire* si son coefficient dominant vaut 1, i.e. si $\deg(f) = n$ et $a_n = 1$.

(5.) Le coefficient a_0 s'appelle le *terme constant* du polynôme.

(6.) Les polynômes de degré 0 s'appellent les *polynômes constants*.

Proposition A.0.1. (i) $K[t]$ est une K -algèbre unitaire commutative.

(ii) L'élément neutre pour la multiplication est le polynôme $f(t) = 1$.

(iii) $\deg(fg) = \deg(f) + \deg(g)$.

Definition. On dit que $\alpha \in K$ est un *zéro* ou une *racine* du polynôme $f \in K[t]$ si $f(\alpha) = 0$.

Proposition A.0.2. (La division euclidienne dans $K[t]$) Soient $f(t), g(t) \in K[t]$, $g(t) \neq 0$. Alors il existe $q(t), r(t) \in K[t]$ tels que

$$f(t) = g(t)q(t) + r(t) \quad \text{et} \quad \deg(r) < \deg(g).$$

De plus les deux polynômes $q(t)$ et $r(t)$ sont uniquement déterminés par ces propriétés.

Definition. Dans la proposition précédente, $q(t)$ s'appelle le *quotient* et $r(t)$ le *reste* de la division de $f(t)$ par $g(t)$.

Preuve. On montre d'abord l'existence de la division euclidienne : Si $\deg(g) > \deg(f)$, alors on pose $q = 0$ et $r = f$. La preuve est finie dans ce cas. Supposons donc $m = \deg(g) \leq n = \deg(f)$ et faisons une preuve par récurrence sur n . Si $n = 0$, alors f et g sont des constantes et on pose $r = 0$ et $q = f/g$ (rappelons qu'on suppose $g \neq 0$). On a démontré le théorème pour $n = 0$.

Supposons le théorème démontré pour les polynômes de degré $< n$ et notons

$$\begin{aligned} f(t) &= a_n t^n + a_{n-1} t^{n-1} + \cdots + a_0 \\ g(t) &= b_m t^m + b_{m-1} t^{m-1} + \cdots + b_0. \end{aligned}$$

On pose alors $q_1(t) = \frac{a_n}{b_m} t^{n-m}$ et on définit $r_1(t)$ par

$$f(t) = q_1(t)g(t) + r_1(t).$$

Alors $\deg(r_1) < n$ et par récurrence on peut écrire

$$r_1(t) = q_2(t)g(t) + r_2(t)$$

avec $\deg(r_2) < m$. On a alors

$$f(t) = (q_1(t) + q_2(t))g(t) + r_2(t)$$

avec $\deg(r_2) < m$. L'existence est démontrée.

Pour prouver l'unicité, on suppose que

$$f(t) = q_1(t)g(t) + r_1(t) = q_2(t)g(t) + r_2(t)$$

avec $\deg(r_1) < \deg(g)$ et $\deg(r_2) < \deg(g)$ Alors

$$(q_1 - q_2)g = r_1 - r_2,$$

et comme $\deg(r_1 - r_2) < \deg(g)$ cette relation doit être $0 = 0$.

□

Corollaire A.0.3. Soient $f(t) \in K[t]$ et $\alpha \in K$. Alors α est une racine de $f(t)$ si et seulement si $f(t) = (t - \alpha)q(t)$ pour un polynôme $q(t) \in K[t]$.

Preuve. On applique la division euclidienne à $f(t)$ et $g(t) = (t - \alpha)$. Donc il existe $q(t)$ et $r(t)$ dans $K[t]$ tels que $f(t) = q(t)(t - \alpha) + r(t)$ avec $\deg(r) < \deg(g) = \deg(t - \alpha) = 1$. Donc $r(t)$ est un polynôme constant, i.e. $r(t) = c$ et on a

$$c = r(\alpha) = f(\alpha) - q(\alpha)(\alpha - \alpha) = 0.$$

Ainsi $r(t)$ est le polynôme nul et on a $f(t) = q(t)(t - \alpha)$.

□

Exercice. Soient $c \in K$ et $f(t) \in K[t]$. Montrer que le reste de la division de $f(t)$ par le polynôme $(t - c)$ est le polynôme constant $f(c)$.

Definitions.

- 1.) Soient $f(t), g(t) \in K[t]$. On dit que g *divise* f s'il existe un polynôme $q(t) \in K[t]$ tel que $f(t) = q(t)g(t)$.
- 2.) Soient $f(t), g(t) \in K[t]$. On dit que $f(t)$ et $g(t)$ sont *premiers entre eux* si les seuls diviseurs communs de f et g dans $K[t]$ sont les polynômes constants. Plus généralement, les polynômes $f_1, \dots, f_m \in K[t]$ sont *premiers entre eux* si les seuls diviseurs communs à chaque f_i dans $K[t]$ sont les polynômes constants.
- 3.) On dit que $f(t) \in K[t]$ est *irréductible* si f est non constant et si pour toute décomposition $f(t) = g(t)h(t)$ avec $g(t), h(t) \in K[t]$ on a ou bien g est constante ou bien h est constante. Dans le cas contraire, on dit que $f(t)$ est *réductible*.

Théorème A.0.4. (*Théorème fondamental de l'algèbre*) Tout polynôme de $\mathbb{C}[t]$ possède au moins une racine dans $\mathbb{C}$.

La preuve de ce théorème est habituellement donnée dans le cours d'analyse complexe.

Corollaire A.0.5. Tout polynôme $f \in \mathbb{C}[t]$ non constant peut s'écrire

$$f(t) = c \cdot (t - \alpha_1)(t - \alpha_2) \cdots (t - \alpha_n).$$

Preuve. On sait qu'il existe $\alpha_1 \in \mathbb{C}$ tel que $f(\alpha_1) = 0$. On peut donc écrire $f(t) = (t - \alpha_1)g(t)$ et $\deg(g) = \deg(f) - 1$. On conclut par récurrence sur le degré de f . □

Les racines α_i ne sont pas forcément distinctes, en regroupant les racines qui se répètent, on peut écrire

$$f(t) = c \cdot (t - \beta_1)^{k_1} (t - \beta_2)^{k_2} \cdots (t - \beta_m)^{k_m}, \quad (\text{A.1})$$

avec les β_j distincts. On dit alors que β_j est une racine de f de *multiplicité* k_j .

Définition. Un polynôme $f \in K[t]$ est dit *scindé* si on peut l'écrire sous la forme (A.1). Le théorème fondamental de l'algèbre (ou plus précisément son corollaire) nous dit que tout polynôme de $\mathbb{C}[t]$ est scindé.

Théorème A.0.6. Tout polynôme $f(t) \in K[t]$ avec $\deg(f) \geq 1$ se décompose de façon unique (à l'ordre des facteurs près) comme un produit

$$f(t) = a g_1(t) \cdots g_r(t),$$

où $a \in K$ et $g_i(t) \in K[t]$ est un polynôme irréductible unitaire pour $1 \leq i \leq r$.

Les facteurs $g_i(t)$ dans le théorème s'appellent *les facteurs irréductibles* de $f(t)$.

Exemple. Un exemple important du théorème précédent est que le polynôme $p(t) = (t^d - 1) \in \mathbb{C}[t]$ peut se scinder en produit de polynômes de degrés 1 de la façon suivante :

$$(t^d - 1) = \prod_{k=0}^{d-1} (t - \omega^k), \quad \text{où} \quad \omega = \exp\left(\frac{2i\pi}{d}\right).$$

En effet, ω^k est racine de $p(t)$ pour tout k , donc les deux membres de l'égalité ci-dessus sont des polynômes unitaires de même degré ayant les mêmes racines, ils ont donc les même facteurs irréductibles.

Théorème A.0.7. (*Identité de Bézout*) Si $f_1(t), \dots, f_m(t) \in K[t]$ sont des polynômes premiers entre eux, alors il existe $p_1(t), \dots, p_m(t) \in K[t]$ tels que

$$p_1(t)f_1(t) + \dots + p_m(t)f_m(t) = 1.$$

Preuve. Nous faisons la preuve pour le cas de deux polynômes (i.e. $m = 2$, rien d'essentiel ne change pour le cas général). Soit $J \subset K[t]$ l'ensemble des polynômes tels que

$$J = \{s_1f_1 + s_2f_2 \mid s_1, s_2 \in K[t]\}.$$

Notons m le degré minimal d'un élément non nul de J et choisissons un élément $u \in J$ de degré m . Ecrivons la division euclidienne de f_1 et f_2 par u . On a donc

$$f_1 = q_1u + r_1 \quad \text{et} \quad f_2 = q_2u + r_2$$

avec $\deg(r_i) < \deg(u) = m$. Montrons que les polynômes $r_1(t)$ et $r_2(t)$ appartiennent à J : on sait que $u \in J$, donc on peut écrire $u = s_1f_1 + s_2f_2$ et par conséquent

$$r_1 = f_1 - q_1u = (1 - q_1s_1)f_1 + (-q_1s_2)f_2 \in J.$$

Et de même $r_2 \in J$. Nous avons par ailleurs $\deg(r_1) < \deg(u) = m$; par minimalité de m on conclut que $r_1 = 0$. De même $r_2 = 0$. Mais alors

$$f_1 = q_1u \quad \text{et} \quad f_2 = q_2u,$$

et comme on a supposé que f_1 et f_2 sont premiers entre eux, le polynôme u doit être constant, $u = c$. On a donc montré qu'il existe une constante non nulle $c \in K^*$ qui appartient à J ; il existe donc $s_1, s_2 \in K[t]$ tels que $c = s_1f_1 + s_2f_2$. En posant $p_1 = s_1/c$ et $p_2 = s_2/c$ on obtient

$$p_1(t)f_1(t) + p_2(t)f_2(t) = \frac{1}{c}(s_1f_1 + s_2f_2) = \frac{c}{c} = 1.$$

□

Annexe B

Un petit guide pour la diagonalisation

B.1 Aspects théoriques de la diagonalisation

Rappelons qu'un endomorphisme $f \in \mathcal{L}(V)$ d'un espace vectoriel de dimension finie est *diagonalisable* s'il existe une base de V pour laquelle la matrice de f est une matrice diagonale. Cette base est donc formée de vecteurs propres. Une matrice $A \in M_n(K)$ est *diagonalisable sur le corps K* si elle est semblable à une matrice diagonale, i.e. s'il existe une matrice inversible $P \in GL_n(K)$ telle que $P^{-1}AP$ est une matrice diagonale.

Les conditions de diagonalisabilité peuvent être formulées dans le théorème suivant :

Théorème B.1.1. *Soit $f \in \mathcal{L}(V)$ un endomorphisme d'un K -espace vectoriel de dimension n , notons $\sigma(f) = \{\lambda_1, \dots, \lambda_r\}$ l'ensemble de ses valeurs propres. Alors les conditions suivantes sont équivalentes.*

- (a) *f est diagonalisable.*
- (b) *Il existe une base de V formée de vecteurs propres pour f*
- (c) *L'espace vectoriel V s'écrit comme somme directe des sous-espaces propres de f :*

$$V = E_{\lambda_1}(f) \oplus \dots \oplus E_{\lambda_r}(f).$$

- (d) *Le polynôme caractéristique $\chi_f(t)$ de f est scindé et la multiplicité algébrique de chaque valeur propre est égale à sa multiplicité géométrique :*

$$\text{multgeom}_\lambda(f) = \text{multalg}_\lambda(f).$$

Preuve du théorème. Nous prouvons les implications dans l'ordre suivant : (a) $\Leftrightarrow$ (b) $\Leftrightarrow$ (c) $\Leftrightarrow$ (d).

La proposition 8.6.1 nous donne l'équivalence (a) $\Leftrightarrow$ (b) et l'équivalence (b) $\Leftrightarrow$ (c) s'obtient en regroupant les bases des différents espaces propres E_{λ_i} pour obtenir une base de V .

L'équivalence (c) $\Leftrightarrow$ (d) découle immédiatement des quatre faits suivants :

- Le théorème de décomposition primaire, qui dit que $V = N_{\lambda_1}(f) \oplus \dots \oplus N_{\lambda_r}(f)$, où $N_{\lambda_i}(f)$ est l'espace caractéristique associé à la valeur propre λ_i .
- $E_{\lambda_i}(f) \subset N_{\lambda_i}(f)$.
- $\text{multgeom}_\lambda(f) = \dim(E_\lambda(f))$.
- $\text{multalg}_\lambda(f) = \dim(N_{\lambda_i}(f))$.

□

Rappelons que des vecteurs propres associés à des valeurs propres distinctes sont linéairement indépendants, par conséquent on a :

Corollaire B.1.2. *Si $f \in \mathcal{L}(V)$ possède n valeurs propres distinctes deux à deux ($n = \dim(V)$), alors f est diagonalisable.*

Méthode.

Voici les étapes concrètes pour décider si une matrice $A \in M_n(K)$ est diagonalisable ou non.

1. Analyser le polynôme caractéristique $\chi_A(t)$ de la façon suivante :
 - i.) Calculer le polynôme caractéristique $\chi_A(t)$.
 - ii.) Chercher toutes les racines de ce polynôme (i.e. résoudre l'équation polynomiale $\chi_A(\lambda) = 0$ avec $\lambda \in K$, ces racines sont les valeurs propres de A).
 - iii.) Si la matrice $A \in M_n(K)$ possède n valeurs propres deux à deux distinctes dans le corps K alors A est diagonalisable sur K .
 - iv.) Sinon, on décompose $\chi_A(t)$ en facteurs irréductibles, on constate alors si ce polynôme est scindé ou non.
 - v.) Si $\chi_A(t)$ n'est pas scindé, la matrice A n'est pas diagonalisable et le procédé s'arrête ici.
2. Si $\chi_A(t)$ est scindé, i.e. $\chi_A(t) = \prod_{i=1}^r (t - \lambda_i)^{m_i}$, alors on vérifie pour chaque valeur propre si sa multiplicité algébrique (= la multiplicité dans le polynôme caractéristique) est égale à $\dim(E_{\lambda_i})$. Si c'est le cas la matrice est diagonalisable.
3. A la place de l'étape précédente, on peut préférer vérifier directement si le polynôme $\nu_A(t)$ annule la matrice A , c'est à dire si

$$\nu_A(A) = (A - \lambda_1 \cdot I_n)(A - \lambda_2 \cdot I_n) \cdots (A - \lambda_r \cdot I_n) = 0.$$

Remarques :

- (1) Le procédé décrit ci-dessus présente un problème majeur dans la deuxième étape car il n'y a pas de méthode pour trouver les racines d'un polynôme général de degré ≥ 5 . Toutefois il fonctionne bien pour les petites matrices et pour des matrices de grande tailles assez spéciales.
- (2) On peut trouver la multiplicité géométrique $\dim(E_{\lambda_i})$ en utilisant la formule du rang :

$$\dim(E_{\lambda_i}) = \dim(\text{Ker}(\lambda_i I_n - A)) = n - \text{rang}(\lambda_i I_n - A);$$

puis utiliser l'algorithme de Gauss-Jordan pour trouver le rang de $(\lambda_i I_n - A)$.

- (3) Ne pas oublier que sur le corps $K = \mathbb{C}$, tout polynôme est scindé.
- (4) Il arrive souvent qu'une matrice $A \in M_n(\mathbb{R})$ soit non diagonalisable sur les réels, mais qu'elle soit diagonalisable sur les complexes. Dans ce cas la matrice modale P est une matrice de $GL_n(\mathbb{C})$.

Exemples. Considérons les matrices

$$A = \begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix} \quad \text{et} \quad B = \begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix}$$

On a $\chi_A(t) = t^2 + 1$, donc A n'est pas diagonalisable dans $M_2(\mathbb{R})$ mais A est diagonalisable dans $M_2(\mathbb{C})$ car dans $\mathbb{C}$ on a $\chi_A(t) = (t + i)(t - i)$.

Pour la matrice B , on a $\chi_B(t) = (t - 1)^2$, ce polynôme est scindé, mais $\dim(E_1) = 1 \neq 2 = m_1$ donc B n'est pas diagonalisable dans $M_2(\mathbb{C})$.

Nous concluons cette section en mentionnant que si une matrice réelle $A \in M_n(\mathbb{R})$ est symétrique, i.e. si $A^\top = A$, alors elle est diagonalisable. Plus généralement, si une matrice complexe $A \in M_n(\mathbb{C})$ vérifie $A^\top = \overline{A}$ (on dit alors que A est une matrice *hermitienne*), alors ses valeurs propres sont réelles et la matrice est diagonalisable. Ces résultats seront démontrés au semestre prochain.

B.2 Aspects pratiques de la diagonalisation

Pour diagonaliser une matrice, on cherche une base propre de $\mathcal{B} = \{v_1, \dots, v_n\} \subset K^n$ pour $A \in M_n(K)$, i.e. $Av_j = \lambda_j v_j$. Puis on note $P \in M_n(K)$ la matrice dont la $j^{\text{ème}}$ colonne est formée des coordonnées de v_j ; c'est donc la matrice de transition de la base canonique vers la base propre $\mathcal{B}$. Notons aussi $D = \text{Diag}(\lambda_1, \dots, \lambda_n)$, alors on a :

$$A = PDP^{-1} \quad \text{et} \quad D = P^{-1}AP.$$

Rappelons que la matrice $D = \text{Diag}(\lambda_1, \dots, \lambda_n)$ est la *forme diagonale* de A et P est la *matrice modale* (c'est la matrice de passage de la base canonique de K^n vers la base propre de A).

Mise en oeuvre.

Si on souhaite diagonaliser une matrice $A \in M_n(K)$, on vérifie d'abord si elle est diagonalisable avec les critères décrits plus haut. Si c'est le cas on continue ainsi :

1. On calcule son spectre $\{\lambda_1, \dots, \lambda_r\}$ et pour chaque valeur propre on note $m_i = \text{multalg}_{\lambda_i}(A)$ sa multiplicité algébrique.
2. Pour chaque valeur propre λ_i on cherche une base $\mathcal{B}_i$ de E_{λ_i} en trouvant m_i solutions *linéairement indépendantes* du système linéaire

$$AX = \lambda_i X \quad \Leftrightarrow \quad (A - \lambda_i I)X = 0.$$

3. On note $\mathcal{B} = \mathcal{B}_1 \cup \dots \cup \mathcal{B}_r$ la réunion des bases partielles obtenues à l'étape précédente. On vérifie que $\text{Card}(\mathcal{B}) = n$ (ça doit être le cas si la matrice est diagonalisable). On a ainsi obtenu une base propre $\mathcal{B}$ de K^n .
4. Chaque élément de $\mathcal{B}$ est un vecteur-colonne de K^n . La matrice modale est la matrice P dont la $j^{\text{ème}}$ colonne est formée par les coordonnées du $j^{\text{ème}}$ vecteur de $\mathcal{B}$.
5. On note D la matrice diagonale dont les m_1 premiers coefficients diagonaux sont λ_1 , les m_2 coefficients diagonaux suivants sont λ_2 etc.

$$D = \text{Diag}(\underbrace{\lambda_1, \dots, \lambda_1}_{m_1}, \dots, \underbrace{\lambda_r, \dots, \lambda_r}_{m_r})$$

Cette matrice est la forme diagonale de A .

6. On vérifie que $D = P^{-1}AP$, ou si on préfère, que $AP = PD$.

Exemple. Soit

$$R_\theta = \begin{pmatrix} \cos(\theta) & -\sin(\theta) \\ \sin(\theta) & \cos(\theta) \end{pmatrix}$$

la matrice de rotation d'angle θ dans le plan $\mathbb{R}^2$. Son polynôme caractéristique est

$$\chi_{R_\theta}(t) = t^2 - 2t \cos(\theta) + 1.$$

Ce polynôme est irréductible dans $\mathbb{R}[t]$ si $\cos(\theta) \neq \pm 1$, c'est-à-dire si θ n'est pas un multiple de π . En particulier $\chi_{R_\theta}(t)$ n'est donc pas scindé et R_θ n'est pas diagonalisable dans $M_2(\mathbb{R})$ si $0 < \theta < \pi$. En fait il n'y a aucun vecteur propre dans $\mathbb{R}^2$, ce qui est évident géométriquement pour une rotation.

Mais si on regarde la matrice R_θ dans $M_2(\mathbb{C})$, le polynôme caractéristique se factorise en

$$\chi_{R_\theta}(t) = (t - e^{i\theta})(t - e^{-i\theta}).$$

Il y a donc deux valeurs propres distinctes qui sont $e^{i\theta}$ et $e^{-i\theta}$. Par conséquent R_θ est diagonalisable comme matrice à coefficients complexes. On cherche une base propre en cherchant des solutions non nulles de

$$R_\theta X = e^{i\theta} X \quad \text{et} \quad R_\theta Y = e^{-i\theta} Y.$$

On trouve $X = \begin{pmatrix} 1 \\ -i \end{pmatrix}$ et $Y = \begin{pmatrix} 1 \\ i \end{pmatrix}$. La matrice modale est donc $P = \begin{pmatrix} 1 & 1 \\ -i & i \end{pmatrix}$ et on peut vérifier directement que

$$P^{-1}R_\theta P = \begin{pmatrix} e^{i\theta} & 0 \\ 0 & e^{-i\theta} \end{pmatrix}.$$

Annexe C

Sur les puissances d'une matrice

Cette annexe présente quelques techniques de calcul des puissances d'une matrice carrée. Rappelons que si A est une matrice carrée et $k \in \mathbb{N}$ est un nombre entier naturel, alors on définit

$$A^0 = \mathbf{I}_n, \quad A^k = \underbrace{A \cdot A \cdots A}_k.$$

C.1 Matrices infra-périodiques

Une première classe de matrices pour lesquelles il est facile de calculer les puissances est la classe des matrices *infra-périodiques*. On appelle ainsi une matrice $A \in M_n(k)$ pour lesquelles il existe un entier $m \geq 2$ et un scalaire $\alpha \in K$ tels que

$$A^{m+1} = \alpha \cdot A.$$

Si $m > 0$ est le plus petit entier positif ayant cette propriété, on dit que A est *infra-périodique d'ordre m* . Lorsque $\alpha = 1$, on dit que la matrice est *périodique*.

Le polynôme $p(t) = t^{m+1} - \alpha t = t(t^m - \alpha)$ est alors un polynôme annulateur de la matrice infra-périodique A . Les valeurs propres de A sont des racines de ce polynôme dans K , par conséquent les valeurs propres vérifient $\lambda^m = \alpha$ ou $\lambda = 0$.

Un cas particulier est lorsque $m = \alpha = 1$. Une matrice telle que $A^2 = A$ s'appelle une *matrice idempotente*.

Les puissances des matrices infra-périodiques sont faciles à calculer, voyons quelques exemples :

Exemples.

- (a) La matrice $R = \begin{pmatrix} 4 & 6 \\ -2 & -3 \end{pmatrix}$ est idempotente, i.e. elle vérifie $R^2 = R$.
- (b) Une matrice carrée S est dite involutive si $S^2 = \mathbf{I}_n$. Une symétrie centrale ou une symétrie à travers un hyperplan sont des exemples d'involutions.
- (c) De façon plus générale, une matrice S vérifiant $S^k = \mathbf{I}_n$ et $S^r \neq \mathbf{I}_n$ pour tout $0 < r < k$, est périodique d'ordre k . Un exemple est

$$S = \begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix}.$$

On a

$$S^2 = \begin{pmatrix} -1 & 0 \\ 0 & -1 \end{pmatrix}, \quad S^3 = \begin{pmatrix} 0 & 1 \\ -1 & 0 \end{pmatrix}, \quad S^4 = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}.$$

S est donc périodique d'ordre 4. On peut alors facilement calculer S^n pour tout n , par exemple

$$S^{19} = S^{16+3} = S^3 = \begin{pmatrix} 0 & 1 \\ -1 & 0 \end{pmatrix}.$$

- (d) La matrice A est *nilpotente* s'il existe m tel que $A^m = 0$. Une telle matrice est donc infra-périodique avec $\alpha = 0$. Un exemple est la matrice $T = \begin{pmatrix} 0 & 2 & 5 \\ 0 & 0 & 3 \\ 0 & 0 & 0 \end{pmatrix}$. On vérifie que

$$T^2 = \begin{pmatrix} 0 & 0 & 6 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix} \quad \text{et} \quad T^3 = \begin{pmatrix} 0 & 0 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix}$$

On en déduit que $T^n = 0$ pour tout $n \geq 3$. De manière générale, toute matrice triangulaire strictement supérieure est nilpotente.

- (e) Toute matrice carrée de rang 1 est infra-périodique. En effet, on sait qu'une telle matrice s'écrit $A = X \cdot Y^\top$ où X, Y sont des vecteur-colonnes de K^n . On a donc

$$A^2 = (XY^\top)(XY^\top) = X(Y^\top X)Y^\top = (Y^\top X)XY^\top = \gamma A,$$

où γ est le scalaire $Y^\top X$.

Une formule générale pour les puissances d'une matrice infra-périodique est donnée dans la proposition suivante :

Proposition C.1.1. *Soit $A \in M_n(K)$ une matrice telle que $A^{m+1} = \gamma \cdot A$. Donnons-nous un entier $k \geq m$ et écrivons $k = pm + q$ avec $p, q \in \mathbb{N}$ et $1 \leq q \leq m$. Alors*

$$A^k = A^{pm+q} = \gamma^p A^q.$$

Preuve. On observe que

$$A^{2m+1} = A^{m+1}A^m = \gamma AA^m = \gamma A^{m+1} = \gamma^2 A,$$

et on vérifie par récurrence que pour tout entier $p \geq 1$ on a $A^{pm+1} = \gamma^p A$. Donc pour tout entier $q \geq 1$ on a

$$A^{pm+q} = A^{pm+1}A^{q-1} = \gamma^p AA^{q-1} = \gamma^p A^q.$$

□

Cette proposition nous dit en particulier que toutes les puissances de A sont déterminées par $A, A^2, \dots, A^m$.

Lemme C.1.2. *Toute matrice complexe infra-périodique qui n'est pas nilpotente est diagonalisable.*

Preuve. Soit $A \in M_n(\mathbb{C})$ une matrice complexe telle que $A^{m+1} = \gamma \cdot A$ avec $\gamma \neq 0$. Cette matrice est annulée par le polynôme $p(t) = t(t^m - \gamma)$, or ce polynôme se scinde en produit de termes du premier degré car

$$(t^m - \gamma) = \prod_{k=0}^{m-1} (t - \alpha \cdot \omega^k) = (t - \alpha)(t - \alpha\omega)(t - \alpha\omega^2) \cdots (t - \alpha\omega^{m-1}),$$

avec $\alpha = \sqrt[m]{|\gamma|}$ et $\omega = \exp\left(\frac{2i\pi}{m}\right)$. On sait que toute matrice ayant cette propriété est diagonalisable.

□

C.2 Puissances d'une matrice diagonalisable

Le cas des matrice diagonalisables à été discuté au paragraphe 8.8. Rappelons brièvement que $A \in M_n(K)$ est diagonalisable si elle admet une base propre $\mathcal{B} = \{v_1, \dots, v_n\}$, i.e. $Av_j = \lambda_j v_j$, alors

$$A = PDP^{-1},$$

où $D = \text{Diag}(\lambda_1, \dots, \lambda_n)$ est la forme diagonale de A et $P \in \text{GL}_n(K)$ est la matrice modale, i.e. la matrice de transition de la base canonique vers la base propre propre $\mathcal{B}$ (c'est la matrice dont la $j^{\text{ème}}$ colonne est formée des coordonnées de v_j). On a alors

$$A^k = PD^k P^{-1}, \quad \text{avec} \quad D^k = \text{Diag}(\lambda_1^k, \dots, \lambda_n^k).$$

C.3 Polynôme annulateur et puissances d'une matrice

Soit $p(t) = b_0 + b_1 t + \dots + b_{m-1} t^{m-1} + t^m$ un polynôme (que nous supposons unitaire) qui annule une matrice $M \in M_n(K)$. Alors, par définition, on a

$$M^m = -b_0 I_n - b_1 M - \dots - b_{m-1} M^{m-1}, \quad (\text{C.1})$$

ce qui entraîne en particulier que toutes les puissances de A sont déterminées par $M, M^2, \dots, M^{m-1}$. L'identité (C.1) peut alors être vue comme une récurrence linéaire qu'il faut résoudre.

On a vu plus haut que la récurrence (C.1) se résout en calculant les puissances de la matrice compagnon

$$B = \begin{pmatrix} 0 & 1 & \dots & 0 \\ \vdots & \ddots & \ddots & \vdots \\ 0 & & \ddots & 1 \\ -b_0 & \dots & \dots & -b_{m-1} \end{pmatrix}.$$

A priori calculer les puissances de B n'est pas plus simple que de calculer celles de la matrice originale M . Mais lorsque le degré du polynôme annulateur $p(t)$ est plus petit que la taille de la matrice M , il y a un gain d'efficacité.

Supposons par exemple que $p(t) = (t - \alpha)(t - \beta)$ annule la matrice M , alors on peut appliquer la formule de récurrence (8.7), et on a donc

$$M^k = \left(\frac{\alpha^k \beta - \alpha \beta^k}{\beta - \alpha} \right) I_n + \left(\frac{\beta^k - \alpha^k}{\beta - \alpha} \right) M.$$

Une autre façon d'utiliser un polynôme annulateur est en recourant à la division euclidienne. Plus précisément, supposons que $p(A) = 0$ et que

$$t^k = q(t)p(t) + r(t) \quad \text{avec } \deg(r(t)) < \deg(p(t)).$$

Alors on a clairement

$$A^k = r(A).$$

Exemple. Supposons que $p(t) = t^3 - t^2 + 1$ annule la matrice A . On la division avec reste

$$t^7 = (t^4 + t^3 + t^2 - 1) \cdot p(t) + (-2t^2 + 1),$$

donc $A^7 = -2A^2 + 1$.

C.4 Le cas d'une somme commutative de deux matrices

Supposons que $M \in M_n(K)$ s'écrive

$$M = A + B,$$

avec $AB = BA$. Supposons en outre que les puissances A^k et B^k sont connues. Alors on peut calculer les puissances de M par la formule du binôme de Newton :

$$M^k = (A + B)^k = \sum_{j=1}^k \binom{k}{j} A^j B^{k-j}. \quad (\text{C.2})$$

Attention : cette formule est fausse lorsque A et B ne commutent pas.

Exemple. On considère la matrice

$$M = \begin{pmatrix} a & 1 \\ 0 & a \end{pmatrix} = \begin{pmatrix} a & 0 \\ 0 & a \end{pmatrix} + \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix}.$$

Il est clair que les deux matrices de cette somme commutent. On obtient donc

$$M^k = \begin{pmatrix} a^k & ka^{k-1} \\ 0 & a^k \end{pmatrix}$$

On démontrera au prochain semestre que toute matrice $M \in M_n(\mathbb{C})$ s'écrit $M = D + N$ où D est une matrice diagonalisable, N est une matrice nilpotente et $DN = ND$ (c'est le théorème de Dunford). La formule (C.2) permet alors de calculer en principe toute puissance de n'importe quelle matrice carrée à coefficients complexes.

Annexe D

La notion de groupe quotient

Définition. Soit $(G, \cdot)$ un groupe et $H \subset G$ un sous-groupe. Pour tout $a \in G$, on appelle *classe à gauche de a modulo H* le sous-ensemble de G défini par

$$a \cdot H = \{x = a \cdot h \mid h \in H\}.$$

On définit de même la *classe à droite de a modulo H* le sous-ensemble

$$H \cdot a = \{y = h \cdot a \mid h \in H\}.$$

Lorsque $(G, +)$ est un groupe abélien, dont la loi de groupe est notée additivement, alors la classe à gauche et la classe à droites d'un élément a coïncident et on note cet ensemble $a + H$. Ainsi

$$a + H = \{x = a + h \mid h \in H\}.$$

Proposition. L'ensemble des classes à gauche modulo H forme une partition de G .

Rappelons qu'une *partition* d'un ensemble E est une collection $\mathcal{A} = \{A_i\}_{i \in I} \subset \mathcal{P}(E)$ de sous-ensembles vérifiant qui sont deux-à-deux disjoints et dont la réunion est l'ensemble E tout entier. De façon équivalente, $\mathcal{A}$ est une partition de E si tout élément de E appartient à un et un seul parmi les ensembles A_i .

Preuve de la proposition. Notons

$$\mathcal{Q} = \{aH \mid a \in G\} \subset \mathcal{P}(G)$$

l'ensemble des classes à gauche modulo H . Il est clair que $a = a \cdot e \in a \cdot H$ (où on note $e \in G$ l'élément neutre du groupe), donc la réunion des classes à droites recouvre bien G tout entier, ce que l'on peut écrire

$$G = \bigcup_{a \in G} a \cdot H.$$

Nous devons prouver que les éléments de $\mathcal{Q}$ sont des ensembles deux-à-deux disjoints. Supposons que $x \in aH \cap bH$, alors il existe $h_1, h_2 \in H$ tels que $x = ah_1 = bh_2$. Mais ceci implique que b appartient à la classe aH car

$$b = xh_1^{-1} = ah_1h_2^{-1} = ah,$$

où on a posé $h = h_1h_2^{-1} \in H$ (on sait que $h \in H$ car H est un sous-groupe de G). Mais il est facile de voir que la condition $b \in aH$ entraîne que $bH = aH$. En effet, si $b = ah$ pour un certain élément $h \in H$, alors on a

$$x \in bH \quad \Rightarrow \quad x = bh' \text{ avec } h' \in H \quad \Rightarrow \quad x = (ah)h' = a(hh') \in aH.$$

On a donc $bH \subset aH$ et on montre de façon symétrique que $aH \subset bH$. □

Exemple. La classe à gauche de l'élément neutre $e \in G$ est le sous-groupe H lui-même.

Définition. On dit que deux éléments $x, y \in G$ sont *équivalents à gauche modulo H* et on note $x \sim_H y$ si x et y appartiennent à une même classe à gauche modulo H .

Proposition. La relation $\sim_H$ est une relation d'équivalence sur G et $\mathcal{Q} = G / \sim_H$ est l'ensemble quotient pour cette relation d'équivalence. De plus on a

$$x \sim_H y \quad \Rightarrow \quad y \in xH \quad \Rightarrow \quad x^{-1}y \in H.$$

Nous laissons la preuve en exercice. □

Définition. Le quotient pour cette relation d'équivalence se note $G/H = \mathcal{Q}$. On note $\pi : G \rightarrow G/H$ l'application qui associe à un élément $a \in H$ sa classe d'équivalence (donc $\pi(a) = aH$). Cette application s'appelle la *projection canonique de G sur le quotient G/H* .

Proposition (Formule des classes). Soit G un groupe fini et $H \subset G$ un sous-groupe. Alors on a

$$\text{Card}(G/H) = \frac{\text{Card}(G)}{\text{Card}(H)}.$$

Preuve. On prouve d'abord que toutes les classes ont le même cardinal. En effet, pour tout $a \in G$, l'application

$$\rho_a : H \rightarrow aH \quad \text{définie par} \quad \rho_a(h) = ah$$

est une bijection (son inverse est $\rho_{a^{-1}}$). On a donc pour $a, b \in G$ quelconque

$$\text{Card}(aH) = \text{Card}(H) = \text{Card}(bH).$$

Le nombre de classes à gauche modulo G est fini puisque le groupe G est fini. On peut donc écrire

$$G/H = \{a_1H \cup a_2H \cup \dots \cup a_mH\} \subset \mathcal{P}(G),$$

et supposer que ces classes sont deux-à-deux disjointes. Cela signifie que G s'écrit comme réunion disjointe

$$G = a_1H \cup a_2H \cup \dots \cup a_mH,$$

et donc

$$\text{Card}(G) = \sum_{i=1}^m \text{Card}(a_iH) = m \text{Card}(H) = \text{Card}(G/H) \text{Card}(H).$$

□

En théorie des groupes, le cardinal d'un groupe fini s'appelle l'*ordre* de ce groupe ('ordre' et 'cardinal' sont donc synonymes, mais dans le contexte restreint de la théorie des groupes).

Corollaire. L'ordre de tout sous-groupe d'un groupe fini G divise l'ordre de G .

Exemples. 1. Si l'ordre de G est un nombre premier, alors G ne contient aucun sous-groupe propre (i.e. autre que G lui-même et le groupe trivial $\{e\}$).

2. Un groupe G d'ordre 10 peut posséder des sous-groupes d'ordre 2 et 5 mais aucun sous-groupe d'ordre 3.

Sous-groupe normal.

Dans ce paragraphe, nous nous intéressons au problème suivant : Soient G un groupe et $H \subset G$ un sous-groupe. A quelle condition peut-on définir sur le quotient G/H une structure naturelle de groupe ? Pour répondre à cette question nous avons besoin d'une définition.

Définition. On dit qu'un sous-groupe H d'un groupe G est *normal* (on dit aussi qu'il est *distingué*) si la classe à droite de tout élément coïncide avec sa classe à gauche, c'est-à-dire $aH = Ha$ pour tout $a \in G$.

Remarques. (1) La condition peut se récrire $a^{-1}Ha = H$. Donc $H \subset G$ est un sous-groupe normal si c'est un sous groupe tel que

$$\forall a \in G; \quad [h \in H \Rightarrow a^{-1}ha \in H]. \quad (\text{D.1})$$

(2) Il est clair que tout sous-groupe d'un groupe abélien G est un sous-groupe normal.

Lemme. Si $f : G \rightarrow G'$ est un homomorphisme de groupes, alors $H = \text{Ker}(f) \subset G$ est un sous-groupe normal.

Preuve. Vérifions la condition (D.1). Pour tous $a \in G$ et $h \in H = \text{Ker}(f)$ on a donc $f(h) = e'$ (on note e l'élément neutre de G et e' l'élément neutre de G'). Par conséquent

$$f(aha^{-1}) = f(a)f(h)f(a^{-1}) = f(a)e'f(a^{-1}) = f(a)f(a^{-1}) = f(aa^{-1}) = f(e) = e'.$$

Ainsi $aha^{-1} \in \text{Ker}(f) = H$, ce qui prouve que H est un sous-groupe normal. □

Exemple. Pour tout corps K , le groupe

$$SL_n(K) = \{A \in GL_n(K) \mid \det(A) = 1\}$$

est un sous-groupe de $GL_n(K)$; car c'est le noyau de l'homomorphisme $\det : GL_n(K) \rightarrow K$.

Le théorème suivant répond au problème posé plus haut. Il donne aussi une réciproque au lemme précédent.

Théorème. Soit $(G, \cdot)$ un groupe et $H \subset G$ un sous-groupe. Alors il existe une structure de groupe sur le quotient G/H pour laquelle la projection canonique $\pi : G \rightarrow G/H$ est un homomorphisme si et seulement si H est un sous-groupe normal. Dans ce cas on a $H = \text{Ker}(\pi)$.

Preuve Supposons que $H \subset G$ est un sous-groupe normal et notons

$$[a] = \pi(a) = aH \in G/H$$

la classe de l'élément $a \in G$. On définit une multiplication sur G/H par

$$[a] \cdot [b] := [a \cdot b].$$

Cependant il faut s'assurer que cette opération est bien définie (i.e. qu'elle ne dépend pas du choix des représentants). Soient donc $a' \sim_H a$ et $b' \sim_H b$ (ou de façon équivalente $a' \in [a]$ et $b' \in [b]$). Alors il existe $h_1, h_2 \in H$ tels que $a' = ah_1$ et $b' = bh_2$. On a donc

$$a'b' = ah_1bh_2 = abh'_1h_2 = abh,$$

où on a posé $h'_1 = b^{-1}h_1b$ et $h = h'_1h_2$. Il est important de remarquer que $h'_1 \in H$ par (D.1, car on suppose que H est un sous-groupe normal, on a alors aussi $h \in H$.

On a donc montré que si $a' \sim_H a$ et $b' \sim_H b$, alors $a'b' \sim_H ab$ et donc $[a'b'] = [ab]$. La multiplication des classes est donc bien définie.

La vérification que G/H , avec la multiplication que nous venons de définir, est un groupe est facile :

- (i) Associativité : $([a][b])[c] = [ab][c] = [(ab)c] = [a(bc)] = [a][bc] = [a]([b][c])$.
- (ii) Element neutre : $[e][a] = [ea] = [a]$ and $[a][e] = [ae] = [a]$.
- (iii) Inverse : $[a^{-1}][a] = [a^{-1}a] = [e]$.

En particulier, nous voyons que l'élément neutre de G/H est la classe $[e]$, et donc

$$\text{Ker}(\pi) = \{g \in G \mid \pi(g) = [e]\} = \{g \in G \mid g \sim_h e\} = H.$$

□

Exemple 1. Soit $H = 2\pi\mathbb{Z} \subset \mathbb{R}$ l'ensemble des multiples entiers de 2π . C'est un sous-groupe de $\mathbb{R}$ et en particulier un sous-groupe normal (car $\mathbb{R}$ est abélien). Le groupe quotient $\mathbb{R}/2\pi\mathbb{Z}$ est le groupe des "nombres réels modulo 2π ". On peut y penser comme l'ensemble des arguments (= variable angulaire) du cercle unité dans le plan. Un élément de ce groupe représente donc un angle et on observe que l'application

$$\mathbb{R}/2\pi\mathbb{Z} \rightarrow U(1) = \{z \in \mathbb{C} \mid |z| = 1\}, \quad t \mapsto e^{it}$$

est un isomorphisme de groupes.

Exemple 2. Pour tout $m \in \mathbb{N}$ le sous-groupe $m\mathbb{Z} \subset \mathbb{Z}$ est un sous-groupe normal et le groupe quotient $\mathbb{Z}/m\mathbb{Z}$ est le groupe additif de l'arithmétique modulo m :

$$\mathbb{Z}/m\mathbb{Z} = \{[0], [1], \dots, [m-1]\}, \quad \text{avec l'addition modulo } m.$$

Rappelons que $\mathbb{Z}/m\mathbb{Z}$ est non seulement un groupe abélien, mais c'est aussi un anneau pour la multiplication modulo m (et c'est un corps si et seulement si $m = p$ est un nombre premier, dans ce cas on note $\mathbb{F}_p = \mathbb{Z}/p\mathbb{Z}$).

Exemple 3. Le groupe quotient $GL_n(K)/SL_n(K)$ est isomorphe au groupe multiplicatif $K^* = \{x \in K \mid x \neq 0\}$. L'isomorphisme est donné par l'application déterminant :

$$\det : GL_n(K)/SL_n(K) \rightarrow K^*.$$

Exemple 4. Si V est un espace vectoriel sur un corps K et $W \subset V$ est un sous-espace vectoriel, alors W est un sous-groupe normal pour l'addition (car $(V, +)$ est un groupe abélien). Le quotient s'identifie à l'ensemble des sous-espaces affines de V qui sont parallèles à W :

$$V/W = \{E = a + W \mid a \in V\}.$$

On démontre facilement que le quotient V/W est non seulement un groupe, mais qu'il admet une structure de K -espace vectoriel. De plus, si $U \subset V$ est un sous-espace vectoriel supplémentaire de W dans V , alors V/W est isomorphe à U .