

UNIVERSITÉS FRANCOPHONES



ALGÈBRE

PREMIER CYCLE

MP₁

Saliou Touré



EDICEF/AUPELF

UNIVERSITÉS FRANCOPHONES



ALGÈBRE

PREMIER CYCLE

MP_1

Saliou Touré

EDICEF

58, rue Jean-Bleuzen
92178 VANVES Cedex

Diffusion EDICEF ou ELLIPSES selon pays

© EDICEF, 1991

ISBN 2-850-69697-8

ISSN 0993-3948

La loi du 11 mars 1957 n'autorise, aux termes des alinéas 2 et 3 de l'article 41, que « les copies ou reproductions strictement réservées à l'usage privé du copiste et non destinées à une utilisation collective » d'une part, et, d'autre part, que « les analyses et les courtes citations dans un but d'exemple et d'illustration », toute représentation ou reproduction, intégrale ou partielle, faite sans le consentement de l'auteur, ou de ses ayants-droit ou ayants-cause, est illicite (loi du 11 mars 1957, alinéa 1^{er} de l'article 40). Cette représentation ou reproduction, par quelque procédé que ce soit, constituerait une contrefaçon sanctionnée par les articles 425 et suivants du Code pénal.

SOMMAIRE

INTRODUCTION	7
Chapitre 1 : ENSEMBLES – APPLICATIONS – RELATIONS	9
1.1. Notion de logique	9
1.1.1. Propositions – Connecteurs	9
1.1.2. Quantificateurs	11
1.2. Ensembles	12
1.2.1. Définitions et notations	12
1.2.2. Parties d'un ensemble – Complémentaire	13
1.2.3. Intersection et réunion de deux ensembles	14
1.2.4. Produit d'ensembles	16
1.3. Applications	17
1.3.1. Définitions – Exemples	17
1.3.2. Compositions des applications	18
1.3.3. Applications injectives, surjectives, bijectives	19
1.3.4. Images directes et images réciproques	21
1.3.5. Familles	23
1.3.6. Fonctions de plusieurs variables	24
1.4. Relations dans un ensemble	25
1.4.1. Définitions – Exemples	25
1.4.2. Relations d'équivalence	26
1.4.3. Relations d'ordre	29
1.4.4. Applications monotones – Applications dans un ensemble ordonné	31
1.4.5. Intervalles	32
1.5. Entiers naturels – Ensembles finis	33
1.5.1. L'ensemble des entiers naturels	33
1.5.2. Ensembles finis – Cardinaux	34
1.6. Ensembles dénombrables	37
1.6.1. Définition – Exemples	37
1.6.2. Propriétés élémentaires	38
1.7. Analyse combinatoire	40
1.7.1. Arrangements avec répétition	40
1.7.2. Arrangements sans répétition – Permutations	42
1.7.3. Combinaisons sans répétition	43
Chapitre 2 : LOIS DE COMPOSITION	46
2.1. Généralités	46
2.1.1. Définitions – Notations – Exemples	46
2.1.2. Parties stables – Lois induites	47
2.1.3. Composé de deux parties	47
2.1.4. Translations	48
2.2. Propriétés des lois de composition internes	48
2.2.1. Lois associatives	48
2.2.2. Lois commutatives	49
2.2.3. Élément neutre	50
2.2.4. Éléments symétrisables	52
2.2.5. Distributivité	54
2.2.6. Loi quotient	54
2.3. Morphismes	55
2.3.1. Définition – Exemples	55
2.4. Lois de composition externes	56
2.4.1. Définition – Notation	56
2.4.2. Parties stables – Lois induites	56
2.4.3. Restriction du domaine d'opérateurs	56

COURS D'ALGÈBRE

Chapitre 3 : GROUPE	57
3.1. Généralités	57
3.1.1. Définitions – Exemples	57
3.2. Sous-groupes d'un groupe	59
3.2.1. Définition et caractérisation d'un sous-groupe	59
3.2.2. Sous-groupe engendré par une partie	60
3.3. Morphismes de groupes	62
3.3.1. Définitions – Exemples	62
3.3.2. Propriétés des morphismes de groupes	63
3.4. Groupes-quotients	65
3.4.1. Classes modulo un sous-groupe	65
3.4.2. Groupes-quotients	67
3.4.3. Décomposition canonique d'un homomorphisme	68
3.4.4. Applications aux groupes cycliques	70
3.5. Groupes symétriques	72
3.5.1. Généralités	72
3.5.2. Transpositions	73
3.5.3. Signature d'une permutation	74
3.6. Groupes opérant sur un ensemble	77
3.6.1. Définitions – Exemples	77
3.6.2. Sous-groupe d'isotropie – Orbites	78
Chapitre 4 : ANNEAUX ET CORPS	80
4.1. Structure d'anneaux	80
4.1.1. Définitions – Exemples	80
4.1.2. Règles de calcul dans un anneau	83
4.1.3. Propriétés élémentaires des anneaux	85
4.2. Sous-anneaux – Idéaux – Anneaux-quotients	87
4.2.1. Sous-anneaux	87
4.2.2. Idéaux	88
4.2.3. Anneaux-quotients	90
4.3. Morphismes d'anneaux	91
4.3.1. Définition et propriétés des morphismes d'anneaux	91
4.3.2. Décomposition canonique d'un morphisme d'anneaux	92
4.3.3. Caractéristique d'un anneau	93
4.4. Divisibilité dans un anneau	94
4.4.1. Généralités	94
4.4.2. Plus grand commun diviseur	95
4.4.3. Plus petit commun multiple	98
4.5. Corps	99
4.5.1. Définitions – Propriétés fondamentales	99
4.5.2. Sous-corps – Idéaux d'un corps – Morphismes de corps	100
4.5.3. Corps des fractions d'un anneau commutatif intègre	102
Chapitre 5 : POLYNÔMES ET FRACTIONS RATIONNELLES	104
5.1. Définitions générales	104
5.2. Structure d'anneau de $K[X]$	105
5.2.1. Addition de deux polynômes	105
5.2.2. Multiplication de deux polynômes	106
5.3. Notation définitive	109
5.3.1. Immersion de K dans $K[X]$	109
5.3.2. Notion d'indéterminée	110
5.4. Propriétés arithmétiques de $K[X]$	110
5.4.1. Division euclidienne dans $K[X]$	111
5.4.2. Idéaux de $K[X]$	113
5.4.3. Plus grand commun diviseur	114
5.4.4. Plus petit commun multiple	117
5.4.5. Polynômes irréductibles	118
5.5. Division suivant les puissances croissantes	120
5.6. Fonctions polynômes – Racines d'un polynôme	123

5.6.1. Fonctions polynômes	123
5.6.2. Racines d'un polynôme	124
5.7. Étude de $\mathbb{C}[X]$ et de $\mathbb{R}[X]$	126
5.7.1. Corps algébriquement clos	126
5.7.2. Polynômes de $\mathbb{C}[X]$	128
5.7.3. Polynômes de $\mathbb{R}[X]$	129
5.7.4. Relations entre les coefficients et les racines d'un polynôme	130
5.8. Dérivation des polynômes	131
5.8.1. Dérivée d'un polynôme	131
5.8.2. Formule de Taylor	133
5.9. Polynômes à plusieurs indéterminées	135
5.9.1. Définitions générales	136
5.9.2. Isomorphisme canonique de $K[X][Y]$ sur $K[X,Y]$	137
5.9.3. Degrés d'un polynôme à deux indéterminées	138
5.9.4. Fonctions polynômes	138
5.9.5. Dérivation partielle des polynômes	139
5.10. Définition du corps des fractions rationnelles	141
5.10.1. Fractions rationnelles	142
5.10.2. Fonction rationnelle	144
5.11. Décomposition d'une fraction rationnelle en éléments simples	146
5.11.1. Théorèmes généraux	146
5.11.2. Décomposition en éléments simples d'une fraction rationnelle sur $\mathbb{C}$	149
5.11.3. Décomposition en éléments simples d'une fraction rationnelle sur $\mathbb{R}$	152
Chapitre 6 : ESPACES VECTORIELS	155
6.1. Définition d'un espace vectoriel	155
6.1.1. Définitions	155
6.1.2. Règles de calcul dans un espace vectoriel	156
6.1.3. Exemples d'espaces vectoriels	157
6.2. Sous-espaces vectoriels	159
6.2.1. Définitions – Exemples	159
6.2.2. Intersection de sous-espaces vectoriels. Sous-espace engendré par une partie d'un espace vectoriel	160
6.2.3. Espaces vectoriels quotients	161
6.2.4. Somme de sous-espaces vectoriels	161
6.3. Familles génératrices – Familles libres – Bases	165
6.3.1. Familles génératrices	165
6.3.2. Familles libres	167
6.3.3. Bases d'un espace vectoriel	168
6.3.4. Familles infinies	170
Chapitre 7 : APPLICATIONS LINÉAIRES	172
7.1. Généralités	172
7.1.1. Définitions	172
7.1.2. Exemples	173
7.2. Propriétés des applications linéaires	174
7.2.1. Composée de deux applications linéaires	174
7.2.2. Image et noyau d'une application linéaire	175
7.2.3. Applications linéaires et familles de vecteurs	176
7.2.4. Décomposition canonique d'une application linéaire	180
7.3. L'espace vectoriel $\mathcal{L}(E, F)$	181
7.3.1. Addition dans $\mathcal{L}(E, F)$	182
7.3.2. Produit d'une application linéaire par un scalaire	182
7.3.3. Cas particulier : $E = F$	183
7.4. Projecteurs	184
7.4.1. Définition	184
7.4.2. Propriétés des projecteurs	185
Chapitre 8 : ESPACES VECTORIELS DE DIMENSION FINIE	187
8.1. Le théorème de la dimension	187

COURS D'ALGÈBRE

8.1.1.	Existence de bases en dimension finie	187
8.1.2.	Dimension	188
8.1.3.	Dimension d'un sous-espace vectoriel	192
8.2.	Applications linéaires en dimension finie	195
8.2.1.	Dimension de $\mathcal{L}(E, F)$	195
8.2.2.	Espaces vectoriels isomorphes	197
8.2.3.	Rang d'une application linéaire, d'une famille de vecteurs	198
8.3.	Dualité	201
8.3.1.	Dual d'un espace vectoriel	201
8.3.2.	Bidual d'un espace vectoriel	203
8.3.3.	Orthogonalité	204
8.3.4.	Transposée d'une application linéaire	206
Chapitre 9 : MATRICES		209
9.1.	Généralités	209
9.2.	Matrice d'une application linéaire	209
9.2.1.	Définitions, exemples et théorèmes	209
9.3.	Opérations sur les matrices	215
9.3.1.	Égalité de deux matrices	216
9.3.2.	Addition des matrices	216
9.3.3.	Multiplication d'une matrice par un scalaire	217
9.3.4.	Produit de deux matrices	218
9.4.	Matrices inversibles – Changement de bases	220
9.4.1.	Matrices inversibles	220
9.4.2.	Changement de bases	222
9.4.3.	Matrices équivalentes	227
Chapitre 10 : DÉTERMINANTS		230
10.1.	Applications et formes bilinéaires	230
10.1.1.	Applications et formes bilinéaires alternées	230
10.1.2.	Cas où $\dim(E) = 2$	232
10.1.3.	Déterminant d'ordre 2	233
10.1.4.	Déterminant d'un endomorphisme	233
10.1.5.	Déterminant d'une matrice carrée d'ordre 2	235
10.2.	Applications et formes multilinéaires	236
10.2.1.	Applications et formes multilinéaires alternées	237
10.2.2.	Propriétés des applications et des formes multilinéaires alternées	238
10.3.	Déterminants	242
10.3.1.	Déterminant d'un système de vecteurs	242
10.3.2.	Déterminant d'un endomorphisme	243
10.3.3.	Déterminant d'une matrice carrée	245
10.3.4.	Calculs des déterminants	248
Chapitre 11 : APPLICATIONS DES DÉTERMINANTS		254
11.1.	Calcul de l'inverse d'une matrice carrée	254
11.2.	Détermination du rang	256
11.3.	Systèmes d'équations linéaires	260
11.3.1.	Définitions	260
11.3.2.	Interprétations d'un système d'équations linéaires	260
11.4.	Systèmes de Cramer	261
11.4.1.	Définition	262
11.4.2.	Formules de Cramer	263
11.5.	Résolution d'un système linéaire quelconque	265
11.5.1.	Équations principales – Inconnues principales	265
11.5.2.	Condition de compatibilité et résolution	266
11.6.	Systèmes homogènes	269
PROBLÈMES		271
BIBLIOGRAPHIE		287

Introduction

L'algèbre générale et l'algèbre linéaire sont des outils fondamentaux dans les disciplines mathématiques modernes (analyse, analyse fonctionnelle, probabilité, physique mathématique, etc.). Elles constituent par conséquent des éléments essentiels du bagage mathématique indispensable aux mathématiciens, physiciens, ingénieurs et autres scientifiques.

Le cours d'algèbre que nous soumettons aujourd'hui au public s'adresse aux étudiants en mathématiques du premier cycle des universités et aux étudiants préparant l'entrée dans les grandes écoles scientifiques. Il peut également être utile aux scientifiques qui désirent se recycler en mathématiques et à tous ceux qui veulent acquérir de bonnes connaissances de base en algèbre.

L'expérience montre que le passage de la classe terminale à la première année de faculté constitue pour la majorité des étudiants une difficulté quasi insurmontable. C'est pourquoi nous nous sommes efforcés, dans la rédaction de ce livre, de répondre à une double exigence : d'une part, présenter chaque notion nouvelle à partir du début sans supposer que les étudiants en ont entendu parler et, d'autre part, faire un exposé qui, tout en couvrant l'ensemble des programmes, ne s'y cantonne pas strictement, et soit assez rigoureux et assez riche pour servir de base à une solide formation mathématique. Le livre est ainsi conçu en fonction des besoins immédiats et futurs des étudiants. Sans jamais abandonner la rigueur des démonstrations, nous avons voulu illustrer notre exposé par de nombreux exemples et remarques qui, nous l'espérons, aideront le lecteur à mieux assimiler les notions introduites.

Nous voudrions insister une fois de plus sur le fait que « faire des mathématiques » c'est d'abord et avant tout « résoudre des problèmes ». **Il faut donc, après avoir appris et compris son cours, l'assimiler en résolvant des exercices.** De nombreux exercices et problèmes placés à la fin du livre permettront au lecteur de contrôler l'acquisition de ses connaissances. Certains d'entre eux sont des applications immédiates d'un résultat du cours et doivent servir de test d'assimilation. D'autres, plus difficiles, présentent des compléments qu'il peut être utile de connaître.

Le livre comprend deux grandes parties.

Les chapitres 1 à 5 sont consacrés aux notions d'algèbre générale. Après une étude générale des structures algébriques les plus courantes (relations d'équivalence et d'ordre, groupes, anneaux, corps), le chapitre 5 traite l'anneau des polynômes et le corps des fractions rationnelles.

Les chapitres 6 à 11 sont consacrés à l'algèbre linéaire et abordent l'étude des espaces vectoriels de dimension finie, des applications linéaires, des matrices, des déterminants et des systèmes d'équations linéaires.

COURS D'ALGÈBRE

Il nous paraît bon de signaler que cet ouvrage repose en grande partie sur le cours que nous dispensons à la Faculté des Sciences et Techniques d'Abidjan. Les réactions de plusieurs générations d'étudiants nous ont grandement guidés dans sa mise en forme définitive.

Nous avons beaucoup utilisé plusieurs livres d'Algèbre parus au cours des dernières décennies. On en trouvera une liste non exhaustive dans la bibliographie.

Des erreurs et des imperfections, il y en a certainement. Nous espérons que de nombreux lecteurs et collègues voudront bien nous soumettre critiques et suggestions afin de nous permettre d'apporter les améliorations qui s'imposent à l'occasion des prochaines éditions. D'avance nous les en remercions.

Je remercie vivement Madame Nadine BELLAMY dont les critiques, les suggestions et le goût du travail bien fait m'ont permis d'améliorer la présentation de nombreux points délicats de l'exposé.

Je remercie Monsieur Henri DICI d'avoir lu le manuscrit et de m'avoir signalé quelques imprécisions.

Mes remerciements vont également à Monsieur N'Cho ADOU pour la compétence et le dévouement avec lesquels il a assuré la dactylographie.

Saliou TOURÉ

Chapitre 1 : ENSEMBLES APPLICATIONS RELATIONS

La plupart des notions introduites dans ce chapitre sont déjà connues des étudiants. Notre but est donc de préciser le vocabulaire et les notations qui seront utilisés dans tout le cours.

Nous aurons toujours à considérer des propositions exprimées dans un langage formalisé, mais il ne s'agit pas pour nous de construire une logique formalisée. Nous resterons toujours sur un plan « naïf » et utiliserons le langage usuel, qui doit cependant être précisé.

Le lecteur intéressé par l'étude axiomatique de la logique et de la théorie des ensembles pourra consulter les ouvrages spécialisés consacrés à cette question.

1.1. Notion de logique

1.1.1. PROPOSITIONS. CONNECTEURS LOGIQUES

1.1.1.1. Définition

*On appelle **proposition** un énoncé qui est vrai dans certaines conditions, faux dans d'autres, mais dont on peut toujours dire s'il est vrai ou s'il est faux.*

La propriété essentielle d'une proposition P est donc d'être dotée de l'une des valeurs de vérité V (vrai) ou F (faux).

1.1.1.2. Exemple

« n est un nombre entier et n est multiple de 2 » est une proposition vraie pour les nombres pairs mais fausse pour les nombres impairs.

Nous appellerons **assertion** une proposition qui est toujours vraie ou qui est toujours fausse.

Par exemple, « 10 est un nombre premier » est une assertion fausse.

On appelle **axiome**, dans une théorie mathématique, toute proposition à laquelle on attribue, par convention, la valeur vrai.

On appelle **théorème**, toute proposition dont on démontre qu'elle a la valeur vrai.

A partir des propositions P et Q , on peut former d'autres propositions à l'aide des liaisons **et**, **ou**, **non**, ... appelées **connecteurs logiques**.

Les principaux connecteurs sont :

1.1.1.3. Le connecteur de négation

Si P est une proposition, on note $\neg P$, et on lit «non P », la négation de P .

Par définition, «non P » est vraie si P est fausse, fausse si P est vraie.

La valeur de vérité de $\neg P$ en fonction de celle de P est donnée par un tableau appelé **table de vérité de $\neg$** :

P	$\neg P$
V	F
F	V

1.1.1.4. La conjonction et la disjonction

La **conjonction** est le connecteur logique noté $\wedge$, qui associe à tout couple (P, Q) de propositions, la proposition $(P \wedge Q)$, vraie si et seulement si P et Q sont vraies simultanément.

De même $P \vee Q$, qu'on lit « P ou Q », est vraie si l'une au moins des propositions P, Q est vraie, fausse si P et Q le sont.

Le signe $\vee$ s'appelle le **connecteur de disjonction** ; il se lit «ou».

Les tables de vérité de $\vee$ et $\wedge$ sont :

P	Q	$P \wedge Q$	$P \vee Q$
V	V	V	V
V	F	F	V
F	V	F	V
F	F	F	F

1.1.1.5. L'implication

L'implication est le connecteur logique qui, à tout couple (P, Q) de propositions, associe la proposition $(P \Rightarrow Q)$ (lue « P implique Q ») fausse lorsque P est vraie et Q fausse, vraie dans les autres cas.

En mathématiques il est d'usage de n'écrire que des propositions vraies ; c'est pourquoi, pour traduire « $P \Rightarrow Q$ », on emploie souvent l'expression «si P , alors Q ». Mais la relation $(P \Rightarrow Q)$ peut fort bien être vraie alors que ni P ni Q ne le sont ; par exemple la proposition $(4 \text{ est impair}) \Rightarrow (7 \text{ est pair})$ est vraie.

Les propositions suivantes sont également vraies :
 $(4 \text{ est pair}) \implies (7 \text{ est impair}), (4 \text{ est impair}) \implies (7 \text{ est impair}).$

1.1.1.6. Le connecteur d'équivalence

noté $\iff$:

Si P et Q sont des propositions, on note $P \iff Q$, et on lit « P est équivalente à Q », la proposition :

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

Les tables de vérité des connecteurs $\implies$ et $\iff$ sont :

P	Q	$P \implies Q$	$P \iff Q$
V	V	V	V
V	F	F	F
F	V	V	F
F	F	V	V

En comparant les tables de vérité, on montre facilement que deux propositions P et Q étant données, $(P \implies Q)$ est équivalente à $((\neg Q) \implies (\neg P))$. On dit que $(\neg Q) \implies (\neg P)$ est la **contraposée** de $(P \implies Q)$ ou encore qu'elle est obtenue par **contraposition** de $(P \implies Q)$.

1.1.2. QUANTIFICATEURS

Soit $P(x)$ une proposition contenant un objet x appelé **variable** assujetti à appartenir à un ensemble E appelé **réfèrentiel**.

On convient d'écrire :

$$(\forall x \in E) P(x) \quad \text{ou} \quad (\forall x) P(x)$$

pour exprimer que lorsque x appartient au réfèrentiel E , la proposition P est toujours vraie. On lit «pour tout x , $P(x)$ » ou «quel que soit x , $P(x)$ ».

Le symbole $\forall$ s'appelle le **quantificateur universel**.

Pour exprimer l'assertion «il existe au moins un objet x du réfèrentiel pour lequel $P(x)$ est vraie» on convient d'écrire

$$(\exists x \in E) P(x) \quad \text{ou} \quad (\exists x) P(x).$$

ce qui se lit «il existe au moins un élément x de E tel que " $P(x)$ " ».

Le symbole $\exists$ s'appelle le **quantificateur existentiel**.

Enfin l'expression $\exists! x, < P(x)$ signifie «il existe un et un seul élément x tel que l'assertion $P(x)$ soit vraie»

On utilise très souvent les équivalences logiques suivantes :

$$(\neg((\exists x) P(x))) \iff ((\forall x) (\neg P(x)))$$

$$(\neg((\forall x) P(x))) \iff ((\exists x) (\neg P(x)))$$

$$(\neg((\forall x) (P(x) \implies Q(x)))) \iff (\exists x) (P(x) \text{ et } \neg Q(x))$$

1.1.2.1. Exemple

$$\begin{aligned} &(\forall x \text{ réel}) ((x+1)^2 = x^2 + 2x + 1) \\ &(\exists x \text{ réel}), (x^2 + 3x - 1 = 0) \end{aligned}$$

1.2. Ensembles

1.2.1. DÉFINITIONS ET NOTATIONS

La théorie axiomatique des ensembles est trop délicate pour être exposée au niveau élémentaire auquel nous nous plaçons. Intuitivement, un ensemble est une collection d'objets; ces objets s'appellent les **éléments** ou les **points** de l'ensemble.

Nous désignerons en général les ensembles par des lettres majuscules : E, F, A, B , etc. Les éléments d'un ensemble seront désignés en général par des lettres minuscules : a, b, x, y , etc.

Si a est un élément d'un ensemble E , on écrit $a \in E$ et on lit « a appartient à E » ou « a est élément de E ».

Pour exprimer que a n'est pas un élément de E , on écrit :

$$a \notin E \text{ et on lit «} a \text{ n'appartient pas à } E \text{»}.$$

Nous admettons l'existence d'un ensemble noté $\emptyset$, appelé **ensemble vide**, qui ne contient aucun élément.

Un ensemble réduit à un seul élément a est noté $\{a\}$. Plus généralement, un ensemble qui ne contient que les éléments $x_1, \dots, x_n$ est noté $\{x_1, \dots, x_n\}$.

Si E est un ensemble et P une propriété vraie pour certains éléments de E , l'ensemble des éléments de E qui vérifient la propriété P est souvent noté :

$$\{x : x \text{ vérifie } P\}$$

1.2.1.1. Exemples

Les exemples suivants sont déjà bien connus du lecteur.

$\mathbb{N} = \{0, 1, 2, \dots\}$ est l'ensemble des entiers naturels ;
 $\mathbb{Z} = \{\dots, -3, -2, -1, 0, 1, 2, \dots\}$ est l'ensemble des entiers relatifs ;
 $\mathbb{Q}$ est l'ensemble des nombres rationnels ;
 $\mathbb{R}$ est l'ensemble des nombres réels ;
 $\mathbb{R}^*$ est l'ensemble des nombres réels non nuls ;
 $\mathbb{R}_+$ est l'ensemble des nombres réels positifs ou nuls ;
 $\mathbb{R}_+^*$ désigne l'ensemble des nombres réels strictement positifs ;
 $\mathbb{C}$ est l'ensemble des nombres complexes ;
 $\mathbb{C}^*$ est l'ensemble des nombres complexes non nuls.

1.2.2. PARTIES D'UN ENSEMBLE – COMPLÉMENTAIRE

1.2.2.1. Définition

On dit que l'ensemble E est inclus ou est contenu dans l'ensemble F si tout élément de E est élément de F .

On dit aussi que E est une partie ou un sous-ensemble de F .

On écrit

$$E \subset F \text{ ou } F \supset E$$

Par définition,

$$(E \subset F) \iff (\forall x, x \in E \implies x \in F)$$

Il est immédiat que :

$$E \subset E \text{ quel que soit } E,$$

$$(E \subset F \text{ et } F \subset X) \implies (E \subset X)$$

On dit que l'ensemble E est égal à l'ensemble F , et on note $E = F$, si on a $E \subset F$ et $F \subset E$.

Nous admettons que pour tout ensemble E , il existe un nouvel ensemble appelé ensemble des parties de E , noté $\mathcal{P}(E)$, et dont les éléments sont tous les sous-ensembles de E , y compris l'ensemble vide et E lui-même. Ainsi

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

On note souvent 2^E l'ensemble de parties de E car si E possède n éléments, $\mathcal{P}(E)$ possède 2^n éléments.

1.2.2.2. Définition

Soient E un ensemble et A une partie de E . On appelle complémentaire de A dans E l'ensemble des éléments de E qui n'appartiennent pas à A .

Le complémentaire de A dans E se note

$$E - A \text{ ou } {}^E C_A \text{ ou } C_A$$

s'il n'y a pas de confusion à craindre.

On a donc

$$C_E A = \{x : x \in E \text{ et } x \notin A\}.$$

Les égalités suivantes sont évidentes :

$$C_E (C_E A) = A; C_E E = \emptyset; C_E \emptyset = E.$$

De même, on verrait facilement que si A et B sont des parties d'un ensemble E , et si $A \subset B$, alors

$$C_E B \subset C_E A.$$

1.2.3. INTERSECTION ET RÉUNION DE DEUX ENSEMBLES

On appelle **intersection** de deux ensembles E et F , et on note $E \cap F$, l'ensemble des éléments x tels que $x \in E$ et $x \in F$.

On a donc

$$E \cap F = \{x : x \in E \text{ et } x \in F\}.$$

Si $E \cap F = \emptyset$, on dit que E et F sont **disjoints**.

On a évidemment

$$E \cap \emptyset = \emptyset.$$

On appelle **réunion** de deux ensembles E et F , et on note $E \cup F$, l'ensemble des éléments x tels que $x \in E$ ou $x \in F$.

On a donc

$$E \cup F = \{x : x \in E \text{ ou } x \in F\}.$$

Il est évident que

$$E \cup \emptyset = E \text{ et } E \cup F = \emptyset \implies (E = \emptyset \text{ et } F = \emptyset)$$

Voici quelques propriétés de l'intersection et de la réunion.

- La réunion et l'intersection sont **associatives** :

$$(A \cup B) \cup C = A \cup (B \cup C) \text{ et } (A \cap B) \cap C = A \cap (B \cap C)$$

quels que soient les ensembles A , B et C .

- Elles sont **commutatives** :

$$A \cup B = B \cup A \text{ et } A \cap B = B \cap A$$

quels que soient les ensembles A et B .

- Elles sont **distributives** l'une par rapport à l'autre :

$$\begin{aligned} A \cup (B \cap C) &= (A \cup B) \cap (A \cup C) \\ A \cap (B \cup C) &= (A \cap B) \cup (A \cap C). \end{aligned}$$

- Elles sont **idempotentes** :

$$A \cup A = A \text{ et } A \cap A = A.$$

Les démonstrations sont simples et sont laissées au lecteur.

Il existe des relations simples entre le complémentaire, la réunion et l'intersection. Ces relations sont données par le théorème suivant.

1.2.3.1. Théorème

Soient A et B deux parties d'un ensemble E . Alors on a les égalités suivantes :

$$\mathbf{C}_{(A \cap B)} = (\mathbf{C}_A) \cup (\mathbf{C}_B) \text{ et } \mathbf{C}_{(A \cup B)} = (\mathbf{C}_A) \cap (\mathbf{C}_B).$$

Démonstration. Il suffit de démontrer la première égalité. La deuxième s'en déduit en posant $A_1 = \mathbf{C}_A$, $B_1 = \mathbf{C}_B$ et en utilisant l'égalité

$$\mathbf{C}(\mathbf{C}_A) = A.$$

Soit $x \in \mathbf{C}_{(A \cap B)}$. Alors $x \in E$ et $x \notin A \cap B$; donc $x \in \mathbf{C}_A$ ou bien $x \in \mathbf{C}_B$. Donc $x \in (\mathbf{C}_A) \cup (\mathbf{C}_B)$, d'où l'inclusion

$$\mathbf{C}_{(A \cap B)} \subset (\mathbf{C}_A) \cup (\mathbf{C}_B).$$

Réciproquement, soit x un élément de $(\mathbf{C}_A) \cup (\mathbf{C}_B)$.

Si $x \in \mathbf{C}_A$, $x \notin A$, donc $x \notin A \cap B$ et par suite $x \in \mathbf{C}_{(A \cap B)}$.

De même si $x \in \mathbf{C}_B$, $x \notin B$, donc $x \notin A \cap B$ et par suite $x \in \mathbf{C}_{(A \cap B)}$.

Dans les deux cas, $x \in \mathbf{C}_{(A \cap B)}$ d'où l'inclusion

$$(\mathbf{C}_A) \cup (\mathbf{C}_B) \subset \mathbf{C}_{(A \cap B)}$$

La première égalité est donc démontrée.

1.2.4. PRODUIT D'ENSEMBLES

Soient x et y deux objets. Nous admettrons qu'il est possible de former un troisième objet que l'on note (x, y) et qu'on appelle le couple (x, y) , tel qu'on ait l'équivalence :

$$((x, y) = (x', y')) \iff ((x = x') \text{ et } (y = y')).$$

On dit que x est le **premier élément** et y le **deuxième élément** du couple (x, y) .

L'équivalence précédente montre que l'ordre dans lequel on écrit les deux éléments figurant dans un couple est essentiel. Il ne faut donc pas confondre (x, y) et (y, x) , sauf si $x = y$; de même le couple (x, y) est différent de l'ensemble $\{x, y\}$.

1.2.4.1. Définition

Soient E et F deux ensembles. On appelle **produit cartésien** de E et F , et on note $E \times F$, l'ensemble des couples (x, y) tels que $x \in E$ et $y \in F$.

$$E \times F = \{(x, y) : x \in E \text{ et } y \in F\}.$$

On vérifie facilement les propriétés suivantes :

a) Si A, B, E et F sont quatre ensembles tels que

$$A \subset E \text{ et } B \subset F, \text{ alors } A \times B \subset E \times F.$$

b) Si A, B et C sont des ensembles quelconques, alors

$$A \times (B \cup C) = (A \times B) \cup (A \times C),$$

$$A \times (B \cap C) = (A \times B) \cap (A \times C).$$

c) $(E \times F = \emptyset) \iff (E = \emptyset \text{ ou } F = \emptyset)$
 $(E \times F \neq \emptyset) \iff (E \neq \emptyset \text{ et } F \neq \emptyset).$

Lorsque $E = F$, $E \times E$ se note E^2 et on appelle **diagonale** de E^2 l'ensemble des couples (x, x) avec $x \in E$.

Plus généralement, le produit cartésien de n ensembles $E_1, \dots, E_n$ est l'ensemble

$$E_1 \times \dots \times E_n \text{ encore noté } \prod_{j=1}^n E_j$$

de toutes les suites ordonnées $(x_1, \dots, x_n)$ telles que $x_1 \in E_1, \dots, x_n \in E_n$.

On dit que $(x_1, \dots, x_n)$ est un **n-uple**.

Si $E_1 = \dots = E_n = E$, on note E^n au lieu de $E \times E \times \dots \times E$.

Par exemple, $\mathbb{R} \times \mathbb{R} \times \mathbb{R}$ se note $\mathbb{R}^3$.

1.3. Applications

1.3.1. DÉFINITIONS. EXEMPLES

Soient E et F deux ensembles. On appelle **graphe** de E vers F , toute partie non vide Γ de $E \times F$. Autrement dit, tout élément de Γ est un couple ordonné (x, y) où $x \in E$ et $y \in F$.

L'étude générale des graphes n'est pas notre objectif immédiat. Nous allons appliquer cette notion de graphe à l'étude des **applications d'un ensemble dans un autre ensemble**.

1.3.1.1. Définition

Soient E et F deux ensembles. On appelle **fonction** ou **application** de E dans F , tout triplet $f = (\Gamma, E, F)$ vérifiant les conditions suivantes :

- a) Γ est un graphe de E vers F .
- b) Pour tout $x \in E$, il existe un élément y et un seul de F tel que $(x, y) \in \Gamma$.

On dit que Γ est le **graphe de la fonction** f .

Si $x \in E$, l'unique élément y de F tel que $(x, y) \in \Gamma$ s'appelle l'**image** ou la **valeur de la fonction** f en x ; on la note $f(x)$.

Le graphe Γ de la fonction f est donc l'ensemble des couples $(x, f(x))$ où $x \in E$.

Si $f = (\Gamma, E, F)$ est une fonction, on dit que E est l'**ensemble de départ** (ou de **définition**) et F l'**ensemble d'arrivée**.

Pour exprimer que f est une application de E dans F on utilise les notations suivantes :

$$f : E \longrightarrow F, \quad E \xrightarrow{f} F \quad \text{ou} \quad x \longmapsto f(x).$$

L'ensemble des applications de E dans F se note

$$\mathcal{F}(E, F) \text{ ou } F^E.$$

1.3.1.2. Exemples

a) Si $F = \mathbb{R}$, on dit que f est une **fonction réelle**. Si $E \subset \mathbb{R}$, on dit que f est une **fonction d'une variable réelle**. Par exemple $x \mapsto \sin x$ est une fonction réelle d'une variable réelle.

b) On appelle **application identique** d'un ensemble E , et on note Id_E ou 1_E , l'application qui à tout $x \in E$ fait correspondre x lui-même. On a donc par définition :

$$Id_E(x) = x \text{ pour tout } x \in E.$$

c) Soient E et F deux ensembles. Les applications $(x, y) \mapsto x$ de $E \times F$ dans E et $(x, y) \mapsto y$ de $E \times F$ dans F s'appellent respectivement la **première projection** et la **deuxième projection**. On les note pr_1 et pr_2 respectivement.

d) On dit qu'une application $f : E \longrightarrow F$ est **constante** si l'on a $f(x) = f(y)$ quels que soient $x, y \in E$.

e) Soient $f = (\Gamma, E, F)$ et $g = (\Gamma', E', F')$ deux fonctions. On dit que ces fonctions sont égales si les trois conditions suivantes sont vérifiées :

$$E = E', \quad F = F' \text{ et } f(x) = g(x) \text{ pour tout } x \in E.$$

f) Soit $f = (\Gamma, E, F)$ une fonction et soit A une partie de E . On appelle **restriction** de f à A , et on note $f|_A$, l'application h de A dans F telle que $h(x) = f(x)$ pour tout $x \in A$.

Inversement, étant donné deux fonctions $f = (\Gamma, E, F)$ et $g = (\Gamma', E', F')$, on dit que f est un **prolongement** de g si l'on a les relations

$$E' \subset E, \quad F' \subset F \text{ et } f(x) = g(x) \text{ pour tout } x \in E'.$$

g) Soit E un ensemble. On appelle **fonction caractéristique** de E , la fonction χ_E à valeurs réelles définie par :

$$\chi_E(x) = \begin{cases} 1 & \text{si } x \in E \\ 0 & \text{si } x \notin E \end{cases}$$

1.3.2. COMPOSITIONS DES APPLICATIONS

Soient E, F, G trois ensembles et $f = (\Gamma, E, F)$, $g = (\Gamma', F, G)$ deux applications de E dans F et de F dans G respectivement.

Pour tout $x \in E$, $f(x) \in F$, donc $g(f(x)) \in G$. L'application $x \mapsto g(f(x))$ de E dans G s'appelle l'**application composée** de f et g et se note $g \circ f$ ou gf s'il n'y a pas de confusion possible.

On a donc par définition

$$(g \circ f)(x) = g(f(x)) \text{ pour tout } x \in E.$$

On notera bien que dans l'écriture $g \circ f$, on effectue d'abord l'opération $x \mapsto f(x)$ puis l'opérateur $f(x) \mapsto g(f(x))$.

On définit de même la composée d'un nombre fini d'applications. En particulier, si f est une application de E dans E , on peut former $f \circ f$, $f \circ f \circ f$, etc. Ces applications sont notées f^2 , f^3 , ...

1.3.2.1. Exemples

Prenons $E = F = G = \mathbb{R}$ et $f(x) = \cos x$, $g(x) = x^2 + 1$. On a

$$\begin{aligned}(gof)(x) &= g(f(x)) = \cos^2 x + 1 \\ (fog)(x) &= f(g(x)) = \cos(x^2 + 1),\end{aligned}$$

ce qui montre que $fog \neq gof$ en général.

1.3.2.2. Théorème

Quelles que soient les applications

$$f : E \longrightarrow F, \quad g : F \longrightarrow G, \quad h : G \longrightarrow H,$$

on a

$$(I.3.1) \quad (hog)of = ho(gof).$$

En calculant la valeur du premier membre en un point x quelconque de E , on trouve $h(g(f(x)))$. De même la valeur du second membre au même point x est $h(g(f(x)))$, d'où l'égalité à établir.

Le théorème permet de noter simplement $hogof$ l'application $(hog)of = ho(gof)$.

1.3.3. APPLICATIONS INJECTIVES, SURJECTIVES, BIJECTIVES

1.3.3.1. Définition

Soient E et F deux ensembles et f une application de E dans F .

a) *On dit que f est injective (ou est une injection) si quels que soient $x, y \in E$, la relation $f(x) = f(y)$ entraîne $x = y$, ou encore si la relation $x \neq y$ implique $f(x) \neq f(y)$.*

b) *On dit que f est surjective (ou est une surjection) ou applique E sur F si pour tout $y \in F$ il existe au moins un $x \in E$ tel que $y = f(x)$.*

c) *On dit que f est bijective (ou est une bijection) si elle est à la fois injective et surjective, c'est-à-dire si, pour tout $y \in F$, il existe un et un seul $x \in E$ tel que $y = f(x)$.*

1.3.3.2. Exemples

a) Soit A une partie d'un ensemble E .

L'application $j : A \longrightarrow E$ définie par $j(x) = x$ pour tout $x \in A$ est injective; on l'appelle l'injection canonique de A dans E .

b) On appelle **permutation** d'un ensemble E , toute application bijective de E sur E . L'ensemble des permutations de E se note $\mathcal{S}(E)$.

Si $E = \{1, \dots, n\}$, on écrit $\mathcal{S}_n$ au lieu de $\mathcal{S}(E)$. Cet ensemble sera étudié au Chapitre 3.

c) Si a et b sont des nombres réels et si $a \neq 0$, l'application $f : \mathbb{R} \longrightarrow \mathbb{R}$ définie par $f(x) = ax + b$ pour tout $x \in \mathbb{R}$, est bijective (vérification facile laissée au lecteur).

1.3.3.3. Définition

Soient E et F deux ensembles et $f : E \longrightarrow F$ une application. On dit que f est **inversible** s'il existe une application $g : F \longrightarrow E$ telle que

$$(1.3.3.1) \quad g \circ f = Id_E \quad \text{et} \quad f \circ g = Id_F.$$

Voici une caractérisation des applications inversibles.

1.3.3.4. Théorème

Soit f une application d'un ensemble E dans un ensemble F . Pour que f soit inversible, il faut et il suffit qu'elle soit bijective.

Démonstration. Si f est bijective, alors pour tout $y \in F$, il existe un x unique de E tel que $y = f(x)$. On peut donc définir une application $g : F \longrightarrow E$ en associant à tout $y \in F$ l'unique $x \in E$ tel que $y = f(x)$. Ainsi, si $y = f(x)$, on a $x = g(y)$. Par suite, pour tout $y \in F$ et pour tout $x \in E$, on a

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

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

Donc f est inversible.

Réciproquement, supposons que f soit inversible; donc il existe une application $g : F \longrightarrow E$ telle que $g \circ f = Id_E$ et $f \circ g = Id_F$.

Soient $x \in E$, $x' \in E$ tels que $f(x) = f(x')$. On en déduit $g(f(x)) = g(f(x'))$ d'où, puisque $g \circ f = Id_E$, $x = x'$ et f est injective.

Soit maintenant z un élément de F . Comme $f \circ g = Id_F$, on a $f(g(z)) = z$. Il existe donc au moins un $x \in E$ (à savoir $x = g(z)$) tel que $f(x) = z$, ce qui montre que f est surjective, donc bijective.

1.3.3.5. Remarque

On montrerait de même que l'application g est bijective. De plus g est unique car si g_1 et g_2 sont deux applications de F dans E vérifiant (1.3.3.1), on a

$$g_1 = g_1 \circ Id_F = g_1 \circ (f \circ g_2) = (g_1 \circ f) \circ g_2 = Id_E \circ g_2 = g_2.$$

Ainsi, si $f : E \longrightarrow F$ est bijective, il existe une bijection $g : F \longrightarrow E$ et une seule telle que l'on ait

$$g \circ f = Id_E \text{ et } f \circ g = Id_F.$$

On dit que g est l'**application réciproque** de f et on la note f^{-1} .

1.3.3.6. Théorème

Soient $f : E \longrightarrow F$ et $g : F \longrightarrow G$ deux applications. Si f et g sont injectives (resp. surjectives) il en est de même de $g \circ f$. Si f et g sont bijectives alors $g \circ f$ est bijective et l'on a

$$(1.3.3.2) \quad (g \circ f)^{-1} = f^{-1} \circ g^{-1}.$$

Enfin, si f est bijective, f^{-1} est bijective et on a

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

Démonstration. Supposons que f et g soient injectives et soient $x \in E$, $x' \in E$ tels que $g(f(x)) = g(f(x'))$. Alors $f(x) = f(x')$ puisque g est injective, donc $x = x'$ puisque f est injective; ceci prouve que $g \circ f$ est injective.

Si f et g sont surjectives, pour tout $z \in G$, il existe un $y \in F$ tel que $z = g(y)$, puis un $x \in E$ tel que $y = f(x)$, d'où $z = g(f(x))$, ce qui montre que $g \circ f$ est surjective.

On déduit de ce qui précède que si f et g sont bijectives, il en est de même de $g \circ f$. De plus, on a

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

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

ce qui montre que l'application réciproque de $g \circ f$ est bien $f^{-1} \circ g^{-1}$ (bien noter l'ordre des facteurs).

Enfin, si f est bijective, les relations

$$f \circ f^{-1} = Id_F \text{ et } f^{-1} \circ f = Id_E$$

montrent que f^{-1} est bijective et admet f pour bijection réciproque.

1.3.3.7. Remarque

D'après le Théorème 1.3.3.4, toute application $f : E \longrightarrow E$ telle que $f \circ f = Id_E$ est bijective et $f^{-1} = f$. Une telle bijection s'appelle une **involution** de E .

1.3.4. IMAGES DIRECTES ET IMAGES RÉCIPROQUES

1.3.4.1. Définition

Soient E et F deux ensembles, $f : E \longrightarrow F$ une application, $A \subset E$ et $B \subset F$.

a) On appelle **image directe** de A par f , et l'on note $f(A)$, l'ensemble des $f(x)$ pour $x \in A$.

On a

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

b) On appelle **image réciproque** de B par f , et l'on note $f^{-1}(B)$, l'ensemble des $x \in E$ tels que $f(x) \in B$.

On a

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

L'image $f(E)$ de E s'appelle l'**image** de f et se note $Im(f)$.

On remarquera que la notation $f^{-1}(B)$ ne signifie nullement que l'application réciproque de f existe : il s'agit simplement d'une notation « abusive ».

On a les relations évidentes suivantes :

$$f(\emptyset) = \emptyset, f^{-1}(\emptyset) = \emptyset, f^{-1}(F) = E.$$

$$A \subset f^{-1}(f(A)) \text{ pour toute partie } A \text{ de } E,$$

$$f(f^{-1}(B)) \subset B \text{ pour toute partie } B \text{ de } F.$$

Soit f une application d'un ensemble E dans lui même et soit A une partie de E . On dit que A est **stable par f** si l'on a $f(A) \subset A$. On dit que A est **invariant par f** si $f(A) = A$. L'application $h : A \longrightarrow A$ qui coïncide avec f sur A s'appelle l'**application induite** par f sur A .

Nous donnons ci-après quelques propriétés de l'image directe et de l'image réciproque ; ces propriétés ne doivent pas être apprises par coeur mais doivent être retrouvées rapidement en cas de besoin.

1.3.4.2. Théorème

Soit $f : E \longrightarrow F$ une application.

a) Soient A et A' deux parties de E .

1. Si $A \subset A'$ alors $f(A) \subset f(A')$.

2. Si A et A' sont quelconques, on a

$$f(A \cup A') = f(A) \cup f(A');$$

$$f(A \cap A') \subset f(A) \cap f(A') \text{ avec égalité si } f \text{ est injective.}$$

b) Soient B et B' deux parties de F .

1. Si $B \subset B'$, on a $f^{-1}(B) \subset f^{-1}(B')$.

2. Si B et B' sont quelconques, on a

$$f^{-1}(B \cup B') = f^{-1}(B) \cup f^{-1}(B');$$

$$f^{-1}(B \cap B') = f^{-1}(B) \cap f^{-1}(B').$$

$$3. f^{-1}\left(\bigcap_{F} B\right) = \bigcap_{E} f^{-1}(B).$$

On vérifiera toutes ces propriétés à titre d'exercice.

1.3.5. FAMILLES

Soient E un ensemble et I un autre ensemble non vide appelé **ensemble d'indices**. On appelle **famille d'éléments** de E indexée par I , toute application de I dans E .

Si $x : i \mapsto x(i)$ est une telle famille, on utilise la notation indicielle $i \mapsto x_i$ et on parle de la famille $(x_i)_{i \in I}$ d'éléments de E .

Si I est un ensemble fini, on dit que la famille est **finie**.

Si J est une partie non vide de I , on dit que la famille $(x_i)_{i \in J}$ est une **sous-famille** ou une **famille extraite** de la famille $(x_i)_{i \in I}$.

Si nous prenons $I = \mathbb{N}$, une famille d'éléments de E indexée par $\mathbb{N}$ s'appelle une **suite d'éléments** de E et se note $(x_0, x_1, \dots)$ ou $(x_n)_{n \geq 0}$.

Si $(k_1, k_2, \dots)$ est une suite strictement croissante d'entiers et si on pose $x_{k_1} = z_1, x_{k_2} = z_2, \dots$, on dit que la suite $z_1, z_2, \dots$ est une **suite extraite** de la suite $(x_1, x_2, \dots)$.

On appelle **famille d'ensembles** $(A_i)_{i \in I}$, une famille telle que chaque A_i soit un ensemble.

Nous aurons souvent à considérer une famille $(A_i)_{i \in I}$ de parties d'un ensemble E . La relation $A_i \subset E$ est donc équivalente à $A_i \in \mathcal{P}(E)$; alors la famille $(A_i)_{i \in I}$ de parties de E est la famille des éléments A_i de $\mathcal{P}(E)$.

Pour une famille d'ensembles, on peut généraliser les notions d'intersection, de réunion et de produit de la façon suivante :

Soit $(A_i)_{i \in I}$ une famille d'ensembles.

a) On appelle **intersection** de cette famille, et on note $\bigcap_{i \in I} A_i$, l'ensemble des éléments x tels que $x \in A_i$ pour tout $i \in I$.

b) On appelle **réunion** de cette famille, et on note $\bigcup_{i \in I} A_i$, l'ensemble des éléments x qui appartiennent à l'un au moins des A_i .

Rappelons que le produit cartésien $A_1 \times \dots \times A_n$ des ensembles $A_1, \dots, A_n$ est l'ensemble des n -uples $(x_1, \dots, x_n)$ tels que $x_1 \in A_1, \dots, x_n \in A_n$, ou encore l'ensemble des suites ordonnées $(x_i)_{1 \leq i \leq n}$ telles que $x_i \in A_i$ pour $1 \leq i \leq n$.

Cette remarque va nous permettre de généraliser la notion d'ensemble produit.

Si $A = \bigcup_{i \in I} A_i$ est la réunion de la famille $(A_i)_{i \in I}$, on appelle **produit** de cette famille, et on note $\prod_{i \in I} A_i$, l'ensemble des familles $(x_i)_{i \in I}$ d'éléments de A telles que $x_i \in A_i$ pour tout $i \in I$.

Donc $x \in \prod_{i \in I} A_i$ si et seulement si $x = (x_i)_{i \in I}$ avec $x_i \in A_i$ pour tout $i \in I$.

On dit que x_i est la **composante** ou la **coordonnée**, ou encore la **projection d'indice i** de x . L'ensemble A_i est appelé **facteur d'indice i** du produit $\prod_{i \in I} A_i$.

L'application $(x_i)_{i \in I} \mapsto x_i$ de $\prod_{i \in I} A_i$ dans A_i s'appelle la **projection d'indice i** et se note pr_i .

Pour terminer ce numéro, donnons encore deux définitions.

On dit qu'une famille d'ensembles $(A_i)_{i \in I}$, I non vide, est un **recouvrement** d'un ensemble E si

$$E \subset \bigcup_{i \in I} A_i.$$

On appelle **partition d'un ensemble E** une famille de parties de E , non vides, deux à deux disjointes, dont la réunion est E .

Par exemple si $A \subset E$ ($A \neq \emptyset$ et $A \neq E$) alors A et $\complement_E A$ forment une partition de E .

1.3.6. FONCTIONS DE PLUSIEURS VARIABLES

Soient $E_1, \dots, E_n$ des ensembles, $E_1 \times \dots \times E_n$ leur produit cartésien et $E \subset E_1 \times \dots \times E_n$. Si f est une application de E dans un ensemble F , on dit que f est une **fonction de n variables**.

Si $x = (x_1, \dots, x_n)$ est un point de E , la valeur de f au point x devrait s'écrire $f((x_1, \dots, x_n))$ mais on préfère écrire $f(x_1, \dots, x_n)$.

Il arrive souvent qu'on ait à considérer des fonctions dont les ensembles d'arrivée sont des produits d'ensembles. Considérons par exemple une application

$$f : E \longrightarrow F_1 \times F_2 \times \dots \times F_n$$

où $E, F_1, \dots, F_n$ sont des ensembles.

Pour tout $x \in E$, on a $f(x) \in F_1 \times \dots \times F_n$, donc $f(x)$ est un n -uplet. Si nous désignons par $f_1(x), \dots, f_n(x)$ les coordonnées de $f(x)$, on peut écrire pour chaque $x \in E$,

$$f(x) = (f_1(x), \dots, f_n(x)).$$

On voit donc que l'application f détermine les applications

$$f_i = pr_i \circ f : E \longrightarrow F_i \quad (1 \leq i \leq n),$$

et réciproquement, la donnée des fonctions f_i définit une application f de E dans $F_1 \times \dots \times F_n$ en posant

$$f : x \mapsto (f_1(x), \dots, f_n(x)).$$

1.4. Relations dans un ensemble

Nous rappelons dans ce paragraphe les notions élémentaires de relations, de relation d'équivalence et de relation d'ordre. Nous nous limitons volontairement aux développements indispensables à la compréhension du reste du cours.

1.4.1. DÉFINITIONS. EXEMPLES

1.4.1.1. Définition

Soit E un ensemble. On appelle **relation binaire sur E** , tout couple $\mathcal{R} = (E, \Gamma)$, où Γ est une partie de $E \times E$ que l'on appelle **graphe de la relation $\mathcal{R}$** .

Si $(x, y) \in \Gamma$, on dit que x est **en relation avec y** ; on note $x\mathcal{R}y$, sinon on note $\neg(x\mathcal{R}y)$.

1.4.1.2. Exemples

a) Soit E un ensemble et soient $x, y \in E$. $x\mathcal{R}y$ si et seulement si $x = y$ est une relation binaire sur E .

b) La relation d'inclusion est une relation binaire dans l'ensemble $\mathcal{P}(E)$ des parties d'un ensemble E .

1.4.1.3. Définition

Soient E un ensemble et $\mathcal{R}$ une relation binaire sur E . On dit que :

- a) $\mathcal{R}$ est **réflexive** si pour tout $x \in E$, on a $x\mathcal{R}x$;
- b) $\mathcal{R}$ est **symétrique** si pour tout $x \in E$ et pour tout $y \in E$, $x\mathcal{R}y \implies y\mathcal{R}x$;
- c) $\mathcal{R}$ est **antisymétrique** si pour tout $x \in E$ et pour tout $y \in E$, $(x\mathcal{R}y \text{ et } y\mathcal{R}x) \implies x = y$;
- d) $\mathcal{R}$ est **transitive** si quels que soient $x, y, z \in E$, $(x\mathcal{R}y \text{ et } y\mathcal{R}z) \implies x\mathcal{R}z$.

1.4.1.4. Exemples

a) La relation d'égalité de l'exemple 1.4.1.2 est réflexive: pour tout $x \in E$, on a $x = x$.

b) Dans l'ensemble $\mathcal{P}(E)$ des parties non vides d'un ensemble E , l'inclusion est réflexive, antisymétrique et transitive.

c) Dans $\mathbb{Z}^* = \mathbb{Z} - \{0\}$, la relation $x\mathcal{R}y \iff x$ divise y est réflexive et transitive mais elle n'est ni symétrique ni antisymétrique, ce qui montre au passage que la propriété d'antisymétrie n'est pas la négation de la propriété de symétrie.

1.4.2. RELATIONS D'ÉQUIVALENCE

1.4.2.1. Définition

On dit qu'une relation binaire $\mathcal{R}$ dans un ensemble E est une relation d'équivalence si elle est réflexive, symétrique et transitive.

On note $x\mathcal{R}y$ ou $x \equiv y(\text{mod } \mathcal{R})$ qui se lit « x est équivalent à y modulo $\mathcal{R}$ ».

1.4.2.2. Exemples

a) L'égalité dans E , $(x\mathcal{R}y \iff x = y)$ est une relation d'équivalence.

b) Dans $\mathbb{Z} \times \mathbb{Z}^*$, $(p, q)\mathcal{R}(p', q') \iff pq' = p'q$ est une relation d'équivalence.

c) Soit p un entier ≥ 1 . Dans $\mathbb{Z}$ la relation

$$x \equiv y(\text{mod } p) \iff p \text{ divise } x - y$$

est une relation d'équivalence. En effet si $n \in \mathbb{Z}$, on a $n - n = 0.p$, donc la relation est réflexive. Si n, m et k sont des éléments de $\mathbb{Z}$ tels que $n - m = kp$, on a $m - n = (-k)p$ et la relation est symétrique. Si $n, n', m, k, k' \in \mathbb{Z}$ sont tels que $n - n' = kp$ et $n' - m = k'p$, on a $n - m = (n - n') + (n' - m) = (k + k')p$, donc la relation est transitive.

1.4.2.3. Définition

Soit $\mathcal{R}$ une relation d'équivalence sur un ensemble E . On appelle **classe d'équivalence** d'un élément x de E , et on note $cl(x)$ ou $\overline{x}$ ou $\dot{x}$, l'ensemble des $y \in E$ qui sont équivalents à x modulo $\mathcal{R}$.

L'ensemble des classes d'équivalence s'appelle **l'ensemble quotient** de E par $\mathcal{R}$ et se note $E/\mathcal{R}$. Tout élément d'une classe d'équivalence s'appelle un **représentant** de cette classe.

Par définition de $E/\mathcal{R}$, l'application $\pi : x \longmapsto \dot{x}$ de E dans $E/\mathcal{R}$ est surjective ; on l'appelle **l'application canonique** ou la **surjection canonique**.

1.4.2.4. Exemple

Dans l'exemple 1.4.2.2 c) la classe d'équivalence d'un entier n est l'ensemble

$$\{\dots, n - 2p, n - p, n, n + p, n + 2p, \dots\}$$

qu'on appelle la **classe de congruence de n modulo p** ; une classe de congruence modulo p est aussi appelée un **entier modulo p** .

L'ensemble quotient $E/\mathcal{R}$ se note ici

$$\mathbb{Z}/p\mathbb{Z}.$$

1.4.2.5. Théorème

Soit $\mathcal{R}$ une relation d'équivalence sur un ensemble E . L'ensemble des classes d'équivalence modulo $\mathcal{R}$ forme une partition de E . Réciproquement, toute partition de E définit une relation d'équivalence dont les classes sont les éléments de la partition donnée.

Démonstration. On a $x \in \dot{x}$, puisque la relation est réflexive; donc pour tout $x \in E$, $\dot{x} \neq \emptyset$. Puisque $x \in \dot{x}$, on a

$$E = \bigcup_{x \in E} \{x\} \subset \bigcup_{x \in E} \dot{x} \subset E.$$

Les classes ne sont donc pas vides et leur réunion est l'ensemble E .

Montrons que deux classes d'équivalence $\dot{x}$ et $\dot{y}$ sont identiques si et seulement si $x\mathcal{R}y$.

Supposons d'abord que $x\mathcal{R}y$ et soit $z \in \dot{x}$; alors on a $x\mathcal{R}z$, donc $y\mathcal{R}z$ et par suite $z \in \dot{y}$, ce qui montre que $\dot{x} \subset \dot{y}$. Par symétrie, on verrait de même que $\dot{y} \subset \dot{x}$. Donc $\dot{x} = \dot{y}$.

Inversement supposons que $\dot{x} = \dot{y}$. On a toujours $x \in \dot{x} = \dot{y}$, donc $x\mathcal{R}y$.

Supposons maintenant que $\dot{x} \neq \dot{y}$. S'il existait $z \in E$ tel que $z \in \dot{x} \cap \dot{y}$, alors $z\mathcal{R}x$ et $z\mathcal{R}y$. Comme la relation est symétrique et transitive, on en déduirait $x\mathcal{R}y$, d'où $\dot{x} = \dot{y}$ contrairement à l'hypothèse. Donc $\dot{x} \cap \dot{y} = \emptyset$.

Réciproquement, si $(A_i)_{i \in I}$ est une partition de E , la relation $x\mathcal{R}y$ si et seulement si x et y appartiennent au même A_i est une relation d'équivalence. En effet puisque $\bigcup_{i \in I} A_i = E$, pour tout $x \in E$, il existe un indice i tel que

$x \in A_i$. Donc $x\mathcal{R}x$. La relation $\mathcal{R}$ est évidemment symétrique. Soient x, y, z des éléments de E tels que $x\mathcal{R}y$ et $y\mathcal{R}z$. Alors il existe $i \in I$ tel que $x \in A_i$ et $y \in A_i$, et il existe $j \in I$ tel que $y \in A_j$ et $z \in A_j$. Comme $y \in A_i \cap A_j$, on a $A_i \cap A_j \neq \emptyset$, donc $A_i = A_j$, et par suite $x\mathcal{R}z$, d'où la transitivité de $\mathcal{R}$.

Il est clair que les classes d'équivalence sont les éléments A_i de la partition donnée.

Le résultat suivant sera souvent utilisé dans la suite du cours.

1.4.2.6. Théorème (Décomposition canonique d'une application.)

Soient E et F deux ensembles et $f : E \longrightarrow F$ une application.

a) La relation binaire $\mathcal{R}$ définie sur E par : $x\mathcal{R}y$ si et seulement si $f(x) = f(y)$ est une relation d'équivalence dans E dite associée à f .

b) Soient π la surjection canonique de E sur $E/\mathcal{R}$ et j l'injection canonique de $f(E)$ dans F . Alors il existe une application bijective unique $\bar{f} : E/\mathcal{R} \longrightarrow f(E)$ telle que $f = j \circ \bar{f} \circ \pi$.

Démonstration.

a) On vérifie facilement que $\mathcal{R}$ est une relation d'équivalence dans E .

b) Définissons une application $\bar{f} : E/\mathcal{R} \longrightarrow f(E)$ en posant, pour toute classe $\dot{x} \in E/\mathcal{R}$,

$$\bar{f}(\dot{x}) = f(x)$$

où x est un représentant de la classe $\dot{x}$; autrement dit, on a

$$f = \bar{f} \circ \pi.$$

L'application $\bar{f}$ ne dépend que des classes modulo $\mathcal{R}$ et non du représentant de la classe $\dot{x}$. Si en effet y est un autre représentant de la classe $\dot{x}$, on a $x\mathcal{R}y$, i.e. $f(x) = f(y) = \bar{f}(\dot{x})$.

Montrons que $\bar{f}$ est bijective.

Soient $\dot{x}, \dot{y} \in E/\mathcal{R}$ tels que $\bar{f}(\dot{x}) = \bar{f}(\dot{y})$; si x est un représentant de $\dot{x}$ et y est un représentant de $\dot{y}$, on a

$$f(x) = \bar{f}(\dot{x}) = \bar{f}(\dot{y}) = f(y),$$

donc $x\mathcal{R}y$ et par suite $\dot{x} = \dot{y}$, ce qui prouve que $\bar{f}$ est injective.

Pour tout $y \in f(E)$, il existe un $x \in E$ tel que $y = f(x)$. Donc $y = \bar{f}(\dot{x})$ et $\bar{f}$ est surjective, donc bijective.

S'il existait une autre application $g : E/\mathcal{R} \longrightarrow F$ telle que $f = g \circ \pi$, on aurait pour tout $x \in E$,

$$\bar{f}(\dot{x}) = f(x) = g(\dot{x})$$

d'où $\bar{f} = g$, ce qui prouve l'unicité de $\bar{f}$.

Si enfin j est l'injection canonique de $f(E)$ dans F , il est clair que pour tout $x \in E$, on a

$$j(\bar{f}(\pi(x))) = \bar{f}(\pi(x)) = \bar{f}(\dot{x}) = f(x).$$

Donc $f = j \circ \bar{f} \circ \pi$ et on a le diagramme suivant

$$\begin{array}{ccc} E & \xrightarrow{f} & F \\ \pi \downarrow & & \uparrow j \\ E/\mathcal{R} & \xrightarrow{\bar{f}} & f(E) \end{array}$$

1.4.2.7. Définition

La décomposition $f = j \circ \bar{f} \circ \pi$ s'appelle la décomposition canonique ou la factorisation canonique de f .

La bijection $\bar{f}$ s'appelle l'application induite par f ou encore l'application déduite de f par passage au quotient.

1.4.3. RELATIONS D'ORDRE

1.4.3.1. Définition

On dit qu'une relation binaire dans un ensemble E est une relation d'ordre si elle est réflexive, antisymétrique et transitive.

En général, une relation d'ordre sera notée $\leq$ et on dit que $(E, \leq)$ est un ensemble ordonné.

$x \leq y$ se lit « x est inférieur ou égal à y » ou « x est plus petit que y ».

Si $x \leq y$ et $x \neq y$, on écrit $x < y$ ou $y > x$ et on dit que « x est strictement inférieur à y » ou « y est strictement supérieur à x ».

Soit $(E, \leq)$ un ensemble ordonné. On dit que deux éléments x et y de E sont comparables si $x \leq y$ ou $y \leq x$. Si deux éléments quelconques de E sont comparables, on dit que la relation d'ordre $\leq$ est une relation d'ordre total; $(E, \leq)$ est alors dit totalement ordonné. Dans le cas contraire on dit que $\leq$ est une relation d'ordre partiel et $(E, \leq)$ est dit partiellement ordonné.

1.4.3.2. Exemples

a) Dans $\mathbb{N}$, $\mathbb{Z}$, $\mathbb{Q}$ et $\mathbb{R}$, l'ordre usuel est un ordre total.

b) Soit E un ensemble. La relation d'inclusion est une relation d'ordre partiel (en général) entre éléments de $\mathcal{P}(E)$. On dit que $\mathcal{P}(E)$ est ordonné par inclusion.

c) Dans $\mathbb{N}$, la relation de divisibilité est une relation d'ordre partiel. Par exemple 3 ne divise pas 8 et 8 ne divise pas 3.

Ici encore, notre but n'est pas de faire la théorie générale des ensembles ordonnés. Nous n'exposerons que les notions dont nous aurons besoin ultérieurement.

Soit $(E, \leq)$ un ensemble ordonné et soit A une partie de E . La relation $x \leq y$ entre éléments de A est évidemment une relation d'ordre sur A , appelée relation d'ordre induite sur A par celle de E .

1.4.3.3. Définition

Soient $(E, \leq)$ un ensemble ordonné et A une partie non vide de E .

a) On dit qu'un élément $a \in E$ est un **majorant** (resp. un **minorant**) de A , si $x \leq a$ (resp. $a \leq x$) pour tout $x \in A$.

b) On dit qu'un élément $a \in E$ est un **plus grand** (resp. un **plus petit**) élément de A , si $a \in A$ et si a est un majorant (resp. un minorant) de A .

On dit que A est **majorée** (resp. **minorée**) si A admet des majorants (resp. des minorants). Si A est majorée et minorée, on dit que A est une **partie bornée**.

Remarquons qu'une partie A d'un ensemble ordonné n'admet pas nécessairement un plus grand ou un plus petit élément. Par exemple si $E = \mathbb{N}$ et si A est l'ensemble des nombres pairs, A n'a pas de plus grand élément. Toutefois, si A admet un plus grand (ou un plus petit) élément, celui-ci est unique. En effet, si $a, a' \in A$ sont tels que $x \leq a$ et $x \leq a'$ quel que soit $x \in A$, alors, en particulier, on a : $a \leq a'$ et $a' \leq a$, d'où $a = a'$. On pourra donc parler du plus grand (ou du plus petit) élément de A lorsqu'il existe.

Le plus petit des majorants (resp. le plus grand des minorants) de A , s'il existe, s'appelle la **borne supérieure** (resp. la **borne inférieure**) de A et se note $\sup(A)$ (resp. $\inf(A)$).

Une partie A d'un ensemble ordonné E n'admet pas nécessairement une borne supérieure (resp. inférieure). Toutefois, si A admet une borne supérieure (resp. inférieure), elle est unique mais elle peut ne pas appartenir à A .

Par exemple, si $E = \mathbb{Q}$, ordonné par l'inégalité habituelle, et si $A = \{x \in \mathbb{Q} : x > 0 \text{ et } x^2 < 2\}$, alors l'ensemble A est minoré par n'importe quel nombre rationnel négatif ou nul. On a $\inf(A) = 0$ mais A n'admet pas de borne supérieure dans $\mathbb{Q}$ puisque $\sqrt{2} \notin \mathbb{Q}$.

Si A possède un plus grand (resp. un plus petit) élément a , alors a , qui appartient à A , est la borne supérieure (resp. la borne inférieure) de A . Soit en effet a le plus grand élément de A ; alors a est un majorant de A et si a' est un autre majorant de A , on a : $a \leq a'$ car $a \in A$. Donc a est le plus petit des majorants de A , c'est-à-dire la borne supérieure de A .

On démontrerait de même l'assertion concernant la borne inférieure.

Il faut bien noter qu'une partie peut très bien avoir une borne supérieure (resp. inférieure) sans avoir de plus grand (resp. plus petit) élément. Si en effet nous reprenons $E = \mathbb{Q}$ et $A = \{x \in \mathbb{Q} : x > 0 \text{ et } x^2 < 2\}$, 0 est bien la borne inférieure de A mais A n'admet pas de plus petit élément. De même A n'admet pas de borne supérieure, donc pas de plus grand élément.

Voici une caractérisation de la borne supérieure.

Théorème

Soient $(E, \leq)$ un ensemble totalement ordonné, et A une partie de E . Pour qu'un élément b de E soit la borne supérieure de A , il faut et il suffit que b vérifie les deux conditions.

a) *Pour tout $x \in A$, on a : $x \leq b$.*

b) *Pour tout élément $c \in E$ tel que $c < b$, il existe $x \in A$ tel que $c < x$.*

Démonstration. Si b est la borne supérieure de A , b est un majorant de A . La condition b) est vérifiée car sinon c serait un majorant de A strictement inférieur à b .

Réciproquement, si les deux conditions sont vérifiées, alors b est un majorant de A et tout élément de E strictement inférieur à b n'est pas un majorant de A . Donc b est le plus petit des majorants de A , i.e., la borne supérieure de A .

Soit $(E, \leq)$ un ensemble ordonné. S'il existe un élément $a \in E$ tel que la relation $a \leq x$ entraîne $x = a$ (resp. la relation $x \leq a$ entraîne $x = a$) pour tout $x \in E$, on dit que a est un **élément maximal** (resp. **minimal**) de E .

On dit qu'un élément de E est **extrémal** s'il est ou maximal ou minimal.

Un élément maximal (resp. minimal) de E , s'il existe, est noté $\max E$ (resp. $\min E$).

Par exemple, pour la relation de divisibilité dans $\mathbb{N} - \{0, 1\}$, les éléments minimaux sont les nombres premiers mais il n'y a pas d'élément maximal, ni de plus petit élément.

De même pour la relation d'ordre naturel $\leq$ dans $\mathbb{R}$, il n'y a ni élément minimal, ni élément maximal.

Remarquons que si E est totalement ordonné, les notions d'élément maximal et de plus grand élément coïncident ; mais lorsque E est partiellement ordonné, il n'en est pas ainsi, c'est donc dans ce dernier cas que la notion d'élément maximal présente un intérêt (voir l'exemple de $\mathbb{N} - \{0, 1\}$ ci-dessus).

1.4.4. APPLICATIONS MONOTONES. APPLICATIONS DANS UN ENSEMBLE ORDONNÉ

Soient E et F deux ensembles ordonnés (on notera $\leq$ ou $<$ la relation d'ordre dans E et dans F) et soit f une application de E dans F .

On dit que f est **croissante** (resp. **strictement croissante**) si la relation $x \leq y$ (resp. $x < y$) dans E entraîne la relation $f(x) \leq f(y)$ (resp. $f(x) < f(y)$) dans F .

On dit que f est **décroissante** (resp. **strictement décroissante**) si la relation $x \leq y$ (resp. $x < y$) dans E entraîne la relation $f(x) \geq f(y)$ (resp. $f(x) > f(y)$) dans F .

On dira que f est **monotone** (resp. **strictement monotone**) si f est croissante ou décroissante (resp. strictement croissante ou strictement décroissante).

Soit maintenant f une application d'un ensemble quelconque A dans un ensemble ordonné E . On dit que f est **majorée** (resp. **minorée**) si $f(A)$ est une partie majorée (resp. minorée) dans E . Si f est majorée et minorée dans A , on dit que f est **bornée** dans A .

On appelle **borne supérieure** (resp. **borne inférieure**) de f dans A la borne supérieure (resp. la borne inférieure) dans E (si elle existe) de l'ensemble $f(A)$.

On les note respectivement

$$\sup_{x \in A} f(x) \quad , \quad \inf_{x \in A} f(x)$$

Remarquons que les notations et les définitions précédentes s'appliquent aux familles d'éléments de E . Soit $(x_i)_{i \in I}$ une famille d'éléments de E . Si les bornes supérieure et inférieure de la famille existent, on les note

$$\sup_{i \in I} x_i \quad \text{et} \quad \inf_{i \in I} x_i.$$

En particulier, pour une suite $(x_n)_{n \geq 1}$, on écrit

$$\sup_{n \geq 1} x_n \quad , \quad \inf_{n \geq 1} x_n.$$

1.4.5. INTERVALLES

Soient $(E, \leq)$ un ensemble totalement ordonné, a, b des éléments de E tels que $a \leq b$. On appelle **intervalle fermé** $[a, b]$ l'ensemble des $x \in E$ tels que $a \leq x \leq b$.

Soient a et b des éléments de E tels que $a < b$. On appelle **intervalle ouvert** $]a, b[$, l'ensemble des $x \in E$ tels que $a < x < b$.

On définit de même les **intervalles semi-ouverts**

$$]a, b] = \{x \in E : a < x \leq b\}$$

et

$$[a, b[= \{x \in E : a \leq x < b\}.$$

Par extension, on définit les **demi-droites fermées**

$$[a, \longrightarrow [= \{x \in E : x \geq a\}$$

$$] \longleftarrow, a] = \{x \in E : x \leq a\}$$

et les **demi-droites ouvertes**

$$]a, \longrightarrow [= \{x \in E : x > a\}$$

$$] \longleftarrow, a[= \{x \in E : x < a\},$$

pour tout $a \in E$.

Il est clair que si I est un intervalle ou une demi-droite de E , alors pour tout couple (x, y) d'éléments de I , tout élément compris entre x et y appartient encore à I . On vérifie cette propriété en utilisant la transitivité de la relation d'ordre.

Soit E un ensemble totalement ordonné et soit $x \in E$. On appelle **prédécesseur** $'x$ de x tout élément tel que $]x, x[= \emptyset$. On notera que si $'x$ existe, il est unique.

De même le **successeur** x' de x est un élément tel que $]x, x'[= \emptyset$. Si x' existe, il est unique.

1.5. Entiers naturels. Ensembles finis

Dans ce paragraphe, nous allons donner quelques propriétés des entiers naturels. Nous n'exposerons pas ici les détails de la construction axiomatique de l'ensemble $\mathbb{N}$ des entiers naturels ce qui déborderait largement le cadre de ce cours. Nous admettons donc l'existence de $\mathbb{N}$ et nous supposons que le lecteur est familiarisé avec l'addition, la multiplication des entiers et les propriétés usuelles de ces opérations.

Nous étudierons ensuite quelques propriétés élémentaires des ensembles finis.

1.5.1. L'ENSEMBLE DES ENTIERS NATURELS

Nous admettrons qu'il existe un ensemble non vide et ordonné, noté $\mathbb{N}$, appelé **ensemble des entiers naturels**, et vérifiant les axiomes suivants :

- (N_1) Toute partie non vide de $\mathbb{N}$ admet un plus petit élément.
 - (N_2) Toute partie non vide et majorée de $\mathbb{N}$ admet un plus grand élément.
 - (N_3) $\mathbb{N}$ n'a pas de plus grand élément.
- Le plus petit élément de $\mathbb{N}$ est noté 0.

Conséquences de la définition

a) Toute partie $\{n, m\}$ à deux éléments de $\mathbb{N}$ admet un plus petit élément, donc $\mathbb{N}$ est totalement ordonné.

b) Tout élément $a \in \mathbb{N} - \{0\}$ admet un prédécesseur.

Posons

$$N(a) = \{n \in \mathbb{N} : n < a\}.$$

On a $N(a) \neq \emptyset$ car $0 \in N(a)$. L'ensemble non vide $N(a)$ est majoré par a ; il admet donc un plus grand élément que nous noterons α . On a évidemment $\alpha < a$ et $\alpha, a[= \emptyset$, donc α est le prédécesseur de a .

c) Tout élément $a \in \mathbb{N}$ admet un successeur.

Posons en effet

$$N'(a) = \{n \in \mathbb{N} : a < n\}.$$

$N'(a)$ est non vide car si $N'(a)$ était vide, $\mathbb{N}$ admettrait un plus grand élément contrairement à l'axiome (N_3). D'après (N_1), $N'(a)$ admet un plus petit élément que nous noterons β . Il est clair que $a < \beta$ et $a, \beta[= \emptyset$, ce qui montre que β est le successeur de a .

Le successeur de 0 est noté 1, celui de 1 est noté 2, etc.

On pose

$$\mathbb{N}^* = \mathbb{N} - \{0\}.$$

Ces remarques montrent que l'application qui, à tout élément de $\mathbb{N}$ associe son successeur, est une bijection croissante de $\mathbb{N}$ sur $\mathbb{N}^*$.

Nous allons introduire le raisonnement par récurrence qui joue un rôle essentiel en mathématiques. Il est fondé sur le résultat suivant appelé **théorème de récurrence**.

1.5.1.1. Théorème

Soit $P(n)$ une propriété dépendant de l'entier n . Supposons que

a) *$P(0)$ est vraie.*

b) *Pour tout $n \in \mathbb{N}$, la relation $P(n)$ vraie implique $P(n+1)$ est vraie. Alors $P(n)$ est vraie pour tout $n \in \mathbb{N}$.*

Démonstration. Considérons l'ensemble A des éléments n de $\mathbb{N}$ tels que $P(n)$ soit vraie. C'est une partie de $\mathbb{N}$, non vide puisque $0 \in A$. Il s'agit de montrer que $A = \mathbb{N}$, ce qui revient au même, que $\mathbb{N} - A = \emptyset$. Supposons que $\mathbb{N} - A \neq \emptyset$. Alors d'après l'axiome (N_1) , $\mathbb{N} - A$ possède un plus petit élément m et comme $0 \notin \mathbb{N} - A$, on a $m \in \mathbb{N}^*$. Soit a le prédécesseur de m . Puisque m est le plus petit élément de $\mathbb{N} - A$, $a \in A$. Alors d'après l'hypothèse **b)**, $a \in A$ entraîne $m \in A$ contrairement au fait que $m \in \mathbb{N} - A$. Donc $\mathbb{N} - A = \emptyset$ et par suite $A = \mathbb{N}$.

1.5.1.2. Exemple

Soient x un nombre réel positif et n un entier. Montrons que $(1+x)^n \geq 1+nx$.

Nous allons démontrer cette propriété par récurrence sur n . Soit $P(n)$ la propriété

$$(\forall x \in \mathbb{R}_+) ((1+x)^n \geq 1+nx).$$

$P(0)$ est vraie car $(\forall x \in \mathbb{R}_+) ((1+x)^0 \geq 1)$.

Montrons que $P(n) \implies P(n+1)$. Supposons $P(n)$ vraie, i.e. $(1+x)^n \geq 1+nx$. Multiplions les deux membres de cette inégalité par $1+x$ qui est positif; on obtient

$$(1+x)^{n+1} \geq (1+nx)(1+x) = 1 + (n+1)x + nx^2.$$

Comme $1 + (n+1)x + nx^2 \geq 1 + (n+1)x$, il vient

$$(1+x)^{n+1} \geq 1 + (n+1)x,$$

et la propriété $P(n+1)$ est vérifiée.

1.5.2. ENSEMBLES FINIS. CARDINAUX

1.5.2.1. Définition

*Soient E et F deux ensembles. On dit que E et F sont **équipotents** s'il existe une bijection de E sur F .*

On vérifie facilement que l'équipotence est une relation réflexive, symétrique et transitive entre ensembles. Toutefois, il ne s'agit pas d'une relation binaire sur un ensemble dont les éléments sont tous les ensembles ; on montre qu'un tel ensemble n'existe pas. On ne peut donc parler de relation d'équivalence ni *a fortiori* de classes d'équivalence.

1.5.2.2. Définition

On dit qu'un ensemble E est fini s'il est vide ou s'il existe un entier naturel $n \geq 1$ tel que E soit équipotent à l'intervalle $[1, n]$ de $\mathbb{N}$.

Un ensemble qui n'est pas fini est dit infini.

1.5.2.3. Théorème

Soient n et m deux entiers ≥ 1 . Pour qu'il existe une injection de $[1, n]$ dans $[1, m]$, il faut et il suffit que $n \leq m$.

Démonstration. Si $n \leq m$, on a $[1, n] \subset [1, m]$. Alors l'injection canonique de $[1, n]$ dans $[1, m]$ est injective.

Pour démontrer la réciproque, nous allons raisonner par l'absurde. Supposons qu'il existe une injection de $[1, n]$ dans $[1, m]$ avec $m < n$. Alors l'ensemble E des entiers $n \geq 1$ tels qu'il existe un entier $m \geq 1$ et une injection de $[1, n]$ dans $[1, m]$ avec $m < n$ est non vide ; donc E admet un plus petit élément que nous noterons n_0 . Soit m_0 un entier non nul tel que $m_0 < n_0$ et soit f une injection de $[1, n_0]$ dans $[1, m_0]$. Comme $1 \leq m_0 < n_0$, on a $n_0 \neq 1$, d'où $f(n_0) \neq f(1)$ puisque f est injective. On en déduit que $m_0 > 1$ car si on avait $m_0 = 1$, f serait constante et égale à 1. La restriction f_1 de f à $[1, n_0 - 1]$ est injective comme restriction d'une application injective.

Définissons une application $h : [1, m_0] - f(n_0) \longrightarrow [1, m_0 - 1]$ en posant

$$h(n) = \begin{cases} n & \text{si } n < f(n_0) \\ n - 1 & \text{si } n > f(n_0). \end{cases}$$

Alors h est injective. Donc l'application $h \circ f_1$, composée de deux injections, est une injection de $[1, n_0 - 1]$ dans $[1, m_0 - 1]$; ce résultat contredit la définition de n_0 car $m_0 - 1 < n_0 - 1$ puisque $m_0 < n_0$. Le théorème est donc démontré.

1.5.2.4. Corollaire

Soient n et m deux entiers ≥ 1 . Pour qu'il existe une bijection de $[1, n]$ sur $[1, m]$, il faut et il suffit que $n = m$.

On en déduit que si un ensemble E est équipotent à $\{1, \dots, n\}$ et à $\{1, \dots, m\}$, alors $n = m$.

Nous pouvons donc poser la définition suivante :

1.5.2.5. Définition

Soit E un ensemble fini non vide. On appelle **cardinal de E ou nombre d'éléments de E** , l'unique entier naturel $n \geq 1$ tel que E soit équipotent à $[1, n]$.

On écrit $\text{Card}(E) = n$ ou $|E| = n$.

Par définition, $\text{Card}(\emptyset) = 0$.

On démontre qu'une partie non vide de $\mathbb{N}$ est finie si et seulement si elle est majorée. On en déduit que l'ensemble $\mathbb{N}$ est infini.

Nous allons énoncer quelques propriétés élémentaires des ensembles finis. Ces propriétés se démontrent facilement à partir des définitions et des résultats précédents.

1.5.2.6. Théorème

a) Si E et F sont deux ensembles finis équipotents, on a $\text{Card}(E) = \text{Card}(F)$.

b) Si F est une partie d'un ensemble fini E , alors F est fini et $\text{Card}(F) \leq \text{Card}(E)$. Si de plus, $\text{Card}(F) = \text{Card}(E)$, alors $F = E$.

c) L'intersection d'une famille finie ou infinie d'ensembles finis est finie.

d) Si E et F sont deux ensembles finis, alors $E \cup F$ est fini et

$$\text{Card}(E \cup F) + \text{Card}(E \cap F) = \text{Card}(E) + \text{Card}(F).$$

e) Si E et F sont deux ensembles finis, alors $E \times F$ est fini et

$$\text{Card}(E \times F) = \text{Card}(E) \cdot \text{Card}(F).$$

Plus généralement, si $E_1, \dots, E_p$ sont des ensembles finis, on

$$\text{Card}(E_1 \times \dots \times E_p) = \text{Card}(E_1) \cdot \text{Card}(E_2) \dots \text{Card}(E_p).$$

Voici encore quelques propriétés des applications d'un ensemble fini dans un autre ensemble fini.

1.5.2.7. Théorème

Soient E et F deux ensembles finis et f une application de E dans F . Alors :

a) $\text{Card}(f(E)) \leq \inf(\text{Card}(E), \text{Card}(F))$.

b) $\text{Card}(f(E)) = \text{Card}(E)$ si et seulement si f est injective.

c) $\text{Card}(f(E)) = \text{Card}(F)$ si et seulement si f est surjective.

d) Si $\text{Card}(E) = \text{Card}(F)$, alors f est injective si et seulement si f est surjective.

Démonstration

a) Il est clair que $\text{Card}(f(E)) \leq \text{Card}(F)$ puisque $f(E) \subset F$. Donc $f(E)$ est un ensemble fini. Posons

$$f(E) = \{x_1, \dots, x_n\}.$$

On peut écrire :

$$E = \bigcup_{i=1}^n f^{-1}(\{x_i\}).$$

Ainsi E apparaît comme une réunion d'ensembles disjoints deux à deux. Comme $f^{-1}(\{x_i\}) \neq \emptyset$, on a (Théorème d), $\text{Card}(E) \geq n$, c'est-à-dire $\text{Card}(E) \geq \text{Card}(f(E))$.

b) Conservons les notations du a). Alors f est injective si et seulement si pour tout $i \in \{1, \dots, n\}$, on a $\text{Card}(f^{-1}(\{x_i\})) = 1$, ce qui signifie que $\text{Card}(E) = n = \text{Card}(f(E))$.

c) résulte du Théorème 1.5.2.6, b).

d) découle immédiatement des résultats précédents. $\square$

1.6. Ensembles dénombrables

1.6.1. DÉFINITION. EXEMPLES

1.6.1.1. Définition

On dit qu'un ensemble E est dénombrable s'il est équipotent à $\mathbb{N}$.

On dit qu'un ensemble E est au plus dénombrable s'il est fini ou s'il est dénombrable.

Si un ensemble E est dénombrable, il est infini comme $\mathbb{N}$. L'existence d'une bijection $f : \mathbb{N} \longrightarrow E$ permet de numérotier les éléments de E ; en notant x_n au lieu de $f(n)$, on peut donc ranger les éléments de E en une suite $x_0, x_1, \dots, x_n$.

1.6.1.2. Exemple

L'ensemble P des entiers pairs est dénombrable.

En effet, l'application $f : \mathbb{N} \longrightarrow P$ définie par $f(n) = 2n$ est bijective.

1.6.1.3. Exemple

L'ensemble $\mathbb{N}^*$ est dénombrable car l'application $f : \mathbb{N} \longrightarrow \mathbb{N}^*$ définie par $f(n) = n + 1$ est bijective par définition même de $\mathbb{N}$.

1.6.1.4. Exemple

$\mathbb{Z}$ est dénombrable.

En effet, l'application $f : \mathbb{N} \longrightarrow \mathbb{Z}$ définie par

$$f(n) = \begin{cases} \frac{n}{2} & \text{si } n \text{ est pair} \\ -\frac{n+1}{2} & \text{si } n \text{ est impair} \end{cases}$$

est bijective.

1.6.1.5. Exemple

$\mathbb{N} \times \mathbb{N}$ est dénombrable.

Ecrivons en effet les éléments de $\mathbb{N} \times \mathbb{N}$ dans un tableau infini à double entrée :

	0	1	2	...	p	...
0	(0, 0)	(0, 1)	(0, 2)	...	(0, p)	...
1	(1, 0)	(1, 1)	(1, 2)	...	(1, p)	...
2	(2, 0)	(2, 1)	(2, 2)	...	(2, p)	...
$\vdots$	$\vdots$					
q	(q , 0)	(q , 1)	(q , 2)	...	(q , p)	...
$\vdots$	$\vdots$					

On obtient la suite :

$$x_0 = (0, 0), \quad x_1 = (1, 0), \quad x_2 = (0, 1), \quad x_3 = (2, 0), \quad x_4 = (1, 1), \quad \dots$$

d'où une bijection de $\mathbb{N}$ sur $\mathbb{N} \times \mathbb{N}$.

1.6.2. PROPRIÉTÉS ÉLÉMENTAIRES**1.6.2.1. Théorème**

Toute partie A d'un ensemble dénombrable E est au plus dénombrable.

Démonstration. Par hypothèse, il existe une bijection f de E sur $\mathbb{N}$. La restriction de f à A est encore une bijection de A sur une partie $f(A)$ de $\mathbb{N}$. Il suffit donc de démontrer que toute partie infinie B de $\mathbb{N}$ est dénombrable.

Pour cela, définissons par récurrence une application $k \mapsto n_k$ de $\mathbb{N}$ dans B de la façon suivante ; n_0 est le plus petit élément de B et pour $k \geq 1$, n_k est le plus petit élément de $B - \{n_0, n_1, \dots, n_{k-1}\}$ qui est une partie non vide de B puisque B est infinie par hypothèse.

Cette application est injective car si $k < p$, alors par construction, $n_p \notin \{n_0, n_1, \dots, n_k\}$, et donc $n_p \neq n_k$. Pour montrer qu'elle est surjective, nous allons raisonner par l'absurde et supposer qu'il existe un élément $a \in B$ tel que $a \neq n_k$ pour tout $k \in \mathbb{N}$. Alors $a \in B - \{n_0, n_1, \dots, n_{k-1}\}$ pour tout $k \in \mathbb{N}^*$ et, par définition de n_k , on a $n_k \leq a$. Comme d'autre part $n_0 \leq a$, on a : $n_k \leq a$ pour tout $k \in \mathbb{N}$. Ainsi l'image de $\mathbb{N}$ par l'application injective $k \mapsto n_k$ est contenue dans l'intervalle $[n_0, a]$ de $\mathbb{N}$ et $\mathbb{N}$ est équipotent à cet intervalle. Comme $[n_0, a]$ est un ensemble fini, cela contredit le fait que $\mathbb{N}$ est infini.

Ainsi l'application $k \mapsto n_k$ est surjective ; elle est donc bijective d'où notre assertion. $\square$

1.6.2.2. Lemme

Si E est un ensemble dénombrable et si f est une application bijective de E sur un ensemble F , alors F est dénombrable.

C'est évident, car si φ est une bijection de $\mathbb{N}$ sur E , alors $f \circ \varphi$ est une bijection de $\mathbb{N}$ sur F .

1.6.2.3. Lemme

Soient E un ensemble dénombrable et f une application surjective de E sur un ensemble F . Alors F est au plus dénombrable.

Démonstration. Puisque E est équipotent à $\mathbb{N}$, on peut supposer que $E = \mathbb{N}$. Comme f est surjective, pour tout $y \in F$, la partie $f^{-1}(\{y\})$ de $\mathbb{N}$ est non vide. Soit $m(y)$ le plus petit élément de $f^{-1}(\{y\})$; on définit ainsi une application $m : F \rightarrow \mathbb{N}$. Comme $f \circ m = Id_F$, l'application m est injective. L'application $y \mapsto m(y)$ est une bijection de F sur la partie $m(F)$ de $\mathbb{N}$. $m(F)$ étant au plus dénombrable d'après le Théorème 1.6.2.1, F est au plus dénombrable.

1.6.2.4. Corrolaire

Tout ensemble quotient d'un ensemble dénombrable est dénombrable.

1.6.2.5. Théorème

Si $E_1, E_2, \dots, E_n$ sont des ensembles dénombrables, l'ensemble $E_1 \times E_2 \times \dots \times E_n$ est dénombrable.

Démonstration. Il suffit de démontrer le théorème si $n = 2$; le cas général s'en déduit par récurrence.

Soit $m \mapsto x_m$ (resp. $k \mapsto y_k$) une bijection de $\mathbb{N}$ sur E_1 , (resp. sur E_2). Alors $(m, k) \mapsto (x_m, y_k)$ est une bijection de $\mathbb{N} \times \mathbb{N}$ sur $E_1 \times E_2$ donc, puisque $\mathbb{N} \times \mathbb{N}$ est dénombrable, $E_1 \times E_2$ est dénombrable (Lemme 1.6.7).

1.6.2.6. Théorème

Soit $(E_1, E_2, \dots)$ une suite finie ou infinie d'ensembles dénombrables. Alors $E = \bigcup_{i=1}^{\infty} E_i$ est un ensemble dénombrable.

Démonstration. Pour tout n fixé, il existe une bijection $f_n : \mathbb{N} \longrightarrow E_n$. Définissons une application $g : \mathbb{N} \times \mathbb{N} \longrightarrow E$ en posant $g(n, m) = f_n(m)$. Cette application est surjective. En effet, pour tout $x \in E$, il existe au moins un $i \in \mathbb{N}$ tel que $x \in E_i$; si on pose $k = f_i^{-1}(x)$, alors $g(i, k) = x$.

D'après le Lemme 1.6.2.3, E est au plus dénombrable; mais puisque un des ensembles E_n est infini, E est lui-même infini.

1.6.2.7. Corollaire

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

Démonstration. L'ensemble A des couples (p, q) d'entiers $p > 0$ et $q > 0$ est dénombrable comme sous-ensemble infini de $\mathbb{N} \times \mathbb{N}$. L'ensemble X des nombres rationnels positifs, qui est l'image de A par l'application $(p, q) \longmapsto p/q$, est donc dénombrable d'après le Lemme 1.6.2.3. De même l'ensemble Y des nombres rationnels négatifs est dénombrable (considérer l'application $(p, q) \longmapsto -p/q$). Comme $\mathbb{Q} = \{0\} \cup X \cup Y$, on voit que $\mathbb{Q}$ est dénombrable. $\square$

On démontre que l'ensemble $\mathbb{R}$ des nombres réels n'est pas dénombrable (Théorème de Cantor) et même plus précisément, que si $(a, b) \in \mathbb{R}^2$ avec $a < b$, alors $]a, b[$ n'est pas dénombrable.

1.7. Analyse combinatoire

Dans ce paragraphe, nous allons aborder les problèmes de dénombrement; il s'agira de calculer les cardinaux d'ensembles finis.

1.7.1. ARRANGEMENTS AVEC RÉPÉTITION

1.7.1.1. Définition

Soit E un ensemble fini non vide et soit p un entier naturel non nul. On appelle **arrangement de p éléments de E avec répétition** toute application de $\{1, 2, \dots, p\}$ dans E .

Le théorème suivant donne le nombre d'arrangements avec répétition des éléments d'un ensemble fini.

1.7.1.2. Théorème

Soient F un ensemble fini de cardinal p et E un ensemble fini de cardinal n . L'ensemble $\mathcal{F}(F, E)$ des applications de F dans E est fini et a pour cardinal n^p .

Démonstration. Soient $F = \{b_1, \dots, b_p\}$ et $E = \{a_1, \dots, a_n\}$.

Pour définir une application f de F dans E , il suffit de se donner les éléments $f(b_1), \dots, f(b_p)$ de E .

Nous allons calculer le nombre d'applications de F dans E par récurrence sur p , n étant fixé.

Si $p = 1$, on a $F = \{b_1\}$. Les applications de F dans E sont alors définies par

$$b_1 \mapsto f_i(b_1) = a_i, \quad (1 \leq i \leq n);$$

donc $\mathcal{F}(F, E)$ est fini et a pour cardinal n .

Supposons la propriété vraie pour $\text{Card}(F) = p - 1$ et démontrons-la pour $\text{Card}(F) = p$. Soit b un élément de F ; posons

$$F = F' \cup \{b\}, \quad \text{où } \text{Card}(F') = p - 1.$$

Une application de F dans E est définie de façon unique par sa restriction à F' et par l'image de b . Il y a n images possibles pour b et, d'après l'hypothèse de récurrence, il y a n^{p-1} restrictions possibles à F' . Le nombre d'applications de F dans E est donc $n \cdot n^{p-1} = n^p$.

1.7.1.3. Corollaire

Si E est un ensemble fini de cardinal $n \geq 1$, alors l'ensemble $\mathcal{P}(E)$ des parties de E est fini et a pour cardinal 2^n .

Démonstrations. A toute partie A de E associons sa fonction caractéristique $\chi_A : A \longrightarrow \{0, 1\}$ définie par

$$\chi_A(x) = \begin{cases} 1 & \text{si } x \in A \\ 0 & \text{si } x \in E - A \end{cases}$$

Réciproquement, à toute fonction $\varphi \in \mathcal{F}(E, \{0, 1\})$ associons la partie A de E définie par $A = \varphi^{-1}(\{1\})$. Alors on a $\varphi = \chi_A$.

On a ainsi établi une bijection $A \longrightarrow \chi_A$ de $\mathcal{P}(E)$ sur $\mathcal{F}(E, \{0, 1\})$ et il suffit d'appliquer le Théorème 1.7.1.2.

1.7.1.4. Remarque

C'est à cause de la formule donnée au Théorème 1.7.1.2 que l'ensemble $\mathcal{F}(F; E)$ de toutes les applications de F dans E se note souvent E^F .

De même, en raison de la formule du Corollaire 1.7.1.3, $\mathcal{P}(E)$ est souvent noté 2^E .

1.7.1.5. Théorème

a) Soient E un ensemble fini et $(A_i)_{1 \leq i \leq n}$ une partition finie de A . Alors

$$\text{Card}(A) = \sum_{i=1}^n \text{Card}(A_i);$$

b) (Principe des bergers) Soient E et F des ensembles finis et f une application surjective de E sur F . On suppose que pour tout $y \in F$, $\text{Card}(f^{-1}(\{y\})) = p$. Alors $\text{Card}(E) = p \cdot \text{Card}(F)$.

Démonstration. a) Comme $A = \bigcup_{i=1}^n A_i$ et $A_i \cap A_j = \emptyset$ si $i \neq j$, on a par récurrence sur n (le cas $n = 2$ étant évident) :

$$\text{Card}(A) = \sum_{i=1}^n \text{Card}(A_i).$$

b) Il suffit d'appliquer a) à la partition de E définie par la relation d'équivalence associée à $f : x \mathcal{R} y$ si et seulement si $f(x) = f(y)$.

1.7.2. ARRANGEMENTS SANS RÉPÉTITION. PERMUTATIONS**1.7.2.1. Définition**

Soit E un ensemble fini de cardinal n et soit p un entier naturel ≥ 1 . On appelle **arrangement de p éléments de E sans répétition**, toute application injective de $\{1, \dots, p\}$ dans E .

Le nombre de ces arrangements est noté A_n^p .

On remarquera qu'il ne peut exister d'injection de $A = \{1, \dots, p\}$ dans E que si $p \leq n$; en effet, si f est une injection de A dans E , c'est une bijection de A sur $f(A) \subset E$, donc (Théorème 1.5.2.7, b)) $\text{Card}(f(A)) = p$. On a donc $p \leq n$ (Théorème 1.5.2.7, a)). Nous supposons donc que $p \leq n$.

1.7.2.2. Théorème

Soient F un ensemble fini de cardinal p et E un ensemble fini de cardinal n , avec $p \leq n$. L'ensemble des applications injectives de F dans E est fini et possède $n(n-1) \dots (n-p+1)$ éléments.

Démonstration. Soient $F = \{b_1, \dots, b_p\}$ et $E = \{a_1, \dots, a_n\}$. Pour définir une injection f de F dans E , il suffit de se donner les éléments distincts $f(b_1), \dots, f(b_p)$ de E . Nous allons calculer le nombre A_n^p d'injections de F dans E par récurrence sur p , n étant fixé.

Si $p = 1$, le résultat est évident. Supposons le théorème vrai pour $\text{Card}(F) = p - 1$ et prouvons-le si $\text{Card}(F) = p$.

Soit b un élément de F ; on a $F = \{b\} \cup F'$, où $\text{Card}(F') = p - 1$. Une application f de F dans E est définie de façon unique par sa restriction f' à F' et par la donnée de $f(b)$. Si f est injective, il est clair que f' est injective. Si f' est injective, f est injective si et seulement si $f(b)$ est choisi parmi les éléments de $E - f'(F')$ qui sont au nombre de $n - (p - 1) = n - p + 1$ puisque $\text{Card } f'(F') = p - 1$.

Toute injection de F' dans E se prolonge ainsi en $n - p + 1$ injections de F dans E . Comme d'après l'hypothèse de récurrence il y a $n(n - 1) \dots (n - p + 2)$ injections de F' dans E , on voit qu'il y a $n(n - 1) \dots (n - p + 2) \times (n - p + 1)$ injections de F dans E et le théorème est démontré.

Notation. Soit n un entier > 0 . On pose

$$1 \cdot 2 \cdot 3 \dots n = n!,$$

qui se lit «factorielle n » avec par convention $0! = 1$.

La formule du Théorème 1.7.2.2 peut alors s'écrire, en multipliant et divisant par $(n - p)!$

$$A_n^p = n(n - 1) \dots (n - p + 1) = \frac{n!}{(n - p)!}.$$

1.7.2.3. Corollaire

Le nombre des permutations d'un ensemble E à n éléments est $n!$

En effet, si $F = E$, l'ensemble des applications injectives de E dans E possède $n!$ éléments. Mais d'après le Théorème 1.5.2.7 d), les injections de E dans E sont aussi les permutations de E .

1.7.3. COMBINAISONS SANS RÉPÉTITION

1.7.3.1. Définition

*Soit E un ensemble fini de cardinal n et soit p un entier $\leq n$. On appelle **combinaison sans répétition** (ou **combinaison**) des n éléments de E p à p , toute partie de E ayant p éléments.*

On dit aussi **combinaison de n objets p à p** .

Une combinaison est formée d'objets distincts; deux combinaisons diffèrent par la nature de leurs objets, et non par l'ordre de ces objets.

On note C_n^p ou $\binom{n}{p}$ le nombre de combinaisons de n objets p à p .

1.7.3.2. Théorème

Soit E un ensemble fini à n éléments et soit p un entier $\leq n$. Le nombre des parties de E à p éléments est :

$$C_n^p = \frac{n!}{p!(n-p)!} = \frac{n(n-1) \dots (n-p+1)}{p!}.$$

Démonstration. Supposons formé le tableau $\mathcal{I}$ des combinaisons des n éléments de E , p à p . Considérons une combinaison de $\mathcal{I}$ et effectuons sur les p éléments qui la composent toutes les permutations possibles ; nous obtenons $p!$ groupes. En procédant de même avec chaque combinaison du tableau $\mathcal{I}$, nous obtenons le tableau $\mathcal{I}'$ des arrangements des n éléments de E , p à p , sans omission ni répétition.

On a donc

$$A_n^p = p! C_n^p.$$

D'où

$$C_n^p = \frac{1}{p!} \cdot A_n^p = \frac{n!}{p!(n-p)!}. \quad \square$$

Convention. On pose $C_n^p = 0$ si $p > n$.

Donnons ci-dessous quelques propriétés des entiers C_n^p qu'on appelle les **coefficients du binôme** pour des raisons que nous verrons bientôt.

1.7.3.3. Théorème

On a les propriétés suivantes :

- a) $C_n^0 = 1, \quad C_n^1 = n \quad \text{pour tout } n \in \mathbb{N};$
- b) $C_n^p = C_n^{n-p};$
- c) $C_n^p = \frac{n}{p} C_{n-1}^{p-1} \quad \text{pour tout } n \in \mathbb{N}^* \text{ et pour tout } p \in \mathbb{N}^*.$
- d) $C_n^p = C_{n-1}^p + C_{n-1}^{p-1} \quad \text{pour tout } n \in \mathbb{N}^* \text{ et pour tout } p \in \mathbb{N}^*.$

Démonstration. Tous ces résultats peuvent se vérifier à l'aide de calculs simples. Mais à titre d'exercice, nous allons donner une démonstration sans calcul de b) et d).

b) Si E est un ensemble à n éléments, notons $\mathcal{P}_p$ l'ensemble des parties de E ayant p élément. L'application qui à toute partie associe son complémentaire est une bijection de $\mathcal{P}_p$ sur $\mathcal{P}_{n-p}$; donc

$$C_n^p = C_n^{n-p}$$

d) Soit E un ensemble à n éléments et soit $a \in E$. Considérons les ensembles :

$$\mathcal{P}'_p = \{A \in \mathcal{P}_p : a \in A\}$$

$$\mathcal{P}''_p = \{A \in \mathcal{P}_p : a \notin A\}.$$

On a évidemment

$$\mathcal{P}'_p \cup \mathcal{P}''_p = \mathcal{P}_p \quad \text{et} \quad \mathcal{P}'_p \cap \mathcal{P}''_p = \emptyset,$$

d'où

$$C_n^p = \text{Card}(\mathcal{P}_p) = \text{Card}(\mathcal{P}'_p) + \text{Card}(\mathcal{P}''_p).$$

Pour calculer $\text{Card}(\mathcal{P}'_p)$, on remarque que l'application qui à $A \in \mathcal{P}'_p$ associe $A - \{a\}$ est une bijection de $\mathcal{P}'_p$ sur l'ensemble des parties à $p - 1$ éléments de $E - \{a\}$; comme $\text{Card}(E - \{a\}) = n - 1$, on a $\text{Card}(\mathcal{P}'_p) = C_{n-1}^{p-1}$.

D'autre part, $\text{Card}(\mathcal{P}''_p)$ est le nombre de parties de $E - \{a\}$ ayant p éléments, c'est-à-dire C_{n-1}^p . On en déduit la formule.

Triangle de Pascal

La formule de récurrence $C_n^p = C_{n-1}^p + C_{n-1}^{p-1}$ permet de construire un tableau triangulaire, appelé **triangle de Pascal**, dont les éléments sont les C_n^p . On écrit sur une même ligne les valeurs des C_n^p pour n fixé :

$$n = 0 \quad C_0^0$$

$$n = 1 \quad C_1^0 \quad C_1^1$$

$$n = 2 \quad C_2^0 \quad C_2^1 \quad C_2^2$$

$$n = 3 \quad C_3^0 \quad C_3^1 \quad C_3^2 \quad C_3^3$$

$$n - 1 \quad C_{n-1}^0 \quad C_{n-1}^1 \quad \dots \quad C_{n-1}^{p-1} \quad C_{n-1}^p \quad \dots \quad C_{n-1}^{n-1}$$

$$n \quad C_n^0 \quad C_n^1 \quad \dots \quad C_n^p \quad \dots \quad C_n^{n-1} \quad C_n^n$$

D'après la formule $C_n^p = C_{n-1}^p + C_{n-1}^{p-1}$, chaque terme du tableau qui n'est pas un terme extrême est la somme du terme situé au-dessus et du terme à gauche de ce dernier.

Chapitre 2 : LOIS DE COMPOSITION

Dans ce chapitre, nous allons étudier les opérations algébriques permettant de composer entre eux, deux éléments quelconques d'un ensemble donné. Autrement dit, à tout couple d'éléments d'un ensemble E , nous ferons correspondre un élément bien défini de E .

Si l'ensemble considéré et la loi possèdent certaines propriétés, on obtient ce qu'on appelle une **structure algébrique**.

On notera qu'il est souvent possible de définir plusieurs structures algébriques sur un même ensemble E . Par exemple, l'ensemble des endomorphismes d'un espace vectoriel est muni d'une structure d'espace vectoriel et d'une structure d'anneau.

Nous avons regroupé ici les propriétés générales des lois de composition et certaines démonstrations afin que le lecteur puisse s'y reporter aisément.

2.1. Généralités

2.1.1. DÉFINITIONS. NOTATIONS. EXEMPLES

2.1.1.1. Définition

Soit E un ensemble. On appelle loi de composition interne sur E , toute application de $E \times E$ dans E .

*Un ensemble muni d'une loi de composition interne s'appelle un **magma**.*

Notations. On note de plusieurs manières les lois de composition. Voici quelques notations utilisées fréquemment :

$$\begin{aligned}(x, y) &\longmapsto x + y ; & (x, y) &\longmapsto x \cdot y \\(x, y) &\longmapsto x \top y ; & (x, y) &\longmapsto x \perp y.\end{aligned}$$

Dans ce chapitre, nous utiliserons souvent le signe $\top$ ou le signe $\perp$ pour noter les lois de composition, mais ce sont les notations **additive** $(x, y) \longmapsto x + y$ et **multiplicative** $(x, y) \longmapsto x \cdot y$ qui sont le plus fréquemment utilisées dans les applications.

L'image $x \top y$ du couple $(x, y) \in E \times E$ par la loi de composition $\top$ s'appelle le **composé** de x et de y pris dans cet ordre.

2.1.1.2. Exemple

Dans $\mathbb{R}$, l'addition $(x, y) \mapsto x + y$ et la soustraction $(x, y) \mapsto x - y$ sont des lois de composition internes.

2.1.1.3. Exemple

Soit E un ensemble. Dans $\mathcal{P}(E)$, $(A, B) \mapsto A \cup B$ et $(A, B) \mapsto A \cap B$ sont des lois de composition internes.

2.1.1.4. Exemple

Soit $\mathcal{F}(E, E)$ l'ensemble des applications d'un ensemble E dans lui-même. L'application qui, aux éléments f et g de $\mathcal{F}(E, E)$, associe $f \circ g$ (composition des applications) est une loi de composition interne sur $\mathcal{F}(E, E)$.

2.1.2. PARTIES STABLES. LOIS INDUITES

2.1.2.1. Définition

Soient $(E, \top)$ un magma, et A une partie de E . On dit que A est **stable** pour la loi $\top$ si les relations $x \in A$ et $y \in A$ entraînent $x \top y \in A$.

L'application $(x, y) \mapsto x \top y$ de $A \times A$ dans A est donc une loi de composition interne sur A . On l'appelle la **loi induite** sur A par la loi $\top$ définie sur E .

2.1.3. COMPOSÉ DE DEUX PARTIES

Soit $(E, \top)$ un magma et soient A et B deux parties de E . On note $A \top B$ l'ensemble des éléments de la forme $x \top y$, où $x \in A$ et $y \in B$.

Si la loi $\top$ est notée multiplicativement, on a

$$A \cdot B = \{xy : x \in A, y \in B\}.$$

Si la loi est notée additivement, on a

$$A + B = \{x + y : x \in A, y \in B\}.$$

Si par exemple A se réduit à un élément x et si B est une partie quelconque de E , on écrit $x \top B$, soit xB en notation multiplicative et $x + B$ en notation additive.

2.1.4. TRANSLATIONS

2.1.4.1. Définition

Soient $(E, \top)$ un magma et a un élément de E . On appelle **translation à gauche** (resp. **à droite**) définie par a , l'application L_a (resp. R_a) de E dans E définie par :

$$x \longmapsto L_a(x) = a \top x \quad (\text{resp. } x \longmapsto R_a(x) = x \top a).$$

En notation multiplicative, on écrit :

$$x \longmapsto L_a(x) = ax \quad \text{et} \quad x \longmapsto R_a(x) = xa$$

pour tout $x \in E$.

En notation additive, on écrit :

$$x \longmapsto L_a(x) = a + x \quad \text{et} \quad x \longmapsto R_a(x) = x + a.$$

2.2. Propriétés des lois de composition internes

2.2.1. LOIS ASSOCIATIVES

2.2.1.1. Définition

Soit $(E, \top)$ un magma. On dit que la loi $\top$ est **associative** si l'on a

$$(x \top y) \top z = x \top (y \top z)$$

quels que soient $x, y, z \in E$.

On écrit alors :

$$(x \top y) \top z = x \top (y \top z) = x \top y \top z,$$

et on dit que $(E, \top)$ est un **magma associatif**.

Par exemple, l'addition et la multiplication dans $\mathbb{R}$ sont des lois associatives.

Soit $(x_1, \dots, x_n)$ une suite d'éléments d'un ensemble E muni d'une loi de composition associative $\top$. On définit par récurrence sur n , le composé de ces éléments en posant

$$x_1 \top x_2 \top \dots \top x_n = (x_1 \top \dots \top x_{n-1}) \top x_n.$$

En notation additive, on écrit

$$x_1 + x_2 + \dots + x_n = \sum_{i=1}^n x_i.$$

En notation multiplicative, on écrit

$$x_1 \cdot x_2 \dots x_n = \prod_{i=1}^n x_i$$

Si $x_1 = \dots = x_n = x$, $\sum_{i=1}^n x_i$ se note nx et le produit $x \cdot x \dots x$ (n facteurs) se note x^n . On dit que x^n est la **puissance** $n^{\text{ème}}$ de x .

On vérifie facilement que pour tout entier p tel que $1 \leq p \leq n$, on a la relation

$$x_1 \top x_2 \top \dots \top x_n = (x_1 \top \dots \top x_p) \top (x_{p+1} \top \dots \top x_n).$$

Si la loi de composition est notée multiplicativement, cette égalité s'écrit :

$$x_1 \cdot x_2 \dots x_n = (x_1 \cdot x_2 \dots x_p) \cdot (x_{p+1} \dots x_n).$$

On en déduit que, quels que soient les entiers positifs m et n et quel que soit $x \in E$, on a :

$$x^m \cdot x^n = x^{m+n} \quad \text{et} \quad (x^m)^n = x^{mn}.$$

En notation additive, ces formules s'écrivent :

$$mx + nx = (m + n)x \quad \text{et} \quad n(mx) = (nm)x.$$

2.2.2. LOIS COMMUTATIVES

2.2.2.1. Définition

Soit $(E, \top)$ un magma. On dit que la loi $\top$ est **commutative** si l'on a
 $x \top y = y \top x$ quels que soient $x, y \in E$.

Il peut arriver qu'une loi $\top$ n'étant pas commutative, il existe cependant des éléments x et y de E tels que $x \top y = y \top x$. On dit alors que ces éléments **commutent** ou encore qu'ils sont **permutables**.

On dit qu'un élément x de E est **central** si tout élément de E est permutable avec x . On appelle **centre** de E l'ensemble des éléments centraux.

2.2.2.2. Remarque

Soit $\top$ une loi de composition associative et commutative sur un ensemble E , et soit $(x_1, \dots, x_n)$ une suite d'éléments de E . On démontre par récurrence sur n , que le composé $x_1 \top x_2 \top \dots \top x_n$ est indépendant de l'ordre des facteurs.

Soit E un ensemble muni d'une addition et d'une multiplication associatives et commutatives. Soit x_{ij} le terme général d'une famille d'éléments de E telle que $1 \leq i \leq m$ et $1 \leq j \leq n$, i, j, m et n étant des entiers positifs. On suppose que i et j varient indépendamment l'un de l'autre.

Disposons les éléments x_{ij} sous la forme d'un tableau rectangulaire :

$$\begin{array}{ccccccc} x_{11} & x_{12} & \dots & \dots & \dots & x_{1n} \\ x_{21} & x_{22} & \dots & \dots & \dots & x_{2n} \\ \dots & \dots & \dots & \dots & \dots & \dots \\ x_{m1} & x_{m2} & \dots & \dots & \dots & x_{mn} \end{array}$$

Il est clair que, quel que soit l'ordre de sommation, la somme S des x_{ij} sera la même. On a donc

$$\begin{aligned} S &= (x_{11} + x_{12} + \dots + x_{1n}) + (x_{21} + x_{22} + \dots + x_{2n}) + \dots + \\ & (x_{m1} + x_{m2} + \dots + x_{mn}) = \sum_{j=1}^n x_{ij} + \sum_{j=1}^n x_{2j} + \dots + \sum_{j=1}^n x_{mj} = \sum_{i=1}^m \left(\sum_{j=1}^n x_{ij} \right). \end{aligned}$$

D'autre part, cette même somme vaut

$$\begin{aligned} & (x_{11} + x_{21} + \dots + x_{m1}) + (x_{12} + x_{22} + \dots + x_{m2}) + \dots + \\ & (x_{1n} + x_{2n} + \dots + x_{mn}) = \sum_{i=1}^m x_{i1} + \sum_{i=1}^m x_{i2} + \dots + \sum_{i=1}^m x_{in} = \sum_{j=1}^n \left(\sum_{i=1}^m x_{ij} \right). \end{aligned}$$

Par conséquent, on a

$$S = \sum_{i=1}^m \left(\sum_{j=1}^n x_{ij} \right) = \sum_{j=1}^n \left(\sum_{i=1}^m x_{ij} \right)$$

que l'on écrit

$$S = \sum_{i=1}^m \sum_{j=1}^n x_{ij}.$$

2.2.3. ÉLÉMENT NEUTRE

2.2.3.1. Définition

Soit $(E, \top)$ un magma. On appelle **élément neutre** pour la loi $\top$, tout élément $e \in E$ tel que l'on ait

$$x \top e = e \top x$$

pour tout $x \in E$.

On appelle magma **unifère**, un magma dont la loi de composition possède un élément neutre.

Pour un tel magma, on pose par convention

$$\overset{\circ}{\top} x = e, \quad \text{pour tout } x \in E.$$

2.2.3.2. Exemples

a) Dans $\mathbb{R}$, $e = 0$ est un élément neutre pour l'addition car quel que soit $x \in \mathbb{R}$, on a : $x + 0 = 0 + x = x$.

b) Dans $\mathbb{R}$, $e = 1$ est un élément neutre pour la multiplication car quel que soit $x \in \mathbb{R}$, on a :

$$x \cdot 1 = 1 \cdot x = x.$$

c) Id_E est un élément neutre pour la composition des applications dans $\mathcal{F}(E, E)$.

Une loi de composition peut ne pas admettre d'élément neutre. Toutefois, s'il existe, un tel élément est **unique**. En effet, si e et e' sont des éléments neutres pour la loi $\top$, alors on a $e \top e' = e'$ car e est élément neutre ; de même e' étant élément neutre, on a $e \top e' = e$. Donc $e = e'$.

Cette remarque nous permet de parler de l'élément neutre.

2.2.3.3. Définition

Soit $(E, \top)$ un magma. On dit qu'un élément $a \in E$ est **régulier** (ou **simplifiable**) à gauche (resp. à droite) si la relation

$$\begin{array}{ll} a \top x = a \top y & \text{entraîne } x = y \\ (\text{resp. } x \top a = y \top a) & \text{entraîne } x = y \end{array}$$

quels que soient $x, y \in E$.

On dit que a est **régulier** ou **simplifiable** s'il est régulier à gauche et à droite.

2.2.3.4. Exemples

a) Dans un magma unifié $(E, \top)$, l'élément neutre est régulier.

b) Dans $\mathbb{N}$ tout élément est régulier pour l'addition et tout élément non nul est régulier pour la multiplication.

2.2.4. ÉLÉMENTS SYMÉTRISABLES

2.2.4.1. Définition

Soit $(E, \top)$ un magma unifié d'élément neutre e . On appelle **symétrique à gauche** (resp. **symétrique à droite**) d'un élément x de E , tout élément $y \in E$ tel que l'on ait

$$y \top x = e \quad (\text{resp. } x \top y = e).$$

On appelle **symétrique** de x tout élément y tel que

$$y \top x = x \top y = e.$$

On dit que x est **symétrisable** s'il admet un symétrique.

Notation. Le symétrique d'un élément symétrisable x sera noté x' .

2.2.4.2. Théorème

Soient $(E, \top)$ un magma associatif et unifère, d'élément neutre e , et a un élément de E . Alors :

- a) Si a admet un symétrique à gauche (resp. un symétrique à droite), a est régulier à gauche (resp. à droite).
- b) Si a admet un symétrique à gauche b et un symétrique à droite c , on a : $b = c$.
- c) Si a est symétrisable, son symétrique a' est unique ; a' est aussi symétrisable et a' admet a pour symétrique.
- d) Si deux éléments x et y sont symétrisables, il en est de même de $x \top y$, et on a

$$(x \top y)' = y' \top x'.$$

Démonstration.

a) Supposons qu'il existe $b \in E$ tel que $b \top a = e$. Alors la relation $a \top x = a \top y$ implique $b \top (a \top x) = b \top (a \top y)$ et puisque la loi $\top$ est associative, on a $(b \top a) \top x = (b \top a) \top y$, c'est-à-dire $e \top x = e \top y$; d'où $x = y$.

b) Supposons qu'il existe $b \in E$ et $c \in E$ tels que $b \top a = e$ et $a \top c = e$. On en déduit, en utilisant l'associativité de $\top$:

$$\begin{aligned} b \top a \top c &= (b \top a) \top c = e \top c = c \\ b \top a \top c &= b \top (a \top c) = b \top e = b. \end{aligned}$$

D'où $b = c$.

c) Si b et c sont deux symétriques de a , le calcul précédent montre que $b = c$. Soit a' le symétrique de a . Les relations

$$a' \top a = a \top a' = e$$

montrent que a' est symétrisable et que $(a')' = a$.

d) Si x et y admettent pour symétriques x' et y' respectivement, on a :

$$\begin{aligned} (y' \top x') \top (x \top y) &= y' \top (x' \top x) \top y = y' \top e \top y = y' \top y = e, \\ (x \top y) \top (y' \top x') &= x \top (y \top y') \top x' = x \top e \top x' = x \top x' = e \end{aligned}$$

ce qui montre que $x \top y$ admet pour symétrique $y' \top x'$.

Remarquons que l'élément neutre e est toujours symétrisable et est égal à son symétrique.

Lorsque la loi de composition est notée multiplicativement, on parle d'**inverse** et d'**élément inversible** ; l'inverse d'un élément inversible x se note alors x^{-1} .

En notation additive, le symétrique de x s'appelle l'**opposé** de x et se note $-x$.

2.2.4.3. Théorème

Soit $(E, \top)$ un magma associatif unifère et soit a un élément symétrisable de E . Alors pour tout $b \in E$, l'équation

$$a \top x = b \quad (\text{resp. } x \top a = b)$$

admet la solution unique $a' \top b$ (resp. $b \top a'$), a' désignant le symétrique de a .

Démonstration. La relation $a \top x = b$ implique $a' \top (a \top x) = a' \top b$, c'est-à-dire

$$a' \top b = (a' \top a) \top x = e \top x = x.$$

Réciproquement, si $x = a' \top b$, alors

$$a \top x = a \top (a' \top b) = (a \top a') \top b = e \top b = b.$$

La deuxième équation se traite de la même manière.

2.2.4.4. Théorème

Soit $(E, \top)$ un magma associatif unifié et soit a un élément symétrisable de E . Les translations à gauche L_a et $L_{a'}$ sont bijectives et inverses l'une de l'autre.

De même les translations à droite R_a et $R_{a'}$ sont bijectives et inverses l'une de l'autre.

Démonstration. Pour tout $x \in E$, on a

$$\begin{aligned} L_a \circ L_{a'}(x) &= L_a(L_{a'}(x)) = L_a(a' \top x) = a \top (a' \top x) \\ &= (a \top a') \top x = e \top x = x \\ L_{a'} \circ L_a(x) &= L_{a'}(L_a(x)) = L_{a'}(a \top x) = a' \top (a \top x) \\ &= (a' \top a) \top x = e \top x = x. \end{aligned}$$

Donc $L_a \circ L_{a'} = L_{a'} \circ L_a = Id_E$.

La démonstration pour les translations à droite se fait de la même manière.

2.2.5. DISTRIBUTIVITÉ

2.2.5.1. Définition

Soit E un ensemble muni de deux lois de composition internes notées $\top$ et $$. On dit que la loi $\top$ est **distributive à gauche** (resp. **à droite**) par rapport à la loi $*$ si l'on a :*

$$\begin{aligned} x \top (y * z) &= (x \top y) * (x \top z) \\ (\text{resp. } (y * z) \top x &= (y \top x) * (z \top x)) \end{aligned}$$

quels que soient $x, y, z \in E$.

Si la loi $\top$ est distributive à gauche et à droite par rapport à la loi $*$, on dit que $\top$ est **distributive par rapport à $*$** .

Par exemple, sur $\mathcal{P}(E)$, les lois $\cap$ et $\cup$ sont distributives l'une par rapport à l'autre. De même dans $\mathbb{N}$, la multiplication est distributive par rapport à l'addition.

Si la loi $\top$ est commutative les trois notions de distributivité sont identiques.

2.2.6. LOI QUOTIENT

2.2.6.1. Définition

Soit $(E, \top)$ un magma et soit $\mathcal{R}$ une relation binaire sur E . On dit que $\mathcal{R}$ est compatible avec $\top$ si quels que soient $x, x', y, y' \in E$, les relations $x\mathcal{R}x'$ et $y\mathcal{R}y'$ impliquent

$$(x \top y)\mathcal{R}(x' \top y').$$

2.2.6.2. Exemple

Dans $\mathbb{Z}$, considérons la relation binaire: $x\mathcal{R}y$ si et seulement si, il existe $k \in \mathbb{Z}$ tel que $x - y = 2k$. Alors $\mathcal{R}$ est compatible avec l'addition de $\mathbb{Z}$.

2.2.6.3. Théorème

Soient $(E, \top)$ un magma et $\mathcal{R}$ une relation d'équivalence sur E . Si $\mathcal{R}$ est compatible avec $\top$, il existe une loi de composition interne $\dot{\top}$ sur l'ensemble quotient $E/\mathcal{R}$ telle que, quels que soient $x, y \in E$, on ait

$$\dot{x} \dot{\top} \dot{y} = cl(x \top y),$$

où $cl(x \top y)$ désigne la classe d'équivalence de l'élément $x \top y$. On dit que $\dot{\top}$ est la loi quotient de $\top$ par $\mathcal{R}$.

Démonstration. Soient $\dot{x}$ et $\dot{y}$ des éléments de $E/\mathcal{R}$. Définissons le composé $\dot{x} \dot{\top} \dot{y}$ en posant

$$\dot{x} \dot{\top} \dot{y} = cl(x \top y)$$

où x et y sont des représentants des classes $\dot{x}$ et $\dot{y}$ respectivement. Pour montrer que

$$(\dot{x}, \dot{y}) \longmapsto cl(x \top y)$$

est une application de $(E/\mathcal{R}) \times (E/\mathcal{R})$ dans $E/\mathcal{R}$, il faut s'assurer que $cl(x \top y)$ ne dépend que des classes $\dot{x}$ et $\dot{y}$ mais pas des représentants x de $\dot{x}$ et y de $\dot{y}$. Soient donc x' un autre représentant de $\dot{x}$ et y' un autre représentant de $\dot{y}$; on a $x\mathcal{R}x'$ et $y\mathcal{R}y'$. Comme la relation $\mathcal{R}$ est compatible avec la loi $\top$, on a $(x \top y)\mathcal{R}(x' \top y')$, d'où $cl(x \top y) = cl(x' \top y')$, et le théorème est démontré.

2.3. Morphismes

2.3.1. DÉFINITION. EXEMPLES

2.3.1.1. Définition

Soient E et F des ensembles, $\top$ et $*$ des lois de composition internes sur E et F respectivement. On dit qu'une application $f : E \longrightarrow F$ est un **morphisme** (ou un **homomorphisme**) de $(E, \top)$ dans $(F, *)$ si on a

$$f(x \top y) = f(x) * f(y)$$

quels que soient $x, y \in E$.

Si $(E, \top) = (F, *)$, on dit que f est un **endomorphisme**.

Attention. Si $E = F$ mais si les lois $\top$ et $*$ sont différentes, on ne peut parler d'endomorphisme.

Si f est un morphisme bijectif de $(E, \top)$ sur $(F, *)$, on dit que f est un **isomorphisme**; on dit alors que E et F sont **isomorphes**.

Un isomorphisme de $(E, \top)$ sur lui-même s'appelle un **automorphisme** de $(E, \top)$.

2.3.1.2. Exemples

1) Prenons $(E, \top) = (\mathbb{R}, +)$ et $(F, *) = (\mathbb{R}, \times)$ et soit a un nombre réel > 0 . L'application $x \mapsto f(x) = a^x$ est un homomorphisme car on a

$$f(x + y) = a^{x+y} = a^x \cdot a^y = f(x) f(y)$$

quels que soient $x, y \in \mathbb{R}$.

2) $\mathbb{R}_+^*$ désignant l'ensemble des nombres réels strictement positifs, l'application $x \mapsto \ln(x)$ ($\ln(x)$ désignant la fonction logarithme népérien) est un morphisme de $(\mathbb{R}_+^*, \times)$ dans $(\mathbb{R}, +)$. En effet

$$\ln(xy) = \ln(x) + \ln(y) \quad \text{quels que soient } x, y \in \mathbb{R}_+^*.$$

2.3.1.3. Théorème

a) Soient $(E, \top)$, $(F, *)$ et $(G, \perp)$ trois magmas, f un morphisme de $(E, \top)$ dans $(F, *)$ et g un morphisme de $(F, *)$ dans $(G, \perp)$. Alors $g \circ f$ est un morphisme de $(E, \top)$ dans $(G, \perp)$.

b) Si f est un morphisme bijectif de $(E, \top)$ sur $(F, *)$, l'application réciproque f^{-1} est un morphisme bijectif de $(F, *)$ sur $(E, \top)$.

Démonstration.

a) On a quels que soient $x, y \in E$

$$\begin{aligned} (g \circ f)(x \top y) &= g[f(x \top y)] = g[f(x) * f(y)] = g(f(x)) \perp g(f(y)) \\ &= (g \circ f)(x) \perp (g \circ f)(y), \end{aligned}$$

en utilisant la définition de $g \circ f$ et le fait que f et g sont des morphismes; cela prouve que $g \circ f$ est un morphisme de $(E, \top)$ dans $(G, \perp)$.

b) On sait déjà (Théorème 1.3.3.6) que f^{-1} est bijective si f est bijective. Soient $u, v \in F$; l'application f étant bijective, il existe des éléments x et y uniques de E tels que

$$f(x) = u \quad \text{et} \quad f(y) = v.$$

Alors $f^{-1}(u) = x$ et $f^{-1}(v) = y$.

f étant un homomorphisme, on a $f(x \top y) = f(x) * f(y) = u * v$.

Donc $f^{-1}(u * v) = x \top y = f^{-1}(u) \top f^{-1}(v)$

ce qui montre que f^{-1} est un morphisme.

2.4. Lois de composition externes

2.4.1. DÉFINITION. NOTATION

2.4.1.1. Définition

Soient E et Ω des ensembles. On appelle loi de composition externe à gauche sur E , ayant pour domaine d'opérateurs l'ensemble Ω , toute application de $\Omega \times E$ dans E .

On appelle loi de composition externe à droite sur E , de domaine d'opérateurs Ω , toute application de $E \times \Omega$ dans E .

Notation. Une loi de composition externe se note $(\alpha, x) \mapsto \alpha \top x$ ou $(x, \alpha) \mapsto x \top \alpha$ suivant le cas, le signe $\top$ pouvant être remplacé par un point (notation multiplicative).

Nous dirons en abrégé «soit $(E, \top)_\Omega$ une loi externe» au lieu de «soit $\top$ une loi de composition externe à gauche sur E , de domaine d'opérateurs l'ensemble Ω »

2.4.2. PARTIES STABLES. LOIS INDUITES

2.4.2.1. Définition

Soient $(E, \top)_\Omega$ une loi externe, et A une partie de E . On dit que A est stable pour la loi $\top$ si l'on a :

$$\alpha \top x \in A$$

quel que soit $(\alpha, x) \in \Omega \times A$.

L'application $(\alpha, x) \mapsto \alpha \top x$ de $\Omega \times A$ dans A est donc une loi de composition externe sur A , de domaine d'opérateurs l'ensemble Ω . On l'appelle la loi induite sur A par $\top$.

2.4.3. RESTRICTION DU DOMAINE D'OPÉRATEURS

Soit $(E, \top)_\Omega$ une loi externe, et Ω' une partie de Ω . L'application $(\alpha, x) \mapsto \alpha \top x$ de $\Omega' \times E$ dans E est une loi de composition externe sur E , ayant pour domaine d'opérateurs l'ensemble Ω' . On dit qu'on a restreint le domaine d'opérateurs.

Par exemple, un $\mathbb{C}$ -espace vectoriel peut toujours être considéré comme un $\mathbb{R}$ -espace vectoriel.

Nous verrons des exemples importants de lois externes lorsque nous étudierons les groupes opérant dans un ensemble et les espaces vectoriels.

Chapitre 3 : GROUPES

La théorie des groupes occupe une place très importante en mathématiques, avec des applications dans de nombreuses branches de la science : physique, chimie, sciences de l'ingénieur, cristallographie, etc.

Il n'est évidemment pas question de faire voir ici tous les développements mathématiques auxquels conduit la théorie des groupes ; nous nous proposons, simplement, d'exposer quelques notions élémentaires de cette théorie.

3.1. Généralités

3.1.1. DÉFINITIONS. EXEMPLES

3.1.1.1. Définition

*On appelle groupe, un ensemble G muni d'une loi de composition interne $(x, y) \longmapsto x * y$ possédant les propriétés suivantes :*

a) *Elle est associative :*

$$x * (y * z) = (x * y) * z \text{ quels que soient } x, y, z \in G.$$

b) *Elle admet un élément neutre $e \in G$.*

c) *Tout élément de G admet un symétrique : pour tout $x \in G$, il existe un élément x' de G , tel que*

$$x * x' = x' * x = e.$$

Si de plus, la loi de composition est commutative, le groupe est dit **commutatif** ou **abélien**. Dans ce cas la loi de composition est souvent notée additivement, l'élément neutre est désigné par 0 et le symétrique d'un élément x est noté $-x$.

Un groupe peut être fini ou infini. On appelle **ordre** d'un groupe fini le nombre de ses éléments.

Dans ce qui suit, nous noterons multiplicativement la loi de composition d'un groupe, sauf dans certains exemples particuliers ; l'élément neutre sera noté e .

Rappelons les résultats suivants que nous avons démontrés au Chapitre 2 et qui restent vrais pour les groupes (Théorème 2.2.4.2) :

— Tous les éléments d'un groupe sont réguliers ;

— Si x et y sont deux éléments d'un groupe G , on a $(xy)^{-1} = y^{-1} x^{-1}$.

3.1.1.2. Exemples

L'ensemble $\mathbb{R}$ des nombres réels est un groupe abélien pour l'addition. On l'appelle le **groupe additif des nombres réels**.

De même l'ensemble $\mathbb{Z}$ des entiers relatifs et l'ensemble $\mathbb{Q}$ des nombres rationnels sont des groupes abéliens pour l'addition.

3.1.1.3. Exemple

$\mathbb{R}^* = \mathbb{R} - \{0\}$ est un groupe abélien pour la multiplication.

De même $\mathbb{Q}^* = \mathbb{Q} - \{0\}$ et $\mathbb{C}^* = \mathbb{C} - \{0\}$ sont des groupes abéliens pour la multiplication.

3.1.1.4. Exemple

Soient E un ensemble non vide, et $S(E)$ l'ensemble des bijections de E sur E . Muni de la loi de composition interne $(f, g) \mapsto fog$ (composition des applications), $S(E)$ est un groupe, en général non abélien. On l'appelle le **groupe des permutations de E** ou le **groupe symétrique de E** . L'élément neutre est l'application identique Id_E . L'inverse de $f \in S(E)$ est la bijection réciproque f^{-1} de f .

Lorsque E est l'ensemble $\{1, 2, \dots, n\}$ des n premiers nombres entiers positifs, le groupe $S(E)$ se note S_n et s'appelle le **groupe symétrique d'ordre n** ; c'est un groupe fini et on a $|S_n| = n!$ (voir le Corollaire 1.7.2.3).

3.1.1.5. Exemple

Soient G_1 et G_2 deux groupes dont les lois sont notées multiplicativement. On définit sur l'ensemble produit $G = G_1 \times G_2$ une structure de groupe en posant, si $(x_1, x_2) \in G$ et $(y_1, y_2) \in G$:

$$(x_1, x_2)(y_1, y_2) = (x_1 y_1, x_2 y_2).$$

On dit alors que G est le **produit direct** des groupes G_1 et G_2 .

On vérifiera à titre d'exercice que la multiplication définie par (3.1.1) est associative, qu'elle admet pour élément neutre $e = (e_1, e_2)$ où e_1 et e_2 sont les éléments neutres de G_1 et G_2 respectivement, et que tout élément $x = (x_1, x_2)$ de G admet pour inverse $x^{-1} = (x_1^{-1}, x_2^{-1})$.

De plus le groupe produit $G = G_1 \times G_2$ est abélien si et seulement si G_1 et G_2 sont abéliens.

On définirait de même le produit direct de n groupes $G_1, G_2, \dots, G_n$. En particulier, lorsque $G_i = G$, $1 \leq i \leq n$, le groupe G^n est le groupe produit:

$$G^n = G \times G \times \dots \times G, \quad n \text{ facteurs.}$$

Par exemple, l'ensemble $\mathbb{R}^n = \mathbb{R} \times \mathbb{R} \times \dots \times \mathbb{R}$ (n facteurs), muni de la loi de composition:

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

est un groupe abélien.

3.2. Sous-groupes d'un groupe

3.2.1. DÉFINITION ET CARACTÉRISATION D'UN SOUS-GROUPE

3.2.1.1. Définition

On dit qu'une partie H d'un groupe G est un sous-groupe de G si H vérifie les deux conditions suivantes :

- a) H est non vide.
- b) Les relations $x \in H$ et $y \in H$ entraînent $xy^{-1} \in H$.

Le théorème suivant donne une caractérisation des sous-groupes.

3.2.1.2. Théorème

Soient G un groupe, et H une partie de G . Les propriétés suivantes sont équivalentes :

- a) H est un sous-groupe de G .
- b) $e \in H$ et quels que soient $x, y \in H$, on a $x^{-1} \in H$ et $xy \in H$.

Démonstration. a) $\implies$ b). Comme H est non vide, il existe au moins un $x \in H$. Les relations $x \in H$ et $x \in H$ impliquent alors $xx^{-1} = e \in H$.

D'autre part, si $x \in H$, comme $e \in H$, on a $ex^{-1} = x^{-1} \in H$.

Enfin si $x, y \in H$, on a $y^{-1} \in H$ d'après ce qui précède et donc $x(y^{-1})^{-1} = xy \in H$.

b) $\implies$ a) Si la condition b) est vérifiée, H est non vide car $e \in H$. D'autre part, si $x, y \in H$, H contient x et y^{-1} , donc aussi xy^{-1} ce qui montre que H est un sous-groupe de G .

Le théorème est donc démontré.

3.2.1.3. Remarques

a) Le sous-groupe H muni de la loi induite par la loi de composition définie sur G est un groupe.

Mais une partie stable d'un groupe n'est pas nécessairement un sous-groupe. Par exemple $\mathbb{N}$ est une partie stable de $\mathbb{Z}$ pour l'addition, mais ce n'est pas un sous-groupe de $\mathbb{Z}$.

b) Pour démontrer qu'un ensemble muni d'une loi de composition est un groupe, il est souvent recommandé de montrer que c'est un sous-groupe d'un groupe connu, ce qui abrège les démonstrations.

3.2.1.4. Exemple

Soit G un groupe. Alors G et $\{e\}$ sont des sous-groupes de G . Tout sous-groupe de G , distinct de G et $\{e\}$ s'appelle un **sous-groupe propre** de G .

3.2.1.5. Exemple

$\mathbb{Z}$ et $\mathbb{Q}$ sont des sous-groupes propres du groupe additif $\mathbb{R}$.

3.2.1.6. Exemple

Soit $G = \mathbb{Z}$ et soit n un entier fixé. L'ensemble $n\mathbb{Z}$ des multiples de n est un sous-groupe de $\mathbb{Z}$. Réciproquement, si H est un sous-groupe de $\mathbb{Z}$, nous allons montrer qu'il existe $n \in \mathbb{N}$ tel que $H = n\mathbb{Z}$.

Si $H = \{0\}$ l'assertion est évidente car H est alors l'ensemble des multiples de 0. Supposons donc que $H \neq \{0\}$. Soit $x \in H \cap \mathbb{Z}^+$; alors $-x \in H$, donc H contient des éléments strictement positifs. $H \cap \mathbb{N}^*$, partie non vide de $\mathbb{N}$, contient un plus petit élément que nous notons n . Nous allons montrer que $H = n\mathbb{Z}$.

La relation $n \in H$ entraîne $n\mathbb{Z} \subset H$ puisque H est stable pour l'addition et le passage à l'opposé.

Montrons que, réciproquement, tout élément de H appartient à $n\mathbb{Z}$. Pour tout $x \in H$, effectuons la division euclidienne de x par n :

$$x = nq + r \quad \text{avec} \quad 0 \leq r < n \quad \text{et} \quad q \in \mathbb{Z}.$$

On a $r = x - nq \in H$ puisque $x \in H$ et $nq \in n\mathbb{Z} \subset H$. Comme n est le plus petit élément positif de H et puisque $0 \leq r < n$, on a nécessairement $r = 0$ et $x = nq \in n\mathbb{Z}$, c'est-à-dire $H \subset n\mathbb{Z}$, ce qui achève la démonstration.

3.2.2. SOUS-GROUPE ENGENDRÉ PAR UNE PARTIE**3.2.2.1. Théorème**

Soient G un groupe et $(H_i)_{i \in I}$ une famille de sous-groupes de G . Alors $H = \bigcap_{i \in I} H_i$ est un sous-groupe de G .

Démonstration. H n'est pas vide puisque les sous-groupes H_i contiennent tous l'élément neutre e de G et par suite $e \in H$. Soient x et y deux éléments de H . Pour tout $i \in I$, $x \in H_i$ et $y \in H_i$, donc $xy^{-1} \in H_i$. Par suite $xy^{-1} \in H$ et H est bien un sous-groupe de G .

3.2.2.2. Exemple

La réunion de deux sous-groupes n'est un sous-groupe que si l'un d'eux est inclus dans l'autre. En effet si $x \in H_1 - H_2$ et $y \in H_2 - H_1$, alors $xy \notin H_1 \cup H_2$.

Soient G un groupe et A une partie de G . Il existe des sous-groupes de G contenant A (par exemple G lui-même). L'intersection de tous ces sous-groupes

est, d'après le Théorème 3.2.2.1, un sous-groupe de G contenant A et c'est le plus petit, au sens de l'inclusion. On l'appelle le **sous-groupe de G engendré par A** et on le note $[A]_G$ ou $[A]$.

3.2.2.3. Théorème

Soient G un groupe et A une partie de G . Alors $[A]$ est l'ensemble des produits finis d'éléments de G dont les termes ou leurs inverses sont dans A .

Démonstrations. Désignons par H l'ensemble des éléments x de G qui peuvent s'écrire

$$x = x_1 x_2 \dots x_p, \quad \text{avec } x_i \in A \text{ ou } x_i^{-1} \in A \text{ pour tout } i,$$

en convenant que $H = \{e\}$ si $A = \emptyset$.

Nous allons montrer que H est un sous-groupe de G contenant A et que tout sous-groupe de G qui contient A contient nécessairement H .

En prenant $p = 1$, on voit que $A \subset H$.

Si $A \neq \emptyset$ et si $a \in A$, $aa^{-1} = e \in H$, donc $H \neq \emptyset$ si $A \neq \emptyset$ et $H = \{e\}$ si $A = \emptyset$.

Si $x = x_1 x_2 \dots x_p$ et $y = y_1 y_2 \dots y_s$ appartiennent à H , avec $x_i \in A$ ou $x_i^{-1} \in A$ pour tout i et $y_j \in A$ ou $y_j^{-1} \in A$ pour tout j , on a

$$xy^{-1} = x_1 x_2 \dots x_p y_s^{-1} \dots y_1^{-1} \in H.$$

Donc H est un sous-groupe de G contenant A .

Soit K un sous-groupe de G contenant A . Pour toute famille finie $(x_1, \dots, x_p)$ d'éléments de A , $x_1^{-1}, \dots, x_p^{-1}$ sont aussi des éléments de K ; il en est de même de tout produit de la forme $x_1 \cdot x_2 \dots x_p$, avec $x_i \in A$ ou $x_i^{-1} \in A$.

Par conséquent tout sous-groupe de G qui contient A contient aussi H qui est bien le plus petit sous-groupe de G contenant A .

Par exemple si $A = \{x\}$ où $x \in G$, le sous-groupe de G engendré par x est le sous-groupe

$$H = \{\dots, x^{-2}, x^{-1}, e, x, x^2, \dots\},$$

où x^p est défini pour tout $x \in G$ et pour tout entier relatif p par :

$$x^p = x \dots x \text{ (} p \text{ fois)}, \quad x^{-p} = (x^{-1})^p = x^{-1} \dots x^{-1} \text{ (} p \text{ fois)} \text{ et } x^0 = e.$$

On appelle **ordre de x** , l'ordre du sous-groupe H engendré par x .

Le sous-groupe H engendré par x est abélien car $x^m \cdot x^n = x^{m+n}$ quels que soient $m, n \in \mathbb{Z}$.

3.2.2.4. Définition

*Soit G un groupe. On dit qu'un sous-ensemble A de G est une **partie génératrice de G** lorsque $[A] = G$.*

*Si G admet une partie génératrice finie, on dit que G est un **groupe de type fini**.*

On dit qu'un groupe G est **cyclique** lorsqu'il peut être engendré par un seul élément x ; on dit alors que x est un **générateur** de G .

D'après ce qui précède, tout groupe cyclique est abélien et de type fini.

3.3. Morphismes de groupes

3.3.1. DÉFINITIONS. EXEMPLES

3.3.1.1. Définition

Soient G et G' deux groupes. On appelle **morphisme** ou **homomorphisme** de G dans G' , toute application $f : G \longrightarrow G'$ telle que, pour tout $x \in G$ et pour tout $y \in G$, on ait :

$$f(xy) = f(x) f(y).$$

Si de plus f est bijectif, on dit que f est un **isomorphisme** et on dit alors que les groupes G et G' sont **isomorphes**.

Si $G = G'$ on dit que f est un **endomorphisme** ; un endomorphisme bijectif s'appelle un **automorphisme**.

3.3.1.2. Exemples

a) L'ensemble $\mathbb{R}_+^*$ des nombres réels strictement positifs est un groupe multiplicatif. L'application $x \longmapsto e^x$ est un homomorphisme du groupe additif $\mathbb{R}$ dans le groupe multiplicatif $\mathbb{R}_+^*$, car $e^{x+y} = e^x \cdot e^y$.

On sait d'après le cours d'analyse, que la fonction exponentielle est un isomorphisme de $\mathbb{R}$ sur $\mathbb{R}_+^*$, l'isomorphisme réciproque étant la fonction logarithmique $x \longmapsto \ln(x)$.

b) Soient G un groupe et a un élément de G . L'application $n \longmapsto a^n$, de $\mathbb{Z}$ dans G , est un morphisme du groupe additif $\mathbb{Z}$ dans G car $a^{m+n} = a^m a^n$.

3.3.1.3. Exemple

Soit G un groupe. Pour tout $a \in G$, l'application f_a de G dans G définie par

$$f_a(x) = axa^{-1}$$

est un automorphisme de G . On a en effet

$$f_a(xy) = axya^{-1} = (axa^{-1})(aya^{-1}) = f_a(x) f_a(y)$$

quels que soient $x, y \in G$ et f_a est un homomorphisme.

Montrons que f_a est bijectif. Pour tout $y \in G$, on a $f_a(x) = y$ avec $x = a^{-1}ya$, donc f_a est surjectif. f_a est aussi injectif car l'équation $axa^{-1} = aya^{-1}$ admet la solution unique $x = y$.

Les automorphismes de la forme f_a s'appellent les **automorphismes intérieurs** de G .

Les automorphismes intérieurs de G se réduisent à l'application identique si G est abélien.

3.3.2. PROPRIÉTÉS DES MORPHISMES DE GROUPES

Nous allons énoncer quelques propriétés des morphismes de groupes. Certaines ne sont que des rappels du Chapitre 2.

3.3.2.1. Théorème

Soient G, G' et K trois groupes, f un morphisme de G dans G' et g un morphisme de G' dans K . Alors

a) $g \circ f$ est un morphisme de G dans K .

b) Si f est un isomorphisme de G sur G' , l'application réciproque f^{-1} est un isomorphisme de G' sur G .

La démonstration est identique à celle du Théorème 2.3.2.1.

3.3.2.2. Théorème

Soit f un homomorphisme d'un groupe G dans un groupe G' . Alors

a) Si e est l'élément neutre de G et e' l'élément neutre de G' , on a $f(e) = e'$.

b) Si $x \in G$, on a $f(x^{-1}) = (f(x))^{-1}$.

c) L'image par f de tout sous-groupe de G est un sous-groupe de G' .

d) L'image réciproque par f de tout sous-groupe de G' est un sous-groupe de G .

Démonstration.

a) Pour tout $x \in G$, on a

$$f(x) = f(xe) = f(x) f(e),$$

d'où $f(e) = (f(x))^{-1} f(x) = e'$

b) De même pour tout $x \in G$, on a

$$f(x) f(x^{-1}) = f(xx^{-1}) = f(e) = e'$$

$$f(x^{-1}) f(x) = f(x^{-1}x) = f(e) = e';$$

donc $f(x^{-1}) = (f(x))^{-1}$.

c) Soit H un sous-groupe de G et soit $H' = f(H)$. H' est non vide car $e' = f(e) \in f(H)$.

Si $u, v \in H'$, il existe $x, y \in H$ tels que $u = f(x)$ et $v = f(y)$; alors

$$uv^{-1} = f(x) (f(y))^{-1} = f(x) f(y^{-1}) = f(xy^{-1}).$$

Comme $xy^{-1} \in H$, puisque H est un sous-groupe de G , on voit que $uv^{-1} \in H'$ et H' est un sous-groupe de G' .

d) Soit B' un sous-groupe de G' et soit $B = f^{-1}(B')$.

On a $e \in B$, car $f(e) = e' \in B'$, donc $B \neq \emptyset$. Si $x \in B$ et $y \in B$, on a par définition de l'image réciproque $f(x) \in B'$ et $f(y) \in B'$. On en déduit :

$$f(xy^{-1}) = f(x) f(y^{-1}) = f(x) (f(y))^{-1} \in B'$$

car B' est un sous-groupe de G' . Donc $xy^{-1} \in B$, ce qui achève la démonstration du théorème.

Le théorème 3.3.2.2 montre que si f est un homomorphisme de G dans G' , alors $f(G)$ est un sous-groupe de G' ; ce sous-groupe s'appelle l'**image** de f et on le note $\text{Im}(f)$.

De même l'ensemble $f^{-1}(\{e'\})$ formé des $x \in G$ tels que $f(x) = e'$ est un sous-groupe de G ; on l'appelle le **noyau** de f et on le note $\text{Ker}(f)$.

3.3.2.3. Remarque

D'après le Théorème 3.3.2.2, c), l'image aHa^{-1} d'un sous-groupe H par un automorphisme intérieur f_a est un sous-groupe de G . Tout sous-groupe de la forme aHa^{-1} , avec $a \in G$, est dit **conjugué** de H .

3.3.2.4. Théorème

Soient G et G' deux groupes et soit f un homomorphisme de G dans G' . Pour que f soit injectif, il faut et il suffit que $\text{Ker}(f) = \{e\}$.

Démonstration. Supposons que f soit injectif. Comme $e \in \text{Ker}(f)$, s'il existait un autre élément $x \in \text{Ker}(f)$, on aurait $f(x) = e' = f(e)$, d'où $x = e$ puisque f est injectif; donc $\text{Ker}(f) = \{e\}$.

Réciproquement, supposons que $\text{Ker}(f) = \{e\}$ et soient $x, y \in G$ tels que $f(x) = f(y)$. Cette relation entraîne

$$f(x) f(y)^{-1} = f(x) f(y^{-1}) = f(xy^{-1}) = f(y) f(y)^{-1} = e';$$

donc $xy^{-1} \in \text{Ker}(f) = \{e\}$, d'où $xy^{-1} = e$, c'est-à-dire $x = y$ et f est injectif.

3.4. Groupes-quotients

3.4.1. CLASSES MODULO UN SOUS-GROUPE

3.4.1.1. Théorème

Soient G un groupe et H un sous-groupe de G . Alors :

a) *La relation $x\mathcal{R}y$ si et seulement si $x^{-1}y \in H$ est une relation d'équivalence sur G .*

b) *La classe de x modulo $\mathcal{R}$ est l'ensemble xH .*

Démonstration.

a) Pour tout $x \in G$, on a $x^{-1}x = e \in H$; donc la relation $\mathcal{R}$ est réflexive.

Si $x^{-1}y \in H$, alors $(x^{-1}y)^{-1} \in H$, c'est-à-dire $y^{-1}x \in H$. Donc $x\mathcal{R}y$ implique $y\mathcal{R}x$ et $\mathcal{R}$ est symétrique.

Si $x^{-1}y \in H$ et $y^{-1}z \in H$, alors $(x^{-1}y)(y^{-1}z) \in H$, c'est-à-dire $x^{-1}z \in H$. Donc $x\mathcal{R}y$ et $y\mathcal{R}z$ impliquent $x\mathcal{R}z$ et la relation $\mathcal{R}$ est transitive, ce qui démontre a).

b) Par définition, la classe de $x \in G$ est l'ensemble de $y \in G$ tels que $x^{-1}y \in H$. Posons $x^{-1}y = z$; alors $y = xz$ avec $z \in H$, donc $y \in xH$.

Réciproquement, si $y \in xH$, on a $y = xz$ avec $z \in H$, donc $x^{-1}y = x^{-1}xz = z \in H$.

La classe de x est bien l'ensemble xH .

3.4.1.2. Remarque

Si on remplace la relation $x^{-1}y \in H$ par la relation $yx^{-1} \in H$, on obtient un théorème analogue au Théorème 3.4.1 mais l'ensemble xH est remplacé par Hx . Cette remarque nous amène à poser la définition suivante.

3.4.1.3. Définition

L'ensemble xH s'appelle une classe à gauche modulo H ; l'ensemble Hx s'appelle une classe à droite modulo H .

L'ensemble des classes xH (resp. Hx) modulo H se note G/H (resp. $H \backslash G$).

Donc $G/H = \{xH : x \in G\}$ et $H \backslash G = \{Hx : x \in G\}$.

Lorsque G est abélien, les classes à gauche coïncident avec les classes à droite. Dans ce cas on note $x + H$ la classe de x .

3.4.1.4. Exemple

Prenons pour G le groupe additif $\mathbb{Z}$ et pour H le sous-groupe $n\mathbb{Z}$, où n est un entier strictement positif fixé. La relation d'équivalence s'écrit ici

$$y - x \in n\mathbb{Z} \quad \text{ou} \quad x \equiv y \pmod{n}.$$

La classe d'équivalence d'un élément $x \in \mathbb{Z}$ est

$$\dot{x} = \{\dots, x - 2n, x - n, x, x + n, \dots\}.$$

L'ensemble quotient est noté $\mathbb{Z}/n\mathbb{Z}$.

3.4.1.5. Théorème

Soient G un groupe fini et H un sous-groupe de G . Alors l'ordre de H divise l'ordre de G .

Démonstration. Remarquons d'abord que pour toute classe à gauche xH modulo H , l'application $y \mapsto xy$ de H dans xH est bijective. Cette application est injective car si $xy_1 = xy_2$ avec $y_1, y_2 \in H$ alors puisque x^{-1} existe, on a $x^{-1}(xy_1) = x^{-1}(xy_2)$ d'où $y_1 = y_2$; elle est surjective par définition même des classes d'équivalence.

Donc H étant fini puisque G l'est, chaque classe xH possède autant d'éléments que H . Comme les classes xH forment une partition de G , le nombre d'éléments de G est égal au produit du nombre d'éléments de H par le nombre de classes xH distinctes.

3.4.1.6. Corollaire

Soit G un groupe fini d'ordre p premier. Alors les seuls sous-groupes de G sont $\{e\}$ et G lui-même.

3.4.2. GROUPES-QUOTIENTS

3.4.2.1. Définition

On dit qu'un sous-groupe H d'un groupe G est distingué, ou normal, ou invariant si pour tout $x \in G$, on a $xH = Hx$.

Autrement dit H est un sous-groupe distingué de G si et seulement si la classe à gauche de tout $x \in G$ coïncide avec sa classe à droite.

Il est clair que dans un groupe abélien tout sous-groupe est distingué.

Le résultat suivant donne une caractérisation des sous-groupes distingués.

3.4.2.2. Théorème

Soient G un groupe et H un sous-groupe de G . Les conditions suivantes sont équivalentes.

a) H est distingué.

b) Pour tout $x \in G$, on a $xHx^{-1} = H$.

c) Pour tout $x \in G$, on a $xHx^{-1} \subset H$.

Démonstration. Il est clair que a) $\iff$ b) et b) $\implies$ c). Il reste donc à montrer que c) $\implies$ a). Supposons la condition c) vérifiée. Pour tout $x \in G$, on a alors

$$xH(x^{-1}x) \subset Hx, \quad \text{i.e. } xH \subset Hx.$$

En changeant x en x^{-1} dans la relation $xHx^{-1} \subset H$, il vient $x^{-1}Hx \subset H$, d'où comme précédemment, $Hx \subset xH$; par suite on a $xH = Hx$ et le théorème est démontré.

Soient G un groupe et H un sous-groupe distingué de G . Nous allons voir que l'ensemble quotient G/H peut être muni d'une structure de groupe.

3.4.2.3. Théorème

*Soient G un groupe et H un sous-groupe distingué de G . La relation d'équivalence $x\mathcal{R}y$ si et seulement si $x^{-1}y \in H$ est compatible avec la loi de G et l'ensemble quotient G/H , muni de la loi quotient, est un groupe appelé **groupe-quotient** de G par H . Si, de plus, G est abélien, alors G/H est abélien.*

Démonstration. Montrons tout d'abord que la relation $\mathcal{R}$ est compatible avec la multiplication de G , c'est-à-dire que les relations $x\mathcal{R}x'$ et $y\mathcal{R}y'$ impliquent $(xy)\mathcal{R}(x'y')$, ou encore que $x^{-1}x' \in H$ et $y^{-1}y' \in H$ impliquent $(xy)^{-1}(x'y') \in H$.

Des relations $x^{-1}x' \in H$ et $y^{-1}y' \in H$, on déduit qu'il existe $h, k \in H$ tels que $x' = xh$ et $y' = yk$.

$$\text{Alors } (xy)^{-1}(x'y') = y^{-1}x^{-1}x'y' = y^{-1}x^{-1}xhyk = y^{-1}hyk.$$

Comme H est un sous-groupe distingué, $y^{-1}hy \in H$ et puisque $k \in H$, $y^{-1}hyk \in H$ car H est stable pour la loi de G .

On peut donc définir une loi quotient (notée encore multiplicativement) dans G/H , en posant

$$(xH)(yH) = (xy)H$$

quels que soient $\dot{x} = xH \in G/H$ et $\dot{y} = yH \in G/H$, le résultat ne dépendant pas des représentants choisis dans $\dot{x}$ et dans $\dot{y}$ (voir le Théorème 2.2.6.3).

Il est clair que si G est abélien, alors

$$(xH)(yH) = (yH)(xH).$$

Montrons que G/H , muni de la loi quotient est un groupe. On a

$$\begin{aligned} [(xH)(yH)](zH) &= ((xy)H)(zH) = [(xy)z]H = [x(yz)]H \\ &= (xH)[(yz)H] = (xH)[(yH)(zH)] \end{aligned}$$

quels que soient $x, y, z \in G$, d'où l'associativité de la loi quotient.

COURS D'ALGÈBRE

La classe $eH = H$ est l'élément neutre car pour tout $x \in G$, on

$$(xH)(eH) = (xe)H = xH = (ex)H = (eH)(xH).$$

Enfin, l'inverse de la classe xH est la classe $x^{-1}H$. En effet, on a

$$(xH)(x^{-1}H) = (xx^{-1})H = eH = (x^{-1}x)H = (x^{-1}H)(xH),$$

ce qui achève la démonstration du théorème.

Notons que, par définition de la multiplication dans G/H , l'application canonique $\pi : G \longrightarrow G/H$ est un homomorphisme de groupes car

$$\pi(xy) = (xy)H = (xH)(yH) = \pi(x)\pi(y);$$

c'est pourquoi, on dit parfois que π est l'**homomorphisme canonique**. On notera également que H est le noyau de π car pour tout $x \in H$, on a $\pi(x) = xH = H$.

3.4.2.4. Remarque

Il est clair que H est un sous-groupe distingué d'un groupe G si et seulement si on a $\sigma_a(H) = H$ pour tout automorphisme intérieur σ_a de G , ce qui justifie la terminologie «sous-groupe invariant».

3.4.2.5. Théorème

Soit f un morphisme du groupe G dans le groupe G' . Alors $N = \text{Ker}(f)$ est un sous-groupe distingué de G .

Démonstration. Nous savons déjà (Théorème 3.3.2.2) que N est un sous-groupe de G . Quel que soit $x \in G$ et quel que soit $h \in N$, on a

$$f(xhx^{-1}) = f(x)f(h)f(x^{-1}) = f(x)e'f(x)^{-1} = e',$$

où e' désigne l'élément neutre de G' .

Donc $xNx^{-1} \subset N$ et par suite, N est un sous-groupe distingué de G .

3.4.3. DÉCOMPOSITION CANONIQUE D'UN HOMOMORPHISME

Nous avons vu au Chapitre 1 (Théorème 1.4.2.6) comment décomposer une application. Dans le cas des morphismes de groupes, on peut préciser un peu mieux les choses.

3.4.3.1. Théorème

Soient G et G' deux groupes, f un homomorphisme de G dans G' , π l'homomorphisme canonique de G sur $G/\text{Ker}(f)$ et j l'injection canonique

de $f(G)$ dans G' . Alors il existe un isomorphisme unique $\bar{f}$ du groupe-quotient $G/\text{Ker}(f)$ sur le sous-groupe $f(G)$ de G' tel que $f = j \circ \bar{f} \circ \pi$.

Démonstration. La décomposition canonique se fait, dans le cas général, en considérant la relation d'équivalence associée à $f : x \mathcal{R} y$ si et seulement si $f(x) = f(y)$. Ici, cette relation devient

$$\begin{aligned} x \mathcal{R} y &\iff f(x) = f(y) \\ &\iff e' = f(x)^{-1} f(x) = f(x)^{-1} f(y) = f(x^{-1}) f(y) = f(x^{-1} y). \end{aligned}$$

Donc $x \mathcal{R} y \iff x^{-1} y \in \text{Ker}(f)$.

On reconnaît la relation d'équivalence modulo le sous-groupe $\text{Ker}(f)$. On obtient donc la bijection $\bar{f} : G/\text{Ker}(f) \longrightarrow f(G)$ définie par

$$\bar{f}(\bar{x}) = f(x)$$

pour tout $\bar{x} \in G/\text{Ker}(f)$, où $\bar{x}$ désigne la classe de x (Théorème 1.4.2.6).

Comme $\text{Ker}(f)$ est un sous-groupe distingué de G , $G/\text{Ker}(f)$ est un groupe. L'application canonique π et l'injection canonique j sont des homomorphismes de groupes. $\bar{f}$ est aussi un homomorphisme car on a, quels que soient $\bar{x}, \bar{y} \in G/\text{Ker}(f)$,

$$\bar{f}(\bar{x} \bar{y}) = \bar{f}(\overline{xy}) = f(xy) = f(x) f(y) = \bar{f}(\bar{x}) \bar{f}(\bar{y}).$$

Ainsi $\bar{f}$ est un homomorphisme bijectif, c'est-à-dire un isomorphisme de $G/\text{Ker}(f)$ sur $f(G)$. La factorisation canonique de f s'écrit

$$f = j \circ \bar{f} \circ \pi$$

et on a le diagramme

$$\begin{array}{ccc} G & \xrightarrow{f} & G' \\ \pi \downarrow & & \uparrow j \\ G/\text{Ker}(f) & \xrightarrow{\bar{f}} & f(G) \end{array}$$

3.4.3.2. Exemple

Soit n un entier ≥ 1 . Alors, comme $n\mathbb{Z}$ est un sous-groupe abélien, donc distingué de $\mathbb{Z}$, $\mathbb{Z}/n\mathbb{Z}$ est un groupe abélien pour la loi quotient $\bar{x} + \bar{y} = \overline{x + y}$, où $\bar{x}$ désigne la classe de x .

Montrons que la restriction π_n de l'homomorphisme canonique π à l'ensemble $\{0, 1, \dots, n-1\}$ est une bijection.

COURS D'ALGÈBRE

La relation $\pi_n(x) = \pi_n(y)$ avec $0 \leq x \leq n-1$ et $0 \leq y \leq n-1$ implique

$$\pi(x-y) = 0 \quad \text{et} \quad x-y \in n\mathbb{Z}.$$

Or le seul multiple de n dans l'intervalle $[0, n-1]$ est 0 ; donc $x-y=0$, i.e. $x=y$, donc π_n est injective.

Montrons que π_n est surjective. Soit $\bar{x} \in \mathbb{Z}/n\mathbb{Z}$ et soit x un représentant de $\bar{x}$. La division euclidienne de x par n donne :

$$x = nq + r \quad \text{avec} \quad 0 \leq r < n.$$

D'où, puisque $n\mathbb{Z}$ est le noyau de π , et que $nq \in n\mathbb{Z}$,

$$\pi(x) = \pi(nq) + \pi(r) = \pi(r) = \pi_n(r),$$

ce qui montre que π_n est surjective, donc bijective.

Le groupe quotient $\mathbb{Z}/n\mathbb{Z}$ contient donc n éléments, à savoir $\pi(0), \pi(1), \dots, \pi(n-1)$. On dit que $\mathbb{Z}/n\mathbb{Z}$ est le **groupe des entiers modulo n** .

Pour abrégé, nous noterons $\bar{0}, \bar{1}, \dots, \overline{n-1}$ les éléments de $\mathbb{Z}/n\mathbb{Z}$, $n \geq 1$.

Dressons par exemple la table d'addition du groupe $\mathbb{Z}/4\mathbb{Z}$.

	$\bar{0}$	$\bar{1}$	$\bar{2}$	$\bar{3}$
$\bar{0}$	$\bar{0}$	$\bar{1}$	$\bar{2}$	$\bar{3}$
$\bar{1}$	$\bar{1}$	$\bar{2}$	$\bar{3}$	$\bar{0}$
$\bar{2}$	$\bar{2}$	$\bar{3}$	$\bar{0}$	$\bar{1}$
$\bar{3}$	$\bar{3}$	$\bar{0}$	$\bar{1}$	$\bar{2}$

A l'intersection de la ligne $\bar{n}$ et de la colonne $\bar{m}$ figure l'élément $\bar{n} + \bar{m}$.

3.4.3.3. Exemple

Soit $\mathbf{U}$ le groupe multiplicatif des nombres complexes de module 1. L'application $f \mapsto e^{ix}$ de $\mathbb{R}$ dans $\mathbf{U}$ est un homomorphisme surjectif dont le noyau est le groupe $2\pi\mathbb{Z}$. Le groupe quotient $\mathbb{R}/2\pi\mathbb{Z} = \mathbf{T}$ s'appelle le **tore à une dimension** ou encore le **groupe des nombres réels modulo 2π** . D'après le Théorème 3.4.3.1, $\mathbb{R}/2\pi\mathbb{Z}$ est isomorphe à $\mathbf{U}$.

3.4.4. APPLICATION AUX GROUPES CYCLIQUES

3.4.4.1. Théorème

Soit G un groupe cyclique.

a) Si G est infini, il est isomorphe à $\mathbb{Z}$.

b) Si G est fini d'ordre n , il est isomorphe au groupe additif $\mathbb{Z}/n\mathbb{Z}$.

Démonstration.

a) Soit x un générateur de G ; tout élément de G est donc de la forme x^k avec $k \in \mathbb{Z}$. Il en résulte que l'homomorphisme $f : \mathbb{Z} \longrightarrow G$ défini par $f(n) = x^n$ est surjectif. Son noyau est un sous-groupe de $\mathbb{Z}$, donc de la forme $a\mathbb{Z}$ avec $a \geq 0$ (Exemple 3.2.1.7). D'après le Théorème 3.4.3.1, G est isomorphe au groupe quotient $\mathbb{Z}/\text{Ker}(f)$.

Si G est infini, il en est de même de $\mathbb{Z}/\text{Ker}(f)$, ce qui exige $a = 0$; alors (Théorème 3.3.2.4), f est injectif, donc f est bijectif, et le groupe G est isomorphe à $\mathbb{Z}$.

b) Si G est fini, d'ordre n , $\mathbb{Z}/\text{Ker}(f)$ est aussi d'ordre n , ce qui exige $a = n$, donc G est isomorphe à $\mathbb{Z}/n\mathbb{Z}$.

Dans ce cas, on a $G = \{x^0 = e, x, x^2, \dots, x^{n-1}\}$.

3.4.4.2. Corollaire

Soient G un groupe et x un élément de G . Les propriétés suivantes sont équivalentes :

a) x est d'ordre fini m .

b) m est le plus petit entier naturel non nul tel que $x^m = e$.

Démonstration. a) $\implies$ b) Considérons le morphisme $f : \mathbb{Z} \longrightarrow G$ utilisé précédemment. Soit $H = \text{Im}(f)$ le sous-groupe de G engendré par x . Par hypothèse H est fini. Le noyau de f est un sous-groupe de $\mathbb{Z}$, donc de la forme $p\mathbb{Z}$, p étant le plus petit élément positif non nul de $\text{Ker}(f)$. D'après le Théorème 3.4.3.1, H est isomorphe à $\mathbb{Z}/p\mathbb{Z}$, ce qui exige $p = m$ et m est le plus petit élément positif non nul de $\text{Ker}(f)$, c'est-à-dire le plus petit entier naturel non nul tel que $x^m = e$.

b) $\implies$ a) Par hypothèse m est le plus petit élément positif non nul de $\text{Ker}(f)$. Donc $\text{Ker}(f) = m\mathbb{Z}$ et H est isomorphe à $\mathbb{Z}/m\mathbb{Z}$, donc est d'ordre fini m .

3.4.4.3. Corollaire

Soit G un groupe fini d'ordre n . Alors, on a $x^n = e$ pour tout $x \in G$.

Démonstration. Soit $x \in G$ et soit H le sous-groupe de G engendré par x . L'ordre m de H est fini puisque H est un sous-groupe du groupe fini G . Donc (Théorème 3.4.1.5) il existe un entier p tel que $n = mp$. Comme $x^m = e$, il vient $x^n = (x^m)^p = e^p = e$.

3.4.4.4. Corollaire

Tout groupe G d'ordre p premier est cyclique ; il est engendré par l'un quelconque de ses éléments autre que l'élément neutre e .

Démonstrations. Soient x un élément de G distinct de e et H le sous-groupe de G engendré par x . Alors (Corollaire 3.4.1.6), $H = \{e\}$ ou $H = G$; comme $x \neq e$, on a nécessairement $H = G$. Donc G est cyclique. D'après le Théorème 3.4.4.1, G est isomorphe à $\mathbb{Z}/p\mathbb{Z}$.

3.5. Groupes symétriques

3.5.1. GÉNÉRALITÉS

Rappelons que si E est un ensemble ayant n éléments, l'ensemble $\mathcal{S}(E)$ des permutations de E , muni de la composition des applications est un groupe ayant $n!$ éléments.

3.5.1.1. Théorème

Soient E et F deux ensembles, et f une bijection de E sur F . Alors l'application $\Psi : \mathcal{S}(E) \longrightarrow \mathcal{S}(F)$ définie par $\Psi(h) = f \circ h \circ f^{-1}$ est un isomorphisme de groupes.

Démonstrations. Il est clair que $\Psi(h) \in \mathcal{S}(F)$ pour tout $h \in \mathcal{S}(E)$. D'autre part, on a

$$\Psi(h \circ h') = f \circ (h \circ h') \circ f^{-1} = (f \circ h \circ f^{-1}) \circ (f \circ h' \circ f^{-1}) = \Psi(h) \circ \Psi(h').$$

Enfin Ψ est bijective car tout élément h' de $\mathcal{S}(F)$ est l'image d'un élément et d'un seul de $\mathcal{S}(E)$, à savoir $f^{-1} \circ h' \circ f$.

Si n est un entier positif, nous noterons $\mathbb{N}_n$ l'intervalle $[1, n]$. Donc si E est un ensemble fini, de cardinal n , il existe une bijection de E sur $\mathbb{N}_n$ et d'après le théorème précédent, $\mathcal{S}(E)$ est isomorphe au groupe symétrique $\mathcal{S}_n$ d'ordre $n!$. Cette remarque permet de ne s'intéresser qu'au groupe $\mathcal{S}_n$.

Notation. Si $\sigma \in \mathcal{S}_n$, on convient d'écrire

$$(3.5.1) \quad \sigma = \begin{pmatrix} 1 & 2 & \cdots & n \\ \sigma(1) & \sigma(2) & \cdots & \sigma(n) \end{pmatrix}.$$

On appelle **permutation circulaire** la permutation

$$\tau = \begin{pmatrix} 1 & 2 & \cdots & n \\ 2 & 3 & \cdots & 1 \end{pmatrix}$$

où chaque élément sauf n , a pour image par τ le suivant, l'image de n étant 1.

3.5.1.2. Théorème

Le groupe $\mathcal{S}_n$ n'est pas commutatif si $n \geq 3$.

Démonstration. Il suffit d'exhiber deux éléments de $\mathcal{S}_n$ qui ne commutent pas.

Considérons les deux permutations σ et σ' suivantes qui laissent invariants tous les éléments de $[1, n]$ autres que a, b et c .

$$\sigma = \begin{pmatrix} a & b & c & x & \dots & z \\ b & a & c & x & \dots & z \end{pmatrix}, \quad \sigma' = \begin{pmatrix} a & b & c & x & \dots & z \\ c & b & a & x & \dots & z \end{pmatrix}$$

On a

$$\sigma\sigma\sigma' = \begin{pmatrix} a & b & c & x & \dots & z \\ c & a & b & x & \dots & z \end{pmatrix}$$

$$\sigma'\sigma\sigma = \begin{pmatrix} a & b & c & x & \dots & z \\ b & c & a & x & \dots & z \end{pmatrix},$$

donc $\sigma\sigma\sigma' \neq \sigma'\sigma\sigma$.

3.5.2. TRANSPOSITIONS

3.5.2.1. Définition

Soit n un entier naturel ≥ 2 . On dit qu'une permutation $\tau \in S_n$ est une **transposition** s'il existe deux entiers distincts i et j de $\mathbb{N}_n$ tels que

$$\tau(i) = j, \quad \tau(j) = i \quad \text{et} \quad \tau(k) = k \quad \text{pour} \quad k \neq i \quad \text{et} \quad k \neq j.$$

On note τ_{ij} la transposition qui échange i et j et qui laisse fixes les autres éléments de $\mathbb{N}_n$.

On remarquera que si τ est une transposition alors $\tau\sigma\tau = Id_{\mathbb{N}_n}$, i.e. $\tau = \tau^{-1}$.

3.5.2.2. Théorème

Si $n \geq 2$, tout élément de S_n peut s'écrire comme composé de transpositions.

Démonstrations. Nous allons démontrer ce résultat par récurrence sur n .

Si $n = 2$, S_2 contient les deux éléments

$$e = \begin{pmatrix} 1 & 2 \\ 1 & 2 \end{pmatrix} \quad \text{et} \quad \tau = \begin{pmatrix} 1 & 2 \\ 2 & 1 \end{pmatrix}$$

e est l'application identique de $\mathbb{N}_2$, τ est une transposition et on a

$$e = \tau\sigma\tau.$$

Supposons le théorème vrai pour tout élément de S_{n-1} et démontrons-le pour S_n .

Soit σ une permutation de $\mathbb{N}_n$ et soit i un élément quelconque de $\mathbb{N}_n$. Posons

$$A = \mathbb{N}_n - \{i\}.$$

A contient $n - 1$ éléments. Il y a deux éventualités, ou $\sigma(i) = i$, ou $\sigma(i) \neq i$.

a) $\sigma(i) = i$.

Alors la restriction σ' de σ à A est une permutation de A . D'après l'hypothèse de récurrence, il existe des transpositions $\sigma_1, \sigma_2, \dots, \sigma_p$ dans $\mathcal{S}_{n-1}$ telles que

$$\sigma' = \sigma_1 \circ \sigma_2 \circ \dots \circ \sigma_p.$$

Posons

$$\tau_j(x) = \begin{cases} \sigma_j(x) & \text{si } x \in A \\ i & \text{si } x = i \end{cases}$$

Comme les σ_j sont des éléments de $\mathcal{S}_{n-1}$, les τ_j sont des transpositions, éléments de $\mathcal{S}_n$. On a alors

$$\sigma = \tau_1 \circ \tau_2 \circ \dots \circ \tau_p$$

et σ est bien la composée de transpositions $\tau_j \in \mathcal{S}_n$.

b) $\sigma(i) = j \neq i$.

Soit $\tau \in \mathcal{S}_n$ la transposition de $\mathbb{N}_n$ qui échange i et j et qui laisse fixes tous les autres éléments de $\mathbb{N}_n$. Alors $\tau \circ \sigma \in \mathcal{S}_n$ et on a

$$(\tau \circ \sigma)(i) = \tau(\sigma(i)) = \tau(j) = i.$$

D'après le premier cas, il existe des transpositions $\tau_1, \tau_2, \dots, \tau_p$ dans $\mathcal{S}_n$ telles que

$$\tau \circ \sigma = \tau_1 \circ \tau_2 \circ \dots \circ \tau_p.$$

En composant à gauche par $\tau = \tau^{-1}$, il vient

$$\sigma = \tau \circ \tau_1 \circ \dots \circ \tau_p \quad \text{et le théorème est démontré.}$$

3.5.2.3. Remarque

La décomposition de $\sigma \in \mathcal{S}_n$ en produit de transpositions n'est pas unique. Par exemple si

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

on a $\sigma = \tau_{2,3} \circ \tau_{1,3} \circ \tau_{2,3} = \tau_{1,2} \circ \tau_{2,3}$.

Cependant, nous allons voir que si σ est fixé, la parité du nombre de ces transpositions dans une décomposition quelconque de σ est entièrement définie par σ .

3.5.3. SIGNATURE D'UNE PERMUTATION

3.5.3.1. Définition

Soit σ un élément de $\mathcal{S}_n$. On dit qu'un couple (i, j) d'éléments de $\mathbb{N}_n$ est une *inversion* pour σ ou une σ -*inversion* si $i < j$ et si $\sigma(i) > \sigma(j)$.

3.5.3.2. Exemple

Soit

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

un élément de $\mathcal{S}_4$. Les couples $(1,2)$, $(1,3)$ et $(1,4)$ sont des σ -inversions.

Soit σ une permutation de $\mathbb{N}_n$. Désignons par $I(\sigma)$ le nombre total de σ -inversions. Notons

$$\Delta = \prod_{i < j} (j - i)$$

le produit de toutes les différences $j - i$ avec $i < j$. On a

$$i \in \{1, 2, \dots, n-1\} \quad \text{et} \quad j \in \{2, 3, \dots, n\}.$$

Notons de même

$$\Delta_\sigma = \prod_{i < j} (\sigma(j) - \sigma(i))$$

le produit de toutes les différences $\sigma(j) - \sigma(i)$ avec $i < j$. Puisque σ est une bijection, chaque facteur de Δ se retrouve au signe près une fois et une seule dans Δ_σ et on a

$$(3.5.2) \quad \Delta_\sigma = (-1)^{I(\sigma)} \Delta.$$

L'application $\sigma \mapsto \varepsilon(\sigma) = (-1)^{I(\sigma)}$ de $\mathcal{S}_n$ dans le groupe multiplicatif $\{-1, 1\}$ s'appelle la **signature de la permutation** σ .

Si $\varepsilon(\sigma) = +1$, on dit que la permutation σ est **paire**.

Si $\varepsilon(\sigma) = -1$, on dit que la permutation σ est **impaire**.

3.5.3.3. Exemple

Parité d'une transposition.

Soit τ la transposition qui échange les éléments i et j de $\mathbb{N}_n$ et laisse fixes les autres. On peut l'écrire, en supposant $i < j$:

$$\tau = \begin{pmatrix} 1 & 2 & \dots & i-1 & i & i+1 & \dots & j-1 & j & j+1 & \dots & n \\ 1 & 2 & \dots & i-1 & j & i+1 & \dots & j-1 & i & j+1 & \dots & n \end{pmatrix}$$

Les couples (a, b) tels que $a < b$ et $\tau(a) > \tau(b)$ sont :

$$(i, i+1), (i, i+2), \dots, (i, j-1), (i, j)$$

$$(i+1, j), (i+2, j), \dots, (j-1, j).$$

Le nombre total d'inversions pour τ est donc

$$I(\tau) = 2(j - i) - 1$$

qui est un nombre impair, d'où $\varepsilon(\tau) = -1$.

Le théorème suivant donne les propriétés essentielles de la signature d'une permutation.

3.5.3.4. Théorème

Soient σ et π des éléments de S_n . On a :

$$\varepsilon(Id) = 1, \quad \varepsilon(\sigma\pi) = \varepsilon(\sigma) \varepsilon(\pi), \quad \varepsilon(\sigma^{-1}) = \varepsilon(\sigma).$$

Démonstration. On a

$$I(Id) = 0, \quad \text{d'où} \quad \varepsilon(Id) = (-1)^{I(Id)} = 1.$$

Si σ et π sont des permutations, nous avons d'après (3.5.2) :

$$\Delta_{\sigma\pi} = (\Delta_\pi)_\sigma = \varepsilon(\sigma) \Delta_\pi = \varepsilon(\sigma) \varepsilon(\pi) \Delta.$$

Comme d'autre part,

$$\Delta_{\sigma\pi} = \varepsilon(\sigma \circ \pi) \Delta,$$

il vient

$$\varepsilon(\sigma\pi) = \varepsilon(\sigma) \varepsilon(\pi).$$

On en déduit $\varepsilon(\sigma) \varepsilon(\sigma^{-1}) = \varepsilon(\sigma \circ \sigma^{-1}) = \varepsilon(Id) = 1$.

Comme $(\varepsilon(\sigma))^2 = 1$, on a

$$\varepsilon(\sigma^{-1}) = 1/\varepsilon(\sigma) = (\varepsilon(\sigma))^2 / \varepsilon(\sigma) = \varepsilon(\sigma).$$

3.5.3.5. Remarque

Le Théorème 3.5.3.4 signifie que l'application $\sigma \mapsto \varepsilon(\sigma)$ est un homomorphisme de S_n sur le groupe multiplicatif $\{-1, 1\}$. Le noyau de cet homomorphisme est un sous-groupe de S_n appelé **groupe alterné d'ordre n** . On le note $\mathcal{A}_n$.

3.5.3.6. Remarque

Nous savons que toute permutation σ est un produit de transpositions et que cette décomposition n'est pas unique. Si

$$\sigma = \tau_1 \circ \tau_2 \circ \dots \circ \tau_s = t_1 \circ t_2 \circ \dots \circ t_r$$

où les τ_i et τ_j sont des transpositions, alors d'après l'Exemple 3.5.3.3 et le Théorème 3.5.3.4, nous avons

$$\varepsilon(\sigma) = (-1)^s \quad \text{et} \quad \varepsilon(\sigma) = (-1)^r,$$

d'où $(-1)^s = (-1)^r$. Donc s et r sont tous les deux pairs ou tous les deux impairs.

Ainsi, lorsqu'on décompose une permutation en produits de transpositions, le nombre de ces transpositions est toujours soit pair soit impair.

3.6. Groupes opérant sur un ensemble

3.6.1. DÉFINITION. EXEMPLES

3.6.1.1. Définition

Soient G un groupe et X un ensemble. On dit que G opère à gauche sur X , si l'on s'est donné une application $(g, x) \mapsto g \cdot x$ de $G \times X$ dans X satisfaisant aux conditions suivantes :

- a) $(g_1 g_2) \cdot x = g_1 \cdot (g_2 \cdot x)$ quels que soient $g_1, g_2 \in G$ et $x \in X$;
- b) $e \cdot x = x$ quel que soit $x \in X$.

On définirait de même la notion de groupe opérant à droite sur l'ensemble X .

Pour tout $g \in G$, soit $\varphi(g)$ l'application $x \mapsto g \cdot x$ de X dans X .

D'après a), on a $\varphi(g_1 g_2) = \varphi(g_1) \circ \varphi(g_2)$; d'après b), on a $\varphi(e) = Id_X$. Il en résulte que $\varphi(g) \circ \varphi(g^{-1}) = \varphi(g^{-1}) \circ \varphi(g) = Id_X$, donc que $\varphi(g)$ est bijective pour tout $g \in G$. Donc φ est un homomorphisme de G dans le groupe des permutations de X .

Réciproquement, soit G un groupe, X un ensemble et φ un homomorphisme de G dans $\mathcal{S}(X)$. Pour $g \in G$ et $x \in X$, posons $g \cdot x = \varphi(g)(x)$. Alors les conditions a) et b) ci-dessus sont vérifiées.

En effet, puisque φ est un homomorphisme, on a :

$$\begin{aligned} (gg') \cdot x &= \varphi(gg')(x) = [\varphi(g) \circ \varphi(g')](x) \\ &= \varphi(g) [\varphi(g')(x)] = g \cdot (g' \cdot x) \end{aligned}$$

quels que soient $g, g' \in G$ et $x \in X$.

De même on a :

$$e \cdot x = \varphi(e)(x) = Id_X(x) \quad \text{pour tout } x \in X.$$

On peut donc définir l'action d'un groupe G dans un ensemble X par la donnée d'un homomorphisme de G dans $\mathcal{S}(X)$.

Dans la suite de ce paragraphe, nous nous intéressons uniquement aux groupes opérant à gauche sur un ensemble.

On dit que G opère transitivement sur X , si quels que soient x et y dans X , il existe $g \in G$ tel que $y = g \cdot x$. On dit alors que X est un espace homogène.

3.6.1.2. Exemple

On peut faire opérer G sur lui-même à l'aide des translations à gauche

$$(s, x) \mapsto sx,$$

où à l'aide des translations à droite.

$$(s, x) \mapsto xs.$$

3.6.2. SOUS-GROUPE D'ISOTROPIE. ORBITES

Soit G un groupe opérant à gauche sur l'ensemble X , et soit $x_0 \in X$. L'ensemble H_{x_0} des $g \in G$ tels que $g \cdot x_0 = x_0$ est un sous-groupe de G . En effet si $g_1, g_2 \in H_{x_0}$, on a

$$(g_1 g_2) \cdot x_0 = g_1 \cdot (g_2 \cdot x_0) = g_1 \cdot x_0 = x_0,$$

donc $g_1 g_2 \in H_{x_0}$.

Si $g \in H_{x_0}$, on a

$$g^{-1} \cdot x_0 = g^{-1} \cdot (g \cdot x_0) = (g^{-1} g) \cdot x_0 = e \cdot x_0 = x_0,$$

donc $g^{-1} \in H_{x_0}$. Enfin $H_{x_0} \neq \emptyset$ car $e \in H_{x_0}$ d'après b).

3.6.2.1. Définition

Soient G un groupe opérant à gauche sur l'ensemble X et soit $x_0 \in X$. On appelle stabilisateur de x_0 dans G ou sous-groupe d'isotropie de x_0 , le sous-groupe H_{x_0} des $g \in G$ tels que $g \cdot x_0 = x_0$.

3.6.2.2. Théorème

Soit G un groupe opérant à gauche sur l'ensemble X . Pour $x, y \in X$, la relation binaire «il existe $g \in G$ tel que $y = g \cdot x$ » est une relation d'équivalence dans X . On appelle orbite ou trajectoire de x suivant G , et on note $G \cdot x$, la classe d'équivalence de $x \in X$, pour cette relation.

Démonstration. Pour tout $x \in X$, on a $x = e \cdot x$, donc la relation est réflexive.

Si $y = g \cdot x$, alors on a :

$$g^{-1}y = g^{-1} \cdot (g \cdot x) = (g^{-1}g) \cdot x = e \cdot x = x,$$

donc la relation est symétrique.

Enfin, si $y = g \cdot x$ et $z = g' \cdot y$ ($g \in G, g' \in G$), on a :

$$z = g' \cdot (g \cdot x) = (g'g) \cdot x,$$

donc la relation est transitive.

Notons que le groupe G opère transitivement sur l'ensemble X si et seulement si la seule orbite suivant G est X .

3.6.2.3. Exemple

Soient G un groupe, H un sous-groupe de G et G/H l'ensemble des classes à gauche $\dot{x} = xH$ modulo H . On vérifie facilement que G opère transitivement à gauche sur G/H par l'application

$$(s, xH) \longmapsto (sx)H.$$

Autrement dit, G/H est un espace homogène. Inversement, on a le résultat suivant.

3.6.2.4. Théorème

Soient G un groupe et X un espace homogène de G . Il existe un sous-groupe H de G et une application bijective f de G/H sur X tels que, quels que soient $s, g \in G$, on ait :

$$sf(gH) = f((sg)H).$$

Démonstration. Soit x_0 un point fixé de X et soit H le stabilisateur de x_0 . Définissons une application f de G/H dans X en posant

$$f(sH) = sx_0, \quad s \in G.$$

Cette définition est justifiée car si s_1 et s_2 appartiennent à la même classe modulo H , i.e. si $s_2^{-1}s_1 \in H$, alors $s_1x_0 = s_2x_0$.

Si $sx_0 = gx_0$, alors $g^{-1}sx_0 = x_0$ et $g^{-1}s \in H$, donc $sH = gH$, ce qui prouve que f est injective.

D'autre part, pour tout $x \in X$, il existe par hypothèse un $s \in G$ tel que $x = sx_0 = f(sH)$; donc f est surjective, donc f est bijective et on a :

$$sf(gH) = f((sg)H).$$

Chapitre 4 : ANNEAUX ET CORPS

Dans les classes antérieures, on a étudié les ensembles $\mathbb{Z}$, $\mathbb{Q}$, $\mathbb{R}$ et $\mathbb{C}$ munis de l'addition et de la multiplication ordinaires. On a montré que $\mathbb{Z}$ possédait une structure d'anneau et que $\mathbb{Q}$, $\mathbb{R}$ et $\mathbb{C}$ possédaient une structure de corps.

Dans ce chapitre, nous nous proposons de décrire les structures d'anneau et de corps d'une manière générale mais nous ne donnons ici que quelques définitions et quelques résultats élémentaires de la théorie des anneaux et des corps.

4.1. Structure d'anneaux

4.1.1. DÉFINITIONS. EXEMPLES

4.1.1.1. Définition

On appelle anneau un ensemble A muni de deux lois de composition internes : une addition $(x, y) \mapsto x + y$ et une multiplication $(x, y) \mapsto xy$, satisfaisant aux axiomes suivants :

- (A₁) L'addition est une loi de groupe abélien.*
- (A₂) La multiplication est associative et admet un élément neutre, noté 1_A ou 1 , et appelé élément unité.*
- (A₃) La multiplication est distributive par rapport à l'addition.*

Si de plus la multiplication est commutative, c'est-à-dire si on a $xy = yx$ quels que soient $x, y \in A$, on dit que l'anneau est **commutatif**.

Si A est un anneau quelconque, on dit que deux éléments x et y de A **commutent** ou sont **permutables** si l'on a $xy = yx$.

On appelle **pseudo-anneau**, un ensemble A muni d'une addition et d'une multiplication satisfaisant aux axiomes des anneaux mais tel que la multiplication n'ait pas d'élément neutre.

L'élément neutre pour l'addition dans un anneau A est noté 0 et est appelé **l'élément nul** de A .

4.1.1.2. Remarque

Certains auteurs appellent anneaux les objets que nous avons appelés pseudo-anneaux ; ils appellent anneaux unitaires, ou unifiés, les triplets que nous appelons anneaux. Notre point de vue est justifié par le fait que tout pseudo-anneau peut être plongé dans un anneau avec élément unité.

4.1.1.3. Exemples

a) Munis de l'addition et de la multiplication ordinaires, $\mathbb{Z}$, $\mathbb{Q}$, $\mathbb{R}$ et $\mathbb{C}$ sont des anneaux commutatifs.

b) Soit A un ensemble réduit à un seul élément, que l'on peut toujours noter 0 : $A = \{0\}$. Muni de l'addition $0 + 0 = 0$ et de la multiplication $0 \times 0 = 0$, A est un anneau qu'on appelle l'anneau nul. Dans cet anneau, on a : $1 = 0$.

Un anneau est dit **non nul** s'il n'est pas réduit à $\{0\}$.

Si A est un anneau non nul on note $A^* = A - \{0\}$.

4.1.1.4. Exemple

Soient A un anneau et E un ensemble non vide ; soit $\mathcal{F}(E, A)$ l'ensemble des applications de E dans A .

Pour $f, g \in \mathcal{F}(E, A)$, définissons la somme $f + g$ et le produit fg par les équations

$$\begin{aligned}(f + g)(x) &= f(x) + g(x) \\ (fg)(x) &= f(x)g(x)\end{aligned}$$

quel que soit $x \in E$.

Alors l'ensemble $\mathcal{F}(E, A)$, muni des deux lois de composition

$$(f, g) \longmapsto f + g \quad \text{et} \quad (f, g) \longmapsto fg$$

est un anneau (commutatif si et seulement si A est commutatif). L'élément unité de $\mathcal{F}(E, A)$ est l'application constante égale à 1.

4.1.1.5. Exemple

Soit G un groupe abélien non réduit à $\{0\}$ (noté additivement). L'ensemble $\text{End}(G)$ des endomorphismes de G , muni des deux lois de composition

$$(f, g) \longmapsto f + g \quad \text{et} \quad (f, g) \longmapsto fog$$

définies par

$$\begin{aligned}(f + g)(x) &= f(x) + g(x) \\ (fog)(x) &= f(g(x))\end{aligned}$$

pour tout $x \in G$, est un anneau en général non commutatif. L'élément unité de $\text{End}(G)$ est l'application identique de G .

4.1.1.6. Exemple

Soient A et B deux anneaux. L'ensemble produit $A \times B$, muni des lois définies par :

$$\begin{aligned}(a, b) + (a', b') &= (a + a', b + b') \\ (a, b)(a', b') &= (aa', bb')\end{aligned}$$

quels que soient $a, a' \in A$ et $b, b' \in B$, est un anneau appelé **anneau produit** des anneaux A et B .

On vérifiera, à titre d'exercice, qu'on définit bien ainsi une structure d'anneau sur $A \times B$, l'élément unité étant $(1,1)$.

De plus, l'anneau produit $A \times B$ est commutatif si et seulement si A et B le sont.

4.1.1.7. Exemple

Soit n un entier naturel > 0 . Nous savons déjà, d'après l'Exemple 3.4.3.2, que $\mathbb{Z}/n\mathbb{Z}$ est un groupe abélien pour l'addition $(\bar{x}, \bar{y}) \mapsto \bar{x} + \bar{y} = \overline{x+y}$. Montrons que $\mathbb{Z}/n\mathbb{Z}$ est un anneau commutatif. Pour cela, remarquons d'abord que la congruence modulo n est compatible avec la multiplication de $\mathbb{Z}$.

En effet, si x, y, x' et y' sont des nombres entiers relatifs tels que

$$x \equiv x' \pmod{n} \quad \text{et} \quad y \equiv y' \pmod{n},$$

il existe des éléments k et k' de $\mathbb{Z}$ tels que

$$x - x' = kn \quad \text{et} \quad y - y' = k'n.$$

On en tire
$$x = x' + kn, \quad y = y' + k'n.$$

D'où
$$xy = x'y' + n(ky' + k'x' + kk'n).$$

On en déduit que
$$xy \equiv x'y' \pmod{n}.$$

On peut donc définir une loi de composition interne dans $\mathbb{Z}/n\mathbb{Z}$, appelée multiplication, en posant :

$$\bar{x} \times \bar{y} = \overline{x \times y}.$$

Alors le groupe $\mathbb{Z}/n\mathbb{Z}$, muni de l'addition $(\bar{x}, \bar{y}) \mapsto \bar{x} + \bar{y}$ et de la multiplication $(\bar{x}, \bar{y}) \mapsto \bar{x} \times \bar{y}$ ainsi définies, est un anneau commutatif, l'élément unité étant $\bar{1}$ (à vérifier).

L'anneau ainsi défini s'appelle l'**anneau des entiers modulo n** .

Nous avons déjà dressé la table d'addition de $\mathbb{Z}/4\mathbb{Z}$. Dressons sa table de multiplication (la classe de n étant notée $\dot{n}$).

	$\dot{0}$	$\dot{1}$	$\dot{2}$	$\dot{3}$
$\dot{0}$	$\dot{0}$	$\dot{0}$	$\dot{0}$	$\dot{0}$
$\dot{1}$	$\dot{0}$	$\dot{1}$	$\dot{2}$	$\dot{3}$
$\dot{2}$	$\dot{0}$	$\dot{2}$	$\dot{0}$	$\dot{2}$
$\dot{3}$	$\dot{0}$	$\dot{3}$	$\dot{2}$	$\dot{1}$

4.1.2. RÈGLES DE CALCULS DANS UN ANNEAU

Soit A un anneau. Toutes les règles de calcul valables dans un groupe abélien s'appliquent évidemment au **groupe additif** de A qui est l'ensemble A considéré comme groupe abélien. Par exemple, l'**opposé** d'un élément $x \in A$ se note $-x$ et on note $x + (-y) = x - y$.

1° Pour tout élément x d'un anneau A , on a :

$$(4.1.2.1) \quad x \cdot 0 = 0 \cdot x = 0.$$

En effet, on a :

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

d'où puisque tout élément de A est régulier pour l'addition, $x \cdot 0 = 0$.

On montre de même que $0 \cdot x = 0$.

2° Pour tout $x \in A$ et tout $y \in A$, on a :

$$(4.1.2.2) \quad x(-y) = (-x)y = -(xy).$$

En effet, pour tout $y \in A$, on a $y + (-y) = 0$; donc

$$x(y + (-y)) = xy + x(-y) = x \cdot 0 = 0.$$

Par suite $x(-y)$ est l'opposé de xy .

On démontrerait de même que $(-x)y$ est l'opposé de xy .

On en déduit ;

$$(4.1.2.3) \quad (-x)(-y) = -((-x)y) = -(-(xy)) = xy.$$

Notons que si A n'est pas l'anneau nul, alors $1 \neq 0$. En effet, la relation (4.1.2.1) montre que s'il existe $x \in A$ tel que $x \neq 0$, alors $x \cdot 0 = 0 \neq x$, donc 0 ne peut être l'élément neutre de la multiplication. L'anneau nul est donc le seul anneau dans lequel on a : $1 = 0$.

3° Pour tout élément $x \in A$, on définit par récurrence sur l'entier $n \in \mathbb{N}$ les éléments x^n et $n \cdot x$, en posant :

$$\begin{aligned} x^0 &= 1, & x^n &= x^{n-1} \cdot x \\ 0 \cdot x &= 0, & n \cdot x &= (n-1)x + x. \end{aligned}$$

On a alors les propriétés suivantes que l'on vérifie aisément par récurrence :

$$(4.1.2.4) \quad x^m \cdot x^n = x^{m+n}$$

$$(4.1.2.5) \quad (m+n)x = m \cdot x + n \cdot x$$

quels que soient $m, n \in \mathbb{N}$ et $x \in A$.

4° Pour tout $x \in A$ et pour tout $n \in \mathbb{N}$, on a :

$$(4.1.2.6) \quad n \cdot x = (n \cdot 1)x = x \cdot (n \cdot 1).$$

Pour $n = 0$, les égalités sont vraies d'après (4.1.2.1) et la définition de $0 \cdot x$. Supposons (4.1.2.6) vraie pour l'entier n . On a

$$\begin{aligned}(n+1) \cdot x &= nx + x = (n \cdot 1)x + 1 \cdot x = (n \cdot 1 + 1)x \\ &= (n \cdot 1 + 1 \cdot 1)x = ((n+1) \cdot 1)x.\end{aligned}$$

On montrerait de même que

$$(n+1)x = x((n+1) \cdot 1).$$

Ces règles de calcul nous permettent de développer les produits de sommes d'éléments d'un anneau A en tenant compte de l'ordre des termes. Si A est commutatif, on peut procéder à des simplifications.

5° Formule du binôme

4.1.2.1. Théorème

Soit A un anneau et soient a et b deux éléments permutables de A . Pour tout entier $n \geq 1$, on a la formule dite du binôme :

$$(4.1.2.7) \quad (a+b)^n = \sum_{k=0}^n C_n^k a^{n-k} b^k.$$

Démonstration. Raisonnons par récurrence sur n . Pour $n = 1$, le résultat est trivial car la formule se réduit alors à $(a+b)^1 = 1 \cdot a + 1 \cdot b$.

Supposons (4.1.2.7) vraie pour l'entier $n \geq 1$, et montrons qu'elle est vraie pour $n+1$. On a donc d'après l'hypothèse de récurrence

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

En multipliant cette relation par $a+b$, il vient :

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

Or, on montre facilement que, puisque a et b sont permutables, il en est de même de a^p et b^q quels que soient $p, q \in \mathbb{N}$.

Donc :

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

Le coefficient de $a^{n+1-k}b^k$ dans le second membre de cette relation est

$$C_n^k + C_n^{k-1} = C_{n+1}^k \quad \text{pour } 1 \leq k \leq n.$$

Comme d'autre part

$$C_n^0 = C_{n+1}^0 = 1 \quad \text{et} \quad C_n^n = C_{n+1}^{n+1} = 1,$$

on obtient, après regroupements :

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

Donc (4.1.2.7) est vraie pour $n+1$, donc pour tout entier $n \geq 1$.

4.1.3. PROPRIÉTÉS ÉLÉMENTAIRES DES ANNEAUX

• Diviseurs de zéro

Soit A un anneau et soit $a \in A$. Alors nous savons que $a \cdot 0 = 0 \cdot a = 0$.

On voit ainsi que dans un anneau, le produit de deux facteurs est nul lorsque l'un des facteurs est nul. La réciproque est inexacte comme le montre l'exemple suivant.

4.1.3.1. Exemple

Prenons $A = \mathbb{R} \times \mathbb{R}$. Pour $(a, b) \in A$ et $(c, d) \in A$, posons :

$$(a, b) + (c, d) = (a + c, b + d), \quad (a, b)(c, d) = (ac, bd).$$

Alors A est un anneau commutatif, l'élément unité étant $(1, 1)$ et l'élément nul étant $0 = (0, 0)$. On a $(1, 0)(0, 1) = (0, 0) = 0$ et pourtant $(1, 0) \neq 0$ et $(0, 1) \neq 0$. Cette remarque permet de poser la définition suivante.

4.1.3.2. Définition

Soit A un anneau non réduit à $\{0\}$. On dit qu'un élément $a \in A$ est un **diviseur de zéro à gauche** (resp. à droite) si $a \neq 0$ et s'il existe un élément non nul b de A tel que $ab = 0$ (resp. $ba = 0$).

Si A est commutatif, les notions de diviseur de zéro à gauche et à droite coïncident.

Dire que a est un diviseur de zéro à gauche, revient à dire que $a \neq 0$ et que a n'est pas régulier à gauche pour la multiplication.

En effet, si a est non nul et est un diviseur de zéro à gauche, il existe b non nul dans A tel que $ab = 0$; alors la relation $ab = 0 = a \cdot 0$ montre que a n'est pas régulier à gauche pour la multiplication de A .

Réciproquement, si $ax = ay$ avec $x \neq y$, alors on a

$$a(x - y) = 0 \quad \text{avec} \quad x - y \neq 0;$$

a est un diviseur de zéro à gauche.

De même, a est un diviseur de zéro à droite si et seulement si $a \neq 0$ et n'est pas régulier à droite pour la multiplication de A .

4.1.3.3. Définition

On dit qu'un anneau A est intègre ou est un anneau d'intégrité s'il est non nul, commutatif et s'il ne possède pas de diviseurs de zéro.

Autrement dit l'anneau A est intègre si la relation $ab = 0$ implique $a = 0$ ou $b = 0$.

• Éléments nilpotents

4.1.3.4. Définition

Soit A un anneau. On dit qu'un élément $x \in A$ est nilpotent s'il existe un entier $n \geq 1$ tel que $x^n = 0$.

On notera que si A possède un élément nilpotent a non nul, alors A possède des diviseurs de zéro car alors $a \cdot a^{n-1} = 0 = a^{n-1} \cdot a$.

• Éléments inversibles

4.1.3.5. Définition

Soit A un anneau et soit $a \in A$. On dit que a est un élément inversible de A si a possède un symétrique pour la multiplication.

Nous noterons $A^\times$ l'ensemble des éléments inversibles de A .

4.1.3.6. Théorème

Soit A un anneau non nul. L'ensemble $A^\times$ des éléments inversibles de A est un groupe pour la multiplication de A .

Démonstrations. D'après le Théorème 2.2.4.2 d), si $x \in A^\times$ et $y \in A^\times$, alors $xy \in A^\times$; on peut donc définir sur $A^\times$ la multiplication $(x, y) \mapsto xy$. La multiplication est associative dans $A^\times$ puisqu'elle l'est déjà dans A .

On a évidemment $1 \in A^\times$ et 1 est l'élément neutre pour la multiplication dans $A^\times$. Enfin si $a \in A^\times$, on a

$$aa^{-1} = a^{-1}a = 1,$$

ce qui montre que $a^{-1} \in A^\times$. Donc $A^\times$ est un groupe pour la multiplication.

L'ensemble $A^\times$, muni de la multiplication $(x, y) \mapsto xy$ s'appelle le groupe multiplicatif de l'anneau A ou encore le groupe des éléments inversibles de l'anneau A .

Par exemple, dans l'anneau $\mathbb{Z}$ des entiers relatifs, on a $\mathbb{Z}^\times = \{-1, 1\}$.

4.2. Sous-anneaux. Idéaux. Anneaux quotients

4.2.1. SOUS-ANNEAUX

4.2.1.1. Définition

Soient A un anneau et B une partie non vide de A . On dit que B est un sous-anneau de A si les conditions suivantes sont vérifiées :

- a) B est un sous-groupe du groupe additif A .
- b) Les relations $x \in B$ et $y \in B$ impliquent $xy \in B$.
- c) L'élément unité 1 de A appartient à B .

On vérifie facilement que l'ensemble B , muni des deux lois de composition

$$(x, y) \longmapsto x + y \quad \text{et} \quad (x, y) \longmapsto xy$$

induites par celles de A , est un anneau.

Le théorème suivant donne une caractérisation des sous-anneaux.

4.2.1.2. Théorème

Soient A un anneau et B une partie de A . Les conditions suivantes sont équivalentes :

- a) B est un sous-anneau de A .
- b) $1 \in B$ et quels que soient $x, y \in B$, on a $x - y \in B$ et $xy \in B$.

Démonstrations. Si B est un sous-anneau de A , il est clair que la condition b) est vérifiée.

Réciproquement, si la condition b) est vérifiée, $B \neq \emptyset$ car $1 \in B$ et B est un sous-groupe du groupe additif A (puisque les relations $x \in B$ et $y \in B$ entraînent $x - y \in B$). D'autre part, comme les relations $x \in B$ et $y \in B$ impliquent $xy \in B$, on voit que B est bien un sous-anneau de A .

On démontre aisément que toute intersection de sous-anneaux de A est un sous-anneau de A . On peut donc parler du plus petit sous-anneau contenant une partie non vide H de A ; c'est l'intersection de tous les sous-anneaux de A contenant H . On l'appelle le sous-anneau engendré par H .

4.2.1.3. Exemples

- a) $\mathbb{R}$ est un sous-anneau de $\mathbb{C}$.
- b) Soit A un anneau; A est un sous-anneau de A mais $\{0\}$ n'est pas un sous-anneau de A si $A \neq \{0\}$.

Nous allons introduire maintenant la notion d'idéal dont le rôle en théorie des anneaux est l'analogue de celui des sous-groupes distingués en théorie des groupes.

4.2.2. IDÉAUX

4.2.2.1. Définition

Soit A un anneau et I une partie de A . On dit que I est un **idéal à gauche** (resp. **à droite**) de A si :

a) I est un sous-groupe du groupe additif A .

b) Quel que soit $a \in A$ et quel que soit $x \in I$, on a $ax \in I$ (resp. $xa \in I$).
On dit que I est un **idéal bilatère** ou simplement un **idéal** de A si I est à la fois un idéal à gauche et un idéal à droite de A .

Notons que dans un anneau commutatif, tous les idéaux sont bilatères.

4.2.2.2. Exemple

Dans tout anneau A , les sous-groupes triviaux A et $\{0\}$ sont des idéaux. Tout idéal de A autre que A et l'idéal nul $\{0\}$ s'appelle un **idéal propre** de A .

4.2.2.3. Exemple

Soit A un anneau et soit $a \in A$. Alors $Aa = \{xa : x \in A\}$ est un idéal à gauche de A . En effet, d'une part $Aa \neq \emptyset$ car $a = 1 \cdot a \in Aa$, et d'autre part si x et y appartiennent à Aa , il existe x' et y' dans A tels que $x = x'a$ et $y = y'a$, d'où :

$$x - y = (x' - y')a \in Aa;$$

donc Aa est un sous-groupe du groupe additif A . Enfin pour tout $z \in A$, on a :

$$zx = z(x'a) = (zx')a \in Aa$$

et la condition b) de la définition d'un idéal est vérifiée.

De même aA est un idéal à droite de A .

Désormais, lorsqu'on parlera d'idéal sans préciser, il s'agira toujours d'idéal bilatère.

4.2.2.4. Définition

Soit I un idéal de l'anneau A . On dit que I est un **idéal maximal** si $I \neq A$ et si, pour tout idéal J différent de I , $I \subset J$ implique $J = A$.

4.2.2.5. Remarques

a) Si I est un idéal de l'anneau A et si $1 \in I$, alors $I = A$. En effet, quel que soit $a \in A$, on a $1 \cdot a = a \in I$, donc $A \subset I$. Comme on a toujours $I \subset A$, alors $I = A$.

b) Si I est un idéal propre de l'anneau A , aucun élément de I n'est inversible. En effet, s'il existe $a \in I$ tel que a^{-1} existe, alors $a^{-1}a = 1 \in I$, et $I = A$ d'après a) ce qui est contraire à l'hypothèse, donc a n'est pas inversible.

4.2.2.6. Théorème

Soient A un anneau et $(I_\lambda)_{\lambda \in L}$ une famille d'idéaux à gauche (resp. à droite) de A . Alors $I = \bigcap_{\lambda \in L} I_\lambda$ est un idéal à gauche (resp. à droite) de A .

Démonstration. Supposons que $(I_\lambda)_{\lambda \in L}$ soit une famille d'idéaux à gauche de A . On sait déjà (Théorème 3.2.2.1) que I est un sous-groupe du groupe additif A . Soit $x \in I$; on a $x \in I_\lambda$ pour tout $\lambda \in L$. Comme I_λ est un idéal à gauche, on a $bx \in I_\lambda$ pour tout $b \in A$ et pour tout $\lambda \in L$, d'où $bx \in I$; donc I est un idéal à gauche de A .

On démontrerait de même que toute intersection d'idéaux à droite (resp. d'idéaux bilatères) de A est un idéal à droite (resp. bilatère) de A .

Soient A un anneau et X une partie de A . Il existe des idéaux à gauche de A contenant X (par exemple A lui-même). L'intersection de tous ces idéaux à gauche est un idéal à gauche de A contenant X et c'est le plus petit, au sens de l'inclusion. On l'appelle l'idéal à gauche engendré par X .

On définirait de même l'idéal à droite engendré par X et l'idéal bilatère engendré par X .

Par exemple, si a est un élément fixé de A , l'idéal à gauche engendré par a est l'ensemble

$$Aa = \{xa : x \in A\}.$$

Nous savons déjà que Aa est un idéal à gauche de A ; cet idéal contient a car $a = 1 \cdot a \in Aa$. De plus, si I est un idéal à gauche de A contenant a , I contient xa pour tout $x \in A$; donc on a $Aa \subset I$ et par suite Aa est l'idéal à gauche engendré par a .

On montrerait de même que l'idéal à droite engendré par a est l'ensemble aA .

Plus généralement on démontrera le théorème suivant à titre d'exercice.

4.2.2.7. Théorème

Soient A un anneau et X une partie non vide de A . L'ensemble $\mathcal{J}$ des éléments x de A ayant la propriété suivante : il existe un entier $n > 0$, une suite $x_1, \dots, x_n$ de n éléments de X et une suite $a_1, \dots, a_n$ de n éléments de A tels que $x = a_1x_1 + \dots + a_nx_n$ est l'idéal à gauche de A engendré par X .

4.2.2.8. Définition

Soit A un anneau. On dit qu'un idéal I de A est un idéal principal s'il existe $a \in A$ tel que $I = Aa = aA$.

On dit qu'un anneau A est principal s'il est commutatif, intègre, et si tout idéal de A est principal.

4.2.2.9. Exemple

Les idéaux de $\mathbb{Z}$ sont les ensembles de la forme $n\mathbb{Z}$, $n \in \mathbb{N}$. On sait déjà que tout sous-groupe de $\mathbb{Z}$ est de la forme $n\mathbb{Z}$, $n \in \mathbb{N}$ (Exemple 3.2.1.6) et il est clair que si $a \in \mathbb{Z}$ et $x \in n\mathbb{Z}$, alors $ax \in n\mathbb{Z}$. Donc $\mathbb{Z}$ est un anneau principal.

4.2.3. ANNEAUX-QUOTIENTS

4.2.3.1. Théorème

Soient A un anneau et I un idéal bilatère de A . Alors la relation définie par $x\mathcal{R}y \iff x - y \in I$ est une relation d'équivalence sur A , compatible avec les deux lois de A . L'ensemble quotient, noté A/I , muni des deux lois quotients est un anneau appelé anneau-quotient de A par I .

Si, de plus, A est commutatif, l'anneau A/I est commutatif.

Démonstration. Comme I est un sous-groupe du groupe additif A , $\mathcal{R}$ est une relation d'équivalence (Théorème 3.4.2.3). On peut donc définir le groupe additif quotient A/I .

Montrons que la relation $\mathcal{R}$ est compatible avec la multiplication de A .

On doit démontrer que les relations $x - x' \in I$ et $y - y' \in I$ impliquent $xy - x'y' \in I$. Or si $x - x' = u \in I$ et $y - y' = v \in I$, on a :

$$x = x' + u \quad \text{et} \quad y = y' + v,$$

d'où

$$xy = x'y' + x'v + uy' + uv.$$

Comme I est un idéal, $x'v \in I$, $uy' \in I$ et $uv \in I$; on en déduit :

$$xy - x'y' = x'v + uy' + uv \in I.$$

On peut donc définir la loi quotient de la multiplication en posant

$$(x + I)(y + I) = xy + I$$

quels que soient $\dot{x} = x + I \in A/I$ et $\dot{y} = y + I \in A/I$.

Si A est commutatif, il est clair que cette multiplication de A/I est commutative.

On vérifie facilement que la multiplication de A/I est associative et distributive par rapport à l'addition de A/I . Donc A/I , muni des deux lois quotient, est un anneau.

Par exemple $\mathbb{Z}/n\mathbb{Z}$, $n \in \mathbb{N}$, est un anneau commutatif (voir l'Exemple 4.1.1.7) car $n\mathbb{Z}$ est un idéal de $\mathbb{Z}$.

4.3. Morphismes d'anneaux

4.3.1. DÉFINITION ET PROPRIÉTÉS DES MORPHISMES D'ANNEAUX

4.3.1.1. Définition

Soient A et B deux anneaux. On dit qu'une application f de A dans B est un **morphisme** ou un **homomorphisme d'anneaux** si :

- a) $f(1) = 1$.
- b) $f(x + y) = f(x) + f(y)$ et $f(xy) = f(x)f(y)$ quels que soient $x, y \in A$.

On définit comme pour les groupes, les notions d'endomorphisme, d'isomorphisme et d'automorphisme d'anneaux.

Le **noyau** d'un morphisme d'anneaux est le noyau du morphisme des groupes sous-jacents $(A, +)$ et $(B, +)$.

Le noyau d'un homomorphisme d'anneaux f est noté $\text{Ker}(f)$.

4.3.1.2. Exemple

Si I est un idéal bilatère d'un anneau A , l'application canonique $\pi : A \longrightarrow A/I$ est, par définition de l'anneau A/I , un morphisme d'anneaux. On l'appelle l'**homomorphisme canonique**.

4.3.1.3. Théorème

Soient A, B et C trois anneaux, f un morphisme de A dans B et g un morphisme de B dans C . Alors :

- a) $g \circ f$ est un morphisme de A dans C .
- b) Si f est un isomorphisme de A sur B , l'application réciproque f^{-1} est un isomorphisme de B sur A .

La démonstration est évidente et est laissée au lecteur.

4.3.1.4. Théorème

Soit $f : A \longrightarrow B$ un morphisme d'anneaux. Alors :

- a) $f(0) = 0$ et $f(-x) = -f(x)$ pour tout $x \in A$.
- b) Si x est un élément inversible de A , on a $f(x^{-1}) = (f(x))^{-1}$.
- c) Si A' est un sous-anneau de A , $f(A')$ est un sous-anneau de B .

d) Si B' est un sous-anneau de B , $f^{-1}(B')$ est un sous-anneau de A .

e) Si I est un idéal de B , $f^{-1}(I)$ est un idéal de A , et f est injectif si et seulement si $\text{Ker}(f) = \{0\}$.

Démonstration. a) Puisque le morphisme d'anneaux $f : A \longrightarrow B$ possède les propriétés d'un morphisme de groupes abéliens, les propriétés du a) sont vérifiées.

b) se démontre comme le b) du Théorème 3.3.2.2.

c) On sait que $f(A')$ est un sous-groupe du groupe additif B . On a d'abord $1 \in f(A')$ car A' est un sous-anneau et f est un morphisme d'anneaux. Soient d'autre part $f(x)$ et $f(y)$ deux éléments de $f(A')$. On a $f(x)f(y) = f(xy) \in f(A')$; donc $f(A')$ est bien un sous-anneau de B .

d) Nous savons déjà que $f^{-1}(B')$ est un sous-groupe du groupe additif A (Théorème 3.3.2.2). On a $1 \in f^{-1}(B')$ car $f(1) = 1 \in B'$. Si $x \in f^{-1}(B')$ et $y \in f^{-1}(B')$, on a $f(x) \in B'$ et $f(y) \in B'$. D'où $f(xy) = f(x)f(y) \in B'$ car B' est un sous-anneau de B . Donc $xy \in f^{-1}(B')$ et $f^{-1}(B')$ est bien un sous-anneau de A .

e) I étant un sous-groupe du groupe additif B , $f^{-1}(I)$ est un sous-groupe du groupe additif A d'après le Théorème 3.3.2.2. Soient $a \in A$ et $x \in f^{-1}(I)$. Les relations $f(a) \in B$ et $f(x) \in I$ impliquent

$f(ax) = f(a)f(x) \in I$ et $f(xa) = f(x)f(a) \in I$
 puisque I est un idéal de B . Donc $ax \in f^{-1}(I)$ et $xa \in f^{-1}(I)$, ce qui prouve que $f^{-1}(I)$ est un idéal de A .

En particulier, $\text{Ker}(f) = f^{-1}(\{0\})$ est un idéal de A et d'après le Théorème 3.3.2.4, f est injectif si et seulement si $\text{Ker}(f) = \{0\}$.

4.3.2. DÉCOMPOSITION CANONIQUE D'UN MORPHISME D'ANNEAUX

Soient A et B deux anneaux et $f : A \longrightarrow B$ un morphisme d'anneaux. Considérons la relation d'équivalence $x \mathcal{R} y$ si et seulement si $f(x) = f(y) \iff f(x - y) = 0 \iff x - y \in \text{Ker}(f)$. Comme $\text{Ker}(f)$ est un idéal de A , on peut former l'anneau-quotient $A/\text{Ker}(f)$. L'application canonique π de A sur $A/\text{Ker}(f)$ est un homomorphisme d'anneaux (Exemple 4.3.1.2); l'injection canonique $j : f(A) \longrightarrow B$ est un homomorphisme d'anneaux. On peut montrer enfin comme pour les groupes, que la bijection $\bar{f}$ de $A/\text{Ker}(f)$ sur $f(A)$ définie par $\bar{f}(\bar{x}) = f(x)$ où $x \in \bar{x}$ est un isomorphisme d'anneaux.

On peut donc énoncer :

4.3.2.1. Théorème

Soient $f : A \longrightarrow B$ un morphisme d'anneaux, π le morphisme canonique de A sur $A/\text{Ker}(f)$ et j l'injection canonique de $f(A)$ dans B . Alors il existe un isomorphisme unique $\bar{f}$ de l'anneau-quotient $A/\text{Ker}(f)$ sur le sous-anneau $f(A)$ de B tel que $f = j \circ \bar{f} \circ \pi$.

4.3.3. CARACTÉRISTIQUE D'UN ANNEAU

Soit A un anneau non nul. On vérifie immédiatement que l'application $f : \mathbb{Z} \longrightarrow A$ définie par

$$f(n) = n \cdot 1$$

pour tout $n \in \mathbb{Z}$ est un morphisme d'anneaux tel que $f(1) = 1$. Si g est un morphisme de $\mathbb{Z}$ dans A alors, $g(1) = 1$; on en déduit $g(n) = n \cdot 1$ pour tout $n \in \mathbb{Z}$, donc $g = f$. Il existe donc un morphisme unique f de $\mathbb{Z}$ dans A tel que $f(1) = 1$. Le noyau de f , qui est un idéal de $\mathbb{Z}$, est de la forme $p\mathbb{Z}$, pour un entier $p \in \mathbb{N}$.

Ces considérations nous amènent à poser la définition suivante :

4.3.3.1. Définition

On appelle **caractéristique** de l'anneau non nul A , l'entier $p \geq 0$ tel que $p\mathbb{Z}$ soit le noyau du morphisme $f : \mathbb{Z} \longrightarrow A$ défini par $f(n) = n \cdot 1$.

D'après le Théorème 4.3.2.1, l'image $f(\mathbb{Z})$ de f est un sous-anneau de l'anneau A , isomorphe à $\mathbb{Z}/p\mathbb{Z}$. Donc si le morphisme f est injectif, le seul entier p tel que $p \cdot 1 = 0$ est $p = 0$. On dit que l'anneau A est de **caractéristique nulle**. Dans ce cas, $f(\mathbb{Z})$ est isomorphe à $\mathbb{Z}$, donc A est un ensemble infini.

Si le morphisme f n'est pas injectif, il existe un plus petit entier $p > 0$ tel que $p \cdot 1 = 0$. On dit alors que l'anneau A est de **caractéristique** $p > 0$. Alors tout entier n tel que $n \cdot 1 = 0$ est un multiple de p et pour tout $a \in A$, on a $pa = (p \cdot 1)a = 0 \cdot a = 0$.

4.3.3.2. Exemples

a) Les anneaux $\mathbb{Z}$, $\mathbb{Q}$, $\mathbb{R}$ et $\mathbb{C}$ sont de caractéristique nulle.

b) Soit n un entier ≥ 2 . La caractéristique de l'anneau $\mathbb{Z}/n\mathbb{Z}$ est n . En effet, on voit facilement par récurrence sur $m \in \mathbb{N}$, que $m\bar{x} = \overline{m} \bar{x}$ pour tout $\bar{x} \in \mathbb{Z}/n\mathbb{Z}$. Donc pour tout $\bar{x} \in \mathbb{Z}/n\mathbb{Z}$, on a $n\bar{x} = \overline{n} \bar{x} = \overline{0} \bar{x} = \overline{0 \cdot x} = \overline{0}$.

4.3.3.3. Théorème

Soit A un anneau d'intégrité et soit p sa caractéristique. Alors $p = 0$ ou p est un nombre premier.

Démonstration. Supposons $p > 0$. Si p n'est pas premier, on a $p = mn$ avec $0 < m < p$ et $0 < n < p$. Donc

$$0 = p \cdot 1 = (mn)1 = (m \cdot 1)(n \cdot 1).$$

Comme A est intègre, on a nécessairement $m \cdot 1 = 0$ ou $n \cdot 1 = 0$ ce qui contredit le fait que p est le plus petit entier positif tel que $p \cdot 1 = 0$. Donc p est un nombre premier.

4.4. Divisibilité dans un anneau

Dans ce paragraphe, nous considérons uniquement des anneaux commutatifs et intègres. Dans un tel anneau A l'idéal engendré par l'élément a de A est $aA = Aa$ et est noté (a) .

Notre objectif principal est l'étude sommaire de la divisibilité dans un cadre plus général que celui déjà bien connu de l'anneau $\mathbb{Z}$.

4.4.1. GÉNÉRALITÉS

4.4.1.1. Définition

Soit A un anneau commutatif et intègre et soient a et b deux éléments de A . On dit que a divise b ou que b est divisible par a , et on écrit $a|b$, s'il existe $q \in A$ tel que $b = aq$.

On dit aussi que a est un diviseur de b ou que b est un multiple de a .

Par exemple, tout $a \in A$ est multiple de tout élément inversible u de A car $a = u(u^{-1}a)$ et tout $a \in A$ divise 0 car $0 = a \cdot 0$.

On dit que deux éléments a et b de l'anneau A sont **associés** s'il existe un élément inversible $u \in A$ tel que $b = au$. Alors $a = bv$ avec $v = u^{-1}$.

4.4.1.2. Définition

*Soit A un anneau principal. On dit qu'un élément $p \neq 0$ de A est **premier** ou **irréductible** ou **extrémal** s'il n'est pas inversible et si ses seuls diviseurs sont les éléments inversibles de A et les éléments qui lui sont associés.*

Dans $\mathbb{Z}$, un élément extrémal est appelé un **nombre premier**.

Le théorème suivant montre l'importance des idéaux principaux dans l'étude de la divisibilité.

4.4.1.3. Théorème

Soit A un anneau commutatif et intègre et soient a et b deux éléments de A . Alors :

a) a divise b si et seulement si $(b) \subset (a)$.

b) $(a) = (b)$ si et seulement s'il existe un élément inversible u de A tel que $b = au$.

Démonstration. **a)** Si $a|b$, il existe $q \in A$ tel que $b = aq$; alors pour tout $x \in A$, on a $bx = a(qx)$, donc $(b) \subset (a)$.

Réciproquement, si $(b) \subset (a)$, on a $b = b \cdot 1 \in aA$, donc $b \in aA$ et il existe $c \in A$ tel que $b = ac$, ce qui montre que $a|b$.

b) Si $(a) = (b)$, il existe $q, u \in A$ tels que $a = bq$ et $b = au$; d'où $b = bqu$. Si $b = 0$, $a = b = 0$ et $b = a \cdot 1$. Si $b \neq 0$, comme A est intègre, $qu = 1$ et u est inversible.

Réciproquement, s'il existe un élément inversible u de A tel que $b = au$, alors $a|b$, donc $(b) \subset (a)$. La relation $b = au$ implique $a = bu^{-1}$, donc $b|a$ et par suite $(a) \subset (b)$. Finalement, on a bien $(a) = (b)$.

Dans tout ce qui suit, lorsque nous dirons «soit A un anneau principal», il s'agira toujours d'anneau d'intégrité commutatif dans lequel tous les idéaux sont principaux. Nous nous proposons d'établir certaines propriétés arithmétiques des anneaux principaux; le cas de l'anneau principal $\mathbb{Z}$ est déjà bien connu.

4.4.2. PLUS GRAND COMMUN DIVISEUR

Soit A un anneau principal et soient $a_1, \dots, a_n$ des éléments non nuls de A . Nous nous proposons d'étudier les éléments p de A qui divisent chacun des a_i , c'est-à-dire les éléments p tels qu'il existe des éléments q_i de A vérifiant

$$a_i = pq_i \quad (1 \leq i \leq n).$$

Il est clair que tout diviseur commun à $a_1, \dots, a_n$ divise l'élément $a_1 u_1 + \dots + a_n u_n$ quels que soient les n éléments $u_1, \dots, u_n$ de A . Cela nous conduit à étudier l'ensemble $\mathcal{J}$ des éléments de A de la forme

$$(4.4.2.1) \quad a_1 u_1 + \dots + a_n u_n$$

où $u_1, \dots, u_n$ sont des éléments arbitraires de A .

4.4.2.1. Théorème

soit A un anneau principal et soient $a_1, \dots, a_n$ des éléments non nuls de A . Alors

a) *Il existe un élément $d \in A$, unique à la multiplication près par un élément inversible de A , tel que l'ensemble des multiples de d dans A soit l'ensemble des éléments de la forme $z = a_1 u_1 + \dots + a_n u_n$.*

b) *Il existe des éléments $u_1, \dots, u_n$ de A tels que*

$$(4.4.2.2) \quad d = a_1 u_1 + \dots + a_n u_n.$$

c) *Pour qu'un élément de A divise simultanément $a_1, \dots, a_n$, il faut et suffit qu'il divise d .*

Démonstrations. a) Considérons l'ensemble $\mathcal{J}$ des éléments de A de la forme

$$z = a_1 u_1 + \dots + a_n u_n$$

où $u_1, \dots, u_n \in A$.

Il est clair que $\mathcal{J}$ est un idéal de A . C'est même l'idéal, noté $(a_1, \dots, a_n)$, engendré par $a_1, \dots, a_n$. Comme l'anneau A est principal, $\mathcal{J}$ est engendré par un élément d de A ; autrement dit, on a $\mathcal{J} = (d)$. L'unicité de d résulte du Théorème 4.4.1.3 b).

b) Puisque $d \in (d)$, il existe $u_1, \dots, u_n \in A$ tels que

$$d = a_1 u_1 + \dots + a_n u_n.$$

c) Soit p un diviseur commun à $a_1, \dots, a_n$. Écrivons

$$a_i = p q_i \quad (1 \leq i \leq n).$$

En portant dans 4.4.2.2, il vient

$$d = p(q_1 u_1 + \dots + q_n u_n)$$

ce qui montre que p divise d .

Réciproquement, soit c un diviseur de d . La relation $\mathcal{J} = (d)$ montre que l'idéal (d) contient les a_i qui sont donc des multiples de d ; par suite c est un diviseur commun à $a_1, \dots, a_n$.

4.4.2.2. Définition

On appelle plus grand commun diviseur (P.G.C.D.) de $a_1, \dots, a_n$, et on note P.G.C.D. $(a_1, a_2, \dots, a_n)$, tout élément d de A tel que

$$dA = a_1 A + \dots + a_n A.$$

La relation 4.4.2.2 s'appelle l'identité de Bezout.

4.4.2.3. Définition

On dit que les éléments $a_1, \dots, a_n$ de A sont premiers entre eux dans leur ensemble s'ils admettent 1 pour P.G.C.D.

On dit que les a_i sont premiers entre eux deux à deux si $\text{PGCD}(a_i, a_j) = 1$ pour tous les i, j tels que $i \neq j$.

On remarque que si $n \geq 3$ et si $a_1, \dots, a_n$ sont premiers entre eux deux à deux, ils sont premiers entre eux dans leur ensemble ; cependant si $a_1, \dots, a_n$ sont premiers entre eux dans leur ensemble, ils ne sont pas nécessairement premiers entre eux deux à deux. Par exemple, dans $\mathbb{Z}$, les nombres 12, 15 et 16 sont premiers entre eux mais il ne sont pas premiers entre eux deux à deux puisque $\text{PGCD}(12, 15) = 3$.

4.4.2.4. Théorème (Bezout)

Soient A un anneau principal et $a_1, \dots, a_n$ des éléments de A . Les propriétés suivantes sont équivalentes :

a) $a_1, \dots, a_n$ sont premiers entre eux dans leur ensemble.

b) Il existe des éléments $u_1, \dots, u_n$ de A tels que

$$a_1 u_1 + \dots + a_n u_n = 1.$$

Démonstration. a) $\implies$ b) : Si $a_1, \dots, a_n$ sont premiers entre eux dans leur ensemble, 1 est un PGCD de $a_1, \dots, a_n$, donc un élément de l'idéal engendré par $a_1, \dots, a_n$. Donc il existe $u_1, \dots, u_n \in A$ tels qu'on ait

$$a_1 u_1 + \dots + a_n u_n = 1.$$

b) $\implies$ a) : S'il existe des éléments $u_1, \dots, u_n$ de A tels que $a_1 u_1 + \dots + a_n u_n = 1$, alors $1 \in (a_1, \dots, a_n)$, donc $(a_1, \dots, a_n) = A$ et par suite 1 est un PGCD de $a_1, \dots, a_n$.

4.4.2.5. Corollaire

Soient a, b et c des éléments de l'anneau principal A . Si a est premier séparément avec b et c , alors il est premier avec le produit bc .

Démonstration. D'après l'identité de Bezout, il existe des éléments u_1, v_1, x, y de A tels que

$$au_1 + bv_1 = 1 \quad \text{et} \quad ax + cy = 1.$$

D'où, en multipliant membre à membre :

$$au + bcv = 1$$

avec $u = au_1 x + cu_1 y + bv_1 x$ et $v = v_1 y$, ce qui montre que a et bc sont premiers entre eux.

4.4.2.6. Théorème (Gauss)

Soient a et b des éléments non nuls de A et soit d un diviseur du produit ab . Si d est premier avec a , alors d divise b .

Démonstration. Si d et a sont premiers entre eux, il existe $u, v \in A$ tels que

$$du + av = 1.$$

En multipliant les deux membres de cette identité par b , il vient

$$bdu + abv = b;$$

puisque d divise à la fois bdu et ab (par hypothèse), il divise b .

4.4.2.7. Corollaire

Soient $b_1, \dots, b_n$ des éléments de A premiers entre eux deux à deux. Si un élément a de A est divisible séparément par $b_1, \dots, b_n$, alors a est divisible par le produit $b_1 b_2 \dots b_n$.

Démonstration. Il suffit de démontrer le résultat dans le cas où $n = 2$; le théorème s'obtient en faisant une récurrence facile sur n .

Par hypothèse, il existe $c_1 \in A$ tel que $a = b_1 c_1$; b_2 , divisant $b_1 c_1$ et étant premier avec b_1 , divise c_1 . Il existe donc $c_2 \in A$ tel que $c_1 = b_2 c_2$; d'où $a = b_1 b_2 c_2$.

4.4.3. PLUS PETIT COMMUN MULTIPLE

Soient $a_1, \dots, a_n$ des éléments non nuls de l'anneau principal A . Tout multiple de a_i est un élément de l'idéal (a_i) ; donc les multiples communs aux a_i sont les éléments de l'idéal $(a_1) \cap \dots \cap (a_n)$. Cet idéal étant principal, il existe un élément m de A , unique à la multiplication près par un élément inversible quelconque de A , tel que

$$(4.4.3.1) \quad (a_1) \cap \dots \cap (a_n) = (m)$$

et tout multiple commun aux a_i est un multiple de m .

On est ainsi amené à poser la définition suivante.

4.4.3.1. Définition

On appelle **plus petit commun multiple (P.P.C.M.)** de $a_1, \dots, a_n$, tout élément m de A tel que

$$(a_1) \cap \dots \cap (a_n) = (m).$$

4.4.3.2. Théorème

Soient A un anneau principal, a et b des éléments non nuls de A , d un P.G.C.D., m un P.P.C.M. de a et b . Alors $ab = md$ à la multiplication près par un élément inversible.

Démonstration. Posons $a = a'd$ et $b = b'd$; alors a' et b' sont premiers entre eux. En effet, il existe $u, v \in A$ tels que $au = bv = d$, d'où en simplifiant par d , la relation $a'u + b'v = 1$, qui montre que P.G.C.D. $(a', b') = 1$.

On peut écrire $a'b'd = ab' = a'b$, donc $a'b'd \in (m)$.

Réciproquement, en écrivant $m = xa$ et $m = yb$, on a: $xa = yb$ ou $xa'd = yb'd$. On en déduit $xa' = yb'$ puisque d est un élément non nul de l'anneau intègre A . Comme P.G.C.D. $(a', b') = 1$, on peut écrire (Théorème de

Gauss), $x = zb'$, d'où $m = zb'a = za'b'd$. Donc $m \in (a'b'd)$ et par suite, on a $(m) = (a'b'd)$. D'après le Théorème 4.4.1.3b), il existe un élément inversible u de A tel que $m = a'b'du$, d'où $md = a'b'dud = abu$.

4.5. Corps

4.5.1. DÉFINITIONS. PROPRIÉTÉS FONDAMENTALES

4.5.1.1. Définition

On appelle **corps** tout anneau K non nul dans lequel tout élément non nul est inversible.

On dit qu'un corps est **commutatif** si sa multiplication est commutative.

Ainsi pour un corps K on a, $K^\times = K^*$.

Si 1 est l'élément unité du groupe multiplicatif $K^\times$, alors 1 est l'élément unité de K . Ainsi un corps possède toujours au moins les deux éléments 0 et 1.

Les notions définies pour les anneaux (intégrité, morphisme, idéal, caractéristique) s'appliquent également aux corps qui sont des anneaux particuliers, mais certains résultats prendront ici des formes particulières.

4.5.1.2. Exemple

Les anneaux $\mathbb{Q}$, $\mathbb{R}$ et $\mathbb{C}$ sont des corps commutatifs de caractéristique 0.

4.5.1.3. Exemple

L'ensemble $\mathbb{Q}[\sqrt{2}] = \{a + b\sqrt{2} : a, b \in \mathbb{Q}\}$ muni de l'addition et de la multiplication ordinaires est un corps commutatif.

Dans un corps commutatif, on écrit souvent $xy^{-1} = y^{-1}x = x/y$. On vérifie facilement que toutes les règles de calculs habituelles dans $\mathbb{R}$ et $\mathbb{C}$ sont valables dans un corps commutatif.

Le résultat suivant fournit un exemple important de corps.

4.5.1.4. Théorème

a) Soit $a \in \mathbb{Z}$. Dans l'anneau $\mathbb{Z}/n\mathbb{Z}$, l'élément $\bar{a}$ est inversible si et seulement si a et n sont premiers entre eux.

b) L'anneau $\mathbb{Z}/n\mathbb{Z}$ est un corps si et seulement si n est un nombre premier.

Démonstration. a) On a les équivalences :

a inversible $\iff$ il existe $u \in \mathbb{Z}$ tel que $\bar{a} \bar{u} = \bar{1}$

$\iff au \equiv 1 \pmod{n}$

$\iff$ il existe $v \in \mathbb{Z}$ tel que $au - 1 = vn$ soit $au - vn = 1$, ce qui, d'après l'identité de Bezout, signifie que a et n sont premiers entre eux.

b) Rappelons que $\mathbb{Z}/n\mathbb{Z} = \{\bar{0}, \bar{1}, \dots, \overline{n-1}\}$.

Supposons que n soit premier et soit $p \in \mathbb{N}$ tel que $0 < p < n$; alors n et p sont premiers entre eux, donc d'après a), $\bar{p}$ est inversible et $\mathbb{Z}/n\mathbb{Z}$ est un corps.

Inversement, si $\mathbb{Z}/n\mathbb{Z}$ est un corps, tout élément $\bar{m}$ ($0 < m < n$) de $\mathbb{Z}/n\mathbb{Z}$ est inversible; donc d'après a), m et n sont premiers entre eux. n n'admet donc pas d'autres diviseurs positifs que 1 et lui-même. Donc n est un nombre premier.

4.5.1.5. Théorème

a) Tout corps K est intègre.

b) Dans un corps K , tout élément non nul est régulier pour la multiplication de K .

Démonstration. a) Soient $a, b \in K$ tels que $ab = 0$ et $a \neq 0$. Alors a^{-1} existe et $a^{-1} \in K$. La relation $ab = 0$ implique

$$0 = a^{-1}0 = a^{-1}(ab) = (a^{-1}a)b = 1 \cdot b = b.$$

b) On sait que $K^* = K - \{0\}$ est un groupe multiplicatif; donc tout élément de K^* est régulier pour la multiplication de K d'après le Théorème 2.2.4.2.

4.5.1.6. Remarque

Le théorème 4.3.3.3 montre que la caractéristique d'un corps est, soit 0, soit un nombre premier.

4.5.2. SOUS-CORPS. IDÉAUX D'UN CORPS. MORPHISMES DE CORPS

4.5.2.1. Définition

Soit K un corps. On dit qu'une partie K' de K est un sous-corps de K si :

a) K' est un sous-anneau de K .

b) Les relations $x \neq 0$ et $x \in K'$ impliquent $x^{-1} \in K'$.

On vérifie alors que l'ensemble K' , muni des lois de composition induites par celles de K est un corps. On dit aussi que K est un sur-corps ou une extension du sous-corps K' .

Par exemple $\mathbb{Q}$ est un sous-corps de $\mathbb{R}$ et $\mathbb{R}$ est un sous-corps de $\mathbb{C}$.

On montre facilement que toute intersection de sous-corps d'un corps K est un sous-corps de K ; on peut donc définir le sous-corps engendré par une partie A de K comme étant l'intersection de tous les sous-corps contenant A .

4.5.2.2. Théorème

Soit M un idéal d'un anneau commutatif A . Alors M est maximal si et seulement si l'anneau quotient A/M est un corps.

Démonstration. Supposons M maximal et soit a un élément de A n'appartenant pas à M . L'ensemble $aA + M$ est un idéal de A contenant M et différent de M ; donc $aA + M = A$. Alors, il existe un élément z de A et un élément m de M tels que $az + m = 1$. En prenant les classes modulo M , on obtient $\overline{az} = \overline{a} \overline{z} = \overline{1}$, ce qui montre que la classe de a est inversible dans A/M ; autrement dit l'anneau quotient A/M est un corps.

Réciproquement, supposons que A/M soit un corps. Soit I un idéal contenant M . Prenons un élément quelconque a de I , n'appartenant pas à M . La classe de a n'est pas la classe nulle et est donc inversible dans A/M . Il existe un élément u de A tel que $\overline{1} = \overline{a} \overline{u} = \overline{au}$, ou $\overline{1 - au} = \overline{0}$; donc $1 - au \in M$ et par suite $1 = (1 - au) + au \in I$. On en déduit que $I = A$ d'après la Remarque 4.2.2.5a). Donc l'idéal I est maximal.

4.5.2.3. Théorème

Soit K un anneau commutatif non nul. Pour que K soit un corps il faut et il suffit que les seuls idéaux de K soient $\{0\}$ et K .

Démonstration. Supposons que K soit un corps et soit I un idéal de K distinct de $\{0\}$. Soit x un élément non nul de I . Comme K est un corps, x^{-1} existe et $x^{-1} \in K$; donc $x^{-1}x = 1 \in I$ puisque I est un idéal. Ainsi, pour tout $a \in K$, $1 \cdot a = a \in I$, ce qui prouve que $K \subset I$; donc $I = K$.

Réciproquement, supposons que les seuls idéaux de K soient $\{0\}$ et K . Soit a un élément non nul de K et considérons l'idéal principal (a) engendré par a :

$$(a) = \{ax : x \in K\}.$$

$(a) \neq \{0\}$ puisque $a = a \cdot 1 \in (a)$ et $a \neq 0$. On a donc nécessairement $(a) = K$. Comme $1 \in (a)$, il existe un élément $x \in K$ tel que $ax = xa = 1$; autrement dit x est l'inverse de a . Ainsi tout élément non nul de K est inversible dans K ; donc K est un corps.

Le théorème précédent est très utile dans l'étude des morphismes de corps.

4.5.2.4. Théorème

Soient K et K' deux corps commutatifs et $f : K \rightarrow K'$ un morphisme de corps. Alors f est une application injective.

Démonstration. f étant un morphisme d'anneaux, $\text{Ker}(f)$ est un idéal de K ; donc $\text{Ker}(f) = \{0\}$ ou $\text{Ker}(f) = K$. Si $\text{Ker}(f) = K$, f est nulle, ce qui est impossible puisque $f(1) = 1 \neq 0 \in K'$; donc $\text{Ker}(f) = \{0\}$ et par suite, f est injective.

4.5.3. CORPS DES FRACTIONS D'UN ANNEAU COMMUTATIF INTÈGRE

Étant donné un anneau commutatif et intègre A , nous nous proposons de trouver un corps commutatif K contenant A comme sous-anneau.

4.5.3.1. Définition

On dit qu'un corps K est un corps des fractions de l'anneau intègre A si les deux conditions suivantes sont vérifiées :

a) A est un sous-anneau du corps K .

b) Pour tout $x \in K$, il existe dans A des éléments a et b tels que $x = ab^{-1}$.

4.5.3.2. Théorème

Tout anneau commutatif intègre A admet un corps des fractions. Si K et L sont des corps des fractions de l'anneau A , alors K et L sont isomorphes.

Démonstration. Soit E l'ensemble des couples (p, q) où $p \in A$, $q \in A$ et $q \neq 0$. Sur E , la relation $\mathcal{R}$ définie par $(p, q)\mathcal{R}(p', q') \iff pq' = qp'$ est une relation d'équivalence.

Il est clair que la relation $\mathcal{R}$ est réflexive et symétrique. Montrons qu'elle est transitive.

Si $(p, q)\mathcal{R}(p', q')$ et $(p', q')\mathcal{R}(p'', q'')$, on a $pq' = qp'$ et $p'q'' = q'p''$, donc $pq'q'' = qp'q'' = qq'p''$. Comme A est intègre et $q' \neq 0$, on en déduit $pq'' = qp''$, donc $(p, q)\mathcal{R}(p'', q'')$. Ainsi $\mathcal{R}$ est bien une relation d'équivalence.

Soit K l'ensemble quotient $E/\mathcal{R}$; notons φ l'application canonique de E sur $E/\mathcal{R}$.

On définit deux lois internes sur E en posant :

Addition : $(p, q) + (r, s) = (ps + qr, qs)$

Multiplication : $(p, q) \cdot (r, s) = (pr, qs)$.

On vérifie sans peine que l'addition et la multiplication ainsi définies sont associatives, commutatives, admettent pour élément neutre $(0,1)$ et $(1,1)$ respectivement, et que la multiplication est distributive par rapport à l'addition.

On vérifie également qu'elles sont compatibles avec la relation d'équivalence $\mathcal{R}$; par exemple, pour l'addition si

$$(p, q)\mathcal{R}(p', q') \quad \text{et} \quad (r, s)\mathcal{R}(r', s')$$

la relation

$$((p, q) + (r, s)) \mathcal{R} ((p', q') + (r', s'))$$

qui s'écrit

$$(ps + qr)q's' - (p's' + q'r')qs = 0$$

ou

$$ss'(pq' - p'q) + qq'(rs' - r's) = 0$$

est vérifiée.

Dans l'ensemble quotient notons encore $+$ et $\bullet$ les lois quotients. Ces lois sont associatives, commutatives et la multiplication est distributive par rapport à l'addition.

Pour l'addition, $\varphi(0, 1)$ est l'élément neutre et $\varphi(-p, q)$ est l'opposé de $\varphi(p, q)$; donc $(K, +)$ est un groupe abélien.

Pour la multiplication, $\varphi(1, 1)$ est l'élément neutre. Donc $(K, +, \bullet)$ est un anneau commutatif.

En outre, si $\varphi(p, q)$ est différent de $\varphi(0, 1)$, c'est-à-dire si $p \neq 0$, $\varphi(p, q)$ admet pour inverse $\varphi(q, p)$. En conclusion $(K, +, \bullet)$ est un corps commutatif.

Considérons l'application $\psi : A \longrightarrow K$, définie par $\psi(x) = \varphi(x, 1)$.

On a

$$\psi(x + y) = \varphi(x + y, 1) = \varphi(x, 1) + \varphi(y, 1) = \psi(x) + \psi(y)$$

$$\psi(xy) = \varphi(xy, 1) = \varphi(x, 1) \cdot \varphi(y, 1) = \psi(x) \psi(y)$$

$$\psi(1) = \varphi(1, 1) = 1_K.$$

ψ est donc un morphisme d'anneaux. Ce morphisme est injectif car si $\varphi(x, 1) = 0$, alors $x = 0$. Donc l'anneau A s'identifie au sous-anneau $\psi(A)$ de K , puisque ψ est un isomorphisme de A sur le sous-anneau $\psi(A)$ de K .

Supposons maintenant qu'il existe un corps L et un morphisme injectif $\theta : A \longrightarrow L$. La relation $(p, q)\mathcal{R}(p', q')$ qui s'écrit $pq' = p'q$ entraîne $\theta(p)\theta(q') = \theta(p')\theta(q)$. Comme $q \neq 0$ et $q' \neq 0$, on a $\theta(q) \neq 0$ et $\theta(q') \neq 0$ puisque θ est injectif; ainsi $\theta(p)\theta(q)^{-1}$ ne change pas lorsqu'on remplace (p, q) par un autre représentant de $\varphi(p, q)$. Alors l'application $\varphi(p, q) \longmapsto \theta(p)\theta(q)^{-1}$ de K dans L est un morphisme injectif de corps.

Le corps K qui vient d'être construit répond donc à la question et c'est, à un isomorphisme près, le seul corps répondant à la question.

4.5.3.3. Remarque

Si l'on identifie A au sous-anneau $\psi(A)$ de K , ce qui revient à remplacer l'élément $\psi(p)$ de $\psi(A)$ par p , alors l'élément générique $\psi(p, q)$ de K prend la forme pq^{-1} que l'on écrit aussi p/q .

4.5.3.4. Remarque

Si l'on applique la construction du corps des fractions en partant de l'anneau $\mathbb{Z}$, on obtient le corps $\mathbb{Q}$ des nombres rationnels. Tout nombre rationnel s'écrit donc p/q , avec $p \in \mathbb{Z}$ et $q \in \mathbb{Z} - \{0\}$.

Nous étudierons au Chapitre 5 un autre exemple important des corps de fractions : le corps $K(X)$ des fractions rationnelles à une indéterminée à coefficients dans un anneau commutatif intègre K .

Chapitre 5 : POLYNÔMES ET FRACTIONS RATIONNELLES

Dans ce chapitre, nous nous proposons de définir, à partir d'un anneau commutatif K donné, un nouvel anneau commutatif noté $K[X]$ et contenant K comme sous-anneau. Cet anneau $K[X]$ possède de nombreuses «propriétés arithmétiques» de l'anneau $\mathbb{Z}$; nous étudierons quelques-unes de ces propriétés.

Les polynômes, c'est-à-dire les éléments de l'anneau $K[X]$ seront introduits comme expressions formelles.

Si K est un anneau commutatif intègre, nous verrons que l'anneau $K[X]$ est intègre; cet anneau n'est pas un corps mais il possède un corps des fractions. Si en particulier K est un corps commutatif, le corps des fractions de l'anneau $K[X]$ est appelé l'anneau des fractions rationnelles à une indéterminée à coefficients dans K . Ce corps sera étudié dans la deuxième partie.

Polynômes

5.1. Définitions générales

5.1.0.1. Définition

*Soit K un anneau commutatif. On appelle **polynôme à une indéterminée à coefficients dans K** , toute suite $(a_0, a_1, \dots)$ d'éléments de K nuls à partir d'un certain rang.*

Un tel polynôme est noté

$$P = (a_0, a_1, \dots, a_n, \dots) \quad \text{ou} \quad P = (a_n)_{n \in \mathbb{N}}.$$

Les a_n sont appelés les **coefficients** du polynôme P ; on dit que a_n est le **coefficient d'indice n** ; a_0 s'appelle le **terme constant**.

Si tous les coefficients a_n du polynôme P sont nuls, on dit que P est le **polynôme nul** et on le note 0. On appelle **monôme** un polynôme dont tous les coefficients sont nuls sauf, au plus, l'un d'entre eux.

L'ensemble des polynômes à une indéterminée à coefficients dans l'anneau commutatif K se note $K[X]$.

Étant donné un polynôme $P = (a_0, a_1, \dots, a_n \dots)$ à coefficients dans K , on appelle **degré** de P , et on note $\deg(P)$, le plus grand entier n tel que $a_n \neq 0$. Cette définition s'applique à tout polynôme de $K[X]$ sauf au polynôme nul 0. En effet, tous les coefficients de 0 étant nuls, l'ensemble des indices des coefficients non nuls de 0 est vide, donc cet ensemble n'admet pas de plus grand élément.

De même, le plus petit entier k tel que $a_k \neq 0$ s'appelle la **valuation** du polynôme P et se note $\text{val}(P)$. Tout polynôme P sauf 0 admet une valuation.

On convient de noter $\deg(0) = -\infty$ et $\text{val}(0) = +\infty$ où les symboles $-\infty$ et $+\infty$ sont assujettis à vérifier les relations suivantes :

Pour tout $n \in \mathbb{Z}$, $-\infty < n < +\infty$

Pour tout $n \in \mathbb{Z}$, $n + (+\infty) = +\infty$; $n + (-\infty) = -\infty$.

$$(+\infty) + (+\infty) = +\infty; (-\infty) + (-\infty) = -\infty.$$

Si P est un monôme non nul, c'est-à-dire si un seul coefficient a_n est non nul, alors $\deg(P) = \text{val}(P) = n$.

5.1.0.2. Exemple

Prenons $K = \mathbb{Z}$ et $P = (0, 1, 0, 4, 0, 0, \dots)$. On a $\deg(P) = 3$ et $\text{val}(P) = 1$.

5.1.0.3. Définition

On dit que deux polynômes $P = (a_n)_{n \in \mathbb{N}}$ et $Q = (b_n)_{n \in \mathbb{N}}$ de $K[X]$ sont **égaux** si pour tout entier n , on a $a_n = b_n$.

En particulier P est le polynôme nul si et seulement si, pour tout entier n , $a_n = 0$.

5.2. Structure d'anneau de $K[X]$

Dans ce paragraphe K désigne un anneau commutatif sauf mention expresse du contraire.

5.2.1. ADDITION DE DEUX POLYNÔMES

5.2.1.1. Définition

Soient $P = (a_n)_{n \geq 0}$ et $Q = (b_n)_{n \geq 0}$ deux polynômes de $K[X]$. On appelle **somme** de P et Q , et on note $P + Q$, le polynôme dont le coefficient d'indice n est égal à $a_n + b_n$:

$$(5.2.1.1) \quad P + Q = (a_0 + b_0, \dots, a_n + b_n, \dots).$$

5.2.1.2. Théorème

Le couple $(K[X], +)$ est un groupe abélien.

Démonstration. Comme $(K, +)$ est un groupe abélien, pour tout $n \in \mathbb{N}$, $a_n + b_n \in K$; donc $(a_0 + b_0, \dots, a_n + b_n, \dots)$ est une suite d'éléments de K . Si l'un des polynômes est nul, $P + Q$ est égal à l'autre polynôme, donc la suite $(a_n + b_n)_{n \geq 0}$ possède un nombre fini d'éléments non nuls et $P + Q \in K[X]$. Si aucun des polynômes n'est nul, soient n et m les degrés de P et Q respectivement; si $k > \max(n, m)$, on aura $a_k = b_k = 0$, donc $a_k + b_k = 0$, ce qui montre que $P + Q \in K[X]$.

Les propriétés de l'addition dans $K[X]$ se déduisent facilement de celles de l'addition dans K . Ainsi, pour tout $n \in \mathbb{N}$, $(a_n + b_n) + c_n = a_n + (b_n + c_n)$, donc $(P + Q) + R = P + (Q + R)$ et pour tout $n \in \mathbb{N}$, $a_n + b_n = b_n + a_n$, donc $P + Q = Q + P$. Le polynôme 0 est l'élément neutre pour l'addition dans $K[X]$ et tout polynôme $P = (a_n)_{n \geq 0}$ a pour opposé le polynôme, noté $-P$, tel que $-P = (-a_n)_{n \geq 0}$.

5.2.1.3. Théorème

Posons $K[X]^ = K[X] - \{0\}$. Alors si $P, Q \in K[X]^*$ et si $Q \neq -P$, on a*

$$(5.2.1.2) \quad \deg(P + Q) \leq \max(\deg(P), \deg(Q))$$

$$(5.2.1.3) \quad \text{val}(P + Q) \geq \min(\text{val}(P), \text{val}(Q)).$$

Démonstration. Il existe deux possibilités: $\deg(P) = \deg(Q)$ et $\deg(P) \neq \deg(Q)$.

1^{er} cas : $\deg(P) = \deg(Q) = n$.

- Si $a_n + b_n = 0$, alors $\deg(P + Q) < \deg(P)$, donc (5.2.1.2) est vraie, l'inégalité stricte ayant lieu.

- Si $a_n + b_n \neq 0$, alors $\deg(P + Q) = \deg(P) = \deg(Q)$, donc (5.2.1.2) est encore vraie, l'égalité ayant lieu dans ce cas.

2^e cas : $\deg(P) \neq \deg(Q)$. Supposons par exemple que $\deg(P) < \deg(Q)$. Alors, par définition du degré et de la somme, on a $\deg(P + Q) = \deg(Q) = \max(\deg(P), \deg(Q))$. Donc (5.2.1.2) est vérifiée.

On démontrerait de même la relation (5.2.1.3), en raisonnant sur a_k et b_r , où $\text{val}(P) = k$ et $\text{val}(Q) = r$ et en considérant les deux cas $k = r$ et $k \neq r$.

5.2.2. MULTIPLICATION DE DEUX POLYNÔMES

5.2.2.1. Définition

Soient $P = (a_n)_{n \geq 0}$ et $Q = (b_n)_{n \geq 0}$ deux polynômes de $K[X]$. On appelle produit de P et Q , et on note PQ , le polynôme dont le coefficient d'indice n est défini par :

$$(5.2.2.1) \quad c_n = \sum_{i+j=n} a_i b_j = \sum_{k=0}^n a_k b_{n-k}.$$

Montrons que l'on définit bien ainsi un polynôme à coefficients dans K . Si $a_i = 0$ pour $i > n_0$ et $b_j = 0$ pour $j > m_0$, on a $c_n = 0$ pour $n > n_0 + m_0$. En effet, si $n > n_0 + m_0$, dans chaque terme $a_i b_j$ de c_n , on a soit $i > n_0$ et alors $a_i = 0$, soit $j > m_0$ et alors $b_j = 0$; donc chaque terme de c_n est nul, et par suite $PQ \in K[X]$.

5.2.2.2. Théorème

Le triplet $(K[X], +, \cdot)$ est un anneau commutatif.

Démonstration. Nous savons déjà que $(K[X], +)$ est un groupe abélien et que la multiplication est une loi interne sur $K[X]$. Il suffit donc de vérifier que la multiplication des polynômes est commutative, associative, distributive par rapport à l'addition et qu'elle admet un élément neutre.

Soit $1 = (1, 0, 0, \dots)$ le polynôme constant dont tous les coefficients sont nuls sauf a_0 qui vaut 1. On vérifie facilement que pour tout $P \in K[X]$, $1 \cdot P = P \cdot 1 = P$, donc le polynôme 1 est l'élément unité de l'anneau $K[X]$.

Montrons que la multiplication est associative.

Soient $P = (a_n)_{n \geq 0}$, $Q = (b_n)_{n \geq 0}$ et $R = (c_n)_{n \geq 0}$ trois polynômes de $K[X]$.

On a

$$PQ = (d_n)_{n \geq 0} \quad \text{avec} \quad d_n = \sum_{i+j=n} a_i b_j$$

$$QR = (e_n)_{n \geq 0} \quad \text{avec} \quad e_n = \sum_{r+s=n} b_r c_s$$

Écrivons le terme u_n de rang n du produit $(PQ)R$:

$$\begin{aligned} u_n &= d_0 c_n + \dots + d_k c_{n-k} + \dots + d_n c_0 \\ &= a_0 b_0 c_n + (a_0 b_1 + a_1 b_0) c_{n-1} + \dots + (a_0 b_n + a_1 b_{n-1} + \dots + a_n b_0) c_0. \end{aligned}$$

De même le terme général v_n du produit $P(QR)$ est :

$$\begin{aligned} v_n &= a_0 e_n + \dots + a_k e_{n-k} + \dots + a_n e_0 \\ &= a_0 (b_0 c_n + b_1 c_{n-1} + \dots + b_n c_0) + \dots + a_n b_0 c_0. \end{aligned}$$

Si on utilise la distributivité de la multiplication par rapport à l'addition dans K et l'associativité de ces lois dans K , on trouve $u_n = v_n$ pour tout $n \in \mathbb{N}$. Donc les deux polynômes $(PQ)R$ et $P(QR)$ sont égaux.

On démontrerait par une méthode analogue que

$$P(Q + R) = PQ + PR \quad \text{et} \quad (Q + R)P = QP + RP.$$

5.2.2.3. Théorème

a) L'anneau $K[X]$ est intègre si et seulement si l'anneau K est intègre.

b) Si $P, Q \in K[X]$, on a

$$(5.2.2.2) \quad \deg(PQ) \leq \deg(P) + \deg(Q)$$

$$(5.2.2.3) \quad \text{val}(PQ) \geq \text{val}(P) + \text{val}(Q)$$

Dans les relations (5.2.2.1) et (5.2.2.2) il y a égalité si l'anneau K est intègre.

Démonstration. a) Si l'anneau K possède des diviseurs de zéro, il existe des éléments $a \neq 0$ et $b \neq 0$ de K tels que $ab = 0$. Alors dans $K[X]$, on a

$$(a, 0, \dots)(b, 0, \dots) = (ab, 0, \dots) = 0,$$

donc $K[X]$ possède des diviseurs de zéro.

Si au contraire K est intègre, montrons que $K[X]$ est intègre. Soient $P = (a_n)_{n \geq 0}$ et $Q = (b_n)_{n \geq 0}$ deux polynômes non nuls; posons $k = \text{val}(P)$ et $r = \text{val}(Q)$.

$$P = (0, \dots, 0, a_k, a_{k+1}, \dots) \quad \text{avec } a_k \neq 0.$$

$$Q = (0, \dots, 0, b_r, b_{r+1}, \dots) \quad \text{avec } b_r \neq 0.$$

Alors par définition de la multiplication,

$$PQ = (0, \dots, 0, a_k b_r, \dots).$$

Comme K est intègre, les relations $a_k \neq 0$ et $b_r \neq 0$ impliquent $a_k b_r \neq 0$; donc $PQ \neq 0$.

b) Les inégalités sont évidentes si l'un des deux polynômes est nul. Écartons ce cas. Supposons que $\deg(P) = n$ et $\deg(Q) = m$. La formule (5.2.2.1) montre que $c_p = 0$ si $p > n + m$ et $c_{n+m} = a_n b_m$. Si l'anneau K possède des diviseurs de zéro, rien ne dit que $c_{n+m} \neq 0$. On a donc en général

$$\deg(PQ) \leq n + m = \deg(P) + \deg(Q).$$

Si l'anneau K est intègre, les relations $a_n \neq 0$ et $b_m \neq 0$ entraînent $c_{n+m} = a_n b_m \neq 0$. Donc

$$\deg(PQ) = \deg(P) + \deg(Q).$$

On démontrerait de même l'inégalité (5.2.2.3) en posant $\text{val}(P) = k$ et $\text{val}(Q) = r$ et en remarquant que le premier coefficient non nul de PQ est $a_k b_r$ si l'anneau K est intègre.

5.2.2.4. Théorème

Soit K un corps commutatif. Le groupe des éléments inversibles de l'anneau $K[X]$ est l'ensemble des polynômes de degré 0.

Démonstration. Si le polynôme $P \in K[X]$ est inversible et si Q est son inverse, on a $PQ = 1$; d'où $\deg(PQ) = \deg(P) + \deg(Q) = \deg(1) = 0$. Donc $\deg(P) = \deg(Q) = 0$, ce qui, par définition du degré, montre que P et Q sont de la forme

$$P = (a_0, 0, \dots) \quad \text{et} \quad Q = (b_0, 0, \dots).$$

Réciproquement dans $K[X]$, tout élément $\lambda \in K^*$ est inversible, son inverse étant le polynôme constant λ^{-1} (inverse de λ dans K).

5.3. Notation définitive

Dans ce paragraphe on désigne par K un anneau commutatif.

5.3.1. IMMERSION DE K DANS $K[X]$

Considérons l'application $f : a \mapsto f(a) = (a, 0, \dots)$ de K dans $K[X]$; elle est évidemment injective. C'est aussi un morphisme d'anneaux car les formules (5.2.1.1) et (5.2.2.1) montrent que

$$f(a + b) = f(a) + f(b), \quad f(ab) = f(a) f(b)$$

et de plus,

$$f(0) = 0, \quad f(1) = 1.$$

Par suite, f est un isomorphisme de K sur $f(K)$. On identifiera désormais, grâce à cet isomorphisme, tout élément de K avec son image dans $K[X]$ en posant, pour tout $a \in K$:

$$(5.3.1.1) \quad a = (a, 0, \dots).$$

Si $\lambda \in K$ et si $P = (a_0, a_1, \dots)$ est un polynôme, on a

$$(5.3.1.2) \quad \lambda P = (\lambda, 0, \dots) (a_0, a_1, \dots) = (\lambda a_0, \lambda a_1, \dots).$$

On vérifie aisément que si $\lambda, \mu \in K$ et $P, Q \in K[X]$, alors

$$(5.3.1.3) \quad (\lambda + \mu)P = \lambda P + \mu P, \quad \lambda(P + Q) = \lambda P + \lambda Q.$$

$$(5.3.1.4) \quad (\lambda \mu)P = \lambda(\mu P), \quad 1 \cdot P = P.$$

On exprime ces propriétés en disant que $K[X]$, muni des opérations $(P, Q) \mapsto P + Q$ et $(\lambda, P) \mapsto \lambda P$ possède une structure de K -module (nous y reviendrons).

Remarquons en outre que, quels que soient $P, Q \in K[X]$ et quel que soit $\lambda \in K$, on a

$$(5.3.1.5) \quad \lambda(PQ) = (\lambda P)Q = P(\lambda Q).$$

5.3.2. NOTION D'INDÉTERMINÉE

5.3.2.1. Définition

On appelle **indéterminée** le polynôme X dont tous les coefficients sont nuls, sauf le coefficient d'indice $1 \in \mathbb{N}$ qui est égal à $1 \in K$:

$$(5.3.2.1) \quad X = (0, 1, 0, \dots).$$

La formule (5.2.2.1) donne facilement

$$X^2 = (0, 0, 1, 0, \dots)$$

$$X^3 = (0, 0, 0, 1, 0, \dots)$$

$$-----$$

$$X^n = (0, 0, \dots, 1, 0, \dots)$$

où le coefficient 1 de X^n se trouve au $(n + 1)^{\text{ème}}$ rang.

En utilisant la formule (5.3.1.2), on voit que pour tout $a_n \in K$,

$$a_n X^n = (0, \dots, 0, a_n, 0, \dots)$$

d'où, si $a_0, a_1, \dots$, sont des éléments de K et si $a_k = 0$ pour $k > n$,

$$(a_0, a_1, \dots) = a_0 + a_1 X + \dots + a_n X^n.$$

Nous écrivons donc désormais le polynôme $P = (a_0, a_1, \dots)$ de degré n sous la forme

$$(5.3.2.2) \quad a_0 + a_1 X + \dots + a_n X^n = \sum_{k=0}^n a_k X^k.$$

Quand on écrit P sous la forme $a_0 + a_1 X + \dots + a_n X^n$, on dit que P est **ordonné suivant les puissances croissantes de X** ; si on écrit : $P = a_n X^n + \dots + a_0$, on dit que P est **ordonné suivant les puissances décroissantes de X** .

Le coefficient a_n est appelé **coefficient dominant** de P ; lorsque $a_n = 1$, on dit que le polynôme P est **unitaire** ou mieux **normalisé**.

Nous verrons plus tard que si K est un corps commutatif, alors $K[X]$ est un K -espace vectoriel ; de même si on désigne par $K_n[X]$ l'ensemble des polynômes à coefficients dans K de degré $\leq n$, $K_n[X]$ est un K -espace vectoriel dont une base est $(1, X, \dots, X^n)$ (cf. Chapitre 6).

5.4. Propriétés arithmétiques de $K[X]$

Dans tout ce paragraphe, on désigne par K un corps commutatif. Nous allons étendre à $K[X]$ un certain nombre de propriétés de l'anneau $\mathbb{Z}$.

5.4.1. DIVISION EUCLIDIENNE DANS $K[X]$

5.4.1.1. Définition

Soient A et B deux polynômes de $K[X]$. On dit que B divise A , et l'on note $B|A$ s'il existe un polynôme Q de $K[X]$ tel que $A = BQ$.

On dit aussi que A est un multiple de B , ou que B est un diviseur de A . Q s'appelle le quotient dans la division de A par B .

5.4.1.2. Remarques

a) Si A et B sont des polynômes non nuls et si B divise A , on a nécessairement $\deg(B) \leq \deg(A)$; si de plus $\deg(B) = \deg(A)$, alors la relation $A = BQ$ donne $\deg(B) = \deg(A) - \deg(B) = 0$, ce qui montre que Q est une constante non nulle de K . On dit que A et B sont **proportionnels** ou **associés**. Réciproquement, toute constante non nulle divise A .

b) La relation de divisibilité dans $K[X]$ est réflexive et transitive mais elle n'est ni symétrique, ni antisymétrique.

5.4.1.3. Théorème

Soient A et B deux polynômes de $K[X]$ tels que $B \neq 0$. Alors il existe un couple unique (Q, R) de polynômes de $K[X]$ tels que

$$(5.4.1.1) \quad A = BQ + R \quad \text{avec} \quad \deg(R) < \deg(B).$$

Q s'appelle le quotient et R le reste de la division euclidienne de A par B . A s'appelle le dividende, B le diviseur.

Démonstration. a) Existence d'une solution

Posons

$$A = a_n X^n + a_{n-1} X^{n-1} + \dots + a_0 \quad \text{avec} \quad a_n \neq 0$$

$$B = b_m X^m + b_{m-1} X^{m-1} + \dots + b_0 \quad \text{avec} \quad b_m \neq 0.$$

Si $n < m$, alors on peut écrire $A = 0 \cdot B + A$ avec $\deg(A) < \deg(B)$, donc $Q = 0$ et $R = A$ conviennent.

Supposons donc $n \geq m$. Posons

$$q_1 = \frac{a_n}{b_m} X^{n-m} \quad \text{et} \quad R_1 = A - Bq_1.$$

Comme Bq_1 admet $a_n X^n$ pour terme de plus haut degré, on a $\deg(R_1) < \deg(A)$.

Si $\deg(R_1) < \deg(B)$, on peut prendre $Q = q_1$ et $R = R_1$.

Si $\deg(R_1) \geq \deg(B)$, on répète sur le couple (R_1, B) ce qui vient d'être fait sur le couple (A, B) ; en désignant par q_2 le quotient du terme de plus haut degré de R_1 par $b_m X^m$, on effectue la différence

$$R_1 - Bq_2 = R_2 \quad \text{avec} \quad \deg(R_2) < \deg(R_1).$$

Si $\deg(R_2) < \deg(B)$, on peut écrire $A = Bq_1 + R_1 = Bq_1 + Bq_2 + R_2 = B(q_1 + q_2) + R_2$ avec $\deg(R_2) < \deg(B)$; donc on peut prendre comme solution $Q = q_1 + q_2$ et $R = R_2$.

Si $\deg(R_2) \geq \deg(B)$, on continue le processus avec R_2 et B . En répétant ce processus on forme une suite de polynômes R_k tels que

$$\deg(A) > \deg(R_1) > \dots > \deg(R_k) > \dots$$

La suite des degrés des R_k étant décroissante et majorée, est nécessairement finie; au bout d'un nombre fini d'itérations, on arrivera à une égalité de la forme.

$$R_{k-1} - Bq_k = R_k \quad \text{avec} \quad \deg(R_k) < \deg(B).$$

Si on ajoute membre à membre les k égalités

$$A - Bq_1 = R_1, \quad R_1 - Bq_2 = R_2, \dots, R_{k-1} - Bq_k = R_k,$$

on obtient :

$$A - B(q_1 + \dots + q_k) = R_k.$$

On peut donc prendre $Q = q_1 + \dots + q_k$ et $R = R_k$.

b) Unicité

Supposons qu'il existe un autre couple (Q_1, R_1) tel que $A = BQ_1 + R_1$ avec $\deg(R_1) < \deg(B)$. On en déduit $B(Q - Q_1) = R_1 - R$. Par conséquent,

$$(5.4.1.2) \quad \deg(R_1 - R) = \deg(B) + \deg(Q - Q_1).$$

Or puisque $\deg(R) < \deg(B)$ et $\deg(R_1) < \deg(B)$, on a

$$(5.4.1.3) \quad \deg(R_1 - R) \leq \max(\deg(R_1), \deg(R)) < \deg(B).$$

Les deux relations (5.4.1.2) et (5.4.1.3) sont incompatibles si $\deg(Q - Q_1) \geq 0$ c'est-à-dire si $Q - Q_1 \neq 0$. Par suite $Q = Q_1$ et $R = R_1$ ce qui établit l'unicité et achève la démonstration du théorème.

5.4.1.4. Corollaire

Si $B \neq 0$, pour que le polynôme A soit divisible par le polynôme B , il faut et suffit que le reste de la division euclidienne de A par B soit nul.

5.4.1.5. Exemple

Diviser $A = X^4 + X^2 - 4X + 2$ par $B = X^2 + 2X + 1$.

Voici comment il faut disposer les calculs.

A	$:$	X^4	$+X^2$	$-4X$	$+2$	$X^2 + 2X + 1$	$:$	B
$-Bq_1$	$:$	$-X^4$	$-2X^3$	$-X^2$		$X^2 - 2X + 4$	$:$	Q
R_1	$:$	$-2X^3$	$-4X$	$+2$		q_1	q_2	q_3
$-Bq_2$	$:$	$2X^3$	$+4X^2$	$+2X$				
R_2	$:$	$4X^2$	$-2X$	$+2$				
$-Bq_3$	$:$	$-4X^2$	$-8X$	-4				
R	$:$	$-10X$	-2					

5.4.2. IDÉAUX DE $K[X]$

5.4.2.1. Théorème

L'anneau $K[X]$ est principal.

Démonstration. Nous savons déjà que $K[X]$ est un anneau commutatif et intègre ; il suffit donc de montrer que tout idéal I peut être engendré par un élément.

Si $I = \{0\}$, alors I est engendré par 0 et le théorème est évident. Supposons $I \neq \{0\}$. L'ensemble des degrés des éléments non nuls de I est une partie non vide de $\mathbb{N}$. Cet ensemble admet donc un plus petit élément que nous noterons n_0 . Soit P un polynôme non nul de I tel que $\deg(P) = n_0$. Tout multiple de P appartient à I . Réciproquement, pour tout élément non nul A de I , on a $\deg(A) \geq \deg(P)$. La division euclidienne de A par P donne alors

$$A = PQ + R \quad \text{avec} \quad \deg(R) < \deg(P).$$

Des relations $A \in I$ et $P \in I$ on déduit $PQ \in I$ puis $R = A - PQ \in I$. Si $R \neq 0$, l'inégalité $\deg(R) < \deg(P)$ contredit le choix de P . Donc $R = 0$ et A est un multiple de P . L'idéal I est donc bien l'ensemble des multiples de P .

On notera (P) l'idéal engendré par le polynôme P .

5.4.2.2. Remarques

a) Le générateur P n'est pas unique car quel que soit $\lambda \in K^*$, $\lambda P \in I$ et a même degré que P .

b) Si A et B sont deux polynômes de $K[X]$ tels que $B \neq 0$, alors $B|A$ si et seulement si $(A) \subset (B)$.

Le Théorème 5.4.2.1 nous permet d'établir des propriétés qui, vraies dans tous les anneaux principaux, généralisent celles de $\mathbb{Z}$.

5.4.3. PLUS GRAND COMMUN DIVISEUR

5.4.3.1. Théorème

Soient $A_1, \dots, A_n$ des polynômes non nuls de $K[X]$. Alors :

a) Tout diviseur commun à $A_1, \dots, A_n$ divise le polynôme $A_1 P_1 + \dots + A_n P_n$ quels que soient les polynômes $P_1, \dots, P_n$.

b) Il existe un polynôme D dans $K[X]$ tel que :

1) D divise chaque polynôme A_i .

2) Il existe des polynômes $U_1, \dots, U_n$ de $K[X]$ tels que

$$(5.4.3.1) \quad D = A_1 U_1 + \dots + A_n U_n,$$

et l'ensemble des diviseurs communs aux polynômes A_i est égal à l'ensemble des diviseurs de D .

Démonstration. a) Soit Q un diviseur commun à $A_1, \dots, A_n$; il existe donc des polynômes $C_1, \dots, C_n$ tels que $A_i = C_i Q$, $1 \leq i \leq n$. Alors

$$A_1 P_1 + \dots + A_n P_n = C_1 Q P_1 + \dots + C_n Q P_n = Q(C_1 P_1 + \dots + C_n P_n),$$

ce qui prouve a).

b) Considérons l'ensemble I des polynômes de la forme $A_1 P_1 + \dots + A_n P_n$, où les P_i sont des polynômes quelconques de $K[X]$. I est évidemment un idéal de $K[X]$ contenant $A_1, \dots, A_n$. I étant un idéal principal d'après le Théorème 5.4.2.1, il existe un polynôme D de I de degré minimum tel que I soit l'ensemble des multiples de D . Comme $A_i \in I$ pour tout i , on a $A_i = D Q_i$ et D divise chaque polynôme A_i . D'autre part, puisque $D \in I$, il existe des polynômes $U_1, \dots, U_n$ tels que

$$D = A_1 U_1 + \dots + A_n U_n.$$

D'après a) tout polynôme qui divise chaque polynôme A_i divise D . Inversement tout diviseur de D divise tous les éléments de I (qui sont des multiples de D) et en particulier les polynômes $A_1, \dots, A_n$.

Remarquons que le polynôme D n'est défini qu'à une constante multiplicative non nulle près. En effet, si D' est un autre polynôme vérifiant la condition b), alors D divise D' et D' divise D , donc $D = \lambda D'$ avec $\lambda \in K^*$.

5.4.3.2. Définition

Avec les notations du Théorème 5.4.3.1, le polynôme D est appelé un **plus grand commun diviseur** (en abrégé P.G.C.D.) des polynômes $A_1, A_2, \dots, A_n$ et se note P.G.C.D. $(A_1, \dots, A_n)$.

La relation (5.4.3.1) s'appelle l'**identité de Bezout**.

D est un diviseur commun à $A_1, \dots, A_n$ de degré maximum, d'où son nom. Si on choisit D normalisé, on dit que D est le P.G.C.D. de $A_1, \dots, A_n$.

5.4.3.3. Définition

On dit que les polynômes $A_1, \dots, A_n$ sont **premiers entre eux dans leur ensemble** si P.G.C.D. $(A_1, \dots, A_n) = 1$.

Il revient au même de dire que l'idéal engendré par les A_i dans $K[X]$ est $K[X]$ tout entier.

D'après la remarque précédant la Définition 5.4.3.2, la relation P.G.C.D. $(A_1, \dots, A_n) = 1$ équivaut à P.G.C.D. $(A_1, \dots, A_n) = \lambda$, où λ est un élément quelconque de K^* .

On dit que les polynômes $A_1, \dots, A_n$ sont **premiers entre eux deux à deux** si P.G.C.D. $(A_i, A_j) = 1$ pour tous les indices i et j tels que $i \neq j$.

Il est évident que si $n \geq 3$ et si les polynômes $A_1, \dots, A_n$ sont premiers entre eux deux à deux ils sont premiers entre eux dans leur ensemble. Mais si $A_1, \dots, A_n$ sont premiers entre eux dans leur ensemble, ils ne sont pas nécessairement premiers entre eux deux à deux. En effet, dans $\mathbb{Q}[X]$, les trois polynômes $A = X(X+1)$, $B = X(X+2)$, $C = (X+1)(X+2)$ sont premiers entre eux dans leur ensemble mais ils ne sont pas premiers entre eux deux à deux puisque

$$\text{P.G.C.D.}(A, B) = X, \quad \text{P.G.C.D.}(B, C) = X+2, \quad \text{P.G.C.D.}(A, C) = X+1.$$

Recherche pratique du P.G.C.D. : Algorithme d'Euclide

Soient A et B deux polynômes de $K[X]$ et supposons que $\deg(B) \leq \deg(A)$. La division euclidienne de A par B donne $A = BQ_1 + R_1$ avec $\deg(R_1) < \deg(B)$.

Tout polynôme qui divise A et B divise B et $R_1 = A - BQ_1$. Réciproquement, tout polynôme qui divise B et R_1 divise $A = BQ_1 + R_1$ et B . Le calcul d'un P.G.C.D. de A et B est donc équivalent au calcul d'un P.G.C.D. de B et R_1 .

1^{er} cas : Si $R_1 = 0$, alors B est un P.G.C.D. de A et B .

2^e cas : Si $R_1 \neq 0$, soit R_2 le reste de la division euclidienne de B par R_1 . Si $R_2 = 0$, B est multiple de R_1 et R_1 est donc un P.G.C.D. de A et B . Si $R_2 \neq 0$ un P.G.C.D. de R_1 et R_2 est un P.G.C.D. de B et R_1 donc un P.G.C.D. de A et B , et ainsi de suite.

En effectuant la suite des divisions euclidiennes

$$A = BQ_1 + R_1 \quad , \quad \deg(R_1) < \deg(B)$$

$$B = R_1Q_2 + R_2 \quad , \quad \deg(R_2) < \deg(R_1)$$

$$R_1 = R_2Q_3 + R_3 \quad , \quad \deg(R_3) < \deg(R_2)$$

$$R_{j-1} = R_jQ_{j+1} + R_{j+1} \quad , \quad \deg(R_{j+1}) < \deg(R_j)$$

on obtient nécessairement un reste $R_{h+1} = 0$, puisque les degrés des restes sont strictement décroissants. Le dernier reste précédant le reste nul est donc un P.G.C.D. de A et B .

Dans le cas de n polynômes $A_1, \dots, A_n$, les communs diviseurs de A_1 et A_2 étant les diviseurs de leur plus grand commun diviseur D_2 nous avons

$$\text{P.G.C.D.}(A_1, \dots, A_n) = \text{P.G.C.D.}(D_2, A_3, \dots, A_n).$$

Les diviseurs communs de D_2 et A_3 étant les diviseurs de leur plus grand commun diviseur D_3 , on a

$$\text{P.G.C.D.}(A_1, \dots, A_n) = \text{P.G.C.D.}(D_3, A_4, \dots, A_n).$$

De proche en proche, on voit que l'on peut ramener le calcul du P.G.C.D. de n polynômes au calcul du P.G.C.D. de deux polynômes.

Les propriétés démontrées au Chapitre 4 pour le P.G.C.D. dans un anneau principal restent valables dans l'anneau $K[X]$ et se démontrent de la même manière. Énonçons quelques propriétés du P.G.C.D. à titre d'exemples.

5.4.3.4. Théorème

Soient $A_1, \dots, A_n$ des éléments de $K[X]$. Si on multiplie (resp. divise, si cela est possible) $A_1, \dots, A_n$ par un même polynôme $P \neq 0$, le P.G.C.D. de $A_1, \dots, A_n$ est multiplié (resp. divisé) par P .

Démonstration. Démontrons la propriété relative à la multiplication par P ; l'autre propriété se démontre de la même manière.

Posons $D = \text{P.G.C.D.}(A_1, \dots, A_n)$. Il existe des polynômes $U_1, \dots, U_n$ tels que

$$D = A_1U_1 + \dots + A_nU_n.$$

Donc si un polynôme Q divise chacun des polynômes $A_1P, \dots, A_nP$, il divise $A_1PU_1 + \dots + A_nPU_n = (A_1U_1 + \dots + A_nU_n)P = DP$. Réciproquement, DP divise chacun des polynômes $A_1P, \dots, A_nP$; donc tout diviseur de DP est un diviseur de $A_1P, \dots, A_nP$. On en déduit

$$\text{P.G.C.D.}(A_1P, \dots, A_nP) = DP.$$

5.4.3.5. Théorème (Bezout)

Pour que les polynômes $A_1, A_2, \dots, A_n$ de $K[X]$ soient premiers entre eux dans leur ensemble, il faut et il suffit qu'il existe des polynômes $U_1, \dots, U_n$ de $K[X]$ tels que

$$A_1 U_1 + \dots + A_n U_n = 1.$$

La démonstration est la même que celle du Théorème 4.4.2.4.

5.4.3.6. Corollaire

Soient A, B et C des polynômes de $K[X]$. Si A est premier séparément avec B et C , alors A est premier avec le produit BC .

Du Théorème de Bezout, on déduit également le théorème suivant connu sous le nom de Théorème de Gauss.

5.4.3.7. Théorème (Gauss)

Soient A, B et C des polynômes de $K[X]$. Si A divise le produit BC et si A est premier avec B , alors A divise C .

Même démonstration que celle du Théorème 4.4.2.6.

5.4.3.8. Corollaire

Soient $B_1, \dots, B_n$ des polynômes de $K[X]$ premiers entre eux deux à deux. Si un polynôme A de $K[X]$ est divisible par chaque polynôme B_i , alors A est divisible par le produit $B_1 \dots B_n$.

5.4.4. PLUS PETIT COMMUN MULTIPLE

Soient $A_1, A_2, \dots, A_n$ des éléments non nuls de $K[X]$. Les multiples communs à $A_1, \dots, A_n$ dans $K[X]$ sont les éléments de l'idéal $(A_1) \cap (A_2) \cap \dots \cap (A_n)$. Cet idéal étant principal, il existe un élément M de $K[X]$, unique à la multiplication près par une constante non nulle, tel que

$$(A_1) \cap \dots \cap (A_n) = (M)$$

et tout multiple commun à $A_1, A_2, \dots, A_n$ est un multiple de M ; réciproquement tout multiple de M est un multiple commun à $A_1, \dots, A_n$.

5.4.4.1. Définition

On appelle **plus petit commun multiple (P.P.C.M.)** des polynômes $A_1, \dots, A_n$, et on note P.P.C.M. $(A_1, \dots, A_n)$, tout polynôme M de $K[X]$ tel que

$$(A_1) \cap (A_2) \cap \dots \cap (A_n) = (M).$$

Notons que M est un multiple commun à $A_1, \dots, A_n$ de degré minimum, d'où son nom.

On a le résultat suivant qui ramène le calcul du P.P.C.M. de deux polynômes non nuls à celui de leur P.G.C.D. La démonstration est exactement la même que celle du Théorème 4.4.3.2.

5.4.4.2. Théorème

Soient A et B deux polynômes non nuls de $K[X]$, D un P.G.C.D., M un P.P.C.M. de A et B . Alors $AB = MD$ (à une constante multiplicative non nulle près).

5.4.5. POLYNÔMES IRRÉDUCTIBLES

5.4.5.1. Définition

*On dit qu'un polynôme P de $K[X]$ est **premier** ou **irréductible** sur le corps K s'il n'est pas constant et si ses seuls diviseurs dans $K[X]$ sont les polynômes associés à P et les éléments non nuls de K .*

Il faut bien noter que cette définition dépend essentiellement du corps K : il se peut très bien qu'un polynôme donné soit irréductible sur un corps et réductible sur un autre. Par exemple, $X^2 + 1$ est irréductible sur $\mathbb{R}$ mais est réductible sur $\mathbb{C}$.

5.4.5.2. Remarque

Dire qu'un polynôme de $K[X]$ est irréductible revient à dire qu'il est impossible de l'écrire comme produit de deux polynômes de $K[X]$ de degrés positifs.

5.4.5.3. Exemple

Tout polynôme P de $K[X]$, du premier degré, est irréductible. Si en effet $P = AB$, avec $A \in K[X]$ et $B \in K[X]$, on a $1 = \deg(P) = \deg(A) + \deg(B)$ donc, nécessairement, l'un des polynômes A ou B est de degré 0 et l'autre de degré un.

5.4.5.4. Exemple

Le polynôme $X^2 - 2$ est irréductible sur $\mathbb{Q}[X]$. Si en effet $X^2 - 2$ n'était pas irréductible, on pourrait écrire

$$X^2 - 2 = (aX + b)(cX + d) = acX^2 + (ad + bc)X + bd$$

avec $a, b, c, d \in \mathbb{Q}$ et $a \neq 0, c \neq 0$. D'où

$$ac = 1, \quad ad + bc = 0, \quad bd = -2.$$

Il résulte de là que

$$c = \frac{1}{a}, \quad d = -\frac{2}{b}.$$

et en substituant dans $ad + bc = 0$, il vient $b^2 - 2a^2 = 0$, d'où $\left(\frac{b}{a}\right)^2 = 2$, ce qui est impossible puisque $\sqrt{2}$ n'est pas un nombre rationnel.

$$X^2 - 2 \text{ est réductible sur } \mathbb{R} \text{ car } X^2 - 2 = (X - \sqrt{2})(X + \sqrt{2}).$$

5.4.5.5. Théorème

Soient P un polynôme irréductible de $K[X]$ et Q un polynôme de $K[X]$. Alors P et Q sont premiers entre eux si et seulement si P ne divise pas Q .

Démonstration. Si P et Q sont premiers entre eux, les seuls diviseurs communs à P et Q sont les constantes, donc P ne divise pas Q .

Réciproquement, si P ne divise pas Q , un P.G.C.D. de P et Q étant un diviseur de P ne peut être qu'une constante non nulle ou un polynôme associé à P . Le second cas étant exclu par hypothèse, la conclusion en résulte. $\square$

5.4.5.6. Corollaire

Si un polynôme irréductible divise un produit de polynômes, il divise au moins l'un des facteurs.

Démonstration. Soit P un polynôme irréductible et supposons que P divise le produit AB sans diviser A . D'après le théorème précédent, P est premier avec A , donc (Théorème 5.4.3.7), P divise B .

Par récurrence, on démontrerait l'assertion lorsqu'on a un nombre fini quelconque de facteurs.

5.4.5.7. Théorème

Tout polynôme P non nul de $K[X]$ peut s'écrire d'une manière unique sous la forme

$$(5.4.5.1) \quad P = \lambda P_1 \dots P_n$$

où λ est constante non nulle et où $P_1, \dots, P_n$ sont des polynômes irréductibles normalisés.

Démonstration. Démontrons l'existence d'une factorisation de P par récurrence sur le degré de P .

Si $\deg(P) = 1$, le résultat est trivial car d'après l'Exemple 5.4.5.3, si $P = aX + b = a(X + a^{-1}b)$, $X + a^{-1}b$ est un polynôme irréductible normalisé.

Supposons le théorème établi pour les polynômes de degré inférieur à n . Soit P un polynôme de degré n . Alors, ou bien P est irréductible et le théorème est évident, ou bien P est décomposable en un produit de deux polynômes de degré inférieur à n : $P = Q_1 Q_2$ avec $\deg(Q_1) < n$, $\deg(Q_2) < n$.

En appliquant l'hypothèse de récurrence à Q_1 et à Q_2 on obtient une décomposition de P sous la forme (5.4.5.1) L'existence est établie.

Démontrons l'unicité de la décomposition. Supposons que

$$(5.4.5.2) \quad P = \lambda P_1 \dots P_n = \mu Q_1 \dots Q_m.$$

où λ et μ sont des constantes, et où les facteurs P_i et Q_j sont des polynômes irréductibles et normalisés. On a nécessairement $\lambda = \mu$ car chacune de ces constantes est le coefficient dominant de P . D'autre part, le polynôme irréductible P_1 divise le premier membre de l'égalité (5.4.5.2). Il divise donc le second membre et, d'après le Corollaire 5.4.5.8, il en divise l'un des facteurs. Appelons Q_1 ce facteur en changeant au besoin la numérotation des facteurs du second membre. Comme P_1 et Q_1 sont irréductibles et normalisés, on a $P_1 = Q_1$. Puisque $K[X]$ est intègre, on peut simplifier par P_1 . En recommençant ce raisonnement avec P_2 puis avec P_3 et ainsi de suite, on pourra identifier les n facteurs du premier produit avec autant de facteurs du second. On en déduit $n \leq m$. On verrait de même que $m \leq n$, d'où $n = m$. Les deux décompositions sont donc identiques à l'ordre près.

5.4.5.8. Remarque

Dans la décomposition de P , il est possible que certains polynômes premiers soient égaux. En regroupant les polynômes normalisés, on peut écrire

$$(5.4.5.3) \quad P = k P_1^{\alpha_1} P_2^{\alpha_2} \dots P_r^{\alpha_r}$$

$P_1, \dots, P_r$ désignant des polynômes irréductibles normalisés distincts et k une constante.

On dit que (5.4.5.3) est la décomposition de P en facteurs irréductibles.

5.5. Division suivant les puissances croissantes

K désigne toujours un corps commutatif. Nous présentons dans ce paragraphe la division suivant les puissances croissantes. Il s'agit, étant donnés deux polynômes A et B , de trouver un polynôme Q tel que $\text{val}(A - BQ)$ soit supérieure à un entier naturel fixé à l'avance.

5.5.0.1. Théorème

Soient n un entier naturel, A un polynôme quelconque et B un polynôme tel que $\text{val}(B) = 0$. Il existe un couple unique (Q, R) de polynômes de $K[X]$ tels que

$$(5.5.0.1) \quad A = BQ + X^{n+1}R \quad \text{avec} \quad \deg(Q) \leq n.$$

Démonstration. a) Existence d'une solution

Ordonnons les polynômes A et B suivant les puissances croissantes de X :

$$A = a_r X^r + a_{r+1} X^{r+1} + \dots + a_m X^m \quad \text{avec} \quad a_r \neq 0 \quad \text{et} \quad a_m \neq 0$$

$$B = b_0 + b_1 X + \dots + b_p X^p \quad \text{avec} \quad b_0 \neq 0 \quad \text{et} \quad b_p \neq 0.$$

Si $A = 0$ ou si $r = \text{val}(A) > n$, on peut écrire $A = X^{n+1}R$, d'où la solution évidente $Q = 0$ et R .

Supposons $A \neq 0$ et $r = \text{val}(A) \leq n$. Posons $q_1 = \frac{a_r}{b_0} X^r$ et

$$(5.5.0.2) \quad R_1 = A - Bq_1.$$

Comme A et Bq_1 ont le même terme de plus bas degré, on a soit $R_1 = 0$ soit $\text{val}(R_1) > \text{val}(A)$.

Si $R_1 = 0$ ou si $\text{val}(R_1) > n$, nous pouvons écrire $R_1 = X^{n+1}R$ et nous pouvons prendre comme solution $Q = q_1$ et R .

Si $\text{val}(R_1) \leq n$, répétons sur le couple (R_1, B) ce qui vient d'être fait sur le couple (A, B) : on calcule le quotient q_2 du premier terme de R_1 par b_0 et on forme la différence

$$(5.5.0.3) \quad R_2 = R_1 - Bq_2.$$

Remarquons que $\deg(q_1) < \deg(q_2) = \text{val}(R_1) \leq n$.

Les premiers termes de R_1 et Bq_2 étant égaux, nous avons soit $R_2 = 0$, soit $\text{val}(R_2) > \text{val}(R_1)$. Si $R_2 = 0$ ou si $\text{val}(R_2) > n$ on peut écrire $R_2 = X^{n+1}R$.

Si nous ajoutons membre à membre les égalités (5.5.0.2) et (5.5.0.3) nous obtenons $A - B(q_1 + q_2) = R_2$, d'où la solution $Q = q_1 + q_2$ et R .

Si $\text{val}(R_2) \leq n$, on recommence l'opération en remplaçant R_1 par R_2 , et ainsi de suite. Comme $\text{val}(A) < \text{val}(R_1) < \text{val}(R_2) < \dots$, nous obtiendrons au bout d'un nombre fini d'itérations, un polynôme R_k tel que $R_k = 0$ ou $\text{val}(R_k) > n$, donc tel que $R_k = X^{n+1}R$.

En ajoutant membre à membre les égalités

$$A - Bq_1 = R_1, \quad R_1 - Bq_2 = R_2, \dots, R_{k-1} - Bq_k = R_k,$$

nous obtenons $A - B(q_1 + \dots + q_k) = R_k$.

Comme $\deg(q_1) < \deg(q_2) < \dots < \deg(q_k) = \text{val}(R_{k-1}) \leq n$, il suffit de prendre $Q = q_1 + q_2 + \dots + q_k$ et R tel que $R_k = X^{n+1}R$, R_k étant le premier reste dont la valuation est strictement supérieure à n .

b) Unicité

Supposons que l'on ait deux décompositions

$$A = BQ + X^{n+1}R = BQ_1 + X^{n+1}R_1 \quad \text{avec} \quad \deg(Q) \leq n \quad \text{et} \quad \deg(Q_1) \leq n.$$

Par différence, nous obtenons $B(Q - Q_1) = X^{n+1}(R_1 - R)$.

Si on avait $Q - Q_1 \neq 0$, le terme de plus bas degré de $B(Q - Q_1)$ serait le degré n au plus puisque $b_0 \neq 0$ et $\deg(Q - Q_1) \leq n$. Comme le terme de plus bas degré de $X^{n+1}(R_1 - R)$ est de degré $(n + 1)$ au moins, on aboutit à une contradiction. Donc $Q = Q_1$ et alors $R = R_1$.

5.5.0.2. Définition

Pour un entier n donné, l'écriture 5.5.0.1 s'appelle la **division suivant les puissances croissantes de A par B l'ordre n** . Dans cette division Q est le quotient à l'ordre n et $X^{n+1}R$ le reste à l'ordre n .

5.5.0.3. Exemple

Diviser $A = 1 + X$ par $B = 1 - X + X^2$ suivant les puissances croissantes de X à l'ordre 2.

La disposition pratique est la même que celle de la division euclidienne mais on écrit les monômes de A et B par ordre de degré croissant, ce qui explique le nom donné à cette division.

$1 + X$	$1 - X + X^2$
$-Bq_1 : -1 + X - X^2$	$1 + 2X + X^2$
$R_1 : 2X - X^2$	$q_1 \quad q_2 \quad q_3$
$-Bq_2 : -2X + 2X^2 - 2X^3$	
$R_2 : X^2 - 2X^3$	
$-Bq_3 : -X^2 + X^3 - X^4$	
$R_3 : -X^3 - X^4$	

$$\text{Donc } 1 + X = (1 - X + X^2)(1 + 2X + X^2) - X^3(1 + X).$$

5.6. Fonction polynômes. Racines d'un polynôme

5.6.1. FONCTIONS POLYNÔMES

5.6.1.1. Définition

Soient K un anneau commutatif et $P = a_0 + a_1 X + \dots + a_n X^n$ un polynôme de $K[X]$. On appelle **fonction polynôme** (ou **fonction polynomiale**) associée au polynôme P , l'application $\tilde{P}$ de K dans K , associant à tout x de K l'élément

$$\tilde{P}(x) = a_0 + a_1 x + \dots + a_n x^n.$$

5.6.1.2. Remarque

Il est essentiel de ne pas confondre le polynôme $P = (a_0, a_1, \dots, a_n, 0, \dots)$ qui est une suite d'éléments de K dont un nombre fini sont non nuls et qui s'écrit, si $X = (0, 1, 0, \dots)$, $P = a_0 + a_1 X + \dots + a_n X^n$, et la fonction polynôme $\tilde{P}$. La fonction polynôme associée au polynôme P se note souvent à l'aide du même symbole P mais cette notation est dangereuse à cause des confusions possibles.

5.6.1.3. Théorème

K^K désignant l'ensemble des applications de K dans K , l'application φ de $K[X]$ dans K^K qui associe à un polynôme P la fonction polynôme $\tilde{P}$ est un morphisme d'anneaux.

Démonstration. On sait (voir l'Exemple 4.1.1.4) que K^K est un anneau. Il faut démontrer que, quels que soient les polynômes P et Q de $K[X]$, on a $\varphi(P + Q) = \varphi(P) + \varphi(Q)$, $\varphi(PQ) = \varphi(P) \varphi(Q)$ et $\varphi(1) = 1$.

Soient $P = (a_0, a_1, \dots, a_n)$, $Q = (b_0, b_1, \dots, b_n, \dots)$.

On a

$$P + Q = (a_0 + b_0, a_1 + b_1, \dots).$$

$\varphi(P + Q)$ est l'application de K^K définie, pour tout $x \in K$, par

$$(\varphi(P + Q))(x) = a_0 + b_0 + (a_1 + b_1)x + \dots + (a_I + b_i)x^i + \dots$$

ce qui, dans l'anneau K , est égal à

$$\begin{aligned} (a_0 + a_1 x + \dots + a_i x^i + \dots) + (b_0 + b_1 x + \dots + b_i x^i + \dots) \\ = (\varphi(P))(x) + (\varphi(Q))(x). \end{aligned}$$

On vérifierait de même que quel que soit $x \in K$,

$$(\varphi(PQ))(x) = (\varphi(P))(x) \cdot (\varphi(Q))(x).$$

Les images de tout $x \in K$ par $\varphi(P + Q)$ et par $\varphi(P) + \varphi(Q)$ étant égales, ces deux applications de K^K sont égales. On a de même $\varphi(PQ) = \varphi(P)\varphi(Q)$.

Enfin, si $1 = (1, 0, 0, \dots)$ est l'élément unité de $K[X]$ on a, pour tout $x \in K$, $(\varphi(1))(x) = 1$, d'où $\varphi(1) = 1$.

5.6.1.4. Remarque

Il faut noter que le morphisme φ défini dans le Théorème 5.6.1.3 n'est pas nécessairement injectif. Ainsi, si $K = \mathbb{Z}/3\mathbb{Z}$, le polynôme $P = X^3 - X$ a pour fonction polynôme associée la fonction nulle (c'est le théorème de Fermat) mais P n'est pas le polynôme nul. Le morphisme φ n'est pas surjectif en général car si $K = \mathbb{R}$, il existe des applications de $\mathbb{R}$ dans $\mathbb{R}$ (par exemple les fonctions exponentielles) qui ne sont pas polynomiales.

5.6.2. RACINES D'UN POLYNÔME

Désormais et jusqu'à la fin de ce paragraphe, on suppose que K est un corps commutatif.

5.6.2.1. Définition

Soient P un polynôme de $K[X]$ et a un élément de K . On dit que a est une racine ou un zéro de P si $\tilde{P}(a) = 0$.

5.6.2.2. Théorème

Soient $P \in K[X]$ et $a \in K$. Pour que a soit racine de P , il faut et il suffit que P soit divisible par $X - a$.

Démonstration. Si P est divisible par $X - a$, alors il existe $Q \in K[X]$ tel que $P = (X - a)Q$; donc d'après le Théorème 5.6.1.3, pour tout $x \in K$, on a $\tilde{P}(x) = (x - a)\tilde{Q}(x)$. Si en particulier $x = a$, on a $\tilde{P}(a) = 0$. $\tilde{Q}(a) = 0$, donc a est racine de P .

Réciproquement, supposons que $\tilde{P}(a) = 0$. Effectuons la division euclidienne de P par $X - a$. On obtient $P = (X - a)Q + R$ avec $\deg(R) < \deg(X - a) = 1$, donc R est un polynôme constant. En prenant les valeurs des fonctions polynômes associées au point $x = a$, il vient $0 = \tilde{P}(a) = 0 \cdot \tilde{Q}(a) + R$. Donc P est divisible par $X - a$.

5.6.2.3. Définition

Soient P un polynôme de $K[X]$, a un élément de K et α un entier ≥ 1 . On dit que a est racine d'ordre α (ou de multiplicité α) de P si P est divisible par $(X - a)^\alpha$ sans l'être par $(X - a)^{\alpha+1}$.

On dit que l'entier α est la **multiplicité** ou l'**ordre de multiplicité** de la racine a .

Une racine d'ordre 1 est dite **racine simple**, une racine d'ordre 2 est dite **racine double**, etc.

5.6.2.4. Théorème

Soit P un élément de $K[X]$. Soient $a_1, a_2, \dots, a_r$ les r racines distinctes de P dans K , et $\alpha_1, \alpha_2, \dots, \alpha_r$ leurs ordres de multiplicité respectifs. Il existe un polynôme $Q \in K[X]$ qui n'admet pas de racine dans K tel que

$$P = (X - a_1)^{\alpha_1} (X - a_2)^{\alpha_2} \dots (X - a_r)^{\alpha_r} Q.$$

Démonstration. Pour $1 \leq i \leq r$ et $1 \leq j \leq r$, avec $i \neq j$, les deux polynômes $(X - a_i)$ et $(X - a_j)$ sont premiers entre eux puisque $a_i \neq a_j$. Donc $(X - a_i)^{\alpha_i}$ et $(X - a_j)^{\alpha_j}$ sont premiers entre eux ; ainsi les polynômes $(X - a_i)^{\alpha_i}$ pour $1 \leq i \leq r$ sont premiers entre eux deux à deux et divisent P . Il en résulte que P est divisible par le produit $(X - a_1)^{\alpha_1} \dots (X - a_r)^{\alpha_r}$; il existe donc un polynôme Q tel que $P = (X - a_1)^{\alpha_1} \dots (X - a_r)^{\alpha_r} Q$. D'autre part, si Q admettait la racine a_i , Q serait divisible par $(X - a_i)$ et P serait divisible par $(X - a_i)^{\alpha_i+1}$, ce qui est absurde.

5.6.2.5. Remarque

Le théorème ne dit pas que le polynôme Q n'admet pas d'autres racines ; il peut en admettre dans un sur-corps de K . Par exemple, si $P = X^4 - 2X^3 - X^2 + 4X - 2$ est considéré comme polynôme de $\mathbb{Q}[X]$, on a $P = (X - 1)^2 (X^2 - 2)$, donc $Q = X^2 - 2$ et ce n'admet pas de racines dans $\mathbb{Q}$, mais admet $\sqrt{2}$ et $-\sqrt{2}$ comme racines dans $\mathbb{R}$.

5.6.2.6. Corollaire

Tout polynôme $P \in K[X]$ non nul de degré n admet au plus n racines dans K en convenant de compter α fois une racine multiple d'ordre α . Si P admet n racines distinctes, elles sont toutes simples.

Démonstration. Soient $a_1, \dots, a_r$ les racines distinctes de P , $\alpha_1, \dots, \alpha_r$ leurs ordres de multiplicité. On a

$$P = (X - a_1)^{\alpha_1} \dots (X - a_r)^{\alpha_r} Q;$$

comme $P \neq 0$, on a $Q \neq 0$, donc

$$\deg(P) = \alpha_1 + \dots + \alpha_r + \deg(Q) \quad \text{avec} \quad \deg(Q) \geq 0.$$

On en déduit $\alpha_1 + \dots + \alpha_r \leq \deg(P) = n$.

Si P admet n racines distinctes, on a d'après ce qui précède, $\alpha_1 + \dots + \alpha_n \leq n$; or comme $\alpha_i \geq 1$ pour tout $i \in \{1, \dots, n\}$, on a $\alpha_1 + \dots + \alpha_n \geq n$. Donc $\alpha_1 + \dots + \alpha_n = n$ et par suite $\alpha_i = 1$ pour tout $i \in \{1, \dots, n\}$.

5.6.2.7. Corollaire

Soit K un corps infini. Alors l'application $P \mapsto \tilde{P}$ de $K[X]$ dans K^K est injective.

Démonstration. Soient P et Q deux polynômes de $K[X]$ tels que $\tilde{P} = \tilde{Q}$. Si $P - Q \neq 0$, soit n le degré de $P - Q$. Considérons $n + 1$ éléments distincts de K , ce qui est possible puisque K est infini. La fonction $\tilde{P} - \tilde{Q}$ s'annule pour ces valeurs; donc le polynôme $P - Q$ admet $(n + 1)$ racines, ce qui est impossible d'après le Corollaire 5.6.2.6. Donc $P = Q$. $\square$

Lorsque K est infini (par exemple, si $K = \mathbb{Q}, \mathbb{R}$, ou $\mathbb{C}$) on peut identifier sans inconvénient un polynôme P à la fonction polynôme $\tilde{P}$ qui lui est associée; c'est ce que nous ferons désormais lorsque K est $\mathbb{Q}, \mathbb{R}$ ou $\mathbb{C}$. La valeur $\tilde{P}(a)$ de P en un point a de K sera donc notée $P(a)$.

5.6.2.8. Définition

On dit qu'un polynôme P de $K[X]$ est scindé sur K si $P = 0$, ou, dans le cas contraire, si P est décomposable en un produit de facteurs du premier degré (distincts ou non) de $K[X]$.

5.7. Étude de $\mathbb{C}[X]$ et de $\mathbb{R}[X]$

5.7.1. CORPS ALGÈBRIQUEMENT CLOS

5.7.1.1. Définition

On dit qu'un corps commutatif K est algébriquement clos si tout polynôme non constant de $K[X]$ possède au moins une racine dans K .

Un tel corps est nécessairement infini car si K est fini et possède n éléments $a_1, \dots, a_n$, le polynôme $(X - a_1)(X - a_2) \dots (X - a_n) + 1$ n'a pas de racine dans K .

5.7.1.2. Théorème

Soit K un corps commutatif. Les propriétés suivantes sont équivalentes :

- a) K est algébriquement clos.
- b) Tout polynôme P non nul de $K[X]$ de degré $n > 0$ admet n racines dans K .

c) Les seuls polynômes irréductibles de $K[X]$ sont les polynômes de degré 1.

Démonstration. a) $\implies$ b) Supposons K algébriquement clos. Soit P un polynôme de $K[X]$ de degré $n \geq 1$. Alors P possède au moins une racine a_1 dans K ; il existe donc un polynôme $Q_1 \in K[X]$ tel que $P = (X - a_1)Q_1$ avec $\deg(Q_1) = \deg(P) - 1 = n - 1$. Par récurrence, on voit que

$$P = (X - a_1)(X - a_2) \dots (X - a_n) Q_n$$

avec $\deg(Q_n) = 0$, i.e. $Q_n \in K$, ce qui montre que P admet n racines dans K et prouve b).

Il est évident que b) $\implies$ a), donc a) $\iff$ b).

b) $\implies$ c) Supposons la condition b) vérifiée. Soit P un polynôme irréductible de $K[X]$. Alors P est de degré $n \geq 1$ puisqu'il est non inversible dans $K[X]$; donc P admet n racines dans K . Par suite, il existe au moins un élément $a \in K$ tel que P soit multiple de $X - a$. Comme P est irréductible il est proportionnel à $X - a$, donc P est de degré 1.

c) $\implies$ b) Supposons que tout polynôme irréductible de $K[X]$ est de degré 1. Soit P un polynôme non nul de $K[X]$. D'après le théorème 5.4.4.9, P admet une décomposition unique en facteurs irréductibles :

$$P = \lambda P_1^{\alpha_1} \dots P_r^{\alpha_r}$$

$P_1, \dots, P_r$ désignant des polynômes irréductibles normalisés distincts et λ une constante non nulle. Par hypothèse, les P_i sont des polynômes de degré 1; donc P s'écrit

$$P = \lambda (X - a_1)^{\alpha_1} \dots (X - a_r)^{\alpha_r}$$

les a_i étant deux à deux distincts. Donc P possède n racines dans K (en comptant chaque racine avec son ordre de multiplicité) et b) est démontré.

On en déduit immédiatement les corollaires suivants :

5.7.1.3. Corollaire

Soit K un corps algébriquement clos. Alors tout polynôme de $K[X]$ est scindé sur K .

5.7.1.4. Corollaire

Soient K un corps algébriquement clos, A et B deux polynômes de $K[X]$. alors B divise A si et seulement si toute racine de B , d'ordre α , est racine de A avec un ordre au moins égal à α .

Les corps $\mathbb{Q}$ et $\mathbb{R}$ ne sont pas algébriquement clos car $X^2 - 2$ n'a pas de racine dans $\mathbb{Q}$ et $X^2 + 1$ n'a pas de racine dans $\mathbb{R}$.

Le théorème suivant, dont la démonstration dépasse de loin le niveau de ce cours, fournit un exemple important de corps algébriquement clos.

5.7.1.5. Théorème (d'Alembert-Gauss)

Le corps $\mathbb{C}$ des nombres complexes est algébriquement clos.

$\mathbb{Q}$ et $\mathbb{R}$ sont des sous-corps du corps algébriquement clos $\mathbb{C}$. On démontre, plus généralement (Théorème de Steinitz) que tout corps commutatif K peut être plongé dans un corps commutatif L , algébriquement clos, unique à un isomorphisme près. L est appelé *clôture algébrique* de K .

5.7.2. POLYNÔMES DE $\mathbb{C}[X]$

Le corps $\mathbb{C}$ étant algébriquement clos, tous les résultats du n° précédent s'appliquent. En particulier le Théorème 5.6.2.4 peut être amélioré si K est un corps algébriquement clos, donc si $K = \mathbb{C}$.

5.7.2.1. Théorème

Soit P un élément de $\mathbb{C}[X]$ de degré n . Soient $a_1, \dots, a_r$ les racines distinctes de P , et $\alpha_1, \dots, \alpha_r$ leurs ordres de multiplicité. On a alors

$$(5.7.2.1) \quad P(X) = \lambda(X - a_1)^{\alpha_1} \dots (X - a_r)^{\alpha_r}$$

où λ est le coefficient dominant de P . En outre

$$(5.7.2.2) \quad \alpha_1 + \alpha_2 + \dots + \alpha_r = n.$$

Démonstration. D'après le Théorème 5.6.2.4, on a

$$P(X) = (X - a_1)^{\alpha_1} \dots (X - a_r)^{\alpha_r} Q(X)$$

où le polynôme Q n'admet plus $a_1, \dots, a_r$ pour racines ; comme toute racine de Q est racine de P (Corollaire 5.7.1.4), Q est constant. Le monôme de plus haut degré dans $\lambda(X - a_1)^{\alpha_1} \dots (X - a_r)^{\alpha_r}$ est $\lambda X^{\alpha_1 + \dots + \alpha_r}$; comme P est de degré n , on a bien $\alpha_1 + \dots + \alpha_r = n$. $\square$

Le corps $\mathbb{C}$ étant algébriquement clos, les seuls polynômes irréductibles de $\mathbb{C}[X]$ sont ceux de degré 1. Donc la formule (5.7.2.1) est la décomposition de P en facteurs irréductibles normalisés.

5.7.2.2. Définition

*Soit $P = a_0 + a_1 X + \dots + a_n X^n$ un élément de $\mathbb{C}[X]$. On appelle **polynôme conjugué de P** , et on note $\overline{P}$, le polynôme $\overline{P} = \overline{a_0} + \overline{a_1} X + \dots + \overline{a_n} X^n$.*

On vérifie facilement que si P et Q sont des éléments de $\mathbb{C}[X]$ et si $\lambda \in \mathbb{C}$, on a

$$\overline{P + Q} = \overline{P} + \overline{Q}, \quad \overline{PQ} = \overline{P} \overline{Q}, \quad \overline{\lambda P} = \overline{\lambda} \overline{P}, \quad \overline{\overline{P}} = P.$$

Autrement dit, l'application $P \mapsto \overline{P}$ de $\mathbb{C}[X]$ dans $\mathbb{C}[X]$ (évidemment bijective) est un automorphisme de l'anneau $\mathbb{C}[X]$.

On en déduit que le polynôme B divise le polynôme A dans $\mathbb{C}[X]$ si, et seulement si $\overline{B}$ divise $\overline{A}$. En effet, l'égalité $A = BQ$ est équivalente à $\overline{A} = \overline{B} \overline{Q}$.

5.7.2.3. Théorème

Soient P un élément de $\mathbb{C}[X]$ et a un élément de $\mathbb{C}$. Alors :

a) a est racine d'ordre α de P si et seulement si $\bar{a}$ (conjugué de a) est racine d'ordre α de $\overline{P}$.

b) Si $P \in \mathbb{R}[X]$, a est racine d'ordre α de P si, et seulement si $\bar{a}$ est racine d'ordre α de P .

Démonstration. D'après la remarque précédente, $(X - a)^\alpha$ divise P si, et seulement si, $(X - a)^\alpha = (X - \bar{a})^\alpha$ divise $\overline{P}$, ce qui prouve a). b) résulte immédiatement de a) car si $P \in \mathbb{R}[X]$, on a $P = \overline{P}$.

5.7.3. POLYNÔMES DE $\mathbb{R}[X]$

Dans ce numéro, nous nous proposons de déterminer tous les polynômes irréductibles de $\mathbb{R}[X]$.

5.7.3.1. Théorème

Les polynômes irréductibles de $\mathbb{R}[X]$ sont les polynômes de degré un et les polynômes de degré deux sans racine réelle.

Démonstration. Les polynômes de degré 1 sont irréductibles (Exemple 5.4.5.3). De même, un polynôme P du second degré qui n'a aucune racine réelle est irréductible dans $\mathbb{R}[X]$. En effet, si P n'était pas irréductible dans $\mathbb{R}[X]$, il aurait un diviseur de degré 1, donc une racine réelle.

Réciproquement, montrons qu'un polynôme $P \in \mathbb{R}[X]$ de degré $n > 2$ n'est pas irréductible. Si P admet une racine réelle, le résultat est évident. Supposons que P n'ait aucune racine réelle. Considéré comme élément de $\mathbb{C}[X]$, P admet une racine complexe a . D'après le Théorème 5.7.2.3 b), $\bar{a}$ est aussi racine de P ; donc P est divisible par $X - a$ et par $X - \bar{a}$. Les deux polynômes $X - a$ et $X - \bar{a}$ étant premiers entre eux puisque $a \neq \bar{a}$, P est divisible par $A = (X - a)(X - \bar{a})$ d'après le Corollaire 5.4.3.8. Mais

$$A = X^2 - (a + \bar{a})X + a\bar{a}$$

est un polynôme à coefficients réels puisque

$$a + \bar{a} = 2\operatorname{Re}(a) \in \mathbb{R} \quad \text{et} \quad a\bar{a} = |a|^2 \in \mathbb{R}.$$

donc P n'est pas irréductible.

5.7.3.2. Corollaire

Tout polynôme P non nul de $\mathbb{R}[X]$ de degré $n \geq 1$ et de coefficient dominant b_n s'écrit, d'une façon unique, sous la forme :

$$(5.7.3.1) \quad P = b_n \prod_{i=1}^r (X - a_i)^{\alpha_i} \prod_{j=1}^m (X^2 + p_j X + q_j)^{\beta_j}$$

où les a_i ($1 \leq i \leq r$), les p_j et les q_j ($1 \leq j \leq m$) sont des nombres réels, où les polynômes $X^2 + p_j X + q_j$ vérifient les relations $p_j^2 - 4q_j < 0$ et où les entiers $\alpha_1, \dots, \alpha_r, \beta_1, \dots, \beta_m$ vérifient la relation :

$$(5.7.3.2) \quad \alpha_1 + \dots + \alpha_r + 2(\beta_1 + \dots + \beta_m) = n.$$

Démonstration. Soient $a_1, \dots, a_r$ les racines réelles distinctes de P , $\alpha_1, \dots, \alpha_r$ leurs ordres de multiplicité et soient $z_1, \overline{z_1}, \dots, z_m, \overline{z_m}$ les racines deux à deux imaginaires conjuguées de P avec les ordres de multiplicité $\beta_1, \dots, \beta_m$.

D'après le Théorème 5.7.2.1, on a

$$P = b_n (X - a_1)^{\alpha_1} \dots (X - a_r)^{\alpha_r} (X - z_1)^{\beta_1} (X - \overline{z_1})^{\beta_1} \dots (X - z_m)^{\beta_m} (X - \overline{z_m})^{\beta_m}$$

$$\text{et } \alpha_1 + \dots + \alpha_r + 2(\beta_1 + \dots + \beta_m) = n.$$

D'autre part, pour $1 \leq j \leq m$, on a

$$(X - z_j)^{\beta_j} (X - \overline{z_j})^{\beta_j} = (X^2 + p_j X + q_j)^{\beta_j}$$

$$\text{avec } z_j + \overline{z_j} = -p_j \in \mathbb{R}, \quad z_j \overline{z_j} = q_j \in \mathbb{R} \quad \text{et} \quad p_j^2 - 4q_j < 0.$$

On en déduit la formule (5.7.3.1) appelée **décomposition de Gauss du polynôme P** .

L'unicité de cette décomposition résulte du fait que les facteurs $X - a_i$ ($1 \leq i \leq r$) et $X^2 + p_j X + q_j$ ($1 \leq j \leq m$) sont irréductibles et normalisés.

5.7.4. RELATIONS ENTRE LES COEFFICIENTS ET LES RACINES D'UN POLYNÔME

On sait que si $P = aX^2 + bX + c$ est un polynôme de $\mathbb{C}[X]$ et si x_1 et x_2 sont les racines de P , alors

$$x_1 + x_2 = -\frac{b}{a} \quad \text{et} \quad x_1 x_2 = \frac{c}{a}.$$

Nous nous proposons de généraliser ces relations entre les coefficients et les racines d'un polynôme du second degré au cas d'un polynôme de degré n quelconque.

5.7.4.1. Théorème

Soit $P = a_n X^n + a_{n-1} X^{n-1} + \dots + a_0$ un élément de $\mathbb{C}[X]$ de degré $n \geq 1$. Soient $x_1, \dots, x_n$ les n racines distinctes ou non de P . Soient $\sigma_1, \dots, \sigma_n$ les fonctions symétriques élémentaires des racines :

$$\sigma_1 = \sum_i x_i = x_1 + \dots + x_n,$$

$$\sigma_2 = \sum_{i < j} x_i x_j = x_1 x_2 + x_1 x_3 + \dots + x_{n-1} x_n,$$

$$\sigma_3 = \sum_{i < j < k} x_i x_j x_k = x_1 x_2 x_3 + x_1 x_2 x_4 + \dots + x_{n-2} x_{n-1} x_n,$$

$$\sigma_n = x_1 x_2 \dots x_n.$$

On a

$$\sigma_1 = -\frac{a_{n-1}}{a_n}, \quad \sigma_2 = \frac{a_{n-2}}{a_n}, \quad \dots, \quad \sigma_k = (-1)^k \frac{a_{n-k}}{a_n}, \quad \dots, \quad \sigma_n = (-1)^n \frac{a_0}{a_n}.$$

Démonstration. D'après le Théorème 5.7.2.1, on peut écrire :

$$P = a_n (X - x_1) \dots (X - x_n).$$

Développons et ordonnons le produit $(X - x_1) \dots (X - x_n)$ suivant les puissances décroissantes de X . On obtient

$$P = a_n (X^n - \sigma_1 X^{n-1} + \sigma_2 X^{n-2} + \dots + (-1)^k \sigma_k X^{n-k} + \dots + (-1)^n \sigma_n).$$

En égalant les coefficients de X^{n-k} ($0 \leq k \leq n$) dans les deux expressions de P , il vient $a_{n-k} = (-1)^k \sigma_k a_n$, d'où les formules annoncées.

5.8. Dérivation des polynômes

Dans ce paragraphe, nous allons définir une notion algébrique de dérivée qui coïncide avec la notion de dérivée des fonctions polynômes dans $\mathbb{R}$. On suppose que l'anneau K est de caractéristique 0.

5.8.1. DÉRIVÉE D'UN POLYNÔME

5.8.1.1. Définition

Soit K un anneau commutatif et $P = a_n X^n + a_{n-1} X^{n-1} + \dots + a_0$ un polynôme de $K[X]$. On appelle polynôme dérivé de P le polynôme $P' \in K[X]$ défini par :

$$(5.8.1.1) \quad P' = na_n X^{n-1} + (n-1) a_{n-1} X^{n-2} + \dots + 2a_2 X + a_1.$$

En particulier, si P est une constante, P' est le polynôme nul.

Si $\deg(P) = n$, c'est-à-dire si $a_n \neq 0$, $na_n \neq 0$ si la caractéristique de K est nulle ou ne divise pas n . Par suite, si la caractéristique de K est nulle, on a $\deg(P') = \deg(P) - 1$.

5.8.1.2. Théorème

Soient P et Q deux éléments de $K[X]$, $\lambda \in K$ et n un entier > 0 . On a

$$\begin{aligned} (P + Q)' &= P' + Q', & (\lambda P)' &= \lambda P' \\ (PQ)' &= P'Q + PQ', & (P^n)' &= nP^{n-1}P'. \end{aligned}$$

Démonstration. Soient

$$\begin{aligned} P &= a_m X^m + a_{m-1} X^{m-1} + \dots + a_0, & a_m &\neq 0 \\ Q &= b_p X^p + b_{p-1} X^{p-1} + \dots + b_0, & b_p &\neq 0, \end{aligned}$$

deux polynômes de $K[X]$. Supposons $p \leq m$; pour $p < k \leq m$, b_k est nul. Donc

$$P + Q = \sum_{k=0}^m (a_k + b_k) X^k.$$

Prenons les dérivées des polynômes P , Q et $P + Q$:

$$P' = \sum_{k=1}^m K a_k X^{k-1}, \quad Q' = \sum_{k=1}^p k b_k X^{k-1}, \quad (P+Q)' = \sum_{k=1}^m k(a_k + b_k) X^{k-1}.$$

Par suite, on a

$$(P + Q)' = P' + Q'.$$

On montrerait de même que pour tout $\lambda \in K$, on a :

$$(\lambda P)' = \lambda P'.$$

Pour démontrer la troisième formule, supposons d'abord que P et Q soient des monômes :

$$P = \lambda X^p, \quad Q = \mu X^q, \quad \lambda, \mu \in K.$$

En écartant le cas trivial $p = 0$ et $q = 0$, on a

$$(PQ)' = (\lambda \mu X^{p+q})' = (p+q) \lambda \mu X^{p+q-1},$$

$$P'Q + PQ' = p \lambda \mu X^{p+q-1} + q \lambda \mu X^{p+q-1} = (p+q) \lambda \mu X^{p+q-1} = (PQ)'.$$

La propriété est donc vraie dans ce cas.

Soient alors $P = \sum_{r=0}^n a_r X^r$ et $Q = \sum_{k=0}^m b_k X^k$ deux polynômes de $K[X]$.

On a :

$$PQ = \sum_r \sum_k a_r b_k X^r X^k.$$

D'où, d'après ce qui précède :

$$\begin{aligned} (PQ)' &= \sum_r \sum_k a_r b_k (X^r X^k)' = \sum_r \sum_k a_r b_k ((X^r)' X^k + X^r (X^k)') \\ &= \sum_r r a_r X^{r-1} \sum_k b_k X^k + \sum_r a_r X^r \sum_k k b_k X^{k-1} = P'Q + PQ'. \end{aligned}$$

D'après la troisième formule, le polynôme dérivé de P^2 est $2P'P$. Par récurrence sur n , on voit facilement que $(P^n)' = nP^{n-1}P'$.

5.8.1.3. Définition

Soient $P \in K[X]$ et n un entier. On définit le polynôme dérivé d'ordre n de P par récurrence en posant $P^{(n)} = (P^{(n-1)})'$.

On pose par convention, $P^{(0)} = P$.

5.8.2. FORMULE DE TAYLOR

On suppose maintenant que K est un corps commutatif de caractéristique 0.

5.8.2.1. Théorème

Soient P un polynôme de $K[X]$ de degré $n \geq 1$, et a un élément de K . On a (formule de Taylor) :

$$\begin{aligned} (5.8.2.1) \quad P &= P(a) + (X - a) P'(a) + \dots + \frac{(X - a)^k}{K!} P^{(k)}(a) \\ &\quad + \dots + \frac{(X - a)^n}{n!} P^{(n)}(a) \end{aligned}$$

où $P^{(k)}(a)$ désigne la valeur au point a de la fonction associée au polynôme $P^{(k)}$.

Démonstration. Soit

$$P = a_0 + a_1 X + \dots + a_n X^n, \quad a_n \neq 0.$$

Pour simplifier l'écriture, nous noterons encore P la fonction polynôme $\tilde{P}$ associée à P .

Nous allons démontrer la formule (5.8.2.1), d'abord lorsque $a = 0$. On a

$$P(0) = a_0$$

$$P' = a_1 + 2a_2X + \dots + na_nX^{n-1} \implies P'(0) = a_1$$

$$P^{(k)} = k(k-1) \dots 3.2.1 a_k + \dots + n(n-1) \dots (n-K+1) a_n X^{n-k},$$

d'où $P^{(k)}(0) = K! a_k$,

$$P^{(n)} = n(n-1) \dots 3.2.1 a_n \implies P^{(n)}(0) = n! a_n.$$

On en déduit

$$a_0 = P(0), \quad a_1 = P'(0), \quad a_2 = \frac{1}{2!} P''(0), \quad \dots, \quad a_k = \frac{1}{k!} P^{(k)}(0), \quad \dots$$

$$\dots \quad a_n = \frac{1}{n!} P^{(n)}(0).$$

En remplaçant les a_i ($0 \leq i \leq n$) par ces expressions, on obtient la formule

$$P = P(0) + X P'(0) + \frac{X^2}{2!} P''(0) + \dots + \frac{X^k}{k!} P^{(k)}(0) + \dots + \frac{X^n}{n!} P^{(n)}(0)$$

appelée **formule de Mac-Laurin**.

Avec les notations précédentes, posons

$$Q(Y) = P(Y + a).$$

On a d'après la formule de Mac-Laurin,

$$Q = Q(0) + Y Q'(0) + \dots + \frac{Y^k}{k!} Q^{(k)}(0) + \dots + \frac{Y^n}{n!} Q^{(n)}(0).$$

Comme

$$Q(Y) = P(Y + a) = a_0 + a_1(Y + a) + \dots + a_K(Y + a)^k + \dots + a_n(Y + a)^n,$$

on a

$$Q(0) = P(a), \quad Q'(0) = P'(a), \quad \dots, \quad Q^{(k)}(0) = P^{(k)}(a), \quad \dots, \quad Q^{(n)}(0) = P^{(n)}(a).$$

On en déduit la formule de Taylor

$$(5.8.2.2) \quad P(Y + a) = P(a) + Y P'(a) + \dots + \frac{Y^k}{k!} P^{(k)}(a) + \dots + \frac{Y^n}{n!} P^{(n)}(a)$$

ou, en posant $Y + a = X$,

$$P = P(a) + (X - a) P'(a) + \dots + \frac{(X - a)^k}{k!} P^{(k)}(a) + \dots + \frac{(X - a)^n}{n!} P^{(n)}(a). \quad \square$$

La formule de Taylor va nous permettre d'obtenir une condition nécessaire et suffisante pour qu'un élément a de K soit une racine d'ordre k d'un polynôme. Démontrons d'abord un Lemme.

5.8.2.2. Lemme

Soient P un élément de $K[X]$, a un élément de K et k un entier ≥ 2 . Alors, si a est racine d'ordre k de P , a est racine d'ordre $k - 1$ de P' .

Démonstration. Si a est racine d'ordre k de P , il existe un polynôme Q de $K[X]$ tel que

$$P = (X - a)^k Q \quad \text{avec} \quad Q(a) \neq 0.$$

Donc

$$P' = k(X - a)^{k-1} Q + (X - a)^k Q' = (X - a)^{k-1} H$$

où $H = kQ + (X - a) Q'$. On a $H(a) = kQ(a) \neq 0$, car $Q(a) \neq 0$ et $k \cdot 1 \neq 0$ dans K puisque K est de caractéristique 0. Donc a est bien racine d'ordre $k - 1$ de P' .

5.8.2.3. Théorème

Soient K un corps de caractéristique 0, P un polynôme de $K[X]$, $a \in K$ et k un entier ≥ 1 . Pour que a soit racine d'ordre k de P , il faut et il suffit que :

$$(5.8.2.3) \quad P(a) = 0, \quad P'(a) = 0, \quad \dots, \quad P^{(k-1)}(a) = 0, \quad P^{(k)}(a) \neq 0.$$

Démonstration. Si a est racine d'ordre k de P , d'après le Lemme 5.8.2.2, a est racine d'ordre $k - 1, k - 2, \dots, 1$ de $P', P'', \dots, P^{(k-1)}$ respectivement et $P^{(k)}(a) \neq 0$. Les relations (5.8.2.3) sont donc vérifiées.

Réciproquement, si les relations (5.8.2.3) sont vérifiées, le degré n de P est $\geq k$.

La formule de Taylor donne :

$$P = (X - a)^k \left[\frac{P^{(k)}(a)}{k!} + \frac{X - a}{(k + 1)!} P^{(k+1)}(a) + \dots + \frac{(X - a)^{n-k}}{n!} P^{(n)}(a) \right]$$

$$= (X - a)^k Q$$

avec $Q(a) = \frac{1}{k!} P^{(k)}(a) \neq 0$. Donc a est racine d'ordre k de P .

5.9. Polynômes à plusieurs indéterminées

Nous allons définir les polynômes à plusieurs indéterminées en nous inspirant du cas d'une indéterminée. Nous ne ferons pas une étude détaillée des propriétés de l'anneau des polynômes à plusieurs indéterminées mais nous invitons le lecteur intéressé à le faire lui-même (c'est très instructif). Les notions fondamentales seront exposées dans le cas de deux indéterminées.

5.9.1. DÉFINITIONS GÉNÉRALES

5.9.1.1. Définition

Soit K un anneau commutatif. On appelle **polynôme à deux indéterminées à coefficients dans K** , toute suite double $(a_{ij})_{i \geq 0, j \geq 0}$ d'éléments de K dont seuls un nombre fini de termes sont non nuls.

On note

$$P = (a_{00}, a_{10}, a_{01}, \dots) \quad \text{ou} \quad P = (a_{ij})_{i \geq 0, j \geq 0}.$$

Les a_{ij} sont appelés **coefficients** du polynôme P ; a_{00} s'appelle le **terme constant**.

L'ensemble des polynômes à deux indéterminées à coefficients dans l'anneau commutatif K se note $K[X, Y]$.

Soient $P = (a_{ij})$ et $Q = (b_{ij})$ deux polynômes de $K[X, Y]$. On définit la somme et le produit de P et Q par les formules :

$$(5.9.1.1) \quad \begin{aligned} (a_{ij} + b_{ij}) &= (a_{ij}) + (b_{ij}) \\ (a_{ij})(b_{ij}) &= (c_{ij}) \end{aligned}$$

où

$$(5.9.1.2) \quad c_{ij} = \sum_{\substack{i' + i'' = i \\ j' + j'' = j}} a_{i'j'} b_{i''j''}.$$

On vérifie facilement que l'on définit bien ainsi sur $K[X, Y]$ une structure d'anneau commutatif (intègre si K est intègre).

On dit que deux polynômes $P = (a_{ij})$ et $Q = (b_{ij})$ sont égaux si on a $a_{ij} = b_{ij}$ pour tout élément (i, j) de $\mathbb{N} \times \mathbb{N}$. On identifie K à un sous-anneau de $K[X, Y]$ en identifiant tout élément $a \in K$ au polynôme (a_{ij}) tel que $a_{00} = a$ et $a_{ij} = 0$ si $i \geq 1$ ou $j \geq 1$.

Soit X (resp. Y) le polynôme dont tous les coefficients sont nuls sauf a_{10} (resp. a_{01}) qui est égal à 1. En raisonnant comme dans le cas de $K[X]$, on vérifie facilement que tout polynôme $P = (a_{ij})_{i \geq 0, j \geq 0}$ s'écrit de façon unique sous la forme :

$$(5.9.1.3) \quad P = \sum_{\substack{i \geq 0 \\ j \geq 0}} a_{ij} X^i Y^j.$$

On dit que X et Y sont les indéterminées et que P est un polynôme en X et Y .

On définirait de même l'anneau $K[X_1, X_2, \dots, X_n]$ des polynômes à n indéterminées $X_1, X_2, \dots, X_n$ à coefficients dans K . Un élément de cet anneau s'écrit

$$\sum_{(i_1, \dots, i_n) \in \mathbb{N}^n} a_{i_1 \dots i_n} X_1^{i_1} \dots X_n^{i_n}$$

où un nombre fini de coefficients $a_{i_1 \dots i_n}$ sont non nuls.

5.9.2. ISOMORPHISME CANONIQUE DE $K[X][Y]$ SUR $K[X, Y]$

Soit K un anneau commutatif. $A = K[X]$ est un anneau commutatif. Pour tout polynôme $P \in A[Y]$, on a

$$P = \sum_j P_j Y^j$$

avec $P_j \in K[X]$, l'ensemble des indices j tels que $P_j \neq 0$ étant fini. Posons, pour tout indice j :

$$P_j = \sum_i a_{ij} X^i, \quad \text{où } a_{ij} \in K.$$

Alors

$$Q = \sum_j \left(\sum_i a_{ij} X^i \right) Y^j = \sum_{i,j} a_{ij} X^i Y^j$$

est un élément de $K[X, Y]$. On peut donc associer à P l'élément $Q = \sum_{i,j} a_{ij} X^i Y^j$ de $K[X, Y]$.

5.9.2.1. Théorème

L'application $P \mapsto Q$ est un isomorphisme de $A[Y]$ sur $K[X, Y]$.

Démonstration. Notons φ l'application qui fait correspondre à P l'élément Q . Elle est surjective car les coefficients a_{ij} de $\varphi(P) = Q = \sum_{i,j} a_{ij} X^i Y^j$ sont arbitraires.

Montrons qu'elle est injective.

Soit $A = \sum_j A_j Y^j$ avec $A_j = \sum_i \alpha_{ij} X^i$; on a

$\varphi(A) = B = \sum_{i,j} \alpha_{ij} X^i Y^j$. Si $\varphi(A) = \varphi(P)$, on a $a_{ij} = \alpha_{ij}$ quels que soient i et j , donc $A_j = P_j$ quel que soit j ; par suite $A = P$ et φ est injective.

On verrait de même qu'il existe un isomorphisme de $K[X_1, \dots, X_{n-1}][X_n]$ sur $K[X_1, \dots, X_{n-1}, X_n]$.

5.9.3. DEGRÉ D'UN POLYNÔME A DEUX INDÉTERMINÉES

5.9.3.1. Définition

Soit $P = \sum_{i,j} a_{ij} X^i Y^j$ un polynôme de $K[X, Y]$. On appelle **degré partiel** de P par rapport à X , le degré de P considéré comme élément de $K[Y][X]$, c'est-à-dire le plus grand des entiers i tels que $a_{ij} \neq 0$.

Ce degré partiel se note $\deg_X(P)$. On a les inégalités suivantes :

$$\deg_X(P + Q) \leq \sup(\deg_X(P), \deg_X(Q))$$

$$\deg_X(P \cdot Q) \leq \deg_X(P) + \deg_X(Q)$$

cette dernière inégalité étant une égalité si K est intègre.

On définit de même le degré partiel de P par rapport à Y .

On a $\deg_X(P) = -\infty$ si $P = 0$.

On appelle **degré total**, ou simplement **degré** de P , et on note $\deg(P)$, le plus grand des entiers $i + j$ pour tout les couples (i, j) tels que $a_{ij} \neq 0$.

Si $P, Q \in K[X, Y]$, on a aussi les inégalités :

$$\deg(P + Q) \leq \sup(\deg(P), \deg(Q))$$

$$\deg(P \cdot Q) \leq \deg(P) + \deg(Q).$$

Si par exemple $P = X^4 + X^2 Y^3$, $\deg_X(P) = 4$, $\deg_Y(P) = 3$, et $\deg(P) = 5$.

On dit que le polynôme $P = \sum_{i,j} a_{ij} X^i Y^j$ est **homogène de degré d** si $a_{ij} = 0$ pour $i + j \neq d$. Autrement dit, P est homogène de degré d si chacun des monômes qu'il contient est de degré total d .

En regroupant ensemble, dans l'expression de $P = \sum_{i,j} a_{ij} X^i Y^j$ les termes pour lesquels $i + j$ a une valeur donnée, on voit que tout polynôme $P \in K[X, Y]$ de degré total d s'écrit d'une manière unique sous la forme

$$(5.9.3.1) \quad P = P_0 + P_1 + \dots + P_d$$

où chaque polynôme P_r est homogène de degré r et où $P_d \neq 0$. On dit que P_r est la **composante homogène de degré r** de P .

Toutes les notions qui viennent d'être définies se généralisent aisément au cas de n indéterminées.

5.9.4. FONCTIONS POLYNÔMES

Nous traitons comme précédemment le cas de deux indéterminées X et Y , la généralisation étant facile à faire.

5.9.4.1. Définition

Soit $P = \sum_{i,j} a_{ij} X^i Y^j$ un polynôme de $K[X, Y]$. On appelle fonction polynôme de deux variables associée au polynôme P , l'application $\tilde{P}$ de $K \times K$ dans K , associant à tout élément (x, y) de $K \times K$, l'élément $\sum_{i,j} a_{ij} x^i y^j$.

On démontre, comme dans le cas de $K[X]$, que quels que soient les polynômes P et Q de $K[X, Y]$ et l'élément $\lambda \in K$, on a

$$(P + Q) = \tilde{P} + \tilde{Q}, \quad (P \cdot Q) = \tilde{P} \cdot \tilde{Q}, \quad (\lambda P) = \lambda \tilde{P}.$$

On démontre également que si K est un anneau intègre infini, l'application qui à un polynôme P de $K[X, Y]$ fait correspondre la fonction polynôme $\tilde{P}$ est un isomorphisme de $K[X, Y]$ sur l'anneau des fonctions polynômes à deux variables sur K . Ceci permet d'identifier un polynôme $P \in K[X, Y]$ à la fonction polynôme $\tilde{P}$ qui lui est associée lorsque $K = \mathbb{Q}, \mathbb{R}$ ou $\mathbb{C}$.

Toutes ces considérations se généralisent sans peine au cas de n indéterminées $X_1, \dots, X_n$.

5.9.5. DÉRIVATION PARTIELLE DES POLYNÔMES

5.9.5.1. Définition

Soient K un anneau commutatif et P un polynôme de $K[X, Y]$. On appelle polynôme dérivé partiel de P par rapport à l'indéterminée X (resp. Y), et on note P'_X (resp. P'_Y) ou $\frac{\partial P}{\partial X}$ (resp. $\frac{\partial P}{\partial Y}$) le polynôme dérivé de P considéré comme élément de $K[Y][X]$ (resp. $K[X][Y]$).

Si $K = \mathbb{R}$, on retrouve les dérivées partielles définies en analyse. On a des règles de calcul analogues à celles qui ont été établies au paragraphe 5.8 dans le cas d'une indéterminée.

Si $P = \sum a_{ij} X^i Y^j$ est un polynôme de $K[X, Y]$, on a

$$P''_{XY} = P''_{YX}.$$

Il suffit de vérifier cette formule lorsque $P = X^i Y^j$. Dans ce cas, on a :

$$P''_{XY} = (X^{i-1} Y^j)'_Y = i X^{i-1} j Y^{j-1} = ij X^{i-1} Y^{j-1}$$

$$P''_{YX} = (X^i j Y^{j-1})'_X = i X^{i-1} j Y^{j-1} = ij X^{i-1} Y^{j-1}.$$

Cette remarque permet de calculer les dérivées partielles successives de P sans se préoccuper de l'ordre des dérivations.

On note $P^{(\alpha+\beta)}_{X^\alpha Y^\beta}$ le polynôme obtenu en dérivant α fois par rapport à X et β fois par rapport à Y .

5.9.5.2. Théorème (Formule d'Euler)

Soient K un corps de caractéristique 0, et P un élément de $K[X, Y]$. Pour que P soit homogène de degré n , il faut et il suffit que $XP'_X + YP'_Y = nP$.

Démonstration. Supposons P homogène de degré n . P s'écrit donc

$$P = \sum_{i=0}^n a_i X^i Y^{n-i}.$$

Alors :

$$XP'_X + YP'_Y = \sum_{i=0}^n i a_i X^i Y^{n-i} + \sum_{i=0}^n (n-i) a_i X^i Y^{n-i} = nP.$$

Réciproquement, supposons qu'on ait

$$XP'_X + YP'_Y = nP.$$

Nous savons que tout polynôme P de $K[X, Y]$ s'écrit de façon unique sous la forme $P = \sum_i P_i$ où chaque P_i est homogène de degré i . Alors d'après la première partie, on a

$$nP = XP'_X + YP'_Y = \sum_i (X(P_i)'_X + Y(P_i)'_Y) = \sum_i i P_i.$$

D'où

$$\sum_i n P_i = \sum_i i P_i.$$

En vertu de l'unicité de la décomposition d'un polynôme en somme de polynômes homogènes, nous obtenons $(n-i)P_i = 0$ pour tout i ; donc $P_i = 0$ pour $i \neq n$, puisque K est de caractéristique 0. Donc $P = P_n$ est homogène de degré n .

5.9.5.3. Théorème (Formule de Taylor)

Soient K un corps de caractéristique 0, a et b des éléments de K , et P un élément de $K[X, Y]$. On a

$$(5.9.5.1) \quad P(X, Y) = \sum_{i \geq 0, j \geq 0} \frac{1}{i!j!} P_{X^i Y^j}^{(i+j)}(a, b) (X-a)^i (Y-b)^j.$$

Démonstration. Notons d'abord que 5.9.5.1 a un sens car la somme du second membre est finie puisque les dérivées $P^{(i+j)}$ sont nulles pour $i+j$ assez grand.

D'autre part, tout élément de $K[X, Y]$ est somme d'éléments de la forme $\lambda(X-a)^r(Y-b)^s$, où $\lambda \in K$, et où r et s sont dans $\mathbb{N}$. En effet, pour un monôme $X^\alpha Y^\beta$, on a

$$X^\alpha Y^\beta = ((X-a) + a)^\alpha ((Y-b) + b)^\beta.$$

La formule du binôme donne

$$X^\alpha Y^\beta = \sum_{k=0}^{\alpha} C_{\alpha}^k a^{\alpha-k} (X-a)^k \sum_{q=0}^{\beta} C_{\beta}^q b^{\beta-q} (Y-b)^q,$$

expression que l'on peut toujours écrire sous la forme indiquée.

Par linéarité, il suffit donc de démontrer (5.9.5.1) lorsque $P(X, Y) = (X-a)^m(Y-b)^n$.

On a

$$P_{X,Y}^{(i+j)}(a, b) = 0 \text{ pour } (i, j) \neq (m, n) \text{ et } P_{X,Y}^{m+n}(a, b) = m! n!$$

Le second membre de (5.9.5.1) se réduit donc à $\frac{1}{m! n!} m! n!$
 $(X-a)^m(Y-b)^n$, d'où le théorème.

Fractions rationnelles

Soit K un corps commutatif. Nous savons que l'ensemble $K[X]$ des polynômes à une indéterminée à coefficients dans K est un anneau commutatif intègre dans lequel les seuls éléments inversibles sont les polynômes de degré 0. $K[X]$ n'est donc pas un corps, mais on peut construire le corps des fractions de $K[X]$ (cf. Chapitre 4, § 4.5, n° 4.5.3).

Ce corps s'appelle le corps des fractions rationnelles à une indéterminée à coefficients dans K . On le note $K(X)$.

5.10. Définition du corps des fractions rationnelles

Nous allons rappeler dans ce paragraphe les principales étapes de la construction du corps $K(X)$ des fractions rationnelles à une indéterminée à coefficients dans K .

5.10.1. FRACTIONS RATIONNELLES

On pose comme d'habitude $K[X]^* = K[X] - \{0\}$.

• $K(X)$ est l'ensemble quotient de $K[X] \times K[X]^*$ par la relation d'équivalence $\mathcal{R}$:

$$(A, B)\mathcal{R}(A_1, B_1) \iff AB_1 = A_1B.$$

Une fraction rationnelle F de $K(X)$ est donc une classe d'équivalence représentée par un couple (A, B) d'élément de $K[X]$ dans lequel $B \neq 0$; un autre couple (A_1, B_1) représente la même fraction rationnelle F si, et seulement si $AB_1 = A_1B$.

Si (A, B) est un représentant quelconque de F , on convient d'écrire $F = \frac{A}{B}$ ou $F = AB^{-1}$; on dit que A est le **numérateur** et que B est le **dénominateur** de la fraction rationnelle F .

• Dans $K[X] \times K[X]^*$, on définit l'addition et la multiplication en posant :

$$(A, B) + (C, D) = (AD + BC, BD)$$

$$(A, B) \cdot (C, D) = (AC, BD).$$

Les deux lois de $K(X)$ sont les lois quotients et on les note aussi, $+$ et $\cdot$; alors le triplet $(K(X), +, \cdot)$ est un corps commutatif. L'élément neutre pour l'addition est la fraction rationnelle nulle 0 qui est la classe des couples $(0, B)$ tels que $B \neq 0$. L'élément neutre pour la multiplication, appelée **fraction rationnelle unité**, et notée 1 , est la classe des couples (B, B) avec $B \neq 0$. Pour la multiplication, l'inverse de la fraction rationnelle non nulle F , classe de (A, B) avec $A \neq 0$ et $B \neq 0$ est la fraction rationnelle notée $\frac{1}{F}$.

L'application qui, à $A \in K[X]$, associe la fraction rationnelle dont un représentant est le couple $(A, 1)$, est un morphisme injectif de l'anneau $K[X]$ dans l'anneau $K(X)$. On peut donc identifier $K[X]$ au sous-anneau de $K(X)$ constitué par les fractions rationnelles dont les représentants sont de la forme $(A, 1)$.

5.10.1.1. Définition

Soit $F \in K(X) - \{0\}$. On appelle **représentant irréductible de F** ou **forme irréductible de F** toute représentation de F sous la forme $\frac{A}{B}$, où A et B sont des éléments de $K[X]$ premiers entre eux.

Montrons que toute fraction rationnelle admet une forme irréductible.

5.10.1.2. Théorème

Soit F un élément de $K(X) - \{0\}$. Alors :

a) F possède des représentants irréductibles.

b) Si (A, B) est un représentant irréductible de F , les autres représentants irréductibles de F sont de la forme $(\lambda A, \lambda B)$, où $\lambda \in K^*$, et les représentants de F sont de la forme (AC, BC) , où $C \in K[X]^*$.

Démonstration. a) Soit (A, B) un représentant de F . Par hypothèse, on a $A \neq 0$ et $B \neq 0$. Soit D le PGCD de A et B . On a $A = A_1 D$ et $B = B_1 D$, avec $\text{PGCD}(A_1, B_1) = 1$, $B_1 \neq 0$ puisque B n'est pas nul. Par suite $\frac{A}{B} = \frac{A_1}{B_1} = F$ et $\frac{A_1}{B_1}$ est bien un représentant irréductible de F .

b) Soit (A, B) un représentant irréductible de F et soit (A_1, B_1) un autre représentant irréductible de F . On a $AB_1 = A_1 B$. Donc B divise AB_1 et est premier avec A ; d'après le Théorème de Gauss, B divise B_1 . De même B_1 divise B . Ainsi, il existe $C \in K[X]$ tel que $B_1 = BC$. Comme l'anneau $K[X]$ est intègre et puisque $B \neq 0$, l'égalité $AB_1 = A_1 B$ qui s'écrit $A_1 B = ABC$, montre que $A_1 = AC$. On a ensuite $B_1 = BC$. De plus, P.G.C.D. $(A_1, B_1) = C$, d'où P.G.C.D. $(A_1, B_1) = 1$ si et seulement si $C \in K^*$. Si (A, B) est un représentant irréductible de F , il est clair que, pour tout $C \in K[X]^*$, (AC, BC) est un représentant de F .

Notation. Par abus de langage, nous dirons et nous noterons :

- $\frac{A}{B}$ la fraction rationnelle dont un représentant est le couple (A, B) .
- $\frac{A}{B} = \frac{A_1}{B_1}$ si (A, B) et (A_1, B_1) représentent la même fraction rationnelle.
- La fraction irréductible $\frac{A}{B}$ au lieu de la fraction de représentant irréductible (A, B) .
- Si B divise A , $A = BQ$. La fraction $\frac{A}{B}$ sera donc identifiée au polynôme Q .

Les opérations dans $K(X)$ s'expriment par

$$\frac{A}{B} + \frac{C}{D} = \frac{AD + BC}{BD} \quad \text{et} \quad \frac{A}{B} \cdot \frac{C}{D} = \frac{AC}{BD}.$$

5.10.1.3. Théorème

Soit $F \in K(X) - \{0\}$. Si $\frac{A}{B}$ est un représentant quelconque de F , l'entier $\deg(A) - \deg(B)$ ne dépend que de F . On l'appelle le **degré de la fraction rationnelle F** , et on le note $\deg(F)$.

Démonstration. Soient en effet $\frac{A}{B}$ et $\frac{A_1}{B_1}$ deux représentants de F . On a $AB_1 = A_1 B$, donc

$$\deg(A) + \deg(B_1) = \deg(A_1) + \deg(B).$$

Comme $\deg(B)$ et $\deg(B_1)$ sont des éléments de $\mathbb{N}$ puisque $B \neq 0$ et $B_1 \neq 0$, on en déduit

$$\deg(A) - \deg(B) = \deg(A_1) - \deg(B_1).$$

Comme pour les polynômes, on convient de poser $\deg(0) = -\infty$.

5.10.1.4. Théorème

Si $\frac{A}{B}$ et $\frac{C}{D}$ sont des éléments non nuls de $K(X)$, on a :

$$\begin{aligned} \deg\left(\frac{A}{B} + \frac{C}{D}\right) &\leq \sup\left(\deg\left(\frac{A}{B}\right), \deg\left(\frac{C}{D}\right)\right) \\ \deg\left(\frac{A}{B} \cdot \frac{C}{D}\right) &= \deg\left(\frac{A}{B}\right) + \deg\left(\frac{C}{D}\right). \end{aligned}$$

La vérification est immédiate et est laissée au lecteur.

5.10.2. FONCTION RATIONNELLE

Soit K un corps commutatif. Nous allons associer à toute fraction rationnelle sur K une fonction de K dans K , comme cela a été fait pour les polynômes (voir § 5.6).

5.10.2.1. Définition

Soit $F = \frac{P}{Q}$ une fraction rationnelle de $K(X)$ écrite sous forme irréductible.

On appelle **pôle** de F toute racine de son dénominateur. On dit que $a \in K$ est un **pôle d'ordre** α de F si a est un zéro d'ordre α de Q .

Toute racine de multiplicité k du polynôme P est dite **racine d'ordre** k de F .

5.10.2.2. Remarque

Si $\frac{P}{Q}$ et $\frac{P_1}{Q_1}$ sont deux formes irréductibles de F , P et P_1 sont associés et ont donc les mêmes zéros ; de même Q et Q_1 étant associés ont les mêmes racines. Les notions de pôle et de racine sont donc indépendantes du représentant irréductible choisi. Mais il est indispensable, dans la définition précédente, de prendre un représentant irréductible ; ainsi $+1$ n'est pas un pôle de la fraction rationnelle $\frac{X^3 - 1}{X^2 - 1}$.

5.10.2.3. Définition

Soit F une fraction rationnelle écrite sous forme irréductible $\frac{P}{Q}$. On appelle **domaine de définition** de F , et on note D_F , le complémentaire dans K de l'ensemble des pôles de F .

On appelle **fonction rationnelle associée** à F , l'application $\tilde{F}$ de D_F dans K définie par $\tilde{F}(x) = \frac{\tilde{P}(x)}{\tilde{Q}(x)}$ pour tout $x \in D_F$.

Cette définition est justifiée par le fait que si $\frac{P}{Q}$ et $\frac{P_1}{Q_1}$ sont deux représentants irréductibles de F , $\frac{P(x)}{Q(x)} = \frac{P_1(x)}{Q_1(x)}$ et $\tilde{F}$ ne dépend donc pas du représentant irréductible choisi.

5.10.2.4. Définition

Soient F et G deux fractions rationnelles sur K , de domaines de définition D_F et D_G respectivement. On définit les applications $F \dot{+} G$ et $F \dot{\cdot} G$ de $D_F \cap D_G$ dans K en posant :

$$\begin{aligned}(F \dot{+} G)(x) &= \tilde{F}(x) + \tilde{G}(x) \\ (F \dot{\cdot} G)(x) &= \tilde{F}(x) \cdot \tilde{G}(x)\end{aligned}$$

pour tout $x \in D_F \cap D_G$.

5.10.2.5. Théorème

Soient K un corps commutatif infini, et F et G deux fractions rationnelles sur K telles que $\tilde{F}(x) = \tilde{G}(x)$ pour tout $x \in D_F \cap D_G$. Alors $F = G$.

Démonstration. Soient $\frac{P}{Q}$ et $\frac{P_1}{Q_1}$ des représentants irréductibles de F et G respectivement. Pour tout $x \in D_F \cap D_G$, on a $\frac{\tilde{P}(x)}{\tilde{Q}(x)} = \frac{\tilde{P}_1(x)}{\tilde{Q}_1(x)}$. Donc

$$\tilde{P}(x) \tilde{Q}_1(x) - \tilde{Q}(x) \tilde{P}_1(x) = 0.$$

Autrement dit la fonction polynôme associée au polynôme $PQ_1 - QP_1$ de $K[X]$ s'annule en tout point de l'ensemble $D_F \cap D_G$ qui est infini puisque le nombre de pôles est fini. Ainsi le polynôme $PQ_1 - QP_1$ s'annule pour une infinité de valeurs, donc (Corollaire 5.6.2.6), $PQ_1 - QP_1 = 0$, i.e. $F = G$.

5.11. Décomposition d'une fraction rationnelle en éléments simples

5.11.1. THÉORÈMES GÉNÉRAUX

5.11.1.1. Lemme

Soit F un élément de $K(X)$. Il existe un unique polynôme E tel que l'on ait $F = E + R$ où R est une fraction rationnelle de degré strictement négatif. On dit que E est la partie entière de F .

Démonstration. Soit $\frac{A}{B}$ un représentant quelconque de F . Comme $B \neq 0$, on peut effectuer la division euclidienne de A par B . On obtient

$$A = EB + D \text{ avec } D = 0 \text{ ou } \deg(D) < \deg(B).$$

Si $D = 0$, $F = E + 0$ et l'existence est établie. Sinon, on a

$$\frac{A}{B} = \frac{EB + D}{B} = E + \frac{D}{B}.$$

Posons $R = \frac{D}{B}$; on a $\deg(R) < 0$.

L'écriture est unique car si E et E_1 sont des polynômes tels que $F = E + R = E_1 + R_1$ avec $\deg(R) < 0$ et $\deg(R_1) < 0$, on a

$$\deg((F - E) + (E_1 - F)) \leq \sup(\deg(F - E), \deg(E_1 - F)).$$

On en déduit $\deg(E_1 - E) < 0$, d'où $E_1 = E$.

5.11.1.2. Lemme

Soit $F = \frac{P}{Q_1 Q_2 \dots Q_n}$ une fraction rationnelle, où les polynômes $Q_1, \dots, Q_n$ sont premiers entre eux deux à deux et $\deg(P) < \deg(Q_1 Q_2 \dots Q_n)$. Il existe une famille et une seule de polynôme P_i , $1 \leq i \leq n$, tels que

$$\frac{P}{Q_1 Q_2 \dots Q_n} = \sum_{i=1}^n \frac{P_i}{Q_i}; \quad \deg\left(\frac{P_i}{Q_i}\right) < 0 \quad \text{pour tout } i \in \{1, \dots, n\}.$$

Démonstration. Raisonnons par récurrence sur n .

a) Existence d'une décomposition lorsque $Q = Q_1 Q_2$:

Q_1 et Q_2 étant deux polynômes premiers entre eux, il existe, d'après le Théorème de Bezout, deux polynômes R_1 et R_2 tels que $R_2 Q_1 + R_1 Q_2 = 1$, soit $P = (P R_2) Q_1 + (P R_1) Q_2$.

On en déduit

$$\frac{P}{Q_1 Q_2} = \frac{P R_1}{Q_1} + \frac{P R_2}{Q_2}.$$

D'après le Lemme 5.11.1.1, il existe des polynômes E_1, P_1, E_2, P_2 tels que :

$$\frac{P R_1}{Q_1} = E_1 + \frac{P_1}{Q_1}, \quad \frac{P R_2}{Q_2} = E_2 + \frac{P_2}{Q_2}, \quad \deg\left(\frac{P_1}{Q_1}\right) < 0, \quad \deg\left(\frac{P_2}{Q_2}\right) < 0.$$

D'où

$$\frac{P}{Q_1 Q_2} = E_1 + E_2 + \frac{P_1}{Q_1} + \frac{P_2}{Q_2}, \quad \deg\left(\frac{P_1}{Q_1}\right) < 0, \quad \deg\left(\frac{P_2}{Q_2}\right) < 0.$$

On en déduit $\deg\left(\frac{P_1}{Q_1} + \frac{P_2}{Q_2}\right) < 0$, donc $E_1 + E_2$ est la partie entière de $\frac{P}{Q_1 Q_2}$. Comme cette partie entière est nulle d'après l'hypothèse sur les degrés de P et $Q_1 Q_2$, on a finalement

$$\frac{P}{Q_1 Q_2} = \frac{P_1}{Q_1} + \frac{P_2}{Q_2} \quad \text{avec} \quad \deg\left(\frac{P_1}{Q_1}\right) < 0 \quad \text{et} \quad \deg\left(\frac{P_2}{Q_2}\right) < 0.$$

b) Unicité de la décomposition dans le cas de $Q = Q_1 Q_2$:

Supposons qu'on ait trouvé deux couples (P_1, P_2) et (P'_1, P'_2) de polynômes vérifiant les conditions du théorème.

On obtient par différence :

$$\frac{P_1 - P'_1}{Q_1} = \frac{P'_2 - P_2}{Q_2} \quad \text{ou} \quad Q_2(P_1 - P'_1) = Q_1(P'_2 - P_2).$$

Le polynôme Q_2 est premier avec Q_1 et il divise $Q_1(P'_2 - P_2)$; d'après le Théorème de Gauss, Q_2 divise $P'_2 - P_2$. Or $\deg(P'_2 - P_2) \leq \sup(\deg(P'_2), \deg(P_2)) < \deg(Q_2)$, donc $P'_2 - P_2 = 0$. Q_2 n'étant pas nul et $K[X]$ étant intègre, on a $P_1 - P'_1 = 0$, d'où l'unicité de la décomposition.

c) Existence et unicité dans le cas général :

Raisonnons par récurrence sur n , le résultat étant vrai pour $n = 2$. Supposons le résultat démontré pour $(n-1)$ facteurs et soit $Q = Q_1 Q_2 \dots Q_n$ un produit de n polynômes premiers entre eux deux à deux. Q_1 étant premier avec $Q_2, \dots, Q_n$ est premier avec le produit $Q_2 Q_3 \dots Q_n$. En appliquant la première partie de la démonstration, on conclut qu'il existe un couple unique (P_1, R_1) de polynômes vérifiant $\deg(P_1) < \deg(Q_1)$, $\deg(R_1) < \deg(Q_2 \dots Q_n)$ et tels que :

$$\frac{P}{Q} = \frac{P_1}{Q_1} + \frac{R_1}{Q_2 Q_3 \dots Q_n}.$$

Appliquons alors l'hypothèse de récurrence à la fraction rationnelle $\frac{R_1}{Q_2 \dots Q_n}$; on obtient d'une manière et d'une seule, $n - 1$ polynômes $P_2, \dots, P_n$ tels que :

$$\frac{P}{Q} = \frac{P_1}{Q_1} + \frac{P_2}{Q_2} + \dots + \frac{P_n}{Q_n}$$

avec $\deg(P_i) < \deg(Q_i)$ pour tout $i \in \{1, \dots, n\}$ ce qui achève la preuve.

5.11.1.3. Lemme

Soit $F = \frac{P}{Q^n}$, $n \in \mathbb{N}^*$, une fraction rationnelle telle que $\deg(P) < \deg(Q^n)$.
Il existe une famille unique de polynômes $P_1, P_2, \dots, P_n$ telle que :

$$F = \sum_{i=1}^n \frac{P_i}{Q^i} \text{ et } \deg(P_i) < \deg(Q) \text{ pour tout } i \in \{1, \dots, n\}.$$

Démonstration. Raisonnons par récurrence sur l'entier n .

Le lemme est évident si $n = 1$. Supposons-le démontré jusqu'à l'ordre $n - 1$ et démontrons-le à l'ordre n .

Effectuons la division euclidienne de P par Q : il existe des polynômes B et P_n dans $K[X]$ tels que

$$P = QB + P_n \text{ avec } \deg(P_n) < \deg(Q).$$

$$\text{Alors } \frac{P}{Q^n} = \frac{B}{Q^{n-1}} + \frac{P_n}{Q^n}.$$

Cette décomposition est unique à cause de l'unicité de la division euclidienne dans $K[X]$. Comme $\frac{B}{Q^{n-1}} = \frac{P - P_n}{Q^n}$ on a, d'après le Théorème 5.10.1.4, $\deg(B) < \deg(Q^{n-1})$. On peut donc appliquer l'hypothèse de récurrence à $\frac{B}{Q^{n-1}}$, ce qui achève la démonstration du lemme.

5.11.1.4. Définition

Toute fraction rationnelle de la forme $\frac{A}{B^\alpha}$, où B est un polynôme irréductible de $K[X]$, α un entier supérieur ou égal à 1 et où $\deg(A) < \deg(B)$, s'appelle un **élément simple**.

Le résultat fondamental de ce paragraphe est le théorème suivant qui s'obtient en utilisant les lemmes précédents.

5.11.1.5. Théorème

Soit F une fraction rationnelle de $K(X)$ écrite sous sa forme irréductible $\frac{P}{Q}$, Q étant un polynôme de degré au moins égal à 1. Si $Q = \lambda A^\alpha B^\beta \dots L^\gamma$ est la décomposition de Q en facteurs irréductibles, il existe une famille unique $E, A_1, \dots, A_\alpha, B_1, \dots, B_\beta, \dots, L_1, \dots, L_\gamma$ de polynômes de $K[X]$ tels que

$$\begin{aligned} \frac{P}{Q} = E + \frac{A_1}{A} + \frac{A_2}{A^2} + \dots + \frac{A_\alpha}{A^\alpha} \\ + \frac{B_1}{B} + \frac{B_2}{B^2} + \dots + \frac{B_\beta}{B^\beta} \\ + \frac{L_1}{L} + \frac{L_2}{L^2} + \dots + \frac{L_\gamma}{L^\gamma} \end{aligned}$$

et $\deg(A_i) < \deg(A)$, $\deg(B_i) < \deg(B)$, $\dots$, $\deg(L_i) < \deg(L)$ pour tout i .

Calculer ces polynômes c'est décomposer la fraction rationnelle F en éléments simples.

Démonstration. Puisque les polynômes $A, B, \dots, L$ sont premiers entre eux deux à deux il en est de même de $A^\alpha, B^\beta, \dots, L^\gamma$. D'après les lemmes 5.11.1.1 et 5.11.1.2, il existe une famille unique de polynômes $E, C_1, \dots, C_\gamma$ tels que

$$\frac{P}{Q} = E + \frac{C_1}{A^\alpha} + \dots + \frac{C_\gamma}{L^\gamma}$$

avec $\deg\left(\frac{C_1}{A^\alpha}\right) < 0, \dots, \deg\left(\frac{C_\gamma}{L^\gamma}\right) < 0$. Il suffit maintenant d'appliquer le

Lemme 5.11.1.3 à chacun des termes $\frac{C_1}{A^\alpha}, \dots, \frac{C_\gamma}{L^\gamma}$ pour obtenir la décomposition annoncée.

5.11.2. DÉCOMPOSITION EN ÉLÉMENTS SIMPLES D'UNE FRACTION RATIONNELLE SUR $\mathbb{C}$

Nous savons que les polynômes irréductibles de $\mathbb{C}[X]$ sont de la forme $X - a$, où $a \in \mathbb{C}$. Le Théorème 5.11.1.5 se simplifie de la manière suivante :

5.11.2.1. Théorème

Soit F une fraction rationnelle de $\mathbb{C}(X)$ écrite sous forme irréductible $\frac{P}{Q}$ telle que $\deg(Q) \geq 1$. Si $Q(X) = \lambda(X - a_1)^{\alpha_1} \dots (X - a_n)^{\alpha_n}$ est la décomposition de Q en facteurs irréductibles, il existe un polynôme unique E et une unique famille de scalaires $(b_{ij})_{1 \leq i \leq n, 1 \leq j \leq \alpha_i}$ tels que :

$$F = E + \sum_{i=1}^n \left(\sum_{j=1}^{\alpha_i} \frac{b_{ij}}{(X - a_i)^j} \right).$$

Démonstration. Les polynômes $A, B, \dots, L$ du Théorème 5.11.1.5 sont ici $(X - a_1), (X - a_2), \dots, (X - a_n)$; comme ces polynômes sont du premier degré, les polynômes $A_i, B_j \dots L_k$ de la théorie générale, qui vérifient $\deg(A_i) < \deg(A)$, sont des constantes.

5.11.2.2. Définition

Les notations étant celles du Théorème 5.11.2.1, $\sum_{j=1}^{\infty} \frac{b_{ij}}{(X - a_i)^j}$ s'appelle la **partie polaire (ou partie principale) de F relative au pôle a_i .**

Méthode pratique de calcul des b_{ij}

Étant donné une fraction rationnelle F de $\mathbb{C}(X)$, le Théorème 5.11.2.1 affirme l'existence et l'unicité de la décomposition en éléments simples de F . La partie entière (qui est non nulle si et seulement si $\deg(F) \geq 0$) s'obtient en effectuant la division euclidienne du numérateur de F par son dénominateur. Nous supposons dans la suite que cette opération a été faite et nous ne considérons plus que des fractions rationnelles de degré strictement négatif.

a) Partie polaire relative à un pôle multiple :

Soit $F(X) = \frac{P(X)}{(X - a)^k Q(X)}$ une fraction rationnelle de $\mathbb{C}(X)$ de degré strictement négatif admettant a pour pôle d'ordre k (on a donc $P(a) \neq 0$ et $Q(a) \neq 0$).

D'après le Théorème 5.11.2.1, la décomposition de F s'écrit :

$$F(X) = \frac{b_1}{X - a} + \frac{b_2}{(X - a)^2} + \dots + \frac{b_k}{(X - a)^k} + \frac{P_1(X)}{Q(X)}.$$

D'où

$$P(X) = (b_k + b_{k-1}(X - a) + \dots + b_1(X - a)^{k-1}) Q(X) + (X - a)^k P_1(X).$$

Prenons alors le polynôme $X - a = Y$ comme nouvelle indéterminée; on obtient

$$P(a + Y) = (b_k + b_{k-1}Y + \dots + b_1Y^{k-1}) Q(a + Y) + Y^k P_1(a + Y).$$

La partie polaire relative au pôle a apparaît ainsi comme le quotient de la division suivant les puissances croissantes de $P(a + Y)$ par $Q(a + Y)$ à l'ordre $k - 1$.

b) Partie polaire relative à un pôle simple :

Si $\frac{P}{Q}$ est une fraction rationnelle irréductible de $\mathbb{C}(X)$ admettant le nombre a pour pôle simple, la partie polaire de cette fraction rationnelle relative au pôle a est $\frac{P(a)}{Q'(a)} \cdot \frac{1}{X - a}$.

En posant $Q(X) = (X - a) Q_1(X)$, on a une décomposition de la forme

$$\frac{P(X)}{(X - a)Q_1(X)} = \frac{A}{X - a} + \frac{R(X)}{Q_1(X)} \quad \text{avec } Q_1(a) \neq 0.$$

Multiplions les deux membres par $X - a$ puis faisons $X = a$; on obtient $A = \frac{P(a)}{Q_1(a)}$.

De l'égalité $Q(X) = (X - a) Q_1(X)$, on obtient en prenant les polynômes dérivés, $Q'(X) = (X - a) Q_1'(X) + Q_1(X)$. D'où $Q'(a) = Q_1(a)$ et $A = \frac{P(a)}{Q'(a)}$.

5.11.2.3. Exemple

Décomposer la fraction rationnelle $F(X) = \frac{1}{X(X + 1)(X - 1)^3}$ en éléments simples sur $\mathbb{C}$.

Le degré du numérateur étant strictement inférieur à celui du dénominateur, la partie entière de la décomposition est nulle.

La décomposition est de la forme :

$$F(X) = \frac{A}{X} + \frac{B}{X + 1} + \frac{a}{X - 1} + \frac{b}{(X - 1)^2} + \frac{c}{(X - 1)^3}.$$

Multiplions les deux membres de cette égalité par X puis faisons $X = 0$; il vient $A = -1$.

En multipliant de même les deux membres de l'égalité par $X + 1$ puis faisant $X = -1$, nous obtenons $B = \frac{1}{8}$.

Pour obtenir la partie principale relative au pôle triple 1, nous posons $Y = X - 1$. D'où $F(1 + Y) = \frac{1}{(2 + 3Y + Y^2)Y^3}$.

La division de 1 par $2 + 3Y + Y^2$ suivant les puissances croissantes à l'ordre 2 donne :

$$1 = (2 + 3Y + Y^2) \left(\frac{1}{2} - \frac{3}{4} Y + \frac{7}{8} Y^2 \right) + Y^3 \left(-\frac{15}{8} - \frac{7}{8} Y \right),$$

soit en revant à $F(1 + Y)$:

$$\frac{1}{(2 + 3Y + Y^2) Y^3} = \frac{\frac{1}{2}}{Y^3} - \frac{\frac{3}{4}}{Y^2} + \frac{\frac{7}{8}}{Y} + \frac{-\frac{15}{8} - \frac{7}{8} Y}{2 + 3Y + Y^2}$$

et en revant enfin à X ,

$$F(X) = \frac{1}{2(X - 1)^3} - \frac{3}{4(X - 1)^2} + \frac{7}{8(X - 1)} + \frac{-\frac{7}{8} X - 1}{X(X + 1)}.$$

La décomposition cherchée est donc

$$F(X) = -\frac{1}{X} + \frac{1}{8(X+1)} + \frac{7}{8(X-1)} - \frac{3}{4(X-1)^2} + \frac{1}{2(X-1)^3}.$$

5.11.3. DÉCOMPOSITION EN ÉLÉMENTS SIMPLES D'UNE FRACTION RATIONNELLE SUR $\mathbb{R}$

Dans $\mathbb{R}[X]$, les polynômes irréductibles sont les polynômes du premier degré et ceux du second degré à discriminant négatif. Le Théorème 5.11.1.5 prend alors la forme suivante :

5.11.3.1. Théorème

Soit F une fraction rationnelle de $\mathbb{R}(X)$ admettant un représentant irréductible $\frac{P}{Q}$ tel que $\deg(Q) \geq 1$.

Si $Q(X) = \lambda \prod_{i=1}^n (X - a_i)^{\alpha_i} \prod_{j=1}^m (X^2 + p_j X + q_j)^{\beta_j}$ est la décomposition de Q en polynômes irréductibles, il existe un polynôme unique E et des familles uniques de nombres réels $(A_{ij})_{1 \leq i \leq n, 1 \leq j \leq \alpha_i}$, $(B_{kr})_{1 \leq k \leq m, 1 \leq r \leq \beta_k}$ et $(C_{kr})_{1 \leq k \leq m, 1 \leq r \leq \beta_k}$ tels que :

$$F = E + \sum_{i=1}^n \left(\sum_{j=1}^{\alpha_i} \frac{A_{ij}}{(X - a_i)^j} \right) + \sum_{k=1}^m \left(\sum_{r=1}^{\beta_k} \frac{B_{kr} X + C_{kr}}{(X^2 + p_k X + q_k)^r} \right).$$

5.11.3.2. Définition

Dans la décomposition en éléments simples d'une fraction rationnelle de $\mathbb{R}(X)$, une fraction de la forme $\frac{A_{ij}}{(X - a)^j}$ s'appelle un élément simple de première espèce, une fraction de la forme $\frac{B_{kr} X + C_{kr}}{(X^2 + p_k X + q_k)^r}$ s'appelle un élément simple de deuxième espèce.

Méthode pratique de décomposition

a) Pour la recherche de la partie entière et les éléments simples de première espèce, tout ce qui a été dit au n° 5.11.2 reste valable.

b) Pour les éléments simples de deuxième espèce, les méthodes suivantes peuvent être utilisées :

- On écrit la décomposition de F à l'aide de coefficients indéterminés et on détermine ces coefficients par des considérations numériques particulières ;

l'examen de la parité de la fraction rationnelle considérée peut simplifier les calculs.

- On utilise la décomposition dans $\mathbb{C}(X)$ puis en regroupant les parties polaires relatives aux pôles conjugués, on obtient la décomposition dans $\mathbb{R}(X)$.

- Si F n'admet que deux pôles complexes conjugués, on procède par divisions successives.

Pour conclure ce chapitre, examinons quelques exemples qui mettent en œuvre les méthodes qui viennent d'être expliquées.

5.11.3.3. Exemple

Décomposer la fraction rationnelle $F(X) = \frac{1}{(X^2 - 1)(X^2 + 1)^2}$ en éléments simples sur $\mathbb{R}$.

On a $P = 1$, $Q = (X^2 - 1)(X^2 + 1)^2$; $\deg(P) < \deg(Q)$, donc la partie entière est nulle. F s'écrit

$$F(X) = \frac{A}{X-1} + \frac{B}{X+1} + \frac{aX+b}{X^2+1} + \frac{cX+d}{(X^2+1)^2}.$$

Nous remarquons que F est paire, c'est-à-dire $F(-X) = F(X)$.

Comme

$$F(-X) = -\frac{A}{X+1} - \frac{B}{X-1} + \frac{-aX+b}{X^2+1} + \frac{-cX+d}{(X^2+1)^2},$$

l'unicité de la décomposition nous permet d'affirmer que $B = -A$, $a = -a$, $c = -c$, donc $a = c = 0$.

Pour déterminer A , multiplions les deux membres par $X-1$ et après simplification, donnons à X la valeur 1; nous trouvons $A = \frac{1}{8}$. On utilise une méthode analogue pour déterminer d ; on trouve $d = -\frac{1}{2}$.

Il ne reste plus qu'un coefficient inconnu; il suffit de donner à X une valeur particulière, par exemple 0, pour obtenir $b = -\frac{1}{4}$.

La décomposition cherchée est donc

$$\frac{1}{(X^2 - 1)(X^2 + 1)^2} = \frac{1}{8(X-1)} - \frac{1}{8(X+1)} - \frac{1}{4(X^2 + 1)} - \frac{1}{2(X^2 + 1)^2}.$$

5.11.3.4. Exemple

Décomposer la fraction rationnelle $F(X) = \frac{X}{(X+1)(X^2+1)}$ en éléments simples sur $\mathbb{R}$.

La partie entière étant nulle, la décomposition sur $\mathbb{C}$ est de la forme

$$F(X) = \frac{A}{X+1} + \frac{B}{X+i} + \frac{C}{X-i}.$$

Pour obtenir A , on multiplie les deux membres de cette égalité par $X+1$ puis on fait $X = -1$, d'où $A = \frac{-1}{2}$. On obtient de même $B = \frac{1+i}{4}$, $C = \frac{1-i}{4}$, ce qui donne

$$F(X) = \frac{-1}{2(X+1)} + \frac{1+i}{4(X+i)} + \frac{1-i}{4(X-i)}.$$

En regroupant les parties polaires relatives aux pôles conjugués et en réduisant au même dénominateur, nous obtenons la décomposition sur $\mathbb{R}$:

$$\begin{aligned} F(X) &= \frac{-1}{2(X+1)} + \frac{(1+i)(X-i) + (1-i)(X+i)}{4(X^2+1)} \\ &= \frac{-1}{2(X+1)} + \frac{X+1}{2(X^2+1)}. \end{aligned}$$

5.11.3.5. Exemple

Décomposer la fraction rationnelle $\frac{2X^4 + X^3 + 1}{(X^2 + X + 1)^3}$ en éléments simples sur $\mathbb{R}$.

La partie entière est nulle. Le dénominateur admet deux racines complexes conjuguées.

La division euclidienne du numérateur par $X^2 + X + 1$ donne :

$$2X^4 + X^3 + 1 = (X^2 + X + 1)(2X^2 - X - 1) + 2X + 2,$$

d'où

$$\frac{2X^4 + X^3 + 1}{(X^2 + X + 1)^3} = \frac{2X^2 - X - 1}{(X^2 + X + 1)^2} + \frac{2X + 2}{(X^2 + X + 1)^3}.$$

La division euclidienne du quotient $2X^2 - X - 1$ par $X^2 + X + 1$ s'écrit :

$$2X^2 - X - 1 = 2(X^2 + X + 1) - 3X - 3.$$

On en déduit :
$$\frac{2X^2 - X - 1}{(X^2 + X + 1)^2} = \frac{2}{X^2 + X + 1} + \frac{-3X - 3}{(X^2 + X + 1)^2}.$$

La décomposition cherchée est donc :

$$\frac{2X^4 + X^3 + 1}{(X^2 + X + 1)^3} = \frac{2}{X^2 + X + 1} + \frac{-3X - 3}{(X^2 + X + 1)^2} + \frac{2X + 2}{(X^2 + X + 1)^3}.$$

Chapitre 6 : ESPACES VECTORIELS

La notion d'espace vectoriel est l'une des notions les plus importantes en mathématiques et dans les applications des mathématiques aux autres sciences. En physique et dans les sciences de l'ingénieur, les espaces vectoriels constituent un outil indispensable pour représenter certaines quantités : forces, vitesses, état d'un système en mécanique quantique, etc.

Dans ce chapitre, nous allons introduire la notion abstraite d'espace vectoriel puis nous étudierons les principales conséquences des axiomes définissant la structure d'espace vectoriel. Dans tout ce qui suit, K désignera toujours un corps commutatif.

6.1. Définition d'un espace vectoriel

6.1.1. DÉFINITIONS

6.1.1.1. Définition

*Soient E un groupe commutatif noté additivement, et K un corps commutatif. On dit que E est un **espace vectoriel** sur K ou un **K -espace vectoriel**, s'il existe une loi externe, de domaine K , associant à tout élément (λ, x) de $K \times E$ l'élément de E noté $\lambda \cdot x$ ou λx avec les propriétés suivantes :*

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

$$(\lambda + \mu) \cdot x = \lambda \cdot x + \mu \cdot x$$

$$\lambda \cdot (\mu \cdot x) = (\lambda\mu) \cdot x$$

$$1 \cdot x = x$$

quels que soient $x, y \in E$ et $\lambda, \mu \in K$, 1 désignant l'élément unité du corps K .

Les éléments de E s'appellent les **vecteurs** ; ceux de K s'appellent les **scalaires**. Nous désignerons en général les vecteurs par des lettres latines minuscules, et les scalaires par des lettres grecques minuscules. On parlera d'**espace vectoriel réel** si $K = \mathbb{R}$, d'**espace vectoriel complexe** si $K = \mathbb{C}$.

L'élément neutre du groupe $(E, +)$ est noté 0_E ou simplement 0 et est appelé le **vecteur nul** ou l'**origine** de E . Si $E = \{0\}$, on dit que E est un **espace vectoriel nul**.

6.1.1.2. Remarque

Si au lieu de prendre un corps K , on prend un anneau A (non commutatif), on dit que E , muni des deux lois de composition précédentes est un **module à gauche** sur l'anneau A ou encore un **A -module à gauche** (l'ordre d'écriture λx pour le produit externe de $(\lambda, x) \in K \times E$ doit être respecté).

On définit de même un **A -module à droite**. C'est un ensemble E sur lequel on a défini :

a) une loi de groupe abélien, notée $(x, y) \mapsto x + y$;

b) une application $(x, \lambda) \mapsto x \cdot \lambda$ de $E \times A$ dans E , vérifiant les propriétés suivantes :

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

$$x \cdot (\lambda + \mu) = x \cdot \lambda + x \cdot \mu$$

$$(x \cdot \lambda) \cdot \mu = x \cdot (\lambda \mu)$$

$$x \cdot 1 = x$$

quels que soient $x, y \in E$ et $\lambda, \mu \in A$.

Si A est un anneau commutatif, ces deux notions de module à gauche et de module à droite coïncident. On dit que E est un **A -module**.

6.1.2. RÈGLES DE CALCUL DANS UN ESPACE VECTORIEL

L'opposé du vecteur x dans le groupe $(E, +)$ est noté $-x$; la somme du vecteur x et de l'opposé du vecteur y est notée $x - y$ et est appelée **différence** de x et y .

a) On a

$$(6.1.2.1) \quad \lambda(x - y) = \lambda x - \lambda y$$

quels que soient $x, y \in E$ et $\lambda \in K$.

En effet,

$$\lambda(x - y) + \lambda y = \lambda[(x - y) + y] = \lambda x.$$

D'où la formule (6.1.2.1) Si dans cette formule, on fait $x = y$, il vient

$$\lambda \cdot 0 = 0$$

pour tout $\lambda \in K$.

b) On a

$$(6.1.2.2) \quad (\lambda - \mu)x = \lambda x - \mu x$$

quels que soient $x \in E$ et $\lambda, \mu \in K$.

En effet,

$$(\lambda - \mu)x + \mu x = [(\lambda - \mu) + \mu]x = \lambda x.$$

On en déduit la formule (6.1.2.2). En faisant $\lambda = \mu$ dans cette formule, on obtient

$$0 \cdot x = 0$$

pour tout vecteur x .

c) Quel que soit $x \in E$ et quel que soit $\lambda \in K$, on a

$$(6.1.2.3) \quad (-\lambda)x = \lambda(-x) = -(\lambda x).$$

On a en effet

$$\lambda x + \lambda(-x) = \lambda(x - x) = \lambda \cdot 0 = 0$$

et

$$\lambda x + (-\lambda)x = (\lambda + (-\lambda))x = (\lambda - \lambda)x = 0 \cdot x = 0.$$

d) Pour tout $x \in E$ et pour tout $\lambda \in K$,

$$\lambda x = 0 \iff \lambda = 0 \text{ ou } x = 0.$$

D'après b) et c), on a $\lambda \cdot 0 = 0 \cdot x = 0$.

Réciproquement, supposons que $\lambda \cdot x = 0$ et $\lambda \neq 0$. Alors λ^{-1} existe et on a

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

Or

$$\lambda^{-1}(\lambda x) = (\lambda^{-1}\lambda)x = 1 \cdot x = x,$$

donc $x = 0$.

6.1.3. EXEMPLE D'ESPACES VECTORIELS

6.1.3.1. Exemple

L'ensemble E des vecteurs libres du plan (ou de l'espace) de la géométrie élémentaire est un espace vectoriel réel pour les lois de composition $(\vec{v}_1, \vec{v}_2) \mapsto \vec{v}_1 + \vec{v}_2$ de $E \times E$ dans E et $(\lambda, \vec{v}) \mapsto \lambda \vec{v}$ de $\mathbb{R} \times E$ dans E . Cet exemple est à l'origine de l'étude abstraite des espaces vectoriels.

6.1.3.2. Exemple

Soient K un corps commutatif et K' un sous-corps de K . Muni de la loi de groupe abélien et de la loi externe $(\lambda, \mu) \mapsto \lambda \cdot \mu$ de $K' \times K$ dans K , K est un K' -espace vectoriel. Si on prend en particulier $K' = K$, on voit que tout corps K est un espace vectoriel sur lui-même, la loi externe étant la multiplication interne. Ainsi $\mathbb{R}$ est un espace vectoriel réel. $\mathbb{C}$ est un espace vectoriel réel ou complexe.

6.1.3.3. Exemple

Soit E un espace vectoriel sur K et soit K' un sous-corps de K . En considérant la restriction de l'opération externe de $K \times E$ à $K' \times E$, on obtient une structure d'espace vectoriel sur le corps K' . Ainsi, tout $\mathbb{C}$ -espace vectoriel est aussi un $\mathbb{R}$ -espace vectoriel. Ce procédé est appelé **restriction du corps des scalaires**. Le procédé inverse, appelé **extension du corps des scalaires** ne sera pas exposé dans ce cours.

6.1.3.4. Exemple

Soient $E_1, E_2, \dots, E_n, n$ espaces vectoriels sur le même corps K . Nous allons munir l'ensemble produit $E_1 \times E_2 \times \dots \times E_n$ d'une structure d'espace vectoriel sur K . Pour cela, si $(x_1, x_2, \dots, x_n) \in E_1 \times E_2 \times \dots \times E_n$, $(y_1, y_2, \dots, y_n) \in E_1 \times E_2 \times \dots \times E_n$ et $\lambda \in K$, posons :

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

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

On vérifie facilement qu'on obtient ainsi une structure de K -espace vectoriel.

On dit que $E_1 \times E_2 \times \dots \times E_n$ est l'**espace vectoriel produit** des n espaces $E_1, E_2, \dots, E_n$.

Si en particulier $E_1 = E_2 = \dots = E_n = E$, on obtient le K -espace vectoriel E^n . Par exemple $K^n, n \geq 2$, est un espace vectoriel sur K puisque K est un espace vectoriel sur lui-même.

6.1.3.5. Exemple

Soient X un ensemble quelconque F un K -espace vectoriel et soit E l'ensemble de toutes les applications $f : X \longrightarrow F$.

Si $f, g \in E$ et $\lambda \in K$, définissons $f + g$ et λf en posant

$$(f + g)(x) = f(x) + g(x)$$

$$(\lambda f)(x) = \lambda f(x)$$

pour tout $x \in X$.

On vérifie facilement que E est un K -espace vectoriel lorsqu'on le munit des deux lois de composition $(f, g) \longmapsto f + g$ et $(\lambda, f) \longmapsto \lambda f$ ainsi définies.

6.1.3.6. Exemple

L'ensemble E des fonctions réelles continues d'une variable réelle, muni des lois de composition définies dans l'Exemple 6.1.3.5 est un espace vectoriel réel. On démontre en effet dans le cours d'analyse que si f et g sont partout continues et si $\lambda \in \mathbb{R}$, alors $f + g$ et λf sont continues partout, donc sont dans E .

De même l'ensemble des fonctions réelles d'une variable réelle, dérivables en un point $x_0 \in \mathbb{R}$ (ou dérivables en tout point) est un espace vectoriel réel.

6.1.3.7. Exemple

Soit K un corps commutatif. L'ensemble $K[X]$ des polynômes à une indéterminée X , à coefficients dans K , est un K -espace vectoriel lorsqu'on le munit de l'addition des polynômes : $(P, Q) \mapsto P + Q$ et de la multiplication par un scalaire $(\lambda, P) \mapsto \lambda \cdot P$, définies au Chapitre 5.

6.2. Sous-espaces vectoriels

6.2.1. DÉFINITIONS. EXEMPLES

6.2.1.1. Définition

Soit E un K -espace vectoriel. On dit qu'une partie F de E est un **sous-espace vectoriel** ou simplement un **sous-espace** de E (en abrégé : s. e. v.) si F vérifie les conditions suivantes :

- a) F n'est pas vide.
- b) Les relations $x \in F$ et $y \in F$ entraînent $\lambda x + \mu y \in F$ quels que soient $\lambda, \mu \in K$.

Faisons $\lambda = 1$ et $\mu = -1$ dans la condition b) ; alors les relations $x \in E$ et $y \in F$ impliquent $x - y \in F$; donc, puisque $F \neq \emptyset$, F est un sous-groupe du groupe additif E . Si dans la condition b) on fait maintenant $\mu = 0$, on voit que les relations $x \in F$ et $\lambda \in K$ entraînent $\lambda x \in F$. On peut donc définir une application $(\lambda, x) \mapsto \lambda x$ de $K \times F$ dans F et F , muni des lois induites par les lois $(x, y) \mapsto x + y$ et $(\lambda, x) \mapsto \lambda x$ est un K -espace vectoriel, ce qui justifie la locution « sous-espace vectoriel ».

6.2.1.2. Exemple

Soit E un K -espace vectoriel. Les parties $\{0\}$ et E sont des sous-espaces vectoriels de E , dits **triviaux**. Tout sous-espace différent de ces deux sous-espaces est dit **propre**.

6.2.1.3. Exemple

Soit $n \in \mathbb{N}$. L'ensemble $K_n[X]$ des polynômes à une indéterminée X , à coefficients dans le corps commutatif K , de degré $\leq n$, est un sous-espace vectoriel de $K[X]$.

En effet, quels que soient les polynômes P et Q et les scalaires λ et μ , on a

$$\deg(\lambda P + \mu Q) \leq \sup(\deg(P), \deg(Q)) .$$

Donc si $\deg(P) \leq n$ et $\deg(Q) \leq n$, alors $\deg(\lambda P + \mu Q) \leq n$; par suite $K_n[X]$ est bien un sous-espace vectoriel de $K[X]$.

6.2.1.4. Exemple

L'ensemble $\mathcal{C}(I, \mathbb{C})$ des fonctions continues sur un intervalle I de la droite réelle et à valeurs complexes, est un sous-espace vectoriel de l'espace vectoriel $\mathcal{F}(I, \mathbb{C})$ de toutes les fonctions définies sur I et à valeurs complexes.

De même l'ensemble $\mathcal{D}(I, \mathbb{C})$ des fonctions dérivables sur I et à valeurs complexes est un sous-espace vectoriel de $\mathcal{F}(I, \mathbb{C})$.

6.2.1.5. Exemple

Soit E un K -espace vectoriel et soit x un vecteur non nul de E . L'ensemble Kx des vecteurs de la forme λx , où $\lambda \in K$, est un sous-espace vectoriel de E . On l'appelle la **droite vectorielle engendrée par x** .

6.2.1.6. Exemple

Soit E un K -espace vectoriel et soient x et y deux vecteurs de E tels que $x \neq 0$ et $y \notin Kx$. L'ensemble des vecteurs de la forme $\lambda x + \mu y$, où $\lambda, \mu \in K$, est un sous-espace vectoriel de E . On l'appelle le **plan vectoriel engendré par x et y** .

6.2.2. INTERSECTION DE SOUS-ESPACES VECTORIELS. SOUS-ESPACE ENGENDRÉ PAR UNE PARTIE D'UN ESPACE VECTORIEL

6.2.2.1. Théorème

Soit $(F_i)_{i \in I}$ une famille de sous-espaces vectoriels d'un K -espace vectoriel E . Alors $F = \bigcap_{i \in I} F_i$ est un sous-espace vectoriel de E .

Démonstration. Les F_i étant des sous-groupes du groupe additif E , on a $0 \in F_i$ pour tout $i \in I$; donc $0 \in F$ et $F \neq \emptyset$. Soient x et y des éléments de F ; alors pour tout $i \in I$, x et y appartiennent à F_i . Donc quels que soient les scalaires λ et μ de K , $\lambda x + \mu y \in F_i$ puisque F_i est un sous-espace vectoriel de E . Par suite $\lambda x + \mu y \in F$ et F est bien un sous-espace vectoriel de E .

On notera que la réunion d'une famille de sous-espaces n'est pas toujours un sous-espace (considérer par exemple deux droites distinctes dans le plan).

Mais si la réunion de deux sous-espaces quelconques de la famille est toujours contenue dans un élément de la famille, alors cette réunion est un sous-espace vectoriel.

Soit A une partie d'un K -espace vectoriel E . Il existe des sous-espaces vectoriels de E contenant A (par exemple E lui-même). L'intersection de tous ces sous-espaces vectoriels est le plus petit sous-espace vectoriel (pour l'inclusion) contenant A . On l'appelle le sous-espace vectoriel engendré par A et on le note $\text{Vect}(A)$.

Si $\text{Vect}(A) = E$, on dit que A est une partie génératrice de E .

6.2.3. ESPACES VECTORIELS QUOTIENTS

Soit E un espace vectoriel sur le corps K et soit F un sous-espace vectoriel de E . La relation

$$x \equiv y \pmod{F} \iff x - y \in F$$

est une relation d'équivalence sur E (à vérifier!) compatible avec la loi de groupe de E , c'est-à-dire :

$$x \equiv y \text{ et } x' \equiv y' \pmod{F} \implies x + x' \equiv y + y' \pmod{F}.$$

Dans l'ensemble quotient, noté E/F , des classes d'équivalence $\bar{x} = x + F$ modulo F , définissons une loi interne en posant

$$\bar{x} + \bar{y} = \overline{x + y}.$$

Cette loi confère à E/F une structure de groupe abélien (voir le Théorème 3.4.2.3).

Posons de plus

$$\lambda \bar{x} = \overline{\lambda x}$$

quels que soient $\bar{x} \in E/F$ et $\lambda \in K$.

Cette définition sera justifiée si nous montrons que la classe obtenue ne dépend que de $\bar{x}$ et de λ et non du représentant x choisi. Soient donc x et y deux représentants d'une même classe ; on a $x - y \in F$. Alors $\lambda(x - y) = \lambda x - \lambda y \in F$ puisque F est un sous-espace vectoriel de E . On a donc bien $\lambda x \equiv \lambda y \pmod{F}$, c'est-à-dire $\lambda \bar{x} = \overline{\lambda y}$.

On vérifie facilement que E/F , muni de ces deux lois, est un K -espace vectoriel. On l'appelle l'espace vectoriel quotient de E par F .

6.2.4. SOMME DE SOUS-ESPACES VECTORIELS

Nous avons observé que la réunion de deux sous-espaces vectoriels d'un K -espace vectoriel E n'est pas, en général, un sous-espace vectoriel de E . On est donc amené à introduire une autre opération : la somme d'une famille de sous-espaces vectoriels.

6.2.4.1. Définition

Soit $E_1, \dots, E_n$ une famille finie de sous-espaces vectoriels d'un K -espace vectoriel E . On appelle **somme des sous-espaces** $E_1, \dots, E_n$, et on note $E_1 + \dots + E_n$, ou $\sum_{i=1}^n E_i$, l'ensemble des éléments x de E qui sont de la forme

$$x = x_1 + \dots + x_n \text{ où } x_1 \in E_1, \dots, x_n \in E_n.$$

Plus généralement, si $(E_i)_{i \in I}$ est une famille de sous-espaces vectoriels de E , on appelle **somme des sous-espaces** E_i et on note $\sum_{i \in I} E_i$, l'ensemble des sommes finies de la forme $\sum_{i \in J} x_i$ où J est une partie finie de I et où pour tout $i \in J, x_i \in E_i$.

6.2.4.2. Théorème

Soient E un K -espace vectoriel et $E_1, \dots, E_n$ ($n \geq 2$), n sous-espaces vectoriels de E . La somme $E_1 + \dots + E_n$ est le sous-espace vectoriel engendré par $E_1 \cup E_2 \cup \dots \cup E_n$.

Démonstration. Montrons d'abord que $\mathcal{S} = E_1 + \dots + E_n$ est un sous-espace vectoriel de E .

Il est clair que $0 \in \mathcal{S}$ donc $\mathcal{S} \neq \emptyset$. Soient x et y deux éléments de $\mathcal{S}$; on a :

$$x = x_1 + \dots + x_n \text{ et } y = y_1 + \dots + y_n$$

avec $x_i \in E_i$ et $y_i \in E_i$ ($1 \leq i \leq n$).

Quels que soient les scalaires λ et μ ,

$$\lambda x + \mu y = (\lambda x_1 + \mu y_1) + \dots + (\lambda x_n + \mu y_n).$$

Chaque E_i étant un sous-espace vectoriel, $\lambda x_i + \mu y_i \in E_i$ pour tout indice $i \in \{1, \dots, n\}$. Donc $\lambda x + \mu y \in \mathcal{S}$ et $\mathcal{S}$ est bien un sous-espace vectoriel de E .

Montrons maintenant que $\mathcal{S}$ est le plus petit sous-espace vectoriel de E (au sens de l'inclusion) contenant la réunion $E_1 \cup \dots \cup E_n$.

Soit $x_i \in E_i$; on a $x_i = 0 + \dots + 0 + x_i + 0 + \dots + 0$ donc $E_i \subset \mathcal{S}$ pour tout i et par suite

$$E_1 \cup \dots \cup E_n \subset \mathcal{S}.$$

De plus, si F est un sous-espace vectoriel de E contenant $E_1 \cup \dots \cup E_n$, alors pour tout $i \in \{1, \dots, n\}, x_i \in E_i \implies x_i \in F$, donc $x_1 + \dots + x_n \in F$ puisque F est un sous-espace vectoriel de E , et par suite $\mathcal{S} \subset F$.

Ainsi, $\mathcal{S}$ est le plus petit sous-espace vectoriel de E contenant $E_1 \cup \dots \cup E_n$.

Si $x = x_1 + \dots + x_n$ est un vecteur de $E_1 + \dots + E_n$, il n'est pas dit que les vecteurs $x_1, \dots, x_n$ sont uniques. Nous sommes ainsi conduits à examiner les conditions pour que la décomposition soit unique.

Conservons les notations précédentes. Tout élément x de $E_1 + \dots + E_n$ s'écrit sous la forme

$$(6.2.4.1) \quad x = x_1 + \dots + x_n \text{ avec } x_i \in E_i \text{ pour tout } i.$$

Considérons l'application $f : E_1 \times \dots \times E_n \longrightarrow E$ définie par

$$(6.2.4.2) \quad f(x_1, \dots, x_n) = x_1 + \dots + x_n.$$

f est un homomorphisme de groupes abéliens et par définition de la somme $\sum_{i=1}^n E_i$, on a $\text{Im}(f) = \sum_{i=1}^n E_i$.

• Si $\text{Ker}(f) \neq \{0\}$, c'est-à-dire si f n'est pas injectif, il existe un élément $(y_1, \dots, y_n)$ non nul de $E_1 \times \dots \times E_n$ tel que $y_1 + \dots + y_n = 0$. Alors on peut écrire

$$x = (x_1 + y_1) + \dots + (x_n + y_n) \text{ avec } x_i + y_i \in E_i \text{ pour tout } i,$$

ce qui montre que la décomposition (6.2.4.1) n'est pas unique.

• Si $\text{Ker}(f) = \{0\}$ c'est-à-dire si f est injectif, et si x admet une autre décomposition

$$x = x'_1 + \dots + x'_n \text{ avec } x'_i \in E_i \text{ pour tout } i,$$

on a par soustraction :

$$(x_1 - x'_1) + (x_2 - x'_2) + \dots + (x_n - x'_n) = 0,$$

d'où $(x_1 - x'_1, x_2 - x'_2, \dots, x_n - x'_n) \in \text{Ker}(f) = \{0\}$.

Par suite $x_1 = x'_1, x_2 = x'_2, \dots, x_n = x'_n$ et tout $x \in E_1 + \dots + E_n$ s'écrit de manière unique sous la forme $x = x_1 + \dots + x_n$ avec $x_i \in E_i$ pour tout $i \in \{1, \dots, n\}$.

Ces remarques nous amènent à poser la définition suivante :

6.2.4.3. Définition

Soit $(E_i)_{1 \leq i \leq n}$ une famille finie de sous-espaces d'un K -espace vectoriel E . On dit que les E_i sont linéairement indépendants si l'application f est injective. On dit alors que la somme $E_1 + \dots + E_n$ est directe et on écrit $E_1 \oplus \dots \oplus E_n$.

On a le théorème suivant :

6.2.4.4. Théorème

Soit $E_1, \dots, E_n$ une famille finie de sous-espaces d'un K -espace vectoriel E . Les conditions suivantes sont équivalentes :

a) Les sous-espaces $E_1, \dots, E_n$ sont linéairement indépendants.

b) $E_i \cap \sum_{j \neq i} E_j = \{0\}$ pour tout $i \in I = \{1, \dots, n\}$.

Démonstration. a) $\implies$ b) Fixons un entier $i \in I$ et soit $x \in E_i \cap \sum_{j \neq i} E_j$.

On peut écrire x sous la forme $x = \sum_{j \neq i} x_j$, avec $x_j \in E_j$ pour $j \neq i$. Soit

$y = (y_1, \dots, y_n)$ l'élément de $E_1 \times \dots \times E_n$ tel que $y_i = x$ et $y_j = -x_j$ si $j \neq i$. f désignant l'application définie par (6.2.4.2), on a $f(y_1, \dots, y_n) = 0$. Donc $y = 0$, d'où $x = 0$ et b) est démontré.

b) $\implies$ a) Supposons la condition b) vérifiée et soit $x = (x_1, \dots, x_n)$ un élément de $E_1 \times \dots \times E_n$ tel que $f(x_1, \dots, x_n) = x_1 + x_2 + \dots + x_n = 0$. Fixons un entier $i \in I$. On a $x_i = \sum_{j \in I, j \neq i} (-x_j)$. Comme $x_i \in E_i$ et $\sum_{j \in I, j \neq i} (-x_j) \in \sum_{j \neq i} E_j$, on obtient $x_i = 0$; ceci étant vrai pour tout $i \in I$, on a $x = 0$, c'est-à-dire $\text{Ker}(f) = \{0\}$. $\square$

6.2.4.5. Corollaire

Soit E un K -espace vectoriel et soient E_1 et E_2 deux sous-espaces vectoriels de E . Les conditions suivantes sont équivalentes :

a) E_1 et E_2 sont linéairement indépendants.

b) $E_1 \cap E_2 = \{0\}$.

6.2.4.6. Définition

On dit que deux sous-espaces vectoriels d'un K -espace vectoriel E sont supplémentaires dans E s'ils sont linéairement indépendants i.e. si $E = E_1 \oplus E_2$.

D'après ce qui précède, E_1 et E_2 sont supplémentaires dans E si et seulement si $E = E_1 + E_2$ et $E_1 \cap E_2 = \{0\}$.

6.2.4.7. Remarques

a) Un sous-espace vectoriel donné peut admettre plusieurs sous-espaces supplémentaires. Ainsi si $E = \mathbb{R}^2$, considérons les sous-espaces

$$E_1 = \{(x, 0) : x \in \mathbb{R}\}$$

$$E_2 = \{(0, x) : x \in \mathbb{R}\}$$

$$E_3 = \{(x, x) : x \in \mathbb{R}\}$$

On vérifie facilement que E_2 et E_3 sont des sous-espaces supplémentaires de E_1 dans E .

b) On démontre que tout sous-espace vectoriel F d'un espace vectoriel E admet (au moins) un sous-espace supplémentaire. Nous démontrerons ce résultat dans le cas des espaces vectoriels de dimension finie (voir le Théorème 8.1.3.1).

6.2.4.8. Exemple

Soit E l'espace vectoriel formé par les fonctions réelles d'une variable réelle. Soient E_1 l'ensemble des fonctions paires et E_2 l'ensemble des fonctions impaires. Alors E_1 et E_2 sont des sous-espaces vectoriels supplémentaires dans E . En effet, pour toute fonction f , $f(x)$ s'écrit de manière unique sous la forme : $f(x) = h(x) + g(x)$, où

$$h(x) = \frac{f(x) + f(-x)}{2} \in E_1 \text{ et } g(x) = \frac{f(x) - f(-x)}{2} \in E_2.$$

6.2.4.9. Exemple

Soient E_1 et E_2 deux espaces vectoriels sur le même corps K et soit $E = E_1 \times E_2$ l'espace vectoriel produit (voir l'Exemple 6.1.3.4). On peut considérer E_1 et E_2 comme des sous-espaces de E de la façon suivante : on identifie l'élément x_1 de E_1 et l'élément $(x_1, 0)$ de E ; de même λx_1 et $(\lambda x_1, 0)$ seront identifiés pour tout $\lambda \in K$.

Pour E_2 , on aurait de même

$$x_2 = (0, x_2), \quad \lambda x_2 = (0, \lambda x_2) \quad \text{pour tout } \lambda \in K.$$

Si $(x_1, x_2) \in E$, on peut écrire $(x_1, x_2) = (x_1, 0) + (0, x_2)$ avec $(x_1, 0) \in E_1$ et $(0, x_2) \in E_2$. Donc $E = E_1 + E_2$.

D'autre part, si $(x_1, x_2) \in E_1 \cap E_2$, alors la relation $(x_1, x_2) \in E_1$, entraîne $x_2 = 0$ et la relation $(x_1, x_2) \in E_2$ entraîne $x_1 = 0$; donc $(x_1, x_2) = 0$ et par suite $E_1 \cap E_2 = \{0\}$. Par conséquent, E_1 et E_2 sont deux sous-espaces supplémentaires dans E .

6.3. Familles génératrices. Familles libres. Bases

6.3.1. FAMILLES GÉNÉRATRICES

Si E est un K -espace vectoriel, on appelle **système de vecteurs** ou **famille finie de vecteurs** de E , toute partie finie de E .

6.3.1.1. Définition

Soient E un K -espace vectoriel et $\{x_1, \dots, x_n\}$ une famille finie de n vecteurs de E . On appelle **combinaison linéaire** de $x_1, \dots, x_n$ tout vecteur $x \in E$ de la forme :

$$x = \lambda_1 x_1 + \dots + \lambda_n x_n.$$

où $\lambda_1, \dots, \lambda_n$ sont des scalaires appelés **coefficients** de la combinaison linéaire.

Plus généralement, soit A une partie d'un K -espace vectoriel E . On appelle **combinaison linéaire des éléments de A** tout vecteur $x \in E$ possédant la

propriété suivante : il existe un entier naturel n , une famille $x_1, \dots, x_n$ de n vecteurs de A et une famille $\lambda_1, \dots, \lambda_n$ de n scalaires de K tels que

$$x = \lambda_1 x_1 + \dots + \lambda_n x_n.$$

6.3.1.2. Exemple

Dans K^n ($n \geq 2$), considérons les n vecteurs

$$e_1 = (1, 0, \dots, 0), \quad e_2 = (0, 1, 0, \dots, 0), \dots, \quad e_n = (0, \dots, 0, 1).$$

Tout vecteur $x = (\lambda_1, \dots, \lambda_n)$ de l'espace vectoriel K^n est combinaison linéaire de la famille $\{e_1, \dots, e_n\}$ car

$$\begin{aligned} (\lambda_1, \dots, \lambda_n) &= (\lambda_1, 0, \dots, 0) + (0, \lambda_2, \dots, 0) + \dots + (0, \dots, \lambda_n) \\ &= \lambda_1 e_1 + \lambda_2 e_2 + \dots + \lambda_n e_n. \end{aligned}$$

Cette notion de combinaison linéaire va nous permettre de caractériser le sous-espace vectoriel engendré par une partie A d'un espace vectoriel E .

6.3.1.3. Théorème

Soient E un K -espace vectoriel et A une partie non vide de E . Alors, le sous-espace vectoriel engendré par A est l'ensemble des combinaisons linéaires des éléments de A .

Démonstration. Notons A' l'ensemble des combinaisons linéaires des éléments de A . si $x \in A$, on a $0 = x - x \in A'$, donc $A' \neq \emptyset$. Soient

$$x = \lambda_1 x_1 + \dots + \lambda_n x_n \text{ et } y = \mu_1 y_1 + \dots + \mu_m y_m$$

deux vecteurs de A' .

Alors pour tous scalaires λ et μ de K ,

$$\lambda x + \mu y = \lambda \lambda_1 x_1 + \dots + \lambda \lambda_n x_n + \mu \mu_1 y_1 + \dots + \mu \mu_m y_m$$

est une combinaison linéaire des vecteurs $x_1, \dots, x_n, y_1, \dots, y_m$ de A . Donc $\lambda x + \mu y \in A'$ et par suite A' est un sous-espace vectoriel de E . En outre, $A \subset A'$ car pour tout $x \in A$, $x = 1 \cdot x \in A'$.

Soit F le sous-espace vectoriel engendré par A . F étant le plus petit sous-espace vectoriel contenant A , on a $F \subset A'$. Comme F est un sous-espace vectoriel contenant A , F contient toutes les combinaisons linéaires des éléments de A ; donc $A' \subset F$ et par suite $F = A'$.

6.3.1.4. Corollaire

Soient E un K -espace vectoriel, $x_1, \dots, x_n$ des vecteurs de E . Le sous-espace vectoriel engendré par $x_1, \dots, x_n$ est l'ensemble des combinaisons linéaires des x_i .

En particulier, pour tout vecteur non nul $x \in E$, le sous-espace engendré par x est la droite vectorielle $Kx = \{\lambda x : \lambda \in K\}$.

6.3.1.5. Définition

Soit E un K -espace vectoriel. On dit que les vecteurs $x_1, \dots, x_n$ de E forment un système de générateurs (ou une famille génératrice de E) si le sous-espace vectoriel engendré par $x_1, \dots, x_n$ est égal à E , c'est-à-dire si pour tout x de E , il existe des scalaires $\lambda_1, \dots, \lambda_n$ tels que $x = \lambda_1 x_1 + \dots + \lambda_n x_n$.

On dit que E est de dimension finie, s'il existe dans E une famille finie de générateurs de E . Dans le cas contraire, on dit que E est de dimension infinie.

6.3.2. FAMILLES LIBRES

6.3.2.1. Définition

Soit E un K -espace vectoriel. On dit que les vecteurs $x_1, \dots, x_n$ de E sont linéairement indépendants, ou que la famille $\{x_1, \dots, x_n\}$ est libre si la relation

$$\lambda_1 x_1 + \dots + \lambda_n x_n = 0$$

entraîne $\lambda_1 = \lambda_2 = \dots = \lambda_n = 0$.

On dit que les vecteurs $x_1, \dots, x_n$ sont linéairement dépendants, ou que la famille $\{x_1, \dots, x_n\}$ est liée s'il existe des scalaires $\lambda_1, \dots, \lambda_n$ non tous nuls tels que

$$\lambda_1 x_1 + \dots + \lambda_n x_n = 0.$$

6.3.2.2. Remarques

- a) Toute sous-famille d'une famille libre est une famille libre.
- b) Les éléments d'une famille libre sont non nuls.
- c) Toute sur-famille d'une famille liée est une famille liée. Ces assertions résultent immédiatement des définitions.

6.3.2.3. Exemple

Dans l'espace vectoriel K^n , les n vecteurs $e_1, \dots, e_n$ de l'Exemple 6.3.1.2 sont linéairement indépendants car la relation

$$\lambda_1 e_1 + \dots + \lambda_n e_n = 0 \text{ entraîne } \lambda_1 = \dots = \lambda_n = 0.$$

6.3.2.4. Exemple

Soient E un K -espace vectoriel, $x \in E$. Alors $\{x\}$ est libre si et seulement si $x \neq 0$. En effet, si la famille $\{x\}$ n'est pas libre, il existe un scalaire λ non nul tel que $\lambda x = 0$; λ étant inversible dans K (car $\lambda \neq 0$), on a

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

Réciproquement, si $x = 0$, on a $1 \cdot x = 0$ et la famille $\{x\}$ n'est pas libre.

6.3.2.5. Exemple

Dans l'espace vectoriel $\mathbb{R}^3$, les vecteurs $x_1 = (1, 0, 0)$, $x_2 = (0, 1, 0)$ et $x_3 = (1, 1, 0)$ sont linéairement dépendants car

$$x_1 + x_2 - x_3 = 0.$$

6.3.2.6. Théorème

Soient E un K -espace vectoriel et $x_1, \dots, x_n$ des vecteurs de E . Alors pour que la famille $\{x_1, \dots, x_n\}$ soit liée, il faut et il suffit que l'un des vecteurs x_i soit combinaison linéaire des autres.

Démonstration. Supposons que la famille $\{x_1, \dots, x_n\}$ soit liée; alors il existe des scalaires $\lambda_1, \dots, \lambda_n$ non tous nuls tels que

$$\lambda_1 x_1 + \dots + \lambda_n x_n = 0.$$

Si par exemple $\lambda_1 \neq 0$, on peut écrire

$$x_1 = -\frac{\lambda_2}{\lambda_1} x_2 - \dots - \frac{\lambda_n}{\lambda_1} x_n,$$

donc x_1 est combinaison linéaire de $x_2, \dots, x_n$.

Réciproquement, si par exemple $x_1 = \alpha_2 x_2 + \dots + \alpha_n x_n$ avec $\alpha_2, \dots, \alpha_n \in K$, on a

$$1 \cdot x_1 + (-\alpha_2)x_2 + \dots + (-\alpha_n)x_n = 0$$

ce qui montre que la famille $\{x_1, \dots, x_n\}$ est liée puisque le coefficient de x_1 est 1, donc non nul.

6.3.3. BASES D'UN ESPACE VECTORIEL

Nous venons d'étudier deux classes importantes de vecteurs d'un espace vectoriel : les familles génératrices et les familles libres. Les familles qui sont à la fois génératrices et libres vont jouer un rôle fondamental dans toute la suite du cours.

6.3.3.1. Définition

Soit E un K -espace vectoriel. On dit qu'une famille $\{e_1, \dots, e_n\}$ de vecteurs de E est une **base** de E si les vecteurs $e_1, \dots, e_n$ sont linéairement indépendants et engendrent E .

Le résultat suivant caractérise les bases de E .

6.3.3.2. Théorème

Soient $e_1, \dots, e_n$ des vecteurs d'un K -espace vectoriel E . Pour que la famille $\{e_1, \dots, e_n\}$ soit une base de E , il faut et il suffit que tout vecteur $x \in E$ s'exprime de façon unique comme combinaison linéaire des e_i .

Démonstration. Supposons que la famille $\{e_1, \dots, e_n\}$ soit une base de E . Alors tout vecteur $x \in E$ est combinaison linéaire des e_i ; il existe donc des scalaires $\lambda_1, \dots, \lambda_n$ tels que

$$x = \lambda_1 e_1 + \dots + \lambda_n e_n.$$

S'il existe des scalaires $\mu_1, \dots, \mu_n$ tels qu'on ait aussi

$$x = \mu_1 e_1 + \dots + \mu_n e_n$$

alors par soustraction, on obtient $(\lambda_1 - \mu_1)e_1 + \dots + (\lambda_n - \mu_n)e_n = 0$.

Comme la famille $\{e_1, \dots, e_n\}$ est libre, on a

$$\lambda_1 - \mu_1 = \dots = \lambda_n - \mu_n = 0$$

c'est-à-dire $\lambda_1 = \mu_1, \dots, \lambda_n = \mu_n$. La décomposition de x est bien unique.

Réciproquement supposons que tout vecteur de E s'écrit de façon unique comme combinaison linéaire de $e_1, \dots, e_n$; la famille $\{e_1, \dots, e_n\}$ est donc génératrice. Montrons qu'elle est libre. Par hypothèse, le vecteur nul s'écrit de façon unique sous la forme

$$0 \cdot e_1 + \dots + 0 \cdot e_n = 0.$$

Alors la relation $\lambda_1 e_1 + \dots + \lambda_n e_n = 0$,

montre que $\lambda_1 = \lambda_2 = \dots = \lambda_n = 0$, d'où le résultat.

6.3.3.3. Définition

Soient E un K -espace vectoriel et $(e_1, \dots, e_n)$ une base de E . Pour tout vecteur $x \in E$, il existe une famille unique $\{\lambda_1, \dots, \lambda_n\}$ de scalaires telle que

$$x = \lambda_1 e_1 + \dots + \lambda_n e_n.$$

Les scalaires $\lambda_1, \dots, \lambda_n$ s'appellent les **coordonnées** ou les **composantes** de x par rapport à la base $(e_1, \dots, e_n)$.

6.3.3.4. Remarque

On démontre que tout espace vectoriel sur un corps K possède au moins une base. Nous démontrerons plus tard ce résultat dans le cas des espaces vectoriels de dimension finie.

6.3.3.5. Exemple

Dans l'espace vectoriel K^n ($n \geq 2$), les n vecteurs $e_1, \dots, e_n$ de l'Exemple 6.3.1.2 sont linéairement indépendants et engendrent K^n . La famille $\{e_1, \dots, e_n\}$ est une base de K^n . On l'appelle la **base canonique** de K^n .

6.3.3.6. Exemple

Soit n un entier ≥ 0 . Notons $K_n[X]$ l'ensemble des polynômes à une indéterminée X , à coefficients dans un corps commutatif K , de degré $\leq n$. Nous savons déjà (Exemple 6.2.1.3) que $K_n[X]$ est un sous-espace vectoriel de $K[X]$; et c'est donc un K -espace vectoriel.

Montrons que la famille

$$(6.3.3.1) \quad \mathcal{B} = \{1, X, X^2, \dots, X^n\}$$

est une base de $K_n[X]$.

Par définition, tout polynôme $P = a_0 + a_1X + \dots + a_nX^n$ de $K_n[X]$ est combinaison linéaire d'éléments de $\mathcal{B}$ donc $\mathcal{B}$ est une famille génératrice de $K_n[X]$. Par conséquent $K_n[X]$ est de dimension finie.

D'autre part, la relation $\lambda_0 \cdot 1 + \lambda_1 X + \dots + \lambda_n X^n = 0$ implique $\lambda_0 = \lambda_1 = \dots = \lambda_n = 0$

puisque'un polynôme est identiquement nul si et seulement si tous ses coefficients sont nuls; donc $\mathcal{B}$ est une famille libre.

6.3.4. FAMILLES INFINIES

Les notions de familles génératrices, de familles libres et de bases peuvent se généraliser au cas des familles infinies de vecteurs mais la situation est plus délicate.

6.3.4.1. Définition

Soit I un ensemble fini ou non et soit $(\lambda_i)_{i \in I}$ une famille d'éléments du corps K . On dit que la famille $(\lambda_i)_{i \in I}$ est à **support fini** ou que les λ_i sont **presque tous nuls**, si l'ensemble des indices i tels que $\lambda_i \neq 0$ est fini.

Bien entendu, si I est un ensemble fini, cette définition n'a pas d'intérêt.

On dit que la famille $(\lambda_i)_{i \in I}$ est **triviale** si $\lambda_i = 0$ pour tout $i \in I$.

6.3.4.2. Définition

Soit $(x_i)_{i \in I}$ une famille de vecteurs d'un K -espace vectoriel E . On dit qu'un vecteur $x \in E$ est **combinaison linéaire** de la famille $(x_i)_{i \in I}$ s'il existe une famille $(\lambda_i)_{i \in I}$ de scalaires à support fini telle que

$$(6.3.4.1) \quad x = \sum_{i \in I} \lambda_i x_i.$$

La famille $(\lambda_i)_{i \in I}$ étant à support fini, on remarquera qu'un nombre fini seulement de termes $\lambda_i x_i$ sont non nuls dans la somme (6.3.4.1). Par définition, $\sum_{i \in I} \lambda_i x_i$ désigne la somme de ces termes non nuls.

Si $(x_i)_{i \in I}$ est une famille d'éléments d'un K -espace vectoriel E , on démontre facilement (cf. démonstration du Théorème 6.3.1.3) que l'ensemble des combinaisons linéaires des x_i est le plus petit sous-espace vectoriel de E contenant tous les x_i pour tout $i \in I$. On l'appelle le **sous-espace vectoriel engendré par la famille $(x_i)_{i \in I}$** .

On dit que la famille $(x_i)_{i \in I}$ est une **famille génératrice** si le sous-espace vectoriel qu'elle engendre est égal à E .

La notion d'indépendance linéaire se généralise comme suit.

6.3.4.3. Définition

Soit $(x_i)_{i \in I}$ une famille de vecteurs d'un K -espace vectoriel E . On appelle **relation linéaire entre les x_i** , toute famille $(\lambda_i)_{i \in I}$ de scalaires presque tous nuls tels que

$$\sum_{i \in I} \lambda_i x_i = 0.$$

On appelle **relation triviale**, toute famille triviale de scalaires (qui est toujours une relation linéaire entre les x_i).

On dit qu'une famille $(x_i)_{i \in I}$ de vecteur d'un K -espace vectoriel est **libre** ou que les x_i ($i \in I$) sont **linéairement indépendants**, si la seule relation linéaire entre les x_i est la relation triviale.

Autrement dit, la famille $(x_i)_{i \in I}$ est libre si pour toute partie finie J de I , la relation $\sum_{i \in J} \lambda_i x_i = 0$ entraîne $\lambda_i = 0$ pour tout $i \in J$.

Une famille qui n'est pas libre est dite **liée**.

On appelle **base** de E , toute famille $(e_i)_{i \in I}$ d'éléments de E à la fois libre et génératrice.

Le théorème suivant se démontre exactement comme le théorème 6.3.3.2.

6.3.4.4. Théorème

Soit $(e_i)_{i \in I}$ une famille de vecteurs d'un K -espace vectoriel E . Pour que la famille $(e_i)_{i \in I}$ soit une base de E , il faut et il suffit que tout vecteur de E s'exprime de façon unique comme combinaison linéaire des e_i .

Si $(e_i)_{i \in I}$ est une base de E , pour tout vecteur $x \in E$, il existe une famille unique $(\lambda_i)_{i \in I}$ de scalaires presque tous nuls telle que

$$x = \sum_{i \in I} \lambda_i x_i.$$

Les scalaires λ_i s'appellent les **coordonnées** ou les **composantes** de x par rapport à la base $(e_i)_{i \in I}$.

Chapitre 7 : APPLICATIONS LINÉAIRES

En mathématiques, on est souvent amené à associer des «morphismes» à une structure donnée. Dans la théorie des espaces vectoriels, les morphismes sont des applications qui sont compatibles avec la structure de K -espace vectoriel.

Après avoir défini les applications linéaires d'un K -espace vectoriel dans un autre (ou dans lui-même) et donné quelques exemples élémentaires, nous examinons quelques propriétés générales de ces applications. En particulier, nous montrons que si E et F sont des K -espaces vectoriels l'ensemble $\mathcal{L}(E, F)$ des applications linéaires de E dans F est un K -espace vectoriel.

Nous étudions enfin un exemple important d'endomorphisme : les projecteurs.

7.1. Généralités

7.1.1. DÉFINITIONS

7.1.1.1. Définition

Soient E et F deux espaces vectoriels sur le même corps K . On dit qu'une application $u : E \longrightarrow F$ est K -linéaire ou est un **homomorphisme** si l'on a :

a) $u(x + y) = u(x) + u(y)$

b) $u(\lambda x) = \lambda u(x)$

quels que soient $x, y \in E$ et $\lambda \in K$.

Lorsqu'il n'y a pas d'ambiguïté sur le corps K , on dit simplement que u est une **application linéaire** ou un **morphisme d'espaces vectoriels**.

Une application linéaire de E dans lui-même s'appelle un **endomorphisme** de E ou un **opérateur linéaire** dans E . On appelle **automorphisme** de E tout endomorphisme bijectif de E .

Une application linéaire de E dans K s'appelle une **forme linéaire** sur E .

Une application linéaire bijective $u : E \longrightarrow F$ s'appelle un **isomorphisme** d'espaces vectoriels. S'il existe un isomorphisme de E sur F on dit alors que les espaces vectoriels E et F sont **isomorphes** et on écrit $E \approx F$.

L'ensemble des applications K -linéaires de E dans F se note $\mathcal{L}_K(E, F)$ ou simplement $\mathcal{L}(E, F)$ si aucune confusion n'est à craindre. Cet ensemble est aussi noté $\text{Hom}(E, F)$.

L'ensemble des endomorphismes de E est noté $\mathcal{L}_K(E)$ ou $\mathcal{L}(E)$. On le note aussi $\text{End}(E)$.

L'ensemble des automorphismes de E est noté $GL_K(E)$ ou $GL(E)$ et s'appelle le **groupe linéaire général** de E (nous justifierons plus loin au Corollaire 7.2.1.2 cette dénomination).

7.1.1.2. Remarques

a) Si dans la condition a) de la Définition 7.1.1.1, on pose $x = 0$, on obtient $u(0) = 0$. De même, en prenant $\lambda = -1$ dans la condition b), on obtient $u(-x) = -u(x)$.

b) Une application $u : E \longrightarrow F$ est linéaire si et seulement si $u(\lambda x + \mu y) = \lambda u(x) + \mu u(y)$ quels que soient $x, y \in E$, $\lambda, \mu \in K$.

7.1.2. EXEMPLES

7.1.2.1. Exemple

Soient E et F deux K -espaces vectoriels. L'application $0 : E \longrightarrow F$ définie par $0(x) = 0$ pour tout $x \in E$ est linéaire. On l'appelle l'**application nulle** de E dans F .

7.1.2.2. Exemple

L'application identique, notée 1_E ou Id_E , d'un K -espace vectoriel E est linéaire. C'est même un automorphisme de E .

7.1.2.3. Exemple

Soit E un K -espace vectoriel et soit α un élément non nul de K . On appelle **homothétie** de rapport α l'application $h_\alpha : E \longrightarrow E$ telle que $h_\alpha(x) = \alpha \cdot x$ pour tout $x \in E$. h_α est un élément de $GL(E)$.

7.1.2.4. Exemple

Soient E un K -espace vectoriel et F un sous-espace de E . L'application canonique $\pi : E \longrightarrow E/F$ qui, à tout vecteur x de E associe sa classe $\bar{x}$ dans l'espace vectoriel quotient E/F , est K -linéaire d'après la définition de la structure vectorielle de E/F . On l'appelle le **morphisme canonique**.

7.1.2.5. Exemple

Soient $E_1, \dots, E_n$ des espaces vectoriels sur le même corps K et $E = E_1 \times \dots \times E_n$ leur produit (cf. Exemple 6.1.3.4).

Pour $1 \leq i \leq n$, l'application

$$pr_i : E_1 \times \dots \times E_n \longrightarrow E_i$$

définie par $pr_i(x_1, \dots, x_n) = x_i$ est linéaire. On l'appelle la $i^{\text{ème}}$ projection.

7.2. Propriétés des applications linéaires

Nous donnons dans ce paragraphe quelques propriétés générales des applications linéaires. D'autres propriétés classiques seront exposées lorsque nous étudierons les espaces vectoriels de dimension finie.

7.2.1. COMPOSÉE DE DEUX APPLICATIONS LINÉAIRES

7.2.1.1. Théorème

Soient E, F et G trois espaces vectoriels sur le corps K . Soient $u \in \mathcal{L}(E, F)$ et $v \in \mathcal{L}(F, G)$. Alors :

a) L'application composée $v \circ u$ est une application linéaire de E dans G .

b) Si u est un isomorphisme de E sur F , l'application réciproque u^{-1} est un isomorphisme de F sur E .

Démonstration.

a) De la linéarité de u et de v , nous déduisons, quels que soient $x, y \in E$, $\lambda, \mu \in K$,

$$\begin{aligned} (v \circ u)(\lambda x + \mu y) &= v[u(\lambda x + \mu y)] = v[\lambda u(x) + \mu u(y)] \\ &= \lambda v(u(x)) + \mu v(u(y)) \\ &= \lambda(v \circ u)(x) + \mu(v \circ u)(y). \end{aligned}$$

b) L'application réciproque d'une application bijective étant bijective, il suffit de montrer que u^{-1} est linéaire.

Quels que soient les éléments x et y de F , il existe des éléments x' et y' uniques de E tels que $u(x') = x$ et $u(y') = y$. Quels que soient $\lambda, \mu \in K$, on a alors

$$\begin{aligned} u^{-1}(\lambda x + \mu y) &= u^{-1}(\lambda u(x') + \mu u(y')) \\ &= u^{-1}(u(\lambda x' + \mu y')) \\ &= \lambda x' + \mu y' \quad (\text{car } u^{-1} \circ u = \text{Id}_E) \\ &= \lambda u^{-1}(x) + \mu u^{-1}(y) \end{aligned}$$

d'où la linéarité de u^{-1} .

7.2.1.2. Corollaire

L'ensemble $GL(E)$ des automorphismes d'un K -espace vectoriel E est un groupe pour la loi de composition des applications.

On sait que la composée de deux bijections est un bijection et que la composée de deux applications linéaires est linéaire.

Donc, si $u, v \in GL(E)$, alors $uov \in GL(E)$.

D'autre part, d'après le Théorème 7.2.1.1 b), si $u \in GL(E)$, alors $u^{-1} \in GL(E)$.

Nous avons ainsi démontré que $GL(E)$ (qui est non vide puisque l'application identique $1_E \in GL(E)$) est un sous-groupe du groupe des permutations de E .

7.2.1.3. Remarque

Si $u, v \in GL(E)$, on a d'après le Théorème 2.2.4.2 d) :

$$(v \circ u)^{-1} = u^{-1} \circ v^{-1}$$

7.2.2. IMAGE ET NOYAU D'UNE APPLICATION LINÉAIRE

7.2.2.1. Théorème

Soient E et F deux K -espaces vectoriels et soit $u \in \mathcal{L}(E, F)$. Alors :

a) *L'image par u d'un sous-espace vectoriel de E est un sous-espace vectoriel de F .*

b) *L'image réciproque par u d'un sous-espace vectoriel de F est un sous-espace vectoriel de E .*

Démonstration. a) Soit A un sous-espace vectoriel de E . Comme $0 \in A$, $0 = u(0) \in u(A)$ et $u(A) \neq \emptyset$. Soient $y, z \in u(A)$, $\lambda, \mu \in K$. Il existe x et x' dans A tels que $y = u(x)$ et $z = u(x')$. Comme u est linéaire, on a

$$\lambda y + \mu z = \lambda u(x) + \mu u(x') = u(\lambda x + \mu x').$$

Or $\lambda x + \mu x' \in A$ puisque A est un sous-espace vectoriel de E . Donc $\lambda y + \mu z \in u(A)$ et $u(A)$ est un sous-espace vectoriel de F .

Démontrons b). Soit B un sous-espace vectoriel de F . $u^{-1}(B)$ contient 0 , puisque $u(0) = 0 \in B$. Donc $u^{-1}(B) \neq \emptyset$.

Soient x et y deux éléments de $u^{-1}(B)$, λ et μ deux scalaires. On a par définition $u(x) \in B$ et $u(y) \in B$. D'où

$$u(\lambda x + \mu y) = \lambda u(x) + \mu u(y) \in B$$

car B est un sous-espace vectoriel de F . Donc $\lambda x + \mu y \in u^{-1}(B)$.

7.2.2.2. Définition

Soient E et F deux K -espaces vectoriels et soit $u \in \mathcal{L}(E, F)$.

- a) On appelle **image** de u , et on note $\text{Im}(u)$, le sous-espace vectoriel $u(E)$ de F .
- b) On appelle **noyau** de u , et on note $\text{Ker}(u)$, le sous-espace vectoriel $u^{-1}(\{0\})$ de E .

7.2.2.3. Théorème

Soient E et F deux K -espaces vectoriels et $u \in \mathcal{L}(E, F)$.

- a) u est injective si et seulement si $\text{Ker}(u) = \{0\}$.
- b) u est surjective si et seulement si $\text{Im}(u) = F$.

Démonstration. a) découle aussitôt du Théorème 3.3.2.4 puisque u est un morphisme de groupes.

b) résulte immédiatement de la définition de $\text{Im}(u)$.

7.2.3. APPLICATIONS LINÉAIRES ET FAMILLES DE VECTEURS

7.2.3.1. Théorème

Soient E et F deux K -espaces vectoriels et $u \in \mathcal{L}(E, F)$.

- a) Si $(x_i)_{i \in I}$ est une famille génératrice de E , la famille $(u(x_i))_{i \in I}$ est une famille génératrice de $\text{Im}(u)$.
- b) Si $(x_i)_{i \in I}$ est une famille liée de vecteurs de E , la famille $(u(x_i))_{i \in I}$ est une famille liée de vecteurs de F .

Démonstration. a) Soit $y \in \text{Im}(u)$; il existe $x \in E$ tel que $y = u(x)$. Puisque $(x_i)_{i \in I}$ est une famille génératrice, il existe une famille $(\lambda_i)_{i \in I}$ de scalaires presque tous nuls telle que

$$x = \sum_{i \in I} \lambda_i x_i.$$

On en déduit (puisque u est linéaire et la famille de scalaires à support fini) :

$$u(x) = \sum_{i \in I} \lambda_i u(x_i),$$

donc $u(x)$ est combinaison linéaire des $u(x_i)$.

b) Soit $(x_i)_{i \in I}$ une famille liée de vecteurs de E . Il existe donc une famille non triviale de scalaires $(\lambda_i)_{i \in I}$ telle que $\sum_{i \in I} \lambda_i x_i = 0$. D'où $u \left(\sum_{i \in I} \lambda_i x_i \right) = \sum_{i \in I} \lambda_i u(x_i) = 0$ car u est linéaire et la famille de scalaires à support fini. Par suite la famille de vecteurs $(u(x_i))_{i \in I}$ est liée.

Attention. L'image d'une famille génératrice de E n'est une famille génératrice de F que si u est surjective.

L'image d'une famille libre de E n'est pas, en général, une famille libre de F . Toutefois on a le résultat suivant.

7.2.3.2. Théorème

Soient E et F deux K -espaces vectoriels et soit $u \in \mathcal{L}(E, F)$. Les propositions suivantes sont équivalentes :

a) u est injective.

b) L'image par u de toute famille libre de E est une famille libre de F .

Démonstration. Démontrons que a) $\implies$ b). Soit $(x_i)_{i \in I}$ une famille libre de E et soit $(\lambda_i)_{i \in I}$ une famille de scalaires presque tous nuls telle que

$$\sum_{i \in I} \lambda_i u(x_i) = 0,$$

ou encore, puisque u est linéaire :

$$u \left(\sum_{i \in I} \lambda_i x_i \right) = 0.$$

u étant injective, on en déduit

$$\sum_{i \in I} \lambda_i x_i = 0,$$

puisque $\text{Ker}(u) = \{0\}$.

La famille $(x_i)_{i \in I}$ étant libre, la relation linéaire $(\lambda_i)_{i \in I}$ est la relation triviale. La famille $(u(x_i))_{i \in I}$ est donc libre.

b) $\implies$ a) Soit $x \in \text{Ker}(u)$. L'image de la famille $\{x\}$ est la famille $\{0\}$ qui est liée ; la famille $\{x\}$ est donc liée, i.e. $x = 0$ et u est injective.

7.2.3.3. Théorème

Soient E et F deux K -espaces vectoriels et $u \in \mathcal{L}(E, F)$. Les propositions suivantes sont équivalentes :

a) u est un isomorphisme de E sur F .

b) L'image par u de toute base de E est une base de F .

Démonstration. a) $\implies$ b) Supposons que u soit un isomorphisme de E sur F et soit $(e_i)_{i \in I}$ une base de E . Alors, d'après les Théorèmes 7.2.3.1 et 7.2.3.2, $(u(e_i))_{i \in I}$ est une base de F puisque u est à la fois injectif et surjectif, donc b) est vérifiée.

b) $\implies$ a) Soit $(e_i)_{i \in I}$ une base de E dont l'image par u est une base de F . Montrons que u est injective. Soit $x \in \text{Ker}(u)$; on a $x = \sum_{i \in I} x_i e_i$ où $(x_i)_{i \in I}$ est une famille de scalaires presque tous nuls. Par linéarité, on a $u(x) = \sum_{i \in I} x_i u(e_i) = 0$.

La famille $(u(e_i))_{i \in I}$ étant libre, la famille $(x_i)_{i \in I}$ est la famille triviale, donc $x = 0$ et u est bien injective.

Montrons que u est surjective. Soit $y \in F$. $(u(e_i))_{i \in I}$ étant une base de F , y s'écrit de façon unique $y = \sum_{i \in I} \lambda_i u(e_i)$ où $(\lambda_i)_{i \in I}$ est une famille de scalaires presque tous nuls. Alors $y = u(x)$ avec $x = \sum_{i \in I} \lambda_i e_i$, donc u est bien surjective.

7.2.3.4. Théorème

Soient E et F deux K -espaces vectoriels. Soit $(e_i)_{i \in I}$ une base de E et soit $(y_i)_{i \in I}$ une famille de vecteurs de F indexée par le même ensemble d'indices I .

a) Il existe une application linéaire et une seule u de E dans F telle que $u(e_i) = y_i$ pour tout $i \in I$.

b) u est injective si et seulement si $(y_i)_{i \in I}$ est une famille libre de F .

c) u est surjective si et seulement si $(y_i)_{i \in I}$ est une famille génératrice de F .

d) u est un isomorphisme de E sur F si et seulement si $(y_i)_{i \in I}$ est une base de F .

Démonstration. a) Tout vecteur x de E s'écrit de façon unique sous la forme

$$x = \sum_{i \in I} \lambda_i e_i$$

où les scalaires $(\lambda_i)_{i \in I}$ sont presque tous nuls.

En posant

$$u(x) = \sum_{i \in I} \lambda_i y_i,$$

on définit une application u de E dans F et il est clair que $u(e_i) = y_i$ pour tout i , puisque $e_i = \sum_{j \in I} \lambda_j e_j$ avec $\lambda_i = 1$ et $\lambda_j = 0$ si $j \neq i$.

Montrons que u est linéaire.

Soient $x = \sum_{i \in I} \lambda_i e_i$ et $y = \sum_{i \in I} \mu_i e_i$ deux vecteurs de E . Quels que soient les scalaires α et λ , on a, par définition de u :

$$\begin{aligned} u(\alpha x + \lambda y) &= u\left(\sum_{i \in I} (\alpha \lambda_i + \lambda \mu_i) e_i\right) = \sum_{i \in I} (\alpha \lambda_i + \lambda \mu_i) y_i \\ &= \alpha \sum_{i \in I} \lambda_i y_i + \lambda \sum_{i \in I} \mu_i y_i = \alpha u(x) + \lambda u(y). \end{aligned}$$

On a ainsi établi l'existence et la linéarité de u .

Soit $v \in \mathcal{L}(E, F)$ telle que $v(e_i) = y_i$ pour tout $i \in I$. Pour tout vecteur $x = \sum_{i \in I} \lambda_i e_i$, on a

$$v(x) = \sum_{i \in I} \lambda_i v(e_i) = \sum_{i \in I} \lambda_i y_i = u(x).$$

Donc $v = u$, d'où l'unicité de u .

b) Supposons u injective. Soit $(\lambda_i)_{i \in I}$ une famille de scalaires presque tous nuls telle que

$$\sum_{i \in I} \lambda_i y_i = 0.$$

Cette relation s'écrit $\sum_{i \in I} \lambda_i u(e_i) = 0$, ou encore, puisque u est linéaire

$u\left(\sum_{i \in I} \lambda_i e_i\right) = 0$. Comme u est injective, on en déduit $\sum_{i \in I} \lambda_i e_i = 0$. Comme $(e_i)_{i \in I}$ est une famille libre, la relation linéaire $(\lambda_i)_{i \in I}$ est la relation triviale, donc la famille $(y_i)_{i \in I}$ est libre.

Réciproquement, supposons la famille $(y_i)_{i \in I}$ libre et montrons que u est injective. Soit $x = \sum_{i \in I} \lambda_i e_i$ un élément de $\text{Ker}(u)$, où la famille $(\lambda_i)_{i \in I}$ est à support fini. On a

$$u(x) = \sum_{i \in I} \lambda_i y_i = 0.$$

Comme $(y_i)_{i \in I}$ est une famille libre, la relation linéaire $(\lambda_i)_{i \in I}$ est triviale, donc $x = 0$ et u est injective.

c) On a les équivalences suivantes :

u est surjective $\iff$ pour tout $y \in F$, il existe $x \in E$ tel que $y = u(x)$ $\iff$ pour tout $y \in F$, il existe une famille $(\lambda_i)_{i \in I}$ à support fini telle que $y =$

$$u\left(\sum_{i \in I} \lambda_i e_i\right) = \sum_{i \in I} \lambda_i y_i, \text{ puisque } u \text{ est linéaire.}$$

Ainsi u est surjective si et seulement si la famille $(y_i)_{i \in I}$ est génératrice.

d) résulte immédiatement de b) et c).

7.2.4. DÉCOMPOSITION CANONIQUE D'UNE APPLICATION LINÉAIRE

7.2.4.1. Théorème

Soient E et F deux K -espaces vectoriels, $u \in \mathcal{L}(E, F)$ et N le noyau de u .

La relation $x \mathcal{R} y \iff u(x) = u(y)$ est une relation d'équivalence dans E . Soit $\pi : E \longrightarrow E/N$ le morphisme canonique et soit $j : u(E) \longrightarrow F$ l'injection canonique. Alors, il existe un isomorphisme unique $\bar{u}$ de l'espace vectoriel quotient E/N sur le sous-espace $u(E)$ de F tel que :

$$j \circ \bar{u} \circ \pi = u.$$

Démonstration. On vérifie facilement que $\mathcal{R}$ est une relation d'équivalence. De plus

$$\begin{aligned} x \mathcal{R} y &\iff u(x) = u(y) \iff u(x) - u(y) = u(x - y) = 0 \\ &\iff x - y \in N. \end{aligned}$$

Comme N est un sous-espace vectoriel de E , E/N est un K -espace vectoriel.

D'après l'Exemple 7.1.2.4, l'application canonique π est linéaire.

Rappelons que si X et Y sont deux ensembles tels que $X \subset Y$, l'injection canonique $j : X \longrightarrow Y$ est définie par $j(x) = x$ pour tout $x \in X$. Si on prend pour X , le sous-espace $u(E)$ de F , et pour Y l'espace F , alors il est évident que l'injection canonique est linéaire.

Rappelons que l'application $\bar{u} : E/N \longrightarrow u(E)$ est définie en posant :

$$\bar{u}(\bar{x}) = u(x)$$

pour tout $\bar{x} \in E/N$, où x est un représentant de la classe $\bar{x}$. $\bar{u}$ est une bijection ; c'est l'application déduite de u par passage au quotient (voir le Théorème 1.4.2.6).

Montrons que $\bar{u}$ est K -linéaire.

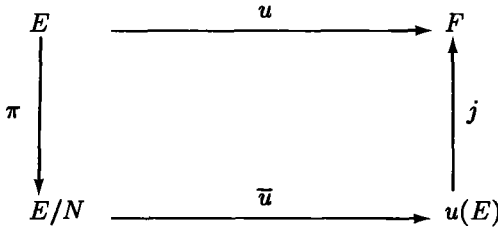
Soient $\bar{x}, \bar{y} \in E/N$. Quels que soient les scalaires λ et μ , on a :

$$\begin{aligned} \bar{u}(\lambda \bar{x} + \mu \bar{y}) &= \bar{u}(\overline{\lambda x + \mu y}) = u(\lambda x + \mu y) = \lambda u(x) + \mu u(y) \\ &= \lambda \bar{u}(\bar{x}) + \mu \bar{u}(\bar{y}). \end{aligned}$$

Ainsi $\bar{u}$ est une application linéaire bijective, c'est-à-dire un isomorphisme de E/N sur $u(E)$. La factorisation canonique de u s'écrit :

$$u = j \circ \bar{u} \circ \pi$$

et on a le diagramme suivant



appelé diagramme de décomposition canonique de u .

La bijection $\bar{u}$ s'appelle l'isomorphisme canonique.

7.2.4.2. Théorème

Soient E un K -espace vectoriel et E_1 un sous-espace de E . Si E_2 et E_3 sont deux sous-espaces supplémentaires de E_1 dans E , c'est-à-dire si $E = E_1 \oplus E_2 = E_1 \oplus E_3$, alors E_2 et E_3 sont isomorphes.

Démonstration. Par hypothèse, pour tout $x \in E$, il existe (x_1, x_2) unique dans $E_1 \times E_2$ tel que

$$x = x_1 + x_2.$$

De même il existe (x'_1, x_3) unique dans $E_1 \times E_3$ tel que $x = x'_1 + x_3$.

Soient u et v les endomorphismes de E définis par

$$u(x) = x_2 \text{ et } v(x) = x_3.$$

On a

$$\text{Ker}(u) = \text{Ker}(v) = E_1$$

et

$$\text{Im}(u) = E_2, \text{ Im}(v) = E_3.$$

La décomposition canonique de u et v montre que E_2 et E_3 sont isomorphes à l'espace vectoriel quotient E/E_1 . Donc E_2 et E_3 sont isomorphes.

7.3. L'espace vectoriel $\mathcal{L}(E, F)$

Soient E et F deux espaces vectoriels sur un corps commutatif K . Nous allons munir l'ensemble $\mathcal{L}(E, F)$ des applications linéaires de E dans F d'une structure de K -espace vectoriel. Pour cela, il nous faut définir la somme de deux applications linéaires et le produit d'une application linéaire par un scalaire.

7.3.1. ADDITION DANS $\mathcal{L}(E, F)$ **7.3.1.1. Définition**

Soient u et v deux éléments de $\mathcal{L}(E, F)$. On appelle **somme** de u et v , et on note $u + v$, l'application de E dans F définie par :

$$(u + v)(x) = u(x) + v(x) \quad \text{pour tout } x \in E.$$

Montrons que $u + v$ est linéaire.

Quels que soient les vecteurs x, y de E et quels que soient les scalaires α et β , on a :

$$\begin{aligned} (u + v)(\alpha x + \beta y) &= u(\alpha x + \beta y) + v(\alpha x + \beta y) \\ &= \alpha u(x) + \beta u(y) + \alpha v(x) + \beta v(y) \\ &= \alpha (u + v)(x) + \beta (u + v)(y). \end{aligned}$$

Muni de cette addition $(u, v) \mapsto u + v$, $\mathcal{L}(E, F)$ est un groupe abélien ; on vérifie en effet facilement les relations suivantes :

$$\begin{aligned} (u + v) + w &= u + (v + w), \\ u + v &= v + u, \\ 0 + u &= u + 0, \\ u + (-u) &= 0 \end{aligned}$$

l'application $-u$ étant définie par $(-u)(x) = -u(x)$ pour tout $x \in E$ et 0 étant l'application nulle (cf. Exemple 7.1.2.1).

7.3.2. PRODUIT D'UNE APPLICATION LINÉAIRE PAR UN SCALAIRE**7.3.2.1. Définition**

Soit $u \in \mathcal{L}(E, F)$. On appelle **produit** de u par le scalaire λ , et on note λu , l'application de E dans F , définie par

$$(\lambda u)(x) = \lambda u(x) \quad \text{pour tout } x \in E.$$

Montrons que λu est linéaire.

Soient x et y deux vecteurs de E et soient α et β deux scalaires ; on a :

$$\begin{aligned} (\lambda u)(\alpha x + \beta y) &= \lambda u(\alpha x + \beta y) = \lambda (\alpha u(x) + \beta u(y)) \\ &= \lambda \alpha u(x) + \lambda \beta u(y) = \alpha (\lambda u)(x) + \beta (\lambda u)(y). \end{aligned}$$

On vérifie facilement que quels que soient les éléments u et v de $\mathcal{L}(E, F)$ et quels que soient les scalaires λ et μ ,

$$\begin{aligned} \lambda(u + v) &= \lambda u + \lambda v, \quad (\lambda + \mu)u = \lambda u + \mu u \\ (\lambda \mu)u &= \lambda(\mu u), \quad 1 \cdot u = u. \end{aligned}$$

On peut donc énoncer le

7.3.2.2. Théorème

Soient E et F deux K -espaces vectoriels. L'ensemble $\mathcal{L}(E, F)$ des applications linéaires de E dans F , muni des lois de compositions

$$(u, v) \mapsto u + v \text{ et } (\lambda, u) \mapsto \lambda u$$

est un K -espace vectoriel.

7.3.3. CAS PARTICULIER : $E = F$

Si $E = F$, le Théorème 7.3.2.2 montre que $\mathcal{L}(E)$ est un espace vectoriel sur K . D'après le Théorème 7.2.1.1, la loi de composition des applications est une loi interne dans $\mathcal{L}(E)$.

Si u, v, w sont des éléments de $\mathcal{L}(E)$, on a :

$$\begin{aligned} (u \circ v) \circ w &= u \circ (v \circ w), \\ u \circ (v + w) &= u \circ v + u \circ w, \\ (v + w) \circ u &= v \circ u + w \circ u. \end{aligned}$$

Vérifions par exemple le deuxième résultat. Soit $x \in E$; on a, en utilisant la linéarité de u et la définition de $v + w$:

$$\begin{aligned} [u \circ (v + w)](x) &= u((v + w)(x)) = u(v(x) + w(x)) \\ &= u(v(x)) + u(w(x)) \\ &= (u \circ v)(x) + (u \circ w)(x) = (u \circ v + u \circ w)(x). \end{aligned}$$

D'où

$$u \circ (v + w) = u \circ v + u \circ w.$$

Ainsi, $\mathcal{L}(E)$, muni des lois de composition $(u, v) \mapsto u + v$ et $(u, v) \mapsto u \circ v$ est un anneau (non commutatif en général), l'élément unité étant l'application identique de E .

On a de plus, pour tout $\lambda \in K$:

$$\lambda(u \circ v) = (\lambda u) \circ v = u \circ (\lambda v).$$

En effet, pour tout $x \in E$, on a :

$$(\lambda(u \circ v))(x) = \lambda(u \circ v)(x) = \lambda(u(v(x))) = ((\lambda u) \circ v)(x),$$

d'où

$$\lambda(u \circ v) = (\lambda u) \circ v.$$

On démontrerait de même que $\lambda(u \circ v) = u \circ (\lambda v)$.

Ces propriétés de $\mathcal{L}(E)$ nous amènent à poser la définition suivante.

7.3.3.1. Définition

Soit K un corps commutatif. On dit qu'un ensemble A est une **algèbre sur K** , ou une **K -algèbre**, si les conditions suivantes sont vérifiées :

- a) A est muni d'une structure de K -espace vectoriel.
- b) Il existe une loi interne dans A , appelée **multiplication** et notée $(x, y) \mapsto x \cdot y$ telle que $(A, +, \cdot)$ soit un anneau.
- c) Quels que soient les vecteurs x et y de A et pour tout scalaire λ on a :

$$\lambda(x \cdot y) = (\lambda x) \cdot y = x \cdot (\lambda y).$$

Si la loi $(x, y) \mapsto x \cdot y$ est commutative, on dit que l'algèbre A est **commutative**.

Ainsi $\mathcal{L}(E)$ est un algèbre sur K . On l'appelle l'**algèbre des endomorphismes de E** .

7.4. Projecteurs

7.4.1. DÉFINITION

7.4.1.1. Définition

Soit E un espace vectoriel sur le corps K . On dit qu'un endomorphisme p de E est un **projecteur** si $p \circ p = p$.

Soit I un ensemble non vide. On dit qu'une famille $(p_i)_{i \in I}$ de projecteurs est **orthogonale** si $p_i \circ p_j = 0$ pour $i \neq j$.

Soit E un K -espace vectoriel. Si E est somme directe d'une famille finie $(E_i)_{1 \leq i \leq n}$ de sous-espaces, tout $x \in E$ s'écrit d'une manière unique sous la forme

$$x = x_1 + \dots + x_n \text{ avec } x_i \in E_i \text{ pour } 1 \leq i \leq n.$$

Définissons l'application p_i de E dans E en posant :

$$(7.4.1.1) \quad p_i(x) = x_i.$$

Pour tout $x \in E$, $p_i(x)$ est donc l'unique vecteur de E_i tel que $x - p_i(x)$ appartienne au sous-espace engendré par $E_1, \dots, E_{i-1}, E_{i+1}, \dots, E_n$. On dit que $p_i(x)$ est la **projection de x sur E_i** .

Nous allons montrer que $(p_i)_{1 \leq i \leq n}$ est une famille orthogonale de projecteurs. On dit que p_i est le **projecteur de E sur E_i parallèlement à $\bigoplus_{j \neq i} E_j$** .

7.4.2. PROPRIÉTÉS DES PROJECTEURS

7.4.2.1. Théorème

Conservons les notations qui viennent d'être introduites. Alors :

a) p_i est un endomorphisme de E et $p_1 + \dots + p_n = \text{Id}_E$.

b) $\text{Ker}(p_i) = \bigoplus_{j \neq i} E_j$ et $\text{Im}(p_i) = E_i$.

c) $p_i \circ p_j = \begin{cases} p_i & \text{si } i = j \\ 0 & \text{si } i \neq j \end{cases}$

Démonstration. a) Soient $x = x_1 + \dots + x_n$ et $y = y_1 + \dots + y_n$ deux éléments de E avec $x_i, y_i \in E_i$ pour $1 \leq i \leq n$. On a, par définition de p_i :

$$p_i(x) = x_i \text{ et } p_i(y) = y_i.$$

Pour tout couple (λ, μ) de scalaires, on a :

$$\lambda x + \mu y = (\lambda x_1 + \mu y_1) + \dots + (\lambda x_n + \mu y_n).$$

D'où

$$p_i(\lambda x + \mu y) = \lambda x_i + \mu y_i = \lambda p_i(x) + \mu p_i(y),$$

et p_i est linéaire.

Comme tout vecteur x de E s'écrit de façon unique

$$x = x_1 + \dots + x_n \text{ avec } x_i \in E_i \text{ pour } 1 \leq i \leq n,$$

on a $\text{Id}_E(x) = p_1(x) + \dots + p_n(x) = (p_1 + \dots + p_n)(x)$, d'où

$$\text{Id}_E = p_1 + \dots + p_n.$$

Démontrons b) Si $x \in E$ est tel que $p_i(x) = 0$, on a

$$x = x_1 + \dots + x_{i-1} + x_{i+1} + \dots + x_n$$

et par suite

$$\text{Ker}(p_i) \subset \bigoplus_{j \neq i} E_j.$$

Réciproquement tout vecteur $x \in \bigoplus_{j \neq i} E_j$ s'écrit de façon unique sous la forme

$$x = x_1 + \dots + x_{i-1} + 0 + x_{i+1} + \dots + x_n$$

avec $0 \in E_i$ et $x_j \in E_j$ pour $j \neq i$.

On en déduit $p_i(x) = 0$, d'où $\bigoplus_{j \neq i} E_j \subset \text{Ker}(p_i)$.

Par suite, on a l'égalité $\text{Ker}(p_i) = \bigoplus_{j \neq i} E_j$.

De la définition de p_i , il résulte que $p_i(x) \in E_i$ pour tout $x \in E$; d'où $\text{Im}(p_i) \subset E_i$.

Inversement, tout $x \in E_i$ s'écrit de façon unique

$$x = 0 + \dots + 0 + x + 0 + \dots + 0$$

avec $x \in E_i$ et $0 \in E_j$ pour $j = 1, \dots, i-1, i+1, \dots, n$.

On en déduit: $p_i(x) = x$ pour tout $x \in E_i$, d'où $E_i \subset \text{Im}(p_i)$. Finalement, on a $\text{Im}(p_i) = E_i$.

Démontrons enfin c).

Nous venons de voir que si $x \in E_i$, alors $p_i(x) = x$. Réciproquement, soit $x \in E$ tel que $p_i(x) = x$. Comme $p_i(x) \in E_i$ nécessairement $x \in E_i$.

Comme $p_i(x) \in E_i$ pour tout $x \in E$, on a $p_i(p_i(x)) = p_i(x)$, c'est-à-dire

$$p_i \circ p_i = p_i.$$

D'autre part, pour tout $x \in E$, $p_j(x) \in E_j$; donc $p_i(p_j(x)) = 0$ si $i \neq j$, c'est-à-dire

$$p_i \circ p_j = 0 \quad \text{si } i \neq j.$$

Le théorème suivant établit la réciproque du résultat précédent.

7.4.2.2. Théorème

Soit $(p_i)_{1 \leq i \leq n}$ une famille finie orthogonale de projecteurs d'un K -espace vectoriel E , telle que $p_1(x) + \dots + p_n(x) = x$ pour tout $x \in E$. Alors E est somme directe des sous-espaces $E_i = p_i(E)$.

Démonstration. La relation

$$p_1(x) + \dots + p_n(x) = x$$

pour tout $x \in E$, montre que E est somme des sous-espaces E_i .

Pour montrer que la somme est directe, il suffit de montrer que la relation $x_1 + \dots + x_n = 0$ où $x_i \in E_i$ pour $1 \leq i \leq n$, entraîne $x_1 = \dots = x_n = 0$.

Considérons donc des $x_i \in E_i$ tels que

$$(7.4.2.1) \quad x_1 + \dots + x_n = 0.$$

Puisque $x_j \in E_j = p_j(E)$, il existe $z_j \in E$ tel que $x_j = p_j(z_j)$; on a alors :

$$p_i(x_j) = p_i(p_j(z_j)) = \begin{cases} p_i(z_i) = x_i & \text{si } i = j \\ 0 & \text{si } i \neq j \end{cases}$$

On en déduit de là, en appliquant p_i à la relation (7.4.2.1), $x_i = 0$ pour $i = 1, \dots, n$, ce qui prouve que E est somme directe des E_i .

Chapitre 8 : ESPACES VECTORIELS DE DIMENSION FINIE

Les espaces vectoriels ont été définis dans le Chapitre 6 et nous avons mis en évidence quelques-unes de leurs principales propriétés.

Dans ce chapitre, nous complétons cette étude par l'examen des espaces vectoriels de dimension finie. Nous verrons en particulier qu'un espace vectoriel de dimension finie peut être complètement définie par sa dimension. En fait, un K -espace vectoriel de dimension n est isomorphe à K^n .

En plus des propriétés générales des applications linéaires que nous avons étudiées au Chapitre 7, nous obtiendrons d'autres propriétés valables lorsque les espaces sont de dimension finie. Ainsi nous verrons que si E et F sont des K -espaces vectoriels de dimension finie, alors le K -espace vectoriel $\mathcal{L}(E, F)$ est de dimension finie.

8.1. Le théorème de la dimension

Rappelons qu'un K -espace vectoriel E est dit de **dimension finie** s'il possède une famille génératrice finie. Dans le cas contraire, on dit que E est de **dimension infinie**.

Pour l'instant, nous n'avons défini que la locution « dimension finie » mais nous verrons bientôt ce qu'est la dimension d'un K -espace vectoriel.

8.1.1. EXISTENCE DE BASES EN DIMENSION FINIE

8.1.1.1. Théorème

Soient E un K -espace vectoriel non nul de dimension finie, G une partie génératrice finie de E et L une partie libre contenue dans G . Alors, il existe une base B de E telle que $L \subset B \subset G$.

Démonstration. Soit $\mathcal{L}$ l'ensemble des parties libres X de E telles que $L \subset X \subset G$. $\mathcal{L}$ n'est pas vide car $L \in \mathcal{L}$ et $\mathcal{L}$ est fini puisque, G étant fini, l'ensemble $\mathcal{P}(G)$ des parties de G est finie et $\mathcal{L} \subset \mathcal{P}(G)$. En outre tout élément de $\mathcal{L}$ possède un nombre fini d'éléments. Choisissons dans $\mathcal{L}$ une partie B ayant le plus grand nombre possible d'éléments. Soit p ce nombre et montrons que B est une base

de E . Il suffit de montrer que B est une partie génératrice de E puisque par construction B est une partie libre de E .

Comme G engendre E , il suffit de voir que tout élément de G est combinaison linéaire des éléments de B .

Si $\text{card}(G) = p$, alors on a $B = G$. La partie génératrice G de E est donc libre. Par conséquent G est une base de E et le théorème est démontré.

Supposons maintenant que $p < \text{card}(G)$. Si x est un élément de G n'appartenant pas à B , l'ensemble $B \cup \{x\}$ est contenu dans G et possède $(p+1)$ éléments. Donc la famille $B \cup \{x\}$ est liée. Si $B = \{x_1, \dots, x_p\}$, il existe des scalaires $\lambda_1, \dots, \lambda_p, \lambda$ non tous nuls tels que $\lambda_1 x_1 + \dots + \lambda_p x_p + \lambda x = 0$.

On a nécessairement $\lambda \neq 0$, car si $\lambda = 0$ les éléments de B vérifieraient une combinaison linéaire nulle à coefficients non tous nuls et B ne serait pas libre. On en déduit

$$x = -\frac{\lambda_1}{\lambda} x_1 - \dots - \frac{\lambda_p}{\lambda} x_p.$$

Ainsi tout élément $x \in G$ est combinaison linéaire des éléments de B et le théorème est démontré. $\square$

Le théorème 8.1.1.1 admet des conséquences importantes que nous allons examiner.

8.1.1.2. Corollaire

Tout espace vectoriel E de dimension finie, non réduit à $\{0\}$, admet une base finie.

En effet, par définition, E admet une partie génératrice finie $G = \{x_1, \dots, x_n\}$ et il existe un $x_i \in G$ non nul puisque $E \neq \{0\}$. Il suffit de poser $L = \{x_i\}$ et d'appliquer le Théorème 8.1.1.1 à L et à G .

8.1.1.3. Corollaire (Théorème de la base incomplète)

Soit E un espace vectoriel de dimension finie sur un corps K . Pour toute partie libre $\{x_1, \dots, x_p\}$ de E , il existe des vecteurs $y_1, y_2, \dots, y_q$ de E tels que $(x_1, x_2, \dots, x_p, y_1, \dots, y_q)$ soit une base de E .

Démonstration. Posons $L = \{x_1, \dots, x_p\}$ et soit G une famille génératrice finie de E . Alors $L \cup G$ est une famille génératrice finie de E et on a $L \subset L \cup G$. D'après le Théorème 8.1.1.1, il existe une base B telle que $L \subset B \subset L \cup G$; on peut mettre B sous la forme $B = L \cup H$, où H est une partie de G .

8.1.2. DIMENSION

Nous venons de voir que dans un espace vectoriel de dimension finie, il existe toujours une base finie. Nous allons démontrer maintenant que toutes les

bases ont le même nombre d'éléments. Pour cela, nous aurons besoin d'utiliser le lemme suivant.

8.1.2.1. Lemme

Soient E un K -espace vectoriel et n vecteurs $x_1, \dots, x_n$ de E . Si les vecteurs $y_1, \dots, y_{n+1}$ sont des combinaisons linéaires de $x_1, \dots, x_n$, alors la famille $\{y_1, \dots, y_{n+1}\}$ est liée.

Démonstration. Raisonnons récurrence sur n .

Pour $n = 1$, il existe $\alpha_1, \alpha_2 \in K$ tels que $y_1 = \alpha_1 x_1$ et $y_2 = \alpha_2 x_1$.

On en déduit, par élimination de x_1 :

$$\alpha_1 y_2 - \alpha_2 y_1 = 0.$$

Si les deux scalaires α_1 et α_2 ne sont pas nuls en même temps, la famille $\{y_1, y_2\}$ est liée. Si $\alpha_1 = \alpha_2 = 0$, on a $y_1 = y_2 = 0$ et la famille $\{y_1, y_2\}$ est encore liée.

Le lemme est donc vrai si $n = 1$. Supposons-le établi pour $n - 1$ vecteurs et démontrons-le pour n vecteurs.

Soient $y_1, y_2, \dots, y_{n+1}$ des combinaisons linéaires de $x_1, \dots, x_n$. On peut écrire

$$(8.1.2.1) \quad y_1 = z_1 + \alpha_1 x_n, \dots, y_{n+1} = z_{n+1} + \alpha_{n+1} x_n$$

où les vecteurs $z_1, \dots, z_{n+1}$ sont des combinaisons linéaires de $x_1, x_2, \dots, x_{n-1}$.

Envisageons les deux cas suivants :

1^{er} cas : $\alpha_1 = \alpha_2 = \dots = \alpha_{n+1} = 0$.

Alors $y_1, y_2, \dots, y_{n+1}$ sont des combinaisons linéaires de $x_1, \dots, x_{n-1}$. D'après l'hypothèse de récurrence, les n vecteurs $y_1, \dots, y_n$ sont linéairement dépendants ; *à fortiori*, les vecteurs $y_1, \dots, y_{n+1}$ sont linéairement dépendants et le lemme est démontré dans ce cas.

2^e cas : L'un au moins des α_i est non nul.

En changeant au besoin l'ordre de la numérotation, on peut supposer que $\alpha_{n+1} \neq 0$. Alors x_n est combinaison linéaire de $y_{n+1}, x_1, \dots, x_{n-1}$:

$$x_n = \frac{1}{\alpha_{n+1}} (y_{n+1} - z_{n+1}).$$

En prenant dans les n premières relations (8.1.2.1), on voit que $y_1, y_2, \dots, y_n$ sont des combinaisons linéaires de $y_{n+1}, x_1, \dots, x_{n-1}$:

$$y_1 = z_1 + \frac{\alpha_1}{\alpha_{n+1}} (y_{n+1} - z_{n+1}), \dots, y_n = z_n + \frac{\alpha_n}{\alpha_{n+1}} (y_{n+1} - z_{n+1}),$$

d'où, pour $1 \leq i \leq n$,

$$y_i - \frac{\alpha_i}{\alpha_{n+1}} y_{n+1} = z_i - \frac{\alpha_i}{\alpha_{n+1}} z_{n+1}$$

ce qui montre que les n vecteurs $y_i - \frac{\alpha_i}{\alpha_{n+1}} y_{n+1}$, $1 \leq i \leq n$, sont des combinaisons linéaires des $n - 1$ vecteurs $x_1, \dots, x_{n-1}$. D'après l'hypothèse de récurrence, il existe des scalaires $\lambda_1, \dots, \lambda_n$ non tous nuls tels que

$$\lambda_1 \left(y_1 - \frac{\alpha_1}{\alpha_{n+1}} y_{n+1} \right) + \dots + \lambda_n \left(y_n - \frac{\alpha_n}{\alpha_{n+1}} y_{n+1} \right) = 0.$$

Si on pose

$$\lambda_{n+1} = -\frac{\lambda_1 \alpha_1 + \dots + \lambda_n \alpha_n}{\alpha_{n+1}},$$

on obtient

$$\lambda_1 y_1 + \dots + \lambda_{n+1} y_{n+1} = 0$$

où les scalaires $\lambda_1, \dots, \lambda_{n+1}$ ne sont pas tous nuls. La famille $\{y_1, \dots, y_{n+1}\}$ est donc liée et le lemme est démontré. $\square$

8.1.2.2. Théorème

Dans un espace vectoriel E de dimension finie, toutes les bases ont le même nombre d'éléments.

Démonstration. On sait déjà que tout espace vectoriel E de dimension finie possède au moins une base finie. Soient $B = (x_1, \dots, x_n)$ et $B' = (e_1, \dots, e_m)$ deux bases de E .

Comme $e_1, \dots, e_m$ sont des combinaisons linéaires de $x_1, \dots, x_n$, si on avait $m > n$, la suite $(e_1, \dots, e_m)$ serait liée d'après le Lemme 8.1.2.1, ce qui est impossible puisque $(e_1, \dots, e_m)$ est une base. Donc $m \leq n$.

On montrerait de même que $n \leq m$, d'où $n = m$.

Le théorème précédent nous permet de définir la dimension d'un espace vectoriel.

8.1.2.3. Définition

*Soit E un K -espace vectoriel de dimension finie. On appelle **dimension** de E , et on note $\dim_K(E)$ ou $\dim(E)$, le nombre d'éléments d'une base quelconque de E .*

On pose par définition $\dim(\{0\}) = 0$.

8.1.2.4. Remarque

La dimension d'un espace vectoriel dépend du corps de base. Ainsi $\dim_{\mathbb{C}}(\mathbb{C}) = 1$ mais $\dim_{\mathbb{R}}(\mathbb{C}) = 2$. Plus généralement tout $\mathbb{C}$ -espace vectoriel de dimension n est un $\mathbb{R}$ -espace vectoriel de dimension $2n$.

8.1.2.5. Exemple

La base canonique de K^n contient n vecteurs. Donc

$$\dim_K(K^n) = n.$$

8.1.2.6. Exemple

$\dim_K(K_n[X]) = n + 1$ car la base $B = \{1, X, X^2, \dots, X^n\}$ de $K_n[X]$ contient $n + 1$ éléments.

8.1.2.7. Théorème

Soit E un espace vectoriel de dimension finie n sur un corps K . Alors :

a) *Toute partie libre L de E possède au plus n éléments. Si ce nombre d'éléments est égal à n , L est une base de E .*

b) *Toute famille génératrice G de E possède au moins n éléments. Si ce nombre d'éléments est égal à n , G est une base de E .*

Démonstration. a) D'après le Théorème 8.1.1.1, il existe une base B de E telle que $L \subset B$. Comme B possède n éléments, L possède au plus n éléments. Si L possède n éléments, alors $L = B$ et L est une base de E .

b) De même, il existe une base B de E telle que $B \subset G$. Comme B possède n éléments, G possède au moins n éléments. Si G possède n éléments, on a $G = B$ et G est une base de E . $\square$

Donnons maintenant une caractérisation des bases en dimension finie. Ce résultat permet de montrer le plus souvent, qu'une partie B est une base.

8.1.2.8. Théorème

Soit E un K -espace vectoriel de dimension finie n et soit B une partie de E . Les conditions suivantes sont équivalentes :

a) *B est une base de E .*

b) *B est une partie libre de E et B possède n éléments.*

c) *B est une partie génératrice de E ayant n éléments.*

Démonstration. a) $\implies$ b) est évident.

b) $\implies$ c) Si B n'était pas une partie génératrice, on pourrait la compléter pour obtenir une base (Corollaire 8.1.1.3) et cette base aurait au moins $n + 1$ éléments, ce qui est absurde puisque toutes les bases possèdent n éléments.

c) $\implies$ a) Si B n'était pas libre, on pourrait en extraire une base de E , possédant au plus $(n - 1)$ éléments, ce qui contredirait la définition de la dimension.

8.1.3. DIMENSION D'UN SOUS-ESPACE VECTORIEL

8.1.3.1. Théorème

Soit E un espace vectoriel de dimension finie n sur un corps K et soit F un sous-espace vectoriel de E .

a) F est de dimension finie et $\dim_K(F) \leq \dim_K(E)$.

b) $\dim_K(F) = \dim_K(E)$ si et seulement si $F = E$.

c) F admet un supplémentaire F_1 et on a

$$\dim_K(E) = \dim_K(F) + \dim_K(F_1).$$

Démonstration. a) Toute partie libre L de F est une partie libre de E ; donc L possède au plus n éléments. Choisissons une partie libre $B = \{x_1, \dots, x_p\}$ de F dont le nombre p d'éléments est le plus grand possible. On a $p \leq n$.

Montrons que B est une base de F . Pour cela, il suffit de montrer que B est une partie génératrice de F puisqu'elle est libre. Il s'agit de montrer que le sous-espace vectoriel B' de F engendré par B est égale à F . Si on avait $B' \neq F$, il existerait un vecteur x de F qui ne serait pas combinaison linéaire des éléments de B . Montrons que les vecteurs $x, x_1, \dots, x_p$ sont linéairement indépendants. Considérons en effet une relation de la forme

$$\lambda x + \lambda_1 x_1 + \dots + \lambda_p x_p = 0.$$

Si $\lambda \neq 0$, λ^{-1} existe et on a

$$x = -\frac{\lambda_1}{\lambda} x_1 - \dots - \frac{\lambda_p}{\lambda} x_p$$

ce qui est absurde; donc $\lambda = 0$. Il reste $\lambda_1 x_1 + \dots + \lambda_p x_p = 0$, donc $\lambda_1 = \lambda_2 = \dots = \lambda_p = 0$, puisque $x_1, \dots, x_p$ sont linéairement indépendants. Donc la relation $\lambda x + \lambda_1 x_1 + \dots + \lambda_p x_p = 0$ entraîne $\lambda = \lambda_1 = \dots = \lambda_p = 0$ et la famille $\{x, x_1, \dots, x_p\}$ est bien libre. On aboutit à une contradiction puisque par hypothèse p est le nombre maximum de vecteurs linéairement indépendants de F . Donc $B' = F$ et par suite F est de dimension $p \leq n = \dim(E)$.

b) Si $\dim_K(F) = \dim_K(E)$, toute base B de F est une partie libre de E ayant n éléments; c'est donc une base de E d'après le Théorème 8.1.2.8. Alors tout vecteur de E est combinaison linéaire des éléments de B , donc appartient à F . Par conséquent $F = E$.

c) Soit $(x_1, \dots, x_p)$ une base de F . Comme E est de dimension n , il existe $n - p$ vecteurs $x_{p+1}, \dots, x_n$ tels que $(x_1, \dots, x_p, x_{p+1}, \dots, x_n)$ soit une base de E (Corollaire 8.1.1.3).

La famille $\{x_{p+1}, \dots, x_n\}$, sous-famille d'une famille libre de E , est libre. Le sous-espace vectoriel F_1 de E qu'elle engendre est de dimension $n - p$.

D'autre part, tout vecteur $x \in E$ s'écrit de façon unique sous la forme :

$$x = \lambda_1 x_1 + \dots + \lambda_p x_p + \lambda_{p+1} x_{p+1} + \dots + \lambda_n x_n = y + z$$

où $y = \lambda_1 x_1 + \dots + \lambda_p x_p \in F$ et $z = \lambda_{p+1} x_{p+1} + \dots + \lambda_n x_n \in F_1$, ce qui montre que $E = F \oplus F_1$ et

$$\dim_K(E) = \dim_K(F) + \dim_K(F_1).$$

8.1.3.2. Théorème

Soit E un espace vectoriel de dimension finie n sur un corps K et soit F un sous-espace vectoriel de E . Alors l'espace vectoriel quotient E/F est de dimension finie et on a :

$$\dim(E) = \dim(F) + \dim(E/F).$$

On appelle **codimension** de F dans E et on note $\text{codim}_K(F)$ ou $\text{codim}(F)$, la dimension de E/F .

Démonstration. D'après le théorème précédent, F est de dimension finie et $\dim(F) \leq \dim(E)$. Soit $(x_1, \dots, x_m)$ une base de F . D'après le Théorème de la base incomplète, il existe des vecteurs $y_1, \dots, y_p$ de E tels que $(x_1, \dots, x_m, y_1, \dots, y_p)$ soit une base de E . tout vecteur $x \in E$ s'écrit alors de façon unique :

$$x = \lambda_1 x_1 + \dots + \lambda_m x_m + \alpha_1 y_1 + \dots + \alpha_p y_p$$

où les λ_i et les α_i sont des scalaires.

D'où

$$\begin{aligned} x + F &= \lambda_1 x_1 + \dots + \lambda_m x_m + \alpha_1 y_1 + \dots + \alpha_p y_p + F \\ &= \alpha_1 y_1 + \dots + \alpha_p y_p + F \end{aligned}$$

puisque $\lambda_1 x_1 + \dots + \lambda_m x_m \in F$.

Or, d'après la définition de la structure vectorielle de E/F ,

$$\alpha_1 y_1 + \dots + \alpha_p y_p + F = \alpha_1 (y_1 + F) + \dots + \alpha_p (y_p + F),$$

donc les éléments $y_1 + F, \dots, y_p + F$ engendrent l'espace vectoriel quotient E/F . Il reste à montrer que ces classes sont linéairement indépendantes.

Supposons que

$$c_1(y_1 + F) + \dots + c_p(y_p + F) = 0 + F$$

où $0 + F = F$ est l'élément 0 de l'espace E/F . Alors $c_1 y_1 + \dots + c_p y_p \in F$ et le vecteur $c_1 y_1 + \dots + c_p y_p$ est combinaison linéaire des vecteurs $x_1, \dots, x_m$:

$$c_1 y_1 + \dots + c_p y_p = \mu_1 x_1 + \dots + \mu_m x_m.$$

Comme les vecteurs $x_1, \dots, x_m, y_1, \dots, y_p$ sont linéairement indépendants, on en déduit

$$c_1 = c_2 = \dots = c_p = \mu_1 = \dots = \mu_m = 0.$$

Ainsi, l'espace vectoriel quotient E/F admet pour base les p classes $y_1 + F, \dots, y_p + F$ et on a

$$\dim(E/F) = p = (m + p) - m = \dim(E) - \dim(F).$$

8.1.3.3. Remarque

Les Théorèmes 8.1.3.1 et 8.1.3.2 montrent que la codimension d'un sous-espace vectoriel F est égale à la dimension d'un supplémentaire de F .

8.1.3.4. Théorème

Soit E un K -espace vectoriel. Si $E_1, \dots, E_n$ sont des sous-espaces vectoriels de dimension finie, et si $E = E_1 \oplus \dots \oplus E_n$ alors E est de dimension finie et on a

$$\dim(E) = \sum_{i=1}^n \dim(E_i).$$

Démonstration. Posons $p_i = \dim(E_i)$. Soit $B_i = (e_{i1}, \dots, e_{ip_i})$ une base de E_i pour $1 \leq i \leq n$. Comme E est somme directe des E_i , on a (Théorème 6.2.4.4) $E_i \cap \sum_{j \neq i} E_j = \{0\}$ pour tout $i \in \{1, \dots, n\}$ donc $E_i \cap E_j = \{0\}$ si $i \neq j$ et par

suite $B_i \cap B_j = \emptyset$ si $i \neq j$ car le vecteur nul n'appartient à aucune famille libre. D'après l'hypothèse, tout vecteur $x \in E$ s'écrit de façon unique sous la forme :

$$x = (\lambda_{11}e_{11} + \dots + \lambda_{1p_1}e_{1p_1}) + \dots + (\lambda_{n1}e_{n1} + \dots + \lambda_{np_n}e_{np_n}) ;$$

donc (Théorème 6.3.3.2) la suite $(e_{11}, \dots, e_{1p_1}, \dots, e_{n1}, \dots, e_{np_n})$ est une base de E qui est donc de dimension $p_1 + \dots + p_n$.

8.1.3.5. Théorème

Soient $E_1, \dots, E_p$ des K -espaces vectoriels de dimension finie $n_1, \dots, n_p$ respectivement. Alors l'espace vectoriel $E_1 \times \dots \times E_p$ a pour dimension $n_1 + \dots + n_p$.

Démonstrations. a) cas où $p = 2$: Soient E_1 et E_2 deux K -espaces vectoriels de dimension finie n et m respectivement. Soient $(e_1, \dots, e_n)$ une base de E_1 , $(f_1, \dots, f_m)$ une base de E_2 .

Considérons la famille $\mathcal{B}$ de vecteurs de $E_1 \times E_2$:

$$\mathcal{B} = \{(e_1, 0), \dots, (e_n, 0), (0, f_1), \dots, (0, f_m)\}.$$

Tout élément de $E_1 \times E_2$ s'écrit :

$$\left(\sum_{i=1}^n x_i e_i, \sum_{j=1}^m y_j f_j \right)$$

ou, en utilisant les lois définies sur $E_1 \times E_2$:

$$\left(\sum_{i=1}^n x_i e_i, 0 \right) + \left(0, \sum_{j=1}^m y_j f_j \right) = \sum_{i=1}^n x_i (e_i, 0) + \sum_{j=1}^m y_j (0, f_j),$$

ce qui montre que la famille $\mathcal{B}$ engendre l'espace vectoriel $E_1 \times E_2$.

De plus, la famille $\mathcal{B}$ est libre dans $E_1 \times E_2$. Soit en effet

$$\alpha_1 (e_1, 0) + \dots + \alpha_n (e_n, 0) + \lambda_1 (0, f_1) + \dots + \lambda_m (0, f_m) = 0$$

une combinaison linéaire nulle des vecteurs de $\mathcal{B}$. On a, en remontant les calculs précédents :

$$(\alpha_1 e_1 + \dots + \alpha_n e_n, \lambda_1 f_1 + \dots + \lambda_m f_m) = (0, 0).$$

On en déduit

$$\alpha_1 e_1 + \dots + \alpha_n e_n = 0 \text{ et } \lambda_1 f_1 + \dots + \lambda_m f_m = 0,$$

d'où $\alpha_1 = \dots = \alpha_n = \lambda_1 = \dots = \lambda_m = 0$.

La famille $\mathcal{B}$ des $n + m$ vecteurs de $E_1 \times E_2$ est donc une base de $E_1 \times E_2$, donc $\dim(E_1 \times E_2) = \dim(E_1) + \dim(E_2)$.

b) Cas général : Raisonnons par récurrence sur p . Nous venons de voir que la formule est vraie pour $p = 2$. Supposons-la vraie pour $p - 1$. Si on pose $F = E_1 \times \dots \times E_{p-1}$ on a

$$\dim(E_1 \times \dots \times E_{p-1} \times E_p) = \dim(F \times E_p) = \dim(F) + \dim(E_p).$$

Il suffit d'appliquer l'hypothèse de récurrence pour obtenir la formule annoncée.

8.2. Applications linéaires en dimension finie

8.2.1. DIMENSION DE $\mathcal{L}(E, F)$

8.2.1.1. Théorème

Soient E et F deux espaces vectoriels de dimension finie sur le même corps K . Alors $\mathcal{L}(E, F)$ est de dimension finie et on a :

$$\dim_K (\mathcal{L}(E, F)) = \dim_K(E) \cdot \dim_K(F).$$

Démonstration. Notons n la dimension de E et p la dimension de F . Soient $(e_1, \dots, e_n)$ une base de E et $(f_1, \dots, f_p)$ une base de F .

Nous savons (Théorème 7.2.3.4) qu'une application linéaire est parfaitement déterminée par les images des vecteurs d'une base de l'espace de départ. On peut donc définir une famille de np application linéaires de E dans F , indexée par $\{1, \dots, n\} \times \{1, \dots, p\}$ en posant :

$$u_{ij}(e_k) = \begin{cases} f_j & \text{si } k = i \\ 0 & \text{si } k \neq i \end{cases}$$

Montrons que la famille $(u_{ij})_{1 \leq i \leq n, 1 \leq j \leq p}$ est une base de $\mathcal{L}(E, F)$.

a) La famille (u_{ij}) engendre $\mathcal{L}(E, F)$: soit $u \in \mathcal{L}(E, F)$. Pour tout $x \in E$, on peut écrire de manière unique $x = \sum_{i=1}^n x_i e_i$, et $u(e_i) = \sum_{j=1}^p a_{ij} f_j$ pour $1 \leq i \leq n$, où les a_{ij} sont des scalaires.

D'où

$$\begin{aligned} u(x) &= \sum_{i=1}^n x_i u(e_i) = \sum_{i=1}^n x_i \left(\sum_{j=1}^p a_{ij} f_j \right) \\ &= \sum_{i=1}^n \sum_{j=1}^p a_{ij} x_i f_j. \end{aligned}$$

Comme $u_{ij}(x) = \sum_{k=1}^n x_k u_{ij}(e_k) = x_i f_j$, on a pour tout vecteur $x \in E$:

$$u(x) = \sum_{i=1}^n \sum_{j=1}^p a_{ij} u_{ij}(x),$$

c'est-à-dire

$$u = \sum_{i=1}^n \sum_{j=1}^p a_{ij} u_{ij}$$

et $(u_{ij})_{1 \leq i \leq n, 1 \leq j \leq p}$ est une famille génératrice de $\mathcal{L}(E, F)$.

b) La famille (u_{ij}) est libre : soit (λ_{ij}) une famille de scalaires telle que

$$\sum_{i=1}^n \sum_{j=1}^p \lambda_{ij} u_{ij} = 0.$$

Alors, pour tout $x \in E$, on a

$$\sum_{i=1}^n \sum_{j=1}^p \lambda_{ij} u_{ij}(x) = 0.$$

En particulier, pour tout $k \in [1, n]$, on a

$$\sum_{i=1}^n \sum_{j=1}^p \lambda_{ij} u_{ij}(e_k) = 0,$$

ce qui donne, d'après la définition des u_{ij} : $\sum_{j=1}^p \lambda_{kj} f_j = 0$.

Comme $(f_j)_{1 \leq j \leq p}$ est une famille libre de F , on obtient $\lambda_{kj} = 0$ pour tout $k \in [1, n]$ et tout $j \in [1, p]$, donc $(u_{ij})_{1 \leq i \leq n, 1 \leq j \leq p}$ est une famille libre de $\mathcal{L}(E, F)$.

La famille (u_{ij}) est donc une base de $\mathcal{L}(E, F)$. Comme elle comporte np éléments, le théorème est démontré. $\square$

8.2.1.2. Corollaire

Soit E un K -espace vectoriel de dimension finie n . Alors l'espace vectoriel $\mathcal{L}(E)$ des endomorphismes de E est de dimension finie n^2 .

C'est un cas particulier du Théorème 8.2.1.1 pour $E = F$.

8.2.2. ESPACES VECTORIELS ISOMORPHES

8.2.2.1. Théorème

Soient E et F deux espaces vectoriels de dimension finie sur le corps K . Alors E et F sont isomorphes si et seulement si $\dim_K(E) = \dim_K(F)$.

Démonstration. Supposons que E et F soient isomorphes. Par hypothèse, il existe une application linéaire bijective u de E sur F . Soit $(e_i)_{1 \leq i \leq n}$ une base de E ; alors $(u(e_i))_{1 \leq i \leq n}$ est une base de F (Théorème 7.2.3.3). Donc F et E ont même dimension.

Réciproquement, supposons que $\dim(E) = \dim(F) = n$. Soient (e_i) une base de E et (f_j) une base de F . Par hypothèse elles ont le même nombre d'éléments. Soit u l'application linéaire définie par $u(e_i) = f_i$ pour tout $i \in [1, n]$. D'après le Théorème 7.2.3.4, u est bijective, donc E et F sont isomorphes.

8.2.2.2. Corollaire

Soit E un espace vectoriel de dimension finie n sur un corps K . Alors E est isomorphe à K^n .

Il suffit d'appliquer le Théorème 8.2.2.1 en se souvenant que $\dim_K(K^n) = n$.

8.2.2.3. Corollaire

Soit E un espace vectoriel de dimension finie sur un corps K et soit F un sous-espace vectoriel de E . Alors tout sous-espace supplémentaire de F dans E est isomorphe à l'espace quotient E/N .

En effet, si F_1 est un supplémentaire de F dans E , on sait (Théorèmes 8.1.3.1) et 8.1.3.2 que

$$\dim(E) = \dim(F) + \dim(F_1) = \dim(F) + \dim(E/F)$$

d'où $\dim(F_1) = \dim(E/F)$.

8.2.3. RANG D'UNE APPLICATION LINÉAIRE,
D'UNE FAMILLE DE VECTEURS

8.2.3.1. Définition

*Soient E un K -espace vectoriel de dimension finie et F un K -espace vectoriel quelconque. Soit $u \in \mathcal{L}(E, F)$. On appelle **rang** de u , et on note $rg(u)$, la dimension sur K de $\text{Im}(u)$.*

Si $(e_1, \dots, e_n)$ est une base de E , les vecteurs $u(e_1), \dots, u(e_n)$ engendrent $\text{Im}(u)$. D'après le Théorème 8.1.2.7, il existe une base de $\text{Im}(u)$ comportant $rg(u)$ éléments. Ainsi le rang de u est le nombre maximum de vecteurs linéairement indépendants extraits de la suite $u(e_1), \dots, u(e_n)$. On a donc $rg(u) \leq n = \dim(E)$.

Plus généralement, nous poserons la définition suivante :

8.2.3.2. Définition

*On appelle **rang** d'une famille $\{x_1, \dots, x_n\}$ de vecteurs d'un K -espace vectoriel E de dimension finie ou non, la dimension du sous-espace vectoriel E' de E engendré par ces vecteurs, c'est-à-dire le nombre maximum de vecteurs linéairement indépendants que l'on peut extraire de la suite $(x_1, \dots, x_n)$.*

Le théorème suivant appelé **théorème noyau image** ou encore **théorème du rang** est d'un usage courant.

8.2.3.3. Théorème

Soient E un K -espace vectoriel de dimension finie, F un K -espace vectoriel quelconque, et $u \in \mathcal{L}(E, F)$. Alors :

$$\dim(E) = \dim(\text{Ker}(u)) + \dim(\text{Im}(u)).$$

Démonstration. D'après le Théorème 8.1.3.1, $\text{Ker}(u)$ admet un sous-espace supplémentaire W dans E . Soit

$$v : W \longrightarrow \text{Im}(u)$$

la restriction de u à W . Il est clair que v est linéaire. Montrons qu'elle est bijective.

Par définition, si $x \in \text{Ker}(v)$, on a $u(x) = 0$, donc $x \in \text{Ker}(u)$, d'où $x \in W \cap \text{Ker}(u) = \{0\}$; donc $x = 0$. L'application v est donc injective.

Montrons que v est surjective. Soit $y \in \text{Im}(u)$. Il existe $x \in E$ tel que $u(x) = y$. Or x s'écrit de façon unique sous la forme :

$$x = x_1 + x_2 \text{ avec } x_1 \in \text{Ker}(u) \text{ et } x_2 \in W.$$

D'où, puisque $u(x_1) = 0$, $u(x) = u(x_1 + x_2) = u(x_1) + u(x_2) = u(x_2)$.

Par suite $y = u(x_2) = v(x_2)$, et v est une application linéaire surjective sur $\text{Im}(u)$.

On a ainsi démontré que tout supplémentaire de $\text{Ker}(u)$ est isomorphe à $\text{Im}(u)$. D'après le Théorème 8.2.2.1, W a même dimension que $\text{Im}(u)$. Comme $\dim(E) = \dim(W) + \dim(\text{Ker}(u))$, on a bien $\dim(E) = \dim(\text{Im}(u)) + \dim(\text{Ker}(u))$.

8.2.3.4. Corollaire

Soient E et F deux K -espaces vectoriels de dimension finie et soit $u \in \mathcal{L}(E, F)$. Alors :

- a) $\text{rg}(u) \leq \dim(E)$; $\text{rg}(u) = \dim(E)$ si et seulement si u est injective.
- b) $\text{rg}(u) \leq \dim(F)$; $\text{rg}(u) = \dim(F)$ si et seulement si u est surjective.

Démonstration. a) On a (Théorème du rang) :

$$\text{rg}(u) = \dim(E) - \dim(\text{Ker}(u)) \leq \dim(E).$$

$\text{rg}(u) = \dim(E)$ si et seulement si $\dim(\text{Ker}(u)) = 0$, c'est-à-dire si et seulement si $\text{Ker}(u) = \{0\}$; donc $\text{rg}(u) = \dim(E)$ si et seulement si u est injective (Théorème 7.2.2.3).

b) De même $\text{rg}(u) = \dim(u(E)) \leq \dim(F)$ et $\text{rg}(u) = \dim(F)$ si et seulement si $\dim(u(E)) = \dim(F)$, c'est-à-dire $u(E) = F$. Finalement, $\text{rg}(u) = \dim(F)$ si et seulement si u est surjective.

8.2.3.5. Corollaire

Soient E et F deux K -espaces vectoriels de même dimension finie n et soit u une application linéaire de E dans F . Les propriétés suivantes sont équivalentes :

- a) u est un isomorphisme de E sur F .
- b) u est injective.

c) u est surjective.

d) u est de rang n .

Démonstration. Les conditions b) et d) sont équivalentes d'après le Corollaire 8.2.3.4 ; de même c) et d) sont équivalentes. Donc b), c) et d) sont équivalentes.

Il est clair que a) $\implies$ b).

Montrons que b) $\implies$ a). Supposons u injective ; alors $\text{Ker}(u) = \{0\}$, donc $\dim(\text{Ker}(u)) = 0$. D'après le Théorème 8.2.3.3, $\dim(\text{Im}(u)) = \dim(E) = \dim(F)$, d'où $\text{Im}(u) = F$ et u est surjective, donc bijective. $\square$

8.2.3.6. Remarque

La proposition précédente est très importante dans la pratique car elle permet d'affirmer (en particulier) qu'un endomorphisme est bijectif dès qu'on sait qu'il est injectif ou surjectif. Mais ce n'est plus vrai en dimension infinie. Ainsi l'application linéaire $D : K[X] \longrightarrow K[X]$ qui, à tout polynôme, associe sa dérivée est surjective mais non injective.

8.2.3.7. Corollaire

Soient E_1 et E_2 deux sous-espaces vectoriels d'un K -espace vectoriel E de dimension finie. On a

$$\dim(E_1 + E_2) = \dim(E_1) + \dim(E_2) - \dim(E_1 \cap E_2).$$

Démonstration. Considérons l'application $u : E_1 \times E_2 \longrightarrow E$ définie par

$$u(x_1, x_2) = x_1 + x_2, \quad x_1 \in E_1, \quad x_2 \in E_2.$$

Elle est évidemment linéaire et $\text{Im}(u) = E_1 + E_2$ par définition de la somme de deux sous-espaces vectoriels. D'autre part :

$$\begin{aligned} \text{Ker}(u) &= \{(x, y) \in E_1 \times E_2 : x + y = 0\} \\ &= \{(x, -x) : x \in E_1 \cap E_2\}. \end{aligned}$$

L'application $x \longmapsto (x, -x)$ de $E_1 \cap E_2$ sur $\text{Ker}(u)$ est un isomorphisme d'espaces vectoriels (pourquoi ?), d'où $\dim(\text{Ker}(u)) = \dim(E_1 \cap E_2)$. Alors, d'après le Théorème 8.1.3.5 et le Théorème 8.2.3.3 appliqué à u , on a

$$\begin{aligned} \dim(E_1) + \dim(E_2) &= \dim(E_1 \times E_2) \\ &= \dim(\text{Ker}(u)) + \dim(\text{Im}(u)) \\ &= \dim(E_1 \cap E_2) + \dim(E_1 + E_2), \end{aligned}$$

d'où notre assertion. $\square$

8.3. Dualité

8.3.1. DUAL D'UN ESPACE VECTORIEL

Rappelons qu'une forme linéaire sur un K -espace vectoriel E est une application linéaire de E dans K .

D'après le Théorème 7.3.2.2, l'ensemble $\mathcal{L}(E, K)$ des formes linéaires sur E est un espace vectoriel. On l'appelle l'espace vectoriel dual ou simplement le dual de E et on le note E^* .

8.3.1.1. Notation

Si $x \in E$ et $u \in E^*$, posons $u(x) = \langle u, x \rangle$. Alors, comme u est linéaire et E^* est un espace vectoriel, on a

$$\begin{aligned} \langle u, x + y \rangle &= \langle u, x \rangle + \langle u, y \rangle \\ (8.3.1.1) \quad \langle u, \lambda x \rangle &= \lambda \langle u, x \rangle \\ \langle u + v, x \rangle &= \langle u, x \rangle + \langle v, x \rangle \\ \langle \lambda u, x \rangle &= \lambda \langle u, x \rangle. \end{aligned}$$

quels que soient $u, v \in E^*$, $x, y \in E$ et $\lambda \in K$.

8.3.1.2. Exemple

Soit E l'espace vectoriel des fonctions continues sur un intervalle $[a, b]$ de $\mathbb{R}$ et à valeurs réelles. En posant

$$u(f) = \int_a^b f(x) dx$$

pour toute $f \in E$, on obtient une forme linéaire sur E .

Si E est un K -espace vectoriel de dimension finie, le Théorème 8.2.1.1 montre que E^* est de dimension finie égale à la dimension de E puisque $\dim_K(K) = 1$. Nous allons construire une base de E^* .

Supposons E de dimension finie n et soit $(e_1, \dots, e_n)$ une base de E . Donnons-nous n scalaires $\alpha_1, \dots, \alpha_n$ quelconques ; d'après le Théorème 7.2.3.4, il existe une forme linéaire u et une seule telle que

$$u(e_1) = \alpha_1, u(e_2) = \alpha_2, \dots, u(e_n) = \alpha_n.$$

On peut donc définir n formes linéaires $e_1^*, e_2^*, \dots, e_n^*$ en posant

$$\langle e_i^*, e_j \rangle = 1 \text{ si } i = j \text{ et } \langle e_i^*, e_j \rangle = 0 \text{ si } i \neq j$$

c'est-à-dire $\langle e_i^*, e_j \rangle = \delta_{ij}$ où le symbole δ_{ij} appelé symbole de Kronecker est défini par

$$\delta_{ij} = \begin{cases} 1 & \text{si } i = j \\ 0 & \text{si } i \neq j \end{cases}$$

8.3.1.3. Théorème

Soit E un K -espace vectoriel de dimension n et soit $(e_1, \dots, e_n)$ une base de E . Les n formes linéaires $e_1^*, \dots, e_n^*$ de E^* définie par $\langle e_i^*, e_j \rangle = \delta_{ij}$ forment une base de E^* , donc E^* est de dimension n .

Démonstration. a) Les e_i^* engendrent E^* . Soit u une forme linéaire quelconque. Il existe n scalaires $\alpha_1, \alpha_2, \dots, \alpha_n$ tels que

$$u(e_1) = \alpha_1, u(e_2) = \alpha_2, \dots, u(e_n) = \alpha_n.$$

Alors pour tout $x = \sum_{j=1}^n \lambda_j e_j \in E$, on a

$$u(x) = \lambda_1 u(e_1) + \dots + \lambda_n u(e_n) = \lambda_1 \alpha_1 + \dots + \lambda_n \alpha_n.$$

Or par définition de e_i^* , on a :

$$e_i^*(x) = \sum_{j=1}^n \lambda_j e_i^*(e_j) = \lambda_i,$$

donc

$$u(x) = \alpha_1 e_1^*(x) + \dots + \alpha_n e_n^*(x) = (\alpha_1 e_1^* + \dots + \alpha_n e_n^*)(x)$$

et par suite

$$u = \alpha_1 e_1^* + \dots + \alpha_n e_n^*$$

ce qui montre que les e_i^* engendrent E^* .

b) Les e_i^* sont linéairement indépendantes dans E^* .

Soit $(a_1, \dots, a_n)$ une suite de scalaires tels que :

$$a_1 e_1^* + a_2 e_2^* + \dots + a_n e_n^* = 0.$$

Alors pour tout $x \in E$, on a

$$a_1 e_1^*(x) + a_2 e_2^*(x) + \dots + a_n e_n^*(x) = 0.$$

En particulier pour chaque $j \in [1, n]$ nous aurons

$$\sum_{i=1}^n a_i e_i^*(e_j) = \sum_{i=1}^n a_i \delta_{ij} = a_j = 0,$$

d'où $a_1 = a_2 = \dots = a_n = 0$.

Ainsi la relation $a_1 e_1^* + \dots + a_n e_n^* = 0$ entraîne $a_1 = \dots = a_n = 0$ et la famille $\{e_1^*, \dots, e_n^*\}$ est libre. Elle forme donc une base de E^* .

8.3.1.4. Définition

La base $(e_1^*, \dots, e_n^*)$ de E^* associée à la base $(e_1, \dots, e_n)$ de E s'appelle la **base duale de la base** $(e_1, \dots, e_n)$ de E .

8.3.1.5. Remarques

a) Si $x = \sum_{j=1}^n \lambda_j e_j \in E$, on a

$$e_i^*(x) = \sum_{j=1}^n \lambda_j e_i^*(e_j) = \lambda_i.$$

Pour cette raison, on dit que e_i^* est la $i^{\text{ème}}$ **forme coordonnée** relative à la base $(e_1, \dots, e_n)$.

b) Si $u \in E^*$, alors $u = \alpha_1 e_1^* + \dots + \alpha_n e_n^*$, où $\alpha_i = \langle u, e_i \rangle$.

c) $\langle u, x \rangle = 0$ pour tout $u \in E^*$ $\iff x = 0$.

En effet, l'hypothèse montre que $\langle e_i^*, x \rangle = 0$ pour tout indice $i \in \{1, 2, \dots, n\} \iff$ (les coordonnées de x sont nulles).

8.3.2. BIDUAL D'UN ESPACE VECTORIEL

8.3.2.1. Définition

Soit E un K -espace vectoriel. On appelle **bidual** de E et on note E^{**} l'espace vectoriel dual de E^* .

Si E est de dimension finie, on a d'après le Théorème 8.3.1.3.

$$\dim(E^{**}) = \dim(E^*) = \dim(E).$$

Nous allons définir une application de E dans E^{**} de la manière suivante : soit $x \in E$; à tout $u \in E^*$ associons sa valeur au point x . On définit ainsi une application T_x de E^* dans K et d'après les formules (8.3.1.1), T_x est une forme linéaire sur E^* , autrement dit $T_x \in E^{**}$.

Posons

$$J(x) = T_x \text{ et } \langle u, J(x) \rangle = \langle u, x \rangle.$$

D'après les formules (8.3.1.1), l'application $J : E \longrightarrow E^{**}$ est linéaire. On l'appelle le **morphisme canonique** de E dans E^{**} .

Nous allons montrer que si E est de dimension finie, alors le morphisme canonique est un isomorphisme.

8.3.2.2. Théorème

Soit E un K -espace vectoriel de dimension finie n . Alors le morphisme canonique $x \mapsto T_x$ de E dans E^{**} est un isomorphisme.

Démonstration. Nous savons déjà que $\dim(E) = \dim(E^{**})$. Il suffit donc, d'après le Corollaire 8.2.3.5, de montrer que J est injectif. On a les équivalences suivantes :

$$\begin{aligned} x \in \text{Ker}(J) &\iff J(x) = 0 \iff \langle u, J(x) \rangle = 0 \text{ pour tout } u \in E^* \\ &\iff \langle u, x \rangle = 0 \text{ pour tout } u \in E^* \iff x = 0. \end{aligned}$$

La dernière équivalence logique résulte des Remarques 8.3.1.5.

Ainsi J est injective, donc bijective.

En général, on identifie un espace vectoriel E de dimension finie à son bidual E^{**} à l'aide de l'isomorphisme J . On identifie le vecteur x et la forme linéaire T_x définie à partir de x . Ainsi x peut être considérée comme une forme linéaire sur E^* , ce qui permet d'écrire indifféremment :

$$u(x) = \langle u, x \rangle = \langle x, u \rangle.$$

Ainsi un espace vectoriel de dimension finie et son bidual jouent des rôles parfaitement symétriques.

8.3.2.3. Remarque

Si E n'est pas de dimension finie, on démontre que le morphisme canonique est toujours injectif, mais pas nécessairement surjectif et E est seulement un sous-espace vectoriel de E^{**} .

8.3.3. ORTHOGONALITÉ

8.3.3.1. Définition

Soient E un K -espace vectoriel, $x \in E$ et $u \in E^*$. On dit que u et x sont orthogonaux si $\langle u, x \rangle = 0$.

On dit qu'une partie A de E est orthogonale à une partie B de E^* si tout élément de A est orthogonal à tout élément de B .

Si A est une partie de E , on appelle **orthogonal de A dans E^*** , et on note A° , l'ensemble des formes linéaires u sur E qui vérifient $\langle u, x \rangle = 0$ pour tout $x \in A$.

De même si B est une partie de E^* , l'orthogonal de B est l'ensemble

$$(8.3.3.1) \quad B_1 = \{x \in E : \langle u, x \rangle = 0 \text{ pour tout } u \in B\}.$$

Les formules (8.3.1.1) montrent immédiatement que pour toute partie A de E , l'ensemble A° est un sous-espace vectoriel de E^* et pour toute partie B de E^* , l'ensemble B_1 est un sous-espace vectoriel de E .

8.3.3.2. Définition

Soient E un K -espace vectoriel de dimension finie n . On appelle **hyperplan vectoriel** de E tout sous-espace de dimension $(n - 1)$.

Le théorème suivant donne une caractérisation des hyperplans à l'aide des formes linéaires.

8.3.3.3. Théorème

Soient E un K -espace vectoriel de dimension finie n et soit u une forme linéaire non nulle sur E . Alors le noyau de u est un hyperplan vectoriel de E . Réciproquement, si H est un hyperplan vectoriel de E , il existe une forme linéaire non nulle u dont le noyau est H , et toute forme linéaire v dont le noyau est H est proportionnelle à u .

Démonstration. Le sous-espace vectoriel $\text{Im}(u)$ de K est de dimension 1. D'après le Théorème 8.2.3.3, $\dim(\text{Ker}(u)) = n - 1$, donc $\text{Ker}(u)$ est un hyperplan.

Réciproquement, soit H un hyperplan vectoriel de E et soit $(e_1, \dots, e_{n-1})$ une base de H que nous complétons en une base $(e_1, \dots, e_n)$ de E . Considérons la forme linéaire u définie par

$$\begin{aligned} u(e_i) &= 0 \text{ pour } i = 1, 2, \dots, n-1 \\ u(e_n) &= 1. \end{aligned}$$

Il est clair que $\text{Ker}(u) = H$. Si v est une forme linéaire telle que $\text{Ker}(v) = H$, posons $\alpha = v(e_n)$ et considérons la forme linéaire $v - \alpha u$. α est non nul puisque $H \neq E$.

On a d'une part $(v - \alpha u)(e_n) = 0$ et, d'autre part $(v - \alpha u)(e_i) = 0$ pour $i = 1, \dots, n-1$.

Donc $v - \alpha u = 0$, soit $v = \alpha u$.

8.3.3.4. Théorème

Soient E un K -espace vectoriel de dimension finie n , F un sous-espace vectoriel de dimension m . Alors :

a) F° est de dimension $n - m$; autrement dit :

$$\dim(F) + \dim(F^\circ) = \dim(E) ;$$

b) $(F^\circ)^\circ = F$.

Démonstration. a) Soit $(e_1, \dots, e_m)$ une base de F . D'après le théorème de la base incomplète, il existe des vecteurs $e_{m+1}, \dots, e_n$ de E tels que $(e_1, \dots, e_n)$ soit une base de E . Soit $(e_1^*, \dots, e_n^*)$ la base duale de E^* .

Pour qu'un vecteur $f = \alpha_1 e_1^* + \dots + \alpha_n e_n^*$ de E^* soit orthogonal à F , il faut et il suffit que pour tout $x \in F$, on ait $f(x) = 0$, et en particulier $f(e_i) = 0$ pour $i = 1, 2, \dots, m$. Comme $f(e_i) = \alpha_i$, on voit que $f = \alpha_{m+1} e_{m+1}^* + \dots + \alpha_n e_n^*$. Réciproquement, tout élément f de E^* de la forme précédente vérifie $f(x) = 0$ pour tout $x \in F$, donc est dans F° . Donc F° admet pour base $(e_{m+1}^*, \dots, e_n^*)$, d'où le a).

b) Soit x un élément de F ; x est orthogonal à F° , donc $x \in (F^\circ)^\circ$, et par suite $F \subset (F^\circ)^\circ$. D'après le a), on a $\dim [(F^\circ)^\circ] = n - (n - m) = m = \dim(F)$. Le Théorème 8.1.3.1, b), appliqué à F et $(F^\circ)^\circ$ montre alors que $F = (F^\circ)^\circ$.

8.3.4. TRANSPOSÉE D'UNE APPLICATION LINÉAIRE

Soient E et F deux K -espaces vectoriels, et u une application linéaire de E dans F . Nous nous proposons de définir une application linéaire de F^* dans E^* . Pour cela, considérons une forme linéaire y' sur F . Alors l'application composée $y' \circ u$ est une forme linéaire sur E , c'est-à-dire un élément de E^* .

Ainsi, étant donnée une application linéaire u de E dans F , à tout élément y' de F^* il correspond un élément $y' \circ u$ de E^* . On peut donc définir une application ${}^t u : F^* \longrightarrow E^*$ en posant

$${}^t u(y') = y' \circ u \quad \text{quel que soit } y' \in F^*.$$

8.3.4.1. Définition

Soient E et F deux K -espaces vectoriels E^* et F^* leurs espaces duaux et $u : E \longrightarrow F$ une application linéaire. On appelle **transposée** de u , et on note ${}^t u$, l'application linéaire ${}^t u : F^* \longrightarrow E^*$ définie par ${}^t u(y') = y' \circ u$ pour tout $y' \in F^*$.

On a par définition

$$({}^t u(y'))(x) = (y' \circ u)(x) = y'(u(x))$$

pour tout $x \in E$, ou encore

$$(8.3.4.1) \quad \langle {}^t u(y'), x \rangle = \langle y', u(x) \rangle$$

quels que soient $x \in E$ et $y' \in F^*$.

La formule (8.3.4.1) s'appelle la **formule fondamentale de la transposition**.

Montrons que l'application ${}^t u$ est linéaire. Soient $y'_1, y'_2 \in F^*$. Pour tout $x \in E$, on a

$$\langle {}^t u(y'_1 + y'_2), x \rangle = \langle y'_1 + y'_2, u(x) \rangle;$$

d'après (8.3.4.1) on obtient

$$\begin{aligned} \langle y'_1 + y'_2, u(x) \rangle &= \langle y'_1, u(x) \rangle + \langle y'_2, u(x) \rangle \\ &= \langle {}^t u(y'_1), x \rangle + \langle {}^t u(y'_2), x \rangle = \langle {}^t u(y'_1) + {}^t u(y'_2), x \rangle. \end{aligned}$$

D'où

$${}^t u(y'_1 + y'_2) = {}^t u(y'_1) + {}^t u(y'_2).$$

On vérifierait de même que ${}^t u(\lambda y') = \lambda {}^t u(y')$ pour tout $y' \in F^*$ et pour tout $\lambda \in K$.

8.3.4.2. Théorème

L'application $u \mapsto {}^t u$ possède les propriétés suivantes :

a) Soient E, F deux K -espaces vectoriels, $u, v \in \mathcal{L}(E, F)$ et $\lambda \in K$; on a

$${}^t(u + v) = {}^t u + {}^t v ; {}^t(\lambda u) = \lambda {}^t u.$$

b) Si E et F sont de dimension finie, et si $u \in \mathcal{L}(E, F)$ on a ${}^t({}^t u) = u$.

c) Soient E, F, G trois K -espaces vectoriels, $u \in \mathcal{L}(E, F)$ et $v \in \mathcal{L}(F, G)$. On a ${}^t(vou) = {}^t u {}^t v$.

d) Soit E un K -espace vectoriel et soit $u \in \mathcal{L}(E)$. Si u est un automorphisme de E , ${}^t u$ est aussi un automorphisme de E^* et $({}^t u)^{-1} = {}^t(u^{-1})$. La transposée de l'application identique de E est l'application identique de E^* .

Démonstration.

a) Pour tout $y' \in F^*$ et pour tout $x \in E$, on a

$$\begin{aligned} \langle {}^t(u + v)(y'), x \rangle &= \langle y', (u + v)(x) \rangle = \langle y', u(x) + v(x) \rangle \\ &= \langle y', u(x) \rangle + \langle y', v(x) \rangle \\ &= \langle {}^t u(y'), x \rangle + \langle {}^t v(y'), x \rangle \\ &= \langle ({}^t u + {}^t v)(y'), x \rangle, \end{aligned}$$

d'où

$${}^t(u + v)(y') = ({}^t u + {}^t v)(y')$$

et par suite

$${}^t(u + v) = {}^t u + {}^t v.$$

On vérifie de même que ${}^t(\lambda u) = \lambda {}^t u$.

b) Puisque $u \in \mathcal{L}(E, F)$, on a ${}^t u \in \mathcal{L}(F^*, E^*)$, donc ${}^t({}^t u) \in \mathcal{L}(E^{**}, F^{**})$.

Soient donc $x \in E^{**} = E$ et $y' \in F^*$. On a

$$\langle {}^t({}^t u)(x), y' \rangle = \langle x, {}^t u(y') \rangle = \langle u(x), y' \rangle,$$

d'où ${}^t({}^t u)(x) = u(x)$ pour tout $x \in E$ et par suite ${}^t({}^t u) = u$.

c) Soient $x \in E$ et $z' \in G^*$. On a, en appliquant la formule (8.3.4.1),

$$\begin{aligned} \langle {}^t(vou)(z'), x \rangle &= \langle z', v(u(x)) \rangle = \langle {}^t v(z'), u(x) \rangle \\ &= \langle {}^t u({}^t v(z')), x \rangle, \end{aligned}$$

d'où ${}^t(vou)(z') = {}^t u({}^t v(z'))$ et par suite ${}^t(vou) = {}^t u {}^t v$.

d) Pour établir (iv) montrons d'abord que si u est l'application identique de E , alors ${}^t u$ est l'application identique de E^* . En effet, pour toute forme linéaire y' sur E et pour tout $x \in E$, on a

$$\langle {}^t u(y'), x \rangle = \langle y', u(x) \rangle = \langle y', x \rangle,$$

d'où ${}^t u(y') = \text{Id}_{E^*}(y')$ et par suite ${}^t u = \text{Id}_{E^*}$.

Si maintenant u est un automorphisme de E , soit u^{-1} l'automorphisme réciproque. On a $u \circ u^{-1} = u^{-1} \circ u = \text{Id}_E$.

On en déduit, en appliquant c) :

$${}^t(u^{-1}) \circ {}^t u = {}^t u \circ {}^t(u^{-1}) = {}^t(\text{Id}_E) = \text{Id}_{E^*},$$

ce qui montre que ${}^t u$ est bien un automorphisme de E^* et que $({}^t u)^{-1} = {}^t(u^{-1})$.

8.3.4.3. Théorème

Soient E et F deux K -espaces vectoriels de dimension finie, $u : E \longrightarrow F$ une application linéaire et ${}^t u$ sa transposée. Alors :

a) Le noyau de ${}^t u$ est égal à l'orthogonal de $u(E)$ dans F^* .

b) Le noyau de u est égal à l'orthogonal de ${}^t u(F^*)$ dans E .

c) Le rang de u est égal au rang de ${}^t u$.

Démonstration. a) Soit $y' \in F^*$. On a les équivalences logiques suivantes :

$$\begin{aligned} y' \in \text{Ker}({}^t u) &\iff {}^t u(y') = 0 \iff \langle {}^t u(y'), x \rangle = 0 \text{ pour tout } x \in E \\ &\iff \langle y', u(x) \rangle = 0 \text{ pour tout } x \in E \iff y' \in (u(E))^\circ. \end{aligned}$$

Donc

$$\text{Ker}({}^t u) = (u(E))^\circ.$$

L'assertion b) se démontre de la même manière.

Démontrons c). On sait (Théorème 8.2.3.3), que

$$\text{rg}(u) = \dim(E) - \dim(\text{Ker}(u)).$$

Comme $\dim(E) = \dim(E^*)$ et $\text{Ker}(u) = ({}^t u(F^*))^\circ$, il vient, en appliquant le Théorème 8.3.3.4,

$$\text{rg}(u) = \dim(E^*) - \dim(({}^t u(F^*))^\circ) = \dim({}^t u(F^*)) = \text{rg}({}^t u).$$

Chapitre 9 : MATRICES

Si E et F sont des K -espaces vectoriels de dimension finie, il résulte de la linéarité, qu'une application linéaire u de E dans F est entièrement définie par la donnée des images par u des vecteurs d'une base B de E .

Si E et F sont munis de bases B et C respectivement, on est ainsi conduit à la notion de matrice d'une application linéaire relativement aux bases B et C . Alors la structure de l'ensemble des applications linéaires de E dans F détermine les opérations sur les matrices : addition, multiplication par un scalaire, et multiplication. En particulier, nous obtiendrons l'algèbre des matrices carrées associé à l'algèbre des endomorphismes d'un K -espace vectoriel E .

Ce chapitre est essentiel non seulement pour la suite du cours, mais également pour les applications aux autres sciences.

Dans tout ce chapitre, on désigne par K un corps commutatif quelconque.

9.1. Généralités

9.1.0.1. Définition

Soit K un corps commutatif. On appelle **matrice à m lignes et n colonnes** ou **matrice de type (m, n) à coefficients dans K** , un tableau rectangulaire d'éléments de K de la forme

$$M = \begin{pmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \dots & \dots & \dots & \dots \\ a_{m1} & a_{m2} & \dots & a_{mn} \end{pmatrix}.$$

Les a_{ij} s'appellent les **éléments** de la matrice ; l'élément a_{ij} est situé à l'intersection de la $i^{\text{ème}}$ ligne et de la $j^{\text{ème}}$ colonne de M . On note souvent

$$M = (a_{ij})_{1 \leq i \leq m, 1 \leq j \leq n}$$

ou simplement $M = (a_{ij})$ s'il n'y a pas de confusion possible.

Si $K = \mathbb{R}$, on dit que la matrice M est **réelle** ; si $K = \mathbb{C}$, elle est dite **complexe**.

Si $m = n$, on dit que M est une **matrice carrée d'ordre n** ; les éléments $a_{11}, a_{22}, \dots, a_{nn}$ forment alors la **diagonale principale** de la matrice.

Si $n = 1$, on obtient la **matrice colonne** :

$$M = \begin{pmatrix} a_{11} \\ a_{21} \\ \vdots \\ a_{m1} \end{pmatrix}$$

Si $m = 1$, on obtient la **matrice ligne** :

$$M = (a_{11} \ a_{12} \ \dots \ a_{1n}).$$

Si $M = (a_{ij})$ est une matrice de type (m, n) on appelle **transposée** de M , et on note tM , la matrice ${}^tM = (b_{ij})$ de type (n, m) définie par $b_{ij} = a_{ji}$.

tM est obtenue en permutant les lignes et les colonnes de M . On a : ${}^t({}^tM) = M$.

Une matrice carrée $M = (a_{ij})$ est dite **symétrique** si $a_{ij} = a_{ji}$ quels que soient i et j , c'est-à-dire si $M = {}^tM$.

Une matrice carrée $M = (a_{ij})$ est dite **antisymétrique** si ${}^tM = -M$.

La relation ${}^tM = -M$ implique $a_{ji} = -a_{ij}$, d'où $a_{ii} = -a_{ii}$ et par suite $a_{ii} = 0$ si le corps K est de caractéristique $\neq 2$.

Si $M = (a_{ij})$ est une matrice complexe de type (m, n) la matrice $\overline{M} = (\overline{a_{ij}})$ est appelée la **matrice conjuguée** de M , $\overline{a_{ij}}$ étant le nombre complexe conjugué de a_{ij} . On appelle **adjointe** de M , et on note M^* , la matrice $M^* = {}^t(\overline{M}) = \overline{{}^tM}$. Si M est une matrice carrée et si $M^* = M$, on dit que M est **hermitienne**.

On dit qu'une matrice carrée $M = (a_{ij})$ est **triangulaire supérieure** si les éléments situés au-dessous de la diagonale principale sont nuls, c'est-à-dire si $a_{ij} = 0$ pour $i > j$. On définit de même une matrice **triangulaire inférieure**.

Une matrice **diagonale** est une matrice carrée $M = (a_{ij})$ telle que $a_{ij} = 0$ si $i \neq j$. Si $\lambda_1, \dots, \lambda_n$ désignent les éléments diagonaux d'une matrice diagonale D , nous écrirons

$$D = \text{diag}(\lambda_1, \dots, \lambda_n).$$

On appelle **matrice scalaire** une matrice diagonale dont tous les éléments de la diagonale principale sont égaux.

On appelle **matrice unité d'ordre n** , et on note I_n , la matrice scalaire d'ordre n dont les éléments de la diagonale principale sont égaux à un.

9.1.0.2. Exemple

$$I_3 = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix}.$$

Ces quelques définitions montrent déjà l'importance des matrices carrées.

L'ensemble des matrices de type (m, n) à coefficients dans le corps K se note $M_{m,n}(K)$ et l'ensemble des matrices carrées d'ordre n à coefficients dans K se note $M_n(K)$.

9.2. Matrice d'une application linéaire

9.2.1. DÉFINITIONS, EXEMPLES ET THÉORÈMES

9.2.1.1. Définition

Soient E et F deux K -espaces vectoriels de dimension n et m respectivement, $B = (e_1, \dots, e_n)$ une base de E , $C = (f_1, \dots, f_m)$ une base de F et $u \in \mathcal{L}(E, F)$. On appelle **matrice de u par rapport aux bases B et C** , et on note $\text{Mat}_{B,C}(u)$, la matrice (a_{ij}) de type (m, n) dont la $j^{\text{ème}}$ colonne est constituée par les coordonnées du vecteur $u(e_j)$ par rapport à la base $(f_1, \dots, f_m)$.

On a donc

$$(9.2.1.1) \quad u(e_j) = a_{1j}f_1 + a_{2j}f_2 + \dots + a_{mj}f_m.$$

Si $E = F$ et si $B = C$, on écrit $\text{Mat}_{B,B}(u) = \text{Mat}_B(u)$, et on dit que $\text{Mat}_B(u)$ est la **matrice de l'endomorphisme u par rapport à la base B** .

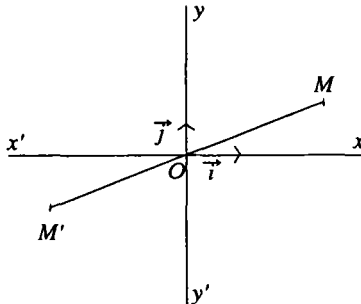
On notera que $\text{Mat}_{B,C}(u)$ dépend en général de B et C même lorsque $B = C$.

9.2.1.2. Remarque

Le nombre de lignes de la matrice de l'application linéaire $u : E \longrightarrow F$ est égale à la dimension de F ; le nombre de ses colonnes est égal à la dimension de E .

9.2.1.3. Exemple

Dans le plan rapporté à des axes rectangulaires, soient $\vec{i} = (1, 0)$ et $\vec{j} = (0, 1)$ la base canonique. Cherchons la matrice de la symétrie u par rapport à 0.



On a $u(\vec{i}) = -\vec{i} = (-1, 0)$; $u(\vec{j}) = -\vec{j} = (0, -1)$.

La matrice M de u dans la base $(\vec{i}, \vec{j})$ est donc :

$$M = \begin{pmatrix} -1 & 0 \\ 0 & -1 \end{pmatrix}.$$

9.2.1.4. Théorème

Soient E et F deux K -espaces vectoriels de dimension n et m respectivement, $B = (e_1, \dots, e_n)$ une base de E , $C = (f_1, \dots, f_m)$ une base de F . L'application qui, à tout $u \in \mathcal{L}(E, F)$, fait correspondre sa matrice $\text{Mat}_{B,C}(u)$ par rapport aux bases B et C est une application bijective de $\mathcal{L}(E, F)$ sur l'ensemble $M_{m,n}(K)$ des matrices de type (m, n) .

Démonstration. Si u et u' sont deux éléments de $\mathcal{L}(E, F)$ tels que $\text{Mat}_{B,C}(u) = \text{Mat}_{B,C}(u')$, alors on a pour tout indice j , $u(e_j) = u'(e_j)$ et donc par linéarité $u(x) = u'(x)$ pour tout $x \in E$, i.e. $u = u'$ et l'application est injective.

Soit $M = (a_{ij})$ une matrice de type (m, n) à coefficients dans K ; considérons les n vecteurs y_j de F :

$$y_j = a_{1j}f_1 + a_{2j}f_2 + \dots + a_{mj}f_m, \quad 1 \leq j \leq n.$$

On sait (Théorème 7.2.3.4), qu'il existe une application linéaire unique $u : E \rightarrow F$ telle que $u(e_j) = y_j$ pour tout indice j , donc telle que $\text{Mat}_{B,C}(u) = M$; donc l'application $u \mapsto \text{Mat}_{B,C}(u)$ est surjective, ce qui achève la démonstration du théorème.

D'après ce théorème, toute matrice M de type (m, n) peut être considérée comme la matrice d'une unique application linéaire u_M de K^n dans K^m relativement aux bases canoniques de ces espaces. On dit que u_M est l'application linéaire associée à M . On peut donc définir l'image, le noyau et le rang de M comme étant respectivement l'image, le noyau et le rang de u_M . En particulier, l'image de u_M est le sous-espace vectoriel de K^m engendré par les images des éléments de la base canonique de K^n , c'est-à-dire par les colonnes de M . Ces remarques nous permettent de définir le rang de M comme étant la dimension du sous-espace vectoriel de K^m engendré par les n colonnes de M ; c'est aussi le nombre de colonnes de M qui sont linéairement indépendantes.

Nous verrons plus tard comment calculer le rang de M .

9.2.1.5. Théorème

Soient E et F deux K -espaces vectoriels de dimension n et m respectivement, $(e_1, \dots, e_n)$ une base de E , $(f_1, \dots, f_m)$ une base de F , et $u \in \mathcal{L}(E, F)$. Alors le rang de u est égal au rang de la matrice de u par rapport aux bases $(e_1, \dots, e_n)$ et $(f_1, \dots, f_m)$.

Démonstration. Soit $A = (a_{ij})$ la matrice de u par rapport aux bases $(e_1, \dots, e_n)$ et $(f_1, \dots, f_m)$. Il s'agit de comparer $\text{rg}(A)$ et $\text{rg}(u)$.

Soient $(e'_1, \dots, e'_n)$ la base canonique de K^n , $(f'_1, \dots, f'_m)$ la base canonique de K^m et u_A l'application linéaire de K^n dans K^m associée à A . Soit φ l'unique isomorphisme de E sur K^n défini par $\varphi(e_j) = e'_j$ pour $1 \leq j \leq n$ et soit ψ l'unique isomorphisme de F sur K^m défini par $\psi(f_i) = f'_i$ pour $1 \leq i \leq m$.

Montrons d'abord que le diagramme

$$\begin{array}{ccc} E & \xrightarrow{u} & F \\ \varphi \downarrow & & \downarrow \psi \\ K^n & \xrightarrow{u_A} & K^m \end{array}$$

est commutatif.

Pour tout $j \in [1, n]$, on a

$$\begin{aligned} (\psi \circ u)(e_j) &= \psi(u(e_j)) = \psi\left(\sum_{i=1}^m a_{ij} f_i\right) = \sum_{i=1}^m a_{ij} \psi(f_i) = \sum_{i=1}^m a_{ij} f'_i \\ &= u_A(e'_j) = u_A(\varphi(e_j)) = (u_A \circ \varphi)(e_j). \end{aligned}$$

Par linéarité, on en déduit que $\psi \circ u = u_A \circ \varphi$. D'où

$$\text{rg}(\psi \circ u) = \text{rg}(u_A \circ \varphi),$$

soit

$$\dim(\psi(u(E))) = \dim(u_A(\varphi(E)))$$

ou, puisque φ et ψ sont bijectives,

$$\dim(u(E)) = \dim(u_A(K^n)).$$

Par suite, on a $\text{rg}(u) = \text{rg}(u_A) = \text{rg}(A)$.

9.2.1.6. Théorème

Soient E un K -espace vectoriel de dimension finie n , $B = (e_1, \dots, e_n)$ une base de E et $\{x_1, \dots, x_p\}$ une famille de p vecteurs de E définie par

$$x_j = \sum_{i=1}^n a_{ij} e_i \quad (1 \leq j \leq p).$$

Alors le rang de la famille $\{x_1, \dots, x_p\}$ est égal au rang de la matrice $M = (a_{ij})_{1 \leq i \leq n, 1 \leq j \leq p}$.

Démonstration. Soit u l'application linéaire de K^p dans E qui envoie les éléments $(e'_1, \dots, e'_p)$ de la base canonique de K^p sur les vecteurs x_i donnés :

$$u(e'_j) = x_j \quad (1 \leq j \leq p).$$

La matrice de l'application linéaire u par rapport aux bases $(e'_1, \dots, e'_p)$ et $(e_1, \dots, e_n)$ est justement $M = (a_{ij})$ et d'après le Théorème 9.2.1.4, on a $rg(u) = rg(M)$.

Or, le rang de la famille $\{x_1, \dots, x_p\}$ est la dimension du sous-espace vectoriel engendré par cette famille et d'après la construction de u , ce sous-espace est identique à $\text{Im}(u)$. Donc le rang de la famille $\{x_1, \dots, x_p\}$ est égal au rang de u , c'est-à-dire au rang de la matrice M .

9.2.1.7. Théorème

Conservons les notations du Théorème 9.2.1.4. Soient $u \in \mathcal{L}(E, F)$ et $M = (a_{ij})$ la matrice de u par rapport aux bases $(e_1, \dots, e_n)$ et $(f_1, \dots, f_m)$. Alors si x est un vecteur de E défini par ses coordonnées $(x_1, \dots, x_n)$ par rapport à la base $(e_1, \dots, e_n)$ les coordonnées $(y_1, \dots, y_m)$ de $y = u(x)$ par rapport à la base $(f_1, \dots, f_m)$ sont données par les formules

$$\begin{aligned} y_1 &= a_{11}x_1 + a_{12}x_2 + \dots + a_{1n}x_n \\ y_2 &= a_{21}x_1 + a_{22}x_2 + \dots + a_{2n}x_n \\ &\vdots \\ y_m &= a_{m1}x_1 + a_{m2}x_2 + \dots + a_{mn}x_n. \end{aligned}$$

Démonstration. On a

$$y = u(x) = u \left(\sum_{j=1}^n x_j e_j \right) = \sum_{j=1}^n x_j u(e_j).$$

D'après la formule (9.2.1.1), cette expression vaut

$$\sum_{j=1}^n x_j \left(\sum_{i=1}^m a_{ij} f_i \right).$$

En regroupant les termes par rapport aux f_i , il vient :

$$y = (a_{11}x_1 + \dots + a_{1n}x_n)f_1 + \dots + (a_{m1}x_1 + \dots + a_{mn}x_n)f_m.$$

D'où

$$y_i = a_{i1}x_1 + \dots + a_{in}x_n$$

pour tout $i \in [1, m]$.

9.2.1.8. Théorème

Soient E et F deux K -espaces vectoriels, $(e_1, \dots, e_n)$ une base de E , $(f_1, \dots, f_m)$ une base de F , $(e_1^*, \dots, e_n^*)$ et $(f_1^*, \dots, f_m^*)$ les bases de E^* et F^* respectivement duales des bases données. Alors si $u \in \mathcal{L}(E, F)$ et si A est la matrice de u par rapport aux bases $(e_1, \dots, e_n)$ et $(f_1, \dots, f_m)$, la matrice de la transposée ${}^t u$ par rapport aux bases $(f_1^*, \dots, f_m^*)$ et $(e_1^*, \dots, e_n^*)$ est ${}^t A$.

Démonstration. Posons $A = (a_{ij})_{1 \leq i \leq m, 1 \leq j \leq n}$. On a

$$u(e_j) = \sum_{k=1}^m a_{kj} f_k.$$

Par définition de la base duale $(f_1^*, \dots, f_m^*)$, on a

$$\langle f_i^*, f_j \rangle = \delta_{ij} \quad (1 \leq i, j \leq m).$$

Donc

$$\langle f_i^*, u(e_j) \rangle = a_{ij}.$$

Soit (b_{ij}) la matrice de ${}^t u \in \mathcal{L}(F^*, E^*)$ par rapport aux bases $(f_1^*, \dots, f_m^*)$ et $(e_1^*, \dots, e_n^*)$. On a

$${}^t u(f_j^*) = \sum_{k=1}^n b_{kj} e_k^*.$$

D'où, puisque $\langle e_k^*, e_j \rangle = \delta_{kj}$,

$$\langle {}^t u(f_j^*), e_i \rangle = b_{ij}.$$

Comme

$$\langle {}^t u(f_j^*), e_i \rangle = \langle f_j^*, u(e_i) \rangle,$$

d'après la formule (8.3.4.1), on a bien $b_{ij} = a_{ji}$.

9.2.1.9. Remarque

D'après le Théorème 8.3.4.3, toute matrice A de type (m, n) a même rang r que sa transposée ${}^t A$ et r est aussi le nombre maximal de lignes de A linéairement indépendantes.

9.3. Opérations sur les matrices

Nous définirons les opérations sur les matrices à partir des opérations correspondantes sur les applications linéaires.

Soient E et F deux espaces vectoriels sur K , $(e_1, \dots, e_n)$ une base de E , $(f_1, \dots, f_m)$ une base de F , u et v des éléments de $\mathcal{L}(E, F)$. Soient $M(u)$ et $M(v)$ les matrices de u et v respectivement par rapport aux bases $(e_1, \dots, e_n)$ et $(f_1, \dots, f_m)$.

9.3.1. ÉGALITÉ DE DEUX MATRICES

Les deux applications linéaires u et v sont égales si et seulement si $u(e_j) = v(e_j)$ pour tout $j \in [1, n]$, c'est-à-dire

$$a_{1j}f_1 + a_{2j}f_2 + \dots + a_{mj}f_m = b_{1j}f_1 + b_{2j}f_2 + \dots + b_{mj}f_m.$$

Donc $u = v$ si et seulement si $a_{ij} = b_{ij}$, pour $i = 1, \dots, n$ et $j = 1, \dots, m$, ce qui nous conduit à la définition suivante :

9.3.1.1. Définition

On dit que deux matrices $A = (a_{ij})$ et $B = (b_{ij})$ de type (m, n) sont égales si on a

$$a_{ij} = b_{ij}$$

quels que soient i et j .

9.3.2. ADDITION DES MATRICES

Conservons les notations précédentes et notons $S = (s_{ij})$ la matrice de l'application linéaire $u + v$. On a

$$\begin{aligned}(u + v)(e_j) &= u(e_j) + v(e_j) = a_{1j}f_1 + \dots + a_{mj}f_m + b_{1j}f_1 + \dots + b_{mj}f_m \\ &= (a_{1j} + b_{1j})f_1 + \dots + (a_{mj} + b_{mj})f_m.\end{aligned}$$

D'où

$$s_{ij} = a_{ij} + b_{ij}.$$

On peut donc poser la définition suivante :

9.3.2.1. Définition

Si $A = (a_{ij})$ et $B = (b_{ij})$ sont deux matrices de type (m, n) , on appelle **somme de A et de B** , et on note $A + B$, la matrice de type (m, n) dont les éléments sont $a_{ij} + b_{ij}$ quel que soit $(i, j) \in [1, m] \times [1, n]$.

On a donc $M(u + v) = M(u) + M(v)$.

Il est clair que si A et B sont des matrices de type (m, n) , on a

$${}^t(A + B) = {}^tA + {}^tB.$$

On vérifie facilement que l'addition des matrices est une opération associative, commutative, et qu'elle possède un élément neutre noté 0 : c'est la matrice dont tous les éléments sont nuls. Enfin toute matrice $A = (a_{ij})$ possède une matrice opposée notée $-A$ et définie par $-A = (-a_{ij})$. Autrement dit, l'ensemble des matrices de type (m, n) est un groupe abélien pour l'addition des matrices.

9.3.3. MULTIPLICATION D'UNE MATRICE PAR UN SCALAIRE

Avec les notations précédentes, soit $M(\lambda u) = (a'_{ij})$, $\lambda \in K$, la matrice de λu par rapport aux bases $(e_1, \dots, e_n)$ et $(f_1, \dots, f_m)$. On a pour tout $j \in [1, n]$,

$$(\lambda u)(e_j) = \lambda(u(e_j)) = \lambda \left(\sum_{i=1}^m a_{ij} f_i \right) = \sum_{i=1}^m (\lambda a_{ij}) f_i.$$

Donc, quel que soit $(i, j) \in [1, m] \times [1, n]$,

$$a'_{ij} = \lambda a_{ij}$$

d'où la définition

9.3.3.1. Définition

Si $A = (a_{ij})$ est une matrice de type (m, n) , le produit de A par le scalaire λ est la matrice λA obtenue en multipliant tous les éléments de A par λ .

On a donc $M(\lambda u) = \lambda M(u)$.

Il est clair que si A est une matrice de type (m, n) , on a

$${}^t(\lambda A) = \lambda {}^t A.$$

9.3.3.2. Théorème

L'ensemble $M_{m,n}(K)$ des matrices de type (m, n) , muni des opérations $(A, B) \mapsto A + B$ et $(\lambda, A) \mapsto \lambda A$, est un espace vectoriel de dimension mn .

Démonstration. Il est immédiat que si A et B sont des matrices de type (m, n) et si λ et μ sont des scalaires, alors

$$\lambda(A + B) = \lambda A + \lambda B; (\lambda + \mu)A = \lambda A + \mu A$$

$$\lambda(\mu A) = (\lambda \mu)A; 1 \cdot A = A.$$

Comme $M_{m,n}(K)$ est un groupe abélien pour l'addition des matrices, c'est un espace vectoriel sur K .

Si E et F sont des K -espaces vectoriels de dimension n et m respectivement, rapportés à des bases, alors l'application $u \mapsto M(u)$ est un isomorphisme de l'espace vectoriel $\mathcal{L}(E, F)$ sur l'espace vectoriel $M_{m,n}(K)$. Comme $\mathcal{L}(E, F)$ est de dimension mn d'après le Théorème 8.2.1.1, $M_{m,n}(K)$ est aussi de dimension mn .

Pour trouver une base de l'espace vectoriel $M_{m,n}(K)$, considérons les mn matrices E_{ij} de type (m, n) qui contiennent un 1 à l'intersection de la $i^{\text{ème}}$ ligne et de la $j^{\text{ème}}$ colonne et des zéros partout ailleurs.

Si $A = (a_{ij}) \in M_{m,n}(K)$, on a

$$A = \sum_{i=1}^m \sum_{j=1}^n a_{ij} E_{ij},$$

donc les matrices E_{ij} engendrent l'espace vectoriel $M_{m,n}(K)$. Ces matrices sont linéairement indépendantes car la relation $\sum_{i=1}^m \sum_{j=1}^n a_{ij} E_{ij} = 0$ implique $a_{ij} = 0$ quels que soient $i \in [1, m]$ et $j \in [1, n]$.

Donc les $(E_{ij})_{1 \leq i \leq m, 1 \leq j \leq n}$ forment une base de $M_{m,n}(K)$ appelée **base canonique** de $M_{m,n}(K)$.

9.3.4. PRODUIT DE DEUX MATRICES

Soient E, F, G trois espaces vectoriels sur K , $(e_1, \dots, e_n)$, $(f_1, \dots, f_m)$, $(g_1, \dots, g_r)$ des bases de E, F et G respectivement. Soient u un élément de $\mathcal{L}(E, F)$, v un élément de $\mathcal{L}(F, G)$, $M(u) = (a_{ij})$ et $M(v) = (b_{ij})$ les matrices de u et v par rapport aux bases données.

Nous nous proposons de chercher la matrice de vou par rapport aux bases $(e_1, \dots, e_n)$ et $(g_1, \dots, g_r)$.

Posons

$$M(vou) = (c_{ij}).$$

On a quel que soit $j \in [1, n]$,

$$\begin{aligned} (vou)(e_j) &= v(u(e_j)) = v\left(\sum_{k=1}^m a_{kj} f_k\right) = \sum_{k=1}^m a_{kj} v(f_k) \\ &= \sum_{k=1}^m a_{kj} \left(\sum_{i=1}^r b_{ik} g_i\right) = \sum_{k=1}^m \sum_{i=1}^r a_{kj} b_{ik} g_i \\ &= \sum_{i=1}^r \left(\sum_{k=1}^m b_{ik} a_{kj}\right) g_i. \end{aligned}$$

Donc, par définition de c_{ij} , on a

$$c_{ij} = \sum_{k=1}^m b_{ik} a_{kj},$$

d'où la définition suivante :

9.3.4.1. Définition

Soient $A = (a_{ij})$ une matrice de type (m, n) et $B = (b_{ij})$ une matrice de type (n, r) à coefficients dans K . On appelle **produit de A par B** , et on note AB , la matrice de type (m, r) dont l'élément c_{ij} situé à l'intersection de la $i^{\text{ème}}$ ligne et de la $j^{\text{ème}}$ colonne est donnée par :

$$(9.3.4.1) \quad c_{ij} = \sum_{k=1}^n a_{ik} b_{kj}.$$

On a donc $M(vou) = M(v) \cdot M(u)$.

9.3.4.2. Remarque

Le produit AB n'est défini que si le nombre de colonnes de A est égal au nombre de lignes de B .

L'expression de c_{ij} s'obtient à partir de la ligne d'indice i de A et de la colonne d'indice j de B . On dit que l'on fait le produit AB «lignes par colonnes».

9.3.4.3. Exemples

$$\begin{pmatrix} 1 & 1 & 2 \\ 3 & 5 & 1 \end{pmatrix} \begin{pmatrix} 2 & 4 \\ 1 & 6 \\ 3 & 2 \end{pmatrix} = \begin{pmatrix} 9 & 14 \\ 14 & 44 \end{pmatrix}.$$

Propriétés du produit

a) Si A, B, C sont des matrices telles que les différents produits ci-dessous soient définis et si λ est un scalaire quelconque, on a

$$A(B \cdot C) = (A \cdot B)C$$

$$A(B + C) = AB + AC; (B + C)A = BA + CA$$

$$A(\lambda B) = (\lambda A)B = \lambda(AB).$$

Ces propriétés se démontrent en utilisant des propriétés bien connues des applications linéaires. Si par exemple $u : K^r \rightarrow K^s$, $v : K^m \rightarrow K^n$ et $w : K^n \rightarrow K^m$ sont des applications linéaires dont les matrices sont A, B et C respectivement par rapport aux bases canoniques, alors la matrice de $u \circ (v \circ w)$ est $A(BC)$ et la matrice de $(u \circ v) \circ w$ est $(AB)C$. Comme $(u \circ v) \circ w = u \circ (v \circ w)$, on a l'égalité $A(BC) = (AB)C$.

Les autres égalités se démontrent de même.

b) Le produit de deux matrices n'est pas commutatif en général.

c) Puisque $\mathcal{L}(E)$ est un anneau pour l'addition et la composition des applications, l'ensemble $M_n(K)$ est un anneau pour l'addition et la multiplication des matrices, l'élément unité étant la matrice unité d'ordre n . Cet anneau possède des diviseurs de zéro car on a par exemple, si $n = 2$,

$$\begin{pmatrix} 0 & 0 \\ 1 & 0 \end{pmatrix} \begin{pmatrix} 0 & 0 \\ 0 & 1 \end{pmatrix} = \begin{pmatrix} 0 & 0 \\ 0 & 0 \end{pmatrix}.$$

$M_n(K)$ est même une K -algèbre et l'application $u \mapsto M(u)$ de $\mathcal{L}(E)$ dans $M_n(K)$ est un isomorphisme d'algèbres.

d) Si le produit AB est défini, ${}^t B^t A$ l'est aussi et on a

$${}^t(AB) = {}^t B^t A.$$

C'est immédiat.

e) En revenant au Théorème 9.2.1.7, les équations (9.2.1.2) montrent que si

$$X = \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix} \quad \text{et} \quad Y = \begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_m \end{pmatrix}$$

sont les matrices colonnes des composantes de x dans la base $(e_1, \dots, e_n)$ de E et des composantes de $y = u(x)$ dans la base $(f_1, \dots, f_m)$ de F , la $i^{\text{ème}}$ composante de y est le produit de la $i^{\text{ème}}$ ligne de M par la matrice colonne X . On obtient ainsi la traduction matricielle de la relation $y = u(x) : Y = MX$, soit

$$\begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_m \end{pmatrix} = \begin{pmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \dots & \dots & \dots & \dots \\ a_{m1} & a_{m2} & \dots & a_{mn} \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix}$$

9.4. Matrices inversibles. Changement de bases

9.4.1. MATRICES INVERSIBLES

9.4.1.1. Définition

On dit qu'une matrice $M \in M_n(K)$ est **inversible** ou est **régulière** s'il existe une matrice $A \in M_n(K)$ telle que

$$(9.4.1.1) \quad MA = AM = I_n.$$

La matrice A est alors unique ; on la note M^{-1} et on l'appelle l'inverse de M .

L'ensemble des matrices carrées inversibles d'ordre n à coefficients dans K se note $GL(n, K)$. On vérifie facilement que $GL(n, K)$, muni de la multiplication des matrices, est un groupe non commutatif en général. On l'appelle le **groupe linéaire à n variable sur le corps K** ; ce groupe n'est rien d'autre que le groupe multiplicatif des éléments inversibles de l'anneau $M_n(K)$.

Les groupes $GL(n, K)$ et leurs sous-groupes s'appellent les **groupes classiques**. Leur étude a constitué une étape importante le développement de la théorie des groupes dont les applications à l'algèbre, à l'analyse moderne et à la physique sont nombreuses.

9.4.1.2. Théorème

Soient E un K -espace vectoriel de dimension finie n , $(e_1, \dots, e_n)$ une base de E et u un endomorphisme de E . Pour que la matrice $M(u)$ de u par rapport à la base $(e_1, \dots, e_n)$ soit inversible, il faut et il suffit que u soit un automorphisme de E . Alors on a

$$M(u^{-1}) = (M(u))^{-1}.$$

Démonstration. Posons $A = M(u)$. Si A est inversible il existe $B \in M_n(K)$ telle que

$$AB = BA = I_n.$$

L'application $u \mapsto M(u)$ de $\mathcal{L}(E)$ sur $M_n(K)$ étant bijective, on peut considérer l'application linéaire v associée à B . Elle est telle que

$$uov = vov = Id_E,$$

donc u est bijective et v est sa bijection réciproque : $v = u^{-1}$.

Réciproquement, si u est bijective, u^{-1} existe et

$$M(u) M(u^{-1}) = M(uou^{-1}) = M(Id_E) = I_n,$$

$$M(u^{-1}) M(u) = M(u^{-1}ou) = M(Id_E) = I_n.$$

Par suite, $M(u)$ est inversible et $(M(u))^{-1} = M(u^{-1})$.

9.4.1.3. Théorème

Soit $A \in M_n(K)$. Les conditions suivantes sont équivalentes :

- a) A est inversible.
- b) ${}^t A$ est inversible.

c) Les lignes de A (considérées comme vecteurs de K^n) sont linéairement indépendantes.

d) Les colonnes de A (considérées comme vecteurs de K^n) sont linéairement indépendantes.

e) A est de rang n .

Démonstration. a) $\implies$ b) Les relations $AA^{-1} = A^{-1}A = I_n$ impliquent par transposition, ${}^t(A^{-1})^t A = {}^t A {}^t(A^{-1}) = I_n$, donc ${}^t A$ est inversible et $({}^t A)^{-1} = {}^t(A^{-1})$.

b) $\implies$ a) Il suffit d'échanger les rôles de A et ${}^t A$ et d'observer que ${}^t({}^t A) = A$.

On a a) $\iff$ e) d'après le Corollaire 8.2.3.5 car cela revient à dire qu'un endomorphisme u d'un K -espace vectoriel E de dimension finie n est inversible si et seulement si $\text{rg}(u) = n$.

On a c) $\iff$ e) d'après la Remarque 9.2.1.9. Enfin d) $\iff$ e) d'après les remarques qui suivent le Théorème 9.2.1.4, ce qui achève la démonstration du théorème.

9.4.2. CHANGEMENT DE BASES

Soit E un K -espace vectoriel de dimension n muni d'une première base $B = (e_1, \dots, e_n)$ appelée « ancienne base ». Tout vecteur x de E s'écrit de façon unique

$$x = x_1 e_1 + \dots + x_n e_n.$$

On est souvent conduit à calculer les coordonnées $x'_1, x'_2, \dots, x'_n$ du vecteur x par rapport à une autre base $B' = (f_1, \dots, f_n)$ de E appelée « nouvelle base ».

Avant d'examiner ce problème, nous poserons la définition suivante.

9.4.2.1. Définition

Soient E un K -espace vectoriel de dimension n , $B = (e_1, \dots, e_n)$ et $B' = (f_1, \dots, f_n)$ deux bases de E . On appelle **matrice de passage de la base B à la base B'** la matrice carrée P d'ordre n dont la $j^{\text{ème}}$ colonne est formée par les composantes de f_j par rapport à la base B .

Autrement dit, si $P = (a_{ij})$, on a

$$(9.4.2.1) \quad f_j = \sum_{i=1}^n a_{ij} e_i.$$

9.4.2.2. Remarque

Il existe un endomorphisme u de E , et un seul, tel que pour $j \in [1, n]$, on ait :

$$u(e_j) = f_j.$$

Les formules (9.4.2.1) montrent que P est la matrice de u par rapport à la base B . Comme u est un automorphisme de E (puisque'il transforme une base de E en une base de E), P est inversible :

$$P \in GL(n, K).$$

9.4.2.3. Remarque

Les relations (9.4.2.1) et les relations $Id_E(f_j) = f_j$ ($j = 1, 2, \dots, n$) montrent que P est la matrice de l'application identique de E par rapport aux bases B' et B (bien noter l'ordre dans lequel doivent être prises les bases).

9.4.2.4. Théorème

Conservons les notations précédentes. Alors la matrice de passage de la base B' à la base B est P^{-1} .

Démonstration. Désignons par P' la matrice de passage de la base B' à la base B . Considérons le diagramme

$$\begin{array}{ccccc} E & & Id_E & & E \\ \downarrow & \xrightarrow{\quad} & \downarrow & \xrightarrow{\quad} & \downarrow \\ B & & B' & & B \end{array}$$

En passant aux matrices associées relativement aux bases indiquées, ce diagramme se traduit par $PP' = I_n$. De même, on a $P'P = I_n$, donc

$$P' = P^{-1}.$$

9.4.2.5. Remarque

Nous venons de voir qu'une matrice de changement de base est inversible. Nous allons montrer que, réciproquement, une matrice inversible P et une base B d'un espace vectoriel E étant données, P est la matrice de passage de la base B à une base unique B' de E .

Soient en effet $P = (a_{ij})$ une matrice carrée inversible d'ordre n et $B = (e_1, \dots, e_n)$ une base d'un K -espace vectoriel E de dimension n . Si $f_1, \dots, f_n$ sont les n vecteurs de E définis par

$$f_j = a_{1j}e_1 + \dots + a_{nj}e_n,$$

alors l'endomorphisme u de E tel que $u(e_j) = f_j$ admet P pour matrice par rapport à la base B , donc est un automorphisme de E . Donc (Théorème 7.2.3.4 d)), $B' = (f_1, \dots, f_n)$ est une base de E et P est évidemment la matrice de passage de la base B à la base B' .

• Influence d'un changement de bases sur les composantes d'un vecteur.

9.4.2.6. Théorème

Soient E un K -espace vectoriel de dimension n , $B = (e_1, \dots, e_n)$ et $B' = (f_1, \dots, f_n)$ deux bases de E , $P = (a_{ij})$ la matrice de passage de la base B à la base B' . Soient x un vecteur de E , X (resp X') la matrice colonne de ses coordonnées dans l'ancienne base B (resp. dans la nouvelle base B'). On a la relation $X = PX'$.

Démonstration. Posons

$$X = \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix} \quad \text{et} \quad X' = \begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_n \end{pmatrix}.$$

On peut écrire

$$x = \sum_{j=1}^n y_j f_j = \sum_{j=1}^n y_j \left(\sum_{i=1}^n a_{ij} e_i \right) = \sum_{i=1}^n \left(\sum_{j=1}^n a_{ij} y_j \right) e_i.$$

D'autre part, on peut écrire x sous la forme

$$x = \sum_{i=1}^n x_i e_i.$$

En identifiant ces deux dernières expressions de x , on voit que

$$x_i = \sum_{j=1}^n a_{ij} y_j, \quad i = 1, 2, \dots, n,$$

ce qui s'écrit matriciellement :

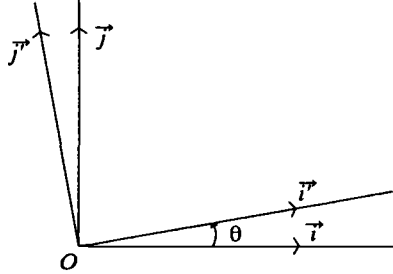
$$X = PX'.$$

Insistons sur le fait que ce sont les « anciennes » composantes $(x_j)_{1 \leq j \leq n}$ que l'on calcule en fonction des « nouvelles » composantes $(y_j)_{1 \leq j \leq n}$ par l'intermédiaire de la matrice de passage. Naturellement, on peut également exprimer X' en fonction de X :

$$X' = P^{-1}X.$$

9.4.2.7. Exemple

Dans le plan vectoriel euclidien, soit $B = (\vec{i}, \vec{j})$ une base orthonormée. Soit $B' = (\vec{i}', \vec{j}')$ la base obtenue en effectuant sur la base B une rotation R_θ d'angle θ .



Soit $\vec{v}$ un vecteur de coordonnées $\begin{pmatrix} x \\ y \end{pmatrix}$ dans la base B et $\begin{pmatrix} x' \\ y' \end{pmatrix}$ dans la base B' . On a

$$\vec{i}' = R_\theta(\vec{i}) = (\cos \theta)\vec{i} + (\sin \theta)\vec{j}$$

$$\vec{j}' = R_\theta(\vec{j}) = R_{\theta+\pi/2}(\vec{i}) = (-\sin \theta)\vec{i} + (\cos \theta)\vec{j}.$$

La matrice de passage de la base B à la base B' est donc

$$P = \begin{pmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{pmatrix}.$$

D'où

$$\begin{pmatrix} x \\ y \end{pmatrix} = P \begin{pmatrix} x' \\ y' \end{pmatrix},$$

soit

$$x = x' \cos \theta - y' \sin \theta$$

$$y = x' \sin \theta + y' \cos \theta.$$

• Influence d'un changement de bases sur la matrice d'une application linéaire.

Un problème important lié aux changements de bases dans les espaces vectoriels, est de savoir ce que devient la matrice d'une application linéaire $u : E \longrightarrow F$ lorsqu'on change de base dans E et dans F .

La solution de ce problème est donnée par le théorème suivant :

9.4.2.8. Théorème

Soient E et F deux K -espaces vectoriels de dimension n et m respectivement, et $u \in \mathcal{L}(E, F)$. Soient $B = (e_1, \dots, e_n)$ et $B' = (e'_1, \dots, e'_n)$ deux bases de E ,

$\mathcal{C} = (f_1, \dots, f_m)$ et $\mathcal{C}' = (f'_1, \dots, f'_m)$ deux bases de F , et soit P (resp. S). La matrice de passage de la base $\mathcal{B}$ à la base $\mathcal{B}'$ (resp. de la base $\mathcal{C}$ à la base $\mathcal{C}'$). Si A est la matrice de u par rapport aux bases $\mathcal{B}$ et $\mathcal{C}$, alors la matrice M de u par rapport aux bases $\mathcal{B}'$ et $\mathcal{C}'$ est donnée par la formule

$$M = S^{-1}AP.$$

Démonstration. Soient $x \in E$ et $y = u(x)$, X et X' les matrices colonnes des coordonnées de x dans les bases $\mathcal{B}$ et $\mathcal{B}'$, Y et Y' les matrices colonnes des coordonnées de y dans les bases $\mathcal{C}$ et $\mathcal{C}'$.

D'après le Théorème 9.4.2.6, on a $X = PX'$ et $Y = SY'$. Mais, par définition de A , on a $Y = AX$, donc

$$SY' = APX', \text{ d'où } Y' = S^{-1}APX'.$$

Comme $Y' = MX'$, on a bien

$$M = S^{-1}AP.$$

9.4.2.9. Remarque

En utilisant la Remarque 9.4.2.2 et le Théorème 9.4.2.4, on peut retrouver le Théorème 9.4.2.8. Le faire à titre d'exercice en considérant le diagramme

$$\begin{array}{ccccccc} E & \xrightarrow{Id_E} & E & \xrightarrow{u} & F & \xrightarrow{Id_F} & F \\ \mathcal{B}' & & \mathcal{B} & & \mathcal{C} & & \mathcal{C}' \end{array}$$

9.4.2.10. Corollaire

Soient E un K -espace vectoriel de dimension n , u un endomorphisme de E et $\mathcal{B} = (e_1, \dots, e_n)$, $\mathcal{C} = (f_1, \dots, f_n)$ deux bases de E . Soient P la matrice de passage de la base $\mathcal{B}$ à la base $\mathcal{C}$ et A la matrice de u par rapport à la base $\mathcal{B}$. Alors la matrice de u par rapport à la base $\mathcal{C}$ est $M = P^{-1}AP$.

Si $F = E$ il suffit, dans le Théorème 9.4.2.8, de prendre la base $\mathcal{C}$ identique à la base $\mathcal{B}$ et la base $\mathcal{C}'$ identique à la base $\mathcal{B}'$. Les deux matrices S et P sont alors identiques et on a bien la relation

$$M = P^{-1}AP.$$

9.4.2.11. Définition

On dit que deux matrices A et M de $M_n(K)$ sont **semblables** s'il existe une matrice $P \in GL(n, K)$ telle que

$$M = P^{-1}AP.$$

La relation de similitude est une relation d'équivalence sur $M_n(K)$.

La définition et le Corollaire 9.4.2.10 montrent que deux matrices carrées A et M d'ordre n sont semblables si et seulement si, étant donné un K -espace vectoriel E de dimension n , A et M représentent le même endomorphisme u de E par rapport à deux bases différentes de E . On pourra donc déterminer une base de E par rapport à laquelle la matrice de u soit aussi simple que possible. Nous examinerons ce problème ultérieurement.

9.4.3. MATRICES ÉQUIVALENTES

9.4.3.1. Définition

On dit que deux matrices A et B de type (m, n) sont équivalentes, s'il existe une matrice inversible P d'ordre m et une matrice inversible Q d'ordre n telle que

$$A = PBQ.$$

Il est clair que l'équivalence des matrices est une relation d'équivalence sur $M_{m,n}(K)$.

9.4.3.2. Théorème

Soient E et F deux K -espaces vectoriels de dimension n et m respectivement, $B = (e_1, \dots, e_n)$ une base de E , $C = (f_1, \dots, f_m)$ une base de F et $u \in \mathcal{L}(E, F)$. Soient B la matrice de u par rapport aux bases B et C et A une matrice de type (m, n) . Pour que A et B soient équivalentes, il faut et il suffit que A soit la matrice de u par rapport à des bases B' de E et C' de F .

Démonstration. Supposons que A et B soient équivalentes. Il existe donc $P \in GL(m, K)$ et $Q \in GL(n, K)$ telles que $A = PBQ$. La matrice Q étant inversible, c'est la matrice de passage de la base B à une base B' de E . De même P^{-1} est la matrice de passage de la base C à une base C' de F . Alors d'après le Théorème 9.4.2.8, la matrice de u par rapport aux bases B' et C' est $(P^{-1})^{-1}BQ = PBQ = A$.

Réciproquement, supposons que A et B soient les matrices de la même application linéaire u par rapport aux bases B' et C' d'une part, B et C d'autre part. Soient P la matrice de passage de la base B à la base B' et S la matrice de passage de la base C à la base C' . D'après le Théorème 9.4.2.8, on a $A = S^{-1}BP$. Les matrices S^{-1} et P étant inversibles, A et B sont équivalentes.

9.4.3.3. Théorème

Soient E et F deux K -espaces vectoriels de dimension n et m respectivement et u une application linéaire de E dans F , de rang r . Alors il existe une base de E et une base de F par rapport auxquelles la matrice de u est de la forme

$$A(m, n, r) = \begin{pmatrix} 1 & 0 & \dots & 0 & : & 0 & 0 & \dots & 0 \\ 0 & 1 & \dots & 0 & : & 0 & 0 & \dots & 0 \\ \dots & \dots & \dots & \dots & : & \dots & \dots & \dots & \dots \\ 0 & 0 & \dots & 1 & : & 0 & 0 & \dots & 0 \\ \dots & \dots & \dots & \dots & : & \dots & \dots & \dots & \dots \\ 0 & 0 & \dots & 0 & : & 0 & 0 & \dots & 0 \\ \dots & \dots & \dots & \dots & : & \dots & \dots & \dots & \dots \\ 0 & 0 & \dots & 0 & : & 0 & \dots & \dots & 0 \end{pmatrix}$$

$\begin{matrix} \updownarrow r \\ \updownarrow m-r \end{matrix}$

$\begin{matrix} \leftarrow r & \leftarrow n-r \end{matrix}$

Démonstration. Soit E' un supplémentaire de $\text{Ker}(u)$ dans E . Soit $(e_1, \dots, e_n)$ une base de E obtenue en complétant une base $(e_1, \dots, e_r)$ de E' par une base $(e_{r+1}, \dots, e_n)$ de $\text{Ker}(u)$. Posons $f_j = u(e_j)$ si $1 \leq j \leq r$. Montrons que les vecteurs $f_1, \dots, f_r$ forment une base de $\text{Im}(u)$. La famille $\{f_1, \dots, f_r\}$ est libre dans $\text{Im}(u)$ car :

$$\sum_{i=1}^r \lambda_i u(e_i) = 0 \implies u \left(\sum_{i=1}^r \lambda_i e_i \right) = 0.$$

Le vecteur $\sum_{i=1}^r \lambda_i e_i$ appartenant à $\text{Ker}(u)$ et à E' est nul, donc $\lambda_1 = \dots = \lambda_r = 0$.

Comme $u(e_i) = 0$ pour tout $i \in \{r+1, \dots, n\}$, les vecteurs $f_1, \dots, f_r$ engendrent $\text{Im}(u)$, d'où notre assertion.

Complétons les vecteurs $f_1, \dots, f_r$ par des vecteurs arbitraires $f_{r+1}, \dots, f_m$ pour obtenir une base $\mathcal{C} = (f_1, \dots, f_r, f_{r+1}, \dots, f_m)$ de F . La matrice de u par rapport aux bases $\mathcal{B}$ et $\mathcal{C}$ a bien la forme indiquée puisque

$$u(e_1) = f_1, \dots, u(e_r) = f_r, u(e_{r+1}) = \dots = u(e_n) = 0.$$

9.4.3.4. Corollaire

Toute matrice de type (m, n) de rang r est équivalente à la matrice $A(m, n, r)$. En particulier, deux matrices de type (m, n) sont équivalentes si et seulement si elles ont le même rang.

Démonstration. La première partie du Corollaire n'est qu'une traduction du Théorème 9.4.3.3.

Si deux matrices A et B de type (m, n) sont équivalentes, elles représentent la même application linéaire u par rapport à des bases différentes (Théorème 9.4.3.2). A et B ont donc le même rang que u .

Réciproquement, si A et B ont pour rang r , elles sont équivalentes à la matrice $A(m, n, r)$ du Théorème 9.4.3.3 ; elles sont donc équivalentes entre elles.

9.4.3.5. Corollaire

Soit A une matrice de type (m, n) . Alors A et tA ont même rang.

Soit r le rang de A . D'après le Corollaire 9.4.3.4, il existe deux matrices inversibles P et Q de type (m, m) et (n, n) respectivement telles que

$$A = PA(m, n, r)Q.$$

En transposant, il vient

$${}^tA = {}^tQ {}^tA(m, n, r) {}^tP.$$

Comme tQ et tP sont inversibles d'après le Théorème 9.4.1.3, tA est équivalente à ${}^tA(m, n, r)$. Mais ${}^tA(m, n, r) = A(n, m, r)$; donc ${}^tA(m, n, r)$ est de rang r et on a $r = rg(A) = rg({}^tA)$.

Chapitre 10 : DÉTERMINANTS

Le lecteur a déjà rencontré la notion de déterminant dans les classes antérieures à propos de l'étude des systèmes d'équations linéaires dans $\mathbb{R}^2$ et dans $\mathbb{R}^3$.

Dans ce chapitre, nous définissons d'abord le déterminant d'un couple de vecteurs, le déterminant d'un endomorphisme relativement à une base d'un espace vectoriel de dimension deux, puis le déterminant d'une matrice carrée d'ordre deux.

Les applications et les formes multilinéaires sont alors introduites, ce qui nous permet de définir les déterminants dans un espace vectoriel de dimension finie quelconque. Nous étudions en particulier, les propriétés des déterminants qui découlent du fait qu'un déterminant est une forme multilinéaire alternée.

Dans ce chapitre, K désignera un corps commutatif de caractéristique différente de 2. Dans les applications, ce corps sera toujours $\mathbb{R}$ ou $\mathbb{C}$.

10.1. Applications et formes bilinéaires

10.1.1. APPLICATIONS ET FORMES BILINÉAIRES ALTERNÉES

10.1.1.1. Définition

*Soient E_1, E_2, F trois espaces vectoriels sur le même corps K . On dit qu'une application f de $E_1 \times E_2$ dans F est **bilinéaire** si, quels que soient les vecteurs x et x' de E_1 , y et y' de E_2 et le scalaire λ , on a*

$$f(x + x', y) = (f(x, y) + f(x', y))$$

$$f(x, y + y') = f(x, y) + f(x, y')$$

$$f(\lambda x, y) = f(x, \lambda y) = \lambda f(x, y).$$

Autrement dit, l'application $(x, y) \mapsto f(x, y)$ est, pour y fixé, linéaire par rapport à la variable x et, pour x fixé, linéaire par rapport à la variable y .

Si $E_1 = E_2 = E$, on dit que f est une **application bilinéaire** sur E . Si, de plus, $F = K$, on dit que f est une **forme bilinéaire** sur E .

10.1.1.2. Remarque

Il résulte immédiatement des définitions que quels que soient les vecteurs x et y de E et le scalaire λ , on a

$$f(x, 0) = 0, \quad f(0, y) = 0, \quad f(\lambda x, \lambda y) = \lambda^2 f(x, y).$$

10.1.1.3. Définition

Soit E un K -espace vectoriel. On dit qu'une forme bilinéaire f sur E est :

- **symétrique** si $f(x, y) = f(y, x)$ quels que soient les vecteurs $x, y \in E$,
- **antisymétrique** si $f(y, x) = -f(x, y)$ quels que soient les vecteurs $y, x \in E$,
- **alternée** si $f(x, x) = 0$ quel que soit le vecteur x de E .

On a la caractérisation suivante des formes bilinéaires alternées.

10.1.1.4. Théorème

Soit E un K -espace vectoriel. Une forme bilinéaire f sur E est alternée si et seulement si elle est antisymétrique.

Démonstration. Si f est alternée, on a quels que soient $x, y \in E$,

$$f(x + y, x + y) = 0.$$

En développant, il vient

$$f(x, x) + f(x, y) + f(y, x) + f(y, y) = 0,$$

d'où, puisque f est alternée, $f(x, y) = -f(y, x)$.

Réciproquement, si f est une forme bilinéaire antisymétrique, en posant $y = x$ dans l'égalité $f(y, x) = -f(x, y)$ on obtient :

$$f(x, x) = -f(x, x),$$

soit $2f(x, x) = 0$, d'où $f(x, x) = 0$ puisque la caractéristique de K est différente de 2.

10.1.1.5. Exemple

Prenons $K = \mathbb{R}$ et $E = \mathbb{R}^2$. L'application $f : E \times E \longrightarrow K$ définie par

$$f((x, x'), (y, y')) = xy' - x'y$$

est une forme bilinéaire alternée.

La vérification est immédiate et est laissée au lecteur.

10.1.2. CAS OÙ $\dim(E) = 2$

Jusqu'à présent, nous n'avons pas supposé que les espaces vectoriels considérés étaient de dimension finie. Nous allons étudier maintenant les formes bilinéaires alternées sur un K -espace vectoriel E de dimension 2.

10.1.2.1. Théorème

Soit E un K -espace vectoriel de dimension 2. L'ensemble des formes bilinéaires alternées sur E est un K -espace vectoriel de dimension 1.

Démonstration. On démontre sans peine que l'ensemble $A_2(E)$ des formes bilinéaires alternées sur E est un sous-espace vectoriel de l'espace vectoriel $\mathcal{F}(E \times E, K)$ des applications de $E \times E$ dans K .

Soient $\mathcal{B} = (e_1, e_2)$ une base de E et (x, y) un couple de vecteurs de E . Si f est une forme bilinéaire alternée sur E , on a

$$\begin{aligned} f(x, y) &= f(x_1 e_1 + x_2 e_2, y_1 e_1 + y_2 e_2) \\ &= x_1 y_1 f(e_1, e_1) + x_1 y_2 f(e_1, e_2) + x_2 y_1 f(e_2, e_1) + x_2 y_2 f(e_2, e_2). \end{aligned}$$

Comme f est alternée, on a

$$\begin{aligned} f(e_1, e_1) &= f(e_2, e_2) = 0 \\ f(e_2, e_1) &= -f(e_1, e_2). \end{aligned}$$

Par suite, il reste

$$f(x, y) = (x_1 y_2 - x_2 y_1) f(e_1, e_2).$$

D'après l'Exemple 10.1.1.5, l'application $h_{\mathcal{B}}$ définie par

$$h_{\mathcal{B}}(x, y) = x_1 y_2 - x_2 y_1$$

est une forme bilinéaire alternée vérifiant :

$$h_{\mathcal{B}}(e_1, e_2) = 1 \cdot 1 - 0 \cdot 0 = 1.$$

On a donc, dans $A_2(E)$, l'égalité

$$f = f(e_1, e_2) h_{\mathcal{B}},$$

ce qui montre que $\{h_{\mathcal{B}}\}$ est une base de $A_2(E)$.

10.1.2.2. Corollaire

Soit E un K -espace vectoriel de dimension 2 et soit $\mathcal{B} = (e_1, e_2)$ une base de E . Il existe une forme bilinéaire alternée f et une seule sur E telle que $f(e_1, e_2) = 1$.

Soit f une forme bilinéaire alternée telle que $f(e_1, e_2) = 1$. On a d'après le Théorème 10.1.2.1, $f = 1 \cdot h_{\mathcal{B}} = h_{\mathcal{B}}$, ce qui montre que f coïncide avec la forme bilinéaire alternée $h_{\mathcal{B}}$.

10.1.3. DÉTERMINANT D'ORDRE 2

10.1.3.1. Définition

Soient E un espace vectoriel sur K de dimension 2 et $B = (e_1, e_2)$ une base de E . Soient x et y des vecteurs de E . On appelle **déterminant de x et y par rapport à la base B** , le scalaire $\det_B(x, y)$. On le note $\det_B(x, y)$.

Nous avons vu que si $x = x_1 e_1 + x_2 e_2$ et $y = y_1 e_1 + y_2 e_2$, alors $\det_B(x, y) = x_1 y_2 - x_2 y_1$. On convient d'écrire

$$(10.1.3.1) \quad \det_B(x, y) = x_1 y_2 - x_2 y_1 = \begin{vmatrix} x_1 & y_1 \\ x_2 & y_2 \end{vmatrix}.$$

10.1.3.2. Remarque

L'application $(x, y) \mapsto \det_B(x, y)$ vérifie évidemment toutes les propriétés des formes bilinéaires alternées. Mais l'intérêt essentiel des déterminants vient du théorème suivant.

10.1.3.3. Théorème

Soient E un K -espace vectoriel de dimension 2 et B une base de E . Alors pour que deux vecteurs x et y de E soient linéairement indépendants, il faut et il suffit que $\det_B(x, y) \neq 0$.

Démonstration. Supposons que x et y soient linéairement dépendants. Alors, on a par exemple $x = \lambda y$ et

$$\det_B(x, y) = \det_B(\lambda y, y) = \lambda \det_B(y, y) = 0.$$

Supposons maintenant que x et y soient linéairement indépendants. Ils forment alors une base B' de E . Puisque $\det_{B'}$ est une forme bilinéaire alternée, il existe d'après le Théorème 10.1.2.1, un scalaire $\lambda \in K$ tel que $\det_{B'}(x, y) = \lambda \det_B(x, y)$ et puisque d'après le Corollaire 10.1.2.2, $\det_{B'}(x, y) = 1$, on obtient $\det_B(x, y) \neq 0$.

10.1.4. DÉTERMINANT D'UN ENDOMORPHISME

Soient E un K -espace vectoriel de dimension 2 et f une forme bilinéaire alternée non nulle sur E . Pour tout endomorphisme u de E , considérons l'application $g : E \times E \rightarrow K$ définie par

$$g(x, y) = f(u(x), u(y))$$

quels que soient x, y dans E .

Il est clair que g est une forme bilinéaire alternée sur E ; de plus on vérifie facilement que l'application φ_u qui à f associe g est un endomorphisme de $A_2(E)$. Comme $\dim(A_2(E)) = 1$, tout endomorphisme de $A_2(E)$ est une homothétie. Il existe donc un scalaire unique λ indépendant de f tel que, quels que soient les vecteurs x et y de E , on ait

$$f(u(x), u(y)) = \lambda f(x, y).$$

Cette remarque nous amène à poser la définition suivante :

10.1.4.1. Définition

Soit u un endomorphisme d'un K -espace vectoriel E de dimension 2. On appelle **déterminant** de u , et on note $\det(u)$, l'unique scalaire tel que pour toute forme bilinéaire alternée f sur E et pour tout couple $(x, y) \in E^2$, on ait

$$(10.1.4.1) \quad f(u(x), u(y)) = (\det(u)) f(x, y).$$

10.1.4.2. Remarque

Si on applique la relation (10.1.4.1) pour une base $\mathcal{B} = (e_1, e_2)$ de E et si on prend $(x, y) = (e_1, e_2)$, on obtient une expression du déterminant de u :

$$(10.1.4.2) \quad \det(u) = \det_{\mathcal{B}}(u(e_1), u(e_2)).$$

On a les propriétés suivantes du déterminant d'un endomorphisme.

10.1.4.3. Théorème

Soit E un K -espace vectoriel de dimension 2. Alors :

- a) $\det(Id_E) = 1$.
- b) Pour tout endomorphisme u de E et pour tout $\lambda \in K$, on a $\det(\lambda u) = \lambda^2 \det(u)$.
- c) Quels que soient les endomorphismes u et v de E , on a $\det(vu) = \det(v) \det(u)$.
- d) Un endomorphisme u de E est inversible si et seulement si $\det(u) \neq 0$ et dans ce cas

$$\det(u^{-1}) = (\det(u))^{-1}.$$

Démonstration. Soit $\mathcal{B} = (e_1, e_2)$ une base de E .

a) On a

$$\det(Id_E) = \det_{\mathcal{B}}(Id_E(e_1), Id_E(e_2)) = \det_{\mathcal{B}}(e_1, e_2) = 1.$$

b) On a de même pour tout scalaire λ ,

$$\det(\lambda u) = \det_B(\lambda u(e_1), \lambda u(e_2)) = \lambda^2 \det_B(u(e_1), u(e_2)) = \lambda^2 \det(u).$$

c) En appliquant plusieurs fois la relation (10.1.4.2), on obtient

$$\begin{aligned} \det(vou) &= \det_B(vou(e_1), vou(e_2)) \\ &= \det_B(v(u(e_1)), v(u(e_2))) \\ &= (\det(v)) \det_B(u(e_1), u(e_2)) = (\det(v)) \cdot (\det(u)). \end{aligned}$$

d) L'endomorphisme u est inversible si et seulement si $(u(e_1), u(e_2))$ est une base de E , c'est-à-dire si et seulement si $\det_B(u(e_1), u(e_2)) \neq 0$ (Théorème 10.1.3.3), donc si et seulement si $\det(u) \neq 0$ d'après (10.1.4.2).

Si u est inversible, de la relation $uou^{-1} = Id_E$ on déduit en appliquant a) et c) :

$$\det(uou^{-1}) = 1 = (\det(u)) (\det(u^{-1})).$$

$$\text{D'où} \quad \det(u^{-1}) = (\det(u))^{-1}.$$

10.1.5. DÉTERMINANT D'UNE MATRICE CARRÉE D'ORDRE 2

Soit A une matrice carrée d'ordre 2 à coefficients dans K . Nous savons que A peut être considérée comme la matrice d'un endomorphisme unique u de K^2 rapporté à sa base canonique. Cela permet de définir le déterminant de A comme étant le déterminant de u .

10.1.5.1. Définition

Soit A une matrice carrée d'ordre 2 à coefficients dans K :

$$A = \begin{pmatrix} a & b \\ c & d \end{pmatrix}.$$

On appelle **déterminant** de A , et on note $\det(A)$, le scalaire défini par

$$\det(A) = \det(u) = ad - bc.$$

On a les propriétés suivantes du déterminant d'une matrice carrée d'ordre 2.

10.1.5.2. Théorème

a) $\det(I_2) = 1$.

b) Si $A \in M_2(K)$, on a $\det(A) = \det({}^t A)$ et $\det(\lambda A) = \lambda^2 \det(A)$ pour tout $\lambda \in K$.

c) Soient A et B deux matrices carrées d'ordre 2. On a

$$\det(AB) = \det(A) \cdot \det(B).$$

d) Une matrice $A \in M_2(K)$ est inversible si et seulement si $\det(A) \neq 0$ et dans ce cas, on a

$$\det(A^{-1}) = (\det(A))^{-1}.$$

e) Si A et B sont deux matrices carrées d'ordre 2 semblables, on a

$$\det(A) = \det(B).$$

Démonstration. Ces propriétés se vérifient aisément par le calcul. Elles découlent également des propriétés analogues du déterminant d'un endomorphisme.

Pour démontrer e), on remarque que si A et B sont semblables, il existe une matrice carrée inversible P d'ordre 2 telle que $B = P^{-1}AP$. Donc, d'après c), on a

$$\begin{aligned} \det(B) &= \det(P^{-1}AP) = \det(P^{-1}) \cdot \det(A) \cdot \det(P) \\ &= \det(P)^{-1} \cdot \det(P) \cdot \det(A) = \det(A). \end{aligned}$$

10.1.5.3. Corollaire

Pour tout endomorphisme u de E , on a

$$\det({}^t u) = \det(u).$$

10.2. Applications et formes multilinéaires

Dans ce paragraphe, nous nous proposons d'étendre la notion d'application bilinéaire définie au paragraphe précédent.

Soit p un nombre entier ≥ 1 et soient $E_1, \dots, E_p$ et F des K -espaces vectoriels. L'ensemble $E_1 \times E_2 \times \dots \times E_p$ sera muni de sa structure d'espace vectoriel produit.

10.2.1. APPLICATIONS ET FORMES MULTILINÉAIRES ALTERNÉES

10.2.1.1. Définition

On dit qu'une application $f : E_1 \times E_2 \times \dots \times E_p \longrightarrow F$ est p -linéaire si pour tout indice j tel que $1 \leq j \leq p$ et pour tout système de vecteurs $x_i \in E_i$ ($i \neq j$), l'application de E_j dans F définie par :

$$x_j \longmapsto f(x_1, \dots, x_{j-1}, x_j, x_{j+1}, \dots, x_p)$$

est linéaire.

Une application p -linéaire de $E_1 \times E_2 \times \dots \times E_p$ dans K s'appelle une forme p -linéaire.

Si $E_1 = \dots = E_p = E$ une application p -linéaire sur E^p à valeurs dans F s'appelle une application p -linéaire sur E .

Une forme p -linéaire sur E^p s'appelle une forme p -linéaire sur E .

Si $p = 1$, on retrouve les notions d'application linéaire et de forme linéaire. Si $p = 2$, on a les notions d'application bilinéaire et de forme bilinéaire.

Les propriétés suivantes résultent immédiatement des définitions :

a) Si f est une application p -linéaire de $E_1 \times \dots \times E_p$ dans F , alors $f(x_1, \dots, x_p) = 0$ dès que l'un des x_i est le vecteur nul.

b) Quels que soient les scalaires $\lambda_1, \dots, \lambda_p$, on a

$$f(\lambda_1 x_1, \dots, \lambda_p x_p) = \lambda_1 \dots \lambda_p f(x_1, \dots, x_p).$$

c) L'ensemble des applications p -linéaires de $E_1 \times \dots \times E_p$ dans F est un sous-espace vectoriel de l'espace vectoriel des applications de $E_1 \times \dots \times E_p$ dans F . Ce sous-espace est noté $\mathcal{L}_p(E_1, \dots, E_p; F)$ et $\mathcal{L}_p(E; F)$ si $E_1 = \dots = E_p = E$. $\mathcal{L}_p(E)$ désignera l'espace des formes p -linéaires sur E .

10.2.1.2. Définition

Soient E et F des espaces vectoriels sur K , et $p \geq 1$ un nombre entier. On dit qu'une application p -linéaire $f : E^p \longrightarrow F$ est :

– symétrique si

$$f(x_{\sigma(1)}, \dots, x_{\sigma(p)}) = f(x_1, \dots, x_p)$$

quels que soient les $x_i \in E$ et la permutation $\sigma \in s_p$,

– antisymétrique si

$$f(x_{\sigma(1)}, \dots, x_{\sigma(p)}) = \varepsilon(\sigma) f(x_1, \dots, x_p)$$

quels que soient les $x_i \in E$ et la permutation $\sigma \in S_p$ ($\varepsilon(\sigma)$ désigne la signature de σ),

– alternée si

$$f(x_1, \dots, x_p) = 0$$

chaque fois que deux des vecteurs x_i sont égaux.

On vérifie facilement que l'ensemble des applications p -linéaire symétriques (resp. alternées) $f : E^p \longrightarrow F$ est un sous-espace vectoriel de $\mathcal{L}_p(E; F)$. On le note $\mathcal{S}_p(E; F)$ (resp. $\mathcal{A}_p(E; F)$).

Lorsqu'il s'agit de formes p -linéaires, ces espaces seront notés respectivement $\mathcal{S}_p(E)$ et $\mathcal{A}_p(E)$.

10.2.2. PROPRIÉTÉS DES APPLICATIONS ET DES FORMES MULTILINÉAIRES ALTERNÉES

10.2.2.1. Théorème

Soient E et F des K -espaces vectoriels. Une application p -linéaire f de E^p dans F est alternée si et seulement si elle est antisymétrique.

Démonstration. Supposons f alternée. Soit σ une permutation de $\{1, 2, \dots, p\}$. Cette permutation est un produit de s transpositions et sa signature est $\varepsilon(\sigma) = (-1)^s$. Il suffit donc de montrer que pour la transposition qui échange les entiers i et j , on a

$$f(x_1, \dots, x_j, \dots, x_i, \dots, x_p) = -f(x_1, \dots, x_i, \dots, x_j, \dots, x_p).$$

En utilisant le fait que f est alternée et multilinéaire, nous obtenons

$$\begin{aligned} 0 &= f(x_1, \dots, x_i + x_j, \dots, x_i + x_j, \dots, x_p) \\ &= f(x_1, \dots, x_i, \dots, x_i, \dots, x_p) + \\ &\quad f(x_1, \dots, x_i, \dots, x_j, \dots, x_p) + \\ &\quad f(x_1, \dots, x_j, \dots, x_i, \dots, x_p) + \\ &\quad f(x_1, \dots, x_j, \dots, x_j, \dots, x_p). \end{aligned}$$

Puisque f est alternée, il reste

$$f(x_1, \dots, x_i, \dots, x_j, \dots, x_p) + f(x_1, \dots, x_j, \dots, x_i, \dots, x_p) = 0,$$

d'où l'égalité à démontrer.

Supposons f antisymétrique. Soit $(x_1, \dots, x_p)$ un élément de E^p tel que $x_i = x_j$ avec $i \neq j$. La signature de la transposition τ qui échange i et j étant égale à -1 , nous avons :

$$f(x_1, \dots, x_i, \dots, x_i, \dots, x_p) = -f(x_1, \dots, x_i, \dots, x_i, \dots, x_p).$$

On en déduit $2f(x_1, \dots, x_p) = 0$ et puisque K n'est pas de caractéristique 2, il vient $f(x_1, \dots, x_p) = 0$, ce qui achève la démonstration du théorème.

Le théorème suivant et ses corollaires joueront un rôle important dans la suite de ce chapitre.

10.2.2.2. Théorème

Soient E et F des K -espaces vectoriels et f une application p -linéaire alternée de E^p dans F . Si les vecteurs $x_1, \dots, x_p$ de E sont linéairement dépendants, on a $f(x_1, \dots, x_p) = 0$.

Démonstration. Si les vecteurs $x_1, \dots, x_p$ sont linéairement dépendants, l'un d'eux, par exemple x_1 , est combinaison linéaire des autres :

$$x_1 = \lambda_2 x_2 + \dots + \lambda_p x_p,$$

où les λ_i sont des scalaires. Donc

$$f(x_1, \dots, x_p) = \sum_{j=2}^p \lambda_j f(x_j, x_2, \dots, x_p).$$

Puisque f est alternée, on a pour tout $j \in \{2, \dots, p\}$, $f(x_j, x_2, \dots, x_p) = 0$ car x_j est égal à l'un des vecteurs $x_2, \dots, x_p$.

10.2.2.3. Corollaire

Les notations étant celles du Théorème 10.2.2.2, $f(x_1, \dots, x_p)$ ne change pas si on ajoute à l'un des vecteurs x_i une combinaison linéaire des autres vecteurs.

Démonstration. Supposons que le vecteur y soit une combinaison linéaire des vecteurs x_j pour $j \neq i$:

$$y = \lambda_1 x_1 + \dots + \lambda_{i-1} x_{i-1} + \lambda_{i+1} x_{i+1} + \dots + \lambda_p x_p,$$

où les λ_j sont des scalaires.

D'après le Théorème 10.2.2.2, on a

$$f(x_1, \dots, x_{i-1}, y, x_{i+1}, \dots, x_p) = 0.$$

Donc

$$\begin{aligned} f(x_1, \dots, x_{i-1}, x_i + y, x_{i+1}, \dots, x_p) &= \\ f(x_1, \dots, x_{i-1}, x_i, x_{i+1}, \dots, x_p) + f(x_1, \dots, x_{i-1}, y, x_{i+1}, \dots, x_p) &= \\ f(x_1, \dots, x_{i-1}, x_i, x_{i+1}, \dots, x_p). \end{aligned}$$

10.2.2.4. Corollaire

Soit E un K -espace vectoriel de dimensions finie r . Alors quel que soit le K -espace vectoriel F , toute application p -linéaire alternée de E^p dans F est nulle pour $r < p$.

Démonstration. Soit $f : E^p \longrightarrow F$ une application p -linéaire alternée. Si $\dim(E) = r < p$, toute famille $\{x_1, \dots, x_p\}$ de p vecteurs de E est liée. Donc quels que soient les vecteurs $x_1, \dots, x_p$ de E , on a

$$f(x_1, \dots, x_p) = 0,$$

et par suite f est l'application nulle de E^p dans F .

Lorsque E est un K -espace vectoriel de dimension finie p , il existe effectivement des applications p -linéaires alternées non identiquement nulles sur E à valeurs dans un K -espace vectoriel F quelconque. Nous allons démontrer ce résultat lorsque $F = K$.

10.2.2.5. Théorème

Soient E un K -espace vectoriel de dimension finie p et $B = (e_1, \dots, e_p)$ une base de E . Alors :

a) L'espace vectoriel $A_p(E)$ des formes p -linéaires alternées sur E est de dimension 1.

b) Il existe une unique forme p -linéaire alternée D sur E telle que

$$D(e_1, \dots, e_p) = 1.$$

Démonstration. a) Soit f une forme p -linéaire alternée sur E . Pour tout système de p vecteurs $x_1, \dots, x_p$ de E , posons

$$x_i = \sum_{j=1}^p \lambda_{ji} e_j \quad (1 \leq i \leq p).$$

f étant p -linéaire, nous avons

$$(10.2.2.1) \quad f \left(\sum_{j_1=1}^p \lambda_{j_1,1} e_{j_1} \right), \dots, \sum_{j_p=1}^p (\lambda_{j_p,p} e_{j_p}) = \sum \lambda_{j_1,1} \dots \lambda_{j_p,p} f(e_{j_1}, \dots, e_{j_p})$$

la sommation étant étendue aux p^p systèmes $\{j_1, \dots, j_p\}$ d'entiers pris dans $\{1, 2, \dots, p\}$ indépendamment les uns des autres.

Puisque f est alternée, on a $f(e_{j_1}, \dots, e_{j_p}) = 0$ si deux des entiers $j_1, \dots, j_p$ sont égaux. Dans la somme (10.2.2.1), nous pouvons donc retenir les seuls termes pour lesquels les entiers $j_1, \dots, j_p$ sont tous différents, c'est-à-dire constituent une permutation des p premiers entiers naturels ; il existe alors une permutation $\sigma \in S_p$ telle que

$$j_1 = \sigma(1), \dots, j_p = \sigma(p);$$

et puisque f est antisymétrique, nous avons

$$f(e_{j_1}, \dots, e_{j_p}) = f(e_{\sigma(1)}, \dots, e_{\sigma(p)}) = \varepsilon(\sigma) f(e_1, \dots, e_p).$$

Tous les systèmes d'indices distincts sont obtenus lorsque σ parcourt le groupe symétrique s_p . Finalement, la somme (10.2.2.1) s'écrit

$$(10.2.2.2) \quad f(x_1, \dots, x_p) = f(e_1, \dots, e_p) \left(\sum_{\sigma \in S_p} \varepsilon(\sigma) \lambda_{\sigma(1),1} \dots \lambda_{\sigma(p),p} \right).$$

Soit D l'application de E^p dans K définie par

$$(10.2.2.3) \quad D(x_1, \dots, x_p) = \sum_{\sigma \in S_p} \varepsilon(\sigma) \lambda_{\sigma(1),1} \dots \lambda_{\sigma(p),p}.$$

Nous avons, quels que soient les vecteurs $x_1, \dots, x_p$ de E

$$f(x_1, \dots, x_p) = f(e_1, \dots, e_p) D(x_1, \dots, x_p),$$

c'est-à-dire $f = f(e_1, \dots, e_p)D$.

Pour achever la démonstration du a), il suffit donc de montrer que D est une forme p -linéaire alternée non nulle sur E . Chaque terme de la somme (10.2.2.3) dépend linéairement de chacun des vecteurs $x_1, \dots, x_p$, donc D est une forme p -linéaire sur E . Montrons qu'elle est alternée. Soient i et j deux entiers quelconques mais distincts appartenant à $\{1, \dots, p\}$ et supposons que $x_i = x_j$. Nous devons montrer que $D(x_1, \dots, x_p) = 0$.

Soit τ la transposition qui échange i et j . On a pour toute permutation $\sigma \in S_p$, $\varepsilon(\sigma\tau) = \varepsilon(\sigma) \varepsilon(\tau) = -\varepsilon(\sigma)$. D'autre part, A_p désignant l'ensemble des permutations paires, on a $S_p = A_p \cup A_p\tau$, d'où

$$\begin{aligned} D(x_1, \dots, x_p) &= \sum_{\sigma \in A_p} \varepsilon(\sigma) \lambda_{\sigma(1),1} \dots \lambda_{\sigma(p),p} + \sum_{\sigma \in A_p} \varepsilon(\sigma\tau) \lambda_{\sigma\tau(1),1} \dots \lambda_{\sigma\tau(p),p} \\ &= \sum_{\sigma \in A_p} \varepsilon(\sigma) \lambda_{\sigma(1),1} \dots \lambda_{\sigma(p),p} - \sum_{\sigma \in A_p} \varepsilon(\sigma) \lambda_{\sigma(1),1} \dots \lambda_{\sigma(p),p} \end{aligned}$$

car

$$\sigma(\tau(1)) = \sigma(1), \dots, \sigma(\tau(i)) = \sigma(j), \dots, \sigma(\tau(j)) = \sigma(i), \dots, \sigma(\tau(p)) = \sigma(p).$$

Comme $x_i = x_j$ par hypothèse, on a

$$\lambda_{\sigma(i),i} = \lambda_{\sigma(i),j} \text{ et } \lambda_{\sigma(j),j} = \lambda_{\sigma(j),i}.$$

Donc $D(x_1, \dots, x_p) = 0$ et D est bien alternée.

Montrons que D est non nulle. Pour cela prenons $x_1 = e_1, \dots, x_p = e_p$. On a alors $\lambda_{ij} = \delta_{ij}$ (symbole de Kronecker). D'où d'après la formule (10.2.2.3),

$$D(e_1, \dots, e_p) = \sum_{\sigma \in S_p} \varepsilon(\sigma) \delta_{\sigma(1),1} \dots \delta_{\sigma(p),p}.$$

Tous les termes de cette dernière somme sont nuls sauf lorsque $\sigma(1) = 1, \dots, \sigma(p) = p$. Dans ce cas σ est la permutation identique dont la signature est $+1$. On en déduit que

$$(10.2.2.4) \quad D(e_1, \dots, e_p) = 1.$$

b) Nous venons de voir qu'il existe une forme p -linéaire alternée D sur E telle que $D(e_1, \dots, e_p) = 1$ et telle que

$$f = f(e_1, \dots, e_p)D$$

pour toute forme p -linéaire alternée f sur E . Donc si $f(e_1, \dots, e_p) = 1$, alors $f = D$, d'où l'unicité de D .

10.2.2.6. Corollaire

Soient E un K -espace vectoriel de dimension finie p , $\mathcal{B} = (e_1, \dots, e_p)$ une base de E et f une forme p -linéaire alternée sur E . Pour que f soit nulle il faut et il suffit que $f(e_1, \dots, e_p) = 0$.

Démonstration. Si f est nulle, il est clair que $f(e_1, \dots, e_p) = 0$.

Réciproquement, si $f(e_1, \dots, e_p) = 0$, on a d'après la formule (10.2.2.2), $f(x_1, \dots, x_p) = 0$ pour toute suite $(x_1, \dots, x_p)$ d'éléments de E , donc $f = 0$.

10.3. Déterminants

Les résultats du paragraphe précédent permettent de définir le déterminant d'un système de vecteurs, d'un endomorphisme et d'une matrice carrée.

10.3.1. DÉTERMINANT D'UN SYSTÈME DE VECTEURS

10.3.1.1. Définition

*Soient E un espace vectoriel de dimension p sur K , $\mathcal{B} = (e_1, \dots, e_p)$ une base de E et p vecteurs $x_1, \dots, x_p$ de E . on appelle **déterminant de $x_1, \dots, x_p$ par rapport à la base $\mathcal{B}$** , le scalaire $D(x_1, \dots, x_p)$ défini par la relation (10.2.2.3).*

On note $\det_{\mathcal{B}}(x_1, \dots, x_p)$ le déterminant de $x_1, \dots, x_p$ par rapport à la base $\mathcal{B}$.

L'application $\det_{\mathcal{B}}$ possède les propriétés suivantes :

10.3.1.2. Théorème

Soient E un K -espace vectoriel de dimension finie p et $\mathcal{B} = (e_1, \dots, e_p)$ une base de E .

a) *L'application déterminant est une forme p -linéaire alternée et on a*

$$\det_{\mathcal{B}}(x_{\sigma(1)}, \dots, x_{\sigma(p)}) = \varepsilon(\sigma) \det_{\mathcal{B}}(x_1, \dots, x_p)$$

pour toute suite $(x_1, \dots, x_p)$ de p vecteurs de E et pour toute permutation $\sigma \in S_p$.

b) Le déterminant de p vecteurs ne change pas si l'on ajoute à l'un d'entre eux une combinaison linéaire des autres vecteurs.

c) Si $B' = (e'_1, \dots, e'_p)$ est une autre base de E , on a pour toute suite $(x_1, \dots, x_p)$ de p vecteurs de E ,

$$(10.3.1.1) \quad \det_{B'}(x_1, \dots, x_p) = \det_B(x_1, \dots, x_p) \det_{B'}(e_1, \dots, e_p).$$

d) p vecteurs $x_1, \dots, x_p$ de E sont linéairement indépendants si et seulement si $\det_B(x_1, \dots, x_p) \neq 0$.

Démonstration. La propriété a) résulte du Théorème 10.2.2.5. et la propriété b) est une propriété générale des formes multilinéaires alternées (voir le Corollaire 10.2.2.3).

c) D'après le Théorème 10.2.2.5, il existe un scalaire λ tel que

$$\det_{B'}(x_1, \dots, x_p) = \lambda \det_B(x_1, \dots, x_p).$$

En prenant $x_1 = e_1, \dots, x_p = e_p$, il vient d'après la formule (10.2.2.4) :

$$\det_{B'}(x_1, \dots, x_p) = \lambda \det_B(e_1, \dots, e_p) = \lambda,$$

d'où l'égalité à démontrer.

d) Nous savons déjà que si les vecteurs $x_1, \dots, x_p$ sont linéairement dépendants, alors $\det_B(x_1, \dots, x_p) = 0$ (Théorème 10.2.2.2).

Réciproquement, si la famille $(x_1, \dots, x_p)$ est libre, c'est une base B' de E . Puisque $\det_{B'}$ est une forme p -linéaire alternée, il existe d'après le Théorème 10.2.2.5 un scalaire λ tel que $\det_{B'}(x_1, \dots, x_p) = \lambda \det_B(x_1, \dots, x_p)$ et puisque $\det_{B'}(x_1, \dots, x_p) = 1$, on a bien $\det_B(x_1, \dots, x_p) \neq 0$.

10.3.2. DÉTERMINANT D'UN ENDORMORPHISME

10.3.2.1. Théorème

Soient E un K -espace vectoriel de dimension finie p et u un endomorphisme de E . Il existe un scalaire et un seul, appelé **déterminant de u** , et noté $\det(u)$, tel que pour toute forme p -linéaire alternée f sur E et pour toute suite $(x_1, \dots, x_p)$ de vecteurs de E , on ait :

$$(10.3.2.1) \quad f(u(x_1), \dots, u(x_p)) = \det(u) f(x_1, \dots, x_p).$$

Démonstration. Pour tout $f \in A_p(E)$, considérons l'application $\varphi_u(f)$ de E^p dans K donnée par

$$\varphi_u(f)(x_1, \dots, x_p) = f(u(x_1), \dots, u(x_p)).$$

On vérifie facilement que $\varphi_u(f)$ est une forme p -linéaire alternée sur E et que l'application $f \mapsto \varphi_u(f)$ est un endomorphisme de $A_p(E)$. Comme d'après

le Théorème 10.2.2.5, $\dim(A_p(E)) = 1$, tout endomorphisme de $A_p(E)$ est une homothétie. Le rapport d'une homothétie dans un espace vectoriel non nul étant déterminé de façon unique, on voit qu'il existe bien un scalaire unique, noté $\det(u)$, indépendant de f et tel que l'on ait la formule (10.3.2.1).

10.3.2.2. Remarque

Soit $\mathcal{B} = (e_1, \dots, e_p)$ une base quelconque de E . Si on applique la relation (10.3.2.1), en prenant $f = \det_{\mathcal{B}}$, alors pour toute suite $(x_1, \dots, x_p)$ de E et pour tout endomorphisme u de E , on a

$$(10.3.2.2) \quad \det_{\mathcal{B}}(u(x_1), \dots, u(x_p)) = \det(u) \det_{\mathcal{B}}(x_1, \dots, x_p).$$

Si en particulier on prend $(x_1, \dots, x_p) = (e_1, \dots, e_p)$, on obtient une expression du déterminant de u :

$$(10.3.2.3) \quad \det(u) = \det_{\mathcal{B}}(u(e_1), \dots, u(e_p)).$$

Rassemblons quelques propriétés du déterminant d'un endomorphisme dans un théorème.

10.3.2.3. Théorème

Soit E un K -espace vectoriel de dimension p . Alors :

a) $\det(\text{Id}_E) = 1$.

b) Pour tout endomorphisme u de E et pour tout $\lambda \in K$, on a

$$\det(\lambda u) = \lambda^p \det(u).$$

c) Quels que soient les endomorphismes u et v de E , on a

$$\det(vu) = (\det(v)) (\det(u)).$$

d) Un endomorphisme de E est inversible si et seulement si $\det(u) \neq 0$ et dans ce cas on a :

$$\det(u^{-1}) = (\det(u))^{-1}.$$

La démonstration est la même que celle du Théorème 10.1.4.3.

10.3.2.4. Corollaire

Soit E un K -espace vectoriel de dimension finie. L'application $u \mapsto \det(u)$ est un homomorphisme du groupe linéaire $GL(E)$ sur le groupe multiplicatif K^ des éléments non nuls de K . Le noyau de cet homomorphisme, constitué des endomorphismes de E dont le déterminant est égal à 1 est un sous-groupe*

distingué de $GL(E)$, appelé **groupe spécial linéaire** ou **groupe unimodulaire** de E , et noté $SL(E)$.

C'est immédiat.

10.3.3. DÉTERMINANT D'UNE MATRICE CARRÉE

Soit $A = (a_{ij})_{1 \leq i, j \leq p}$ une matrice carrée d'ordre p à coefficients dans K . Les colonnes de A représentent des vecteurs $x_1, \dots, x_p$ de K^p rapporté à sa base canonique. Nous savons aussi que A peut être considérée comme la matrice d'un endomorphisme unique u de K^p rapporté à sa base canonique. Ceci nous amène à poser la définition suivante :

10.3.3.1. Définition

Soit $A = (a_{ij})$ une matrice carrée d'ordre p à éléments dans K . On appelle **déterminant** de A , et on note $\det(A)$, le déterminant de ses vecteurs colonnes dans la base canonique de K^p .

Soit $\mathcal{B}$ la base canonique de K^p . Avec les notations précédentes, on a donc

$$(10.3.3.1) \quad \det(A) = \det_{\mathcal{B}}(x_1, \dots, x_p) = \sum_{\sigma \in S_p} \varepsilon(\sigma) a_{\sigma(1),1} \dots a_{\sigma(p),p} = \det(u).$$

Le déterminant de la matrice $A = (a_{ij})$ se note aussi

$$\det(A) = \begin{vmatrix} a_{11} & \dots & a_{1p} \\ \dots & \dots & \dots \\ a_{p1} & \dots & a_{pp} \end{vmatrix}$$

Les propriétés des déterminants des matrices carrées se déduisent des propriétés des déterminants des endomorphismes en utilisant l'isomorphisme canonique de l'algèbre $\mathcal{L}(K^p)$ des endomorphismes de l'espace vectoriel K^p sur l'algèbre $M_p(K)$ des matrices carrées d'ordre p à coefficients dans K . Nous obtenons ainsi le théorème suivant :

10.3.3.2. Théorème

- a) Le déterminant de la matrice unité I_p est égal à 1.
- b) Quelle que soit la matrice $A \in M_p(K)$ et quel que soit $\lambda \in K$, on a

$$\det(\lambda A) = \lambda^p \det(A).$$
- c) Pour tout couple (A, B) d'éléments de $M_p(K)$, on a

$$\det(AB) = \det(A) \det(B).$$

d) Pour qu'une matrice carrée A soit inversible, il faut et il suffit que son déterminant soit non nul ; on a alors :

$$\det(A^{-1}) = (\det(A))^{-1}.$$

e) Si A et B sont deux matrices carrées d'ordre p semblables, on a

$$\det(A) = \det(B).$$

10.3.3.3. Corollaire

L'application $A \mapsto \det(A)$ est un homomorphisme du groupe linéaire $GL(p, K)$ sur le groupe multiplicatif K^* . Le noyau de cet homomorphisme, qui est l'ensemble des matrices carrées d'ordre p dont le déterminant est égal à 1, est un sous-groupe distingué de $GL(p, K)$, appelé groupe spécial linéaire ou groupe unimodulaire ; on le note $SL(p, K)$.

La démonstration est immédiate.

Voici encore un résultat important pour les applications.

10.3.3.4. Théorème

Le déterminant de la transposée d'une matrice est égal au déterminant de cette matrice.

Démonstration. Soit $A = (a_{ij})$ une matrice carrée d'ordre p . Posons ${}^tA = (b_{ij})$; on a $b_{ij} = a_{ji}$.

Nous avons par définition :

$$\det({}^tA) = \sum_{\sigma \in S_p} \varepsilon(\sigma) b_{\sigma(1),1} \dots b_{\sigma(p),p} = \sum_{\sigma \in S_p} \varepsilon(\sigma) a_{1,\sigma(1)} \dots a_{p,\sigma(p)}.$$

Si on pose, pour tout $j \in \{1, \dots, p\}$, $i = \sigma(j)$, alors $a_{j,\sigma(j)} = a_{\sigma^{-1}(i),i}$ et $\varepsilon(\sigma^{-1}) = \varepsilon(\sigma)$.

Mais puisque l'application $\sigma \mapsto \sigma^{-1}$ est une bijection de S_p sur lui-même, on peut écrire, en posant $\tau = \sigma^{-1}$,

$$\begin{aligned} \det({}^tA) &= \sum_{\sigma^{-1} \in S_p} \varepsilon(\sigma^{-1}) a_{\sigma^{-1}(1),1} \dots a_{\sigma^{-1}(p),p} \\ &= \sum_{\tau \in S_p} \varepsilon(\tau) a_{\tau(1),1} \dots a_{\tau(p),p} = \det(A). \end{aligned}$$

10.3.3.5. Corollaire

Soit E un K -espace vectoriel de dimension finie et u un endomorphisme de E . On a $\det({}^tu) = \det(u)$.

Ceci résulte du Théorème 10.3.3.4 en considérant les matrices de u et de ${}^t u$.

10.3.3.6. Remarque

Le théorème 10.3.3.4 montre que le déterminant des vecteurs colonnes d'une matrice A est égal au déterminant des vecteurs lignes de A . Il en résulte que toute propriété démontrée pour les colonnes d'un déterminant est également valable pour les lignes de ce déterminant. Nous parlerons de **rangées** pour désigner les lignes ou les colonnes d'un déterminant.

Voici maintenant quelques propriétés qui découlent du fait que le déterminant est une forme p -linéaire alternée.

10.3.3.7. Théorème

Soit A une matrice carrée d'ordre p .

- a) *Si deux rangées parallèles de A sont identiques ou proportionnelles, on a $\det(A) = 0$.*
- b) *Si l'on échange deux rangées parallèles de A , le déterminant de A change de signe ; plus généralement, si l'on fait subir aux rangées de A une permutation σ , le déterminant de A est multiplié par $\varepsilon(\sigma)$.*
- c) *Si l'on multiplie par le même scalaire λ tous les éléments d'une rangée de A , le déterminant de A est multiplié par λ .*
- d) *$\det(A)$ ne change pas si l'on ajoute à l'une de ses rangées une combinaison linéaire des autres rangées parallèles.*
- e) *$\det(A) = 0$ si et seulement si les vecteurs colonnes (resp. les vecteurs lignes) de A sont linéairement dépendants.*

10.3.3.8. Théorème

Soient E un K -espace vectoriel de dimension p et $B = (e_1, \dots, e_p)$ une base de E .

- a) *Soient $x_1, \dots, x_p$ des vecteurs de E . Notons $a_{1j}, a_{2j}, \dots, a_{pj}$ les coordonnées du vecteur x_j par rapport à la base B et A la matrice (a_{ij}) de type (p, p) . On a $\det_B(x_1, \dots, x_p) = \det(A)$.*
- b) *Soient u endomorphisme de E et $M(u)$ sa matrice par rapport à la base B . On a $\det(u) = \det(M(u))$.*

Démonstration. L'assertion a) résulte immédiatement des formules (10.2.2.3) et (10.3.3.1).

b) On a d'après (10.3.2.3) $\det(u) = \det_B(u(e_1), \dots, u(e_p))$ et d'après a) ce déterminant n'est autre que le déterminant de la matrice associée à u relativement à la base B .

Nous disposons, grâce à ce théorème, d'un moyen pour calculer effectivement le déterminant d'un endomorphisme.

10.3.4. CALCULS DES DÉTERMINANTS

Nous allons, dans ce numéro, apprendre à développer un déterminant suivant une colonne ou une ligne. Mais auparavant, nous allons établir les formules qui donnent le développement des déterminants de certaines matrices remarquables.

10.3.4.1. Théorème

Soit M une matrice carrée d'ordre p de la forme :

$$M = \begin{pmatrix} A & : & B \\ & : & \\ \dots & \dots & \dots \\ 0 & : & C \\ & : & \end{pmatrix} \begin{matrix} \uparrow r \\ \\ \\ \downarrow p-r \end{matrix}$$

$$\begin{matrix} \leftarrow r & \quad p-r \rightarrow \end{matrix}$$

où $A \in M_r(K)$, $B \in M_{r,p-r}(K)$, $C \in M_{p-r}(K)$ et où 0 est la matrice nulle de type $(p-r, r)$. Alors

$$(10.3.4.1) \quad \det(M) = \det(A) \det(C).$$

Démonstrations. Soient $B = (e_1, \dots, e_p)$ la base canonique de K^p , E le sous-espace vectoriel, de K^p engendré par $B' = (e_1, \dots, e_r)$ et F le sous-espace vectoriel de K^p engendré par $B'' = (e_{r+1}, \dots, e_p)$. Notons u l'endomorphisme de K^p admettant M pour matrice par rapport à la base B , v l'endomorphisme de E dont la matrice par rapport à la base B' est A et enfin w l'endomorphisme de F canoniquement associé à C .

Par définition du déterminant d'une matrice, on a

$$\det(M) = \det_B(u(e_1), \dots, u(e_p)).$$

On a évidemment

$$u(e_1) = v(e_1), \dots, u(e_r) = v(e_r),$$

d'où

$$\det(M) = \det_B(v(e_1), \dots, v(e_r), u(e_{r+1}), \dots, u(e_p)).$$

Il est clair que l'application $f : E^r \longrightarrow K$ définie par

$$f(x_1, \dots, x_r) = \det_B(x_1, \dots, x_r, u(e_{r+1}), \dots, u(e_p)),$$

est une forme r -linéaire alternée sur E . Par définition même du déterminant de l'endomorphisme v , nous avons :

$$f(v(e_1), \dots, v(e_r)) = \det(v) f(e_1, \dots, e_r),$$

d'où puisque $\det(v) = \det(A)$,

$$\det(M) = \det(A) \det_B(e_1, \dots, e_r, u(e_{r+1}), \dots, u(e_p)).$$

Or

$$u(e_{r+1}) = z_{r+1} + w(e_{r+1}), \dots, u(e_p) = z_p + w(e_p)$$

où les vecteurs $z_{r+1}, \dots, z_p$ sont dans E ; donc

$$\begin{aligned} \det_B(e_1, \dots, e_r, u(e_{r+1}), \dots, u(e_p)) &= \\ \det_B(e_1, \dots, e_r, z_{r+1} + w(e_{r+1}), \dots, z_p + w(e_p)). \end{aligned}$$

Ce dernier déterminant ne change si l'on retranche du vecteur colonne d'indice k ($r+1 \leq k \leq p$), le vecteur z_k qui est une combinaison linéaire de $e_1, e_2, \dots, e_r$. D'où

$$\begin{aligned} \det_B(e_1, \dots, e_r, u(e_{r+1}), \dots, u(e_p)) &= \\ \det_B(e_1, \dots, e_r, w(e_{r+1}), \dots, w(e_p)). \end{aligned}$$

En introduisant l'application $h : F^{p-r} \longrightarrow K$ définie par :

$$h(x_{r+1}, \dots, x_p) = \det_B(e_1, \dots, e_r, x_{r+1}, \dots, x_p),$$

et en utilisant un raisonnement analogue au précédent, on voit que

$$\begin{aligned} \det_B(e_1, \dots, e_r, w(e_{r+1}), \dots, w(e_p)) &= \\ \det(w) \det_B(e_1, \dots, e_r, e_{r+1}, \dots, e_p) &= \det(w) \end{aligned}$$

puisque $\det_B(e_1, \dots, e_p) = 1$. Comme $\det(w) = \det(C)$, on obtient finalement

$$\det(M) = \det(A) \det(C).$$

10.3.4.2. Corollaire

Soit A une matrice carrée de la forme

$$A = \begin{pmatrix} A_{11} & A_{12} & \dots & A_{1p} \\ 0 & A_{22} & \dots & A_{2p} \\ \dots & \dots & \dots & \dots \\ 0 & 0 & \dots & A_{pp} \end{pmatrix}$$

où les matrices A_{ii} sont des matrices carrées et où les matrices A_{ij} telles que $i > j$ sont nulles. Alors

$$(10.3.4.2) \quad \det(A) = \det(A_{11}) \det(A_{22}) \dots \det(A_{pp}).$$

En particulier, si

$$A = \begin{pmatrix} A_1 & & 0 \\ & A_2 & \\ 0 & & A_p \end{pmatrix}$$

est une matrice diagonale de matrices carrées, on a

$$(10.3.4.3) \quad \det(A) = \det(A_1) \det(A_2) \dots \det(A_p).$$

La formule (10.3.4.2) découle du Théorème 10.3.4.1 en raisonnant par récurrence sur le nombre p de blocs diagonaux. La formule (10.3.4.3) s'en déduit aisément.

10.3.4.3. Corollaire

Le déterminant d'une matrice triangulaire est égal au produit des éléments diagonaux. En particulier, le déterminant d'une matrice diagonale est égal au produit des éléments diagonaux.

Il s'agit là de deux cas particulier de ceux traités dans le Corollaire 10.3.4.2.

10.3.4.4. Définition

Soit $A = (A_{ij})$ une matrice carrée d'ordre p . On appelle **mineur relatif à l'élément a_{ij}** , le déterminant de la matrice carrée A_{ij} , d'ordre $p - 1$, déduite de A en supprimant la $i^{\text{ème}}$ ligne et la $j^{\text{ème}}$ colonne de A .

Le scalaire $\Delta_{ij} = (-1)^{i+j} \det(A_{ij})$ s'appelle le **cofacteur de a_{ij}** .

10.3.4.5. Théorème

Soient $A = (a_{ij})$ une matrice carrée d'ordre $p > 1$ et Δ_{ij} le cofacteur de a_{ij} . Alors quels que soient les entiers i et j dans $\{1, \dots, p\}$, on a

$$(10.3.4.4) \quad \det(A) = \sum_{i=1}^p a_{ij} \Delta_{ij}$$

et

$$(10.3.4.5) \quad \det(A) = \sum_{j=1}^p a_{ij} \Delta_{ij}.$$

La formule (10.3.4.4) s'appelle le **développement de $\det(A)$ suivant la $i^{\text{ème}}$ colonne de A** ; la formule (10.3.4.5) s'appelle le **développement de $\det(A)$ suivant la $i^{\text{ème}}$ ligne de A** .

Démonstrations. Nous allons démontrer la formule (10.3.4.4) Comme $\det({}^t A) = \det(A)$, on en déduit la formule (10.3.4.5).

Soit $\mathcal{B} = (e_1, \dots, e_p)$ la base canonique de K^p ; désignons par $x_1, \dots, x_p$ les vecteurs colonnes de A . On a :

$$x_j = \sum_{i=1}^p a_{ij} e_i, \quad 1 \leq j \leq p,$$

et par définition

$$\det(A) = \det_{\mathcal{B}}(x_1, \dots, x_p).$$

Comme $\det_{\mathcal{B}}$ est une forme p -linéaire alternée, on a

$$\begin{aligned} \det(A) &= (-1)^{j-1} \det_{\mathcal{B}}(x_j, x_1, \dots, x_{j-1}, x_{j+1}, \dots, x_p) \\ &= (-1)^{j-1} \sum_{i=1}^p a_{ij} \det_{\mathcal{B}}(e_i, x_1, \dots, x_{j-1}, x_{j+1}, \dots, x_p) \\ &= (-1)^{j-1} \sum_{i=1}^p a_{ij} \det_{\mathcal{B}}(A'_{ij}), \end{aligned}$$

où A'_{ij} est la matrice ayant pour vecteurs colonnes

$$e_i, x_1, \dots, x_{j-1}, x_{j+1}, \dots, x_p.$$

Dans la matrice A'_{ij} , amenons la ligne de rang i à la première place en l'échangeant successivement avec les $i-1$ premières lignes ; il y a $i-1$ transpositions, d'où

$$\det(A'_{ij}) = (-1)^{i-1} \det(B_{ij})$$

où la matrice B_{ij} est de la forme

$$B_{ij} = \begin{pmatrix} 1 & * \\ 0 & A_{ij} \end{pmatrix},$$

A_{ij} étant la matrice déduite de A en supprimant la $i^{\text{ème}}$ ligne et la $j^{\text{ème}}$ colonne de A .

D'après le Théorème 10.3.4.1 et le fait que $\det(1) = 1$, on a

$$\det(B_{ij}) = \det(A_{ij}).$$

D'où

$$\begin{aligned} \det(A) &= (-1)^{j-1} \sum_{i=1}^p a_{ij} \det(A'_{ij}) = (-1)^{j-1} \sum_{i=1}^p a_{ij} (-1)^{i-1} \det(B_{ij}) \\ &= \sum_{i=1}^p (-1)^{i+j} a_{ij} \det(A_{ij}) = \sum_{i=1}^p a_{ij} \Delta_{ij}. \end{aligned}$$

10.3.4.6. Remarque

Les formules (10.3.4.4) et (10.3.4.5) ramènent le calcul d'un déterminant d'ordre n au calcul d'un déterminant d'ordre $n - 1$ donc, par itérations successives, au calcul des déterminants d'ordre 3 et 2. Mais les calculs devenant assez pénibles lorsque l'ordre du déterminant est élevé, on a toujours intérêt, avant de développer un déterminant, à faire apparaître des zéros en ajoutant à une ligne (ou une colonne) une combinaison linéaire des autres lignes (ou des autres colonnes).

10.3.4.7. Exemples

Le développement suivant la première colonne du déterminant

$$D = \begin{vmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{vmatrix},$$

donne

$$D = a_{11}a_{22} - a_{21}a_{12}.$$

De même le développement suivant la première colonne du déterminant

$$D = \begin{vmatrix} a & b & c \\ a' & b' & c' \\ a'' & b'' & c'' \end{vmatrix}$$

donne

$$\begin{aligned} D &= a \begin{vmatrix} b' & c' \\ b'' & c'' \end{vmatrix} - a' \begin{vmatrix} b & c \\ b'' & c'' \end{vmatrix} + a'' \begin{vmatrix} b & c \\ b' & c' \end{vmatrix} \\ &= a(b'c'' - b''c') - a'(bc'' - b''c) + a''(bc' - b'c) \\ &= ab'c'' - ab''c' - a'bc'' + a'b''c + a''bc' - a''b'c. \end{aligned}$$

On dispose d'un moyen mnémotechnique – la règle de Sarrus – pour retrouver ce développement. On écrit

$$\begin{vmatrix} a & b & c \\ a' & b' & c' \\ a'' & b'' & c'' \end{vmatrix} \begin{vmatrix} a & b \\ a' & b' \\ a'' & b'' \end{vmatrix}$$

– Les éléments dont le produit doit être affecté du signe + sont ceux qui sont situés sur des «parallèles à la diagonale principale» (traits continus), et ceux dont le produit doit être affecté du signe – sont situés sur des «parallèles à la diagonale non principale» (traits pointillés).

– Le déterminant est égal à la somme de ces six termes.

10.3.4.8. Exemple

Soit $(\alpha_1, \dots, \alpha_n)$ une suite de n éléments de K ($n \geq 2$). On appelle **déterminant de VANDERMONDE** associé à la suite $(\alpha_1, \dots, \alpha_n)$ le scalaire

$$V(\alpha_1, \dots, \alpha_n) = \begin{vmatrix} 1 & \alpha_1 & \alpha_1^2 & \dots & \alpha_1^{n-1} \\ 1 & \alpha_2 & \alpha_2^2 & \dots & \alpha_2^{n-1} \\ \dots & \dots & \dots & \dots & \dots \\ 1 & \alpha_n & \alpha_n^2 & \dots & \alpha_n^{n-1} \end{vmatrix}$$

On suppose évidemment que les α_i sont distincts deux à deux car sinon le déterminant est nul.

$V(\alpha_1, \dots, \alpha_n)$ est un polynôme en α_1 , de degré $n - 1$, dont les zéros sont $\alpha_2, \alpha_3, \dots, \alpha_n$, comme on le voit facilement en remplaçant dans V , α_1 par $\alpha_2, \alpha_3, \dots, \alpha_n$. Donc il existe un scalaire λ non nul tel que :

$$V(\alpha_1, \dots, \alpha_n) = \lambda(\alpha_2 - \alpha_1)(\alpha_3 - \alpha_1) \dots (\alpha_n - \alpha_1).$$

En égalant les coefficients de α_1^{n-1} dans les deux membres de cette dernière égalité, on obtient

$$(-1)^{n+1} D = (-1)^{n-1} \lambda$$

où D est le mineur relatif à α_1^{n-1} dans $V(\alpha_1, \dots, \alpha_n)$. D est le déterminant de Vandermonde associé à la suite $(\alpha_2, \dots, \alpha_n)$. Donc

$$V(\alpha_1, \dots, \alpha_n) = V(\alpha_2, \dots, \alpha_n) \prod_{i=2}^n (\alpha_i - \alpha_1).$$

Par récurrence, on en déduit

$$V(\alpha_1, \dots, \alpha_n) = \prod_{1 \leq i < j \leq n} (\alpha_j - \alpha_i).$$

Chapitre 11 : APPLICATIONS DES DÉTERMINANTS

Dans ce chapitre, nous allons appliquer la théorie des déterminants à la recherche de l'inverse d'une matrice carrée inversible, à la détermination du rang d'une application linéaire, puis à la résolution des systèmes d'équations linéaires. Cette étude montre l'importance fondamentale des déterminants dans les applications et justifie amplement l'intérêt qui leur est accordé.

Nous désignerons toujours par K un corps commutatif qui, dans la plupart des cas sera $\mathbb{R}$ ou $\mathbb{C}$.

11.1. Calcul de l'inverse d'une matrice carrée

D'après le Théorème 10.3.3.2 d), une matrice carrée A est inversible si et seulement si $\det(A) \neq 0$. Nous nous proposons de montrer comment utiliser les déterminants pour calculer la matrice inverse d'une matrice inversible.

11.1.0.1. Définition

Soit $A = (a_{ij})$ une matrice carrée d'ordre n . On appelle matrice complémentaire de A , et on note $\tilde{A}$, la transposée de la matrice des cofacteurs de A .

Autrement dit, si $\tilde{A} = (b_{ij})$, on a

$$b_{ij} = (-1)^{i+j} \det(A_{ji})$$

où A_{ij} est la matrice déduite de A par suppression de la $i^{\text{ème}}$ ligne et de la $j^{\text{ème}}$ colonne.

11.1.0.2. Théorème

Quelle que soit la matrice $A \in M_p(K)$, on a

$$A\tilde{A} = \tilde{A}A = \det(A) \cdot I_p$$

où I_p désigne la matrice unité d'ordre p .

Par suite, si A est inversible, son inverse est donné par :

$$A^{-1} = \frac{1}{\det(A)} \cdot \tilde{A}$$

Démonstration. Soit $A = (a_{ij})$. Posons

$$\tilde{A} = (b_{ij}) \text{ et } A\tilde{A} = (c_{ij}).$$

On a

$$b_{ij} = (-1)^{i+j} \det(A_{ji}).$$

Par définition du produit de deux matrices, on a

$$c_{ij} = \sum_{k=1}^p a_{ik} b_{kj} = \sum_{k=1}^p (-1)^{k+j} a_{ik} \det(A_{jk}).$$

Si $i = j$, la formule (10.3.4.5) montre que c_{ii} est égal à $\det(A)$. Si $i \neq j$, on voit que c_{ij} n'est autre que le déterminant de la matrice obtenue en remplaçant dans A la $j^{\text{ème}}$ ligne par la $i^{\text{ème}}$ sans toucher aux autres. Cette matrice ayant deux lignes identiques, son déterminant est nul : $c_{ij} = 0$ si $i \neq j$.

On a donc

$$A\tilde{A} = \det(A) \cdot I_p.$$

On montrerait de même que

$$\tilde{A}A = \det(A) \cdot I_p.$$

Si la matrice A est inversible, on a $\det(A) \neq 0$. Alors

$$A \cdot \frac{\tilde{A}}{\det(A)} = \frac{\tilde{A}}{\det(A)} \cdot A = I_p,$$

d'où

$$A^{-1} = \frac{1}{\det(A)} \cdot \tilde{A}.$$

Règle pratique :

Posons $A^{-1} = (\alpha_{ij})$. Alors si Δ_{ij} est le cofacteur de a_{ij} dans A , la formule

$$A^{-1} = \frac{1}{\det(A)} \cdot \tilde{A}$$

montre que

$$\alpha_{ij} = \frac{\Delta_{ji}}{\det(A)}.$$

Ainsi, on obtient l'inverse d'une matrice inversible A d'ordre $p > 1$, en formant la transposée $\tilde{A}$ de la matrice des cofacteurs de A , puis en divisant tous les éléments de $\tilde{A}$ par $\det(A)$.

11.1.0.3. Exemple

Considérons la matrice

$$A = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$$

Si $\det(A) = ad - bc \neq 0$, A est inversible. On a

$$\tilde{A} = \begin{pmatrix} d & -b \\ -c & a \end{pmatrix};$$

d'où

$$A^{-1} = \frac{1}{ad - bc} \begin{pmatrix} d & -b \\ -c & a \end{pmatrix}.$$

11.2. Détermination du rang

Nous allons appliquer la théorie des déterminants à la recherche pratique du rang d'une application linéaire, d'une matrice ou d'un système de vecteurs. Le calcul du rang d'une application linéaire ou d'une famille de vecteurs peut d'ailleurs se ramener à celui du rang d'une matrice puisque le rang d'une application linéaire est celui de sa matrice dans des bases données, qui est aussi celui des vecteurs colonnes.

Posons la définition suivante :

11.2.0.1. Définition

Soit A une matrice de type (m, n) . On appelle **matrice extraite** de A , toute matrice obtenue en supprimant un certain nombre de lignes et un certain nombre de colonnes de A .

On appelle **déterminant extrait** de A , tout déterminant d'une matrice carrée extraite de A .

Nous ne considérerons que les matrices carrées extraites et nous avons le résultat suivant.

11.2.0.2. Théorème

Soit $A = (a_{ij})$ une matrice de type (m, n) à éléments dans K . Alors le rang de A est le plus grand entier r tel que l'on puisse extraire de A au moins une matrice carrée inversible d'ordre r .

Démonstration. Soient E un K -espace vectoriel de dimension m et $\mathcal{B} = (e_1, \dots, e_m)$ une base de E . Soient $x_1, \dots, x_n$ les vecteurs colonnes de la matrice A :

$$x_j = \sum_{i=1}^m a_{ij} e_i \quad (1 \leq j \leq n).$$

D'après le Théorème 9.2.1.5 le rang de la famille $\{x_1, \dots, x_n\}$ est égal au rang de la matrice A .

Soit ρ le rang de la matrice A et soit r le plus grand entier tel que l'on puisse extraire de A une matrice carrée inversible d'ordre r . On a évidemment

$$r \leq \inf(m, n).$$

Nous allons montrer que $\rho = r$.

Démontrons d'abord que $\rho \geq r$. Soit M une matrice carrée inversible d'ordre r extraite de A . On peut supposer, en changeant au besoin la numérotation des lignes et des colonnes de A , que M est la matrice formée par les r premières lignes et les r premières colonnes de A . Posons :

$$A = \begin{pmatrix} a_{11} & \dots & a_{1r} & \dots & a_{1n} \\ \vdots & & \vdots & & \\ \vdots & & \vdots & & \\ \vdots & M & \vdots & & \\ \vdots & & \vdots & & \\ a_{r1} & \dots & a_{rr} & \dots & a_{rn} \\ \vdots & & \vdots & & \\ \vdots & & \vdots & & \\ a_{m1} & \dots & a_{mr} & \dots & a_{mn} \end{pmatrix}$$

Les r premiers vecteurs colonnes $x_1, \dots, x_r$ de A sont linéairement indépendants ; sinon l'un d'eux, x_1 par exemple, serait une combinaison linéaire des autres :

$$x_1 = \lambda_2 x_2 + \dots + \lambda_r x_r.$$

Dans la matrice M , les éléments de la première colonne seraient des combinaisons linéaires des éléments correspondants des autres colonnes et le déterminant de M serait nul contrairement au fait que M est inversible.

Comme parmi les vecteurs $x_1, \dots, x_n$ il existe r vecteurs linéairement indépendants, la dimension du sous-espace vectoriel engendré par $x_1, \dots, x_n$, c'est-à-dire le rang ρ de A , est supérieure ou égale à r .

Démontrons maintenant que $r \geq \rho$. Pour cela, nous allons extraire de la matrice A , une matrice carrée inversible d'ordre ρ .

Par définition du rang de la matrice A , il existe ρ vecteurs colonnes de A qui sont linéairement indépendants. Soit A' la matrice de type (m, ρ) obtenue en supprimant $n - \rho$ colonnes de A . On peut extraire de A' , ρ vecteurs lignes

linéairement indépendants puisque le rang d'une matrice est le rang de ses vecteurs colonnes qui est aussi le rang de ses vecteurs lignes. Soit A'' la matrice carrée d'ordre ρ obtenue en supprimant $m - \rho$ lignes de A' . Alors A'' est une matrice inversible de rang ρ extraite de A . Donc on a $r \geq \rho$.

On en déduit $r = \rho$ et le théorème est démontré. $\square$

Nous déduisons de ce théorème un critère général pour que des vecteurs soient linéairement indépendants.

11.2.0.3. Théorème

Soient E un K -espace vectoriel de dimension n , $B = (e_1, \dots, e_n)$ une base de E et $x_1, \dots, x_p$ des vecteurs de E définis par :

$$x_j = \sum_{i=1}^n a_{ij} e_i \quad (1 \leq j \leq p).$$

Soit $M = (a_{ij})_{1 \leq i \leq n, 1 \leq j \leq p}$ la matrice des coordonnées de ces vecteurs par rapport à la base B . Alors pour que $x_1, \dots, x_p$ soient linéairement indépendants, il faut et il suffit que l'on puisse extraire de M une matrice carrée inversible d'ordre p .

Démonstration. Nous savons que si $p > n$, les vecteurs $x_1, \dots, x_p$ sont linéairement dépendants. Supposons donc $p \leq n$. La famille $\{x_1, \dots, x_p\}$ est libre si et seulement si le sous-espace vectoriel de E engendré par cette famille est de dimension p , c'est-à-dire si et seulement si le rang de $\{x_1, \dots, x_p\}$ est égal à p . D'après le Théorème 9.2.15, le rang de la famille $\{x_1, \dots, x_p\}$ est égal à celui de la matrice M . Le Théorème 11.2.0.2 montre alors que les vecteurs $x_1, \dots, x_p$ sont linéairement indépendants si et seulement si on peut extraire de M une matrice carrée inversible d'ordre p . $\square$

Caractérisation du sous-espace engendré par un système libre.

Le résultat suivant nous sera fort utile dans l'étude des systèmes d'équations linéaires.

11.2.0.4. Théorème

Soient E un K -espace vectoriel de dimension n , $B = (e_1, \dots, e_n)$ une base de E et $x_1, \dots, x_r$, r vecteurs linéairement indépendants de E ($r < n$) définis par

$$x_j = \sum_{i=1}^n a_{ij} e_i \quad (1 \leq j \leq r).$$

Alors pour qu'un vecteur v de composantes $b_1, \dots, b_n$ de E appartienne au sous-espace vectoriel engendré par $x_1, \dots, x_r$, il faut et il suffit que les $n - r$ déterminants

$$\Delta_k = \begin{vmatrix} a_{11} & \dots & a_{1r} & b_1 \\ \vdots & & \vdots & \\ a_{r1} & \dots & a_{rr} & b_r \\ a_{k1} & \dots & a_{kr} & b_k \end{vmatrix}$$

soient nuls pour tout $k \in \{r+1, \dots, n\}$.

Démonstration. Comme les vecteurs $x_1, \dots, x_r$ sont linéairement indépendants, on peut extraire de la matrice $M = (a_{ij})$ qui est de type (n, r) une matrice carrée P inversible d'ordre r . En changeant au besoin l'ordre de numérotation des lignes et des colonnes de M , nous pouvons supposer que

$$P = \begin{pmatrix} a_{11} & \dots & a_{1r} \\ \vdots & & \vdots \\ a_{r1} & \dots & a_{rr} \end{pmatrix}.$$

Considérons la matrice

$$M' = \begin{pmatrix} a_{11} & \dots & a_{1r} & b_1 \\ \vdots & & \vdots & \vdots \\ a_{r1} & \dots & a_{rr} & b_r \\ \vdots & & \vdots & \vdots \\ a_{n1} & \dots & a_{nr} & b_n \end{pmatrix}$$

dont les vecteurs colonnes sont les vecteurs $x_1, \dots, x_r, v$.

Si le vecteur v appartient au sous-espace vectoriel engendré par $x_1, \dots, x_r$, on a $\text{rg}(M') < r+1$ et, puisque P est inversible, le rang de M' est égal à r . Tous les déterminants d'ordre $r+1$ extraits de M' sont donc nuls. En particulier, tout déterminant d'ordre $r+1$ extrait de M' et dont $\det(P)$ est un déterminant extrait est nul. Autrement dit, tous les déterminants Δ_k sont nuls pour $k = r+1, \dots, n$.

Réciproquement, supposons que $\Delta_k = 0$ pour tout $k \in \{r+1, \dots, n\}$.

Si $k \leq r$, Δ_k est évidemment nul car ce déterminant a deux lignes identiques.

Si $k > r$, en développant Δ_k par rapport à la ligne d'indice $r+1$, on obtient

$$\sum_{j=1}^r a_{kj} \Delta_{r+1,j} + \det(P) b_k = 0$$

où les $\Delta_{r+1,j}$ sont les cofacteurs des éléments de la ligne d'indice $r+1$ de Δ_k .

On en déduit :

$$x_1 \Delta_{r+1,1} + x_2 \Delta_{r+1,2} + \dots + x_r \Delta_{r+1,r} + \det(P) v = 0,$$

ce qui achève la démonstration du théorème. $\square$

Posons

$$X = \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix}, \quad B = \begin{pmatrix} b_1 \\ b_2 \\ \vdots \\ b_m \end{pmatrix}.$$

Alors le système (11.3.1.1) est équivalent à l'équation matricielle

$$(11.3.2.1) \quad AX = B$$

dans laquelle A et B sont des matrices connues et X une matrice inconnue. On dit que (11.3.2.1) est la **forme matricielle** du système 11.3.1.1.

11.3.2.1. Définition

On appelle rang du système (11.3.1.1) le rang de la matrice A de ce système.

b) Soient $A_1, \dots, A_n$ les vecteurs colonnes de la matrice A ; ce sont des vecteurs de K^m définis par leurs composantes dans la base canonique :

$$A_1 = \begin{pmatrix} a_{11} \\ \vdots \\ a_{m1} \end{pmatrix}, \dots, \quad A_n = \begin{pmatrix} a_{1n} \\ \vdots \\ a_{mn} \end{pmatrix}.$$

Notons encore B le vecteur de K^m ayant pour composantes $b_1, \dots, b_m$.

Le système (11.3.1.1) est alors équivalent à l'équation vectorielle dans K^m :

$$(11.3.2.2) \quad x_1 A_1 + \dots + x_n A_n = B.$$

On dit que (11.3.2.2) est la **forme vectorielle** du système (11.3.1.1).

Résoudre le système (11.3.1.1) revient donc à chercher si B est combinaison linéaire des vecteurs $A_1, \dots, A_n$, i.e. si B appartient au sous-espace vectoriel de K^m engendré par les vecteurs $A_1, \dots, A_n$ et, dans l'affirmative, à trouver toutes les décompositions possibles de B sur la famille $\{A_1, \dots, A_n\}$.

11.4. Système de Cramer

Dans ce paragraphe, nous supposons que $m = n$; autrement dit, le système considéré contient autant d'équations que d'inconnues.

11.4.1. DÉFINITION

11.4.1.1. Définition

On dit qu'un système de n équations linéaires à n inconnues est un système de Cramer si la matrice A de ce système est inversible.

Un système de Cramer peut toujours s'écrire sous la forme matricielle.

$$(11.4.1.1) \quad AX = B.$$

Nous allons voir qu'un tel système possède toujours une solution unique donnée par :

$$X = A^{-1}B.$$

On appelle **déterminant du système** (11.4.1.1) le déterminant de la matrice A de ce système.

Le théorème suivant donne plusieurs caractérisations des système de Cramer.

11.4.1.2. Théorème

Soit

[illegible]

un système de n équations linéaires à n inconnues écrit sous forme matricielle : $AX = B$. Les propriétés suivantes sont équivalentes :

- a) Quel que soit B , le système (11.4.1.2) admet une solution et une seule.
- b) Quel que soit B , le système (11.4.1.2) admet au moins une solution.
- c) Quel que soit B , le système (11.4.1.2) admet au plus une solution.
- d) Le système homogène associé au système (11.4.1.2) n'admet que la solution triviale.
- e) La matrice A du système (11.4.1.2) est inversible.
- f) $\det(A) \neq 0$.

La solution unique du système (11.4.1.2) est alors

$$X = A^{-1}B.$$

Démonstration. Soit u l'endomorphisme de K^n dont la matrice par rapport à la base canonique de K^n est A . L'assertion a) signifie que u est bijectif, b) que u

est surjectif, c) que u est injectif, d) que le noyau de u est réduit à $\{0\}$, e) que u est inversible, et f) que $\det(u) \neq 0$. Les assertions a), b), c), d) et e) sont donc équivalentes d'après le Corollaire 8.2.3.5 et le Théorème 10.3.2.3 d).

En outre, comme la matrice A est inversible, on obtient la solution unique du système (11.4.1.2) en multipliant à gauche les deux membres par A^{-1} . On obtient

$$X = A^{-1}B. \quad \square$$

11.4.2. FORMULES DE CRAMER

Si on a déjà calculé la matrice inversée A^{-1} , la formule $X = A^{-1}B$ permet d'obtenir facilement la solution unique d'un système de Cramer. Sinon, on peut utiliser les formules de Cramer que nous allons établir.

11.4.2.1. Théorème

Soit

$$(11.4.2.1) \quad \sum_{i=1}^n a_{ji}x_i = b_j, \quad (1 \leq j \leq n)$$

un système de Cramer.

La solution unique $(x_1, \dots, x_n)$ de ce système est donnée par les formules de Cramer :

$$(11.4.2.2) \quad x_i = \frac{\Delta_i}{\Delta} \quad (1 \leq i \leq n),$$

où Δ est le déterminant du système et où Δ_i est le déterminant déduit de Δ en remplaçant la $i^{\text{ème}}$ colonne par la colonne des termes constants $b_1, \dots, b_n$.

Démonstration. Le système (11.4.2.1) est équivalent à l'équation

$$(11.4.2.3) \quad x_1A_1 + \dots + x_nA_n = B$$

où les vecteurs $A_1, \dots, A_n$ de K^n sont linéairement indépendants car $\det(A) \neq 0$. Ces vecteurs forment une base de K^n et par suite, il existe un système unique de scalaires $x_1, \dots, x_n$ (les composantes du vecteur B dans la base $(A_1, \dots, A_n)$) vérifiant (11.4.2.3).

Calculons le déterminant :

$$\Delta_i = \det(A_1, \dots, A_{i-1}, B, A_{i+1}, \dots, A_n)$$

dans la base canonique de K^n .

Puisque l'application déterminant est une forme n -linéaire alternée, on a

$$\begin{aligned}\Delta_i &= \det \left(A_1, \dots, A_{i-1}, \sum_{k=1}^n x_k A_k, A_{i+1}, \dots, A_n \right) \\ &= \sum_{k=1}^n \det(A_1, \dots, A_{i-1}, A_k, A_{i+1}, \dots, A_n) \\ &= x_i \det(A_1, \dots, A_{i-1}, A_i, A_{i+1}, \dots, A_n) = x_i \Delta.\end{aligned}$$

D'où, puisque $\Delta \neq 0$,

$$x_i = \frac{\Delta_i}{\Delta}.$$

11.4.2.2. Remarque

Les formules de Cramer conduisent souvent à des calculs pénibles de déterminants. Il est possible dans certains cas de s'en passer en recourant à une succession de combinaisons linéaires des lignes du système.

11.4.2.3. Exemple

Résoudre le système

$$(S) = \begin{cases} y + z = 1 & (1) \\ x + z = 2 & (2) \\ x + y = 0 & (3) \end{cases}$$

Le déterminant du système (S) est

$$\Delta = \begin{vmatrix} 0 & 1 & 1 \\ 1 & 0 & 1 \\ 1 & 1 & 0 \end{vmatrix} = 2.$$

D'où, d'après les formules de Cramer :

$$x = \frac{\begin{vmatrix} 1 & 1 & 1 \\ 2 & 0 & 1 \\ 0 & 1 & 0 \end{vmatrix}}{\Delta} = \frac{1}{2}, \quad y = \frac{\begin{vmatrix} 0 & 1 & 1 \\ 1 & 2 & 1 \\ 1 & 0 & 0 \end{vmatrix}}{\Delta} = -\frac{1}{2}$$

$$z = \frac{\begin{vmatrix} 0 & 1 & 1 \\ 1 & 0 & 2 \\ 1 & 1 & 0 \end{vmatrix}}{\Delta} = \frac{3}{2}.$$

On peut retrouver ces résultats plus rapidement en procédant de la façon suivante :

Additionnons les équations (1), (2) et (3) ; le système (S) est équivalent au système

$$\begin{cases} y + z &= 1 \\ x + z &= 2 \\ x + y + z &= \frac{3}{2} \end{cases}$$

En retranchant la 2^e équation de la 3^e, on obtient $y = -\frac{1}{2}$, d'où $x = \frac{1}{2}$ d'après la 3^e équation de (S). Enfin $z = 1 - y = 1 + \frac{1}{2} = \frac{3}{2}$.

11.5. Résolution d'un système linéaire quelconque

11.5.1. ÉQUATIONS PRINCIPALES. INCONNUES PRINCIPALES

Revenons maintenant au cas d'un système linéaire quelconque de m équations à n inconnues :

$$(11.5.1.1) \quad \sum_{i=1}^n a_{ji} x_i = b_j, \quad (1 \leq j \leq m).$$

Notons r le rang du système (11.5.1.1).

On sait qu'il existe une matrice carrée M inversible d'ordre r , extraite de la matrice A de ce système.

En changeant au besoin la numérotation des équations et des inconnues, on peut supposer que

$$M = \begin{pmatrix} a_{11} & \dots & a_{1r} \\ \dots & \dots & \dots \\ a_{r1} & \dots & a_{rr} \end{pmatrix}.$$

On dit que le déterminant Δ de la matrice M est le **déterminant principal** du système. Les r premières équations sont les **équations principales** et les inconnues $x_1, \dots, x_r$ sont les **inconnues principales**, les autres inconnues étant les **inconnues non principales**.

11.5.2. CONDITION DE COMPATIBILITÉ ET RÉOLUTION

Comme nous l'avons vu, le système (11.5.1.1) s'écrit sous forme vectorielle :

$$x_1 A_1 + \cdots + x_n A_n = B.$$

Les vecteurs colonnes $A_1, \dots, A_n$ engendrent un sous-espace vectoriel E de K^m et le système n'est possible que si le vecteur B appartient à E . On est donc ramené à trouver une base de E ou, ce qui revient au même, à déterminer le rang r du système (11.5.1.1).

Le cas où $r = m = n$ ayant été traité au paragraphe 11.4, nous allons examiner les cas où $r = m < n$ puis $r < m$.

1^{er} cas : $r = m < n$. Si le rang du système (11.5.1.1) est égal au nombre d'équations, on a $E = K^m$ donc B , qui est un vecteur K^m , appartient à E et par suite le système est compatible quel que soit B .

Les vecteurs colonnes $A_1, \dots, A_m$ forment une base de K^m et on peut écrire le système sous la forme :

$$x_1A_1 + \cdots + x_mA_m = B - x_{m+1}A_{m+1} - \cdots - x_nA_n = B'.$$

Quelles que soient les valeurs données aux inconnues non principales $x_{m+1}, \dots, x_n$, le vecteur B' appartient à K^m . On peut donc l'écrire de manière unique comme combinaison linéaire des vecteurs $A_1, \dots, A_m$. On obtient $x_1, \dots, x_m$ en résolvant par les formules des Cramer les m équations principales par rapport aux inconnues principales :

[illegible]

La résolution de ce système donne une solution unique en $x_1, \dots, x_m$ mais cette solution dépend des inconnues non principales $x_{m+1}, \dots, x_n$ qui ont des valeurs arbitraires dans K . On dit qu'il y a une **indétermination d'ordre $n - m$** .

2^e cas : $r < m$. D'après le choix de l'ordre des équations et de l'ordre des inconnues, les vecteurs $A_1, \dots, A_r$ sont linéairement indépendants; ils forment donc une base du sous-espace vectoriel E de K^m .

Donc si $r < m$, le système (11.5.1.1) est compatible si et seulement si le vecteur B appartient au sous-espace vectoriel E , donc (Théorème 11.2.0.4), si et seulement si les $m - r$ déterminants

$$\Delta_k = \begin{vmatrix} a_{11} & \dots & a_{1r} & b_1 \\ \dots & \dots & \dots & \dots \\ a_{r1} & \dots & a_{rr} & b_r \\ a_{k1} & \dots & a_{kr} & b_k \end{vmatrix}$$

sont nuls pour tout $k \in \{r + 1, \dots, m\}$.

Les déterminants $\Delta_k (r + 1 \leq k \leq m)$ s'appellent les **déterminants caractéristiques** du système (11.5.1.1).

Supposons que le système Σ soit compatible.

a) Si $r = n$, les vecteurs $A_1, \dots, A_n$ forment une base du sous-espace vectoriel E . Le vecteur B s'exprime donc, d'une manière et d'une seule, comme combinaison linéaire de $A_1, \dots, A_n$:

$$x_1 A_1 + \cdots + x_n A_n = B.$$

On obtient $x_1, \dots, x_n$ par les formules de Cramer en résolvant les n équations principales par rapport aux n inconnues principales.

b) Si $r < n$, on peut écrire :

$$x_1A_1 + \cdots + x_rA_r = B - x_{r+1}A_{r+1} - \cdots - x_nA_n = B''.$$

Quelles que soient les valeurs données aux inconnues non principales $x_{r+1}, \dots, x_n$ le vecteur B'' appartient au sous-espace vectoriel E engendré par les vecteurs $A_1, \dots, A_n$. On peut donc l'exprimer de manière unique comme combinaison linéaire des vecteurs de la base $(A_1, \dots, A_r)$ de E . On obtient $x_1, \dots, x_r$ en résolvant, par les formules de Cramer, les r équations principales par rapport aux r inconnues principales :

[illegible]

La solution dépend alors des $n - r$ inconnues arbitraires $x_{r+1}, \dots, x_n$.

c) Lorsque le système (11.5.1.1) est compatible, toute solution est également solution du système (11.5.1.2). Ce qui précède montre que, réciproquement, toute solution de (11.5.1.2) est solution du système (11.5.1.1). Ces deux systèmes sont donc équivalents dans le cas où le système (11.5.1.1) admet des solutions.

Résumons l'étude qui vient d'être faite en énonçant le théorème suivant :

11.5.2.1. Théorème (Théorème de Rouché-Fontené)

Soit (S) un système de m équations linéaires à n inconnues, de rang r .

a) Si $r = m = n$, le système admet une solution unique.

b) Si $r = m < n$, le système admet des solutions que l'on obtient en donnant des valeurs arbitraires aux $n - m$ inconnues non principales et en résolvant le système de Cramer aux m inconnues principales.

c) Si $r < m$, le système admet des solutions si et seulement si les $m - r$ déterminants caractéristiques sont tous nuls. Lorsque cette condition est vérifiée, le système se réduit aux r équations principales que l'on résout comme au b).

11.5.2.2. Exemple

Résoudre le système

$$(S) \begin{cases} \alpha x & + 2z & = & 2 \\ x & + 2y & & = & 1 \\ x & - 2y & + \beta z & = & 1 \end{cases}$$

où α et β sont des paramètres réels.

Le déterminant du système est

$$\Delta = \begin{vmatrix} \alpha & 0 & 2 \\ 1 & 2 & 0 \\ 1 & -2 & \beta \end{vmatrix} = 2(\alpha\beta - 4).$$

1^{er} cas : Si $\alpha\beta \neq 4$, on a un système de Cramer. Les formules de Cramer (11.4.2.2) donnent

$$x = \frac{2(\beta - 2)}{\alpha\beta - 4}, \quad y = \frac{\beta(\alpha - 2)}{2(\alpha\beta - 4)}, \quad z = \frac{2(\alpha - 2)}{\alpha\beta - 4}.$$

2^e cas : Si $\alpha\beta = 4$, on a

$$\begin{vmatrix} \alpha & 0 \\ 1 & 2 \end{vmatrix} = 2\alpha.$$

Ce déterminant étant non nul puisque $\alpha\beta = 4$, le système est de rang 2. Il y a un seul déterminant caractéristique qui doit être nul

$$\Delta_3 = \begin{vmatrix} \alpha & 0 & 2 \\ 1 & 2 & 1 \\ 1 & -2 & 1 \end{vmatrix} = 0$$

On a $\Delta_3 = 4(\alpha - 2)$.

a) Si $\alpha \neq 2$, le système est impossible.

b) Si $\alpha = 2$, alors $\beta = 2$ puisque $\alpha\beta = 4$. Le système admet des solutions ; les inconnues principales sont x et y . Nous sommes ramenés à résoudre le système :

$$\begin{cases} 2x = 2 - 2z \\ x + 2y = 1 \end{cases}$$

qui est un système de Cramer. Les solutions sont

$$x = 1 - z, \quad y = \frac{z}{2} \quad \text{et } z, \text{ où } z \in \mathbb{R}.$$

11.6. Systèmes homogènes

Considérons un système linéaire homogène de la forme

$$\sum_{i=1}^n a_{ji}x_i = 0 \quad (1 \leq j \leq m).$$

Avec les notations habituelles, ce système s'écrit :

$$x_1 A_1 + \dots + x_n A_n = 0.$$

Nous savons déjà qu'un tel système admet toujours au moins la solution banale $x_1 = 0, \dots, x_n = 0$.

Soit r le rang du système. La résolution d'un système homogène se ramène toujours à celle du système des r équations principales ; tous les déterminants caractéristiques sont nuls car ils ont une colonne de zéros.

a) Si $r = n$, les vecteurs $A_1, \dots, A_n$ sont linéairement indépendants. Le système n'admet que la solution triviale.

b) Si $r < n$, les $n - r$ inconnues non principales sont arbitraires et le système est indéterminé. On le résout par la méthode générale du paragraphe 11.5. En résumé, nous avons :

11.6.0.1. Théorème

Pour qu'un système linéaire homogène admette des solutions autres que la solution triviale, il faut et il suffit que le rang de la matrice du système soit inférieur au nombre des inconnues.

En particulier, un système linéaire homogène de n équations à n inconnues admet des solutions autres que la solution banale si, et seulement si, le déterminant de la matrice du système est nul.

Il est clair que si $x_1, \dots, x_n$ est une solution non nulle d'un système linéaire homogène, $\lambda x_1, \dots, \lambda x_n$ où λ est un scalaire, est également une solution de ce système. Nous sommes donc amenés à chercher une solution non nulle quelconque d'un système homogène

$$(S) \quad a_{j1}x_1 + \dots + a_{jn}x_n = 0 \quad (1 \leq j \leq m).$$

Soit $A = (a_{ij})$ la matrice du système (S).

Supposons par exemple que : $m = n$ et $r = n - 1$.

Il y a des solutions non nulles si $\det(A) = 0$. Par exemple, supposons que la matrice extraite de A par suppression de la dernière ligne et de la dernière colonne soit inversible. Soient $\Delta_{n1}, \dots, \Delta_{nn}$ les cofacteurs des éléments de la dernière ligne de A .

On sait (Théorème 11.1.0.2), que

$$a_{n1}\Delta_{n1} + \cdots + a_{nn}\Delta_{nn} = \det(A) = 0.$$

Donc $(\Delta_{n1}, \dots, \Delta_{nn})$ est une solution non nulle du système (S) . D'après la remarque précédente, $x_1 = \lambda\Delta_{n1}, \dots, x_n = \lambda\Delta_{nn}$, où λ est un scalaire quelconque, est une autre solution du système.

Supposons maintenant que : $m = n - 1$ et $r = n - 1$.

On se ramène au cas précédent en adjoignant à la matrice $A = (a_{ij})$ une $n^{\text{ème}}$ ligne arbitraire. Les cofacteurs $\Delta_{n1}, \dots, \Delta_{nn}$ des éléments de la dernière ligne de la nouvelle matrice obtenue ne sont pas tous nuls puisque le système est de rang $n - 1$. On montre comme précédemment que $(\Delta_{n1}, \dots, \Delta_{nn})$ est une solution non nulle du système considéré.

11.6.0.2. Exemple

Résoudre le système

$$\begin{cases} 2x - y + 3z = 0 \\ x + y + 2z = 0 \end{cases}$$

La matrice du système est

$$A = \begin{pmatrix} 2 & -1 & 3 \\ 1 & 1 & 2 \end{pmatrix};$$

elle est de rang 2, donc le système admet des solutions autres que la solution nulle.

En appliquant la méthode précédente, on voit qu'une solution non nulle est :

$$x_0 = \begin{vmatrix} -1 & 3 \\ 1 & 2 \end{vmatrix} = -5, \quad y_0 = -\begin{vmatrix} 2 & 3 \\ 1 & 2 \end{vmatrix} = -1, \quad z_0 = \begin{vmatrix} 2 & -1 \\ 1 & 1 \end{vmatrix} = 3,$$

d'où la solution générale

$$x = -5\lambda, \quad y = -\lambda, \quad z = 3\lambda, \quad \lambda \in \mathbb{R}.$$

PROBLÈMES

1

I. Dans tout le problème on désigne par E un espace vectoriel réel. On note I l'application identique de E .

1°) E_1 et E_2 étant deux sous-espaces vectoriels supplémentaires de E , on définit une application f de E dans E de la façon suivante :

Si $x = x_1 + x_2$ où $x_1 \in E_1$ et $x_2 \in E_2$, on pose

$$f(x) = x_1 - x_2.$$

a) Montrer que f est une application linéaire et que l'on a $f^2 = I$.

b) Déterminer les noyaux des applications linéaires

$$u = \frac{I - f}{2} \quad \text{et} \quad v = \frac{I + f}{2}.$$

et les sous-espaces vectoriels $u(E)$ et $v(E)$. Evaluer u^2 , v^2 , uv et vu .

2°) Réciproquement soit f un endomorphisme de E tel que $f^2 = I$. On pose

$$u = \frac{I - f}{2} \quad \text{et} \quad v = \frac{I + f}{2}.$$

Montrer que $u(E)$ et $v(E)$ sont stables par f et sont deux sous-espaces supplémentaires de E .

Montrer que f peut être obtenu de la manière considérée ci-dessus au 1°).

3°) On suppose maintenant que E est de dimension finie n .

a) Trouver la matrice A de l'application linéaire f . En déduire que l'application transposée de f s'obtient de la même manière que f .

b) Soit I_n la matrice unité d'ordre n et soit λ un scalaire. Former le polynôme $P(\lambda) = \text{Dét}(A - \lambda I_n)$ et trouver ses racines.

II. Soient a et b deux nombres réels distincts.

1°) E_1 et E_2 étant deux sous-espaces vectoriels supplémentaires de E , on définit une application de E dans E de la façon suivante :

Si $x = x_1 + x_2$ où $x_1 \in E_1$ et $x_2 \in E_2$, on pose $f(x) = ax_1 + bx_2$.

a) Montrer que f est une application linéaire et que l'on a :

$$f^2 - (a + b)f + abI = 0.$$

b) Déterminer les noyaux des applications linéaires

$$u = \frac{f - aI}{b - a} \text{ et } v = \frac{f - bI}{a - b}$$

et les sous-espaces vectoriels $u(E)$ et $v(E)$.

2°) Calculer u^2, v^2, uv et vu .

Calculer f en fonction de u et v ; en déduire f^n pour n entier > 0 .

3°) Montrer que si $ab \neq 0$, alors f est inversible. Exprimer alors f^{-1} à l'aide de u et v .

4°) Inversement, soit f un endomorphisme de E tel que :

$$f^2 - (a + b)f + abI = 0.$$

a) Montrer qu'il existe deux sous-espaces supplémentaires F et G tels que f puisse être défini comme il a été dit au II, 1°) à partir de F et G .

N.B. Les deux parties de ce problème sont indépendantes mais il est conseillé de résoudre d'abord la partie I.

DUES MP1, Abidjan, Juin 1969.

2

I. Soit $x \in \mathbb{R}^4$. On pose :

$$x = \begin{pmatrix} x_0 \\ x_1 \\ x_2 \\ x_3 \end{pmatrix}, \quad tx = (x_0, x_1, x_2, x_3)$$

$$J = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & -1 & 0 & 0 \\ 0 & 0 & -1 & 0 \\ 0 & 0 & 0 & -1 \end{pmatrix}.$$

On désigne par G le sous-ensemble de $M_4(\mathbb{R})$ formé des matrices M qui laissent invariante la forme quadratique

$$f(x) = x_0^2 - x_1^2 - x_2^2 - x_3^2.$$

(on dit que $f(x)$ est invariante par la matrice M , si la relation $y = Mx$ implique $y_0^2 - y_1^2 - y_2^2 - y_3^2 = x_0^2 - x_1^2 - x_2^2 - x_3^2$).

1°) Montrer que $f(x) = {}^t x J x$ et que si $M \in G$, on a ${}^t M J M = J$ où ${}^t M$ est la transposée de M . En déduire que si $M \in G$, $M = (a_{ij})$, on a $a_{11}^2 \geq 1$ et $\det(M) = \pm 1$.

2°) Montrer que G est un groupe multiplicatif.

3°) Soit $M \in G$. Calculer l'expression $(MJ)({}^t M J M)(M^{-1}J)$ et montrer que ${}^t M \in G$.

II. On rappelle qu'une matrice carrée A à éléments complexes est dite hermitienne si $A^* = A$ où A^* est la matrice ${}^t \bar{A}$ ($\bar{A}$ est la matrice complexe conjuguée de A). Quelle est la forme générale des matrices carrées hermitiennes d'ordre 2?

1°) On associe à tout point $(x_0, x_1, x_2, x_3) \in \mathbb{R}^4$, la matrice

$$X = \begin{pmatrix} x_0 + x_1 & x_2 + ix_3 \\ x_2 - ix_3 & x_0 - x_1 \end{pmatrix}.$$

Montrer que X est hermitienne et que réciproquement, toute matrice hermitienne est de cette forme.

2°) On note $SL(2, \mathbb{C})$ le sous-ensemble de $M_2(\mathbb{C})$ formé des matrices dont le déterminant est égal à 1. Si $g \in SL(2, \mathbb{C})$, on pose $X' = gXg^*$.

Montrer que X' peut se mettre sous la forme

$$X' = \begin{pmatrix} x'_0 + x'_1 & x'_2 + ix'_3 \\ x'_2 - ix'_3 & x'_0 - x'_1 \end{pmatrix}$$

avec x'_0, x'_1, x'_2, x'_3 réels et $i = \sqrt{-1}$.

3°) On considère la matrice $T(g)$ qui, à (x_0, x_1, x_2, x_3) associe le vecteur (x'_0, x'_1, x'_2, x'_3) .

Montrer que $T(g) \in G$ et que l'application $g \mapsto T(g)$ est un homomorphisme de $SL(2, \mathbb{C})$ dans G . Quel est le noyau de cet homomorphisme?

DUES MP1, Abidjan, Juin 1971

3

I. On note G l'ensemble des applications de $\mathbb{R}$ dans $\mathbb{R}$, définies par $f_{a,b}(x) = ax + b$, où $a > 0$ et b sont des nombres réels.

COURS D'ALGÈBRE

1°) Montrer que $f_{a,b}$ est une bijection de $\mathbb{R}$ sur $\mathbb{R}$. Quelle est sa bijection réciproque ?

2°) Montrer que G est un groupe pour la composition des applications. Est-il abélien ?

3°) Soient H l'ensemble des éléments de G de la forme $f_{a,0}$ et K l'ensemble des éléments de G de la forme $f_{1,b}$.

Montrer que :

a) Tout élément de G peut s'écrire comme composé d'un élément de H par un élément de K .

b) H est un sous-groupe de G isomorphe à $\mathbb{R}_+^* = \{x \in \mathbb{R} : x > 0\}$.

c) K est un sous-groupe distingué de G isomorphe à $\mathbb{R}$.

II. Soit K un corps commutatif.

1°) Montrer qu'il existe un homomorphisme unique f de l'anneau $\mathbb{Z}$ dans l'anneau K tel que $f(1) = 1$.

2°) Montrer que le noyau de f est un idéal I de $\mathbb{Z}$ tel que, quels que soient $x, y \in \mathbb{Z}$, la relation $xy \in I$ implique $x \in I$ ou $y \in I$.

En déduire que $I = \{0\}$ ou $I = p\mathbb{Z}$ où p est un nombre premier.

3°) On pose $K' = f(\mathbb{Z})$. Démontrer que :

a) Si $I = \{0\}$, alors K' est un sous-anneau de K ;

b) Si $I = p\mathbb{Z}$, alors K' est un sous-corps de K .

4°) Montrer que si A est un sous-anneau de K , K' est inclus dans A .

DUES MP1, Abidjan, Janvier 1986

4

I. Soient E et F deux espaces vectoriels sur un même corps commutatif. On dit qu'une application linéaire $f : E \longrightarrow F$ est triviale si $f(x) = 0$ pour tout $x \in E$.

1°) Soit $f \in \mathcal{L}(E, F)$. Montrer que les conditions suivantes sont équivalentes :

a) f est triviale.

b) $\text{Ker}(f) = E$.

c) $\text{Im}(f) = \{0\}$.

2°) Etant donnée une suite

$$E_1 \xrightarrow{f_1} E_2 \xrightarrow{f_2} E_3 \longrightarrow \dots \longrightarrow E_n \xrightarrow{f_n} E_{n+1}$$

d'espaces vectoriels sur un même corps commutatif et d'applications linéaires $f_k : E_k \longrightarrow E_{k+1}$, on dit que cette suite est exacte si pour $1 \leq k < n$, on a

$$\text{Im}(f_k) = \text{Ker}(f_{k+1}).$$

On note $\{0\}$ l'espace vectoriel réduit à 0.

a) Montrer que si $E_1 = \{0\}$ (resp. $E_{n+1} = \{0\}$), f_1 (resp. f_n) est l'application triviale que l'on n'écrit pas dans la suite.

b) Donner les conditions nécessaires et suffisantes pour que les suites suivantes soient exactes

$$\begin{aligned} \{0\} &\longrightarrow E \xrightarrow{f} F \\ E &\xrightarrow{f} F \longrightarrow \{0\} \\ \{0\} &\longrightarrow E \xrightarrow{f} F \longrightarrow \{0\} \end{aligned}$$

3°) Soit E_1 un sous-espace vectoriel de E . Montrer que la suite

$$\{0\} \longrightarrow E_1 \xrightarrow{j} E \xrightarrow{\pi} E/E_1 \longrightarrow \{0\}$$

(où j est l'injection canonique et π la surjection canonique) est exacte.

4°) Soit

$$\{0\} \longrightarrow E \xrightarrow{\pi} E/\text{Ker}(f) \xrightarrow{\bar{f}} f(E) \xrightarrow{j} F \longrightarrow \{0\},$$

la suite de décomposition canonique de l'application linéaire $f : E \longrightarrow F$. Est-elle exacte ?

II. Soit :

$$A = \begin{pmatrix} 1 & 1 & 1 \\ 0 & 1 & 1 \\ 0 & 0 & 1 \end{pmatrix}$$

1°) Calculer A^n pour n entier ≥ 1 . (On pourra écrire $A = I_3 + B$, où I_3 désigne la matrice unité d'ordre 3.)

2°) Soient (x_n) , (y_n) et (z_n) trois suite définies par les relations

$$x_n = x_{n-1} + y_{n-1} + z_{n-1}$$

$$y_n = y_{n-1} + z_{n-1}$$

$$z_n = z_{n-1}$$

et la donnée de x_0, y_0, z_0 .

Calculer x_n, y_n, z_n en fonction de x_0, y_0, z_0 et de n .

DUES MP1, Abidjan, Mars 1986

5

I. Soit A un anneau (unitaire), 1 son élément unité. On appelle dérivation de A toute application d de A dans A telle que

a) $d(x + y) = d(x) + d(y)$

b) $d(xy) = d(x)y + xd(y)$ quels que soient $x, y \in A$.

1°) Montrer que la dérivation ordinaire des polynômes à coefficients réels est une dérivation de $\mathbb{R}[X]$.

2°) Soient d_1 et d_2 deux dérivations de A .

a) Montrer que $d_1 + d_2$ est une dérivation de A .

b) L'application $d_1 \circ d_2$ est-elle une dérivation de A ?

Montrer que l'application $[d_1, d_2] = d_1 \circ d_2 - d_2 \circ d_1$ est une dérivation de A .

3°) On suppose désormais que A est commutatif. On pose $d^0 = 1_A$, $d^1 = d$ et par récurrence sur n , $d^n = d \circ d^{n-1}$. Montrer que l'on a

$$d^n(xy) = \sum_{k=0}^n C_n^k d^k(x) d^{n-k}(y)$$

quels que soient $x, y \in A$ et $n \in \mathbb{N}$

4°) On suppose désormais que A est intègre et de caractéristique p . Montrer que d^p est une dérivation de A .

DUES MP1, Abidjan, Juin 1986 (extrait)

6

I. On note $\mathbb{C}$ le corps des nombres complexes et $\mathbb{R}$ celui des nombres réels.

1°) Montrer que l'ensemble G des matrices $g = \begin{pmatrix} a & b \\ \bar{b} & \bar{a} \end{pmatrix}$ à coefficients complexes tels que $|a|^2 - |b|^2 = 1$ est un groupe multiplicatif.

Montrer que $G = \{g \in M_2(\mathbb{C}) : g^* J g = J \text{ et } \det(g) = 1\}$ où $J = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}$ et où $g^* = {}^t \bar{g}$ est l'adjointe de g .

2°) Soit $D = \{z \in \mathbb{C} : |z| < 1\}$. Pour tout $g \in G$ et pour tout $z \in D$, on pose

$$\psi_g(z) = \frac{az + b}{\bar{b}z + \bar{a}}.$$

Montrer que ψ_g est une application de D dans D et que l'application $g \mapsto \psi_g$ est un homomorphisme. Trouver le sous-groupe K de G tel que $\psi_g(0) = 0$ pour tout $g \in K$.

3°) On pose

$$C = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 & -i \\ 1 & i \end{pmatrix}$$

et on considère l'automorphisme intérieur de $GL(2, \mathbb{C})$ défini par $\sigma(g) = CgC^{-1}$.

Montrer que $\sigma(g) \in G$ si et seulement si g appartient au groupe multiplicatif G' formé des matrices $\begin{pmatrix} \alpha & \beta \\ \gamma & \delta \end{pmatrix}$, où $\alpha, \beta, \gamma, \delta$ sont réels et vérifient $\alpha\delta - \gamma\beta = 1$.

Quelle est l'image par σ du sous-groupe N de G' formé des matrices de la forme $\begin{pmatrix} 1 & \xi \\ 0 & 0 \end{pmatrix}$, $\xi \in \mathbb{R}$?

II. On se donne trois nombres réels différents a, b, c et le polynôme en X :

$$P(X) = X^3 + X^2z + Xy + x.$$

En considérant le système d'équations

$$(1) \quad \begin{cases} x + ay + a^2z + a^3 = 0 \\ x + by + b^2z + b^3 = 0 \\ x + cy + c^2z + c^3 = 0 \end{cases}$$

montrer que $P(X)$ est divisible par le polynôme $(X - a)(X - b)(X - c)$.

En déduire la solution du système (1). Retrouver cette solution par une méthode directe.

7

I. Soit A un anneau commutatif unitaire. Soient I et J deux idéaux de A . On pose

$$I + J = \{x + y : x \in I \text{ et } y \in J\}.$$

1°) Montrer que :

a) $I + J$ et $I \cap J$ sont des idéaux de A .

b) $I + J = A$ si et seulement si, il existe $a \in I$, il existe $b \in J$ tels que $a + b = 1$.

2°) On note

$$IJ = \left\{ a \in A / \exists n \in \mathbb{N}^*, \exists x_1, \dots, x_n \in I, \exists y_1, \dots, y_n \in J \text{ et } a = \sum_{i=1}^n x_i y_i \right\}.$$

Montrer que IJ est un idéal de A et que $IJ \subset I \cap J$.

II. Soit A un anneau. Quels que soient $x, y \in A$, on pose

$$[x, y] = xy - yx.$$

1°) Démontrer l'identité :

$$[x, [y, z]] + [y, [z, x]] + [z, [x, y]] = 0.$$

2°) On considère des éléments x, y, h vérifiant les relations

$$[h, x] = 2x, \quad [h, y] = -2y, \quad [x, y] = h.$$

Etablir les formules $[h, x^n] = 2nx^n, \quad [h, y^n] = -2ny^n$.

3°) Soit a un élément de A . On considère l'application u de A dans A donnée par

$$u(x) = [a, x] \quad \text{pour tout } x \in A.$$

Montrer que si $a^2 = 0$, on a $u^3(x) = 0$ pour tout $x \in A$, et que si $a^3 = 0$, on a $u^5(x) = 0$ pour tout $x \in A$.

Montrer que l'on a :

$$u^n(x) = \sum_{k=0}^n (-1)^k C_n^k a^{n-k} x a^k \quad \text{pour tout } x \in A.$$

8

I. Soit $M_n(\mathbb{R})$ l'algèbre des matrices carrées d'ordre n à coefficients réels. On note $Tr(A)$ la trace de la matrice $A \in M_n(\mathbb{R})$.

1°) a) Montrer que l'application $A \mapsto Tr(A)$ est une forme linéaire sur $M_n(\mathbb{R})$ et que $Tr(AB) = Tr(BA)$ quelles que soient A et B dans $M_n(\mathbb{R})$.

b) Montrer que $Tr(P^{-1}AP) = Tr(A)$ si $A \in M_n(\mathbb{R})$ et $P \in GL(n, \mathbb{R})$. En déduire que si E est un $\mathbb{R}$ -espace vectoriel de dimension finie n , B une base de E et $u \in \mathcal{L}(E)$, la trace de la matrice de u par rapport à la base B est indépendante de cette base.

2°) Si A et B sont deux éléments de $M_n(\mathbb{R})$, on pose $[A, B] = AB - BA$.

Montrer que

a) $[A, A] = 0$ et $[B, A] = -[A, B]$.

b) $[A, [B, C]] + [B, [C, A]] + [C, [A, B]] = 0$.

c) Que vaut $Tr([A, B])$?

3°) Soit A une matrice fixée dans $M_n(\mathbb{R})$. On pose

$$f(X) = [A, X] \text{ pour toute matrice } X \in M_n(\mathbb{R}).$$

Montrer que f est un endomorphisme de $M_n(\mathbb{R})$ et que

$$f([X, Y]) = [f(X), Y] + [X, f(Y)]$$

quelles que soient X et Y dans $M_n(\mathbb{R})$. (On pourra utiliser 2°, b)).

4°) a) Montrer que l'ensemble F des matrices carrées réelles d'ordre 2 de trace nulle est un espace vectoriel de dimension 3 dont une base est formée des matrices

$$I = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}, \quad J = \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix}, \quad K = \begin{pmatrix} 0 & 0 \\ 1 & 0 \end{pmatrix}.$$

b) Montrer que si A et B sont dans F , alors $[A, B] \in F$.

c) Ecrire la matrice de l'endomorphisme f du 3°) par rapport à la base (I, J, K) .

II. Soit z un nombre complexe non nul. On pose $x = z + \frac{1}{z}$.

1°) Montrer que pour tout entier $n \geq 0$, $z^n + \frac{1}{z^n}$ est un polynôme $P_n(x)$ de degré n en x à coefficients entiers et que pour $n \geq 2$, on a la relation

$$P_n - xP_{n-1} + P_{n-2} = 0.$$

2°) On suppose que $x = 2 \cos \theta$, où θ est un nombre réel. Calculer la valeur de $P_n(X)$. En déduire que pour $-2 \leq x \leq 2$ et pour tout $n \geq 2$, on a

$$(4 - x^2)P_n'' - xP_n' + n^2P_n = 0.$$

DUES MP1, Abidjan, Juin 1987

9

I. On considère le polynôme P à coefficients dans $\mathbb{R}$ défini par :

$$P(x) = x^4 - 2x^2 \cos 2\alpha + 1, \text{ où } \alpha \in \mathbb{R}.$$

1°) Déterminer les racines de $P(x)$ dans $\mathbb{C}$. En déduire la décomposition de $P(x)$ en produits de polynômes irréductibles de $\mathbb{C}[x]$.

2°) Donner la décomposition de $P(x)$ en produits de polynômes irréductibles de $\mathbb{R}[x]$. (On pourra discuter, suivant les valeurs de α , le nombre des racines réelles de $P(x)$).

3°) On considère $F(x)$ définie par : $F(x) = \frac{x^2}{x^4 - 2x^2 \cos 2\alpha + 1}$. Donner la décomposition de $F(x)$ en éléments simples dans $\mathbb{R}$. (Plusieurs cas sont à considérer. On utilisera la question 2° /).

4°) Calculer le P.G.C.D. des polynômes P et Q , où Q est défini par :

$$Q(x) = x^8 - 2x^4 \cos 4\alpha + 1.$$

II. On considère l'équation $(E_1) : x^2 - bx + c = 0$, où x est l'inconnue, b et c étant deux éléments de $\mathbb{Z}$ vérifiant $b^2 - 4c < 0$.

α étant un élément de $\mathbb{C}$, on définit Z_α , ensemble des nombres complexes de la forme $z = p + q\alpha$, où p et q sont des éléments de $\mathbb{Z}$.

1°) Montrer que si α est racine de (E_1) alors Z_α est un sous-anneau de $\mathbb{C}$ muni des opérations usuelles et que G_α , ensemble des éléments de Z_α inversibles pour la multiplication, est un groupe multiplicatif.

2°) Si α et $\bar{\alpha}$ sont les 2 racines de (E_1) montrer que $Z_\alpha = Z_{\bar{\alpha}}$.

3°) On considère l'application f de Z_α dans $\mathbb{N}$ définie par : $\forall z \in Z_\alpha, f(z) = |z|^2$ (où $|z|$ désigne le module de z). Soit α une racine de (E_1) .

Montrer que si $z = p + q \propto$, alors $f(z) = p^2 + bpq + cq^2$. Quelle est l'image par f du groupe $G_\propto$?

En déduire que si $z = p + q \propto$ est un élément de $G_\propto$, alors on a :

$$0 \leq q^2(4c - b^2) \leq 4$$

(on pourra considérer l'équation $(E_2) : x^2 + bqx + cq^2 - 1 = 0$ qui admet p pour racine dans $\mathbb{Z}$).

4°) Quelles sont les valeurs possibles de $b^2 - 4c$? En déduire la détermination des éléments du groupe $G_\propto$ suivant les valeurs de $b^2 - 4c$.

DUES MP1, Abidjan, Septembre 1987 (extrait)

10

I. Soient G un groupe noté multiplicativement et A une partie de G . On appelle **centralisateur** de A dans G , l'ensemble $C(A)$ des $x \in G$ tels que $xa = ax$ pour tout $a \in A$. On appelle **normalisateur** de A dans G , l'ensemble $N(A)$ des $x \in G$ tels que $xAx^{-1} = A$.

1°) Montrer que $C(A)$ est un sous-groupe distingué de G et que $N(A)$ est un sous-groupe de G .

2°) Pour tout a fixé dans G , on considère l'application $f_a : G \longrightarrow G$ définie par $f_a(x) = axa^{-1}$. On note G' l'ensemble des applications de la forme f_a .

a) Montrer que G' est un groupe pour la composition des applications.

b) Montrer que $a \longmapsto f_a$ est un homomorphisme de G sur G' . Quel est son noyau ?

c) En déduire que G' est isomorphe à G/C où C est le centralisateur de G .

II. Soit A un anneau. On dit qu'un élément $x \in A$ est nilpotent s'il existe un entier $n \geq 1$ tel que $x^n = 0$.

1°) Montrer que, si A ne possède pas de diviseur de zéro, le seul élément nilpotent de A est 0.

2°) Montrer que, si deux éléments nilpotents x et y de A commutent, alors $x + y$ et xy sont nilpotents (pour $x + y$, on utilisera la formule du binôme avec un exposant convenable).

3°) Montrer que si deux éléments a et b de A commutent, on a

$$a^n - b^n = (a - b)(a^{n-1} + a^{n-2}b + \dots + ab^{n-2} + b^{n-1})$$

pour tout entier $n \geq 1$.

En déduire que si x est un élément nilpotent de A , alors $1_A - x$ est inversible dans A et déterminer $(1_A - x)^{-1}$.

DUES MP1, Abidjan, Février 1988 (extrait)

11

I. Soit (e_1, e_2, e_3) la base canonique de $\mathbb{R}^3$ et soit u l'endomorphisme de $\mathbb{R}^3$ défini par :

$$u(e_1) = -e_2 - e_3 \sin \theta$$

$$u(e_2) = e_1 + e_3 \cos \theta$$

$$u(e_3) = -e_1 \sin \theta + e_2 \cos \theta$$

où θ est un nombre réel donné.

1°) Ecrire la matrice A de u par rapport à la base (e_1, e_2, e_3) . Calculer A^3 et en déduire que $u^3 = 0$.

2°) A tout nombre réel t , on associe l'endomorphisme f_t de $\mathbb{R}^3$ défini par

$$f_t = Id + tu + \frac{t^2}{2} u^2$$

où Id désigne l'application identique de $\mathbb{R}^3$.

Montrer que pour tout couple (t, t') de nombres réels, on a $f_t \circ f_{t'} = f_{t+t'}$.

Que peut-on en déduire pour la structure de l'ensemble des endomorphismes f_t , lorsque t parcourt $\mathbb{R}$?

3°) On pose $e'_1 = e_1 \cos \theta + e_2 \sin \theta$, $e'_2 = u(e_1)$, $e'_3 = u(e_2)$.

a) Montrer que (e'_1, e'_2, e'_3) est une base de $\mathbb{R}^3$.

b) Déterminer la matrice M de u par rapport à la base (e'_1, e'_2, e'_3) .

II. On pose
$$P_n(X) = \left(1 + \frac{X}{n}\right)^n - \left(1 - \frac{X}{n}\right)^n$$

où n est un entier ≥ 0 .

Factoriser $P_n(X)$ dans $\mathbb{C}[X]$.

DUES MP1, Abidjan, Juin 1988 (extrait)

12

I. Soient $f : E \longrightarrow F$ et $g : F \longrightarrow G$ deux applications.

1°) Montrer les implications :

$$g \circ f \text{ injective} \implies f \text{ injective}$$

$$g \circ f \text{ surjective} \implies g \text{ surjective.}$$

2°) En déduire que, s'il existe deux applications s et h de F dans E telles que

$$s \circ f = id_E \quad \text{et} \quad f \circ h = id_F \quad \text{alors} \quad s = h.$$

II. Soit G l'ensemble des couples (a, b) de nombres rationnels non simultanément nuls.

On définit la loi $*$ de la façon suivante : pour $(a, b) \in G$ et $(a', b') \in G$ on pose

$$(a, b) * (a', b') = (aa' - 3bb', ab' + a'b).$$

1°) Montrer que $*$ est une loi de composition interne dans G .

2°) Etudier les propriétés de $*$ dans G . Quelle est la structure de $(G, *)$?

3°) Soit $K = \{x \in \mathbb{R} / \exists (a, b) \in \mathbb{Q} \times \mathbb{Q}, (a, b) \neq (0, 0) \text{ et } x = a + b\sqrt{3}\}$. K est muni de la multiplication des réels.

Quelle est la structure de $(K, \times)$?

4°) On considère $f : (G, *) \longrightarrow (K, \times)$ définie par

$$f((a, b)) = a + b\sqrt{3}. \quad f \text{ est-il un morphisme ?}$$

III. Soit A un anneau commutatif et soit ξ l'ensemble des idéaux de A . Pour $I \in \xi$ et $J \in \xi$, on pose :

$$I * J = \{x \in A / \forall j \in J, \quad j \cdot x \in I\}.$$

Le problème se propose d'étudier quelques propriétés de la loi $*$.

— A —

1°) Montrer : $\forall I \in \xi, \quad \forall J \in \xi, \quad I \subset I * J.$

COURS D'ALGÈBRE

2°) Pour $I \in \xi$ déterminer

a) $I * I$.

b) $I * A$.

c) $I * (0)$.

3°) Montrer que $*$ est une loi de composition interne dans ξ .

– B –

Dans toute la partie – B –, on se placera dans le cas particulier où l'anneau A est égal à $\mathbb{Z}$.

1°) Soient p et q deux nombres premiers entre eux. Montrer que :

$$p\mathbb{Z} * a\mathbb{Z} = p\mathbb{Z}.$$

2°) Soient m et n deux entiers naturels, non nuls, tels que m soit un multiple de n . Montrer que $n\mathbb{Z} * m\mathbb{Z} = A$.

3°) a) La loi $*$ est-elle commutative ?

b) La loi $*$ est-elle associative ?

– C –

Dans la partie – C –, A est un anneau commutatif quelconque.

Montrer que $\forall I \in \xi, \quad \forall J \in \xi, \quad \forall K \in \xi$

1°) $I * (J + K) = (I * J) \cap (I * K)$.

2°) $(I \cap J) * K = (I * K) \cap (J * K)$.

DUES MP1, Abidjan, Février 1989

13

I. On note $\mathbb{C}[X]$ l'algèbre des polynômes à coefficients complexes. Si

$$P = a_0 + a_1 X + \cdots + a_p X^p$$

est un élément de $\mathbb{C}[X]$ et si A est une matrice carrée d'ordre n à coefficients complexes, on pose

$$P(A) = a_0 I + a_1 A + \cdots + a_p A^p,$$

où I désigne la matrice unité d'ordre n . On dit qu'un polynôme P de $\mathbb{C}[X]$ est annulateur de la matrice A si $P(A)$ est la matrice nulle.

1°) Montrer que l'ensemble $M_n(\mathbb{C})$ des matrices carrées d'ordre n à coefficients complexes est un espace vectoriel de dimension n^2 .

Montrer que toute matrice A de $M_n(\mathbb{C})$ admet un polynôme annulateur non nul. (On pourra considérer les $n^2 + 1$ matrices $I, A, A^2, \dots, A^{n^2}$.)

2°) Soit A un élément de $M_n(\mathbb{C})$. Montrer que l'application $\varphi : \mathbb{C}[X] \longrightarrow M_n(\mathbb{C})$ définie par $\varphi(P) = P(A)$ est un homomorphisme d'algèbres dont l'image est une sous-algèbre commutative de $M_n(\mathbb{C})$ et dont le noyau est l'ensemble des polynômes annulateurs de A .

3°) On appelle polynôme minimal de la matrice $A \in M_n(\mathbb{C})$ le polynôme unitaire annulateur de A de plus bas degré. Montrer qu'il est unique.

4°) Soit A une matrice de $M_n(\mathbb{C})$ et soit P le polynôme minimal de A . Montrer que A est inversible si et seulement si $P(0) \neq 0$.

Montrer qu'un polynôme $Q \in \mathbb{C}[X]$ est annulateur de A si et seulement si P divise Q .

DUES MP1, Abidjan, Juin 1989 (extrait)

14

I. Rappeler la définition d'un anneau principal. Soit K un corps commutatif; montrer que l'anneau $K[X]$ est principal.

II. Dans tout le problème, on note $H = \mathbb{C} \times \mathbb{C}$, où $\mathbb{C}$ désigne le corps des nombres complexes. Si $h = (a, \alpha) \in H$ et $h' = (b, \beta) \in H$, on pose :

$h + h' = (a + b, \alpha + \beta)$, $h \cdot h' = (a, \alpha) \cdot (b, \beta) = (ab - \alpha \bar{\beta}, a\beta + \alpha \bar{b})$ où $\bar{z}$ désigne le nombre complexe conjugué de z .

1°) Montrer que $(H, +)$ est un groupe abélien, que $(1, 0)$ est l'élément neutre pour la multiplication et que $(H, +, \cdot)$ est un anneau. Est-il commutatif?

2°) Montrer que l'application $\Phi : \mathbb{C} \longrightarrow H$ définie par $\Phi(a) = (a, 0)$ est un homomorphisme injectif d'anneaux. En déduire que l'ensemble des éléments de H de la forme $(a, 0)$ est un corps isomorphe à $\mathbb{C}$. Dans la suite on identifiera le nombre complexe z et l'élément $(z, 0)$ de H .

Calculer $(\alpha, 0) \cdot (0, 1)$ et en déduire que tout élément (a, α) de H s'écrit de façon unique $(a, \alpha) = a + \alpha \omega$, avec $\omega = (0, 1)$.

Calculer ω^2 puis montrer que $\omega a = \bar{a} \omega$ pour tout $a \in \mathbb{C}$.

3°) Soit $Z(H)$ le centre de H . Montrer qu'un élément $h = a + \alpha \omega$ de H appartient à $Z(H)$ si et seulement si $a \in \mathbb{R}$ et $\alpha = 0$.

4°) On définit le conjugué d'un élément $h = a + \alpha \omega$ de H en posant : $\bar{h} = \bar{a} - \alpha \omega$.

Montrer que, quel que soit $(h, h') \in H^2$, on a : $\overline{h + h'} = \bar{h} + \bar{h'}$; $\overline{\bar{h}} = h$ et $\overline{h \cdot h'} = \bar{h'} \cdot \bar{h}$.

Pour $h \in H$, on pose $N(h) = h \cdot \bar{h}$. Montrer que : $N(h) = \bar{h} \cdot h = a\bar{a} + \alpha \bar{\alpha}$ et que $N(h) \geq 0$ pour tout h dans H . A quelle condition a-t-on $N(h) = 0$? En déduire que H est un corps non commutatif.

5°) En utilisant la relation : $\forall (h, h') \in H^2, \overline{h \cdot h'} = \bar{h'} \cdot \bar{h}$, montrer que $N(h \cdot h') = N(h) \cdot N(h')$. En déduire que l'ensemble G des éléments de H tels que $N(h) = 1$ est un groupe multiplicatif.

III. P est un plan affine rapporté à un repère $R = (0, \vec{i}, \vec{j})$. On rappelle que l'équation cartésienne d'une droite D de P est de la forme $D : ux + vy + h = 0$ où u et v ne sont pas tous les deux nuls.

Δ est la droite d'équation $x = 0$ et on note $P' = P - \Delta$. Dans P on définit la loi de composition interne, notée $*$ par : si M a pour coordonnées (x, y) , M' a pour coordonnées (x', y') , alors $M * M'$ est le point de P de coordonnées $(xx', xy' + y)$.

1°) a) Montrer que la loi $*$ est associative dans P .

b) Montrer qu'il existe un point E de P élément neutre pour $*$.

2°) Déterminer l'ensemble des points M de P qui admettent un symétrique pour la loi $*$. Ce symétrique, s'il existe sera noté M^{-1} . Déterminer les coordonnées de M^{-1} en fonction de celles de M .

3°) Quelle est la structure de $(P', *)$?

4°) Soit A un point de P fixé, de coordonnées (a, b) et soit H_A l'ensemble des points de P' qui commutent avec A , c'est-à-dire : $H_A = \{M \in P' / A * M = M * A\}$.

a) Déterminer l'équation cartésienne de H_A .

b) Dans quel cas a-t-on $H_A = P'$?

c) On suppose désormais que A est tel que $H_A \neq P'$. Quelle est la nature de H_A ? Montrer que $(H_A, *)$ est un sous-groupe commutatif de $(P', *)$.

Bibliographie

- (1) *Cours de mathématiques—1, algèbre*, J.M. ARNAUDIES et H. FRAYSSE, Dunod Université, Paris, 1987.
- (2) *Cours de mathématiques du premier cycle, 1^{ère} et 2^e années*, J. DIXMIER, Gauthier-Villars, Paris 1976.
- (3) *Nouveau cours de mathématiques tomes 1, 2 et 3*, A. DONEDDU, Vuibert, Paris, 1982.
- (4) *Cours d'algèbre*, R. GODEMENT, Hermann, Paris, 1973.
- (5) *Fundamental Structures of Algebra*, G.D. MOSTOW, J.H. SAMPSON, J.-P. MEYER, Mc Graw-Hill Book Company, New-York, 1963.
- (6) *Algèbre, M.P. et Spéciales AA'*, M. QUEYSANNE, Armand Colin, Paris, 1964.
- (7) *Cours d'algèbre*, J.L. ROQUE, Chr. LEBOEUF, G. CHASSARD et J. GUEGAND, Edition Marketing, Paris, 1980.

Imprimé en France par I.M.E. - 25-Baume-les-Dames
Dépôt légal n° 4655-10/1991
Collection n° 49 - Edition n° 01
59/4306/3

Universités francophones est la collection de l'Université des réseaux d'expression française (UREF). Cette dernière, qui fonctionne au sein de l'AUPELF comme une Université sans murs, a été choisie par le Sommet des Chefs d'État et de gouvernement des pays ayant en commun l'usage du français comme l'opérateur privilégié du Sommet en matière d'enseignement supérieur et de recherche.

Cette collection de manuels universitaires et d'ouvrages de référence s'adresse à tous les étudiants francophones. Elle est appelée à constituer une bibliothèque universitaire en langue française dont les ouvrages sont proposés à des prix modérés.

Cet ouvrage de base a pour but d'exposer le plus simplement possible, mais de façon rigoureuse, les principaux résultats d'algèbre générale et d'algèbre linéaire. Il s'adresse aux étudiants en mathématiques du premier cycle des Universités et aux étudiants préparant l'entrée dans les grandes écoles scientifiques. Il peut également être utile aux scientifiques qui désirent se recycler en mathématiques et à tous ceux qui veulent acquérir de bonnes connaissances de base en algèbre.

L'auteur s'est efforcé de faire un exposé qui soit assez rigoureux et assez riche pour servir de base à une solide formation mathématique.

Une série de problèmes posés aux examens de MP_1 sont placés à la fin de l'ouvrage.

★
★ ★

Saliou TOURÉ est né à Kolia (Côte d'Ivoire) en 1937. Docteur ès Sciences mathématiques, il a enseigné dans de nombreuses Universités africaines, européennes et américaines. Depuis 1976, il est professeur titulaire de mathématiques à la Faculté des Sciences et Techniques de l'Université d'Abidjan. Il est également Directeur de l'Institut de Recherches Mathématiques d'Abidjan (IRMA) depuis 1975, Président de la Société Mathématique de Côte-d'Ivoire depuis 1977, Secrétaire Général de l'Union Mathématique Africaine depuis 1986, Vice-Président de l'Union Panafricaine de la Science et de la Technologie depuis 1990. Son domaine de recherche est l'analyse harmonique sur les groupes de Lie.

Prix France : 100 FF • Prix préférentiel UREF (Afrique, Asie, Amérique latine, Moyen-Orient, Haïti) : 50 FF



I.S.S.N. 0993-3948
Diffusion EDICEF ou ELLIPSES selon pays
Distribution Canada D.P.L.U.

59/4306/3
Imprimé en France
S.S.Q.I. - PARIS