

République Algérienne Démocratique et Populaire

Ministère de l'Enseignement Supérieur et
de la Recherche Scientifique



Centre Universitaire Abd Elhafid Boussouf Mila

Institut des Mathématiques et Informatique

Département d'Informatique

Cours : Algèbre 1

Par : Dr. Kecies Mohamed

Life is good only for two things, discovering mathematics and teaching mathematics.

SIMÉON DENIS POISSON

Année Universitaire 2024/2025

Table des matières

Introduction Générale	2
1 Notions de logique	4
1.1 Assertion et table de vérité	4
1.2 Quantificateurs mathématiques	5
1.2.1 Quantificateur universel	5
1.2.2 Quantificateur existentiel	6
1.3 Assertions quantifiées et leurs négations	7
1.4 Connecteurs logiques usuels	10
1.4.1 Négation	10
1.4.2 Conjonction	11
1.4.3 Disjonction	11
1.4.4 Implication logique	12
1.4.5 Equivalence logique	13
1.5 Règles logiques	15
1.6 Différents types de raisonnement mathématiques	17
1.6.1 Raisonnement direct	17
1.6.2 Raisonnement par Contraposée	18
1.6.3 Raisonnement par l’Absurde	19
1.6.4 Raisonnement par Disjonction des cas	21
1.6.5 Raisonnement par Contre-exemple	22
1.6.6 Raisonnement par Récurrence	22
2 Ensembles et Relations binaires sur un ensemble	25
2.1 Théorie des ensembles	25
2.1.1 Définitions et notations	25
2.1.2 Relations entre ensembles	28
2.1.2.1 Inclusion d’ensembles	28
2.1.2.2 Égalité d’ensembles	29
2.1.3 Opérations sur les ensembles	30
2.1.3.1 Intersection	31
2.1.3.2 Réunion	31
2.1.3.3 Différence ensembliste	32
2.1.3.4 Différence symétrique	33
2.1.3.5 Complémentaire	34
2.1.4 Propriétés des opérations élémentaires	35
2.1.5 Union et intersection d’une famille de sous-ensembles	37
2.1.6 Partition d’un ensemble	38
2.1.7 Produit cartésien	39
2.1.7.1 Couples et n -uplets	40
2.2 Relations Binaires	42

2.2.1	Généralités	43
2.2.2	Représentation d'une relation binaire	44
2.2.3	Propriétés d'une relation binaire	49
2.2.4	Relations d'équivalence	51
2.2.4.1	Classe d'équivalence et ensemble quotient	52
2.2.5	Relations d'ordre	56
2.2.5.1	Comparaison dans un ensemble ordonné, ordre total et partiel	57
2.2.5.2	Éléments remarquables d'un ensemble ordonné	59
2.2.5.2.0.1	Minorants, majorants	59
2.2.5.2.0.2	Parties majorées, minorées, bornées	60
2.2.5.2.0.3	Plus grand élément, plus petit élément	60
2.2.5.2.0.4	Borne supérieure, Borne inférieure	61
3	Fonctions et Applications	64
3.1	Fonctions	64
3.2	Applications	66
3.2.1	Applications particulières	69
3.2.2	Partie stable par une application	72
3.2.3	Prolongements et restrictions	72
3.2.3.1	Restriction d'une application	73
3.2.3.2	Prolongement d'une application	73
3.2.4	Composition des applications	74
3.2.5	Image directe et image réciproque	77
3.2.5.1	Propriétés de l'image directe et réciproque	80
3.2.6	Types d'applications	84
3.2.6.1	Applications injectives	84
3.2.6.2	Applications surjectives	88
3.2.6.3	Applications bijectives	91
3.2.7	Application réciproque d'une application bijective	94
3.2.8	Injection, surjection, bijection sur des ensembles finis	98
3.2.9	Ordre et applications	99
4	Structures algébriques	101
4.1	Lois de composition interne	101
4.1.1	Partie stable pour une loi interne	102
4.1.2	Propriétés d'une loi de composition interne	103
4.1.2.1	Commutativité	103
4.1.2.2	Associativité	103
4.1.2.3	Distributivité	104
4.1.3	Éléments particuliers	104
4.1.3.1	Élément neutre à gauche, neutre à droite, élément neutre	104
4.1.3.2	Symétrique à gauche, symétrique à droite, symétrique	106
4.2	Demi-groupe et Monoïde	107
4.3	Groupes	108
4.3.1	Généralités sur les groupes	108
4.3.2	Puissances entières d'un élément dans un groupe	109
4.3.3	Ordre d'un Groupe	110
4.3.4	Sous-groupes	111
4.4	Groupe $\mathbb{Z}/n\mathbb{Z}$	114
4.5	Généralités sur le groupe symétrique $\mathcal{S}_n$	116
4.5.1	Rappels et notations	116

4.5.2	Transpositions et cycles	121
4.6	Morphismes de Groupes	124
4.6.1	Définitions et exemples	124
4.6.2	Image et noyau d'un morphisme de groupes	127
4.6.3	Image directe et réciproque de sous-groupes par un morphisme	129
4.7	Anneaux	131
4.7.1	Définitions et exemples	131
4.7.2	Règles de calculs dans un anneau	134
4.7.3	Sous-anneaux d'un anneau	137
4.7.4	Diviseurs de zéro et Intégrité	139
4.7.5	Morphismes d'anneaux	140
4.7.5.1	Propriétés des morphismes d'anneaux	142
4.7.6	Idéaux	143
4.7.6.1	Notions d'idéal à gauche, à droite, bilatère	143
4.7.7	Opérations sur les idéaux d'un anneau	146
4.7.7.1	Somme d'idéaux	146
4.7.7.2	Intersection d'idéaux	147
4.7.7.3	Produit d'idéaux	148
4.7.7.4	Propriétés des opérations sur les idéaux	150
4.8	Corps	151
4.8.1	Sous-corps	152
5	Anneau des polynômes	155
5.1	Polynôme à une indéterminée à coefficients dans $\mathbb{K}$	155
5.2	Opérations sur les polynômes et structures	158
5.2.1	Addition de deux polynômes	158
5.2.1.1	Structure de groupe commutatif sur $\mathbb{K}[X]$	159
5.2.2	Produit de deux polynômes	160
5.2.2.1	Structure d'anneau sur $\mathbb{K}[X]$	162
5.2.3	Multiplication d'un polynôme par un scalaire	164
5.3	Degré d'un polynôme et Intégrité de $\mathbb{K}[X]$	166
5.3.1	Propriétés fondamentales des degrés des polynômes	167
5.4	Composition des polynômes	169
5.5	Dérivation des polynômes	172
5.6	Fonctions polynomiales	176
5.7	Arithmétique dans $\mathbb{K}[X]$	177
5.7.1	Divisibilité dans l'anneau $\mathbb{K}[X]$	177
5.7.2	Division euclidienne dans $\mathbb{K}[X]$	180
5.7.3	Division suivant les puissances croissantes	183
5.7.4	Plus grand commun diviseur	184
5.7.4.1	Calcul du pgcd (Algorithme d'Euclide)	185
5.7.5	Plus petit commun multiple	188
5.7.6	Polynômes premiers entre eux	190
5.8	Racines des polynômes	191
5.8.1	Généralités	191
5.8.2	Ordre de multiplicité des racines d'un polynôme	193
5.8.3	Utilisation des dérivées successives pour le calcul d'une multiplicité	196
5.8.4	Polynôme conjugué et racines complexes	198
5.9	Polynômes Irréductibles et Factorisation	202
5.9.1	Polynômes Irréductibles	202
5.9.2	Factorisation dans $\mathbb{C}[X]$	204

5.9.3	Factorisation dans $\mathbb{R}[X]$	205
Bibliographie		206

Table des figures

2.1	Intersection de deux ensembles	31
2.2	Réunion de deux ensembles	32
2.3	Différence entre deux ensembles	33
2.4	Différence symétrique de deux ensembles	34
2.5	Complémentaire d'un ensemble	35
2.6	Représentation au moyen d'un tableau	45
2.7	Représentation au moyen d'une matrice binaire	46
2.8	Représentation au moyen d'un diagramme sagittal	46
2.9	Représentation au moyen d'un diagramme cartésien	46
2.10	Représentation au moyen d'un graphe	47
2.11	Composée de deux relations Binaires	48
2.12	Représentation graphique d'une relation d'équivalence. Partition en classes d'équivalence.	56
3.1	Fonction, non Fonction	65
3.2	Application, non Application	67
3.3	Image directe, Image réciproque	79
3.4	Application injective, non injective	86
3.5	Application surjective, non surjective	89
3.6	Application bijective, non bijective	92

RÉSUMÉ

Ce manuscrit est issu du cours **d'Algèbre 1** destiné aux étudiants de première année L.M.D en mathématiques et informatique. Il couvre un large éventail de sujets fondamentaux, notamment la logique mathématique, la théorie des ensembles, les relations binaires, les applications, et les structures algébriques (groupes, anneaux, corps). La dernière section est consacrée à l'anneau des polynômes, avec des explications approfondies sur les opérations, la factorisation, et les propriétés des polynômes. Le contenu est structuré de manière à introduire progressivement les concepts de base, tout en fournissant des détails et des démonstrations mathématiques importantes pour les étudiants débutants dans ces domaines. Chaque section est riche en explications théoriques et en exemples pratiques, ce qui permet aux étudiants d'acquérir une compréhension solide des fondements de la logique et de l'algèbre.

ABSTRACT

This manuscript is derived from the **Algebra 1** course intended for first-year L.M.D. students in mathematics and computer science. It covers a wide range of fundamental topics, including mathematical logic, set theory, binary relations, functions, and algebraic structures (groups, rings, fields). The final section is devoted to the ring of polynomials, with detailed explanations on operations, factorization, and polynomial properties. The content is structured to gradually introduce basic concepts, while providing important mathematical details and proofs for students new to these fields. Each section is rich in theoretical explanations and practical examples, enabling students to gain a solid understanding of the foundations of logic and algebra."

Introduction Générale

L'algèbre, branche fondamentale des mathématiques, explore les structures abstraites et les opérations associées. Elle offre des outils puissants pour résoudre des problèmes mathématiques, modéliser des phénomènes du monde réel, et développer des théories complexes. L'algèbre utilise des symboles, souvent des lettres, pour représenter des nombres et des relations générales, facilitant ainsi la résolution d'équations et l'analyse des relations mathématiques de manière plus abstraite. Ses applications s'étendent à de nombreux domaines, notamment la science, la technologie et l'ingénierie.

Historiquement, le terme «algèbre» provient du mot arabe al-jabr, qui signifie «restauration» ou «réunion des parties brisées». Ce terme a été employé pour la première fois par le mathématicien perse Al-Khawarizmi dans son ouvrage Kitab al-Jabr wa-l-Muqabala au 9^e siècle, souvent traduit par «Abrégé de calcul par la restauration et la comparaison». Cet ouvrage est considéré comme l'une des premières contributions majeures à l'algèbre. De nombreux historiens des sciences considèrent Muhammad Ibn Musa al-Khwarizmi comme le «père de l'algèbre».

Aujourd'hui, l'algèbre se divise en plusieurs branches, chacune abordant des aspects et des niveaux d'abstraction différents : l'algèbre élémentaire, l'algèbre linéaire, et l'algèbre abstraite.

Ce manuscrit s'inscrit dans le cadre du module d'algèbre générale, qui constitue une introduction approfondie aux concepts fondamentaux de l'algèbre. Il est issu de l'enseignement du module **Algèbre 1**, principalement destiné aux étudiants de première année du système L.M.D en mathématiques et informatique, ainsi qu'à toute personne désireuse d'acquérir des bases solides en algèbre. Ce document aborde les fondements de la logique mathématique, la théorie des ensembles, les fonctions, les structures algébriques, et les polynômes, fournissant ainsi une base essentielle pour des études ultérieures en mathématiques et en informatique.

L'objectif de ce document est d'offrir aux étudiants une compréhension claire et structurée des notions élémentaires d'algèbre générale, afin de les préparer à des études plus avancées dans le domaine. La structure des chapitres suit rigoureusement le programme récemment révisé, tel qu'établi dans le canevas de formation. Le contenu est organisé autour de cinq chapitres principaux :

Le document s'ouvre sur une introduction à la logique mathématique, qui traite du raisonnement formel. Il présente les concepts fondamentaux, tels que les assertions, qui sont des propositions pouvant être vraies ou fausses, ainsi que les tables de vérité, utilisées pour analyser les combinaisons de valeurs logiques. Les quantificateurs mathématiques, comme le quantificateur universel et le quantificateur existentiel, sont ensuite abordés afin d'exprimer des affirmations généralisées ou particulières. La section se penche également sur les connecteurs logiques : négation, conjonction, implication, équivalence qui permettent de construire des propositions complexes à partir de propositions simples. Différents types de raisonnements sont explorés, notamment le raisonnement direct, par contraposée, par l'absurde, et par récurrence, chacun étant illustré par des exemples concrets.

Le deuxième chapitre est consacré à la théorie des ensembles et aux relations binaires. Il introduit les opérations de base sur les ensembles, telles que l'union, l'intersection, le complémentaire, la différence et la différence symétrique, ainsi que des notions plus avancées comme les familles de sous-ensembles

et les partitions d'ensembles. Les relations binaires sont également abordées, avec un accent particulier sur les relations d'équivalence, qui permettent de partitionner un ensemble en classes d'équivalence, et les relations d'ordre, avec des concepts clés comme les bornes supérieures et inférieures, essentiels en mathématiques.

Le troisième chapitre se concentre sur les fonctions et applications. Après avoir défini ce qu'est une fonction, le document décrit les différentes catégories d'applications, notamment les applications injectives, surjectives, et bijectives. Il explore également des notions telles que l'image directe et image réciproque, ainsi que la composition des applications. Les concepts de prolongement et de restriction d'une application sont aussi abordés, montrant comment ajuster des applications à des domaines spécifiques, un aspect fondamental dans l'étude des fonctions en mathématiques.

Le document poursuit avec un chapitre consacré aux structures algébriques, un domaine clé de l'algèbre abstraite. Il introduit les définitions de groupes, anneaux et corps, ainsi que les propriétés fondamentales de ces structures, telles que la commutativité, l'associativité, et la distributivité. Le concept de morphisme, qui établit des liens entre différentes structures algébriques, est également traité, avec un accent particulier sur les notions d'image et de noyau d'un morphisme. Des concepts comme les sous-groupes, sous-anneaux, et idéaux sont également discutés pour montrer comment des sous-structures peuvent émerger à partir de structures algébriques plus grandes.

Enfin, le dernier chapitre porte sur les polynômes et l'anneau des polynômes. Il commence par exposer les opérations de base sur les polynômes, telles que l'addition, la multiplication, et la division euclidienne. Le document aborde ensuite les notions de degré d'un polynôme, de racines, et de factorisation. Les polynômes irréductibles, qui ne peuvent pas être décomposés davantage, sont étudiés en profondeur, ainsi que les méthodes de calcul du plus grand commun diviseur (pgcd) et du plus petit commun multiple (ppcm) dans un anneau de polynômes.

Nous espérons que ce polycopié répondra aux attentes des étudiants et qu'il les aidera à réussir dans leurs études.

Avertissement : Pour toute remarque ou suggestion visant à améliorer la rédaction de ce cours, n'hésitez pas à me contacter à l'adresse e-mail suivante : m.kecies@centre-univ-mila.dz.

Chapitre 1

Notions de logique

La logique mathématique est une discipline fondamentale des mathématiques, apparue à la fin du 19^e et au début du 20^e siècle. Elle repose sur plusieurs concepts clés tels que les assertions ou propositions, qui sont des énoncés pouvant être vrais ou faux, les connecteurs logiques comme "et", "ou", "non", "implication", les quantificateurs universel et existentiel, ainsi que la théorie de la démonstration. Ces notions constituent les bases de la logique mathématique et permettent de formaliser et d'analyser les raisonnements en mathématiques et en informatique, tout en assurant la validité des démonstrations.

Ce chapitre aborde les fondements de la logique, en incluant les assertions, les tables de vérité et les quantificateurs universels et existentiels. Il explore également les connecteurs logiques tels que la négation, la conjonction, l'implication et l'équivalence, ainsi que divers types de raisonnements mathématiques, notamment le raisonnement direct, par contraposée, par l'absurde, entre autres.

Les notions présentées dans ce chapitre sont largement utilisées en mathématiques et en informatique, constituant une base solide pour les développements à venir dans le reste de l'ouvrage.

1.1 Assertion et table de vérité

Définition 1.1.1 Une assertion (ou bien proposition) en mathématiques est une phrase P qui est soit vraie, soit fausse mais pas les deux simultanément. On utilise en général les lettres P, Q, R etc. pour désigner les assertions.

Définition 1.1.2 Une table de vérité est un tableau qui indique si une assertion P , construite à partir des assertions $P_1, P_2, P_3, \dots$, est vraie ou fausse suivant les valeurs de vérités de $P_1, P_2, P_3, \dots$

- Une table de vérité associe à l'assertion P les deux possibilités vrai (1) ou faux (0)

P	1	0
-----	---	---

Table de vérité d'une assertion

- Une table de vérité a une colonne pour chaque assertion donnée (par exemple, P et Q), et une dernière colonne affichant tous les résultats possibles de l'opération logique qui nous intéresse (par exemple, $P \wedge Q$).

Remarque 1.1.1 (*Intérêts principaux des tables de vérité*)

1. Les tables de vérité sont des outils fondamentaux en logique mathématique. Elles permettent de

représenter toutes les combinaisons possibles des valeurs de vérité (vrai ou faux) pour des assertions logiques et de déterminer la valeur de vérité d'une expression logique en fonction de ces propositions.

2. Permettre de comparer facilement deux assertions, et parfois comparer qu'elles sont vraies en même temps, et fausses en même temps.

3. Les tables de vérité sont aussi largement utilisées en informatique, notamment dans la conception de circuits logiques (AND, OR, NOT, etc.), où chaque composant de circuit peut être représenté par une table de vérité qui montre comment il se comporte selon les entrées.

Exemple 1.1.1

1. « π est un nombre entier » est une assertion fausse.
2. « 18 est divisible par 3 » est une assertion vraie.
3. « $3 > 1$ » est une assertion vraie.
4. « $\sqrt{2} > 2$ » est une assertion fausse.
5. « Pour tout réel x , on a $x^3 - 1 = (x - 1)(x^2 + x + 1)$ » est une assertion vraie.
6. « L'entier 2 est plus grand que l'entier 1 » est une assertion vraie.
7. « Il existe un réel x tel que $x^2 < 0$ » est une assertion fausse.

Remarque 1.1.2 Une assertion P peut dépendre d'un paramètre x , par exemple « $x^2 \geq 1$ », l'assertion $P(x)$ est vraie ou fausse selon la valeur de x .

1.2 Quantificateurs mathématiques

Les quantificateurs mathématiques sont des symboles utilisés pour exprimer des propositions concernant un ensemble d'éléments. Il en existe principalement deux types : Quantificateur universel et existentiel. Ces quantificateurs sont essentiels en logique et dans les démonstrations mathématiques.

1.2.1 Quantificateur universel

Le quantificateur universel est un symbole utilisé en logique et en mathématiques pour indiquer qu'une assertion est vraie pour tous les éléments d'un ensemble. Le symbole utilisé pour le quantificateur universel est $\forall$. Ainsi un quantificateur permet de préciser le domaine de validité d'une assertion.

Définition 1.2.1 Le symbole $\forall$ qui signifie « **quel que soit** » ou « **pour tout** » représente le quantificateur universel. Ce symbole représente la lettre « A » renversée qui est l'initiale du mot anglais « All ».

Exemple 1.2.1

1. Soit l'assertion

$$\forall x \in \mathbb{R} : x^2 \geq 0.$$

Cela signifie, pour tout x dans l'ensemble des nombres réels, x^2 est positif ou nul.

2. Soit l'assertion

$$\forall p \in \mathbb{P} : p > 1.$$

Cela signifie que pour tout p appartenant à l'ensemble des nombres premiers $\mathbb{P}$, p est strictement supérieur à 1.

1.2.2 Quantificateur existentiel

Le quantificateur existentiel est un symbole utilisé en logique et en mathématiques pour indiquer qu'une assertion est vraie pour au moins un élément d'un ensemble. Le symbole pour le quantificateur existentiel est $\exists$.

Définition 1.2.2 *Le symbole $\exists$ qui signifie «**il existe au moins un ... tel que** » représente le quantificateur existentiel. Ce symbole représente la lettre «E» renversée qui est l'initiale du mot anglais «**exist** ».*

Exemple 1.2.2

1. Soit l'assertion

$$\exists x \in \mathbb{Z} : x \text{ est pair.}$$

Cela signifie qu'il existe au moins un nombre entier x qui est pair.

2. Soit l'assertion

$$\exists x \in \mathbb{R} : x^2 - 4 = 0..$$

Cela signifie qu'il existe un nombre réel x tel que $x^2 - 4 = 0$.

Notation 1.2.1 (Quantificateur d'existence et d'unicité)

1. On rencontre aussi parfois l'écriture $\exists!$ pour signifier l'existence et l'unicité. On a alors $\exists!$ qui signifie «**il existe un unique ... tel que**», ou «**il existe un et un seul ... tel que**». La notation $\exists!$ s'appelle le quantificateur d'existence et d'unicité qui est utilisé en mathématiques et en logique pour indiquer qu'il existe un et un seul élément dans un ensemble qui satisfait une certaine propriété.
2. Il convient également de signaler le quantificateur «d'existence et d'unicité», d'usage moins courant que les deux précédents.

Exemple 1.2.3 Soit l'assertion

$$\exists! x \in [0, 1] : x^2 + 4x + 1 = 0.$$

Signifie qu'il existe un unique x appartenant à l'intervalle $[0, 1]$ tel que $x^2 + 4x + 1 = 0$.

Remarque 1.2.1 *L'ordre des quantificateurs est très important. Autrement dit l'ordre dans lequel on écrit les quantificateurs change la signification.*

1. Par exemple les deux phrases logiques

$$\forall x \in \mathbb{R}, \exists y \in \mathbb{R} : x + y > 0 \text{ et } \exists y \in \mathbb{R}, \forall x \in \mathbb{R} : x + y > 0,$$

sont différentes. La première est vraie, la seconde est fausse. En effet une phrase logique se lit de gauche à droite, ainsi la première phrase affirme « Pour tout réel x , il existe un réel y (qui peut donc dépendre de x) tel que $x + y > 0$. » (par exemple on peut prendre $y = |x| + 1$). C'est donc une phrase vraie. Par contre la deuxième se lit « Il existe un réel y , tel que pour tout réel x , $x + y > 0$ ». Cette phrase est fausse, cela ne peut pas être le même y qui convient pour tous les x .

2. Soit l'assertion

$$\forall x \in \mathbb{R}, \exists y \in \mathbb{R} : y > x,$$

se lit « Quel que soit le réel x , il existe au moins un réel y tel que y soit supérieur à x ». On peut toujours trouver un nombre supérieur à un nombre réel donné car l'ensemble $\mathbb{R}$ n'est pas borné.

L'assertion est vraie.

Inversons maintenant les quantificateurs

$$\exists x \in \mathbb{R}, \forall y \in \mathbb{R} : y > x,$$

se lit « Il existe au moins un réel x tel que pour tout réel y , y soit supérieur à x ». Cette assertion cette fois est fausse car on ne peut trouver un réel inférieur à tous les autres. En effet l'ensemble $\mathbb{R}$ n'a pas de borne inférieure.

3. On retrouve la même différence dans les phrases en français suivantes. Voici une phrase vraie « Pour toute personne, il existe un numéro de téléphone », bien sûr le numéro dépend de la personne. Par contre cette phrase est fausse « Il existe un numéro, pour toutes les personnes ». Ce serait le même numéro pour tout le monde.

1.3 Assertions quantifiées et leurs négations

Les assertions quantifiées sont des propositions logiques dans lesquelles on utilise des quantificateurs : universel et existentiel pour exprimer des affirmations sur des ensembles d'éléments. Ces quantificateurs permettent de formuler des énoncés concernant tous les éléments d'un ensemble (quantificateur universel) ou l'existence d'au moins un élément qui satisfait une certaine propriété (quantificateur existentiel). On peut alors construire deux nouvelles assertions.

Définition 1.3.1 (Assertions quantifiées)

Soient E un ensemble et $P(x)$ une proposition dépendant d'une variable x de E . Alors .

1. Assertion quantifiée universelle : Le quantificateur universel $\forall$, permet de former l'assertion

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

qui est vraie lorsque $P(x)$ est vraie pour tous les éléments x de E , et qui est fausse si $P(x)$ est fausse pour au moins un élément x de E .

• *L'assertion « $\forall x \in E : P(x)$ » se lit «pour tout x dans E la proposition $P(x)$ est vérifiée», ou «tout élément x de E vérifie P » ou «quel que soit x dans E l'assertion $P(x)$ est vérifiée».*

2. Assertion quantifiée existentielle : Le quantificateur existentiel, $\exists$, permet de former l'assertion

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

qui est vraie lorsque $P(x)$ est vraie pour au moins un élément x de E , et qui est fausse si $P(x)$ est fausse pour tous les éléments de E .

• *L'assertion « $\exists x \in E : P(x)$ » se lit «il existe au moins x dans E tel que la proposition $P(x)$ est vérifiée», ou «il existe au moins x dans E qui vérifie P ».*

3. *On définit aussi l'assertion quantifiée « $\exists ! x \in E : P(x)$ » qui est vraie lorsque $P(x)$ est vraie pour un seul élément x de E .*

Remarque 1.3.1

1. *Remplacer une variable quantifiée par une autre (indépendante des autres variables intervenant dans la formule) ne change pas le sens de l'assertion. Ainsi « $\forall x, P(x)$ », et « $\forall y, P(y)$ » sont équivalentes.*

2. *Par convention, l'assertion quantifiée « $\forall x \in \phi : P(x)$ » est toujours vraie (il n'y a rien à vérifier puisque l'ensemble vide n'a pas d'éléments) et l'assertion quantifiée « $\exists x \in \phi : P(x)$ » est toujours fausse (il n'existe aucun élément dans l'ensemble vide).*

Exemple 1.3.1 *Par exemple*

1.
 - $\langle \forall x \in [1, +\infty[: x^2 \geq 1 \rangle$ est une assertion vraie.
 - $\langle \forall x \in \mathbb{R} : x^2 \geq 1 \rangle$ est une assertion fausse.
 - $\langle \forall n \in \mathbb{N} : n(n+1) \text{ est divisible par } 2 \rangle$ est une assertion vraie.
2.
 - $\langle \exists x \in \mathbb{R} : x(x-1) < 0 \rangle$ est une assertion vraie (par exemple $x = \frac{1}{2}$ vérifie bien la propriété).
 - $\langle \exists n \in \mathbb{N} : n^2 - n > n \rangle$ est une assertion vraie (il y a plein de choix, par exemple $n = 3$).
 - $\langle \exists x \in \mathbb{R} : x^2 = -1 \rangle$ est une assertion fausse (aucun réel au carré ne donnera un nombre négatif).

Définition 1.3.2 (*Négation des Assertions quantifiées*)

1. **Négation d'une assertion universelle.** Une assertion universelle s'énonce « Pour tout élément x d'un ensemble E , x possède la propriété P ». Sa négation sera « il existe au moins un élément x de l'ensemble E qui ne possède pas la propriété P ». Autrement dit, la négation de

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

est

$$\exists x \in E : \text{non}(P(x)).$$

2. **Négation d'une assertion existentielle.** Une assertion existentielle s'énonce « Il existe au moins un élément x de l'ensemble E qui possède la propriété P ». Sa négation sera « Pour tout élément x de l'ensemble E , x ne vérifie pas P ». Autrement dit, la négation de

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

est

$$\forall x \in E : \text{non}(P(x)).$$

Remarque 1.3.2 Quand on nie une assertion contenant des quantificateurs, le quantificateur universel $\forall$ devient existentiel $\exists$, et le quantificateur existentiel $\exists$ devient universel $\forall$. De plus, la propriété contenue dans l'assertion est niée elle aussi.

Exemple 1.3.2

1. La négation de

$$\forall x \in [1, +\infty[: x^2 \geq 1,$$

est l'assertion

$$\exists x \in [1, +\infty[: x^2 < 1.$$

En effet, la négation de $(x^2 \geq 1)$ est $\text{non}(x^2 \geq 1)$ mais s'écrit plus simplement $x^2 < 1$.

2. La négation de

$$\exists x \in \mathbb{C} : x^2 + x + 1 = 0,$$

est

$$\forall x \in \mathbb{C} : x^2 + x + 1 \neq 0.$$

3. La négation de

$$\forall x \in \mathbb{R} : x + 1 \in \mathbb{Z},$$

est

$$\exists x \in \mathbb{R} : x + 1 \notin \mathbb{Z}.$$

Dans le cas général d'une assertion qui renferme un certain nombre de quantificateurs, on a le résultat suivant.

Corollaire 1.3.1 *La négation d'une assertion contenant un certain nombre de fois les quantificateurs $\forall, \exists$, et ensuite l'énoncé d'une propriété P , s'obtient en remplaçant chaque quantificateur $\forall$ par le quantificateur $\exists$ et vice versa, et la propriété P par sa négation $\text{non}(P)$. Autrement dit*

- Remplacer chaque quantificateur universel $\forall$ par le quantificateur existentiel $\exists$.
- Remplacer chaque quantificateur existentiel $\exists$ par le quantificateur universel $\forall$.
- Nier la propriété ou la proposition P , c'est-à-dire remplacer P par sa négation $\text{non}(P)$.

Exemple 1.3.3 *Soit une assertion $P(x, y)$ dépendant de deux variables, alors*

1. *La négation de*

$$\forall x \in E, \exists y \in F : P(x, y), \quad (1.1)$$

est

$$\exists x \in E, \forall y \in F : \text{non}(P(x, y)).$$

2. *La négation de*

$$\exists x \in E, \forall y \in F : P(x, y), \quad (1.2)$$

est

$$\forall x \in E, \exists y \in F : \text{non}(P(x, y)).$$

3. *Soit l'assertion avec plusieurs quantificateurs*

$$\forall x \in \mathbb{R}, \exists y \in \mathbb{R} : y > x,$$

qui signifie "pour tout réel x , il existe un réel y tel que y est supérieur à x ". La négation de cette assertion est

$$\exists x \in \mathbb{R}, \forall y \in \mathbb{R} : y \leq x,$$

Cela se lit "il existe un réel x tel que pour tout réel y , y est inférieur ou égal à x ", ce qui signifie que x est un maximum global, ce qui est faux dans $\mathbb{R}$.

Remarque 1.3.3

1. *L'assertion (1.1) signifie que pour tout x il existe une valeur y (qui dépend a priori de x) telle que $P(x, y)$ est vérifiée, alors que l'assertion (1.2) signifie qu'il existe au moins un élément x dans l'ensemble E tel que, pour tous les éléments y de l'ensemble F , la propriété $P(x, y)$ est vraie. En d'autres termes, on peut trouver un x spécifique pour lequel P est vrai pour n'importe quel y dans F .*

2. *Si l'on utilise deux fois le même quantificateur, l'ordre n'a pas d'importance. On peut permuter les quantificateurs de même nature dans des écritures du type*

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

$$\exists x \in E, \exists y \in E : P(x, y).$$

Mais si les quantificateurs sont de natures différentes, leur ordre est important.

1.4 Connecteurs logiques usuels

Les connecteurs logiques sont des opérations qui relient des assertions afin d'en former de nouvelles. On distingue les connecteurs unaires, qui s'appliquent à une seule assertion, et les connecteurs binaires, qui relient deux assertions. Leur utilisation vise à établir des relations entre plusieurs assertions pour en obtenir une nouvelle. Parmi ces symboles, le premier est un connecteur unaire. Cela signifie qu'il occupe une seule place et permet de former une assertion à partir d'une unique assertion. Par exemple, avec P , ce connecteur permet de construire l'assertion $\bar{P}$. En revanche, les quatre autres connecteurs sont binaires, nécessitant deux assertions pour en produire une troisième.

Nous avons donc les cinq symboles suivants : $\neg$, $\wedge$, $\vee$, $\implies$ et $\iff$. Voici leur désignation :

- $\neg$, est le symbole de négation, que l'on prononce "non".
- $\wedge$ est le symbole de conjonction, que l'on prononce "et".
- $\vee$ est le symbole de disjonction, que l'on prononce "ou".
- $\implies$ est le symbole d'implication, que l'on prononce "implique".
- $\iff$ est le symbole de double implication, ou d'équivalence, que l'on prononce "si et seulement si" ou "équivalent à".

1.4.1 Négation

Le connecteur de négation est un opérateur logique unaire fondamental qui inverse la valeur de vérité d'une assertion.

Définition 1.4.1 (*Négation*)

La négation de l'assertion P est l'assertion notée $\text{non}(P)$ ou $\bar{P}$ qui est vraie lorsque P est fausse et fausse lorsque P est vraie.

On représente les valeurs de vérités de $\bar{P}$ en fonction de celles de P dans une table de vérité, qui est un tableau de la forme suivante

P	1	0
$\bar{P}$	0	1

Table de vérité de l'opérateur logique de négation.

Remarque 1.4.1

1. L'assertion $\text{non}(P)$ exprime le contraire de l'assertion P .
2. Il résulte de cette définition que $\text{non}(\text{non}(P))$ et P ont la même valeur logique, c'est-à-dire sont vraies simultanément ou fausses simultanément.
3. Lorsque l'assertion P s'exprime par un signe spécifique, l'assertion $\text{non}(P)$ se note par ce même signe barré par le signe $/$.
4. En général, une double négation vient souvent renforcer la négation tel que : voulez-vous sortir ? non, non. En mathématique, une double négation est considérée comme une affirmation.

Exemple 1.4.1

1. Soit P l'assertion « $x \in E$ » qui se lit x appartient à E . La négation de P peut se noter aussi bien par « $x \notin E$ » qui se lit x n'appartient pas à E , ou $\text{non}(x \in E)$. De la même façon, la négation de

« $A \subset B$ » qui se lit A est inclus dans B , peut s'écrire $\text{non}(A \subset B)$ ou $\overline{(A \subset B)}$, alors que la négation de $(A = B)$ qui se lit A est égal à B notée $(A \neq B)$ (se lit A est différent de B) ou $\text{non}(A = B)$ ou $\overline{(A = B)}$.

2. Considérons l'assertion P définie par «24 est un multiple de 2». Il s'agit d'une assertion vraie. Sa négation est $\text{non}(P)$ définie par «24 n'est pas un multiple de 2». Il s'agit donc d'une assertion fausse.

3. L'assertion « $\forall x \in \mathbb{R} : x > 0$ » est fausse et sa négation « $\exists x \in \mathbb{R} : x \leq 0$ » est vraie.

1.4.2 Conjonction

Après avoir défini un connecteur appliqué à une seule assertion dans le paragraphe précédent, nous allons maintenant examiner un connecteur qui relie deux assertions : la conjonction logique, exprimée par "et". Ce connecteur, également appelé opérateur logique binaire, relie deux assertions et exprime l'idée que les deux doivent être vraies simultanément.

Définition 1.4.2 (*Conjonction*)

Soient P et Q deux assertions, on appelle conjonction de P et Q , l'assertion notée $(P \wedge Q)$ ou encore $(P \text{ et } Q)$, qui est vraie lorsque les deux assertions P et Q sont vraies simultanément, et qui est fausse dans tous les autres cas.

- La table de vérité comporte quatre cas distincts puisque les deux assertions P et Q peuvent être vraies ou fausses de façon indépendante. On résume ces quatre cas de figure dans la table de vérité ci-dessous.

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

Table de vérité de l'opérateur logique de conjonction.

- Il est évident que l'assertion $(P \text{ et } Q)$ d'une part et l'assertion $(Q \text{ et } P)$ ont même valeur de vérité.

Exemple 1.4.2

1. Soient les assertions P : « k est un multiple de 2 » et Q : « k est un multiple de 3 », alors l'assertion $(P \wedge Q)$ est « k est un multiple de 6».

2. L'assertion « π n'est pas un entier et $\pi > 2$ » est vraie.

3. L'assertion «la fonction exponentielle $\exp(\cdot)$ est décroissante et pour tout $x \in \mathbb{R} : \exp'(x) = \exp(x)$ » est fausse.

1.4.3 Disjonction

Définissons maintenant la notion cousine à celle de la conjonction : la disjonction logique, exprimée par "ou". Le connecteur logique de disjonction est un opérateur binaire qui relie deux assertions et exprime l'idée que l'une des deux, ou les deux, peuvent être vraies.

Définition 1.4.3 (*Disjonction*)

Soient P et Q deux assertions, on appelle disjonction de P et Q , l'assertion notée $(P \vee Q)$ ou encore $(P \text{ ou } Q)$, qui est vraie si au moins une des deux assertions P et Q est vraie, et qui est fausse quand les deux assertions P et Q sont fausses simultanément.

- La table de vérité du connecteur $\vee$ sera donc

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

Table de vérité de l'opérateur logique de disjonction.

- Il est évident que l'assertion $(P \text{ ou } Q)$ d'une part et l'assertion $(Q \text{ ou } P)$ ont même valeur de vérité.

Exemple 1.4.3

1. L'assertion « $\pi \geq 2$ ou $e \geq 1$ » est vraie.
2. Soient les assertions P : « $x < 2$ » et Q : « $x > 10$ », alors l'assertion $(P \vee Q)$ est « $x \in]-\infty, 2[\cup]10, +\infty[$ ».
3. Soient les assertions P : « n multiple de 3 inférieur à 10 » et Q : « n pair inférieur à 10», alors l'assertion $(P \vee Q)$ est « $n \in \{0, 2, 3, 4, 6, 8, 9\}$ ».

1.4.4 Implication logique

L'implication logique est un opérateur dérivé, construit à partir des opérateurs fondamentaux "non" et "ou". Ce connecteur logique constitue la base de la plupart des raisonnements que nous rencontrerons par la suite.

Définition 1.4.4

1. Soient P et Q deux assertions, l'implication de Q par P est l'assertion notée « $P \implies Q$ » ou « P implique Q » qui est fausse seulement si P est vraie et Q est fausse. L'implication est vraie dans tous les autres cas.
2. Le symbole « $\implies$ » s'appelle la flèche d'implication et suit la direction de l'implication logique.

- La table de vérité de l'implication logique est la suivante :

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

Table de vérité de l'opérateur logique de l'implication.

Proposition 1.4.1 On peut exprimer le connecteur logique $\implies$ à l'aide des connecteurs ou et non par

$$\text{non}(P) \vee Q.$$

Autrement dit l'assertion $(P \implies Q)$ est l'assertion $(\text{non}(P) \vee Q)$.

Preuve. Il suffit de comparer les deux assertions $(P \implies Q)$ et $(\bar{P} \vee Q)$ pour toutes les valeurs possibles de P et Q . Pour s'en convaincre voici la table de vérité montrant ceci

P	Q	$\bar{P}$	$\bar{P} \vee Q$	$P \implies Q$
1	1	0	1	1
1	0	0	0	0
0	1	1	1	1
0	0	1	1	1

Comme on peut le voir dans la table de vérité, les deux expressions $(P \implies Q)$ et $(\bar{P} \vee Q)$ ont exactement les mêmes valeurs de vérité dans toutes les situations, ce qui prouve leur équivalence. $\square$

Remarque 1.4.2 *L'implication est l'outil de base du raisonnement mathématique. Il est essentiel de bien l'assimiler, et de comprendre toutes ses formulations. Alors on peut exprimer $(P \implies Q)$ de l'une des façons suivantes :*

- Si P alors Q .
- P implique Q .
- P entraîne Q .
- P est une condition suffisante (CS) de Q (pour que Q soit vraie, il suffit que P le soit).
- Q est une condition nécessaire (CN) de P (si P est vraie, nécessairement Q l'est).
- Pour que P , il faut que Q .
- Pour que Q , il suffit que P .

Exemple 1.4.4

1. L'assertion « $x^2 \leq 1 \implies x \geq -1$ » est vraie.
2. L'assertion « $x^2 \leq 1 \implies 1 \geq x \geq -1$ » est vraie.
3. L'assertion « $x \leq 1 \implies x^2 \leq 1$ » est fausse (prendre par exemple $x = -2$).
4. L'assertion « $0 \leq x \leq 25 \implies x \leq 5$ » est vraie (prendre la racine carrée).
5. L'assertion « $x \in]-\infty, -4[\implies x^2 + 3x - 4 > 0$ » est vraie (étudier le binôme).
6. L'assertion « $\sin x \implies x = 0$ » est fausse (regarder pour $x = 2\pi$ par exemple).
7. L'assertion « $2 + 2 = 5 \implies \sqrt{2} = 2$ » est vraie (si P est fausse alors l'assertion $P \implies Q$ est toujours vraie).

Définition 1.4.5 Soient P et Q deux assertions, alors

1. Dans l'implication $P \implies Q$, P s'appelle l'hypothèse et Q la conclusion.
2. L'assertion $Q \implies P$ s'appelle l'implication réciproque de $P \implies Q$.
3. L'assertion $\bar{Q} \implies \bar{P}$ s'appelle l'implication contraposée de $P \implies Q$.

Exemple 1.4.5

1. Soient $x, y \in \mathbb{R}$, l'implication

$$x.y = 0 \implies (x = 0) \vee (y = 0),$$

sa contraposée est

$$(x \neq 0) \wedge (y \neq 0) \implies x.y \neq 0.$$

2. Soit $n \in \mathbb{N}$, l'implication

$$n^2 \text{ est pair} \implies n \text{ est pair},$$

sa contraposée est

$$n \text{ est impair} \implies n^2 \text{ est impair}.$$

1.4.5 Équivalence logique

Le connecteur logique d'équivalence, noté $\iff$, relie deux assertions en indiquant qu'elles sont vraies ou fausses ensemble. Cela signifie que, dans toutes les situations possibles, les deux assertions doivent être soit toutes les deux vraies, soit toutes les deux fausses.

Définition 1.4.6 On dit que deux assertions P et Q sont équivalentes si on a, à la fois $(P \implies Q)$ et $(Q \implies P)$. En termes logiques, P et Q sont équivalentes si elles ont les mêmes valeurs de vérité. On note dans ce cas

$$P \iff Q.$$

- La table de vérité de l'équivalence logique est la suivante

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

Table de vérité de l'opérateur logique de l'équivalence.

- La table ci-dessus montre que l'assertion P est équivalente à l'assertion Q si et seulement si elles sont simultanément vraies ou simultanément fausses.

Remarque 1.4.3 On peut traduire $(P \iff Q)$ par

- P équivaut à Q .
- P entraîne Q et réciproquement.
- Si P est vraie alors Q est vraie et réciproquement.
- P est vraie si et seulement si (ssi) Q est vraie.
- Pour que P soit vraie il faut et il suffit que Q le soit.
- P est une condition nécessaire et suffisante (CNS) pour Q .

Exemple 1.4.6

1. Pour deux réels x et y , l'équivalence

$$x.y = 0 \iff x = 0 \vee y = 0,$$

est vraie.

2. Pour deux réels x et y , l'équivalence

$$x.y = 0 \iff x = 0 \wedge y = 0,$$

est fausse.

3. Pour deux réels x et y , l'équivalence

$$|x + y| = |x| + |y| \iff x.y \geq 0,$$

est vraie.

4. Soit $x \in \mathbb{R}$, l'assertion « $x \in [0, 1]$ » n'est pas équivalente à l'assertion « $x \in \mathbb{R}_+$ ». En effet l'implication directe

$$x \in [0, 1] \implies x \in \mathbb{R}_+,$$

est vraie. Par contre l'implication réciproque

$$x \in \mathbb{R}_+ \implies x \in [0, 1],$$

est fausse.

5. Soit $x \in \mathbb{R}$, l'assertion « $x^2 = 1$ » est équivalente à l'assertion « $x = +1$ ou $x = -1$ ». En effet les deux implications $\implies$ et $\impliedby$ sont vraies.

6. Soit $x \in \mathbb{R}_+$, l'assertion « $x^2 = 1$ » est équivalente à l'assertion « $x = +1$ ». En effet ici on s'est limité à un paramètre x positif, donc la solution $x = -1$ n'est pas à prendre en compte.

Proposition 1.4.2 *Etant données deux assertions logiques P et Q , alors les assertions $(P \implies Q)$ et $(\bar{Q} \implies \bar{P})$ sont logiquement équivalentes. En d'autres termes, elles ont même valeur de vérité.*

Preuve. On établit la preuve de ce résultat en donnant les valeurs de vérité des assertions logiques correspondantes. Alors

P	Q	$\bar{P}$	$\bar{Q}$	$P \implies Q$	$\bar{Q} \implies \bar{P}$
1	1	0	0	1	1
1	0	0	1	0	0
0	1	1	0	1	1
0	0	1	1	1	1

On voit que les assertions $(P \implies Q)$ et $(\bar{Q} \implies \bar{P})$ sont simultanément vraies ou simultanément fausses. Elles sont donc bien équivalentes. $\square$

Conclusion 1.4.1 *Voici un tableau qui résume les valeurs de vérités des divers connecteurs que nous avons introduits. Vrai (resp. Faux) est symbolisé avec 1 (resp. 0)*

P	Q	$\bar{P}$	$P \wedge Q$	$P \vee Q$	$P \implies Q$	$P \iff Q$
1	1	0	1	1	1	1
1	0	0	0	1	0	0
0	1	1	0	1	1	0
0	0	1	0	0	1	1

Table de vérité de des principaux connecteurs.

1.5 Règles logiques

En logique, il existe plusieurs règles qui établissent un calcul des assertions, semblable à celui du calcul algébrique. L'application de l'une de ces règles sur une assertion permet d'obtenir une assertion équivalente, c'est-à-dire ayant la même valeur de vérité. Les principales propriétés des connecteurs sont résumées dans la proposition suivante.

Proposition 1.5.1 *Pour toutes assertions P, Q , et R , on a*

1. *Loi de non contradiction*

$$P \wedge \text{non}(P) \text{ est une assertion fausse.}$$

2. *Double négation*

$$\text{non}(\text{non}(P)) \iff P.$$

3. *Lois de Morgan*

$$\begin{cases} \text{non}(P \wedge Q) \iff \text{non}(P) \vee \text{non}(Q) \text{ (Négation d'une conjonction),} \\ \text{non}(P \vee Q) \iff \text{non}(P) \wedge \text{non}(Q) \text{ (Négation d'une disjonction).} \end{cases}$$

4. *Commutativité de $\wedge$ et de $\vee$*

$$\begin{cases} (P \wedge Q) \iff (Q \wedge P), \\ (P \vee Q) \iff (Q \vee P). \end{cases}$$

5. Idempotence de $\wedge$ et de $\vee$

$$\begin{cases} (P \wedge P) \iff P, \\ (P \vee P) \iff P. \end{cases}$$

6. Associativité de $\wedge$ et de $\vee$

$$\begin{cases} P \wedge (Q \wedge R) \iff (P \wedge Q) \wedge R, \\ P \vee (Q \vee R) \iff (P \vee Q) \vee R. \end{cases}$$

7. Distributivité de $\wedge$ et de $\vee$

$$\begin{cases} P \wedge (Q \vee R) \iff (P \wedge Q) \vee (P \wedge R) \text{ (Distributivité de } \wedge \text{ par rapport à } \vee), \\ P \vee (Q \wedge R) \iff (P \vee Q) \wedge (P \vee R) \text{ (Distributivité de } \vee \text{ par rapport à } \wedge). \end{cases}$$

8. Loi de contraposition

$$(P \implies Q) \iff (\text{non}(Q) \implies \text{non}(P)).$$

9. Négation de l'implication

$$\text{non}(P \implies Q) \iff (P \wedge \text{non}(Q)).$$

10. Transitivité de l'implication

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

11. Transitivité de l'équivalence

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

Preuve.

• Ces résultats peuvent s'obtenir à partir des tables de vérités des assertions étudiées. Par exemple, pour montrer (1), il suffit de tracer la table de vérité de $(P \wedge \text{non}(P))$ et de vérifier qu'il n'y a que des 0 dans la colonne finale.

P	$\text{non}(P)$	$(P \wedge \text{non}(P))$
1	0	0
0	1	0

• Montrons maintenant (7) (Distributivité de $\wedge$ par rapport à $\vee$). La table de vérité de $(P \wedge Q) \vee (P \wedge R)$ est définie de la manière suivante :

Tout d'abord on s'inquiète de savoir quel est le nombre d'assertions qu'il y a dans $(P \wedge Q) \vee (P \wedge R)$. Dans notre cas il y en a 3, cela induit une table de vérité avec $(1 + 8)$ lignes. Si la formule avait n assertions, cela donnerait une table avec $(1 + 2^n)$ lignes. Les n premières colonnes sont réservées aux assertions, la $(n + 1)^{\text{ème}}$ est pour l'assertion elle-même. On obtient donc

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

On voit que les assertions logiques $P \wedge (Q \vee R)$ et $(P \wedge Q) \vee (P \wedge R)$ ont les mêmes valeurs de vérité, donc elles sont équivalentes. De même pour la distributivité de $\vee$ par rapport à $\wedge$. Nous laissons au lecteur le soin de vérifier de même chacune des autres équivalences. $\square$

1.6 Différents types de raisonnement mathématiques

Les raisonnements mathématiques sont des processus formels utilisés pour démontrer des vérités ou établir des conclusions à partir d'axiomes, de théorèmes et de règles logiques. Plusieurs types de raisonnements mathématiques existent, et nous allons examiner ici les plus courants : le raisonnement direct, le raisonnement par contraposée, le raisonnement par l'absurde, le raisonnement par récurrence, la disjonction des cas et le raisonnement par contre-exemple.

1.6.1 Raisonnement direct

Le raisonnement direct est une méthode de démonstration où l'on part d'hypothèses pour arriver directement à une conclusion grâce à une succession d'implications.

Définition 1.6.1 (*Principe*)

Soient P et Q deux assertions données. Pour montrer directement l'implication $P \implies Q$, on suppose que P est vraie et on démontre que Q est vraie. La démonstration commence par «supposons que P est vraie» et se termine par « Q est vraie».

Exemple 1.6.1

1. Montrer que si $x, y \in \mathbb{Q}$ alors $x + y \in \mathbb{Q}$.

Solution. Prenons $x, y \in \mathbb{Q}$. Alors $x = \frac{m_1}{n_1}$ pour un certain $m_1 \in \mathbb{Z}$ et un certain $n_1 \in \mathbb{N}^*$. De même $y = \frac{m_2}{n_2}$ pour un certain $m_2 \in \mathbb{Z}$ et un certain $n_2 \in \mathbb{N}^*$. Maintenant

$$x + y = \frac{m_1}{n_1} + \frac{m_2}{n_2} = \frac{m_1 n_2 + n_1 m_2}{n_1 n_2}.$$

Or le numérateur $m_1 n_2 + n_1 m_2$ est bien un élément de $\mathbb{Z}$, le dénominateur $n_1 n_2$ est lui un élément de $\mathbb{N}^*$. Donc $x + y$ s'écrit bien de la forme

$$x + y = \frac{m_3}{n_3},$$

avec $m_3 \in \mathbb{Z}$, $n_3 \in \mathbb{N}^*$. Ainsi $x + y \in \mathbb{Q}$. Avec des notations plus symboliques, on vient de montrer que

$$x, y \in \mathbb{Q} \implies x + y \in \mathbb{Q}.$$

2. Montrer que si $x < 1$, alors $|x - 4| > 3$.

Solution. Si $x < 1$, alors

$$x - 4 < -3.$$

Par décroissance de la fonction valeur absolue sur $\mathbb{R}^-$, on en déduit que

$$|x - 4| > |-3|,$$

Ceci veut bien dire que $|x - 4| > 3$. Avec des notations plus symboliques, on vient de montrer que

$$x < 1 \implies |x - 4| > 3.$$

1.6.2 Raisonnement par Contraposée

Le raisonnement par contraposée est une méthode qui consiste à établir l'implication « si $\text{non}(Q)$, alors $\text{non}(P)$ » à partir de l'implication « si P , alors Q ». L'implication « si $\text{non}(Q)$ alors $\text{non}(P)$ » est appelée contraposée de « si P alors Q ». Ce mode de raisonnement est particulièrement utile lorsque prouver directement $P \implies Q$ s'avère complexe, tandis que la démonstration de la contraposée $\bar{Q} \implies \bar{P}$ est plus aisée. Cette technique repose sur le principe selon lequel $P \implies Q$ est logiquement équivalent à $\bar{Q} \implies \bar{P}$.

Définition 1.6.2 (*Principe*)

1. Le raisonnement par contraposée (ou *contraposition*) est basé sur l'équivalence d'une implication et de sa «contraposée» : P et Q étant deux assertions quelconques,

$$(P \implies Q) \iff (\bar{Q} \implies \bar{P}).$$

2. Démontrer par contraposition l'implication $(P \implies Q)$ consiste à démontrer l'implication $(\bar{Q} \implies \bar{P})$. Pour cela, on suppose $\bar{Q}$ vraie et on en déduit $\bar{P}$ vraie.

Remarque 1.6.1 Il est difficile de donner une règle générale d'utilisation de ce raisonnement. Un bon conseil avant de se lancer dans la démonstration d'une implication, est d'écrire d'abord sa contraposée. Avec un peu d'expérience, on arrive vite à sentir laquelle des deux est la plus facile à démontrer. Si le résultat désiré est R , on cherche les conséquences de $\text{non}(R)$ pour arriver aux bonnes hypothèses. Notre premier exemple est un résultat facile, mais très utile.

Exemple 1.6.2 Soit $n \in \mathbb{N}$, prouver que n^2 est pair si et seulement si n est pair.

Solution. Soit $n \in \mathbb{N}$, montrons l'équivalence demandée en procédant par double implication.

• L'implication réciproque $\Leftarrow$. On suppose que n est pair, alors il existe un entier $k \in \mathbb{N}$ tel que $n = 2k$. Donc

$$n^2 = 4k^2 = 2k', k' = 2k^2 \in \mathbb{N}.$$

Ce qui prouve que n^2 est pair. On a ainsi prouvé que si n est pair alors n^2 est pair.

• L'implication directe $\Rightarrow$. Il ne paraît pas simple de prouver directement que si n^2 est pair, alors n est pair. Essayons donc de montrer la contraposée. Pour cela, montrons que

$$\text{non}(n \text{ est pair}) \implies \text{non}(n^2 \text{ est pair}).$$

Autrement dit

$$n \text{ est impair} \implies n^2 \text{ est impair}.$$

Supposons n impair. Il existe $k \in \mathbb{N}$ tel que $n = 2k + 1$. On a alors

$$n^2 = (2k + 1)^2 = 2(2k^2 + 2k) + 1 = 2m + 1.$$

On a ainsi écrit n^2 sous la forme

$$n^2 = 2m + 1,$$

avec $m = 2k^2 + 2k \in \mathbb{N}$, et donc n^2 est impair. On a ainsi prouvé que si n est impair alors n^2 est impair.

Par contraposition, on en déduit que si n^2 est pair, alors n est pair.

Conclusion : Par double implication, nous avons montré que pour tout $n \in \mathbb{N}$,

$$n^2 \text{ est pair} \iff n \text{ est pair}.$$

Exemple 1.6.3 Soit $x \in \mathbb{R}$, montrons en utilisant un raisonnement par contraposée l'assertion suivante

$$(\forall \varepsilon > 0 : |x| < \varepsilon) \implies x = 0.$$

Solution. On doit donc montrer que

$$\text{non}(x = 0) \implies \text{non}(\forall \varepsilon > 0 : |x| < \varepsilon).$$

Autrement dit

$$x \neq 0 \implies (\exists \varepsilon > 0 : |x| \geq \varepsilon).$$

Soit $x \in \mathbb{R}$, supposons que $x \neq 0$ et montrons qu'il existe $\varepsilon > 0$ tel que $|x| \geq \varepsilon$. C'est immédiat. Il suffit en effet de prendre

$$\varepsilon = \frac{|x|}{2},$$

puisque $\frac{|x|}{2}$ est strictement positif et

$$|x| \geq \frac{|x|}{2}.$$

D'une manière générale,

$$\varepsilon = \frac{|x|}{\alpha},$$

avec $\alpha \geq 1$ convient aussi puisque, pour tout réel $\alpha \geq 1$, on a

$$\begin{cases} \frac{|x|}{\alpha} > 0, \\ \text{et} \\ |x| \geq \frac{|x|}{\alpha}. \end{cases}$$

1.6.3 Raisonnement par l'Absurde

Le raisonnement par l'absurde, également connu sous le nom de raisonnement par contradiction, est une méthode de démonstration qui consiste à supposer que l'assertion que l'on souhaite prouver est fausse. On montre ensuite que cette supposition conduit à une contradiction. En établissant qu'une contradiction découle de l'hypothèse initiale, on conclut que l'assertion originale doit nécessairement être vraie. Cela présente quelques similitudes avec le raisonnement par contraposée, mais ce n'est pas tout à fait la même chose.

Définition 1.6.3 (*Principe*)

Le principe du raisonnement par l'absurde est le suivant : pour démontrer qu'une assertion P est vraie, on suppose le contraire (c'est-à-dire P est fausse), et on essaye d'arriver à un résultat faux (souvent sous forme de contradiction). Ceci montre que l'hypothèse de départ est fausse, donc que P est vraie.

Exemple 1.6.4

1. Montrons que le nombre 0 n'a pas d'inverse dans $\mathbb{R}$.

Solution. On suppose que 0 a un inverse dans $\mathbb{R}$. On le note a . Par définition de l'inverse,

$$0 \cdot a = 1.$$

Or pour tout réel x ,

$$0 \cdot x = 0.$$

On en déduit que $0 = 1$. Ce qui est faux. On en déduit que 0 n'a pas d'inverse dans $\mathbb{R}$.

2. On souhaite démontrer que $\sqrt{2}$ est irrationnel. La preuve classique de l'irrationalité de $\sqrt{2}$ est une preuve par l'absurde. Pour cela, on suppose qu'il est rationnel et on aboutit à une contradiction.

Solution. Supposons que $\sqrt{2}$ soit un nombre rationnel. Il existe un couple $(m, n) \in \mathbb{N} \times \mathbb{N}^*$ avec m et n sont premiers entre eux (c'est-à-dire sans diviseurs commun autre que 1) tel que

$$\sqrt{2} = \frac{m}{n} \text{ (la forme d'une fraction irréductible).}$$

On obtient

$$m^2 = 2n^2,$$

d'où p^2 est pair. D'après l'exemple 1.6.2 précédent, on en déduit que m est pair. Il existe donc $k \in \mathbb{N}$ tel que

$$m = 2k.$$

On obtient

$$n^2 = 2k^2,$$

ce qui implique que n^2 est pair. On en déduit alors que n est pair (voir l'exemple 1.6.2 précédent), ce qui signifie qu'il existe $l \in \mathbb{N}^*$ tel que

$$n = 2l.$$

On a donc fait apparaître un diviseur commun à m et n (à savoir le nombre 2), ce qui est contraire à notre hypothèse (m et n sont premiers entre eux). Donc m et n n'existent pas et la supposition $\sqrt{2}$ est rationnel est fausse. $\sqrt{2}$ est un nombre irrationnel.

3. Soient $a, b > 0$. Montrons que si

$$\frac{a}{1+b} = \frac{b}{1+a},$$

alors

$$a = b.$$

Solution. Nous raisonnons par l'absurde en supposant que $\frac{a}{1+b} = \frac{b}{1+a}$ et $a \neq b$. Comme

$$\frac{a}{1+b} = \frac{b}{1+a},$$

alors

$$a(a+1) = b(b+1),$$

donc

$$a^2 + a = b^2 + b,$$

d'où

$$a^2 - b^2 = b - a.$$

Cela conduit à

$$(a-b)(a+b) = -(a-b),$$

Comme $a \neq b$ alors $a-b \neq 0$ et donc en divisant par $a-b$ on obtient

$$a+b = -1.$$

La somme des deux nombres positifs a et b ne peut être négative. Nous obtenons une contradiction.

Conclusion. si $\frac{a}{1+b} = \frac{b}{1+a}$ alors $a = b$.

1.6.4 Raisonnement par Disjonction des cas

Lorsqu'on souhaite démontrer une assertion, il peut être plus simple de la décomposer en sous-cas particuliers. Si l'ensemble de ces cas couvre toutes les possibilités, alors l'assertion est prouvée dans toute sa généralité. Cette approche correspond à un raisonnement par disjonction des cas ou raisonnement cas par cas. Ce mode de raisonnement consiste à décomposer l'assertion que l'on cherche à établir en un nombre fini de cas (sous-propositions), qui sont vérifiés de manière indépendante.

Définition 1.6.4 (Principe)

Soient P et Q deux assertions. Pour montrer que « $P \implies Q$ », on sépare l'hypothèse P de départ en différents cas possibles et on montre que l'implication est vraie dans chacun des cas.

Exemple 1.6.5

1. Montrons que pour tout $x \in \mathbb{R}$,

$$|x - 1| - |x - 2| + 1 \geq 0.$$

Solution. Les valeurs absolues sont gênantes à manipuler, on va donc raisonner par disjonction des cas pour s'en débarrasser. On sépare le problème en trois cas selon que

$$x \leq 1, 1 < x \leq 2, x > 2.$$

Soit $x \in \mathbb{R}$. Nous distinguons les cas suivants :

• Premier cas : Si $x \leq 1$, alors dans ce cas

$$|x - 1| = 1 - x \text{ et } |x - 2| = 2 - x,$$

ainsi

$$|x - 1| - |x - 2| + 1 = 1 - x - (2 - x) + 1 = 0 \geq 0.$$

• Deuxième cas : Si $1 < x \leq 2$, alors dans ce cas

$$|x - 1| = x - 1 \text{ et } |x - 2| = 2 - x,$$

ainsi

$$|x - 1| - |x - 2| + 1 = x - 1 - (2 - x) + 1 = 2(x - 1) \geq 0.$$

• Troisième cas : Si $x > 2$, alors dans ce cas

$$|x - 1| = x - 1 \text{ et } |x - 2| = x - 2,$$

ainsi

$$|x - 1| - |x - 2| + 1 = x - 1 - (x - 2) + 1 = 0 \geq 0.$$

Conclusion. Dans tous les cas $|x - 1| - |x - 2| + 1 \geq 0$.

2. Montrer que pour tout $n \in \mathbb{N}$, $n(n + 1)$ est pair.

Solution. Soit $n \in \mathbb{N}$. Nous distinguons deux cas :

• Premier cas : Si n est pair, alors on peut écrire $n = 2p$ avec $p \in \mathbb{N}$. Ainsi

$$n(n + 1) = 2p(2p + 1) = 2k, k = p(2p + 1) \in \mathbb{N},$$

est pair.

• Deuxième cas : Si n est impair, alors on peut écrire $n = 2p + 1$ avec $p \in \mathbb{N}$. Ainsi

$$n(n + 1) = (2p + 1)(2p + 2) = 2m + 1, m = (p + 1)(2p + 1) \in \mathbb{N},$$

est impair.

Conclusion. Dans tous les cas, on a $n(n + 1)$ est pair.

1.6.5 Raisonnement par Contre-exemple

Un contre-exemple est un exemple particulier et concret qui contredit une affirmation, un énoncé, une conjecture, une règle générale ou une loi. Trouver un contre-exemple pour prouver qu'une assertion est fausse peut parfois s'avérer aussi difficile que de démontrer que l'assertion est vraie. Ce type de raisonnement est essentiel pour établir les limites de certaines propositions et pour clarifier des énoncés qui peuvent sembler vrais dans des cas généraux mais qui échouent dans des cas spécifiques.

Définition 1.6.5 (Principe)

Soit $P(x)$ une propriété dépendant de la variable x dans E . Pour montrer que l'assertion « $\forall x \in E, P(x)$ » est fausse, il suffit de trouver un x de E tel que $P(x)$ soit fausse. Trouver un tel x c'est trouver un contre-exemple à l'assertion « $\forall x \in E, P(x)$ ».

Remarque 1.6.2 Ce mode de raisonnement repose sur le principe que la négation de l'assertion

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

est

$$\exists x \in E : \text{non}(P(x)).$$

Trouver tel x c'est trouver un contre-exemple à l'assertion « $\forall x \in E, P(x)$ ».

Exemple 1.6.6

1. Montrer que l'assertion «Pour tout entier naturel n , si n est divisible par 4 et par 6, alors n est divisible par 24» est fausse.

Solution. Un contre-exemple est $n = 12$, il est divisible par 4 et par 6, mais il n'est pas divisible par 24.

2. Montrer que l'assertion suivante est fausse «Tout entier positif est somme de trois carrés». Par exemple

$$14 = 1^2 + 2^2 + 3^2.$$

Solution. Un contre-exemple est 7 : les carrés inférieurs à 7 sont 0, 1, 4 mais avec trois de ces nombres on ne peut faire 7 ($0 + 1 + 4 \neq 7$).

1.6.6 Raisonnement par Récurrence

Pour conclure cette liste de stratégies de raisonnement, évoquons brièvement le raisonnement par récurrence, une méthode que vous connaissez probablement bien depuis le lycée. Le raisonnement par récurrence est une technique mathématique utilisée pour démontrer une propriété qui s'applique à tous les entiers naturels, ou, plus généralement, à une infinité d'entiers naturels.

Ce type de raisonnement est pertinent lorsque l'on souhaite prouver une assertion $\mathcal{P}(n)$ pour tout $n \in \mathbb{N}$. On commence par établir la base de la récurrence en prouvant l'assertion pour $\mathcal{P}(0)$. Ensuite, on procède à l'hérédité, qui consiste à démontrer que l'implication $(\mathcal{P}(n) \implies \mathcal{P}(n+1))$ est vraie pour tout $n \in \mathbb{N}$. En établissant ces deux étapes, on conclut que la propriété est vraie pour tous les entiers naturels.

Définition 1.6.6 (Principe)

Soit $\mathcal{P}(n)$ une propriété dépendant d'un entier $n \in \mathbb{N}$. Pour montrer que $\mathcal{P}(n)$ est vraie pour tout entier $n \geq 0$, il suffit de démontrer que :

- La proposition $\mathcal{P}(0)$ est vraie (initialisation).
 - Pour n'importe quel entier $n \geq 0$ fixé, si $\mathcal{P}(n)$ est vraie, alors $\mathcal{P}(n+1)$ est vraie également (hérédité). C'est-à-dire $\mathcal{P}(n) \implies \mathcal{P}(n+1)$ est vraie.
- La propriété $\mathcal{P}(n)$ supposée vraie est appelée **hypothèse de récurrence**.

Exemple 1.6.7**1. Inégalité de Bernoulli.**

Soit $x \in \mathbb{R}^+$, montrer pour tout entier naturel $n \in \mathbb{N}$ l'inégalité (de Bernoulli) suivante

$$(1+x)^n \geq 1+nx.$$

Solution. Soit $x \in \mathbb{R}^+$, on note $\mathcal{P}(n)$ l'assertion suivante

$$(1+x)^n \geq 1+nx.$$

Nous allons montrer par récurrence que $\mathcal{P}(n)$ est vraie pour tout $n \in \mathbb{N}$.

- Pour $n = 0$, nous avons

$$(1+x)^0 = 1 \geq 1+0x.$$

Donc $\mathcal{P}(0)$ est vraie.

- Maintenant fixons $n \in \mathbb{N}$ et supposons que $\mathcal{P}(n)$ est vraie, c'est-à-dire

$$(1+x)^n \geq 1+nx.$$

Comme x est positif, on a aussi $1+x \geq 0$, on peut donc multiplier l'inégalité par $(1+x)$ sans en changer le sens. On obtient

$$(1+x)^n(1+x) \geq (1+nx)(1+x).$$

C'est-à-dire

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

Ce qui donne

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

Comme on a $nx^2 \geq 0$, alors

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

On en déduit en particulier qu'on a

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

$\mathcal{P}(n+1)$ est donc vraie.

Conclusion. Par le principe de récurrence $\mathcal{P}(n)$ est vraie pour tout $n \in \mathbb{N}$, c'est-à-dire on a $(1+x)^n \geq 1+nx$ pour tout $n \in \mathbb{N}$ et tout $x \geq -1$.

Remarque 1.6.3 Comme l'illustre l'exemple suivant, un raisonnement par récurrence peut être utilisé si la propriété $\mathcal{P}(n)$ n'est vraie qu'à partir d'un certain rang n_0 où n_0 désigne un entier naturel non nécessairement égal à 0. Dans ce cas précis, la première étape consistera à montrer que la propriété $\mathcal{P}(n_0)$ est vraie.

Exemple 1.6.8 Montrons que pour tout entier naturel $n \in \mathbb{N}^*$ non nul,

$$1+2+\dots+n = \frac{n(n+1)}{2}.$$

Solution. Soit $n \in \mathbb{N}$, notons $\mathcal{P}(n)$ l'assertion suivante

$$1 + 2 + \dots + n = \frac{n(n+1)}{2}.$$

Vérifions que $\mathcal{P}(1)$ est vraie. Si $n = 1$, le membre de gauche vaut 1. Le membre de droite vaut, $\frac{1(1+1)}{2}$ soit 1 aussi, on a bien

$$1 = \frac{1(1+1)}{2},$$

autrement dit $\mathcal{P}(1)$ est vraie.

Soit n un entier naturel non nul, si $\mathcal{P}(n)$ est vraie, alors

$$1 + 2 + \dots + n = \frac{n(n+1)}{2},$$

il s'agit d'en déduire que $\mathcal{P}(n+1)$ est vraie. Alors, en ajoutant $(n+1)$ à chaque membre,

$$1 + 2 + \dots + n + (n+1) = \frac{n(n+1)}{2} + (n+1),$$

alors, en factorisant le membre de droite

$$1 + 2 + \dots + n + (n+1) = \frac{(n+1)(n+2)}{2},$$

donc

$$1 + 2 + \dots + n + (n+1) = \frac{(n+1)((n+1)+1)}{2},$$

alors $\mathcal{P}(n+1)$ est vraie.

Conclusion. Par le principe de récurrence, $\mathcal{P}(n)$ est vraie pour tout entier n supérieur ou égal à 1.

Chapitre 2

Ensembles et Relations binaires sur un ensemble

Ce chapitre présente les concepts fondamentaux de la théorie des ensembles, incluant les opérations d'union, d'intersection, et de complémentaire. Elle aborde aussi les relations binaires et leurs propriétés, telles que les relations d'équivalence et d'ordre.

2.1 Théorie des ensembles

La théorie des ensembles constitue une branche fondamentale des mathématiques, dédiée à l'étude des ensembles, c'est-à-dire des collections d'objets distincts considérés comme un tout. Ces objets peuvent être des nombres, des points, des fonctions ou toute autre entité mathématique. Cette théorie est cruciale pour de nombreux sous-domaines des mathématiques, en particulier pour la probabilité, qui est une branche des statistiques. Cette connexion est essentielle, car la probabilité s'appuie sur des concepts issus de la théorie des ensembles pour définir et analyser les événements, les variables aléatoires et les distributions.

2.1.1 Définitions et notations

Définition 2.1.1

1. *Un ensemble est une collection (ou un groupement) d'objets, ces objets s'appellent les éléments ou les points de l'ensemble, et ces objets sont considérés sans ordre et sans répétition.*
2. *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.*

Exemple 2.1.1 *Les exemples suivants sont déjà bien connus du lecteur.*

1.
 - $\mathbb{N}$ l'ensemble des entiers naturels.
 - $\mathbb{Z}$ l'ensemble des entiers relatifs.
 - $\mathbb{Q}$ l'ensemble des nombres rationnels.
 - $\mathbb{R}$ l'ensemble des nombres réels.
 - $\mathbb{C}$ l'ensemble des nombres complexes.
2. *Les ensembles $\mathbb{N}, \mathbb{Z}, \mathbb{Q}, \mathbb{R}, \mathbb{C}$ privés du nombre 0 sont souvent utilisés, on les dénote par la notation étoilée $\mathbb{N}^*, \mathbb{Z}^*, \mathbb{Q}^*, \mathbb{R}^*, \mathbb{C}^*$, ou la notation ensembliste usuelle $\mathbb{N} \setminus \{0\}, \mathbb{Z} \setminus \{0\}, \mathbb{Q} \setminus \{0\}, \mathbb{R} \setminus \{0\}, \mathbb{C} \setminus \{0\}$.*

Définition 2.1.2 (Relation d'appartenance)

1. Lorsqu'un objet x est un élément d'un ensemble E , on dit que x appartient à E . L'appartenance de x à E est notée $x \in E$.
2. Pour exprimer que x n'est pas un élément de E ou l'objet x n'appartient pas à l'ensemble E , on écrit et $x \notin E$.

Exemple 2.1.2 Si $E = \{a, b, c\}$, alors

- $a \in E$ (l'élément a appartient à l'ensemble E).
- $d \notin E$ (l'élément d n'appartient pas à E).

Remarque 2.1.1 Deux points importants sont à noter :

- a. **L'ordre** dans lequel les éléments apparaissent dans la liste n'a aucune importance. Autrement dit les ensembles sont uniquement définis par la présence ou l'absence de leurs éléments, et non par l'ordre dans lequel ces éléments sont listés. Par exemple, soit $A = \{1, 2, 3\}$. Cet ensemble est exactement le même que $B = \{3, 2, 1\}$, $C = \{2, 1, 3\}$ et $D = \{1, 3, 2\}$
- b. Si un ou plusieurs éléments sont **répétés** dans la liste de définition d'un ensemble, l'ensemble ne change pas. En d'autres termes, chaque élément d'un ensemble est unique. Ainsi si $A = \{1, 2, 2, 3, 3, 3\}$, alors A est en réalité l'ensemble $A = \{1, 2, 3\}$, car les répétitions des éléments 2 et 3 n'ont pas d'importance.

En théorie des ensembles, il existe principalement deux manières fondamentales de définir un ensemble, en extension et en compréhension.

Définition 2.1.3

1. **Définition en extension.** La définition en extension consiste à donner la liste explicite de tous les éléments d'un ensemble. Cette liste est écrite entre accolades et les éléments sont séparés par des virgules. Les accolades $\{$ et $\}$ servent à marquer le début et la fin de la description de l'ensemble.
2. **Définition en compréhension.** La définition en compréhension consiste à décrire un ensemble en spécifiant une propriété ou une condition que ses éléments doivent satisfaire. Au lieu de lister explicitement les éléments, on utilise une formule qui indique la règle à respecter pour faire partie de l'ensemble. L'ensemble est noté sous la forme $\{x \in E : \mathcal{P}(x)\}$. Ce qui se lit "l'ensemble des éléments x de l'ensemble E qui vérifient la propriété $\mathcal{P}$."

• La définition en extension est utile pour les ensembles avec un petit nombre d'éléments, tandis que la définition en compréhension est plus puissante et pratique pour les ensembles infinis ou définis par une propriété spécifique.

Exemple 2.1.3

1. L'ensemble des jours de la semaine peut être défini en extension par

$$S = \{\text{vendredi, samedi, dimanche, lundi, mardi, mercredi, jeudi}\}.$$

2. Si A est l'ensemble des trois premiers nombres naturels, on peut le définir en extension comme suit

$$A = \{0, 1, 2\}.$$

3. L'ensemble des nombres pairs peut être défini en compréhension comme suit

$$\mathbb{P} = \{x \in \mathbb{Z} : x \text{ est pair}\}.$$

4. L'ensemble des solutions d'une équation est, par nature, défini en compréhension. Ainsi

$$E = \{x \in \mathbb{R} : x^2 - 2x + 1 = 0\},$$

est l'ensemble des nombres réels x qui vérifient la condition $x^2 - 2x + 1 = 0$.

5. Voici d'autres exemples d'ensemble sous les définitions par extension et en compréhension :

- $E = \{\text{farine, eau, levure, huile, sel, mzzarella, tomates, champignons, olives, basilic}\}$ ou $\{x : x \text{ est un ingrédient de préparation de pizza}\}$.
- $E = \{1, 3, 5, 7, 9\}$ ou $\{x : x \text{ est un entier impair compris entre 1 et 9}\}$.

Remarque 2.1.2

- Une erreur d'écriture fréquente $A = \{x \in \mathbb{R} \mid x \geq 0\}$ (incorrect) au lieu de $A = \{x \in \mathbb{R} : x \geq 0\}$ (correct). La condition $\mathcal{P}$, ici $x \geq 0$, doit se trouver à l'intérieur des accolades et non à l'extérieur.
- Il est important de préciser explicitement à quel ensemble E s'applique la propriété $\mathcal{P}$ qui caractérise l'ensemble que l'on veut décrire, sinon la définition risque d'être ambiguë. Par exemple, $\{x \in \mathbb{N} : x^2 = 1\} = \{1\}$ et $\{x \in \mathbb{Z} : x^2 = 1\} = \{1, -1\}$. On voit ainsi que l'écriture $\{x : x^2 = 1\}$ serait incomplète pour définir un ensemble bien déterminé.
- Un même ensemble peut admettre plusieurs descriptions en compréhension différentes. Par exemple,

$$E = \{x \in \mathbb{N} : x > 1 \text{ et } x \leq 3\} = \{x \in \mathbb{R} : x^2 - 5x + 6 = 0\}.$$

- Il est généralement facile de déterminer si un objet donné appartient ou non à un ensemble défini en compréhension : il suffit de vérifier si cet objet satisfait ou non la condition qui définit l'ensemble. En revanche, il n'est pas toujours facile de savoir quels sont exactement tous les éléments d'un tel ensemble. Par exemple, soit $E = \{x \in \mathbb{R} : x^5 + x^4 - 5x^3 + 4x^2 + x - 2 = 0\}$. On vérifie facilement que $1 \in E$, mais il est plus difficile de déterminer tous les éléments de E , autrement dit de résoudre l'équation qui définit E .

Deux ensembles particuliers vont jouer un rôle spécial en théorie des ensembles. Le premier est l'ensemble référentiel, appelé également l'univers. On le note E . Le second ensemble particulier est l'ensemble vide, noté ϕ , qui ne comporte aucun élément.

Définition 2.1.4 (Ensemble vide)

On appelle ensemble vide l'ensemble qui ne contient aucun élément. Autrement dit, c'est un ensemble sans aucun objet ou membre. On le note généralement par le symbole ϕ ou par deux accolades vides. On a donc $\phi = \{\}$. Il est donc défini par la propriété

$$\forall x, x \notin \phi.$$

Exemple 2.1.4 Si l'on définit un ensemble avec une condition impossible à satisfaire, cet ensemble sera vide. Par exemple

$$A = \{x \in \mathbb{R} : x^2 = -1\}.$$

Comme il n'existe aucun nombre réel dont le carré est égal à -1 , cet ensemble est vide, donc $A = \phi$.

Définition 2.1.5 (Singletons, paires)

1. Un singleton est un ensemble qui a exactement un élément. Un singleton s'écrit donc $\{x\}$, où x est l'unique élément du singleton.
2. Une paire est un ensemble qui a exactement deux éléments. Une paire s'écrit donc $\{x, y\}$, où x et y sont les deux éléments (distincts) de la paire.

Exemple 2.1.5

1. $\{1\}$ et $\{\pi\}$ sont des singletons.
2. $\{1, \pi\}$ est une paire. En revanche, $\{1, 1\}$ est un singleton, car la répétition d'un élément n'a pas d'importance dans un ensemble.

Définition 2.1.6 (Ensemble fini, infini)

1. Un ensemble E est dit fini lorsque le nombre d'éléments qui le composent est un entier naturel. Dans ce cas, le nombre d'éléments est appelé le cardinal de l'ensemble. On le note $\text{Card}(E)$ ou $|E|$.
2. Un ensemble qui n'est pas fini est dit infini.
3. Par convention $\text{Card}(\emptyset) = 0$.

Exemple 2.1.6

1. L'ensemble $E = \{a, b, c\}$ est fini.
2. Les ensembles $\mathbb{N}, \mathbb{Z}, \mathbb{Q}, \mathbb{R}$ et $[0, 1]$ sont infinis.

2.1.2 Relations entre ensembles

Les relations entre ensembles sont des concepts fondamentaux en théorie des ensembles, permettant de comparer et de relier différents ensembles. On donne les principales relations entre ensembles.

2.1.2.1 Inclusion d'ensembles

À partir du signe d'appartenance $\in$ on introduit une abréviation notée $\subset$ et appelée signe d'inclusion.

Définition 2.1.7 (Inclusion)

1. Soient A et B deux ensembles. On dit que A est inclus dans B (ou que A est un sous-ensemble de B , ou encore que A est une partie de B ou que B contient A) et on note $A \subset B$ si et seulement si tout élément de A est aussi un élément de B . Formellement, cela se traduit par

$$(A \subset B) \iff \forall x, (x \in A \implies x \in B).$$

2. On utilise le symbole $\subset$ pour désigner l'inclusion $A \subset B$ signifie A est inclus dans B et $A \not\subset B$ signifie A n'est pas inclus dans B , c'est à dire il existe au moins un élément de A qui n'appartient pas à B . On a donc

$$A \not\subset B \iff \exists x : (x \in A \text{ et } x \notin B).$$

Remarque 2.1.3

1. Soit E un ensemble, alors x est un élément de E si et seulement si le singleton $\{x\}$ est inclus dans E ,

$$x \in E \iff \{x\} \subset E.$$

2. Il est important de ne pas confondre appartenance (symbole $\in$) et inclusion (symbole $\subset$). Par exemple, $a \in \{a, 2, 4\}$ mais $a \not\subset \{a, 2, 4\}$.
3. On emploie aussi la notation $B \supset A$ au lieu de $A \subset B$, et l'on dit alors parfois que B est un sur-ensemble de A .
4. Si A et B sont deux ensembles finis et si $A \subset B$ alors $\text{Card}(A) \leq \text{Card}(B)$.

Exemple 2.1.7

1. Soient $A = \{1, 2\}$ et $B = \{1, 2, 3\}$, on a

$$A \subset B \iff \forall x, (x \in \{1, 2\} \implies x \in \{1, 2, 3\})$$

Cela signifie que pour chaque élément de A , cet élément appartient aussi à B .

2. $\mathbb{N} \subset \mathbb{Z} \subset \mathbb{Q} \subset \mathbb{R} \subset \mathbb{C}$.

3. On a $\{1, 2, 4\} \subset \{1, 2, 3, 4\}$ et $\{1, 2, 4\} \not\subset \{1, 2, 3\}$.

4. $\mathbb{Q} \not\subset \mathbb{Z}$ car $\frac{1}{3} \in \mathbb{Q}$ et $\frac{1}{3} \notin \mathbb{Z}$.

Corollaire 2.1.1

1. L'ensemble vide ϕ est inclus dans tout ensemble, y compris lui-même. Pour tout ensemble A , on a $\phi \subset A$.

2. Un ensemble est toujours inclus dans lui-même. Cela signifie que, pour tout ensemble A , on a toujours $A \subset A$.

2.1.2.2 Égalité d'ensembles

Une autre relation fondamentale entre les ensembles est en effet l'égalité d'ensembles.

Définition 2.1.8 (*Egalité de deux ensembles*)

Soient E et F deux ensembles.

1. On dit que E et F sont égaux (ou identiques) et on note $E = F$, s'ils ont les mêmes éléments. Cela signifie que tout élément de E est élément de F et tout élément de F est élément de E . On a alors

$$E = F \iff \forall x, (x \in E \iff x \in F).$$

En d'autres termes,

$$E = F \iff (E \subset F \text{ et } F \subset E). \quad (\text{Principe de double-inclusion}).$$

2. Dans le cas contraire, lorsque les ensembles E et F ne contiennent pas exactement les mêmes éléments, on dit qu'ils sont distincts et on note $E \neq F$. Cela se traduit par

$$E \neq F \iff \exists x, (x \in E \text{ et } x \notin F) \text{ ou } \exists x, (x \notin E \text{ et } x \in F).$$

3. L'égalité $E = F$ se lit « E est égal à F » ou « E est identique à F ».

4. La notation $E \neq F$ se lit « E est différent de F » ou « E est distinct de F ».

Remarque 2.1.4 Pour montrer que deux ensembles E et F sont égaux, on procède souvent par double inclusion, de la manière suivante : on montre que tout élément de E est élément de F (première inclusion) puis que tout élément de F est élément de E (deuxième inclusion).

Exemple 2.1.8

1. $E = \{0, 1, 2, 3, 4\} \neq \{n \in \mathbb{N} : n \text{ divise } 4\}$.

2. Si $E = \{0, 2, 4, 6, 8, 10\}$ et $F = \{6, 8, 10, 0, 2, 4, 6, 8\}$, alors on a $E = F$.

3. Si $E = \{x \in \mathbb{R} : x^2 - 3x + 2 = 0\}$ et $F = \{1, 2\}$, alors on a $E = F$.

Nous admettrons 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 .

Définition 2.1.9 (*Ensemble des parties d'un ensemble*)

Soit E un ensemble, l'ensemble des parties de E est l'ensemble, noté $\mathcal{P}(E)$, dont les éléments sont tous les sous-ensembles (les parties) de E .

$$\mathcal{P}(E) = \{X : X \subset E\}.$$

Autrement dit $X \in \mathcal{P}(E)$ signifie que $X \subset E$.

Remarquons que les éléments de $\mathcal{P}(E)$ sont des sous-ensembles de E et non pas des éléments de E . De plus, contrairement à l'ensemble E qui peut être vide, l'ensemble $\mathcal{P}(E)$ n'est, lui, jamais vide puisqu'il contient au moins les ensembles ϕ et E ,

$$\phi, E \in \mathcal{P}(E).$$

Exemple 2.1.9

1. Si $E = \{a\}$, alors les parties de E sont l'ensemble vide ϕ et E lui-même. Donc

$$\mathcal{P}(\{a\}) = \{\phi, \{a\}\}.$$

2. Si E a deux éléments, disons $E = \{a, b\}$ avec $a \neq b$, alors

$$\mathcal{P}(\{a, b\}) = \{\phi, \{a\}, \{b\}, \{a, b\}\}.$$

3. Si E a trois éléments, disons $E = \{a, b, c\}$ avec $a \neq b \neq c$, alors

$$\mathcal{P}(\{a, b, c\}) = \{\phi, \{a\}, \{b\}, \{c\}, \{a, b\}, \{a, c\}, \{b, c\}, \{a, b, c\}\}.$$

En comptant les éléments de $\mathcal{P}(E)$, on remarque que l'on a

$$\text{Card}(\mathcal{P}(E)) = 2^3 = 8.$$

Proposition 2.1.1 Si E est un ensemble fini de cardinal n alors l'ensemble $\mathcal{P}(E)$ est fini et

$$\text{Card}(\mathcal{P}(E)) = 2^n.$$

2.1.3 Opérations sur les ensembles

Tout comme on effectue des opérations mathématiques sur des nombres (addition, soustraction, division, multiplication) pour obtenir un nouveau nombre, on peut appliquer certaines opérations sur des ensembles pour en créer de nouveaux. Ces opérations ensemblistes incluent la détermination de leur union, de leur intersection, de leur complémentaire, ainsi que leur différence et leur différence symétrique. Dans cette section, nous allons présenter les opérations ensemblistes les plus fondamentales.

2.1.3.1 Intersection

L'intersection de deux ensembles peut être illustrée par une métaphore : si l'on imagine deux patates qui se superposent, la région commune entre elles représente l'intersection de ces deux ensembles. Les éléments contenus dans cette région sont précisément ceux qui appartiennent aux deux ensembles. L'intersection est l'une des opérations les plus fréquemment utilisées dans la vie quotidienne.

Définition 2.1.10 Soient E un ensemble et A, B deux parties de E .

1. L'intersection de A et B est l'ensemble formé des éléments appartenant à A et à B . L'intersection de A et de B se note $A \cap B$ et se lit A intersection B ou A inter B . On a donc

$$A \cap B = \{x \in E : x \in A \wedge x \in B\}.$$

2. Si $A \cap B = \emptyset$, alors nous dirons que les deux ensembles A et B sont disjoints.

Exemple 2.1.10

1. Si $A = \{1, 2, 3, 7\}$ et $B = \{0, 2, 7\}$, alors $A \cap B = \{2, 7\}$ puisque 2 et 7 sont les deux éléments communs aux deux ensembles.
2. Si $A = \{1, 2, 3\}$ et $B = \{5, 8\}$, alors $A \cap B = \emptyset$ puisque les deux ensembles n'ont aucun élément en commun.
3. Soient $A = \{\text{Dr. Smith}, \text{Dr. Johnson}, \text{Dr. Lee}\}$ (professeurs du département A) $B = \{\text{Dr. Lee}, \text{Dr. Brown}\}$ (professeurs du département B), alors $A \cap B = \{\text{Dr. Lee}\}$.

- On peut visualiser l'intersection de deux ensembles A et B par le diagramme suivant

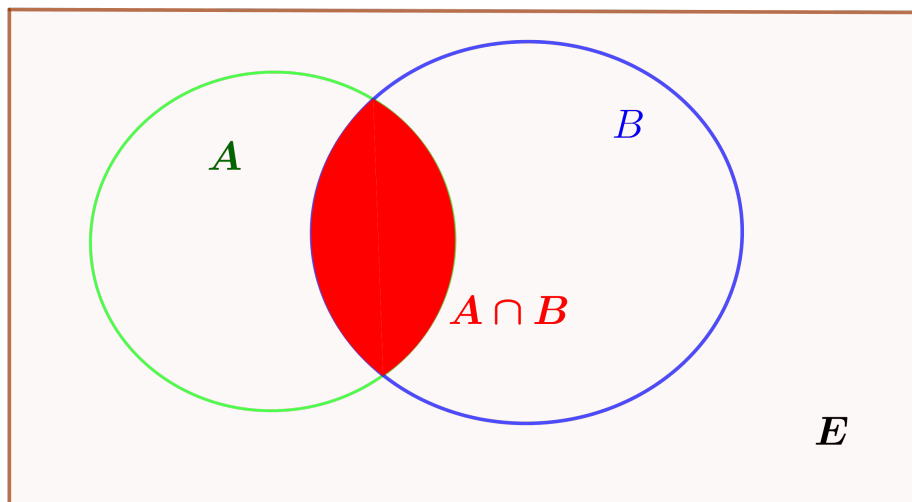


FIGURE 2.1: Intersection de deux ensembles

2.1.3.2 Réunion

L'opération duale de l'intersection est la réunion. Au lieu de former un sous-ensemble contenant uniquement les éléments communs à A et B , l'union crée un sous-ensemble qui comprend tous les éléments présents dans A ou dans B . L'union ensembliste est une opération courante dans la vie quotidienne, utilisée pour combiner des ensembles d'éléments.

Définition 2.1.11 Soient E un ensemble et A, B deux sous-ensembles de E . La réunion de A et B , notée $A \cup B$ (lire A union de B), est l'ensemble constitué par les éléments de E appartenant à A ou à B , c'est-à-dire

$$A \cup B = \{x \in E : x \in A \vee x \in B\}.$$

Remarque 2.1.5 Il est évident que A et B sont des sous-ensembles de $A \cup B$, c'est-à-dire que $A \subset A \cup B$ et $B \subset A \cup B$. De plus, si $A \subset B$, alors $A \cup B = B$. On vérifie que

$$A \cup \emptyset = A, A \cup A = A, A \cup E = E.$$

Exemple 2.1.11 Si $A = \{1, 3, 5\}$ et $B = \{2, 3, 4\}$, alors $A \cup B = \{1, 2, 3, 4, 5\}$. On notera que l'élément 3 qui se trouve dans A et dans B n'est pas répété dans $A \cup B$.

Définition 2.1.12 (Réunion disjointe)

Soient E un ensemble et A, B deux sous-ensembles de E . On dit que la réunion $A \cup B$ de A et B est une réunion disjointe lorsque l'intersection des deux ensembles est vide.

- On peut visualiser l'union de deux ensembles A et B par le diagramme suivant

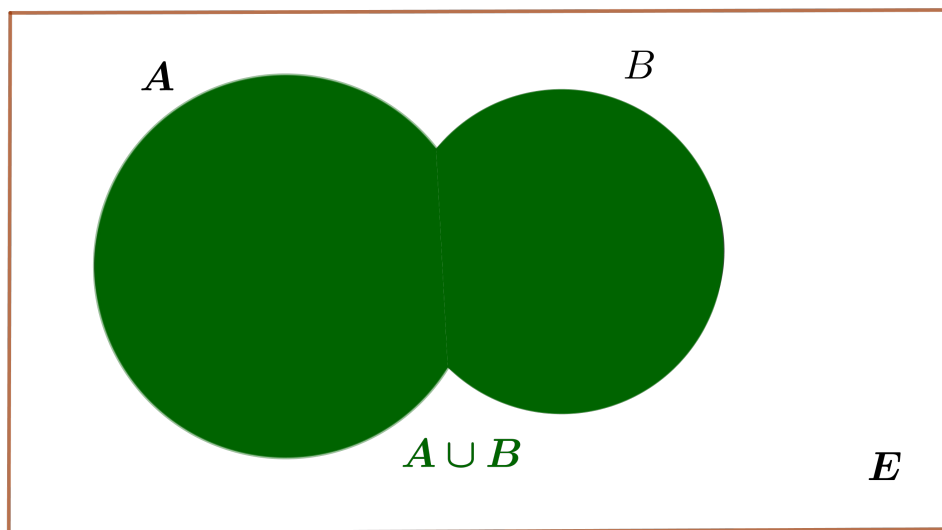


FIGURE 2.2: Réunion de deux ensembles

2.1.3.3 Différence ensembliste

La réunion est comparable à l'addition, car l'ensemble $A \cup B$ contient tous les éléments présents dans A ainsi que ceux qui se trouvent dans B . Mais quelle opération pourrait être considérée comme l'équivalent de la soustraction pour les ensembles ? Pour illustrer, cette opération donnerait, à partir de A et B , l'ensemble des éléments qui sont dans A mais pas dans B . Cette opération est appelée la différence ensembliste. Elle peut être définie de plusieurs manières, toutes équivalentes.

Définition 2.1.13 Soient E un ensemble et A, B deux sous-ensembles de E .

1. La différence de A et B est l'ensemble des éléments qui appartiennent à A mais qui n'appartiennent

pas à B .

2. La différence de A et de B (dans cet ordre) se note $A \setminus B$ ou $A - B$, ce qui se lit « A moins B » ou « A différence B ». On a donc

$$A \setminus B = \{x \in E : x \in A \wedge x \notin B\}.$$

Remarque 2.1.6 La différence entre deux ensembles A et B consiste simplement à prendre les éléments de A et à en retirer les éléments qui sont également dans B .

Exemple 2.1.12

1. Un cas particulier intéressant est celui où on considère une différence $A \setminus B$, mais avec $B \subset A$. $\mathbb{R} \setminus \mathbb{Q}$ désigne donc l'ensemble des nombres réels qui ne sont pas rationnels, c'est-à-dire l'ensemble des irrationnels.

2. Considérons les ensembles $A = \{1, 2, 3, 4, 5\}$ et $B = \{3, 4, 5, 6, 7, 8\}$. Pour trouver la différence $A \setminus B$ de ces deux ensembles, nous commençons par écrire tous les éléments de A , puis enlevons tous les éléments de A qui est aussi un élément de B . Puisque A partage les éléments 3, 4 et 5 avec B , cela nous donne la différence d'ensemble $A \setminus B = \{1, 2\}$.

- On peut visualiser la différence entre deux ensembles A et B par le diagramme suivant

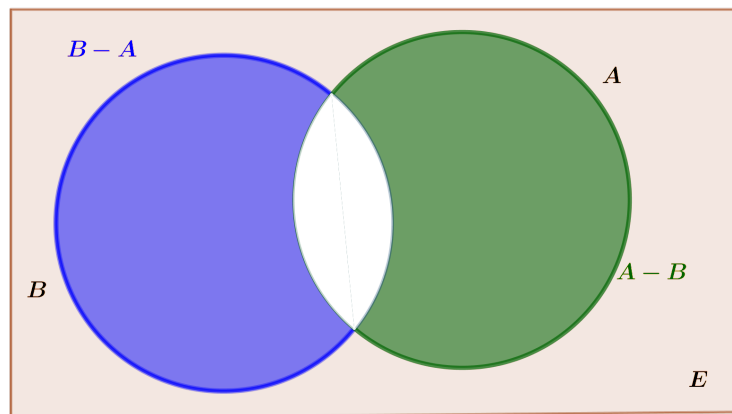


FIGURE 2.3: Différence entre deux ensembles

2.1.3.4 Différence symétrique

La différence symétrique entre deux ensembles est une opération ensembliste qui renvoie un ensemble contenant tous les éléments présents dans l'un des deux ensembles, mais pas dans les deux à la fois. Autrement dit, la différence symétrique entre deux ensembles A et B est constituée des éléments qui appartiennent soit à A , soit à B , mais pas aux deux.

Définition 2.1.14

1. La différence symétrique de deux sous-ensembles A et B d'un ensemble E est l'ensemble formé des éléments de E qui appartiennent soit à A , soit à B , mais pas à leur intersection.

2. La différence symétrique de A et de B (dans cet ordre) se note $A\Delta B$, ce qui se lit A delta B . On a donc

$$A\Delta B = \{x \in E : x \in A \setminus B \vee x \in B \setminus A\}.$$

Autrement dit

$$A\Delta B = (A \setminus B) \cup (B \setminus A).$$

Exemple 2.1.13 Considérons les ensembles $A = \{1, 2, 3, 4, 5\}$ et $B = \{3, 4, 5, 6, 7, 8\}$. Alors la différence symétrique de A et B est l'ensemble

$$A\Delta B = (A \setminus B) \cup (B \setminus A) = \{1, 2\} \cup \{6, 7, 8\} = \{1, 2, 6, 7, 8\}.$$

• On peut visualiser la différence symétrique de deux sous-ensembles A et B par le diagramme suivant

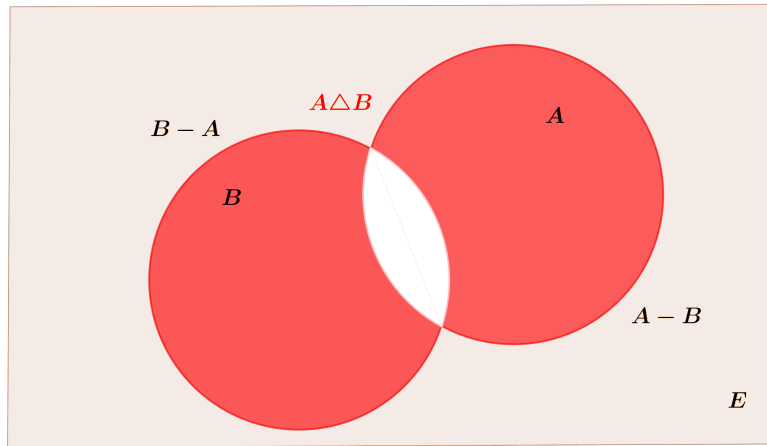


FIGURE 2.4: Différence symétrique de deux ensembles

2.1.3.5 Complémentaire

Les quatre opérations précédentes sont binaires, car elles agissent sur deux ensembles. Nous nous intéresserons maintenant à une opération unaire, qui n'a qu'un seul argument. Le complémentaire d'un ensemble est une opération unaire qui, à partir d'un ensemble A inclus dans un ensemble de référence E , renvoie l'ensemble des éléments qui appartiennent à E mais qui ne sont pas dans A .

Définition 2.1.15 Soit A une partie de l'ensemble E . On appelle complémentaire de A dans E le sous-ensemble de E , constitué des éléments de E qui n'appartiennent pas à A . Cet ensemble est représenté de différentes manières C_E^A , ou encore cA , ou encore $\overline{A}$. On a donc

$$C_E^A = \{x \in E : x \notin A\} = E \setminus A.$$

Remarque 2.1.7 Le complémentaire de A est le plus grand (au sens de l'inclusion) sous-ensemble de E qui ne contient aucun élément en commun avec A .

Exemple 2.1.14

1. Le complémentaire dans l'ensemble $\mathbb{N}$ des entiers naturels du sous-ensemble $\mathbb{P}$ des entiers pairs est l'ensemble des entiers impairs $\mathbb{I}$,

$$C_{\mathbb{N}}^{\mathbb{P}} = \mathbb{I}.$$

2. Si $E = \{a, b, c, d\}$ et $A = \{a, c\}$, alors

$$C_E^A = \{b, d\}.$$

- On peut visualiser le complémentaire de A dans E par le diagramme suivant

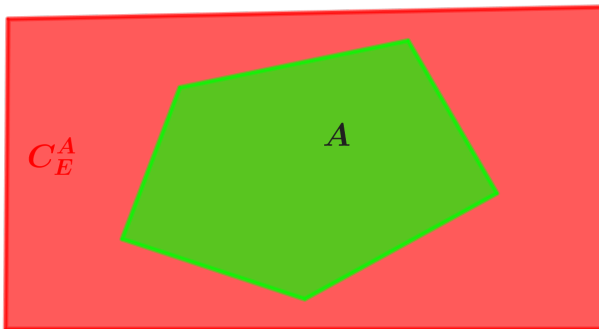


FIGURE 2.5: Complémentaire d'un ensemble

Cela complète l'introduction des cinq opérations de base sur les ensembles, opérations qui sont pratiquement omniprésentes en mathématiques.

2.1.4 Propriétés des opérations élémentaires

Avec la proposition qui suit, on résume les résultats essentiels relatifs à ces opérateurs ensemblistes. Le lecteur pourra établir, à titre d'exercice, les propriétés suivantes.

Proposition 2.1.2 (Règles de calculs)

Soient A, B et C des parties d'un ensemble E , alors

1. $A \cup B = B \cup A$ (Commutativité de la réunion).
2. $A \cap B = B \cap A$ (Commutativité de l'intersection).
3. $(A \cup B) \cup C = A \cup (B \cup C)$ (Associativité de la réunion).
4. $(A \cap B) \cap C = A \cap (B \cap C)$ (Associativité de l'intersection).
5. $A \cup (B \cap C) = (A \cup B) \cap (A \cup C)$ (Distributivité de la réunion sur l'intersection).
6. $A \cap (B \cup C) = (A \cap B) \cup (A \cap C)$ (Distributivité de l'intersection sur la réunion).
7. $A \cup \phi = A$ et $\phi \cup A = A$ (l'ensemble vide est un élément neutre pour la réunion).
8. $A \cap \phi = \phi$ et $\phi \cap A = \phi$ (l'ensemble vide est un élément absorbant pour l'intersection).
9. $A \cap E = A$ et $E \cap A = A$ (l'ensemble E est un élément neutre pour l'intersection).

10. $A \cup E = E$ et $E \cup A = E$ (l'ensemble E est un élément absorbant pour la réunion).
11. $C_E^\phi = E$ et $C_E^E = \phi$ (E et ϕ sont symétriques par complémentation).
12. $A \cap C_E^A = \phi$, $A \cup C_E^A = E$ et $C_E^{C_E^A} = A$ (Double passage au complémentaire).
13. $C_E^{A \cap B} = C_E^A \cup C_E^B$ et $C_E^{A \cup B} = C_E^A \cap C_E^B$ (Lois de De Morgan).
14. $A \setminus (B \cup C) = (A \setminus B) \cap (A \setminus C)$ et $A \setminus (B \cap C) = (A \setminus B) \cup (A \setminus C)$.
15. $A \triangle B = B \triangle A$ (Commutativité de la différence symétrique).
16. $(A \triangle B) \triangle C = A \triangle (B \triangle C)$ (Associativité de la différence symétrique).
17. $A \cap (B \triangle C) = (A \cap B) \triangle (A \cap C)$ (Distributivité de $\cap$ par rapport à $\triangle$).
19. $A \triangle \phi = A$ et $\phi \triangle A = A$ (l'ensemble vide est un élément neutre pour la différence symétrique).
20. $A \triangle B = (A \cup B) \setminus (A \cap B)$ (Caractérisation de la différence symétrique).

Corollaire 2.1.2 Pour A et B ensembles, on a les propriétés suivantes

1. Si $A \subset B$, alors $\mathcal{P}(A) \subset \mathcal{P}(B)$.
2. $\mathcal{P}(A \cap B) = \mathcal{P}(A) \cap \mathcal{P}(B)$.
3. $\mathcal{P}(A) \cup \mathcal{P}(B) \subset \mathcal{P}(A \cup B)$.

Preuve.

1. • Supposons que $A \subset B$. Cela signifie que tout élément de A est aussi un élément de B . Soit $X \in \mathcal{P}(A)$, cela signifie que X est un sous-ensemble de A , donc $X \subset A$. Puisque $A \subset B$, ce qui implique que $X \subset B$. Donc, $X \in \mathcal{P}(B)$. Cela montre que $\mathcal{P}(A) \subset \mathcal{P}(B)$.

2. • Montrons l'inclusion

$$\mathcal{P}(A \cap B) \subset \mathcal{P}(A) \cap \mathcal{P}(B).$$

Soit $X \in \mathcal{P}(A \cap B)$, c'est-à-dire que $X \subset A \cap B$. Cela implique que $X \subset A$ et $X \subset B$, donc $X \in \mathcal{P}(A)$ et $X \in \mathcal{P}(B)$. Ainsi,

$$X \in \mathcal{P}(A) \cap \mathcal{P}(B),$$

donc

$$\mathcal{P}(A \cap B) \subset \mathcal{P}(A) \cap \mathcal{P}(B).$$

• De la même manière montrons l'inclusion

$$\mathcal{P}(A) \cap \mathcal{P}(B) \subset \mathcal{P}(A \cap B).$$

Par conséquent, $\mathcal{P}(A \cap B) = \mathcal{P}(A) \cap \mathcal{P}(B)$.

3. Soit $X \in \mathcal{P}(A) \cup \mathcal{P}(B)$, c'est-à-dire que X est un sous-ensemble de A (donc $X \subset A$) ou de B (donc $X \subset B$).

• Si $X \subset A$, alors $X \subset A \cup B$ (car $A \subset A \cup B$), donc $X \in \mathcal{P}(A \cup B)$.

• Si $X \subset B$, alors de la même manière, $X \subset A \cup B$, puisque car $B \subset A \cup B$, donc $X \in \mathcal{P}(A \cup B)$.

Dans les deux cas, $X \in \mathcal{P}(A \cup B)$, cela montre que

$$\mathcal{P}(A) \cup \mathcal{P}(B) \subset \mathcal{P}(A \cup B).$$

□

Remarque 2.1.8 Remarquons que l'inclusion du point 3 n'est pas en général une égalité, comme on le voit sur le contre-exemple $A = \{1\}$ et $B = \{2\}$ pour lequel

$$\mathcal{P}(A \cup B) = \{\phi, \{1\}, \{2\}, \{1, 2\}\},$$

et

$$\mathcal{P}(A) \cup \mathcal{P}(B) = \{\phi, \{1\}, \{2\}\}.$$

Alors

$$\mathcal{P}(A \cup B) \not\subset \mathcal{P}(A) \cup \mathcal{P}(B).$$

2.1.5 Union et intersection d'une famille de sous-ensembles

Il est possible de généraliser la réunion à un nombre fini d'ensembles : on se ramène au cas de deux ensembles par réunion binaires successives et l'associativité de la réunion assure que l'ordre n'a pas d'importance. De même pour l'intersection. Mais il est possible également de généraliser ces opérations à une famille non nécessairement finie d'ensembles.

Définition 2.1.16 Si $(A_i)_i$ est une famille d'ensembles indexés par $i \in I$ (les éléments de I sont les indices i). Alors

1. L'intersection de la famille $(A_i)_i$ est notée $\bigcap_{i \in I} A_i$. Cette intersection est donc définie explicitement par

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

C'est-à-dire que l'intersection de la famille d'ensembles indexés comprend tous les x qui se trouvent dans chaque ensemble de tous les ensembles de la famille.

2. La réunion de la famille $(A_i)_i$ est notée $\bigcup_{i \in I} A_i$. Cette réunion est donc définie explicitement par

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

C'est-à-dire que la réunion de la famille d'ensembles indexés comprend tous les x pour lesquels il existe un ensemble indexé par i tel que x soit appartient à cet ensemble A_i .

3. Ces deux ensembles se lisent respectivement réunion (ou union) des A_i et intersection des A_i . (On dirait aussi réunion et intersection de la famille $(A_i)_{i \in I}$ ou encore de la famille des A_i).

Remarque 2.1.9

1. La définition 2.1.16 s'applique bien sûr au cas d'un ensemble fini d'indices, $I = \{1, 2, \dots, n\}$. On a alors

$$\bigcap_{i \in I} A_i = \bigcap_{i=1}^n A_i = A_1 \cap A_2 \cap \dots \cap A_n,$$

et

$$\bigcup_{i \in I} A_i = \bigcup_{i=1}^n A_i = A_1 \cup A_2 \cup \dots \cup A_n.$$

2. Dans le cas où les indices parcourent les naturels, l'usage a retenu l'emploi du symbole ∞ pour indiquer que ces indices ne sont pas bornés supérieurement. Étant donné la famille infinie $\{A_1, A_2, \dots\}$, on désignerait par exemple la réunion des ensembles A_n par les notations

$$\bigcup_{n \in \mathbb{N}^*} A_n \text{ ou encore } \bigcup_{n=1}^{\infty} A_n.$$

3. À noter qu'il serait toujours possible de décortiquer de telles opérations finies par étape

$$\begin{aligned} \bigcap_{i=1}^n A_i &= \left(\bigcap_{i=1}^{n-1} A_i \right) \cap A_n. \\ \bigcup_{i=1}^n A_i &= \left(\bigcup_{i=1}^{n-1} A_i \right) \cup A_n. \end{aligned}$$

Exemple 2.1.15

1. Soit $I = \{1, 2, 3, 4\}$ et soit une famille quelconque d'ensembles $(A_i)_{i \in I}$. On a alors

$$\bigcap_{i \in I} A_i = \bigcap_{i=1}^4 A_i = A_1 \cap A_2 \cap A_3 \cap A_4.$$

2. Soit $I = \mathbb{N}^*$ et pour tout $n \in \mathbb{N}^*$, posons $A_n = \{1, 2, 3, \dots, n\} \subset \mathbb{N}^*$. Alors

$$\bigcup_{n \in \mathbb{N}^*} A_n = \bigcup_{n=1}^{\infty} A_n = \mathbb{N}^* \text{ et } \bigcap_{n \in \mathbb{N}^*} A_n = \bigcap_{n=1}^{\infty} A_n = \{1\}.$$

On aurait de même

$$\bigcup_{n \in \mathbb{N}} B_n = \bigcup_{n=0}^{\infty} B_n = \mathbb{N} \text{ et } \bigcap_{n \in \mathbb{N}} B_n = \bigcap_{n=0}^{\infty} B_n = \{0\},$$

avec

$$B_n = \{0, 1, 2, 3, \dots, n\} \subset \mathbb{N}.$$

Un certain nombre de propriétés vues dans le paragraphe précédent se généralisent aux unions et intersections quelconques.

Proposition 2.1.3 (Propriétés liées aux unions et intersections quelconques)

Étant donné $(A_i)_{i \in I}$ une famille d'ensembles, et B un ensemble, tous inclus dans un ensemble E , alors

1. $B \cap \left(\bigcup_{i \in I} A_i \right) = \bigcup_{i \in I} (B \cap A_i)$ (Distributivité)
2. $B \cup \left(\bigcap_{i \in I} A_i \right) = \bigcap_{i \in I} (B \cup A_i)$ (Autre Distributivité)
3. $C_E \left(\bigcup_{i \in I} A_i \right) = \bigcap_{i \in I} (C_E^{A_i})$ (Loi de De Morgan)
4. $C_E \left(\bigcap_{i \in I} A_i \right) = \bigcup_{i \in I} (C_E^{A_i})$ (Loi de De Morgan).

Preuve. Éléments de preuve. Les deux premières égalités découlent de propriétés logiques similaires (ou par double-inclusion), les deux dernières de la négation des quantificateurs et des connecteurs logiques. $\square$

2.1.6 Partition d'un ensemble

Une partition d'un ensemble est une décomposition de cet ensemble en sous-ensembles non vides et disjoints, tels que l'union de ces sous-ensembles reconstitue l'ensemble d'origine.

Définition 2.1.17 Soit E , un ensemble non vide. Une partition de E est une famille $\mathcal{F} = (A_i)_{i \in I}$ de sous-ensembles de E satisfaisant aux trois conditions suivantes

1. Pour tout $i \in I$, on a $A_i \neq \emptyset$.
2. Pour tout $i, j \in I$ tels que $i \neq j$, on a $A_i \cap A_j = \emptyset$.
3. $\bigcup_{i \in I} A_i = E$.

Remarque 2.1.10

1. Un ensemble $A_i \in \mathcal{F}$ est appelé un bloc de la partition $\mathcal{F}$.
2. On exprimera habituellement :
 - La condition 1. en disant que les blocs A_i sont tous non vides.
 - La condition 2. en disant que les blocs A_i sont deux à deux disjoints.
 - La condition 3. en disant que les blocs A_i recouvrent entièrement E .

Exemple 2.1.16

1. $\mathbb{R}_+^*$ et $\mathbb{R}_-^*$ ne constituent pas une partition de $\mathbb{R}$ car l'union $\mathbb{R}_+^* \cup \mathbb{R}_-^*$ ne couvre pas tout $\mathbb{R}$, car $0 \notin \mathbb{R}_+^* \cup \mathbb{R}_-^*$.
2. $\mathbb{R}_+^*$, $\mathbb{R}_-^*$ et $\{0\}$ constituent une partition de $\mathbb{R}$.
4. Soit $E = \{1, 2, 3\}$. Alors
 - il y a une partition de E avec un seul bloc

$$\mathcal{F}_1 = \{A_1 = \{1, 2, 3\}\}$$

- les partitions de E avec deux blocs sont

$$\mathcal{F}_1 = \{A_1 = \{1\}, A_2 = \{2, 3\}\}, \mathcal{F}_2 = \{A_1 = \{2\}, A_2 = \{1, 3\}\}, \mathcal{F}_3 = \{A_1 = \{3\}, A_2 = \{1, 2\}\}.$$

- il y a une partition de E avec trois blocs

$$\mathcal{F} = \{A_1 = \{1\}, A_2 = \{2\}, A_3 = \{3\}\}.$$

5. Voici des exemples de familles qui ne sont pas des partitions.

- $\mathcal{F}_1$ n'est pas une partition de $E = \{1, 2, 3, 4, 5\}$ où

$$\mathcal{F}_1 = \{A_1 = \{1, 2\}, A_2 = \{3, 5\}\}.$$

- $\mathcal{F}_2$ n'est pas une partition de $E = \{1, 2\}$ où

$$\mathcal{F}_2 = \{A_1 = \{1, 2\}, A_2 = \emptyset\}.$$

- $\mathcal{F}_3$ n'est pas une partition de $E = \{1, 2, 3, 4\}$ où

$$\mathcal{F}_3 = \{A_1 = \{1, 2\}, A_2 = \{1, 4\}, A_3 = \{3\}\}.$$

6. Voici deux partitions différentes de $\mathbb{Z}$,

-

$$\mathcal{F}_1 = \{\mathbb{P}, \mathbb{I}\}, \mathbb{P} = \{2n, n \in \mathbb{Z}\}, \mathbb{I} = \{2n + 1, n \in \mathbb{Z}\}.$$

-

$$\mathcal{F}_2 = \{\dots, \{-3\}, \{-2\}, \{-1\}, \{0\}, \{1\}, \{2\}, \{3\}, \dots\}.$$

2.1.7 Produit cartésien

Nous allons maintenant introduire une nouvelle opération binaire sur les ensembles : le produit cartésien. Le produit cartésien est une opération qui prend deux ensembles et crée un nouvel ensemble composé de toutes les paires d'éléments possibles.

2.1.7.1 Couples et n -uplets

Dans la section précédente, nous avons défini l'ensemble $\{a, b\}$, appelé paire, dont les seuls éléments sont a et b . Le couple (a, b) est un nouvel objet qui, à la différence de la paire $\{a, b\}$, prend en compte l'ordre des éléments. Ainsi, (a, b) et (b, a) sont deux couples distincts, tandis que $\{a, b\}$ et $\{b, a\}$ sont identiques. De plus, (a, a) est un couple, alors que $\{a, a\} = \{a\}$ est un singleton. La notion de couple joue ici un rôle clé.

Les couples et les n -uplets sont des concepts fondamentaux en mathématiques et en informatique, utilisés pour organiser des objets dans une séquence ordonnée. Ils généralisent la notion de "paire" (ou "couple") à des ensembles de plus de deux éléments.

Définition 2.1.18 (*Couples*)

1. On appelle couple (ou paire ordonnée) une collection de deux éléments dans un ordre spécifique. Si a et b sont deux éléments, le couple formé par a et b est noté (a, b) , où
 - a est le premier élément (ou la première composante) du couple.
 - b est le deuxième élément (ou la seconde composante) du couple.
2. Le symbole (a, b) se lit couple a, b .

Remarque 2.1.11 La caractéristique fondamentale d'un couple est l'ordre des éléments. Autrement dit, la façon dont les éléments sont disposés dans le couple est cruciale et détermine son identité. Par conséquent, Deux couples sont égaux si et seulement si leurs éléments respectifs sont égaux dans le même ordre. Autrement dit,

$$(a, b) = (c, d) \iff \begin{cases} a = c \\ \text{et} \\ b = d. \end{cases}$$

De plus l'inversion de l'ordre des éléments produit un couple différent, sauf si les deux éléments sont identiques. On a donc $(a, b) \neq (b, a)$ sauf si $a = b$.

Exemple 2.1.17

1. Si $a = 1$ et $b = 2$, alors le couple (a, b) est $(1, 2)$. Dans ce cas, $(1, 2) \neq (2, 1)$, car l'ordre des éléments est différent.

La notion de couple se généralise de manière naturelle à celle de liste ordonnée (finie) de longueur quelconque. On parlera alors de n -uplet pour désigner l'objet résultant de la donnée de n éléments avec un ordre déterminé, avec $n \in \mathbb{N}^*$.

Définition 2.1.19

1. Soit $n \in \mathbb{N}^*$. On appelle n -uplet une collection ordonnée de n éléments avec un ordre déterminé.
2. Si nous notons a_1 le premier élément, a_2 le deuxième élément, ..., a_n le $n^{\text{ème}}$ élément, le n -uplet s'écrit $(a_1, a_2, a_3, \dots, a_n)$. En particulier, un 2-uplet s'appelle un couple et un 3-uplet s'appelle un triplet et un 4-uplet s'appelle un quadruplet.

La notion de produit cartésien joue un rôle majeur dans la suite de ce cours. En effet, cette opération ensembliste fournit le cadre conceptuel pour l'un des concepts mathématiques les plus riches : celui de relation binaire. À partir de la notion de couple, nous pouvons définir le produit cartésien. Pour tous ensembles E et F , nous pouvons créer l'ensemble $E \times F$, qui est constitué de tous les couples (a, b) tels que le premier élément a appartient à E et le second élément b appartient à F .

Définition 2.1.20 Soient E, F deux ensembles.

1. On appelle produit cartésien de E et F l'ensemble formé des couples (x, y) dont la première composante x appartient à E , et la seconde y , à F .
2. Le produit cartésien de E et de F se note $E \times F$ et se lit E fois F , ou E croix F , ou E produit cartésien F . On a donc

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

3. Si $E = F$, on note cet ensemble E^2 .

Remarque 2.1.12

- En pratique, au lieu d'écrire $\forall (x, y) \in E^2, \dots$ on peut noter $\forall x, y \in E, \dots$
- Il est à noter que les couples sont ordonnés, c'est-à-dire que $(x, y) = (y, x)$ si et seulement si $x = y$. De manière plus générale, on a $(x, y) = (x', y')$ dans $E \times F$ si, et seulement si $x = x'$ et $y = y'$.
- En général, $E \times F \neq F \times E$ car les couples (x, y) et (y, x) ne sont pas les mêmes..

Exemple 2.1.18

1. Soit $E = \{1, 2\}$ et $F = \{2, 3\}$, alors

$$E \times F = \{(1, 2), (1, 3), (2, 2), (2, 3)\},$$

et

$$F \times E = \{(2, 1), (2, 2), (3, 1), (3, 2)\},$$

et

$$E^2 = E \times E = \{(1, 1), (1, 2), (2, 1), (2, 2)\}.$$

On voit donc que le produit cartésien n'est pas une opération commutative.

2. Pour $E = \{1, 2\}$ et $F = \{x, y\}$, alors

$$E \times F = \{(1, x), (1, y), (2, x), (2, y)\}.$$

Sous les mêmes hypothèses, on a

$$F \times E = \{(x, 1), (x, 2), (y, 1), (y, 2)\}.$$

- 3.

$$\mathbb{R}^2 = \mathbb{R} \times \mathbb{R} = \{(x, y) : x \in \mathbb{R} \text{ et } y \in \mathbb{R}\}.$$

3. Soit $E = [0, 1] \subset \mathbb{R}$, l'intervalle unité sur la droite réelle. Alors $E^2 = E \times E$ est le carré unité du plan cartésien.

Nous voulons maintenant envisager le produit cartésien de plusieurs ensembles. La notion de produit cartésien se généralise au cas de plus de deux ensembles.

Définition 2.1.21 Soient $E_1, E_2, \dots, E_n$ des ensembles non vides.

1. On appelle produit cartésien des ensembles $E_1, E_2, \dots, E_n$ l'ensemble noté $E_1 \times E_2 \times \dots \times E_n$, constitué des n -uplets $(x_1, x_2, \dots, x_n)$ avec $x_i \in E_i$ pour tout $i \in \{1, 2, \dots, n\}$. En d'autres termes,

$$E_1 \times E_2 \times \dots \times E_n = \{(x_1, x_2, \dots, x_n) : x_1 \in E_1, x_2 \in E_2, \dots, x_n \in E_n\}.$$

2. Deux n -uplets $(x_1, x_2, \dots, x_n)$ et $(x'_1, x'_2, \dots, x'_n)$ sont égaux (ou identiques) si

$$x_i = x'_i, \forall i \in \{1, 2, \dots, n\}.$$

On écrit alors

$$(x_1, x_2, \dots, x_n) = (x'_1, x'_2, \dots, x'_n).$$

Exemple 2.1.19 *Exemple avec trois ensembles.*

1. Si

$$A = \{1, 2\}, B = \{x, y\}, C = \{a\}.$$

Alors

$$A \times B \times C = \{(1, x, a), (1, y, a), (2, x, a), (2, y, a)\}.$$

2. Si

$$A = \{1, 2\}, B = \{x, y\}, C = \{a, b\}.$$

Alors

$$A \times B \times C = \{(1, x, a), (1, x, b), (1, y, a), (1, y, b), (2, x, a), (2, x, b), (2, y, a), (2, y, b)\}.$$

Notation 2.1.1 *On convient de la notation suivante*

$$\prod_{i=1}^n E_i = E_1 \times E_2 \times \dots \times E_n.$$

En particulier, pour tout $n \in \mathbb{N}^*$, on note

$$E^n = \underbrace{E \times E \times \dots \times E}_{n \text{ fois}}.$$

Corollaire 2.1.3 1. Soient E et F deux ensembles finis. Leur produit cartésien $E \times F$ est fini avec

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

2. Plus généralement, soient $E_1, \dots, E_n$ des ensembles finis. Leur produit cartésien est fini avec

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

En notation condensée

$$\text{Card}\left(\prod_{i=1}^n E_i\right) = \prod_{i=1}^n \text{Card}(E_i).$$

On pourrait s'intéresser à diverses propriétés du produit cartésien en lien avec les opérations ensemblistes préalablement introduites. La proposition suivante est immédiate.

Proposition 2.1.4 (Propriétés du produit cartésien)

Soient E, F, G des ensembles, alors

1. $E \times F = \phi \iff E = \phi \text{ ou } F = \phi.$
2. $(E \cup F) \times G = (E \times G) \cup (F \times G).$
3. $(E \cap F) \times G = (E \times G) \cap (F \times G).$
4. $(A \setminus B) \times C = (A \times C) \setminus (B \times C).$

2.2 Relations Binaires

Après les notions d'ensemble introduites dans les sections précédentes, cette section aborde les notions de relations sur un ensemble. Nous y présentons deux types de relations fondamentales, les relations d'équivalence qui permettent dans un ensemble de regrouper des éléments qui ont une même propriété, et les relations d'ordre qui servent à comparer des éléments au sein d'un ensemble.

2.2.1 Généralités

Une relation binaire est une règle qui associe un élément d'un ensemble donné à un élément d'un autre ensemble ou du même ensemble. Formellement elle est définie comme un sous-ensemble du produit cartésien de deux ensembles. Elles sont utilisées pour définir des concepts importants comme les relations d'équivalence et les relations d'ordre, qui permettent de regrouper ou de comparer les éléments dans un ensemble selon certaines propriétés ou règles.

Définition 2.2.1

1. On appelle relation binaire d'un ensemble E vers un ensemble F toute correspondance (ou propriété) notée $\mathfrak{R}$, qui lie des éléments de E à des éléments de F .
2. Si $x \in E$ est lié à $y \in F$ par la relation $\mathfrak{R}$, on dit que x est en relation avec y , ou (x, y) vérifie la relation $\mathfrak{R}$ et on écrit $x\mathfrak{R}y$. Dans le cas contraire on écrit $x\not\mathfrak{R}y$.
3. E s'appelle l'ensemble de départ (ou la source) de $\mathfrak{R}$, F s'appelle l'ensemble d'arrivée (ou le but) de la relation $\mathfrak{R}$, et le sous-ensemble Γ de $E \times F$ s'appelle le graphe de $\mathfrak{R}$ avec

$$\Gamma = \{(x, y) \in E \times F : x\mathfrak{R}y\}.$$

Remarque 2.2.1

1. Un cas particulièrement important de relation binaire est obtenu lorsque les deux ensembles E et F sont les mêmes, $E = F$. Alors on dit que $\mathfrak{R}$ est une relation binaire sur E . La plupart des relations avec lesquelles nous travaillerons dans ce cours seront de ce type.
2. Dans une relation binaire, un élément de l'ensemble de départ peut être en relation :
 - Soit avec plusieurs éléments de l'ensemble d'arrivée.
 - Soit avec un seul élément de l'ensemble d'arrivée.
 - Soit avec aucun seul élément de l'ensemble d'arrivée.

Exemple 2.2.1 Voici quelques exemples fondamentaux de relations.

1. Les relations $<, >, \leq, \geq$ (inférieur à, supérieur à, inférieur ou égal à, supérieur ou égal à) définissent des relations binaires sur $\mathbb{N}, \mathbb{Z}, \mathbb{Q}, \mathbb{R}$.
2. Pour tout ensemble E , la relation d'**égalité** sur E notée $=$ est une relation binaire sur E .

$$\forall x, y \in E : x\mathfrak{R}y \iff x = y.$$

3. Si X est un ensemble, alors l'**inclusion** notée $\subset$ est une relation binaire sur l'ensemble $E = \mathcal{P}(X)$ des parties de X .

$$\forall A, B \in E : A\mathfrak{R}B \iff A \subset B.$$

4. La **divisibilité** sur $\mathbb{Z}$ notée $|$ est une relation binaire sur $\mathbb{Z}$,

$$\forall x, y \in \mathbb{Z} : x\mathfrak{R}y \iff x \text{ divise } y \iff x | y.$$

5. Soit $E = \mathbb{Z}$ et $n \in \mathbb{N}^*$ fixé. La propriété « $x - y$ **est multiple de n** » est une relation binaire sur $\mathbb{Z}$ appelée la relation de **congruence modulo n** . On écrit $x \equiv y \pmod{n}$ ou encore $x \equiv y[n]$ (se lit x est congru à y modulo n) pour signifier que $x\mathfrak{R}y$. On a donc

$$\forall x, y \in \mathbb{Z} : x\mathfrak{R}y \iff x \equiv y \pmod{n} \iff n | (x - y)$$

$$\iff \exists k \in \mathbb{Z} : x - y = kn.$$

Cet entier y correspond au reste de la division euclidienne de x par n . Alors dire que x est congru à y modulo n équivaut à dire que x et y ont le même reste dans la division euclidienne par n .

6. La relation de **parallélisme** notée $\parallel$ sur les droites du plan ou de l'espace E est une relation binaire sur E ,

$$\forall D_1, D_2 \in E : D_1 \mathcal{R} D_2 \iff D_1 \parallel D_2.$$

7. La relation de l'**orthogonalité** notée $\perp$ est une relation binaire sur l'ensemble des droites du plan ou de l'espace E ,

$$\forall D_1, D_2 \in E : D_1 \mathcal{R} D_2 \iff D_1 \perp D_2.$$

8. Soit la relation $\mathcal{R}$ « **est l'ami de** » et E l'ensemble des étudiants en première année MI d'une université donnée. Dans ce cas, tous les couples $(x, y) \in E^2$ qui satisfont cette relation, (c'est-à-dire tels que $x \mathcal{R} y$) nous indiquent que x est l'ami de y , avec x et y appartenant au même ensemble.

Exemple 2.2.2 La correspondance $\mathcal{R}$ qui lie les chiffres aux voyelles utilisées pour écrire les chiffres en toutes lettres est une relation entre deux ensembles : l'ensemble des chiffres et l'ensemble des voyelles. Cette correspondance peut être vue comme une relation binaire entre ces deux ensembles, où chaque chiffre est mis en relation avec les voyelles qui apparaissent lorsqu'il est écrit en toutes lettres. Soient $E = \{0, 1, 2, 3\}$ et $F = \{a, e, i, o, u, y\}$. On a donc

$$0\mathcal{R}e, 0\mathcal{R}o, 1\mathcal{R}u, 2\mathcal{R}e, 2\mathcal{R}u, 3\mathcal{R}o, 3\mathcal{R}i, 0\mathcal{R}a, 1\mathcal{R}y, 2\mathcal{R}o, 3\mathcal{R}u.$$

Définition 2.2.2 Soit $\mathcal{R}$ une relation d'un ensemble E vers un ensemble F .

1. On appelle domaine de $\mathcal{R}$ (ou ensemble de définition de $\mathcal{R}$) et on note $Dom(\mathcal{R})$ l'ensemble des éléments de E qui sont en relation avec (au moins) des éléments de F , c'est-à-dire

$$Dom(\mathcal{R}) = \{x \in E, \exists y \in F : x \mathcal{R} y\} \subset E.$$

2. On appelle image de $\mathcal{R}$ (ou ensemble des valeurs de $\mathcal{R}$) et on note $Im(\mathcal{R})$ l'ensemble des éléments (valeurs) de F mis en relation avec le domaine de définition, c'est-à-dire l'ensemble

$$Im(\mathcal{R}) = \{y \in F, \exists x \in E : x \mathcal{R} y\} \subset F.$$

Exemple 2.2.3 Reprenons la relation $\mathcal{R}$ donnée par l'exemple 2.2.2 précédent, alors l'ensemble de définition est

$$Dom(\mathcal{R}) = \{0, 1, 2, 3\},$$

l'ensemble de valeurs est

$$Im(\mathcal{R}) = \{e, i, o, u\},$$

le graphe est

$$\Gamma = \{(x, y) \in E \times F : x \mathcal{R} y\} = \{(0, e), (0, o), (1, u), (2, e), (2, u), (3, o), (3, i)\}.$$

2.2.2 Représentation d'une relation binaire

La représentation d'une relation $\mathcal{R}$ d'un ensemble E vers un ensemble F dépend de la nature des ensembles E et F . Les représentations les plus courantes sont les suivantes :

1. Représentation par un tableau, une matrice binaire ou un diagramme sagittal/cartésien (utile lorsque E et F sont finis) :

- Matrice : Dans cette représentation, une matrice est utilisée pour indiquer les connexions entre les éléments des deux ensembles. Voici comment cela fonctionne :

- Les éléments de l'ensemble E sont placés sur les lignes.
- Les éléments de l'ensemble F sont placés sur les colonnes.
- Si un élément $x \in E$ est en relation avec un élément $y \in F$, alors on place 1 à l'intersection de la ligne de x et de la colonne de y .
- Si x n'est pas en relation avec y , on place un 0 à cette position.
- Tableau : Dans la représentation par un tableau, les 1 et 0 peuvent être remplacés par des cases noires et blanches pour indiquer l'existence d'une relation. Une case noire signifie qu'il existe une relation entre $x \in E$ et $y \in F$, tandis qu'une case blanche indique l'absence de relation.
- Diagramme sagittal : On dessine les éléments des ensembles E et F sous forme de points, et on trace une flèche partant d'un élément de E vers un élément de F pour indiquer la relation (réseau de flèches).
- Diagramme cartésien : Les éléments des ensembles E et F sont placés sur deux axes perpendiculaires : l'axe des abscisses (horizontal) pour E et l'axe des ordonnées (vertical) pour F . Les relations entre les éléments de E et F sont représentées par des points dans le plan cartésien. Par exemple, si un élément $x \in E$ est en relation avec un élément $y \in F$, un point est placé à l'intersection correspondante dans le plan.

2. Représentation par une formule : Cette approche implique de décrire la relation $\mathfrak{R}$ entre deux ensembles E et F à l'aide d'une expression mathématique ou logique. Cette formule spécifie les conditions sous lesquelles un élément de E est en relation avec un élément de F .

3. Représentation par un graphe : Cette méthode est efficace pour représenter des relations binaires entre les éléments de deux ensembles. Elle utilise des graphes pour illustrer comment les éléments d'un ensemble sont reliés entre eux ou avec ceux d'un autre ensemble.

Exemple 2.2.4

1. Représentation au moyen d'un tableau, d'une matrice ou d'un diagramme. La relation $\mathfrak{R}$ donnée par l'exemple 2.2.2 précédent peut être représentée par les quatre façons suivantes :

	<i>a</i>	<i>e</i>	<i>i</i>	<i>o</i>	<i>u</i>	<i>y</i>
0						
1						
2						
3						

FIGURE 2.6: Représentation au moyen d'un tableau

	<i>a</i>	<i>e</i>	<i>i</i>	<i>o</i>	<i>u</i>	<i>y</i>
0	0	1	0	1	0	0
1	0	0	0	0	1	0
2	0	1	0	0	1	0
3	0	0	1	1	0	0

FIGURE 2.7: Représentation au moyen d'une matrice binaire

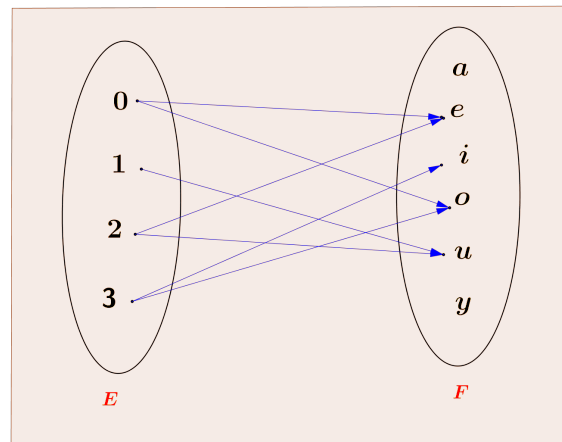


FIGURE 2.8: Représentation au moyen d'un diagramme sagittal

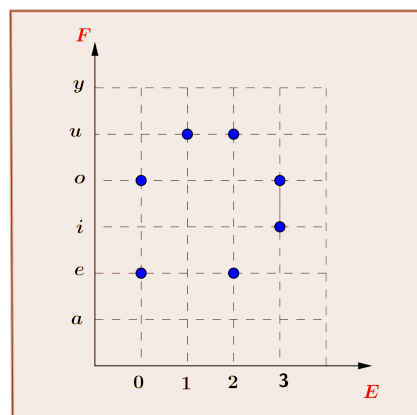


FIGURE 2.9: Représentation au moyen d'un diagramme cartésien

2. Représentation au moyen d'une formule. Soit la relation $\mathfrak{R}$ sur $\mathbb{R}$ telle que

$$\forall x, y \in \mathbb{R} : x \mathfrak{R} y \iff x^2 = y^2.$$

3. Représentation au moyen d'un graphe. La relation $\mathfrak{R}$ donnée par le graphe suivant

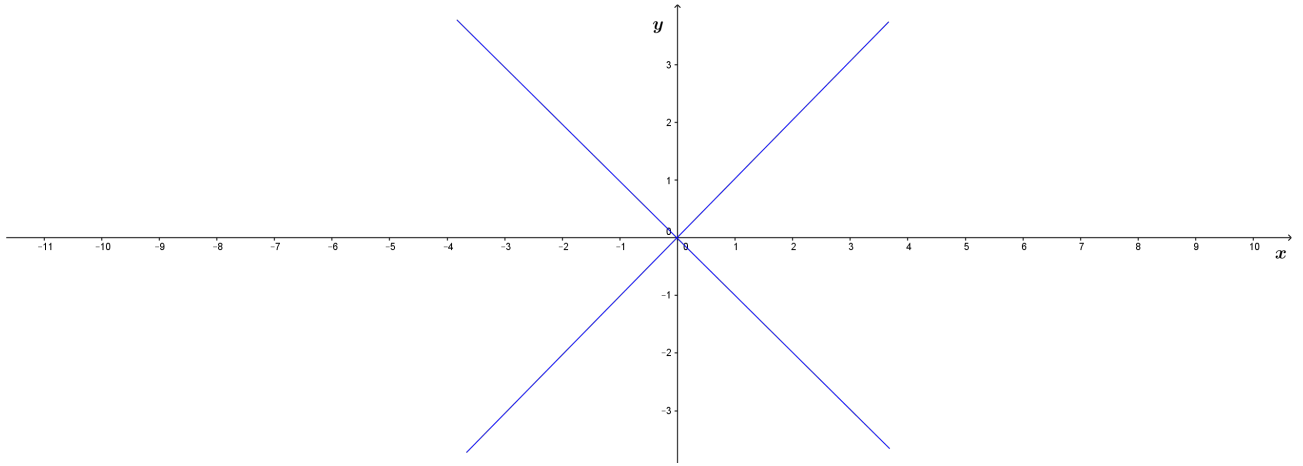


FIGURE 2.10: Représentation au moyen d'un graphe

Définition 2.2.3 Soit $\mathfrak{R}$ une relation binaire d'un ensemble E vers un ensemble F .

1. On appelle relation réciproque de $\mathfrak{R}$ et on note $\mathfrak{R}^{-1}$ la relation binaire de F vers E définie par

$$\forall (x, y) \in E \times F : x\mathfrak{R}y \iff y\mathfrak{R}^{-1}x.$$

2. Autrement dit, $\mathfrak{R}^{-1}$ est la relation où les couples (x, y) de $\Gamma_{\mathfrak{R}}$ sont inversés pour devenir (y, x) . On a donc

$$\forall (x, y) \in E \times F : (x, y) \in \Gamma_{\mathfrak{R}} \iff (y, x) \in \Gamma_{\mathfrak{R}^{-1}}.$$

Exemple 2.2.5 Si on reprend la relation $\mathfrak{R}$ donnée par l'exemple 2.2.2 précédent, on aura

$$\Gamma_{\mathfrak{R}^{-1}} = \{(e, 0), (o, 0), (u, 1), (e, 2), (u, 2), (o, 3), (i, 3)\}.$$

Définition 2.2.4 Soient $\mathfrak{R}$ une relation binaire d'un ensemble E vers un ensemble F et S une relation binaire de l'ensemble F vers un ensemble G .

1. La composée de $\mathfrak{R}$ et S est une relation binaire de l'ensemble E vers l'ensemble G notée $S \circ \mathfrak{R}$ ou $S\mathfrak{R}$ et définie par

$$\forall (x, y) \in E \times G : x(S \circ \mathfrak{R})y \iff \exists z \in F : \begin{cases} x\mathfrak{R}z \\ \text{et} \\ z\mathfrak{R}y. \end{cases}$$

2. Dans le cas particulier où $E = F = G$ on note $\mathfrak{R}^2 = \mathfrak{R} \circ \mathfrak{R}$.

Remarque 2.2.2 Pour que (x, y) appartienne à la relation composée $S \circ \mathfrak{R}$, il doit exister un élément intermédiaire $z \in F$ tel que $(x, z) \in \Gamma_{\mathfrak{R}}$ et $(z, y) \in \Gamma_S$.

Exemple 2.2.6

1. Dans un plan affine euclidien, la relation composée de l'orthogonalité avec elle-même est effectivement le parallélisme. En effet : soient $E = F = G$ l'ensemble des droites d'un plan affine euclidien et $\mathfrak{R} = S = \perp$ est la relation d'orthogonalité. Alors

$$S \circ \mathfrak{R} = \parallel,$$

est la relation de parallélisme, car

$$\forall D_1, D_2 \in E : D_1 \parallel D_2 \iff \exists D_3 \in E : \begin{cases} D_1 \perp D_3 \\ \text{et} \\ D_3 \perp D_2. \end{cases}$$

On peut donc écrire ici

$$\perp \circ \perp = \parallel.$$

En effet, en géométrie, si deux droites sont orthogonales à une même droite, alors elles sont parallèles entre elles.

2. Soient

$$E = \{\alpha, \beta, \gamma, \delta\}, F = \{1, 2, 3, 4, 5\}, G = \{a, b, c, d\}.$$

Telles que

$$\Gamma_{\mathfrak{R}} = \{(\beta, 1), (\beta, 2), (\gamma, 2), (\delta, 4)\},$$

$$\Gamma_S = \{(2, a), (3, b), (4, b), (4, c)\},$$

Alors

$$\Gamma_{S \circ \mathfrak{R}} = \{(\beta, a), (\gamma, a), (\delta, b), (\delta, c)\}.$$

On peut visualiser la composée de deux relations $\mathfrak{R}$ et S par le diagramme suivant

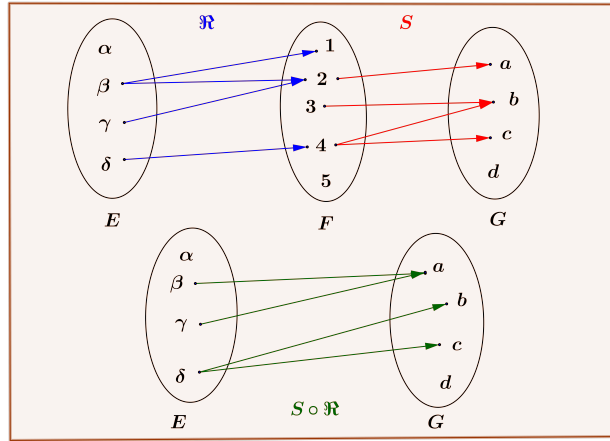


FIGURE 2.11: Composée de deux relations Binaires

Proposition 2.2.1 (Associativité de la composition des relations)

Soient E, F, G, H des ensembles, $\mathfrak{R}_1$ (resp. $\mathfrak{R}_2$, resp $\mathfrak{R}_3$) une relation de E vers F (resp. F vers G , resp. G vers H). On a alors

$$(\mathfrak{R}_3 \circ \mathfrak{R}_2) \circ \mathfrak{R}_1 = \mathfrak{R}_3 \circ (\mathfrak{R}_2 \circ \mathfrak{R}_1).$$

Autrement dit, la composition des relations est associative.

Preuve. Soit $(x, t) \in E \times H$. On a

$$\begin{aligned} x (\mathfrak{R}_3 \circ \mathfrak{R}_2) \circ \mathfrak{R}_1 t &\iff \exists y \in F : \begin{cases} x \mathfrak{R}_1 y \\ et \\ y \mathfrak{R}_3 \circ \mathfrak{R}_2 z \end{cases} \iff \exists y \in F, \exists z \in G : \begin{cases} x \mathfrak{R}_1 y \\ et \\ y \mathfrak{R}_2 z \\ et \\ z \mathfrak{R}_3 t \end{cases} \\ &\iff \exists z \in G : \begin{cases} x \mathfrak{R}_2 \circ \mathfrak{R}_1 z \\ et \\ z \mathfrak{R}_3 t \end{cases} \\ &\iff x \mathfrak{R}_3 \circ (\mathfrak{R}_2 \circ \mathfrak{R}_1) t. \end{aligned}$$

Cela montre que $(\mathfrak{R}_3 \circ \mathfrak{R}_2) \circ \mathfrak{R}_1 = \mathfrak{R}_3 \circ (\mathfrak{R}_2 \circ \mathfrak{R}_1)$. □

Proposition 2.2.2

1. **Inverse double d'une relation.** Pour toute relation binaire $\mathfrak{R}$, on a

$$(\mathfrak{R}^{-1})^{-1} = \mathfrak{R}.$$

2. **Inverse de la composition des relations.** Soient E, F, G des ensembles, $\mathfrak{R}$ (resp. S) une relation de E vers F (resp. F vers G). On a alors

$$(S \circ \mathfrak{R})^{-1} = \mathfrak{R}^{-1} \circ S^{-1}.$$

Preuve.

1. D'après la définition de $\mathfrak{R}^{-1}$, nous savons que

$$\forall (x, y) \in E \times F : (y, x) \in \Gamma_{\mathfrak{R}^{-1}} \iff (x, y) \in \Gamma_{\mathfrak{R}}.$$

Donc

$$\forall (x, y) \in E \times F : (x, y) \in \Gamma_{(\mathfrak{R}^{-1})^{-1}} \iff (y, x) \in \Gamma_{\mathfrak{R}},$$

ce qui équivaut à $(x, y) \in \Gamma_{\mathfrak{R}}$. Donc $(\mathfrak{R}^{-1})^{-1} = \mathfrak{R}$.

2. Pour tout $(x, z) \in E \times G$, on a

$$\begin{aligned} z(S \circ \mathfrak{R})^{-1}x &\iff x(S \circ \mathfrak{R})z \iff \exists y \in F : \begin{cases} x\mathfrak{R}y \\ \text{et} \\ ySz \end{cases} \\ &\iff \exists y \in F : \begin{cases} y\mathfrak{R}^{-1}x \\ \text{et} \\ zS^{-1}y \end{cases} \\ &\iff \exists y \in F : \begin{cases} zS^{-1}y \\ \text{et} \\ y\mathfrak{R}^{-1}x \end{cases} \\ &\iff z(\mathfrak{R}^{-1} \circ S^{-1})x. \end{aligned}$$

Ainsi, on obtient

□

2.2.3 Propriétés d'une relation binaire

On s'intéresse maintenant à une relation binaire dont la source coïncide avec le but. On se retrouve donc avec une relation sur un ensemble E donné. On définit maintenant un certain nombre de propriétés susceptibles d'être satisfaites par une relation d'un ensemble E dans lui-même.

Définition 2.2.5 Soit $\mathfrak{R}$ une relation binaire sur un ensemble E . On dit que $\mathfrak{R}$ est

1. **Réflexive** si tout élément de E est en relation avec lui-même. Formellement,

$$\forall x \in E : x\mathfrak{R}x.$$

2. **Symétrique** si x est lié à y entraîne que y est lié à x . Formellement,

$$\forall x, y \in E : x\mathfrak{R}y \implies y\mathfrak{R}x.$$

Cela signifie que la relation ne change pas si les deux éléments du couple sont inversés.

3. Transitive si x est lié à y et y à z alors x est lié à z . Formellement,

$$\forall x, y, z \in E : \begin{cases} x\mathcal{R}y \\ \text{et} \\ y\mathcal{R}z \end{cases} \implies x\mathcal{R}z.$$

Cela signifie que la relation "transmet" la connexion à travers un élément intermédiaire.

Remarque 2.2.3

1. On voit que la réflexivité exprime le fait que tout élément est en relation avec lui-même.
2. La symétrie revient au fait que la présence d'un couple donné dans la relation implique que le couple « symétrique », obtenu en intervertissant les deux composantes, est lui aussi dans la relation.
3. La transitivité permet de « sauter par-dessus » un élément mitoyen, pour ainsi dire, de manière à établir une relation entre deux éléments reliés à celui-ci.

Sur un diagramme sagittal à une seule courbe, ces propriétés se manifestent respectivement comme suit :

- La réflexivité d'une relation est représentée par la présence d'une boucle à chaque élément de l'ensemble.
- Symétrie : pour chaque flèche « aller » d'un élément x vers un élément y , il y a une flèche « retour ».
- Transitivité : pour chaque paire de flèches consécutives, de x à y puis de y à z , il y a un « raccourci » de x à z . (On a donc une sorte de triangle de flèches, mais avec des directions particulières.)

Une situation nettement plus porteuse concerne la propriété de symétrie. Dans le cas d'une relation non symétrique, il existe (au moins) un couple (x, y) tel que $x\mathcal{R}y$ et $y\not\mathcal{R}x$ " autrement dit, une flèche aller sans flèche retour. Une propriété plus forte serait qu'aucune flèche allée ne s'accompagne d'une flèche retour". Et l'expérience montre que cette situation joue un rôle fondamental. C'est cette idée que cherche à cerner la propriété d'antisymétrie, qui vient compléter notre définition des principales propriétés d'une relation.

Définition 2.2.6 Soit $\mathcal{R}$ une relation binaire sur un ensemble E . On dit que $\mathcal{R}$ est Anti-symétrique si x lié à y et y lié à x entraînent $x = y$. Formellement,

$$\forall x, y \in E : \begin{cases} x\mathcal{R}y \\ \text{et} \\ y\mathcal{R}x \end{cases} \implies x = y.$$

Remarque 2.2.4 Par définition, l'antisymétrie exprime le fait que la présence d'un couple et de son symétrique dans la relation entraîne que les deux composantes sont forcément égales. En d'autres termes, l'existence de deux flèches réciproques entre deux éléments distincts est impossible.

Exemple 2.2.7 (Exemples basiques)

1. La relation donnée par l'égalité sur $\mathbb{R}$,

$$\forall x, y \in \mathbb{R} : x\mathcal{R}y \iff x = y,$$

est réflexive, symétrique, transitive, et antisymétrique.

2. La relation d'**inclusion** $\subset$ dans l'ensemble $E = \mathcal{P}(X)$ des parties de X est réflexive, antisymétrique et transitive mais n'est pas symétrique ($A = \{1\}$, $B = \{1, 2\}$).

3. La relation **de divisibilité** sur l'ensemble $\mathbb{N}^*$ est réflexive, antisymétrique et transitive mais n'est pas symétrique ($x = 2, y = 4$). De plus cette relation est réflexive et transitive mais n'est pas symétrique ($x = 2, y = 4$) et n'est pas antisymétrique ($x = 2, y = -2$) sur l'ensemble $\mathbb{Z}^*$.
4. Pour tout $n \in \mathbb{N}^*$, la relation de **congruence** sur l'ensemble $\mathbb{Z}$,

$$\forall x, y \in \mathbb{Z} : x \mathcal{R} y \iff x \equiv y \pmod{n} \iff n \mid (x - y),$$

est réflexive, symétrique et transitive mais n'est pas antisymétrique ($n = 2, x = 8, y = -8$).

5. La relation d'**inégalité** $\leq$ dans les ensembles $\mathbb{N}, \mathbb{Z}, \mathbb{Q}$ et $\mathbb{R}$ est réflexive, antisymétrique et transitive mais n'est pas symétrique ($x = 2, y = 3$).
6. La relation « **être parallèle** » sur les droites du plan E est réflexive, symétrique et transitive mais n'est pas antisymétrique.
7. La relation « **être perpendiculaire** » (Ou relation d'orthogonalité « **est orthogonal à** ») est symétrique mais n'est pas réflexive, n'est pas transitive et n'est pas antisymétrique.
8. La relation « **être du même âge** » dans l'ensemble des personnes E d'une ville,

$$\forall x, y \in E : x \mathcal{R} y \iff \text{âge}(x) = \text{âge}(y),$$

est réflexive, symétrique et transitive mais n'est pas antisymétrique.

2.2.4 Relations d'équivalence

Les quatre propriétés introduites aux définitions précédentes jouent un rôle fondamental dans de nombreuses situations mathématiques. C'est en effet sur ces quatre propriétés que reposent les concepts-clés de relation d'équivalence (relation à la fois réflexive, symétrique et transitive) et de relation d'ordre (relation à la fois réflexive, transitive et antisymétrique), dont il sera maintenant question.

Définition 2.2.7 Soit $\mathcal{R}$ une relation binaire sur un ensemble E .

1. On dit que $\mathcal{R}$ est une relation d'équivalence si elle vérifie les trois propriétés suivantes

(a) **Réflexivité** :

$$\forall x \in E : x \mathcal{R} x.$$

(b) **Symétrie** :

$$\forall x, y \in E : x \mathcal{R} y \implies y \mathcal{R} x.$$

(c) **Transitivité** :

$$\forall x, y, z \in E : \begin{cases} x \mathcal{R} y \\ \text{et} \\ y \mathcal{R} z \end{cases} \implies x \mathcal{R} z.$$

2. Etant donné x, y tels que $x \mathcal{R} y$, on dit alors que x est équivalent à y pour la relation d'équivalence $\mathcal{R}$ ou x et y sont équivalents pour la relation $\mathcal{R}$ ou encore x et y sont équivalents modulo $\mathcal{R}$.

Exemple 2.2.8

1. La relation d'**égalité** dans un ensemble quelconque E est une relation d'équivalence.
2. La relation $\mathcal{R}$ donnée sur $\mathbb{R}$ par la formule

$$\forall x, y \in \mathbb{R} : x \mathcal{R} y \iff x^2 = y^2,$$

est une relation d'équivalence.

3. La relation « **être parallèle** » est une relation d'équivalence pour l'ensemble E des droites affines

du plan.

4. La relation «**être du même âge**» dans l'ensemble des personnes E est une relation d'équivalence.
5. Pour tout $n \in \mathbb{N}^*$, la relation de **congruence modulo** n est une relation d'équivalence sur l'ensemble $\mathbb{Z}$.
6. La relation $\mathcal{R}$ définie sur $\mathbb{Z} \times \mathbb{Z}^*$ par

$$\forall (x, y), (z, t) \in \mathbb{Z} \times \mathbb{Z}^* : (x, y) \mathcal{R} (z, t) \iff xt = yz,$$

est une relation d'équivalence sur l'ensemble $\mathbb{Z} \times \mathbb{Z}^*$.

7. La relation «**être perpendiculaire**» n'est pas une relation d'équivalence pour l'ensemble E des droites affines du plan (ni la réflexivité, ni la transitivité ne sont vérifiées).
8. La relation de **divisibilité** sur l'ensemble $\mathbb{N}^*$ n'est pas une relation d'équivalence (elle n'est pas symétrique).
9. La relation d'**inégalité** $\leq$ dans les ensembles $\mathbb{N}, \mathbb{Z}, \mathbb{Q}$ et $\mathbb{R}$ n'est pas une relation d'équivalence (elle n'est pas symétrique).

2.2.4.1 Classe d'équivalence et ensemble quotient

Plus généralement, une relation d'équivalence $\mathcal{R}$ sur un ensemble E nous permet de regrouper les éléments de E en paquets, ces paquets étant formés par des éléments de E qui sont équivalents entre eux. Nous introduisons maintenant la notion de classe d'équivalence.

Définition 2.2.8 Soit $\mathcal{R}$ une relation d'équivalence dans un ensemble E .

1. Pour tout $x \in E$, on appelle classe d'équivalence de x (modulo $\mathcal{R}$) l'ensemble, noté $\dot{x}$ (ou $\bar{x}$, ou $\hat{x}$, ou $cl(x)$) de tous les éléments de E équivalents à x (selon la relation $\mathcal{R}$). On a donc

$$\dot{x} = \{y \in E : y \mathcal{R} x\}.$$

Autrement dit $\dot{x}$ est la partie de E formée des éléments qui sont en relation avec x , c'est-à-dire

$$y \in \dot{x} \iff x \mathcal{R} y.$$

2. Tout élément t appartenant à $\dot{x}$ est appelé un représentant de la classe $\dot{x}$.

Remarque 2.2.5

- Une classe d'équivalence regroupe les éléments d'un ensemble qui sont "équivalents" selon une relation d'équivalence donnée.
- Notons que, grâce à la propriété de transitivité, nous pouvons affirmer que tous les éléments d'une classe d'équivalence sont équivalents entre eux.
- On a évidemment $\dot{x} \subset E$ et par symétrie de la relation d'équivalence, on a aussi

$$\dot{x} = \{y \in E : x \mathcal{R} y\}.$$

Le résultat suivant fournit des caractérisations utiles de la relation d'équivalence en termes de classes d'équivalence.

Proposition 2.2.3 (Propriétés de base des classes d'équivalence)

Soit $\mathcal{R}$ une relation d'équivalence sur un ensemble E . Alors les classes d'équivalence de E ont les

propriétés de base suivantes :

1. Les classes d'équivalence ne sont jamais vides,

$$\forall x \in E : \dot{x} \neq \phi.$$

2. Deux classes d'équivalence sont soit identiques, soit disjointes

$$(i) \forall x, y \in E : x \mathcal{R} y \iff \dot{x} = \dot{y},$$

et

$$(ii) \forall x, y \in E : x \not\mathcal{R} y \iff \dot{x} \cap \dot{y} = \phi.$$

3. Les classes d'équivalence recouvrent E ,

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

Preuve.

1. Soit $\dot{x}, x \in E$, l'une des classes d'équivalence. Comme $\mathcal{R}$ est réflexive, on a $x \mathcal{R} x$, ce qui donne $x \in \dot{x}$. Donc pour tout $x \in E$,

$$\dot{x} \neq \phi.$$

2. (i) Supposons que $x \mathcal{R} y$, on procède par double inclusion

$$\dot{x} \subset \dot{y} \text{ et } \dot{y} \subset \dot{x}.$$

• Soit $z \in \dot{x}$, alors on a

$$z \mathcal{R} x.$$

En utilisant la transitivité de $\mathcal{R}$, on obtient

$$z \mathcal{R} y,$$

et par suite

$$z \in \dot{y},$$

ce qui montre que

$$\dot{x} \subset \dot{y}.$$

• L'inclusion réciproque $\dot{y} \subset \dot{x}$ se montre de la même manière (ou est immédiate par symétrie). Donc

$$\dot{x} = \dot{y}.$$

Inversement supposons que $\dot{x} = \dot{y}$. On a toujours

$$x \in \dot{x} = \dot{y},$$

donc

$$x \mathcal{R} y.$$

(ii) • Soient $x, y \in E$ tel que $x \not\mathcal{R} y$. Supposons que $\dot{x} \cap \dot{y} \neq \phi$, alors il existe au moins $z \in E$ tel que

$$z \in \dot{x} \cap \dot{y}.$$

Alors

$$z \mathcal{R} x \text{ et } z \mathcal{R} y.$$

Comme la relation est symétrique et transitive, on en déduirait $x\mathfrak{R}y$. Ce qui contredit notre hypothèse que $x\not\mathfrak{R}y$. Par conséquent $\dot{x} \cap \dot{y} = \phi$.

• Réciproquement, soient $x, y \in E$ telle que $\dot{x} \cap \dot{y} = \phi$. Supposons que $x\mathfrak{R}y$, alors $\dot{x} = \dot{y}$ et ainsi

$$\dot{x} \cap \dot{y} = \dot{x} \neq \phi,$$

puisque $\dot{x}$ est non vide. Ce qui contredit notre hypothèse que $\dot{x} \cap \dot{y} = \phi$. Par conséquent $x\not\mathfrak{R}y$.

3. Il découle des observations faites en 1. que

$$\{x\} \subset \dot{x}.$$

Mais par ailleurs on a bien sûr que

$$\dot{x} \subset E.$$

Il s'ensuit alors que

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

et donc

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

Ce qui termine la preuve. □

On a le corollaire utile suivant.

Corollaire 2.2.1 *Soit $\mathfrak{R}$ une relation d'équivalence sur un ensemble E et soit $y \in \dot{x}$, alors*

$$\dot{x} = \dot{y}.$$

Exemple 2.2.9 *Donnons à présent quelques exemples de classe équivalence pris parmi les exemples de relation donnés en (2.2.8).*

1. *Pour la relation d'égalité dans un ensemble quelconque E , la classe d'équivalence d'un élément x de E est réduite au singleton $\{x\}$,*

$$\forall x \in E : \dot{x} = \{y \in E : y\mathfrak{R}x\} = \{y \in E : y = x\} = \{x\}.$$

2. *Pour la relation $\mathfrak{R}$ donnée sur $\mathbb{R}$ par la formule*

$$\forall x, y \in \mathbb{R} : x\mathfrak{R}y \iff x^2 = y^2,$$

la classe d'équivalence d'un élément x de $\mathbb{R}$ est

$$\begin{aligned} \dot{x} &= \{y \in \mathbb{R} : y\mathfrak{R}x\} = \{y \in \mathbb{R} : y^2 = x^2\} = \{y \in \mathbb{R} : (y - x)(y + x) = 0\} \\ &= \begin{cases} \{x\}, & \text{si } x = 0, \\ \{-x, x\}, & \text{si } x \neq 0. \end{cases} \end{aligned}$$

3. *Pour la relation « être du même âge », la classe d'équivalence d'une personne est l'ensemble des personnes E ayant le même âge,*

$$\forall p \in E : \dot{p} = \{q \in E : q\mathfrak{R}p\} = \{q \in E : \text{âge}(q) = \text{âge}(p)\}.$$

4. *Pour la relation « être parallèle », la classe d'équivalence d'une droite est l'ensemble des droites parallèles à cette droite,*

$$\forall X \in E : \dot{X} = \{D \in E : D\mathfrak{R}X\} = \{D \in E : D \parallel X\}.$$

5. Pour la relation de congruence modulo n sur l'ensemble $\mathbb{Z}$, ($n \in \mathbb{N}^*$), on a la classe d'équivalence d'un élément x de $\mathbb{Z}$ est donnée par

$$\dot{x} = \{y \in \mathbb{Z} : y \mathcal{R} x\} = \{y \in \mathbb{Z} : y \equiv x \pmod{n}\} = \{y \in \mathbb{Z} : y = x + kn, k \in \mathbb{Z}\} = x + n\mathbb{Z}.$$

Les classes d'équivalence de cette relation sont

$$\dot{0}, \dot{1}, \dot{2}, \dots, \widehat{\dot{n-1}}.$$

Telles que

$$\blacktriangleright \dot{0} = \{y = kn, k \in \mathbb{Z}\} = n\mathbb{Z}$$

$\forall m \in \dot{0}$, le reste de la division euclidienne de m par n est égal à 0.

0 est un représentant de la classe $\dot{0}$.

$$\blacktriangleright \dot{1} = \{y = kn + 1, k \in \mathbb{Z}\} = n\mathbb{Z} + 1$$

$\forall m \in \dot{1}$, le reste de la division euclidienne de m par n est égal à 1.

1 est un représentant de la classe $\dot{1}$.

$\vdots$

$$\blacktriangleright \widehat{\dot{n-1}} = \{y = kn + (n-1), k \in \mathbb{Z}\} = n\mathbb{Z} + (n-1).$$

$\forall m \in \widehat{\dot{n-1}}$, le reste de la division euclidienne de m par n est égal à $(n-1)$.

$(n-1)$ est un représentant de la classe $\widehat{\dot{n-1}}$.

alors la relation de congruence modulo n sur $\mathbb{Z}$ possède exactement n classes d'équivalence distinctes.

Une relation d'équivalence peut servir à trier les éléments de E selon un critère et nous obtenons des sous-ensembles (les classes d'équivalence) constitués des éléments équivalents vis-à-vis de ce critère. Nous pouvons alors, d'une certaine façon, identifier les éléments équivalents en considérant l'ensemble des classes d'équivalence. Ce nouvel ensemble appelé **ensemble quotient** de E par $\mathcal{R}$. C'est un outil très puissant en algèbre, permettant par exemple de construire de nouveaux ensembles ($\mathbb{Z}$ à partir de $\mathbb{N}$, puis $\mathbb{Q}$ à partir de $\mathbb{Z}$ etc).

Définition 2.2.9 Soit E un ensemble muni d'une relation d'équivalence $\mathcal{R}$. On appelle ensemble quotient de E par la relation d'équivalence $\mathcal{R}$ l'ensemble formé de toutes les classes d'équivalence déterminées par la relation $\mathcal{R}$. On représente cet ensemble par la notation $E/\mathcal{R}$. On a donc

$$E/\mathcal{R} = \{\dot{x} : x \in E\} \subset \mathcal{P}(E).$$

Exemple 2.2.10

1. Pour la relation d'égalité dans un ensemble quelconque E , pour tout x de E , on a

$$E/\mathcal{R} = \{\dot{x} : x \in E\} = \{\{x\} : x \in E\}.$$

2. Pour la relation de congruence modulo n sur l'ensemble $\mathbb{Z}$, ($n \in \mathbb{N}^*$), les classes d'équivalences sont $\dot{0}, \dot{1}, \dot{2}, \dots, \widehat{\dot{n-1}}$. L'ensemble quotient de $\mathbb{Z}$ par la relation de congruence est donné par

$$E/\mathcal{R} = \left\{ \dot{0}, \dot{1}, \dot{2}, \dots, \widehat{\dot{n-1}} \right\}.$$

Ce nouvel ensemble est noté $\mathbb{Z}/n\mathbb{Z}$ (lire $\mathbb{Z}$ sur $n\mathbb{Z}$), ainsi

$$\mathbb{Z}/n\mathbb{Z} = \left\{ \dot{0}, \dot{1}, \dot{2}, \dots, \widehat{\dot{n-1}} \right\}.$$

La famille constituée de ces n sous-ensembles de $\mathbb{Z}$ forme donc une partition de $\mathbb{Z}$. C'est un sous-ensemble fini de $\mathcal{P}(\mathbb{Z})$ contient n éléments.

Proposition 2.2.4 (Partition associée à une relation d'équivalence)

Soit $\mathcal{R}$ une relation d'équivalence sur un ensemble E . L'ensemble des classes d'équivalence distinctes forme une partition de E .

Preuve. Ce résultat découle immédiatement de la proposition (2.2.3). Il s'agit de montrer que toutes les trois conditions de la définition 2.1.17 soient vérifiées.

(i) Chacun des blocs de la partition est non vide.

(ii) Les blocs sont deux à deux disjoints.

(iii) La réunion de tous les blocs donne E . □

On peut visualiser la partition associée à une relation d'équivalence par le diagramme suivant

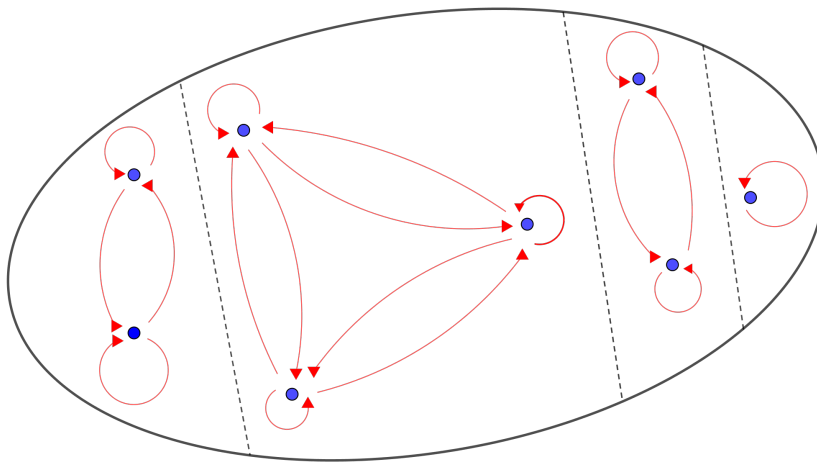


FIGURE 2.12: Représentation graphique d'une relation d'équivalence. Partition en classes d'équivalence.

2.2.5 Relations d'ordre

L'autre type de relation binaire que nous allons explorer est la relation d'ordre. Elle se distingue de la relation d'équivalence par l'une de ses propriétés : au lieu d'être symétrique, elle est antisymétrique. Une relation d'ordre dans un ensemble est une relation binaire qui permet de comparer les éléments de cet ensemble de manière cohérente. Elle organise les éléments selon un certain critère de comparaison, généralement appelé ordre.

Définition 2.2.10

1. Une relation binaire $\mathcal{R}$ sur un ensemble E est une relation d'ordre si et seulement si elle est réflexive, antisymétrique et transitive. Autrement dit, une relation binaire $\mathcal{R}$ sur E est une relation d'ordre lorsque, on a

a. **Réflexivité :**

$$\forall x \in E : x\mathcal{R}x.$$

b. **Antisymétrie :**

$$\forall x, y \in E : \left\{ \begin{array}{l} x\mathcal{R}y \\ \text{et} \\ y\mathcal{R}x \end{array} \right\} \implies x = y.$$

c. **Transitivité** :

$$\forall x, y, z \in E : \begin{cases} x \mathcal{R} y \\ \text{et} \\ y \mathcal{R} z \end{cases} \implies x \mathcal{R} z.$$

On dit alors que E est un ensemble ordonné (par $\mathcal{R}$).

Notation 2.2.1 Une relation d'ordre est souvent notée $\leq$ au lieu de $\mathcal{R}$ et $x \leq y$ se lit " **x est plus petit que y** " ou " **y est plus grand que x** ".

Voici quelques exemples de relation d'ordre.

Exemple 2.2.11

1. Les relations usuelles égalité, supériorité, infériorité ($=, \geq, \leq$) sont des relations d'ordre sur $\mathbb{N}, \mathbb{Z}, \mathbb{Q}$ et $\mathbb{R}$.
2. La relation d'inclusion dans l'ensemble $\mathcal{P}(E)$ des parties de E est une relation d'ordre.
3. La relation de divisibilité dans $\mathbb{N}^*$ est une relation d'ordre dans cet ensemble.
4. La relation $\mathcal{R}$ de puissances définie par

$$\forall x, y \in \mathbb{N}^* : x \mathcal{R} y \iff \exists n \in \mathbb{N}^* : y = x^n,$$

est une relation d'ordre dans $\mathbb{N}^*$. Cette relation peut s'énoncer aussi " y est une puissance entière non nulle de x ".

Définition 2.2.11 (Ensemble ordonné)

Un ensemble E muni d'une relation d'ordre $\mathcal{R}$ est appelé un ensemble ordonné. On note alors $(E, \mathcal{R})$ pour signifier que l'on considère cet ensemble avec la relation d'ordre $\mathcal{R}$.

Exemple 2.2.12

1. $(\mathbb{N}, \mathbb{Z}, \mathbb{Q}, \mathbb{R}, \leq)$ sont des ensembles ordonnés.
2. $(\mathbb{N}^*, |)$ est un ensemble ordonné.
3. $(\mathcal{P}(E), \subset)$ est un ensemble ordonné.

2.2.5.1 Comparaison dans un ensemble ordonné, ordre total et partiel

La comparaison de deux éléments d'un ensemble ordonné vise à déterminer lequel des deux « précède » l'autre en vertu de l'ordre en cause. Mais il arrivera que de telles comparaisons ne soient pas concluantes.

Définition 2.2.12 (Eléments comparables, incomparables)

Soit $\mathcal{R}$ une relation d'ordre sur un ensemble E .

1. Deux éléments x, y de E sont dits comparables par la relation $\mathcal{R}$ si l'une des deux affirmations suivantes est vérifiée

$$x \mathcal{R} y \text{ ou } y \mathcal{R} x.$$

2. Deux élément x, y de E sont dits incomparables par la relation $\mathcal{R}$ si les deux des affirmations suivantes est vérifiée

$$x \mathcal{R} y \text{ et } y \mathcal{R} x.$$

Autrement dit, x et y sont incomparables s'il n'existe aucune relation entre eux, ni dans un sens, ni dans l'autre.

Exemple 2.2.13

1. Dans l'ensemble ordonné $(\mathbb{N}^*, |)$, les éléments 3 et 24 de $\mathbb{N}^*$ sont comparables par la relation de divisibilité,

$$3 \mid 24.$$

tandis que 3 et 25 ne le sont pas car

$$3 \nmid 25 \text{ et } 25 \nmid 3.$$

2. Dans l'ensemble ordonné $(\mathcal{P}(\mathbb{N}), \subset)$, les deux parties de $\mathbb{N}$ suivantes $X_1 = \{2, 7\}$ et $Y_1 = \{1, 2, 4, 7\}$ sont comparables par la relation de l'inclusion

$$X_1 = \{2, 7\} \subset Y_1 = \{1, 2, 4, 7\}.$$

Mais $X_2 = \{5\}$ et $Y_2 = \{0, 4, 11\}$ sont incomparables

$$X_2 = \{5\} \not\subset Y_2 = \{0, 4, 11\} \text{ et } Y_2 = \{0, 4, 11\} \not\subset X_2 = \{5\}.$$

Dans $\mathbb{R}$ muni de l'ordre usuel, on peut comparer deux à deux tous les éléments, mais ce n'est pas toujours le cas (cas de la divisibilité par exemple). Cette remarque motive la définition suivante.

Définition 2.2.13 (Ordre total, ordre partiel)

Soit $\mathcal{R}$ une relation d'ordre sur un ensemble E .

1. On dit que $\mathcal{R}$ est une relation d'ordre total (ou E est totalement ordonné par $\mathcal{R}$) si et seulement si les éléments de E sont tous comparables deux à deux par $\mathcal{R}$, c'est-à-dire

$$\forall x, y \in E : x \mathcal{R} y \text{ ou } y \mathcal{R} x.$$

2. Dans le cas contraire, on dit que $\mathcal{R}$ est une relation d'ordre partiel (on dit aussi que l'ensemble E est partiellement ordonné). Il s'agit donc d'une relation d'ordre telle qu'il existe au moins deux éléments incomparables par $\mathcal{R}$.

Remarque 2.2.6 Un ensemble totalement ordonné est un ensemble ordonné dans lequel deux éléments quelconques sont toujours comparables. Un ensemble partiellement ordonné est un ensemble muni d'une relation d'ordre qui indique que pour certains couples d'éléments, l'un est plus petit que l'autre. Tous les éléments ne sont pas forcément comparables, contrairement au cas d'un ensemble muni d'un ordre total.

Exemple 2.2.14

1. La relation d'inégalité usuelle $\leq$ dans les ensembles $\mathbb{N}, \mathbb{Z}, \mathbb{Q}$ et $\mathbb{R}$, est une relation d'ordre total.
 2. La relation de divisibilité sur l'ensemble $\mathbb{N}^*$ est une relation d'ordre partiel car 3 et 4 ne sont pas comparables

$$3 \nmid 4 \text{ et } 4 \nmid 3.$$

3. La relation $\mathcal{R}$ définie sur l'ensemble $\mathbb{N}^*$ par

$$\forall x, y \in \mathbb{N}^* : x \mathcal{R} y \iff \exists n \in \mathbb{N}^* : y = x^n,$$

est une relation d'ordre partiel car 3 et 4 ne sont pas comparables.

4. La relation d'inclusion $\subset$ dans l'ensemble $\mathcal{P}(E)$ des parties de E est une relation d'ordre partiel car si E contient au moins deux éléments, alors pour tous a et b éléments distincts ($a \neq b$) de E , on a

$$\{a\} \not\subset \{b\} \text{ et } \{b\} \not\subset \{a\}.$$

5. La relation d'inégalité usuelle $\leq$ dans l'ensemble $\mathcal{F}(\mathbb{R}, \mathbb{R})$ des applications de $\mathbb{R}$ vers $\mathbb{R}$ est une relation d'ordre partiel car, par exemple $f = \sin$ et $g = \cos$ ne sont pas comparables,

$$\sin \not\leq \cos \text{ et } \cos \not\leq \sin,$$

telle que

$$\forall f, g \in \mathcal{F}(\mathbb{R}, \mathbb{R}) : f \leq g \iff \forall x \in \mathbb{R} : f(x) \leq g(x).$$

2.2.5.2 Éléments remarquables d'un ensemble ordonné

Nous allons maintenant donner la définition de quelques éléments remarquables d'un ensemble ou d'une partie d'un ensemble ordonné.

2.2.5.2.0.1 Minorants, majorants Dans le cadre d'un ensemble partiellement ordonné E muni d'une relation d'ordre $\mathfrak{R}$, les notions de minorants et majorants sont définies comme suit.

Définition 2.2.14 (*Minorants, majorants*)

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

1. On dit que $x \in E$ est un *minorant* de A (ou x *minore* A) si et seulement si tout élément de A est plus grand que x pour $\leq$,

$$\forall a \in A : x \leq a.$$

2. On dit que $x \in E$ est un *majorant* de A (ou x *majore* A) si et seulement si tout élément de A est plus petit que x pour $\leq$,

$$\forall a \in A : a \leq x.$$

Notation 2.2.2 On peut noter $\text{Maj}_E(A)$ (resp. $\text{Min}_E(A)$) l'ensemble des majorants (resp. minorants) de A dans E .

Donnons maintenant quelques exemples pour illustrer notre propos.

Exemple 2.2.15

1. Si $(E, \leq) = (\mathbb{R}, \leq)$ et si $A =]0, 1[$, alors tout réel $x \geq 1$ est un majorant de A , tout réel $x \leq 0$ est un minorant de A . C'est-à-dire

$$\text{Maj}_{\mathbb{R}}(A) = [1, +\infty[\text{ et } \text{Min}_{\mathbb{R}}(A) =]-\infty, 0].$$

2. Si $(E, \leq) = (\mathbb{N}^*, |)$ et si $A = \{1, 2, 3\}$, alors les majorants de A sont tous les entiers naturels simultanément divisibles par 1, 2, 3 donc divisibles par $\text{PPCM}(1, 2, 3) = 6$ (tel que $\text{PPCM}(1, 2, 3)$ désigne le Plus Petit Commun Multiple des nombres 1, 2, 3). D'autre part, un minorant de A est un entier naturel qui divise 1, 2 et 3. En particulier il doit diviser 1, ce qui n'est le cas que de 1. Le seul minorant de A est donc 1.

$$\text{Min}_{\mathbb{N}^*}(A) = \{1\}.$$

3. Soit $E = \mathcal{P}(\{1, 2\}) = \{\phi, \{1\}, \{2\}, \{1, 2\}\}$ ordonné par la relation de l'inclusion $\subset$ et soit $A = \{\phi, \{1\}, \{2\}\}$ un sous-ensemble de E . Alors $\{1, 2\}$ (resp. ϕ) est un majorant (resp. minorant) de A . En effet,

• On a $\{1, 2\}$ est le plus grand ensemble dans E pour la relation d'inclusion, car tout sous-ensemble de $\{1, 2\}$ est inclus dans $\{1, 2\}$. Il est donc vérifié que

$$\phi \subset \{1, 2\}, \{1\} \subset \{1, 2\}, \{2\} \subset \{1, 2\}.$$

Donc, $\{1, 2\}$ est bien un majorant de A .

• On a ϕ est le plus petit ensemble dans E pour la relation d'inclusion, et on a

$$\phi \subset \phi, \phi \subset \{1\}, \phi \subset \{2\}.$$

Donc, ϕ est un minorant de A .

2.2.5.2.0.2 Parties majorées, minorées, bornées**Définition 2.2.15 (*Parties majorées, minorées, bornées*)**

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

1. On dit que A est minorée dans E si et seulement si A admet au moins un minorant dans E , c'est-à-dire

$$\exists x \in E, \forall a \in A : x \leq a.$$

2. On dit que A est majorée dans E si et seulement si A admet au moins un majorant dans E , c'est-à-dire

$$\exists x \in E, \forall a \in A : a \leq x.$$

3. La partie A est dite bornée dans E si et seulement si A est à la fois minorée et majorée. Autrement dit

$$A \text{ est bornée} \iff \exists x, y \in E, \forall a \in A : x \leq a \leq y.$$

Exemple 2.2.16

1. On munit $\mathbb{R}$ de son ordre usuel $\leq$. La partie

$$A = \left\{ \frac{1}{n}, n \in \mathbb{N}^* \right\},$$

est majorée par 1 et minorée par 0.

2. On munit $\mathcal{P}(\mathbb{R})$ de la relation d'inclusion $\subset$. Soit A et B deux parties de $\mathbb{R}$. On a

$$\begin{aligned} A \cap B &\subset A \subset A \cup B, \\ A \cap B &\subset B \subset A \cup B, \end{aligned}$$

donc $X = \{A, B\}$ est majorée par $A \cup B$ et minorée par $A \cap B$.

3. Reprenons l'exemple avec $E = \mathcal{P}(\{1, 2\}) = \{\emptyset, \{1\}, \{2\}, \{1, 2\}\}$ et $A = \{\emptyset, \{1\}, \{2\}\}$, nous avons vu que $\emptyset$ est un minorant de A . Ainsi, A est minoré dans E puisqu'il admet $\emptyset$ comme minorant et que $\{1, 2\}$ est un majorant de A . Ainsi, A est majorée dans E puisqu'il admet $\{1, 2\}$ comme majorant. Ainsi, la partie A est bornée dans E car elle admet à la fois un minorant $\emptyset$ et un majorant $\{1, 2\}$.

2.2.5.2.0.3 Plus grand élément, plus petit élément Les concepts de plus grand élément et de plus petit élément sont des outils essentiels pour étudier les propriétés des ensembles ordonnés. Ils permettent de caractériser certains types d'ensembles et de simplifier de nombreuses démonstrations.

Définition 2.2.16 (*Plus grand élément, plus petit élément*)

Soit $(E, \leq)$ un ensemble muni d'une relation d'ordre, et soit A un sous-ensemble non vide de E .

1. On dit que $\alpha \in E$ est un plus grand élément de A (ou un maximum de A) noté $\max(A)$ si et seulement si α est un majorant de A , et $\alpha \in A$. C'est-à-dire

$$\alpha = \max(A) \iff \begin{cases} \alpha \in A, \\ \text{et} \\ \forall x \in A : x \leq \alpha. \end{cases}$$

2. On dit que $\beta \in E$ est un plus petit élément de A (ou un minimum de A) noté $\min(A)$ si et seulement si β est un minorant de A , et $\beta \in A$. C'est-à-dire

$$\beta = \min(A) \iff \begin{cases} \beta \in A, \\ \text{et} \\ \forall x \in A : \beta \leq x. \end{cases}$$

3. Un extremum est un élément qui est un minimum ou un maximum.

Remarque 2.2.7

- Un ensemble n'a pas nécessairement de minimum ou de maximum.
- Un majorant de A qui appartient à A est appelé plus grand élément de A .
- Un minorant de A qui appartient à A est appelé plus petit élément de A .
- Il est important de noter que le maximum d'un ensemble doit nécessairement appartenir à cet ensemble. De la même manière, si un ensemble a un minimum, cet élément appartient à l'ensemble.

Exemple 2.2.17

1. Si $(E, \leq) = (\mathbb{R}, \leq)$ (avec la relation d'ordre usuelle), alors
 - Si $A = [0, 1]$, alors $\max(A) = 1$ et $\min(A) = 0$.
 - Si $A =]0, 1]$, alors A n'admet pas de minimum, mais admet 1 comme maximum.
 - Si $A =]0, 1[$, alors A n'admet ni maximum ni minimum.
 - Si $A =]-\infty, 1]$, alors A est majorée, mais non minorée (donc non bornée). L'ensemble de ses majorants est $[1, +\infty[$. Le maximum de A n'existe pas.
 - Si $A = [-2, +\infty[$, alors A est minorée, mais non majorée (donc non bornée). L'ensemble de ses minorants est $]-\infty, -2]$. Le minimum de A est -2 .
2. Si $(E, \leq) = (E, \subset)$, alors le minimum de $\mathcal{P}(E)$ pour la relation d'inclusion est l'ensemble vide ϕ et E est le maximum,

$$\max(\mathcal{P}(E)) = E \text{ et } \min(\mathcal{P}(E)) = \phi.$$

3. Si $(E, \leq) = (\mathbb{N}^*, |)$ (ordonné par la relation de divisibilité $|$) et si $A = \{1, 2, 3\}$, alors A a comme minimum 1, mais elle n'a pas de maximum (car $2 \nmid 3$).

Théorème 2.2.1 (Unicité de minimum, maximum)

Soit $(E, \leq)$ un ensemble ordonné et A une partie de A . Si A possède un plus petit élément ou un plus grand élément, alors chacun de ces éléments est unique.

Preuve. Supposons au contraire que β_1 et β_2 sont deux minimums de A . On a donc que β_1 et β_2 sont tous deux des éléments de A . Comme β_1 est un minimum de A , il s'ensuit par définition que $\beta_1 \leq \beta_2$. Par un raisonnement similaire on a que $\beta_2 \leq \beta_1$. La propriété d'antisymétrie de la relation entraîne alors que $\beta_1 = \beta_2$. Le même raisonnement s'applique pour montrer l'unicité du plus grand élément de A . $\square$

2.2.5.2.0.4 Borne supérieure, Borne inférieure Les notions de borne supérieure et de borne inférieure sont des concepts fondamentaux en théorie des ensembles ordonnés et en analyse mathématique. Elles permettent de généraliser les idées de plus grand et de plus petit élément dans le contexte d'ensembles potentiellement infinis, ou dans des situations où le plus grand ou le plus petit élément n'existe pas au sein de l'ensemble.

Définition 2.2.17 (Borne supérieure, Borne inférieure)

Soit $(E, \leq)$ un ensemble muni d'une relation d'ordre, et soit A un sous-ensemble non vide de E . Alors

1. Si l'ensemble $\text{Maj}_E(A)$ des majorants de A dans E admet un plus petit élément M , alors M est appelé la borne supérieure de A dans E (ou le supremum de A) et est noté $\sup(A)$ ou $\sup_E(A)$.
2. Si l'ensemble $\text{Min}_E(A)$ des minorants de A dans E admet un plus grand élément m , alors m est appelé la borne inférieure de A dans E (ou l'infimum de A) et est noté $\inf(A)$ ou $\inf_E(A)$.

On notera que les bornes (supérieure ou inférieure) d'un ensemble A n'appartiennent pas nécessairement à A .

Remarque 2.2.8 Soit $(E, \leq)$ un ensemble muni d'une relation d'ordre, soient A un sous-ensemble non vide de E et $m, M \in E$. Alors

1. Pour que M soit la borne supérieure de A dans E , il faut et il suffit que

$$\begin{cases} \forall a \in A : a \leq M, \\ \text{et} \\ \forall x \in E, ((\forall a \in A : a \leq x) \implies M \leq x). \end{cases}$$

Autrement dit un élément $M \in E$ est appelé une borne supérieure de A dans E si et seulement si

$$\begin{cases} M \text{ est un majorant de } A, \\ \text{et} \\ \text{pour tout majorant } x \text{ de } A, \text{ on a } M \leq x. \end{cases}$$

2. Pour que m soit la borne inférieure de A dans E , il faut et il suffit que

$$\begin{cases} \forall a \in A : m \leq a, \\ \text{et} \\ \forall x \in E, ((\forall a \in A : x \leq a) \implies x \leq m). \end{cases}$$

Autrement dit un élément $m \in E$ est appelé une borne inférieure de A dans E si et seulement si

$$\begin{cases} m \text{ est un minorant de } A, \\ \text{et} \\ \text{pour tout minorant } x \text{ de } A, \text{ on a } x \leq m. \end{cases}$$

3. Un majorant de A plus petit que tous les autres s'appelle borne supérieure de A .

4. Un minorant de A plus grand que tous les autres s'appelle borne inférieure de A .

Exemple 2.2.18

1. Pour la relation d'ordre usuelle $\leq$ sur $\mathbb{R}$. Soient a et b des réels, avec $a < b$. Soit $A = [a, b[$. Alors les majorants de A sont les réels supérieurs ou égaux à b . Les minorants de A sont les réels inférieurs ou égaux à a . L'ensemble A n'a pas de plus grand élément, mais il a une borne supérieure $\sup(A) = b$. L'ensemble A a un plus petit élément $\min(A) = a$. De ce fait $\inf(A)$ existe et vaut a .

2. Soit $E = \mathcal{P}(\mathbb{N})$ ordonné par la relation de l'inclusion $\subset$ et soit

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

un sous-ensemble de E . Alors :

► Il a une borne supérieure. En effet, les majorants de A sont les parties de $\mathbb{N}$ contenant tous les ensembles

$$\{1, 2\}, \{2, 3\}, \{1, 4\}, \{1, 2, 3\},$$

c'est à dire les parties de $\mathbb{N}$ contenant leur **union**. Leur union est donc un majorant de A , et plus petite que tous les majorants de A . On en déduit que A a une borne supérieure qui est

$$\sup(A) = \{1, 2, 3, 4\}.$$

En revanche, Il n'a donc pas de plus grand élément car

$$\{1, 2, 3, 4\} \notin A.$$

► Il a une borne inférieure. En effet, les minorants de A sont les parties de $\mathbb{N}$ incluses dans tous les ensembles

$$\{1, 2\}, \{2, 3\}, \{1, 4\}, \{1, 2, 3\},$$

donc incluse dans leur **intersection**. On en déduit que

$$\inf_{\mathbb{N}}(A) = \{1\}.$$

En revanche, A n'a pas de plus petit élément car $\{1\} \notin A$.

3. Soit $E = \mathbb{N}^* - \{1\}$ ordonné par la relation de la divisibilité $|$. Pour tous $x, y \in E$, soit $A = \{x, y\}$, alors

$$\begin{aligned}\sup_E(A) &= \text{PPCM}(x, y), \\ \inf_E(A) &= \text{PGCD}(x, y).\end{aligned}$$

4. Si on considère l'ensemble ordonné $(\mathbb{N}^*, |)$ et $A = \{1, 2, 4, 6\}$, alors

$$\inf_E(A) = \min(A) = 1, \sup_E(A) = 12,$$

et $\max(A)$ n'existe pas (sinon on aurait $\sup_E(A) = \max(A) \in A$).

Corollaire 2.2.2 Pour toute partie A d'un ensemble ordonné $(E, \leq)$, on a

1. La borne inférieure ou supérieure de A quand elle existe n'est pas nécessairement un élément de A .
2. Si A admet un plus grand élément M , alors M est la borne supérieure de A et

$$\sup_E(A) = \max_E(A) = M.$$

La réciproque est fausse.

3. Si A admet un plus petit élément m , alors m est la borne inférieure de A et

$$\inf_E(A) = \min_E(A) = m.$$

La réciproque est fausse.

4. **Cas particulier de l'ensemble ordonné $(\mathbb{R}, \leq)$:**

- Toute partie non vide et majorée de $\mathbb{R}$ a une borne supérieure.
- Toute partie non vide et minorée de $\mathbb{R}$ a une borne inférieure.

Proposition 2.2.5 (Unicité de la borne supérieure, inférieure)

La borne supérieure (resp. inférieure) si elle existe est unique.

Preuve. Supposons que A admette deux bornes supérieures M_1 et M_2 , alors les deux sont des des majorants de A . Puisque M_1 est le plus petit des majorants de A , alors

$$M_1 \leq M_2$$

Par le même argument, on a

$$M_2 \leq M_1$$

La propriété d'antisymétrie de la relation entraîne alors que $M_1 = M_2$. Le même raisonnement s'applique pour montrer l'unicité de la borne inférieure de A . $\square$

Chapitre 3

Fonctions et Applications

Dans ce chapitre, nous explorons un type particulier de relations entre deux ensembles : les fonctions et les applications, qui sont des concepts étroitement liés en théorie des ensembles. Ce chapitre se divise en deux parties principales. La première partie introduit les notions fondamentales des fonctions en tant que relations spécifiques entre deux ensembles. La section consacrée aux applications approfondit des concepts tels que les applications injectives, surjectives et bijectives. Elle aborde également les notions d'image directe et réciproque, ainsi que les propriétés associées aux applications, sans oublier les prolongements et restrictions de celles-ci. Ce chapitre est essentiel pour comprendre les relations entre ensembles et les diverses manières de transformer ou de relier ces ensembles par le biais des fonctions et des applications.

3.1 Fonctions

Une fonction désigne un type de correspondance particulière entre deux ensembles, dans laquelle chaque élément $x \in X$ (appartenant à l'ensemble de départ, appelé domaine) est associé à au plus un élément dans un autre ensemble Y . Cela signifie qu'il existe au plus une flèche partant de chaque élément x de X vers un élément de Y .

Définition 3.1.1 Soient E, F deux ensembles non vides.

1. On appelle fonction d'un ensemble E vers un ensemble F , toute correspondance (ou relation) notée f , qui, à chaque élément x de E , fait correspondre **au plus** un élément y de F .
2. Si x est un élément de E en relation avec un élément y de F par une fonction f , on écrit $y = f(x)$ et on dit que y est l'image de x par la fonction f , mais aussi que x est un antécédent de y par la fonction f .

Notation 3.1.1

1. Une fonction se note

$$\begin{aligned} f : E &\longrightarrow F \\ x &\longmapsto y = f(x), \end{aligned}$$

ou parfois plus simplement $f : x \longmapsto f(x)$ ou encore plus simplement f . Mais une fonction ne se note pas $f(x)$ car la notation $f(x)$ désigne l'image de l'élément x par la fonction f .

2. $f : E \longrightarrow F$ se lit " f est une fonction de E vers F " ou encore " f est une fonction définie sur E à valeurs dans F " et $x \longmapsto f(x)$ se lit " à x est associé son image $f(x)$ ".
3. Parmi les autres symboles fréquemment employés pour désigner une fonction on retrouve $g, h, s, i, \dots$, mais aussi des lettres grecques telles que ϕ ou ψ .

Définition 3.1.2 Soient E, F deux ensembles non vides et f une fonction de E vers F .

1. L'ensemble E s'appelle l'ensemble de départ (ou source) et F l'ensemble d'arrivée (ou but) de la fonction f .
2. L'ensemble des couples $(x, f(x))$ où x décrit E s'appelle le graphe de la fonction f noté $\Gamma(f)$ où

$$\Gamma(f) = \{(x, f(x)) : x \in E\},$$

qui est un sous-ensemble de $E \times F$.

Remarque 3.1.1 Une fonction est une relation, mais une relation n'est pas forcément une fonction. Considérons l'exemple suivant de la relation « **fils ou fille de** ». Nous pouvons voir que notre relation de « **fils ou fille de** » n'est pas une fonction, car chaque personne a deux parents.

Grphe : Sur un diagramme sagittal cela donne une présentation de type

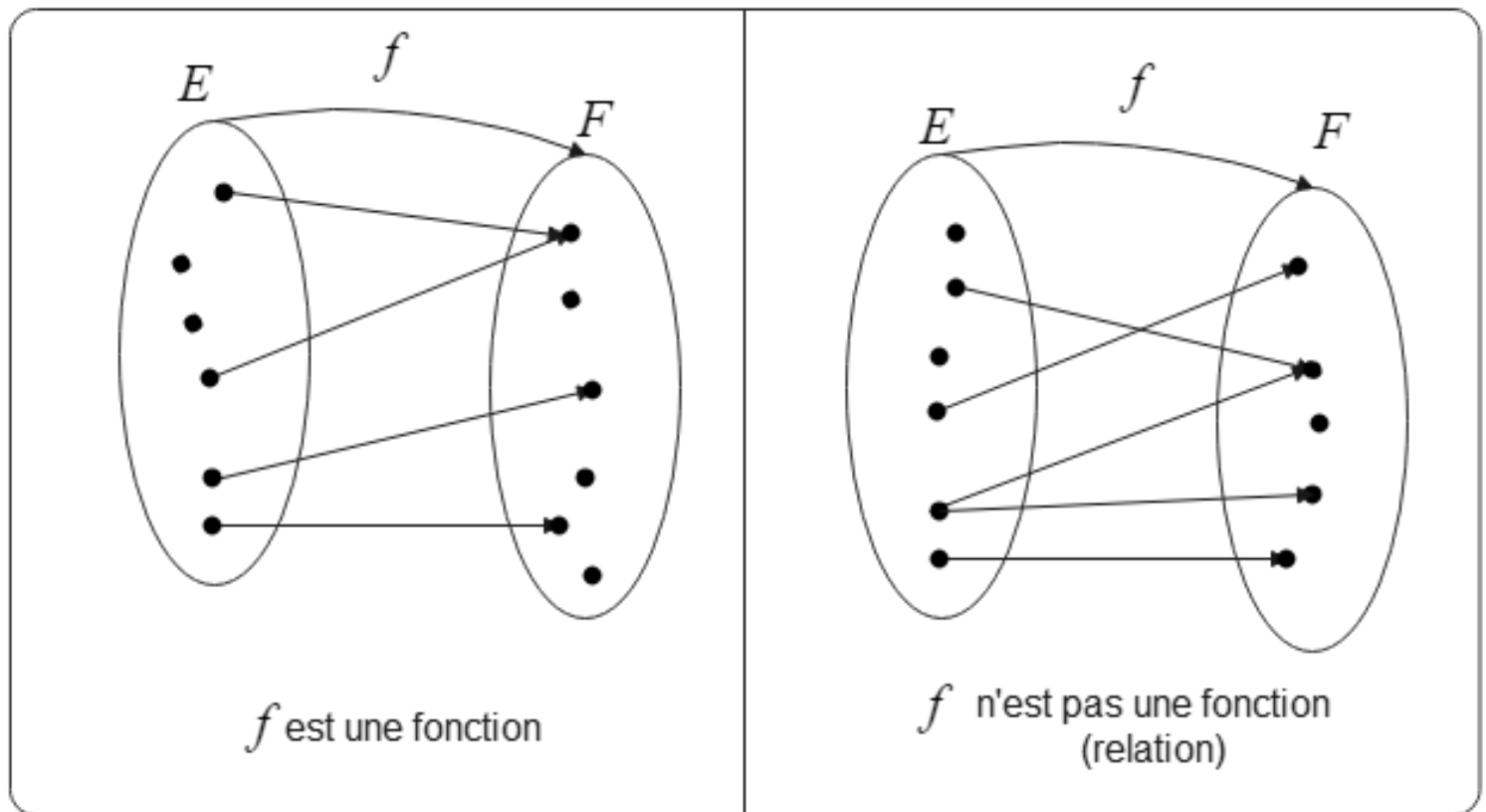


FIGURE 3.1: Fonction, non Fonction

Pour une fonction $f : E \longrightarrow F$, les notions d'ensemble de définition et d'ensemble image sont importantes pour comprendre son domaine d'application et les valeurs qu'elle peut prendre.

Définition 3.1.3 Soient E, F deux ensembles non vides et f une fonction de E vers F .

1. On appelle ensemble de définition (ou domaine de définition) de la fonction f , et on note $D(f)$, l'ensemble

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

Autrement dit, $D(f)$ est l'ensemble des éléments de E ayant une image par f ou l'ensemble des valeurs pour lesquelles la fonction est définie. C'est-à-dire l'ensemble de tous les x pour lesquels $f(x)$ existe et prend une valeur.

2. On appelle ensemble image de f , et on note $Im(f)$, ou $F(E)$ l'ensemble

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

Autrement dit, $Im(f)$ est l'ensemble des valeurs effectivement prises par la fonction f .

Exemple 3.1.1

1. Soit la fonction (Fonction valeur absolue)

$$\begin{aligned} f : \mathbb{R} &\longrightarrow \mathbb{R} \\ x &\longmapsto f(x) = |x - 1|, \end{aligned}$$

alors

• **Ensemble de définition.** La fonction valeur absolue est définie pour tous les réels, alors $D(f) = \mathbb{R}$.

• **Ensemble Image.** Pour tout x réel, $|x - 1|$ est toujours positive ou nulle, alors $Im(f) = \mathbb{R}^+$.

• **Graphe.** Le graphe de f est une ligne en forme de $\vee$ avec son sommet à $(1, 0)$.

2. La correspondance f qui associe à chaque entier naturel le mois correspondant est une fonction de $E = \mathbb{N}$ dans l'ensemble

$$F = \{\text{Janvier, Février, Mars, Avril, Mai, Juin, Juillet, Août, Septembre, Octobre, Novembre, Décembre}\}.$$

On a dans ce cas $f(2) = \text{Novembre}$, et $f(17)$ n'existe pas.

$$D(f) = \{1, 2, 3, 4, 5, 6, 7, 8, 9, 10, 11, 12\}.$$

et

$$Im(f) = F.$$

Ainsi, $Im(f)$ est l'ensemble des mois de l'année, car chaque entier de 1 à 12 correspond précisément à l'un des 12 mois.

3. La correspondance qui associe à chaque mois le nombre possible de jours du mois n'est une fonction de l'ensemble $\mathbb{N}$ de l'exemple précédent dans F , car elle fait associer à Février, les deux éléments 28 et 29.

3. Dans un commerce au détail, le commerçant choisi, sous certaines contraintes, le prix de chacun des articles en vente. L'ensemble de départ est ici les articles vendus, l'ensemble d'arrivée est par exemple $\mathbb{R}$, l'ensemble des nombres réels. À chaque article correspond un et un seul prix. Cette correspondance « **a le prix de** » est donc une fonction.

4. Considérons la fonction $f : \mathbb{R} \longrightarrow \mathbb{R}$ définie par $f(x) = x^2$. Pour tout $x \in \mathbb{R}$, x^2 est toujours un nombre réel positif ou nul. Donc ensemble image de f est $Im(f) = \mathbb{R}^+$.

3.2 Applications

Un cas particulier de relations est celui où, pour un x donné, il est relié à un seul y . Lorsque cette relation associe précisément un y à chaque x , on parle alors d'une application. Autrement dit, dans une application, chaque élément du domaine est associé à un unique élément de l'ensemble d'arrivée, garantissant une correspondance bien définie.

Définition 3.2.1 Soient E, F deux ensembles non vides. Une application $f : E \longrightarrow F$ est une correspondance qui à tout élément $x \in E$ associe **un unique** élément noté $f(x)$ de l'ensemble F . Avec des quantificateurs, cela donne

$$f \text{ est une application} \iff \forall x \in E, \exists ! y \in F : f(x) = y.$$

En d'autres termes, une fonction $f : E \longrightarrow F$ est une application si et seulement si $D(f) = E$.

Notation 3.2.1

1. L'ensemble des applications de E dans F se note habituellement $\mathcal{F}(E, F)$ ou plus fréquemment F^E (faire attention à l'ordre des lettres). Ainsi, l'ensemble des suites réelles se note $\mathbb{R}^{\mathbb{N}}$ (car une suite réelle est une application de $\mathbb{N}$ dans $\mathbb{R}$).

2. Si $E = F$, l'ensemble des applications de E dans E se note plus simplement E^E ou $\mathcal{F}(E)$.

Remarque 3.2.1

- La différence entre fonction et application ne concerne donc que l'ensemble de départ.
- On transforme facilement une fonction en une application en prenant son ensemble de définition pour ensemble de départ.
- À ce point, il est important de souligner la différence entre fonction et application. Pour une fonction f , chaque élément de l'ensemble de départ admet au plus une image par f . En particulier, il peut arriver qu'un élément n'admette pas d'image par f , c'est-à-dire que la fonction f n'est pas définie en ce point. En revanche, pour une application, chaque élément de l'ensemble de départ admet exactement une image par f .

Graphe : Sur un diagramme sagittal cela donne une présentation de type

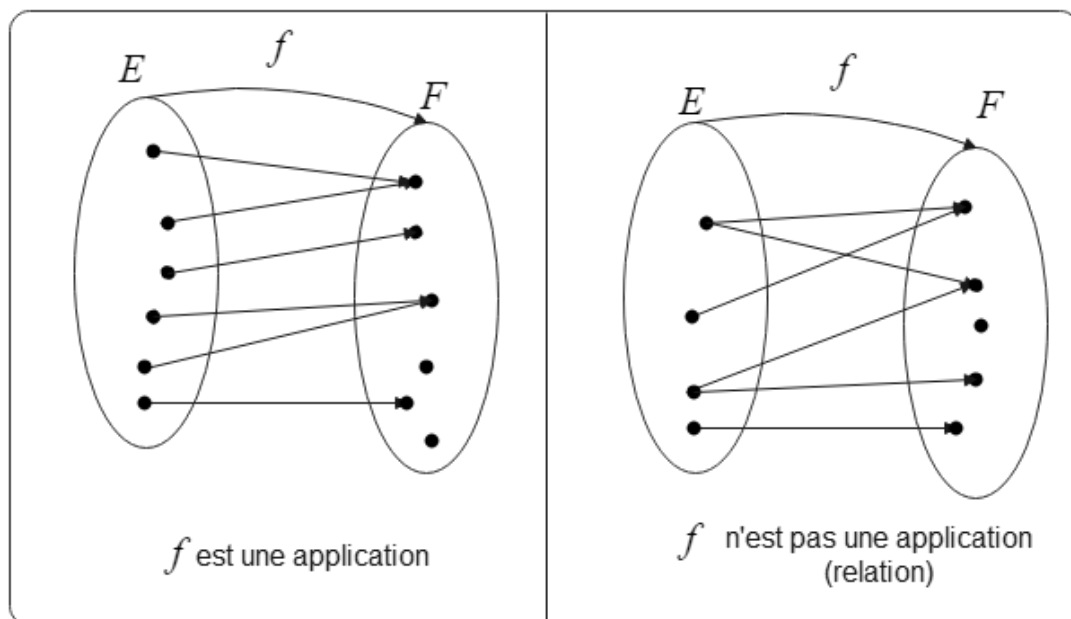


FIGURE 3.2: Application, non Application

Exemple 3.2.1

1. Considérons l'ensemble des êtres humains comme ensemble de départ et d'arrivée. La relation « **a pour père** » est bien une application, puisqu'à chaque être humain est associé un et un seul père. Par contre, la relation « **a pour enfant** » n'est pas une application (n'est pas une fonction), puisque chaque être humain peut avoir zéro, un, deux, etc. enfants.

2. Soit la relation $f : \mathbb{R} \longrightarrow \mathbb{R}, x \longmapsto f(x) = x^2$, à chaque nombre réel x est associé son carré. Il s'agit bien d'une application puisque chaque nombre réel n'a qu'un seul carré, mais $(-2)^2 = 2^2 = 4$, deux nombres réels peuvent avoir le même carré.

3. On définit une application f de $\mathbb{N}$ dans $\{0, 1, \dots, 9\}$ en posant

$$f(n) = \text{le chiffre des unités de } n.$$

On peut le noter

$$\begin{aligned} f : \mathbb{N} &\longrightarrow \{0, 1, \dots, 9\} \\ n &\longmapsto \text{le chiffre des unités de } n. \end{aligned}$$

4. On définit une application f de $\{1, 2, 3, 4\}$ dans $\{0, 1\}$ en précisant que f associe à n son reste modulo 2, ou de façon équivalente, en posant

$$f(n) = \begin{cases} 0, & \text{si } n \text{ est pair,} \\ 1, & \text{si } n \text{ est impair.} \end{cases}$$

On peut le noter

$$\begin{aligned} f : \{1, 2, 3, 4\} &\longrightarrow \{0, 1\} \\ n &\longmapsto (n) = \begin{cases} 0, & \text{si } n \text{ est pair,} \\ 1, & \text{si } n \text{ est impair.} \end{cases} \end{aligned}$$

Définition 3.2.2 (Égalité des applications)

Deux applications $f : E \longrightarrow F$ et $g : E' \longrightarrow F'$ sont dites égales si les trois propriétés suivantes sont vérifiées.

(i). $E = E'$ (même ensemble de départ).

(ii). $F = F'$ (même ensemble d'arrivée).

(iii). Pour tout x , $f(x) = g(x)$.

En particulier, si $E = E'$ et $F = F'$ alors $f \neq g$ si et seulement si

$$\exists x \in E : f(x) \neq g(x).$$

Exemple 3.2.2

1. Soient

$$\begin{aligned} f_1 : \mathbb{R} &\longrightarrow \mathbb{R} & f_2 : \mathbb{R} &\longrightarrow \mathbb{R} \\ x &\longmapsto f_1(x) = \cos x & x &\longmapsto f_2(x) = 2 \cos^2\left(\frac{x}{2}\right) - 1. \end{aligned} \quad \text{et}$$

Alors, on a $f_1 = f_2$. En effet, l'identité trigonométrique suivante

$$\cos 2\alpha = 2 \cos^2 \alpha - 1,$$

nous donne en posant $\alpha = \frac{x}{2}$,

$$f_1(x) = \cos x = 2 \cos^2\left(\frac{x}{2}\right) - 1.$$

C'est exactement la forme de $f_2(x)$.

2. Les trois applications

$$\begin{aligned} g_1 : \mathbb{R} &\longrightarrow \mathbb{R} & g_2 : \mathbb{R} &\longrightarrow \mathbb{R}_+ & g_3 : \mathbb{R}_+ &\longrightarrow \mathbb{R} \\ x &\longmapsto g_1(x) = x^2, & x &\longmapsto g_2(x) = x^2 & x &\longmapsto g_3(x) = x^2. \end{aligned} \quad \text{et}$$

sont deux à deux distinctes.

3. Les applications f et g définies de $\mathbb{N}$ dans $\mathbb{Z}$ par $f(x) = \cos(\pi x)$ et $g(x) = (-1)^n$ sont égales. C.à.d $f = g$.

3.2.1 Applications particulières

Définissons maintenant une famille d'applications remarquables : application identité, nulle, constante et indicatrice.

Définition 3.2.3 (*Application Identique*)

Soit E un ensemble quelconque. L'application identité de E (ou application identique de E) est l'application de E dans E , notée Id_E , définie par

$$\forall x \in E : Id_E(x) = x.$$

En d'autres termes, l'application identité associe à chaque élément de son ensemble de départ le même élément dans l'ensemble d'arrivée.

Nous verrons plus loin qu'elle joue le rôle d'élément neutre pour la loi de composition.

Définition 3.2.4 (*Application constante*)

1. Soient E et F deux ensembles non vides. Une application f de E dans F est dite constante s'il existe un élément c de F , tel que pour tout x de E , on ait $f(x) = c$, c'est-à-dire si

$$\exists c \in F, \forall x \in E : f(x) = c.$$

- On dit alors que f est constante égale à c et est souvent simplement notée c .
- 2. Une application constante associe à chaque élément de son ensemble de départ un unique élément fixe c dans l'ensemble d'arrivée.

Définition 3.2.5 (*Application nulle*)

Pour $f : E \longrightarrow \mathbb{R}$, l'application constante f définie par $f(x) = 0$ est appelée application nulle notée o . C'est-à-dire l'application qui, à tout élément de E associe le nombre 0.

$$\forall x \in E : f(x) = 0.$$

Une application indicatrice est une application définie sur un ensemble E qui explicite l'appartenance ou non à un sous-ensemble A de E de tout élément de E . Elle est paramétrée par un sous-ensemble, disons A , et qui ne peut prendre que deux valeurs : la valeur 1 si la variable de l'application est élément de A , et la valeur 0 sinon.

Définition 3.2.6 (*Application indicatrice d'une partie d'un ensemble*)

Soit A une partie d'un ensemble E . On appelle application indicatrice de A dans E , ou parfois application caractéristique de A dans E et on note φ_A , l'application définie par

$$\begin{array}{ccc} \varphi_A : E & \longrightarrow & \{0, 1\} \\ x & \longmapsto & \varphi_A(x) = \begin{cases} 1, & \text{si } x \in A \\ 0, & \text{si } x \notin A \end{cases} \end{array} \iff \begin{array}{ccc} \varphi_A : E & \longrightarrow & \{0, 1\} \\ x & \longmapsto & \varphi_A(x) = \begin{cases} 1, & \text{si } x \in A \\ 0, & \text{si } x \in \overline{A}. \end{cases} \end{array}$$

Remarque 3.2.2 1. Il est clair que

$$\varphi_\emptyset = 0, \varphi_E = 1,$$

où 0 désigne l'application nulle, c'est-à-dire l'application qui, à tout élément de E associe le nombre 0 et 1 désigne l'application qui, à tout élément de E associe le nombre 1.

2. D'autres notations souvent employées pour l'application caractéristique de A sont 1_A , χ_A , et I_A (i majuscule).

3. Soit E un ensemble, et A un sous-ensemble de E , alors on peut écrire

$$A = \{x \in E : \varphi_A(x) = 1\}.$$

Exemple 3.2.3

1. Il est clair que

$$\varphi_{\mathbb{R}^+}(2) = 1, \varphi_{\mathbb{R}^-}(2) = 0, \varphi_{\mathbb{Q}}(\pi) = 0.$$

2. Considérons l'application

$$f(x) = \begin{cases} e^{-x}, & \text{si } x \geq 0 \\ 0, & \text{si } x < 0. \end{cases}$$

En utilisant une application indicatrice, f peut s'écrire en une seule ligne

$$f(x) = e^{-x} \varphi_{[0, +\infty[}(x),$$

en tout point x .

3. Considérons maintenant l'application

$$f(x) = \begin{cases} 0, & \text{si } x \leq -1 \\ \frac{x+1}{4}, & \text{si } -1 < x < 3 \\ 1, & \text{si } x \geq 3. \end{cases}$$

En utilisant deux applications indicatrices, on peut écrire

$$f(x) = \left(\frac{x+1}{4} \right) \varphi_{]-1, 3[}(x) + \varphi_{[3, +\infty[}(x),$$

en tout point x .

4. Considérons la partie

$$A = 2\mathbb{N} + 1 = \{2n + 1, n \in \mathbb{N}\},$$

de $\mathbb{N}$. Alors son indicatrice est l'application

$$\begin{aligned} \varphi_A : \mathbb{N} &\longrightarrow \{0, 1\} \\ n &\longmapsto \varphi_A(n) = \begin{cases} 0, & \text{si } n \text{ est pair} \\ 1, & \text{si } n \text{ est impair.} \end{cases} \end{aligned}$$

Il s'agit donc de l'application qui à un entier associe son reste modulo 2.

La proposition suivante est un simple dictionnaire traduisant les opérations sur les ensembles en termes des applications indicatrices.

Proposition 3.2.1 Soient E un ensemble, A et B deux parties de E . Alors l'application indicatrice vérifie les propriétés suivantes,

- | | |
|---|---|
| 1. $\varphi_A^2 = \varphi_A$ | 2. $A \subset B \iff \varphi_A \leq \varphi_B$ |
| 3. $A = B \iff \varphi_A = \varphi_B$ | 4. $\varphi_{C_E^A} = 1 - \varphi_A$ |
| 5. $\varphi_{A \cap B} = \varphi_A \cdot \varphi_B$ | 6. $\varphi_{A B} = \varphi_A(1 - \varphi_B)$ |
| 7. $\varphi_{A \cup B} = \varphi_A + \varphi_B - \varphi_A \cdot \varphi_B$ | 8. $\varphi_{A \Delta B} = \varphi_A + \varphi_B - 2 \cdot \varphi_A \varphi_B = (\varphi_A - \varphi_B)^2 = \varphi_A - \varphi_B $. |

Preuve.

1. Montrons que

$$\forall x \in E : \varphi_A^2(x) = \varphi_A(x).$$

Observons déjà que φ_A^2, φ_A sont toutes les deux des applications de E dans $\{0, 1\}$. Soit $x \in E$, nous avons deux cas possibles :

- Si $x \in A$, alors $\varphi_A(x) = 1$, donc $\varphi_A^2(x) = \varphi_A(x) \cdot \varphi_A(x) = 1 = \varphi_A(x)$.

• Si $x \notin A$, alors $\varphi_A(x) = 0$, donc $\varphi_A^2(x) = \varphi_A(x) \cdot \varphi_A(x) = 0 = \varphi_A(x)$.

Dans les deux cas, nous avons $\varphi_A^2(x) = \varphi_A(x)$, ce qui prouve que $\varphi_A^2 = \varphi_A$.

2. • Si $A \subset B$, alors tout élément x qui appartient à A appartient également à B . Par conséquent, pour chaque x de E ,

Si $x \in A$, alors $\varphi_A(x) = 1$ et $\varphi_B(x) = 1$, donc $\varphi_A(x) \leq \varphi_B(x)$.

Si $x \notin A$, alors $\varphi_A(x) = 0$ et puisque $\varphi_B(x)$ peut être soit 0 (si $x \notin B$) soit 1 (si $x \in B$), donc on a toujours $\varphi_A(x) \leq \varphi_B(x)$.

• Réciproquement. Si $\varphi_A(x) \leq \varphi_B(x)$ pour tout $x \in E$, alors cela signifie que chaque fois que x appartient à A ($\varphi_A(x) = 1$), x doit aussi appartenir à B ($\varphi_B(x) = 1$). Cela implique que A est un sous-ensemble de B .

3. • L'implication directe est évidente.

• Supposons que $\varphi_A = \varphi_B$ et montrons que $A = B$ par double inclusion.

(i) Soit $x \in A$, alors $\varphi_A(x) = 1$ et donc $\varphi_B(x) = 1$, d'où $x \in B$. Ceci prouve que $A \subset B$.

(i) Soit $x \in B$, alors $\varphi_B(x) = 1$ et donc $\varphi_A(x) = 1$, d'où $x \in A$. Ceci prouve que $B \subset A$. Donc $A = B$.

4. Observons déjà que $\varphi_{C_E^A}$, $1 - \varphi_A$ sont toutes les deux des applications de E dans $\{0, 1\}$. Soit $x \in E$, traitons deux cas :

• Si $x \in A$, alors on a $\varphi_A(x) = 1$ et $\varphi_{\bar{A}}(x) = 0$ (i.e. $x \notin \bar{A}$). Par conséquent

$$\varphi_A(x) + \varphi_{\bar{A}}(x) = 1.$$

• Si $x \notin A$, alors on a $\varphi_A(x) = 0$ et $\varphi_{\bar{A}}(x) = 1$ (i.e. $x \in \bar{A}$). Par conséquent

$$\varphi_A(x) + \varphi_{\bar{A}}(x) = 1.$$

Donc les applications $\varphi_{C_E^A}$ et $1 - \varphi_A$ sont égales.

5. Observons déjà que $\varphi_{A \cap B}$, φ_A , et φ_B sont toutes les deux des applications de E dans $\{0, 1\}$. Soit $x \in E$, on distingue deux cas :

• Si $x \in A \cap B$, alors on a $\varphi_{A \cap B}(x) = 1$, et par ailleurs

$$x \in A \cap B \iff x \in A \wedge x \in B,$$

et donc

$$\varphi_A(x) = \varphi_B(x) = 1 \implies \varphi_A(x) \cdot \varphi_B(x) = 1.$$

Donc finalement

$$\varphi_{A \cap B}(x) = \varphi_A(x) \cdot \varphi_B(x).$$

• Si $x \notin A \cap B$, alors on a $\varphi_{A \cap B}(x) = 0$, et par ailleurs

$$x \notin A \cap B \iff x \notin A \vee x \notin B,$$

et donc

$$\varphi_A(x) = 0 \vee \varphi_B(x) = 0 \implies \varphi_A(x) \cdot \varphi_B(x) = 0.$$

Donc

$$\varphi_{A \cap B}(x) = \varphi_A(x) \cdot \varphi_B(x).$$

Donc les applications $\varphi_{A \cap B}$ et $\varphi_A \cdot \varphi_B$ sont égales.

6. Observons déjà que $\varphi_{A|B}$ et $\varphi_A(1 - \varphi_B)$ sont toutes les deux des applications de E dans $\{0, 1\}$. Utilisons ce qui précède, on obtient

$$\varphi_{A|B} = \varphi_{A \cap \bar{B}} = \varphi_A \cdot \varphi_{\bar{B}} = \varphi_A(1 - \varphi_B).$$

7. En passant par des complémentaires et en utilisant des résultats précédents, on a

$$\varphi_{A \cup B} + \varphi_{\overline{A \cup B}} = 1 \implies \varphi_{A \cup B} + \varphi_{\bar{A} \cap \bar{B}} = 1$$

$$\begin{aligned} &\implies \varphi_{A \cup B} + \varphi_{\bar{A}} \cdot \varphi_{\bar{B}} = 1 \\ &\implies \varphi_{A \cup B} + (1 - \varphi_A) \cdot (1 - \varphi_B) = 1, \end{aligned}$$

ce qui correspond à

$$\varphi_{A \cup B} = \varphi_A + \varphi_B - \varphi_A \cdot \varphi_B.$$

8. En utilisant des résultats précédents

$$\begin{aligned} \varphi_{A \Delta B} &= \varphi_{(A \cup B) \setminus A \cap B} = \varphi_{(A \cup B)}(1 - \varphi_{\overline{A \cap B}}) = \varphi_{(A \cup B)}(1 - \varphi_{\overline{A \cup B}}) = \\ &= (\varphi_A + \varphi_B - \varphi_A \cdot \varphi_B)(1 - \varphi_{\bar{A}} - \varphi_{\bar{B}} + \varphi_{\bar{A}} \cdot \varphi_{\bar{B}}) \\ &= (\varphi_A + \varphi_B - \varphi_A \cdot \varphi_B)(1 - (1 - \varphi_A) - (1 - \varphi_B) + (1 - \varphi_A) \cdot (1 - \varphi_B)) \\ &= \varphi_A + \varphi_B - 2 \cdot \varphi_A \varphi_B. \end{aligned}$$

□

3.2.2 Partie stable par une application

Une partie stable par une application (ou fonction) est un sous-ensemble d'un ensemble qui, lorsqu'il est soumis à cette application, reste invariant sous cette application. En d'autres termes, une partie A d'un ensemble E est stable par une application $f : E \longrightarrow E$ si, en appliquant f à tout élément de A , on obtient toujours un élément qui appartient également à A .

Définition 3.2.7 Soit $f : E \longrightarrow E$ une application d'un ensemble E dans lui-même et soit A une partie de E .

1. On dit que A est stable par f si et seulement si

$$\forall x \in A : f(x) \in A.$$

Ceci est équivalent à

$$f(A) \subset A.$$

2. On dit que A est invariante par f si $f(A) = A$. L'application $g : A \longrightarrow A$ qui coïncide avec f sur A s'appelle l'application induite par f sur A .

Exemple 3.2.4 Pour l'application $f : x \longmapsto x^2$, étudions si les parties suivantes sont stables :

$$A = \mathbb{R}^+, B = [0, 2], C = [2, +\infty[.$$

- Si $x \in A = \mathbb{R}^+$, alors on a $x^2 \in \mathbb{R}^+$ et donc A est stable par f .
- La partie B n'est pas stable par f car il existe $x = 2 \in B$ avec $f(x) = 4 \notin B$.
- Si $x \geq 2$, alors $f(x) \geq 4$ donc $f(x) > 2$. Ainsi C est stable par f .

3.2.3 Prolongements et restrictions

À partir d'une application donnée, nous pouvons créer d'autres applications en remplaçant simplement l'ensemble de départ ou d'arrivée par un sous-ensemble ou un sur-ensemble de cet ensemble.

3.2.3.1 Restriction d'une application

Définition 3.2.8 (*Restriction d'une application*)

Soient E et F deux ensembles non vides et f une application de E dans F . Soit E' une partie de E ($E' \subset E$).

1. On appelle restriction de f à E' , l'application g de E' dans F qui à tout x de E' associe $f(x)$, i.e. telle que

$$\forall x \in E' : g(x) = f(x).$$

Autrement dit g coïncide avec f sur E' . Cette application g est habituellement notée $f|_{E'}$.

Remarque 3.2.3

- L'application $f|_{E'}$ est la même que f , mais elle est seulement appliquée aux éléments de E' .
- La restriction d'une application consiste donc à limiter son domaine de définition à un sous-ensemble plus petit tout en conservant la même règle de correspondance pour les éléments de ce sous-ensemble.

Exemple 3.2.5

1. La restriction d'une fonction à son domaine de définition est une application. Ceci fait que l'on identifie parfois les notions de fonction et d'application.
2. Les applications suivantes

$$\begin{aligned} f_1 : [-1, 1] &\longrightarrow \mathbb{R} \\ x &\longmapsto f_1(x) = 1 - x^2, \end{aligned}$$

et

$$\begin{aligned} f_2 : [1, +\infty[&\longrightarrow \mathbb{R} \\ x &\longmapsto f_2(x) = x^2 - 1, \end{aligned}$$

sont des restrictions de l'application

$$\begin{aligned} f : \mathbb{R} &\longrightarrow \mathbb{R} \\ x &\longmapsto f_2(x) = |x^2 - 1| \end{aligned}$$

à $[-1, 1]$ et $[1, +\infty[$ respectivement. C'est-à-dire

$$f|_{[-1, 1]} = f_1, f|_{[1, +\infty[} = f_2.$$

3.2.3.2 Prolongement d'une application

À l'inverse de l'opération de restriction, on peut prolonger une application f de E dans F à un ensemble E' contenant E .

Définition 3.2.9 (*Prolongements d'une application*)

Soient E et F deux ensembles non vides et f une application de E dans F . Soit E' un ensemble contenant E ($E \subset E'$).

1. On appelle prolongement de f à E' , toute application $\tilde{f}$ de E' dans F dont la restriction à E est égale à f , i.e. telle que

$$\forall x \in E : \tilde{f}(x) = f(x).$$

Remarque 3.2.4

1. Le prolongement $\tilde{f}$ doit coïncider avec f sur E , mais est aussi défini sur tout l'ensemble E' .
2. Le prolongement d'une application consiste à étendre son domaine de définition à un ensemble plus grand, tout en gardant la même règle de correspondance là où elle était initialement définie.

Exemple 3.2.6

1. L'application

$$\begin{aligned}\tilde{f} : \mathbb{R} &\longrightarrow \mathbb{R} \\ x &\longmapsto \tilde{f}(x) = \begin{cases} \frac{\sin x}{x}, & \text{si } x \neq 0 \\ 1, & \text{si } x = 0, \end{cases}\end{aligned}$$

est un prolongement de l'application f à $\mathbb{R}$ telle que

$$\begin{aligned}f : \mathbb{R}^* &\longrightarrow \mathbb{R} \\ x &\longmapsto f(x) = \frac{\sin x}{x}.\end{aligned}$$

2. Considérons la fonction

$$\begin{aligned}f : [-1, 1] &\longrightarrow \mathbb{R} \\ x &\longmapsto f(x) = x^2,\end{aligned}$$

on peut prolonger cette fonction à tout $\mathbb{R}$ en définissant

$$\begin{aligned}\tilde{f} : \mathbb{R} &\longrightarrow \mathbb{R} \\ x &\longmapsto \tilde{f}(x) = \begin{cases} x^2, & \text{si } x \in [-1, 1] \\ 0, & \text{si } x \notin [-1, 1]. \end{cases}\end{aligned}$$

Remarque 3.2.5 Il existe en général plusieurs prolongements d'une même application. Par exemple les applications

$$\begin{aligned}\tilde{f}_1 : \mathbb{R} &\longrightarrow \mathbb{R} \\ x &\longmapsto \tilde{f}_1(x) = \begin{cases} \frac{\sin x}{x}, & \text{si } x \neq 0 \\ 1, & \text{si } x = 0, \end{cases}\end{aligned}$$

et

$$\begin{aligned}\tilde{f}_2 : \mathbb{R} &\longrightarrow \mathbb{R} \\ x &\longmapsto \tilde{f}_2(x) = \begin{cases} \frac{\sin x}{x}, & \text{si } x \neq 0 \\ 0, & \text{si } x = 0, \end{cases}\end{aligned}$$

sont des prolongements de la même application

$$\begin{aligned}f : \mathbb{R}^* &\longrightarrow \mathbb{R} \\ x &\longmapsto f(x) = \frac{\sin x}{x}.\end{aligned}$$

3.2.4 Composition des applications

La composition d'applications est une opération fondamentale qui combine deux applications pour former une nouvelle application. Cette opération est largement utilisée dans divers domaines tels que l'analyse, l'algèbre, et la topologie.

Définition 3.2.10 Soient E, F et G trois ensembles non vides. Soient f une application de E vers F et g une application de F vers G .1. La composée des applications g et f (ou f suivie de g), notée $g \circ f$ (se lit « g rond f »), est l'application de E vers G définie par

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

2. Ainsi, pour composer f et g , on applique d'abord f à x pour obtenir un élément de F , puis on applique g à ce résultat pour obtenir un élément de G .

Remarque 3.2.6 1. La définition précédente a bien un sens, car si on a $x \in E$, on a bien $f(x) \in F$ et il est donc loisible de considérer $g(f(x))$.

2. Pour que la composition $g \circ f$ soit bien définie, il est nécessaire que l'ensemble d'arrivée de f soit le même que l'ensemble de départ de g .

3. En général, $g \circ f \neq f \circ g$ où f et g sont deux applications de $\mathcal{F}(E)$.

4. On peut schématiser la composition de deux applications $f : E \longrightarrow F$ et $g : F \longrightarrow G$ de la manière suivante :

$$\begin{array}{ccccc} E & \xrightarrow{f} & F & \xrightarrow{g} & G \\ x & \mapsto & f(x) = y & \mapsto & g(y) \end{array}$$

$$\begin{array}{ccc} E & \xrightarrow{g \circ f} & G \\ x & \mapsto & g(f(x)) \end{array}$$

Exemple 3.2.7

1. Soient $f : \{1, 2, 3, 4\} \longrightarrow \{a, b, c, d\}$ et $g : \{a, b, c, d\} \longrightarrow \{1, 2, 3, 4\}$ les applications définies par les relations suivantes

$$f(1) = d \quad f(3) = a \quad g(a) = 3 \quad g(c) = 4$$

$$f(2) = b \quad f(4) = c \quad g(b) = 2 \quad g(d) = 1.$$

Comme l'ensemble d'arrivée de f est égal à l'ensemble de départ de g , l'application $g \circ f$ est bien définie. De plus, on laisse soin au lecteur de vérifier l'égalité suivante,

$$\forall x \in \{1, 2, 3, 4\} : g \circ f(x) = x.$$

2. Considérons les deux applications

$$\begin{array}{ccc} f : \mathbb{R} & \longrightarrow & \mathbb{R} \\ x & \mapsto & f(x) = x + 1, \end{array}$$

et

$$\begin{array}{ccc} g : \mathbb{R} & \longrightarrow & \mathbb{R} \\ x & \mapsto & g(x) = x^2. \end{array}$$

Pour tout réel x , on a

$$(g \circ f)(x) = g(f(x)) = (f(x))^2 = (x + 1)^2,$$

alors que

$$(f \circ g)(x) = f(g(x)) = g(x) + 1 = x^2 + 1,$$

et en particulier, $(g \circ f)(-1) = 0$ et $(f \circ g)(-1) = 2$. Donc ici, $g \circ f \neq f \circ g$.

3. Considérons un ensemble E , l'application

$$\begin{array}{ccc} f : \mathcal{P}(E) & \longrightarrow & \mathcal{P}(E) \\ A & \mapsto & f(A) = C_E^A, \end{array}$$

et

$$\begin{array}{ccc} g : \mathcal{P}(E) \times \mathcal{P}(E) & \longrightarrow & \mathcal{P}(E) \\ (A, B) & \mapsto & f(A, B) = A \cup B. \end{array}$$

Sur cet exemple, la composée $g \circ f$ n'a pas de sens, car la source de g qui est $\mathcal{P}(E) \times \mathcal{P}(E)$ ne coïncide pas avec le but de f qui est $\mathcal{P}(E)$. Par contre la source de f coïncide bien avec le but de g (dans les deux cas, il s'agit de $\mathcal{P}(E)$). La composée $f \circ g$ a donc un sens. Il s'agit de l'application

$$f \circ g : (A, B) \mapsto C_E^{A \cup B}.$$

D'après les lois de Morgan ensemblistes, et par définition de l'égalité entre applications, on peut aussi écrire

$$f \circ g : (A, B) \longmapsto C_E^A \cap C_E^B.$$

4. Considérons les deux applications

$$\begin{aligned} f : \mathbb{R} &\longrightarrow \mathbb{R}^+ \\ x &\longmapsto f(x) = x^2, \end{aligned}$$

et

$$\begin{aligned} g : \mathbb{R}^+ &\longrightarrow \mathbb{R} \\ x &\longmapsto g(x) = \sqrt{x}. \end{aligned}$$

Les deux composées $f \circ g$ et $g \circ f$ ont un sens. On a $f \circ g \in \mathcal{F}(\mathbb{R}^+, \mathbb{R}^+)$ et $g \circ f \in \mathcal{F}(\mathbb{R}, \mathbb{R})$. Plus précisément, on a

$$\forall x \in \mathbb{R}^+ : f(g(x)) = (\sqrt{x})^2 = x.$$

$$\forall x \in \mathbb{R} : g(f(x)) = \sqrt{x^2} = |x|.$$

Remarquons que la première ligne peut se lire $f \circ g = Id_{\mathbb{R}^+}$, mais qu'on a $g \circ f \neq Id_{\mathbb{R}}$.

5. Considérons les deux applications

$$\begin{aligned} f : \mathbb{N} &\longrightarrow \mathbb{N} \\ n &\longmapsto f(n) = n^2, \end{aligned}$$

et

$$\begin{aligned} g : \mathbb{N} &\longrightarrow \mathbb{N} \\ n &\longmapsto g(n) = n + 1. \end{aligned}$$

Les deux composées $f \circ g$ et $g \circ f$ ont un sens, elles appartiennent toutes les deux à $\mathcal{F}(\mathbb{N}, \mathbb{N})$. Plus précisément, on a

$$\begin{aligned} f \circ g : \mathbb{N} &\longrightarrow \mathbb{N} \\ n &\longmapsto (f \circ g)(n) = n^2 + 2n + 1, \end{aligned}$$

et

$$\begin{aligned} g \circ f : \mathbb{N} &\longrightarrow \mathbb{N} \\ n &\longmapsto (g \circ f)(n) = n^2 + 1. \end{aligned}$$

Définition 3.2.11 Soit $f : E \longrightarrow E$ une application, nous convenons de poser

$$f^0 = Id_E.$$

et pour tout $n \in \mathbb{N}^*$, la composée de f par elle-même n fois se note

$$f^n = \underbrace{f \circ f \circ \dots \circ f}_{n \text{ fois}} = f \circ f^{n-1} = f^{n-1} \circ f.$$

On a immédiatement

$$\forall n, m \in \mathbb{N}^* : \begin{cases} f^n \circ f^m = f^{n+m} \\ (f^n)^m = f^{nm} \end{cases}.$$

Exemple 3.2.8 Considérons l'application $f : \mathbb{R} \longrightarrow \mathbb{R}$ définie par $f(x) = x^2$. Alors

$$f^1(x) = f(x).$$

$$f^2(x) = f \circ f(x) = f(f(x)) = (x^2)^2 = x^4.$$

$$f^3(x) = f \circ f \circ f(x) = f(f^2(x)) = (x^4)^2 = x^8.$$

En général, on a

$$f^n(x) = \underbrace{f \circ f \circ \dots \circ f}_{n \text{ fois}}(x) = x^{2^n}.$$

2. Considérons l'application $f : \mathbb{R} \longrightarrow \mathbb{R}$ définie par $f(x) = x + 1$. Alors

$$f^1(x) = f(x) = x + 1$$

$$f^2(x) = f \circ f(x) = f(f(x)) = x + 2.$$

$$f^3(x) = f \circ f \circ f(x) = f(f^2(x)) = x + 3.$$

En général, on a

$$f^n(x) = \underbrace{f \circ f \circ \dots \circ f}_{n \text{ fois}}(x) = x + n.$$

Proposition 3.2.2

1. Soient E, F, G et H quatre ensembles non vides. Pour toutes applications $f : E \longrightarrow F$, $g : F \longrightarrow G$, $h : G \longrightarrow H$, on a

• L'opération $\circ$ est associative,

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

•

$$f \circ Id_E = Id_F \circ f = f.$$

Preuve.

1. • Les applications $(h \circ g) \circ f$ et $h \circ (g \circ f)$ ont le même ensemble de départ E , et le même ensemble d'arrivée H . D'autre part, Soit $x \in E$, alors

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

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

• Ces trois applications $f \circ Id_E$, $Id_F \circ f$, et f vont de E dans F . De plus, pour tout x dans E ,

$$(Id_F \circ f)(x) = (Id_F)(f(x)) = f(x),$$

et

$$(f \circ Id_E)(x) = f(Id_E(x)) = f(x).$$

On obtient $f \circ Id_E = Id_F \circ f = f$. □

3.2.5 Image directe et image réciproque

Définition 3.2.12 (Image directe)

1. Soient E et F deux ensembles non vides et f une application de E vers F . Soit A une partie de E . L'image directe de la partie A par l'application f , notée $f(A)$, est l'ensemble des images des éléments de A par f . Alors

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

Avec des quantificateurs, cela donne

$$\forall y \in F : (y \in f(A) \iff \exists x \in A : f(x) = y).$$

2. L'image de l'ensemble de départ, c'est-à-dire l'ensemble $f(E)$ de tous les éléments de la forme $f(x)$, s'appelle l'image de f (ou ensemble image de f), notée $Im(f)$,

$$Im(f) = f(E) = \{f(x) : x \in E\} \subset F.$$

Remarque 3.2.7 *L'image d'une partie de l'ensemble de départ est une partie de l'ensemble d'arrivée. Il ne faut pas confondre :*

- *L'image d'un élément x : c'est l'élément $f(x)$ appartenant à l'ensemble d'arrivée. En fait, on a*

$$f(\{x\}) = \{f(x)\}.$$

- *L'image d'une partie A : c'est une partie de l'ensemble d'arrivée.*

Exemple 3.2.9

1. *Si $A = \{a, b, c\}$, alors*

$$f(A) = \{f(a), f(b), f(c)\}.$$

2. *Soit*

$$\begin{aligned} f : \mathbb{Z} &\longrightarrow \mathbb{Z} \\ x &\longmapsto f(x) = x + 2. \end{aligned}$$

Prenons $A = \{-1, 0, 1, 2\} \subset \mathbb{Z}$. Alors l'image directe de A par f , notée $f(A)$, est l'ensemble des valeurs prises par f pour chaque élément de A . On obtient

$$f(A) = \{f(-1), f(0), f(1), f(2)\} = \{1, 2, 3, 4\}.$$

3. *Si*

$$\begin{aligned} f : \mathbb{R} &\longrightarrow \mathbb{R} \\ x &\longmapsto f(x) = x^2 \end{aligned}$$

et $A = [-1, 1]$, alors

$$f(A) = f([-1, 1]) = \{x^2 : x \in [-1, 1]\} = [0, 1],$$

et

$$\text{Im}(f) = f(\mathbb{R}) = \mathbb{R}^+.$$

4. *Soit*

$$\begin{aligned} f : \mathbb{N} &\longrightarrow \mathbb{N} \\ x &\longmapsto f(x) = x \bmod (3) \text{ (le reste de la division de } n \text{ par } 3) \end{aligned}$$

Prenons $A = \{3, 4, 5, 6, 7\}$ une partie de $\mathbb{N}$. L'image directe de A par f est

$$f(A) = \{f(3), f(4), f(5), f(6), f(7)\} = \{0, 1, 2, 0, 1\}.$$

Comme les ensembles n'ont pas de répétition, on peut écrire

$$f(A) = \{0, 1, 2\}.$$

Définition 3.2.13 (Image réciproque)

Soient E et F deux ensembles non vides et f une application de E vers F . Soit B une partie de F . L'image réciproque de la partie B par l'application f , notée $f^{-1}(B)$, est l'ensemble des antécédents des éléments de B par f .

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

Avec des quantificateurs, cela donne

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

Remarque 3.2.8 Attention, la notation f^{-1} est également utilisée pour décrire l'application réciproque (ou inverse) étudiée en analyse. L'application réciproque de f n'existe pas toujours. En revanche l'image réciproque d'une partie B par une application f existe toujours.

Exemple 3.2.10

1. Si

$$\begin{aligned} f : \mathbb{N} &\longrightarrow \mathbb{N} \\ x &\longmapsto f(x) = 2x + 1. \end{aligned}$$

Soit $B = \{5\}$, alors

$$f^{-1}(B) = \{x \in \mathbb{N} : f(x) \in B\} = \{x \in \mathbb{N} : f(x) = 5\} = \{2\}.$$

2. Si

$$\begin{aligned} f : \mathbb{R} &\longrightarrow \mathbb{R} \\ x &\longmapsto f(x) = x^2. \end{aligned}$$

Alors, on a pour $B = \{4\}$,

$$f^{-1}(B) = \{-2, 2\}.$$

En effet, $x \in f^{-1}(B)$ si et seulement si $f(x) = 4$, c'est-à-dire si et seulement si $x^2 = 4$. Il est alors aisé de montrer que les solutions de cette équation du second degré sont -2 et 2 .

Graphes : Sur un diagramme sagittal cela donne une présentation de type

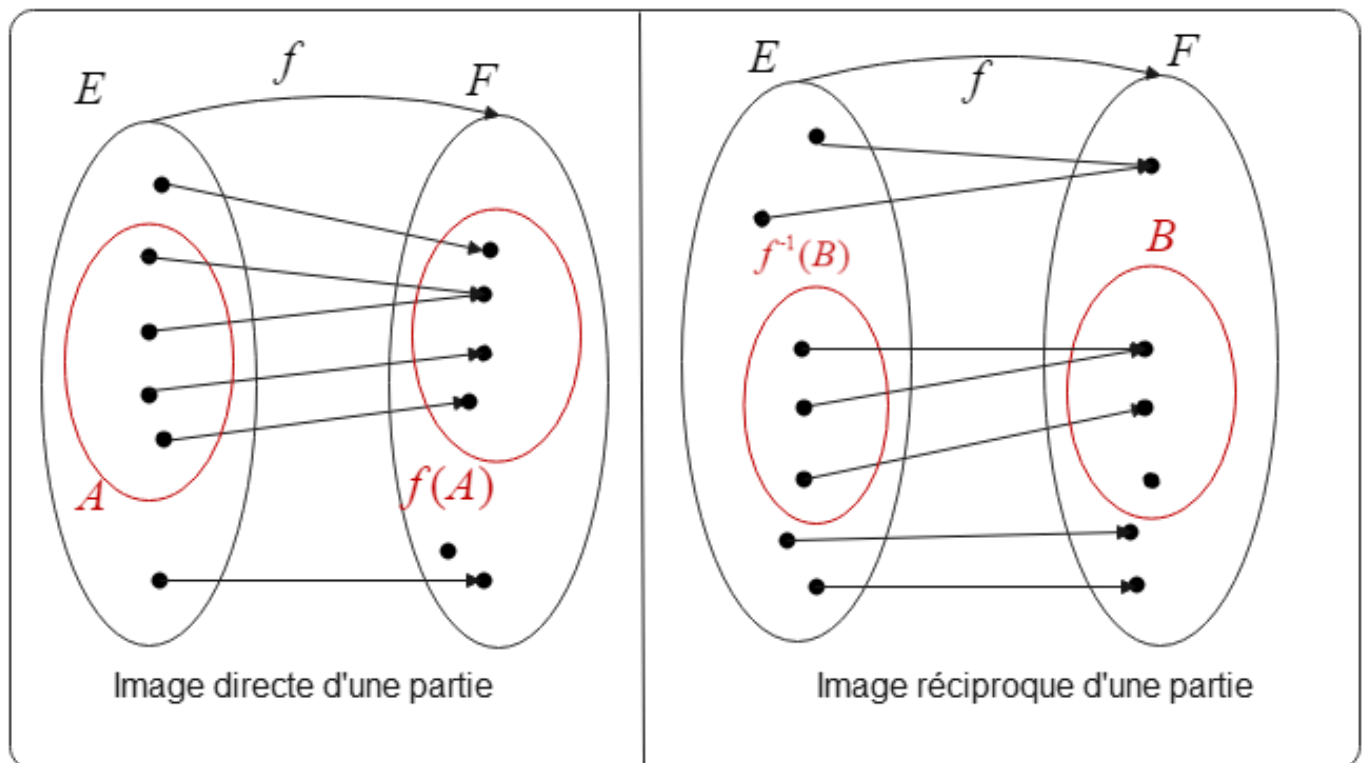


FIGURE 3.3: Image directe, Image réciproque

3.2.5.1 Propriétés de l'image directe et réciproque

Les concepts d'image directe et d'image réciproque d'une partie d'un ensemble par une application possèdent plusieurs propriétés importantes qui aident à comprendre comment une fonction agit sur des sous-ensembles. Examinons maintenant quelques propriétés immédiates et utiles de l'image directe.

Proposition 3.2.3 *Soient E et F deux ensembles non vides et f une application de E vers F . Alors*

1. *Image directe de l'ensemble vide et de l'ensemble de départ,*

$$f(\emptyset) = \emptyset, f(E) \subset F.$$

2. *Monotonie,*

$$\forall A, B \in \mathcal{P}(E) : A \subset B \implies f(A) \subset f(B).$$

3. *Image directe de l'Union,*

$$\forall A, B \in \mathcal{P}(E) : f(A \cup B) = f(A) \cup f(B).$$

4. *Image directe de l'Intersection,*

$$\forall A, B \in \mathcal{P}(E) : f(A \cap B) \subset f(A) \cap f(B).$$

5. *Image directe du complémentaire,*

$$\forall A \in \mathcal{P}(E) : f(C_E^A) \subset C_F^{f(A)}.$$

Preuve.

1. • On sait que

$$f(\emptyset) = \{f(x) : x \in \emptyset\}.$$

Puisque $\emptyset$ ne contient aucun élément, il n'y a aucun x pour lequel $f(x)$ peut être calculé. Par conséquent

$$f(\emptyset) = \{f(x) : x \in \emptyset\} = \emptyset.$$

• $f(E) \subset F$ est immédiat.

2. Soient A et B deux parties de E telles que $A \subset B$. Si $f(A) = \emptyset$, alors $f(A) \subset f(B)$. Sinon, soit $y \in f(A)$. Il existe $x \in A$ tel que $y = f(x)$. Puisque $A \subset B$, on a $x \in B$ et donc y est l'image par f d'un élément de B ou encore $y \in f(B)$. Ceci montre dans tous les cas que $f(A) \subset f(B)$.

3. Soient A et B deux parties de E . Soit $y \in F$, alors

$$\begin{aligned} y \in f(A \cup B) &\iff \exists x \in A \cup B : f(x) = y \\ &\iff \exists x \in E, (x \in A \text{ ou } x \in B : f(x) = y) \\ &\iff \exists x \in E, ((x \in A : f(x) = y) \text{ ou } (x \in B : f(x) = y)) \\ &\iff (\exists x \in E, x \in A : f(x) = y) \text{ ou } (\exists x \in E, x \in B : f(x) = y) \\ &\iff (\exists x \in A : f(x) = y) \text{ ou } (\exists x \in B : f(x) = y) \\ &\iff y \in f(A) \text{ ou } y \in f(B) \\ &\iff y \in f(A) \cup f(B). \end{aligned}$$

Donc, $f(A \cup B) = f(A) \cup f(B)$. (Cette démonstration est correcte même si $f(A \cup B) = \phi$).

4. Soient A et B deux parties de E . Alors

$$A \cap B \subset A$$

et donc

$$f(A \cap B) \subset f(A).$$

De même,

$$f(A \cap B) \subset f(B).$$

et donc

$$f(A \cap B) \subset f(A) \cap f(B).$$

5. Soit $y \in F$, alors

$$\begin{aligned} y \in f(C_E^A) &\implies \exists x \in C_E^A : f(x) = y \\ &\implies \exists x \notin A : f(x) = y \\ &\implies f(x) \notin f(A) \\ &\implies f(x) \in C_F^{f(A)}, \end{aligned}$$

ce qui prouve que $f(C_E^A) \subset C_F^{f(A)}$. □

Remarque 3.2.9

1. La réciproque de l'assertion $A \subset B \implies f(A) \subset f(B)$ n'est pas toujours vraie. Prenons $E = \{1, 2\}$, $F = \{1\}$ et $f : E \longrightarrow F$ définie par $f(1) = f(2) = 1$. Soient $A = \{1\}$ et $B = \{2\}$ deux parties de E . Alors $f(A) = f(B) = \{1\}$ alors que pourtant A n'est pas inclus dans B .

2. Il est possible que $f(A \cap B) \neq f(A) \cap f(B)$ (c'est-à-dire $f(A) \cap f(B) \not\subset f(A \cap B)$). Considérons par exemple, l'application $f : x \longmapsto x^2$ puis $A = \{-1\}$ et $B = \{1\}$. On a alors $A \cap B = \phi$ et donc $f(A \cap B) = \phi$ puis $f(A) \cap f(B) = \{1\} \cap \{1\} = \{1\}$. Dans ce cas $f(A \cap B) \neq f(A) \cap f(B)$.

3. On rappelle que quand on écrit

$$\exists x \in A : y = f(x) \text{ et } \exists x \in B : y = f(x),$$

la même lettre x utilisée dans chacun des deux morceaux de phrase ne désigne pas forcément un seul et même élément. Une écriture plus explicite est

$$\exists x_1 \in A : y = f(x_1) \text{ et } \exists x_2 \in B : y = f(x_2),$$

Dans ce cas, l'élément $y = f(x_1) = f(x_2)$ est dans $f(A) \cap f(B)$ mais pas nécessairement dans $f(A \cap B)$.

On a vu que l'image directe avait un comportement un peu délicat. Ce n'est pas le cas de l'image réciproque avec laquelle tout se passe bien.

Proposition 3.2.4 Soient E et F deux ensembles non vides et f une application de E vers F . Alors l'application image réciproque possède les propriétés suivantes :

1. Image réciproque du vide et de l'ensemble d'arrivée,

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

2. Monotonie,

$$\forall C, D \in \mathcal{P}(F) : C \subset D \implies f^{-1}(C) \subset f^{-1}(D).$$

3. Image réciproque de l'union,

$$\forall C, D \in \mathcal{P}(F) : f^{-1}(C \cup D) = f^{-1}(C) \cup f^{-1}(D).$$

4. Image réciproque de l'intersection,

$$\forall C, D \in \mathcal{P}(F) : f^{-1}(C \cap D) = f^{-1}(C) \cap f^{-1}(D).$$

5. Image réciproque du complémentaire,

$$\forall D \in \mathcal{P}(F) : f^{-1}(C_F^D) = C_E^{f^{-1}(D)}.$$

Autrement dit, l'image réciproque d'un complémentaire est égale au complémentaire de l'image réciproque.

Preuve.

1. • On sait que

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

La condition $f(x) \in \phi$ est impossible à satisfaire, car l'ensemble vide ne contient aucun élément. Par conséquent, il n'existe aucun $x \in E$ tel que $f(x) \in \phi$. Par conséquent

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

• On a

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

Par définition de l'application f , chaque $x \in E$ est envoyé par f dans un élément de F . Autrement dit, $f(x) \in F$ pour tout $x \in E$. Puisque cette condition est vraie pour tout $x \in E$, cela signifie que l'ensemble des x qui appartiennent à $f^{-1}(F)$ est tout E . Donc

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

2. On suppose que $C \subset D$. Prenons un élément $x \in f^{-1}(C)$, par définition, cela signifie que $f(x) \in C$ et comme $C \subset D$, si $f(x) \in C$, alors $f(x) \in D$. On obtient $x \in f^{-1}(D)$. On a donc montré que $f^{-1}(C) \subset f^{-1}(D)$.

3. Soient C et D deux parties de F . Soit $x \in E$. Alors

$$\begin{aligned} x \in f^{-1}(C \cup D) &\iff f(x) \in C \cup D \\ &\iff (f(x) \in C) \text{ ou } (f(x) \in D) \\ &\iff (x \in f^{-1}(C)) \text{ ou } (x \in f^{-1}(D)) \\ &\iff x \in f^{-1}(C) \cup f^{-1}(D). \end{aligned}$$

Ceci montre que $f^{-1}(C \cup D) = f^{-1}(C) \cup f^{-1}(D)$.

4. Soient C et D deux parties de F . Soit $x \in E$. Alors

$$\begin{aligned} x \in f^{-1}(C \cap D) &\iff f(x) \in C \cap D \\ &\iff (f(x) \in C) \text{ et } (f(x) \in D) \\ &\iff (x \in f^{-1}(C)) \text{ et } (x \in f^{-1}(D)) \\ &\iff x \in f^{-1}(C) \cap f^{-1}(D). \end{aligned}$$

Ceci montre que $f^{-1}(C \cap D) = f^{-1}(C) \cap f^{-1}(D)$.

5. Procédons par double inclusion.

- Si $f^{-1}(C_F^D) = \phi$, on a nécessairement l'inclusion $f^{-1}(C_F^D) \subset C_E^{f^{-1}(D)}$.
- Supposons que $f^{-1}(C_F^D) \neq \phi$ et soit $x \in f^{-1}(C_F^D)$, par définition de l'image réciproque, il existe donc $y \in C_F^D$ tel que $y = f(x)$. Mais puisque $y \notin D$, on a $f(x) \notin D$. Ainsi, l'ensemble des antécédents des éléments de D ne contient pas x ,

$$x \notin f^{-1}(D).$$

Ce qui signifie $x \in C_E^{f^{-1}(D)}$. On a donc montré que $f^{-1}(C_F^D) \subset C_E^{f^{-1}(D)}$.

Réciproquement,

- Si $C_E^{f^{-1}(D)} = \phi$, on a alors l'inclusion $C_E^{f^{-1}(D)} \subset f^{-1}(C_F^D)$.
- Supposons que $C_E^{f^{-1}(D)} \neq \phi$ et soit $x \in C_E^{f^{-1}(D)}$. Donc $x \notin f^{-1}(D)$, autrement dit aucun élément de D n'a x pour antécédent. Donc

$$f(x) \notin D.$$

Cela signifie que $f(x) \in C_F^D$. On a donc

$$x \in f^{-1}(C_F^D).$$

D'où l'inclusion $C_E^{f^{-1}(D)} \subset f^{-1}(C_F^D)$. On en déduit bien l'égalité souhaitée. $\square$

Remarque 3.2.10 *Ainsi, une partie A de E et la partie "aller-retour" $f^{-1}(f(A))$ ne coïncident pas nécessairement. On a cependant toujours une inclusion dans un sens comme le montre le théorème suivant.*

Théorème 3.2.1 *Soient E et F deux ensembles non vides et f une application de E vers F . Soient $A \in \mathcal{P}(E)$ et $B \in \mathcal{P}(F)$. Alors*

- $A \subset f^{-1}(f(A))$.
- $f(f^{-1}(B)) \subset B$.

Preuve.

- Soit $A \in \mathcal{P}(E)$. Soit $x \in E$, alors

$$x \in A \implies f(x) \in f(A)$$

$$\implies x \in f^{-1}(f(A)),$$

car les éléments de A sont des antécédents d'éléments de $f(A)$. Ceci montre que $A \subset f^{-1}(f(A))$.

- Soit $B \in \mathcal{P}(F)$. Choisissons un élément $y \in f(f^{-1}(B))$, par définition, cela signifie qu'il existe un $x \in E$ tel que $f(x) = y$ et $x \in f^{-1}(B)$. Cela signifie que $f(x) = y \in B$. Donc $f(f^{-1}(B)) \subset B$. $\square$

Les concepts d'images directes et réciproques peuvent être généralisés à une famille de parties d'un ensemble. Voici comment ces concepts s'appliquent dans ce cadre plus général.

Proposition 3.2.5 (Images directes et réciproques d'une famille des parties)

Soient E et F deux ensembles non vides et f une application de E vers F , et $(E_i)_{i \in I}$ une famille de sous-ensembles de E , $(F_i)_{i \in I}$ une famille de sous-ensembles de F . Alors,

1. *L'image directe de l'union de la famille est égale à l'union des images directes des ensembles individuels,*

$$f\left(\bigcup_{i \in I} E_i\right) = \bigcup_{i \in I} f(E_i).$$

2. L'image directe de l'intersection de la famille est incluse dans l'intersection des images directes des ensembles individuels,

$$f \left(\bigcap_{i \in I} E_i \right) \subset \bigcap_{i \in I} f(E_i).$$

3. L'image réciproque de l'union de la famille est égale à l'union des images réciproques des ensembles individuels,

$$f^{-1} \left(\bigcup_{i \in I} F_i \right) = \bigcup_{i \in I} f^{-1}(F_i).$$

4. L'image réciproque de l'intersection de la famille est égale à l'intersection des images réciproques des ensembles individuels,

$$f^{-1} \left(\bigcap_{i \in I} F_i \right) = \bigcap_{i \in I} f^{-1}(F_i).$$

3.2.6 Types d'applications

Les applications entre ensembles peuvent être classées en différentes catégories selon la manière dont les éléments de l'ensemble de départ sont associés à ceux de l'ensemble d'arrivée. Les trois types principaux sont : les injections, les surjections et les bijections. Ces concepts sont essentiels dans l'étude des applications.

3.2.6.1 Applications injectives

L'injectivité est une propriété particulièrement intéressante, souvent essentielle pour définir certains modèles. Par exemple, lors de l'établissement d'une communication téléphonique, on cherche a priori à exclure la possibilité que deux numéros de téléphone soient associés à la même carte SIM (Subscriber Identity Module). La table de correspondance utilisée par les opérateurs n'est rien d'autre qu'une application injective. De la même manière, chaque numéro de sécurité sociale distinct est associé à un individu unique, illustrant encore cette notion d'injectivité.

Définition 3.2.14 (*Application injective ou injection*)

Soient E et F deux ensembles non vides.

1. On dit qu'une application f de E dans F est injective ou une injection si tout élément de F admet au plus un antécédent dans E .
2. On peut aussi formuler cela de la façon suivante : deux éléments distincts de l'ensemble de départ E ont des images distinctes par f dans l'ensemble d'arrivée F , ce qui s'écrit avec des quantificateurs

$$\forall x_1, x_2 \in E : x_1 \neq x_2 \implies f(x_1) \neq f(x_2).$$

Par passage à la contraposée, cette définition est équivalente à

$$\forall x_1, x_2 \in E : f(x_1) = f(x_2) \implies x_1 = x_2. \quad (3.1)$$

3. De manière équivalente, pour tout $y \in F$, l'équation $y = f(x)$ d'inconnue x admet au plus une solution. Autrement dit, chaque élément $y \in F$ admet au plus un antécédent.

- On note $\text{Inj}(E, F)$ l'ensemble des injections de E vers F .

Remarque 3.2.11*1. La première équivalence*

$$f \text{ injective} \iff \forall x_1, x_2 \in E : x_1 \neq x_2 \implies f(x_1) \neq f(x_2)$$

est simple à comprendre. Elle dit que deux éléments différents ont des images différentes. Elle présente néanmoins un défaut : elle utilise la relation $\neq$. La relation $\neq$ a de nombreuses propriétés désagréables : elle n'est pas transitive (on a $0 \neq 1$ et $1 \neq 0$ mais on n'a pas $0 \neq 0$), elle n'est pas compatible avec l'addition (on a $0 \neq 1$ et $1 \neq 0$ mais on n'a pas $0 + 1 \neq 1 + 0$). Pour cette raison, on préfère utiliser la contraposée de l'implication précédente

$$f \text{ injective} \iff \forall x_1, x_2 \in E : f(x_1) = f(x_2) \implies x_1 = x_2,$$

pour travailler avec la relation $=$ beaucoup plus facile à manipuler. La phrase obtenue dit que si deux éléments ont même image, ils ne peuvent qu'être égaux. Cette phrase est peut-être moins claire mais beaucoup plus agréable à manipuler dans la pratique.

2. Dans la plupart des cas, pour démontrer qu'une application est injective on utilise plutôt l'assertion 3.1 ci-dessus. C'est-à-dire, on considère x_1 et x_2 dans E vérifiant $f(x_1) = f(x_2)$, et on cherche à prouver que $x_1 = x_2$.

3. Dans le cas d'un diagramme sagittal, une fonction n'est pas injective si deux flèches arrivent sur le même élément de l'ensemble d'arrivée. Lorsqu'on a fait un tableau, la fonction n'est pas injective lorsqu'il y a deux étoiles dans la même colonne.

4. Dans le cas d'une application donnée par une formule, on résout dans E , pour un y quelconque dans F , l'équation d'inconnue x , $y = f(x)$. Si, pour tout y dans F , elle admet au plus une solution dans E , alors f est une injection de E dans F .

5. Graphiquement, une application est injective si et seulement si toute droite horizontale coupe la courbe représentative de cette application en au plus un point.

6. Notons que dans la définition d'une injection, on aurait pu écrire

$$\forall x_1, x_2 \in E : f(x_1) = f(x_2) \iff x_1 = x_2.$$

Mais on ne l'a pas fait car seule l'implication de gauche à droite fait de f une injection (même si l'implication de droite à gauche est vraie).

7. Pour une preuve de la non injectivité d'une application f , on montre la négation de la définition en exhibant un exemple

$$\exists x_1, x_2 \in E : x_1 \neq x_2 \text{ et } f(x_1) = f(x_2).$$

8. L'implication $(x_1 = x_2 \implies f(x_1) = f(x_2))$ est vérifiée par n'importe quelle application.

Grphe : Sur un diagramme sagittal cela donne une présentation de type

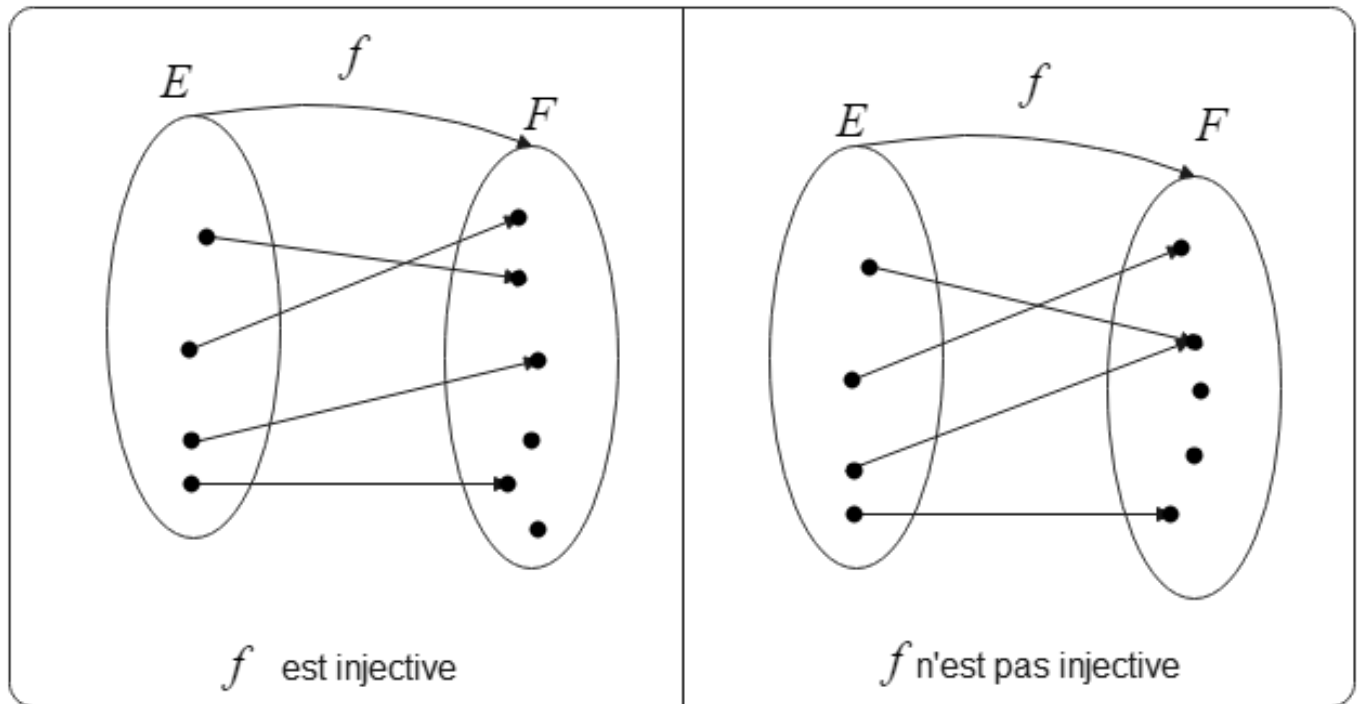


FIGURE 3.4: Application injective, non injective

Exemple 3.2.11

1. Pour tout ensemble E , l'application identique Id_E (aussi appelée application identité) est injective. En effet, Soient $x_1, x_2 \in E$ tels que $Id_E(x_1) = Id_E(x_2)$. Par définition de l'application identité, cela revient à $x_1 = x_2$. Donc, Id_E est injective.

2. Montrons que l'application de $\mathbb{C} \setminus \{i\}$ dans $\mathbb{C}$, définie par $f(z) = \frac{z-1}{z-i}$ est injective.

Solution.

Il est clair que f est bien une application. Soit $(z_1, z_2) \in (\mathbb{C} \setminus \{i\})^2$. Alors

$$\begin{aligned} f(z_1) = f(z_2) &\implies \frac{z_1 - 1}{z_1 - i} = \frac{z_2 - 1}{z_2 - i} \\ &\implies (1 - i)z_1 = (1 - i)z_2 \\ &\implies z_1 = z_2. \end{aligned}$$

On a montré que

$$\forall z_1, z_2 \in \mathbb{C} \setminus \{i\} : f(z_1) = f(z_2) \implies z_1 = z_2.$$

L'application f est donc injective.

3. Montrons que $f : \mathbb{R} \longrightarrow \mathbb{R}$ définie par $f(x) = x^2$ n'est pas injective sur $\mathbb{R}$, et $g : \mathbb{R}^+ \longrightarrow \mathbb{R}$ définie par $g(x) = x^2$ est injective sur $\mathbb{R}^+$.

Solution.

a. Pour f c'est évident puisque par exemple

$$f(-1) = 1 = f(1) \text{ et } -1 \neq 1.$$

Ceci prouve que f n'est pas injective sur $\mathbb{R}$.

b. Pour montrer que g est injective on se donne $x_1, x_2 \in \mathbb{R}^+$ tel que $g(x_1) = g(x_2)$. On a donc

$$x_1^2 = x_2^2$$

alors

$$x_1 = x_2 \text{ ou } x_1 = -x_2.$$

Comme x_1 et x_2 sont des réels positifs, on peut en conclure que $x_1 = x_2$. Ceci prouve que g est injective sur $\mathbb{R}^+$.

4. Montrons que l'application de $\mathbb{R}$ dans $\mathbb{R}$, définie par $f(x) = 2x + 1$ est injective.

Solution.

Première méthode. Soit $y \in \mathbb{R}$, résolvons l'équation (d'inconnue x) $y = 2x + 1$. On trouve une (unique) solution

$$x = \frac{y - 1}{2}.$$

L'application f est donc injective.

Deuxième méthode. Soient $x_1, x_2 \in \mathbb{R}$ tel que $f(x_1) = f(x_2)$. On a donc

$$2x_1 + 1 = 2x_2 + 1,$$

qui conduit bien sûr à $x_1 = x_2$. L'injectivité est donc établie.

5. Montrons que l'application de $\mathbb{R}$ dans $\mathbb{R}$ définie par $f(x) = \sin x$ n'est pas injective.

Solution.

Prenons $x_1 = 0$ et $x_2 = 2\pi$. On vérifie les relations $x_1 \neq x_2$ et $f(x_1) = f(x_2) = 0$. Alors le nombre 0 a deux antécédents. L'application f n'est pas injective.

Un autre exemple fondamental des applications injectives est l'injection canonique.

Définition 3.2.15 (Injection canonique)

Soit E un ensemble et X une partie de E . L'application

$$\begin{aligned} i_X : X &\longrightarrow E \\ x &\longmapsto i(x) = x, \end{aligned}$$

qui à x associe x est injective, appelée l'injection canonique i_X (ou inclusion canonique) de X dans E .

Lorsque $X = E$, l'injection canonique n'est autre que l'application identité de E ($i_X = Id_E$).

Proposition 3.2.6 (Composée d'injections). Soient E, F, G trois ensembles non vides, f une application de E vers F , g une application de F vers G . Si f et g sont injectives, alors $g \circ f$ est injective. Autrement dit la composée de deux injections est une injection.

Preuve. Soit $(x_1, x_2) \in E^2$ tels que

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

On a donc

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

Comme g est injective, alors

$$f(x_1) = f(x_2),$$

mais f est aussi injective, donc

$$x_1 = x_2.$$

D'où la conclusion. □

La réciproque n'est que partiellement vraie.

Proposition 3.2.7 Soient E, F, G trois ensembles non vides, f une application de E vers F et g une application de F vers G . Si $g \circ f$ est injective, alors f est injective.

Preuve. Supposons que $g \circ f$ est injective. Soit $(x_1, x_2) \in E^2$. Alors

$$\begin{aligned} f(x_1) = f(x_2) &\implies g(f(x_1)) = g(f(x_2)) \text{ (car } g \text{ est une application)} \\ &\implies (g \circ f)(x_1) = (g \circ f)(x_2) \\ &\implies x_1 = x_2 \text{ (car } g \circ f \text{ est injective)}. \end{aligned}$$

On a montré que

$$\forall (x_1, x_2) \in E^2 : f(x_1) = f(x_2) \implies x_1 = x_2,$$

f est donc injective. □

Proposition 3.2.8 Soit f une application de E vers F . Alors

$$f \text{ est injective} \iff \forall A \in \mathcal{P}(E) : A = f^{-1}(f(A)).$$

Preuve.

1. Montrons l'implication directe.

• On sait que pour tout $A \in \mathcal{P}(E)$, on a

$$A \subset f^{-1}(f(A)).$$

Donc il suffit de montrer que $f^{-1}(f(A)) \subset A$.

• Soit $x \in f^{-1}(f(A))$, alors $f(x) \in f(A)$ et par suite il existe $a \in A$ tel que $f(x) = f(a)$. Comme f est injective alors $x = a$, donc $x \in A$, par suite

$$f^{-1}(f(A)) \subset A.$$

2. Montrons l'implication réciproque.

Soient x_1 et x_2 deux éléments de E tels que $f(x_1) = f(x_2)$, alors avec $A = \{x_1\}$ on a $f(A) = \{f(x_1)\} = \{f(x_2)\}$, et par suite $x_2 \in f^{-1}(f(A))$. Comme $A = f^{-1}(f(A))$, alors il résulte que $x_2 \in A = \{x_1\}$, par suite $x_1 = x_2$. □

3.2.6.2 Applications surjectives

La surjectivité est une propriété particulièrement intéressante, notamment dans l'étude des applications entre ensembles. Elle garantit que chaque élément de l'ensemble d'arrivée F est atteint par au moins un élément de l'ensemble de départ E . Autrement dit, pour chaque $y \in F$, il existe au moins un $x \in E$ tel que $f(x) = y$. En d'autres termes, la surjectivité assure qu'il n'y a pas de "trous" dans l'ensemble d'arrivée F .

Définition 3.2.16 Soient E et F deux ensembles non vides.

1. On dit qu'une application f de E dans F est surjective ou une surjection de E dans F , si tout élément de F admet au moins un antécédent dans E , ce qui s'écrit avec des quantificateurs

$$\forall y \in F, \exists x \in E : f(x) = y.$$

2. De manière équivalente, f est surjective si et seulement si, son image est égale à son ensemble d'arrivée,

$$f \text{ est surjective} \iff \text{Im}(f) = f(E) = F.$$

3. On peut aussi formuler cela de la façon suivante, pour tout $y \in F$, l'équation $y = f(x)$ d'inconnue x admet au moins une solution. Autrement dit, chaque élément $y \in F$ admet au moins un antécédent.

- Nous notons $\text{Surj}(E, F)$ l'ensemble des surjections de E vers F .

Graphe : Sur un diagramme sagittal cela donne une présentation de type

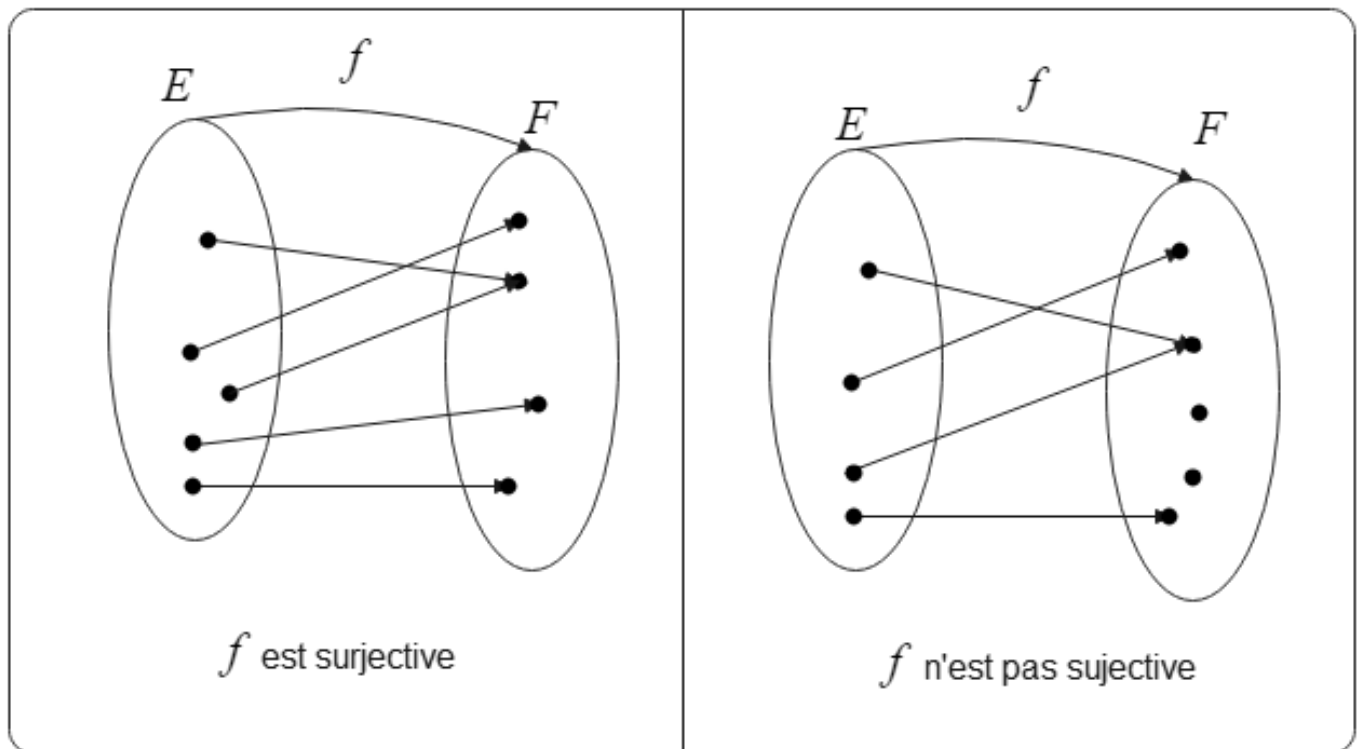


FIGURE 3.5: Application surjective, non surjective

Remarque 3.2.12

1. La surjectivité dépend essentiellement de l'ensemble d'arrivée. Alors une application f devient surjective si on limite l'ensemble d'arrivée à $\text{Im}(f)$.
3. On doit prendre garde à ne pas transformer la définition correcte d'une surjection (à savoir $\forall y \in F, \exists x \in E : f(x) = y$) en $\forall x \in F, \exists y \in E : f(x) = y$. La conclusion est la même mais cette dernière phrase ne signifie pas que f est surjective. En fait, elle ne signifie même pas que f est une application.
3. De manière générale, montrer qu'une application $f : E \longrightarrow F$ est surjective consiste à montrer que pour tout y de F , il existe au moins un élément x de E tel que $y = f(x)$. On se donne donc un élément y de F et on montre que l'équation $y = f(x)$, d'inconnue x , a au moins une solution x dans E , en résolvant cette équation par exemple ou autrement.
4. Pour montrer qu'une application $f : E \longrightarrow F$ n'est pas surjective, il suffit de trouver un élément dans l'ensemble d'arrivée qui n'est l'image d'aucun élément de E ,

$$\exists y \in F, \forall x \in E : f(x) \neq y.$$

Exemple 3.2.12

1. Pour tout ensemble E , l'application identique Id_E (aussi appelée application identité) est surjective. En effet, prenons un élément $y \in E$. Par définition de Id_E , on a $\text{Id}_E(y) = y$. Ainsi, pour chaque élément $y \in E$, il existe un élément $x = y$ dans E tel que $\text{Id}_E(x) = y$.
2. Soit E un ensemble et $A \in \mathcal{P}(E)$ tel que $A \notin \{\emptyset, E\}$, alors l'application φ_A indicatrice de A est

surjective. Pour trouver un antécédent de 1, il suffit de considérer $x \in A$ (possible car $A \neq \emptyset$). Pour un antécédent de 0 il suffit de considérer $x \in C_E^A$ (possible car $A \neq E$).

3. Soit $f : \mathbb{Z} \rightarrow \mathbb{N}$ définie par $f(x) = x^2$, alors f n'est pas surjective. En effet il existe des éléments $y \in \mathbb{N}$ qui n'ont aucun antécédent, par exemple, $y = 2$. Supposons que 2 possède un antécédent x par f . Alors on a

$$\begin{aligned} f(x) = 2 &\implies x^2 = 2 \\ &\implies x = \sqrt{2} \text{ ou } x = -\sqrt{2}. \end{aligned}$$

Mais alors x n'est pas un entier de $\mathbb{Z}$. Donc $y = 2$ n'a pas d'antécédent et f n'est pas surjective.

4. L'application $f : x \in \mathbb{R} \mapsto x^2 \in \mathbb{R}^+$ est surjective. En effet si $y \in \mathbb{R}^+$, l'équation $y = x^2$ d'inconnue x admet au moins une solution $x = \sqrt{y}$. En revanche, $x \in \mathbb{R} \mapsto x^2 \in \mathbb{R}$ n'est pas surjective puisque pour $y < 0$, l'équation $y = x^2$ n'admet pas de solution dans $\mathbb{R}$.

5. Considérons l'application $f : \mathbb{R}^2 \rightarrow \mathbb{R}$ définie par $f(x, y) = x + y$. Cette application est surjective car pour tout $z \in \mathbb{R}$, on peut trouver $(x, y) \in \mathbb{R}^2$ tel que $x + y = z$. Par exemple, $x = z$ et $y = 0$ est une solution.

6. Soit $A \in \mathcal{P}(E)$ et soit l'application

$$\begin{aligned} f_A : \mathcal{P}(E) &\longrightarrow \mathcal{P}(E) \\ X &\longmapsto f_A(X) = X \cap A. \end{aligned}$$

Montrons que si $A \neq E$, alors l'application f_A n'est pas surjective.

Solution.

- Supposons que $A = E$, alors f_A est l'identité de $\mathcal{P}(E)$, elle est donc surjective.
- Supposons que $A \neq E$, alors, pour tout $X \in \mathcal{P}(E)$, on a $f_A(X) \subset A$ et donc $f_A(X) \neq E$. Ainsi E n'a pas d'antécédent par f_A , et cette application n'est pas surjective.

Un autre exemple fondamental des applications surjectives est la surjection canonique dans le cadre de l'ensemble des parties d'un ensemble donné. La surjection canonique ou projection canonique donc est la surjection associée à une relation d'équivalence sur un ensemble.

Définition 3.2.17 (Surjection canonique)

Soit $\mathcal{R}$ une relation d'équivalence sur un ensemble E .

1. L'application

$$\begin{aligned} s : E &\longrightarrow E/\mathcal{R} \\ x &\longmapsto s(x) = \dot{x} \end{aligned}$$

est surjective, appelée surjection canonique de E dans $E/\mathcal{R}$ associée à $\mathcal{R}$ où $\dot{x}$ désigne la classe d'équivalence de x .

2. $s(\cdot)$ est l'application définie de E dans l'ensemble quotient $E/\mathcal{R}$, qui à chaque élément de E associe sa classe d'équivalence.

Proposition 3.2.9 (Composée de surjections). Soient E, F, G trois ensembles non vides, f une application de E vers F g une application de F vers G . Si f et g sont surjectives, alors $g \circ f$ est surjective. Autrement dit la composée de deux surjections est une surjection.

Preuve. Soit $z \in G$. Puisque g est surjective, il existe un élément y de F tel que $g(y) = z$. Puis, comme f est aussi surjective, il existe donc un élément x de E tel que $f(x) = y$. Mais alors,

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

et x est un antécédent de z par $g \circ f$. Ce qui montre que $g \circ f$ est surjective. □

La réciproque n'est que partiellement vraie.

Proposition 3.2.10 Soient E, F, G trois ensembles non vides, f une application de E vers F et g une application de F vers G . Si $g \circ f$ est surjective, alors g est surjective.

Preuve. Supposons que $g \circ f$ est surjective. Soit $z \in G$. Il existe un élément x de E tel que $(g \circ f)(x) = z$. Posons $y = f(x)$, alors y est un élément de F tel que $g(y) = z$. On a montré que

$$\forall z \in G : \exists y \in F : g(y) = z,$$

g est donc surjective. □

Proposition 3.2.11 Soit f une application de E vers F . Alors

$$f \text{ est surjective} \iff \forall B \in \mathcal{P}(F) : B = f(f^{-1}(B)).$$

Preuve.

1. Montrons l'implication directe.

• On sait que pour tout $B \in \mathcal{P}(F)$, on a

$$f(f^{-1}(B)) \subset B.$$

Donc il suffit de montrer que $B \subset f(f^{-1}(B))$.

• Soit $y \in B$, alors comme f est surjective il existe $x \in E$ tel que $f(x) = y$. On a $x \in f^{-1}(B)$ et ainsi $f(x) = y \in f(f^{-1}(B))$. Par conséquent

$$B \subset f(f^{-1}(B)).$$

2. Montrons l'implication réciproque.

Soit $y \in F$, alors avec $B = \{y\}$ on a par hypothèse $f(f^{-1}(\{y\})) = \{y\}$. Donc, $f^{-1}(\{y\})$ n'est pas vide, et y admet un antécédent par f . L'application f est alors surjective. □

3.2.6.3 Applications bijectives

La bijectivité est une propriété fondamentale en mathématiques, car elle permet d'établir une correspondance parfaite entre deux ensembles. Elle combine les deux caractéristiques essentielles d'une application : l'injectivité et la surjectivité. Ainsi, chaque élément de l'ensemble de départ est associé à un unique élément de l'ensemble d'arrivée, et vice versa, garantissant ainsi une relation réciproque complète entre les deux ensembles.

Définition 3.2.18 Soient E et F deux ensembles non vides, et f une application de E vers F .

1. On dit que f est une application bijective ou une bijection de E dans F , si tout élément de F admet exactement un antécédent dans E , ce qui s'écrit avec des quantificateurs

$$\forall y \in F, \exists ! x \in E : f(x) = y.$$

2. De cette définition, il découle que f est bijective si et seulement si f est à la fois injective et surjective.

3. On peut aussi formuler cela de la façon suivante, pour tout $y \in F$, l'équation $y = f(x)$ d'inconnue x admet exactement une solution.

• Nous notons $\text{Bij}(E, F)$ l'ensemble des bijections de E vers F .

Grphe : Sur un diagramme sagittal cela donne une présentation de type

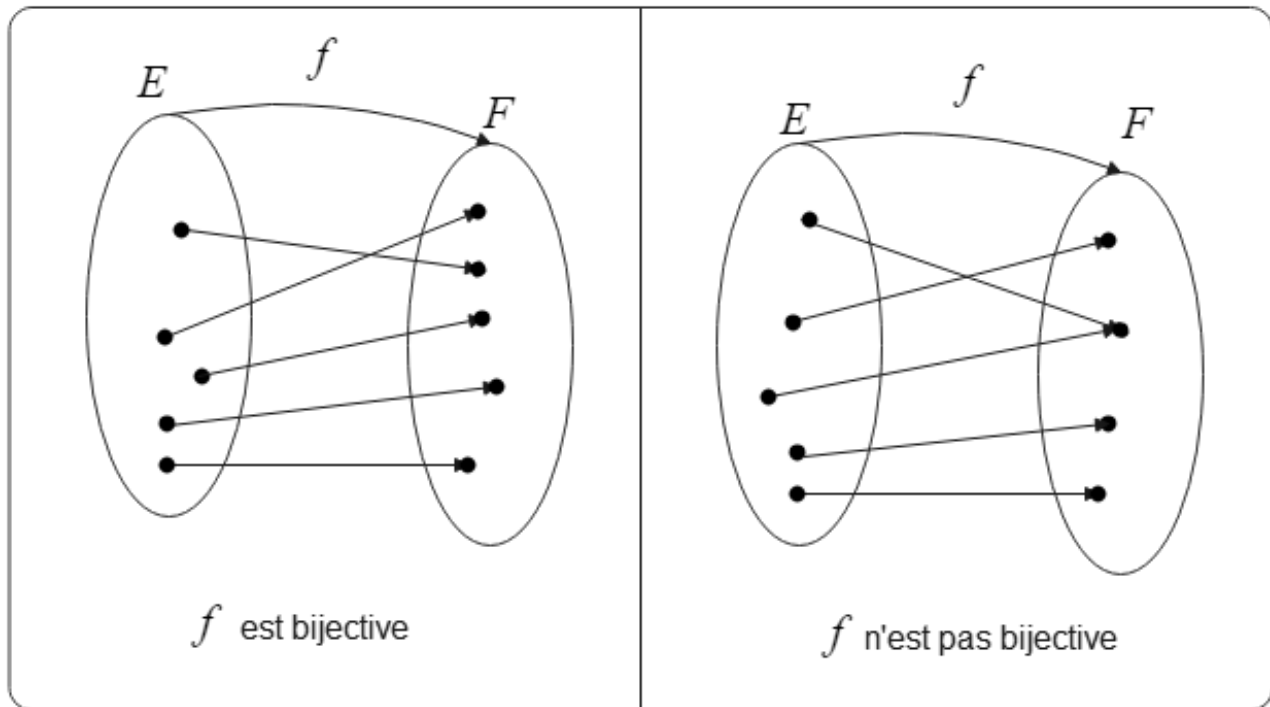


FIGURE 3.6: Application bijective, non bijective

Remarque 3.2.13

1. Ainsi, f est bijective si et seulement si tout élément de F admet exactement un antécédent. Pour cette raison, une bijection est parfois aussi appelée correspondance 1 à 1 (issu de la terminologie anglaise « one-to-one correspondance »).
2. Pour qu'une application ne soit pas bijective, il suffit qu'elle ne soit pas injective ou qu'elle ne soit pas surjective.

Exemple 3.2.13

1. Pour tout ensemble E , l'application identité $Id_E : E \longrightarrow E$ est bijective. Elle est injective car $Id_E(x_1) = Id_E(x_2)$ implique $x_1 = x_2$. Elle est surjective car pour tout $y \in E$, il existe $x = y$ tel que $Id_E(y) = y$.
2. L'application $f : \mathbb{R} \longrightarrow \mathbb{R}$ définie par $f(x) = 2x + 1$ est bijective, puisque pour tout réel y , il existe exactement une solution réelle de l'équation $y = 2x + 1$ d'inconnue x , à savoir $x = \frac{1}{2}(y - 1)$.
3. L'application exponentielle $f : x \longmapsto \exp(x)$ est une application bijective de $\mathbb{R}$ dans $\mathbb{R}_+^*$. En effet,
 - **Injectivité :** Si $\exp(x_1) = \exp(x_2)$, alors $x_1 = x_2$.
 - **Surjectivité :** Pour tout $y > 0$, il existe $x \in \mathbb{R}$ tel que $y = \exp(x)$, donc $x = \ln(y)$.
 Par contre si on définit l'application de $\mathbb{R}$ dans $\mathbb{R}$, l'application ainsi définie n'est pas surjective (un réel négatif n'a pas d'antécédent), elle n'est pas bijective.
4. L'application logarithme népérien définie par

$$\begin{aligned} f : \mathbb{R}_+^* &\longrightarrow \mathbb{R} \\ x &\longmapsto f(x) = \ln x, \end{aligned}$$

est bijective.

5. Soit E un ensemble et $\mathcal{P}(E)$ l'ensemble de toutes les parties de E . Alors l'application complémentaire définie par

$$\begin{aligned} f : \mathcal{P}(E) &\longrightarrow \mathcal{P}(E) \\ X &\longmapsto f(X) = C_E^X \end{aligned}$$

est bijective. En effet,

• **Injectivité :** Supposons que $f(X_1) = f(X_2)$ pour deux parties X_1, X_2 de E . Cela signifie que $C_E^{X_1} = C_E^{X_2}$, alors

$$C_E(C_E^{X_1}) = C_E(C_E^{X_2}) \implies X_1 = X_2,$$

ce qui prouve que f est injective.

• **Surjectivité :** Pour prouver que f est surjective, considérons une partie quelconque $Y \subset E$. Nous devons montrer qu'il existe une partie X telle que $f(X) = Y$. Prenons simplement $X = C_E^Y$. Alors,

$$f(X) = C_E(C_E^Y) = Y.$$

Ainsi, chaque partie de E est l'image d'une autre partie par f , ce qui montre que f est surjective.

6. Montrons que l'application

$$\begin{aligned} f : \mathbb{R} &\longrightarrow]-1, 1[\\ x &\longmapsto f(x) = \frac{x}{1+|x|} \end{aligned}$$

est une bijection.

Solution.

Tout d'abord, pour tout réel x , $1 + |x| \geq 1 > 0$, et en particulier, $1 + |x| \neq 0$. Ainsi, pour tout réel x , $f(x)$ existe. De plus, pour $x \in \mathbb{R}$,

$$|f(x)| = \frac{|x|}{1 + |x|} < \frac{1 + |x|}{1 + |x|} = 1$$

et donc, pour tout réel x , $-1 < f(x) < 1$. Par suite, f est bien une application de $\mathbb{R}$ vers $] -1, 1[$.

• **Montrons maintenant que f est bijective.** C'est-à-dire pour tout réel y de $] -1, 1[$, il existe un et un seul réel $x \in \mathbb{R}$ tel que $f(x) = y$. Soit $y \in] -1, 1[$. Pour x réel, si $f(x) = y$, alors, puisque

$$\frac{1}{1 + |x|} > 0,$$

il est nécessaire que x et y aient même signe (Dire que les deux réels x et y sont de même signe signifie que x et y sont simultanément strictement négatifs, simultanément nuls et simultanément strictement positifs). Nous avons donc deux cas.

Cas 1. Soit $y \in [0, 1[$. Pour $x \in \mathbb{R}$,

$$\begin{aligned} f(x) = y &\iff \begin{cases} x \geq 0 \\ \text{et} \\ \frac{x}{1+|x|} = y \end{cases} \\ &\iff \begin{cases} x \geq 0 \\ \text{et} \\ x = y(1+x) \quad (\text{car } 1+x \neq 0 \text{ pour } x \geq 0) \end{cases} \\ &\iff \begin{cases} x \geq 0 \\ \text{et} \\ x(1-y) = y \end{cases} \\ &\iff \begin{cases} x \geq 0 \\ \text{et} \\ x = \frac{y}{1-y} \quad (\text{car } 1-y \neq 0) \end{cases} \end{aligned}$$

$$\Longleftrightarrow x = \frac{y}{1-y} \text{ (car } \frac{y}{1-y} \geq 0).$$

Ainsi, pour tout réel $y \in [0, 1[$, il existe un et un seul réel x tel que $f(x) = y$, à savoir $x = \frac{y}{1-y}$.

Cas 2. Soit $y \in]-1, 0[$. Pour $x \in \mathbb{R}$,

$$\begin{aligned} f(x) = y &\Longleftrightarrow \begin{cases} x < 0 \\ \text{et} \\ \frac{x}{1-x} = y \end{cases} \\ &\Longleftrightarrow \begin{cases} x < 0 \\ \text{et} \\ x = y(1-x) \text{ (car } 1-x \neq 0 \text{ pour } x < 0) \end{cases} \\ &\Longleftrightarrow \begin{cases} x < 0 \\ \text{et} \\ x(1+y) = y \end{cases} \\ &\Longleftrightarrow \begin{cases} x < 0 \\ \text{et} \\ x = \frac{y}{1+y} \text{ (car } 1+y \neq 0) \end{cases} \\ &\Longleftrightarrow x = \frac{y}{1+y} \text{ (car } \frac{y}{1+y} < 0). \end{aligned}$$

Ainsi, pour tout réel $y \in]-1, 0[$, il existe un et un seul réel x tel que $f(x) = y$, à savoir $x = \frac{y}{1+y}$.

En résumé, pour tout réel y élément de $]-1, 1[$, il existe un et un seul réel x tel que $f(x) = y$, à savoir

$$x = \frac{y}{1-|y|}.$$

Par suite, f est une bijection de $\mathbb{R}$ sur $]-1, 1[$.

Proposition 3.2.12 (Composée de bijections) Soient E, F, G trois ensembles non vides, f une application de E vers F et g une application de F vers G . Si f et g sont bijectives. Alors $g \circ f$ est bijective. Autrement dit la composée de deux bijections est une bijection.

Preuve. Puisque f et g sont bijectives, alors elles sont injectives et surjectives. Donc, d'après ce qui précède, $g \circ f$ est injective et surjective c'est-à-dire bijective. $\square$

3.2.7 Application réciproque d'une application bijective

Lorsqu'une application $f : E \longrightarrow F$ est bijective, il est possible de définir une nouvelle application appelée application réciproque ou inverse, notée $f^{-1} : F \longrightarrow E$. Cette application inverse permet de revenir de F vers E . L'existence d'une application réciproque est garantie uniquement lorsque l'application originale est bijective. Dans notre représentation graphique, il suffit d'inverser le sens de la flèche symbolisant l'application.

Définition 3.2.19 (Application réciproque d'une application bijective)

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

1. L'application de F dans E , qui à tout élément $y \in F$ de l'ensemble d'arrivée de f , associe son unique antécédent $x \in E$ tel que $y = f(x)$ s'appelle application réciproque de f et se note f^{-1} .
2. La réciproque de f est définie par

$$\forall (x, y) \in E \times F : y = f(x) \Longleftrightarrow x = f^{-1}(y).$$

Exemple 3.2.14

1. *L'application*

$$\begin{aligned} f : \mathbb{R}^+ &\longrightarrow \mathbb{R}^+ \\ x &\longmapsto f(x) = x^2 \end{aligned}$$

est bijective. Sa réciproque est l'application

$$\begin{aligned} f^{-1} : \mathbb{R}^+ &\longrightarrow \mathbb{R}^+ \\ x &\longmapsto f^{-1}(x) = \sqrt{x}. \end{aligned}$$

car

$$\forall (x, y) \in \mathbb{R}^{+2} : y = x^2 \iff x = \sqrt{y}.$$

2. *L'application*

$$\begin{aligned} f : \mathbb{R} &\longrightarrow]0, +\infty[\\ x &\longmapsto f(x) = e^x \end{aligned}$$

est bijective. Sa réciproque est l'application

$$\begin{aligned} f^{-1} :]0, +\infty[&\longrightarrow \mathbb{R} \\ x &\longmapsto f^{-1}(x) = \ln(x). \end{aligned}$$

car

$$\forall (x, y) \in \mathbb{R} \times]0, +\infty[: y = e^x \iff x = \ln(y).$$

3. *L'application*

$$\begin{aligned} f : \left[-\frac{\pi}{2}, \frac{\pi}{2}\right] &\longrightarrow [-1, 1] \\ x &\longmapsto f(x) = \sin x, \end{aligned}$$

est bijective, et son application réciproque est notée $f^{-1} = \text{Arcsin}$.

4. *Soit E un ensemble. Alors l'application*

$$\begin{aligned} f : \mathcal{P}(E) &\longrightarrow \mathcal{P}(E) \\ X &\longmapsto f(X) = C_E^X \end{aligned}$$

est bijective et $f^{-1} = f$.

5. *L'application*

$$\begin{aligned} f : \mathbb{R} &\longrightarrow]-1, 1[\\ x &\longmapsto f(x) = \frac{x}{1+|x|}, \end{aligned}$$

est une bijection de $\mathbb{R}$ sur $]-1, 1[$ et sa réciproque est l'application

$$\begin{aligned} f^{-1} :]-1, 1[&\longrightarrow \mathbb{R} \\ x &\longmapsto f^{-1}(x) = \frac{x}{1-|x|}. \end{aligned}$$

6. *Soit f l'application de $\{1, 2, 3\}$ dans $\{1, 5, 7\}$ telle que*

$$f(1) = 5, f(2) = 1, f(3) = 7,$$

elle est bijective. Sa réciproque f^{-1} est l'application de $\{1, 5, 7\}$ dans $\{1, 2, 3\}$ donnée

$$f^{-1}(1) = 2, f^{-1}(5) = 1, f^{-1}(7) = 3.$$

Théorème 3.2.2 (Caractérisation d'une bijection et de sa réciproque)

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

1. *f est bijective si et seulement s'il existe une application g de F dans E telle que*

$$g \circ f = \text{Id}_E \quad \text{et} \quad f \circ g = \text{Id}_F.$$

2. *Si f est bijective, alors l'application g est unique et elle est aussi bijective. De plus*

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

Autrement dit $g = f^{-1}$.

Preuve.

1. (i) Supposons que f est bijective. Alors pour tout élément y de F , il existe un unique élément x de E tel que $y = f(x)$. Posons $g(y) = x$, g est donc une application de F dans E . On a donc pour tout $y \in F$

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

Autrement dit

$$f \circ g = Id_F.$$

D'autre part, en composant à droite avec f , il vient

$$f \circ g \circ f = Id_F \circ f,$$

et par suite pour tout $x \in E$, on a

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

Comme f est injective, alors

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

Ainsi

$$g \circ f = Id_E.$$

En résumé, il existe $g : F \longrightarrow E$ telle que $g \circ f = Id_E$ et $f \circ g = Id_F$.

(ii) Réciproquement, supposons qu'il existe $g : F \longrightarrow E$ telle que $g \circ f = Id_E$ et $f \circ g = Id_F$, et montrons que f est bijective.

• **Injectivité :** Soient $x_1, x_2 \in E$ tels que $f(x_1) = f(x_2)$. En appliquant g il vient alors

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

Cela donne

$$Id_E(x_1) = Id_E(x_2),$$

ou encore

$$x_1 = x_2.$$

Ainsi, f est injective.

• **Surjectivité :** Soit $y \in F$, alors

$$f(g(y)) = (f \circ g)(y) = Id_F(y) = y,$$

donc y admet par f un antécédent $x = g(y) \in E$, et donc f est bien surjective.

2. (i) Montrons maintenant que g est unique. Soit $h : F \longrightarrow E$ une autre application telle que $h \circ f = Id_E$ et $f \circ h = Id_F$. On a donc

$$f \circ h = Id_F = f \circ g.$$

Autrement dit pour tout $y \in F$, on a

$$f(h(y)) = f(g(y)),$$

et comme f est injective alors

$$h(y) = g(y).$$

Par suite $h = g$.

(ii) Montrons la bijectivité de g . On a f est bijective, alors

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

En appliquant la réciproque de (1), on en déduit que g est bijective (il suffit de permuter les lettres g et f). Ainsi $g^{-1} = f$. □

Exemple 3.2.15 Montrer que l'application

$$\begin{aligned} f : \mathbb{R} - \{-1\} &\longrightarrow \mathbb{R}^* \\ x &\longmapsto f(x) = \frac{x}{1+x}, \end{aligned}$$

est une bijection et préciser sa réciproque f^{-1} .

Solution.

Pour tout réel $x \in \mathbb{R} - \{-1\}$, $f(x)$ existe et est un réel non nul. Soit alors

$$\begin{aligned} g : \mathbb{R}^* &\longrightarrow \mathbb{R} - \{-1\} \\ x &\longmapsto g(x) = -1 + \frac{1}{x} \end{aligned}$$

Puisque pour tout réel x non nul, $g(x)$ existe et est différent de -1 , g est bien une application. De plus, pour tout réel x distinct de -1 ,

$$(g \circ f)(x) = g(f(x)) = -1 + \frac{1}{f(x)} = -1 + \frac{1}{\frac{x}{1+x}} = -1 + 1 + x = x,$$

et pour tout réel x distinct de 0 ,

$$(f \circ g)(x) = f(g(x)) = \frac{1}{1 + g(x)} = \frac{1}{1 - 1 + \frac{1}{x}} = x.$$

Ainsi, $g \circ f = Id_{\mathbb{R} - \{-1\}}$ et $f \circ g = Id_{\mathbb{R}^*}$. On a montré que f est bijective et que $f^{-1} = g$.

Remarque 3.2.14 Une seule des deux égalités $g \circ f = Id_E$ et $f \circ g = Id_F$ ne suffit pas à ce que f soit bijective. Considérons par exemple l'application

$$\begin{aligned} f : \mathbb{N} &\longrightarrow \mathbb{N} \\ n &\longmapsto f(n) = n + 1, \end{aligned} \quad \text{et} \quad \begin{aligned} g : \mathbb{N} &\longrightarrow \mathbb{N} \\ n &\longmapsto g(n) = \begin{cases} 0, & \text{si } n = 0 \\ n - 1, & \text{si } n \geq 1. \end{cases} \end{aligned}$$

L'application f n'est pas bijective car l'élément 0 de $\mathbb{N}$ n'a pas d'antécédent par f . Pourtant, pour tout entier naturel n ,

$$(g \circ f)(n) = g(f(n)) = (n + 1) - 1 = n \quad (\text{car } n + 1 \geq 1),$$

et donc $g \circ f = Id_{\mathbb{N}}$. On note que

$$(f \circ g)(0) = f(g(0)) = f(0) = 1 \neq 0,$$

et donc $f \circ g \neq Id_{\mathbb{N}}$.

Corollaire 3.2.1 Soient E et F deux ensembles et $f : E \longrightarrow F$ une bijection. Alors l'application réciproque f^{-1} de f vérifie

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

Preuve.

- Soit x un élément de E . Soit $y = f(x)$. Alors $x = f^{-1}(y) = f^{-1}(f(x))$. Ainsi pour tout $x \in E$, $f^{-1}(f(x)) = x$ et donc $f^{-1} \circ f = Id_E$.
- Soit y un élément de F . Soit $x = f^{-1}(y)$. Alors $y = f(x) = f(f^{-1}(y))$. Ainsi pour tout $y \in F$, $f(f^{-1}(y)) = y$ et donc $f \circ f^{-1} = Id_F$. $\square$

Exemple 3.2.16 On sait que ,

$$\forall x \in \mathbb{R} : \ln(e^x) = x \text{ et } \forall x > 0 : e^{\ln x} = x.$$

Proposition 3.2.13 Soient E, F, G trois ensembles non vides, f une application de E vers F et g une application de F vers G . Si f et g sont bijectives. Alors la réciproque de la composée $g \circ f$ de deux bijections est donnée par la formule

$$(g \circ f)^{-1} = f^{-1} \circ g^{-1}. \quad (3.2)$$

Preuve. L'application $f^{-1} \circ g^{-1}$ est bien définie et c'est une application de G dans E qui vérifie, par associativité

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

et

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

On en déduit d'après le théorème 3.2.2 que la réciproque de $g \circ f$ est $f^{-1} \circ g^{-1}$. □

3.2.8 Injection, surjection, bijection sur des ensembles finis

Lorsqu'on considère des applications entre des ensembles finis, les concepts d'injection, de surjection et de bijection prennent des caractéristiques spécifiques. Il existe un cas particulier dans lequel, pour établir la bijectivité de $f : E \longrightarrow F$, il suffit de démontrer soit la surjectivité, soit l'injectivité. Ce cas se produit lorsque E et F sont des ensembles finis de même cardinalité. Le théorème suivant est alors admis.

Théorème 3.2.3 Soient E et F des ensembles finis de même cardinal ($Card(E) = Card(F)$). Si f une application de E dans F , alors les propriétés suivantes sont équivalentes :

1. f est injective.
2. f est surjective.
3. f est bijective.

Ce théorème signifie que, pour des ensembles finis de même cardinal, il suffit de vérifier l'une des deux propriétés, injectivité ou surjectivité, pour conclure que l'application est bijective. En général, ce résultat n'est pas vrai pour des ensembles infinis. Pour les ensembles infinis, une application peut être injective sans être surjective (ou vice versa), et donc la bijectivité ne peut pas être déduite uniquement à partir d'une seule de ces propriétés. Par exemple,

- L'application

$$\begin{aligned} f_1 : \mathbb{N} &\longrightarrow \mathbb{N} \\ n &\longmapsto 2n, \end{aligned}$$

est injective et non surjective ($y = 1$).

- L'application

$$\begin{aligned} f_2 : \mathbb{N} &\longrightarrow \mathbb{N} \\ n &\longmapsto \begin{cases} \frac{n}{2}, & \text{si } n \text{ est pair} \\ 0, & \text{si } n \text{ est impair,} \end{cases} \end{aligned}$$

est surjective et non injective ($n_1 = 0, n_2 = 2$).

Remarque 3.2.15

1. Pour des ensembles finis, les concepts d'injection, de surjection et de bijection sont liés à la taille des ensembles, on a

- Une injection implique que l'ensemble de départ est de taille inférieure ou égale à celui d'arrivée ($\text{Card}(E) \leq \text{Card}(F)$).
- Une surjection implique que l'ensemble de départ est de taille supérieure ou égale à celui d'arrivée ($\text{Card}(E) \geq \text{Card}(F)$).
- Une bijection implique que les deux ensembles ont la même taille ($\text{Card}(E) = \text{Card}(F)$).

3.2.9 Ordre et applications

L'étude des relations d'ordre dans le contexte des applications (ou fonctions) est un sujet fondamental en mathématiques, particulièrement en théorie des ensembles, en algèbre et en analyse. Voici une présentation des concepts liés aux ordres et aux applications.

Définition 3.2.20 Soient $(E, \leq)$, $(F, \leq)$ deux ensembles ordonnés et $f : E \longrightarrow F$ une application, alors on dit que

1. f est croissante si et seulement si

$$\forall x, y \in E : x \leq y \implies f(x) \leq f(y).$$

2. f est décroissante si et seulement si

$$\forall x, y \in E : x \leq y \implies f(y) \leq f(x).$$

3. f est monotone si et seulement si f est croissante ou décroissante.

4. f est strictement croissante si et seulement si

$$\forall x, y \in E : \begin{cases} x \leq y \\ \text{et} \\ x \neq y \end{cases} \implies \begin{cases} f(x) \leq f(y) \\ \text{et} \\ f(x) \neq f(y). \end{cases}$$

5. f est strictement décroissante si et seulement si

$$\forall x, y \in E : \begin{cases} x \leq y \\ \text{et} \\ x \neq y \end{cases} \implies \begin{cases} f(y) \leq f(x) \\ \text{et} \\ f(x) \neq f(y). \end{cases}$$

6. f est strictement monotone si et seulement si f est strictement croissante ou strictement décroissante.

Exemple 3.2.17

1. L'application

$$\begin{aligned} f : (\mathbb{N}^*, |) &\longrightarrow (\mathbb{N}^*, |) \\ x &\longmapsto f(x) = x^2 \end{aligned}$$

est strictement croissante pour la relation de la divisibilité $|$. En effet, f est strictement croissante pour la relation de la divisibilité si et si

$$\forall x, y \in \mathbb{N}^* : \begin{cases} x | y \\ \text{et} \\ x \neq y \end{cases} \implies \begin{cases} f(x) | f(y) \\ \text{et} \\ f(x) \neq f(y). \end{cases}$$

Soient $x, y \in \mathbb{N}^*$, alors

$$\begin{cases} x \mid y \\ \text{et} \\ x \nmid y \end{cases} \implies \begin{cases} y = m.x, m \in \mathbb{N}^* \\ \text{et} \\ m \nmid 1. \end{cases}$$

Par ailleurs, on a

$$\begin{aligned} f(y) = y^2 = m^2 x^2 = k f(x), k = m^2 \in \mathbb{N}^*. \\ \implies f(x) \mid f(y). \end{aligned}$$

Or $m \nmid 1$, alors $k \nmid 1$. Par suite, $f(x) \nmid f(y)$. Par conséquent, f est strictement croissante pour la relation de la divisibilité.

2. Pour tout ensemble E , l'application

$$\begin{array}{ccc} g : (\mathcal{P}(E), \subset) & \longrightarrow & (\mathcal{P}(E), \subset) \\ X & \longmapsto & g(X) = C_E^X \end{array}$$

est strictement décroissante pour la relation de l'inclusion $\subset$.

Chapitre 4

Structures algébriques

Les structures algébriques constituent des concepts fondamentaux en mathématiques, en particulier en algèbre abstraite. Elles offrent un cadre pour l'étude d'ensembles dotés d'opérations et des propriétés qui en résultent. Présentes dans de nombreux domaines, telles que la physique, l'informatique et l'économie, ces structures ont des applications variées. Parmi les exemples essentiels de structures algébriques, on trouve les groupes, les anneaux et les corps.

Ce chapitre examine en détail les différentes structures algébriques, notamment les groupes, les anneaux et les corps. Il aborde les propriétés des lois de composition interne, les morphismes et les opérations associées à ces structures. De plus, des concepts tels que les sous-groupes, les sous-anneaux et les idéaux sont également introduits.

4.1 Lois de composition interne

Les lois de composition interne sont des concepts fondamentaux en algèbre. Elles interviennent dans la définition de structures algébriques telles que les groupes, les anneaux et les corps. Une loi de composition interne est une opération mathématique qui associe à deux éléments d'un ensemble un troisième élément appartenant également à cet ensemble. Cette opération joue un rôle central non seulement en algèbre, mais aussi en théorie des groupes.

Définition 4.1.1

1. Soit E un ensemble non vide. On appelle loi de composition interne (en abrégé : **l.c.i**) ou opération interne sur E toute application de $E \times E$ dans E .
2. Si $*$ est le symbole désignant cette **l.c.i**, alors $*$ désigne l'application

$$\begin{aligned} * : E \times E &\longrightarrow E \\ (x, y) &\longmapsto x * y. \end{aligned}$$

3. L'élément $x * y$ est appelé composé de x par y via $*$.
4. Les lois de composition interne sont généralement notées $*, +, \cdot, \times, \Delta, \circ, \dots$.
5. Traditionnellement, on utilise la notation $x * y$ pour désigner l'image d'un couple $(x, y) \in E^2$ par une loi $*$ plutôt que la notation fonctionnelle $*(x, y)$.

Exemple 4.1.1

1. L'addition $+$ et la multiplication usuelles $\cdot$ sont des lois de composition interne sur $\mathbb{N}, \mathbb{Z}, \mathbb{Q}, \mathbb{R}$ et $\mathbb{C}$.

2. Si E est un ensemble, la composition des applications $\circ$ est une loi interne sur l'ensemble $\mathcal{F}(E, E)$ des applications de E dans E .
3. Si E est un ensemble, l'intersection, la réunion, la différence et la différence symétrique sont des lois internes sur l'ensemble $\mathcal{P}(E)$ des parties de E .
4. La division (des nombres) peut être ou non une loi selon l'ensemble de nombres considéré : Elle n'est pas interne dans $\mathbb{Z}^*$, mais elle l'est dans $\mathbb{R}^*$.
5. La soustraction n'est pas une loi interne dans $\mathbb{N}$, mais elle l'est dans $\mathbb{Z}$.

Définition 4.1.2 On appelle magma tout ensemble E non vide muni d'une loi de composition interne $*$. On le note $(E, *)$.

Un magma donc est l'une des structures algébriques les plus simples définies à partir d'une loi de composition interne. Un magma se compose d'un ensemble et d'une opération binaire qui associe à chaque paire d'éléments de l'ensemble un élément du même ensemble.

Exemple 4.1.2 $(\mathbb{C}, +)$, $(\mathbb{C}, \cdot)$, $(\mathcal{F}(E, E), \circ)$, $(\mathcal{P}(E), \cap)$, $(\mathcal{P}(E), \cup)$ sont des magmas usuels.

4.1.1 Partie stable pour une loi interne

Une partie stable (ou sous-ensemble stable) pour une loi de composition interne est un sous-ensemble qui est fermé sous cette loi. Cela signifie que l'application de la loi de composition interne à des éléments de ce sous-ensemble produit toujours des éléments qui appartiennent encore à ce sous-ensemble.

Définition 4.1.3 Soit $*$ une loi de composition interne sur un ensemble non vide E .

1. Une partie A de E est dite stable pour $*$ si et seulement si

$$\forall x, y \in A : x * y \in A.$$

2. Si A est une partie stable de E pour $*$, alors la loi de composition interne dans A définie par

$$\begin{array}{ccc} A \times A & \longrightarrow & A \\ (x, y) & \longmapsto & x * y \end{array}$$

est appelée loi de composition interne induite sur A par $*$ de E , et encore notée $*$ (Autrement dit, la restriction de la loi $*$ à A est une loi de composition interne dans A).

Exemple 4.1.3

1. $\mathbb{N}, \mathbb{Z}, \mathbb{Q}$ et $\mathbb{R}$ sont des parties stables de $(\mathbb{C}, +)$ et $(\mathbb{C}, \cdot)$.
2. $\mathbb{R}^-$ et $\mathbb{R}^+$ sont deux parties stables de $\mathbb{R}$, pour la loi $+$ usuelle.
3. Pour la loi $\cdot$, $\mathbb{R}^+$ est encore une partie stable, mais ce n'est pas le cas de $\mathbb{R}^-$.
4. $\mathbb{N}$ n'est pas stable pour la soustraction, mais $\mathbb{Z}$ l'est.
5. La partie $\mathbb{P}$ des nombres pairs de $\mathbb{N}$ est stable pour l'addition, puisque la somme de deux nombres pairs est paire. Elle est également stable pour la multiplication.
6. De même, la partie $\mathbb{I}$ des nombres impairs de $\mathbb{N}$ est stable pour la multiplication.

4.1.2 Propriétés d'une loi de composition interne

4.1.2.1 Commutativité

La commutativité est une propriété importante dans la théorie des structures algébriques, car elle simplifie les calculs et les propriétés des ensembles munis d'une loi de composition. En mathématiques et en informatique, la commutativité est souvent utilisée pour simplifier les algorithmes et les preuves théoriques.

Définition 4.1.4 Soit $*$ une loi interne sur un ensemble non vide E .

1. **Eléments qui commutent pour une loi.** On dit que deux éléments $x, y \in E$ commutent (ou : sont permutables) si et seulement si

$$x * y = y * x.$$

2. **Commutativité d'une loi interne.** On dit que la loi $*$ est commutative si et seulement si

$$\forall x, y \in E : x * y = y * x.$$

Autrement dit tous les éléments de E commutent deux à deux pour cette loi. Le magma $(E, *)$ est alors dit commutatif.

Exemple 4.1.4

1. L'addition $+$ et la multiplication usuelles $\cdot$ sont des lois commutatives sur $\mathbb{N}, \mathbb{Z}, \mathbb{Q}, \mathbb{R}$ et $\mathbb{C}$.
2. Soit E un ensemble, les opérations $\cup$ (union), $\cap$ (intersection) et Δ (différence symétrique) sur $\mathcal{P}(E)$ sont toutes trois commutatives.
3. La loi $\circ$ (composition des applications) n'est pas commutative sur $\mathcal{F}(E, E)$.
4. La soustraction n'est pas commutative dans $\mathbb{R}$.

4.1.2.2 Associativité

L'associativité est une propriété clé en algèbre. Elle permet de simplifier les calculs en garantissant que l'ordre des parenthèses n'affecte pas le résultat d'une opération interne. Cette propriété est essentielle dans de nombreuses structures algébriques, telles que les groupes, les anneaux et les corps.

Définition 4.1.5 (Associativité d'une loi interne)

Soit $*$ une loi interne sur un ensemble non vide E . On dit que la loi $*$ est associative si et seulement si

$$\forall x, y, z \in E : (x * y) * z = x * (y * z).$$

On dit aussi que le magma $(E, *)$, ou simplement E , est associatif.

Exemple 4.1.5

1. L'addition et la multiplication des nombres sont associatives, ce qui permet d'écrire par exemple (deux parenthésages différents)

$$1 + 9 + 3 + 7 = (1 + 9) + (3 + 7) = 10 + 10 = 20,$$

et

$$3 + 6 + 14 + 7 = 3 + ((6 + 14) + 7) = (20 + 7) + 3 = 20 + (7 + 3) = 20 + 10 = 30.$$

2. Dans l'ensemble des parties d'un ensemble, la réunion, l'intersection et la différence symétrique sont associatives.
3. La loi $\circ$ (composition des applications) est associative sur $\mathcal{F}(E, E)$.
4. La soustraction (sur $\mathbb{Z}, \mathbb{Q}, \mathbb{R}, \mathbb{C}$) n'est pas associative.

Remarque 4.1.1

1. Si la loi $*$ est associative, les parenthèses n'étant plus nécessaires à la compréhension de l'opération, on note $x * y * z$, au lieu de $(x * y) * z$ ou $x * (y * z)$.
2. Lorsque E est muni d'une loi associative, on peut effectuer les opérations dans l'ordre que l'on veut, à condition de respecter la position respective des éléments les uns par rapport aux autres. Si la loi est commutative, on peut échanger la position respective des éléments (mais pas nécessairement faire les opérations dans l'ordre qu'on veut si la loi n'est pas associative).

4.1.2.3 Distributivité

La distributivité est une propriété fondamentale des opérations internes qui lie deux opérations différentes, comme l'addition et la multiplication. Cette propriété est essentielle dans de nombreuses structures algébriques, telles que les anneaux, et les corps. Il existe deux types principaux de distributivité que l'on rencontre fréquemment : la distributivité à gauche et à droite.

Cette fois-ci, on va travailler avec deux lois de composition interne définies sur un même ensemble.

Définition 4.1.6 (*Distributivité*).

Soient $*$ et T deux lois internes sur un ensemble non vide E . On dit que :

1. $*$ est distributive à gauche par rapport à la loi T si et seulement si

$$\forall x, y, z \in E : x * (yTz) = (x * y) T(x * z).$$

2. $*$ est distributive à droite par rapport à la loi T si et seulement si

$$\forall x, y, z \in E : (yTz) * x = (y * x) T(z * x).$$

3. $*$ est distributive par rapport à la loi T si et seulement si loi $*$ est distributive à gauche et à droite par rapport à la loi T ,

$$\forall x, y, z \in E : \begin{cases} x * (yTz) = (x * y) T(x * z) \\ \text{et} \\ (yTz) * x = (y * x) T(z * x). \end{cases}$$

Remarque 4.1.2 Si la loi $*$ est commutative, chacune des deux distributivités (à gauche ou à droite) implique l'autre. Cela démontre que, dans le contexte d'une loi interne commutative, il est suffisant de vérifier l'une des formes de distributivité pour garantir l'autre.

Exemple 4.1.6

1. Dans $\mathcal{P}(E)$, les lois $\cap$ et $\cup$ sont distributives l'une par rapport à l'autre.
2. Dans $\mathbb{N}, \mathbb{Z}, \mathbb{Q}, \mathbb{R}$, et $\mathbb{C}$, la multiplication $\cdot$ est distributive par rapport à l'addition $+$.

4.1.3 Eléments particuliers**4.1.3.1 Élément neutre à gauche, neutre à droite, élément neutre**

Un élément neutre est un élément d'un ensemble E qui, lorsqu'il est combiné avec n'importe quel autre élément de l'ensemble à l'aide d'une loi interne, laisse cet autre élément inchangé. Les éléments neutres jouent un rôle central dans de nombreuses structures algébriques telles que les groupes, les anneaux et les corps.

Définition 4.1.7 Soit $*$ une loi interne sur un ensemble non vide E . Soit $e \in E$, alors on dit que e est un élément :

1. Neutre à gauche pour $*$ et on note e_g si et seulement si

$$\forall x \in E : e * x = x.$$

En d'autres termes, lorsque e est utilisé à gauche dans l'opération $*$, le résultat est toujours x , quel que soit x .

2. Neutre à droite pour $*$ et on note e_d si et seulement si

$$\forall x \in E : x * e = x.$$

En d'autres termes, lorsque e est utilisé à droite dans l'opération $*$, le résultat est toujours x , quel que soit x .

3. Neutre si et seulement si e est neutre à gauche et à droite,

$$\forall x \in E : e * x = x * e = x.$$

Remarque 4.1.3

1. Si la loi $*$ est commutative, les notions d'élément neutre à gauche, neutre à droite, et élément neutre coïncident. Cela montre que les notions d'élément neutre à gauche et d'élément neutre à droite sont équivalentes dans ce contexte. C'est-à-dire $e \in E$ est un élément neutre de E si et seulement si $x * e = x$ pour tout $x \in E$.
2. Si la loi $*$ n'est pas commutative, il doit être vérifié les deux égalités. Autrement dit, si $*$ n'est pas commutative, un élément peut être un élément neutre à gauche sans être un élément neutre à droite et vice versa.

Notation 4.1.1

1. **Notation additive** : Si la loi $*$ est notée additivement ($*$ = +), alors l'élément neutre s'il existe est noté 0_E .
2. **Notation multiplicative** : Si la loi $*$ est notée multiplicativement ($*$ = $\cdot$), alors l'élément neutre s'il existe est noté 1_E .

Exemple 4.1.7

1. Dans $\mathbb{R}$, l'élément neutre de l'addition est 0 et celui de la multiplication est 1.
2. Dans $\mathcal{P}(E)$, l'ensemble vide ϕ est élément neutre pour la réunion $\cup$ et la différence symétrique Δ , et E est élément neutre pour l'intersection $\cap$.
3. Quand E est un ensemble, l'application Id_E est élément neutre pour la composition $\circ$ dans $\mathcal{F}(E, E)$.
4. 0 est l'élément neutre à droite mais pas à gauche de $(\mathbb{Z}, -)$, car pour tout $x \in \mathbb{Z}$, on a

$$x - 0 = x \text{ mais } 0 - x = -x.$$

Proposition 4.1.1 (Unicité de l'élément neutre)

L'élément neutre de E pour la loi $*$, s'il existe, est unique.

Preuve. supposons que la loi $*$ possède deux éléments neutres e_1 et e_2 dans E , alors par définition on a

$$\forall x \in E : \begin{cases} e_1 * x = x * e_1 = x, \\ e_2 * x = x * e_2 = x. \end{cases}$$

En particulier

$$e_1 * e_2 = e_2 \text{ et } e_1 * e_2 = e_1,$$

car e_1 (resp. e_2) est le neutre pour $*$. On en déduit que $e_1 = e_2$. □

4.1.3.2 Symétrie à gauche, symétrie à droite, symétrie

Un élément symétrique ou élément inverse dans certains contextes est un élément d'un ensemble qui, lorsqu'il est combiné avec un autre élément à l'aide d'une loi interne, produit l'élément neutre de cette loi. Ce concept est fondamental en algèbre, notamment dans les groupes et les autres structures algébriques

Définition 4.1.8 Soit E un ensemble non vide muni d'une loi interne $*$. On suppose que $(E, *)$ possède un élément neutre $e \in E$. Soit x un élément de E .

1. On dit qu'un élément $x' \in E$ est un symétrique à gauche de x , et on note x'_g (ou $\text{sym}_g(x)$) si et seulement si

$$x' * x = e.$$

En d'autres termes, un élément symétrique à gauche de x est un élément qui, lorsqu'il est combiné avec x à gauche selon la loi interne, produit l'élément neutre.

2. On dit qu'un élément $x' \in E$ est un symétrique à droite de x , et on note x'_d (ou $\text{sym}_d(x)$) si et seulement si

$$x * x' = e.$$

En d'autres termes, un élément symétrique à droite de x est un élément qui, lorsqu'il est combiné avec x à droite selon la loi interne, produit l'élément neutre.

3. On dit qu'un élément $x' \in E$ est un symétrique de x (on le note x' ou $\text{sym}(x)$), si et seulement si x' est un symétrique à gauche et à droite de x ,

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

4. On dit que $x \in E$ est symétrisable (resp. symétrisable à gauche, resp. symétrisable à droite) si x admet au moins un symétrique (resp. un symétrique à gauche, resp. un symétrique à droite).

Remarque 4.1.4

1. Si la loi $*$ est commutative, les notions de symétrique à droite, symétrique à gauche et symétrique coïncident. C'est-à-dire $x' \in E$ est un symétrique de l'élément x si et seulement si $x' * e = e$.

2. Si la loi $*$ est commutative, il doit être vérifié les deux égalités. C'est-à-dire un élément peut avoir un symétrique à gauche sans avoir un symétrique à droite, et vice versa. Les symétriques à gauche et à droite peuvent donc exister indépendamment l'un de l'autre dans des contextes non commutatifs.

Notation 4.1.2

1. Notation additive : Lorsque la loi est notée additivement $+$, le symétrique s'il existe, est appelé « opposé » et est noté $x' = -x$. On dit alors que x est opposable.

2. Notation multiplicative : Lorsque la loi est notée multiplicativement $\cdot$, le symétrique, s'il existe, est appelé « inverse » et est noté $x' = x^{-1}$. On dit alors que x est inversible.

Définition 4.1.9

1. Dans $(\mathbb{N}, +)$ seul 0 est symétrisable. Mais tous les éléments de $(\mathbb{Z}, +), (\mathbb{Q}, +), (\mathbb{R}, +), (\mathbb{C}, +)$ sont symétrisables (opposables).

2. Les éléments inversibles de $(\mathbb{Q}, \mathbb{R}, \mathbb{C}, \cdot)$ sont les éléments non nuls.

3. Le seul élément inversible de $(\mathbb{N}, \cdot)$ est 1. Ceux de $(\mathbb{Z}, \cdot)$ sont -1 et 1 .

4. Dans $(\mathcal{F}(E, E), \circ)$, une application est inversible si et seulement si elle est bijective.

5. Dans $(\mathcal{P}(E), \cup)$ (resp. $(\mathcal{P}(E), \cap)$) seul $\emptyset$ (resp. E), est symétrisable.

Remarque 4.1.5 *L'élément neutre e de $(E, *)$ s'il existe est symétrisable pour $*$ et il est son propre symétrique,*

$$e' = e$$

Proposition 4.1.2 (Unicité du symétrique)

Soit E un ensemble non vide muni d'une loi interne $$ associative et possédant un élément neutre $e \in E$. Si $x \in E$ est symétrisable pour $*$, alors x admet un et un seul symétrique x' pour $*$.*

Preuve. Si $x \in E$ est symétrisable pour $*$ qui admet deux éléments symétriques x_1 et x_2 , alors

$$\begin{cases} x * x_1 = x_1 * x = e, \\ x * x_2 = x_2 * x = e. \end{cases}$$

Ce qui implique que

$$\begin{cases} x_2 * x * x_1 = x_2 * e, \\ x_2 * x = e. \end{cases}$$

Par conséquent

$$e * x_1 = x_2 * e.$$

D'où

$$x_1 = x_2.$$

On en déduit l'unicité du symétrique. □

4.2 Demi-groupe et Monoïde

Un monoïde est une structure algébrique essentielle en mathématiques et en informatique. Cette structure algébrique composée d'un ensemble muni d'une loi interne qui est associative et possède un élément neutre. Un demi-groupe est une structure algébrique plus simple qu'un monoïde. C'est un ensemble muni d'une loi interne qui est associative, mais qui n'a pas nécessairement d'élément neutre. Il est plus général qu'un monoïde puisqu'il n'a pas besoin d'un élément neutre.

Définition 4.2.1 (Demi-groupe et Monoïde)

1. *Un demi-groupe est un ensemble non vide $(E, *)$ muni d'une loi interne associative.*
2. *Un monoïde est un ensemble non vide $(E, *)$ muni d'une loi interne associative et possédant un élément neutre. Autrement dit, un monoïde est un demi-groupe possédant un élément neutre.*
3. *Si de plus la loi $*$ est commutative, le demi-groupe (resp. le monoïde) est dit commutatif.*

Exemple 4.2.1

1. $(\mathbb{N}, +), (\mathbb{Z}, +), (\mathbb{Q}, +), (\mathbb{R}, +), (\mathbb{C}, +)$ sont des monoïdes commutatifs d'élément neutre 0.
2. $(\mathbb{N}, \cdot), (\mathbb{Z}, \cdot), (\mathbb{Q}, \cdot), (\mathbb{R}, \cdot), (\mathbb{C}, \cdot)$ sont des monoïdes commutatifs d'élément neutre 1.
3. $(\mathcal{P}(E), \cup)$ est un monoïde commutatif d'élément neutre $\emptyset$.
4. $(\mathcal{P}(E), \cap)$ est un monoïde commutatif d'élément neutre E .
5. $(\mathcal{F}(E, E), \circ)$ est un monoïde d'élément neutre Id_E .

Proposition 4.2.1 *Soient $(E, *)$ un monoïde de neutre e , x et y deux éléments de E .*

1. *Si x et y sont symétrisables pour $*$, alors x' et $(x * y)$ sont symétrisables pour $*$ et*

$$\begin{cases} (x')' = x \\ (x * y)' = y' * x' \end{cases} \iff \begin{cases} \text{sym}(\text{sym}(x)) = x \\ \text{sym}(x * y) = \text{sym}(y) * \text{sym}(x). \end{cases}$$

Preuve.

1. Si x est symétrisable, alors il existe x' tel que

$$x' * x = x * x' = e,$$

donc x' est symétrisable pour $*$ et $(x')' = x$.

2. Soient x et y deux éléments de E symétrisables pour $*$. Posons

$$z = y' * x'.$$

Alors, par associativité de la loi $*$, on a

$$(x * y) * z = x * (y * z) = x * (y * (y' * x')) = x * ((y * y') * x') = x * (e * x') = x * x' = e,$$

et de même

$$z * (x * y) = e.$$

On a ainsi montré que $(x * y)$ est symétrisable pour $*$ et $(x * y)' = y' * x'$. □

Remarque 4.2.1

1. Traduisons pour une loi multiplicative $\cdot$: si x et y sont inversibles, d'inverses x^{-1} et y^{-1} , alors $x \cdot y$ aussi, et

$$(x \cdot y)^{-1} = y^{-1} \cdot x^{-1}.$$

Dans le cas d'une loi additive $+$, on obtient

$$-(x + y) = (-y) + (-x).$$

2. $(x * y)' = y' * x'$, attention à l'ordre des facteurs si la loi $*$ n'est pas commutative.

Exemple 4.2.2 Si f et g sont deux bijections d'un ensemble E dans lui-même, alors $g \circ f$ est aussi bijective et sa réciproque $(g \circ f)^{-1}$ est $f^{-1} \circ g^{-1}$.

4.3 Groupes

Les groupes constituent des structures algébriques essentielles et largement étudiées en mathématiques, car ils émergent naturellement dans de nombreux domaines, allant de la théorie des nombres à la physique et à l'informatique. Un groupe est défini par un ensemble muni d'une loi de composition interne qui respecte des propriétés spécifiques.

4.3.1 Généralités sur les groupes

Définition 4.3.1

1. On dit qu'un ensemble non vide G muni d'une loi interne $*$ est un groupe si et seulement si

$$\left\{ \begin{array}{ll} (j) & \text{La loi } * \text{ est associative. } \forall x, y, z \in G : (x * y) * z = x * (y * z). \\ (jj) & \text{La loi } * \text{ possède un élément neutre. } \exists e \in G, \forall x \in G : e * x = x * e = x. \\ (jjj) & \text{Tout élément possède un symétrique pour } *. \forall x \in G, \exists x' \in G : x' * x = x * x' = e. \end{array} \right.$$

2. Si, de plus, la loi est $*$ commutative, on dit que $(G, *)$ est un groupe commutatif, ou un groupe abélien (du nom de Niels Henrik ABEL).

Remarque 4.3.1

1. Un groupe est un monoïde dans lequel tout élément admet un symétrique.
2. Un groupe n'est jamais vide, il contient e .

Exemple 4.3.1

1. $(\mathbb{Z}, \mathbb{Q}, \mathbb{R}, \mathbb{C}, +)$ sont des groupes commutatifs.
2. $(\mathbb{Q}^*, \mathbb{R}^*, \mathbb{C}^*, \cdot)$ sont des groupes commutatifs.
3. $(\mathbb{R}, \cdot)$ n'est pas un groupe car, par exemple, 0 n'est pas inversible.

4.3.2 Puissances entières d'un élément dans un groupe

Dans un groupe $(G, *)$, les puissances entières d'un élément $x \in G$ sont définies en utilisant l'opération $*$ de groupe. Ces puissances entières permettent d'étudier la structure et les propriétés du groupe.

Définition 4.3.2 (Puissances entières d'un élément dans un groupe)

Soit $(G, *)$ un groupe, on définit les puissances (ou les itérés n – ièmes) entières d'un élément $x \in G$ de la façon suivante :

1. Les itérés d'ordre n de l'élément x sont donnés par les relations

$$\left\{ \begin{array}{l} x^n = \underbrace{x * x * \dots * x}_{n \text{ fois}}, n \in \mathbb{N}^* \\ x^0 = e_G \\ x^{-n} = (x^{-1})^n = \underbrace{x^{-1} * x^{-1} * \dots * x^{-1}}_{n \text{ fois}}, n \in \mathbb{N}^*. \end{array} \right.$$

En particulier

$$\forall n \in \mathbb{N} : x^{n+1} = x * (x^n).$$

2. Les puissances entières de $x \in G$ satisfont les propriétés suivantes pour tous $n, m \in \mathbb{Z}$, on a

$$\left\{ \begin{array}{l} x^n * x^m = x^{n+m} = x^{m+n} = x^m * x^n. \\ (x^n)^m = x^{nm} = x^{mn} = (x^m)^n. \end{array} \right.$$

Notation 4.3.1

1. Lorsque la loi est notée multiplicativement $*$ = $\cdot$, alors
 - a. Les itérés multiplicatifs d'un élément x sont donnés par les relations

$$\left\{ \begin{array}{l} x^n = \underbrace{x \cdot x \cdot \dots \cdot x}_{n \text{ fois}}, n \in \mathbb{N}^*. \\ x^0 = 1_G. \\ x^{-n} = (x^{-1})^n = \underbrace{x^{-1} \cdot x^{-1} \cdot \dots \cdot x^{-1}}_{n \text{ fois}}, n \in \mathbb{N}^*. \end{array} \right.$$

- b. Pour tous $n, m \in \mathbb{Z}$, on a

$$\left\{ \begin{array}{l} x^n \cdot x^m = x^{n+m} = x^{m+n} = x^m \cdot x^n \\ (x^n)^m = x^{nm} = x^{mn} = (x^m)^n. \end{array} \right.$$

2. Lorsque la loi est notée additivement $*$ = +, alors

a. Les itérés additifs d'un élément x sont donnés par les relations

$$\begin{cases} n.x = \underbrace{x + x + \dots + x}_{n \text{ fois}}, n \in \mathbb{N}^*. \\ 0.x = 0_G. \\ (-n).x = n.(-x) = \underbrace{(-x) + \dots + (-x)}_{n \text{ fois}}, n \in \mathbb{N}^*. \end{cases}$$

b. Pour tous $n, m \in \mathbb{Z}$, on a

$$\begin{cases} (n + m)x = nx + mx \\ n(mx) = (nm)x. \end{cases}$$

Remarque 4.3.2

1. En général

$$(x * y)^n \neq x^n * y^n,$$

tels que $x, y \in G$. En effet :

$$(x * y)^n = \underbrace{(x * y) * (x * y) * \dots * (x * y)}_{n \text{ fois}}$$

et

$$x^n * y^n = \left(\underbrace{x * x * \dots * x}_{n \text{ fois}} \right) * \left(\underbrace{y * y * \dots * y}_{n \text{ fois}} \right).$$

2. Si $x, y \in G$ commutent ($x * y = y * x$), alors

$$(x * y)^n = x^n * y^n,$$

pour tout $n \in \mathbb{Z}$.

4.3.3 Ordre d'un Groupe

L'ordre d'un groupe donne une mesure du "taille" du groupe. Pour les groupes finis, c'est simplement le nombre d'éléments. Pour les groupes infinis, l'ordre est infini.

Définition 4.3.3 (Ordre d'un Groupe)

L'ordre d'un groupe $(G, *)$, noté $\text{ord}(G)$ ou $|G|$, est le nombre d'éléments de G . Si G est un groupe fini, son ordre est un entier positif. Si G est infini, on dit simplement que son ordre est infini.

Exemple 4.3.2

1. $\mathbb{Z}, \mathbb{Q}, \mathbb{R}$ sont des groupes d'ordre infini.
2. Groupe des entiers modulo $n \in \mathbb{N}^*$. Pour tout $n \in \mathbb{N}^*$, le groupe $(\mathbb{Z}/n\mathbb{Z}, +)$ est fini d'ordre n . Par exemple, $\mathbb{Z}/5\mathbb{Z} = \{0, 1, 2, 3, 4\}$ a un ordre de 5.
3. Groupe Symétrique $(\mathcal{S}_n, \circ)$ (le groupe des permutations de n éléments) a un ordre de $n!$.

4.3.4 Sous-groupes

Les sous-groupes sont des concepts essentiels en théorie des groupes, car ils permettent d'explorer des parties plus petites d'un groupe tout en préservant la structure algébrique de celui-ci.

Démontrer qu'un ensemble est un groupe à partir de sa définition peut s'avérer long et complexe. Une approche alternative consiste à établir qu'un sous-ensemble d'un groupe est également un groupe, ce qui constitue la notion de sous-groupe.

Définition 4.3.4 Soit $(G, *)$ un groupe, et H un sous-ensemble non vide de G . Un sous-groupe est une partie H d'un groupe $(G, *)$ qui, munie de la loi $*$ de G , est aussi un groupe.

De manière équivalente, cette définition peut s'écrire comme suit.

Définition 4.3.5 Soit $(G, *)$ un groupe, et H un sous-ensemble non vide de G .

1. On dit que H est un sous-groupe de G , si et seulement si les assertions suivantes sont vérifiées :

$$\left\{ \begin{array}{ll} \text{i)} & e \in H \quad (H \text{ non vide}) \\ \text{ii)} & \forall x, y \in H : x * y \in H \quad (H \text{ est stable pour la loi } *) \\ \text{iii)} & \forall x \in H : x' \in H \quad (H \text{ est stable par passage au symétrique}) \end{array} \right.$$

2. On appelle sous-groupe propre tout sous-groupe de G distinct de G et de $\{e\}$.

Remarque 4.3.3 On vérifie facilement que si H est un sous-groupe de $(G, *)$, alors

1. Les groupes $(H, *)$ et $(G, *)$ partagent le même élément neutre.
2. Le symétrique d'un élément x de H est le même, que l'on considère x comme un élément du groupe $(H, *)$ ou un élément du groupe $(G, *)$.

Exemple 4.3.3

1. Soit $(G, *)$ un groupe, alors G et $\{e\}$ sont deux sous-groupes de G (appelés sous-groupes triviaux).
2. Dans $(\mathbb{Z}, +)$, $(\mathbb{Q}, +)$, $(\mathbb{R}, +)$, $(\mathbb{C}, +)$ chacun est un sous-groupe du suivant. Même chose avec $(\mathbb{Q}^*, \cdot)$, $(\mathbb{R}^*, \cdot)$, $(\mathbb{C}^*, \cdot)$.
3. $\mathbb{R}_+^*$ est un sous-groupe de $(\mathbb{R}^*, \cdot)$.
4. Caractérisation des sous-groupes de $(\mathbb{Z}, +)$: Pour tout $n \in \mathbb{N}^*$, les sous-groupes de $(\mathbb{Z}, +)$ sont les sous-ensembles de $\mathbb{Z}$ suivants

$$n\mathbb{Z} = \{nk : k \in \mathbb{Z}\} \quad (\text{l'ensemble des multiples de } n).$$

5. On note U l'ensemble des nombres complexes de module 1, on a donc

$$U = \{z \in \mathbb{C} : |z| = 1\}.$$

$(U, \cdot)$ est un sous-groupe de $(\mathbb{C}^*, \cdot)$.

6. Pour $n \geq 1$, on définit l'ensemble des racines n -ièmes de l'unité

$$R_n = \{z \in \mathbb{C} : z^n = 1\}.$$

R_n forme un sous-groupe du groupe multiplicatif des nombres complexes non nuls $(\mathbb{C}^*, \cdot)$.

7. L'ensemble

$$H = \{x = a + b\sqrt{2} : a, b \in \mathbb{Z}\},$$

est un sous-groupe du groupe additif $(\mathbb{Z}, +)$.

Un exemple très important de sous-groupe est le centre d'un groupe.

Définition 4.3.6 (Centre d'un groupe)

Le centre d'un groupe $(G, *)$, noté $Z(G)$, est l'ensemble des éléments de G qui commutent avec tous les éléments de G . Formellement,

$$Z(G) = \{x \in G, \forall y \in G : x * y = y * x\}.$$

Remarque 4.3.4 (Propriétés du Centre d'un Groupe)

1. Si $Z(G) = \{e\}$, alors G est dit avoir un centre trivial. Ceci se produit pour de nombreux groupes non-abéliens.
2. Si G est un groupe abélien, alors chaque élément de G commutera avec tous les autres éléments de G . Par conséquent, $Z(G) = G$ dans ce cas.

Proposition 4.3.1 *Le centre d'un groupe $(G, *)$, est un sous-groupe de G .*

Preuve. Il est clair que $Z(G) \subset G$.

- Par définition de l'élément neutre, on a

$$\forall y \in G : e * y = y * e.$$

Alors $e \in Z(G)$, donc $Z(G) \neq \emptyset$.

- Soient $x_1, x_2 \in Z(G)$, alors pour tout $y \in G$, on a

$$\begin{aligned}(x_1 * x_2) * y &= x_1 * (x_2 * y) \text{ (car la loi } * \text{ est associative dans } G) \\ &= x_1 * (y * x_2) \text{ (car } x_2 \text{ est un élément de } Z(G)) \\ &= (x_1 * y) * x_2 \text{ (car la loi } * \text{ est associative dans } G) \\ &= (y * x_1) * x_2 \text{ (car } x_1 \text{ est un élément de } Z(G)) \\ &= y * (x_1 * x_2) \text{ (car la loi } * \text{ est associative dans } G),\end{aligned}$$

et donc $x_1 * x_2 \in Z(G)$.

- Soit $x \in Z(G)$, alors pour tout $y \in G$, on a

$$y * x = x * y,$$

cela donne

$$(x' * y) * x = x' * (y * x) = x' * (x * y) = (x' * x) * y = y,$$

et par suite

$$\forall y \in G : x' * y = y * x',$$

ce qui entraîne que $x' \in Z(G)$. □

Les deux résultats suivants donnent des critères pratiques que l'on utilise pour montrer qu'une partie donnée d'un groupe est (ou n'est pas) un sous-groupe. Les conditions d'un sous-groupe ci-dessus peuvent être remplacées par celles données par la proposition suivante.

Proposition 4.3.2 (Caractérisation des sous-groupes)

Soient $(G, *)$ un groupe et H un sous-ensemble de G . Le sous-ensemble H est un sous-groupe de G si et seulement si

$$\begin{cases} i) & e \in H \\ ii) & \forall x, y \in H : x * y \in H. \end{cases}$$

Corollaire 4.3.1

1. Cas de la notation additive.

$$H \text{ est un sous-groupe de } (G, +) \iff \begin{cases} i) & 0_G \in H \\ ii) & \forall x, y \in H : x + y \in H \\ iii) & \forall x \in H : -x \in H. \end{cases} \iff \begin{cases} i) & 0_G \in H \\ ii) & \forall x, y \in H : x - y \in H. \end{cases}$$

2. Cas de la notation multiplicative.

$$H \text{ est un sous-groupe de } (G, \cdot) \iff \begin{cases} i) & 1_G \in H \\ ii) & \forall x, y \in H : x \cdot y \in H \\ iii) & \forall x \in H : x^{-1} \in H. \end{cases} \iff \begin{cases} i) & 1_G \in H \\ ii) & \forall x, y \in H : x \cdot y^{-1} \in H. \end{cases}$$

Proposition 4.3.3 (Intersection de sous-groupes)

Soient $(G, *)$ un groupe, $(H_i)_{i \in I}$ une famille (finie ou infinie) de sous-groupes de G . Alors $\bigcap_{i \in I} H_i$ est un sous-groupe de G (Une intersection quelconque de sous-groupes est encore un sous-groupe).

Preuve.

i) Posons $H = \bigcap_{i \in I} H_i$. Il est clair que

$$\bigcap_{i \in I} H_i \subset G.$$

Puisque H_i est un sous-groupe de G , alors pour tout $i \in I$, $e \in H_i$ et donc $e \in H$.

ii) Soient x et y deux éléments de H . Alors pour tout $i \in I$, on a $x \in H_i$ et $y \in H_i$, donc $x * y \in H_i$ (puisque pour tout i de I , H_i est un sous-groupe de G). Il vient donc que

$$x * y \in H.$$

iii) Enfin, soit x un élément de H . Pour tout i de I , on a $x \in H_i$, donc pour tout i de I , $x' \in H_i$. Ainsi,

$$x' \in H.$$

Cela prouve que $\bigcap_{i \in I} H_i$ est un sous-groupe de G . □

Remarque 4.3.5

1. La réunion de deux sous-groupes d'un groupe G n'est pas, en général, un sous-groupe de G . En guise de contre-exemple, si l'on considère les sous-groupes $2\mathbb{Z}$ et $3\mathbb{Z}$ de $(\mathbb{Z}, +)$, leur réunion $A = 2\mathbb{Z} \cup 3\mathbb{Z}$ n'est pas un sous-groupe de $\mathbb{Z}$. En effet, on a 2 et 3 sont deux éléments de A (car $2 \in 2\mathbb{Z}$ et $3 \in 3\mathbb{Z}$) mais $2 + 3 = 5 \notin A$ (car 5 n'est ni un multiple de 2 ni un multiple de 3).

2. La réunion de deux sous-groupes H et K d'un groupe G n'est un sous-groupe de G que dans un cas spécial : lorsqu'un des sous-groupes est inclus dans l'autre, c'est-à-dire $H \subset K$ ou $K \subset H$. Dans ce cas, la réunion $H \cup K$ est simplement le sous-groupe plus grand. Prenons le groupe $G = \mathbb{Z}$ (les entiers sous l'addition), le sous-groupe $H = 2\mathbb{Z}$ et le sous-groupe $K = 4\mathbb{Z}$. Ici, K est un sous-groupe de H car tout multiple de 4 est aussi un multiple de 2. Autrement dit, $H \subset K$. La réunion $H \cup K$ est

$$H \cup K = 2\mathbb{Z} \cup 4\mathbb{Z} = 2\mathbb{Z}.$$

Puisque $H \subset K$, la réunion des deux sous-groupes est simplement $2\mathbb{Z}$.

4.4 Groupe $\mathbb{Z}/n\mathbb{Z}$

Le groupe $\mathbb{Z}/n\mathbb{Z}$ est constitué des entiers modulo n , avec l'opération d'addition effectuée modulo n . Il s'agit d'un groupe fini d'ordre n , qui occupe une place centrale en théorie des groupes, en algèbre abstraite et en théorie des nombres. Voici un aperçu de ce groupe et de ses propriétés.

Fixons $n \in \mathbb{N}^*$, rappelons que $\mathbb{Z}/n\mathbb{Z}$ est l'ensemble des classes d'équivalence pour la congruence modulo n défini par

$$\mathbb{Z}/n\mathbb{Z} = \left\{ \dot{0}, \dot{1}, 2, \dots, \widehat{n-1} \right\},$$

où $\dot{x}$ désigne la classe d'équivalence de $x \in \mathbb{Z}$ pour cette relation. Autrement dit

$$\dot{x} = \dot{y} \iff x \equiv y \pmod{n},$$

ou encore

$$\dot{x} = \dot{y} \iff \exists k \in \mathbb{Z} : x - y = kn.$$

On définit une opération d'addition notée $\dot{+}$ sur $\mathbb{Z}/n\mathbb{Z}$ par

$$\forall \dot{x}, \dot{y} \in \mathbb{Z}/n\mathbb{Z} : \dot{x} \dot{+} \dot{y} = \widehat{x + y}.$$

Par exemple dans $\mathbb{Z}/5\mathbb{Z}$, on a

$$\dot{2} \dot{+} \dot{4} = \widehat{2 + 4} = \dot{6} = \dot{1}.$$

Alors $\dot{+}$ définit une loi de composition interne sur $\mathbb{Z}/n\mathbb{Z}$.

La construction de $\mathbb{Z}/n\mathbb{Z}$ peut paraître conceptuellement difficile la première fois qu'on la voit, mais en fait, la manipulation de cet ensemble est très simple en pratique : écrire $\dot{x} \dot{+} \dot{y} = \dot{z}$ est rigoureusement équivalent à écrire $x + y \equiv z \pmod{n}$, par exemple. Pour passer d'une écriture à l'autre, on enlève les barres et on remplace l'égalité par une relation de congruence. Mais l'énorme avantage conceptuel de l'utilisation de $\mathbb{Z}/n\mathbb{Z}$ est, dans le cas de $\mathbb{Z}/5\mathbb{Z}$ par exemple, le fait que $\dot{2}$ et $\dot{7}$ sont un seul et même nombre, et non plus simplement congrus. De plus, $\mathbb{Z}/n\mathbb{Z}$ possède une certaine structure algébrique, qui nous permet de réaliser toutes nos opérations en restant à l'intérieur de $\mathbb{Z}/n\mathbb{Z}$, et donc sans avoir à repasser par les entiers.

Proposition 4.4.1 (*Addition sur $\mathbb{Z}/n\mathbb{Z}$*)

Pour tout entier $n \geq 1$, l'opération $\dot{+}$ d'addition modulo n est une loi interne sur $\mathbb{Z}/n\mathbb{Z}$.

Preuve.

• Nous devons montrer que cette addition est bien définie pour cela nous devons vérifier que le résultat de l'addition ne dépend pas des représentants des classes d'équivalence choisies. C'est-à-dire que pour tous $x, y \in \mathbb{Z}$, $\widehat{x + y}$ ne dépend que de $x + y$, et non pas de x et de y . Soient x_1 un élément de $\dot{x}$ et y_1 un élément de $\dot{y}$, alors

$$\left\{ \begin{array}{l} x \mathcal{R} x_1 \\ \text{et} \\ y \mathcal{R} y_1. \end{array} \right.$$

Par suite

$$\left\{ \begin{array}{l} x - x_1 = k_1 n, k_1 \in \mathbb{Z} \\ \text{et} \\ y - y_1 = k_2 n, k_2 \in \mathbb{Z}. \end{array} \right.$$

Ceci entraîne que

$$(x + y) - (x_1 + y_1) = k_3 n, k_3 = k_1 + k_2 \in \mathbb{Z},$$

et par conséquent

$$\widehat{x + y} = \widehat{x_1 + y_1}.$$

Donc on a aussi

$$\dot{x} \dot{+} \dot{y} = \dot{x}_1 \dot{+} \dot{y}_1.$$

Finalement, pour tous $\dot{x}, \dot{y}$ éléments de $\mathbb{Z}/n\mathbb{Z}$, le résultat de l'opération $\dot{x} \dot{+} \dot{y}$ ne dépend pas des représentants choisis dans $\dot{x}$ et $\dot{y}$.

• Maintenant, pour montrer que cette opération est une loi interne sur $\mathbb{Z}/n\mathbb{Z}$, on doit vérifier que pour tous $\dot{x}, \dot{y}$ de $\mathbb{Z}/n\mathbb{Z}$, leur somme $\dot{x} \dot{+} \dot{y}$ est aussi un élément de $\mathbb{Z}/n\mathbb{Z}$.

Alors, soient $\dot{x}, \dot{y}$ deux éléments de $\mathbb{Z}/n\mathbb{Z}$, cela signifie que x et y sont des entiers ($x, y \in \mathbb{Z}$). Considérons la somme $x + y$, cette somme est un entier. En réduisant $x + y$ modulo n , nous obtenons un entier qui est dans l'ensemble $\{0, 1, 2, \dots, n-1\}$. Par définition, $\left(\dot{x} \dot{+} \dot{y} = \widehat{x + y}\right)$ est dans $\mathbb{Z}/n\mathbb{Z}$. $\square$

Proposition 4.4.2 *Pour tout entier $n \geq 1$, $(\mathbb{Z}/n\mathbb{Z}, \dot{+})$ est un groupe commutatif d'ordre n .*

Preuve.

1. On sait que l'ordre d'un groupe est le nombre d'éléments dans ce groupe. Or $\mathbb{Z}/n\mathbb{Z}$ contient exactement n éléments distincts. Par conséquent, l'ordre de $(\mathbb{Z}/n\mathbb{Z}, \dot{+})$ est n .

2. Montrons maintenant que $(\mathbb{Z}/n\mathbb{Z}, \dot{+})$ est un groupe commutatif.

- La loi $\dot{+}$ est évidemment une loi interne sur $\mathbb{Z}/n\mathbb{Z}$.
- La loi $\dot{+}$ est associative. En effet, pour tous $x, y, z \in \mathbb{Z}$, on a

$$(\dot{x} \dot{+} \dot{y}) \dot{+} \dot{z} = \widehat{\widehat{x + y} + z} = \widehat{(x + y) + z} = \widehat{x + (y + z)} = \dot{x} \dot{+} (\dot{y} \dot{+} \dot{z}).$$

- $\mathbb{Z}/n\mathbb{Z}$ a un élément neutre pour $\dot{+}$, à savoir

$$e = \dot{0} = n\mathbb{Z}.$$

En effet, pour tout $x \in \mathbb{Z}$, on a

$$\dot{x} \dot{+} \dot{0} = \widehat{x + 0} = \dot{x} = \widehat{0 + x} = \dot{0} \dot{+} \dot{x}.$$

- Soit $x \in \mathbb{Z}$, posons $y = -x$ (ainsi, $y \in \mathbb{Z}$). On a alors

$$\dot{x} \dot{+} \dot{y} = \widehat{x + y} = \widehat{x + (-x)} = \dot{0} = \widehat{(-x) + x} = \dot{y} \dot{+} \dot{x}.$$

Donc tout élément $\dot{x}$ de $\mathbb{Z}/n\mathbb{Z}$ admet pour la loi un symétrique

$$\dot{x}' = \widehat{-x} = \widehat{n - x}.$$

- Enfin, $\dot{+}$ est commutative : pour tous $x, y \in \mathbb{Z}$, on a

$$\dot{x} \dot{+} \dot{y} = \widehat{x + y} = \widehat{y + x} = \dot{y} \dot{+} \dot{x}.$$

Donc $(\mathbb{Z}/n\mathbb{Z}, \dot{+})$ est un groupe commutatif. $\square$

Exemple 4.4.1 Table de l'addition dans $\mathbb{Z}/5\mathbb{Z}$,

$\dot{+}$	$\dot{0}$	$\dot{1}$	$\dot{2}$	$\dot{3}$	$\dot{4}$
$\dot{0}$	$\dot{0}$	$\dot{1}$	$\dot{2}$	$\dot{3}$	$\dot{4}$
$\dot{1}$	$\dot{1}$	$\dot{2}$	$\dot{3}$	$\dot{4}$	$\dot{0}$
$\dot{2}$	$\dot{2}$	$\dot{3}$	$\dot{4}$	$\dot{0}$	$\dot{1}$
$\dot{3}$	$\dot{3}$	$\dot{4}$	$\dot{0}$	$\dot{1}$	$\dot{2}$
$\dot{4}$	$\dot{4}$	$\dot{0}$	$\dot{1}$	$\dot{2}$	$\dot{3}$

1. Pour chaque cellule de la table, l'élément à la ligne i et la colonne j est le résultat de l'addition modulo 5 de i et j .
2. La table d'addition ci-dessus montre que l'opération d'addition modulo 5 sur $\mathbb{Z}/5\mathbb{Z}$ est bien définie et vérifie les propriétés d'un groupe commutatif.

4.5 Généralités sur le groupe symétrique $\mathcal{S}_n$

Le groupe des permutations d'un ensemble E , noté $S(E)$, est un concept fondamental en mathématiques, particulièrement en théorie des groupes, combinatoire et algèbre abstraite. Deux notions essentielles dans ce domaine sont le groupe des permutations et le groupe symétrique $\mathcal{S}_n$. Ces groupes permettent de comprendre comment les éléments d'un ensemble peuvent être réarrangés tout en préservant certaines structures et propriétés.

Le groupe des permutations $S(E)$ est défini comme l'ensemble de toutes les bijections de l'ensemble E sur lui-même. Une bijection est une fonction qui associe chaque élément de E à un autre élément de E de manière unique, sans répétition ni omission. Le groupe symétrique $\mathcal{S}_n$ est un cas particulier de $S(E)$ lorsque E est un ensemble fini contenant exactement n éléments, représentant ainsi le groupe des permutations de cet ensemble spécifique.

En d'autres termes, pour un ensemble de n éléments, par exemple $\{1, 2, \dots, n\}$, le groupe $\mathcal{S}_n$ regroupe toutes les bijections de cet ensemble sur lui-même. Ainsi, le groupe $S(E)$ des permutations d'un ensemble E constitue une généralisation du groupe symétrique $\mathcal{S}_n$ pour les ensembles qui peuvent être infinis, tandis que $\mathcal{S}_n$ se limite aux permutations d'un ensemble fini de n éléments.

4.5.1 Rappels et notations

Définition 4.5.1 (*Permutation sur un ensemble E*)

Soit E un ensemble non vide.

1. On appelle permutation de E toute bijection notée σ de E dans lui-même,

$$\sigma : E \longrightarrow E.$$

Cela signifie que chaque élément de E est associé de manière unique à un autre élément de E .

2. On note $S(E)$ ou S_E l'ensemble des permutations de E dans lui-même.

Exemple 4.5.1

1. L'application identité Id_E de E est une permutation de E .
2. Dans le vocabulaire courant, permuter des objets signifie modifier l'ordre dans lequel ils sont rangés.

Cette idée correspond au concept de permutation de la définition 4.5.1. En effet, permuter les nombres 1, 2, 3, 4, 5 pour les disposer dans l'ordre 3, 4, 1, 2, 5, c'est appliquer la bijection

$$\sigma : \begin{cases} 1 \mapsto 3 \\ 2 \mapsto 4 \\ 3 \mapsto 1 \\ 4 \mapsto 2 \\ 5 \mapsto 5. \end{cases}$$

Remarque 4.5.1 Dans le cas où E est réduit à un élément, on peut quand même définir $S(E)$ et il est réduit à $\{Id_E\}$.

Proposition 4.5.1 L'ensemble $(\mathcal{S}(E), \circ)$ des permutations de E muni avec la loi de composition des applications $\circ$ est un groupe, appelé groupe des permutations de E .

Preuve.

- Si p_1 et p_2 sont des permutations de E , alors leur composition $p_1 \circ p_2$ est aussi une permutation de E car la composition de deux bijections est une bijection, donc on a une loi interne.
- La composition est clairement associative (par l'associativité de la composition des applications).
- L'identité Id_E (La permutation identité) est l'élément neutre pour la composition.
- Enfin, tout élément de $\mathcal{S}(E)$ est inversible d'inverse son application réciproque car chaque bijection a une bijection inverse telle que la composition des deux donne l'identité.

Ainsi, $(\mathcal{S}(E), \circ)$ est un groupe. $\square$

On considère un ensemble fini E de cardinal $n \in \mathbb{N}^*$, et en particulier $E = \{1, 2, \dots, n\}$, nous intéressons maintenant au groupe symétrique, c'est-à-dire le groupe des permutations d'un ensemble fini $E = \{1, 2, \dots, n\}$.

Définition 4.5.2 (Groupe symétrique).

Pour tout entier naturel n non nul, on appelle groupe symétrique d'indice n l'ensemble des permutations de l'ensemble $\{1, 2, \dots, n\}$. C'est-à-dire de toutes les bijections de $\{1, 2, \dots, n\}$ vers $\{1, 2, \dots, n\}$. On note ce groupe $\mathcal{S}_n$.

Remarque 4.5.2 Une permutation de n éléments est aussi appelée permutation sans répétition de ces éléments. Signalons qu'autrefois une permutation était appelée substitution.

Notation 4.5.1 On représente généralement une permutation $\sigma \in \mathcal{S}_n$ par un tableau (ou une Matrice) dont la première ligne est constituée par les entiers de 1 à n rangés par ordre croissant (représente l'ensemble de départ) et dont la seconde ligne est constituée de leurs images respectives (représente l'ensemble d'arrivée). On a donc

$$\sigma = \begin{pmatrix} 1 & 2 & \dots & n \\ \sigma(1) & \sigma(2) & \dots & \sigma(n) \end{pmatrix}.$$

En particulier l'application identité, neutre du groupe $\mathcal{S}_n$, se note e ou Id_E

$$Id_E = \begin{pmatrix} 1 & 2 & \dots & n \\ 1 & 2 & \dots & n \end{pmatrix}.$$

Par exemple la permutation de $\mathcal{S}_7$ notée

$$\sigma = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 & 6 & 7 \\ 3 & 7 & 5 & 4 & 6 & 1 & 2 \end{pmatrix},$$

est la bijection

$$\sigma : \{1, 2, 3, 4, 5, 6, 7\} \longrightarrow \{1, 2, 3, 4, 5, 6, 7\},$$

définie par

$$\sigma(1) = 3, \sigma(2) = 7, \sigma(3) = 5, \sigma(4) = 4, \sigma(5) = 6, \sigma(6) = 1, \sigma(7) = 2.$$

C'est bien une bijection car chaque nombre de 1 à 7 apparaît une fois et une seule sur la deuxième ligne.

Exemple 4.5.2

1. Si $n = 1$, alors le groupe $\mathcal{S}_1$ se réduit à l'application identité de $E = \{1\}$ dans lui-même,

$$\mathcal{S}_1 = \{Id_E\}.$$

2. Si $n = 2$, alors $\mathcal{S}_2 = \{Id_E, \sigma\}$, où $E = \{1, 2\}$ et

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

3. Si $n = 3$, alors $E = \{1, 2, 3\}$ et $\mathcal{S}_3$ est formé de six éléments, qui sont

$$\begin{aligned} \sigma_0 = Id &= \begin{pmatrix} 1 & 2 & 3 \\ 1 & 2 & 3 \end{pmatrix}, \sigma_1 = \begin{pmatrix} 1 & 2 & 3 \\ 1 & 3 & 2 \end{pmatrix}, \sigma_2 = \begin{pmatrix} 1 & 2 & 3 \\ 3 & 2 & 1 \end{pmatrix}, \sigma_3 = \begin{pmatrix} 1 & 2 & 3 \\ 2 & 1 & 3 \end{pmatrix} \\ \sigma_4 &= \begin{pmatrix} 1 & 2 & 3 \\ 2 & 3 & 1 \end{pmatrix}, \sigma_5 = \begin{pmatrix} 1 & 2 & 3 \\ 3 & 1 & 4 \end{pmatrix}. \end{aligned}$$

Corollaire 4.5.1 Le cardinal de $\mathcal{S}_n$ est $n!$.

Preuve. Pour l'élément 1, son image appartient à $\{1, 2, \dots, n\}$ donc nous avons n choix. Pour l'image de 2, il ne reste plus que $(n-1)$ choix (1 et 2 ne doivent pas avoir la même image car notre application est une bijection). Ainsi de suite, pour l'image du dernier élément n il ne reste qu'une possibilité. Au final il y a

$$n \cdot (n-1) \cdot \dots \cdot 2 \cdot 1 = n!$$

façons de construire des bijections de $\{1, 2, \dots, n\}$. □

Dans le contexte du groupe symétrique $\mathcal{S}_n$, il est important de comprendre les concepts de support et de point fixe d'une permutation.

Définition 4.5.3 (Support et point fixe d'une permutation)

1. Soit σ un élément de $\mathcal{S}_n$, un point $i \in E = \{1, 2, \dots, n\}$ est dit point fixe de σ (ou invariant par σ) lorsque

$$\sigma(i) = i.$$

2. L'ensemble des points fixes de σ est noté $Fix(\sigma)$.

3. Soit $\sigma \in \mathcal{S}_n$, on appelle support de σ et on note $Supp(\sigma)$, l'ensemble des éléments i de $E = \{1, 2, \dots, n\}$ tels que $\sigma(i) \neq i$,

$$Supp(\sigma) = \{i \in \{1, 2, \dots, n\} : \sigma(i) \neq i\} \subset E.$$

Autrement dit, l'ensemble $Supp(\sigma)$ est le complémentaire dans $E = \{1, 2, \dots, n\}$ de l'ensemble des points fixes de σ ,

$$Supp(\sigma) = C_E^{Fix(\sigma)}.$$

Ou de manière équivalente

$$Fix(\sigma) = C_E^{Supp(\sigma)}.$$

Exemple 4.5.3

1. Le support de l'identité est l'ensemble vide.

$$\text{Supp}(Id_E) = \phi.$$

2. Pour la permutation

$$\sigma = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 & 6 \\ 1 & 3 & 6 & 4 & 2 & 5 \end{pmatrix} \in S_6.$$

On a

$$\begin{aligned} \text{Fix}(\sigma) &= \{1, 4\}, \\ \text{Supp}(\sigma) &= \{2, 3, 5, 6\}. \end{aligned}$$

Corollaire 4.5.2 Pour toute permutation $\sigma \in \mathcal{S}_n$, on a les égalités

$$\sigma(\text{Fix}(\sigma)) = \text{Fix}(\sigma), \sigma(\text{Supp}(\sigma)) = \text{Supp}(\sigma).$$

Les permutations de E sont définies comme des applications de E dans E , il est donc possible de définir leur produit de composition, qui se note $\circ$.

Définition 4.5.4 (Composition de permutations)

La composition (ou produit) de deux permutations σ_1 et σ_2 de S_n est la permutation $\sigma = \sigma_1 \circ \sigma_2$ obtenue en appliquant σ_2 puis σ_1 au résultat :

$$\forall i \in \{1, 2, \dots, n\} : (\sigma_1 \circ \sigma_2)(i) = \sigma_1(\sigma_2(i)).$$

C'est-à-dire si

$$\sigma_1 = \begin{pmatrix} 1 & 2 & \dots & k & \dots & n \\ \sigma_1(1) & \sigma_1(2) & \dots & \sigma_1(k) & \dots & \sigma_1(n) \end{pmatrix}, \sigma_2 = \begin{pmatrix} 1 & 2 & \dots & k & \dots & n \\ \sigma_2(1) & \sigma_2(2) & \dots & \sigma_2(k) & \dots & \sigma_2(n) \end{pmatrix},$$

alors

$$\sigma_1 \circ \sigma_2 = \begin{pmatrix} 1 & 2 & \dots & k & \dots & n \\ \sigma_1(\sigma_2(1)) & \sigma_1(\sigma_2(2)) & \dots & \sigma_1(\sigma_2(k)) & \dots & \sigma_1(\sigma_2(n)) \end{pmatrix}.$$

Il suffit d'écrire les images par σ_1 des images par σ_2 de la permutation initiale.

Exemple 4.5.4

1. Supposons que nous ayons les permutations σ_1 et σ_2 de $\mathcal{S}_3$ données par

$$\sigma_1(1) = 2, \sigma_1(2) = 3, \sigma_1(3) = 1,$$

et

$$\sigma_2(1) = 1, \sigma_2(2) = 3, \sigma_2(3) = 2.$$

Alors

$$\begin{aligned} (\sigma_1 \circ \sigma_2)(1) &= \sigma_1(\sigma_2(1)) = \sigma_1(1) = 2, \\ (\sigma_1 \circ \sigma_2)(2) &= \sigma_1(\sigma_2(2)) = \sigma_1(3) = 1, \\ (\sigma_1 \circ \sigma_2)(3) &= \sigma_1(\sigma_2(3)) = \sigma_1(2) = 3. \end{aligned}$$

Autrement dit,

$$\begin{pmatrix} 1 & 2 & 3 \\ 2 & 3 & 1 \end{pmatrix} \circ \begin{pmatrix} 1 & 2 & 3 \\ 1 & 3 & 2 \end{pmatrix} = \begin{pmatrix} 1 & 2 & 3 \\ \sigma_1(\sigma_2(1)) & \sigma_1(\sigma_2(2)) & \sigma_1(\sigma_2(3)) \end{pmatrix} = \begin{pmatrix} 1 & 2 & 3 \\ 2 & 1 & 3 \end{pmatrix}.$$

2. Soient $\sigma_1, \sigma_2 \in \mathcal{S}_5$ telles que

$$\sigma_1 = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 \\ 2 & 5 & 3 & 1 & 4 \end{pmatrix}, \sigma_2 = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 \\ 3 & 5 & 2 & 1 & 4 \end{pmatrix},$$

alors

$$\sigma_1 \circ \sigma_2 = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 \\ 2 & 5 & 3 & 1 & 4 \end{pmatrix} \circ \begin{pmatrix} 1 & 2 & 3 & 4 & 5 \\ 3 & 5 & 2 & 1 & 4 \end{pmatrix} = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 \\ 3 & 4 & 5 & 2 & 1 \end{pmatrix}.$$

Remarque 4.5.3 En général, le produit de deux permutations n'est pas commutatif, par exemple

$$\sigma_1 = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 \\ 2 & 3 & 4 & 5 & 1 \end{pmatrix}, \sigma_2 = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 \\ 3 & 2 & 5 & 1 & 4 \end{pmatrix} \in \mathcal{S}_5$$

alors

$$\begin{aligned} \sigma_1 \circ \sigma_2 &= \begin{pmatrix} 1 & 2 & 3 & 4 & 5 \\ 2 & 3 & 4 & 5 & 1 \end{pmatrix} \circ \begin{pmatrix} 1 & 2 & 3 & 4 & 5 \\ 3 & 2 & 5 & 1 & 4 \end{pmatrix} = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 \\ 4 & 3 & 1 & 2 & 5 \end{pmatrix} \\ \neq \sigma_2 \circ \sigma_1 &= \begin{pmatrix} 1 & 2 & 3 & 4 & 5 \\ 3 & 2 & 5 & 1 & 4 \end{pmatrix} \circ \begin{pmatrix} 1 & 2 & 3 & 4 & 5 \\ 2 & 3 & 4 & 5 & 1 \end{pmatrix} = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 \\ 2 & 5 & 1 & 4 & 3 \end{pmatrix}. \end{aligned}$$

Définition 4.5.5 (Permutation inverse)

La permutation inverse (ou réciproque) de la permutation σ de $\mathcal{S}_n$ est la permutation σ^{-1} telle que

$$\sigma \circ \sigma^{-1} = \sigma^{-1} \circ \sigma = Id_E.$$

Autrement dit

$$\forall i \in \{1, 2, \dots, n\} : \sigma^{-1}(\sigma(i)) = \sigma(\sigma^{-1}(i)) = i.$$

Exemple 4.5.5 Soit

$$\sigma = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 \\ 3 & 2 & 5 & 1 & 4 \end{pmatrix} \in \mathcal{S}_5,$$

alors

$$\sigma^{-1} = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 \\ 4 & 2 & 1 & 5 & 3 \end{pmatrix},$$

car

$$\begin{cases} (\sigma \circ \sigma^{-1})(1) = 1 \\ (\sigma \circ \sigma^{-1})(2) = 2 \\ (\sigma \circ \sigma^{-1})(3) = 3 \\ (\sigma \circ \sigma^{-1})(4) = 4 \\ (\sigma \circ \sigma^{-1})(5) = 5 \end{cases} \implies \begin{cases} \sigma(\sigma^{-1}(1)) = 1 \\ \sigma(\sigma^{-1}(2)) = 2 \\ \sigma(\sigma^{-1}(3)) = 3 \\ \sigma(\sigma^{-1}(4)) = 4 \\ \sigma(\sigma^{-1}(5)) = 5 \end{cases} \implies \begin{cases} \sigma^{-1}(1) = 4 \\ \sigma^{-1}(2) = 2 \\ \sigma^{-1}(3) = 1 \\ \sigma^{-1}(4) = 5 \\ \sigma^{-1}(5) = 3. \end{cases}$$

Remarque 4.5.4

1. Toutes les permutations sont inversibles.

2. Il est tout aussi facile de calculer l'inverse d'une permutation : Il suffit donc de permuter les deux lignes (du haut et du bas) et de réordonner la première. Par exemple l'inverse de

$$\sigma = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 & 6 & 7 \\ 3 & 7 & 5 & 4 & 6 & 1 & 2 \end{pmatrix} \in \mathcal{S}_7,$$

se note

$$\begin{pmatrix} 3 & 7 & 5 & 4 & 6 & 1 & 2 \\ 1 & 2 & 3 & 4 & 5 & 6 & 7 \end{pmatrix},$$

ou plutôt après réordonnement

$$\sigma^{-1} = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 & 6 & 7 \\ 6 & 7 & 1 & 4 & 3 & 5 & 2 \end{pmatrix}.$$

3. Pour toute permutation $\sigma \in \mathcal{S}_n$, on a les égalités

$$\text{Fix}(\sigma) = \text{Fix}(\sigma^{-1}), \text{Supp}(\sigma) = \text{Supp}(\sigma^{-1}).$$

Proposition 4.5.2 Soit $n \in \mathbb{N}^*$, alors $\mathcal{S}_n$ muni de la composition des applications $\circ$, forme un groupe, appelé le groupe symétrique de $E = \{1, \dots, n\}$.

Le groupe symétrique $\mathcal{S}_n$ est donc un exemple fondamental de groupe en algèbre, avec des applications vastes en mathématiques et dans d'autres domaines.

Preuve. La preuve est simple.

- L'identité Id_E est une permutation de $E = \{1, 2, \dots, n\}$ donc $\mathcal{S}_n$ n'est pas vide.
- La composition de deux bijections de $E = \{1, 2, \dots, n\}$ est une bijection de $E = \{1, 2, \dots, n\}$, donc on a une loi interne.
- La composition est clairement associative (par l'associativité de la composition des applications).
- L'identité Id_E est l'élément neutre pour la composition.
- Enfin, tout élément de $\mathcal{S}_n$ est inversible d'inverse son application réciproque.

Ainsi, $(\mathcal{S}_n, \circ)$ satisfait toutes les propriétés nécessaires pour être un groupe. $\square$

Remarque 4.5.5 Pour $n \geq 3$, le groupe $(\mathcal{S}_n, \circ)$ n'est pas commutatif. Par exemple, considérons le groupe $\mathcal{S}_3$ et les permutations suivantes

$$\sigma_1 = \begin{pmatrix} 1 & 2 & 3 \\ 2 & 1 & 3 \end{pmatrix}, \sigma_2 = \begin{pmatrix} 1 & 2 & 3 \\ 3 & 2 & 1 \end{pmatrix}.$$

Alors

$$\begin{aligned} \sigma_1 \circ \sigma_2 &= \begin{pmatrix} 1 & 2 & 3 \\ 2 & 1 & 3 \end{pmatrix} \circ \begin{pmatrix} 1 & 2 & 3 \\ 3 & 2 & 1 \end{pmatrix} = \begin{pmatrix} 1 & 2 & 3 \\ 3 & 1 & 2 \end{pmatrix} \\ &\neq \sigma_2 \circ \sigma_1 = \begin{pmatrix} 1 & 2 & 3 \\ 3 & 2 & 1 \end{pmatrix} \circ \begin{pmatrix} 1 & 2 & 3 \\ 2 & 1 & 3 \end{pmatrix} = \begin{pmatrix} 1 & 2 & 3 \\ 2 & 3 & 1 \end{pmatrix}, \end{aligned}$$

cela montre que $\mathcal{S}_3$ n'est pas commutatif.

4.5.2 Transpositions et cycles

Une transposition est un type particulier de permutation qui échange exactement deux éléments d'un ensemble tout en laissant tous les autres éléments inchangés. Les transpositions sont fondamentales dans la théorie des permutations car toute permutation peut être décomposée en un produit de transpositions.

Définition 4.5.6

1. Soit $n \geq 2$, pour tout (i, j) de $\{1, 2, \dots, n\}^2$ tel que $i < j$, on appelle transposition sur i, j et on note $\tau_{i,j}$ (ou τ_{ij} , ou (i, j)) toute permutation de $\mathcal{S}_n$ définie par

$$\begin{cases} \tau_{i,j}(i) = j \\ \tau_{i,j}(j) = i \\ \forall k \in \{1, 2, \dots, n\} - \{i, j\} : \tau_{i,j}(k) = k. \end{cases}$$

Une transposition est donc de la forme

$$\tau_{i,j} = \begin{pmatrix} 1 & 2 & \dots & i & \dots & j & \dots & n \\ 1 & 2 & \dots & j & \dots & i & \dots & n \end{pmatrix}.$$

Remarque 4.5.6

1. Une transposition est une permutation qui échange deux éléments distincts et laisse les autres inchangés.
2. Pour une transposition $\tau_{i,j}$, le support est simplement l'ensemble $\{i, j\}$,

$$\text{Supp}(\tau_{i,j}) = \{i, j\}.$$

de plus toute transposition est involutive,

$$\tau_{i,j} \circ \tau_{i,j} = \text{Id}_E,$$

on a donc

$$\tau_{i,j}^{-1} = \tau_{i,j}.$$

3. Le nombre de transpositions dans le groupe symétrique $\mathcal{S}_n$ est donné par le nombre de façons de choisir deux éléments parmi n éléments, puisque chaque transposition échange deux éléments distincts et laisse les autres inchangés. Pour calculer ce nombre, on utilise la combinaison C_n^2 , qui représente le nombre de façons de choisir 2 éléments parmi n . La formule de la combinaison est la suivante

$$C_n^2 = \frac{n(n-1)}{2}.$$

Exemple 4.5.6 Considérons le groupe symétrique $\mathcal{S}_3$, qui est le groupe de toutes les permutations de l'ensemble $\{1, 2, 3\}$.

- La transposition $\tau_{1,2}$ échange les éléments 1 et 2,

$$\tau_{1,2} = \begin{pmatrix} 1 & 2 & 3 \\ 2 & 1 & 3 \end{pmatrix}.$$

- La transposition $\tau_{1,3}$ échange les éléments 1 et 3,

$$\tau_{1,3} = \begin{pmatrix} 1 & 2 & 3 \\ 3 & 2 & 1 \end{pmatrix}.$$

- La transposition $\tau_{2,3}$ échange les éléments 2 et 3,

$$\tau_{2,3} = \begin{pmatrix} 1 & 2 & 3 \\ 1 & 3 & 2 \end{pmatrix}.$$

2. Pour $n = 5$, on a la transposition $\tau_{2,4}$ échange les éléments 2 et 4,

$$\tau_{2,4} = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 \\ 1 & 4 & 3 & 2 & 5 \end{pmatrix}.$$

Introduisons maintenant la notion de cycle. Dans le groupe symétrique $\mathcal{S}_n$ un cycle est une permutation qui déplace un certain nombre d'éléments de manière cyclique, laissant les autres éléments inchangés.

Définition 4.5.7

1. Soit $n \geq 2$ et $p \in \mathbb{N}$ tel que $2 \leq p \leq n$. On appelle cycle de longueur p ou p -cycle ou encore cycle d'ordre p toute permutation notée c de $\mathcal{S}_n$ telle qu'il existe $x_1, \dots, x_p \in \{1, 2, \dots, n\}$, deux à deux distincts, vérifiant

$$\begin{cases} c(x_1) = x_2, c(x_2) = x_3, \dots, c(x_{p-1}) = x_p, c(x_p) = x_1, \\ \forall k \in \{1, 2, \dots, n\} - \{x_1, \dots, x_p\} : c(k) = k. \end{cases}$$

2. L'ensemble $\{x_1, \dots, x_p\}$ est appelé le support de c , et on note $c = (x_1, \dots, x_p)$.

3. Deux cycles c_1 et c_2 sont disjoints si leurs supports sont disjoints.

Remarque 4.5.7

1. Un cycle de longueur 2 est appelé une transposition.

2. L'identité Id_E n'est pas considérée comme un cycle dans le contexte des permutations.

3. Il y a plusieurs façons d'écrire le même p -cycle. Par exemple, les cycles suivants sont identiques :

$$(x_1, \dots, x_p) = (x_2, \dots, x_p, x_1) = \dots = (x_p, x_1, \dots, x_{p-1}).$$

Exemple 4.5.7

1. Dans $\mathcal{S}_5$, le cycle $c = (1, 2, 3)$ est donc la permutation envoyant 1 sur 2, 2 sur 3 et 3 sur 1. Nous avons donc

$$c(1) = 2, c(2) = 3, c(3) = 1.$$

Pour tout $i \notin \{1, 2, 3\}$, nous avons alors $c(i) = i$.

2. Dans $\mathcal{S}_6$, la permutation

$$\begin{pmatrix} 1 & 2 & 3 & 4 & 5 & 6 \\ 1 & 5 & 2 & 4 & 3 & 6 \end{pmatrix},$$

est un cycle de support $\{2, 3, 5\}$ et de longueur 3. On le note plus simplement $(2, 5, 3)$ (ou $(5, 3, 2)$ ou $(3, 2, 5)$).

Définition 4.5.8 (Permutation circulaire)

1. Une permutation circulaire de $\{1, 2, \dots, n\}$ est une permutation qui déplace chaque élément de cet ensemble à une position suivante dans un cycle unique de longueur n .

2. En d'autres termes, une permutation circulaire de $\{1, 2, \dots, n\}$ est un cycle unique de longueur n qui peut être représenté sous la forme $(x_1, \dots, x_n)$, où chaque élément x_i est envoyé à la position de x_{i+1} , et x_n est envoyé à la position de x_1 .

Exemple 4.5.8 Pour $n = 4$, considérons l'ensemble $\{1, 2, 3, 4\}$. Une permutation circulaire de cet ensemble pourrait être

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

Ceci peut être noté comme le cycle $c = (1, 2, 3, 4)$.

Proposition 4.5.3 (Propriétés)

1. Soit $n \geq 2$ et $p \in \mathbb{N}$ tel que $2 \leq p \leq n$. Soit $c = (x_1, \dots, x_p)$ un cycle de longueur p , alors

$$c^p = c \circ c \circ \dots \circ c = \text{Id}_{\{1, 2, \dots, n\}} \text{ et } c^{-1} = c^{p-1} = (x_p, \dots, x_1).$$

2. Si $c_1 = (x_1, \dots, x_p)$ et $c_2 = (y_1, \dots, y_q)$ deux cycles à supports disjoints, alors c_1 et c_2 commutent.

$$\{x_1, \dots, x_p\} \cap \{y_1, \dots, y_q\} = \emptyset \implies (x_1, \dots, x_p) \circ (y_1, \dots, y_q) = (y_1, \dots, y_q) \circ (x_1, \dots, x_p).$$

La décomposition d'une permutation en produit de cycles à supports disjoints est une méthode pour représenter une permutation en termes de cycles indépendants qui ne partagent aucun élément. Le théorème suivant est admis.

Théorème 4.5.1 (*Décomposition en produit de cycles à support disjoints*)

Soit $n \geq 2$ et soit $\sigma \in \mathcal{S}_n$. Il existe un entier naturel k et des cycles $c_1, c_2, \dots, c_k$ de $\{1, 2, \dots, n\}$ à supports disjoints tels que

$$\sigma = c_1 \circ c_2 \circ \dots \circ c_k.$$

De plus ces cycles commutent et cette décomposition est unique à l'ordre près des facteurs.

Autrement dit, toute permutation différente de l'identité se décompose, de manière unique à l'ordre près des facteurs, en un produit de cycles à support deux à deux disjoints.

Remarque 4.5.8 On peut convenir que l'identité $e = Id$ est décomposable en produit vide de cycles.

Exemple 4.5.9

1. Soit

$$\sigma_1 = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 & 6 & 7 & 8 & 9 & 10 \\ 3 & 6 & 5 & 10 & 7 & 2 & 9 & 4 & 1 & 8 \end{pmatrix} \in \mathcal{S}_{10}.$$

Pour décomposer σ_1 en produit de cycles, on commence par chercher l'image de 1 : c'est 3. Ensuite, on regarde quelle est l'image de 3, c'est 5, puis quelle est l'image de 5 et ainsi de suite jusqu'à ce qu'on arrive à 1. On a ainsi formé le cycle $(1, 3, 5, 7, 9)$. Ensuite, on se demande si tous les éléments restant sont des points fixes ou non. Si oui, on a terminé la décomposition, sinon, on cherche quel est le plus petit élément ne faisant pas partie des cycles précédemment trouvés qui n'est pas point fixe. Ici c'est 2 et en regardant quels sont les itérés de 2 on obtient le cycle $(2, 6)$. On itère ce procédé jusqu'à ce qu'il ne reste que des points fixes. Ainsi

$$\sigma_1 = (1, 3, 5, 7, 9) \circ (2, 6) \circ (4, 10, 8).$$

2. Dans $\mathcal{S}_{10}$, on a

$$\sigma_2 = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 & 6 & 7 & 8 & 9 & 10 \\ 4 & 6 & 9 & 1 & 5 & 8 & 2 & 7 & 10 & 3 \end{pmatrix} = (1, 4) \circ (2, 6, 8, 7) \circ (3, 9, 10).$$

4.6 Morphismes de Groupes

Nous allons maintenant étudier un type particulier d'applications qui agissent sur les groupes. Ces applications envoient le produit de deux éléments du groupe de départ sur le produit des images des éléments du groupe d'arrivée et sont connues sous le nom de morphismes de groupes. Alors un morphisme de groupe est une application entre deux groupes qui préserve la structure de groupe.

4.6.1 Définitions et exemples

Définition 4.6.1 Soient $(G_1, *)$ et $(G_2, \top)$ deux groupes. Soit $f : (G_1, *) \longrightarrow (G_2, \top)$ une application de G_1 dans G_2 . On dit que f est un morphisme (ou encore homomorphisme) de groupes de G_1 dans G_2 si et seulement si

$$\forall x, y \in G_1 : f(x * y) = f(x) \top f(y).$$

Cela signifie que l'image du produit (ou de la composition) de deux éléments dans le groupe G_1 est égale au produit (ou de la composition) des images de ces éléments dans le groupe G_2 .

Définition 4.6.2 Soient $(G_1, *)$ et $(G_2, \top)$ deux groupes. Soit $f : (G_1, *) \longrightarrow (G_2, \top)$ un morphisme de groupes.

1. Si $G_1 = G_2$ et $* = \top$, alors on dit f est endomorphisme de groupe.
2. Si f est une application bijective, alors on dit que f est un isomorphisme de groupes.
3. Si $G_1 = G_2$ et $* = \top$ et f est une application bijective, alors on dit que f est un automorphisme de groupes.

Remarque 4.6.1

1. Quand il existe un isomorphisme entre deux groupes $(G_1, *)$ et $(G_2, \top)$, on dit que ces deux groupes sont isomorphes et on note

$$(G_1, *) \simeq (G_2, \top) \text{ ou } G_1 \simeq G_2.$$

2. Lorsque deux groupes $(G_1, *)$ et $(G_2, \top)$ sont isomorphes, cela signifie qu'ils ont essentiellement les mêmes propriétés structurelles et algébriques, même s'ils peuvent être représentés différemment.

Exemple 4.6.1

1. Pour tout groupe $(G, *)$, l'application identité Id_G est un automorphisme de groupe $(G, *)$.
2. L'application exponentielle $f_1 = \exp$ est un isomorphisme de groupe additif $(\mathbb{R}, +)$ dans le groupe multiplicatif $(\mathbb{R}_+^*, \cdot)$,

$$\begin{aligned} f_1 : (\mathbb{R}, +) &\longrightarrow (\mathbb{R}_+^*, \cdot) \\ x &\longmapsto f_1(x) = \exp(x). \end{aligned}$$

En effet, soient $x, y \in \mathbb{R}$, alors

$$f_1(x + y) = \exp(x + y) = \exp(x) \cdot \exp(y) = f_1(x) \cdot f_1(y),$$

donc f_1 est un morphisme de groupes, et comme l'application exponentielle est bijective, f_1 est en fait un isomorphisme de groupes.

3. L'application logarithme népérien $f_2 = \ln$ est un isomorphisme du groupe multiplicatif $(\mathbb{R}_+^*, \cdot)$ dans le groupe additif $(\mathbb{R}, +)$,

$$\begin{aligned} f_2 : (\mathbb{R}_+^*, \cdot) &\longrightarrow (\mathbb{R}, +) \\ x &\longmapsto f_2(x) = \ln x. \end{aligned}$$

En effet, soient $x, y \in \mathbb{R}_+^*$, alors

$$f_2(x \cdot y) = \ln(x \cdot y) = \ln(x) + \ln(y) = f_2(x) + f_2(y),$$

donc f_2 est un morphisme de groupes, et comme l'application logarithme népérien est bijective, f_2 est en fait un isomorphisme de groupes.

4. L'application

$$\begin{aligned} f_3 : (\mathbb{R}^*, \cdot) &\longrightarrow (\mathbb{R}_+^*, \cdot) \\ x &\longmapsto f_3(x) = |x|, \end{aligned}$$

est un morphisme de groupes. En effet, pour tout $x, y \in \mathbb{R}^*$, on a

$$f_3(x \cdot y) = |x \cdot y| = |x| \cdot |y| = f_3(x) \cdot f_3(y).$$

5. Soit $(G, *)$ un groupe d'élément neutre e . L'application constante

$$\begin{aligned} f_4 : (G, *) &\longrightarrow (G, *) \\ x &\longmapsto f(x) = e, \end{aligned}$$

est un endomorphisme de G .

6. L'application

$$\begin{aligned} f_5 : (\mathbb{Z}, +) &\longrightarrow (\mathbb{Z}, +) \\ x &\longmapsto f_5(x) = -x, \end{aligned}$$

est un automorphisme de groupe $(\mathbb{Z}, +)$ car elle est bijective et préserve l'opération d'addition.

7. L'application

$$\begin{aligned} f_6 : (\mathbb{Z}, +) &\longrightarrow (\mathbb{Z}, +) \\ x &\longmapsto f_6(x) = x^2, \end{aligned}$$

n'est pas un morphisme de groupes. En effet,

$$f_6(1 + 2) = f_6(3) = 9 \neq f_6(1) + f_6(2) = 5.$$

Proposition 4.6.1 Soient $(G_1, *)$ et $(G_2, \top)$ deux groupes. Soit $f : (G_1, *) \longrightarrow (G_2, \top)$ un morphisme de groupes. Si e_1 est l'élément neutre de G_1 et e_2 l'élément neutre de G_2 , alors

1. **Préservation de l'élément neutre** : l'image par un morphisme de groupes de l'élément neutre du groupe de départ est égale à l'élément neutre du groupe d'arrivée, on a donc

$$f(e_1) = e_2.$$

2. **Préservation des inverses** : Le symétrique de l'image d'un élément du groupe de départ par un morphisme de groupes est égal à l'image du symétrique de cet élément, on a donc

$$\forall x \in G_1 : (f(x))' = f(x').$$

Preuve.

1. Remarquons que

$$f(e_1) = f(e_1 * e_1) = f(e_1) \top f(e_1).$$

En composant (à droite par exemple) $(f(e_1))'$, on obtient

$$f(e_1) \top (f(e_1))' = f(e_1) \top f(e_1) \top (f(e_1))'.$$

Ce qui implique que

$$e_2 = f(e_1) \top e_2.$$

Ainsi

$$f(e_1) = e_2.$$

2. Soit $x \in G_1$, alors

$$x * x' = e_1,$$

donc

$$f(x * x') = f(e_1).$$

Cela entraîne

$$f(x) \top f(x') = f(e_1).$$

En composant à gauche $(f(x))'$, nous obtenons

$$f(x') = (f(x))'.$$

□

Exemple 4.6.2 Reprenons l'exemple de l'application exponentielle $f = \exp : (\mathbb{R}, +) \longrightarrow (\mathbb{R}_+^*, \cdot)$. Nous avons bien

$$f(0) = \exp(0) = 1,$$

l'élément neutre 0 de $(\mathbb{R}, +)$ a pour image l'élément neutre 1 de $(\mathbb{R}_+^*, \cdot)$. Pour $x \in \mathbb{R}$ son symétrique x' dans $(\mathbb{R}, +)$ est ici son opposé $(-x)$, alors

$$f(x') = f(-x) = e^{-x} = \frac{1}{e^x} = \frac{1}{f(x)} = (f(x))^{-1} = (f(x))',$$

est bien l'inverse (dans $(\mathbb{R}_+^*, \cdot)$) de $f(x)$.

Corollaire 4.6.1 Soient $(G_1, *)$ et $(G_2, \top)$ deux groupes. Soit $f : (G_1, *) \longrightarrow (G_2, \top)$ un morphisme de groupes. Alors

$$\forall x_1, x_2, \dots, x_n \in G_1 : f(x_1 * x_2 * \dots * x_n) = f(x_1) \top f(x_2) \top \dots \top f(x_n),$$

est une généralisation de la propriété de préservation des opérations pour les morphismes de groupes.

Proposition 4.6.2 (Composition de morphismes de groupes)

Soient $(G_1, *)$, $(G_2, \top)$ et (G_3, Δ) trois groupes. Soit $f : (G_1, *) \longrightarrow (G_2, \top)$, $g : (G_2, \top) \longrightarrow (G_3, \Delta)$ deux morphismes de groupes, alors $g \circ f : (G_1, *) \longrightarrow (G_3, \Delta)$ est un morphisme de groupes. Autrement dit la composée de deux morphismes de groupes est un morphisme de groupe.

Preuve. Soient $x_1, x_2 \in G_1$, alors

$$(g \circ f)(x_1 * x_2) = g(f(x_1 * x_2)) = g(f(x_1) \top f(x_2)) = g(f(x_1)) \Delta g(f(x_2)) = (g \circ f)(x_1) \Delta (g \circ f)(x_2).$$

Par conséquent $g \circ f$ est un morphisme de groupes. $\square$

Proposition 4.6.3 Soient $(G_1, *)$ et $(G_2, \top)$ deux groupes. Si $f : (G_1, *) \longrightarrow (G_2, \top)$ est un isomorphisme de groupes, alors $f^{-1} : (G_2, \top) \longrightarrow (G_1, *)$ est un isomorphisme de groupes aussi. Autrement dit la bijection réciproque d'un isomorphisme de groupe est un isomorphisme de groupe.

Preuve. Soit f un morphisme de groupes bijectif. Soient $y_1, y_2 \in G_2$, montrons que

$$f^{-1}(y_1 \top y_2) = f^{-1}(y_1) * f^{-1}(y_2).$$

Comme $f^{-1}(y_1)$ et $f^{-1}(y_2)$ sont dans G_1 et f est un morphisme de groupes, on a

$$f(f^{-1}(y_1) * f^{-1}(y_2)) = f(f^{-1}(y_1)) \top f(f^{-1}(y_2)) = (f \circ f^{-1})(y_1) \top (f \circ f^{-1})(y_2) = y_1 \top y_2.$$

En composant l'équation précédente par f^{-1} , on obtient

$$f^{-1}(f(f^{-1}(y_1) * f^{-1}(y_2))) = f^{-1}(y_1 \top y_2).$$

Par suite

$$f^{-1}(y_1) * f^{-1}(y_2) = f^{-1}(y_1 \top y_2).$$

Alors f^{-1} est un morphisme de groupes et il est de plus bien connu que f^{-1} est bijective. Ainsi f^{-1} est un isomorphisme de groupes. $\square$

4.6.2 Image et noyau d'un morphisme de groupes

Le noyau d'un morphisme de groupes est effectivement un concept central en théorie des groupes. Il joue un rôle crucial dans l'étude des structures de groupe et des propriétés des morphismes. L'image d'un morphisme de groupes est un autre concept fondamental en théorie des groupes. Cette sous-section introduit deux sous groupes particulièrement importants que l'on appelle image et noyau d'un morphisme de groupes.

Définition 4.6.3 Soient $(G_1, *)$ et $(G_2, \top)$ deux groupes. Soit $f : (G_1, *) \longrightarrow (G_2, \top)$ un morphisme de groupes. Si e_1 est l'élément neutre de G_1 et e_2 l'élément neutre de G_2 , alors

1. On appelle noyau de f , et on note $\ker(f)$, l'ensemble des éléments de G_1 qui sont envoyés sur l'élément neutre de G_2 par f , on a donc

$$\ker(f) = \{x \in G_1 : f(x) = e_2\} = f^{-1}(\{e_2\}) \subset G_1.$$

En d'autres termes, le noyau de f est l'ensemble de tous les éléments de G_1 qui sont "annulés" par le morphisme f .

2. On appelle image de f , et on note $\text{Im}(f)$ ou $f(G_1)$, l'ensemble défini par

$$\text{Im}(f) = \{y \in G_2, \exists x \in G_1 : f(x) = y\} = \{f(x) \in G_2 : x \in G_1\} = f(G_1) \subset G_2.$$

En d'autres termes, l'image de f , est l'ensemble des éléments de G_2 qui sont atteints par f .

Remarque 4.6.2

1. La notation $\ker$ vient du mot Allemand **Kern** dont la traduction est noyau.
2. Pour déterminer $\ker(f)$, on résout l'équation $f(x) = e_2$ d'inconnue $x \in G_1$.
3. Pour déterminer $\text{Im}(f)$, on étudie les valeurs prises par f .
4. Comme $f(e_1) = e_2$, le noyau d'un morphisme n'est jamais vide.

Exemple 4.6.3

1. Le noyau (resp. l'image) de l'application exponentielle $f_1 = \exp : (\mathbb{R}, +) \longrightarrow (\mathbb{R}_+^*, \cdot)$ est

$$\ker(f_1) = \{x \in \mathbb{R} : f_1(x) = 1\} = \{0\}.$$

(resp.

$$\text{Im}(f_1) = \{f_1(x) = \exp(x) : x \in \mathbb{R}\} = \mathbb{R}_+^*).$$

2. Le noyau (resp. l'image) de l'application logarithme népérien $f_2 = \ln : (\mathbb{R}_+^*, \cdot) \longrightarrow (\mathbb{R}, +)$ est

$$\ker(f_2) = \{x \in \mathbb{R}_+^* : f_2(x) = 0\} = \{1\}.$$

(resp.

$$\text{Im}(f_2) = \{f_2(x) = \ln(x) : x \in \mathbb{R}_+^*\} = \mathbb{R}).$$

3. Le noyau (resp. l'image) de l'application

$$\begin{aligned} f_3 : (\mathbb{R}^*, \cdot) &\longrightarrow (\mathbb{R}_+^*, \cdot) \\ x &\longmapsto f_3(x) = |x|, \end{aligned}$$

est

$$\ker(f_3) = \{x \in \mathbb{R}^* : f_3(x) = 1\} = \{x \in \mathbb{R}^* : |x| = 1\} = \{-1, 1\}.$$

(resp.

$$\text{Im}(f_3) = \{f_3(x) = |x| : x \in \mathbb{R}^*\} = \mathbb{R}_+^*).$$

4. Pour $k \in \mathbb{Z}$, l'application

$$\begin{aligned} f_4 : (\mathbb{Z}, +) &\longrightarrow (\mathbb{Z}, +) \\ x &\longmapsto f_4(x) = kx, \end{aligned}$$

est un endomorphisme de groupes. Son noyau dépend de la valeur de k , alors

$$\ker(f_4) = \{x \in \mathbb{Z} : f_4(x) = 0\} = \{x \in \mathbb{Z} : kx = 0\} = \begin{cases} \{0\}, & \text{si } k \neq 0 \\ \mathbb{Z}, & \text{si } k = 0. \end{cases}$$

et

$$\text{Im}(f_4) = \{f_4(x) \in \mathbb{Z} : x \in \mathbb{Z}\} = k\mathbb{Z}.$$

4.6.3 Image directe et réciproque de sous-groupes par un morphisme

L'image directe et réciproque de sous-groupes jouent un rôle clé dans l'analyse des morphismes de groupes. L'image directe nous montre comment les sous-groupes du groupe source sont transformés dans le groupe cible, tandis que l'image réciproque nous indique quels éléments du groupe source sont envoyés dans un sous-groupe du groupe cible.

Proposition 4.6.4 (*Image directe et réciproque de sous-groupes par un morphisme*)

Soient $(G_1, *)$ et $(G_2, \top)$ deux groupes. Soit $f : (G_1, *) \longrightarrow (G_2, \top)$ un morphisme de groupes. Alors

1. Si H_1 est un sous-groupe de G_1 , alors $f(H_1)$ est un sous-groupe de G_2 .
2. Si H_2 est un sous-groupe de G_2 , alors $f^{-1}(H_2)$ est un sous-groupe de G_1 .

Preuve.

1. Soit H_1 un sous-groupe de G_1 .

- Comme $e_2 = f(e_1)$ et que $e_1 \in H_1$ alors $e_2 \in f(H_1)$.
- Soient y_1 et y_2 deux éléments de $f(H_1)$, alors il existe $x_1, x_2 \in H_1$ tels que

$$f(x_1) = y_1 \text{ et } f(x_2) = y_2,$$

et par suite

$$y_1 \top y_2 = f(x_1) \top f(x_2) = f(x_1 * x_2) \in f(H_1).$$

car $x_1 * x_2 \in H_1$ puisque H_1 est un sous-groupe.

- Soit $y \in f(H_1)$, alors il existe $x \in H_1$ tel que $y = f(x)$ et par suite on a

$$y' = (f(x))' = f(x') \in f(H_1),$$

car $x' \in H_1$ puisque H_1 est un sous-groupe. Ainsi $f(H_1)$ est un sous-groupe de G_2 .

2. Soit H_2 un sous-groupe de G_2 .

- Comme $e_2 = f(e_1)$ et que $e_2 \in H_2$ alors $e_1 \in f^{-1}(H_2)$.
- Soient $x_1, x_2 \in f^{-1}(H_2)$. On a alors

$$f(x_1), f(x_2) \in H_2,$$

d'où, puisque H_2 est un sous-groupe,

$$f(x_1 * x_2) = f(x_1) \top f(x_2) \in H_2.$$

On a donc bien $x_1 * x_2 \in f^{-1}(H_2)$.

- Soit $x \in f^{-1}(H_2)$. On a alors $f(x) \in H_2$ et donc

$$(f(x))' = f(x') \in H_2,$$

puisque H_2 est un sous-groupe. On a donc bien $x' \in f^{-1}(H_2)$. □

Remarque 4.6.3

1. Le cas particulier où le sous groupe H_2 est le sous groupe trivial de G_2 , à savoir $(\{e_{G_2}\}, \circ)$, joue un rôle particulièrement important dans l'étude d'un morphisme de groupes. Le sous groupe obtenu dans ce cas s'appelle le noyau de f .
2. Le cas H_1 est égal au groupe G_1 définit l'image de f . Comme le verrez plus loin, ce sous groupe joue un rôle particulier dans l'étude de la surjectivité de l'application f .

Corollaire 4.6.2 (*Structure d'un noyau, d'une image*)

Soit $f : (G_1, *) \longrightarrow (G_2, \top)$ un morphisme de groupes, alors

1. $\ker(f)$ est un sous-groupe de G_1 .
2. $\text{Im}(f)$ est un sous-groupe de G_2 .

Preuve. La démonstration est facile et laissée comme un bon exercice d'entraînement pour le lecteur. $\square$

Proposition 4.6.5 (*Caractérisation des morphismes injectifs, ou surjectifs*)

Soient $(G_1, *)$ et $(G_2, \top)$ deux groupes. Soit $f : (G_1, *) \longrightarrow (G_2, \top)$ un morphisme de groupes.

1. L'application f est injective si et seulement si $\ker(f) = \{e_1\}$.
2. L'application f est surjective si et seulement si $\text{Im}(f) = G_2$.

Preuve.

1. • Supposons que f est injective. On a

$$\{e_1\} \subset \ker(f),$$

car $f(e_1) = e_2$. Montrons maintenant que $\ker(f) \subset \{e_1\}$. Soit $x \in \ker(f)$, alors

$$f(x) = e_2 = f(e_1),$$

donc $x = e_1$ puisque f est injective. Ceci prouve que $\ker(f) \subset \{e_1\}$ et par suite

$$\ker(f) = \{e_1\}.$$

- Réciproquement, supposons que $\ker(f) = \{e_1\}$. Soient $x_1, x_2 \in G_1$ tel que $f(x_1) = f(x_2)$, alors on a

$$f(x_1) \top (f(x_2))' = f(x_2) \top (f(x_2))'.$$

Ce qui entraîne que

$$f(x_1) \top (f(x_2))' = f(x_1) \top f(x_2') = e_2.$$

On obtient

$$f(x_1 * x_2') = e_2.$$

Donc

$$x_1 * x_2' \in \ker(f),$$

et comme $\ker(f) = \{e_1\}$, alors on a

$$x_1 * x_2' = e_1,$$

par suite

$$x_1 * x_2' * x_2 = e_1 * x_2.$$

Ce qui entraîne que

$$x_1 = x_2.$$

Par conséquent, f est injective.

2. • Supposons que f est surjective. Par définition, on a

$$\text{Im}(f) \subset G_2.$$

Il reste donc à montrer que

$$G_2 \subset \text{Im}(f),$$

pour établir l'égalité. Soit $y \in G_2$, puisque f est surjective alors il existe $x \in G_1$ tel que $y = f(x)$. Ce qui entraîne que

$$y \in \text{Im}(f).$$

Ceci prouve donc

$$G_2 \subset \text{Im}(f),$$

et par suite

$$\text{Im}(f) = G_2.$$

• Réciproquement, supposons que $\text{Im}(f) = G_2$. Soit $y \in G_2 = \text{Im}(f)$, alors il existe $x \in G_1$ tel que $y = f(x)$, ainsi f est surjective. $\square$

Dans le cas de groupes finis de même cardinal, l'injectivité devient équivalente à la bijectivité. La réduction du noyau du morphisme à l'élément neutre devient alors caractéristique de la bijectivité du morphisme.

Corollaire 4.6.3 *Soient $(G_1, *)$ et $(G_2, \top)$ deux groupes finis de même cardinal. Soit $f : (G_1, *) \longrightarrow (G_2, \top)$ un morphisme de groupes. Alors f est un isomorphisme si et seulement si $\ker(f) = \{e_{G_1}\}$.*

4.7 Anneaux

Les anneaux sont une structure algébrique fondamentale en mathématiques, en particulier en algèbre abstraite. Ils généralisent des concepts tels que les entiers, les polynômes et les matrices, et trouvent des applications dans diverses branches des mathématiques. La structure d'anneau enrichit celle des groupes en introduisant non pas une, mais deux lois de composition interne : l'une est l'addition, notée $+$, et l'autre est la multiplication, notée $\cdot$. Examinons maintenant la définition de la structure d'anneau.

4.7.1 Définitions et exemples

Définition 4.7.1

1. On appelle anneau tout triplet $(A, +, \cdot)$ constitué d'un ensemble A non vide et de deux lois internes sur A : l'addition notée

$$\begin{aligned} + : A \times A &\longrightarrow A \\ (x, y) &\longmapsto x + y, \end{aligned}$$

et la multiplication notée

$$\begin{aligned} \cdot : A \times A &\longrightarrow A \\ (x, y) &\longmapsto x \cdot y, \end{aligned}$$

vérifiant les propriétés suivantes :

(A₁) Le couple $(A, +)$ est un groupe abélien, dont l'élément neutre est noté 0_A , et l'inverse additif (l'opposé) de $x \in A$ est noté $(-x)$, c'est donc dire que

$$x + (-x) = (-x) + x = 0_A.$$

(A₂) La loi $\cdot$ est associative, c'est-à-dire

$$\forall x, y, z \in A : (x \cdot y) \cdot z = x \cdot (y \cdot z).$$

(A₃) La multiplication $\cdot$ est distributive à gauche et à droite par rapport à l'addition $+$, c'est-à-dire

$$\forall x, y, z \in A : \begin{cases} x \cdot (y + z) = (x \cdot y) + (x \cdot z) & (\text{distributivité à gauche}), \\ \text{et} \\ (y + z) \cdot x = (y \cdot x) + (z \cdot x) & (\text{distributivité à droite}). \end{cases}$$

2. Un anneau $(A, +, \cdot)$ est dit unitaire (ou unifère) si la multiplication $\cdot$ admet un élément neutre (appelé un élément unité de A et noté 1_A). Il est dit commutatif si sa multiplication est commutative,

$$\forall x, y \in A : x \cdot y = y \cdot x.$$

3. Si $A = \{0_A\}$, alors $(A, +, \cdot)$ est un anneau appelé anneau nul.

Remarque 4.7.1 De la même façon que pour les groupes, le neutre de $+$ est unique, les symétriques pour $+$ sont uniques, le neutre pour $\cdot$ s'il existe est unique, et les symétriques pour $\cdot$, lorsqu'ils existent, sont uniques.

Exemple 4.7.1 (Exemples fondamentaux d'anneaux)

1. Les structures $(\mathbb{Z}, +, \cdot)$, $(\mathbb{Q}, +, \cdot)$, $(\mathbb{R}, +, \cdot)$, $(\mathbb{C}, +, \cdot)$ sont des anneaux unitaires, commutatifs.
2. L'ensemble $(\mathcal{P}(E), \Delta, \cap)$ des parties d'un ensemble E muni de la différence symétrique Δ et de l'intersection $\cap$ est un anneau commutatif et unitaire dont $1_{\mathcal{P}(E)} = E$.
3. Soit $(\mathbb{R}^{\mathbb{N}}, +, \cdot)$ l'ensemble des suites réelles muni des lois $+$ et $\cdot$ définies par : si $u = (u_n)_n$, $v = (v_n)_n$ alors

$$u + v = (u_n + v_n)_n, u \cdot v = (u_n \cdot v_n)_n.$$

Alors $(\mathbb{R}^{\mathbb{N}}, +, \cdot)$ est un anneau commutatif et unitaire ($1_{\mathbb{R}^{\mathbb{N}}}$ est la suite réelle constante égale à 1).

3. Soit I un intervalle de $\mathbb{R}$. On munit l'ensemble $(\mathcal{F}(I, \mathbb{R}), +, \cdot)$ des applications de I dans $\mathbb{R}$ des lois d'addition $+$ et de multiplication $\cdot$ suivantes : Soit $f, g \in \mathcal{F}(I, \mathbb{R})$,
L'addition $f + g$ est définie par

$$\forall x \in I : (f + g)(x) = f(x) + g(x).$$

La multiplication $f \cdot g$ est définie par

$$\forall x \in I : (f \cdot g)(x) = f(x) \cdot g(x).$$

Alors $(\mathcal{F}(I, \mathbb{R}), +, \cdot)$ est un anneau commutatif et unitaire (de neutre multiplicatif l'application constante égale à 1).

4. $(\mathbb{R}[X], +, \cdot)$ l'ensemble des polynômes à coefficients réels, est un anneau commutatif unitaire (neutre multiplicatif le polynôme constant égal à 1).

Exemple 4.7.2 Parmi les anneaux commutatifs finis, un exemple incontournable est certainement l'anneau des entiers modulo n : noté $\mathbb{Z}/n\mathbb{Z}$. Fixons $n \in \mathbb{N}^*$, rappelons que $\mathbb{Z}/n\mathbb{Z}$ est l'ensemble des classes d'équivalence pour la congruence modulo n défini par

$$\mathbb{Z}/n\mathbb{Z} = \left\{ \dot{0}, \dot{1}, 2, \dots, \widehat{n-1} \right\},$$

où $\dot{x}$ désigne la classe d'équivalence de $x \in \mathbb{Z}$ pour cette relation. On note $+$ et $\cdot$ l'addition et la multiplication usuelles sur $\mathbb{Z}$. On vérifie facilement que pour tout $n \in \mathbb{N}^*$ et pour tous x, y, x', y' appartenant à $\mathbb{Z}$, on a

$$\begin{cases} x \equiv x' \pmod{n} \\ y \equiv y' \pmod{n} \end{cases} \implies \begin{cases} x + y \equiv x' + y' \pmod{n} \\ x \cdot y \equiv x' \cdot y' \pmod{n}. \end{cases}$$

Cela signifie que si $x, y \in \mathbb{Z}$ alors les classes d'équivalence $\widehat{x+y}$ et $\widehat{x \cdot y}$ ne dépendent que des classes d'équivalence $\dot{x}$ et $\dot{y}$ et non du choix de x dans la classe d'équivalence $\dot{x}$ et de y dans la classe d'équivalence $\dot{y}$. Cela nous permet de définir deux lois de composition interne, notées $\oplus$ et $\odot$ et appelées respectivement addition et multiplication sur l'ensemble quotient $\mathbb{Z}/n\mathbb{Z}$, en posant pour tous $\dot{x}, \dot{y} \in \mathbb{Z}/n\mathbb{Z}$,

$$\dot{x} \oplus \dot{y} = \widehat{x+y} \text{ et } \dot{x} \odot \dot{y} = \widehat{x \cdot y}.$$

On montre sans difficulté que pour tout $n \in \mathbb{N}^*$, l'ensemble quotient $(\mathbb{Z}/n\mathbb{Z}, \oplus, \odot)$ muni des deux lois $\oplus$ et $\odot$ possède une structure d'anneau commutatif (unitaire si $n \geq 2$).

Pour tout $n \in \mathbb{N}^*$, on a $0_{\mathbb{Z}/n\mathbb{Z}}$ est la classe d'équivalence $\dot{0}$ et $1_{\mathbb{Z}/n\mathbb{Z}}$ est la classe d'équivalence $\dot{1}$ avec $n \geq 2$. Car pour tout $\dot{x} \in \mathbb{Z}/n\mathbb{Z}$,

$$\begin{cases} \dot{x} \oplus \dot{0} = \widehat{x+0} = \dot{x} = \dot{0} \oplus \dot{x}, \\ \dot{x} \odot \dot{1} = \widehat{x \cdot 1} = \dot{x} = \dot{1} \odot \dot{x}. \end{cases}$$

Ecrivons par exemple les tables d'addition et de multiplication dans les deux ensembles quotients $\mathbb{Z}/2\mathbb{Z}$ et $\mathbb{Z}/3\mathbb{Z}$.

► Dans $\mathbb{Z}/2\mathbb{Z} = \{\dot{0}, \dot{1}\}$, elles s'écrivent

$\oplus$	$\dot{0}$	$\dot{1}$
$\dot{0}$	$\dot{0}$	$\dot{1}$
$\dot{1}$	$\dot{1}$	$\dot{0}$

 et

$\odot$	$\dot{0}$	$\dot{1}$
$\dot{0}$	$\dot{0}$	$\dot{0}$
$\dot{1}$	$\dot{0}$	$\dot{1}$

► Dans $\mathbb{Z}/3\mathbb{Z} = \{\dot{0}, \dot{1}, \dot{2}\}$, elles s'écrivent

$\oplus$	$\dot{0}$	$\dot{1}$	$\dot{2}$
$\dot{0}$	$\dot{0}$	$\dot{1}$	$\dot{2}$
$\dot{1}$	$\dot{1}$	$\dot{2}$	$\dot{0}$
$\dot{2}$	$\dot{2}$	$\dot{0}$	$\dot{1}$

 et

$\odot$	$\dot{0}$	$\dot{1}$	$\dot{2}$
$\dot{0}$	$\dot{0}$	$\dot{0}$	$\dot{0}$
$\dot{1}$	$\dot{0}$	$\dot{1}$	$\dot{2}$
$\dot{2}$	$\dot{0}$	$\dot{2}$	$\dot{1}$

Définition 4.7.2 (Anneau produit)

1. Soient $(A_1, +, \cdot), (A_2, +, \cdot), \dots, (A_n, +, \cdot)$ une famille d'anneaux. On munit l'ensemble produit $A = A_1 \times A_2 \times \dots \times A_n$ de deux lois internes, notée $+$ et $\cdot$ définies respectivement par

$$\begin{aligned} (x_1, \dots, x_n) + (y_1, \dots, y_n) &= (x_1 + y_1, \dots, x_n + y_n), \\ (x_1, \dots, x_n) \cdot (y_1, \dots, y_n) &= (x_1 \cdot y_1, \dots, x_n \cdot y_n). \end{aligned}$$

On a donc $(A, +, \cdot)$ est un anneau de neutres $0_A = (0_{A_1}, \dots, 0_{A_n})$ appelé anneau produit des anneaux $A_1, \dots, A_n$.

2. Si chaque A_i est un anneau unitaire, alors $A_1 \times A_2 \times \dots \times A_n$ est un anneau unitaire avec l'élément neutre pour la multiplication est $1_A = (1_{A_1}, \dots, 1_{A_n})$

3. De plus, si chaque A_i est un anneau commutatif, alors $A_1 \times A_2 \times \dots \times A_n$ est également un anneau commutatif.

Exemple 4.7.3 • Soit l'ensemble

$$\mathbb{Z}^2 = \mathbb{Z} \times \mathbb{Z} = \{(x, y) : x \in \mathbb{Z} \text{ et } y \in \mathbb{Z}\},$$

munie des lois d'addition et de multiplication terme à terme : pour tous $(x_1, y_1), (x_2, y_2)$ dans $\mathbb{Z}^2$

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

$$(x_1, y_1) \cdot (x_2, y_2) = (x_1 \cdot x_2, \dots, y_1 \cdot y_2).$$

On peut facilement vérifier que $(\mathbb{Z}^2, +, \cdot)$ est un anneau commutatif unitaire ($0_{\mathbb{Z}^2} = (0, 0)$ et $1_{\mathbb{Z}^2} = (1, 1)$).

• **Produit de trois anneaux d'entiers modulo** : Considérons $\mathbb{Z}/2\mathbb{Z}$, $\mathbb{Z}/3\mathbb{Z}$, et $\mathbb{Z}/5\mathbb{Z}$. Leur produit $\mathbb{Z}/2\mathbb{Z} \times \mathbb{Z}/3\mathbb{Z} \times \mathbb{Z}/5\mathbb{Z}$ est l'ensemble des triplets $(\dot{x}, \dot{y}, \dot{z})$ où $\dot{x} \in \mathbb{Z}/2\mathbb{Z}$, $\dot{y} \in \mathbb{Z}/3\mathbb{Z}$, et $\dot{z} \in \mathbb{Z}/5\mathbb{Z}$.

Définition 4.7.3 (Élément inversible dans un anneau)

Soit $(A, +, \cdot)$ un anneau unitaire et $a \in A$.

1. On dit que a est inversible pour la loi $\cdot$ s'il existe b de A tel que

$$a \cdot b = b \cdot a = 1_A.$$

Alors, cet b élément est unique et s'appelle l'inverse de a . On le note a^{-1} .

2. L'ensemble des éléments inversibles de A est noté $A^\times$. Cet ensemble, également appelé le groupe des unités de A , est formé par tous les éléments de A qui possèdent un inverse multiplicatif dans A .

Remarque 4.7.2 On ne confondra pas les ensembles A^* et $A^\times$ respectivement ensemble des éléments non nuls de A et ensemble des unités de A .

Proposition 4.7.1 (Groupe des éléments inversibles d'un anneau)

Soit $(A, +, \cdot)$ un anneau unitaire. Alors l'ensemble $A^\times$ des éléments inversibles de A est un groupe pour cette même loi $\cdot$, appelé groupe des inversibles ou groupe des unités de A .

Preuve.

• La loi $\cdot$ interne dans $A^\times$:

Il est clair que $A^\times$ est stable par loi $\cdot$ car si $a, b \in A^\times$, alors

$$(a \cdot b)^{-1} = b^{-1} \cdot a^{-1} \in A^\times.$$

• Existence d'un élément Neutre : L'élément neutre pour la multiplication dans A est 1_A . Par définition, 1_A est dans $A^\times$ car $1_A \cdot 1_A = 1_A$. Donc, 1_A est l'élément neutre dans $A^\times$.

• Existence d'inverses pour chaque élément : Si $a \in A^\times$, alors $a^{-1} \in A^\times$ car $(a^{-1})^{-1} = a$.

• Associativité : La multiplication $\cdot$ est associative dans A , donc elle est également associative dans $A^\times$. Ainsi $(A^\times, \cdot)$ forme un groupe. $\square$

Exemple 4.7.4

1. Le groupe des inversibles de l'anneau $(\mathbb{Z}, +, \cdot)$ se réduit à la paire $\{-1, 1\}$.

2. Le groupe des inversibles de l'anneau $(\mathbb{R}, +, \cdot)$ est l'ensemble de tous les réels non nuls $\mathbb{R}^*$.

3. Pour tout ensemble E , le groupe des inversibles de l'anneau $(\mathcal{P}(E), \Delta, \cap)$ est $(\{E\}, \cap)$.

4.7.2 Règles de calculs dans un anneau

Ces règles résultent des propriétés élémentaires des lois de composition interne, ainsi que des axiomes $(A_1), (A_2), (A_3)$. Soit $(A, +, \cdot)$ un anneau. On a les règles fondamentales de calcul suivantes.

Proposition 4.7.2 (Quelques Propriétés Arithmétiques) Tout anneau $(A, +, \cdot)$ vérifie les propriétés suivantes :

1. Pour tout $x \in A$,

$$x \cdot 0_A = 0_A \cdot x = 0_A,$$

on dit que 0_A est absorbant pour la loi $\cdot$.

2. **Règle de signes :**

$$\forall x, y \in A : \begin{cases} (i) - (x \cdot y) = (-x) \cdot y = x \cdot (-y), \\ (ii) (-x) \cdot (-y) = x \cdot y. \end{cases}$$

3. **Règles de distributivité :**

$$\forall x, y \in A : \begin{cases} (i) x \cdot (y - z) = x \cdot y - x \cdot z \text{ (Distributivité à gauche de } \cdot \text{ par rapport à la soustraction)}, \\ (ii) (y - z) \cdot x = y \cdot x - z \cdot x. \text{ (Distributivité à droite de } \cdot \text{ par rapport à la soustraction)}. \end{cases}$$

Preuve.

1. Soit $x \in A$, alors la distributivité de la loi $\cdot$ par rapport à la loi $+$ permet d'écrire

$$x \cdot 0_A = x \cdot (0_A + 0_A) = x \cdot 0_A + x \cdot 0_A.$$

En ajoutant à chaque membre l'opposé de $x \cdot 0_A$, on trouve

$$x \cdot 0_A - x \cdot 0_A = x \cdot 0_A + x \cdot 0_A - x \cdot 0_A,$$

et donc

$$x \cdot 0_A = 0_A.$$

L'autre égalité se prouve de même.

2. (i) Soient $x, y \in A$, alors on a

$$y + (-y) = 0_A.$$

Par suite

$$x \cdot (y + (-y)) = x \cdot y + x \cdot (-y) = 0_A.$$

En ajoutant à chaque membre $-(x \cdot y)$, on obtient

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

Alors $x \cdot (-y)$ est l'opposé de $(x \cdot y)$. Un raisonnement similaire fournit l'égalité

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

Maintenant, montrons le point (ii), soient $x, y \in A$, alors on a

$$(-x) \cdot (-y) \stackrel{(i)}{=} -(x \cdot (-y)) \stackrel{(i)}{=} -(-(x \cdot y)) = x \cdot y.$$

3. (i) Soient $x, y, z \in A$, alors on a

$$y - z = y + (-z).$$

En appliquant la distributivité à gauche par rapport à l'addition, nous obtenons

$$x \cdot (y - z) = x \cdot y + x \cdot (-z).$$

Par définition de l'opposé dans un anneau, $x \cdot (-z) = -(x \cdot z)$. Donc

$$x \cdot (y - z) = x \cdot y - x \cdot z.$$

Un raisonnement similaire fournit l'égalité (ii). □

Notation 4.7.1 Soit $(A, +, \cdot)$ un anneau unitaire non nécessairement commutatif. Alors

1. Pour tout $x \in A$ et $n \in \mathbb{Z}$, on note

$$nx = \begin{cases} 0_A, & \text{si } n = 0. \\ \underbrace{x + x + \dots + x}_{n \text{ fois}}, & \text{si } n > 0 \\ \underbrace{(-x) + (-x) + \dots + (-x)}_{-n \text{ fois}} = (-n)(-x), & \text{si } n < 0. \end{cases}$$

Par exemple, si $n = -3$, alors

$$-3 \cdot x = (-x) + (-x) + (-x) = 3 \cdot (-x).$$

2. Pour tout $x \in A$ et $n \in \mathbb{N}$, on note

$$x^n = \begin{cases} 1_A, & \text{si } n = 0. \\ \underbrace{x \cdot x \cdot \dots \cdot x}_{n \text{ fois}}, & \text{si } n > 0 \end{cases}$$

Lorsque x est inversible dans A , on peut définir x^n avec $n < 0$ en posant

$$x^n = x^{-1} \cdot x^{-1} \cdot \dots \cdot x^{-1} = (x^{-1})^{-n} \text{ (multiplié } (-n) \text{ fois)}.$$

Par exemple, si $n = -3$ et si x est inversible, alors

$$x^{-3} = x^{-1} \cdot x^{-1} \cdot x^{-1} = (x^{-1})^3.$$

Le lecteur a déjà rencontré le binôme de Newton dans $\mathbb{R}$ ou $\mathbb{C}$ ainsi que dans le calcul matriciel. C'est en fait un résultat général propre à la théorie des anneaux.

Proposition 4.7.3 (Formule du binôme de Newton).

Soit $(A, +, \cdot)$ un anneau unitaire, et soit x, y deux éléments de A tel que $x \cdot y = y \cdot x$ (x et y commutent), alors pour tout $n \in \mathbb{N}$, on a

$$(x + y)^n = \sum_{k=0}^n C_n^k x^k \cdot y^{n-k}. \quad (4.1)$$

En particulier

$$(x + y)^2 = x^2 + 2x \cdot y + y^2.$$

La relation (4.1) est appelée formule du binôme de Newton.

Preuve. Notons $P(n)$ la relation (4.1) écrite pour x et y donnés dans A et n fixé dans $\mathbb{N}$. La formule $P(n)$ se démontre par récurrence sur n . $\square$

Corollaire 4.7.1 Comme 1_A et $x \in A$ commutent, alors

$$(1_A + x)^n = \sum_{k=0}^n C_n^k x^k.$$

Remarque 4.7.3 (Sur l'importance de l'hypothèse $x \cdot y = y \cdot x$)

Dans la proposition 4.7.3 précédente, l'hypothèse selon laquelle x et y commutent est essentielle. Évidemment, cela est automatiquement vérifié dans un anneau commutatif. Si les deux éléments x et y ne commutent pas, le développement de $(x + y)^n$ est beaucoup plus compliqué. On trouve par exemple

$$\begin{cases} (x + y)^2 = x^2 + x \cdot y + y \cdot x + y^2 \neq x^2 + 2x \cdot y + y^2, \\ (x + y)^3 = x^3 + x^2 \cdot y + x \cdot y \cdot x + x \cdot y^2 + y \cdot x^2 + y \cdot x \cdot y + y^2 \cdot x + y^3. \end{cases}$$

Corollaire 4.7.2 Soit $(A, +, \cdot)$ un anneau unitaire, alors on a les propriétés suivantes :

1. Pour tout $x \in A$ et pour tous $n, m \in \mathbb{Z}$, on a

$$\begin{cases} (n + m)x = nx + mx \\ (nm)x = n(mx). \end{cases}$$

2. Pour tout $x \in A$ et pour tous $n, m \in \mathbb{N}$, on a

$$\begin{cases} x^n \cdot x^m = x^{n+m} \\ (x^n)^m = x^{nm}. \end{cases}$$

En particulier, si x est inversible dans A , alors les deux dernières propriétés sont vraies pour tout $(n, m) \in \mathbb{Z}^2$.

3. Pour tout $x \in A$, et pour tout $n \in \mathbb{Z}$, on a

$$n(-x) = (-n)x = -(nx).$$

4. Pour tous $x, y \in A$ et pour tout $n \in \mathbb{Z}$, on a

$$\begin{cases} n(x + y) = nx + ny \\ n(x - y) = nx - ny. \end{cases}$$

5. Pour tous $x, y \in A$ et pour tout $n \in \mathbb{Z}$, on a

$$n(x \cdot y) = (nx) \cdot y = x \cdot (ny).$$

Remarque 4.7.4 Dans un anneau A non commutatif, $(x \cdot y)^n \neq x^n \cdot y^n$ en général. En effet $(x \cdot y)^n$ se comprend

$$\underbrace{(x \cdot y) \cdot \dots \cdot (x \cdot y)}_{n \text{ fois}}.$$

4.7.3 Sous-anneaux d'un anneau

Nous allons examiner deux sous-structures importantes des anneaux : les sous-anneaux et les idéaux. Ces deux concepts jouent un rôle crucial lorsque l'on souhaite effectuer des opérations arithmétiques au sein d'un anneau. De manière analogue à la définition des sous-groupes en théorie des groupes, la notion de sous-anneaux est définie dans le cadre de la théorie des anneaux.

Définition 4.7.4 (*Sous-anneau*)

Un sous-anneau B d'un anneau $(A, +, \cdot)$ est un sous-ensemble de A qui est lui-même un anneau avec les mêmes opérations d'addition et de multiplication que A .

On donne maintenant un critère couramment utilisé pour vérifier qu'un sous-ensemble est un sous-anneau.

Proposition 4.7.4 (*Caractérisation d'un sous-anneau*)

Soit $(A, +, \cdot)$ un anneau et B une partie non vide de A . Alors

$$\begin{aligned} B \text{ est un sous-anneau de } A &\iff \begin{cases} j) (B, +) \text{ est un sous-groupe du groupe additif } (A, +) \\ jj) B \text{ est stable par la multiplication de } A : \forall x, y \in B : x \cdot y \in B. \end{cases} \\ &\iff \begin{cases} SA_1) 0_A \in B \\ SA_2) \forall x, y \in B : x - y \in B. \\ SA_3) \forall x, y \in B : x \cdot y \in B. \end{cases} \end{aligned}$$

La caractérisation d'un sous-anneau en termes de conditions sur l'addition et la multiplication simplifie considérablement le processus de vérification. En utilisant ces critères, on peut rapidement établir si un sous-ensemble est un sous-anneau d'un anneau donné.

Exemple 4.7.5

1. Si $(A, +, \cdot)$ est un anneau non nul, alors $\{0_A\}$ et A lui-même sont des sous-anneaux (triviaux) de A .
2. Dans $(\mathbb{Z}, +, \cdot)$, $(\mathbb{Q}, +, \cdot)$, $(\mathbb{R}, +, \cdot)$, $(\mathbb{C}, +, \cdot)$ chacun est un sous-anneau du suivant.
3. Pour tout entier $n \in \mathbb{N}$, l'ensemble $B = n\mathbb{Z}$ des multiples de n est un sous-anneau de $(\mathbb{Z}, +, \cdot)$ et tout sous-anneau de $\mathbb{Z}$ s'écrit sous cette forme, de façon unique.
4. On note $\mathbb{Z}[i]$ l'ensemble

$$\mathbb{Z}[i] = \{z = a + ib : a, b \in \mathbb{Z}\},$$

où i est l'unité imaginaire, $\mathbb{Z}[i]$ est un sous-anneau de $(\mathbb{C}, +, \cdot)$ qu'on l'appelle le sous-anneau des entiers de **Gauss**.

5. Soit I un intervalle de $\mathbb{R}$, l'ensemble $\mathcal{C}(I, \mathbb{R})$ des fonctions continues sur I , est un sous-anneau de $(\mathcal{F}(I, \mathbb{R}), +, \cdot)$.
6. L'ensemble des suites bornées (resp. convergentes) est un sous-anneau de $\mathbb{R}^{\mathbb{N}}$ des suites de réels.
7. L'ensemble des classes d'équivalence $B = \{0, 2, 4\}$ est un sous-anneau de $(\mathbb{Z}/6\mathbb{Z}, \oplus, \odot)$.

Exemple 4.7.6 Soit $(A, +, \cdot)$ un anneau et S un sous-ensemble de A . Le commutant (ou centralisateur) de S dans A , noté $C_A(S)$ est défini par

$$C_A(S) = \{a \in A, \forall s \in S : a \cdot s = s \cdot a\}.$$

Autrement dit, $C_A(S)$ est l'ensemble de tous les éléments de A qui commutent avec chaque élément de S . On a donc $C_A(S)$ est un sous-anneau de A . Si $S = A$, le commutant est appelé centre de A souvent noté $Z(A)$.

Remarque 4.7.5

1. Etant donné un anneau unitaire A , nous conviendrons d'appeler sous-anneau unitaire de A , tout sous-anneau de A contenant 1_A .
2. Si B est un sous-anneau unitaire d'un anneau unitaire A , alors B est lui-même un anneau unitaire, et l'on a $1_B = 1_A$.
3. Si l'anneau A est commutatif, alors tout sous-anneau de A est commutatif.
4. Un sous-anneau B d'un anneau unitaire A ne contient pas nécessairement l'élément unité 1_A . Par exemple, pour tout entier $n > 1$, $n\mathbb{Z}$ est un sous-anneau de $\mathbb{Z}$ ne contenant pas 1.

Proposition 4.7.5 Dans tout anneau $(A, +, \cdot)$, l'intersection $\bigcap_{i \in I} A_i$ d'une famille quelconque de sous-anneaux $(A_i)_{i \in I}$ de A est un sous-anneau de A .

Preuve. Même démonstration que dans le cas des groupes. En effet, soit $(A_i)_{i \in I}$ une famille de sous-anneaux de A . Puisque 0_A est un élément de chacun des A_i (puisque les A_i sont des sous-anneaux), il est aussi un élément de $C := \bigcap_{i \in I} A_i$. Pour x et y dans C , par définition, x et y sont dans chacun des A_i , et donc on a que $x - y$ et $x \cdot y$ sont dans chacun des A_i . Il s'ensuit que $x - y$ et $x \cdot y$ sont dans C , ce qui montre que C est bien un sous-anneau. $\square$

Remarque 4.7.6 En général, la réunion d'une famille de sous-anneaux de A n'est pas un sous-anneau de A . Par exemple, dans l'anneau $(\mathbb{Z}, +, \cdot)$, la partie $3\mathbb{Z} \cup 5\mathbb{Z}$ n'est pas un sous-anneau de $\mathbb{Z}$, puisque ce n'est pas un sous-groupe de $(\mathbb{Z}, +)$.

4.7.4 Diviseurs de zéro et Intégrité

La notion d'intégrité est liée à celle de diviseurs de zéros que nous introduisons ci-dessous. On prendra garde à ce que l'équation $a \cdot x = 0_A$ dans un anneau $(A, +, \cdot)$ peut admettre d'autres solutions que $x = 0_A$ même lorsque $a \in A - \{0_A\}$. Autrement dit, on ne peut pas simplifier par a en général.

Définition 4.7.5 (*Diviseur de Zéro*)

Soit $(A, +, \cdot)$ un anneau non réduit à $\{0_A\}$. On dit qu'un élément :

1. $a \in A$ est un diviseur de zéro à gauche si $a \neq 0_A$ et s'il existe un élément non nul $b \in A - \{0_A\}$ tel que $a \cdot b = 0_A$.
 2. $a \in A$ est un diviseur de zéro à droite si $a \neq 0_A$ et s'il existe un élément non nul $b \in A - \{0_A\}$ tel que $b \cdot a = 0_A$.
 3. $a \in A$ est un diviseur de zéro si et si a est diviseur de zéro à droite et à gauche.
- Autrement dit, un diviseur de zéro est un élément qui annule d'autres éléments non nuls lorsqu'il est multiplié par eux.

Remarque 4.7.7

1. Par cette définition, on ne considère pas que 0_A soit un diviseur de zéro.
2. Si A est commutatif, les notions de diviseur de zéro à gauche et à droite coïncident. Cette propriété simplifie les preuves concernant les diviseurs de zéro dans les anneaux commutatifs.
3. Les éléments inversibles de A ne sont pas diviseurs de zéro.

Exemple 4.7.7

1. Dans l'anneau (anneau produit) $(\mathbb{R}^2, +, \cdot)$ les diviseurs de zéros sont les points $(x, 0)$ et $(0, x)$ avec $x \neq 0$.
2. $(\mathbb{Z}, +, \cdot), (\mathbb{Q}, +, \cdot), (\mathbb{R}, +, \cdot), (\mathbb{C}, +, \cdot)$ n'ont pas de diviseurs de zéro.
3. Dans l'anneau $(\mathcal{F}(\mathbb{R}, \mathbb{R}), +, \cdot)$ muni des deux lois $+$ et $\cdot$ définies à l'exemple 4.7.1, les deux applications

$$\begin{aligned} f: \mathbb{R} &\longrightarrow \mathbb{R} & g: \mathbb{R} &\longrightarrow \mathbb{R} \\ x &\longmapsto f(x) = \begin{cases} 0, & \text{si } x \leq 0 \\ \sqrt{x}, & \text{si } x > 0, \end{cases} & \text{et} & x \longmapsto g(x) = \begin{cases} x, & \text{si } x \leq 0 \\ 0, & \text{si } x > 0, \end{cases} \end{aligned}$$

sont des diviseurs de zéro dans $\mathcal{F}(\mathbb{R}, \mathbb{R})$ puisqu'elles sont non nulles et vérifient $f \cdot g = 0$. En effet :

$$\forall x \in \mathbb{R} : (f \cdot g)(x) = \begin{cases} 0 \cdot x = 0, & \text{si } x \leq 0, \\ \sqrt{x} \cdot 0 = 0, & \text{si } x > 0. \end{cases}$$

4. Dans l'anneau $(\mathbb{Z}/6\mathbb{Z}, \oplus, \odot)$ tel que

$$\mathbb{Z}/6\mathbb{Z} = \{\dot{0}, \dot{1}, \dot{2}, \dot{3}, \dot{4}, \dot{5}\},$$

les classes $\dot{2}, \dot{3}$ sont des diviseurs de zéro dans $\mathbb{Z}/6\mathbb{Z}$. En effet :

$$\begin{cases} \dot{2} \neq \dot{0} \text{ et } \exists \dot{3} \neq \dot{0} : \dot{2} \odot \dot{3} = \dot{6} = \dot{0}, \\ \dot{3} \neq \dot{0} \text{ et } \exists \dot{2} \neq \dot{0} : \dot{3} \odot \dot{2} = \dot{6} = \dot{0}. \end{cases}$$

Proposition 4.7.6 Soit $(A, +, \cdot)$ un anneau unitaire. Tout diviseur de zéro dans A est nécessairement non inversible.

Preuve. Cela se montre facilement en raisonnant par l'absurde. Supposons que l'élément x de A soit à la fois diviseur à gauche et inversible. Il existe alors un élément $y \in A$ non nul ($y \neq 0_A$) tel que $x \cdot y = 0_A$. En multipliant à gauche cette égalité par x^{-1} (qui existe car x est inversible), on obtient

$$x^{-1} \cdot (x \cdot y) = x^{-1} \cdot 0_A.$$

On en déduit $y = 0_A$ puisque

$$x^{-1} \cdot (x \cdot y) = (x^{-1} \cdot x) \cdot y = 1_A \cdot y = y.$$

et puisque $x^{-1} \cdot 0_A = 0_A$ (0_A est absorbant pour la loi $\cdot$), ce qui est contraire à nos hypothèses. La démonstration dans le cas où x est diviseur de zéro à droite s'effectue suivant le même principe. $\square$

Une propriété intéressante de certains anneaux est l'absence de diviseur de zéro : le produit de deux éléments de l'anneau est nul si et si l'un des deux éléments est nul. Ceci est le cas par exemple dans les ensembles de nombres munis des lois d'addition et de multiplication usuelles. Mais soyez attentifs, cette propriété n'est pas forcément vraie dans toutes les structures algébriques. Lorsqu'un anneau satisfait cette propriété, on dit que c'est un anneau intègre.

Définition 4.7.6 On dit qu'un anneau $(A, +, \cdot)$ est intègre ou est un anneau d'intégrité s'il est non nul et s'il ne possède pas de diviseurs de zéro. En d'autres termes, un anneau $(A, +, \cdot)$ est intègre si $A \neq \{0_A\}$ et si

$$\forall x, y \in A : x \cdot y = 0_A \implies \begin{cases} x = 0_A \\ \text{ou} \\ y = 0_A. \end{cases}$$

Autrement dit, par contraposition, dans un anneau intègre, si $x \neq 0_A$ et si $y \neq 0_A$, alors $x \cdot y \neq 0_A$ aussi.

Exemple 4.7.8

1. $(\mathbb{Z}, +, \cdot), (\mathbb{Q}, +, \cdot), (\mathbb{R}, +, \cdot), (\mathbb{C}, +, \cdot)$ sont des anneaux intègres.
2. $(\mathbb{Z}/6\mathbb{Z}, \oplus, \odot)$ n'est pas un anneau intègre.
3. L'anneau produit $(\mathbb{Z}^2, +, \cdot)$ de l'exemple 4.7.3 n'est pas intègre, car il contient des diviseurs de zéro $(1, 0)$ et $(0, 1)$.
4. Les deux anneaux $(\mathbb{Z}/2\mathbb{Z}, \oplus, \odot)$ et $(\mathbb{Z}/3\mathbb{Z}, \oplus, \odot)$ sont intègres. Cela se vérifie directement à partir des tables d'addition et de multiplication dans $\mathbb{Z}/2\mathbb{Z}$ et $\mathbb{Z}/3\mathbb{Z}$ données dans l'exemple 4.7.1.
5. L'anneau $(\mathbb{Z}/n\mathbb{Z}, \oplus, \odot)$ des congruences modulo n est intègre si et seulement si n est premier.

Remarque 4.7.8 Un anneau intègre n'a pas nécessairement besoin d'être commutatif. Les anneaux intègres peuvent être commutatifs ou non commutatifs. Toutefois, un anneau commutatif intègre est souvent appelé simplement un anneau intègre.

4.7.5 Morphismes d'anneaux

Les morphismes d'anneaux sont des applications entre deux anneaux qui préservent les structures algébriques des anneaux. En d'autres termes, un morphisme d'anneaux respecte les opérations de somme et de produit des anneaux. Ils jouent un rôle crucial dans la théorie des anneaux et sont analogues aux morphismes de groupes dans la théorie des groupes.

Définition 4.7.7 Soient $(A, +, \cdot), (B, +, \cdot)$ deux anneaux. On dit qu'une application $f : (A, +, \cdot) \longrightarrow (B, +, \cdot)$ est un morphisme d'anneaux de $(A, +, \cdot)$ dans $(B, +, \cdot)$ si les assertions suivantes sont vérifiées :

i) **Préservation de l'Addition :**

$$\forall x, y \in A : f(x + y) = f(x) + f(y).$$

ii) **Préservation de la Multiplication :**

$$\forall x, y \in A : f(x \cdot y) = f(x) \cdot f(y).$$

Remarque 4.7.9

1. Par définition, un morphisme d'anneaux $f : (A, +, \cdot) \longrightarrow (B, +, \cdot)$ est un morphisme du groupe additif $(A, +)$ dans le groupe additif $(B, +)$. Ainsi, f possède, en particulier toutes les propriétés d'un morphisme de groupes.
2. Si $(A, +, \cdot)$ et $(B, +, \cdot)$ sont deux anneaux unitaires, et $f : (A, +, \cdot) \longrightarrow (B, +, \cdot)$ un morphisme d'anneaux n'implique pas nécessairement $f(1_A) = 1_B$.

Les notions d'endomorphisme, d'isomorphisme, et d'automorphisme d'anneaux sont analogues à celles des groupes, avec quelques détails spécifiques à la structure des anneaux. Voici une définition de chaque terme dans le contexte des anneaux :

Définition 4.7.8 Soient $(A, +, \cdot), (B, +, \cdot)$ deux anneaux. Soit $f : (A, +, \cdot) \longrightarrow (B, +, \cdot)$ un morphisme d'anneaux.

1. Si $(A, +, \cdot) = (B, +, \cdot)$, alors on dit f est endomorphisme d'anneaux.
2. Si f est une application bijective, alors on dit que f est un isomorphisme d'anneaux.
3. Si $(A, +, \cdot) = (B, +, \cdot)$ et f est une application bijective, alors on dit que f est un automorphisme d'anneaux.

Définition 4.7.9 Deux anneaux $(A, +, \cdot)$ et $(B, +, \cdot)$ sont dits isomorphes, s'il existe un isomorphisme d'anneaux de l'un sur l'autre ; dans ce cas, on écrit : $A \approx B$.

Exemple 4.7.9

1. Soit $a \in \mathbb{Z}$ et définissons $f : \mathbb{Z} \longrightarrow \mathbb{Z}$ en posant $f(x) = ax$. Alors f est un morphismes de groupe $(\mathbb{Z}, +)$, car pour tout $x, y \in \mathbb{Z}$, on a

$$f(x + y) = a(x + y) = ax + ay = f(x) + f(y).$$

Cependant, si $a \neq 1$, ce n'est pas un morphisme d'anneaux car, en général,

$$f(x \cdot y) \neq f(x) \cdot f(y).$$

2. L'application

$$\begin{aligned} f : (\mathbb{C}, +, \cdot) &\longrightarrow (\mathbb{C}, +, \cdot) \\ z = x + iy &\longmapsto f(z) = \bar{z} = x - iy, \end{aligned}$$

est un morphisme d'anneaux. C'est même un automorphisme de l'anneau $(\mathbb{C}, +, \cdot)$. En effet : On sait que pour tout $z_1, z_2 \in \mathbb{C}$, on a

$$\overline{z_1 + z_2} = \bar{z}_1 + \bar{z}_2,$$

et

$$\overline{z_1 \cdot z_2} = \overline{z_1} \cdot \overline{z_2}.$$

Ainsi,

$$f(z_1 + z_2) = f(z_1) + f(z_2), f(z_1 \cdot z_2) = f(z_1) \cdot f(z_2).$$

Par ailleurs, f est une application bijective. Donc, f est un automorphisme d'anneau.

3. Si $(A, +, \cdot), (B, +, \cdot)$ sont des anneaux et $A \times B$ l'anneau produit, alors les projections (sur les composantes)

$$\begin{aligned} p_1 : A \times B &\longrightarrow A \\ (x, y) &\longmapsto p_1(x, y) = x, \end{aligned} \quad \text{et} \quad \begin{aligned} p_2 : A \times B &\longrightarrow B \\ (x, y) &\longmapsto p_2(x, y) = y, \end{aligned}$$

sont des morphismes d'anneaux.

4. Pour $n > 1$ dans $\mathbb{N}$, la surjection canonique

$$\begin{aligned} s : (\mathbb{Z}, +, \cdot) &\longrightarrow (\mathbb{Z}/n\mathbb{Z}, \oplus, \odot) \\ x &\longmapsto s(x) = \dot{x}, \end{aligned}$$

est un morphisme d'anneaux. En effet, soient $x, y \in \mathbb{Z}$, alors

$$s(x + y) = \widehat{\dot{x} + \dot{y}} = \dot{x} \oplus \dot{y} = s(x) + s(y),$$

et

$$s(x \cdot y) = \widehat{\dot{x} \cdot \dot{y}} = \dot{x} \odot \dot{y} = s(x) \cdot s(y).$$

Définition 4.7.10 Soient $(A, +, \cdot)$ et $(B, +, \cdot)$ deux anneaux. Soit $f : (A, +, \cdot) \longrightarrow (B, +, \cdot)$ un morphisme d'anneaux, alors

1. On appelle noyau de f , et on note $\ker(f)$, l'ensemble défini par

$$\ker(f) = \{x \in A : f(x) = 0_B\} = f^{-1}(\{0_B\}) \subset A.$$

2. On appelle image de f , et on note $\text{Im}(f)$ ou $f(A)$, l'ensemble défini par

$$\text{Im}(f) = \{y \in B, \exists x \in A : f(x) = y\} = \{f(x) \in B : x \in A\} = f(A) \subset B.$$

4.7.5.1 Propriétés des morphismes d'anneaux

Les propriétés générales des morphismes d'anneaux sont de fait analogues à celles que nous avons démontrées pour les morphismes de groupes aux sections précédentes. C'est pourquoi nous synthétisons ci-dessous les plus usuelles en laissant au lecteur le soin d'adapter les démonstrations. La méthode de démonstration est semblable à celle utilisée dans le cas des morphismes de groupes. La proposition suivante donne les premières propriétés des morphismes d'anneaux.

Proposition 4.7.7 (Propriétés des morphismes d'anneaux)

1. Soient $(A, +, \cdot), (B, +, \cdot)$ et $(C, +, \cdot)$ trois anneaux. Soit $f : (A, +, \cdot) \longrightarrow (B, +, \cdot), g : (B, +, \cdot) \longrightarrow (C, +, \cdot)$ deux morphismes d'anneaux, alors

(a) $g \circ f : (A, +, \cdot) \longrightarrow (C, +, \cdot)$ est un morphisme d'anneaux. Autrement dit la composée de deux morphismes d'anneaux est un morphisme d'anneaux.

(b) Si f est un isomorphisme d'anneaux, alors $f^{-1} : (B, +, \cdot) \longrightarrow (A, +, \cdot)$ est un isomorphisme d'anneaux aussi. Autrement dit la bijection réciproque d'un isomorphisme d'anneaux est un isomorphisme d'anneaux.

2. Un morphisme d'anneaux $f : (A, +, \cdot) \longrightarrow (B, +, \cdot)$ est injectif (resp. surjectif) si et seulement si $\ker(f) = \{0_A\}$ (resp. $\text{Im}(f) = B$).

3. Si $f : (A, +, \cdot) \longrightarrow (B, +, \cdot)$ est un morphisme d'anneaux, alors
(a)

$$\begin{cases} f(0_A) = 0_B. \\ \forall x \in A : f(-x) = -f(x) \end{cases} .$$

(b)

$$\forall n \in \mathbb{Z}, \forall x \in A : f(nx) = nf(x).$$

(c)

$$\forall n \in \mathbb{N}^*, \forall x \in A : f(x^n) = (f(x))^n .$$

(d) Si N_1 est un sous-anneau de A , alors $f(N_1)$ est un sous-anneau de B . En particulier, l'image de f , $\text{Im}(f)$, est un sous-anneau de B .

(e) Si N_2 est un sous-anneau de B , alors $f^{-1}(N_2)$ est un sous-anneau de A . En particulier, le noyau de f , $\ker(f)$, est un sous-anneau de A .

Preuve. La démonstration, rigoureusement similaire au cas des groupes, est laissée en exercice. $\square$

4.7.6 Idéaux

Les idéaux sont des concepts fondamentaux en algèbre abstraite, en particulier dans l'étude des anneaux. Dans le cadre de la théorie des anneaux non commutatifs, il est important de distinguer deux types d'idéaux : les idéaux à gauche et les idéaux à droite. Un idéal est un sous-ensemble remarquable d'un anneau ; il s'agit d'un sous-groupe du groupe additif de l'anneau qui est également stable par multiplication avec les éléments de l'anneau.

Nous allons maintenant introduire la notion d'idéal, dont le rôle en théorie des anneaux est l'analogue de celui des sous-groupes normaux en théorie des groupes.

4.7.6.1 Notions d'idéal à gauche, à droite, bilatère

Les notions d'idéal à gauche, à droite, et bilatère sont des concepts importants dans la théorie des anneaux, surtout lorsqu'on travaille avec des anneaux non commutatifs.

Définition 4.7.11 (*Notions d'idéal à gauche, à droite, bilatère*)

Soit $(A, +, \cdot)$ un anneau. Alors

1. Une partie non vide L de A est un idéal à gauche de A , si et si

i) L est un sous-groupe de groupe additif $(A, +)$.

ii) Pour tout $x \in L$ et tout $a \in A$, $a \cdot x \in L$ (c'est-à-dire L est absorbant, ce que l'on écrit plus rapidement $AL \subset L$).

Autrement dit

$$L \text{ est un idéal à gauche de } A \iff \begin{cases} 0_A \in L \\ \forall x, y \in L : x - y \in L \\ \forall x \in L, \forall a \in A : a \cdot x \in L. \end{cases}$$

Autrement dit, un idéal à gauche est stable par addition et par multiplication à gauche par n'importe quel élément de A .

2. Une partie non vide R de A est un idéal à droite de A , si et si

- i) R est un sous-groupe de groupe additif $(A, +)$.
 ii) Pour tout $x \in R$ et tout $a \in A$, $x \cdot a \in R$ (ce que l'on écrit plus rapidement $RA \subset R$).
 Autrement dit

$$R \text{ est un idéal à droite de } A \iff \begin{cases} 0_A \in R \\ \forall x, y \in R : x - y \in R \\ \forall x \in R, \forall a \in A : x \cdot a \in R. \end{cases}$$

Ainsi, un idéal à droite est stable par addition et par multiplication à droite par n'importe quel élément de A .

3. Une partie non vide I de A est un idéal bilatère de A (ou simplement un Idéal), si I est à la fois un idéal à gauche et un idéal à droite de A . Autrement dit

$$I \text{ est un idéal bilatère de } A \iff \begin{cases} 0_A \in I \\ \forall x, y \in I : x - y \in I \\ \forall x \in I, \forall a \in A : a \cdot x \in I \text{ et } x \cdot a \in I. \end{cases}$$

4. Tout idéal à gauche (resp. à droite ou bilatère) d'un anneau A , autre que A et l'idéal nul $\{0_A\}$ s'appelle un idéal à gauche (resp. à droite ou bilatère) propre de A .

Remarque 4.7.10

1. Si A est commutatif, les notions d'idéal à gauche, d'idéal à droite et d'idéal bilatère coïncident. Autrement dit, tout idéal est simultanément un idéal à gauche, un idéal à droite et un idéal bilatère.
2. Si l'anneau A n'est pas commutatif, on distingue les idéaux à gauche, c'est-à-dire ceux vérifiant $AI \subset I$ des idéaux à droite qui vérifient quant à eux $IA \subset I$. Un idéal à la fois à gauche et à droite est qualifié de bilatère.
3. Tout idéal d'un anneau A est un sous-anneau de A .

Exemple 4.7.10

1. Si $(A, +, \cdot)$ est un anneau, alors pour tout $c \in A$, on a

$$I_1 = cA = \{c \cdot a : a \in A\},$$

est un idéal à droite de A et

$$I_1 = Ac = \{a \cdot c : a \in A\},$$

est un idéal à gauche de A . En effet,

Pour I_1 :

- (i) L'élément neutre pour l'addition dans A est 0_A . Alors

$$c \cdot 0_A = 0_A.$$

Comme $0_A \in A$, il en résulte que $0_A \in I_1$.

- (ii) Soient $x, y \in I_1$, alors par définition de I_1 , il existe $a_1, a_2 \in A$ tels que

$$x = c \cdot a_1, y = c \cdot a_2.$$

On obtient

$$x - y = c \cdot a_1 - c \cdot a_2 = c \cdot (a_1 - a_2).$$

Comme $(a_1 - a_2) \in A$, il en résulte que $x - y \in I_1$.

- (iii) Maintenant, montrons que I_1 est stable pour la multiplication à droite par des éléments de A . Soient $x \in I_1, a \in A$, alors par définition de I_1 , il existe $a_1 \in A$ tels que

$$x = c \cdot a_1.$$

On obtient

$$x \cdot a = (c \cdot a_1) \cdot a = c \cdot (a_1 \cdot a).$$

Comme $(a_1 \cdot a) \in A$, il en résulte que $x \cdot a \in I_1$. Par conséquent, I_1 est bien un idéal à droite de A . De la même manière, nous pouvons montrer que I_2 est bien un idéal à gauche de A .

2. Dans l'anneau $(\mathcal{F}(\mathbb{R}, \mathbb{R}), +, \cdot)$, l'ensemble des applications qui s'annulent en 0 est un idéal,

$$I = \{f \in \mathcal{F}(\mathbb{R}, \mathbb{R}) : f(0) = 0\}.$$

3. Les idéaux de l'anneau $(\mathbb{Z}, +, \cdot)$ sont les parties de $\mathbb{Z}$ de la forme $I = n\mathbb{Z}$ avec $n \in \mathbb{N}$. En effet : Les sous-groupes de $(\mathbb{Z}, +)$ sont les parties de $\mathbb{Z}$ de la forme $n\mathbb{Z}$ avec $n \in \mathbb{N}$. Le fait que ces parties soient absorbantes est immédiat.

4. Dans tout anneau A , les sous-groupes triviaux A et $\{0_A\}$ sont des idéaux (des idéaux bilatères de A).

Exemple 4.7.11 (Annulateur d'une partie dans un anneau)

1. Soit S une partie non vide d'un anneau $(A, +, \cdot)$. On pose

$$\text{Ann}_g(S) = \{a \in A, \forall s \in S : a \cdot s = 0_A\},$$

$$\text{Ann}_d(S) = \{a \in A, \forall s \in S : s \cdot a = 0_A\},$$

$\text{Ann}_g(S)$ est appelé l'annulateur à gauche de S dans A . Autrement dit, $\text{Ann}_g(S)$ est l'ensemble des éléments de A qui annulent tous les éléments de S lorsqu'ils sont multipliés à gauche.

$\text{Ann}_d(S)$ est appelé l'annulateur à droite de S dans A . Autrement dit, $\text{Ann}_d(S)$ est l'ensemble des éléments de A qui annulent tous les éléments de S lorsqu'ils sont multipliés à droite.

On vérifie que, quelle que soit la partie non vide S de A , $\text{Ann}_g(S)$ (resp. $\text{Ann}_d(S)$) est un idéal à gauche (resp. à droite) de A . En effet :

(i) L'élément neutre pour l'addition dans A est 0_A . Alors pour tout $s \in S$,

$$0_A \cdot s = 0_A.$$

Comme $0_A \in A$, il en résulte que $0_A \in \text{Ann}_g(S)$.

(ii) Soient $a, b \in \text{Ann}_g(S)$, alors pour tout $s \in S$

$$a \cdot s = 0_A \text{ et } b \cdot s = 0_A.$$

Par suite

$$(a - b) \cdot s = a \cdot s - b \cdot s = 0_A.$$

Donc $a - b \in \text{Ann}_g(S)$.

(iii) Soient $r \in \text{Ann}_g(S)$, $a \in A$, alors pour tout $s \in S$,

$$(a \cdot r) \cdot s = a \cdot (r \cdot s) = a \cdot 0_A = 0_A,$$

il en résulte que $a \cdot r \in \text{Ann}_g(S)$. Par conséquent, $\text{Ann}_g(S)$ est bien un idéal à droite de A .

De la même manière, nous pouvons montrer que $\text{Ann}_d(S)$ est bien un idéal à gauche de A .

2. Si l'anneau A est commutatif, on appelle annulateur d'une partie non vide S de A , l'idéal noté $\text{Ann}(S)$ tel que

$$\text{Ann}(S) = \text{Ann}_g(S) = \text{Ann}_d(S).$$

Par exemple, si $A = \mathbb{Z}$ (l'anneau des entiers) et $S = \{2\}$, alors

$$\text{Ann}(S) = \text{Ann}_g(S) = \text{Ann}_d(S) = \{0\}.$$

4.7.7 Opérations sur les idéaux d'un anneau

Les idéaux dans un anneau permettent une variété d'opérations qui aident à explorer et manipuler la structure de l'anneau. Voici un résumé des principales opérations que l'on peut effectuer sur les idéaux, comme l'intersection, la somme, et le produit.

4.7.7.1 Somme d'idéaux

La somme d'idéaux est une opération importante dans la théorie des anneaux et des idéaux. Elle permet de combiner plusieurs idéaux pour obtenir un nouvel idéal.

Définition 4.7.12 Soient $(A, +, \cdot)$ un anneau et I, J deux idéaux de A de même nature (i.e., deux idéaux à gauche (resp. à droite (resp. bilatères)), alors on appelle somme de I et J et on note $I + J$, le sous-ensemble de A suivant

$$I + J = \{x \in A, \exists a \in I \text{ et } \exists b \in J : x = a + b\}.$$

Autrement dit $I + J$ est l'ensemble de toutes les sommes possibles d'un élément de I et d'un élément de J .

Exemple 4.7.12 Considérons l'anneau (commutatif) des entiers $\mathbb{Z}$ et les idéaux $I = 4\mathbb{Z}$ et $J = 6\mathbb{Z}$, où $4\mathbb{Z}$ est l'ensemble des multiples de 4 et $6\mathbb{Z}$ est l'ensemble des multiples de 6. La somme des idéaux I et J est définie par

$$I + J = \{x = 4n + 6m : n, m \in \mathbb{Z}\}.$$

On peut simplifier cette expression en remarquant que

$$x = 4n + 6m = 2(2n + 3m).$$

Cela montre que chaque élément de $I + J$ est un multiple de 2. En fait, on peut vérifier que

$$I + J = 2\mathbb{Z}.$$

Donc, la somme des idéaux $4\mathbb{Z}$ et $6\mathbb{Z}$ dans $\mathbb{Z}$ est l'idéal $2\mathbb{Z}$, l'ensemble des multiples de 2.

Remarque 4.7.11 Plus généralement, si $I_1, I_2, \dots, I_n$ sont des idéaux de A de même nature, il en est de même de leur somme

$$I = I_1 + I_2 + \dots + I_n = \{x_1 + x_2 + \dots + x_n : x_1 \in I_1, x_2 \in I_2, \dots, x_n \in I_n\}.$$

L'idéal $I_1 + I_2 + \dots + I_n$ contient $I_1, I_2, \dots, I_n$.

Proposition 4.7.8 Si I, J sont deux idéaux à gauche (resp. à droite, bilatères) d'un anneau $(A, +, \cdot)$, alors $I + J$ est un idéal à gauche (resp. à droite, bilatère) de A qui contient I et J .

Preuve. Démontrons la propriété dans le cas des idéaux à gauche (celui de droite étant similaire).

1. Il est clair que

$$I \subset I + J \text{ et } J \subset I + J.$$

En effet, pour tout $x \in I$ (resp. J), on a

$$x = x + 0_A, 0_A \in J \text{ (resp. } x = 0_A + x, 0_A \in I).$$

Par suite

$$x \in I + J.$$

2. Montrons maintenant que $I + J$ est un idéal à gauche.

(i) L'ensemble $I + J$ est non vide car il contient $0_A = 0_A + 0_A$.

(ii) Soient $x, y \in I + J$, alors il existe $(x_1, y_1) \in I \times J$ et $(x_2, y_2) \in I \times J$ tels que

$$x = x_1 + y_1 \text{ et } y = x_2 + y_2.$$

On a donc, par commutativité et associativité de la loi $+$,

$$x - y = (x_1 - x_2) + (y_1 - y_2) \in I + J,$$

car $(x_1 - x_2) \in I$ puisque I est un sous groupe de $(A, +)$ et de même pour $(y_1 - y_2) \in J$.

(iii) Soit $a \in A$ et soit $x = x_1 + y_1 \in I + J$, alors par distributivité de $\cdot$ par rapport à $+$,

$$a \cdot x = a \cdot (x_1 + y_1) = a \cdot x_1 + a \cdot y_1,$$

comme $a \cdot x_1 \in I$ puisque I est un idéal à gauche de A et de même $a \cdot y_1 \in J$, on a donc

$$a \cdot x \in I + J.$$

Ceci prouve que $I + J$ est un idéal à gauche. On montre de même que $x \cdot a \in I + J$. $\square$

4.7.7.2 Intersection d'idéaux

L'opération d'intersection d'idéaux est effectivement l'intersection au sens de la théorie des ensembles. Lorsque l'on parle de l'intersection de deux ou plusieurs idéaux dans un anneau A , on prend simplement l'intersection des ensembles correspondants.

Définition 4.7.13 *L'intersection de deux idéaux I et J dans un anneau $(A, +, \cdot)$ est définie comme l'ensemble des éléments qui appartiennent à la fois à I et à J . Formellement, l'intersection $I \cap J$ est*

$$I \cap J = \{x \in A : x \in I \text{ et } x \in J\}.$$

Exemple 4.7.13 *Considérons l'anneau (commutatif) des entiers $\mathbb{Z}$ et les idéaux $I = 4\mathbb{Z}$ et $J = 6\mathbb{Z}$, alors l'intersection*

$$I \cap J = \{x \in \mathbb{Z} : x \text{ est un multiple de 4 et un multiple de 6}\}.$$

Pour un entier x dans $I \cap J$, x doit être divisible à la fois par 4 et par 6. Pour trouver cet ensemble, nous devons déterminer le plus petit commun multiple (PPCM) de 4 et 6. Donc, le PPCM de 4 et 6 est 12. Cela signifie que tout multiple de 12 est à la fois un multiple de 4 et un multiple de 6. Ainsi

$$I \cap J = 12\mathbb{Z}.$$

Proposition 4.7.9 *Soit $(I_i)_{i \in J}$ une famille non vide d'idéaux à gauche (resp. à droite, bilatères) d'un anneau $(A, +, \cdot)$, alors l'intersection $\bigcap_{i \in J} I_i$ est un idéal à gauche (resp. à droite, bilatère) de A .*

Preuve. La preuve est la même que pour les sous-anneaux. Démontrons la propriété dans le cas des idéaux à gauche (celui de droite étant similaire). Soit donc $(I_i)_{i \in J}$ une famille d'idéaux à gauche de A . Posons $I = \bigcap_{i \in J} I_i$ l'intersection de tous les I_i . On sait déjà que I est un sous-groupe additif en tant qu'intersection d'une famille de sous-groupes. Soient $a \in A$ et $x \in I$, on a $a \cdot x \in I_i$ pour tout $i \in J$ puisque I_i est un idéal à gauche, et donc $a \cdot x \in I$. Ce qui prouve que I est un idéal à gauche de A . On prouve de même que $x \cdot a \in \bigcap_{i \in J} I_i$. $\square$

Remarque 4.7.12

1. De manière générale, une réunion d'idéaux n'est pas un idéal, par exemple $2\mathbb{Z} \cup 3\mathbb{Z}$ n'est pas un idéal de $\mathbb{Z}$.
2. Si I et J sont des idéaux, pour que $I \cup J$ soit un idéal, il faut, et il suffit, que $I \subset J$ ou $J \subset I$.

4.7.7.3 Produit d'idéaux

Le produit de deux idéaux dans un anneau est une autre opération importante qui joue un rôle central dans l'étude de la structure des anneaux. Cette opération permet de combiner les éléments de deux idéaux pour former un nouvel idéal.

Soient I, J deux parties non vides d'un anneau $(A, +, \cdot)$. En général, l'ensemble

$$P = \{ij : i \in I, j \in J\},$$

n'est pas un idéal de A et ceci même si I et J sont des idéaux de A .

Prenons l'anneau commutatif $\mathbb{Z}[X]$, l'anneau des polynômes à coefficients entiers. Considérons deux idéaux I et J dans cet anneau suivants

$$I = (2, X) = \{2P_1 + XQ_1 : P_1, Q_1 \in \mathbb{Z}[X]\},$$

$$J = (3, X) = \{3P_2 + XQ_2 : P_2, Q_2 \in \mathbb{Z}[X]\}.$$

Alors

$$\begin{aligned} P &= \{T_1 \cdot T_2 : T_1 \in I, T_2 \in J\} = \{(2P_1 + XQ_1)(3P_2 + XQ_2) : P_1, Q_1, P_2, Q_2 \in \mathbb{Z}[X]\} \\ &= \{6P_1P_2 + (2P_1Q_2 + 3P_2Q_1)X + Q_1Q_2X^2 : P_1, Q_1, P_2, Q_2 \in \mathbb{Z}[X]\}. \end{aligned}$$

Il n'est pas difficile de constater que

$$a_1 = X^2 \in P \text{ car } (Q_1 = 1 = Q_2, P_1 = P_2 = 0),$$

et

$$a_2 = 6 \in P \text{ car } (P_1 = P_2 = 1, Q_1 = 0 = Q_2).$$

Mais d'autre part, on doit avoir $X^2 + 6 \notin P$. Par l'absurde, si $X^2 + 6 \in P$, on doit pouvoir écrire

$$X^2 + 6 = (2P_1 + XQ_1)(3P_2 + XQ_2) = 6P_1P_2 + (2P_1Q_2 + 3P_2Q_1)X + X^2Q_1Q_2,$$

pour des éléments $P_1, Q_1, P_2, Q_2 \in \mathbb{Z}[X]$. On développe alors le membre de droite pour obtenir

$$Q_1Q_2 = 1 = P_1P_2.$$

Il s'ensuit que

$$P_i = \mp 1, Q_i = \mp 1,$$

pour chaque $i \in \{1, 2\}$. Ceci donne $2P_1Q_2 + 3Q_1P_2 \neq 0$, ce qui est contradiction. Cet exemple montre que l'ensemble

$$P = \{T_1 \cdot T_2 : T_1 \in I, T_2 \in J\},$$

n'est pas un idéal de $\mathbb{Z}[X]$, même si I et J sont des idéaux. Le problème réside dans le fait que P n'est pas stable sous l'addition, une propriété essentielle pour être un idéal. Pour obtenir un véritable idéal, on doit prendre le produit des idéaux I et J , qui est défini comme l'ensemble des sommes finies de produits $T_1 \cdot T_2$. Alors, on considère la définition suivante.

Définition 4.7.14 (Produit d'idéaux)

Soient I, J deux idéaux à gauche (resp. à droite, bilatères) d'un anneau $(A, +, \cdot)$. On appelle produit de I et J l'idéal IJ défini par

$$IJ = \left\{ x = \sum_{k=1}^m i_k \cdot j_k : m \in \mathbb{N}^*, i_k \in I, j_k \in J, \forall k = \overline{1, m} \right\}.$$

Autrement dit, IJ est l'ensemble de toutes les combinaisons linéaires finies des produits d'éléments de I avec des éléments de J .

Remarque 4.7.13 Le produit des idéaux dans un anneau $(A, +, \cdot)$ n'est pas généralement commutatif, sauf si l'anneau A lui-même est commutatif. En effet,

- Si A n'est pas commutatif, il n'est pas nécessairement vrai que $i \cdot j = j \cdot i$ pour tous $i \in I$ et $j \in J$.
- Par conséquent, il se peut que $IJ \neq JI$. Le produit des idéaux dépend de l'ordre des facteurs, et donc le produit des idéaux peut ne pas être commutatif.

Corollaire 4.7.3 Si I, J sont deux idéaux de $(A, +, \cdot)$ de même nature. Alors le produit de I et J est un idéal. De plus $IJ \subset I \cap J$.

Preuve. Montrons la propriété dans le cas des idéaux à gauche (celui de droite étant similaire).

1. L'ensemble IJ est non vide car il contient 0_A . En effet

$$0_A \in I, 0_A \in J \implies 0_A \cdot 0_A = 0_A \in IJ.$$

De plus si

$$\begin{aligned} x &= \sum_{k=1}^n i_k \cdot j_k, i_k \in I, j_k \in J, n \in \mathbb{N}^*, \\ y &= \sum_{l=1}^m i'_l \cdot j'_l, i'_l \in I, j'_l \in J, m \in \mathbb{N}^*. \end{aligned}$$

Alors

$$x - y = \sum_{k=1}^n i_k \cdot j_k - \sum_{l=1}^m i'_l \cdot j'_l = \sum_{k=1}^n i_k \cdot j_k + \sum_{l=1}^m (-i'_l) \cdot j'_l.$$

En posant, pour k allant de $(n+1)$ à $(n+m)$,

$$\begin{cases} i_k = -i'_{k-n} \\ j_k = j'_{k-n}. \end{cases}$$

On a

$$x - y = \sum_{k=1}^{n+m} i_k \cdot j_k.$$

Cette somme est une combinaison finie de produits d'éléments de I avec des éléments de J , et donc $x - y$ est aussi un élément de IJ . Par conséquent IJ est un sous groupe de $(A, +)$.

Enfin, pour tout a dans A , on a

$$a \cdot x = a \cdot \sum_{k=1}^n i_k \cdot j_k = \sum_{k=1}^n (a \cdot i_k) \cdot j_k \in IJ,$$

puisque chaque $a \cdot i_k$ (resp. j_k) est un élément de I (resp. de J). Ceci prouve que IJ est un idéal à gauche.

On montre de même que $x \cdot a \in IJ$.

2. Soit $x = \sum_{k=1}^n i_k \cdot j_k$ un élément de IJ . Pour chaque k , on a $i_k \cdot j_k$ est un élément de I puisque I est un idéal à gauche et stable sous la multiplication par des éléments de A ($j_k \in A$ car $J \subset A$). Comme I est de plus stable par la somme, IJ est bien contenu dans I . Par symétrie du rôle joué par I et J , on a aussi IJ est contenu dans J et donc IJ est contenu dans $I \cap J$. $\square$

4.7.7.4 Propriétés des opérations sur les idéaux

Désignons par $\mathcal{I}_g$ (resp. $\mathcal{I}_d, \mathcal{I}$) l'ensemble des idéaux à gauche (resp. à droite, bilatères) d'un anneau $(A, +, \cdot)$. L'étude faite dans les paragraphes précédents montre que l'intersection, la somme et le produit des idéaux définissent des lois de composition dans les ensembles $\mathcal{I}_g, \mathcal{I}_d, \mathcal{I}$; ces lois vérifient les propriétés suivantes :

Proposition 4.7.10 *$(A, +, \cdot)$ étant un anneau, quels que soient I, J, K , dans $\mathcal{I}_g$ (resp. $\mathcal{I}_d, \mathcal{I}$), on a*

$$1. \quad I \cap J = J \cap I, I + J = J + I.$$

2.

$$(I \cap J) \cap K = I \cap (J \cap K).$$

3.

$$(I + J) + K = I + (J + K).$$

4.

$$(IJ)K = I(JK).$$

5.

$$I(J + K) = IJ + IK, (I + J)K = IK + JK.$$

Preuve. La preuve est laissée comme un exercice au lecteur. □

Remarque 4.7.14

1. La proposition 4.7.10 exprime que dans $\mathcal{I}_g$ (resp. $\mathcal{I}_d, \mathcal{I}$), les opérations "intersection" et "somme" sont commutatives et associatives, le produit est associatif et distributif à droite et à gauche par rapport à la "somme". Le "produit" des idéaux n'est commutatif que si l'anneau A est commutatif.
2. L'associativité du produit des idéaux implique que pour tout idéal à gauche, à droite ou bilatère de $(A, +, \cdot)$ et tout $n > 1$ dans $\mathbb{N}$, l'idéal I^n est défini inductivement par

$$I^n = II^{n-1} = I^{n-1}I.$$

Proposition 4.7.11 *Soient $(A, +, \cdot), (B, +, \cdot)$ deux anneaux et $f : (A, +, \cdot) \longrightarrow (B, +, \cdot)$ un morphisme d'anneaux. Alors*

- (1) $\ker(f)$ est un idéal bilatère de A .
- (2) Si J est un idéal à gauche (resp. à droite (resp. bilatère)) de B , alors $f^{-1}(J)$ est un idéal à gauche (resp. à droite (resp. bilatère)) de A .
- (3) Si f est surjective et I un idéal à droite (resp. à gauche, bilatère) de A , alors $f(I)$ est un idéal à droite (resp. à gauche, bilatère) de B .

Preuve. On va prouver les cas des idéaux à gauche. Les cas à droite se traitent par symétrie et les cas bilatères s'en déduisent.

1. On sait déjà que $\ker(f)$ est un sous-groupe de A pour l'addition, il suffit de vérifier qu'il est stable par multiplication à droite et à gauche par un élément de A . Or clairement, si $x \in \ker(f)$ et $a \in A$, on a

$$f(a \cdot x) = f(a) \cdot f(x) = f(a) \cdot 0_A = 0_A.$$

Il en résulte que $a \cdot x$ appartient à $\ker(f)$.

Par ailleurs

$$f(x \cdot a) = f(x) \cdot f(a) = 0_A \cdot f(a) = 0_A,$$

ce qui entraîne

$$x \cdot a \in \ker(f).$$

Par conséquent $\ker(f)$ est un idéal bilatère de A .

2. Sous les hypothèses de (2), $f^{-1}(J)$ est un sous-groupe additif de A comme image réciproque d'un sous-groupe additif de B par un morphisme de groupes.

Par ailleurs, soient $a \in A$ et $x \in f^{-1}(J)$ (i.e, $x \in A$), on a

$$f(a \cdot x) = f(a) \cdot f(x),$$

avec $f(a) \in B$ et $f(x) \in J$, donc $f(a \cdot x) \in J$ puisque J est un idéal à gauche de B , c'est-à-dire $a \cdot x \in f^{-1}(J)$. Ceci prouve que $f^{-1}(J)$ est un idéal à gauche de A .

3. Supposons que f est surjective et soit I un idéal à gauche de A , alors $f(I)$ est un sous-groupe de B . Par ailleurs, soient $y \in f(I)$, $c \in B$, alors

$$\begin{cases} \exists i \in I : f(i) = y \\ \exists a \in A : f(a) = c, \end{cases}$$

donc

$$c \cdot y = f(a) \cdot f(i) = f(a \cdot i).$$

Puisque i appartient à I , $a \cdot i$ appartient à I . Il en résulte que $c \cdot y = f(a \cdot i)$ appartient à $f(I)$ et le résultat s'ensuit. $\square$

Remarque 4.7.15 Notons que si f n'est pas surjective, l'image directe d'un idéal de A par f n'est pas forcément un idéal de B (prendre pour f l'injection canonique de $\mathbb{Z}$ dans $\mathbb{Q}$). Par contre l'image $\text{Im}(f)$ de f est un sous-anneau de B . De plus pour tout idéal I de A , l'image directe $f(I)$ est un idéal de l'anneau $f(A) = \text{Im}(f)$.

4.8 Corps

Les corps sont des structures centrales en algèbre et en mathématiques en général, avec des applications profondes en algèbre linéaire, en théorie des nombres, en géométrie, en analyse, ainsi que dans diverses branches de la physique et de l'informatique. Un corps est défini comme un ensemble muni de deux opérations internes qui permettent d'effectuer des additions, des soustractions, des multiplications et des divisions (sauf par zéro). Plus précisément, un corps se compose de deux opérations : l'addition, notée $+$, et la multiplication, notée $\cdot$, qui doivent satisfaire certaines propriétés. Ces propriétés constituent essentiellement une combinaison des caractéristiques des anneaux et des groupes.

Définition 4.8.1 On considère un ensemble $\mathbb{K}$ non vide muni de deux lois de composition interne, notées $+$ et $\cdot$.

1. On appelle corps $(\mathbb{K}, +, \cdot)$ un anneau unitaire, non réduit à $\{0_A\}$ dans lequel tout élément non nul $x \neq 0_A$ possède un inverse. Ceci peut être exprimé de manière plus modulaire comme suit :

$(\mathbb{K}, +, \cdot)$ est un corps si et si

(i) $(\mathbb{K}, +)$ est un groupe Abélien.

(ii) $(\mathbb{K}^*, \cdot)$ est un groupe.

(iii) La multiplication $\cdot$ est distributive à gauche et à droite par rapport à l'addition $+$.

2. Un corps $(\mathbb{K}, +, \cdot)$ est dit commutatif si et seulement si la loi de multiplication $\cdot$ est commutative. Cela signifie que pour tous les éléments x et y du corps $\mathbb{K}$, on a

$$x \cdot y = y \cdot x.$$

Exemple 4.8.1

1. $(\mathbb{Q}, +, \cdot), (\mathbb{R}, +, \cdot), (\mathbb{C}, +, \cdot)$ sont des corps commutatifs (fameux).
2. L'ensemble $(\mathbb{Z}, +, \cdot)$ n'est pas un corps car la plupart des éléments de $\mathbb{Z}^*$ ne sont pas inversibles : par exemple, il n'existe pas d'entier relatif n tel que $2 \cdot n = 1$ donc 2 n'est pas inversible.

4.8.1 Sous-corps

Un sous-corps d'un corps est un concept fondamental en algèbre, particulièrement dans l'étude des structures algébriques et des extensions de corps. Un sous-corps est un ensemble d'éléments qui satisfait aux mêmes propriétés qu'un corps complet, mais qui est contenu dans un corps plus grand.

Définition 4.8.2 On appelle sous-corps d'un corps $(\mathbb{K}, +, \cdot)$ un sous-ensemble L de $\mathbb{K}$ qui est lui-même un corps par rapport à l'addition et à la multiplication de $\mathbb{K}$. On dit aussi que $\mathbb{K}$ est un sur-corps ou une extension du sous-corps L .

De manière équivalente, cette définition peut s'écrire comme suit.

Définition 4.8.3 Soit $(\mathbb{K}, +, \cdot)$ un corps. On appelle sous-corps de $\mathbb{K}$ tout sous-anneau unitaire L de $\mathbb{K}$ tel que l'inverse de tout élément non-nul de L appartienne à L . On a donc

$$\begin{aligned}
 L \text{ est un sous-corps de } \mathbb{K} &\iff \begin{cases} j) 1_{\mathbb{K}} \in L \\ jj) (L, +, \cdot) \text{ est un sous-anneau de l'anneau } (\mathbb{K}, +, \cdot) \\ jjj) \forall x \in L^* : x^{-1} \in L. \end{cases} \\
 &\iff \begin{cases} SC_1) 0_{\mathbb{K}}, 1_{\mathbb{K}} \in L \\ SC_2) \forall x, y \in L : x - y \in L \\ SC_3) \forall x, y \in L : x \cdot y \in L \\ SC_4) \forall x \in L^* : x^{-1} \in L. \end{cases} \\
 &\iff \begin{cases} SC_1) 0_{\mathbb{K}}, 1_{\mathbb{K}} \in L \\ SC_2) \forall x, y \in L : x - y \in L \\ SC_3) \forall x \in L, \forall y \in L^* : x \cdot y^{-1} \in L. \end{cases}
 \end{aligned}$$

Remarque 4.8.1 Ainsi un sous-ensemble L de $\mathbb{K}$ est un sous-corps si L est un sous-groupe additif de $\mathbb{K}$ (pour l'addition de $\mathbb{K}$) et si $L^* = L - \{0_A\}$ est un sous-groupe multiplicatif de $\mathbb{K}^*$.

Exemple 4.8.2

1. Dans $(\mathbb{Q}, +, \cdot), (\mathbb{R}, +, \cdot), (\mathbb{C}, +, \cdot)$, chacun est un sous-corps du suivant.
2. Le seul sous-corps de $(\mathbb{Q}, +, \cdot)$ est lui-même.
3. L'ensemble

$$L = \left\{ x = r + s\sqrt{2} : r, s \in \mathbb{Q} \right\},$$

est un sous-corps de $(\mathbb{R}, +, \cdot)$ noté $\mathbb{Q}(\sqrt{2})$.

4. L'ensemble

$$L = \left\{ x = r + s\sqrt{2} : r, s \in \mathbb{Z} \right\},$$

n'est pas un sous-corps de $(\mathbb{R}, +, \cdot)$ car $x = 2 + \sqrt{2} \in L$ n'est pas inversible dans L ($x^{-1} = \frac{1}{x} = 1 - \frac{1}{2}\sqrt{2} \notin L$).

Proposition 4.8.1 Soit $(\mathbb{K}, +, \cdot)$ un corps. Si $(L_i)_{i \in I}$ est une famille de sous-corps de $\mathbb{K}$, alors $\bigcap_{i \in I} L_i$ est un sous-corps de $\mathbb{K}$. Autrement dit toute intersection de sous-corps du corps $\mathbb{K}$ est un sous-corps de $\mathbb{K}$.

Preuve. Comme cette propriété est déjà vraie pour les anneaux unitaires, l'intersection des sous-corps de $\mathbb{K}$ est un sous-anneau unitaire de $\mathbb{K}$. Soit x un élément non nul de cette intersection. Dans chacun des sous-corps, cet élément est inversible dont l'inverse correspond à x^{-1} , l'inverse de x dans $\mathbb{K}$. Ainsi x^{-1} est dans chacun des sous-corps et x est inversible dans l'intersection qui est donc un sous-corps de $\mathbb{K}$. $\square$

es corps possèdent une propriété fondamentale énoncée de manière simple : "un produit est nul si et seulement si l'un des facteurs est nul". Cette caractéristique est essentielle et distingue les corps des autres structures algébriques, comme les anneaux, où des diviseurs de zéro peuvent exister. Dans un corps, l'absence de diviseurs de zéro garantit cette propriété essentielle, ce qui assure des comportements multiplicatifs bien définis et prévisibles.

Proposition 4.8.2 Tout corps est un anneau intègre.

Preuve. Soit donc un corps $\mathbb{K}$. Soient $x, y \in \mathbb{K}$ tels que $x \cdot y = 0_{\mathbb{K}}$. Pour satisfaire l'intégrité au sens de la Définition 4.7.6, nous devons montrer que $x = 0_{\mathbb{K}}$ ou $y = 0_{\mathbb{K}}$.

Supposons premièrement que $x \neq 0_{\mathbb{K}}$ et cherchons à montrer que $y = 0_{\mathbb{K}}$. Puisque $\mathbb{K}$ est un corps, l'élément $x \in \mathbb{K}^*$ est inversible dans $\mathbb{K}$, c'est-à-dire qu'il existe un élément, noté x^{-1} , tel que

$$x \cdot x^{-1} = x^{-1} \cdot x = 1_{\mathbb{K}}.$$

Multiplions alors l'égalité $x \cdot y = 0_{\mathbb{K}}$ à gauche par x^{-1} , ce qui nous donne

$$x^{-1} \cdot (x \cdot y) = x^{-1} \cdot 0_{\mathbb{K}}.$$

Par suite

$$(x^{-1} \cdot x) \cdot y = 0_{\mathbb{K}},$$

et donc nous obtenons bien $y = 0_{\mathbb{K}}$.

Deuxièmement, si nous supposons $y \neq 0_{\mathbb{K}}$, le même argument (symétrisé) donne $x = 0_{\mathbb{K}}$. $\square$

Dans un anneau, un élément ne possède pas nécessairement de symétrie par rapport à la deuxième loi (ou d'inverse puisque la deuxième loi est notée multiplicativement). Il est à noter que l'élément zéro est absorbant. Il ne peut donc pas être inversible. Considérons par exemple l'ensemble quotient $\mathbb{Z}/4\mathbb{Z} = \{\dot{0}, \dot{1}, \dot{2}, \dot{3}\}$ muni des lois $\oplus$ et $\odot$, on obtient

$\odot$	$\dot{0}$	$\dot{1}$	$\dot{2}$	$\dot{3}$
$\dot{0}$	$\dot{0}$	$\dot{0}$	$\dot{0}$	$\dot{0}$
$\dot{1}$	$\dot{0}$	$\dot{1}$	$\dot{2}$	$\dot{3}$
$\dot{2}$	$\dot{0}$	$\dot{2}$	$\dot{0}$	$\dot{2}$
$\dot{3}$	$\dot{0}$	$\dot{3}$	$\dot{2}$	$\dot{1}$

Il est immédiat d'après la table de multiplication que les classes d'équivalence $\dot{1}$ et $\dot{3}$ possèdent des inverses

$$(\dot{1})^{-1} = \dot{1} \text{ et } (\dot{3})^{-1} = \dot{3},$$

où $(\dot{1})^{-1}$ et $(\dot{3})^{-1}$ désignent les inverses respectifs de $\dot{1}$ et $\dot{3}$ pour la loi $\odot$. En revanche, la classe d'équivalence $\dot{2}$ ne possède pas d'inverse puisque

$$\forall \dot{p} \in \mathbb{Z}/4\mathbb{Z} : \dot{2} \odot \dot{p} \neq \dot{1}.$$

Le résultat suivant fournit un exemple important de corps.

Théorème 4.8.1 *Pour tout entier $n \geq 2$, On a*

- 1. 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.*
- 2. L'anneau $\mathbb{Z}/n\mathbb{Z}$ est un corps si et seulement si n est un nombre premier.*

Chapitre 5

Anneau des polynômes

Enfin, ce chapitre se concentre sur les polynômes, leurs propriétés, les opérations sur eux, et leur factorisation. Les notions d'irréductibilité et de division euclidienne sont aussi traitées, ainsi que les racines de polynômes et leurs multiplicités. Dans tout ce chapitre, $(\mathbb{K}, +, \cdot)$ désigne un corps (En pratique, le plus souvent, $\mathbb{K} = \mathbb{R}$ ou $\mathbb{C}$).

5.1 Polynôme à une indéterminée à coefficients dans $\mathbb{K}$

Un polynôme à une indéterminée à coefficients dans $\mathbb{K}$ est une expression algébrique formée par une somme de termes de la forme $a_n X^n$, où X est une indéterminée, a_n est un coefficient pris dans un corps $\mathbb{K}$, et n est un entier naturel. Le concept des polynômes à une indéterminée est fondamental en algèbre, en particulier dans l'étude des structures algébriques comme les anneaux et les corps.

Définition 5.1.1 (*Polynôme en termes de suites*)

1. On appelle polynôme à une indéterminée (souvent notée X) à coefficients dans $\mathbb{K}$ toute suite presque nulle d'éléments de $\mathbb{K}$, c'est-à-dire toute suite d'éléments de $\mathbb{K}$ nulle à partir d'un certain rang.
2. Si $(a_n)_{n \in \mathbb{N}}$ est un polynôme à coefficients dans $\mathbb{K}$, alors

$$\exists k \in \mathbb{N}, \forall n > k : a_n = 0_{\mathbb{K}}.$$

On note $P = (a_0, a_1, \dots, a_k, 0_{\mathbb{K}}, 0_{\mathbb{K}}, \dots)$, et les éléments $a_0, a_1, \dots, a_k$ sont appelés coefficients du polynôme P .

Notation 5.1.1 L'ensemble des polynômes à coefficients dans $\mathbb{K}$ est noté $\mathbb{K}[X]$ si on choisit de noter X l'indéterminée.

Définition 5.1.2 (*Polynôme constant, polynôme nul, monôme*)

1. On appelle polynôme constant de $\mathbb{K}[X]$ tout polynôme $P = (\alpha, 0_{\mathbb{K}}, 0_{\mathbb{K}}, \dots)$ avec $\alpha \in \mathbb{K}$. Un tel polynôme sera simplement noté α , c'est-à-dire tout polynôme dont les coefficients sont nuls à partir du rang 1. On a donc

$$\begin{cases} a_0 = \alpha, \\ \forall n \geq 1 : a_n = 0_{\mathbb{K}}. \end{cases}$$

Un polynôme constant est ainsi associé à une suite constante presque nulle, avec uniquement le premier terme non nul, et tous les autres termes étant $0_{\mathbb{K}}$. Ce type de polynôme ne dépend pas de

l'indéterminée X .

2. Le polynôme nul, noté généralement $0_{\mathbb{K}[X]} = 0$, est le polynôme dont tous les coefficients sont nuls. Il est associé à la suite nulle $(0_{\mathbb{K}}, 0_{\mathbb{K}}, 0_{\mathbb{K}}, \dots)$.
3. On appelle monôme la suite $(a_n)_{n \in \mathbb{N}} \in \mathbb{K}^{\mathbb{N}}$ telles que

$$\exists n_0 \in \mathbb{N}, \forall n \in \mathbb{N} : n \neq n_0 \implies a_n = 0_{\mathbb{K}}.$$

Cela signifie qu'un monôme est une suite où un seul élément a_{n_0} est non nul, tandis que tous les autres termes a_n sont nuls.

Définition 5.1.3 (Notation polynomiale, forme d'un polynôme)

1. Tout polynôme P peut s'écrire de façon unique sous forme

$$P = \sum_{n=0}^{+\infty} a_n X^n,$$

où $(a_n)_{n \in \mathbb{N}}$ est une suite presque nulle d'éléments de $\mathbb{K}$, appelée suite des coefficients de P (la suite des coefficients de P est donc unique).

2. Les polynômes de la forme αX^k pour $k \in \mathbb{N}$ et $\alpha \in \mathbb{K}$ sont appelés monômes. Cela signifie que tout polynôme peut être exprimé comme une somme de monômes, où les puissances de l'indéterminée sont écrites de manière explicite, et chaque monôme a un coefficient non nul dans un corps $\mathbb{K}$.

Remarque 5.1.1 Pour un polynôme $P = \sum_{k=0}^m a_k X^k$, on complète parfois la suite $(a_n)_{n \in \mathbb{N}}$ des coefficients de P en posant $a_n = 0_K$ si $n > m$. On peut alors écrire

$$P = \sum_{k=0}^m a_k X^k = \sum_{k=0}^{+\infty} a_k X^k.$$

Exemple 5.1.1

1. Le polynôme constant $P = 5$ est associé à la suite constante de coefficients $(5, 0, 0, 0, \dots)$.
2. Le monôme $2X$ correspond à la suite de coefficients $(0, 2, 0, 0, 0, \dots)$.
3. Le polynôme $P = 3X^3 - 2X^2 + 4X + 7$ est associé à la suite de coefficients $(7, 4, -2, 3, 0, 0, 0, \dots)$.

Définition 5.1.4 (Identification des coefficients)

Deux polynômes $P = (a_n)_{n \in \mathbb{N}} \in \mathbb{K}[X]$ et $Q = (b_n)_{n \in \mathbb{N}} \in \mathbb{K}[X]$ sont dits égaux ou identiques si et seulement si leurs coefficients correspondants sont égaux pour toutes les suites presque nulles $(a_n)_{n \in \mathbb{N}}$ et $(b_n)_{n \in \mathbb{N}}$. En d'autres termes,

$$\sum_{n=0}^{+\infty} a_n X^n = \sum_{n=0}^{+\infty} b_n X^n \iff \forall n \in \mathbb{N} : a_n = b_n.$$

Exemple 5.1.2 Soit $P = 2X^3 - 9X^2 + 13X - 6$. On demande de déterminer des réels a, b et c tels que

$$P = (X - 2)(aX^2 + bX + c).$$

Pour que cette équation soit vraie pour tous les X , les coefficients correspondants doivent être égaux. Alors on développe

$$(X - 2)(aX^2 + bX + c) = aX^3 + (b - 2a)X^2 + (c - 2b)X - 2c,$$

et donc, par identification, on a

$$P = (X - 2)(aX^2 + bX + c) \iff \begin{cases} a = 2 \\ b - 2a = -9 \\ c - 2b = 13 \\ -2c = -6. \end{cases}$$

On résout ce système et on trouve comme unique solution le triplet $(a = 2, b = -5, c = 3)$.

Les concepts de polynôme pair et de polynôme impair sont analogues aux notions de fonctions paires et impaires. Un polynôme peut être classé comme pair ou impair en fonction de la parité de ses puissances dans son expression développée.

Définition 5.1.5 (Parité -Polynôme pair/impair-)

Soit $P = \sum_{n=0}^{+\infty} a_n X^n$ un polynôme de $\mathbb{K}[X]$.

1) On dit que P est pair lorsque

$$P(-X) = P(X).$$

Cela signifie que le polynôme reste inchangé lorsque l'on remplace X par $-X$.

2) On dit que P est impair lorsque

$$P(-X) = -P(X).$$

Cela signifie que le polynôme change de signe lorsque l'on remplace X par $-X$.

Exemple 5.1.3 (Exemples de polynômes pairs et impairs)

1. Considérons le polynôme suivant

$$P(X) = 4X^4 + 3X^2 + 1.$$

Si nous remplaçons X par $-X$, nous obtenons

$$P(-X) = 4(-X)^4 + 3(-X)^2 + 1 = 4X^4 + 3X^2 + 1 = P(X).$$

Ainsi, P est un polynôme pair.

2. Considérons le polynôme suivant

$$Q(X) = 5X^3 - 2X.$$

Si nous remplaçons X par $-X$, nous obtenons

$$Q(-X) = 5(-X)^3 - 2(-X) = -5X^3 + 2X = -Q(X).$$

Ainsi, Q est un polynôme impair.

Comment caractériser la parité (respectivement l'imparité) d'un polynôme à l'aide de ces coefficients ?.

Théorème 5.1.1 (Caractérisation de polynômes pairs/impairs)

Soit $P = (a_n)_{n \in \mathbb{N}}$ un polynôme de $\mathbb{K}[X]$.

1) On dit que P est pair si et seulement si, tous ses coefficients d'indice impair sont nuls, on a donc

$$\forall n \in \mathbb{N} : a_{2n+1} = 0_{\mathbb{K}}.$$

En d'autres termes, un polynôme pair ne contient que des termes avec des puissances paires de X .

2) On dit que P est impair si et seulement si, tous ses coefficients d'indice pair sont nuls, on a donc

$$\forall n \in \mathbb{N} : a_{2n} = 0_{\mathbb{K}}.$$

En d'autres termes, un polynôme impair ne contient que des termes avec des puissances impaires de X .

Preuve. Si $P = \sum_{n=0}^{+\infty} a_n X^n$ un polynôme de $\mathbb{K}[X]$, on a alors

$$P(-X) = \sum_{n=0}^{+\infty} (-1_{\mathbb{K}})^n a_n X^n.$$

Donc par unicité de l'écriture d'un polynôme, on obtient

$$\begin{aligned} P \text{ est pair} &\iff P(X) - P(-X) = 0_{\mathbb{K}[X]} \iff \sum_{n=0}^{+\infty} a_n (1_{\mathbb{K}} - (-1_{\mathbb{K}})^n) X^n = 0_{\mathbb{K}[X]} \\ &\iff \forall n \in \mathbb{N} : a_n (1_{\mathbb{K}} - (-1_{\mathbb{K}})^n) = 0_{\mathbb{K}} \\ &\iff \forall n \in \mathbb{N} : a_{2n+1} = 0_{\mathbb{K}}. \end{aligned}$$

De même

$$\begin{aligned} P \text{ est impair} &\iff P(X) + P(-X) = 0_{\mathbb{K}[X]} \iff \sum_{n=0}^{+\infty} a_n (1_{\mathbb{K}} + (-1_{\mathbb{K}})^n) X^n = 0_{\mathbb{K}[X]} \\ &\iff \forall n \in \mathbb{N} : a_n (1_{\mathbb{K}} + (-1_{\mathbb{K}})^n) = 0_{\mathbb{K}} \\ &\iff \forall n \in \mathbb{N} : a_{2n} = 0_{\mathbb{K}}. \end{aligned}$$

□

Exemple 5.1.4 $P_1 = X^5 + X$ de $\mathbb{R}[X]$ est impair et $P_2 = X^8 + 4X^4 + 3X^2$ de $\mathbb{R}[X]$ est pair. Par contre, $X^3 + X^2$ n'est ni pair, ni impair.

5.2 Opérations sur les polynômes et structures

Dans l'ensemble $\mathbb{K}[X]$, les opérations principales sont effectivement l'addition et les multiplications interne (ou produit de Polynômes) et externe (ou scalaire). Ces opérations sont essentielles pour travailler dans l'anneau de polynômes $\mathbb{K}[X]$ et sont utilisées dans de nombreuses applications mathématiques et algébriques.

5.2.1 Addition de deux polynômes

Pour additionner un polynôme à un autre, il faut additionner chacun des termes semblables du second polynôme à ceux du premier et réduire l'expression algébrique obtenue. On obtient alors un nouveau polynôme correspondant à la somme recherchée.

Définition 5.2.1 Soient $P = (a_n)_{n \in \mathbb{N}}$ et $Q = (b_n)_{n \in \mathbb{N}}$ deux polynômes de $\mathbb{K}[X]$. On appelle somme de P et Q (ou addition de P et Q) le polynôme de $\mathbb{K}[X]$ noté $P + Q$ dont les coefficients sont

$$c_n = a_n + b_n, \forall n \in \mathbb{N}.$$

Autrement dit, la suite des coefficients de $P+Q$ est la somme de leurs suites de coefficients respectives. On a ainsi défini une première loi (de composition) interne sur $\mathbb{K}[X]$ notée $+$.

Remarque 5.2.1 Si $P = \sum_{k=0}^n a_k X^k = \sum_{k=0}^{+\infty} a_k X^k$ et $Q = \sum_{k=0}^m b_k X^k = \sum_{k=0}^{+\infty} b_k X^k$, alors

$$P + Q = \sum_{k=0}^{\max(m,n)} (a_k + b_k) X^k = \sum_{k=0}^{+\infty} (a_k + b_k) X^k.$$

Exemple 5.2.1

1. Soient $P = (1, 1, 1, 0, \dots)$ et $Q = (0, 2, 3, -1, \dots)$ deux polynômes de $\mathbb{K}[X]$. On a

$$P + Q = (1, 1, 1, 0, \dots) + (0, 2, 3, -1, \dots) = (1, 3, 4, -1, 0, 0, \dots).$$

2. Si on a $P = 5X^2 - 3$ et $Q = 2X^3 - 5X^2 + 7X + 1$ deux polynômes de $\mathbb{K}[X]$, alors on a

$$P + Q = (0 + 2) X^3 + (5 - 5) X^2 + (0 + 7) X + (-3 + 1) = 2X^3 + 7X - 2.$$

Remarque 5.2.2 Remarquons que l'addition $+$ est une loi de composition interne sur $\mathbb{K}[X]$. En effet, soient $P = (a_n)_{n \in \mathbb{N}}$ et $Q = (b_n)_{n \in \mathbb{N}}$ deux polynômes de $\mathbb{K}[X]$, alors il existe $N_1, N_2 \in \mathbb{N}$ tels que

$$\forall n \in \mathbb{N} : n > N_1 \implies a_n = 0_{\mathbb{K}}.$$

$$\forall n \in \mathbb{N} : n > N_2 \implies b_n = 0_{\mathbb{K}},$$

donc, en posant $N_3 = \max(N_1, N_2) \in \mathbb{N}$, on a

$$\forall n \in \mathbb{N} : n > N_3 \implies a_n = b_n = 0_{\mathbb{K}},$$

par suite

$$\forall n \in \mathbb{N} : n > N_3 \implies a_n + b_n = 0_{\mathbb{K}}.$$

Il s'ensuit que $P + Q$ définie ci-dessus est bien un polynôme ($P + Q \in \mathbb{K}[X]$). Cela explique en quoi la loi $+$ est bien interne dans $\mathbb{K}[X]$. Ceci montre que $\mathbb{K}[X]$ est une partie de $\mathbb{K}^{\mathbb{N}}$ stable pour $+$.

5.2.1.1 Structure de groupe commutatif sur $\mathbb{K}[X]$

L'ensemble des polynômes $\mathbb{K}[X]$, muni de l'addition de polynômes, forme un groupe additif commutatif (ou groupe abélien). Les éléments du groupe sont les polynômes, l'opération de groupe est l'addition, l'élément neutre est le polynôme nul, et chaque polynôme possède un inverse additif.

Proposition 5.2.1 L'ensemble $\mathbb{K}[X]$ muni de loi $+$ possède une structure de groupe commutatif.

Preuve. L'addition définie sur l'ensemble des polynômes de $\mathbb{K}[X]$ est une loi de composition interne sur $\mathbb{K}[X]$ puisque l'addition $+\mathbb{K}$ définie sur le corps $\mathbb{K}$ est elle-même une loi de composition interne sur $\mathbb{K}$. De plus, l'addition des polynômes possède les propriétés suivantes (qui se déduisent des propriétés de l'addition sur $\mathbb{K}$). Ainsi,

i) Elle est associative, c'est-à-dire que pour tous $P, Q, R \in \mathbb{K}[X]$, on a

$$(P + Q) + R = P + (Q + R).$$

En effet, on note

$$P = \sum_{n=0}^{+\infty} a_n X^n, Q = \sum_{n=0}^{+\infty} b_n X^n, R = \sum_{n=0}^{+\infty} c_n X^n \in \mathbb{K}[X].$$

Alors

$$\begin{aligned}
(P + Q) + R &= \left(\sum_{n=0}^{+\infty} a_n X^n + \sum_{n=0}^{+\infty} b_n X^n \right) + \sum_{n=0}^{+\infty} c_n X^n = \sum_{n=0}^{+\infty} (a_n + b_n) X^n + \sum_{n=0}^{+\infty} c_n X^n \\
&= \sum_{n=0}^{+\infty} ((a_n + b_n) + c_n) X^n = \sum_{n=0}^{+\infty} (a_n + (b_n + c_n)) X^n \\
&= \sum_{n=0}^{+\infty} a_n X^n + \sum_{n=0}^{+\infty} (b_n + c_n) X^n \\
&= \sum_{n=0}^{+\infty} a_n X^n + \left(\sum_{n=0}^{+\infty} b_n X^n + \sum_{n=0}^{+\infty} c_n X^n \right) \\
&= P + (Q + R).
\end{aligned}$$

ii) Elle admet un élément neutre dans $\mathbb{K}[X]$. C'est le polynôme $0_{\mathbb{K}[X]}$ car

$$\forall P \in \mathbb{K}[X] : P + 0_{\mathbb{K}[X]} = 0_{\mathbb{K}[X]} + P = P.$$

iii) Tout polynôme $P = (a_n)_{n \in \mathbb{N}} \in \mathbb{K}[X]$ admet un symétrique dans $\mathbb{K}[X]$ qui est le polynôme $(-P) = (-a_n)_{n \in \mathbb{N}} \in \mathbb{K}[X]$. En effet, on vérifie que l'on a

$$\forall P \in \mathbb{K}[X] : P + (-P) = (-P) + P = 0_{\mathbb{K}[X]}.$$

iiii) Elle est commutative, pour tous $P = (a_n)_{n \in \mathbb{N}}$ et $Q = (b_n)_{n \in \mathbb{N}} \in \mathbb{K}[X]$,

$$P + Q = (a_n + b_n)_{n \in \mathbb{N}} = (b_n + a_n)_{n \in \mathbb{N}} = (b_n)_{n \in \mathbb{N}} + (a_n)_{n \in \mathbb{N}} = Q + P.$$

□

5.2.2 Produit de deux polynômes

Dans le contexte de la multiplication ou produit de polynômes, il existe effectivement deux types de multiplications distinctes : la première est la multiplication externe ou multiplication scalaire. C'est un exemple de multiplication d'un vecteur par un scalaire. La deuxième est la multiplication interne ou produit de polynômes.

Considérons deux monômes $P = a_m X^m$ et $Q = b_k X^k$. Si l'on interprète ces deux monômes comme des fonctions de la variable réelle ou complexe X , il est naturel de définir le produit de P par Q comme étant le monôme $PQ = a_m b_k X^{m+k}$. Plus généralement, on définit le produit de deux polynômes de la façon suivante.

Définition 5.2.2

1. Soient $P = (a_n)_{n \in \mathbb{N}}$ et $Q = (b_n)_{n \in \mathbb{N}}$ deux polynômes de $\mathbb{K}[X]$. On appelle produit de P et Q le polynôme de $\mathbb{K}[X]$, noté $P \cdot Q$ (ou plus simplement PQ) dont le coefficient d'indice n est défini par

$$\forall n \in \mathbb{N} : c_n = \sum_{k=0}^n a_k b_{n-k} = \sum_{i+j=n} a_i b_j.$$

Remarque 5.2.3

1. La définition précédente peut aussi s'écrire

$$\forall n \in \mathbb{N} : c_n = \sum_{k=0}^n a_k b_{n-k} = \sum_{k=0}^n a_{n-k} b_k.$$

2. Attention la multiplication définie dans $\mathbb{K}[X]$ n'est pas la multiplication terme à terme des suites définie dans $\mathbb{K}^{\mathbb{N}}$.

Exemple 5.2.2 Le polynôme P noté $3X^2 + 2X + 5$ de $\mathbb{R}[X]$ est formellement la suite $(a_n)_{n \in \mathbb{N}}$ telle que $a_0 = 5, a_1 = 2, a_2 = 3$ et $a_n = 0$ pour tout $n \geq 3$, c'est-à-dire la suite $(5, 2, 3, 0, 0, 0, \dots)$. Le polynôme Q noté $X^3 - 2X$ est la suite $(b_n)_{n \in \mathbb{N}}$ dont les premiers termes sont $(0, -2, 0, 1, 0, 0, \dots)$. Le polynôme produit PQ , selon la définition, correspond à la suite $(c_n)_{n \in \mathbb{N}}$ telle que

$$c_n = \sum_{k=0}^n a_k b_{n-k}.$$

Telles que

$$\begin{aligned} c_0 &= a_0 \cdot b_0 = 5 \cdot 0 = 0 \\ c_1 &= a_0 \cdot b_1 + a_1 \cdot b_0 = 5 \cdot (-2) + 2 \cdot 0 = -10 \\ c_2 &= a_0 \cdot b_2 + a_1 \cdot b_1 + a_2 \cdot b_0 = 5 \cdot 0 + 2 \cdot (-2) + 3 \cdot 0 = -4 \\ c_3 &= a_0 \cdot b_3 + a_1 \cdot b_2 + a_2 \cdot b_1 + a_3 \cdot b_0 = 5 \cdot 1 + 2 \cdot 0 + 3 \cdot (-2) + 0 \cdot 0 = -1 \\ c_4 &= a_0 \cdot b_4 + a_1 \cdot b_3 + a_2 \cdot b_2 + a_3 \cdot b_1 + a_4 \cdot b_0 = 5 \cdot 0 + 2 \cdot 1 + 3 \cdot 0 + 0 \cdot (-2) + 0 \cdot 0 = 2 \\ c_5 &= a_0 \cdot b_5 + a_1 \cdot b_4 + a_2 \cdot b_3 + a_3 \cdot b_2 + a_4 \cdot b_1 + a_5 \cdot b_0 = 5 \cdot 0 + 2 \cdot 0 + 3 \cdot 1 + 0 \cdot 0 + 0 \cdot (-2) + 0 \cdot 0 = 3 \\ &\vdots \\ c_n &= 0 \text{ pour tout } n \geq 6, \end{aligned}$$

et les coefficients suivants sont tous nuls, donc PQ est ce que l'on note habituellement $3X^5 + 2X^4 - X^3 - 4X^2 - 10X$. On peut vérifier qu'en calculant $(3X^2 + 2X + 5)(X^3 - 2X)$ avec les règles usuelles, on obtient le même résultat.

Le produit de deux polynômes est équivalent au polynôme obtenu en multipliant chaque monôme de l'un par chaque monôme de l'autre et en faisant la somme des monômes obtenus. Le produit de plusieurs monômes s'obtient en effectuant le produit des deux premiers, puis celui du polynôme obtenu par le troisième et ainsi de suite.

Définition 5.2.3 (Puissance d'un Polynôme)

1. Soient $P = (a_n)_{n \in \mathbb{N}}$ un polynôme de $\mathbb{K}[X]$ et m un entier positif. La puissance m -ième de P est définie par

$$P^m = \underbrace{P \cdot P \cdots P}_{m \text{ fois}}.$$

On peut alors observer que X^m désigne le polynôme dont le coefficient de degré m est 1 et les autres sont 0.

2. Cas particuliers,

$$\begin{cases} P^0 = 1_{\mathbb{K}[X]} \text{ (polynôme constant)} \\ P^1 = P. \end{cases}$$

Exemple 5.2.3 Prenons le polynôme $P = X + 2$ de $\mathbb{R}[X]$, calculons P^3 . On a

$$P^3 = P \cdot P \cdot P = (X + 2)(X + 2)(X + 2) = 8 + 12X + 6X^2 + X^3.$$

5.2.2.1 Structure d'anneau sur $\mathbb{K}[X]$

L'ensemble des polynômes à coefficients dans un corps $\mathbb{K}$, noté $\mathbb{K}[X]$, possède une structure d'anneau. Cette structure combine deux opérations fondamentales : l'addition et la multiplication des polynômes, tout en respectant certaines propriétés algébriques.

Proposition 5.2.2 *L'ensemble $(\mathbb{K}[X], +, \cdot)$ muni de l'addition $+$ et de la multiplication $\cdot$ des polynômes est un anneau unitaire (commutatif si $\mathbb{K}$ est un corps commutatif).*

Preuve.

1. Nous avons déjà vérifié que l'ensemble $\mathbb{K}[X]$ muni de loi $+$ possède une structure de groupe commutatif.

2. Maintenant, on vérifie que le produit de deux polynômes est bien un polynôme (la multiplication $\cdot$ est une loi interne sur $\mathbb{K}[X]$). En effet, soient $P = (a_n)_{n \in \mathbb{N}}$ et $Q = (b_n)_{n \in \mathbb{N}}$ deux polynômes de $\mathbb{K}[X]$, alors

(i). Si $P = 0_{\mathbb{K}[X]}$ ou $Q = 0_{\mathbb{K}[X]}$, alors $P \cdot Q = 0_{\mathbb{K}[X]} \in \mathbb{K}[X]$.

(ii). Si $P \neq 0_{\mathbb{K}[X]}$ et $Q \neq 0_{\mathbb{K}[X]}$, alors il existe deux entiers naturels n_1 et n_2 tels que

$$\forall n \in \mathbb{N} : n > n_1 \implies a_n = 0_{\mathbb{K}},$$

$$\forall n \in \mathbb{N} : n > n_2 \implies b_n = 0_{\mathbb{K}}.$$

Posons $n_3 = n_1 + n_2$. Alors pour tout $n > n_3$, on a

$$c_n = \sum_{k=0}^n a_k b_{n-k} = \sum_{k=0}^{n_3-1} a_k b_{n-k} + \sum_{k=n_3}^n a_k b_{n-k}.$$

Par suite, pour tout $k \in \{0, 1, \dots, n\}$, si $k > n_1$, on a

$$a_k = 0_{\mathbb{K}},$$

et si $k \leq n_1$, alors

$$n - k \geq n - n_1 > n_3 - n_1 = n_2,$$

donc

$$b_{n-k} = 0_{\mathbb{K}}.$$

Par conséquent

$$c_n = \sum_{k=0}^{n_3-1} a_k 0_{\mathbb{K}} + \sum_{k=n_3}^n 0_{\mathbb{K}} b_{n-k} = 0_{\mathbb{K}}.$$

ceci prouve que $P \cdot Q = (c_n)_{n \in \mathbb{N}} \in \mathbb{K}[X]$.

3. Il reste alors à établir que la multiplication des polynômes possède les propriétés suivantes : commutative, associative, distributive par rapport à l'addition et qu'elle admet un élément neutre.

(j) Elle est commutative si $\mathbb{K}$ est un corps commutatif : Pour tous $P = (a_n)_{n \in \mathbb{N}}$, $Q = (b_n)_{n \in \mathbb{N}} \in \mathbb{K}[X]$,

$$P \cdot Q = (c_n)_n,$$

telle que

$$\forall n \in \mathbb{N} : c_n = \sum_{k=0}^n a_k b_{n-k} = \sum_{i+j=n} a_i b_j.$$

En utilisant le fait que la loi $\cdot$ de $\mathbb{K}$ est commutative, on obtient

$$\forall n \in \mathbb{N} : c_n = \sum_{j+i=n} b_j a_i = \sum_{k=0}^n b_k a_{n-k},$$

donc

$$P \cdot Q = Q \cdot P.$$

(jj) Elle est associative : Pour tous $P, Q, R \in \mathbb{K}[X]$,

$$(P \cdot Q) \cdot R = P \cdot (Q \cdot R).$$

En effet, on note

$$P = (p_n)_{n \in \mathbb{N}} \in \mathbb{K}[X], Q = (q_n)_{n \in \mathbb{N}} \in \mathbb{K}[X], R = (r_n)_{n \in \mathbb{N}} \in \mathbb{K}[X].$$

Alors

$$P \cdot Q = (d_n)_{n \in \mathbb{N}},$$

où

$$\forall n \in \mathbb{N} : d_n = \sum_{k=0}^n p_k q_{n-k}.$$

Puis

$$(P \cdot Q) \cdot R = (e_n)_{n \in \mathbb{N}},$$

où

$$\forall n \in \mathbb{N} : e_n = \sum_{k=0}^n d_k r_{n-k}.$$

et $(Q \cdot R) = (f_n)_{n \in \mathbb{N}}$, où

$$\forall n \in \mathbb{N} : f_n = \sum_{k=0}^n q_k r_{n-k}.$$

Puis $P \cdot (Q \cdot R) = (g_n)_{n \in \mathbb{N}}$

$$\forall n \in \mathbb{N} : g_n = \sum_{k=0}^n p_k f_{n-k}.$$

On a, pour tout n de $\mathbb{N}$,

$$\begin{aligned} g_n &= \sum_{i+m=n} p_i f_m = \sum_{i+m=n} p_i \left(\sum_{j+k=m} q_j r_k \right) \\ &= \sum_{i+(j+k)=n} p_i (q_j r_k) = \sum_{(i+j)+k=n} (p_i q_j) r_k \\ &= \sum_{s+k=n} \left(\sum_{i+j=s} p_i q_j \right) r_k \\ &= \sum_{s+k=n} d_s r_k = e_n. \end{aligned}$$

Si on utilise la distributivité de la multiplication par rapport à l'addition dans $\mathbb{K}$ et l'associativité de ces lois dans $\mathbb{K}$, on trouve les deux polynômes $(P \cdot Q) \cdot R$ et $P \cdot (Q \cdot R)$ sont égaux. Ceci prouve

$$(P \cdot Q) \cdot R = P \cdot (Q \cdot R).$$

(jjj) Elle est distributive (à gauche et à droite) par rapport à l'addition : Pour tous $P, Q, R \in \mathbb{K}[X]$,

$$P \cdot (Q + R) = (P \cdot Q) + (P \cdot R),$$

et

$$(Q + R) \cdot P = (Q \cdot P) + (R \cdot P).$$

En effet, soient

$$P = (p_n)_{n \in \mathbb{N}} \in \mathbb{K}[X], Q = (q_n)_{n \in \mathbb{N}} \in \mathbb{K}[X], R = (r_n)_{n \in \mathbb{N}} \in \mathbb{K}[X].$$

On note

$$D = (d_n)_{n \in \mathbb{N}} = P \cdot (Q + R) \text{ et } E = (e_n)_{n \in \mathbb{N}} = (P \cdot Q) + (P \cdot R).$$

Alors pour tout n de $\mathbb{N}$,

$$d_n = \sum_{k=0}^n p_k (q_{n-k} + r_{n-k}) = \sum_{k=0}^n p_k q_{n-k} + \sum_{k=0}^n p_k r_{n-k} = e_n.$$

Donc

$$P \cdot (Q + R) = (P \cdot Q) + (P \cdot R).$$

Comme la multiplication est commutative, on en déduit que

$$(Q + R) \cdot P = P \cdot (Q + R) = (P \cdot Q) + (P \cdot R) = (Q \cdot P) + (R \cdot P)$$

(jjjj) Soit $1_{\mathbb{K}[X]} = (1_{\mathbb{K}}, 0_{\mathbb{K}}, 0_{\mathbb{K}}, \dots, 0_{\mathbb{K}}, \dots)$ le polynôme constant dont tous les coefficients sont nuls sauf a_0 qui vaut 1. On vérifie facilement que pour tout $P \in \mathbb{K}[X]$

$$P \cdot 1_{\mathbb{K}[X]} = 1_{\mathbb{K}[X]} \cdot P = P.$$

donc le polynôme $1_{\mathbb{K}[X]}$ est l'élément unité de l'anneau $\mathbb{K}[X]$. □

5.2.3 Multiplication d'un polynôme par un scalaire

La multiplication d'un polynôme par un scalaire dans l'anneau des polynômes $\mathbb{K}[X]$ est une opération simple mais essentielle. Cette opération consiste à multiplier chaque coefficient du polynôme par un élément du corps $\mathbb{K}$, qui est le scalaire.

Définition 5.2.4 (*Proposition*)

Soient $P = (a_n)_{n \in \mathbb{N}}$ un polynôme de $\mathbb{K}[X]$, et λ un scalaire de $\mathbb{K}$. On définit alors le polynôme noté λP de $\mathbb{K}[X]$ par

$$\lambda P = (\lambda a_n)_{n \in \mathbb{N}}.$$

Preuve. Il s'agit de montrer que la suite $(\lambda a_n)_{n \in \mathbb{N}}$ est presque nulle. En utilisant le fait qu'il existe un entier naturel $N \in \mathbb{N}$ tels que

$$\forall n \in \mathbb{N} : n > N \implies a_n = 0_{\mathbb{K}}.$$

Par suite

$$\forall n \in \mathbb{N} : n > N \implies \lambda a_n = 0_{\mathbb{K}}.$$

Ceci prouve que $\lambda P = (\lambda a_n)_{n \in \mathbb{N}} \in \mathbb{K}[X]$. □

La multiplication d'un polynôme de $\mathbb{K}[X]$ par un élément de $\mathbb{K}$ ne définit pas une loi interne sur $\mathbb{K}[X]$ mais une loi externe sur $\mathbb{K}[X]$. On l'appelle produit externe. Alors l'ensemble $\mathbb{K}[X]$ est stable par la loi externe à domaine d'opérateurs $\mathbb{K}$ sur $\mathbb{K}[X]$. Cette loi possède les propriétés suivantes.

Proposition 5.2.3 *L'anneau $(\mathbb{K}[X], +, \cdot)$ vérifie les propriétés supplémentaires suivantes pour tout $(\alpha, \beta) \in \mathbb{K}^2$ et $(P, Q) \in (\mathbb{K}[X])^2$:*

1. $(\alpha + \beta) \cdot P = \alpha \cdot P + \beta \cdot P$.
2. $\alpha \cdot (P + Q) = \alpha \cdot P + \alpha \cdot Q$.
3. $(\alpha \cdot \beta) \cdot P = \alpha \cdot (\beta \cdot P)$.
4. $1_{\mathbb{K}} \cdot P = P$.
5. $\alpha \cdot (P \cdot Q) = (\alpha \cdot P) \cdot Q$.

Preuve. On note

$$P = \sum_{n=0}^{+\infty} a_n X^n \in \mathbb{K}[X], Q = \sum_{n=0}^{+\infty} b_n X^n \in \mathbb{K}[X].$$

1. Soient $(\alpha, \beta) \in \mathbb{K}^2$, Alors

$$\begin{aligned} (\alpha + \beta) \cdot P &= (\alpha + \beta) \cdot \sum_{n=0}^{+\infty} a_n X^n = \sum_{n=0}^{+\infty} (\alpha + \beta) a_n X^n = \sum_{n=0}^{+\infty} (\alpha a_n + \beta a_n) X^n \\ &= \alpha \cdot \sum_{n=0}^{+\infty} a_n X^n + \beta \cdot \sum_{n=0}^{+\infty} a_n X^n = \alpha \cdot P + \beta \cdot P. \end{aligned}$$

2. Soit $\alpha \in \mathbb{K}$, alors

$$\begin{aligned} \alpha \cdot (P + Q) &= \alpha \cdot \left(\sum_{n=0}^{+\infty} a_n X^n + \sum_{n=0}^{+\infty} b_n X^n \right) = \alpha \cdot \sum_{n=0}^{+\infty} (a_n + b_n) X^n \\ &= \sum_{n=0}^{+\infty} (\alpha a_n + \alpha b_n) X^n = \alpha \cdot \sum_{n=0}^{+\infty} a_n X^n + \alpha \cdot \sum_{n=0}^{+\infty} b_n X^n = \alpha \cdot P + \alpha \cdot Q. \end{aligned}$$

3. Soient $(\alpha, \beta) \in \mathbb{K}^2$, alors

$$\begin{aligned} (\alpha \cdot \beta) \cdot P &= (\alpha \cdot \beta) \cdot \left(\sum_{n=0}^{+\infty} a_n X^n \right) = \sum_{n=0}^{+\infty} (\alpha \cdot \beta) \cdot a_n X^n \\ &= \sum_{n=0}^{+\infty} \alpha \cdot (\beta \cdot a_n) X^n = \alpha \cdot \sum_{n=0}^{+\infty} (\beta \cdot a_n) X^n \\ &= \alpha \cdot (\beta \cdot P). \end{aligned}$$

4. Il est clair que

$$1_{\mathbb{K}} \cdot P = 1_{\mathbb{K}} \cdot \sum_{n=0}^{+\infty} a_n X^n = \sum_{n=0}^{+\infty} (1_{\mathbb{K}} \cdot a_n) X^n = \sum_{n=0}^{+\infty} a_n X^n = P.$$

5. Soit $\alpha \in \mathbb{K}$, alors

$$\alpha \cdot (P \cdot Q) = \alpha \cdot \left(\sum_{n=0}^{+\infty} a_n X^n \cdot \sum_{n=0}^{+\infty} b_n X^n \right) = \alpha \cdot \sum_{n=0}^{+\infty} c_n X^n.$$

Telle que

$$\forall n \in \mathbb{N} : c_n = \sum_{k=0}^n a_k b_{n-k} = \sum_{i+j=n} a_i b_j.$$

Par suite

$$\alpha \cdot c_n = \sum_{k=0}^n \alpha \cdot (a_k \cdot b_{n-k}) = \sum_{k=0}^n (\alpha \cdot a_k) \cdot b_{n-k} = \sum_{k=0}^n a'_k \cdot b_{n-k},$$

où

$$\forall n \in \mathbb{N} : a'_n = \alpha \cdot a_n,$$

telle que

$$\alpha \cdot P = (a'_n)_{n \in \mathbb{N}}.$$

Par suite

$$\alpha \cdot (P \cdot Q) = (\alpha \cdot P) \cdot Q.$$

□

Remarque 5.2.4 Nous serons amenés par la suite à additionner des degrés de polynômes. Comme l'application $\deg$ est à valeurs dans $\mathbb{N} \cup \{-\infty\}$, il faut étendre la définition de l'addition. On adopte la convention suivante pour $n \in \mathbb{N} \cup \{-\infty\}$:

$$-\infty + n = -\infty.$$

5.3 Degré d'un polynôme et Intégrité de $\mathbb{K}[X]$

Le degré d'un polynôme est une notion clé en algèbre qui indique le terme de plus haut degré avec un coefficient non nul dans un polynôme. Le degré d'un polynôme fournit des informations sur la taille du polynôme en termes de la plus haute puissance de la variable présente dans l'expression.

Soit $P = (a_n)_{n \in \mathbb{N}} \in \mathbb{K}[X]$ un polynôme non nul. L'ensemble $\{n \in \mathbb{N} : a_n \neq 0_{\mathbb{K}}\}$ est une partie non vide de $\mathbb{N}$, majorée puisque la suite $(a_n)_{n \in \mathbb{N}}$ est nulle à partir d'un certain rang, cet ensemble admet donc un plus grand élément (un maximum). Cela nous conduit à la définition suivante.

Définition 5.3.1 (Degré d'un polynôme, coefficient dominant, polynôme unitaire)

Soit $P = (a_n)_{n \in \mathbb{N}} \in \mathbb{K}[X]$ un polynôme non nul.

1. Le plus grand indice n pour lequel $a_n \neq 0_{\mathbb{K}}$ est appelé le degré de P et noté $\deg(P)$. On a donc

$$\deg(P) = \max \{n \in \mathbb{N} : a_n \neq 0_{\mathbb{K}}\}.$$

2. Le coefficient $a_{\deg(P)}$ de P est appelé coefficient dominant de P et on dit que $a_{\deg(P)}X^{\deg(P)}$ est le monôme dominant de P .

3. Si le coefficient dominant de P vaut 1, alors on dit que P est unitaire ou encore normalisé.

Remarque 5.3.1

1. On adopte parfois la convention $\deg(0_{\mathbb{K}[X]}) = -\infty$. Cela permet de simplifier certains résultats sur le degré et pour éviter certaines contradictions dans les propriétés algébriques. Par exemple P est un polynôme constant si et si $P = 0_{\mathbb{K}[X]}$ ou $(P \neq 0_{\mathbb{K}[X]}$ et $\deg(0_{\mathbb{K}[X]}) = 0$) devient P est un polynôme constant si et si $\deg(0_{\mathbb{K}[X]}) \leq 0$.

2. Le coefficient dominant de P est le terme qui va dominer le comportement du polynôme pour des valeurs de X très grandes ou très petites.

Exemple 5.3.1

1. $P = 7X^4 - X^3 + 2X^2 - 3X - 5$ a pour degré 4 et coefficient dominant 7, tandis que $Q = X^3 - 4X^2 + 3X + 5$ est unitaire.

2. Le coefficient dominant de $3X^2 + 2X + 5$ est 3, ce polynôme n'est donc pas unitaire. Le coefficient dominant de $X^3 - X$ est 1 et ce polynôme est unitaire.

5.3.1 Propriétés fondamentales des degrés des polynômes

La proposition suivante rassemble les principales propriétés du degré.

Proposition 5.3.1 (*Propriétés des degrés des polynômes*) Soient $P, Q \in \mathbb{K}[X]$ et $\lambda \in \mathbb{K}$.

(a) **Degré d'une somme**

$$\deg(P + Q) \leq \max \{ \deg(P), \deg(Q) \}. \quad (5.1)$$

Cette inégalité est une égalité notamment quand $\deg(P) \neq \deg(Q)$, on a donc

$$\deg(P + Q) = \max \{ \deg(P), \deg(Q) \}. \quad (5.2)$$

(b) **Degré d'un produit**

$$\deg(PQ) = \deg(P) + \deg(Q), \quad (5.3)$$

où par convention $(-\infty + m = -\infty)$ pour tout $m \in \mathbb{N} \cup \{-\infty\}$. En particulier, pour $\lambda \neq 0_{\mathbb{K}}$

$$\deg(\lambda P) = \deg(P). \quad (5.4)$$

Preuve.

(a) Considérons deux polynômes $P = (a_n)_{n \in \mathbb{N}}$, $Q = (b_n)_{n \in \mathbb{N}} \in \mathbb{K}[X]$. Le résultat est évident lorsque P ou Q est nul. Supposons-les donc tous deux non nuls et notons p le degré de P et q celui de Q , ainsi que les coefficients a_p et b_q sont nécessairement non nuls. Supposons (sans perte de généralité) que $p \geq q$ et considérons les cas $p > q$ et $p = q$.

• Si $p > q$, alors

$$P + Q = (a_0 + b_0, \dots, a_q + b_q, a_{q+1}, \dots, a_p, 0, 0, \dots).$$

Par hypothèse, $a_p \neq 0_{\mathbb{K}}$. On en déduit alors que $\deg(P + Q) = p$, c'est-à-dire que

$$\deg(P + Q) = \deg(P) = \max \{ \deg(P), \deg(Q) \}.$$

• Si $p = q$, alors

$$P + Q = (a_0 + b_0, \dots, a_p + b_p, 0, 0, \dots).$$

Il existe deux possibilités : si $a_p + b_p = 0_{\mathbb{K}}$, alors

$$\deg(P + Q) < \deg(P) = p,$$

donc, l'inégalité (5.1) est vraie, ainsi l'inégalité stricte ayant lieu.

Par ailleurs, si $a_p + b_p \neq 0_{\mathbb{K}}$, alors

$$\deg(P + Q) = \deg(P) = \deg(Q),$$

donc l'inégalité (5.1) est encore vraie, ainsi l'égalité ayant lieu dans ce cas.

(b) Considérons deux polynômes $P = (a_n)_{n \in \mathbb{N}}$, $Q = (b_n)_{n \in \mathbb{N}} \in \mathbb{K}[X]$. On distingue les cas suivants :

• Si $P = 0_{\mathbb{K}[X]}$ ou $Q = 0_{\mathbb{K}[X]}$, alors $PQ = 0_{\mathbb{K}[X]}$ et

$$\deg(PQ) = \deg(0_{\mathbb{K}[X]}) = -\infty = \deg(P) + \deg(Q),$$

car, avec l'extension des opérations arithmétiques sur $\overline{\mathbb{R}}$, on a pour tout $n \in \mathbb{N} \cup \{-\infty\}$,

$$n + (-\infty) = -\infty + n = -\infty,$$

ainsi que

$$(-\infty) + (-\infty) = (-\infty).$$

• Sinon, posons $\deg(P) = p$ et $\deg(Q) = q$. On a donc $a_n = 0_{\mathbb{K}}$ pour tout entier n strictement supérieure à p , et $b_n = 0_{\mathbb{K}}$ pour tout entier n strictement supérieure à q . Soient $c_n, n \in \mathbb{N}$, les coefficients du polynôme PQ . Pour montrer que $\deg(PQ) = \deg(P) + \deg(Q)$, on doit vérifier d'une part que $c_{p+q} \neq 0_{\mathbb{K}}$, et d'autre part que $c_{p+q+l} = 0_{\mathbb{K}}$ pour tout $l \geq 1$. Commençons par calculer c_{p+q} . On vérifie que

$$\begin{aligned} c_{p+q} &= \sum_{k=0}^{p+q} a_k b_{p+q-k} = \sum_{k=0}^{p-1} a_k b_{p+q-k} + a_p b_q + \sum_{k=p+1}^{p+q} a_k b_{p+q-k} \\ &= \sum_{k=0}^{p-1} a_k \underbrace{(b_{p+q-k})}_{=0_{\mathbb{K}}} + a_p b_q + \sum_{k=p+1}^{p+q} \underbrace{(a_k)}_{=0_{\mathbb{K}}} b_{p+q-k} = a_p b_q. \end{aligned}$$

Car si $k \in \{0, p-1\}$, alors $p+q-k \geq q+1$ donc $b_{p+q-k} = 0_{\mathbb{K}}$, et si $k \geq p+1$, alors $a_k = 0_{\mathbb{K}}$. Par ailleurs, comme $a_p \neq 0_{\mathbb{K}}$ et $b_q \neq 0_{\mathbb{K}}$, forcément $c_{p+q} = a_p b_q \neq 0_{\mathbb{K}}$.

De même, on vérifie que

$$\begin{aligned} c_{p+q} &= \sum_{k=0}^{p+q+1} a_k b_{p+q-k} = a_0 b_{p+q+1} + a_1 b_{p+q} + \dots + a_{p-1} b_{q+2} + a_p b_{q+1} + a_{p+1} b_q + \dots + a_{p+q+1} b_0 \\ &= a_0 \underbrace{b_{p+q+1}}_{=0_{\mathbb{K}}} + a_1 \underbrace{b_{p+q}}_{=0_{\mathbb{K}}} + \dots + a_{p-1} \underbrace{b_{q+2}}_{=0_{\mathbb{K}}} + a_p \underbrace{b_{q+1}}_{=0_{\mathbb{K}}} + \underbrace{a_{p+1}}_{=0_{\mathbb{K}}} b_q + \dots + \underbrace{a_{p+q+1}}_{=0_{\mathbb{K}}} b_0 = 0_{\mathbb{K}}. \end{aligned}$$

De manière similaire, on peut vérifier que $c_{p+q+l} = 0_{\mathbb{K}}$ pour tout $l \geq 1$. Ce qui implique que

$$\deg(PQ) = \deg(P) + \deg(Q).$$

L'assertion (5.4) découle immédiatement de (5.3). En effet, soit $P = \sum_{k=0}^n a_k X^k$ de degré n et $\lambda \in \mathbb{K}$ un scalaire non nul. Par construction de λP , nous avons

$$\lambda P = \sum_{k=0}^n (\lambda a_k) X^k.$$

Comme $\lambda \neq 0_{\mathbb{K}}$ et $a_n \neq 0_{\mathbb{K}}$, il en résulte que $\lambda a_n \neq 0_{\mathbb{K}}$. Ainsi, le polynôme λP est de degré n . $\square$

Exemple 5.3.2

- Le degré de $(X^3 + X) + (X^2 + 1) = X^3 + X^2 + X + 1$ est $3 = \max(3, 2)$.
- Le degré de $(X^3 + X) - (X^3 + X^2) = -X^2 + X$ est $2 \leq \max(3, 3)$.
- Le degré de $(X^3 + X)(X^2 + 1) = X^5 + 2X^3 + X$ est $3 + 2 = 5$.
- Le degré de $(2(1 + 2X)) = (2 + 4X)$ est $1 = \deg(1 + 2X)$.

Notation 5.3.1 Pour tout $n \in \mathbb{N}$, on note $\mathbb{K}_n[X]$ l'ensemble des polynômes à coefficients dans $\mathbb{K}$ et de degré inférieur ou égal à n ,

$$\mathbb{K}_n[X] = \{P \in \mathbb{K}[X] : \deg(P) \leq n\}.$$

En d'autres termes, $\mathbb{K}_n[X]$ est constitué de tous les polynômes de la forme

$$P = a_0 + a_1 X + \dots + a_n X^n,$$

où $a_i \in \mathbb{K}$ pour tout i tel que $0 \leq i \leq n$.

Remarque 5.3.2

1. On remarque que, pour tout $n \in \mathbb{N}$, $0_{\mathbb{K}[X]} \in \mathbb{K}_n[X]$ puisque $(-\infty \leq n)$.
2. L'expression $\mathbb{K} = \mathbb{K}_0[X]$ indique que l'ensemble $\mathbb{K}_0[X]$ des polynômes de degré inférieur ou égal à 0 est exactement le corps $\mathbb{K}$ lui-même. Cela signifie que tout polynôme de $\mathbb{K}_0[X]$ est simplement un élément constant de $\mathbb{K}$, ce qui relie directement la structure des polynômes à celle du corps $\mathbb{K}$. On a donc la chaîne stricte d'inclusions suivante

$$\mathbb{K} = \mathbb{K}_0[X] \subset \mathbb{K}_1[X] \subset \mathbb{K}_2[X] \subset \dots \subset \mathbb{K}_n[X] \subset \dots \subset \mathbb{K}[X].$$

3. $\mathbb{K}_n[X]$ est un sous-groupe de $(\mathbb{K}[X], +)$.

L'intégrité de l'anneau des polynômes $\mathbb{K}[X]$ peut être démontrée en utilisant les propriétés du degré d'un polynôme. Plus précisément, on peut montrer que $\mathbb{K}[X]$ est un anneau intègre, c'est-à-dire un anneau dans lequel le produit de deux polynômes non nuls est non nul, en utilisant la propriété que le degré du produit de deux polynômes est la somme des degrés de ces polynômes.

Proposition 5.3.2 (Intégrité de $\mathbb{K}[X]$)

L'anneau $(\mathbb{K}[X], +, \cdot)$ est intègre. Formellement, cela signifie que

$$\forall P, Q \in \mathbb{K}[X] : P \cdot Q = 0_{\mathbb{K}[X]} \implies \begin{cases} P = 0_{\mathbb{K}[X]} \\ \text{ou} \\ Q = 0_{\mathbb{K}[X]}. \end{cases}$$

Preuve. Pour commencer, $\mathbb{K}[X]$ est un anneau non nul. Soient ensuite $P, Q \in \mathbb{K}[X]$. Si $P \cdot Q = 0_{\mathbb{K}[X]}$, alors

$$\deg(P) + \deg(Q) = \deg(PQ) = -\infty,$$

donc $\deg(P)$ ou $\deg(Q)$ vaut $-\infty$, autrement dit P ou Q est nul. □

Proposition 5.3.3 Les éléments inversibles de $\mathbb{K}[X]$ par rapport au produit interne sont les polynômes de degré 0, c'est-à-dire les polynômes constants non nuls (précisément ceux de $\mathbb{K}^*$).

Preuve. Il suffit de considérer le degré.

- Si P est un polynôme de degré 0, alors P est un polynôme constant non nul de la forme $P = \lambda$, le polynôme constant λ^{-1} est l'inverse de P .
- Réciproquement, si P est un polynôme inversible, c'est-à-dire tel qu'il existe un polynôme Q vérifiant $P \cdot Q = 1_{\mathbb{K}[X]}$, alors P et Q sont non nuls et l'on a

$$0 = \deg(1_{\mathbb{K}[X]}) = \deg(P \cdot Q) = \deg(P) + \deg(Q).$$

Puisque le degré d'un polynôme est un entier non négatif, la seule solution à cette équation est $\deg(P) = \deg(Q) = 0$. □

5.4 Composition des polynômes

En plus de l'addition, de la multiplication, et de la multiplication par un scalaire, il existe plusieurs autres opérations intéressantes que l'on peut effectuer sur les polynômes. Voici quelques-unes de ces opérations : la composition et la dérivation.

La composition des polynômes est souvent interprétée comme une opération de substitution. En d'autres termes, pour deux polynômes P et Q , la composition $P \circ Q$ revient à substituer Q à la variable X dans le polynôme P . On peut définir la composition $P \circ Q$ de P et Q de la façon suivante.

Définition 5.4.1 (Composition des polynômes)

Soient $P = \sum_{n=0}^{+\infty} a_n X^n, Q \in \mathbb{K}[X]$. On appelle composée de Q suivie de P le polynôme $P \circ Q$, qui est défini par

$$P \circ Q = \sum_{n=0}^{+\infty} a_n Q^n.$$

Remarque 5.4.1

- La composition $P \circ Q$ consiste donc à remplacer l'indéterminée X dans P par un polynôme Q et calculer le résultat avec les opérations d'addition et de multiplication des polynômes.
- Dans le cas particulier où $Q = X$, le polynôme $P(Q) = P(X)$ est donc égal à P , c'est pourquoi on utilise aussi bien P que $P(X)$ pour désigner ce dernier polynôme.
- Lorsque $Q = 0_{\mathbb{K}[X]}$, on a $P \circ Q = P(0_{\mathbb{K}[X]}) = a_0$ et si $P = 0_{\mathbb{K}[X]}$, alors $P \circ Q = 0_{\mathbb{K}[X]}(Q) = 0_{\mathbb{K}[X]}$.

Exemple 5.4.1 Pour $Q = X+1$, on a $P \circ Q = \sum_{n=0}^{+\infty} a_n (X+1)^n$, pour $Q = X^3$, on a $P \circ Q = \sum_{n=0}^{+\infty} a_n X^{3n}$.

Plus concrètement, si on a

$$P = 3X^2 - 2X + 1 \text{ et } Q = X - 1.$$

on obtient

$$P \circ Q = 3(X-1)^2 - 2(X-1) + 1 = 3X^2 - 8X + 6.$$

Si on a $Q = 2$, alors

$$P \circ Q = 3(2)^2 - 2(2) + 1 = 9.$$

Les propriétés de la composition sont rassemblées dans la proposition suivante.

Proposition 5.4.1 La composition des polynômes vérifie les propriétés suivantes :

1. Pour tous $\alpha, \beta \in \mathbb{K}$, et $P, Q, R \in \mathbb{K}[X]$, on a

$$(\alpha P + \beta Q) \circ R = \alpha (P \circ R) + \beta (Q \circ R).$$

En particulier, elle est distributive à droite par rapport à l'addition (mais pas à gauche)

$$\forall P, Q, R \in \mathbb{K}[X] : (P + Q) \circ R = P \circ R + Q \circ R.$$

2. Elle est distributive à droite par rapport à la multiplication :

$$\forall P, Q, R \in \mathbb{K}[X] : (P \cdot Q) \circ R = (P \circ R) \cdot (Q \circ R).$$

3. Elle est associative

$$\forall P, Q, R \in \mathbb{K}[X] : (P \circ Q) \circ R = P \circ (Q \circ R).$$

4. X est neutre pour $\circ$,

$$\forall P \in \mathbb{K}[X] : P \circ X = X \circ P = P.$$

Preuve.

1. On note

$$P = \sum_{n=0}^{+\infty} a_n X^n \in \mathbb{K}[X], Q = \sum_{n=0}^{+\infty} b_n X^n \in \mathbb{K}[X], R = \sum_{n=0}^{+\infty} c_n X^n \in \mathbb{K}[X].$$

Alors

$$(\alpha P + \beta Q) \circ R = \left(\sum_{n=0}^{+\infty} (\alpha a_n + \beta b_n) X^n \right) \circ R = \sum_{n=0}^{+\infty} (\alpha a_n + \beta b_n) R^n$$

$$\begin{aligned}
&= \alpha \sum_{n=0}^{+\infty} a_n R^n + \beta \sum_{n=0}^{+\infty} b_n R^n \\
&= \alpha \left[\left(\sum_{n=0}^{+\infty} a_n X^n \right) \circ R \right] + \beta \left[\left(\sum_{n=0}^{+\infty} b_n X^n \right) \circ R \right] \\
&= \alpha (P \circ R) + \beta (Q \circ R).
\end{aligned}$$

2. On a

$$(P \cdot Q) \circ R = \sum_{n=0}^{+\infty} d_n R^n$$

avec pour tout $n \in \mathbb{N}$,

$$d_n = \sum_{k=0}^n a_k b_{n-k}.$$

Par suite

$$\begin{aligned}
(P \cdot Q) \circ R &= \sum_{n=0}^{+\infty} \sum_{k=0}^n a_k b_{n-k} R^n = \sum_{k=0}^{+\infty} a_k R^k \sum_{n=k}^{+\infty} b_{n-k} R^{n-k} \\
&\stackrel{n-k=l}{=} \sum_{k=0}^{+\infty} a_k R^k \sum_{l=0}^{+\infty} b_l R^l \\
&= (P \circ R) \cdot (Q \circ R).
\end{aligned}$$

3. On a

$$(P \circ Q) \circ R = \left(\sum_{n=0}^{+\infty} a_n Q^n \right) \circ R.$$

En appliquant le résultat 1, on obtient

$$\left(\sum_{n=0}^{+\infty} a_n Q^n \right) \circ R = \sum_{n=0}^{+\infty} a_n (Q^n \circ R).$$

En appliquant le résultat 2, on obtient

$$\begin{aligned}
\sum_{n=0}^{+\infty} a_n (Q^n \circ R) &= \sum_{n=0}^{+\infty} a_n (Q(R))^n \\
&= P \circ (Q \circ R).
\end{aligned}$$

4. Il est clair que dans le cas où $Q = X$, le polynôme composé $P \circ Q = P \circ X = P(X)$ est égal à P . De même pour $X \circ P = P$. $\square$

Remarque 5.4.2

1. La composition de polynômes n'est pas commutative, à savoir $P \circ Q \neq Q \circ P$ en général. Par exemple

$$1 \circ (X + 1) = 1,$$

alors que

$$(X + 1) \circ 1 = 2.$$

On a aussi

$$1 \circ (X + 1) = 0 \text{ et } (X + 1) \circ 0 = 1.$$

2. La composition de polynômes n'est pas distributive à gauche par rapport à l'addition. Par exemple

$$1 \circ (X + 1) = 1,$$

alors que

$$1 \circ 1 + 1 \circ X = 2.$$

On a aussi

$$X^2 \circ (X + 1) = (X + 1)^2 \text{ et } X^2 \circ X + X^2 \circ 1 = X^2 + 1.$$

3. La composition de polynômes n'est pas distributive à gauche par rapport à la multiplication. Par exemple

$$(1 + X) \circ (X \cdot X) = (1 + X) \circ X^2 = 1 + X^2,$$

alors que

$$((1 + X) \circ X) \cdot ((1 + X) \circ X) = (1 + X) \cdot (1 + X) = (1 + X)^2.$$

Proposition 5.4.2 *Etant donné P et Q dans $\mathbb{K}[X]$, si P et Q sont non nuls ou si P est nul mais Q n'est pas constant, alors on a*

$$\deg(P \circ Q) = \deg(P) \cdot \deg(Q).$$

Preuve.

- Si P est nul et Q non constant, alors $P \circ Q$ est nul et $\deg(Q) \geq 1$, donc

$$\deg(P \circ Q) = -\infty = -\infty \deg(Q) = \deg(P) \deg(Q) \text{ (car } \deg(Q) \neq 0).$$

En effet, on étend la multiplication à $\mathbb{N}^* \cup \{-\infty\}$ en posant, si $n \in \mathbb{N}^*$, alors

$$n \times (-\infty) = -\infty \text{ et } (-\infty) \times n = -\infty.$$

- Si P et Q sont non nuls, on écrit $P = \sum_{k=0}^n a_k X^k$ avec $n = \deg(P)$. Ainsi, $a_n \neq 0_{\mathbb{K}}$. Par produit $\deg(Q^m) = m \deg(Q)$ pour tout $m \in \{0, 1, \dots, n\}$, alors

$$P \circ Q = \sum_{k=0}^n a_k Q^k,$$

et comme la suite $(\deg(Q^m))_{m \in \{0, 1, \dots, n\}}$ est croissante, alors $P \circ Q$ est la somme de $(n + 1)$ polynômes de degrés 2 à 2 différents. Finalement, par somme et (5.2), on obtient

$$\deg(P \circ Q) = \deg\left(\sum_{k=0}^n a_k Q^k\right) \stackrel{a_n \neq 0_{\mathbb{K}}}{=} \deg(a_n Q^n) = \deg(Q^n) = n \cdot \deg(Q) = \deg(P) \cdot \deg(Q).$$

□

5.5 Dérivation des polynômes

Une autre opération intéressante que l'on peut effectuer sur les polynômes est la dérivation des polynômes. La dérivation des polynômes est une opération algébrique qui consiste à calculer la dérivée d'un polynôme par rapport à sa variable principale. Cette dérivation est une application directe des règles de dérivation classiques utilisées en analyse, mais elle peut être effectuée dans un cadre purement algébrique.

Définition 5.5.1

1. Soit $P = \sum_{k=0}^n a_k X^k$ un polynôme de $\mathbb{K}[X]$, on appelle polynôme dérivé de P , et on note P' , le polynôme défini par

$$P' = \sum_{k=1}^n k a_k X^{k-1}.$$

Autrement dit, chaque coefficient a_k du polynôme original P est multiplié par l'exposant k de la variable X^k dans ce terme, puis on réduit cet exposant de 1 pour obtenir le terme correspondant dans P' .

2. On a également, après changement d'indice

$$P' = \sum_{k=0}^{n-1} (k+1) a_{k+1} X^k.$$

Exemple 5.5.1 Soit $P = 4X^3 - 5X^2 + 3X - 2$ de $\mathbb{R}[X]$, alors le polynôme dérivé P' est

$$P' = 12X^2 - 10X + 3.$$

Définition 5.5.2 (Polynômes dérivés successifs)

Soit P un polynôme de $\mathbb{K}[X]$. On définit les polynômes dérivés successifs de P en posant

$$\begin{cases} P^{(0)} = P, \\ \forall k \in \mathbb{N}^* : P^{(k)} = (P^{(k-1)})'. \end{cases}$$

On dit que $P^{(k)}$ est le polynôme dérivé k -ième de P . On note bien sûr P' et P'' plutôt que $P^{(1)}$ et $P^{(2)}$.

Exemple 5.5.2

• Pour $P = 8X^3 - 5X^2 + 3X + 1 \in \mathbb{R}[X]$, on a

$$P^{(1)} = 24X^2 - 10X + 3, P^{(2)} = 48X - 10, P^{(3)} = 48, \text{ et pour tout } n \geq 4, P^{(n)} = 0_{\mathbb{R}[X]}.$$

• Pour $0 \leq k \leq n$ entiers,

$$(X^n)^{(k)} = n(n-1)(n-2)\dots(n-k+1)X^{n-k} = \frac{n!}{(n-k)!} X^{n-k}. \quad (5.5)$$

En particulier $(X^n)^{(n)} = n!$ et $(X^n)^{(n+1)} = 0_{\mathbb{R}[X]}$.

Remarque 5.5.1 Soient $P \in \mathbb{K}[X]$ et $n \in \mathbb{N}$, si $\deg(P) = m$, alors

$$\begin{cases} \deg(P^{(n)}) = m - n, \forall n \leq m, \\ P^{(n)} = 0_{\mathbb{K}[X]}, \forall n > m. \end{cases}$$

Cette relation montre que chaque dérivation diminue le degré du polynôme de 1, jusqu'à ce que le degré atteigne zéro. Après m dérivées, le polynôme devient constant, et toute dérivée ultérieure est nulle.

Proposition 5.5.1 Soient $P \in \mathbb{K}[X]$ et $n \in \mathbb{N}$, alors

$$1. \quad \deg(P') = \begin{cases} \deg(P) - 1, & \text{si } \deg(P) \geq 1, \\ -\infty, & \text{si } \deg(P) \leq 0. \end{cases} \quad (5.6)$$

2.

$$\deg(P) \leq n \iff P^{(n+1)} = 0_{\mathbb{K}[X]}.$$

Preuve.

1. Soit $P \in \mathbb{K}[X]$, alors si $\deg(P) \geq 1$, on note $P = \sum_{k=0}^n a_k X^k$ avec $\deg(P) = n \geq 1$ donc $a_n \neq 0_{\mathbb{K}}$, alors

$$P' = \sum_{k=1}^n k a_k X^{k-1} = \sum_{j=0}^{n-1} (j+1) a_{j+1} X^j, \quad n a_n \neq 0_{\mathbb{K}},$$

donc

$$\deg(P') = n - 1.$$

Sinon, P est un polynôme constant, $P = a_0$ avec $a_0 \in \mathbb{K}$ et $P' = 0_{\mathbb{K}[X]}$, donc $\deg(P') = -\infty$.

2. • Supposons que $\deg(P) \leq n$, alors d'après (5.6), pour tout $k \in \mathbb{N}^*$, si $\deg(P) = k$ alors, on voit en itérant que $\deg(P^{(k)}) = k - k = 0$ donc $P^{(k+1)}$ est le polynôme nul. Par suite, si $\deg(P) \leq n$, alors $P^{(n+1)} = 0_{\mathbb{K}[X]}$.

• Réciproquement, par contraposée, si $\deg(P) \geq n + 1$, alors

$$\deg(P^{(n+1)}) = \deg(P) - (n + 1) \geq n + 1 - (n + 1) = 0.$$

Donc $P^{(n+1)}$ n'est pas le polynôme nul. □

La proposition suivante rassemble les principales propriétés des polynômes dérivés.

Proposition 5.5.2 Soient P, Q deux polynômes de $\mathbb{K}[X]$ et $\lambda \in \mathbb{K}$. On a

1. $(P + \lambda Q)' = P' + \lambda Q'$.
2. $(PQ)' = P'Q + PQ'$.
3. $(P \circ Q)' = Q' \cdot (P' \circ Q)$.

Preuve. Soient $P = \sum_{k=0}^{+\infty} p_k X^k$ et $Q = \sum_{k=0}^{+\infty} q_k X^k$ deux polynômes de $\mathbb{K}[X]$.

1. Par définition de la somme et du polynôme dérivé, on a pour tout $\lambda \in \mathbb{K}$,

$$\begin{aligned} (P + \lambda Q)' &= \left(\sum_{k=0}^{+\infty} (p_k + \lambda q_k) X^k \right)' = \sum_{k=1}^{+\infty} k (p_k + \lambda q_k) X^{k-1} = \\ &= \sum_{k=1}^{+\infty} k p_k X^{k-1} + \lambda \sum_{k=1}^{+\infty} k q_k X^{k-1} \\ &= P' + \lambda Q'. \end{aligned}$$

2. D'après la définition du produit sur $\mathbb{K}[X]$, on a

$$PQ = \left(\sum_{k=0}^{+\infty} p_k X^k \right) \left(\sum_{k=0}^{+\infty} q_k X^k \right) = \sum_{n=0}^{+\infty} \left(\sum_{k=0}^n p_k q_{n-k} \right) X^n.$$

Ainsi

$$(PQ)' = \sum_{n=1}^{+\infty} n \left(\sum_{k=0}^n p_k q_{n-k} \right) X^{n-1} = \sum_{n=1}^{+\infty} n p_0 q_n X^{n-1} + \sum_{k=1}^{+\infty} \left(\sum_{n=k}^{+\infty} n p_k q_{n-k} X^{n-1} \right)$$

$$\begin{aligned}
&= p_0 Q' + \sum_{k=1}^{+\infty} p_k \left(\sum_{n=k}^{+\infty} (n-k+k) q_{n-k} X^{n-1} \right) \\
&= p_0 Q' + \sum_{k=1}^{+\infty} p_k \left(\sum_{n=k}^{+\infty} (n-k) q_{n-k} X^{n-1} \right) + \sum_{k=1}^{+\infty} k p_k \left(\sum_{n=k}^{+\infty} q_{n-k} X^{n-1} \right) \\
&= p_0 Q' + \sum_{k=1}^{+\infty} p_k \left(\sum_{l=0}^{+\infty} l q_l X^{l-1+k} \right) + \sum_{k=1}^{+\infty} k p_k \left(\sum_{l=0}^{+\infty} q_l X^{l+k-1} \right) \\
&= p_0 Q' + \sum_{k=1}^{+\infty} p_k X^k \left(\sum_{l=1}^{+\infty} l q_l X^{l-1} \right) + \sum_{k=1}^{+\infty} k p_k X^{k-1} \left(\sum_{l=0}^{+\infty} q_l X^l \right) \\
&= p_0 Q' + \sum_{k=1}^{+\infty} p_k X^k Q' + \sum_{k=1}^{+\infty} k p_k X^{k-1} Q \\
&= P Q' + P' Q.
\end{aligned}$$

3. Comme $P \circ Q = \sum_{k=0}^{+\infty} p_k Q^k$, on déduit de la linéarité de la dérivation que

$$(P \circ Q)' = \sum_{k=0}^{+\infty} p_k (Q^k)' = \sum_{k=1}^{+\infty} k p_k Q' Q^{k-1} = Q' \sum_{k=1}^{+\infty} k p_k Q^{k-1} = Q' \cdot (P' \circ Q).$$

□

L'identité de Leibniz $(PQ)' = P'Q + PQ'$ se généralise à un ordre n quelconque de la manière suivante.

Corollaire 5.5.1 (*Formule de Leibniz*)

Soient P, Q deux polynômes de $\mathbb{K}[X]$ et $n \in \mathbb{N}$. Alors l'identité de Leibniz pour la dérivée d'un produit de ces deux polynômes à l'ordre n est donnée par

$$(PQ)^{(n)} = \sum_{k=0}^n C_n^k P^{(k)} Q^{(n-k)}. \quad (5.7)$$

La formule de Leibniz s'en déduit par récurrence sur n .

Exemple 5.5.3 On a $X^{2n} = X^n \cdot X^n$, alors en appliquant la formule (5.5), il s'en suit que

$$(X^n)^{(k)} = n(n-1)(n-2)\dots(n-k+1)X^{n-k} = \frac{n!}{(n-k)!} X^{n-k}.$$

$$\implies (X^{2n})^{(n)} = \frac{(2n)!}{n!} X^n.$$

Par ailleurs, de (5.7), et de (5.5), on peut déduire que

$$\begin{aligned}
(X^{2n})^{(n)} &= (X^n \cdot X^n)^{(n)} = \sum_{k=0}^n C_n^k (X^n)^{(k)} (X^n)^{(n-k)} = \sum_{k=0}^n C_n^k \frac{n!}{(n-k)!} X^{n-k} \cdot \frac{n!}{k!} X^k \\
&= n! \sum_{k=0}^n (C_n^k)^2 X^n.
\end{aligned}$$

Ceci entraîne que

$$\frac{(2n)!}{n!} X^n = n! \sum_{k=0}^n (C_n^k)^2 X^n.$$

Par conséquent,

$$\frac{(2n)!}{n!} = n! \sum_{k=0}^n (C_n^k)^2.$$

Ceci prouve que

$$\sum_{k=0}^n (C_n^k)^2 = C_{2n}^n.$$

5.6 Fonctions polynomiales

Pour définir la notion de racine d'un polynôme, il faut d'abord pouvoir « appliquer » un polynôme à un élément de $\mathbb{K}$, c'est-à-dire transformer un polynôme formel en une fonction polynomiale.

Définition 5.6.1 (*Fonction polynomiale*)

La fonction polynomiale (ou fonction polynôme) associée à un polynôme $P = a_0 + a_1X + a_2X^2 + \dots + a_nX^n$ dans $\mathbb{K}[X]$ est la fonction notée $\tilde{P}$ définie de $\mathbb{K}$ dans $\mathbb{K}$, associant à tout $x \in \mathbb{K}$ l'élément

$$\tilde{P}(x) = a_0 + a_1x + a_2x^2 + \dots + a_nx^n.$$

En particulier, la fonction polynomiale associée à un monôme est appelée fonction monôme.

Remarque 5.6.1 Dans la pratique, on identifie le plus souvent P et $\tilde{P}$, c'est-à-dire qu'on note P la fonction polynomiale associée.

Exemple 5.6.1

1. La fonction polynomiale associée au polynôme constant λ est la fonction constante $x \mapsto \lambda$.
2. La fonction polynomiale associée au polynôme X est la fonction identité $x \mapsto x$ de $\mathbb{K}$ dans $\mathbb{K}$.
3. Soit le polynôme $P = 3X^2 - 2X + 4$, élément de $\mathbb{R}[X]$, la fonction polynomiale associée au polynôme P est $\tilde{P}(x) = 3x^2 - 2x + 4$.
4. Si $P = (1, 0, 0, 5 + i, 0, 0, \dots) \in \mathbb{C}[X]$, alors

$$\forall x \in \mathbb{C} : \tilde{P}(x) = 1 + (5 + i)x^3.$$

5. Si $P = (0, 0, 12i, 20, 1, 0, \dots) \in \mathbb{C}[X]$, alors

$$\forall x \in \mathbb{C} : \tilde{P}(x) = 12ix^2 + 20x^3 + x^4.$$

6. La fonction monôme associée à $P = (0, 0, 5, 0, \dots)$ est $\tilde{P} : x \mapsto 5x^2$.

Proposition 5.6.1 Pour tous α de $\mathbb{K}$ et P, Q de $\mathbb{K}[X]$, la fonction polynomiale vérifie les propriétés suivantes :

1. $\widetilde{(P + \alpha Q)} = \tilde{P} + \alpha \tilde{Q}$.
2. $\widetilde{(PQ)} = \tilde{P} \tilde{Q}$.
3. $\widetilde{(P \circ Q)} = \tilde{P} \circ \tilde{Q}$.

Preuve. On note $P = \sum_{n=0}^{+\infty} a_n X^n \in \mathbb{K}[X]$, $Q = \sum_{n=0}^{+\infty} b_n X^n \in \mathbb{K}[X]$.

1. Pour la linéarité est claire. En effet, pour tout $x \in \mathbb{K}$,

$$(\widetilde{P + \alpha Q})(x) = \sum_{n=0}^{+\infty} (a_n + \alpha b_n) x^n = \sum_{n=0}^{+\infty} a_n x^n + \alpha \sum_{n=0}^{+\infty} b_n x^n = \tilde{P}(x) + \alpha \tilde{Q}(x) = (\tilde{P} + \alpha \tilde{Q})(x).$$

2. Pour le produit, pour tout $x \in \mathbb{K}$,

$$(\widetilde{PQ})(x) = \sum_{n=0}^{+\infty} d_n x^n,$$

avec pour tout $n \in \mathbb{N}$,

$$d_n = \sum_{k=0}^n a_k b_{n-k},$$

d'où

$$\begin{aligned} (\widetilde{PQ})(x) &= \sum_{n=0}^{+\infty} \left(\sum_{k=0}^n a_k b_{n-k} \right) x^n = \sum_{n=0}^{+\infty} \sum_{k=0}^n a_k x^k b_{n-k} x^{n-k} = \sum_{k=0}^{+\infty} a_k x^k \sum_{n=k}^{+\infty} b_{n-k} x^{n-k} \\ &= \sum_{k=0}^{+\infty} a_k x^k \sum_{l=0}^{+\infty} b_l x^l = \tilde{P}(x) \tilde{Q}(x). \end{aligned}$$

Enfin, pour tout $x \in \mathbb{K}$, on a

$$(\widetilde{P \circ Q})(x) = \left(\sum_{n=0}^{+\infty} a_n Q^n \right) (x) = \sum_{n=0}^{+\infty} a_n (\tilde{Q}(x))^n = \tilde{P}(\tilde{Q}(x)) = (\tilde{P} \circ \tilde{Q})(x),$$

d'après les deux premiers points. □

5.7 Arithmétique dans $\mathbb{K}[X]$

Dans l'anneau des polynômes $\mathbb{K}[X]$, où $\mathbb{K}$ est un corps, l'arithmétique est un ensemble de règles et de procédures qui permettent de manipuler, de diviser et de factoriser les polynômes. Ces opérations sont analogues à celles effectuées avec les entiers, mais adaptées aux polynômes. Voici une vue d'ensemble des principales opérations arithmétiques dans $\mathbb{K}[X]$.

5.7.1 Divisibilité dans l'anneau $\mathbb{K}[X]$

Dans l'anneau des polynômes $\mathbb{K}[X]$, où $\mathbb{K}$ est un corps (comme $\mathbb{R}$, $\mathbb{C}$, ou $\mathbb{Q}$), la divisibilité et la division euclidienne sont des concepts fondamentaux qui jouent un rôle crucial dans la structure algébrique de cet anneau. Tout comme $\mathbb{Z}$ qui est un anneau commutatif, on peut définir dans l'anneau $\mathbb{K}[X]$ la notion de diviseur et multiple ainsi que de division euclidienne.

Définition 5.7.1 (*Divisibilité, diviseur, multiple*)

Soient $A, B \in \mathbb{K}[X]$. On dit que B divise A , ou que B est un diviseur de A , ou que A est divisible par B , ou que A est un multiple de B , s'il existe $Q \in \mathbb{K}[X]$ tel que $A = QB$. Cette relation se note $B|A$.

Notation 5.7.1

1. L'ensemble des multiples d'un polynôme $P \in \mathbb{K}[X]$ est noté

$$P\mathbb{K}[X] = \{PQ : Q \in \mathbb{K}[X]\}.$$

Donc

$$B|A \iff A\mathbb{K}[X] \subset B\mathbb{K}[X].$$

2. L'ensemble des diviseurs d'un polynôme P est noté $\mathcal{D}(P)$.

Exemple 5.7.1

1. **Cas particuliers des polynômes** $0_{\mathbb{K}[X]}$ et $1_{\mathbb{K}[X]}$.

a). Pour tout polynôme B , on a $0_{\mathbb{K}[X]} = QB$ avec $Q = 0_{\mathbb{K}[X]}$. Le polynôme nul $0_{\mathbb{K}[X]}$ est donc un multiple de tout polynôme B (ou encore : tout polynôme B de $\mathbb{K}[X]$ divise le polynôme nul). Ainsi

$$\forall B \in \mathbb{K}[X] : B|0_{\mathbb{K}[X]}.$$

En revanche le polynôme nul ne divise que lui-même (car $A = Q0_{\mathbb{K}[X]} \implies A = 0_{\mathbb{K}[X]}$), donc $0_{\mathbb{K}[X]}\mathbb{K}[X] = \{0_{\mathbb{K}[X]}\}$ et

$$\forall A \in \mathbb{K}[X] : 0_{\mathbb{K}[X]}|A \iff A = 0_{\mathbb{K}[X]}.$$

b). L'égalité évidente $A = 1_{\mathbb{K}[X]}A$ dit que le polynôme constant $1_{\mathbb{K}[X]}$ divise tout polynôme A , ou encore que tout polynôme de $\mathbb{K}[X]$ est multiple du polynôme $1_{\mathbb{K}[X]}$. Ainsi $1_{\mathbb{K}[X]}\mathbb{K}[X] = \mathbb{K}[X]$.

2. Dans $\mathbb{R}[X]$, le polynôme $A = X^2 + 3X + 2$ est divisible par $B = X + 1$ car

$$X^2 + 3X + 2 = (X + 1)(X + 2).$$

3. Dans $\mathbb{R}[X]$, le polynôme $A = X^4 + X$ est divisible par $B_1 = X$ et $B_2 = X^3 + 1$ mais également par le polynôme $B_3 = X + 1$ puisque

$$X^4 + X = (X + 1)(X^3 - X^2 + X).$$

4. Dans $\mathbb{C}[X]$ (mais pas dans $\mathbb{R}[X]$), on a $(X - i)$ divise $(X^2 + 1)$.

5. Dans $\mathbb{R}[X]$, le polynôme $(1 - X)$ divise $1 - X^{n+1}$. En effet,

$$1 - X^{n+1} = (1 + X + X^2 + \dots + X^n)(1 - X).$$

Remarque 5.7.1

1. Si $B = 0_{\mathbb{K}[X]}$ dans la définition 5.7.1, alors $A = 0_{\mathbb{K}[X]}$ et Q est quelconque. Si $B \neq 0_{\mathbb{K}[X]}$; le polynôme Q est unique. En effet si $BQ_1 = BQ_2$, alors $B(Q_1 - Q_2) = 0_{\mathbb{K}[X]}$. Par intégrité de $\mathbb{K}[X]$, on a $Q_1 = Q_2$.

2. Soient $A, B \in \mathbb{K}[X]$, si $B|A$ avec B non nul, alors

$$\deg(A) \geq \deg(B).$$

En effet, il existe un polynôme Q non nul tel que $A = BQ$, et donc

$$\deg(A) = \deg(B) + \deg(C) \geq \deg(B).$$

3. La relation de la divisibilité dans $\mathbb{K}[X]$ est réflexive et transitive mais elle n'est pas ni symétrique, ni antisymétrique. En effet,

i) **Réflexivité** : Pour tout $A \in \mathbb{K}[X]$, on a $A = 1_{\mathbb{K}[X]}A = A$, donc $A|A$.

ii) **Transitivité** : Soient $A, B, C \in \mathbb{K}[X]$. Si $A|B$ et $B|C$, alors il existe $D \in \mathbb{K}[X]$ tel que $B = AD$ et il existe $E \in \mathbb{K}[X]$ tel que $C = BE$. D'où

$$C = BE = (AD)E = A(DE).$$

Alors $A|C$.

iii) **Non symétrie** : X divise X^2 tandis que X^2 ne divise pas X .

iiii) **Non antisymétrie** : X et $2X$ se divisent l'un l'autre mais ne sont pas égaux.

Définition 5.7.2 (Polynômes Associés)

Soient $A, B \in \mathbb{K}[X]$ deux polynômes. On dit que les polynômes A et B sont associés si A divise B et B divise A ,

$$A|B \text{ et } B|A.$$

Exemple 5.7.2 Considérons les polynômes suivants $A = 2X^2 - 4$, $B = -6X^2 + 12$ de $\mathbb{R}[X]$. Alors A et B sont associés, car

$$A|B \text{ et } B|A.$$

En effet,

$$B = -6X^2 + 12 = -3(2X^2 - 4) = -3A, Q_1 = -3 \in \mathbb{R}[X],$$

et

$$A = 2X^2 - 4 = \frac{-1}{3}(-6X^2 + 12) = \frac{-1}{3}B, Q_2 = \frac{-1}{3} \in \mathbb{R}[X].$$

Définition 5.7.3 (Normalisé d'un polynôme non nul)

Soit P un polynôme non nul, et soit a_d son coefficient de plus haut degré. Le polynôme unitaire $Q = a_d^{-1}P$ est appelé le normalisé du polynôme P .

Exemple 5.7.3 Si $P = 3X^3 - 6X + 9$, on peut normaliser ce polynôme en le divisant par son coefficient dominant (dans ce cas, 3) pour obtenir une forme plus simple, alors le polynôme normalisé est

$$Q = X^3 - 2X + 3.$$

Proposition 5.7.1 (Caractérisation des couples de polynômes associés)

Soient $A, B \in \mathbb{K}[X]$ deux polynômes. Alors les polynômes A et B sont associés si et seulement s'il existe $\lambda \in \mathbb{K}^*$ vérifiant $A = \lambda B$. Cette relation se traduit par le fait qu'ils sont proportionnels l'un à l'autre par un facteur constant.

Preuve.

• Supposons que A et B sont associés, alors le résultat est évident si A ou B est nul (et dans ce cas $A = B = 0_{\mathbb{K}[X]}$). Soient donc A et B deux polynômes non nuls tels que $A|B$ et $B|A$, donc il existe des polynômes non nuls $Q_1, Q_2 \in \mathbb{K}[X]^*$ tels que

$$A = Q_1B \text{ et } B = Q_2A.$$

Ainsi

$$A = Q_1Q_2A,$$

ou encore

$$A(1 - Q_1Q_2) = 0_{\mathbb{K}[X]},$$

puisque $\mathbb{K}[X]$ est intègre, on a $Q_1Q_2 = 1_{\mathbb{K}[X]}$. Par conséquent, Q_1 et Q_2 sont des polynômes inversibles inverses l'un de l'autre. Appliquant la proposition 5.3.3, il existe $\lambda \in \mathbb{K}^*$ tel que $Q_1 = \lambda$ et $Q_2 = \lambda^{-1}$. On a alors $A = \lambda B$.

• Réciproquement, on suppose qu'il existe $\lambda \in \mathbb{K}^*$ vérifiant $A = \lambda B$. Alors d'une part $B|A$. D'autre part, λ est non nul donc $B = \lambda^{-1}A$ et $A|B$. Donc A et B sont bien associés. $\square$

Corollaire 5.7.1

1. Tout polynôme $P \in \mathbb{K}[X]$, non nul est associé à un unique polynôme unitaire : $Q = c^{-1}P$, où $c \in \mathbb{K}$ est le coefficient dominant de P .

2. La relation d'association dans $\mathbb{K}[X]$ est une relation d'équivalence.
3. La relation de divisibilité restreinte à l'ensemble des polynômes unitaires est une relation d'ordre. En effet, comme deux polynômes unitaires associés sont égaux, l'antisymétrie, réflexivité et transitivité étant immédiates.

Revenons à des propriétés classiques générales sur la divisibilité. L'ensemble des règles de calcul sur la divisibilité dans $\mathbb{K}[X]$ est consigné dans la proposition suivante dont la preuve, laissée au lecteur, est en tout point analogue au cas des entiers relatifs.

Proposition 5.7.2 Soient $A, B, C, D \in \mathbb{K}[X]$ et $\alpha, \beta \in \mathbb{K}$. Alors

1. Stabilité par combinaison linéaire :

$$\text{Si } D|A \text{ et } D|B \text{ alors } D|(\alpha A + \beta B).$$

2. Compatibilité avec produit :

$$\text{Si } B|A \text{ et } D|C \text{ alors } BD|AC.$$

3. En particulier, on a la stabilité par produit :

$$\text{Si } B|A \text{ alors } BD|AD.$$

et pour tout $n \in \mathbb{N}^*$,

$$\text{Si } B|A \text{ alors } B^n|A^n.$$

5.7.2 Division euclidienne dans $\mathbb{K}[X]$

Les polynômes forment une structure d'anneau, ce qui signifie qu'on peut les additionner et les multiplier en suivant les règles usuelles de calcul. Cependant, ils ne forment pas un corps, car la plupart des polynômes n'ont pas d'inverse pour la multiplication. En ce sens, la structure de $\mathbb{K}[X]$ est assez similaire à celle de l'ensemble $\mathbb{Z}$ des entiers. Nous allons maintenant montrer que de nombreuses notions et résultats de l'arithmétique des entiers se transposent bien dans le cas des polynômes.

Dans l'anneau $\mathbb{K}[X]$, il existe une division euclidienne analogue à celle de $(\mathbb{Z}, +, \cdot)$. Cette division est fondamentale pour l'arithmétique des polynômes. Effectuer la division d'un polynôme A par un autre polynôme B non nul revient à écrire A sous la forme $A = BQ + R$ avec $(Q, R) \in \mathbb{K}[X]^2$ et $\deg(R) < \deg(B)$.

Cette division euclidienne permet de généraliser de nombreux résultats de l'arithmétique aux polynômes, tels que le calcul du pgcd (plus grand commun diviseur) et la factorisation des polynômes. Le théorème suivant établit l'existence et l'unicité de ce couple (Q, R) .

Théorème 5.7.1 (Définition) (Division euclidienne)

1. Soient A et B deux polynômes de $\mathbb{K}[X]$ avec $B \neq 0_{\mathbb{K}[X]}$. Il existe un unique couple (Q, R) de polynômes de $\mathbb{K}[X]$ tel que

$$\begin{cases} A = BQ + R, \\ \deg(R) < \deg(B). \end{cases}$$

- Cette opération s'appelle *division euclidienne de A par B* .
- On appelle A le **dividende**, B le **diviseur**, Q le **quotient** et R le **reste** de cette division.

Preuve. Principe de démonstration.

• Pour l'unicité : Elle s'effectue en utilisant un mode de raisonnement par l'absurde. On suppose qu'il existe deux couples (Q_1, R_1) et (Q_2, R_2) de $(\mathbb{K}[X])^2$ vérifiant

$$\begin{cases} A = BQ_1 + R_1, \\ \deg(R_1) < \deg(B). \end{cases} \quad \text{et} \quad \begin{cases} A = BQ_2 + R_2, \\ \deg(R_2) < \deg(B). \end{cases} \quad (5.8)$$

Alors

$$R_1 - R_2 = B(Q_2 - Q_1).$$

Si $Q_1 \neq Q_2$, alors on a, d'une part

$$\deg(R_1 - R_2) = \deg(B(Q_2 - Q_1)) = \deg(B) + \deg(Q_2 - Q_1) \geq \deg(B).$$

et, d'autre part

$$\deg(R_1 - R_2) \leq \max \{ \deg(R_1), \deg(R_2) \} < \deg(B).$$

Cette dernière inégalité (stricte) est en contradiction avec l'inégalité (large) $\deg(R_1 - R_2) \geq \deg(B)$. On en déduit alors que nécessairement $Q_1 = Q_2$. Ceci implique d'après les relations (??) que $R_1 = R_2$.

• Pour l'existence : la démonstration se fait par récurrence sur $n = \deg(A)$. Fixons pour toute la suite $B = b_0 + b_1X + \dots + b_mX_m$ avec $b_m \neq 0_{\mathbb{K}}$ ($\deg(B) = m$) et pour tout $n \in \mathbb{N}$, notons $\mathcal{P}(n)$ l'hypothèse de récurrence suivante :

$$\mathcal{P}(n) : \forall A \in \mathbb{K}[X] : \deg(A) \leq n, \exists (Q, R) \in \mathbb{K}[X]^2 : A = BQ + R \text{ et } \deg(R) < \deg(B).$$

□

Remarque 5.7.2

1. Pour effectuer la division euclidienne de A par B lorsque $\deg(A) \geq \deg(B)$, il a été impératif d'écrire les deux polynômes A et B dans le sens des puissances décroissantes. La division euclidienne est d'ailleurs aussi appelée division suivant les puissances décroissantes.

2. Si $\deg(A) < \deg(B)$, alors la division euclidienne de A par B s'écrit $A = 0_{\mathbb{K}[X]}B + A$.

3. Il ne faut jamais oublier de mentionner la condition $\deg(R) < \deg(B)$ (Cette division s'arrête lorsque le degré de R est strictement inférieur au degré de B).

Lemme 5.7.1 Soient A, B, Q et R des polynômes de $\mathbb{K}[X]$ tels que $A = BQ + R$. Alors

$$\mathcal{D}(A) \cap \mathcal{D}(B) = \mathcal{D}(B) \cap \mathcal{D}(R).$$

Preuve. En effet, la relation $A = BQ + R$ montre que tout diviseur commun à B et R est aussi un diviseur de A (et aussi de B évidemment). L'inclusion inverse provient de même de la relation $R = (-Q)B + A$. □

Algorithme de la division euclidienne :

- On pose la division euclidienne comme la division des entiers, en disposant A à gauche et B à droite, les monômes étant écrits dans l'ordre décroissant des degrés (donc en marquant d'abord les monômes de plus haut degré).
- On trouve le monôme aX^k tel que aX^kB ait même monôme dominant que A , puis on effectue la différence $A_1 = A - aX^kB$, qui a donc un degré strictement plus petit que A .
- On recommence sur A_1 , et on en déduit A_2 .
- On recommence ainsi jusqu'à obtenir A_k de degré strictement plus petit que B . Alors A_k est le reste recherché, et le quotient est la somme des monômes par lesquels on a multiplié B pour obtenir les A_i successifs.

Exemple 5.7.4

1. Effectuons la division euclidienne de $A = X^4 + X^2 + X + 2$ par $B = X^2 - 3$ dans $\mathbb{R}[X]$.

$$\begin{array}{lcl}
 A & = & \overbrace{X^5 + 4X^4 + 2X^3 + X^2 - X - 1}^{\text{Dividende}} \quad \overbrace{X^3 - 2X + 3}^{\text{Diviseur}} = B \\
 A_1 = A - X^2B & = & 4X^4 + 4X^3 - 2X^2 - X - 1 \quad \underbrace{X^2 + 4X + 4}_{\text{Quotient}} = Q \\
 A_2 = A_1 - 4XB & = & 4X^3 + 6X^2 - 13X - 1 \\
 A_3 = A_2 - 4B & = & \underbrace{6X^2 - 5X - 13}_{\text{Reste}} = R
 \end{array}$$

Nous avons arrêté le processus car le degré de reste $(6X^2 - 5X - 13)$ est strictement inférieur au degré du diviseur B . Ainsi

$$A = BQ + R \text{ avec } Q = X^2 + 4X + 4 \text{ et } R = 6X^2 - 5X - 13.$$

2. Effectuons la division euclidienne de $A = iX^3 - X^2 + (1 - i)$ par $B = (1 + i)X^2 - iX + 3$ dans $\mathbb{C}[X]$.

$$\begin{array}{lcl}
 A & = & iX^3 - X^2 + (1 - i) \quad (1 + i)X^2 - iX + 3 = B \\
 A_1 = A - \frac{(1+i)}{2}XB & = & \left(\frac{-3+i}{2}\right)X^2 - \left(\frac{3}{2} + \frac{3i}{2}\right)X + (1 - i) \quad \frac{(1+i)}{2}X + \frac{-1+2i}{2} \\
 A_2 = A_1 - \left(\frac{-1+2i}{2}\right)B & = & -\left(\frac{5}{2} + 2i\right)X + \left(\frac{5}{2} - 4i\right)
 \end{array}$$

Ainsi

$$A = BQ + R \text{ avec } Q = \frac{(1+i)}{2}X + \frac{-1+2i}{2} \text{ et } R = -\left(\frac{5}{2} + 2i\right)X + \left(\frac{5}{2} - 4i\right).$$

Corollaire 5.7.2 (Equivalence entre divisibilité et reste de la division euclidienne)

Soient A et B deux polynômes de $\mathbb{K}[X]$ avec $B \neq 0_{\mathbb{K}[X]}$. Alors $B|A$ si et si le reste de la division euclidienne de A par B est nul.

Cela signifie que la divisibilité entre deux polynômes est étroitement liée au reste de leur division euclidienne.

Preuve.

- Supposons que $B|A$, alors il existe $Q \in \mathbb{K}[X]$ tel que $A = BQ = BQ + 0_{\mathbb{K}[X]}$. Comme $\deg(0_{\mathbb{K}[X]}) < \deg(B)$, par unicité du reste de la division euclidienne pour les polynômes, on en déduit que ce reste est nul (et le quotient vaut Q).
- Réciproquement, on suppose que le reste de la division euclidienne de A par B est nul. Notons Q le quotient de cette division. Alors $A = BQ + 0_{\mathbb{K}[X]} = BQ$, donc on a bien $B|A$. $\square$

Remarque 5.7.3 Ce résultat sert souvent pour démontrer qu'un polynôme B non nul divise un polynôme A .

5.7.3 Division suivant les puissances croissantes

La division selon les puissances croissantes est une méthode semblable à la division classique, mais elle inverse l'ordre des termes. Bien que les deux approches visent à obtenir un quotient et un reste, elles diffèrent dans la manière dont elles manipulent les termes des polynômes. Cette méthode est moins couramment utilisée que la division par les puissances décroissantes, mais elle peut s'avérer utile dans certains contextes mathématiques ou informatiques. Le théorème suivant est admis.

Théorème 5.7.2 (Définition) Soient $n \in \mathbb{N}$, et $A, B \in \mathbb{K}[X]$ tel que $\tilde{B}(0_{\mathbb{K}}) \neq 0_{\mathbb{K}}$ (le terme constant de B n'est pas nul). Il existe un couple unique (Q, R) de $(\mathbb{K}[X])^2$ tels que

$$\begin{cases} A = BQ + X^{n+1}R \\ \deg(Q) \leq n. \end{cases}$$

Le polynôme Q (resp. $X^{n+1}R$) s'appelle le quotient (resp. : reste) de la division de A par B suivant les puissances croissantes jusqu'à l'ordre n .

Exemple 5.7.5

1. Soient $A = 4 + X^2$ et $B = 1 + X + X^2$ deux polynômes de $\mathbb{R}[X]$. Effectuons la division de A par B selon les puissances croissantes à l'ordre 2. Alors

$$\begin{array}{ll} A_1 = A - 4B = -4X - 3X^2 & \begin{array}{l} A = 4 + X^2 \\ 1 + X + X^2 = B \\ 4 - 4X + X^2 = Q \end{array} \end{array}$$

$$A_2 = A_1 + 4XB = X^2 + 4X^3$$

$$A_3 = A_2 - X^2B = X^3R = \underbrace{3X^3 - X^4}_{\text{Reste}}$$

Nous avons arrêté le processus car le degré du quotient est inférieur ou égal à 2 (ici 2 est l'ordre de la division). Ainsi,

$$A = BQ + X^{n+1}R,$$

avec

$$\begin{cases} Q = 4 - 4X + X^2 =, \\ \deg(Q) = 2 \leq 2. \end{cases}$$

Donc on a l'égalité

$$A = 4 + X^2 = (1 + X + X^2)(4 - 4X + X^2) + (3X^3 - X^4).$$

2. Soient $A = X + X^4$ et $B = 1 + X$ deux polynômes de $\mathbb{R}[X]$. Effectuons la division de A par B selon les puissances croissantes aux ordres $n = 1, 2, 3, 4, 5$.

• À l'ordre 0, on a

$$Q_0 = 0_{\mathbb{R}[X]} \text{ et } XR_0 = X + X^4.$$

• Pour l'ordre $n \geq 1$, on a

$$\begin{array}{ll} A_1 = A - XB = -X^2 + X^4 & \begin{array}{l} A = X + X^4 \\ 1 + X = B \\ X - X^2 + X^3 = Q \end{array} \end{array}$$

$$A_2 = A_1 + X^2B = X^3 + X^4$$

$$A_3 = A_2 - X^3B = 0_{\mathbb{R}[X]}.$$

On obtient ainsi,

- À l'ordre 1, on a

$$Q_1 = X \text{ et } X^2 R_1 = -X^2 + X^4.$$

- À l'ordre 2, on a

$$Q_2 = X - X^2 \text{ et } X^3 R_2 = X^3 + X^4,$$

et on en déduit

$$\forall n \geq 3 : Q_n = X - X^2 + X^3 \text{ et } X^{n+1} R_n = 0_{\mathbb{K}[X]}.$$

Remarque 5.7.4 La division suivant les puissances croissantes est surtout utilisée pour :

1. L'obtention de la décomposition en éléments simples d'une fraction rationnelle.
2. Le calcul de développement limité d'un quotient.

5.7.4 Plus grand commun diviseur

Dans cette section, pour tout polynôme $A \in \mathbb{K}[X]$, nous notons $\mathcal{D}(A)$ l'ensemble des diviseurs de A , et pour tout couple $(A, B) \in \mathbb{K}[X]^2$, nous notons $\mathcal{D}(A, B)$ l'ensemble des diviseurs communs à A et B . Il est important de noter que $\mathcal{D}(A, B) = \mathcal{D}(A) \cap \mathcal{D}(B)$. Rappelons également que l'ensemble des diviseurs du polynôme nul est $\mathcal{D}(0_{\mathbb{K}[X]}) = \mathbb{K}[X]$. En revanche, si $A \neq 0_{\mathbb{K}[X]}$, tous les polynômes de $\mathcal{D}(A)$ sont de degré inférieur ou égal à $\deg(A)$. Dans la suite, nous considérons deux polynômes A et B de $\mathbb{K}[X]$, dont au moins l'un est non nul.

- L'ensemble des diviseurs communs à A et B est non vide (puisqu'il contient $1_{\mathbb{K}[X]}$) et ne comprend pas le polynôme nul (étant donné que AA ou B est non nul).

- L'ensemble des degrés des polynômes diviseurs communs de A et de B constitue donc une partie non vide de $\mathbb{N}$, qui est majorée par $\deg(A)$ si $A \neq 0_{\mathbb{K}[X]}$ ou par $\deg(B)$ sinon. Par conséquent, cet ensemble possède un plus grand élément r dans $\mathbb{N}$. Ainsi, il existe un polynôme D de degré r tel que D soit un diviseur commun à A et B . Cela nous conduit aux définitions suivantes.

Définition 5.7.4 (Plus grand commun diviseur) Soient A et B deux polynômes non tous les deux nuls de $\mathbb{K}[X]$. Alors il existe un unique polynôme unitaire non nul D dans $\mathbb{K}[X]$, diviseur commun de A, B et de plus haut degré parmi les diviseurs communs de A et B . Le polynôme D est appelé le **plus grand commun diviseur** de A et B (en abrégé : pgcd) et est noté $\text{pgcd}(A, B)$.

Remarque 5.7.5

- Un pgcd de deux polynômes non tous les deux nuls est toujours non nul par définition, puisque son degré n'est pas $-\infty$.
- Par convention le pgcd de $0_{\mathbb{K}[X]}$ et $0_{\mathbb{K}[X]}$ est $0_{\mathbb{K}[X]}$ (pour des raisons pratiques et pour garder une certaine cohérence dans les calculs algébriques, on considère souvent que le pgcd de deux polynômes nuls est également le polynôme nul).
- On a évidemment $\text{pgcd}(A, B) = \text{pgcd}(B, A)$.
- Si A et B sont des polynômes non nuls, alors comme un polynôme non nul admet le même ensemble de diviseurs qu'un de ses associés, chercher $\text{pgcd}(A, B)$ revient à chercher le pgcd des polynômes unitaires associés à A et B .
- Soient A et B des polynômes unitaires. Le polynôme $\text{pgcd}(A, B)$ est le plus grand élément de l'ensemble des diviseurs de A et B au sens de la relation de divisibilité dans l'ensemble des polynômes unitaires.

Proposition 5.7.3 (Quelques propriétés du pgcd)

Soient A, B, C trois polynômes non nuls de $\mathbb{K}[X]$. Alors

1. Tout $D = \text{pgcd}(A, B)$ vérifie

$$\forall P \in \mathbb{K}[X] : P \mid A \text{ et } P \mid B \iff P \mid D.$$

2. Multiplicativité : pour tout $Q \in \mathbb{K}[X]$, unitaire,

$$\text{pgcd}(Q \cdot A, Q \cdot B) = Q \cdot \text{pgcd}(A, B).$$

3. Pour les puissances,

$$\forall n \in \mathbb{N}^* : \text{pgcd}(A^n, B^n) = (\text{pgcd}(A, B))^n.$$

4. Pour tous $\alpha, \beta \in \mathbb{K}^*$, on a

$$\text{pgcd}(\alpha \cdot A, \beta \cdot B) = \text{pgcd}(A, B).$$

5. Associativité de pgcd,

$$\text{pgcd}(A, \text{pgcd}(B, C)) = \text{pgcd}(\text{pgcd}(A, B), C).$$

5.7.4.1 Calcul du pgcd (Algorithme d'Euclide)

L'algorithme d'Euclide est une méthode efficace permettant de calculer le pgcd de deux polynômes. L'outil de base est la division euclidienne. Cet algorithme repose sur l'idée que le pgcd de deux polynômes ne change pas si on remplace l'un des polynômes par son reste lorsqu'on le divise par l'autre.

Lemme 5.7.2 Soit B un polynôme non nul de $\mathbb{K}[X]^*$, et A un polynôme quelconque. Notons Q et R le quotient et le reste de la division euclidienne de A par B . Alors on a

$$\text{pgcd}(A, B) = \text{pgcd}(B, R).$$

Preuve. Soit D un diviseur commun aux polynômes A et B . Puisque $R = A - BQ$, D est un diviseur aussi de R . Donc D divise $\text{pgcd}(B, R)$. En choisissant $D = \text{pgcd}(A, B)$, on conclut que

$$\text{pgcd}(A, B) \mid \text{pgcd}(B, R).$$

Soit maintenant D un polynôme divisant B et R . Comme $A = BQ + R$, on a aussi $D \mid A$. Donc $D \mid \text{pgcd}(A, B)$. On a donc finalement

$$\text{pgcd}(B, R) \mid \text{pgcd}(A, B).$$

Les deux polynômes $\text{pgcd}(A, B)$ et $\text{pgcd}(B, R)$ sont unitaires et associés. Ils sont donc égaux. $\square$

Algorithme d'Euclide

Pour calculer explicitement le pgcd de deux polynômes A et B non nuls, on utilise l'algorithme d'Euclide. Comme pour les entiers, l'existence du pgcd est donc fondée sur le lemme 5.7.1 précédent et la division euclidienne : si B est non nul, le pgcd de A et B est le même que celui de B et du

reste R de la division euclidienne de A par B . Il suffit donc de remplacer (A, B) par (B, R) et de recommencer. Si $R \neq 0_{\mathbb{K}[X]}$, on peut alors utiliser le résultat suivant,

Soient A, B deux polynômes non nuls tels que $\deg(A) \geq \deg(B)$. Posons $R_0 = A$ et $R_1 = B$ et définissons par récurrence le polynôme R_{k+1} comme reste de la division euclidienne de R_{k-1} par R_k

$$R_{k-1} = Q_k R_k + R_{k+1} \text{ avec } \deg(R_{k+1}) < \deg(R_k).$$

Alors le pgcd de A et B est le dernier reste non nul normalisé de cet algorithme.

Mise en œuvre de l'algorithme d'Euclide

Sans perte de généralité, on peut supposer que $B \neq 0_{\mathbb{K}[X]}$.

1. On pose $R_0 = A$ et $R_1 = B$. Le polynôme R_1 est le premier “reste” et il est non nul.
2. On effectue la division euclidienne de R_0 par R_1 ,

$$R_0 = Q_1 R_1 + R_2, \deg(R_2) < \deg(R_1).$$

- Si $R_2 = 0_{\mathbb{K}[X]}$, alors le processus s'arrête, et R_1 est le dernier reste non nul obtenu.
3. Sinon on effectue la division de R_1 par R_2 ,

$$R_1 = Q_2 R_2 + R_3, \deg(R_3) < \deg(R_2).$$

- Si $R_3 = 0_{\mathbb{K}[X]}$, alors le processus s'arrête et R_2 est le dernier reste non nul obtenu.
4. Sinon on effectue la division de R_2 par R_3 ,

$$R_2 = Q_3 R_3 + R_4, \deg(R_4) < \deg(R_3).$$

À la k -ième étape de cet algorithme, on effectue une division de la forme
est une division

$$R_{k-1} = Q_k R_k + R_{k+1}.$$

À ce stade on a

$$\deg(R_0) > \deg(R_1) > \dots > \deg(R_k) > \deg(R_{k+1}).$$

Ce processus est fini car la suite des $\deg(R_k)$ est strictement décroissante dans $\mathbb{N}$. Il existe donc une n -ième étape où $R_{n-1} = Q_n R_n$, c'est-à-dire $R_{n+1} = 0$. Une fois ces calculs effectués, on observe

$$\text{pgcd}(A, B) = \text{pgcd}(B, R_1) = \text{pgcd}(R_1, R_2) = \dots = \text{pgcd}(R_{n-1}, R_n).$$

Le polynôme R_n est le dernier reste non nul obtenu dans cette méthode. En divisant par le coefficient dominant de R_n , on obtient le normalisé R_n^* du dernier reste non nul, qui est bien le pgcd de A et de B (c'est-à-dire qui satisfait aux conditions de la définition).

Exemple 5.7.6 Soient $A = X^5 + X + 1$ et $B = X^4 - 2X^3 - X + 2$ de $\mathbb{R}[X]$. On utilise l'algorithme d'Euclide pour déterminer le pgcd de A et B . Posons

$$R_0 = A, R_1 = B.$$

Étape 1 : Divisons R_0 par R_1 .

$$\begin{aligned} A &= X^5 + X + 1 & X^4 - 2X^3 - X + 2 &= B \\ A_1 = A - XB &= 2X^4 + X^2 - X + 1 & X + 2 &= Q_1 \\ A_2 = A_1 - 2B &= 4X^3 + X^2 + X - 3. \end{aligned}$$

Donc

$$R_2 = 4X^3 + X^2 + X - 3.$$

Étape 2 : Divisons R_1 par R_2 .

$$\begin{aligned} R_1 &= X^4 - 2X^3 - X + 2 & 4X^3 + X^2 + X - 3 &= R_2 \\ A_1 = R_1 - \frac{1}{4}XR_2 &= -\frac{9}{4}X^3 - \frac{1}{4}X^2 - \frac{1}{4}X + 2 & \frac{1}{4}X - \frac{9}{16} &= Q_2 \\ A_2 = A_1 + \frac{9}{16}R_2 &= \frac{5}{16}X^2 + \frac{5}{16}X + \frac{5}{16}. \end{aligned}$$

Donc

$$R_3 = \frac{5}{16}X^2 + \frac{5}{16}X + \frac{5}{16}.$$

Étape 3 : Divisons R_2 par R_3 .

$$\begin{aligned} R_2 &= 4X^3 + X^2 + X - 3 & R_3 &= \frac{5}{16}X^2 + \frac{5}{16}X + \frac{5}{16} \\ A_1 = R_2 - \frac{64}{5}XR_3 &= -3X^2 - 3X - 3 & \frac{64}{5}X - \frac{48}{5} &= Q_3 \\ A_2 = A_1 + \frac{48}{5}R_3 &= 0. \\ R_4 &= 0. \end{aligned}$$

Donc R_3 est le dernier reste non nul, on le normalise, c'est à dire qu'on le divise par son coefficient dominant $\frac{5}{16}$ pour le rendre unitaire, et on obtient

$$\text{pgcd}(A, B) = X^2 + X + 1.$$

L'identité de Bézout est un résultat fondamental en algèbre, qui affirme qu'il est possible d'exprimer le plus grand commun diviseur de deux polynômes comme une combinaison linéaire de ces polynômes.

Proposition 5.7.4 (Relation de Bézout)

Soient A et B deux polynômes de $\mathbb{K}[X]$ et D un pgcd de A et B . Alors il existe deux polynômes U et V appelés coefficients de Bézout de A et B tels que $D = AU + BV$.

Preuve. Principe de démonstration. La preuve repose sur l'algorithme d'Euclide appliqué aux polynômes. Cette méthode est pratique, car elle permet non seulement de trouver le pgcd, mais aussi de déterminer les coefficients de Bézout, qui sont les polynômes U et V . $\square$

Remarque 5.7.6 Dans la proposition précédente, il n'y a pas unicité du couple (U, V) , puisque si (U_0, V_0) est un couple de coefficients de Bézout de A et B , alors il en est de même du couple $(U_0 + QB, V_0 - QA)$ pour tout polynôme Q .

Exemple 5.7.7 Prenons deux polynômes $A = X^3 + X + 1$ et $B = X^2 + X + 1$ de $\mathbb{R}[X]$. On utilise l'algorithme d'Euclide pour déterminer le pgcd de A et B : $R_0 = A, R_1 = B$. Alors

$$A = X^3 + X + 1 \quad B = X^2 + X + 1$$

$$A_1 = A - X^2B = -X^2 + 1 \quad X - 1 = Q_1$$

$$A_2 = A_1 + B = X + 2 = R_2$$

Donc $R_2 = X + 2$. Par suite

$$R_1 = X^2 + X + 1 \quad R_2 = X + 2$$

$$D_1 = R_1 - X^2R_2 = -2X + 1 \quad X - 1 = Q_1$$

$$D_2 = D_1 + R_2 = 3$$

Donc $R_3 = 3$ est le dernier reste non nul, on le normalise, c'est à dire qu'on le divise par son coefficient dominant pour le rendre unitaire, et on obtient

$$\text{pgcd}(A, B) = 1.$$

Par conséquent

$$A = R_0 = BQ_1 + R_2 = R_1Q_1 + R_2,$$

$$B = R_1 = R_2Q_2 + R_3 = R_2Q_2 + 3.$$

On obtient

$$1 = \frac{1}{3}(B - R_2Q_2) = \frac{1}{3}B - \frac{1}{3}(A - BQ_1)Q_2 = A\left(\frac{-1}{3}Q_2\right) + B\left(\frac{1}{3} + \frac{1}{3}Q_1Q_2\right).$$

Le couple $U = \frac{-1}{3}(X - 1)$ et $V = \frac{1}{3}(1 + (X - 1)^2)$ convient. Ce qui vérifie la relation de Bézout pour ces polynômes.

La caractérisation du plus grand commun diviseur de deux polynômes A et B dans un anneau de polynômes $\mathbb{K}[X]$ repose sur trois conditions essentielles.

Corollaire 5.7.3 (*Caractérisation de pgcd*)

Soient A et B deux polynômes non tous les deux nuls de $\mathbb{K}[X]$, et D un polynôme unitaire, alors $D = \text{pgcd}(A, B)$ si et seulement si $D \mid A$ et $D \mid B$ et il existe deux polynômes U et V tels que $D = AU + BV$.

5.7.5 Plus petit commun multiple

Soient A et B deux polynômes non nuls. L'ensemble des multiples communs non nuls à A et B est une partie non vide de $\mathbb{K}[X]$, car il contient notamment le produit AB . Les degrés de ces multiples forment donc un sous-ensemble non vide de $\mathbb{N}$, qui possède un plus petit élément d . Par conséquent, il existe un polynôme M de degré d , qui est un multiple commun à A et B . Cela nous amène aux définitions suivantes.

Définition 5.7.5 (*Plus petit commun multiple*) Soient A et B deux polynômes non tous les deux nuls de $\mathbb{K}[X]$. Alors il existe un unique polynôme unitaire non nul M dans $\mathbb{K}[X]$, multiple commun de A, B et de plus bas degré parmi les multiples communs de A et B . Le polynôme M est appelé le **plus petit commun multiple** de A et B (en abrégé : ppcm) et est noté $\text{ppcm}(A, B)$.

Remarque 5.7.7

- On a évidemment $\text{ppcm}(A, B) = \text{ppcm}(B, A)$.
- Par définition, un ppcm de deux polynômes non nuls est toujours non nul.
- Un polynôme admet les mêmes multiples que ses polynômes associés. Donc, si A et B sont des polynômes non nuls, chercher $\text{ppcm}(A, B)$ revient à chercher le ppcm des polynômes unitaires associés à A et B .
- Soient A et B des polynômes unitaires. Alors $\text{ppcm}(A, B)$ est le plus petit élément de l'ensemble des multiples communs à A et B au sens de la relation de divisibilité dans l'ensemble des polynômes unitaires.

Le plus petit commun multiple de deux polynômes possède plusieurs propriétés intéressantes qui sont analogues à celles du ppcm des entiers. Voici quelques-unes des propriétés principales.

Proposition 5.7.5 (*Quelques propriétés du ppcm*)

Soient A, B, C trois polynômes non nuls de $\mathbb{K}[X]$. Alors

1. Tout $M = \text{ppcm}(A, B)$ vérifie

$$\forall P \in \mathbb{K}[X] : A \mid P \text{ et } B \mid P \iff M \mid P.$$

2. Multiplicativité : pour tout $Q \in \mathbb{K}[X]$, unitaire,

$$\text{ppcm}(Q \cdot A, Q \cdot B) = Q \cdot \text{ppcm}(A, B).$$

3. Pour les puissances,

$$\forall n \in \mathbb{N}^* : \text{ppcm}(A^n, B^n) = (\text{ppcm}(A, B))^n.$$

4. Pour tous $\alpha, \beta \in \mathbb{K}^*$, on a

$$\text{ppcm}(\alpha \cdot A, \beta \cdot B) = \text{ppcm}(A, B).$$

5. Associativité de ppcm ,

$$\text{ppcm}(A, \text{ppcm}(B, C)) = \text{ppcm}(\text{ppcm}(A, B), C).$$

6. La relation fondamentale entre le produit des polynômes, leur plus grand commun diviseur, et leur plus petit commun multiple peut être exprimée comme,

$$A \cdot B = \text{pgcd}(A, B) \cdot \text{ppcm}(A, B).$$

Où A et B sont deux polynômes unitaires.

Remarque 5.7.8 Si A et B sont deux polynômes quelconques, alors

$$\lambda AB = \text{pgcd}(A, B) \cdot \text{ppcm}(A, B), \lambda \in \mathbb{K}^*.$$

Exemple 5.7.8 Considérons deux polynômes unitaires $A = X^2 - 1$ et $B = X^2 - X$ de $\mathbb{R}[X]$. Alors

- Le produit de A et B est

$$AB = (X^2 - 1)(X^2 - X) = X - X^2 - X^3 + X^4,$$

et nous trouvons que le pgcd de A et B est $(X - 1)$ car c'est le polynôme unitaire de plus haut degré qui divise les deux.

- Pour trouver le ppcm, nous utilisons la relation

$$\text{ppcm}(A, B) = \frac{AB}{\text{pgcd}(A, B)},$$

on obtient

$$\text{ppcm}(A, B) = X^3 - X.$$

Corollaire 5.7.4 Comme dans le cas de l'arithmétique des entiers, on peut généraliser les notions de pgcd et ppcm de deux polynômes à un nombre fini de polynômes.

5.7.6 Polynômes premiers entre eux

Définition 5.7.6 (Couple de polynômes premiers entre eux)

Soient A et B deux polynômes non tous les deux nuls de $\mathbb{K}[X]$, on dit que A et B sont premiers entre eux (ou copremiers) si et seulement si les seuls diviseurs communs sont les polynômes constants non nuls. En d'autres termes, A et B sont premiers entre eux si et seulement si $\text{pgcd}(A, B) = 1$.

Exemple 5.7.9

- Les polynômes $(X - a)$ et $(X - b)$ de $\mathbb{R}[X]$ sont premiers entre eux si et seulement si a est différent de b .
- Soient $A = X^3 - 2X + 1$ et $B = X^2 + X + 1$ dans $\mathbb{R}[X]$, pour déterminer le pgcd de A et B , on peut utiliser l'algorithme d'Euclide ou observer directement que A et B n'ont pas de facteurs communs autres que des constantes. Puisque le pgcd est 1, ces polynômes sont premiers entre eux.

Le théorème de Bézout est un résultat important en algèbre, qui traite de la relation entre les coefficients des polynômes et leurs combinaisons linéaires. Pour les polynômes, ce théorème peut être énoncé comme suit.

Théorème 5.7.3 (de Bézout) Soient A, B deux polynômes non nuls de $\mathbb{K}[X]$. Pour que A, B soient premiers entre eux, il faut et il suffit qu'il existe deux polynômes U et V de $\mathbb{K}[X]$ tel que

$$AU + BV = 1.$$

Preuve.

- Si A et B sont premiers entre eux, alors $\text{pgcd}(A, B) = 1$ et, d'après la proposition 5.7.4, il existe $(U, V) \in \mathbb{K}[X]^2$ tel que $AU + BV = 1$.
- Réciproquement, s'il existe deux polynômes U et V tels que $AU + BV = 1$, alors tout diviseur commun à A et B divisant $AU + BV$ divise 1, il est donc de degré 0. On en déduit que A et B sont premiers entre eux. $\square$

Exemple 5.7.10 Soient $A = X^4 + 1$ et $B = X^3 - 1$ de $\mathbb{R}[X]$. Montrons que A et B sont premiers entre eux.

On effectue les divisions euclidiennes successives.

$$A = X^4 + 1 \quad B = X^3 - 1$$

$$A_1 = A - XB = X + 1 = R_1 \quad X = Q_1.$$

Où

$$A = BQ_1 + R_1, \deg(R_1) < \deg(B),$$

et

$$B = X^3 - 1 \quad X + 1 = R_1$$

$$D_1 = B - X^2R_1 = -X^2 - 1 \quad X^2 - X + 1 = Q_2$$

$$D_2 = D_1 - XB = X - 1$$

$$D_3 = D_2 - B = -2 = R_2.$$

Où

$$B = R_1Q_2 + R_2, \deg(R_2) < \deg(R_1).$$

On obtient

$$R_2 = B - R_1Q_2 = B - (A - BQ_1)Q_2$$

Par suite

$$\begin{aligned} -2 &= (X^3 - 1) - ((X^4 + 1) - X(X^3 - 1))(X^2 - X + 1) \\ &= (X^3 - X^2 + X + 1)(X^3 - 1) - (X^2 - X + 1)(X^4 + 1). \end{aligned}$$

Un couple (U, V) convenant est

$$U = \frac{1}{2}(X^2 - X + 1), V = -\frac{1}{2}(X^3 - X^2 + X + 1).$$

Proposition 5.7.6 (Théorème de Gauss)

$$\forall A, B, C \in \mathbb{K}[X]^* : \begin{cases} A \mid BC \\ \text{et} \\ \text{pgcd}(A, B) = 1 \end{cases} \implies A \mid C.$$

Preuve. Supposons $\text{pgcd}(A, B) = 1$ et $A \mid BC$, d'après la proposition 5.7.4, il existe des polynômes U et V tels que $AU + BV = 1$, ce qui implique $ACU + BCV = C$. Comme A divise ACU et BCV , on a $A \mid ACU + BCV$ et donc $A \mid C$. $\square$

5.8 Racines des polynômes

5.8.1 Généralités

La notion de racine d'un polynôme sera particulièrement importante dans l'étude des polynômes à coefficients dans un corps.

Définition 5.8.1 Soit P un polynôme de $\mathbb{K}[X]$. On dit que l'élément $\alpha \in \mathbb{K}$ est une racine (ou encore un zéro) du polynôme P si $\tilde{P}(\alpha) = 0_{\mathbb{K}}$. Où $\tilde{P}$ désigne la fonction polynomiale associée au polynôme formel P de $\mathbb{K}[X]$.

Parmi les exemples fondamentaux de division euclidienne, on retiendra l'exemple de la division d'un polynôme P par $(X - \alpha)$, où $\alpha \in \mathbb{K}$. La proposition suivante donne une condition nécessaire et suffisante pour qu'un scalaire soit racine d'un polynôme.

Proposition 5.8.1 (Caractérisation des racines avec la divisibilité)

Soient $P \in \mathbb{K}[X]$ et $\alpha \in \mathbb{K}$. Alors α est racine de P si et seulement si P est divisible par $(X - \alpha)$. En d'autres termes

$$\tilde{P}(\alpha) = 0_{\mathbb{K}} \iff (X - \alpha) | P.$$

Preuve. Soit α une racine de P . Alors $\tilde{P}(\alpha) = 0_{\mathbb{K}}$. Effectuons la division euclidienne de P par $(X - \alpha)$. On obtient $P = (X - \alpha)Q + R$ avec $\deg(R) < \deg(X - \alpha) = 1$. On a alors deux possibilités, soit $\deg(R) = 0$, soit $\deg(R) = -\infty$, c'est à dire $R = 0_{\mathbb{K}[X]}$. Montrons que la première n'est pas possible :

Si on avait $\deg(R) = 0$, alors il existerait $c \in \mathbb{K}^*$ tel que $R = c$ et on aurait

$$P = (X - \alpha)Q + c,$$

En prenant les valeurs des fonctions polynômes associées au point α , il vient

$$0_{\mathbb{K}} = \tilde{P}(\alpha) = \tilde{R}(\alpha) = c \neq 0_{\mathbb{K}},$$

ce qui est une contradiction. On a donc bien $R = 0_{\mathbb{K}[X]}$ et $P = (X - \alpha)Q$. Donc P est divisible par $(X - \alpha)$.

Réciproquement, supposons que $(X - \alpha) | P$. Alors il existe $Q \in \mathbb{K}[X]$ tel que $P = (X - \alpha)Q$. Par conséquent $\tilde{P}(\alpha) = 0_{\mathbb{K}}$, ce qui prouve que α est une racine de P . $\square$

Exemple 5.8.1

1. Le polynôme $P = 5X^2 - 25X + 30 \in \mathbb{R}[X]$ admet la factorisation suivante

$$P = 5(X - 2)(X - 3).$$

On en déduit que P admet deux racines $\alpha_1 = 2$ et $\alpha_2 = 3$ puisque

$$\tilde{P}(2) = \tilde{P}(3) = 0.$$

2. Soit le polynôme $P = (X - 3)^2(X + 1)$ de $\mathbb{R}[X]$. En utilisant l'une ou l'autre des conditions de la proposition précédente, il est immédiat que les réels 3 et -1 sont les racines de P .

3. Soit le polynôme $P = X^4 - X^3 + X^2 - X$ de $\mathbb{R}[X]$. Visiblement, on peut mettre X en facteur dans le polynôme P , et par conséquent le réel 0 est une racine de P . D'autre part, un simple calcul prouve que $\tilde{P}(1) = 0$. Donc 1 est une racine de P . Le polynôme P est donc divisible par $(X - 1)$.

Ce résultat se généralise au cas de $m \geq 1$ racines par des arguments d'arithmétique.

Proposition 5.8.2 Si $\alpha_1, \alpha_2, \dots, \alpha_m$ sont des racines deux à deux distinctes de $P \in \mathbb{K}[X]$, alors

$$\prod_{k=1}^m (X - \alpha_k) = (X - \alpha_1)(X - \alpha_2) \dots (X - \alpha_m) | P.$$

Preuve. La démonstration se fait par récurrence sur le nombre m de racines distinctes de P considérées. Pour tout $m \in \mathbb{N}^*$, soit la propriété $\mathcal{P}(m)$: " Si $P \in \mathbb{K}[X]$ admet m racines distinctes $\alpha_1, \alpha_2, \dots, \alpha_m$ de $\mathbb{K}$, alors P est divisible par $\prod_{k=1}^m (X - \alpha_k)$ ".

► La propriété vient d'être prouvée au rang 1 dans la proposition (5.8.1) précédente.

► On suppose que la propriété est vraie au rang m et prouvons-la au rang $m+1$. Soient $\alpha_1, \alpha_2, \dots, \alpha_m, \alpha_{m+1}$ des racines distinctes de P . Alors $\alpha_1, \alpha_2, \dots, \alpha_m$ sont m racines distinctes de P , donc par application de l'hypothèse de récurrence, il existe $Q \in \mathbb{K}[X]$ tel que

$$P = Q \prod_{k=1}^m (X - \alpha_k).$$

Comme α_{m+1} est une racine de P , on a

$$0_{\mathbb{K}} = \tilde{P}(\alpha_{m+1}) = \tilde{Q}(\alpha_{m+1}) \prod_{k=1}^m (\alpha_{m+1} - \alpha_k).$$

Comme pour tout $k \in \{1, \dots, m\}$, $\alpha_i \neq \alpha_{m+1}$, l'élément $(\alpha_{m+1} - \alpha_1)(\alpha_{m+1} - \alpha_2) \dots (\alpha_{m+1} - \alpha_m)$ est non nul et donc nécessairement

$$\tilde{Q}(\alpha_{m+1}) = 0_{\mathbb{K}},$$

c'est-à-dire α_{m+1} est une racine de Q . Appliquant la proposition (5.8.1) précédente, il existe $Q_1 \in \mathbb{K}[X]$ tel que

$$Q = (X - \alpha_{m+1})Q_1.$$

Par suite

$$P = Q \prod_{k=1}^m (X - \alpha_k) = Q_1 (X - \alpha_{m+1}) \prod_{k=1}^m (X - \alpha_k) = Q_1 \prod_{k=1}^{m+1} (X - \alpha_k).$$

Donc $\prod_{k=1}^{m+1} (X - \alpha_k)$ divise P , et $\mathcal{P}(m+1)$ est vraie. La proposition est alors prouvée par application du principe de récurrence. $\square$

Corollaire 5.8.1 *Si un polynôme P de $\mathbb{K}[X]$ s'annule en une infinité d'éléments de $\mathbb{K}$, alors $P = 0_{\mathbb{K}[X]}$.*

5.8.2 Ordre de multiplicité des racines d'un polynôme

Sur les trois exemples précédents, il est possible de faire la remarque suivante : on sait d'avance que si α est une racine dans $\mathbb{K}$ d'un polynôme P de $\mathbb{K}[X]$, alors le polynôme P est divisible par $(X - \alpha)$, mais il peut être divisible par une puissance de strictement supérieure à 1. Cela nous conduit à l'introduction de la notion d'ordre de multiplicité d'une racine d'un polynôme.

Définition 5.8.2 (Ordre de multiplicité des racines d'un polynôme)

Soient $P \in \mathbb{K}[X]$ (non nul), $\alpha \in \mathbb{K}$ et $m \in \mathbb{N}^*$.

1. On dit que α est une racine de multiplicité m (ou racine d'ordre m) de P si et si $(X - \alpha)^m$ divise P et $(X - \alpha)^{m+1}$ ne divise P .

► Si $m = 1$, on parle de racine simple.

► Si $m = 2$, on dit que α est racine double.

► Si $m = 3$, on dit que α est racine triple, et ainsi de suite.

2. On dit que α est une racine d'ordre **au moins** m de P si et si $(X - \alpha)^m$ divise P .

Remarque 5.8.1 *Par convention, on dira que α est racine de multiplicité 0 si α n'est pas racine de P . Lorsque la multiplicité est supérieure ou égale à 2, on parlera de racine multiple.*

Soit P un polynôme non nul et $\alpha \in \mathbb{K}$. Si α est une racine de P , alors l'ensemble

$$\left\{ k \in \mathbb{N}^* : (X - \alpha)^k \mid P \right\},$$

est une partie de $\mathbb{N}$ non vide (elle contient 1), et cet ensemble est fini car majoré par le degré de P d'après la remarque 5.7.1. Donc il admet un plus grand élément. Par conséquent, la multiplicité est toujours bien définie et majorée par le degré du polynôme. Ce qui nous conduit à la définition équivalente de multiplicité suivante.

Définition 5.8.3 *Soient $P \in \mathbb{K}[X]^*$, $\alpha \in \mathbb{K}$ et $m \in \mathbb{N}^*$. L'ordre de multiplicité de α est le plus grand entier m tel que $(X - \alpha)^m$ divise P , i.e.,*

$$m = \max \left\{ k \in \mathbb{N}^* : (X - \alpha)^k \mid P \right\}.$$

Remarque 5.8.2 *Soit P un polynôme non nul.*

- *L'ordre de multiplicité d'une racine de P est inférieur ou égal à $\deg(P)$.*
- *Puisqu'à chaque fois que α est une racine de multiplicité m d'un polynôme, on peut diviser celui-ci par $(X - \alpha)^m$, il en découle que le nombre de racines (comptées avec leur multiplicité) d'un polynôme de degré n est toujours inférieur ou égal à n .*

Exemple 5.8.2

1. Prenons $P = X^3 - 4X^2 + 5X - 2$ de $\mathbb{R}[X]$. On observe que $\alpha = 1$ est racine de P puisque $\tilde{P}(1) = 0$. Si on divise P par $(X - 1)$ on obtient $(X^2 - 3X + 2)$ et on observe que 1 est aussi racine de ce polynôme. Si on divise à nouveau par $(X - 1)$ on obtient $(X - 2)$, dont 1 n'est pas racine. Ainsi 1 est racine double de P (et on observe aussi que 2 est racine simple).

2. Soit $P = X^5 - 9X^4 + 25X^3 - 9X^2 - 54X + 54$ de $\mathbb{R}[X]$, alors $\alpha = 3$ est une racine d'ordre 3 du polynôme P . En effet :

On a

$$P = (X - 3) (X^4 - 6X^3 + 7X^2 + 12X - 18), \quad (X - 3) \mid P.$$

Puis

$$P = (X - 3)^2 (X^3 - 3X^2 - 2X + 6), \quad (X - 3)^2 \mid P.$$

et

$$P = (X - 3)^3 (X^2 - 2), \quad (X - 3)^3 \mid P.$$

mais

$$P = (X - 3)^4 (X + 3) + (189X - 63X^2 + 7X^3 - 189), \quad (X - 3)^4 \nmid P.$$

Proposition 5.8.3 (Caractérisation de l'ordre d'une racine)

Soient $P \in \mathbb{K}[X]^$, $\alpha \in \mathbb{K}$ et $m \in \mathbb{N}^*$. Alors on dit que α est une racine de multiplicité m de P si et seulement s'il existe un polynôme $Q \in \mathbb{K}[X]$ (qui n'admet pas de racine dans $\mathbb{K}$) tel que*

$$\begin{cases} P = (X - \alpha)^m Q \\ \text{et} \\ \tilde{Q}(\alpha) \neq 0_{\mathbb{K}}. \end{cases}$$

Preuve.

• Supposons que α est une racine de P d'ordre m . Comme $(X - \alpha)^m$ divise P , il existe $Q \in \mathbb{K}[X]$ tel que

$$P = (X - \alpha)^m Q.$$

Montrons que $\tilde{Q}(\alpha) \neq 0_{\mathbb{K}}$. Si c'était le cas, alors α serait une racine de Q et il existerait $Q_1 \in \mathbb{K}[X]$ tel que

$$Q = (X - \alpha)Q_1.$$

Par suite, on aurait

$$P = (X - \alpha)^{m+1} Q_1,$$

et $(X - \alpha)^{m+1}$ diviserait P , ce qui n'est, par hypothèse, pas possible. Donc $\tilde{Q}(\alpha) \neq 0_{\mathbb{K}}$.

• Maintenant supposons qu'il existe $Q \in \mathbb{K}[X]$ tel que

$$\begin{cases} P = (X - \alpha)^m Q \\ \text{et} \\ \tilde{Q}(\alpha) \neq 0_{\mathbb{K}}. \end{cases}$$

Pour montrer que α est une racine de multiplicité m de P , il faut montrer que $(X - \alpha)^{m+1}$ ne divise pas P . Par division Euclidienne de Q par $(X - \alpha)$, il existe $A, B \in \mathbb{K}[X]$ tels que

$$\begin{cases} Q = (X - \alpha) A + B \\ \text{et} \\ \deg(B) < \deg(X - \alpha) = 1. \end{cases}$$

Par conséquent $\deg(B) \geq 0$ et comme α n'est pas une racine de Q , B est un polynôme constant non nul. On a alors

$$P = (X - \alpha)^m Q = (X - \alpha)^m ((X - \alpha) A + B) = (X - \alpha)^{m+1} A + (X - \alpha)^m B.$$

Par unicité du couple quotient-reste dans la division Euclidienne de P par $(X - \alpha)^m$, on a $(X - \alpha)^m B$ est le reste de cette division et comme $B \neq 0_{\mathbb{K}[X]}$, ce reste est non nul. Par conséquent, $(X - \alpha)^{m+1}$ ne divise pas P . $\square$

Exemple 5.8.3 Le polynôme $P = X^5 - X^3 - X^2 + 1 \in \mathbb{R}[X]$ admet une racine simple qui est -1 et une racine double qui est 1 . En effet

$$P = (X + 1) \underbrace{(X^4 - X^3 - X + 1)}_{Q_1} \text{ avec } \tilde{Q}_1(-1) \neq 0,$$

et

$$P = (X - 1)^2 \underbrace{(X^3 + 2X^2 + 2X + 1)}_{Q_2} \text{ avec } \tilde{Q}_2(1) \neq 0.$$

Corollaire 5.8.2 $\alpha \in \mathbb{K}$ est une racine d'ordre **au moins** m de $P \in \mathbb{K}[X]$ si et seulement s'il existe un polynôme $Q \in \mathbb{K}[X]$ tel que

$$P = (X - \alpha)^m Q.$$

Proposition 5.8.4 Soit P un polynôme non nul de $\mathbb{K}[X]$. Si $\alpha_1, \alpha_2, \dots, \alpha_m$ de $\mathbb{K}$ sont des racines distinctes de P de multiplicités respectives $r_1, r_2, \dots, r_m$, alors il existe un polynôme $Q \in \mathbb{K}[X]$ tel que

$$P = (X - \alpha_1)^{r_1} (X - \alpha_2)^{r_2} \dots (X - \alpha_m)^{r_m} \cdot Q = \prod_{k=1}^m (X - \alpha_k)^{r_k} \cdot Q$$

avec $\tilde{Q}(\alpha_i) \neq 0_{\mathbb{K}}$ pour tout $i \in \{1, \dots, m\}$.

Preuve. Elle s'effectue par récurrence sur le nombre m de racines distinctes considérées. $\square$

Exemple 5.8.4 Reprenons l'exemple du polynôme $P = X^5 - X^3 - X^2 + 1 \in \mathbb{R}[X]$. Ce polynôme admet -1 pour racine simple et 1 pour racine double et

$$\exists Q \in \mathbb{R}[X] : P = (X + 1)(X - 1)^2 \underbrace{(X^2 + X + 1)}_Q,$$

avec $\tilde{Q}(-1) \neq 0$ et $\tilde{Q}(1) \neq 0$.

Remarque 5.8.3 Il y a deux manières de compter les racines d'un polynôme :

- Soit on dénombre les racines distinctes.
- Soit on compte chaque racine avec son ordre de multiplicité, c'est-à-dire qu'une racine α d'ordre m compte comme m racines.

Exemple 5.8.5

1. Le polynôme $P = (X - 1)(X + 1)^2(X - 2)^3$ de $\mathbb{R}[X]$ possède :
 - 3 racines distinctes : $-1, 1, 2$.
 - 6 racines comptées avec leur ordre de multiplicité : $-1, -1, 1, 2, 2, 2$.
2. Le polynôme $(X - 1)^2(X + 1)^3$ admet deux racines distinctes mais cinq racines si on compte chaque racine avec son ordre de multiplicité, car alors 1 compte comme 2 racines et -1 compte comme trois racines.

5.8.3 Utilisation des dérivées successives pour le calcul d'une multiplicité

L'utilisation des dérivées successives est une méthode efficace pour déterminer la multiplicité d'une racine d'un polynôme. Elle consiste à calculer les dérivées successives du polynôme et à les évaluer en la racine considérée. La multiplicité de la racine correspond au plus grand entier pour lequel les dérivées successives de l'ordre inférieur sont nulles, tandis que la dérivée d'ordre supérieur est non nulle.

Lemme 5.8.1 Soient $P \in \mathbb{K}[X], \alpha \in \mathbb{K}, m \in \mathbb{N}^*$. Si α est une racine d'ordre m de P alors α est une racine d'ordre $(m - 1)$ du polynôme dérivé P' .

- Cette propriété est utile pour déterminer la multiplicité des racines d'un polynôme en utilisant les dérivées successives.

- Le résultat est valable même si $m = 1$, puisqu'on a convenu qu'un élément de $\mathbb{K}$, qui n'est pas racine d'un polynôme, est racine d'ordre 0 de ce polynôme.

Preuve. Soit α une racine d'ordre m de $P \in \mathbb{K}[X]$. On peut écrire

$$P = (X - \alpha)^m Q, \quad \tilde{Q}(\alpha) \neq 0_{\mathbb{K}}.$$

Alors, en dérivant la relation précédente, on obtient

$$P' = (X - \alpha)^{m-1} (mQ + (X - \alpha)Q') = (X - \alpha)^{m-1} R,$$

avec

$$\tilde{R}(\alpha) = m\tilde{Q}(\alpha) \neq 0_{\mathbb{K}},$$

puisque $m \geq 1$ et $\tilde{Q}(\alpha) \neq 0_{\mathbb{K}}$. On a ainsi exhibé un polynôme $R \in \mathbb{K}[X]$ tel que $P' = (X - \alpha)^{m-1} R$ avec $\tilde{R}(\alpha) \neq 0_{\mathbb{K}}$, autrement dit on a montré que α est une racine d'ordre $(m-1)$ du polynôme dérivé P' . $\square$

La proposition suivante donne une condition nécessaire et suffisante pour qu'un scalaire α soit une racine de multiplicité r d'un polynôme.

Proposition 5.8.5 (Caractérisation de l'ordre de multiplicité d'une racine à l'aide des polynômes dérivés)

Soient $P \in \mathbb{K}[X]$, $\alpha \in \mathbb{K}$, $m \in \mathbb{N}^*$, alors

1. Pour que α soit racine de P d'ordre exactement m de P il faut et il suffit que

$$\begin{cases} \tilde{P}(\alpha) = \tilde{P}^{(1)}(\alpha) = \tilde{P}^{(2)}(\alpha) = \dots = \tilde{P}^{(m-1)}(\alpha) = 0_{\mathbb{K}}, \\ \text{et} \\ \tilde{P}^{(m)}(\alpha) \neq 0_{\mathbb{K}}. \end{cases}.$$

Ceci est équivalent à

$$\begin{cases} \forall k \in \{0, \dots, m-1\} : \tilde{P}^{(k)}(\alpha) = 0_{\mathbb{K}}, \\ \text{et} \\ \tilde{P}^{(m)}(\alpha) \neq 0_{\mathbb{K}}. \end{cases}$$

La multiplicité d'une racine de P est donc l'ordre de la première dérivée non nulle en α .

2. Pour que α soit racine de P d'ordre au moins m de P il faut et il suffit que

$$\tilde{P}(\alpha) = \tilde{P}^{(1)}(\alpha) = \tilde{P}^{(2)}(\alpha) = \dots = \tilde{P}^{(m-1)}(\alpha) = 0_{\mathbb{K}}.$$

Ceci est équivalent à

$$\forall k \in \{0, \dots, m-1\} : \tilde{P}^{(k)}(\alpha) = 0_{\mathbb{K}}.$$

Exemple 5.8.6

1. La multiplicité de $\alpha = 1 \in \mathbb{R}$ dans $P = X^4 + 3X^3 - 3X^2 - 7X + 6 \in \mathbb{R}[X]$ est égale à 2. En effet : On a

$$\tilde{P}(1) = 0.$$

En suite

$$P' = 4X^3 + 9X^2 - 6X - 7,$$

donc $\tilde{P}'(1) = 0$. Enfin

$$P'' = 12X^2 + 18X - 6,$$

donc $\tilde{P}''(1) = 24 \neq 0$.

2. Considérons le polynôme $P = X^3 - 3X + 2 \in \mathbb{R}[X]$. On a

$$P' = 3X^2 - 3 \text{ et } P'' = 6X.$$

P admet 1 pour racine double car $\tilde{P}(1) = \tilde{P}'(1) = 0$ et $\tilde{P}''(1) = 6 \neq 0$.

3. Soit $P = X^6 + X^5 + 3X^4 + 2X^3 + 3X^2 + X + 1$ un élément de $\mathbb{C}[X]$. On peut montrer que $\alpha = i \in \mathbb{C}$ est racine de P et déterminer son ordre de multiplicité. En effet, on obtient les résultats suivants

$$P(i) = 0.$$

$$P' = 6X^5 + 5X^4 + 12X^3 + 6X^2 + 6X + 1 \text{ et } \tilde{P}'(i) = 0.$$

$$P'' = 30X^4 + 20X^3 + 34X^2 + 12X + 6 \text{ et } \tilde{P}''(i) = -8i \neq 0,$$

donc i racine de P d'ordre 2.

Remarque 5.8.4

1. α est une racine simple de $P \iff \tilde{P}(\alpha) = 0_{\mathbb{K}}, \tilde{P}'(\alpha) \neq 0_{\mathbb{K}}$.
2. α est une racine double de $P \iff \tilde{P}(\alpha) = \tilde{P}'(\alpha) = 0_{\mathbb{K}}, \tilde{P}''(\alpha) \neq 0_{\mathbb{K}}$.
3. α est une racine triple de $P \iff \tilde{P}(\alpha) = \tilde{P}'(\alpha) = \tilde{P}''(\alpha) = 0_{\mathbb{K}}, \tilde{P}'''(\alpha) \neq 0_{\mathbb{K}}$.

5.8.4 Polynôme conjugué et racines complexes

Le concept de polynôme conjugué se réfère à un polynôme obtenu en prenant le conjugué complexe des coefficients d'un polynôme à coefficients complexes. Nous confondons ici polynôme P de $\mathbb{R}[X]$ et fonction polynomiale $\tilde{P}$.

Définition 5.8.4 (Polynôme conjugué)

Soit $P = \sum_{k=0}^{+\infty} a_k X^k \in \mathbb{C}[X]$ un polynôme à coefficients complexes. On appelle polynôme conjugué de P , le polynôme noté $\bar{P}$ et donné par

$$\bar{P} = \sum_{k=0}^{+\infty} \bar{a}_k X^k \in \mathbb{C}[X].$$

Le conjugué de ce polynôme est obtenu en prenant le conjugué complexe de chacun de ses coefficients.

Remarque 5.8.5

1. Il est clair que tout polynôme à coefficients réels est égal à son conjugué.
2. On vérifie facilement que si P et Q sont des éléments de $\mathbb{C}[X]$ et si $\lambda \in \mathbb{C}$, on a

$$\overline{P+Q} = \bar{P} + \bar{Q}, \overline{P \cdot Q} = \bar{P} \cdot \bar{Q}, \overline{\lambda P} = \bar{\lambda} \cdot \bar{P}$$

ces propriétés sont des conséquences immédiates des propriétés de la conjugaison.

3. Le polynôme B divise le polynôme A dans $\mathbb{C}[X]$ si, et seulement si $\bar{B}$ divise A . En effet, l'égalité $A = B \cdot Q$ est équivalente à $\bar{A} = \bar{B} \cdot \bar{Q}$.

Exemple 5.8.7 Si

$$P = (3 + 2i)X^2 + (1 - i)X + 4 \in \mathbb{C}[X],$$

alors le polynôme conjugué est

$$\bar{P} = (3 - 2i)X^2 + (1 + i)X + 4.$$

Le concept de polynôme conjugué et de racines conjuguées est étroitement lié au contexte des polynômes à coefficients complexes, et parfois à coefficients réels, lorsque les racines du polynôme ne sont pas toutes réelles.

Proposition 5.8.6 Soient P et Q deux polynômes de $\mathbb{C}[X]$, $\alpha \in \mathbb{C}, r \in \mathbb{N}$, on a

$$1. \overline{P(\alpha)} = \bar{P}(\bar{\alpha})$$

$$2. \overline{P^{(r)}} = (\bar{P})^{(r)}$$

$$3. P \in \mathbb{R}[X] \iff \forall z \in \mathbb{C} : \overline{P(z)} = P(\bar{z}).$$

Preuve. Démontrons par exemple les points 1 et 3.

1. Soient $\alpha \in \mathbb{C}$ et $P = \sum_{k=0}^m a_k X^k \in \mathbb{C}[X]$. On a

$$\overline{P(\alpha)} = \overline{\sum_{k=0}^m a_k \alpha^k} = \sum_{k=0}^m \overline{a_k \alpha^k} = \sum_{k=0}^m \bar{a}_k \bar{\alpha}^k = \bar{P}(\bar{\alpha}).$$

3. On a pour tout $z \in \mathbb{C}$

$$\overline{P(z)} - P(\bar{z}) = \sum_{k=0}^m \bar{a}_k \bar{z}^k - \sum_{k=0}^m a_k \bar{z}^k = \sum_{k=0}^m (\bar{a}_k - a_k) \bar{z}^k,$$

on obtient

$$\overline{\left(\overline{P(z)} - P(\bar{z}) \right)} = \sum_{k=0}^m (a_k - \bar{a}_k) z^k$$

d'où

$$\begin{aligned} \left(\forall z \in \mathbb{C} : \overline{P(z)} = P(\bar{z}) \right) &\iff \left(\forall z \in \mathbb{C} : \sum_{k=0}^m (a_k - \bar{a}_k) z^k = 0 \right) \\ &\iff \forall k \in \{0, 1, \dots, m\} : a_k - \bar{a}_k = 0 \\ &\iff \forall k \in \{0, 1, \dots, m\} : a_k = \bar{a}_k \\ &\iff P \in \mathbb{R}[X]. \end{aligned}$$

□

Proposition 5.8.7 Soient $P \in \mathbb{R}[X]$, $r \in \mathbb{N}^*$ et $\alpha \in \mathbb{C}$, alors

1. α est racine de P d'ordre exactement r si et seulement si $\bar{\alpha}$ est racine de P d'ordre exactement r .
2. α est racine de P d'ordre au moins r si et seulement si $\bar{\alpha}$ est racine de P d'ordre au moins r .

Preuve.

1. (i) Supposons que α soit racine de P d'ordre r , alors

$$P(\alpha) = P^{(1)}(\alpha) = \dots = P^{(r-1)}(\alpha) = 0_{\mathbb{R}[X]} \text{ et } P^{(r)}(\alpha) \neq 0_{\mathbb{R}[X]}.$$

Comme $P, P^{(1)}, \dots, P^{(r-1)}, P^{(r)}$ sont dans $\mathbb{R}[X]$, alors d'après la proposition 5.8.6, on a

$$\begin{cases} \forall k \in \{0, \dots, r-1\} : P^{(k)}(\bar{\alpha}) = \overline{P^{(k)}(\alpha)} = 0_{\mathbb{R}[X]}, \\ \text{et} \\ P^{(r)}(\bar{\alpha}) = \overline{P^{(r)}(\alpha)} \neq 0_{\mathbb{R}[X]}. \end{cases}$$

et donc $\bar{\alpha}$ est racine de P d'ordre r .

(ii) La réciproque se déduit de la partie directe en remplaçant α par $\bar{\alpha}$.

2. Un raisonnement analogue permet de conclure dans le cas de l'ordre au moins r .

□

Exemple 5.8.8 Considérons le polynôme

$$P = X^2 + 1 \in \mathbb{R}[X].$$

Les racines de P sont $\alpha_1 = i, \alpha_2 = -i \in \mathbb{C}$. Il est clair que α_1, α_2 sont d'ordre 1. En effet,

$$\begin{cases} P(\alpha_1) = 0, P'(\alpha_1) = 2i \neq 0, \\ P(\alpha_2) = 0, P'(\alpha_2) = -2i \neq 0. \end{cases}$$

Cet exemple montre que si i est une racine de P d'ordre 1, alors son conjugué $(-i)$ est également une racine de P d'ordre 1, confirmant ainsi la propriété.

Exemple 5.8.9 Déterminer les racines du polynôme

$$P = X^6 + X^5 + 3X^4 + 2X^3 + 3X^2 + X + 1 \in \mathbb{R}[X],$$

sachant que i en est une racine multiple.

On a $P(i) = P'(i) = 0$ mais $P''(i) = -8i \neq 0$. Le nombre i est donc une racine de P de multiplicité deux. Puisque P est à coefficients réels, $(-i)$ est également une racine de P de multiplicité deux. P est donc divisible par

$$(X - i)^2 (X + i)^2 = (X + 1)^2.$$

En posant la division euclidienne, on trouve

$$P = (X + 1)^2 (X^2 + X + 1).$$

Comme

$$X^2 + X + 1 = (X - j)(X - j^2)$$

où

$$j = \exp\left(\frac{2\pi i}{3}\right) = \frac{-1}{2} + i\frac{\sqrt{3}}{2} \in \mathbb{C}, j^2 = \frac{1}{j}.$$

Les racines de P sont $i, -i$ (racines doubles) et j, j^2 (racines simples).

Corollaire 5.8.3 Soient $P \in \mathbb{C}[X]$, $r \in \mathbb{N}^*$ et $\alpha \in \mathbb{C}$, alors α est racine de P d'ordre r si et seulement si $\bar{\alpha}$ est racine de $\bar{P}$ d'ordre r .

Preuve.

- Supposons que α soit racine de P d'ordre r , alors par définition on a

$$P(\alpha) = P^{(1)}(\alpha) = \dots = P^{(r-1)}(\alpha) = 0_{\mathbb{C}[X]} \text{ et } P^{(r)}(\alpha) \neq 0_{\mathbb{C}[X]}.$$

Considérons maintenant le polynôme conjugué $\bar{P}$. En prenant le conjugué complexe de chaque équation, nous obtenons

$$\begin{aligned} \bar{P}(\bar{\alpha}) &= \overline{P(\alpha)} = 0_{\mathbb{C}[X]} \\ \bar{P}^{(1)}(\bar{\alpha}) &= \overline{P^{(1)}(\alpha)} = 0_{\mathbb{C}[X]} \\ &\vdots \\ \bar{P}^{(r-1)}(\bar{\alpha}) &= \overline{P^{(r-1)}(\alpha)} = 0_{\mathbb{C}[X]} \\ \bar{P}^{(r)}(\bar{\alpha}) &= \overline{P^{(r)}(\alpha)} \neq 0_{\mathbb{C}[X]}. \end{aligned}$$

Cela prouve que $\bar{\alpha}$ est une racine de $\bar{P}$ d'ordre r .

- Réciproquement. Supposons que $\bar{\alpha}$ soit racine de $\bar{P}$ d'ordre r . Par un raisonnement analogue, on a

$$\bar{P}(\bar{\alpha}) = \bar{P}^{(1)}(\bar{\alpha}) = \dots = \bar{P}^{(r-1)}(\bar{\alpha}) = 0_{\mathbb{C}[X]} \text{ et } \bar{P}^{(r)}(\bar{\alpha}) \neq 0_{\mathbb{C}[X]}.$$

En prenant le conjugué complexe de ces équations, nous obtenons

$$\begin{aligned} P(\alpha) &= \overline{\bar{P}(\bar{\alpha})} = 0_{\mathbb{C}[X]} \\ P^{(1)}(\alpha) &= \overline{\bar{P}^{(1)}(\bar{\alpha})} = 0_{\mathbb{C}[X]} \\ &\vdots \\ P^{(r-1)}(\alpha) &= \overline{\bar{P}^{(r-1)}(\bar{\alpha})} = 0_{\mathbb{C}[X]} \\ P^{(r)}(\alpha) &= \overline{\bar{P}^{(r)}(\bar{\alpha})} \neq 0_{\mathbb{C}[X]}. \end{aligned}$$

Cela prouve que α est une racine de P d'ordre r . □

Exemple 5.8.10 *Considérons le polynôme suivant*

$$P = X^2 - (2 + 2i)X + 2i \in \mathbb{C}[X].$$

Le conjugué de P est $\bar{P}$ avec les coefficients conjugués complexes

$$\bar{P} = X^2 - (2 - 2i)X - 2i.$$

Maintenant, considérons la racine $\alpha = 1 + i \in \mathbb{C}$ et vérifions que α est une racine de P d'ordre 2. Alors

$$P(\alpha) = (1 + i)^2 - (2 + 2i)(1 + i) + 2i = 0.$$

D'autre part, on a

$$\begin{aligned} P' &= 2X - (2 + 2i). \\ P'' &= 2. \end{aligned}$$

On obtient

$$P'(\alpha) = 0 \text{ et } P''(\alpha) = 2 \neq 0.$$

Par conséquent α est une racine de P d'ordre 2.

Maintenant, vérifions que $\bar{\alpha} = 1 - i \in \mathbb{C}$ est une racine de $\bar{P}$ d'ordre 2. Alors

$$\bar{P}(\bar{\alpha}) = (1 - i)^2 - (2 - 2i)(1 - i) - 2i = 0.$$

D'autre part, on a

$$\begin{aligned} \bar{P}' &= 2X - (2 - 2i). \\ \bar{P}'' &= 2. \end{aligned}$$

On obtient

$$\bar{P}'(\bar{\alpha}) = 0 \text{ et } \bar{P}''(\bar{\alpha}) = 2 \neq 0.$$

Par conséquent $\bar{\alpha}$ est une racine de $\bar{P}$ d'ordre 2.

Remarque 5.8.6 *Si $P \in \mathbb{C}[X]$ et α une racine complexe de P , alors $\bar{\alpha}$ n'est pas nécessairement une racine de P .*

Exemple 5.8.11 *Prenons par exemple*

$$P(x) = X^2 - (1 + i)X + i \in \mathbb{C}[X].$$

Les racines sont donc

$$\alpha_1 = \frac{(1 + i) - \sqrt{-2i}}{2}, \alpha_2 = \frac{(1 + i) + \sqrt{-2i}}{2}.$$

Pour simplifier, notons que

$$\sqrt{-2i} = \sqrt{2} \cdot \sqrt{-i} = \sqrt{2} \cdot \left(\frac{1 - i}{\sqrt{2}} \right) = 1 - i.$$

Ainsi, les racines sont

$$\alpha_1 = 1, \alpha_2 = i.$$

Le conjugué de $\alpha_1 = 1$ est toujours 1, qui est une racine de P . Le conjugué de $\alpha_2 = i$ est $\bar{\alpha}_2 = -i$ avec $\bar{\alpha}_2$ n'est pas une racine de P .

Cet exemple montre que pour un polynôme à coefficients complexes, une racine α n'implique pas nécessairement que son conjugué $\bar{\alpha}$ soit aussi une racine de ce polynôme.

5.9 Polynômes Irréductibles et Factorisation

5.9.1 Polynômes Irréductibles

La notion de polynôme irréductible dépend du corps de coefficients dans lequel on travaille. Un polynôme peut être irréductible dans un corps donné, mais devenir réductible lorsqu'on change de corps. Il est donc crucial de préciser le contexte (par exemple, $\mathbb{R}[X]$ ou $\mathbb{C}[X]$) pour déterminer l'irréductibilité d'un polynôme.

Le théorème fondamental de l'algèbre, également connu sous le nom de théorème de d'Alembert-Gauss, joue un rôle central dans l'étude de l'irréductibilité des polynômes. Ce théorème affirme que tout polynôme non constant à coefficients complexes possède au moins une racine complexe. À partir de ce théorème, on peut déduire plusieurs propriétés essentielles concernant la factorisation des polynômes à coefficients réels ou complexes. Le théorème suivant est admis.

Théorème 5.9.1 (de d'Alembert-Gauss).

Tout polynôme de $\mathbb{C}[X]$ de degré $n \geq 1$ admet au moins une racine dans $\mathbb{C}$.

Preuve. Admise. Il existe de nombreuses démonstrations. Aucune n'est assez élémentaire pour être exposée ici. $\square$

Remarque 5.9.1 *Attention ce théorème est faux dans $\mathbb{R}$. Par exemple $P = X^2 + 1$ est non constant mais ne possède aucune racine dans $\mathbb{R}$.*

Une formulation équivalente du théorème fondamental de l'algèbre est la suivante.

Corollaire 5.9.1 *Un polynôme $P \in \mathbb{C}[X]$ de degré n possède n racines (comptées avec leur multiplicité) dans $\mathbb{C}$.*

Définition 5.9.1 (**Polynôme irréductible**).

Soit P un polynôme non constant (de degré ≥ 1) de $\mathbb{K}[X]$.

1. *On dit que P est irréductible dans $\mathbb{K}[X]$ si ses seuls diviseurs dans $\mathbb{K}[X]$ sont :*
 - *Les polynômes constants non nuls (les éléments de $\mathbb{K}^*$).*
 - *Les polynômes associés à P , c'est-à-dire les λP , avec $\lambda \in \mathbb{K}^*$.*
2. *Un polynôme qui n'est pas irréductible est dit réductible.*

Remarque 5.9.2

1. *La notion de polynôme irréductible pour l'arithmétique de $\mathbb{K}[X]$ correspond à la notion de nombre premier pour l'arithmétique de $\mathbb{Z}$.*
2. *Un polynôme irréductible est un polynôme $P \in \mathbb{K}[X]$ non constant et tels que*

$$\forall A, B \in \mathbb{K}[X] : P = AB \implies \deg(A) = 0 \text{ ou } \deg(B) = 0.$$

Un polynôme irréductible sur $\mathbb{K}$ est donc incassable par division dans $\mathbb{K}[X]$.

2. *Un polynôme P est réductible (non irréductible) :*

- *S'il est constant.*
- *Ou s'il peut s'écrire $P = AB$, avec $\deg(A) \geq 1$ et $\deg(B) \geq 1$ ou encore avec $1 \leq \deg(A) \leq \deg(P) - 1$.*

Le lemme d'Euclide pour les polynômes est une version du lemme d'Euclide pour les entiers, adaptée au cadre des polynômes. Il affirme que si un polynôme irréductible divise le produit de deux polynômes, alors il divise l'un des deux polynômes.

Proposition 5.9.1 (Lemme d'Euclide)

Soit $P \in \mathbb{K}[X]$ un polynôme irréductible et soient A, B deux polynômes de $\mathbb{K}[X]$. Si $P \mid AB$ alors $P \mid A$ ou $P \mid B$.

Preuve. Supposons que P soit irréductible et que P divise AB . Alors

- Si P divise A , dans ce cas, il n'y a rien à prouver, car l'énoncé du lemme est vérifié.
- Si P ne divise pas A , alors cela signifie que le plus grand commun diviseur $\text{pgcd}(P, A)$ est un polynôme constant. Plus précisément, comme P est irréductible, $\text{pgcd}(P, A) = 1$. D'après le théorème de Bézout, il existe deux polynômes U et V dans $\mathbb{K}[X]$ tels que

$$UP + VA = 1.$$

En multipliant cette égalité par B , on obtient

$$UPB + VAB = B.$$

Puisque P divise AB , il divise aussi VAB . Il en résulte que P divise UPB , et donc P divise aussi la somme

$$UPB + VAB = B.$$

Par conséquent, si $P \mid AB$ alors $P \mid A$ ou $P \mid B$. □

Dans l'anneau des polynômes à coefficients complexes $\mathbb{C}[X]$, les polynômes irréductibles sont très simples à décrire grâce au théorème fondamental de l'algèbre. Cela implique que les polynômes irréductibles dans $\mathbb{C}[X]$ sont de degré 1. La situation dans $\mathbb{R}[X]$, l'anneau des polynômes à coefficients réels, est plus complexe. Un polynôme réel irréductible peut avoir un degré supérieur à 1.

Les polynômes irréductibles jouent un rôle central dans l'étude de $\mathbb{K}[X]$ puisque tout polynôme se décompose de manière unique (à un facteur inversible près) en un produit de polynômes irréductibles

Proposition 5.9.2 [?](Polynômes irréductibles de $\mathbb{C}[X]$ et $\mathbb{R}[X]$)

1. Les polynômes irréductibles de $\mathbb{C}[X]$ sont les polynômes de degré 1, c'est-à-dire les polynômes $aX + b$ avec $a \neq 0$.
2. Les polynômes irréductibles de $\mathbb{R}[X]$ sont les polynômes de degré 1 et les polynômes de degré 2 (les trinômes $aX^2 + bX + c$ avec $a \neq 0$) de discriminant strictement négatif ($\Delta < 0$).

Exemple 5.9.1

- Le polynôme $P_1 = X - 3$ de degré 1 est irréductible dans $\mathbb{R}[X]$.
- Le polynôme $P_2 = X^2 + 1$ de degré 2 avec discriminant négatif $\Delta = -4$ est irréductible dans $\mathbb{R}[X]$.
- Le polynôme $P_3 = X^2 - 4$ de degré 2 avec discriminant positif $\Delta = 16$ est réductible dans $\mathbb{R}[X]$.
- Le polynôme $P_4 = X - 4$ de degré 1 est irréductible dans $\mathbb{C}[X]$.
- Le polynôme $P_5 = X^2 + 1$ de degré 2 réductible dans $\mathbb{C}[X]$.
- Le polynôme $P_6 = X^2 - 2 = (X - \sqrt{2})(X + \sqrt{2})$ est réductible dans $\mathbb{R}[X]$ mais est irréductible dans $\mathbb{Q}[X]$.

Remarque 5.9.3 La notion de polynôme irréductible dépend du corps $\mathbb{K}$. Ainsi $P = X^2 - 2$ est irréductible dans $\mathbb{Q}[X]$, mais il ne l'est pas dans $\mathbb{R}[X]$. De même $P = X^2 + 1$ est irréductible dans $\mathbb{R}[X]$, mais il ne l'est pas dans $\mathbb{C}[X]$.

5.9.2 Factorisation dans $\mathbb{C}[X]$

La décomposition des polynômes en produits de facteurs irréductibles est une opération fondamentale en algèbre. Elle consiste à exprimer un polynôme comme le produit de polynômes irréductibles, c'est-à-dire des polynômes qui ne peuvent pas être factorisés davantage dans un corps ou un anneau donné. Cette décomposition est analogue à la factorisation d'un nombre entier en produits de nombres premiers.

Dans l'ensemble des polynômes à coefficients complexes, noté $\mathbb{C}[X]$, la décomposition en facteurs irréductibles est garantie par le théorème fondamental de l'algèbre. Ce théorème stipule que tout polynôme non constant à coefficients complexes possède au moins une racine complexe, permettant ainsi sa factorisation en un produit de polynômes de degré 1 (linéaires).

Proposition 5.9.3 (Factorisation dans $\mathbb{C}[X]$)

1. Soit $P \in \mathbb{C}[X]$ un polynôme non nul à coefficients complexes. Soit $\alpha_1, \alpha_2, \dots, \alpha_s$ les racines complexes deux à deux distinctes de P et $r_1, r_2, \dots, r_s$ leurs ordres de multiplicité respectifs. On a alors

$$P = \lambda (X - \alpha_1)^{r_1} \cdot (X - \alpha_2)^{r_2} \cdot \dots \cdot (X - \alpha_s)^{r_s} = \lambda \prod_{k=1}^s (X - \alpha_k)^{r_k}.$$

où $\lambda \in \mathbb{C}^*$ est le coefficient dominant de P , $s \in \mathbb{N}$, $r_i \in \mathbb{N}^*$.

2. Cette écriture est unique à l'ordre près des facteurs.

3. Tout polynôme non constant peut être décomposé de manière unique, à l'ordre près des facteurs, en produit de polynômes irréductibles dans $\mathbb{C}[X]$

Exemple 5.9.2

1. Exemple avec racines complexes conjuguées. Considérons le polynôme

$$P = X^4 + 1.$$

Pour décomposer P dans $\mathbb{C}[X]$, on cherche $z \in \mathbb{C}$ qui vérifie $z^4 = -1$. On pose $z = re^{i\theta}$, et on obtient $r = 1$ et $4\theta = \pi + 2k\pi$ avec $k \in \mathbb{Z}$. Les racines de $X^4 + 1$ sont donc

$$z_0 = e^{i\frac{\pi}{4}}, \bar{z}_0 = e^{-i\frac{\pi}{4}}, z_1 = e^{i\frac{3\pi}{4}}, \bar{z}_1 = e^{-i\frac{3\pi}{4}}.$$

Ainsi P se factorise dans $\mathbb{C}[X]$

$$P = X^4 + 1 = (X - e^{i\frac{\pi}{4}})(X - e^{-i\frac{\pi}{4}})(X - e^{i\frac{3\pi}{4}})(X - e^{-i\frac{3\pi}{4}}).$$

Les facteurs irréductibles sont donc $(X - e^{i\frac{\pi}{4}})$, $(X - e^{-i\frac{\pi}{4}})$, $(X - e^{i\frac{3\pi}{4}})$, et $(X - e^{-i\frac{3\pi}{4}})$.

2. Exemple avec racines multiples. Considérons le polynôme

$$P = X^4 - 2X^2 + 1.$$

Les racines sont $1, -1$ (chaque racine a une multiplicité de 2). On obtient

$$P = (X - 1)^2 (X + 1)^2.$$

Les facteurs irréductibles sont donc $(X - 1)$ et $(X + 1)$ avec leur multiplicité respective.

3. Exemple avec racines complexes. Considérons le polynôme

$$P(X) = X^3 - 1.$$

Les racines sont

$$\alpha_1 = 1, \alpha_2 = e^{2i\frac{\pi}{3}} = \frac{-1 + i\sqrt{3}}{2}, \bar{\alpha}_2 = e^{-2i\frac{\pi}{3}} = -\frac{1 + i\sqrt{3}}{2}.$$

Ainsi P se factorise dans $\mathbb{C}[X]$,

$$P = X^3 - 1 = (X - 1)(X - e^{2i\frac{\pi}{3}})(X - e^{-2i\frac{\pi}{3}}).$$

Les facteurs irréductibles sont donc $(X - 1)$, $(X - e^{2i\frac{\pi}{3}})$, et $(X - e^{-2i\frac{\pi}{3}})$.

5.9.3 Factorisation dans $\mathbb{R}[X]$

Dans $\mathbb{R}[X]$, l'anneau des polynômes à coefficients réels, la décomposition en facteurs irréductibles est plus complexe que dans $\mathbb{C}[X]$ en raison des propriétés des racines réelles et complexes.

Proposition 5.9.4 (Factorisation dans $\mathbb{R}[X]$)

Tout polynôme non nul P de $\mathbb{R}[X]$, peut s'écrire sous la forme

$$P = \lambda \prod_{j=1}^s (X - \alpha_j)^{r_j} \cdot \prod_{k=1}^t (X^2 + b_k X + c_k)^{s_k},$$

où $\lambda \in \mathbb{R}^*$ est le coefficient dominant de P , $s, t \in \mathbb{N}^*$ et, pour tout $j \in \{1, 2, \dots, s\}$ et pour tout $k \in \{1, 2, \dots, t\}$, r_j et s_k des entiers naturels non nuls et α_j, b_k et c_k des réels vérifiant $b_k^2 - 4c_k < 0$. En outre cette écriture est unique à l'ordre des facteurs près.

• Ainsi, dans $\mathbb{R}[X]$ tout polynôme peut se décomposer en produits de facteurs de degré 1 (pour les racines réelles) et de degré 2 (pour les polynômes irréductibles n'ayant pas de racines réelles).

Remarque 5.9.4 Si un polynôme $P \in \mathbb{R}[X]$ possède les racines réelles $\alpha_1, \dots, \alpha_s$ de multiplicités $r_1, r_2, \dots, r_s$ et les racines complexes de partie imaginaire > 0 , $\beta_1, \beta_2, \dots, \beta_t$ de multiplicités $s_1, s_2, \dots, s_t$, il a pour racines de partie imaginaire < 0 les complexes $\bar{\beta}_1, \bar{\beta}_2, \dots, \bar{\beta}_t$ de mêmes multiplicités $s_1, s_2, \dots, s_t$, et il se factorise ainsi dans $\mathbb{C}[X]$,

$$P = \lambda \prod_{j=1}^s (X - \alpha_j)^{r_j} \cdot \prod_{k=1}^t (X - \beta_k)^{s_k} (X - \bar{\beta}_k)^{s_k}.$$

D'où la factorisation sur $\mathbb{R}[X]$,

$$P = \lambda \prod_{j=1}^s (X - \alpha_j)^{r_j} \cdot \prod_{k=1}^t (X^2 - 2\operatorname{Re}(\beta_k)X + |\beta_k|^2)^{s_k}.$$

Exemple 5.9.3

1. Exemple avec des racines réelles. Considérons le polynôme

$$P = X^3 - 6X^2 + 11X - 6.$$

Les racines réelles sont 1, 2, et 3. Ainsi P se factorise dans $\mathbb{R}[X]$,

$$P = X^3 - 6X^2 + 11X - 6 = (X - 1)(X - 2)(X - 3).$$

Les facteurs irréductibles sont donc $(X - 1)$, $(X - 2)$, et $(X - 3)$.

2. Exemple avec des racines complexes. Considérons le polynôme

$$P = X^3 - 1.$$

Pour décomposer $P = X^3 - 1$ dans $\mathbb{R}[X]$, on commence par le décomposer dans $\mathbb{C}[X]$, donc on peut utiliser les racines troisièmes de l'unité 1,

$$P = X^3 - 1 = (X - 1) \left(X - e^{i\frac{2\pi}{3}} \right) \left(X - e^{-i\frac{2\pi}{3}} \right).$$

Puis on regroupe les facteurs ayant des racines conjuguées, cela conduit à un polynôme à coefficients réels,

$$P = (X - 1) \left(X^2 - 2X \cos \left(\frac{2\pi}{3} \right) + 1 \right) = (X - 1) (X^2 + X + 1).$$

qui est la factorisation de P dans $\mathbb{R}[X]$. Les facteurs irréductibles sont donc $(X - 1)$, et $(X^2 + X + 1)$.

3. Exemple avec racines complexes conjuguées. Pour $P = X^4 + 1$, ce polynôme n'a pas de racines réelles. On rappelle sa factorisation dans $\mathbb{C}[X]$ vue dans l'exemple précédent

$$P = X^4 + 1 = (X - e^{i\frac{\pi}{4}}) (X - e^{-i\frac{\pi}{4}}) (X - e^{3i\frac{\pi}{4}}) (X - e^{-3i\frac{\pi}{4}}).$$

Dans la décomposition ci-dessus, on regroupe les facteurs ayant des racines conjuguées, cela conduit à un polynôme à coefficients réels

$$\begin{aligned} P = X^4 + 1 &= \left(X^2 - 2X \cos \left(\frac{\pi}{4} \right) + 1 \right) \left(X^2 - 2X \cos \left(\frac{3\pi}{4} \right) + 1 \right) \\ &= (X^2 - \sqrt{2}X + 1) (X^2 + \sqrt{2}X + 1), \end{aligned}$$

qui est la factorisation de P dans $\mathbb{R}[X]$. Les facteurs irréductibles sont donc $(X^2 - \sqrt{2}X + 1)$, et $(X^2 + \sqrt{2}X + 1)$.

• Une autre méthode consiste plus simplement à écrire

$$P = X^4 + 1 = (X^2 + 1)^2 - 2X^2,$$

qui est une identité remarquable que l'on sait factoriser. Ainsi, on obtient une nouvelle fois,

$$P = (X^2 + 1)^2 - 2X^2 = (X^2 - \sqrt{2}X + 1) (X^2 + \sqrt{2}X + 1).$$

Comme $X^4 + 1$ n'a pas de racines réelles, ces deux polynômes du second degré n'ont pas de racines réelles, et donc leur discriminant est strictement négatif.

Remarque 5.9.5

1. La factorisation des polynômes dans $\mathbb{R}[X]$ implique de trouver et d'extraire les racines réelles et complexes, puis de décomposer le polynôme en produits de facteurs linéaires et quadratiques irréductibles. Les polynômes de degré supérieur à 2 peuvent être irréductibles ou factorisables en produit de polynômes de degré 1 ou 2.

2. Comment passer de $\mathbb{C}$ à $\mathbb{R}$? Pour trouver la factorisation de $P \in \mathbb{R}[X]$ en produit de polynôme irréductibles sur $\mathbb{R}$ à partir de sa factorisation en produit de polynômes irréductibles sur $\mathbb{C}$, on regroupe les racines non réelles du polynôme P par paires de conjugués $(\alpha, \bar{\alpha})$ afin d'obtenir un polynôme à coefficients réels

$$(X - \alpha) (X - \bar{\alpha}) = (X^2 - 2\operatorname{Re}(\alpha)X + |\alpha|^2).$$

3. En conclusion, la factorisation des polynômes dépend des coefficients :

- Pour $\mathbb{C}[X]$, les polynômes se factorisent en produits de polynômes linéaires.
- Pour $\mathbb{R}[X]$, les polynômes se factorisent en produits de polynômes linéaires et quadratiques irréductibles.

Bibliographie

- [1] S. Balac and F. Sturm. *Algèbre et analyse : cours de mathématiques de première année avec exercices corrigés*. EPFL Press, 2003.
- [2] C. Bertrand. *Ensemble, Relations, Applications, Dénombrement - Exercices corrigés avec rappels de cours. L1, L2, L3, classes préparatoires*. Cépaduès, 2009.
- [3] A. Bégyn. *Mathématiques ECS 1re année : L'essentiel du cours avec exemples*. Ellipses, 2016.
- [4] H. Boualem, et al. *Mathématiques L1 : cours complet avec 1000 tests et exercices corrigés*. Pearson education, 2006.
- [5] J. Calais. *Éléments de théorie des anneaux : anneaux commutatifs ; niveau L3*. Ellipses Éd. Marketing, 2006.
- [6] J. F. Caulier. *Fondements mathématiques : Pour l'économie et la gestion*. De Boeck Supérieur, 2014.
- [7] D. Delaunay. *Exercices d'algèbre et de probabilités MPSI*. De Boeck Supérieur, 2017.
- [8] J. Delcourt. *Théorie des groupes : rappels de cours et exercices corrigés*. Dunod, Paris, 2007.
- [9] J. Delfaud and M. Boukhobza. *Tout ce qu'il faut savoir sur les mathématiques en MPSI et MP2I : cours complet avec démonstrations, 174 méthodes, 260 exemples détaillés et 464 exercices d'entraînement corrigés*. Paris : Ellipses, 2021.
- [10] C. Deschamps, et al. *Mathématiques tout-en-un : 5e édition*. Dunod, 2018.
- [11] J. Duparc. *La logique pas à pas*. Polytechniques et universitaires romandes, 2015.
- [12] J. Franchini and J. C. Jacquens. *Cours particuliers de mathématiques pour l'agrégation interne*. Ellipses, 2010.
- [13] O. Halgand and A. Morra. *ExoMaths : MPSI : exercices corrigés pour comprendre et réussir*. Paris : Ellipses, 2018.
- [14] A. Jeanneret and D. Lines. *Invitation à l'algèbre : Théorie des groupes, des anneaux, des corps et des modules*. Cépaduès, 2008.
- [15] G. Maurice. *Exercices et problèmes d'algèbre : Rappels de cours, Exercices et problèmes avec corrigés détaillés, Fiches de révision Broché*. Dunod, 2008.
- [16] J. M. Monier. *Algèbre 1. 1e année MPSI, PCSI, PTSI Cours et 600 exercices corrigés*. Dunod, Paris, 1996.
- [17] S. Rainero. *Les maths en cours - MPSI. Cours complet et détaillé, enrichi de nombreux exemples*. Ellipses, 2015.