

**Licence – Mathématiques**  
**Algèbre 2**

NOTES DE COURS

# 1 Arithmétique élémentaire

## 1.1 Nombres entiers

### 1.1.1 Quelques propriétés élémentaires

Nous ne donnerons pas ici de définition formelle des nombres entiers, nous nous contenterons de dire qu'il existe un ensemble  $\mathbb{N}$  dont les éléments sont appelés *entiers naturels*, et que cet ensemble est muni d'opérations

$$+ : \begin{array}{ccc} \mathbb{N} \times \mathbb{N} & \rightarrow & \mathbb{N} \\ (a, b) & \mapsto & a + b \end{array} \quad (\text{addition}) \qquad \cdot : \begin{array}{ccc} \mathbb{N} \times \mathbb{N} & \rightarrow & \mathbb{N} \\ (a, b) & \mapsto & a.b \end{array} \quad (\text{multiplication})$$

vérifiant les propriétés suivantes :

- (associativité de  $+$ )  $\forall a, b, c \in \mathbb{N}, (a + b) + c = a + (b + c)$  ;
- (associativité de  $\cdot$ )  $\forall a, b, c \in \mathbb{N}, (a.b).c = a.(b.c)$  ;
- (commutativité de  $+$ )  $\forall a, b \in \mathbb{N}, a + b = b + a$  ;
- (commutativité de  $\cdot$ )  $\forall a, b \in \mathbb{N}, a.b = b.a$  ;
- (élément neutre pour  $+$ )  $\exists 0 \in \mathbb{N}, \forall a \in \mathbb{N}, a + 0 = 0 + a = a$  ;
- (élément neutre pour  $\cdot$ )  $\exists 1 \in \mathbb{N}, \forall a \in \mathbb{N}, a.1 = 1.a = a$  ;
- (distributivité de  $\cdot$  sur  $+$ )  $\forall a, b, c \in \mathbb{N}, a.(b + c) = a.b + a.c$  et  $(a + b).c = a.c + b.c$ .

Le nombre 0 joue un rôle très particulier vis-à-vis de la multiplication car il vérifie :

- (absorbance de l'élément 0)  $\forall a \in \mathbb{N}, 0.a = 0$  ;
- (intégrité)  $\forall a, b \in \mathbb{N}, a.b = 0 \Rightarrow a = 0$  ou  $b = 0$ .

De ce dernier point, on déduit le principe de simplification affirmant que si  $a.b = a.c$  avec  $a, b, c \in \mathbb{N}$  et  $a \neq 0$ , alors  $b = c$ .

Par convention, on note  $\mathbb{N}^* := \mathbb{N} \setminus \{0\}$ .

Sur  $\mathbb{N}$ , il existe une relation d'ordre total définie, pour tout  $a, b \in \mathbb{N}$ , par

$$a \geq b \Leftrightarrow \exists c \in \mathbb{N}, a = b + c.$$

Cette relation d'ordre vérifie les propriétés fondamentales suivantes :

- (compatibilité de  $\geq$  avec  $+$ )  $\forall a, b, c \in \mathbb{N}, a \geq b \Rightarrow a + c \geq b + c$  ;
- (compatibilité de  $\geq$  avec  $\cdot$ )  $\forall a, b, c \in \mathbb{N}, a \geq b \Rightarrow a.c \geq b.c$  ;
- (principe du plus petit élément) tout ensemble  $\Omega \subset \mathbb{N}$  non vide possède un plus petit élément ; en particulier, 0 est le plus petit élément de  $\mathbb{N}$  et 1 le plus petit élément de  $\mathbb{N}^*$  ;
- tout ensemble  $\Omega \subset \mathbb{N}$  non vide et majoré possède un plus grand élément ;
- (principe archimédien)  $\forall a \in \mathbb{N}, \forall b \in \mathbb{N}^*, \exists n \in \mathbb{N}, n.b > a$ .

Un corollaire très utile est le suivant :

**Corollaire.** Pour tout  $a, b \in \mathbb{N}$ ,  $a > b \Leftrightarrow a \geq b + 1$ .

Ce corollaire est au cœur du principe de raisonnement suivant :

**Principe** (de récurrence). Si une proposition dépendant d'un paramètre entier  $n$  est vraie pour une valeur  $n_0$  et que sa véracité en  $n$  implique sa véracité en  $n + 1$ , alors elle est vraie pour tout entier plus grand que  $n_0$ . En particulier, si  $n_0 = 0$ , alors elle est vraie pour tout entier naturel.

parfois renforcé en :

**Principe** (de récurrence généralisée). Si une proposition dépendant d'un paramètre entier  $n$  est vraie pour une valeur  $n_0$  et que sa véracité pour tout entier  $n_0 \leq k < n$  implique sa véracité en  $n$ , alors elle est vraie pour tout entier plus grand que  $n_0$ . En particulier, si  $n_0 = 0$ , alors elle est vraie pour tout entier naturel.

On admettra également que  $\mathbb{N}$  peut être étendu en un ensemble  $\mathbb{Z}$ , dont les éléments sont appelés *entiers relatifs* tels que les opérations  $+$  et  $\cdot$  et la relation d'ordre  $\geq$  s'étendent à  $\mathbb{Z}$  et vérifient :

- associativité, commutativité, éléments neutres et distributivité, absorbance de 0, intégrité, compatibilité de  $\geq$  avec  $+$  ;
- (éléments opposés)  $\forall a \in \mathbb{Z}, \exists -a \in \mathbb{Z}, a + (-a) = 0$  ;
- $\forall a, b \in \mathbb{Z}, a \geq b \Rightarrow -b \geq -a$ .

Plus explicitement, on a  $\mathbb{Z} = \mathbb{N} \sqcup \{-a \mid a \in \mathbb{N}^*\}$ . Cela permet de définir l'application valeur absolue suivante :

$$\begin{array}{ccc} \mathbb{Z} & \rightarrow & \mathbb{N} \\ | \cdot | : & a \mapsto & \begin{cases} a & \text{si } a \in \mathbb{N} \\ -a & \text{si } a \notin \mathbb{N} \end{cases} \end{array} .$$

On note encore par convention  $\mathbb{Z}^* := \mathbb{Z} \setminus \{0\}$ .

### 1.1.2 Division euclidienne

Jusqu'ici, nous avons parlé d'addition et de multiplication, implicitement de différence (qui est la somme du premier terme avec l'opposé du second) mais pas de division. Cette dernière n'est en effet pas bien définie sur les entiers comme opération qui à deux entiers en associe un troisième, mais on peut la raffiner en une opération qui, à deux entiers, en associe deux autres.

**Proposition 1.1.1.** Soit  $a \in \mathbb{N}$  et  $b \in \mathbb{N}^*$ , il existe d'unique entiers  $q \in \mathbb{N}$  et  $r \in \llbracket 0, b - 1 \rrbracket$  tels que  $a = q.b + r$ .

*Démonstration.* Commençons par montrer l'existence. Pour cela on considère l'ensemble  $A := \{k \in \mathbb{N} \mid k.b > a\}$  qui, d'après le principe archimédien, n'est pas vide et possède donc un plus petit élément  $k_0$ . On a  $k_0 \neq 0$  car clairement  $0 \notin A$ , et donc  $k_0 \geq 1$ . On pose alors  $q = k_0 - 1 \in \mathbb{N}$  et  $r = a - q.b$ . Par minimalité de  $k_0$ , on a  $q.b \leq a$  et donc  $r \geq 0$ ; et comme  $k_0 \in A$ , on a  $(q + 1).b = k_0.b > a$  donc  $b > a - q.b = r$ , autrement dit  $r \leq b - 1$ . On a donc bien  $r \in \llbracket 0, b - 1 \rrbracket$ .

Montrons maintenant l'unicité. Pour cela on considère  $q_1, q_2 \in \mathbb{N}$  et  $r_1, r_2 \in \llbracket 0, b - 1 \rrbracket$  tels que  $q_1.b + r_1 = q_2.b + r_2$ , et quitte à échanger le rôle des indices, on suppose par l'absurde que  $q_1 > q_2$ , autrement dit que  $q_1 \geq q_2 + 1$ . On a alors  $r_2 - r_1 = (q_1 - q_2).b$ . Or  $r_2 - r_1 \leq r_2 - 0 = r_2 < b - 1$ , et  $(q_1 - q_2).b \geq b$ , on en déduit que  $b - 1 > b$ , ce qui est absurde.  $\square$

**Corollaire 1.1.2.** Soit  $a \in \mathbb{Z}$  et  $b \in \mathbb{Z}^*$ , il existe d'unique entiers  $q \in \mathbb{Z}$  et  $r \in \llbracket 0, |b| - 1 \rrbracket$  tels que  $a = q.b + r$ .

*Démonstration.* Commençons par remarquer que la preuve de l'unicité peut être réutilisée sans aucune modification. Concernant l'existence :

**Si  $a, b \geq 0$  :** c'est la proposition précédente ;

**Si  $a < 0$  et  $b > 0$  :** on utilise la proposition précédente sur  $|a|$  et  $b$  de sorte à écrire  $|a| = q.b + r$ . Si  $r = 0$ , on a alors  $a = -q.b - r = -q.b + r$  ; sinon, on a  $a = -q.b - r = (-q - 1).b + b - r$  avec  $b - r \in \llbracket 1, b - 1 \rrbracket$  ;

**Si  $b < 0$  :** on applique les cas précédents à  $a$  et  $|b|$  pour écrire  $a = q.|b| + r = (-q).b + r$ .

□

**Définition 1.1.3** (division euclidienne). Si  $a \in \mathbb{Z}$  et  $b \in \mathbb{Z}^*$ , on appelle respectivement *quotient de  $a$  par  $b$*  et *reste de  $a$  par  $b$*  les entiers  $q \in \mathbb{Z}$  et  $r \in \llbracket 0, |b| - 1 \rrbracket$  tels que  $a = q.b + r$ .

*Remarque 1.1.4.* La division telle qu'enseignée à l'école primaire n'est rien d'autre qu'une suite de divisions euclidiennes successives.

## 1.2 Primalité

### 1.2.1 Divisibilité

**Définition 1.2.1.** On dit qu'un entier  $b \in \mathbb{Z}^*$  *divise* un entier  $a \in \mathbb{Z}$ , que  $a$  est *divisible* par  $b$ , ou encore que  $a$  est un *multiple* de  $b$  si le reste de la division euclidienne de  $a$  par  $b$  est nul. On note cela  $b \mid a$ .

*Remarques 1.2.2.*

- Par unicité de la division euclidienne,  $b \in \mathbb{Z}^*$  divise  $a \in \mathbb{Z}$  si et seulement s'il existe  $k \in \mathbb{Z}$  tel que  $a = bk$ .
- Un entier divisible par  $n \in \mathbb{Z}^*$  est également divisible par  $-n$  ; et un multiple de  $n \in \mathbb{Z}$  est également un multiple de  $-n$ .
- Un entier  $n \in \mathbb{Z}$  est toujours divisible par 1,  $-1$ ,  $n$  et  $-n$ .
- L'entier 0 est divisible par tous les autres entiers mais ne divise personne.
- Les entiers 1 et  $-1$  sont les deux seuls à ne pas avoir quatre diviseurs distincts. Ils sont également les seuls à diviser tous les entiers.

**Lemme 1.2.3.** Si  $a \in \mathbb{Z}^*$  est divisible par  $b \in \mathbb{Z}^*$ , alors  $|a| \geq |b|$ .

*Démonstration.* Si  $a = q.b$ , alors  $|a| = |q|.|b|$ . De plus,  $|q| \neq 0$  car autrement on aurait  $q = 0$  et donc  $a = 0.b = 0$ , donc  $|q| \geq 1$  et on a  $|a| = |q|.|b| \geq |b|$ . □

**Corollaire 1.2.4.** La relation de divisibilité sur  $\mathbb{N}^*$  est une relation d'ordre. Sur  $\mathbb{Z}^*$ , seules la réflexivité et la transitivité restent vraies.

*Démonstration.* Pour tout  $a \in \mathbb{Z}^*$ , on a  $a = 1.a$  et donc  $a \mid a$ . Si, pour  $a, b, c \in \mathbb{Z}^*$ , on a  $c \mid b$  et  $b \mid a$ , c'est qu'il existe  $k_1, k_2 \in \mathbb{Z}$  tels que  $a = k_1.b$  et  $b = k_2.c$ , on a alors  $a = (k_1.k_2).c$  et donc  $c \mid a$ . On en déduit que sur  $\mathbb{Z}^*$ , la relation de divisibilité est réflexive et transitive.

Soit  $a, b \in \mathbb{N}^*$  tels que  $a \mid b$  et  $b \mid a$ . On a alors  $a \geq b$  et  $b \geq a$  et donc  $a = b$ . La relation de divisibilité est donc anti-symétrique sur  $\mathbb{N}^*$ , ce qui en fait une relation d'ordre. Dans  $\mathbb{Z}^*$ , par contre, les entiers 1 et  $-1$  se divisent l'un l'autre sans être égaux. □

Terminons cette section par un lemme simple mais essentiel.

**Lemme 1.2.5.** Soit  $a \in \mathbb{Z}^*$  et  $b_1, b_2 \in \mathbb{Z}$ . Si  $a$  divise  $b_1$  et  $b_2$ , alors il divise toute combinaison linéaire à coefficients entiers de  $b_1$  et  $b_2$ .

*Démonstration.* Si  $a$  divise  $b_1$  et  $b_2$ , alors il existe  $k_1, k_2 \in \mathbb{Z}$  tels que  $b_1 = k_1.a$  et  $b_2 = k_2.a$ . Dès lors, pour tout  $c_1, c_2 \in \mathbb{Z}$ , on a  $c_1.b_1 + c_2.b_2 = c_1.k_1.a + c_2.k_2.a = (c_1.k_1 + c_2.k_2).a$ , qui est donc divisible par  $a$ . □

### 1.2.2 Primalité d'un nombre entier

**Définition 1.2.6.** On dit que  $p \in \mathbb{N}^*$  est *premier* s'il possède exactement quatre diviseurs distincts. On note  $\mathcal{P}$  l'ensemble des nombres premiers.

*Remarque 1.2.7.* Une façon de reformuler la définition des nombres premiers, est de dire qu'un nombre  $p$  est premier si et seulement si toute écriture de la forme  $p = k_1.k_2$  avec  $k_1, k_2 \in \mathbb{Z}$  implique qu'exactement un des entiers  $k_1$  et  $k_2$  vaut  $\pm 1$ . Nous verrons plus tard que cela peut encore se reformuler en terme "d'éléments inversibles".

*Exemples 1.2.8.*

- L'entier 2 est premier car il est divisible par 1,  $-1$ , 2 et  $-2$ . Tout autre diviseur  $k$  devrait satisfaire  $|k| \leq 2$ , ce qui ne donne plus comme candidat que 0, lequel ne divise personne.
- L'entier 4 n'est pas premier car il est divisible par  $\pm 1$ ,  $\pm 2$  et  $\pm 4$ .
- L'entier 1 n'est pas premier car il ne possède que deux diviseurs, 1 et  $-1$ .

**Lemme 1.2.9.** Tout entier  $n \in \mathbb{Z} \setminus \{\pm 1\}$  est divisible par un nombre premier.

*Démonstration.* On remarque d'abord que 0 est divisible, par exemple, par 2 car  $0 = 0.2$ , et que si  $b \in \mathbb{Z}^*$  divise  $a \in \mathbb{Z}$ , alors il divise également  $-a$ . Il suffit donc de montrer que tout entier strictement plus grand que 1 est divisible par un nombre premier. On travaille par récurrence sur  $n$ .

Pour  $n = 2$ , le résultat est vrai car 2 se divise lui-même et 2 est premier.

Supposons maintenant le résultat vrai pour tout entier  $2 \leq k < n$ . Si  $n$  est premier, alors il suffit de dire qu'il se divise lui-même. Autrement, il possède un diviseur  $k$  positif distinct de 1 et de  $n$ . On a donc notamment  $2 \leq k < n$  et par principe de récurrence généralisée, il existe un diviseur premier  $p$  de  $k$ , lequel divise donc également  $n$  par transitivité.  $\square$

**Corollaire 1.2.10** (Euclide). Il existe une infinité de nombres premiers

*Démonstration.* Supposons par l'absurde que  $\mathcal{P} = \{p_1, \dots, p_\ell\}$  est un ensemble fini, alors  $N := \prod_{i=1}^\ell p_i + 1 > 1$  possède au moins un diviseur premier, disons  $p_{k_0}$ . On a donc  $N = a.p_{k_0}$  avec  $a \in \mathbb{N}^*$  et donc  $1 = \left(a - \prod_{\substack{i=0 \\ i \neq k_0}}^\ell p_i\right).p_{k_0}$ . Cela implique  $|p_{k_0}| \leq 1$ , ce qui est absurde car  $p_{k_0}$  est premier.  $\square$

### 1.2.3 Primalité entre nombres entiers

*Remarque 1.2.11.* Deux entiers  $a, b \in \mathbb{Z}^*$  quelconque ont toujours un diviseur en commun, à savoir 1, ainsi qu'un multiple en commun, à savoir leur produit  $a.b$ . De plus, l'ensemble des diviseurs positifs d'un entier  $a \in \mathbb{Z}^*$  donné est fini puisque les valeurs absolues de ses éléments sont toutes majorées par  $|a|$ .

**Définition 1.2.12.** Soit  $a, b \in \mathbb{Z}^*$ .

- On appelle *plus petit multiple commun* à  $a$  et  $b$ , l'entier  $\text{ppcm}(a, b) := \min\{c \in \mathbb{N}^* \mid c \text{ multiple de } a\} \cap \{c \in \mathbb{N}^* \mid c \text{ multiple de } b\}$ .
- On appelle *plus grand diviseur commun* à  $a$  et  $b$ , l'entier  $\text{pgcd}(a, b) := \max\{c \in \mathbb{N}^* \mid c \text{ diviseur de } a\} \cap \{c \in \mathbb{N}^* \mid c \text{ diviseur de } b\}$ .
- On dit que  $a$  et  $b$  sont *premiers entre eux* si  $\text{pgcd}(a, b) = 1$ .

*Remarque 1.2.13.* Les notions de plus petit multiple et plus grand diviseur communs se généralisent sans difficulté à un nombre quelconque d'entiers. Nous laisserons le soin au lecteur de généraliser les propositions qui suivent.

*Exemple 1.2.14.* Soit  $p \in \mathcal{P}$ . Les seuls diviseurs positifs de  $p$  sont 1 et  $p$ ; ce sont donc les seuls valeurs possibles pour  $\text{pgcd}(p, n)$  avec  $n \in \mathbb{Z}^*$ . Dans le premier cas  $n$  est premier avec  $p$ , et dans le second cas, c'est un multiple de  $p$ . En particulier, deux nombres premiers distincts sont toujours premiers entre eux.

**Lemme 1.2.15.** Soit  $a, b \in \mathbb{Z}^*$ . Tout multiple commun à  $a$  et  $b$  est un multiple de  $\text{ppcm}(a, b)$ .

*Démonstration.* Soit  $c \in \mathbb{Z}$  un multiple commun à  $a$  et  $b$ . Si  $c = 0$ , le résultat est clair. Sinon, par division euclidienne de  $c$  par  $\text{ppcm}(a, b)$ , on a  $q \in \mathbb{Z}$  et  $r \in \llbracket 1, \text{ppcm}(a, b) - 1 \rrbracket$  tels que  $c = q \cdot \text{ppcm}(a, b) + r$ . Mais alors  $r = c - q \cdot \text{ppcm}(a, b)$  est un multiple commun de  $a$  et  $b$ , et par minimalité de  $\text{ppcm}(a, b)$ , on a  $r = 0$ .  $\square$

**Lemme 1.2.16.** Soit  $a, b, r \in \mathbb{Z}^*$  et  $q \in \mathbb{Z}$ .

- Si  $a = b \cdot q$ , alors  $\text{pgcd}(a, b) = |b|$ .
- Si  $a = b \cdot q + r$  alors  $\text{pgcd}(a, b) = \text{pgcd}(b, r)$ .

*Démonstration.* Pour le premier point, il suffit de remarquer que  $|b|$  est un diviseur commun à  $a$  et  $b$  et que tout diviseur de  $b$  est majoré par  $|b|$ .

Montrons maintenant le second point. Tout diviseur commun de  $b$  et  $r$  est également un diviseur commun de  $b$  et  $a = q \cdot b + r$ . Mais réciproquement, tout diviseur commun de  $b$  et  $a$  est également un diviseur commun de  $b$  et  $r = a - q \cdot b$ . On en déduit que les ensembles des diviseurs communs de  $a$  et  $b$ , et de  $b$  et  $r$  sont les mêmes, et qu'il en va donc de même de leurs plus grands éléments  $\text{pgcd}(a, b)$  et  $\text{pgcd}(b, r)$ .  $\square$

*Application 1.2.17.* Par division euclidienne, on peut donc préserver la valeur du  $\text{pgcd}$ , tout en réduisant strictement la valeurs des entiers dont on considère le  $\text{pgcd}$ . En répétant cela jusqu'à obtenir un reste nul, cela permet de calculer le  $\text{pgcd}$  de deux entiers avec juste un nombre fini de divisions euclidiennes. C'est ce qu'on appelle l'*algorithme d'Euclide*.

**Théorème 1.2.18** (Bachet–Bézout). Pour tous  $a, b \in \mathbb{Z}^*$ , il existe  $r, s \in \mathbb{Z}$  tels que  $r \cdot a + s \cdot b = \text{pgcd}(a, b)$ .

*Démonstration.* On commence par considérer le cas  $a, b \in \mathbb{N}^*$ , que l'on montre par récurrence généralisée sur  $b \geq 1$ . Le résultat est vrai pour  $b = 1$  car alors  $\text{pgcd}(a, b) = 1 = 0 \cdot a + 1 \cdot b$ . Supposons maintenant le résultat vrai jusqu'à  $b - 1$ . Par division euclidienne, on a  $a = q \cdot b + r$  avec  $r \in \llbracket 0, b - 1 \rrbracket$ . Si  $r = 0$ , alors  $\text{pgcd}(a, b) = b = 0 \cdot a + 1 \cdot b$ ; sinon, par hypothèse de récurrence, on sait qu'il existe  $u, v \in \mathbb{Z}$  tels que  $\text{pgcd}(b, r) = u \cdot b + v \cdot r$ . Mais alors  $\text{pgcd}(a, b) = \text{pgcd}(b, r) = u \cdot b + v \cdot (a - q \cdot b) = v \cdot a + (u - q \cdot v) \cdot b$ .

Pour  $a, b \in \mathbb{Z}^*$ , il suffit d'écrire  $\text{pgcd}(|a|, |b|) = r' \cdot |a| + s' \cdot |b|$  avec  $r', s' \in \mathbb{Z}$  et d'en déduire  $\text{pgcd}(a, b) = \text{pgcd}(|a|, |b|) = r \cdot a + s \cdot b$  avec  $r = \text{signe}(a) \cdot r'$  et  $s = \text{signe}(b) \cdot s'$ .  $\square$

*Remarque 1.2.19.* La récurrence de la preuve du théorème de Bachet–Bézout donne un algorithme récursif effectif pour trouver des entiers  $r$  et  $s$ . Cela correspond à “remonter” les calculs dans l'algorithme d'Euclide. C'est ce qu'on appelle l'*algorithme d'Euclide étendu*.

**Corollaire 1.2.20.** Pour tous  $a, b \in \mathbb{Z}^*$ , on a  $a$  et  $b$  premiers entre eux si et seulement si il existe  $r, s \in \mathbb{Z}$  tels que  $r \cdot a + s \cdot b = 1$ .

*Démonstration.* Si  $a$  et  $b$  sont premiers entre eux, alors il existe  $r, s \in \mathbb{Z}$  tels que  $1 = \text{pgcd}(a, b) = r \cdot a + s \cdot b$  d'après la proposition précédente. Et réciproquement, s'il existe de tels  $r, s \in \mathbb{Z}$ , alors tout diviseur commun à  $a$  et  $b$  divise 1. On a donc  $\text{pgcd}(a, b) = 1$ .  $\square$

**Corollaire 1.2.21.** Soit  $a, b, c \in \mathbb{Z}^*$  tels que  $a$  soit simultanément premier avec  $b$  et  $c$ , alors il est premier avec  $b \cdot c$ .

*Démonstration.* Par hypothèse, il existe  $r_b, s_b, r_c, s_c \in \mathbb{Z}$  tels que  $r_b \cdot a + s_b \cdot b = 1 = r_c \cdot a + s_c \cdot c$ , mais alors  $1 = 1 \cdot 1 = (r_b \cdot a + s_b \cdot b) \cdot (r_c \cdot a + s_c \cdot c) = (r_b \cdot r_c \cdot a + s_b \cdot r_c \cdot b + r_b \cdot s_c \cdot c) \cdot a + (s_b \cdot s_c) \cdot b \cdot c$  avec  $r_b \cdot r_c \cdot a + s_b \cdot r_c \cdot b + r_b \cdot s_c \cdot c$  et  $s_b \cdot s_c$  dans  $\mathbb{Z}$ .  $\square$

*Application 1.2.22.* Soit  $a, b, c \in \mathbb{Z}^*$ . On cherche les solutions entières de l'équation, alors dite *diophantienne*,  $a.x + b.y = c$ .

Si  $\text{pgcd}(a, b) \nmid c$ , alors l'équation n'a clairement aucune solution. Sinon, l'algorithme d'Euclide donne  $r, s \in \mathbb{Z}$  tels que  $a.r + b.s = \text{pgcd}(a, b)$ . Dès lors,  $(x_0, y_0) = \left( \frac{r.c}{\text{pgcd}(a, b)}, \frac{s.c}{\text{pgcd}(a, b)} \right) \in \mathbb{Z}^2$  donne une solution à l'équation. Toute autre solution  $(x, y) \in \mathbb{Z}^2$  vérifie  $a.(x - x_0) + b.(y - y_0) = 0$ . Mais  $a.(x - x_0) = b.(y_0 - y)$  est alors un multiple commun à  $a$  et  $b$ , et il existe donc  $k \in \mathbb{Z}$  tel que  $a.(x - x_0) = b.(y_0 - y) = k.\text{ppcm}(a, b)$ . On en déduit que  $x = x_0 + k.\frac{\text{ppcm}(a, b)}{a}$  et  $y = y_0 - k.\frac{\text{ppcm}(a, b)}{b}$ . Réciproquement, pour tout  $k \in \mathbb{Z}$ ,  $\left( x_0 + k.\frac{\text{ppcm}(a, b)}{a}, y_0 - k.\frac{\text{ppcm}(a, b)}{b} \right) \in \mathbb{Z}^2$  est clairement solution de l'équation. Au final, on en déduit que les solutions entières sont exactement les couples de la forme

$$\left( \frac{r.c}{\text{pgcd}(a, b)} + k.\frac{\text{ppcm}(a, b)}{a}, \frac{s.c}{\text{pgcd}(a, b)} - k.\frac{\text{ppcm}(a, b)}{b} \right)$$

avec  $k \in \mathbb{Z}$ .

#### 1.2.4 Décomposition en facteurs premiers

**Lemme 1.2.23** (lemme de Gauss). Soit  $a, b \in \mathbb{Z}^*$  et  $c \in \mathbb{Z}$ . Si  $a$  divise  $b.c$ , et  $a$  est premier avec  $b$ , alors  $a$  divise  $c$ .

*Démonstration.* Par le théorème de Bachet–Bézout, on sait qu'il existe  $r, s \in \mathbb{Z}$  tels que  $1 = r.a + s.b$ ; et par hypothèse qu'il existe  $k \in \mathbb{Z}$  tel que  $b.c = k.a$ . Mais alors  $c = c.r.a + c.s.b = c.r.a + s.k.a = (c.r + s.k).a$  et donc  $a$  divise  $c$ .  $\square$

**Corollaire 1.2.24** (lemme d'Euclide). Soit  $a, b \in \mathbb{Z}$  et  $p \in \mathcal{P}$ . Si  $p$  divise  $a.b$ , alors  $p$  divise  $a$  ou  $b$ .

Plus généralement, si  $p$  divise un produit de  $k \in \mathbb{N}^*$  entiers, alors  $p$  divise au moins l'un des  $k$  facteurs.

*Démonstration.* Si  $p$  divise  $a$ , alors le résultat est vrai. Sinon, alors  $p$  est premier avec  $a$  et d'après le lemme de Gauss,  $p$  divise  $b$ .

La généralisation s'obtient par récurrence immédiate.  $\square$

**Théorème 1.2.25.** Pour tout entier  $a \in \mathbb{N}^* \setminus \{1\}$ , il existe une unique écriture

$$a = p_1^{\alpha_1} \cdot p_2^{\alpha_2} \cdots p_k^{\alpha_k}$$

avec  $p_1 < p_2 < \cdots < p_k \in \mathcal{P}$  et  $\alpha_1, \alpha_2, \dots, \alpha_k \in \mathbb{N}^*$ .

*Remarque 1.2.26.* Le théorème peut s'étendre à  $\mathbb{N}^*$  en admettant que, par convention un produit de  $k = 0$  termes vaut 1.

*Démonstration.* L'existence découle du lemme 1.2.9 par récurrence généralisée sur  $a$ . Le résultat est en effet vrai pour  $a = 1$ . Supposons le résultat vrai jusqu'à  $a$ , alors  $a + 1$  possède un diviseur premier  $p$  et par hypothèse de récurrence  $\frac{a}{p} \in \mathbb{N}^*$  possède une telle décomposition à laquelle il suffit de rajouter  $p$ .

Supposons maintenant par l'absurde qu'il n'y a pas unicité d'une telle décomposition. On considère alors  $a_0$  le plus petit entier positif possédant ainsi plusieurs décompositions. On a clairement  $a_0 > 1$  et donc

$$p_1^{\alpha_1} \cdots p_k^{\alpha_k} = a_0 = q_1^{\beta_1} \cdots q_l^{\beta_l}$$

où les produits de chaque côtés sont deux décompositions distinctes telles que dans l'énoncé, avec  $k, l \geq 1$ . Mais alors  $p_1$  divise  $q_1^{\beta_1} \cdots q_l^{\beta_l}$  et par le lemme d'Euclide,  $p_1$  divise l'un des  $q_i$ . Mais par unicité du diviseur strictement plus grand que 1 de  $q_i$ ,  $q_i$  étant premier, on a même  $p_1 = q_i$ . On peut alors diviser de chaque côté par  $p_1 = q_i$  et contredire la minimalité de  $a_0$ .  $\square$

**Définition 1.2.27.** Soit  $p \in \mathcal{P}$ . On appelle *valuation  $p$ -adique* de  $a \in \mathbb{N}^*$  l'exposant de  $p$  dans l'unique décomposition de  $a$  en facteurs premiers. Cela définit une application  $\nu_p : \mathbb{N} \rightarrow \mathbb{N} \cup \{\infty\}$  que l'on étend par convention en zéro par  $\nu_p(0) = \infty$ .

**Proposition 1.2.28.** • Pour tout  $a \in \mathbb{N}^*$ ,  $a = \prod_{p \in \mathcal{P}} p^{\nu_p(a)}$ .

• Pour tout  $p \in \mathcal{P}$  et pour tout  $a \in \mathbb{N}$ ,  $\nu_p(a) = \max\{k \in \mathbb{N} \mid p^k \text{ divise } a\}$ .

*Remarque 1.2.29.* Puisqu'il porte sur l'ensemble  $\mathcal{P}$  qui est infini, le produit du premier point est infini et nécessite donc d'être précisé. D'après le théorème 1.2.25, il n'y a, pour tout  $n \in \mathbb{N}^*$ , qu'un nombre fini de nombres premiers  $p \in \mathcal{P}$  tels que  $\nu_p(n) \neq 0$ ; tous les facteurs dans le produit valent donc 1 sauf un nombre fini. On peut donc définir le produit infini comme limite, lorsque  $k$  tend vers  $+\infty$ , de la suite des produits partiels portant sur les éléments de  $\mathcal{P}$  inférieurs à  $k$ ; cette suite est bien convergente car constante à partir d'un certain rang.

*Démonstration.* Le premier point découle directement de la définition en lien avec le théorème 1.2.25. Ce dernier théorème montre également, concernant le second point lorsque  $a > 0$ , que  $p^{\nu_p(a)}$  divise  $a$ ; mais, réciproquement, aucun  $p^{\nu_p(a)+k}$  avec  $k \in \mathbb{N}^*$  ne peut diviser  $a$  car autrement, en simplifiant par  $p^{\nu_p(a)}$ , on aurait  $p$  qui diviserait  $\prod_{q \in \mathcal{P} \setminus \{p\}} q^{\nu_q(a)}$  et donc un certain  $q \in \mathcal{P} \setminus \{p\}$  d'après le lemme d'Euclide, ce qui est

absurde car deux nombres premiers distincts sont premiers entre eux. Pour  $a = 0$ , toutes les puissances de  $p$  divisent clairement 0.  $\square$

**Proposition 1.2.30.** Pour tout  $p \in \mathcal{P}$ , et tout  $a, b \in \mathbb{N}$ , on a :

- $\nu_p(a) = \infty$  ssi  $a = 0$ ;
- $\nu_p(a.b) = \nu_p(a) + \nu_p(b)$ ;
- $\nu_p(a + b) \geq \min(\nu_p(a), \nu_p(b))$ .

*Remarque 1.2.31.* Ce sont les propriétés que l'on attend d'une "valuation". Cela permet de définir une notion de distance sur  $\mathbb{Z}$ , définie par  $d(a, b) = p^{-\nu_p(|a-b|)}$ , très différente de la distance usuelle. C'est le premier pas vers l'étude des nombres dits  *$p$ -adiques*, étude que nous ne pousserons néanmoins pas plus loin dans ce cours.

*Démonstration.* Les deux premiers points viennent directement de la définition. Pour le troisième point, posons  $k = \min(\nu_p(a), \nu_p(b))$ . On a alors  $p^k$  qui divise  $a$  et  $b$ , et donc  $a + b$ ; on en déduit que  $\nu_p(a + b) \geq k$ .  $\square$

**Proposition 1.2.32.** Pour tout  $a, b \in \mathbb{N}^*$ , on a

$$\text{ppcm}(a, b) = \prod_{p \in \mathcal{P}} p^{\max(\nu_p(a), \nu_p(b))} \quad \text{et} \quad \text{pgcd}(a, b) = \prod_{p \in \mathcal{P}} p^{\min(\nu_p(a), \nu_p(b))}.$$

*Démonstration.* Notons  $m := \text{ppcm}(a, b)$  et  $d = \text{pgcd}(a, b)$ . Pour tout  $p \in \mathcal{P}$ , il est clair d'après le second point de la proposition 1.2.30 que  $\nu_p(m) \geq \nu_p(a), \nu_p(b)$  et que  $\nu_p(d) \leq \nu_p(a), \nu_p(b)$ . On en déduit que  $m \geq \prod_{p \in \mathcal{P}} p^{\max(\nu_p(a), \nu_p(b))}$  et que  $d \leq \prod_{p \in \mathcal{P}} p^{\min(\nu_p(a), \nu_p(b))}$ . Mais réciproquement, ces derniers sont bien, respectivement, un multiple et un diviseur communs à  $a$  et  $b$ .  $\square$

**Corollaire 1.2.33.** Pour tout  $a, b \in \mathbb{N}^*$ , on a  $\text{ppcm}(a, b) \cdot \text{pgcd}(a, b) = a.b$ .

*Démonstration.* Celui provient de l'égalité, pour tout  $k_1, k_2 \in \mathbb{N}$ ,  $\min(k_1, k_2) + \max(k_1, k_2) = k_1 + k_2$ .  $\square$

## 1.3 Congruences

### 1.3.1 Retour les relations d'équivalence

Avant de définir la notion de congruence sur les nombres entiers, faisons quelques rappels sur la notion de relation d'équivalence. Ces dernières sont parfois définies avec un peu de flottement comme une “lien reliant certains éléments entre eux”, vérifiant certains axiomes. Afin de clarifier totalement cette notion, nous allons ici donner un sens plus formel à ce lien en définissant l'ensemble des couples d'éléments étant “en lien”.

**Définition 1.3.1.** Soit  $X$  un ensemble. Une relation d'équivalence sur  $X$ , c'est un sous-ensemble  $R \subset X \times X$  tel que :

- (réflexivité)  $\forall x \in X, (x, x) \in R$ ;
- (symétrie)  $\forall x, y \in X, (x, y) \in R \Rightarrow (y, x) \in R$ ;
- (transitivité)  $\forall x, y, z \in X, (x, y), (y, z) \in R \Rightarrow (x, z) \in R$ .

*Remarque 1.3.2.* Dans la définition précédente, il faut comprendre que deux éléments  $x, y \in X$  sont en relation si et seulement si le couple  $(x, y)$  est dans  $R$ . On note alors  $x \sim_R y$ , voire  $x \sim y$  s'il n'y a pas d'ambiguïté sur la nature de  $R$ .

*Exemple 1.3.3.* Sur  $\mathbb{Z}$ , on peut considérer la relation  $R := \{(a, b) \in \mathbb{Z}^2 \mid ab > 0\} \cup \{(0, 0)\}$ . Cet ensemble  $R$  est :

- réflexif :  $(0, 0) \in R$  et pour tout  $a \in \mathbb{Z}^*$ , on a bien  $a^2 > 0$ ;
- symétrique : si  $(a, b) \in R$ , alors soit  $a = b = 0$  et alors  $(b, a) = (0, 0) \in R$ ; soit  $ab > 0$  et alors  $ba = ab > 0$  donc  $(b, a) \in R$ ;
- transitif : si  $(a, b), (b, c) \in R$ , alors ou bien  $b = 0$ , mais dans ce cas  $a$  et  $c$  valent aussi zéro car ni  $ab$  ni  $bc$  ne peut alors être strictement positif, et donc  $(a, c) = (0, 0) \in R$ ; ou bien  $b \neq 0$ , mais alors  $ab, bc > 0$  et  $b^2 > 0$ , on en déduit que  $ac = ab^2c = ab \cdot bc > 0$  et donc  $(a, c) \in R$ .

Il s'agit donc bien d'une relation d'équivalence, que l'on peut traduire par  $a \sim b$  si et seulement si  $a$  et  $b$  ont le même signe, le signe en question pouvant être positif, négatif ou nul.

**Définition 1.3.4.** Soit  $X$  un ensemble muni d'une relation d'équivalence  $R$ . Pour tout  $x \in X$ , on définit la *classe de  $x$  pour  $R$*  comme le sous-ensemble  $\{y \in X \mid (x, y) \in R\}$  des éléments de  $X$  en relation avec  $x$ . On le note souvent  $\bar{x}^R$ , voire  $\bar{x}$  s'il n'y a pas d'ambiguïté sur la nature de  $R$ . On appelle *classe d'équivalence* pour  $R$  toute classe d'un élément, et on dit de tout élément d'une classe  $\mathfrak{c}$  qu'il est un *représentant* de cette classe.

*Remarque 1.3.5.* Nous n'insisterons jamais assez là-dessus : une classe, c'est donc *un sous-ensemble (non vide) de  $X$ !*

*Exemple 1.3.6.* Dans l'exemple 1.3.3, il y a trois classes d'équivalences :  $\mathfrak{Z} := \bar{0} = \{0\}$ ,  $\mathfrak{P} := \bar{1} = \mathbb{N}^*$  et  $\mathfrak{N} := \bar{-1} = \mathbb{Z} \setminus \mathbb{N}$ . Nous avons pris 0, 1 et  $-1$  comme représentants de chacune des classes, mais ce choix est totalement arbitraire, car tout aussi bien, on a  $\mathfrak{P} = \bar{2}$  ou  $\bar{157}$ , et  $\mathfrak{N} = \bar{-5}$  ou  $\bar{-1032}$ .

De la définition découlent directement un certain nombre de propriétés “évidentes”.

**Lemme 1.3.7.** Soit  $X$  un ensemble muni d'une relation d'équivalence  $R$ .

- Pour toute classe d'équivalence  $\mathfrak{c} \subset X$ , on a :
  - $\mathfrak{c}' = \mathfrak{c}$  pour tout classe  $\mathfrak{c}'$  telle que  $\mathfrak{c} \cap \mathfrak{c}' \neq \emptyset$ ;
  - $\bar{x} = \mathfrak{c}$  pour tout  $x \in \mathfrak{c}$ .
- Pour tout  $x, y \in X$ ,  $x \sim y$  si et seulement si  $\bar{x} = \bar{y}$ .

*Démonstration.* Soit  $\mathfrak{c}$  et  $\mathfrak{c}'$  deux classes. Par définition, il existe  $x_0, x'_0 \in X$  tels que  $\mathfrak{c} = \bar{x}_0$  et  $\mathfrak{c}' = \bar{x}'_0$ . Supposons maintenant qu'elles soient d'intersection non vide, on a alors  $z_0 \in \mathfrak{c} \cap \mathfrak{c}'$ , et par définition des éléments d'une classe, on a  $x_0 \sim z_0$  et  $x'_0 \sim z_0$ . Mais alors,  $x_0 \sim x'_0$  par symétrie et transitivité et, pour



tout  $y \in \mathfrak{c}$ , on a  $y \sim x_0 \sim x'_0$ , donc par transitivité,  $y \in \mathfrak{c}'$ . On en déduit que  $\mathfrak{c} \subset \mathfrak{c}'$  et même, l'inclusion réciproque s'obtenant similairement,  $\mathfrak{c} = \mathfrak{c}'$ .

Soit  $x \in \mathfrak{c}$ , on a alors  $x_0 \sim x$  et donc  $x \in \mathfrak{c} \cap \bar{x}$ . On en déduit que  $\mathfrak{c} \cap \bar{x} \neq \emptyset$ , et donc  $\mathfrak{c} = \bar{x}$  d'après le point précédent.

Soit  $x, y \in X$ . Si  $x \sim y$ , alors  $x \in \bar{x} \cap \bar{y}$ . On a donc  $\bar{x} \cap \bar{y} \neq \emptyset$ , ce qui implique d'après ce qui précède  $\bar{x} = \bar{y}$ . Réciproquement, si  $\bar{x} = \bar{y}$ , alors  $x \in \bar{x} = \bar{y}$  et donc  $x \sim y$ .  $\square$

La notion de relation d'équivalence sur  $X$  est très fortement reliée à la notion de partition de  $X$ .

**Définition 1.3.8.** Soit  $X$  un ensemble. On dit qu'une famille  $(X_i)_{i \in I}$  de sous-ensembles de  $X$  forment une *partition* de  $X$  si

- $\cup_{i \in I} X_i = X$  ;
- $X_i \cap X_j = \emptyset$  pour tout  $i \neq j \in I$ .

Autrement dit, une partition de  $X$ , c'est un découpage de  $X$  en paquets disjoints.

**Proposition 1.3.9.** Soit  $X$  un ensemble muni d'une relation d'équivalence, alors les classes d'équivalence forment une partition de  $X$ .

*Démonstration.* Puisque, pour tout  $x \in X$ ,  $x \in \bar{x}$ , la réunion des classes d'équivalences donne bien  $X$  tout entier. D'autre part, d'après le lemme 1.3.7, deux classes d'équivalence sont soit identiques, soit d'intersection vide.  $\square$

*Remarque 1.3.10.* Réciproquement, on vérifie facilement que toute partition  $(X_i)_{i \in I}$  de  $X$  induit une relation d'équivalence sur  $X$  définie, pour tout  $x_1, x_2 \in X$ , par  $x_1 \sim x_2$  si et seulement s'il existe  $i \in I$  tel que  $x_1, x_2 \in X_i$ . Dans ce cas, les classes d'équivalences sont exactement les  $X_i$ . Il y a en fait deux correspondances, réciproques l'une pour l'autre, entre les notions de relations d'équivalence et de partitions :

| Relations d'équivalence sur $X$                            | Partitions de $X$                        |
|------------------------------------------------------------|------------------------------------------|
| $R$                                                        | $\leadsto$ classes d'équivalences de $R$ |
| $(x_1, x_2) \in R$ ssi $\exists i \in I, x_1, x_2 \in X_i$ | $\nwarrow (X_i)_{i \in I}$               |

Cela permet de réinterpréter toute relation d'équivalence sur  $X$  comme un découpage de  $X$  en paquets, chaque élément d'un même paquet étant dit "équivalents".

*Exemple 1.3.11.* En reprenant encore l'exemple 1.3.3, la relation d'équivalence y sépare les entiers en trois paquets, correspondant à leur signe. On a de fait  $\mathbb{Z} = (\mathbb{Z} \setminus \mathbb{N}) \sqcup \{0\} \sqcup \mathbb{N}^*$ .

**Définition 1.3.12.** Soit  $X$  un ensemble muni d'une relation d'équivalence  $R$ . On appelle *espace quotient* associé, noté  $X/\sim_R$ , l'ensemble des classes d'équivalence pour  $R$ .

*Exemple 1.3.13.* En suivant toujours l'exemple 1.3.3, on a  $Z/\sim = \{\mathbb{Z} \setminus \mathbb{N}, \{0\}, \mathbb{N}^*\}$ , c'est-à-dire  $Z/\sim = \{\mathfrak{N}, \mathfrak{Z}, \mathfrak{P}\}$ . Dans l'espace quotient, chaque paquet est pris comme un élément dans sa globalité; en ce sens, c'est une simplification de  $X$  puisqu'on en réduit le nombre d'éléments en ne les considérant plus que par paquets.

Une question récurrente est de savoir, lorsqu'un ensemble est muni d'une certaine structure (par exemple l'existence d'une opération définie dessus), si cela permet d'induire une opération similaire sur l'espace quotient. La réponse est parfois oui, parfois non, mais dans tous les cas, la stratégie est la même : si on a, par exemple, une notion de somme  $+$  sur  $X$  et qu'on veut définir  $\mathfrak{c}_1 + \mathfrak{c}_2$  pour  $\mathfrak{c}_1, \mathfrak{c}_2 \in X/\sim$ , on considérera  $x_1 \in \mathfrak{c}_1$  et  $x_2 \in \mathfrak{c}_2$  des représentants pour chacune des deux classes et on regardera  $\overline{x_1 + x_2}$ , la classe de leur somme. Toute la question est alors de savoir si le résultat dépend ou non de ces choix (a priori arbitraire) de représentants. Si le résultat est toujours le même, alors on dit que la somme est *compatible* avec la relation d'équivalence et cela permet de donner un sens à  $\mathfrak{c}_1 + \mathfrak{c}_2$ ; si non, on dit que la somme n'est pas compatible et il n'y a pas de somme induite sur l'espace quotient.

*Exemple 1.3.14.* Reprenons une dernière fois l'exemple 1.3.3. Sur  $\mathbb{Z}$ , nous avons une notion d'addition et une autre de multiplication.

- Peut-on en déduire une somme sur  $\mathbb{Z}/\sim$  ?

Essayons de définir  $\mathfrak{P} + \mathfrak{N}$ . Pour cela on considère successivement :

- ▶  $1 \in \mathfrak{P}$  et  $-1 \in \mathfrak{N}$ , cela donnerait  $\mathfrak{P} + \mathfrak{N} = \overline{1-1} = \overline{0} = \mathfrak{Z}$  ;
- ▶  $10 \in \mathfrak{P}$  et  $-5 \in \mathfrak{N}$ , cela donnerait  $\mathfrak{P} + \mathfrak{N} = \overline{10-5} = \overline{5} = \mathfrak{P}$  ;
- ▶  $3 \in \mathfrak{P}$  et  $-7 \in \mathfrak{N}$ , cela donnerait  $\mathfrak{P} + \mathfrak{N} = \overline{3-7} = \overline{-4} = \mathfrak{N}$ .

On voit que, selon le choix que nous faisons, nous obtenons des résultats différents<sup>1</sup>. Il n'est pas possible de définir de façon cohérente une somme sur  $\mathbb{Z}/\sim$ .

- Peut-on en déduire un produit sur  $\mathbb{Z}/\sim$  ?

Essayons de définir  $\mathfrak{P}.\mathfrak{N}$ . Pour cela on considère successivement :

- ▶  $1 \in \mathfrak{P}$  et  $-1 \in \mathfrak{N}$ , cela donnerait  $\mathfrak{P}.\mathfrak{N} = \overline{1.(-1)} = \overline{-1} = \mathfrak{N}$  ;
- ▶  $10 \in \mathfrak{P}$  et  $-5 \in \mathfrak{N}$ , cela donnerait  $\mathfrak{P}.\mathfrak{N} = \overline{10.(-5)} = \overline{-50} = \mathfrak{N}$  ;
- ▶  $3 \in \mathfrak{P}$  et  $-7 \in \mathfrak{N}$ , cela donnerait  $\mathfrak{P}.\mathfrak{N} = \overline{3.(-7)} = \overline{-21} = \mathfrak{N}$ .

On trouve ici toujours le même résultat ! Et même plus généralement, le produit d'un entier strictement positif, quel qu'il soit, par un entier strictement négatif, quel qu'il soit, étant toujours strictement négatif, on peut sans ambiguïté poser que  $\mathfrak{P}.\mathfrak{N} = \mathfrak{N}$ . De même, dans les autres on peut montrer que la multiplication des entiers induit la multiplication suivante sur  $\mathbb{Z}/\sim$  :

| .              | $\mathfrak{N}$ | $\mathfrak{Z}$ | $\mathfrak{P}$ |
|----------------|----------------|----------------|----------------|
| $\mathfrak{N}$ | $\mathfrak{P}$ | $\mathfrak{Z}$ | $\mathfrak{N}$ |
| $\mathfrak{Z}$ | $\mathfrak{Z}$ | $\mathfrak{Z}$ | $\mathfrak{Z}$ |
| $\mathfrak{P}$ | $\mathfrak{N}$ | $\mathfrak{Z}$ | $\mathfrak{P}$ |

### 1.3.2 Congruence modulo $n$

La division euclidienne entre nombres entiers génère deux entiers : un quotient et un reste. Le lecteur attentif aura pu remarquer que, dans les sections précédentes, seul le reste semble importer vraiment, le quotient n'étant que rarement pris en considération. La notion de congruence est un moyen d'achever définitivement le quotient. Dans la suite de cette section, on fixe  $n \in \mathbb{N}^*$ .

**Définition 1.3.15.** On dit que deux entiers  $a, b \in \mathbb{Z}$  sont *congrus modulo  $n$*  si leur différence  $a - b$  est divisible par  $n$  ; on note alors  $a \equiv b[n]$ .

#### Proposition 1.3.16.

- La relation de congruence modulo  $n$  est une relation d'équivalence.
- Un entier est congru modulo  $n$  au reste de sa division euclidienne par  $n$ . Deux entiers sont donc congrus modulo  $n$  si et seulement si leurs divisions euclidiennes par  $n$  ont le même reste ; en particulier, un entier est divisible par  $n$  si et seulement s'il est congru à 0 modulo  $n$ .

*Démonstration.* La relation est clairement réflexive et symétrique. Elle est plus transitive car si, pour  $a, b, c \in \mathbb{Z}$ ,  $n$  divise  $a - b$  et  $b - c$ , alors  $n$  divise également  $a - c = (a - b) + (b - c)$ .

Concernant le second point, on fixe  $a, b \in \mathbb{Z}$  et on pose  $a =: q_a.n + r_a$  et  $b =: q_b.n + r_b$  les divisions euclidiennes de  $a$  et  $b$  par  $n$ . On a  $a - b$  divisible par  $n$  si et seulement si  $r_a - r_b$  est divisible par  $n$ . Or  $r_a - r_b \in \llbracket 1 - n, n - 1 \rrbracket$  et le seul multiple de  $n$  dans cet intervalle est 0. On en déduit que  $a \equiv b[n]$  si et seulement si  $r_a - r_b = 0$ .  $\square$

Chaque relation de congruence étant une relation d'équivalence, on peut donc, pour chacune, considérer l'ensemble quotient associé, dont les éléments sont les classes d'équivalence. La proposition précise même qu'il y a exactement  $n$  classes, avec chacune un unique représentant dans  $\llbracket 0, n - 1 \rrbracket$ .

1. on pourra rapprocher cela de la forme indéterminée  $\infty - \infty$  parfois rencontrée dans les calculs de limites, il est impossible de lui donner un sens absolu car le résultat dépend de ce qui se cache vraiment derrière chaque signe  $\infty$

**Notation 1.3.17.** Pour tout  $a \in \mathbb{Z}$ , on note  $\bar{a}^n$ , et souvent juste  $\bar{a}$  lorsqu'il n'y a pas d'ambiguïté sur  $n$ , la classe de  $a$  dans l'espace quotient associé à la relation de congruence modulo  $n$ .

Les opérations d'addition et de multiplication sont compatibles avec les relations de congruence :

**Proposition 1.3.18.** Soit  $a, b, a', b' \in \mathbb{Z}$  tels que  $a' \equiv a[n]$  et  $b' \equiv b[n]$ . Alors  $a' + b' \equiv a + b[n]$  et  $a'.b' \equiv a.b[n]$ .

*Démonstration.* Par définition, il existe  $q_a, q_b \in \mathbb{Z}$  tels que  $a' - a = q_a.n$  et  $b' - b = q_b.n$ . On a alors  $(a' + b') - (a + b) = a' - a + b' - b = (q_a + q_b).n$  et  $a'.b' - a.b = a'.b' - a'.b + a'.b - a.b = a'.(b' - b) + b.(a' - a) = (a'.q_b + b.q_a).n$ . On a donc bien  $a' + b' \equiv a + b[n]$  et  $a'.b' \equiv a.b[n]$ .  $\square$

Cela permet de définir des opérations d'addition et de multiplication sur les classes de congruence modulo  $n$  en choisissant arbitrairement des représentants pour chaque classe, en associant respectivement, la classe de la somme et du produit de ces représentants. La proposition ci-dessus montre en effet que le résultat ne dépend pas des représentants choisis.

**Définition 1.3.19.** On note<sup>2</sup>  $\mathbb{Z}/n\mathbb{Z} := \{\bar{0}, \bar{1}, \dots, \overline{n-1}\}$  l'ensemble quotient pour la relation de congruence modulo  $n$ . Il est muni d'une addition  $+$  :  $\mathbb{Z}/n\mathbb{Z} \times \mathbb{Z}/n\mathbb{Z} \rightarrow \mathbb{Z}/n\mathbb{Z}$  et d'une multiplication  $.$  :  $\mathbb{Z}/n\mathbb{Z} \times \mathbb{Z}/n\mathbb{Z} \rightarrow \mathbb{Z}/n\mathbb{Z}$  définies, pour tout  $\bar{a}, \bar{b} \in \mathbb{Z}/n\mathbb{Z}$  par  $\bar{a} + \bar{b} = \overline{a+b}$  et  $\bar{a}.\bar{b} = \overline{a.b}$ .

**Proposition 1.3.20.** Sur  $\mathbb{Z}/n\mathbb{Z}$ , l'addition et la multiplication sont associatives et commutatives ; la multiplication est distributive sur l'addition ;  $\bar{0}$  est un élément neutre pour l'addition et un élément absorbant pour la multiplication ; et  $\bar{1}$  est un élément neutre pour la multiplication.

*Démonstration.* Toutes ces propriétés découlent directement des propriétés correspondantes de l'addition et de la multiplication sur  $\mathbb{Z}$ .  $\square$

On peut donc calculer dans  $\mathbb{Z}/n\mathbb{Z}$  comme dans  $\mathbb{Z}$ , et même plus facilement que dans  $\mathbb{Z}$  puisque l'on peut, entre chaque étape de calcul, se ramener à un élément dans  $\llbracket 0, n-1 \rrbracket$ . Couplé à la dernière affirmation de la proposition 1.3.16, cela permet de simplifier grandement les problématiques de divisibilité par  $n$ .

*Avertissement 1.3.21.* Nous avons parlé d'addition et de multiplication dans  $\mathbb{Z}/n\mathbb{Z}$ , également d'opposé, donc de soustraction, mais pas de division. Cette dernière n'est en, effet pas définie. En particulier, dans  $\mathbb{Z}/n\mathbb{Z}$ , l'égalité  $a.b = a.c$  n'implique pas  $b = c$  ! Dans  $\mathbb{Z}/4\mathbb{Z}$ , on a en effet  $\bar{2}.\bar{2} = \bar{2}.\bar{0}$ , mais  $\bar{2} \neq \bar{0}$ . Parfois, la division par un élément  $a \in \mathbb{Z}/n\mathbb{Z}$  pourra être simulée avec la multiplication par un inverse, c'est-à-dire un élément  $a^{-1} \in \mathbb{Z}/n\mathbb{Z}$  tel que  $a.a^{-1} = \bar{1}$ , mais seulement lorsqu'un tel inverse existe, ce qui n'est pas toujours le cas.

Dans cette même perspective, le théorème suivant pourra également se révéler d'une grande efficacité.

**Théorème 1.3.22.** Soit  $n_1, \dots, n_k \in \mathbb{N}^*$ .

- Si  $a, b \in \mathbb{Z}$  vérifient  $a \equiv b[n_i]$  pour tout  $i \in \llbracket 1, k \rrbracket$ , alors  $a \equiv b[\text{ppcm}(n_1, \dots, n_k)]$ .
- (théorème des restes chinois) Si  $n_1, \dots, n_k$  sont deux à deux premiers entre eux, alors pour tout  $a_1, \dots, a_k \in \mathbb{Z}$ , il existe un unique  $a \in \llbracket 0, N-1 \rrbracket$ , avec  $N := \prod_{i=1}^k n_i$ , tel que  $a \equiv a_i[n_i]$  pour tout  $i \in \llbracket 1, k \rrbracket$ .

*Démonstration.* Pour le premier point, il suffit de constater que  $a - b$  est un multiple commun à tous les  $n_i$ . D'après le lemme 1.2.15, c'est donc un multiple de  $\text{ppcm}(n_1, \dots, n_k)$ , ce qui donne le résultat.

Concernant le second point, on commence par montrer l'existence. Pour cela, on remarque que, pour tout  $i \in \llbracket 1, k \rrbracket$ , les entiers  $n_i$  et  $n'_i = \prod_{j \in \llbracket 1, k \rrbracket \setminus \{i\}} n_j$  sont premiers entre eux et qu'il existe donc, par le théorème

2. cette notation trouvera tout son sens dans les chapitres suivants

de Bachet–Bézout, des entiers  $r_i, s_i \in \mathbb{Z}$  tels que  $r_i.n_i + s_i.n'_i = 1$ . En posant  $e_i := s_i.n'_i$ , on a alors  $e_i \equiv 1[n_i]$  et  $e_i \equiv 0[n_j]$  pour tout  $j \in \llbracket 1, k \rrbracket \setminus \{i\}$ . Il ne reste plus qu'à prendre le reste modulo  $N$  de  $\sum_{i=1}^k a_i.e_i$ . L'unicité est une conséquence immédiate du premier point.  $\square$

*Remarque 1.3.23.* A l'occasion de cette preuve, on pourra remarquer l'utilisation du théorème de Bachet–Bézout pour trouver un inverse de  $\bar{k}$  pour la multiplication dans  $\mathbb{Z}/n\mathbb{Z}$  pour tout  $k \in \mathbb{Z}^*$  premier avec  $n \in \mathbb{N}^*$ .

## 2 Groupes

L'objectif de la théorie des groupes n'est pas d'étudier des objets en fonction de leur nature intrinsèque, mais en fonction de leurs comportement les uns par rapport aux autres. Plus que les objets, ce sont donc les fonctions définies sur ces objets que l'on étudie, ce qui permet de créer des modèles applicables dans différents contextes.

### 2.1 Structure de groupe

#### 2.1.1 Groupes et sous-groupes

##### Définition et règles élémentaires

**Définition 2.1.1.** Un *groupe* est un ensemble  $G$ , muni d'une opération  $*$  :  $G \times G \rightarrow G$ , appelée *loi de composition interne*, vérifiant les propriétés suivantes :

- (associativité)  $\forall g_1, g_2, g_3 \in G, (g_1 * g_2) * g_3 = g_1 * (g_2 * g_3)$  ;
- (existence d'un élément neutre)  $\exists e \in G, \forall g \in G, e * g = g * e = g$  ;
- (existence d'inverses)  $\forall g \in G, \exists h \in G, g * h$  et  $h * g$  sont des éléments neutres.

Si, de plus,  $*$  vérifie  $g_1 * g_2 = g_2 * g_1$  pour tous  $g_1, g_2 \in G$ , on dit que le groupe est *commutatif* ou *abélien*.

*Exemples 2.1.2.*

- $(\{e\}, ((e, e) \mapsto e))$  est un groupe (abélien) mais l'ensemble vide n'en est pas un, car il ne possède pas d'élément neutre.
- $(\mathbb{N}, +)$  n'est pas un groupe, mais  $(\mathbb{Z}, +)$  l'est.  $(\mathbb{Z}, \cdot)$  n'est pas un groupe,  $(\mathbb{Q}, \cdot)$  non plus, mais  $(\mathbb{Q}^*, \cdot)$  l'est.  $(\mathbb{R}, +)$ ,  $(\mathbb{C}, +)$ ,  $(\mathbb{R}^*, \cdot)$  et  $(\mathbb{C}, \cdot)$  sont aussi des groupes. Tous sont abéliens.
- $(\mathbb{Z}/n\mathbb{Z}, +)$  est un groupe abélien.  $(\mathbb{Z}/n\mathbb{Z}, \cdot)$  n'en est pas un, et même  $(\mathbb{Z}/n\mathbb{Z}^*, \cdot)$  n'en est pas un lorsque  $n$  n'est pas premier ; par exemple,  $\bar{2} \in \mathbb{Z}/4\mathbb{Z}$  n'a pas d'inverse quand bien même il n'est pas nul.
- Pour tout  $n \in \mathbb{N}^*$  et  $\mathbb{K} = \mathbb{R}$  ou  $\mathbb{C}$ ,  $(\mathcal{M}_n(\mathbb{K}), +)$  est un groupe abélien, mais  $(\mathcal{M}_n(\mathbb{K}), \cdot)$  ne l'est pas. Par contre  $(\mathrm{GL}_n(\mathbb{K}), \cdot)$  est bien un groupe, mais non abélien si  $n \geq 2$ .
- Le produit de matrices correspondant à la composition d'applications linéaires dans une base fixée, on en déduit que pour tout espace vectoriel  $E$ ,  $(\mathrm{Aut}(E), \circ)$ , où  $\mathrm{Aut}(E)$  est l'ensemble des automorphismes linéaires de  $E$ , est un groupe.  
Remarquons au passage que la notion de groupe est une notion qui peut se rencontrer simultanément à différents niveaux. Nous venons en effet de voir que les applications linéaires sur un espace vectoriel  $E$  forment un groupe pour l'addition, mais les éléments de  $E$ , munis de l'addition forment également un groupe.
- Plus généralement, pour tout ensemble  $\Omega$ ,  $(\mathrm{Bij}(\Omega), \circ)$ , où  $\mathrm{Bij}(\Omega)$  est l'ensemble des bijections de  $\Omega$  dans lui-même, est un groupe non abélien si  $|\Omega| \geq 3$ .
- On peut aussi restreindre les éléments considérés tout en conservant une structure de groupe. Ainsi, les isométries du plan ou de l'espace forment un groupe pour la composition, et même les isométries

préservant une certaine figure, par exemple un polygone régulier donné, forment également un groupe pour la composition.

*Avertissement.* Dans tout ce qui suit, les groupes ne seront souvent noté que par l'ensemble sous-jacent. On notera par exemple  $\mathbb{Z}$  pour  $(\mathbb{Z}, +)$ . Il faut néanmoins bien garder en tête que l'ensemble est muni implicitement d'une loi de composition interne. Il faudra notamment faire attention aux ensembles possédant plusieurs structures de groupes ; si un même groupe  $G$  intervient plusieurs fois dans un énoncé, c'est que l'on considère bien entendu à chaque fois la même loi de composition interne.

**Proposition 2.1.3.** Dans tout groupe, l'élément neutre et l'inverse d'un élément sont uniques.

*Démonstration.* Soit  $G$  un groupe et  $e_1, e_2 \in G$  deux éléments neutres. Alors,  $e_1$  étant neutre, on a  $e_1 * e_2 = e_2$  ; mais  $e_2$  étant également neutre, on a également  $e_1 * e_2 = e_1$ . On en déduit que  $e_1 = e_2$  et donc que l'élément neutre est unique. On le note  $e$ .

Soit  $g \in G$  et  $h_1, h_2 \in G$  deux inverses pour  $g$ . On a alors  $h_1 = h_1 * e = h_1 * (g * h_2) = (h_1 * g) * h_2 = e * h_2 = h_2$ .  $\square$

**Notation 2.1.4.** Soit  $G$  un groupe. Plutôt que  $*$ , on notera le plus souvent sa loi de composition interne par  $\cdot$ . C'est ce qu'on appelle la notation *multiplicative*. Par convention, on notera alors

- $e_G$  ou  $e$  son élément neutre ;
- pour tout  $g \in G$ ,  $g^{-1}$  l'inverse de  $g$  ;
- pour tout  $g \in G$  et tout  $n \in \mathbb{N}^*$ ,  $g^n := \underbrace{g \cdot g \cdot \dots \cdot g}_{n \text{ fois}}$ ,  $g^{-n} := (g^{-1})^n$ , ainsi que  $g^0 := e_G$ .

Toutefois, si  $G$  est notoirement abélien, par analogie avec les entiers, sa loi de composition interne sera plutôt notée  $+$ . C'est ce qu'on appelle la notation *additive*. Dans ce cas, on notera plutôt :

- $0_G$  ou  $0$  à la place de  $e_G$  ;
- pour tout  $g \in G$ ,  $-g$  à la place de  $g^{-1}$ , et on l'appellera alors *opposé* de  $g$  ;
- pour tout  $g \in G$  et  $n \in \mathbb{Z}$ ,  $n \cdot g$  à la place de  $g^n$ , en observant que ça reste cohérent, même si  $\mathbb{Z} \subset G$ .

Voyons maintenant quelques règles simples et universelles de calculs.

**Proposition 2.1.5.** Soit  $G$  un groupe. Alors

- pour tout  $g \in G$ ,  $(g^{-1})^{-1} = g$  ;
- pour tous  $g_1, g_2 \in G$ ,  $(g_1 \cdot g_2)^{-1} = g_2^{-1} \cdot g_1^{-1}$  ;
- pour tout  $g \in G$  et tout  $n \in \mathbb{Z}$ ,  $g^{-n} = (g^n)^{-1}$  ;
- pour tout  $g \in G$  et tous  $n_1, n_2 \in \mathbb{Z}$ ,  $g^{n_1} \cdot g^{n_2} = g^{n_1+n_2}$  et  $(g^{n_1})^{n_2} = g^{n_1 \cdot n_2}$ .

*Démonstration.* Par unicité de l'inverse, il suffit de constater que  $g^{-1} \cdot g = g \cdot g^{-1} = e$ ,  $g_2^{-1} \cdot g_1^{-1} \cdot g_1 \cdot g_2 = g_2^{-1} \cdot g_2 = e$  et  $g^{-n} \cdot g^n = g^{-1} \cdot \dots \cdot g^{-1} \cdot g \cdot \dots \cdot g = e$  pour montrer les trois premiers points.

La première partie du quatrième point se montre sans difficulté par récurrence sur  $n_1 + n_2$  lorsque  $n_1 + n_2 \in \mathbb{N}$ . Si  $n_1 + n_2 = 0$ , on a en effet  $g^{n_1} \cdot g^{n_2} = g^{n_1} \cdot g^{-n_1} = e = g^0 = g^{n_1+n_2}$  et, une fois le résultat supposé vrai pour  $n_1 + n_2 \geq 0$ , on a  $g^{n_1+n_2+1} = g^{n_1+n_2} \cdot g = g^{n_1} \cdot g^{n_2} \cdot g = g^{n_1} \cdot g^{n_2+1}$  ainsi que  $g^{n_1+n_2+1} = g \cdot g^{n_1+n_2} = g \cdot g^{n_1} \cdot g^{n_2} = g^{n_1+1} \cdot g^{n_2}$ . Et lorsque  $n_1 + n_2 < 0$ , alors on a  $g^{n_1+n_2} = (g^{-1})^{-n_1-n_2} = (g^{-1})^{-n_1} \cdot (g^{-1})^{-n_2} = g^{n_1} \cdot g^{n_2}$ .

Enfin, le dernier point se démontre également par récurrence lorsque  $n_2 \in \mathbb{N}$ . On a en effet  $(g^{n_1})^0 = e = g^0 = g^{n_1 \cdot 0}$  et, en supposant  $(g^{n_1})^{n_2} = g^{n_1 \cdot n_2}$ ,  $(g^{n_1})^{n_2+1} = (g^{n_1})^{n_2} \cdot g^{n_1} = g^{n_1 \cdot n_2} \cdot g^{n_1} = g^{n_1 \cdot n_2 + n_1} = g^{n_1 \cdot (n_2+1)}$ . Pour  $n_2 \in \mathbb{Z} \setminus \mathbb{N}$ , on écrit  $(g^{n_1})^{n_2} = ((g^{n_1})^{-1})^{-n_2} = (g^{-n_1})^{-n_2} = g^{(-n_1) \cdot (-n_2)} = g^{n_1 \cdot n_2}$ .  $\square$

Par ailleurs la caractérisation suivante de l'élément neutre se révélera utile.

**Lemme 2.1.6.** Dans tout groupe  $G$ ,  $e_G$  est le seul élément égal à son carré.

*Démonstration.* On a clairement  $e_G.e_G = e_G$  et, réciproquement, si  $g \in G$  vérifie  $g.g = g$ , alors  $g = g.(g.g^{-1}) = (g.g).g^{-1} = g.g^{-1} = e_G$ .  $\square$

**Corollaire 2.1.7.** Soit  $G$  un groupe et  $g, h \in G$ . Alors les affirmations suivantes sont équivalentes :

- i.  $h = g^{-1}$  ;
- ii.  $g.h = e_G$  ;
- iii.  $h.g = e_G$ .

*Démonstration.* D'après les règles de calcul précédentes, on a  $h = g^{-1} \Leftrightarrow g = h^{-1}$ , ce qui symétrise les rôles de  $g$  et  $h$ . Par définition de l'inverse, il suffit donc de montrer ii.  $\Rightarrow$  iii.. Supposons donc  $g.h = e_G$ , alors on a  $(h.g).(h.g) = h.(g.h).g = h.e_G.g = h.g$  et donc, d'après le lemme précédent,  $h.g = e_G$ .  $\square$

## Tables de loi

**Définition 2.1.8.** On appelle *ordre* d'un groupe  $G$  son cardinal en tant qu'ensemble ; on le note  $|G|$ . On dit que  $G$  est *fini* si son ordre est fini.

Profitions-en pour définir également l'ordre d'un élément.

**Définition 2.1.9.** Pour tout élément  $g$  d'un groupe  $G$ , on appelle *ordre de  $g$* , noté  $|g|$ , l'entier  $\min\{k \in \mathbb{N}^* \mid g^k = e\} \in \mathbb{N}^* \cup \{\infty\}$  avec la convention que  $\min \emptyset = \infty$ .

Lorsqu'un groupe  $G$  est fini d'ordre  $n \in \mathbb{N}^*$ , on peut indexer arbitrairement<sup>3</sup> ses éléments  $g_1, g_2, \dots, g_n$  ; la structure de groupe est alors intégralement donnée par la *table* de sa loi de composition interne, une matrice carrée de taille  $n$ , dont le coefficient d'indice  $i, j$  vaut  $g_i.g_j$ .

*Exemples 2.1.10.* Dans les exemples précédents, seul  $(\mathbb{Z}/n\mathbb{Z}, +)$  est fini, avec  $|\mathbb{Z}/n\mathbb{Z}| = n$ . En ordonnant les éléments ainsi :  $\bar{0}, \bar{1}, \dots, \overline{n-1}$ , on obtient la matrice circulante suivante comme table de loi :

$$\begin{pmatrix} \bar{0} & \bar{1} & \bar{2} & \dots & \dots & \overline{n-1} \\ \bar{1} & \bar{2} & \dots & \dots & \overline{n-1} & \bar{0} \\ \bar{2} & \dots & \dots & \overline{n-1} & \bar{0} & \bar{1} \\ \vdots & & & & & \vdots \\ \overline{n-1} & \bar{0} & \bar{1} & \dots & \dots & \overline{n-2} \end{pmatrix}.$$

Dans le cas  $n = 4$ , cela donne

$$\begin{pmatrix} \bar{0} & \bar{1} & \bar{2} & \bar{3} \\ \bar{1} & \bar{2} & \bar{3} & \bar{0} \\ \bar{2} & \bar{3} & \bar{0} & \bar{1} \\ \bar{3} & \bar{0} & \bar{1} & \bar{2} \end{pmatrix},$$

mais ça n'est pas le seul groupe (abélien) d'ordre 4, on pourra vérifier que, sur  $\{e, a, b, c\}$ , la table

$$\begin{pmatrix} e & a & b & c \\ a & e & c & b \\ b & c & e & a \\ c & b & a & e \end{pmatrix}$$

---

3. toutefois, on commence en général par l'élément neutre

défini également un groupe, distinct de  $(\mathbb{Z}/4\mathbb{Z}, +)$ . Son caractère abélien se lit par le caractère symétrique de la matrice. Sur l'ensemble  $\{e, \sigma_{12}, \sigma_{23}, \sigma_{31}, \sigma_{123}, \sigma_{321}\}$ , la table

$$\begin{pmatrix} e & \sigma_{12} & \sigma_{23} & \sigma_{31} & \sigma_{123} & \sigma_{321} \\ \sigma_{12} & e & \sigma_{123} & \sigma_{321} & \sigma_{23} & \sigma_{31} \\ \sigma_{23} & \sigma_{321} & e & \sigma_{123} & \sigma_{31} & \sigma_{12} \\ \sigma_{31} & \sigma_{123} & \sigma_{123} & e & \sigma_{12} & \sigma_{23} \\ \sigma_{123} & \sigma_{31} & \sigma_{12} & \sigma_{23} & \sigma_{321} & e \\ \sigma_{321} & \sigma_{23} & \sigma_{31} & \sigma_{12} & e & \sigma_{123} \end{pmatrix}$$

donne une structure de groupe non abélien.

### Sous-groupes

Pour un groupe donné, on peut s'intéresser aux sous-ensembles qui restent des groupes pour l'opération donnée.

**Définition 2.1.11.** Soit  $G$  un groupe. On dit que  $H \subset G$  est un *sous-groupe* de  $G$  si  $H$  n'est pas vide et qu'il est stable par loi de composition interne de  $G$  et prise d'inverse, c'est-à-dire si

- $H \neq \emptyset$ ;
- $\forall h_1, h_2 \in H, h_1.h_2 \in H$ ;
- $\forall h \in H, h^{-1} \in H$ .

Il existe plusieurs caractérisations équivalentes des sous-groupes.

**Proposition 2.1.12.** Soit  $(G, .)$  un groupe. Pour tout  $H \subset G$ , les propositions suivantes sont équivalentes :

- i.  $H$  est un sous-groupe de  $G$ ;
- ii.  $H$  non vide et pour tout  $x, y \in H, x.y^{-1} \in H$ ;
- iii.  $(H, .)$  est un groupe.

*Démonstration.* Il est clair que i.  $\Rightarrow$  ii.. Réciproquement, si ii. est vrai, alors il existe  $h \in H$  et  $e_G = h.h^{-1} \in H$ . Dès lors,  $H$  est stable par prise d'inverse car, pour tout  $h \in H, h^{-1} = e_G.h^{-1} \in H$ . Et cela implique enfin la stabilité par loi de composition interne, car pour tout  $h_1, h_2 \in H$ , on a donc  $h_2^{-1} \in H$  et donc  $h_1.h_2 = h_1.(h_2^{-1})^{-1} \in H$ .

Il est également clair que si  $(H, .)$  est un groupe, alors  $H$  n'est pas vide car il contient  $e_H$  et est stable par  $.$  et par prise d'inverse.

Réciproquement, supposons que  $H$  est un sous-groupe de  $G$ . On peut alors bien considérer  $._|_H : H \times H \rightarrow H$  et cette opération est associative car elle l'était déjà sur  $G$ . Montrons qu'il contient un élément neutre. Puisque  $H$  n'est pas vide, il contient un élément  $h \in H$ . Par stabilité par prise d'inverse, on a alors  $h^{-1} \in H$ , et par stabilité par  $.$  on a  $e_G = h.h^{-1} \in H$ ; ce dernier vérifie bien la propriété de l'élément neutre. Enfin, tout élément de  $H$  possède un inverse dans  $G$ , qui est aussi dans  $H$  par stabilité de  $H$  par prise d'inverse.  $\square$

*Exemples 2.1.13.*

- On a les suites d'inclusions de groupes  $(\mathbb{Z}, +) \subset (\mathbb{Q}, +) \subset (\mathbb{R}, +) \subset (\mathbb{C}, +)$  et

$$\begin{array}{ccccc} & & (\mathbb{Q}_+^*, .) \subset (\mathbb{R}_+^*, .) & & \\ & \subset & \cap & \subset & \\ (\{1\}, .) & & & & (\mathbb{C}^*, .) \\ & \subset & (\mathbb{Q}^*, .) \subset (\mathbb{R}^*, .) & \subset & \end{array}$$

Par contre, bien que stable par multiplication, ni  $(\mathbb{N}^*, .)$  ni  $(\mathbb{Z}^*, .)$  ne sont des sous-groupes de  $(\mathbb{R}^*, .)$  car ils ne sont pas stables par prise d'inverse.

- Pour  $n \in \mathbb{N}^*$ ,  $n\mathbb{Z} := \{\text{multiple de } n\}$  est un sous-groupe de  $(\mathbb{Z}, +)$ . On peut même montrer que tout sous-groupe de  $\mathbb{Z}$  est de cette forme.
- Pour tout  $n \in \mathbb{N}^*$  et  $\mathbb{K} = \mathbb{R}$  ou  $\mathbb{C}$ ,  $\text{Sym}_n(\mathbb{K}) := \{M \in \mathcal{M}_n(\mathbb{K}) \mid M \text{ symétrique}\}$  est un sous-groupe de  $(\mathcal{M}_n(\mathbb{K}), +)$ ; et  $\text{SL}_n(\mathbb{K}) := \{M \in \mathcal{M}_n(\mathbb{K}) \mid \det(M) = 1\}$  est un sous-groupe de  $(\text{GL}_n(\mathbb{K}), \cdot)$ .
- L'ensemble des isométries du plan est un sous-groupe de  $(\text{Bij}(\mathbb{R}^2), \circ)$ , que l'on note  $\text{Isom}(\mathbb{R}^2)$ .
- Pour toute figure  $F \subset \mathbb{R}^2$ , par exemple un triangle ou une ellipse, l'ensemble des isométries envoyant  $F$  sur lui-même est un sous-groupe de  $\text{Isom}(\mathbb{R}^2)$  que l'on nomme groupe de symétrie de  $F$ . Plus généralement, on pourra parler de groupe de symétrie pour tout ensemble de transformations préservant une structure donnée; c'est une notion importante dans beaucoup de domaines en mathématiques, elle est même à l'origine de la notion de groupe.

### Intersections de sous-groupes

Il existe deux opérations fondamentales sur les ensembles, la réunion et l'intersection. Comme l'essentiel des structures algébriques, la structure de groupe est fortement compatible avec la seconde, mais pas avec la première.

**Proposition 2.1.14.** Soit  $G$  un groupe, et  $(G_i)_{i \in I}$  une famille de sous-groupes de  $G$ . Alors  $\cap_{i \in I} G_i$  est un sous-groupe de  $G$ .

*Démonstration.* Il suffit de remarquer que  $\cap_{i \in I} G_i$  est non vide car il contient  $e_G$ , et que, pour tout  $g_1, g_2 \in G$ ,  $(g_1, g_2 \in \cap_{i \in I} G_i) \Leftrightarrow (\forall i \in I, g_1, g_2 \in G_i) \Rightarrow (\forall i \in I, g_1 \cdot g_2^{-1} \in G_i) \Leftrightarrow (g_1 \cdot g_2^{-1} \in \cap_{i \in I} G_i)$ .  $\square$

*Remarque 2.1.15.* A l'inverse, la réunion de deux sous-groupes n'est que très rarement un sous-groupe. Par exemple, dans  $\mathbb{R}^2$  vu comme un  $\mathbb{R}$ -espace vectoriel de dimension 2, les sous ensembles  $E_1 := \{(\lambda, 0) \mid \lambda \in \mathbb{R}\}$  et  $E_2 := \{(0, \lambda) \mid \lambda \in \mathbb{R}\}$  sont chacun des sous-groupes de  $(\mathbb{R}^2, +)$ , mais  $E_1 \cup E_2$  ne l'est pas car  $(1, 0), (0, 1) \in E_1 \cup E_2$  mais  $(1, 1) = (1, 0) + (0, 1) \notin E_1 \cup E_2$ .

**Définition 2.1.16.** Soit  $G$  un groupe et  $X \subset G$ , on note  $\langle X \rangle$  l'intersection de tous les sous-groupes de  $G$  contenant  $X$ . Lorsque  $X$  est décrit explicitement, on omettra les accolades dans l'écriture  $\langle X \rangle$ ;  $\langle \{a, b, c\} \rangle$  sera par exemple écrit  $\langle a, b, c \rangle$ .

**Proposition 2.1.17.** Soit  $G$  un groupe. Pour tout  $X \subset G$ ,  $\langle X \rangle$  est un sous-groupe de  $G$ , contenu dans tout sous-groupe de  $G$  contenant  $X$ .

*Remarque 2.1.18.* Le sous-groupe  $\langle X \rangle$  est donc le plus petit sous-groupe de  $G$  contenant  $X$ ; on l'appelle *sous-groupe engendré par  $X$* .

*Démonstration.* Le fait que  $\langle X \rangle$  soit un groupe est une conséquence directe de la proposition 2.1.14, et la seconde affirmation est une conséquence directe de la définition de  $\langle X \rangle$ .  $\square$

*Exemples 2.1.19.*

- Pour tout groupe  $G$ , on a  $\langle G \rangle = G$  et  $\langle \emptyset \rangle = \langle e \rangle = \{e\}$ .
- Dans  $\mathbb{Z}$ , on a  $\langle 2 \rangle = 2\mathbb{Z}$ ,  $\langle 2, 4 \rangle = 2\mathbb{Z}$  et  $\langle 2, 3 \rangle = \mathbb{Z}$ .
- Dans  $\mathcal{M}_n(\mathbb{R})$ , pour tout  $1 \leq i, j \leq n$ , on note  $E_{i,j}$  la matrice n'ayant qu'un seul coefficient non nul, 1 en position  $(i, j)$ . Dans  $(\mathcal{M}_n(\mathbb{R}), +)$ , on a alors  $\langle E_{i,j} \mid 1 \leq i, j \leq n \rangle = \mathcal{M}_n(\mathbb{R})$ , tandis que dans  $(\text{GL}_n(\mathbb{R}), \cdot)$ ,  $\langle \text{Id} + E_{i,j} \mid 1 \leq i \neq j \leq n \rangle = \text{SL}_n(\mathbb{R})$ .

Il est possible de donner une description plus intrinsèque des éléments d'un sous-groupe engendré.

**Proposition 2.1.20.** Soit  $G$  un groupe et  $X =: \{g_i \mid i \in I\}$  une famille non vide d'éléments de  $G$ . On a alors

$$\langle X \rangle = \{g_{i_1}^{\varepsilon_1} g_{i_2}^{\varepsilon_2} \cdots g_{i_k}^{\varepsilon_k} \mid k \in \mathbb{N}, i_1, i_2, \dots, i_k \in I, \varepsilon_1, \varepsilon_2, \dots, \varepsilon_k \in \{\pm 1\}\},$$

avec la convention qu'un produit vide vaut  $e_G$ .



*Démonstration.* Notons  $H := \{g_{i_1}^{\varepsilon_1} g_{i_2}^{\varepsilon_2} \cdots g_{i_k}^{\varepsilon_k} \mid k \in \mathbb{N}, i_1, i_2, \dots, i_k \in I, \varepsilon_1, \varepsilon_2, \dots, \varepsilon_k \in \{\pm 1\}\}$ . Puisque  $X \subset \langle X \rangle$  et par stabilité du sous-groupe  $\langle X \rangle$  par produit et prise d'inverse, il est clair que  $H \subset \langle X \rangle$ . mais réciproquement,  $H$  contient  $e_G$  comme produit vide, et il est clairement stable par produit et prise d'inverse ; c'est donc un sous-groupe de  $G$ . De plus, il contient  $X$ , on en déduit que  $\langle X \rangle \subset H$ .  $\square$

*Remarque 2.1.21.* Ce dernier résultat permet notamment d'expliquer *a posteriori* pourquoi le cardinal d'un groupe et la plus petite puissance triviale d'un élément s'appelle tous les deux ordres, noté  $|\cdot|$ . En effet, pour tout  $g \in G$  d'ordre fini, cela donne  $\langle g \rangle = \{g^k \mid k \in \mathbb{Z}\}$ , ce qui, par division euclidienne de  $k$  par  $|g|$ , est égal à  $\{g^k \mid k \in \llbracket 0, |g| - 1 \rrbracket\}$  ; on a donc  $|g| = |\langle g \rangle|$ .

**Définition 2.1.22.** On dit que  $X \subset G$  est une *partie génératrice* d'un groupe  $G$  si  $\langle X \rangle = G$ .

Connaitre une partie génératrice d'un groupe, pour peu que celle-ci soit simple, se révélera souvent très utile pour étudier ce groupe.

## 2.1.2 Morphismes de groupes

### Définition et propriétés élémentaires

L'idée fondamentale de la théorie des groupes est de pouvoir comparer des ensembles dont les éléments se comportent de manière similaire vis-à-vis d'opérations données, quand bien même les ensembles sont d'origines très différentes. Il faut, pour cela, avoir une notion d'applications entre groupes qui respectent les dites structures de groupes.

**Définition 2.1.23.** Soit  $G_1, G_2$  deux groupes dont on notera respectivement  $*_1$  et  $*_2$  les lois de composition interne. On dit qu'une application  $f : G_1 \rightarrow G_2$  est un *morphisme de groupes* si, pour tous  $g_1, g_2 \in G_1$ , on a  $f(g_1 *_1 g_2) = f(g_1) *_2 f(g_2)$ .

Plus précisément, on parlera même

- de *monomorphisme (de groupes)* si  $f$  est injective ;
- d'*épimorphisme (de groupes)* si  $f$  est surjective ;
- d'*isomorphisme (de groupes)* si  $f$  est bijective ;
- d'*endomorphisme (de groupe)* si  $G_1 = G_2$  ;
- d'*automorphisme (de groupe)* si  $G_1 = G_2$  et que  $f$  est bijective.

**Définition 2.1.24.** On dit que deux groupes  $G_1$  et  $G_2$  sont *isomorphes* s'il existe un isomorphisme de groupes  $f : G_1 \rightarrow G_2$ . On note alors  $G_1 \cong G_2$ .

*Remarque 2.1.25.* On considère généralement que deux groupes isomorphes représentent le “même” groupe. Un isomorphisme entre eux correspond en effet à un renommage des éléments de  $G_1$  par ceux de  $G_2$ , mais le comportement des éléments entre eux reste le même. De fait, on ne considère souvent les groupes qu'à isomorphisme près.

*Exemples 2.1.26.*

- Pour tout groupe  $G$ , les applications

$$\begin{array}{ccc} G & \longrightarrow & G \\ \text{Id}_G : & & \\ g & \longmapsto & g \end{array} \quad \text{et} \quad \begin{array}{ccc} G & \longrightarrow & \{e\} \\ \text{Triv}_G : & & \\ g & \longmapsto & e \end{array}$$

sont, respectivement un automorphisme et un épimorphisme de groupes.

- Pour tout groupe  $m \in \mathbb{N}^*$ , l'application

$$\begin{array}{ccc} \mathbb{Z} & \longrightarrow & \mathbb{Z} \\ \text{mult}_m : & & \\ k & \longmapsto & m.k \end{array} \quad \text{et} \quad \begin{array}{ccc} \mathbb{Z}/n\mathbb{Z} & \longrightarrow & \mathbb{Z}/n\mathbb{Z} \\ \text{mult}_m^n : & & \\ \bar{k} & \longmapsto & \overline{m.k} \end{array}$$

sont des endomorphismes de groupe.

- Les applications

$$\begin{array}{ccc} \mathbb{R} & \longrightarrow & \mathbb{R}_+^* \\ \exp: x & \longmapsto & e^x \end{array} \quad \text{et} \quad \begin{array}{ccc} \mathbb{R}_+^* & \longrightarrow & \mathbb{R} \\ \ln: x & \longmapsto & \ln(x) \end{array}$$

sont des isomorphismes entre  $(\mathbb{R}, +)$  et  $(\mathbb{R}_+^*, \cdot)$ . De même,

$$\begin{array}{ccc} \mathbb{C} & \longrightarrow & \mathbb{C}^* \\ \exp: z & \longmapsto & e^z \end{array}$$

est un épimorphisme de groupes entre  $(\mathbb{C}, +)$  et  $(\mathbb{C}^*, \cdot)$ .

- Pour  $\mathbb{K} = \mathbb{R}$  ou  $\mathbb{C}$ , l'application

$$\begin{array}{ccc} \mathrm{GL}_n(\mathbb{K}) & \longrightarrow & \mathbb{K}^* \\ \det: M & \longmapsto & \det(M) \end{array}$$

est un épimorphisme de groupes.

**Proposition 2.1.27.** Pour tout morphisme de groupes  $f : G_1 \rightarrow G_2$ , on a  $f(e_{G_1}) = e_{G_2}$  et  $f(g^{-1}) = f(g)^{-1}$  pour tout  $g \in G_1$ . Plus généralement, on a  $f(g^n) = f(g)^n$  pour tout  $n \in \mathbb{Z}$ .

*Démonstration.* On a  $f(e_{G_1}) = f(e_{G_1} \cdot e_{G_1}) = f(e_{G_1}) \cdot f(e_{G_1})$  et donc  $f(e_{G_1}) = e_{G_2}$  d'après le lemme 2.1.6. Et pour tout  $g \in G$ , on a  $f(g) \cdot f(g^{-1}) = f(g \cdot g^{-1}) = f(e_{G_1}) = e_{G_2}$ , ce qui, d'après le corollaire 2.1.7, suffit pour dire que  $f(g^{-1}) = f(g)^{-1}$ . La dernière affirmation se prouve par récurrence pour  $n \in \mathbb{N}$  puis, pour  $n \in \mathbb{Z} \setminus \mathbb{N}$ , en remarquant que  $f(g^n) = f((g^{-1})^{-n}) = f((g^{-1}))^{-n} = (f(g)^{-1})^{-n} = f(g)^n$ .  $\square$

## Noyaux et images

**Définition 2.1.28.** Pour tout morphisme de groupes  $f : G_1 \rightarrow G_2$ , on définit son *noyau* comme  $\mathrm{Ker}(f) = f^{-1}(e_{G_2})$  et son *image* comme l'ensemble  $\mathrm{Im}(f) = f(G_1)$ .

Une propriété essentielle de la théorie des groupes est la suivante :

**Proposition 2.1.29.** Un morphisme de groupes  $f : G_1 \rightarrow G_2$  est injectif si et seulement si  $\mathrm{Ker}(f) = \{e_{G_1}\}$ .

*Démonstration.* Si  $f$  est injectif, on a clairement  $\mathrm{Ker}(f) = \{e_{G_1}\}$ . Réciproquement, supposons  $\mathrm{Ker}(f) = \{e_{G_1}\}$  et considérons  $g_1, g_2 \in G$  tels que  $f(g_1) = f(g_2)$ . On a alors  $f(g_1^{-1} \cdot g_2) = f(g_1)^{-1} \cdot f(g_2) = f(g_1)^{-1} \cdot f(g_1) = e_{G_2}$  et donc  $g_1^{-1} \cdot g_2 \in \mathrm{Ker}(f) = \{e_{G_1}\}$ . On en déduit que  $g_1^{-1} \cdot g_2 = e_{G_1}$  et donc que  $g_1 = g_2$ .  $\square$

*Exemples 2.1.30.* En reprenant les notations des exemples plus haut, on a :

- $\mathrm{Ker}(\mathrm{Id}_G) = \{e_G\}$  et donc  $\mathrm{Id}_G$  injectif, mais  $\mathrm{triv}_G$  ne l'est que si  $G$  est trivial, car  $\mathrm{Ker}(\mathrm{triv}_G) = G$ ;
- $\mathrm{Ker}(\mathrm{mult}_m) = \{0\}$  donc  $\mathrm{mult}_m$  injectif, mais  $\mathrm{mult}_m^n$  ne le sera que si  $\mathrm{pgcd}(n, m) = 1$  car on peut montrer que  $\mathrm{Ker}(\mathrm{mult}_m^n)$  est engendré par la classe de  $\frac{n}{\mathrm{pgcd}(n, k)}$  dans  $\mathbb{Z}/n\mathbb{Z}$ ;
- dans  $\mathbb{R}$ ,  $e^x = 1$  si et seulement si  $x = 0$ , donc  $\exp : \mathbb{R} \rightarrow \mathbb{R}_+^*$  est injective, mais sur  $\mathbb{C}$ ,  $e^z = 1$  si et seulement si  $z \in \{2k\pi i \mid k \in \mathbb{Z}\}$ , donc  $\exp : \mathbb{C} \rightarrow \mathbb{C}^*$  n'est pas injective;
- $\mathrm{Ker}(\det) = \mathrm{SL}_n(\mathbb{K})$ , donc  $\det$  n'est injective que si  $n = 1$ .

Les morphismes de groupes peuvent également être utiles pour construire de nouveaux sous-groupes.

**Proposition 2.1.31.** Soit  $f : G_1 \rightarrow G_2$  un morphisme de groupes. Pour tout sous-groupe  $H_2 \subset G_2$ ,  $f^{-1}(H_2)$  est un sous-groupe de  $G_1$ ; et pour tout sous-groupe  $H_1 \subset G_1$ ,  $f(H_1)$  est un sous-groupe de  $G_2$ . En particulier,  $\mathrm{Ker}(f)$  et  $\mathrm{Im}(f)$  sont respectivement des sous-groupes de  $G_1$  et  $G_2$ .

*Démonstration.* Montrons que  $f^{-1}(H_2) \subset G_1$  est un sous-groupe. Pour cela, on considère  $g_1, g_2 \in f^{-1}(H_2)$ , on a alors  $f(g_1), f(g_2) \in H_2$ , et donc  $f(g_1 \cdot g_2^{-1}) = f(g_1) \cdot f(g_2)^{-1} \in H_2$ . On en déduit que  $g_1 \cdot g_2^{-1} \in f^{-1}(H_2)$ , montrant donc que  $f^{-1}(H_2)$  est bien un sous-groupe de  $G_1$ .

Considérons maintenant  $g_1, g_2 \in f(H_1)$ ; il existe donc  $h_1, h_2 \in H_1$  tels que  $g_1 = f(h_1)$  et  $g_2 = f(h_2)$ . Mais alors  $g_1 \cdot g_2^{-1} = f(h_1) \cdot f(h_2)^{-1} = f(h_1 \cdot h_2^{-1}) \in f(H_1)$  puisque  $h_1 \cdot h_2^{-1} \in H_1$ .  $\square$

## Composition de morphismes

### Proposition 2.1.32.

- Si  $f_1 : G_1 \rightarrow G_2$  et  $f_2 : G_2 \rightarrow G_3$  sont deux morphismes de groupes, alors  $f_2 \circ f_1 : G_1 \rightarrow G_3$  est aussi un morphisme de groupes.
- Si  $f : G_1 \rightarrow G_2$  est un isomorphisme de groupes, alors  $f^{-1} : G_2 \rightarrow G_1$  est aussi un morphisme de groupes.

*Démonstration.* Pour la première affirmation, notons  $*_1, *_2$  et  $*_3$  les lois de composition internes de  $G_1, G_2$  et  $G_3$ , et considérons  $g_1, g_2 \in G_1$ . On a alors

$$(f_2 \circ f_1)(g_1 *_1 g_2) = f_2(f_1(g_1 *_1 g_2)) = f_2(f_1(g_1) *_2 f_1(g_2)) = f_2(f_1(g_1)) *_3 f_2(f_1(g_2)) = (f_2 \circ f_1)(g_1) *_3 (f_2 \circ f_1)(g_2).$$

L'application  $f_2 \circ f_1$  est donc bien un morphisme de groupes.

Pour la seconde, il suffit d'observer que, pour tout  $g_1, g_2 \in G_2$ , on a

$$f(f^{-1}(g_1) *_1 f^{-1}(g_2)) f(f^{-1}(g_1)) *_2 f(f^{-1}(g_2)) = g_1 *_2 g_2 = f(f^{-1}(g_1 *_2 g_2)).$$

Par injectivité de  $f$ , on a bien  $f^{-1}(g_1 *_2 g_2) = f^{-1}(g_1) *_1 f^{-1}(g_2)$ .  $\square$

**Corollaire 2.1.33.** Pour tout groupe  $G$ ,  $(\text{Aut}(G), \circ)$ , où  $\text{Aut}(G) := \{\text{automorphisme de } G\}$ , est un groupe.

*Démonstration.* D'après la proposition précédente, la composition définit bien une application de  $\text{Aut}(G) \times \text{Aut}(G)$  dans  $\text{Aut}(G)$ . Celle-ci est bien associative car, pour tout  $f_1, f_2, f_3 \in \text{Aut}(G)$  et tout  $g \in G$ , on a  $(f_1 \circ (f_2 \circ f_3))(g) = f_1(f_2(f_3(g))) = ((f_1 \circ f_2) \circ f_3)(g)$ . Elle admet un élément neutre, à savoir  $\text{Id}_G$ . Et tout élément  $f \in \text{Aut}(G)$  admet bien un inverse, à savoir  $f^{-1}$ , qui est encore bien dans  $\text{Aut}(G)$  d'après la proposition précédente.  $\square$

Terminons cette partie par une mise en abyme.

**Définition 2.1.34.** Soit  $G$  un groupe. On dit  $h' \in G$  est le *conjugué* de  $h \in G$  par  $g \in G$  si  $h' = g \cdot h \cdot g^{-1}$ . On dit aussi que  $h'$  est obtenu à partir de  $h$  par *conjugaison*.

**Proposition 2.1.35.** Soit  $G$  un groupe. Pour tout  $g \in G$ , l'application

$$\text{conj}_g : \begin{array}{ccc} G & \longrightarrow & G \\ h & \longmapsto & g \cdot h \cdot g^{-1} \end{array}$$

est un automorphisme de  $G$ , et l'application

$$\text{conj} : \begin{array}{ccc} G & \longrightarrow & \text{Aut}(G) \\ g & \longmapsto & \text{conj}_g \end{array}$$

est un morphisme de groupes.

*Démonstration.* Soit  $g \in G$ . Montrons que  $\text{conj}_g$  est un morphisme de groupe. Pour tout  $h_1, h_2 \in G$ , on a en effet

$$\text{conj}_g(h_1.h_2) = g.h_1.h_2.g^{-1} = g.h_1.g^{-1}.g.h_2.g^{-1} = \text{conj}_g(h_1).\text{conj}_g(h_2).$$

Mais par ailleurs, pour tout  $g_1, g_2 \in G$  et tout  $h \in G$ , on a aussi

$$(\text{conj}_{g_1} \circ \text{conj}_{g_2})(h) = \text{conj}_{g_1}(g_2.h.g_2^{-1}) = g_1.g_2.h.g_2^{-1}.g_1^{-1} = (g_1.g_2).h.(g_1.g_2)^{-1} = \text{conj}_{g_1.g_2}(h).$$

On en déduit

- d'une part, que pour tout  $g \in G$ ,  $\text{conj}_g \circ \text{conj}_{g^{-1}} = \text{conj}_{g.g^{-1}} = \text{conj}_{e_G} = \text{Id}_G$ , donc que  $\text{conj}_g$  est inversible, et donc que  $\text{conj}_g \in \text{Aut}(G)$ ;
- d'autre part, que  $\text{conj}$  est en effet un morphisme de groupes.

□

*Remarque 2.1.36.* L'application  $\text{conj} : G \rightarrow \text{Aut}(G)$  n'est en général pas surjective. On note parfois  $\text{Int}(G)$  son image, et ses éléments sont appelés *automorphismes intérieurs* de  $G$ .

### 2.1.3 Groupes quotient

**Proposition 2.1.37.** Soit  $G$  un groupe et  $H \subset G$  un sous-groupe. La relation  $\sim_H$  définie sur  $G$  par  $g_1 \sim_H g_2 \Leftrightarrow g_1.g_2^{-1} \in H$  est une relation d'équivalence.

*Démonstration.* La relation est réflexive car  $H$  contient l'élément neutre et donc, pour tout  $g \in G$ ,  $g.g^{-1} = e \in H$ . Elle est symétrique car  $H$  est stable par inverse et donc, pour tous  $g_1, g_2 \in G$ ,  $g_1.g_2^{-1} \in H \Leftrightarrow (g_1.g_2^{-1})^{-1} \in H \Leftrightarrow g_2.g_1^{-1} \in H$ . Elle est transitive car  $H$  est stable par loi de composition interne et donc, pour tout  $g_1, g_2, g_3 \in G$ ,  $g_1.g_2^{-1}, g_2.g_3^{-1} \Rightarrow g_1.g_2^{-1}.g_2.g_3^{-1} = g_1.g_3^{-1} \in H$ . □

*Remarque 2.1.38.* Plus précisément, on a défini ici la notion de classe à droite d'un élément par rapport à  $H$ . En posant  $g_1H \sim g_2 \Leftrightarrow g_1^{-1}.g_2 \in H$ , on définit la notion de classe à gauche. Ces deux notions sont, en général, différentes mais ont des propriétés très similaires. Pour ce qui suit, nous allons utiliser les classes à droite mais aurions pu tout aussi bien utiliser celles à gauche.

**Proposition 2.1.39.** Soit  $G$  un groupe et  $H \subset G$  un sous-groupe. Pour tout  $g \in G$ , la classe d'équivalence  $\bar{g}$  de  $g$  pour  $\sim_H$  est  $H.g := \{h.g \mid h \in H\}$ . En particulier,  $\bar{e} = H$ .

*Démonstration.* On travaille par double inclusion. Si  $g' \in \bar{g}$ , c'est que  $g' \sim_H g$  et donc que  $h := g'.g^{-1} \in H$ . Mais alors  $g' = h.g \in H.g$ . Réciproquement, si  $g' = h.g$  avec  $h \in H$ , alors  $g'.g^{-1} \in H$  et  $g' \in \bar{g}$ . □

Dans le cas d'un groupe  $G$  fini, cette relation d'équivalence a une très forte conséquence sur les cardinaux possibles de ses sous-groupes.

**Définition 2.1.40.** Pour tout groupe  $G$  et tout sous-groupe  $H \subset G$ , on appelle *indice* de  $H$  dans  $G$  le cardinal, noté  $[G : H]$ , de l'ensemble des classes d'équivalence pour la relation  $\sim_H$ .

**Théorème 2.1.41** (Lagrange). Si  $G$  est un groupe fini et  $H \subset G$  un sous-groupe, alors  $|G| = [G : H].|H|$ . En particulier,  $|H|$  divise  $|G|$ .

*Démonstration.* La relation d'équivalence  $\sim_H$  donne une partition de  $G$  en ses classes d'équivalence. Or d'après la proposition précédente, l'application  $(h \mapsto h.g)$  donne une bijection entre  $H$  et  $\bar{g}$  pour tout  $g \in G$ . On en déduit que chaque classe possède  $|H|$  éléments et on a donc  $|G| = [G : H].|H|$ . □

*Exemples 2.1.42.*

- Les sous-groupes de  $\mathbb{Z}/4\mathbb{Z}$  ne peuvent avoir que 1, 2 ou 4 éléments, mais pas 3. Or un sous-groupe de cardinal 1 ou 4 ne peut être, respectivement, que le sous-groupe trivial ou le groupe tout entier. On en déduit que les seuls sous-groupes non évidents ne peuvent être que de cardinal 2. Un tel sous-groupe contient de fait  $\bar{0}$  et un autre élément  $\bar{k}$ ; on vérifie rapidement que seul  $\{\bar{0}, \bar{2}\}$  est effectivement un sous-groupe.
- Un groupe  $G$  tel que  $|G|$  soit un nombre premier ne possède que deux sous-groupes,  $\{e\}$  et  $G$ .
- Pour tout élément  $g$  d'un groupe fini  $G$ , on a  $|g| = |\langle g \rangle|$  qui divise  $|G|$ .

Maintenant que l'on a une relation d'équivalence, il est tentant de considérer l'espace quotient. Puisque clairement,  $g \sim_H e \Leftrightarrow g \in H$ , cela revient à vouloir "annuler" tous les éléments de  $H$ . La question est de savoir s'il est possible de le faire tout en préservant une structure de groupe. La réponse est en général non.

*Exemple 2.1.43.* Considérons  $\mathrm{SL}_2(\mathbb{Z}) \subset \mathrm{GL}_2(\mathbb{R})$ . Il s'agit bien d'un sous-groupe car  $\mathrm{SL}_2(\mathbb{Z})$  est clairement stable par produit, et la formule

$$\begin{pmatrix} a & b \\ c & d \end{pmatrix}^{-1} = \frac{1}{ad-bc} \begin{pmatrix} d & -b \\ -c & a \end{pmatrix}$$

montre que c'est également stable par prise d'inverse. Supposons maintenant que la multiplication dans  $\mathrm{GL}_2(\mathbb{R})$  soit compatible avec  $\sim_{\mathrm{SL}_2(\mathbb{Z})}$  (que l'on notera dans suite  $\sim$ ) et notons

$$A := \begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix} \text{ et } B := \begin{pmatrix} 1 & 0 \\ 0 & 2 \end{pmatrix}.$$

Puisque  $A \in \mathrm{SL}_2(\mathbb{Z})$ , on a  $A \sim \mathrm{Id}$  et l'on devrait donc avoir  $B.A \sim B$ , c'est-à-dire  $B.A.B^{-1} \in \mathrm{SL}_2(\mathbb{Z})$ .

$$\text{Or } B.A.B^{-1} = \begin{pmatrix} 1 & 0 \\ 0 & 2 \end{pmatrix} \cdot \begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix} \cdot \begin{pmatrix} 1 & 0 \\ 0 & \frac{1}{2} \end{pmatrix} = \begin{pmatrix} 1 & \frac{1}{2} \\ 0 & 1 \end{pmatrix} \notin \mathrm{SL}_2(\mathbb{Z}).$$

En généralisant cet exemple, on voit que, pour que la loi de composition interne soit compatible avec l'espace quotient, il faut que, pour tout  $h \in H$  et tout  $g \in G$ , on ait  $g.h \sim_H g$ , c'est-à-dire  $g.h.g^{-1} \in H$ . Cela motive la définition suivante.

**Définition 2.1.44.** Soit  $G$  un groupe et  $H \subset G$  un sous-groupe. On dit que  $H$  est *distingué*, s'il est stable par conjugaison, c'est-à-dire si pour tout  $h \in H$  et tout  $g \in G$ , on a  $g.h.g^{-1} \in H$ . On note alors  $H \triangleleft G$ .

*Remarque 2.1.45.* Nous avons évoqué, dans la remarque 2.1.38, l'existence de classes à droite et à gauche. Un sous-groupe  $H \subset G$  est en fait distingué ssi, pour tout élément de  $G$ , ses classes à droite et à gauche vis-à-vis de  $H$  sont identiques. Autrement dit, les sous-groupes distingués sont ceux pour lesquels la relation d'équivalence  $\sim_H$  est globalement bien définie.

*Exemples 2.1.46.*

- Pour tout groupe  $G$ ,  $\{e\}$  et  $G$  sont des sous-groupes distingués.
- Si  $G$  est un groupe abélien, alors tout sous-groupe  $H \subset G$  est distingué car, pour tout  $h \in H$  et tout  $g \in G$ , on a  $g.h.g^{-1} = h.g.g^{-1} = h \in H$ .
- D'après ce qui précède,  $\mathrm{SL}_2(\mathbb{Z})$  n'est pas un sous-groupe distingué de  $\mathrm{GL}_2(\mathbb{R})$ , mais  $\mathrm{SL}_2(\mathbb{R})$ , lui, l'est d'après la proposition qui suit.

**Proposition 2.1.47.** Pour tout un morphisme de groupes  $f : G_1 \rightarrow G_2$ , le noyau  $\mathrm{Ker}(f) \subset G_1$  est un sous-groupe distingué.

*Démonstration.* Soit  $h \in \mathrm{Ker}(f)$  et  $g \in G_1$ , on a  $f(g.h.g^{-1}) = f(g).f(h).f(g^{-1}) = f(g).f(g)^{-1} = e$ .  $\square$

*Remarque 2.1.48.* Au contraire du noyau, l'image d'un morphisme de groupe n'est pas forcément distingué, autrement tout sous-groupe serait distingué puisque pour  $H \subset G$ , on peut toujours considérer l'inclusion

$$\begin{array}{ccc} H & \longrightarrow & G \\ \iota: & & \\ h & \longmapsto & h \end{array}$$

qui est évidemment un morphisme de groupes et dont l'image est  $H$ . Or nous avons déjà vu un exemple de sous-groupe non distingué.

**Proposition 2.1.49.** Soit  $G$  un groupe. Pour tout un sous-groupe distingué  $H \triangleleft G$ , la loi de composition interne de  $G$  est compatible avec  $\sim_H$  dans le sens où, si  $g_1, g_2, g'_1, g'_2 \in G$  vérifient  $g_1 \sim_H g'_1$  et  $g_2 \sim_H g'_2$ , alors  $g_1 \cdot g_2 \sim_H g'_1 \cdot g'_2$ .

*Démonstration.* Il suffit de remarquer que

$$(g_1 \cdot g_2) \cdot (g'_1 \cdot g'_2)^{-1} = g_1 \cdot g_2 \cdot g'_2{}^{-1} \cdot g'_1{}^{-1} = g_1 \cdot g'_1{}^{-1} \cdot g'_1 \cdot g_2 \cdot g'_2{}^{-1} \cdot g'_1{}^{-1} = (g_1 \cdot g'_1{}^{-1}) \cdot (g'_1 \cdot (g_2 \cdot g'_2{}^{-1}) \cdot g'_1{}^{-1})$$

car  $g_2 \cdot g'_2{}^{-1} \in H$ , puisque  $g_2 \sim_H g'_2$ , et donc  $g'_1 \cdot (g_2 \cdot g'_2{}^{-1}) \cdot g'_1{}^{-1} \in H$  puisque  $H$  est distingué; et par ailleurs  $g_1 \cdot g'_1{}^{-1} \in H$  puisque  $g_1 \sim_H g'_1$ .  $\square$

**Définition 2.1.50.** Pour tout sous-groupe  $H \triangleleft G$  d'un groupe  $G$ , on définit le *groupe quotient*  $G/H$  comme l'espace quotient pour la relation d'équivalence  $\sim_H$ , muni de la loi de composition interne induite par celle de  $G$  via la formule  $\bar{g}_1 *_{G/H} \bar{g}_2 := \overline{g_1 *_G g_2}$  où  $\bar{g}$  dénote la classe d'équivalence de  $g \in G$ .

**Proposition 2.1.51.** Pour tout sous-groupe  $H \triangleleft G$  d'un groupe  $G$ ,  $G/H$  est un groupe, et la surjection naturelle

$$\begin{array}{ccc} G & \longrightarrow & G/H \\ \pi_H : g & \longmapsto & \bar{g} \end{array}$$

est un morphisme de groupe.

*Démonstration.* Toutes les vérifications sont immédiates. La loi est associative car, pour tout  $g_1, g_2, g_3 \in G$ , on a  $(\bar{g}_1 \cdot \bar{g}_2) \cdot \bar{g}_3 = \overline{(g_1 \cdot g_2) \cdot g_3} = \overline{g_1 \cdot (g_2 \cdot g_3)} = \bar{g}_1 \cdot (\bar{g}_2 \cdot \bar{g}_3)$ . Elle admet  $\bar{e}$  comme élément neutre puisque, pour tout  $g \in G$ ,  $\bar{e} \cdot \bar{g} = \overline{e \cdot g} = \bar{g} = \overline{g \cdot e} = \bar{g} \cdot \bar{e}$ . Et tout élément  $\bar{g} \in G/H$  admet un élément neutre, à savoir  $\bar{g}^{-1}$ , puisque  $\bar{g}^{-1} \cdot \bar{g} = \overline{g^{-1} \cdot g} = \bar{e} = \overline{g \cdot g^{-1}} = \bar{g} \cdot \bar{g}^{-1}$ .

Enfin, l'application  $\pi_H$  est un morphisme de groupe car, pour tout  $g_1, g_2 \in G$ ,  $\pi_H(g_1 \cdot g_2) = \overline{g_1 \cdot g_2} = \bar{g}_1 \cdot \bar{g}_2 = \pi_H(g_1) \cdot \pi_H(g_2)$ .  $\square$

*Exemples 2.1.52.*

- Pour tout groupe  $G$ ,  $G$  et  $\{e\}$  sont distingués. La relation  $\sim_G$  identifie tout le monde, il n'y a donc qu'une seule classe d'équivalence et  $G/G = \{\bar{e}\}$ . Deux éléments ne sont par contre équivalents pour  $\sim_{\{e\}}$  que s'ils sont égaux, il y a donc autant de classe d'équivalence que d'élément dans  $G$  et  $G/\{e\} = G$ .
- Pour tout  $n \in \mathbb{N}^*$ ,  $(\mathbb{Z}/n\mathbb{Z}, +)$  est le groupe quotient de  $\mathbb{Z}$  par son sous-groupe  $n\mathbb{Z}$ . Ce dernier est bien distingué car  $\mathbb{Z}$  est abélien.

Le théorème suivant est sans doute le plus important en théorie des groupes.

**Théorème 2.1.53** (premier théorème d'isomorphisme). Tout morphisme de groupes  $f : G_1 \rightarrow G_2$  induit un isomorphisme  $\bar{f} : G_1/\text{Ker}(f) \rightarrow \text{Im}(f)$  défini, pour tout  $g \in G_1$ , par  $\bar{f}(\bar{g}) = f(g)$ .

*Démonstration.* Montrons déjà que  $\bar{f}$  est bien défini, c'est-à-dire que  $\bar{f}(\bar{g})$  ne dépend pas du choix du représentant  $g \in G_1$ . Considérons donc  $g'$  un autre représentant de la classe de  $g$ . On a alors  $g' \cdot g^{-1} \in \text{Ker}(f)$  et donc  $e = f(g' \cdot g^{-1}) = f(g') \cdot f(g)^{-1}$ , autrement dit  $f(g') = f(g)$ . De plus, par définition de  $\text{Im}(f)$ ,  $\bar{f}$  va bien dans  $\text{Im}(f)$ , et elle est même surjective. Il ne reste donc plus qu'à montrer que  $\bar{f}$  est injective. Mais si  $\bar{g} \in \text{Ker}(\bar{f})$ , c'est que  $f(g) = e_{G_2}$ , donc que  $g \in \text{Ker}(f)$ , et donc que  $\bar{g} = \bar{e}_{G_1}$ .  $\square$

*Exemples 2.1.54.* Reprenons les situations des exemples 2.1.52.

- On a  $\text{Ker}(\text{Id}_G) = \{e\}$  et  $\text{Im}(\text{Id}_G) = G$ , on en déduit donc que  $\bar{\text{Id}}_G : G/\{e\} \rightarrow G$  est un isomorphisme.

- On a  $\text{Ker}(\text{Triv}_G) = G$  et  $\text{Im}(\text{Triv}_G) = \{e\}$ , on en déduit donc que  $\overline{\text{Triv}}_G : G/G \rightarrow \{e\}$  est un isomorphisme.
- Pour  $n \in \mathbb{N}^*$ , considérons l'application

$$\exp_n : \begin{array}{ccc} \mathbb{Z} & \longrightarrow & \mathbb{C}^* \\ k & \longmapsto & e^{\frac{2ik\pi}{n}} \end{array}.$$

On a  $\text{Ker}(\exp_n) = n\mathbb{Z}$  et  $\text{Im}(\exp_n) = \mathbb{U}_n$ , l'ensemble des racines  $n^{\text{ièmes}}$  de l'unité. On en déduit que  $\overline{\exp}_n : \mathbb{Z}/n\mathbb{Z} \rightarrow \mathbb{U}_n$  est un isomorphisme.

## 2.2 Un exemple important : les groupes de permutation

### 2.2.1 Définition et généralités

Nous allons maintenant étudier plus en détail une famille importante de groupes finis.

**Définition 2.2.1.** Pour tout ensemble  $X$ , on note  $\mathfrak{S}(X)$  l'ensemble des bijections de  $X$  dans lui-même ; et pour tout  $n \in \mathbb{N}^*$ , on note  $\mathfrak{S}_n := \mathfrak{S}([1, n])$ . On appelle *permutations* leurs éléments.

**Notation 2.2.2.** Lorsque  $X := \{x_1, \dots, x_n\}$  est un ensemble fini, on peut représenter tout élément  $\sigma \in \mathfrak{S}(X)$  par le double vecteur suivant :

$$\begin{pmatrix} x_1 & x_2 & \cdots & x_n \\ \sigma(x_1) & \sigma(x_2) & \cdots & \sigma(x_n) \end{pmatrix}.$$

Ainsi

$$\begin{pmatrix} a & b & c \\ c & b & a \end{pmatrix} \quad \text{et} \quad \begin{pmatrix} 1 & 2 & 3 & 4 \\ 2 & 1 & 4 & 3 \end{pmatrix}$$

représentent respectivement les bijections

$$\begin{array}{ccc} \{a, b, c\} & \rightarrow & \{a, b, c\} \\ a & \mapsto & c \\ b & \mapsto & b \\ c & \mapsto & a \end{array} \quad \text{et} \quad \begin{array}{ccc} \{1, 2, 3, 4\} & \rightarrow & \{1, 2, 3, 4\} \\ 1 & \mapsto & 2 \\ 2 & \mapsto & 1 \\ 3 & \mapsto & 4 \\ 4 & \mapsto & 3 \end{array}.$$

S'il n'y a pas d'ambiguïté sur l'ensemble sous-jacent  $X$ , on pourra omettre les colonnes des éléments laissés fixes. Le premier exemple s'écrira alors  $\begin{pmatrix} a & c \\ c & a \end{pmatrix}$ .

**Proposition 2.2.3.** Pour tout ensemble non vide  $X$ ,  $(\mathfrak{S}(X), \circ)$  est un groupe, appelé *groupe des permutations de  $X$*  ou *groupe symétrique* si  $X = [1, n]$ . Ce dernier est non abélien si et seulement si  $|X| \geq 3$ .

*Démonstration.* Nous avons déjà vu qu'il s'agissait d'un groupe. Si  $|X| \geq 3$ , alors il existe trois éléments distincts  $a, b, c \in X$  et on peut vérifier directement que

$$\begin{pmatrix} a & b \\ b & a \end{pmatrix} \circ \begin{pmatrix} b & c \\ c & b \end{pmatrix} = \begin{pmatrix} a & b & c \\ b & c & a \end{pmatrix}$$

tandis que

$$\begin{pmatrix} b & c \\ c & b \end{pmatrix} \circ \begin{pmatrix} a & b \\ b & a \end{pmatrix} = \begin{pmatrix} a & b & c \\ c & a & b \end{pmatrix}.$$

Par contre, si  $|X| = 1$ ,  $\mathfrak{S}(X)$  est le groupe trivial ; et si  $|X| = 2$ ,  $\mathfrak{S}(X)$  ne contient que deux éléments, l'identité et la bijection qui échange les deux éléments. Dans les deux cas,  $\mathfrak{S}(X)$  est abélien.  $\square$

**Proposition 2.2.4.** Si  $\varphi : X_1 \rightarrow X_2$  est une bijection, alors l'application

$$F_\varphi : \begin{array}{ccc} \mathfrak{S}(X_1) & \longrightarrow & \mathfrak{S}(X_2) \\ f & \longmapsto & \varphi^{-1} \circ f \circ \varphi \end{array}$$

est un isomorphisme de groupes. En particulier, si  $X$  est un ensemble fini, alors  $\mathfrak{S}(X)$  est isomorphe à  $\mathfrak{S}_{|X|}$ .

*Démonstration.* Pour  $f_1, f_2 \in \mathfrak{S}(X_1)$ , on a

$$F_\varphi(f_1 \circ f_2) = \varphi^{-1} \circ f_1 \circ f_2 \circ \varphi = \varphi^{-1} \circ f_1 \circ \varphi \circ \varphi^{-1} \circ f_2 \circ \varphi = F_\varphi(f_1) \circ F_\varphi(f_2),$$

$F_\varphi$  est donc un morphisme de groupes qui est clairement bijectif, car d'inverse  $F_{\varphi^{-1}}$ .

Pour montrer la seconde affirmation, il suffit de considérer une bijection qui numérote les éléments de  $X$  de 1 à  $|X|$ .  $\square$

Par la suite, nous nous concentrerons donc sur l'étude de  $\mathfrak{S}_n$ . Ces derniers sont des groupes finis :

**Proposition 2.2.5.** Pour tout  $n \in \mathbb{N}^*$ ,  $\mathfrak{S}_n$  est un groupe fini de cardinal  $n!$ .

*Démonstration.* On montre le résultat par récurrence sur  $n \in \mathbb{N}^*$ . Pour  $n = 1$ ,  $\mathfrak{S}_n$  ne contient qu'un seul élément, à savoir l'identité. Supposons maintenant le résultat vrai au rang  $n - 1$  et considérons l'application

$$\Gamma_n : \begin{array}{ccc} \mathfrak{S}_n & \longrightarrow & \mathfrak{S}_{n-1} \\ \left( \begin{array}{cccccccc} 1 & \cdots & i-1 & i & i+1 & \cdots & n-1 & n \\ k_1 & \cdots & k_{i-1} & n & k_{i+1} & \cdots & k_{n-1} & k_n \end{array} \right) & \longmapsto & \left( \begin{array}{cccccccc} 1 & \cdots & i-1 & i & i+1 & \cdots & n-2 & n-1 \\ k_1 & \cdots & k_{i-1} & k_{i+1} & k_{i+1} & \cdots & k_{n-1} & k_n \end{array} \right) \end{array}$$

qui enlève les deux  $n$ , où qu'il soit sur la seconde ligne. Chaque élément de  $\mathfrak{S}_{n-1}$  a exactement  $n$  antécédents, obtenus en rajoutant  $n$  au bout de la première ligne, et en insérant  $n$  quelque part sur la seconde ligne. On en déduit que  $|\mathfrak{S}_n| = n \cdot |\mathfrak{S}_{n-1}|$ , ce qui par hypothèse de récurrence, donne  $|\mathfrak{S}_n| = n \cdot (n-1)! = n!$ .  $\square$

Mais parmi les groupes finis, ils possèdent même un certain caractère universel :

**Théorème 2.2.6** (Cayley). Tout groupe fini  $G$  est isomorphe à un sous-groupe de  $\mathfrak{S}_{|G|}$ .

*Démonstration.* Il suffit de considérer l'application

$$\xi_G : \begin{array}{ccc} G & \longrightarrow & \mathfrak{S}(G) \\ g & \longmapsto & (h \mapsto g.h) \end{array}.$$

Pour tout  $g \in G$ ,  $\xi_G(g)$  est en effet une bijection, d'inverse  $\xi_G(g^{-1})$ ; et on vérifie facilement que, pour tout  $g_1, g_2 \in G$ , on a  $\xi_G(g_1) \circ \xi_G(g_2) = \xi_G(g_1 \cdot g_2)$ . L'application  $\xi_G$  est donc un morphisme de groupe dont le noyau est trivial. En effet, si  $g \in \text{Ker}(\xi_G)$ , alors  $g \cdot e = (\xi_G(g))(e) = \text{Id}_G(e) = e$  et donc  $g = e$ . D'après le premier théorème d'isomorphisme, on en déduit que  $G$  est isomorphe à  $\text{Im}(\xi_G)$  qui est un sous-groupe de  $\mathfrak{S}(G)$ . Pour obtenir le résultat avec  $\mathfrak{S}_{|G|}$ , il suffit de composer  $\xi_G$  avec un isomorphisme entre  $\mathfrak{S}(G)$  et  $\mathfrak{S}_{|G|}$ .  $\square$



### 2.2.2 Description des éléments

Fixons maintenant  $n \in \mathbb{N}^*$  et étudions plus en détail les éléments de  $\mathfrak{S}_n$ .

#### Définition 2.2.7.

- On appelle *support* d'une permutation  $\sigma \in \mathfrak{S}_n$  l'ensemble  $\text{supp}(\sigma)$  des éléments  $k \in \llbracket 1, n \rrbracket$  tels que  $\sigma(k) \neq k$ .
- On dit qu'une permutation  $\sigma \in \mathfrak{S}_n$  est un *cycle* ou une *permutation circulaire* si  $\text{supp}(\sigma) \neq \emptyset$  et, pour tous  $k_1, k_2 \in \text{supp}(\sigma)$ , il existe  $i \in \mathbb{N}$  tel que  $\sigma^i(k_1) = k_2$ . On appelle alors *longueur* du cycle le cardinal de son support.

Exemples 2.2.8.

- La permutation  $\sigma_1 := \begin{pmatrix} 1 & 2 & 3 \\ 2 & 3 & 1 \end{pmatrix}$  est un cycle car  $\text{supp}(\sigma_1) = \{1, 2, 3\}$  et, en itérant  $\sigma_1$ , on a

$$1 \xrightarrow{\sigma_1} 2 \xrightarrow{\sigma_1} 3 \xrightarrow{\sigma_1} 1.$$

- La permutation  $\sigma_2 := \begin{pmatrix} 1 & 2 & 3 & 4 & 5 \\ 3 & 2 & 5 & 1 & 4 \end{pmatrix}$  l'est aussi car  $\text{supp}(\sigma_2) = \{1, 3, 4, 5\}$  et, en itérant  $\sigma_2$ , on a

$$1 \xrightarrow{\sigma_2} 3 \xrightarrow{\sigma_2} 5 \xrightarrow{\sigma_2} 4 \xrightarrow{\sigma_2} 1.$$

- La permutation  $\sigma_3 := \begin{pmatrix} 1 & 2 & 3 & 4 \\ 2 & 1 & 4 & 3 \end{pmatrix}$  n'est, par contre, pas un cycle car  $\text{supp}(\sigma_3) = \{1, 2, 3, 4\}$  et, en itérant  $\sigma_3$ , on obtient

$$1 \xrightarrow{\sigma_3} 2 \xrightarrow{\sigma_3} 1 \quad \text{et} \quad 3 \xrightarrow{\sigma_3} 4 \xrightarrow{\sigma_3} 3 ;$$

impossible donc d'obtenir, par exemple, 4 en partant de 2 ou 1 en partant de 4. On peut toutefois remarquer que  $\sigma_3$  est le produit de deux cycles,  $\begin{pmatrix} 1 & 2 \\ 2 & 1 \end{pmatrix}$  et  $\begin{pmatrix} 3 & 4 \\ 4 & 3 \end{pmatrix}$ , lesquels ont des supports disjoints ; ceci est un fait général !

**Proposition 2.2.9.** Toute permutation est un produit, éventuellement vide, de cycles à supports disjoints.

*Démonstration.* Montrons par récurrence généralisée sur le cardinal  $s \in \mathbb{N}$  du support que toute permutation  $\sigma$  possède une décomposition en produit de cycles dont les supports sont disjoints et contenus dans le support de  $\sigma$ . Seule l'identité a un support vide, et le produit est alors le produit vide. Supposons maintenant le résultat vrai pour les permutations dont le support contient au plus  $s$  éléments et considérons une permutation  $\sigma \in \mathfrak{S}_n$  avec  $|\text{supp}(\sigma)| = s + 1$ . On a alors  $s + 1 \geq 1$ ,  $\text{supp}(\sigma)$  n'est donc pas vide et on peut fixer un élément  $k_0 \in \text{supp}(\sigma)$ , par exemple son plus petit élément. Considérons alors la suite  $(k_i)_{i \in \mathbb{N}} := (\sigma^i(k_0))_{i \in \mathbb{N}}$  des images itérées de  $k_0$  par  $\sigma$ . Cette suite est périodique. En effet,  $\mathbb{N}$  étant infini mais pas  $\llbracket 1, n \rrbracket$ , il existe forcément  $i_1 < i_2 \in \mathbb{N}$  tels que  $k_{i_1} = k_{i_2}$ , c'est-à-dire tels que  $\sigma^{i_1}(k_0) = \sigma^{i_2}(k_0)$ . Mais puisque  $\sigma$  est bijective, on obtient alors que  $k_{i_2 - i_1} = \sigma^{i_2 - i_1}(k_0) = k_0$  en composant par  $\sigma^{-i_1}$  ; on en déduit que  $\{i \in \mathbb{N} \mid k_i = k_0\}$  n'est pas vide et qu'on peut en considérer  $p$ , son plus petit élément. Alors, en composant par  $\sigma^j$  pour tout  $j \in \mathbb{N}$ , on obtient que  $\sigma^{j+p}(k_1) = \sigma^j(k_1)$ , c'est-à-dire que  $k_{j+p} = k_j$ . La suite est donc  $p$ -périodique et, par minimalité de  $p$ , les éléments  $k_0, k_1, \dots, k_{p-1}$  sont tous distincts. On définit alors la permutation  $c_1 := \begin{pmatrix} k_0 & k_1 & \cdots & k_{p-1} \\ k_1 & k_2 & \cdots & k_p \end{pmatrix}$  et on pose  $\sigma' = c_1^{-1} \circ \sigma$  ; on vérifie aisément que  $c_1$  est un cycle et que  $\text{supp}(\sigma') = \text{supp}(\sigma) \setminus \{k_0, \dots, k_{p-1}\}$ . En particulier,  $|\text{supp}(\sigma')| < |\text{supp}(\sigma)|$ , donc  $|\text{supp}(\sigma')| \leq s$  ; par hypothèse de récurrence, il existe donc des cycles  $c_2, \dots, c_r$  à supports disjoints, tous contenus dans  $\text{supp}(\sigma')$  tels que  $c_1^{-1} \circ \sigma = \sigma' = c_2 \circ \cdots \circ c_r$ . On a alors  $\sigma = c_1 \circ c_2 \circ \cdots \circ c_r$  avec toutes les conditions voulues.  $\square$

*Remarque 2.2.10.* Deux cycles à supports disjoints commutent, l'ordre des cycles n'a donc pas d'importance. Hormis cet ordre, la décomposition en produit de cycles à supports disjoints d'une permutation  $\sigma$  donnée est unique. Les supports des cycles sont en effet déterminés comme les classes d'équivalence de la relation d'équivalence définie par  $k_1 \sim k_2 \Leftrightarrow \exists i \in \mathbb{Z}, k_2 = \sigma^i(k_1)$ , et chaque cycle  $c$  est alors défini, pour tout  $k \in \llbracket 1, n \rrbracket$ , par

$$c(k) = \begin{cases} \sigma(k) & \text{si } k \in \text{supp}(c) \\ k & \text{sinon} \end{cases}.$$

**Notation 2.2.11.** En reprenant les idées de la preuve de la proposition 2.2.9, on observe que, pour tout cycle  $c$ , on peut choisir un élément  $k_0$  dans son support et décrire entièrement  $c$  par la suite

$$k_0 \xrightarrow{c} k_1 \xrightarrow{c} \cdots \xrightarrow{c} k_{p-1}$$

des images itérées de  $k_0$  par  $c$ . On peut donc, de manière efficiente, noter  $c$  par  $(k_0 k_1 \cdots k_{p-1})$ . Puisque toute permutation (non triviale)  $\sigma$  s'écrit comme produit de tels cycles, on peut alors noter  $\sigma$  en concaténant les cycles qui la compose en omettant le signe  $\circ$ .

*Exemples 2.2.12.* En reprenant les notations des exemples 2.2.8, on a  $\sigma_1 = (123)$ ,  $\sigma_2 = (1354)$  et  $\sigma_3 = (12)(34)$ .

*Remarque 2.2.13.* La notation ci-dessus pour un cycle dépend du choix d'un élément  $k_0$ , mais choisir un autre élément aura simplement l'effet de permuter cycliquement les éléments. Par exemple  $(123)$ ,  $(231)$  et  $(312)$  représentent tous le même cycle.

Parmi les permutations, les cycles sont des éléments simples, mais suffisamment généraux pour engendrer tout  $\mathfrak{S}_n$ . On peut encore optimiser cela. Pour cela, regardons les cycles qui fixent un maximum d'éléments. A l'évidence, il n'existe pas de permutation dont le support est de cardinal 1, le plus petit cardinal est donc 2.

**Définition 2.2.14.** On appelle *transposition* tout cycle de longueur 2, c'est-à-dire toute permutation  $(k_1 k_2)$  qui ne fait qu'échanger deux éléments  $k_1 \neq k_2 \in \llbracket 1, n \rrbracket$ . Et on dit que la transposition est *élémentaire* si  $k_1$  et  $k_2$  sont deux entiers consécutifs.

**Proposition 2.2.15.** Tout élément de  $\mathfrak{S}_n$  est un produit de transpositions élémentaires, c'est-à-dire

$$\mathfrak{S}_n = \langle \{ (i(i+1)) \mid 1 \leq i < n \} \rangle.$$

*Démonstration.* D'après la proposition 2.2.9,  $\mathfrak{S}_n$  est déjà engendré par les cycles, il suffit donc de montrer que tout cycle s'écrit comme produit de transpositions élémentaires. Mais commençons par montrer qu'ils sont déjà des produits de transpositions par récurrence sur la longueur  $s$  du cycle. Comme remarqué précédemment, la plus petite valeur pour  $s$  est 2, le cycle est alors lui-même une transposition. Supposons maintenant le résultat vrai pour un cycle de longueur  $s-1$  et considérons un cycle  $c = (k_1 \cdots k_s)$  de longueur  $s$ . Un calcul direct montre alors que  $c' := c \circ (k_{s-1} k_s) = (k_1 \cdots k_{s-1})$ . Par hypothèse de récurrence, il existe donc des transpositions  $\tau_1, \dots, \tau_r$  tels que  $c' = \tau_1 \circ \cdots \circ \tau_r$  et on en déduit que  $c = \tau_1 \circ \cdots \circ \tau_r \circ (k_{s-1} k_s)$  après avoir observé qu'une transposition est toujours son propre inverse.

Pour conclure, il ne reste plus qu'à remarquer que, pour tout  $1 \leq k_1 < k_2 \leq n$ , on a

$$(k_1 k_2) = (k_1(k_1+1)) \circ ((k_1+1)(k_1+2)) \circ \cdots \circ ((k_2-2)(k_2-1)) \circ ((k_2-1)k_2) \circ \\ \circ ((k_2-2)(k_2-1)) \circ \cdots \circ ((k_1+1)(k_1+2)) \circ (k_1(k_1+1)).$$

□

*Remarque 2.2.16.* La proposition 2.2.15 est conforme à l'intuition. Permuter les entiers entre 1 et  $n$ , cela correspond à les réarranger, et pour réaliser "à la main" ce réarrangement, nous procédons naturellement à une suite de transpositions. Plus précisément, alors que nous déplaçons un entier de son ancienne à sa nouvelle place, celui-ci échange successivement sa place avec son prédécesseur ou son successeur.

### 2.2.3 Signature

On fixe ici  $n \in \mathbb{N}^* \setminus \{1\}$ .

Nous venons de montrer que les groupes de permutation sont engendrés par les transpositions. Utilisons ce fait pour montrer le résultat suivant.

**Lemme 2.2.17.** Si  $G$  est un groupe abélien, alors tout morphisme de groupes  $f : \mathfrak{S}_n \rightarrow G$  est entièrement déterminé par  $f((12))$ .

*Démonstration.* On commence par remarquer que la valeur de  $f$  en  $(12)$  impose la valeur de  $f$  sur toute transposition. Considérons pour cela  $(k_1 k_2)$ , avec  $1 \leq k_1 < k_2 \leq n$ , et montrons qu'il existe  $\sigma \in \mathfrak{S}_n$  telle que  $(k_1 k_2) = \sigma^{-1} \circ (12) \circ \sigma$  :

- si  $k_1 > 2$ , alors  $\{k_1, k_2\} \cap \{1, 2\} = \emptyset$  et il suffit de prendre  $\sigma = (1k_1)(2k_2)$  ;
- si  $k_1 = 2$ , alors  $k_2 \notin \{1, 2\}$  et on peut prendre  $\sigma = (1k_2)$  ;
- si  $k_1 = 1$  et  $k_2 > 2$ , alors on prend  $\sigma = (2k_2)$  ;
- si  $k_1 = 1$  et  $k_2 = 2$ , alors on peut prendre  $\sigma = \text{Id}$ .

Dès lors, on a

$$f((k_1 k_2)) = f(\sigma^{-1} \circ (12) \circ \sigma) = f(\sigma)^{-1} \cdot f((12)) \cdot f(\sigma) = f(\sigma)^{-1} \cdot f(\sigma) \cdot f((12)) = f((12))$$

puisque  $G$  est abélien.

Pour une permutation  $\sigma$  quelconque, il ne reste plus qu'à l'exprimer comme produit de transpositions  $\tau_1 \circ \dots \circ \tau_r$  à l'aide de la proposition 2.2.15. On obtient alors  $f(\sigma) = f(\tau_1) \cdot \dots \cdot f(\tau_r) = f((12))^r$ .  $\square$

**Corollaire 2.2.18.** Il existe un unique épimorphisme de groupe allant de  $\mathfrak{S}_n$  à  $\mathbb{Z}/2\mathbb{Z}$ .

*Démonstration.* D'après la proposition précédente, il ne peut exister que deux morphismes de groupes entre  $\mathfrak{S}_n$  et  $\mathbb{Z}/2\mathbb{Z}$ , distingués par le fait d'envoyer  $(12)$  sur  $\bar{0}$  ou  $\bar{1}$ . Dans le premier cas, le morphisme est l'application trivial envoyant tout le monde sur  $\bar{0}$  ; seul le second cas est donc susceptible de produire un épimorphisme. Il ne reste plus qu'à montrer qu'un tel morphisme existe. Or le détail de la preuve du lemme 2.2.17 donne même plus : si un tel morphisme existe, on sait qu'il doit envoyer nécessairement  $\sigma = \tau_1 \circ \dots \circ \tau_r$  sur  $\underbrace{\bar{1} + \dots + \bar{1}}_{r \text{ fois}} = \bar{r}$ , c'est-à-dire sur la parité du nombre de transpositions nécessaire pour

décrire  $\sigma$ . Montrons que cette quantité est bien définie, c'est-à-dire que si  $\tau_1 \circ \dots \circ \tau_r = \sigma = \tau'_1 \circ \dots \circ \tau'_{r'}$ , alors  $r' \equiv r \pmod{2}$ . Pour ce faire, nous allons montrer que  $r$  a la même parité que le nombre d'inversions de  $\sigma$ , c'est-à-dire le nombre de couples  $(i, j) \in \llbracket 1, n \rrbracket^2$  tels que  $i < j$  mais  $\sigma(i) > \sigma(j)$  ; cette dernière quantité ne dépendant pas de la décomposition, cela montrera le résultat. Procédons par récurrence sur le nombre  $s$  de transpositions élémentaires nécessaire pour décrire  $\sigma$ . Si  $s = 0$ , alors  $\sigma = \text{Id}$  et le résultat est vrai. Supposons le résultat vrai pour un produit de  $s - 1$  transpositions élémentaires et considérons  $\sigma = \tau_1 \circ \dots \circ \tau_s$ , où les  $\tau_i$  sont des transpositions élémentaires. On pose  $\sigma' = \tau_1 \circ \dots \circ \tau_{s-1}$  et on écrit  $\tau_s = (i_0(i_0 + 1))$  avec  $i_0 \in \llbracket 1, n - 1 \rrbracket$ . On a donc  $\sigma(i_0) = \sigma'(i_0 + 1)$ ,  $\sigma(i_0 + 1) = \sigma'(i_0)$  et  $\sigma(i) = \sigma'(i)$  pour tout  $i \in \llbracket 1, n \rrbracket \setminus \{i_0, i_0 + 1\}$ . On observe alors que

- si  $(i, j)$ , avec  $\{i, j\} \cap \{i_0, i_0 + 1\} = \emptyset$ , est (resp. n'est pas) une inversion pour  $\sigma'$ , alors il l'est aussi (resp. ne l'est pas non plus) pour  $\sigma$  ;
- si  $(i, i_0)$  est (resp. n'est pas) une inversion pour  $\sigma'$ , alors  $(i, i_0 + 1)$  l'est (resp. ne l'est pas) pour  $\sigma$  ;
- si  $(i, i_0 + 1)$ , avec  $i \neq i_0$ , est (resp. n'est pas) une inversion pour  $\sigma'$ , alors  $(i, i_0)$  l'est (resp. ne l'est pas) pour  $\sigma$  ;
- $(i_0, i_0 + 1)$  est une inversion pour  $\sigma'$  si et seulement si il ne l'est pas pour  $\sigma$  ;
- si  $(i_0, j)$ , avec  $j \neq i_0 + 1$ , est (resp. n'est pas) une inversion pour  $\sigma'$ , alors  $(i_0 + 1, j)$  l'est (resp. ne l'est pas) pour  $\sigma$  ;
- si  $(i_0 + 1, j)$  est (resp. n'est pas) une inversion pour  $\sigma'$ , alors  $(i_0, j)$  l'est (resp. ne l'est pas) pour  $\sigma$ .

On en déduit que  $\sigma$  a une inversion de plus ou de moins que  $\sigma'$ , à savoir  $(i_0, i_0 + 1)$ . Or, par hypothèse de récurrence, la parité du nombre d'inversions de  $\sigma'$  est congrue à  $(s - 1)$  modulo 2, on en déduit que celle de  $\sigma$  est congrue à  $s - 1 \pm 1 \equiv s$  modulo 2.  $\square$

**Définition 2.2.19.**

- On note  $\text{sign}$  l'unique épimorphisme de groupes allant de  $\mathfrak{S}_n$  vers  $\mathbb{Z}/2\mathbb{Z}$ ; et pour toute permutation  $\sigma \in \mathfrak{S}_n$ , on appelle *signature de  $\sigma$*  la quantité  $(-1)^{\text{sign}(\sigma)}$ .
- On appelle  $n^{\text{ième}}$  *groupe alterné* le noyau  $\mathcal{A}_n$  du morphisme  $\text{sign}$  sur  $\mathfrak{S}_n$ , c'est-à-dire l'ensemble des permutations de  $\mathfrak{S}_n$  de signature positive.

*Application 2.2.20* (jeu de taquin). Le jeu de taquin est composé de 15 petits carrés numérotés de 1 à 15 disposés sur les cases d'une grille carrée de côté 4. Une des cases de la grille est donc vide et chacun des carrés peut être déplacé horizontalement ou verticalement d'une case lorsqu'il jouxte la case vide (il vient donc se positionner sur la case anciennement vide, et son ancienne case devient vide) :



La position initiale est 

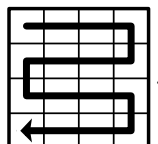
|    |    |    |    |
|----|----|----|----|
| 2  | 1  | 3  | 4  |
| 5  | 6  | 7  | 8  |
| 9  | 10 | 11 | 12 |
| 13 | 14 | 15 |    |

 et le but du jeu est d'arriver à la position 

|    |    |    |    |
|----|----|----|----|
| 1  | 2  | 3  | 4  |
| 5  | 6  | 7  | 8  |
| 9  | 10 | 11 | 12 |
| 13 | 14 | 15 |    |

. A la fin

du XIX<sup>ième</sup> siècle, Sam Loyd proposa une récompense de 1000\$ à celui qui résoudrait le problème. La récompense ne fut jamais réclamée car le problème est impossible. En effet, à toute position de la grille, on peut associer la permutation  $\sigma \in \mathfrak{S}_{15}$  qui envoie  $k$  sur le  $k^{\text{ième}}$  nombre lu en parcourant la grille comme suit, la case vide ne comptant pas :



Un déplacement de case horizontal ne change pas la permutation associée. Selon l'emplacement de la case vide, un déplacement vertical peut ne pas la changer ou bien la multiplier par un cycle de longueur 3, 5 ou 7, lesquels s'écrivent comme produit de 0, 2, 4 ou 6 transpositions. Dans tous les cas, la signature de la permutation n'est pas modifiée. Or la configuration de départ correspond à  $(12)(58)(67)(13\ 15)$ , dont la signature vaut 1, tandis que la position d'arrivée correspond à  $(58)(67)(13\ 15)$ , dont la signature vaut  $-1$ . Il est donc impossible de passer de l'une à l'autre en n'utilisant que les mouvements autorisés.

### 3 Anneaux

Inspirée de plusieurs exemples comme la somme sur les entiers, la multiplication des matrices inversibles ou la composition de bijections, la notion de groupe a permis d'extraire de ces exemples une substantifique moelle commune, permettant d'axiomatiser ces opérations et d'en abstraire certains calculs. La notion d'anneau, quant à elle, vise à généraliser les situations où, comme sur les nombres ou les matrices, deux opérations cohabitent sur un même ensemble.

#### 3.1 Anneaux et sous-anneaux

**Définition 3.1.1.** Un *anneau* est un ensemble  $A$ , muni de deux lois de composition internes, une addition  $+$  :  $A \times A \rightarrow A$  et une multiplication  $\cdot$  :  $A \times A \rightarrow A$ , vérifiant les propriétés suivantes :

- $(A, +)$  est un groupe abélien dont on note  $0_A$  (ou 0 lorsqu'il n'y a pas d'ambiguïté) l'élément neutre ;
- la multiplication est associative, c'est-à-dire  $a_1.(a_2.a_3) = (a_1.a_2).a_3$  pour tous  $a_1, a_2, a_3 \in A$  ;
- il existe un élément neutre pour la multiplication, distinct de  $0_A$ , noté  $1_A$  (ou 1 lorsqu'il n'y a pas d'ambiguïté) et appelé *élément unité* ;
- la multiplication est *distributive* sur l'addition, c'est-à-dire  $a_1.(a_2 + a_3) = a_1.a_2 + a_1.a_3$  et  $(a_1 + a_2).a_3 = a_1.a_3 + a_2.a_3$  pour tous  $a_1, a_2, a_3 \in A$ .

Si, de plus,  $a_1.a_2 = a_2.a_1$  pour tous  $a_1, a_2 \in A$ , on dit que l'anneau est *commutatif*.

*Exemples 3.1.2.*

- $(\mathbb{Z}, +, \cdot)$ ,  $(\mathbb{Q}, +, \cdot)$ ,  $(\mathbb{R}, +, \cdot)$  et  $(\mathbb{C}, +, \cdot)$  sont des anneaux commutatifs.
- Pour tout entier  $n \geq 2$ ,  $(n\mathbb{Z}, +, \cdot)$  n'est pas un anneau car il ne contient pas d'unité.
- Pour tout entier  $n \geq 2$ ,  $(\mathbb{Z}/n\mathbb{Z}, +, \cdot)$  est un anneau commutatif. En particulier  $\{0, 1\} = \mathbb{Z}/2\mathbb{Z}$  sont des anneaux commutatifs.
- Pour tout  $n \in \mathbb{N}^*$  et  $\mathbb{K} = \mathbb{R}$  ou  $\mathbb{C}$ ,  $(\mathcal{M}_n(\mathbb{K}), +, \cdot)$  est un anneau, non commutatif si  $n \geq 2$ . Il en va donc de même pour  $(\text{End}(E), +, \circ)$  pour tout  $\mathbb{K}$ -espace vectoriel de dimension finie (et même de dimension infinie).
- L'ensemble  $\mathbb{R}[X]$  des polynômes réels à une indéterminée, muni de l'addition et de la multiplication des polynômes, est un anneau.
- Pour tout anneau  $A$  et tout ensemble  $X$ ,  $\mathcal{F}(X, A) := \{f : X \rightarrow A\}$  est naturellement muni d'une structure d'anneau, la somme et la multiplication étant définies pour tous  $f_1, f_2 : X \rightarrow A$  par

$$(f_1 + f_2)(x) = f_1(x) + f_2(x) \quad (f_1.f_2)(x) = f_1(x).f_2(x)$$

pour tout  $x \in X$ . Ce dernier est commutatif si et seulement si  $A$  l'est. Par exemple, l'ensemble des applications de  $\mathbb{R}$  dans  $\mathbb{R}$ , ou l'ensemble des suites à valeurs complexes forment des anneaux.

Un anneau est donc un groupe abélien pour l'addition, et à ce titre, tous les résultats du précédent chapitre s'appliquent. La multiplication, quant à elle, n'induit pas une structure de groupe car il manque des inverses pour certains éléments (nous verrons que  $0_A$  ne peut pas avoir d'inverse). Toutefois, en regardant le détail des preuves, on peut observer que certains résultats demeurent.

**Lemme 3.1.3.** Dans tout anneau, l'élément unité est unique.

**Notation 3.1.4.** Pour tout anneau, il est vraiment classique d'utiliser la notation additive pour l'addition et la notation multiplicative pour la multiplication.

À l'instar de la structure de groupe, la structure d'anneau induit naturellement un certain nombre de règles de calculs.

**Proposition 3.1.5.** Soit  $A$  un anneau. On a

- pour tout  $a \in A$ ,  $0_A.a = a.0_A = 0_A$ , on dit que  $0_A$  est *absorbant* ;
- pour tous  $a_1, a_2 \in A$ ,  $-(a_1.a_2) = (-a_1).a_2 = a_1.(-a_2)$ , et notamment  $(-a_1).(-a_2) = a_1.a_2$  ;
- pour tout  $a \in A$  et  $n_1, n_2 \in \mathbb{N}$ ,  $a^{n_1}.a^{n_2} = a^{n_1+n_2}$  et  $(a^{n_1})^{n_2} = a^{n_1.n_2}$  ;
- pour tout  $a \in A$  et  $n \in \mathbb{Z}$ ,  $n.a = (n.1_A).a = a.(n.1_A)$  ;
- (formule du binôme de Newton) si  $A$  est commutatif<sup>4</sup>, pour tous  $a_1, a_2 \in A$  et  $n \in \mathbb{N}$ ,

$$(a_1 + a_2)^n = \sum_{i=0}^n \binom{n}{i} . a_1^i . a_2^{n-i} ;$$

- (identités remarquables) si  $A$  est commutatif<sup>4</sup>, pour tous  $a_1, a_2 \in A$  et  $n \in \mathbb{N}$ ,

$$a_1^n - a_2^n = (a_1 - a_2) . \sum_{i=0}^{n-1} a_1^i . a_2^{n-i-1}.$$

---

4. ou, tout du moins, si  $a_1$  et  $a_2$  commutent

*Démonstration.* Pour le premier point, on observe que  $0_A \cdot a = (0_A + 0_A) \cdot a = 0_A \cdot a + 0_A \cdot a$  et  $a \cdot 0_A = a \cdot (0_A + 0_A) = a \cdot 0_A + a \cdot 0_A$  et on utilise le lemme 2.1.6 dans le groupe  $(A, +)$ . La première partie du deuxième point provient des égalités  $a_1 \cdot a_2 + (-a_1) \cdot a_2 = (a_1 - a_1) \cdot a_2 = 0_A \cdot a_2 = 0_A$  et  $a_1 \cdot a_2 + a_1 \cdot (-a_2) = a_1 \cdot (a_2 - a_2) = a_1 \cdot 0_A = 0_A$ ; le seconde de  $a_1 \cdot a_2 = -(-a_1 \cdot a_2) = -a_1 \cdot (-a_2) = (-a_1) \cdot (-a_2)$ . Le troisième se montre par récurrence sur  $n$  comme dans un groupe.

De même, le quatrième point se montre d'abord par récurrence pour  $n \in \mathbb{N}$ . Le résultat est en effet vrai pour  $n = 0$  d'après le premier point. Et si on suppose le résultat vrai pour  $n - 1$ , on a alors  $n \cdot a = (n - 1) \cdot a + a = ((n - 1) \cdot 1_A) \cdot a + 1_A \cdot a = ((n - 1) \cdot 1_A + 1_A) \cdot a = (n \cdot 1_A) \cdot a$  et de même à droite. Pour  $n \in \mathbb{Z} \setminus \mathbb{N}$ , on a  $n \cdot a = -(-n) \cdot a = -((-n) \cdot 1_A) \cdot a = (-(-n) \cdot 1_A) \cdot a = (n \cdot 1_A) \cdot a$  et de même à droite.

La formule du binôme de Newton se montre par récurrence sur  $n \in \mathbb{N}$ . Le résultat est en effet vrai pour  $n = 0$ . On suppose ensuite le résultat vrai au rang  $n - 1$ , on a alors

$$\begin{aligned} (a_1 + a_2)^n &= (a_1 + a_2) \cdot (a_1 + a_2)^{n-1} = (a_1 + a_2) \cdot \sum_{i=0}^{n-1} \binom{n-1}{i} \cdot a_1^i \cdot a_2^{n-1-i} \\ &= \sum_{i=0}^{n-1} \binom{n-1}{i} \cdot (a_1 + a_2) \cdot a_1^i \cdot a_2^{n-1-i} = \sum_{i=0}^{n-1} \binom{n-1}{i} \cdot (a_1 \cdot a_1^i \cdot a_2^{n-1-i} + a_2 \cdot a_1^i \cdot a_2^{n-1-i}) \\ &= \sum_{i=0}^{n-1} \binom{n-1}{i} \cdot (a_1^{i+1} \cdot a_2^{n-1-i} + a_1^i \cdot a_2^{n-i}) \end{aligned}$$

car  $A$  est commutatif. En regroupant les termes selon la puissance de  $a_1$ , on obtient alors

$$(a_1 + a_2)^n = \sum_{i=0}^n \left( \binom{n-1}{i-1} + \binom{n-1}{i} \right) \cdot a_1^i \cdot a_2^{n-i} = \sum_{i=0}^n \binom{n}{i} \cdot a_1^i \cdot a_2^{n-i},$$

avec la convention que  $\binom{n-1}{-1} = \binom{n-1}{n} = 0$ ; on a en effet, pour tout  $0 < i < n$ ,

$$\begin{aligned} \binom{n-1}{i-1} + \binom{n-1}{i} &= \frac{(n-1)!}{(i-1)!(n-i)!} + \frac{(n-1)!}{i!(n-1-i)!} = \frac{(n-1)!}{(i-1)!(n-1-i)!} \cdot \left( \frac{1}{n-i} + \frac{1}{i} \right) \\ &= \frac{(n-1)!}{(i-1)!(n-1-i)!} \cdot \frac{i+n-i}{i(n-i)} = \frac{n!}{i!(n-i)!} = \binom{n}{i}, \end{aligned}$$

ainsi que  $\binom{n-1}{-1} + \binom{n-1}{0} = 1 = \binom{n}{0}$  et  $\binom{n-1}{n-1} + \binom{n-1}{n} = 1 = \binom{n}{n}$ .

Enfin, les identités remarquables découlent du calcul direct :

$$(a_1 - a_2) \cdot \sum_{i=0}^{n-1} a_1^i \cdot a_2^{n-i-1} = \sum_{i=0}^{n-1} (a_1 - a_2) \cdot a_1^i \cdot a_2^{n-i-1} = \sum_{i=0}^{n-1} a_1^{i+1} \cdot a_2^{n-i-1} - a_1^i \cdot a_2^{n-i}$$

car  $A$  est commutatif, et donc

$$\begin{aligned} (a_1 - a_2) \cdot \sum_{i=0}^{n-1} a_1^i \cdot a_2^{n-i-1} &= \sum_{i=0}^{n-1} a_1^{i+1} \cdot a_2^{n-i-1} - \sum_{i=0}^{n-1} a_1^i \cdot a_2^{n-i} \\ &= \sum_{i=1}^n a_1^i \cdot a_2^{n-i} - \sum_{i=0}^{n-1} a_1^i \cdot a_2^{n-i} \\ &= a_1^n + \sum_{i=1}^{n-1} a_1^i \cdot a_2^{n-i} - \sum_{i=1}^{n-1} a_1^i \cdot a_2^{n-i} - a_2^n = a_1^n - a_2^n. \end{aligned}$$

□

Comme pour les groupes, on peut s'intéresser aux sous-ensembles d'un anneau qui sont eux-mêmes des anneaux.

**Définition 3.1.6.** Soit  $A$  un anneau. On dit que  $B \subset A$  est un *sous-anneau* si  $B$  est un sous-groupe de  $(A, +)$  stable par multiplication, c'est-à-dire si

- $1_A \in B$ ;
- $\forall a_1, a_2 \in A, a_1 - a_2 \in A$ ;
- $\forall a_1, a_2 \in A, a_1 \cdot a_2 \in A$ .

**Proposition 3.1.7.** Soit  $(A, +, \cdot)$  un anneau. Alors pour tout  $B \subset A$ ,  $B$  est un sous-anneau (unitaire) si et seulement si  $(B, +, \cdot)$  est un anneau (unitaire).

*Démonstration.* On sait déjà que  $(B, +)$  est un groupe si et seulement si  $B$  est un sous-groupe de  $A$  pour l'addition. De plus, la multiplication étant déjà associative et distributive dans  $A$ , elle ne peut que le rester dans  $B$ ; la seule question qui reste est donc de savoir si  $B$  est stable par multiplication.  $\square$

*Exemples 3.1.8.*

- Pour tout anneau  $A$ ,  $A$  est un sous-anneau de  $A$ , mais pas  $\{0_A\}$  car il ne contient pas d'unité.
- On a la suite d'inclusion d'anneaux unitaires

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

- On a la suite d'inclusion d'anneaux

$$C^\infty(\mathbb{R}, \mathbb{R}) \subset \dots \subset C^k(\mathbb{R}, \mathbb{R}) \subset \dots \subset C^1(\mathbb{R}, \mathbb{R}) \subset C^0(\mathbb{R}, \mathbb{R}) \subset \mathcal{F}(\mathbb{R}, \mathbb{R}),$$

où  $\mathcal{F}(\mathbb{R}, \mathbb{R})$  est l'ensemble des applications de  $\mathbb{R}$  de  $\mathbb{R}$ ,  $C^0(\mathbb{R}, \mathbb{R})$  le sous-ensemble des fonctions continues,  $C^k(\mathbb{R}, \mathbb{R})$  celui des fonctions  $k$  fois dérivables, et  $C^\infty(\mathbb{R}, \mathbb{R})$  celui des fonctions indéfiniment dérivables.

*Remarque 3.1.9.* Comme pour les groupes, on peut montrer qu'une intersection de sous-anneaux est un sous-anneau et définir une notion de sous-anneau engendré par un sous-ensemble. Dans ce cours, nous ne développerons toutefois ces notions que pour les idéaux, une notion que nous introduirons plus tard.

## 3.2 Inversibilité et intégrité

Dans un anneau, on n'oblige pas tous les éléments à avoir un inverse pour la multiplication, mais cela n'interdit pas certains d'en avoir.

**Définition 3.2.1.** Soit  $A$  un anneau. On dit que  $a \in A$  est *inversible* s'il existe  $b \in A$  tel que  $a \cdot b = b \cdot a = 1_A$ . On note  $A^\times := \{a \in A \mid a \text{ inversible}\}$ .

*Avertissement 3.2.2.* Ne pas confondre  $A^\times$  et  $A^* := A \setminus \{0_A\}$ . On a  $A^\times \subset A^*$ , car l'égalité  $0_A \cdot b = 0_A$  pour tout  $b \in A$  interdit l'existence d'un inverse pour  $0_A$ , mais l'inclusion inverse est en général fausse.

*Exemples 3.2.3.*

- Pour tout anneau  $A$ ,  $1_A \in A^\times$ .
- On a  $\mathbb{Z}^\times = \{\pm 1\}$ , mais  $\mathbb{Q}^\times = \mathbb{Q}^*$ ,  $\mathbb{R}^\times = \mathbb{R}^*$  et  $\mathbb{C}^\times = \mathbb{C}^*$ .
- Pour tout  $n \in \mathbb{N}^*$ ,  $(\mathbb{Z}/n\mathbb{Z})^\times = \{\bar{k} \mid k \in \mathbb{Z}, \text{pgcd}(k, n) = 1\}$ .
- Pour  $\mathbb{K} = \mathbb{R}$  ou  $\mathbb{C}$  et tout  $n \in \mathbb{N}^*$ ,  $(\mathcal{M}_n(\mathbb{K}))^\times = \text{GL}_n(\mathbb{K})$ .

**Proposition 3.2.4.** Pour tout anneau  $A$ ,  $(A^\times, \cdot)$  est un groupe. En particulier, l'élément unité et les inverses sont uniques.

*Démonstration.* Le seul point à vérifier est que  $A^\times$  est bien stable par multiplication. L'associativité et l'existence d'un élément neutre découlent en effet du fait que  $A$  est un anneau, et l'existence d'inverses de la définition même de  $A^\times$ . Or si  $a_1, a_2 \in A^\times$ , alors ils ont des inverses  $b_1, b_2 \in A$  et on a  $(b_2.b_1).(a_1.a_2) = b_2.(b_1.a_1).a_2 = b_2.1_A.a_2 = b_2.a_2 = 1_A$  et pareil pour  $(a_1.a_2).(b_1.b_2)$ . On a donc bien  $a_1.a_2 \in A^\times$ .  $\square$

**Notation 3.2.5.** Si  $A$  est un anneau, alors pour tout  $a \in A^\times$ , on note  $a^{-1}$  l'inverse de  $a$  et, pour tout  $n \in \mathbb{N}$ ,  $a^{-n} := (a^{-1})^n$ . Dès lors, toutes les règles usuelles de calcul dans le groupe  $A^\times$  s'appliquent.

*Remarque 3.2.6.* On appelle *corps* tout anneau  $A$  tel que  $A^\times = A^*$ , c'est-à-dire tel que tout élément non nul soit inversible. Les anneaux  $\mathbb{Q}$ ,  $\mathbb{R}$ ,  $\mathbb{C}$  ou  $\mathbb{F}_p := \mathbb{Z}/p\mathbb{Z}$ , lorsque  $p$  est un nombre premier, sont des corps. L'étude des corps est en soi un domaine important des mathématiques, mais que nous n'aborderons cependant pas dans ce cours.

Nous avons vu que, dans un anneau  $A$ ,  $0_A$  n'est jamais inversible. C'est en réalité un élément dans une famille d'éléments non inversibles potentiellement plus grande.

**Définition 3.2.7.** Soit  $A$  un anneau. On dit que  $a$  est un *diviseur de zéro* s'il existe  $b \in A^*$  tel que  $a.b = 0_A$  ou  $b.a = 0_A$ .

*Exemples 3.2.8.*

- Dans  $\mathbb{Z}$ , il n'y a aucun diviseur de zéro non trivial.
- Si  $n \in \mathbb{N}^*$  n'est pas premier, alors il existe  $a, b \in \llbracket 2, n-1 \rrbracket$  tels que  $n = a.b$  et alors  $\bar{a}$  et  $\bar{b}$  sont des diviseurs de zéro dans  $\mathbb{Z}/n\mathbb{Z}$ .
- Pour  $\mathbb{K} = \mathbb{R}$  ou  $\mathbb{C}$  et tout  $n \in \mathbb{N}^*$ , toute matrice non inversible de  $\mathcal{M}_n(\mathbb{K})$  est un diviseur de zéro.

**Proposition 3.2.9.** Dans un anneau, aucun élément ne peut être simultanément inversible et diviseur de zéro.

*Démonstration.* Supposons par l'absurde qu'il existe  $a$  inversible et  $b \neq 0$  tels que  $b.a = 0$ . On a alors  $0 = b.a.a^{-1} = b.1 = b$ , ce qui contredit l'hypothèse de non trivialité de  $b$ .  $\square$

Plus que la non inversibilité, c'est d'être un diviseur de zéro qui interdit de "simplifier" par un élément.

**Lemme 3.2.10.** Soit  $A$  un anneau et  $a \in A$  un élément qui ne soit pas un diviseur de zéro. Si  $a.b = a.c$  ou  $b.a = c.a$  avec  $b, c \in A$ , alors  $b = c$ .

*Démonstration.* Si  $a.b = a.c$ , alors  $a.(b - c) = a.b - a.c = 0_A$ . Mais puisque  $a$  n'est pas un diviseur de zéro, on a nécessairement  $b - c = 0_A$ , et donc  $b = c$ . On raisonne de même si  $b.a = c.a$ .  $\square$

On comprendra donc l'intérêt de travailler dans un anneau qui ne possède pas d'autre diviseur de zéro que 0 lui-même : dans un tel anneau, on peut, à l'instar de  $\mathbb{Z}$ , simplifier par tout élément non nul.

**Définition 3.2.11.** On dit qu'un anneau  $A$  est *intègre* si  $0_A$  y est le seul diviseur de zéro.

*Exemples 3.2.12.*

- $\mathbb{Z}$  est intègre, mais  $\mathbb{Z}/n\mathbb{Z}$  ne l'est que si  $n \in \mathbb{N}^*$  est premier.
- Tout corps est intègre. En particulier,  $\mathbb{Q}$ ,  $\mathbb{R}$  et  $\mathbb{C}$  le sont.

### 3.3 Morphismes d'anneaux

Nous allons maintenant nous intéresser aux applications qui respectent les structures d'anneaux.

**Définition 3.3.1.** Soit  $A_1, A_2$  deux anneaux. On dit qu'une application  $f : A_1 \rightarrow A_2$  est un *morphisme d'anneaux* si  $f(1_{A_1}) = 1_{A_2}$  et si, pour tous  $a_1, a_2 \in A_1$ , on a  $f(a_1 + a_2) = f(a_1) + f(a_2)$  et  $f(a_1.a_2) = f(a_1).f(a_2)$ . Comme dans le cas des groupes, on pourra préciser la nature de  $f$  en parlant



- de *monomorphisme* (d'anneaux) si  $f$  est injective ;
- d'*épimorphisme* (d'anneaux) si  $f$  est surjective ;
- d'*isomorphisme* (d'anneaux) si  $f$  est bijective ;
- d'*endomorphisme* (d'anneau) si  $A_1 = A_2$  ;
- d'*automorphisme* (d'anneau) si  $A_1 = A_2$  et que  $f$  est bijective.

*Remarque 3.3.2.* Tout morphisme d'anneaux est également un morphisme de groupes pour l'addition. Cela permet déjà de déduire un certain nombre de propriétés, telles que  $f(0_{A_1}) = 0_{A_2}$  ou  $f$  injective si et seulement si  $\text{Ker}(f) := f^{-1}(0_{A_2}) = \{0_{A_1}\}$ .

**Définition 3.3.3.** On dit que deux anneaux  $A_1$  et  $A_2$  sont *isomorphes* en tant qu'anneaux s'il existe un isomorphisme d'anneaux  $f : A_1 \rightarrow A_2$ . On note alors  $A_1 \cong A_2$ .

*Remarque 3.3.4.* Comme pour les groupes, on peut considérer que deux anneaux isomorphes sont deux représentations d'un même anneau abstrait, et on n'étudie en général les anneaux qu'à isomorphisme près. Il est toutefois important de préciser qu'on parle alors d'isomorphisme *d'anneaux*, car si tout isomorphisme d'anneaux induit un isomorphisme de groupes, la réciproque n'est pas vraie.

*Exemples 3.3.5.*

- Si  $A$  est un anneau et  $B \subset A$  un sous-anneau, alors l'injection  $B \hookrightarrow A$  est un monomorphisme d'anneaux.
- Pour tout  $n \in \mathbb{N}^*$ , l'application

$$\begin{array}{ccc} \mathbb{Z} & \longrightarrow & \mathbb{Z}/n\mathbb{Z} \\ k & \longmapsto & \bar{k} \end{array}$$

est un épimorphisme d'anneaux. Plus généralement, pour tout anneau  $A$ , il existe un unique morphisme d'anneaux allant de  $\mathbb{Z}$  vers  $A$ , à savoir

$$\begin{array}{ccc} \mathbb{Z} & \longrightarrow & A \\ k & \longmapsto & k.1_A \end{array}.$$

- Soit  $X$  un ensemble et  $A$  un anneau. Alors pour tout  $x \in X$ , l'application

$$\text{ev}_x : \begin{array}{ccc} \mathcal{F}(X, A) & \longrightarrow & A \\ f & \longmapsto & f(x) \end{array}$$

est un morphisme d'anneaux.

Comme pour les groupes, on a les propriétés suivantes :

**Proposition 3.3.6.**

- La composée de deux morphismes d'anneaux est un morphisme d'anneaux (unitaire).
- Si  $f$  est un isomorphisme d'anneaux, alors  $f^{-1}$  est également un isomorphisme d'anneaux.

*Démonstration.* La démonstration, rigoureusement similaire au cas des groupes, est laissée en exercice.  $\square$

**Corollaire 3.3.7.** Pour tout anneau  $A$ , l'ensemble des automorphismes d'anneaux de  $A$  dans lui-même forme un groupe pour la composition.

Et la notion de morphisme est toujours liée à la notion de sous-anneau par ce qui suit.

**Proposition 3.3.8.** Soit  $A_1, A_2$  deux anneaux, et  $f : A_1 \rightarrow A_2$  un morphisme d'anneaux.

- Pour tout sous-anneau  $B \subset A_2$ ,  $f^{-1}(B)$  est un sous-anneau.
- Pour tout sous-anneau  $B \subset A_1$ ,  $f(B)$  est un sous-anneau.

*Démonstration.* Là encore, la preuve, similaire au cas des groupes, est laissée en exercice.  $\square$

### 3.4 Idéaux et anneaux quotients

Si  $A$  est un anneau et  $B \subset A$  un sous-anneau, alors  $(B, +)$  est un sous-groupe de  $(A, +)$ , et ce sous-groupe est distingué puisque  $(A, +)$  est abélien. On peut donc considérer le groupe quotient  $A/B$  et se demander si la multiplication de  $A$  induit une structure d'anneau sur  $A/B$ . La réponse n'est pas toujours oui. Pour  $\mathbb{Z} \subset \mathbb{Q}$ , par exemple, on a  $0 \sim_{\mathbb{Z}} 1$  et pourtant  $0 \cdot \frac{1}{2} = 0 \not\sim_{\mathbb{Z}} \frac{1}{2} = 1 \cdot \frac{1}{2}$ . Plus généralement, pour que  $A/B$  puisse être un anneau, il faut que  $0_{A/B}$  soit absorbant, et donc que  $a \cdot b \sim_B b \cdot a \sim_B 0_A$  pour tout  $b \in B$  et  $a \in A$ , puisqu'alors  $b \sim_B 0$ .

**Définition 3.4.1.** Soit  $A$  un anneau. On dit que  $I \subset A$  est un *idéal* si  $(I, +)$  est un sous-groupe de  $(A, +)$  et que  $A \cdot I := \{a \cdot x \mid a \in A, x \in I\} \subset I$  et  $I \cdot A := \{x \cdot a \mid a \in A, x \in I\} \subset I$ .

*Remarques 3.4.2.*

- On peut raffiner la définition en parlant d'idéal à droite ou à gauche selon que l'impose seulement  $A \cdot I \subset I$  ou  $I \cdot A \subset I$ . La définition ci-dessus correspond alors à la notion d'idéal *bilatère*. Bien entendu, si  $A$  est commutatif, les notions d'idéal à droite, à gauche ou bilatère sont équivalentes.
- De même que la notion de sous-groupe distingué  $H \triangleleft G$  dépend de  $G$ , la notion d'idéal  $I \subset A$  dépend de  $A$ . L'ensemble  $\mathbb{Z}$  est en effet un idéal de lui-même, mais pas de  $\mathbb{Q}$ .

*Exemples 3.4.3.*

- Pour tout anneau  $A$ ,  $\{0\}$  et  $A$  sont des idéaux.
- Pour tout  $n \in \mathbb{N}^*$ ,  $n\mathbb{Z} \subset \mathbb{Z}$  est un idéal.
- Aucune des inclusions  $\mathbb{Z} \subset \mathbb{Q} \subset \mathbb{R} \subset \mathbb{C}$  n'est un idéal. En fait, puisqu'il doit contenir  $1_A$  et être stable par multiplication par tout élément de  $A$ , le seul sous-anneau d'un anneau  $A$  à être un idéal, c'est  $A$  tout entier. Ainsi, les inclusions  $C^\infty(\mathbb{R}, \mathbb{R}) \subset \dots \subset C^k(\mathbb{R}, \mathbb{R}) \subset \dots \subset C^0(\mathbb{R}, \mathbb{R}) \subset \mathcal{F}(\mathbb{R}, \mathbb{R})$  non plus ne sont pas des idéaux.
- Si  $A$  est un anneau et que  $I \subset A$  contient un élément inversible, alors  $I = A$ . Notamment, les seuls idéaux d'un corps  $\mathbb{K}$  sont  $\{0\}$  et  $\mathbb{K}$ .

La théorie se déroule alors comme dans le cas des sous-groupes distingués.

**Proposition 3.4.4.** Soit  $A$  un anneau (commutatif) et  $I \subset A$  un idéal. Alors la multiplication de  $A$  induit une structure d'anneau (commutatif) sur  $A/I$ .

*Démonstration.* On a déjà vu que  $A/I$  est un groupe abélien pour l'addition. Vérifions maintenant que la multiplication de  $A$  induit, par  $\overline{a} \cdot \overline{b} := \overline{a \cdot b}$ , une opération bien définie sur  $A/I$ . Pour cela, on considère  $a_1, a_2, b_1, b_2 \in A$  tels que  $a_1 \sim_I a_2$  et  $b_1 \sim_I b_2$ , c'est-à-dire tels que  $a_2 - a_1, b_2 - b_1 \in I$ . On a alors

$$a_2 \cdot b_2 - a_1 \cdot b_1 = a_2 \cdot b_2 - a_2 \cdot b_1 + a_2 \cdot b_1 - a_1 \cdot b_1 = a_2 \cdot (b_2 - b_1) + (a_2 - a_1) \cdot b_1 \in I$$

car,  $I$  étant un idéal,  $a_2 \cdot (b_2 - b_1), (a_2 - a_1) \cdot b_1 \in I$ . On a donc bien  $a_2 \cdot b_2 \sim_I a_1 \cdot b_1$ .

L'associativité et la distributivité, ainsi que l'éventuelle commutativité de cette multiplication sur  $A/I$  découlent alors directement des propriétés dans  $A$ .  $\square$

*Exemple 3.4.5.* Pour tout  $n \in \mathbb{N}^*$ ,  $\mathbb{Z}/n\mathbb{Z}$  est un anneau quotient.

La notion d'anneau quotient jouera un rôle important dans le cas des anneaux de polynômes, notamment pour l'étude des extensions de corps, mais ceci sort du cadre de ce cours.

**Proposition 3.4.6.**

- Pour tout morphisme  $f : A_1 \rightarrow A_2$  d'anneaux,  $\text{Ker}(f)$  est un idéal de  $A_1$ .
- Si  $(I_i)_{i \in I}$  est une famille d'idéaux d'un anneau  $A$ , alors  $\bigcap_{i \in I} I_i$  est un idéal de  $A$ .

*Démonstration.* On sait déjà que  $\text{Ker}(f)$  est un sous-groupe de  $A_1$  pour l'addition, il suffit de vérifier qu'il est stable par multiplication à droite ou à gauche par un élément de  $A_1$ . Or clairement, si  $x \in \text{Ker}(f)$  et  $a \in A_1$ , on a  $f(ax) = f(a).f(x) = f(a).0 = 0$  et  $f(xa) = f(x).f(a) = 0.f(a) = 0$ .

De même, on sait que  $\cap_{i \in I} I_i$  est un sous-groupe pour l'addition, et si  $x \in \cap_{i \in I} I_i$  et  $a \in A$ , alors  $x \in I_i$  pour tout  $i \in I$ , donc  $ax, xa \in I_i$  pour tout  $i \in I$ , et de fait  $ax, xa \in \cap_{i \in I} I_i$ .  $\square$

**Définition 3.4.7.** Soit  $A$  un anneau et  $X \subset A$  un sous-ensemble. On appelle *idéal engendré par  $X$*  l'intersection de tous les idéaux de  $A$  contenant  $X$ . On le note  $(X)$  ou  $(x_1, \dots, x_k)$  si  $X = \{x_1, \dots, x_k\}$  est un ensemble fini.

**Proposition 3.4.8.** Si  $A$  est un anneau et  $X \subset A$  un sous-ensemble, alors  $(X)$  est un idéal de  $A$ , minimal pour l'inclusion parmi tous les idéaux contenant  $X$ .

La preuve de ce résultat est rigoureusement identique au cas des groupes.

*Exemple 3.4.9.* Si  $A$  est un anneau commutatif et  $a \in A$  un élément, alors  $(a)$  est égal à l'ensemble  $\{a.b \mid b \in A\}$  des multiples de  $a$ .

Plus généralement, on a :

**Proposition 3.4.10.** Soit  $A$  un anneau et  $X \subset A$  un sous-ensemble. Alors

$$(X) = \left\{ \sum_{i=1}^k a_i.x_i.b_i \mid k \in \mathbb{N}, \forall i \in \llbracket 1, k \rrbracket, x_i \in X, a_i, b_i \in A \right\}.$$

Si  $A$  est commutatif, on a même  $(X) = \{ \sum_{i=1}^k a_i.x_i \mid k \in \mathbb{N}, \forall i \in \llbracket 1, k \rrbracket, x_i \in X, a_i \in A \}$ .

*Démonstration.* Notons  $B := \{ \sum_{i=1}^k a_i.x_i.b_i \mid k \in \mathbb{N}, \forall i \in \llbracket 1, k \rrbracket, x_i \in X, a_i, b_i \in A \}$ . Par stabilité de  $(X)$  par somme et produits, on a clairement  $B \subset (X)$ . Mais réciproquement, on vérifie facilement que  $B$  est un idéal de  $A$  contenant  $X$  ; par minimalité de  $(X)$ , on a donc  $(X) \subset B$ .

Si  $A$  est commutatif, on a alors  $a_i.x_i.b_i = a_i.b_i.x_i = a'_i.x_i$  avec  $a'_i := a_i.b_i$ .  $\square$

On vient de voir que l'intersection permet de définir une opération sur les idéaux. La proposition suivante en donne une seconde.

**Proposition 3.4.11.** Soit  $A$  un anneau et  $I_1, I_2 \subset A$  deux idéaux. Alors  $I_1 + I_2 := \{a_1 + a_2 \mid a_1 \in I_1, a_2 \in I_2\}$  est un idéal.

*Démonstration.* Soit  $a, b \in I_1 + I_2$ , il existe donc  $a_1, b_1 \in I_1$  et  $a_2, b_2 \in I_2$  tels que  $a = a_1 + a_2$  et  $b = b_1 + b_2$ . On a alors

- $a - b = (a_1 - b_1) + (a_2 - b_2) \in I_1 + I_2$  ;
- $a.b = (a_1 + a_2).b = a_1.b + a_2.b \in I_1 + I_2$ , et ce même si  $b$  est un élément quelconque de  $A$  ;
- $a.b = a.(b_1 + b_2) = a.b_1 + a.b_2 \in I_1 + I_2$ , et ce même si  $a$  est un élément quelconque de  $A$ .

$\square$

Nous verrons plus tard que ces deux opérations sur les idéaux, intersection et somme, généralisent en un certain sens les notions de ppcm et de pgcd des entiers.

Terminons maintenant par l'équivalent pour les anneaux du premier théorème d'isomorphisme.

**Théorème 3.4.12.** Tout morphisme d'anneaux  $f : A_1 \rightarrow A_2$  induit un isomorphisme d'anneaux  $\bar{f} : A/\text{Ker}(f) \rightarrow \text{Im}(f)$ .

La preuve de ce résultat est rigoureusement identique au cas des groupes.

On peut, maintenant, revisiter le théorème des restes chinois.

**Théorème 3.4.13.** Soit  $n_1, \dots, n_k \in \mathbb{N}^*$  deux à deux premiers entre eux. Alors  $\mathbb{Z}/n_1\mathbb{Z} \times \dots \times \mathbb{Z}/n_k\mathbb{Z}$  est isomorphe à  $\mathbb{Z}/n_1 \dots n_k\mathbb{Z}$  en tant qu'anneau.

*Démonstration.* Commençons par le cas  $k = 2$ . Pour cela, on considère l'application

$$\begin{aligned} \varphi: \quad \mathbb{Z} &\longrightarrow \mathbb{Z}/n_1\mathbb{Z} \times \mathbb{Z}/n_2\mathbb{Z} \\ k &\longmapsto (k[n_1], k[n_2]) \end{aligned}$$

Il s'agit clairement d'un morphisme d'anneaux. Ce morphisme est de plus surjectif. En effet, pour l'addition,  $(\bar{1}, \bar{0})$  est d'ordre  $n_1$  et  $(\bar{0}, \bar{1})$  d'ordre  $n_2$ ; puisque  $\mathbb{Z}/n_1\mathbb{Z} \times \mathbb{Z}/n_2\mathbb{Z}$  est abélien et que  $\text{pgcd}(n_1, n_2) = 1$ , on a vu en exercice que  $(\bar{1}, \bar{0}) + (\bar{0}, \bar{1}) = (\bar{1}, \bar{1})$  est d'ordre  $n_1 n_2 = |\mathbb{Z}/n_1\mathbb{Z} \times \mathbb{Z}/n_2\mathbb{Z}|$ , c'est-à-dire que  $(\bar{1}, \bar{1})$  est générateur de  $\mathbb{Z}/n_1\mathbb{Z} \times \mathbb{Z}/n_2\mathbb{Z}$ . Or  $(\bar{1}, \bar{1}) = \varphi(1) \in \text{Im}(\varphi)$ , et donc  $\mathbb{Z}/n_1\mathbb{Z} \times \mathbb{Z}/n_2\mathbb{Z} \subset \text{Im}(\varphi)$ , l'inclusion réciproque étant évidente. Enfin, si  $k \in \text{Ker}(\varphi)$ ,  $n_1$  et  $n_2$  divisent  $k$ , et donc  $k \in n_1 n_2 \mathbb{Z}$  puisque  $\text{pgcd}(n_1, n_2) = 1$ . On en conclut que  $\text{Ker}(\varphi) \subset n_1 n_2 \mathbb{Z}$ ; l'inclusion réciproque étant immédiate, on déduit du théorème 3.4.12 que  $\mathbb{Z}/n_1\mathbb{Z} \times \mathbb{Z}/n_2\mathbb{Z} \cong \mathbb{Z}/n_1 n_2 \mathbb{Z}$ .

On travaille ensuite par récurrence sur  $k$  en observant que  $n_k$  est premier avec  $n_1 \dots n_{k-1}$ . On a donc, par hypothèse de récurrence et d'après le cas  $k = 2$ ,

$$\mathbb{Z}/n_1\mathbb{Z} \times \dots \times \mathbb{Z}/n_k\mathbb{Z} \cong \mathbb{Z}/n_1 \dots n_{k-1}\mathbb{Z} \times \mathbb{Z}/n_k\mathbb{Z} \cong \mathbb{Z}/n_1 \dots n_k\mathbb{Z}.$$

□

## 4 Anneaux : compléments

Dans toute cette partie, on fixe  $A$  un anneau commutatif et intègre. Nous avons vu que la notion d'anneau est calquée sur les propriétés algébriques de  $(\mathbb{Z}, +, \cdot)$ , ce dernier exemple vérifie cependant un certain nombre de propriétés supplémentaires, que l'on peut éventuellement demander à un anneau. Nous introduisons maintenant certaines de ces notions, de la plus générale à la plus restrictive.

### 4.1 Anneaux factoriels

Un des éléments les plus centraux de l'étude des nombres entiers est l'unique décomposition en facteurs premiers. Mais avant de généraliser cette notion, il convient de distinguer les notions d'élément premier et d'élément irréductible. Bien entendu, ces notions ne concernent ni 0, qui est absorbant, ni les inversibles qui peuvent toujours être rajoutés artificiellement en facteur.

**Définition 4.1.1.** Soit  $a \in A^* \setminus A^\times$ .

- On dit que  $a$  est *irréductible* s'il n'est pas le produit de deux éléments non inversibles. Cela revient à dire que si  $a = b.c$  avec  $b, c \in A$ , alors  $b$  ou  $c$  est inversible.
- On dit que  $a$  est *premier* si, pour tout  $b_1, b_2 \in A$  tels que  $b_1.b_2 \in (a)$ , on a  $b_1 \in (a)$  ou  $b_2 \in (a)$ . Cela revient à dire que  $a$  vérifie le critère d'Euclide, à savoir que si un produit est un multiple de  $a$ , alors l'un des facteurs au moins est un multiple de  $a$ .

*Exemples 4.1.2.*

- Dans  $\mathbb{Z}$ , les éléments premiers et les éléments irréductibles sont les mêmes, et ils correspondent aux nombres premiers et leurs opposés.
- Dans  $\mathbb{C}[X]$ , les éléments premiers et les éléments irréductibles sont les mêmes, et ils correspondent aux polynômes de degré 1. Dans  $\mathbb{R}[X]$  également, mais il faut alors rajouter les polynômes de degré 2 dont le discriminant est strictement négatif.

- Dans  $\mathbb{Z}[i\sqrt{3}] := \{a + ib\sqrt{3} \mid a, b \in \mathbb{Z}\}$ , on peut montrer que 2 est irréductible mais il n'est pas premier car  $(1 + i\sqrt{3})(1 - i\sqrt{3}) = 4 \in (2)$  mais  $(1 + i\sqrt{3}), (1 - i\sqrt{3}) \notin (2)$ .

**Proposition 4.1.3.** Tout élément premier est irréductible.

*Démonstration.* Soit  $a \in A$  un élément premier. Supposons que  $a = b_1.b_2$  avec  $b_1, b_2 \in A$ . Alors  $b_1.b_2 \in (a)$  donc par primalité, et quitte à échanger  $b_1$  et  $b_2$ , on a  $b_2 \in (a)$ , c'est-à-dire  $b_2 = c.a$  avec  $c \in A$ . Mais alors  $a = b_1.c.a$  et par intégrité de  $A$ ,  $a$  étant non nul,  $1 = b_1.c$ . On en déduit que  $b_1$  est inversible.  $\square$

*Remarque 4.1.4.* L'hypothèse d'intégrité de  $A$  est ici essentielle. En effet, dans  $\mathbb{Z}/6\mathbb{Z}$ , 2 n'est pas irréductible car  $2 = 2.4$  avec 2 et 4 non inversibles, mais 2 est premier car  $(2) = \{0, 2, 4\}$  et un produit d'éléments dans  $\mathbb{Z}/6\mathbb{Z} \setminus (2) = \{1, 3, 5\}$  restera dans  $\{1, 3, 5\}$  par un argument de parité.

**Définition 4.1.5.** On dit que  $A$  est factoriel si tout élément  $a \in A^* \setminus A^\times$  est un produit d'éléments premiers.

Sous cette hypothèse, les notions d'éléments premiers et irréductibles coïncident.

**Proposition 4.1.6.** Si  $A$  est factoriel, alors un élément  $a \in A^* \setminus A^\times$  est premier si et seulement si il est irréductible.

*Démonstration.* On a déjà vu que tout élément premier est irréductible, il suffit donc de montrer la réciproque. Soit donc  $a \in A^* \setminus A^\times$  irréductible. Par factorialité de  $A$ , il existe des éléments  $p_1, \dots, p_k \in A^* \setminus A^\times$  premiers tels que  $a = p_1 \dots p_k$ . Supposons par l'absurde que  $k > 1$ . Puisque  $p_1$  est non inversible, on sait que  $p_2 \dots p_k$  l'est par irréductibilité de  $a$ . Il existe donc  $c \in A$  tel que  $p_2 \dots p_k.c = 1$ . Mais alors  $p_2$  est inversible ce qui contredit sa primalité. On en déduit que  $k = 1$  et  $a = p_1$  est donc premier.  $\square$

On peut maintenant exprimer la factorialité de  $A$  en terme d'unique factorisation.

**Définition 4.1.7.** Deux éléments  $a_1, a_2 \in A$  sont *associés* s'il existe  $b \in A^\times$  tel que  $a_2 = b.a_1$ . On parle alors d'*association* entre  $a_1$  et  $a_2$ .

**Proposition 4.1.8.**

- L'association est une relation d'équivalence.
- Deux éléments  $a_1, a_2 \in A$  sont associés si et seulement si  $(a_1) = (a_2)$ .
- Deux éléments associés sont simultanément irréductibles ou non, et simultanément premiers ou non.

*Démonstration.* Le premier point est clair.

Concernant le second point, supposons d'abord que  $a_1$  et  $a_2$  sont associés. Il existe alors  $b \in A^\times$  tel que  $a_1 = b.a_2$  et on a alors d'une part  $a_1 = b.a_2 \in (a_2)$  et donc  $(a_1) \subset (a_2)$ , et d'autre part  $a_2 = b^{-1}.a_1 \in (a_1)$  et donc  $(a_2) \subset (a_1)$ . Réciproquement, on suppose  $(a_1) = (a_2)$ . On a alors d'une part  $a_1 \in (a_2)$  et il existe donc  $b \in A$  tel que  $a_1 = b.a_2$ , et d'autre part  $a_2 \in (a_1)$  et il existe donc  $c \in A$  tel que  $a_2 = c.a_1$ . On en déduit que  $a_1 = b.a_2 = b.c.a_1$  et, par intégrité de  $A$ , que  $1 = b.c$ . L'élément  $b$  est donc inversible; les éléments  $a_1$  et  $a_2$  sont donc associés.

Enfin, pour le dernier point, l'affirmation est claire concernant la primalité car la définition de cette dernière ne fait intervenir que l'idéal engendré par l'élément en question. Et concernant l'irréductibilité, supposons que  $a_1$  et  $a_2$  soient associés. Il existe donc  $b \in A^\times$  tel que  $a_2 = b.a_1$ . Alors toute décomposition  $a_1 = c_1.c_2$  avec  $c_1$  et  $c_2$  non inversibles induit une décomposition  $a_2 = (b.c_1).c_2$  avec  $b.c_1$  et  $c_2$  non inversibles. En effet, si  $b.c_1$  était inversible, alors  $c_1 = b^{-1}.b.c_1$  le serait aussi. On en déduit que  $a_1$  irréductible implique  $a_2$  irréductible et, par symétrie des rôles de  $a_1$  et  $a_2$ , que  $a_2$  irréductible implique  $a_1$  irréductible.  $\square$

**Proposition 4.1.9.** L'anneau  $A$  est *factoriel* ssi

- pour tout  $a \in A^* \setminus A^\times$  il existe  $p_1, p_2, \dots, p_k$  des éléments irréductibles de  $A$  tels que  $a = p_1 \dots p_k$ ;

- si  $p_1 \cdots p_k = p'_1 \cdots p'_{k'}$  avec tous les  $p_i, p'_i \in A$  irréductibles, alors  $k = k'$  et il existe  $\sigma \in \mathfrak{S}_k$  tels que  $p'_i$  et  $p_{\sigma(i)}$  sont associés pour tout  $i \in \llbracket 1, k \rrbracket$ .

Autrement dit,  $A$  est factoriel si tout élément non nul et non inversible admet une unique décomposition en facteurs irréductibles, à ordre et association près des facteurs.

*Démonstration.* Supposons que  $A$  est factoriel. Alors tout élément non nul et non inversible s'écrit comme produit d'éléments premiers et donc comme produit d'éléments irréductibles. On montre le second point par récurrence sur  $k$ , en supposant, quitte à échanger les rôles des  $p_i$  et des  $p'_i$ , que  $k \leq k'$ . Pour  $k = 1$ , le résultat est clair si  $k' = 1$ . Supposons donc par l'absurde que  $k' > 1$ , alors on a  $p'_1 \cdots p'_{k'} \in (p_1)$  et donc, par application successive de la primalité de  $p_1$ , il existe  $j \in \llbracket 1, k' \rrbracket$  tel que  $p'_j \in (p_1)$ . Par commutativité de  $A$  et quitte à réindexer les  $p'_i$ , on peut supposer que  $j = 1$ . Il existe donc  $b \in A$  tel que  $p'_1 = p_1.b$ , et donc  $p_1 = p_1.b.p'_2 \cdots p'_{k'}$ . Par intégrité de  $A$ , on obtient  $1 = b.p'_2 \cdots p'_{k'}$  et donc  $p'_{k'}$  inversible ce qui contredit sa primalité/son irréductibilité.

Supposons maintenant le résultat vrai pour  $k - 1$ . De l'égalité  $p_1.p_2 \cdots p_k = p'_1.p'_2 \cdots p'_{k'}$ , on en déduit que  $p'_1.p'_2 \cdots p'_{k'} \in (p_1)$  et donc, par application successive de la primalité de  $p_1$ , qu'il existe  $j \in \llbracket 1, k' \rrbracket$  tel que  $p'_j \in (p_1)$ . Par commutativité de  $A$  et quitte à réindexer les  $p'_i$ , on peut supposer que  $j = 1$ . Il existe donc  $b \in A$  tel que  $p'_1 = p_1.b$ . Par irréductibilité de  $p'_1$  l'un des facteurs  $b$  ou  $p_1$  est inversible, et par primalité/irréductibilité de  $p_1$ , ça ne peut être que  $b$ . On en déduit que  $p_1$  et  $p'_1$  sont associés. De plus, par intégrité de  $A$ , on a  $p_2 \cdots p_k = (b.p'_2).p'_3 \cdots p'_{k'}$  avec  $(b.p'_2), p'_3, \dots, p'_{k'}$  irréductibles et on conclut par hypothèse de récurrence.

Réciproquement, supposons maintenant les deux points vérifiés et montrons que  $A$  est factoriel. Pour tout  $a \in A^* \setminus A^\times$ , le premier point donne une factorisation de  $a$  en produits d'éléments irréductibles. Il suffit de montrer que chacun de ces facteurs est en fait premier. Supposons donc par l'absurde qu'il existe  $p \in A^* \setminus A^\times$  irréductible mais non premier. Il existe alors  $b_1, b_2 \in A$  tels que  $b_1.b_2 \in (p)$  mais  $b_1, b_2 \notin (p)$ . On a alors  $b_1, b_2 \neq 0$ , car  $0 \in (p)$ , et  $b_1, b_2 \notin A^\times$  car si, par exemple  $b_2 \in A^\times$ , alors  $b_1 = b_1.b_2.b_2^{-1} \in (p)$ . Par hypothèse, il existe donc des éléments irréductibles  $p_1, \dots, p_k$  et  $q_1, \dots, q_l$  de  $A$  tels que  $b_1 = p_1 \cdots p_k$  et  $b_2 = q_1 \cdots q_l$ . Mais alors  $p = p_1 \cdots p_k.q_1 \cdots q_l$ , ce qui viole le second point car  $k + l > 1$ .  $\square$

*Exemples 4.1.10.*

- L'ensemble des entiers  $\mathbb{Z}$  est factoriel.
- L'ensemble des entiers de Gauss  $\mathbb{Z}[i] := \{a + ib \mid a, b \in \mathbb{Z}\}$  est factoriel.
- L'ensemble  $\mathbb{Z}[i\sqrt{3}] := \{a + ib\sqrt{3} \mid a, b \in \mathbb{Z}\}$  n'est pas factoriel car  $4 = 2.2 = (1 + i\sqrt{3}).(1 - i\sqrt{3})$  or on peut montrer que  $2, 1 + i\sqrt{3}$  et  $1 - i\sqrt{3}$  sont irréductibles et non associés.
- Les ensembles de polynômes  $\mathbb{R}[X]$  et  $\mathbb{C}[X]$  sont factoriels.

Afin de simplifier l'unicité de la décomposition en facteurs irréductibles, on fixe parfois un système de représentants d'éléments irréductibles dans chaque classe d'association. Par exemple, dans  $\mathbb{Z}$ , il existe dans chaque classe d'association un unique représentant positif : cela conduit à l'ensemble  $\mathcal{P}$  des nombres premiers usuels et le théorème d'unique décomposition en facteurs premiers devient donc

$$n = \varepsilon.p_1 \cdots p_k$$

avec  $\varepsilon \in \mathbb{Z}^\times = \{\pm 1\}$  et  $p_i \in \mathcal{P}$ , les  $p_i$  étant uniques à ordre près.

Dans  $\mathbb{C}[X]$ , il existe dans chaque classe d'association un unique représentant unitaire, c'est-à-dire dont le coefficient dominant vaut 1 : cela conduit au théorème d'unique décomposition en facteurs premiers suivant

$$P = \alpha.P_1 \cdots P_k$$

avec  $\alpha \in \mathbb{C}^*$  et  $P_i := X - \beta_i$  avec  $\beta_i \in \mathbb{C}$ , les  $P_i$  étant uniques à ordre près. De même, dans  $\mathbb{R}[X]$ , on peut écrire de manière unique tout polynôme non nul sous la forme

$$P = \alpha.P_1 \cdots P_k$$

avec  $\alpha \in \mathbb{R}^*$  et  $P_i = X - \beta_i$  pour un certain  $\beta_i \in \mathbb{R}$  ou  $P_i = X^2 + \beta_i X + \gamma_i$  pour certains  $\beta_i, \gamma_i \in \mathbb{R}$  tels que  $\beta_i^2 - 4\gamma_i < 0$ , les  $P_i$  étant uniques à ordre près.

Terminons sur une dernière caractérisation des anneaux factoriels.

**Proposition 4.1.11.** L'anneau  $A$  est factoriel ssi

- tout élément irréductible est premier ;
- toute suite croissante d'idéaux engendrés par un élément

$$(a_1) \subset (a_2) \subset (a_3) \subset \cdots$$

est stationnaire.

*Démonstration.* Commençons par supposer que  $A$  est factoriel. Nous avons déjà démontré dans la preuve de la proposition 4.1.9 que tout élément irréductible est premier. Il ne reste donc plus que le second point à montrer. Considérons pour cela une telle suite croissante d'idéaux et supposons par l'absurde qu'elle n'est pas stationnaire. On peut alors en extraire une sous-suite strictement croissante. On supposera donc maintenant que la suite initiale était strictement croissante et, quitte à supprimer le premier terme, que  $a_1 \neq 0$ . Pour tout  $i \in \mathbb{N}$ , il existe donc  $b_i \in A^*$  tel que  $a_i = b_i \cdot a_{i+1}$ , et comme  $(a_i) \neq (a_{i+1})$ , on sait de plus que  $b_i$  n'est pas inversible. Les  $a_i$  sont également non inversibles, car sinon on aurait  $(a_{i_0}) = A$  pour un certain  $i_0 \in \mathbb{N}$  et la suite serait stationnaire à partir de ce terme. Par factorialité de  $A$ , chaque  $a_i$  et chaque  $b_i$  peut donc s'écrire comme produit d'au moins un facteur irréductible. Or par récurrence, on a  $a_1 = b_1 \cdot b_2 \cdots b_k \cdot a_{k+1}$  pour tout  $k \in \mathbb{N}$ . En remplaçant chaque terme par sa décomposition en facteurs irréductibles, cela donne des décompositions de  $a_1$  avec un nombre arbitrairement grand de facteurs irréductibles ; cela contredit l'unicité de la décomposition.

Réciproquement, supposons les deux points vérifiés. Pour montrer que  $A$  est factoriel, il suffit de montrer que tout élément non nul et non inversible s'écrit comme produit de facteurs irréductibles, le premier point permettra alors de conclure. Supposons par l'absurde qu'il existe un élément  $a_0 \in A^* \setminus A^\times$  n'admettant pas de telle décomposition. En particulier, il n'est pas lui-même irréductible, il existe donc  $a_1, b_1 \in A^* \setminus A^\times$  non inversibles tels que  $a_0 = a_1 \cdot b_1$ . Mais parmi ces deux éléments  $a_1$  et  $b_1$ , l'un au moins doit ne pas admettre de décomposition en facteurs irréductibles, autrement la concaténation de ces facteurs donnerait une décomposition pour  $a_0$  ; quitte à les échanger, on peut supposer qu'il s'agit de  $a_1$ . Mais alors, pour les même raisons, il existe  $a_2, b_2 \in A^* \setminus A^\times$  tels que  $a_1 = a_2 \cdot b_2$  et  $a_2$  n'admet pas de décomposition en facteurs irréductibles. Par récurrence, on construit une suite  $(a_i)_{i \in \mathbb{N}}$  d'éléments de  $A^* \setminus A^\times$  telle que, pour tout  $i \in \mathbb{N}^*$ , il existe  $b_i \in A^* \setminus A^\times$  vérifiant  $a_{i-1} = a_i \cdot b_i$ , c'est-à-dire  $(a_{i-1}) \subsetneq (a_i)$ . Mais cela donne donc une suite croissante non stationnaire d'idéaux engendrés chacun par un élément, et cela contredit le second point. Tout élément de  $A^* \setminus A^\times$  admet donc une décomposition en facteurs irréductibles.  $\square$

## 4.2 Anneaux principaux

Tout idéal de  $\mathbb{Z}$  est également un sous-groupe de  $(\mathbb{Z}, +)$ , or nous avons vu que tout sous-groupe de  $\mathbb{Z}$  est de la forme  $n\mathbb{Z}$  avec  $n \in \mathbb{N}$ . On en déduit que tout idéal de  $\mathbb{Z}$  est de la forme  $(n)$  avec  $n \in \mathbb{Z}$ . Cela permet de faire une identification entre les idéaux de  $\mathbb{Z}$  et les classes d'association dans  $\mathbb{Z}$ . Or toutes les notions liées à la divisibilité ne dépendent en réalité que de la classe d'association. De fait, pour tout  $a, b \in \mathbb{Z}^*$ , on a  $a|b \Leftrightarrow b \in (a)$ , et pour tout  $a_1, a_2 \in \mathbb{Z}^*$ ,  $(a_1) \cap (a_2) = (\text{ppcm}(a_1, a_2))$  et  $(a_1) + (a_2) = (\text{pgcd}(a_1, a_2))$ , le second point étant une reformulation du théorème de Bézout. Toutes les propriétés arithmétiques peuvent de fait se réinterpréter en terme d'idéaux. Cela motive les définitions suivantes.

**Définition 4.2.1.** Soit  $a_1, a_2 \in A$ . On dit que :

- $a_1$  *divise*  $a_2$ , ou encore que  $a_2$  est un *multiple* de  $a_1$  si  $a_2 \in (a_1)$  ;
- $m \in A$  est un *plus petit multiple commun* de  $a_1$  et  $a_2$  si  $m$  est un multiple de  $a_1$  et de  $a_2$  et que tout multiple commun à  $a_1$  et  $a_2$  est aussi un multiple de  $m$  ;
- $d \in A$  est un *plus grand diviseur commun* de  $a_1$  et  $a_2$  si  $d$  est un diviseur de  $a_1$  et de  $a_2$  et que tout diviseur commun à  $a_1$  et  $a_2$  est un diviseur de  $d$ .

*Remarques 4.2.2.*

- L'ensemble des multiples communs à deux éléments n'est jamais vide car il contient toujours au moins leur produit, et l'ensemble des diviseurs communs non plus car il contient toujours au moins l'élément unité.
- Comme son nom l'indique, la notion de plus petit multiple commun (resp. plus grand diviseur commun) correspond à un minimum global (resp. maximum global), au sens de la divisibilité, sur l'ensemble des multiples communs (resp. diviseurs communs). Il est toutefois délicat de l'exprimer ainsi car, en général, la relation de divisibilité n'est pas une relation d'ordre. A l'instar du corollaire 1.2.4, elle vérifie en effet la réflexivité et la transitivité, mais pas forcément l'antisymétrie, deux éléments associés étant mutuellement divisibles l'un par l'autre. C'est d'ailleurs pour cela que l'on parle d'un plus petit multiple commun (resp. plus grand diviseur commun) et non *du* car ces derniers ne sont définis qu'à association près. Notons toutefois au passage que, si un choix de représentant a été fait dans chaque classe d'association, on peut alors parler, pour tous  $a_1, a_2 \in A$ , *du* plus grand diviseur commun  $\text{pgcd}(a_1, a_2)$  et *du* plus petit multiple commun  $\text{ppcm}(a_1, a_2)$ . C'est le cas dans  $\mathbb{Z}$ , où l'on choisit l'unique représentant positif, ainsi que dans  $\mathbb{R}[X]$  et  $\mathbb{C}[X]$  où l'on choisit l'unique représentant unitaire.
- Les notions de plus petit multiple et de plus grand diviseurs communs peuvent se réinterpréter en terme d'inclusions d'idéaux. En effet :
  - ▶ – dire que  $m$  est un multiple de  $a_1$  et de  $a_2$ , c'est dire que  $m$  appartient à  $(a_1)$  et à  $(a_2)$ , donc à  $(a_1) \cap (a_2)$ , ou encore dire que  $(m) \subset (a_1) \cap (a_2)$  ;
  - dire que tout multiple commun à  $a_1$  et  $a_2$  est aussi un multiple de  $m$ , c'est dire que  $(a_1) \cap (a_2) \subset (m)$  ;
 donc dire que  $m$  est un plus petit multiple commun de  $a_1$  et  $a_2$ , c'est dire que  $(m) = (a_1) \cap (a_2)$  ;
  - ▶ – dire que  $d$  est un diviseur commun de  $a_1$  et de  $a_2$ , c'est dire que  $a_1, a_2 \in (d)$  et donc que  $(a_1), (a_2) \subset (d)$ , ce qui revient à dire que  $(a_1) + (a_2) \subset (d)$  ;
  - dire que tout diviseur commun à  $a_1$  et  $a_2$  est aussi un diviseur de  $d$ , c'est dire que tout idéal de la forme  $(a)$  contenant  $(a_1)$  et  $(a_2)$ , ce qui est équivalent à contenir  $(a_1) + (a_2)$ , contient également  $(d)$  ;
 donc dire que  $d$  est un plus grand diviseur commun de  $a_1$  et  $a_2$ , c'est dire que  $(d)$  est un minimum global, au sens de l'inclusion, parmi les idéaux engendré par un élément et contenant  $(a_1) + (a_2)$ .

**Définition 4.2.3.** On dit qu'un idéal  $I \subset A$  est *principal* s'il existe  $a \in A$  tel que  $I = (a)$ . On dit que  $A$  est *principal* si tout idéal de  $A$  est principal.

*Exemples 4.2.4.*

- L'anneau des entiers  $\mathbb{Z}$  est principal.
- Tout corps  $\mathbb{K}$  est principal car ses seuls idéaux sont  $\{0\} = (0)$  et  $\mathbb{K} = (1)$ .
- Nous allons voir que les anneaux  $\mathbb{R}[X]$  et  $\mathbb{C}[X]$  sont principaux.
- L'anneau  $\mathbb{Z}[X]$ , par contre, n'est pas principal car  $I = (2, X)$  n'est pas principal. En effet,  $I$  correspond aux polynômes dont le terme constant est pair. S'il était engendré par un élément  $P \in \mathbb{Z}[X]$ , cet élément devrait être de degré 0 car sinon  $I$  ne pourrait pas contenir 2 ; mais s'il était engendré par un élément  $n \in \mathbb{Z}$ , alors  $n$  devrait être pair et tous les coefficients de tous les éléments de  $I$  seraient pairs.

Sous l'hypothèse que  $A$  est principal, deux éléments  $a_1, a_2 \in A$  possèdent toujours des plus petit multiple et plus grand diviseurs communs puisque  $(a_1) \cap (a_2)$  et  $(a_1) + (a_2)$  sont alors principaux et s'écrivent donc respectivement sous la forme  $(a_1) \cap (a_2) = (m)$  et  $(a_1) + (a_2) = (d)$  avec  $m, d \in A$ . Dès lors, le résultat suivant devient totalement tautologique.

**Proposition 4.2.5** (Bachet–Bézout). Si  $A$  est principal, alors pour tout  $a_1, a_2 \in A$ , il existe  $b_1, b_2 \in A$  tels que  $a_1.b_1 + a_2.b_2$  soit un plus grand diviseur commun de  $a_1$  et  $a_2$ .

*Remarques 4.2.6.*



- On peut dès lors développer toute une arithmétique similaire à celle des entiers relatifs. On peut par exemple définir la notion d'éléments  $a_1, a_2 \in A$  *premiers entre eux* par le fait d'avoir 1 comme plus grand diviseur commun. Cela se traduit également par  $(a_1) + (a_2) = (1) = A$ .
- Pour un anneau non principal, on peut également développer une arithmétique, mais celle-ci ne se fera alors pas sur l'ensemble de ses éléments, mais sur l'ensemble de ses idéaux.

Les anneaux principaux sont en fait un cas particulier des anneaux factoriels.

**Proposition 4.2.7.** Tout anneau principal est factoriel.

*Démonstration.* Supposons que  $A$  est principal. Pour montrer qu'il est factoriel, nous allons utiliser la proposition 4.1.11. Commençons donc par montrer que tout élément irréductible  $a \in A$  est premier. Considérons donc  $b_1, b_2 \in A$  tels que  $b_1.b_2 \in (a)$  et montrons que  $b_1$  ou  $b_2$  est dans  $(a)$ . Pour ce faire, on considère un générateur<sup>5</sup>  $d \in A$  de l'idéal  $(a) + (b_1)$ . En particulier, il existe  $c \in A$  tel que  $a = c.d$ , et comme  $a$  est irréductible,  $c$  ou  $d$  doit être inversible. Si  $c$  est inversible, alors  $d$  est associé à  $a$  et  $b_1 \in (d) = (a)$ . Si  $d$  est inversible, alors on écrit  $d = u.a + v.b_1$  avec  $u, v \in A$  à l'aide de la proposition de Bachet-Bézout. On a alors  $b_2 = d^{-1}.d.b_2 = d^{-1}.(u.a + v.b_1).b_2 = d^{-1}.u.a.b_2 + d^{-1}.v.b_1.b_2 \in (a)$  puisque  $d^{-1}.u.a.b_2, d^{-1}.v.b_1.b_2 \in (a)$ . On en déduit donc que  $a$  est premier.

Montrons maintenant que toute suite croissante d'idéaux est stationnaire. On considère donc  $(a_1) \subset (a_2) \subset \dots$ , avec  $a_1, a_2, \dots \in A$ , une suite infinie d'idéaux inclus les uns dans les autres. Une réunion d'idéaux n'est en général pas un idéal, mais elle l'est si les idéaux sont ainsi imbriqués. En effet, considérons  $I = \cup_{i=1}^{\infty} (a_i)$ . Pour tout  $b_1, b_2 \in I$  il existe  $i_1, i_2 \in \mathbb{N}^*$  tels que  $b_1 \in (a_{i_1})$  et  $b_2 \in (a_{i_2})$ , mais alors  $b_1, b_2 \in (a_{\max(i_1, i_2)})$  et donc  $b_1 - b_2 \in (a_{\max(i_1, i_2)}) \subset I$ . Et pour tout  $b \in I$  et  $c \in A$ , il existe  $i \in \mathbb{N}^*$  tel que  $b \in (a_i)$  et donc  $b.c \in (a_i) \subset I$ . L'ensemble  $I$  est donc un idéal, et par principalité de  $A$ , il existe donc  $a \in A$  tel que  $I = (a)$ . En particulier,  $a \in I$  et il existe donc  $i_0 \in \mathbb{N}^*$  tel que  $a \in (a_{i_0})$ . Mais alors, pour tout entier  $i \geq i_0$ , on a  $(a_i) \subset I$  par définition de  $I$ , et  $I = (a) \subset (a_i)$  car  $a \in (a_{i_0}) \subset (a_i)$ . On en déduit que la suite est constante égale à  $I$  à partir du rang  $i_0$ .  $\square$

**Corollaire 4.2.8.** Dans tout anneau principal, un élément non nul et non inversible admet une décomposition en facteurs irréductibles, et cette décomposition est unique à ordre et association des facteurs près.

*Exemple 4.2.9.* Nous avons vu que l'anneau  $\mathbb{Z}[X]$  n'est pas principal car il contient l'idéal  $(2, X)$  qui n'est pas principal. Nous verrons toutefois qu'il est factoriel. Il n'y a donc pas équivalence entre anneaux factoriels et anneaux principaux.

Avec les définitions données dans la section précédente et en reprenant verbatim, lorsque nécessaire, les preuves du cas  $A = \mathbb{Z}$ , on obtient les lemmes usuels d'arithmétique :

**Lemme 4.2.10.** Si  $A$  est principal, alors :

- (lemme d'Euclide) si  $a \in A^* \setminus A^\times$  irréductible divise un produit  $b.c$ , alors  $a$  divise  $b$  ou  $a$  divise  $c$  ;
- (lemme de Gauss) si  $a, b, c \in A^*$  sont tels que  $a$  divise  $b.c$  et que  $a$  et  $b$  sont premiers entre eux, alors  $a$  divise  $c$  ;
- si  $a, b \in A^*$  sont premiers entre eux et divisent  $c \in A$ , alors  $a.b$  divise  $c$ .

### 4.3 Anneaux euclidiens

Dans notre étude  $\mathbb{Z}$  au premier chapitre, la division euclidienne est clairement apparue comme un outil majeur.

**Définition 4.3.1.** On dit que  $A$  est *euclidien* s'il existe une application  $\nu : A \rightarrow \mathbb{N}$ , appelée *stathme*, telle que :

- pour tout  $a \in A$ ,  $\nu(a) = 0$  si et seulement si  $a = 0_A$  ;

---

5. on utilise ici l'hypothèse  $A$  principal

- pour tous  $a, b \in A^*$ ,  $\nu(a.b) \geq \nu(a)$  ;
- pour tout  $a \in A$  et  $b \in A^*$ , il existe  $q, r \in A$  tels que  $a = b.q + r$  et  $\nu(r) < \nu(b)$ .

*Remarques 4.3.2.*

- On n'impose ici aucune condition d'unicité sur  $q$  et  $r$ .
- La seconde condition affirme qu'un stathme est "croissant pour le pré-ordre de la divisibilité", dans le sens où, pour tous  $a, b \in A^*$ , si  $a$  divise  $b$ , alors  $\nu(a) \leq \nu(b)$ .

*Exemples 4.3.3.*

- L'application  $\nu$ : 
$$\begin{array}{ccc} \mathbb{Z} & \longrightarrow & \mathbb{N} \\ n & \longmapsto & |n| \end{array}$$
 définit une structure d'anneau euclidien sur  $\mathbb{Z}$ .
- Pour  $\mathbb{K} = \mathbb{R}$  ou  $\mathbb{C}$ , l'application  $\nu$ : 
$$\begin{array}{ccc} \mathbb{K}[X] & \longrightarrow & \mathbb{N} \\ P & \longmapsto & 1 + \deg(P) \end{array}$$
, avec la convention que  $\deg(0) = -1$ , définit une structure d'anneau euclidien sur  $\mathbb{K}[X]$ .
- L'application  $\nu$ : 
$$\begin{array}{ccc} \mathbb{Z}[i] & \longrightarrow & \mathbb{N} \\ a + ib & \longmapsto & a^2 + b^2 \end{array}$$
 définit une structure d'anneau euclidien sur les entiers de Gauss.

**Proposition 4.3.4.** Tout anneau euclidien est principal.

*Démonstration.* Soit  $A$  un anneau euclidien et  $I \subset A$  un idéal. Si  $I = \{0\}$ , alors  $I = (0)$ . Autrement, l'ensemble  $\nu(I^*)$  est un sous-ensemble non vide de  $\mathbb{N}$ , il possède donc un plus petit élément et on peut considérer  $b \in I^*$  tel que  $\nu(b) = \min\{\nu(a) \mid a \in I^*\}$ . On a alors  $(b) \subset I$ . Et réciproquement, pour tout  $a \in I$ , il existe  $q, r \in A$  tels que  $a = q.b + r$  avec  $\nu(r) < \nu(b)$ . Mais  $r = a - q.b \in I$  et donc, par minimalité de  $b$ , on ne peut avoir que  $r = 0_A$ . On en déduit que  $a = q.b \in (b)$  et donc que  $I = (b)$ .  $\square$

*Exemple 4.3.5.* On peut montrer que l'anneau  $\mathbb{Z}\left[\frac{1+i\sqrt{19}}{2}\right]$  est principal mais pas euclidien.

## 5 Un exemple important : l'anneau des polynômes

La notion de polynôme est une notion fondamentale en mathématiques, apparaissant dans de nombreux contextes, mais sous des formes parfois trompeuses. Les polynômes sont en effet souvent confondus avec les fonctions polynomiales. Il s'ensuit, en général, une grande confusion sur la nature du  $X$  apparaissant couramment dans ce contexte. Cette confusion peut s'expliquer d'une part par les natures très diverses que peut en effet recouvrir ce  $X$ , mais aussi par le fait qu'il ne joue, en réalité, aucun rôle puisque qu'un polynôme est entièrement déterminé par ses coefficients. Nous allons donner ici une définition formelle des polynômes se basant uniquement sur ces derniers.

On fixe maintenant  $\mathbb{K}$  un corps commutatif, par exemple  $\mathbb{K} = \mathbb{R}$  ou  $\mathbb{C}$ .

*Remarque.* On peut définir une notion de polynôme sur n'importe quel anneau. Dans tout ce qui suit, certaines propriétés restent vraies lorsqu'on enlève la propriété de commutativité ou d'intégrité et d'autres deviennent fausses. C'est un très bon exercice de chercher où chaque hypothèse est effectivement utilisée.

### 5.1 Définition

**Notation 5.1.1.** Si  $x$  dénote une suite indexée par  $\mathbb{N}$ , pour tout  $k \in \mathbb{N}$ , on notera par  $x_k$  son  $k^{\text{ième}}$  terme.

**Définition 5.1.2.** On dit qu'une suite  $a$  est à support fini si elle est nulle à partir d'un certain rang, c'est-à-dire s'il existe  $n \in \mathbb{N}$  tel que  $a_k = 0$  pour tout  $k > n$ . On note  $\ell(\mathbb{K})$  l'ensemble des suites à

valeurs dans  $\mathbb{K}$ , et  $\ell_c(\mathbb{K}) \subset \ell(\mathbb{K})$  le sous-ensemble des suites à support fini. On définit également les deux opérations :

$$\begin{array}{ccc} \ell(\mathbb{K}) \times \ell(\mathbb{K}) & \longrightarrow & \ell(\mathbb{K}) \\ +: (a, b) & \longmapsto & (a_k + b_k)_{k \in \mathbb{N}} \end{array} \quad \cdot : \quad \begin{array}{ccc} \ell(\mathbb{K}) \times \ell(\mathbb{K}) & \longrightarrow & \ell(\mathbb{K}) \\ (a, b) & \longmapsto & \left( \sum_{i=0}^k a_i \cdot b_{k-i} \right)_{k \in \mathbb{N}} \end{array}.$$

*Remarque 5.1.3.* La définition du produit  $a.b$  peut sembler disymétrique, mais en faisant le changement de variable  $j = k - i$ , le  $k^{\text{ième}}$  terme de la suite  $a.b$  peut également s'écrire  $\sum_{j=0}^k a_{k-j} \cdot b_j$ . Et en remarquant que  $i + (k - i) = k$ , on peut même symétriser directement la formule en :

$$(a.b)_k = \sum_{\substack{i_1, i_2 \in \mathbb{N} \\ i_1 + i_2 = k}} a_{i_1} \cdot b_{i_2}.$$

**Proposition 5.1.4.** Soit  $a, b \in \ell_c(\mathbb{K})$ , alors  $a + b$  et  $a.b$  sont dans  $\ell_c(\mathbb{K})$ .

*Démonstration.* Puisque  $a$  et  $b$  sont à supports finis, notons  $n_a, n_b \in \mathbb{N}$  tels que  $a_k = 0$  pour  $k > n_a$  et  $b_k = 0$  pour  $k > n_b$ . Alors, pour tout  $k > \max(n_a, n_b)$ ,  $a_k + b_k = 0 + 0 = 0$ , et pour tout  $k > n_a + n_b$ ,

$$\sum_{i=0}^k a_i \cdot b_{k-i} = \sum_{i=0}^{n_a} a_i \cdot b_{k-i} + \sum_{i=n_a+1}^k a_i \cdot b_{k-i} = \sum_{i=0}^{n_a} a_i \cdot 0 + \sum_{i=n_a+1}^k 0 \cdot b_{k-i} = 0.$$

En effet, pour  $i > n_a$ ,  $a_i = 0$  et pour  $i \leq n_a$ , on a  $k - i > n_a + n_b - i \geq n_a + n_b - n_a = n_b$ , et donc  $b_{k-i} = 0$ .  $\square$

**Proposition 5.1.5.** Les opérations  $+$  et  $\cdot$  munissent  $\ell_c(\mathbb{K})$  d'une structure d'anneau commutatif dont le zéro est la suite 0, constante égale à zéro, et l'élément unité est la suite  $1 := (1, 0, 0, \dots)$  dont tous les termes valent zéro sauf celui d'indice 0 qui vaut un.

*Démonstration.* Il est clair que  $(\ell_c, +)$  est un groupe abélien dont la suite 0 est l'élément neutre.

Pour montrer l'associativité du produit, on considère  $a, b, c \in \ell_c(\mathbb{K})$ . Alors, pour tout  $k \in \mathbb{N}$ , on a

$$\begin{aligned} (a.(b.c))_k &= \sum_{\substack{i_1, i_2 \in \mathbb{N} \\ i_1 + i_2 = k}} a_{i_1} \cdot (b.c)_{i_2} = \sum_{\substack{i_1, i'_2, i''_2 \in \mathbb{N} \\ i_1 + i'_2 + i''_2 = k}} a_{i_1} \cdot b_{i'_2} \cdot c_{i''_2} \\ &= \sum_{\substack{i'_1, i''_1, i_2 \in \mathbb{N} \\ i'_1 + i''_1 + i_2 = k}} a_{i'_1} \cdot b_{i''_1} \cdot c_{i_2} = \sum_{\substack{i_1, i_2 \in \mathbb{N} \\ i_1 + i_2 = k}} (a.b)_{i_1} \cdot c_{i_2} \\ &= ((a.b).c)_k. \end{aligned}$$

Avant d'attaquer la distributivité, montrons la commutativité. Pour tous  $a, b \in \ell_c(\mathbb{K})$  et tout  $k \in \mathbb{N}$ , on a

$$(a.b)_k = \sum_{\substack{i_1, i_2 \in \mathbb{N} \\ i_1 + i_2 = k}} a_{i_1} \cdot b_{i_2} = \sum_{\substack{i_1, i_2 \in \mathbb{N} \\ i_1 + i_2 = k}} b_{i_1} \cdot a_{i_2} = (b.a)_k.$$

Concernant la distributivité, pour tous  $a, b, c \in \ell_c(\mathbb{K})$  et tout  $k \in \mathbb{N}$ , on a

$$\begin{aligned} (a.(b+c))_k &= \sum_{i=0}^k a_i \cdot (b+c)_{k-i} = \sum_{i=0}^k a_i \cdot (b_{k-i} + c_{k-i}) = \sum_{i=0}^k a_i \cdot b_{k-i} + a_i \cdot c_{k-i} \\ &= \sum_{i=0}^k a_i \cdot b_{k-i} + \sum_{i=0}^k a_i \cdot c_{k-i} = (a.b)_k + (a.c)_k, \end{aligned}$$

et donc, en utilisant la commutativité,  $(a+b).c = c.(a+b) = c.a + c.b = a.c + b.c$ .

Enfin, pour tout  $a \in \ell_c(\mathbb{K})$  et tout  $k \in \mathbb{N}$ , on a

$$(a.1)_k = (1.a)_k = \sum_{i=0}^k 1_i \cdot a_{k-i} = 1_0 \cdot a_k = a_k$$

et donc  $a.1 = 1.a = a$ . L'élément 1 est donc un élément unité.  $\square$

*Remarque 5.1.6.* De la même manière, on montre que  $(\ell(\mathbb{K}), +, \cdot)$  est un anneau commutatif, dont  $\ell_c(\mathbb{K})$  est un sous-anneau. Par ailleurs, on peut remarquer que  $\ell_0(\mathbb{K}) := \{a \in \ell(\mathbb{K}) \mid \forall k \in \mathbb{N}^*, a_k = 0\} \subset \ell(\mathbb{K})$  est un sous-anneau de  $\ell_c(\mathbb{K})$  qui s'identifie naturellement avec  $\mathbb{K}$ . Plus précisément, l'application  $\psi : \mathbb{K} \rightarrow \ell_0(\mathbb{K})$  qui envoie  $\lambda \in \mathbb{K}$  sur la suite  $(\lambda, 0, \dots)$  est un isomorphisme d'anneaux ; le résultat est clair au niveau de la somme et, concernant le produit, on vérifie directement que  $(\lambda_1, 0, \dots) \cdot (\lambda_2, 0, \dots) = (\lambda_1 \cdot \lambda_2, 0, \dots)$ . Cela permet d'utiliser la multiplication dans  $\ell(\mathbb{K})$  pour définir sur  $\ell(\mathbb{K})$  une notion de multiplication par un scalaire de  $\mathbb{K}$ , à savoir  $\lambda.a := \psi(\lambda).a$ , qui plus prosaïquement parlant, multiplie chaque terme de  $a$  par  $\lambda$ . Puisque cela ne crée pas d'ambiguïté, nous identifierons par la suite  $\mathbb{K}$  et  $\ell_0(\mathbb{K})$  et noterons  $\lambda.a$  pour  $\psi(\lambda).a$ .

**Notation 5.1.7.** On note  $X$  la suite  $(0, 1, 0, \dots)$  dont tous les termes sont nuls sauf celui d'indice 1 qui vaut 1.

**Proposition 5.1.8.** Pour tout  $k \in \mathbb{N}$ ,  $X^k$  est égal à la suite dont tous les termes valent 0 sauf celui d'indice  $k$  qui vaut 1. De fait, pour tout  $a \in \ell_c(\mathbb{K})$ , on a  $a = \sum_{i=0}^n a_i \cdot X^i$  avec  $n \in \mathbb{N}$  tel que  $a_k = 0$  pour tout  $k > n$ .

*Démonstration.* Montrons le résultat par récurrence sur  $k \in \mathbb{N}$ . Pour  $k = 0$ , on a bien (par convention)  $X^0 = 1$ . Supposons maintenant le résultat vrai au rang  $k$ , alors pour tout  $l \in \mathbb{N}$ , on a

$$(X^{k+1})_l = (X.X^k)_l = \sum_{i=0}^l X_i \cdot (X^k)_{l-i} = X_1 \cdot (X^k)_{l-1} = \begin{cases} 1 & \text{si } l-1 = k \text{ c\`ad si } l = k+1 \\ 0 & \text{sinon} \end{cases}.$$

La seconde partie de la proposition en découle directement par linéarité.  $\square$

**Notation 5.1.9.** Il est traditionnel de noter la suite non nulle  $a \in \ell_c(\mathbb{K})^*$  par  $a_0 + a_1X + \dots + a_nX^n$ , avec  $n \in \mathbb{N}$  tel que  $a_n \neq 0$  et  $a_k = 0$  pour tout  $k > n$ . Ici, la lettre  $X$  ne correspond donc ni à une variable, ni à une inconnue, mais juste à une notation pour un élément particulier, que l'on appelle souvent *indéterminée*. La suite nulle, quant à elle, est notée 0. En accord avec ces conventions, l'anneau  $(\ell_c(\mathbb{K}), +, \cdot)$  est noté  $\mathbb{K}[X]$ , et ses éléments sont appelés *polynômes (à coefficients dans  $\mathbb{K}$ )*. Il arrive que, pour certaines raisons,  $X$  soit remplacé par une autre notation, par exemple  $Y$ ,  $Z$  ou  $\alpha$ , il convient alors de modifier la notation  $\mathbb{K}[Y]$ ,  $\mathbb{K}[Z]$  ou  $\mathbb{K}[\alpha]$  en conséquence. Toujours inspiré par cette notation, le terme  $a_k$  est souvent appelé *coefficient (de degré  $k$ )* tandis que  $a_kX^k$  est appelé *terme (de degré  $k$ )*;  $a_0$  est aussi appelé *terme constant*.

*Remarque 5.1.10.* En tant qu'anneau commutatif intègre,  $\mathbb{K}[X]$  vérifie toutes les propriétés de ces derniers. Toutes les règles de calculs restent en particulier valables, notamment le binôme de Newton et les identités remarquables.

## 5.2 Différentes facettes algébriques de $\mathbb{K}[X]$

### 5.2.1 $\mathbb{K}[X]$ vu comme anneau euclidien

**Définition 5.2.1.** Pour tout  $P \in \mathbb{K}[X] \setminus \{0\}$ , on définit  $\deg(P)$  comme le plus grand degré des termes non nuls de  $P$ . Par convention, on pose  $\deg(0) = -\infty$ .

**Définition 5.2.2.**

- Pour tout polynôme non nul  $P \in \mathbb{K}[X]$ , on appelle respectivement *coefficient* et *terme dominant* le coefficient et le terme de degré  $\deg(P)$ , et on dit que le polynôme est *unitaire* si son coefficient dominant vaut 1.
- On appelle *polynôme constant* tout polynôme de degré 0 ou  $-\infty$ .

**Proposition 5.2.3.** Soit  $P_1, P_2 \in \mathbb{K}[X]$ , on a

- $\deg(P_1 + P_2) \leq \max(\deg(P_1), \deg(P_2))$  avec inégalité stricte si et seulement si  $P_1$  et  $P_2$  sont de même degré avec des coefficients dominants opposés ;
- $\deg(P_1 \cdot P_2) = \deg(P_1) + \deg(P_2)$ .

*Démonstration.* Quitte à les permuter par commutativité de la somme et du produit, on peut supposer  $\deg(P_1) \geq \deg(P_2)$ . Au sein de la preuve de la proposition 5.1.4, nous avons déjà montré que  $\deg(P_1 + P_2) \leq \max(\deg(P_1), \deg(P_2))$  et  $\deg(P_1 \cdot P_2) \leq \deg(P_1) + \deg(P_2)$ . Si  $P_2$  est nul, alors  $\deg(P_1 + P_2) = \deg(P_1) = \max(\deg(P_1), \deg(P_2))$  et  $\deg(P_1 \cdot P_2) = \deg(0) = -\infty = \deg(P_1) + \deg(P_2)$ . Sinon, on note  $(a_k)_{k \in \mathbb{N}}$  et  $(b_k)_{k \in \mathbb{N}}$  les coefficients de, respectivement,  $P_1$  et  $P_2$ . Par définition du degré, on a donc  $a_{\deg(P_1)}, b_{\deg(P_2)} \neq 0$ ,  $a_k = 0$  pour  $k > \deg(P_1)$  et  $b_k = 0$  pour  $k > \deg(P_2)$ . On en déduit

$$\begin{aligned}
 (P_1 \cdot P_2)_{\deg(P_1) + \deg(P_2)} &= \sum_{k=0}^{\deg(P_1) + \deg(P_2)} a_k \cdot b_{\deg(P_1) + \deg(P_2) - k} \\
 &= \left( \sum_{k=0}^{\deg(P_1) - 1} a_k \cdot b_{\deg(P_1) + \deg(P_2) - k} \right) + a_{\deg(P_1)} \cdot b_{\deg(P_2)} \\
 &\quad + \left( \sum_{k=\deg(P_1) + 1}^{\deg(P_1) + \deg(P_2)} a_k \cdot b_{\deg(P_1) + \deg(P_2) - k} \right) \\
 &= \left( \sum_{k=0}^{\deg(P_1) - 1} a_k \cdot 0 \right) + a_{\deg(P_1)} \cdot b_{\deg(P_2)} + \left( \sum_{k=\deg(P_1) + 1}^{\deg(P_1) + \deg(P_2)} 0 \cdot b_{\deg(P_1) + \deg(P_2) - k} \right) \\
 &= a_{\deg(P_1)} \cdot b_{\deg(P_2)} \neq 0.
 \end{aligned}$$

De fait  $\deg(P_1 \cdot P_2) \geq \deg(P_1) + \deg(P_2)$ , et donc  $\deg(P_1 \cdot P_2) = \deg(P_1) + \deg(P_2)$ .

Concernant la somme, le terme de degré  $\deg(P_1)$  de  $P_1 + P_2$  vaut  $a_{\deg(P_1)} + b_{\deg(P_1)}$ . Si celui-ci s'annule, alors  $\deg(P_1 + P_2) < \deg(P_1) = \max(\deg(P_1), \deg(P_2))$  ; s'il est non nul, alors  $\deg(P_1 + P_2) \geq \deg(P_1) = \max(\deg(P_1), \deg(P_2))$  et donc  $\deg(P_1 + P_2) = \max(\deg(P_1), \deg(P_2))$ .  $\square$

**Corollaire 5.2.4.** L'anneau  $\mathbb{K}[X]$  est intègre et les inversibles sont les polynômes constants non nuls, c'est-à-dire les polynômes de degré 0.

*Démonstration.* Si  $P_1, P_2 \in \mathbb{K}[X]^*$ , alors  $\deg(P_1), \deg(P_2) \geq 0$  et donc  $\deg(P_1 \cdot P_2) \geq 0$ , c'est-à-dire  $P_1 \cdot P_2 \neq 0$ , ce qui montre par contraposé que  $\mathbb{K}[X]$  est intègre.

Soit  $P \in \mathbb{K}[X]^\times$ . Il existe donc  $P^{-1} \in \mathbb{K}[X]^*$  tel que  $P \cdot P^{-1} = 1$ . Puisque  $P^{-1} \neq 0$ , on a  $\deg(P^{-1}) \geq 0$  et de fait  $\deg(P) \leq \deg(P) + \deg(P^{-1}) = \deg(P \cdot P^{-1}) = \deg(1) = 0$ . On en déduit que  $\deg(P) = 0$ . Réciproquement, il est clair que tout polynôme constant non nul est inversible.  $\square$

*Remarque 5.2.5.* Les éléments de  $\mathbb{K}[X]^* \setminus \mathbb{K}[X]^\times$  sont exactement les polynômes de degré 1 ou plus, c'est-à-dire de degré strictement positif.

**Corollaire 5.2.6.** La classe d'association de tout polynôme non nul possède un unique représentant unitaire ; autrement dit, pour tout  $P \in \mathbb{K}[X]^*$ , il existe un unique polynôme  $\tilde{P} \in \mathbb{K}[X]$  unitaire tel que  $P = \alpha \cdot \tilde{P}$  avec  $\alpha \in \mathbb{K}^*$ .

**Proposition 5.2.7.** L'anneau  $\mathbb{K}[X]$  est euclidien pour le stathme  $\nu: \mathbb{K}[X] \rightarrow \mathbb{N}$  définie par  $\nu := \max(1 + \deg, 0)$ .

Si  $\deg(P) < \deg(Q)$ , alors le résultat est vrai en posant  $S = 0$  et  $R = P$ . Supposons maintenant le résultat vrai en degré  $n \geq \deg(Q) - 1$  et considérons  $P$  de degré  $n + 1 \geq \deg(Q)$ . Alors,  $P$  et  $Q$  étant non nuls, on peut noter  $a_P, a_Q \in \mathbb{K}^*$  leurs coefficients dominants et considérer  $P$  et  $a_P \cdot a_Q^{-1} \cdot X^{\deg(P) - \deg(Q)} \cdot Q$  qui sont deux polynômes de même degré  $n + 1$  et de même coefficient dominant. D'après la proposition 5.2.3,  $\deg(P - a_P \cdot a_Q^{-1} \cdot X^{\deg(P) - \deg(Q)} \cdot Q) \leq n$  et, par hypothèse de récurrence, il existe donc  $R, \tilde{S}$  tels que  $P - a_P \cdot a_Q^{-1} \cdot X^{\deg(P) - \deg(Q)} \cdot Q = \tilde{S} \cdot Q + R$  et  $\nu(R) < \nu(Q)$ . Il ne reste plus qu'à poser  $S := a_P \cdot a_Q^{-1} \cdot X^{\deg(P) - \deg(Q)} + \tilde{S}$  pour conclure.  $\square$

*Remarque 5.2.8.* De cette preuve, on peut extraire une notion de division euclidienne pour les polynôme très similaire à la division des entiers telle que pratiquée dans les écoles primaires. En effet, la démonstration consiste à identifier ce par quoi il faut multiplier  $Q$  pour que son coefficient dominant soit égal à celui de  $P$ , de sorte à ce que, en retranchant ce produit à  $P$  on fasse diminuer le degré, et on répète cela jusqu'à obtenir un reste de degré strictement inférieur à celui de  $Q$ . Sur l'exemple  $P := 2X^3 - X + 1$  et  $Q := X^2 + 2X - 1$ , cela donne :

On pose alors  $P_2 := P - 2X.Q = 2X^3 - X + 1 - 2X.(X^2 + 2X - 1) = -4X^2 + X + 1$ .

On pose alors  $P_3 := P_2 + 4.Q = -4X^2 + X + 1 + 4.(X^2 + 2X - 1) = 9X - 3$ .

On a enfin  $\deg(9X - 3) = 1 < 2 = \deg(X^2 + 2X - 1)$ , et on en déduit que

$$P = P_2 + 2X.Q = P_3 - 4.Q + 2X.Q = (2X - 4).Q + (9X - 3).$$

Ecrit façon *école primaire*, cela donne :

[illegible]

Pour qu'un anneau soit euclidien, il n'est pas nécessaire, *a priori*, que les reste et quotient d'une division euclidienne soient uniques. C'est néanmoins le cas dans le cas d'un anneau de polynômes.

**Proposition 5.2.9.** Pour tout  $P \in \mathbb{K}[X]$  et  $Q \in \mathbb{K}[X]^*$ , les  $R, S \in \mathbb{K}[X]$  tels que  $P = S.Q + R$  et  $\deg(R) < \deg(Q)$  sont uniques.

*Démonstration.* Supposons que  $S_1.Q + R_1 = P = S_2.Q + R_2$  avec  $R_1, R_2, S_1, S_2 \in \mathbb{K}[X]$ ,  $Q \in \mathbb{K}[X]^*$  et  $\deg(R_1), \deg(R_2) < \deg(Q)$ . Alors  $Q.(S_2 - S_1) = R_1 - R_2$  et donc  $\deg(Q.(S_2 - S_1)) = \deg(R_1 - R_2) < \deg(Q)$ . Mais par ailleurs,  $\deg(Q.(S_2 - S_1)) = \deg(Q) + \deg(S_2 - S_1) \geq \deg(Q)$  sauf si  $\deg(S_2 - S_1) = -\infty$ . On en déduit donc que  $S_2 - S_1 = 0$ , c'est-à-dire  $S_2 = S_1$ , et donc  $R_1 = P - S_1.Q = P - S_2.Q = R_2$ .  $\square$

**Corollaire 5.2.10.** Soit  $P, Q \in \mathbb{R}[X] \subset \mathbb{C}[X]$  avec  $Q \neq 0$ . Alors les divisions euclidiennes de  $P$  par  $Q$  dans  $\mathbb{R}[X]$  et dans  $\mathbb{C}[X]$  donnent les mêmes résultats.

*Démonstration.* Soit  $R, S \in \mathbb{R}[X]$  donnés par la division euclidienne de  $P$  par  $Q$  dans  $\mathbb{R}[X]$ . Alors  $P = S.Q + R$  avec  $R, S \in \mathbb{C}[X]$  et  $\deg(R) < \deg(Q)$ . Par unicité, cela correspond donc également à la division euclidienne de  $P$  par  $Q$  dans  $\mathbb{C}[X]$ .  $\square$

### 5.2.2 $\mathbb{K}[X]$ vu comme anneau principal

Nous avons vu dans la section 4.3 que tout anneau euclidien est principal. En corollaire de la section précédente, on a donc :

**Proposition 5.2.11.** L'anneau  $\mathbb{K}[X]$  est principal.

En conséquence, on obtient une notion de multiple et de divisibilité, ainsi que :

**Définition 5.2.12.** Pour tous polynômes  $P_1, P_2 \in \mathbb{K}[X]^*$ , on définit  $\text{pgcd}(P_1, P_2)$  et  $\text{ppcm}(P_1, P_2)$  comme, respectivement, l'unique plus grand diviseur commun unitaire et l'unique plus petit multiple commun unitaire de  $P_1$  et  $P_2$ . Et on dit que  $P_1$  et  $P_2$  sont *premiers entre eux* si  $\text{pgcd}(P_1, P_2) = 1$ .

De la théorie générale, on obtient :

**Théorème 5.2.13.**

- Pour tout  $P_1, P_2 \in \mathbb{K}[X]^*$ , il existe  $Q_1, Q_2 \in \mathbb{K}[X]$  tels que  $P_1.Q_1 + P_2.Q_2 = \text{pgcd}(P_1, P_2)$ .
- Deux polynômes  $P_1, P_2 \in \mathbb{K}[X]^*$  sont premiers entre eux si et seulement si il existe  $Q_1, Q_2 \in \mathbb{K}[X]$  tels que  $P_1.Q_1 + P_2.Q_2 = 1$ .

*Remarque 5.2.14.* Comme pour les entiers, on peut définir un algorithme d'Euclide par divisions euclidiennes successives, lequel permet non seulement de déterminer le plus grand diviseur commun de deux polynômes non nuls, mais aussi une relation de Bézout entre eux.

*Remarque 5.2.15.* Conformément à la remarque 4.3.2, on peut remarquer directement que si  $P_1$  et  $P_2$  sont deux polynômes non nuls tels que  $P_1$  divise  $P_2$ , alors  $\deg(P_1) \leq \deg(P_2)$ .

### 5.2.3 $\mathbb{K}[X]$ vu comme anneau factoriel

Nous avons vu dans la section 4.2 que tout anneau principal est factoriel. En corollaire de la section précédente, on a donc :

**Proposition 5.2.16.** L'anneau  $\mathbb{K}[X]$  est factoriel.

En conséquence, on obtient le résultat suivant :

**Théorème 5.2.17.** A ordre des facteurs près, tout polynôme  $P \in \mathbb{K}[X]$  de degré strictement positif admet une unique décomposition

$$P = \alpha P_1 \cdots P_k$$

avec  $\alpha \in \mathbb{K}^*$  et  $P_1, \dots, P_k \in \mathbb{K}[X]$  des polynômes unitaires irréductibles.

Il est de fait intéressant de savoir quels polynômes sont irréductibles.

**Proposition 5.2.18.** Tout polynôme de degré 1 est irréductible.

*Démonstration.* Soit  $P \in \mathbb{K}[X]$  de degré 1 et  $P_1, P_2 \in \mathbb{K}[X]$  tels que  $P = P_1.P_2$ . Alors  $\deg(P_1), \deg(P_2) \geq 0$  car  $P_1, P_2 \neq 0$  puisque  $P \neq 0$ , et d'autre part  $\deg(P_1) + \deg(P_2) = \deg(P) = 1$ . On déduit que  $(\deg(P_1), \deg(P_2)) = (0, 1)$  ou  $(1, 0)$ , et dans les deux cas on a  $P_1$  ou  $P_2$  inversible.  $\square$

**Définition 5.2.19.** On dit qu'un polynôme non nul et non inversible est *scindé* si tous ses facteurs irréductibles sont de degré un.

Savoir s'il existe des polynômes irréductibles de degré strictement plus grand que 1, autrement dit savoir s'il existe des polynômes non nuls, non inversibles et non scindés, est une question dont la réponse dépend du corps  $\mathbb{K}$ .

**Définition 5.2.20.** On dit qu'un corps  $\mathbb{K}$  est *algébriquement clos* si tout polynôme de degré au moins un dans  $\mathbb{K}[X]$  est scindé, autrement dit si les polynômes de degré 1 sont les seuls à être irréductibles.

**Théorème 5.2.21.** Le corps  $\mathbb{C}$  est algébriquement clos.

*Démonstration.* Ce résultat ne sera prouvé qu'en appendice car sa preuve nécessite un peu d'analyse complexe, c'est-à-dire d'analyse pour les fonctions continues allant de  $\mathbb{C}$  dans  $\mathbb{C}$ .  $\square$

**Corollaire 5.2.22.** Les polynômes irréductibles de  $\mathbb{C}[X]$  sont exactement les polynômes de degré 1. En particulier, pour tout polynôme  $P \in \mathbb{C}[X]$  de degré strictement positif, il existe  $\alpha \in \mathbb{C}^*$  et  $z_1, \dots, z_{\deg(P)} \in \mathbb{C}$  tels que

$$P = \alpha(X - z_1) \cdots (X - z_{\deg(P)}).$$

**Théorème 5.2.23.** Les polynômes irréductibles de  $\mathbb{R}[X]$  sont exactement les polynômes de degré 1 et les polynômes de degré 2 de la forme  $a_0 + a_1X + a_2X^2$  avec  $a_1^2 < 4a_0a_2$ .

*Démonstration.* Considérons un polynôme  $P \in \mathbb{R}[X]$  irréductible et de degré au moins 2. Puisque  $\mathbb{R} \subset \mathbb{C}$ , on peut voir  $P$  comme un élément de  $\mathbb{C}[X]$ , et en tant que tel, il existe  $\alpha \in \mathbb{C}^*$  et  $z_1, \dots, z_k \in \mathbb{C}$  tels que  $P = \alpha(X - z_1) \cdots (X - z_k)$ . Mais en développant ce produit, on observe que  $\alpha$  est égal au coefficient dominant de  $P \in \mathbb{R}[X]$ , on a donc  $\alpha \in \mathbb{R}^*$ . De plus, si  $z_1$  était réel, alors d'après le corollaire 5.2.10, le reste de la division euclidienne de  $P$  par  $X - z_1$  dans  $\mathbb{R}[X]$  serait le même que dans  $\mathbb{C}[X]$ , à savoir 0, et  $P$  serait donc divisible par  $X - z_1$ , ce qui contredirait son irréductibilité. On en déduit que  $z_1 \in \mathbb{C} \setminus \mathbb{R}$ , et donc que  $\bar{z}_1 \neq z_1$ . Or, en notant  $\bar{Q} \in \mathbb{C}[X]$  le polynôme obtenu à partir de  $Q \in \mathbb{C}[X]$  en remplaçant chaque coefficient par son conjugué, on a

$$\alpha(X - \bar{z}_1) \cdots (X - \bar{z}_k) = \bar{\alpha}(X - \bar{z}_1) \cdots (X - \bar{z}_k) = \bar{P} = P = \alpha(X - z_1) \cdots (X - z_k),$$

puisque  $\alpha$  et tous les coefficients de  $P$  sont réels. Par unicité de la décomposition en facteurs irréductibles, on en déduit qu'il existe  $i \in \llbracket 2 \rrbracket k$  tel que  $\bar{z}_1 = z_i$  et, quitte à permuter les racines, on peut supposer  $z_2 = \bar{z}_1$ . On a donc  $P = \alpha(X - z_1)(X - \bar{z}_1) \prod_{k=3}^{\deg(P)} (X - z_k) = \alpha(X^2 - 2\Re(z_1)X + |z_1|^2) \prod_{k=3}^{\deg(P)} (X - z_k)$ . Or  $Q := \alpha(X^2 - 2\Re(z_1)X + |z_1|^2) \in \mathbb{R}[X]$  et, l'algorithme de division euclidienne de  $P$  par  $Q$  étant le même dans  $\mathbb{R}[X]$  et dans  $\mathbb{C}[X]$ , son résultat est le même et on en déduit que  $Q$  divise  $P$ . Cela implique que  $P = Q$  par irréductibilité de  $P$ . De plus, on a  $(-2\alpha\Re(z_1))^2 \leq (-2\alpha\Re(z_1))^2 + (-2\alpha\Im(z_1))^2 = 4\alpha\alpha|z_1|^2$ .

Réciproquement, considérons  $P = a_0 + a_1X + a_2X^2$  avec  $a_1^2 < 4a_0a_2$  et supposons par l'absurde qu'il n'est pas irréductible. Alors ses facteurs irréductibles sont nécessairement de degré 1, et on a  $P =$



$\alpha(X - \beta_1)(X - \beta_2)$  avec  $\alpha, \beta_1, \beta_2 \in \mathbb{R}$ . En développant, on trouve alors  $P = \alpha X^2 - \alpha(\beta_1 + \beta_2)X + \alpha\beta_1\beta_2$  et donc  $a_0 = \alpha\beta_1\beta_2$ ,  $a_1 = -\alpha(\beta_1 + \beta_2)$  et  $a_2 = \alpha$ . Mais alors

$$a_1^2 - 4a_0a_2 = \alpha^2(\beta_1 + \beta_2)^2 - 4\alpha^2\beta_1\beta_2 = \alpha^2(\beta_1^2 + 2\beta_1\beta_2 + \beta_2^2 - 4\beta_1\beta_2) = \alpha^2(\beta_1^2 - 2\beta_1\beta_2 + \beta_2^2) = \alpha^2(\beta_1 - \beta_2)^2 \geq 0,$$

ce qui contredit l'hypothèse de départ. On en déduit donc que  $P$  est irréductible.  $\square$

L'intérêt des polynômes irréductibles va au-delà de la simple décomposition en facteurs.

**Proposition 5.2.24.** Si  $P \in \mathbb{K}[X]$  est irréductible, alors  $\mathbb{K}[X]/(P)$  est un corps.

*Démonstration.* Soit  $\tilde{Q} \in \left(\mathbb{K}[X]/(P)\right)^*$ , alors tout représentant  $Q \in \mathbb{K}[X]$  de  $\tilde{Q}$  est premier avec  $P$ . En effet, si  $D \in \mathbb{K}[X]^* \setminus \mathbb{K}[X]^\times$  est un diviseur commun à  $P$  et  $Q$ , alors  $D = \alpha.P$  avec  $\alpha \in \mathbb{K}[X]^\times$  car  $P$  est irréductible, et donc  $P = \alpha^{-1}.D$  divise  $Q$ , ce qui implique que  $Q \in (P)$ , contredisant la non trivialité de  $\tilde{Q}$ . Par le théorème de Bachet-Bézout, il existe donc  $R, S \in \mathbb{K}[X]$  tels que  $R.Q + S.P = 1$ . On en déduit que  $1 - R.Q \in (P)$  et donc que  $1 - \tilde{R}.\tilde{Q} = 0$ , en notant  $\tilde{R}$  l'image de  $R$  dans  $\mathbb{K}[X]/(P)$ . Cela fournit un inverse  $\tilde{R}$  pour  $\tilde{Q}$ , qui est donc inversible.  $\square$

*Exemple 5.2.25.* En tant que corps, on a  $\mathbb{C} \cong \mathbb{R}[X]/(X^2 + 1)$ . On peut en effet considérer l'application

$$f: \begin{array}{ccc} \mathbb{R}[X] & \longrightarrow & \mathbb{C} \\ \sum_{k=0}^n a_k X^k & \longmapsto & \sum_{k=0}^n a_k i^k \end{array},$$

qui est un épimorphisme d'anneaux. Son noyau  $\text{Ker}(f)$  est un idéal, et comme  $\mathbb{R}[X]$  est principal, il est engendré par un polynôme  $P$  unitaire. Or, clairement,  $X^2 + 1 \in \text{Ker}(f) = (P)$ . Par irréductibilité de  $X^2 + 1$  dans  $\mathbb{R}[X]$ , on en déduit que  $P = X^2 + 1$  ou  $P = 1$ . Mais on ne peut pas avoir  $P = 1$ , car  $1 \notin \text{Ker}(f)$ . On a donc  $P = X^2 + 1$  et, par le théorème 3.4.12, on obtient que  $\mathbb{C} \cong \mathbb{R}[X]/(X^2 + 1)$ .

#### 5.2.4 $\mathbb{K}[X]$ vu comme un espace vectoriel

**Proposition 5.2.26.** La somme et la multiplication par un scalaire munissent  $\mathbb{K}[X]$  d'une structure de  $\mathbb{K}$ -espace vectoriel de dimension infinie.

*Démonstration.* Il a déjà été vérifié que  $(\mathbb{K}[X], +)$  est un groupe abélien. Les propriétés concernant le produit par un scalaire découlent, quant à elles, de la structure d'anneaux de  $\mathbb{K}[X]$ . Rappelons en effet que cette opération, introduite dans la remarque 5.1.6, est définie par  $\lambda.P := \psi(\lambda).P$  où  $\psi$  identifie, en tant qu'anneau,  $\mathbb{K}$  aux polynômes constants. On a alors bien, pour tous  $\lambda, \lambda_1, \lambda_2 \in \mathbb{K}$  et  $P, P_1, P_2 \in \mathbb{K}[X]$ ,

- $\lambda.(P_1 + P_2) := \psi(\lambda).(P_1 + P_2) = \psi(\lambda).P_1 + \psi(\lambda).P_2 =: \lambda.P_1 + \lambda.P_2$ ;
- $(\lambda_1 + \lambda_2).P := \psi(\lambda_1 + \lambda_2).P = (\psi(\lambda_1) + \psi(\lambda_2)).P = \psi(\lambda_1).P + \psi(\lambda_2).P =: \lambda_1.P + \lambda_2.P$ ;
- $\lambda_1.(\lambda_2.P) := \psi(\lambda_1).(\psi(\lambda_2).P) = (\psi(\lambda_1).\psi(\lambda_2)).P = \psi(\lambda_1.\lambda_2).P =: (\lambda_1.\lambda_2).P$ ;
- $\bar{1}.P := \psi(\bar{1}).P = 1.P = P$ , en notant  $\bar{1}$  l'élément unité de  $\mathbb{K}$ .

L'ensemble  $\mathbb{K}[X]$  est donc bien un espace vectoriel.

Concernant sa dimension, montrons que pour tout  $n \in \mathbb{N}$ ,  $\dim(\mathbb{K}[X]) \geq n + 1$ . Pour cela, on montre par récurrence sur  $n$  que la famille  $(1, X, X^2, \dots, X^n)$  est libre. Le résultat est évidemment vrai pour  $n = 0$ . Supposons-le donc vrai au rang  $n$  et considérons une relation  $\sum_{i=0}^{n+1} a_i X^i = 0$  avec  $a_0, \dots, a_{n+1} \in \mathbb{K}$ . Si  $a_{n+1}$  était non nul, on aurait  $X^{n+1} = a_{n+1}^{-1}.a_0 + a_{n+1}^{-1}.a_1 X + \dots + a_{n+1}^{-1}.a_n X^n$ , ce qui n'est pas possible car  $\deg(X^{n+1}) = n + 1$  et  $\deg(a_{n+1}^{-1}.a_0 + a_{n+1}^{-1}.a_1 X + \dots + a_{n+1}^{-1}.a_n X^n) \leq n$ . On en déduit que  $a_{n+1} = 0$  et donc que  $\sum_{i=0}^n a_i X^i = 0$ , ce qui par hypothèse de récurrence implique que tous les  $a_i$  sont nuls. La dimension de  $\mathbb{K}[X]$  en tant qu'espace vectoriel est donc infinie.  $\square$

*Remarque 5.2.27.* Dans la preuve ci-dessus, la famille infinie  $\mathcal{B} := (1, X, X^2, \dots)$  se comporte comme une base infinie, dans le sens que toutes les sous-familles finies sont libres, et que tout élément de  $\mathbb{K}[X]$  s'écrit comme combinaison linéaire (finie) d'éléments de  $\mathcal{B}$ . Il est alors tentant d'étendre tout cela à des combinaisons linéaires infinies, mais il faut pour cela donner un sens à ces sommes infinies. Nous ne traiterons pas de cela ici, mais cela peut se faire de manière abstraite, on parle alors de séries formelles, ou bien avec des considérations de convergence, on met alors le pied dans l'analyse.

Néanmoins, il est souvent agréable de pouvoir raisonner en dimension finie.

**Définition 5.2.28.** Pour tout  $n \in \mathbb{N}^*$ , on définit  $\mathbb{K}_n[X] := \{P \in \mathbb{K}[X] \mid \deg(P) \leq n\}$ .

**Proposition 5.2.29.** Pour tout  $n \in \mathbb{N}$ ,  $\mathbb{K}_n[X]$  est un sous-espace vectoriel de  $\mathbb{K}[X]$  de dimension  $n + 1$ .

*Démonstration.* D'après la proposition 5.2.3,  $\mathbb{K}_n[X]$  est stable par somme, opposé et produit par un scalaire. C'est donc bien un sous-espace vectoriel. Nous avons déjà montré dans la preuve de la proposition 5.2.26 que la famille  $(1, X, \dots, X^n)$  est libre et il est clair qu'elle est génératrice. Il s'agit donc d'une base de cardinal  $n + 1$ , on a donc bien  $\dim(\mathbb{K}_n[X]) = n + 1$ .  $\square$

*Remarques 5.2.30.*

- L'espace  $\mathbb{K}_n[X]$  peut être interprété algébriquement comme l'anneau quotient  $\mathbb{K}[X]/(X^{n+1})$ . En remarquant par ailleurs que  $(X^{n+1}) = \{P \in \mathbb{K}[X] \mid \deg(P) > n\} \cup \{0\}$  est un sous-espace vectoriel de  $\mathbb{K}[X]$ , que l'on notera  $\mathbb{K}_{>n}[X]$ , on peut observer la concomitance des trois opérations suivantes :
  - ▶ projection vectorielle de  $\mathbb{K}[X]$  sur  $\mathbb{K}_n[X]$  selon  $\mathbb{K}_{>n}[X]$  (point de vue “algèbre linéaire”);
  - ▶ projection canonique de  $\mathbb{K}[X]$  vers  $\mathbb{K}[X]/(X^{n+1})$ , (point de vue “algèbre”);
  - ▶ reste de la division euclidienne par  $X^{n+1}$ , (point de vue “arithmétique”).
- Dans le cas  $\mathbb{K} = \mathbb{R}$  (et aussi  $\mathbb{C}$ ), et en empiétant sur la section suivante, tout cela a également fort à voir avec l'analyse. En effet, d'après le théorème de Stone–Weierstrass, toute fonction continue  $f$  sur un compact de  $\mathbb{R}$  peut être approché par une suite  $(P_n)_{n \in \mathbb{N}}$  de polynômes<sup>6</sup>. En chaque degré, le coefficient correspondant de  $P_n$  converge alors lorsque  $n$  tend vers l'infini, et  $f$  peut en ce sens être (un peu abusivement) interprétée comme un “polynôme de degré potentiellement infini”. En projetant ce “polynôme” sur  $\mathbb{K}_n[X]$  par troncation des termes de degré plus grand que  $n$ , on obtient une suite approchant  $f$ ; dans le cas où  $f$  est analytique, c'est la suite de ses développements limités en 0. Pour  $f$  polynomiale, cela donne un quatrième point de vue, étiqueté “analyse”, pour la projection de  $\mathbb{K}[X]$  sur  $\mathbb{K}_n[X]$  évoquée ci-dessus.

### 5.3 Fonctions polynomiales et racines

Commençons par donner une définition purement algébrique de la notion de racine.

**Définition 5.3.1.** On dit que  $\alpha \in \mathbb{K}$  est *racine* de  $P \in \mathbb{K}[X]$  si  $P$  est divisible par  $X - \alpha$ .

De cette définition, il vient directement que :

**Proposition 5.3.2.** Tout polynôme  $P \in \mathbb{K}[X]^*$  non nul admet au plus  $\deg(P)$  racines distinctes.

*Démonstration.* Soit  $\alpha_1 \neq \alpha_2 \in \mathbb{K}$  deux racines de  $P$ . On a alors  $(\alpha_2 - \alpha_1)^{-1}(X - \alpha_1) + (\alpha_1 - \alpha_2)^{-1}(X - \alpha_2) = 1$ , et donc  $\text{pgcd}(X - \alpha_1, X - \alpha_2) = 1$ . Tous les monômes  $X - \alpha$ , avec  $\alpha$  racine de  $P$ , sont donc premiers entre eux et ils divisent tous  $P$ . D'après la proposition 4.2.10 leur produit divise  $P$  et son degré, correspondant au nombre de racines distinctes de  $P$ , est donc inférieur à celui de  $P$ .  $\square$

La définition 5.3.1 ne correspond pas exactement à l'idée, pourtant commune, qu'une racine d'un polynôme, c'est une valeur où le polynôme “s'annule”. Cette approche nécessite de donner un sens à une telle annulation; c'est ici que la distinction polynôme/fonction polynomiale intervient.

6. en vrai, leurs fonctions polynomiales  $f_{P_n}$  associées

**Définition 5.3.3.** Pour tout polynôme  $P =: \sum_{k=0}^n a_k X^k \in \mathbb{K}[X]$ , on définit l'*application polynomiale* associée par

$$f_P: \begin{array}{ccc} \mathbb{K} & \longrightarrow & \mathbb{K} \\ x & \longmapsto & \sum_{k=0}^n a_k x^k \end{array}.$$

**Proposition 5.3.4.** L'application

$$\xi_{\mathbb{K}}: \begin{array}{ccc} \mathbb{K}[X] & \longrightarrow & \mathcal{F}(\mathbb{K}, \mathbb{K}) \\ P & \longmapsto & f_P \end{array},$$

où  $\mathcal{F}(\mathbb{K}, \mathbb{K})$  est l'anneau des applications de  $\mathbb{K}$  dans  $\mathbb{K}$ , est un morphisme d'anneaux. Lorsque  $\mathbb{K} = \mathbb{R}$  ou  $\mathbb{C}$ , on peut même remplacer  $\mathcal{F}(\mathbb{K}, \mathbb{K})$  par l'anneau des applications continues, dérivables, infiniment dérivables ou même analytiques.

*Démonstration.* En notant  $P_1 := a_0 + a_1 X + \dots + a_{n_1} X^{n_1} \in \mathbb{K}[X]$ ,  $a_n := 0$  pour  $n > n_1$ ,  $P_2 := b_0 + b_1 X + \dots + b_{n_2} X^{n_2} \in \mathbb{K}[X]$  et  $b_n := 0$  pour  $n > n_2$ , on a bien, pour tout  $x \in \mathbb{K}$ ,

- $(\xi_{\mathbb{K}}(P_1 + P_2))(x) = \sum_{i=0}^{\max(n_1, n_2)} (a_i + b_i) x^i = \sum_{i=0}^{n_1} a_i x^i + \sum_{i=0}^{n_2} b_i x^i = (\xi_{\mathbb{K}}(P_1) + \xi_{\mathbb{K}}(P_2))(x)$ ;
- $(\xi_{\mathbb{K}}(P_1 \cdot P_2))(x) = \sum_{i=0}^{n_1+n_2} \left( \sum_{j=0}^i a_j b_{i-j} \right) x^i = \sum_{i=0}^{n_1+n_2} \sum_{j=0}^i a_j x^j \cdot b_{i-j} x^{i-j} = \left( \sum_{i=0}^{n_1} a_i x^i \right) \cdot \left( \sum_{i=0}^{n_2} b_i x^i \right) = (\xi_{\mathbb{K}}(P_1) \cdot \xi_{\mathbb{K}}(P_2))(x)$ ;
- $(\xi_{\mathbb{K}}(1))(x) = 1$ .

On en déduit que  $\xi_{\mathbb{K}}(P_1 + P_2) = \xi_{\mathbb{K}}(P_1) + \xi_{\mathbb{K}}(P_2)$ ,  $\xi_{\mathbb{K}}(P_1 \cdot P_2) = \xi_{\mathbb{K}}(P_1) \cdot \xi_{\mathbb{K}}(P_2)$  et  $\xi_{\mathbb{K}}(1) = 1$ , c'est-à-dire que  $\xi_{\mathbb{K}}$  est un morphisme d'anneaux.  $\square$

**Notation 5.3.5.** Par abus de notation, on notera dans la suite  $P(\alpha)$  pour  $f_P(\alpha)$  pour tout  $P \in \mathbb{K}[X]$  et  $\alpha \in \mathbb{K}$ .

**Proposition 5.3.6.** Soit  $P \in \mathbb{K}[X]$  et  $\alpha \in \mathbb{K}$ . Alors  $\alpha$  est racine de  $P$  si et seulement si  $P(\alpha) = 0$ .

*Démonstration.* Si  $\alpha$  est racine de  $P$ , alors il existe  $Q \in \mathbb{K}[X]$  tel que  $P = (X - \alpha)Q$  et on a  $P(\alpha) = (\alpha - \alpha)Q(\alpha) = 0$ .

Réciproquement, supposons que  $P(\alpha) = 0$ . Par division euclidienne de  $P$  par  $X - \alpha$ , il existe  $Q, R \in \mathbb{K}[X]$  tels que  $P = Q \cdot (X - \alpha) + R$  avec  $\deg(R) < \deg(X - \alpha) = 1$ . Le polynôme  $R$  est donc constant égal à  $\beta \in \mathbb{K}$ . Mais on a alors  $R = \beta = Q(\alpha)(\alpha - \alpha) + \beta = P(\alpha) = 0$ . On en déduit que  $P = Q \cdot (X - \alpha)$  et donc que  $X - \alpha$  divise  $P$ .  $\square$

**Corollaire 5.3.7.** Si  $\mathbb{K}$  est de cardinal infini, alors l'application  $\xi_{\mathbb{K}}$  est injective.

*Démonstration.* Soit  $P \in \text{Ker}(\xi_{\mathbb{K}})$ . Alors, pour tout  $x \in \mathbb{K}$ ,  $x$  est racine de  $P$ . Or, d'après la proposition 5.3.2, si  $P$  est non nul, il ne peut posséder qu'un nombre fini de racines. Par contraposée, on en déduit que  $P = 0$  et donc que  $\xi_{\mathbb{K}}$  est injective.  $\square$

*Remarque 5.3.8.* Dans le cas  $\mathbb{K} = \mathbb{R}$  ou  $\mathbb{C}$ , le corollaire 5.3.7 permet d'identifier les notions de polynôme et de fonction polynomiale. C'est un avantage, mais aussi un inconvénient car cela est susceptible de générer un certain nombre de confusions (notamment chez les étudiants). Pour un corps fini, ces deux notions diffèrent en effet nécessairement car  $\mathcal{F}(\mathbb{K}, \mathbb{K})$  est alors de cardinal fini mais pas  $\mathbb{K}[X]$ ; pour  $\mathbb{K} = \mathbb{Z}/2\mathbb{Z}$ , par exemple, l'application  $f_{X(X+1)}$  est triviale sans que  $X(X+1) = X + X^2$  ne le soit.

Plus généralement, on peut montrer que l'application  $\xi_{\mathbb{K}}$  est injective si et seulement si  $\mathbb{K}$  est infini, et qu'elle est surjective si et seulement si  $\mathbb{K}$  est fini. On remarquera de fait que  $\xi_{\mathbb{K}}$  n'est jamais un isomorphisme d'anneaux.

La proposition 5.3.2 peut être raffinée en comptant les racines "avec multiplicité".

**Lemme 5.3.9.** Soit  $P \in \mathbb{K}[X]$  et  $\alpha \in \mathbb{K}$  une racine de  $P$ . Alors il existe un unique  $k_0 \in \mathbb{N}^*$  tel que  $P$  soit divisible par  $(X - \alpha)^{k_0}$  mais pas par  $(X - \alpha)^{k_0+1}$ .

*Démonstration.* Si  $\alpha \in \mathbb{K}$  est racine de  $P \in \mathbb{K}[X]$ , c'est que  $X - \alpha$  divise  $P$ . De plus, si  $(X - \alpha)^{k+1}$  divise  $P$  avec  $k \in \mathbb{N}^*$ , alors  $(X - \alpha)^k$  aussi. Enfin, d'après la remarque 5.2.15,  $(X - \alpha)^{\deg(P)+1}$  ne peut pas diviser  $P$ . On en déduit que l'ensemble  $I_\alpha := \{k \in \mathbb{N}^* \mid (X - \alpha)^k \text{ divise } P\}$  est un ensemble fini non vide de la forme  $\{1, 2, \dots, k_0 - 1, k_0\}$  et que  $k_0 := \max I_\alpha$  est l'unique entier vérifiant les propriétés voulues.  $\square$

**Définition 5.3.10.** Soit  $P \in \mathbb{K}[X]$  et  $\alpha \in \mathbb{K}$  une racine de  $P$ . On appelle *multiplicité* de  $\alpha$  l'unique  $k \in \mathbb{N}^*$  tel que  $P$  soit divisible par  $(X - \alpha)^k$  mais pas par  $(X - \alpha)^{k+1}$ . Si  $\alpha$  est racine de multiplicité 1, on parle de *racine simple*, sinon on parle de *racine multiple*.

*Remarque 5.3.11.* D'après la proposition 5.3.6,  $\alpha$  est donc racine de  $P$  avec multiplicité  $k$  si et seulement si  $P = (X - \alpha)^k Q$  avec  $Q(\alpha) \neq 0$ .

**Proposition 5.3.12.** Tout polynôme  $P \in \mathbb{K}[X]^*$  non nul admet au plus  $\deg(P)$  racines, comptées avec multiplicité, c'est-à-dire

$$\sum_{\alpha \text{ racine de } P} k_\alpha \leq \deg(P),$$

où  $k_\alpha$  est la multiplicité de  $\alpha$  comme racine de  $P$ .

*Démonstration.* On va procéder comme pour la preuve de la proposition 5.3.2, mais pour cela, il faut montrer que, pour  $\alpha \neq \beta \in \mathbb{K}$  racines de  $P$ , les polynômes  $(X - \alpha)^{k_\alpha}$  et  $(X - \beta)^{k_\beta}$  sont premiers entre eux. Supposons donc par l'absurde qu'il existe un facteur commun  $Q$ , que l'on peut supposer irréductible, alors  $Q$  divise  $(X - \alpha)^{k_\alpha}$  et donc, en itérant le lemme d'Euclide, il divise  $X - \alpha$ . De même, il divise  $X - \beta$ . Or nous avons déjà montré que  $X - \alpha$  et  $X - \beta$  sont premiers entre eux, ils ne peuvent donc pas posséder de diviseur commun non trivial.  $\square$

Dans le cas des fonctions polynomiales réelles, les racines multiples se détectent grâce à l'annulation des dérivées successives. Ce principe a un pendant algébrique.

**Définition 5.3.13.** Pour tout  $P =: \sum_{k=0}^{\deg(P)} a_k X^k \in \mathbb{K}[X]$ , on définit la *dérivée* de  $P$  par

$$P' := \sum_{k=1}^{\deg(P)} k a_k X^{k-1}.$$

Pour tout  $n \in \mathbb{N}$ , on définit alors récursivement la  $n^{\text{ième}}$  dérivée de  $P$  par  $P^{(0)} := P$  lorsque  $n = 0$  et  $P^{(n)} = (P^{(n-1)})'$  lorsque  $n > 0$ .

*Remarque 5.3.14.* Cette définition est motivée par le fait que, pour tout  $P \in \mathbb{R}[X]$ ,  $f_{P'} = f'_P$ .

**Proposition 5.3.15.**

- Pour tout  $P \in \mathbb{K}[X]$ ,  $\deg(P') = \deg(P) - 1$  si  $\deg(P) \geq 1$  et  $\deg(P') = -\infty$  si  $\deg(P) \leq 0$ .
- Pour tous  $P_1, P_2 \in \mathbb{K}[X]$ , on a
  - $(P_1 + P_2)' = P'_1 + P'_2$ ;
  - $(P_1 \cdot P_2)' = P'_1 \cdot P_2 + P_1 \cdot P'_2$ .
- Pour tout  $P \in \mathbb{K}[X]$  et tout  $n \in \mathbb{N}^*$ , on a  $(P^n)' = n P' \cdot P^{n-1}$ . En particulier, pour tout  $\alpha \in \mathbb{K}$  et  $n \in \mathbb{N}^*$ , on a  $((X - \alpha)^n)' = n(X - \alpha)^{n-1}$ .

*Remarque 5.3.16.* L'application  $D: \begin{array}{ccc} \mathbb{K}[X] & \longrightarrow & \mathbb{K}[X] \\ P & \longmapsto & P' \end{array}$  est donc un endomorphisme d'espace vectoriel, mais pas un endomorphisme d'anneaux!

*Démonstration.* Le premier point se vérifie directement lorsque  $\deg(P) \leq 1$  et provient de ce que  $\deg(P) \cdot a_{\deg(P)} \neq 0$  lorsque  $\deg(P) \geq 1$ .

Montrons maintenant le second point. Soit  $P_1, P_2 \in \mathbb{K}[X]$  dont on note respectivement  $(a_k)_{k \in \mathbb{N}}$  et  $(b_k)_{k \in \mathbb{N}}$  les coefficients. Par définition, on a

$$\begin{aligned} P'_1 + P'_2 &= \sum_{k=1}^{\deg(P_1)} k a_k X^{k-1} + \sum_{k=1}^{\deg(P_2)} k b_k X^{k-1} = \sum_{k=1}^{\max(\deg(P_1), \deg(P_2))} k(a_k + b_k) X^{k-1} \\ &= \left( \sum_{k=0}^{\max(\deg(P_1), \deg(P_2))} (a_k + b_k) X^k \right)' = (P_1 + P_2)' \end{aligned}$$

ainsi que

$$\begin{aligned} P_1 \cdot P'_2 + P'_1 \cdot P_2 &= \left( \sum_{k=0}^{\deg(P_1)} a_k X^k \right) \cdot \left( \sum_{k=1}^{\deg(P_2)} k b_k X^{k-1} \right) + \left( \sum_{k=1}^{\deg(P_1)} k a_k X^{k-1} \right) \cdot \left( \sum_{k=0}^{\deg(P_2)} b_k X^k \right) \\ &= \left( \sum_{k=0}^{\deg(P_1)} a_k X^k \right) \cdot \left( \sum_{k=0}^{\deg(P_2)-1} (k+1) b_{k+1} X^k \right) + \left( \sum_{k=0}^{\deg(P_1)-1} (k+1) a_{k+1} X^k \right) \cdot \left( \sum_{k=0}^{\deg(P_2)} b_k X^k \right) \\ &= \sum_{k=0}^{\deg(P_1)+\deg(P_2)-1} \left( \sum_{\substack{i_1, i_2 \in \mathbb{N} \\ i_1+i_2=k}} a_{i_1} (i_2+1) b_{i_2+1} \right) X^k + \sum_{k=0}^{\deg(P_1)+\deg(P_2)-1} \left( \sum_{\substack{i_1, i_2 \in \mathbb{N} \\ i_1+i_2=k}} (i_1+1) a_{i_1+1} b_{i_2} \right) X^k \\ &= \sum_{k=0}^{\deg(P_1)+\deg(P_2)-1} \left( \sum_{\substack{i_1, i_2 \in \mathbb{N} \\ i_1+i_2=k+1}} i_2 a_{i_1} b_{i_2} \right) X^k + \sum_{k=0}^{\deg(P_1)+\deg(P_2)-1} \left( \sum_{\substack{i_1, i_2 \in \mathbb{N} \\ i_1+i_2=k+1}} i_1 a_{i_1+1} b_{i_2} \right) X^k \\ &= \sum_{k=0}^{\deg(P_1)+\deg(P_2)-1} \sum_{\substack{i_1, i_2 \in \mathbb{N} \\ i_1+i_2=k+1}} (i_2 + i_1) a_{i_1} b_{i_2} X^k = \sum_{k=0}^{\deg(P_1)+\deg(P_2)-1} \sum_{\substack{i_1, i_2 \in \mathbb{N} \\ i_1+i_2=k+1}} (k+1) a_{i_1} b_{i_2} X^k \\ &= \sum_{k=0}^{\deg(P_1)+\deg(P_2)-1} (k+1) \left( \sum_{\substack{i_1, i_2 \in \mathbb{N} \\ i_1+i_2=k+1}} a_{i_1} b_{i_2} \right) X^k = \sum_{k=1}^{\deg(P_1)+\deg(P_2)} k \left( \sum_{\substack{i_1, i_2 \in \mathbb{N} \\ i_1+i_2=k}} a_{i_1} b_{i_2} \right) X^{k-1} \\ &= (P_1 \cdot P_2)' \end{aligned}$$

Le dernier point provient d'une récurrence immédiate.  $\square$

**Proposition 5.3.17.** Soit  $P_1, P_2 \in \mathbb{K}[X]$  tels que  $P'_1 = P'_2$ , alors il existe  $a \in \mathbb{K}$  tel que  $P_2 = P_1 + a$ .

*Démonstration.* Notons  $P_2 - P_1 =: \sum_{k=0}^n a_k X^k$ . On a alors

$$\sum_{k=1}^n k a_k X^{k-1} = (P_2 - P_1)' = P'_2 - P'_1 = 0.$$

Mais, la famille  $(1, X, \dots, X^{n-1})$  étant libre, on peut en déduire que  $k a_k = 0$  pour tout  $k \in \llbracket 1, n \rrbracket$ , et donc  $a_k = 0$ . En posant  $a := a_0$ , on obtient bien  $P_2 = P_1 + a$ .  $\square$

**Corollaire 5.3.18** (formule de Taylor). Pour tout  $P \in \mathbb{K}[X]^*$  et tout  $\alpha \in \mathbb{K}$ , on a  $P = \sum_{k=0}^{\deg(P)} \frac{P^{(k)}(\alpha)}{k!} (X - \alpha)^k$ .

*Démonstration.* On travaille par récurrence généralisée sur  $\deg(P) \in \mathbb{N}$ . Le résultat est clair si  $\deg(P) = 0$ , c'est-à-dire si  $P$  est un polynôme constant. Supposons maintenant le résultat vrai au rang  $n$  et considérons un polynôme  $P \in \mathbb{K}[X]$  de degré  $n + 1$ . Le polynôme  $P'$  est alors de degré au plus  $n$  et, par hypothèse de récurrence, on a donc  $P' = \sum_{k=0}^{\deg(P)-1} \frac{P^{(k+1)}(\alpha)}{k!} (X - \alpha)^k$ . Mais alors

$$\left( \sum_{k=0}^{\deg(P)} \frac{P^{(k)}(\alpha)}{k!} (X - \alpha)^k \right)' = \sum_{k=1}^{\deg(P)} \frac{P^{(k)}(\alpha)}{(k-1)!} (X - \alpha)^{k-1} = \sum_{k=0}^{\deg(P)-1} \frac{P^{(k+1)}(\alpha)}{k!} (X - \alpha)^k = P',$$

et donc, d'après la proposition précédente,  $\sum_{k=0}^{\deg(P)} \frac{P^{(k)}(\alpha)}{k!} (X - \alpha)^k = P + a$  avec  $a \in \mathbb{K}$ . En évaluant en  $\alpha$ , on obtient enfin  $a = 0$ .  $\square$

**Proposition 5.3.19.** Soit  $P \in \mathbb{K}[X]$ . Un élément  $\alpha \in \mathbb{K}$  est racine de multiplicité  $k_\alpha \in \mathbb{N}^*$  si et seulement si  $P^{(k)}(\alpha) = 0$  pour tout  $k \in \llbracket 0, k_\alpha - 1 \rrbracket$  mais  $P^{(k_\alpha)}(\alpha) \neq 0$ .

*Démonstration.* En utilisant le corollaire 5.3.18, il est clair que, si  $P^{(k)}(\alpha)$  s'annule en  $\alpha$  pour tout  $k \in \llbracket 0, k_0 - 1 \rrbracket$  mais pas pour  $k = k_0$ , alors  $\alpha$  est racine de multiplicité  $k_0$  pour  $P$ .

Pour la réciproque, on commence par remarquer que, sous la convention qu'une racine de multiplicité 0 est une non racine et que  $\llbracket 0, -1 \rrbracket = \emptyset$ , le résultat est vrai pour  $k_\alpha = 0$ . De plus, on peut remarquer que, si  $\alpha \in \mathbb{K}$  est racine de multiplicité  $k_\alpha \in \mathbb{N}^*$  de  $P$ , alors  $P = (X - \alpha)^{k_\alpha} Q$  avec  $Q \in \mathbb{K}[X]$  vérifiant  $Q(\alpha) \neq 0$ , et donc  $P' = k_\alpha (X - \alpha)^{k_\alpha - 1} Q + (X - \alpha)^{k_\alpha} Q' = (X - \alpha)^{k_\alpha - 1} \tilde{Q}$  avec  $\tilde{Q} := (k_\alpha Q + (X - \alpha) Q')$  vérifiant  $\tilde{Q}(\alpha) = k_\alpha Q(\alpha) \neq 0$ . Autrement dit,  $\alpha$  est racine de multiplicité  $k_\alpha - 1$  de  $P'$ . La preuve se fait dès lors par récurrence sur  $k_\alpha \in \mathbb{N}$ . Supposons donc le résultat vrai au rang  $n \in \mathbb{N}$  et considérons une racine  $\alpha$  de multiplicité  $k_\alpha = n + 1$  pour  $P$ . Alors  $\alpha$  est racine de multiplicité  $n$  pour  $P'$  et, par hypothèse de récurrence, on a  $P^{(k)}(\alpha) = (P')^{(k-1)}(\alpha) = 0$  pour tout  $k \in \llbracket 1, n \rrbracket$ , ainsi que  $P^{(n+1)}(\alpha) = (P')^{(n)}(\alpha) \neq 0$ . De plus  $P^{(0)} = P(\alpha) = 0$  car  $\alpha$  est racine de  $P$  puisque  $n + 1 \geq 1$ .  $\square$

## Appendice : Théorème de d'Alembert–Gauss

En supposant acquise la notion de continuité pour les fonctions allant de  $\mathbb{C}$  dans  $\mathbb{C}$ , nous donnons ici une preuve du théorème suivant :

**Théorème (d'Alembert–Gauss).** Tout polynôme  $P \in \mathbb{C}[X]$  non constant admet une racine.

*Remarque.* Cela peut être reformulé en :

- tout polynôme  $P \in \mathbb{C}[X]^* \setminus \mathbb{C}[X]^\times$  est scindé ;
- seuls les polynômes de degré 1 sont irréductibles dans  $\mathbb{C}[X]$  ;
- $\mathbb{C}$  est algebriquement clos.

*Démonstration.* Soit  $P =: \sum_{i=0}^n a_i X^i \in \mathbb{C}[X]$  de degré  $n \in \mathbb{N}^*$ . Par inégalité triangulaire, on a pour tout  $z \in \mathbb{C}$

$$|P(z)| = \left| \sum_{i=0}^n a_i z^i \right| \leq \sum_{i=0}^n |a_i| \cdot |z|^i = |a_n| \cdot |z|^n \left( 1 + \sum_{i=0}^{n-1} \frac{|a_i|}{|a_n| \cdot |z|^{n-i}} \right) \xrightarrow{z \rightarrow \infty} \lim_{z \rightarrow \infty} |a_n| \cdot |z|^n = \infty.$$

En particulier, il existe  $R > 0$  tel que  $|z| > R \Rightarrow |P(z)| \geq |P(0)| + 1$ . On en déduit que  $\min_{\substack{z \in \mathbb{C} \\ |z| \leq R}} |P(z)|$  et que, par compacité de l'ensemble  $\{z \in \mathbb{C} \mid |z| \leq R\}$  et par continuité des applications polynomiales sur  $\mathbb{C}$ , le minimum de  $|P|$  est atteint en un point  $z_0$ . Supposons par l'absurde que  $P(z_0) \neq 0$ .

D'après le corollaire 5.3.18, il existe  $a'_0, \dots, a'_n \in \mathbb{C}$ , avec  $a'_0 = P(z_0) \neq 0$ , tels que

$$P(z) = \sum_{i=0}^n a'_i (z - z_0)^i.$$

Puisque  $P$  est de degré  $n \geq 1$ , l'ensemble  $\{1 \leq i \leq n \mid a'_i \neq 0\}$  est non vide, il possède donc un plus petit élément que l'on note  $i_0$ . On fixe également  $\omega \in \mathbb{C}$  une racine  $i_0^{\text{ième}}$  de  $-\frac{a'_0}{a'_{i_0}} \neq 0$ . Encore par inégalité triangulaire, on a alors pour tout  $0 < t < 1$

$$\begin{aligned} |P(z_0 + t\omega)| &= \left| a'_0 + \sum_{i=i_0}^n a'_i t^i \omega^i \right| = \left| a'_0 + a'_{i_0} t^{i_0} \omega^{i_0} + \sum_{i=i_0+1}^n a'_i t^i \omega^i \right| = \left| a'_0 - a'_0 t^{i_0} + \sum_{i=i_0+1}^n a'_i t^i \omega^i \right| \\ &\leq (1 - t^{i_0}) |a'_0| + \sum_{i=i_0+1}^n |a'_i \omega^i| t^i = (1 - t^{i_0}) |a'_0| + t^{i_0} \varepsilon(t), \end{aligned}$$

avec  $\varepsilon(t) = \sum_{i=1}^{n-i_0} |a'_{i+i_0} \omega^{i+i_0}| t^i \xrightarrow[t \rightarrow 0]{} 0$ . Pour  $t$  suffisamment proche de 0, on a alors  $0 \leq \varepsilon(t) \leq \frac{1}{2} |a'_0|$  et donc

$$|P(z_0 + t\omega)| \leq (1 - t^{i_0}) |a'_0| + \frac{t^{i_0}}{2} |a'_0| \leq \left(1 - \frac{t^{i_0}}{2}\right) |a'_0| < |a'_0| = |P(z_0)|,$$

ce qui contredit la minimalité de  $z_0$ . On en déduit que  $P$  s'annule en  $z_0$ . □