



Faculté des Sciences  
Département de Mathématique

# Mathématiques élémentaires

(du point de vue universitaire)

Pierre Mathonet

Première année d'études de bachelier en Sciences Mathématiques  
Année académique 2019-2020

# Introduction

Le cours “mathématiques élémentaires” est introduit lors de l’année académique 2014-2015 au début de la filière d’études en mathématiques pour faciliter la transition secondaire-université et ainsi favoriser la réussite en première année.

Ses objectifs sont multiples : il s’agit d’abord de revoir certains points de l’enseignement mathématique qui ont été abordés ou utilisés dans votre cursus antérieur, mais sans toujours être approfondis. On adoptera un point de vue assez souvent différent de celui qui a été utilisé auparavant, l’accent étant mis sur la logique et les articulations entre les notions plutôt que sur les méthodes de calcul, qui ne seront pourtant pas oubliées.

Un second objectif important du cours est d’apprendre à rédiger, à comprendre et à étudier des textes mathématiques, en s’exerçant sur des notions qui ont déjà été rencontrées, plutôt que sur des contenus complètement nouveaux.

Quand on entreprend l’étude des mathématiques supérieures, on est assez vite confronté à des problèmes de définitions : en effet, les notions mathématiques introduites dans l’enseignement secondaire reposent souvent sur des idées intuitives, et il est difficile, voire impossible, qu’il en soit autrement. Cependant, quand on veut élever l’édifice mathématique, on doit être sûr de ses fondations, et les idées intuitives, si elles peuvent encore guider la démarche, doivent être remplacées par des définitions claires et cohérentes et par une démarche déductive logique. On part donc d’un nombre minimal de définitions et on tente de démontrer toutes les propriétés des objets que l’on étudie.

Cependant, certaines notions sont encore trop ardues pour être abordées dans ce cours de première année ou n’y trouvent pas naturellement leur place. C’est le cas de la théorie des ensembles dans sa forme actuelle. Nous n’aborderons que la théorie dite naïve des ensembles et ne donnerons pas une définition formelle des ensembles. C’est également le cas de la construction des nombres réels, qui fait souvent intervenir des notions de convergence, et a plus sa place dans un cours d’analyse. Nous conviendrons donc dans ce cours que les exemples que nous prenons sur les nombres réels sont basés sur la compréhension intuitive que vous en avez en sortant de l’enseignement secondaire. En particulier, nous utiliserons les nombres réels pour bâtir les nombres complexes et étudier leurs propriétés.

Dans le premier chapitre, nous rencontrerons tout d’abord la logique et quelques éléments de théorie naïve des ensembles et nous en déduirons quelques techniques de démonstration qui seront utiles dans tous les cours du cursus. Nous reverrons ensuite les nombres complexes, et nous en profiterons pour approfondir l’utilisation du symbole sommatoire.

Nous reverrons ensuite la construction des nombres naturels, entiers (relatifs) et rationnels. Nous élargirons les constructions pour découvrir des structures algébriques omniprésentes dans les mathématiques supérieures, que nous illustrerons par des éléments d’arithmétique modulaire.

Si le temps le permet, nous en profiterons pour revoir des notions importantes d’analyse et de géométrie, avec toujours le souci d’apprendre à rédiger, à comprendre et à étudier des écrits mathématiques, des plus simples aux plus complexes.

J’espère sincèrement que ce cours vous aidera à bien entamer votre cursus universitaire et je vous invite à commencer la lecture sans attendre.

# Chapitre 1

## Logique et ensembles

Ce chapitre présente une introduction à la logique et à la théorie des ensembles. Vous avez déjà rencontré la plupart des notions abordées ici, mais souvent sans les étudier pour elles-mêmes.

Comme dans tous les cours de mathématique, je commence par quelques définitions formelles. Ensuite, nous verrons quelques procédures systématiques qui permettent de n'oublier aucun cas de figure quand on est face à un problème logique, ou faisant appel à des ensembles ; c'est le cas des tables de vérité ou des diagrammes de Venn. Nous passerons ensuite en revue les techniques de démonstration classiques comme la contraposition ou le raisonnement par l'absurde, nous verrons comment démontrer une alternative et nous terminerons par les démonstrations par récurrence.

### 1.1 Logique

Nous commençons par donner une définition des assertions logiques, et des connecteurs qui permettent de former des assertions composées à l'aide d'assertions élémentaires.

**Définition 1.1.1.** On appelle *assertion* ou *proposition logique* toute phrase d'un langage donné dont on peut envisager sans ambiguïté le problème de sa vérité ou de sa fausseté.

Par exemple, on peut considérer les assertions suivantes.

1. "Aujourd'hui, je porte un pull rouge";
2. "3 est un nombre premier";
3. "3 n'est pas divisible par 2";
4. "tout nombre positif est pair";
5. "Il pleut";
6. "J'emporte un parapluie";
7. "Si Berlin est en Suisse, alors je viens de Mars";
8. "Si mon chat aboie, alors je gagne au lotto".

Les assertions sont construites de façon à être compréhensibles sans ambiguïté pour que l'on puisse décider si elles sont vraies ou fausses. Elles admettent donc des valeurs de vérité "vrai" et "faux" que l'on note aussi  $V$  et  $F$  ou 1 et 0.

Les phrases suivantes ne sont donc pas des propositions logiques.

1. "Quelle heure est-il ?"
2. "Paris est-elle la capitale de la France ?"
3. "Cette phrase est fausse."
4. "Je cornifle gauche bien."

En effet, les deux premières sont des questions. Les propositions logiques correspondantes pourraient être “Il est 15 heures”, “Paris n’est pas la capitale de la France”. La troisième est contradictoire, puisque si elle est vraie, elle doit alors être fausse et vice-versa. Enfin, la quatrième n’a pas de sens.

Les règles formatives (ou règles de syntaxe) permettent de construire de nouvelles propositions à partir d’anciennes. L’idée est qu’on déclare dans ces règles qui est une proposition. On commence par admettre qu’il existe des propositions élémentaires dites *propositions atomiques* ou *variables propositionnelles*, notées  $p, q, r, \dots$ , puis on donne les règles stipulant qu’on peut en former de nouvelles, en utilisant les opérations logiques de négation, conjonction (et), disjonction (ou), implication ou bi-implication, encore appelées connecteurs logiques.

A ce point, on peut former des propositions logiques composées, mais on ne peut pas encore décider de la vérité de telles propositions, en fonction de la vérité ou la fausseté des propositions atomiques qui les composent. Cette étude s’appelle la *sémantique*.

On résout ce problème en associant à chaque connecteur une table de vérité qui précise la vérité de toute proposition logique composée à l’aide de ce connecteur. Une assertion composée a alors des valeurs de vérité qui dépendent des valeurs de vérité des assertions qui la composent.

**Définition 1.1.2.** Si  $P$  est une assertion, alors la négation de  $P$  est une assertion. On la note  $\neg P$ . Cette assertion est vraie si  $P$  est fausse et elle est fausse si  $P$  est vraie. La table de vérité de l’opérateur de négation  $\neg$  est donc la suivante.<sup>a</sup>

|     |          |           |     |          |
|-----|----------|-----------|-----|----------|
| $P$ | $\neg P$ |           | $P$ | $\neg P$ |
| 0   | 1        | ou encore | $F$ | $V$      |
| 1   | 0        |           | $V$ | $F$      |

*Remarque 1.1.* Dans la suite, nous adopterons les valeurs 0 pour faux, et 1 pour vrai, mais vous pouvez conserver  $V$  et  $F$  si c’est plus concret pour vous.

Voici quelques exemples qui sont les négations des assertions introduites plus haut.

1. “Aujourd’hui, je ne porte pas un pull rouge”;
2. “3 n’est pas un nombre premier”;
3. “3 est divisible par 2”;
4. “Il existe un nombre positif qui n’est pas pair”;<sup>b</sup>
5. “Il ne pleut pas”;
6. “Je n’emporte pas de parapluie”;
7. “Mon chat aboie et je ne gagne pas au lotto”.

Comme dans le langage courant, nier deux fois revient à ne rien faire. On pourrait écrire  $\neg(\neg P) = P$ , quelle que soit l’assertion  $P$ . Arrêtons-nous un instant sur cette égalité. En effet,  $\neg(\neg P)$  et  $P$  sont des assertions qui ne sont pas écrites de la même manière. On touche ici une définition importante.

**Définition 1.1.3.** Deux propositions logiques  $P$  et  $Q$  (composées à partir de propositions élémentaires) sont *logiquement équivalentes*<sup>c</sup> si elles ont les mêmes tables de vérité (i.e. les mêmes valeurs de vérité, dans tous les cas). On note alors  $P \equiv Q$ .

a. Voyez la construction de la table : la première colonne donne les deux valeurs possibles pour  $P$ , la deuxième donne les valeurs correspondantes de  $\neg P$ .

b. Il est important de remarquer que la négation de  $P$  : “tout nombre positif est pair” **n’est pas** “tout nombre positif est impair”, nous y reviendrons.

c. On dit aussi tautologiquement équivalentes.

Dans le cas de la double négation, on a bien

| $P$ | $\neg P$ | $\neg(\neg P)$ |
|-----|----------|----------------|
| 0   | 1        | 0              |
| 1   | 0        | 1              |

et donc on peut noter  $P \equiv \neg(\neg P)$ , et dire que  $P$  et  $\neg(\neg P)$  sont logiquement équivalentes. Il est important de noter que  $P$  et  $\neg(\neg P)$  ont les mêmes valeurs de vérité, quel que soit le contexte et quel que soit  $P$ .

Dans une expression complexe, on peut toujours remplacer une assertion par une assertion logiquement équivalente sans changer la valeur de vérité globale. Dans le cas présent, on peut remplacer l'assertion  $\neg(\neg P)$  par l'assertion  $P$ . Cela peut sembler fort théorique, mais vous pouvez utiliser ce fait dans le langage courant pour simplifier des phrases compliquées.

**Exemple 1.1.1.** La phrase

“Il n’est pas impossible que ce cours ne soit pas dépourvu de concepts nouveaux.”

est logiquement équivalente à

“Il est possible que ce cours contienne des concepts nouveaux.”

Voici maintenant deux connecteurs logiques bien connus dans la vie de tous les jours, le *et* et le *ou*. Il n’y a pas de grande surprise pour le “et”. Pour le “ou”, il faut juste noter qu’il n’est pas exclusif : en mathématiques, si on vous dit “tu peux avoir pour dessert une glace ou une coupe de fruit” vous pouvez répondre “d’accord, je mangerai les deux”.

**Définition 1.1.4.** Si  $P$  et  $Q$  sont deux assertions, alors la conjonction de  $P$  et  $Q$ , notée  $P \wedge Q$  ou “ $P$  et  $Q$ ” est une assertion qui est vraie quand  $P$  est vraie et  $Q$  est vraie (simultanément) et fausse sinon. La table de vérité du connecteur “et” est donc <sup>d</sup>

| $P$ | $Q$ | $P$ et $Q$ |
|-----|-----|------------|
| 0   | 0   | 0          |
| 0   | 1   | 0          |
| 1   | 0   | 0          |
| 1   | 1   | 1          |

On peut ainsi former les assertions

1. “Il pleut et je porte un pull rouge”;
2. “J’emporte un parapluie et 3 est un nombre premier”.

Il est intéressant de remarquer que la valeur de vérité de  $P \wedge Q$  est le minimum des valeurs de vérités de  $P$  et  $Q$ . Cela permet de faciliter les calculs, et cela justifie l’utilisation de 0 et 1, plutôt que  $F$  et  $V$ .

De la même manière on définit la disjonction de deux assertions.

**Définition 1.1.5.** Si  $P$  et  $Q$  sont deux assertions, alors la disjonction de  $P$  et  $Q$ , notée  $P \vee Q$  ou “ $P$  ou  $Q$ ” est une assertion qui est vraie quand au moins l’une des deux assertions  $P$ ,  $Q$  est vraie et qui est fausse sinon. Sa table de vérité est donc la suivante <sup>e</sup>.

| $P$ | $Q$ | $P$ ou $Q$ |
|-----|-----|------------|
| 0   | 0   | 0          |
| 0   | 1   | 1          |
| 1   | 0   | 1          |
| 1   | 1   | 1          |

d. Ici, les deux premières colonnes permettent d’avoir les quatre valeurs possibles pour le couple  $(P, Q)$ .

e. Notez la différence, dans la dernière ligne de la table, avec le “ou exclusif” souvent utilisé dans le langage courant.

Vous aurez sans doute remarqué que puisque  $P$  et  $Q$  peuvent prendre chacun deux valeurs de vérité, la table contient quatre lignes. Combien y aura-t-il de lignes pour des assertions composées de  $P$ ,  $Q$  et  $R$ , ou encore de  $P$ ,  $Q$ ,  $R$  et  $S$ ? Voici quelques exemples simples.

1. “Il pleut ou je porte un pull rouge”;
2. “J’emporte un parapluie ou 3 est un nombre premier”.

Ici aussi, on peut remarquer que la valeur de vérité de  $P \vee Q$  est le maximum des valeurs de vérités de  $P$  et  $Q$ .

Il est intéressant pour la suite de nos développements de déjà regarder quelques façons de construire des assertions logiques composées avec les trois connecteurs que nous avons vus jusqu’à présent.

**Proposition 1.1.1.** *On a les équivalences logiques suivantes*

1.  $\neg(P \vee Q) \equiv (\neg P) \wedge (\neg Q)$ ;
2.  $\neg(P \wedge Q) \equiv (\neg P) \vee (\neg Q)$ ;

Que veulent dire ces équivalences logiques sur des exemples ?

- La négation de “Il pleut ou je porte un pull rouge” est “il ne pleut pas **et** je ne porte pas de pull rouge”.
- La négation de “Il pleut et nous sommes mardi” est “il ne pleut pas **ou** nous ne sommes pas mardi”.
- Comment obtenir dans  $\mathbb{R}$  une condition équivalente à  $x^2 - 2x \neq 0$  ?
- Comment obtenir dans  $\mathbb{R}^2$  une condition équivalente à  $x^2(y - 4) \geq 0$  ?

Passons maintenant aux implications et bi-implications, ces dernières étant encore appelées équivalences. La définition de l’implication est la première de ce chapitre qui soit un peu surprenante. L’implication est le “si...alors” du français. Pour bien comprendre la table de vérité ci-dessous, on peut se demander quand l’assertion “Si on est vendredi, alors je porte un pull rouge” est fausse. Si vous voulez faire une expérience scientifique pour mettre en défaut cette affirmation, vous ne viendrez certainement pas voir la couleur de mes vêtements un jeudi. Vous viendrez le vendredi. L’assertion est fausse uniquement dans le cas où on est vendredi et où je n’ai pas de pull rouge.

**Définition 1.1.6.** Si  $P$  et  $Q$  sont deux assertions, alors “ $P$  implique  $Q$ ” est une assertion. On la note  $P \Rightarrow Q$ . Elle est toujours vraie sauf si  $P$  est vrai et  $Q$  faux. La table de vérité du connecteur  $\Rightarrow$  est donc

| $P$ | $Q$ | $P \Rightarrow Q$ |
|-----|-----|-------------------|
| 0   | 0   | 1                 |
| 0   | 1   | 1                 |
| 1   | 0   | 0                 |
| 1   | 1   | 1                 |

On peut également dire “si  $P$ , alors  $Q$ ” pour indiquer  $P \Rightarrow Q$ . Il est important de remarquer les deux premières lignes de la table de vérité. Quand  $P$  est faux, alors  $P \Rightarrow Q$  est vrai. Voici quelques exemples :

1. “S’il pleut alors j’emporte un parapluie.”<sup>f</sup>
2. “Si on est vendredi, je porte un pull rouge.”<sup>g</sup>
3. “Si 3 est un nombre premier, alors je porte un pull rouge.”

A titre d’exercice, on pourra démontrer l’équivalence logique entre  $P \Rightarrow Q$  et  $(\neg P) \vee Q$ .

Finalement, on peut définir la bi-implication entre de deux assertions, encore appelée équivalence.

f. Cette implication ne donne aucune indication s’il ne pleut pas.

g. On peut aussi dire “Tous les vendredis, je porte un pull rouge.”

**Définition 1.1.7.** Si  $P$  et  $Q$  sont deux assertions alors “ $P$  bi-implique  $Q$ ”, “ $P$  est équivalent à  $Q$ ” est une assertion. On la note  $P \Leftrightarrow Q$ . Elle est vraie quand  $P$  implique  $Q$  et  $Q$  implique  $P$  sont vraies. La table de vérité du connecteur  $\Leftrightarrow$  est donc

| $P$ | $Q$ | $P \Leftrightarrow Q$ |
|-----|-----|-----------------------|
| 0   | 0   | 1                     |
| 0   | 1   | 0                     |
| 1   | 0   | 0                     |
| 1   | 1   | 1                     |

Vous remarquerez que  $P \Leftrightarrow Q$  est vraie exactement quand  $P$  et  $Q$  ont la même valeur de vérité. Si  $P$  est équivalent à  $Q$ , on dira aussi que  $P$  est vrai si et seulement si  $Q$  est vrai. Voici quelques exemples

1. “J’emporte un parapluie si et seulement si il pleut”;
2. “Je porte un pull rouge si et seulement si on est vendredi”.

Remarquons que, par définition, nous avons  $P \Leftrightarrow Q \equiv (P \Rightarrow Q) \text{ et } (Q \Rightarrow P)$ .

Terminons cette liste d’opérations logiques en y ajoutant les deux *quantificateurs* : Le signe  $\forall$  se lit “pour tout” et le signe  $\exists$  se lit “il existe”. Ainsi, si  $P$  et  $Q$  sont deux assertions, on peut écrire

$$\forall x : P, \exists y : Q$$

pour signifier “pour tout  $x$  tel que  $P$ , il existe un  $y$  tel que  $Q$ ”.

L’ordre des quantificateurs a de l’importance. En effet, dans l’assertion précédente,  $y$  peut varier en fonction de  $x$ , ce qui n’est pas le cas si on l’écrit dans l’autre sens. Par exemple, le lecteur conviendra que les assertions commençant par

“Pour tout garçon dans la salle, il existe une fille dans la salle telle que...”

et

“Il existe une fille dans la salle telle que pour tout garçon dans la salle...”

n’auront sans doute pas les mêmes significations.<sup>h</sup> Remarquez également que dans les expressions ci-dessus, les mots “pour tout” et “il existe” n’ont pas été remplacés par les symboles correspondants, qui ne devraient être utilisés que dans des expressions purement mathématiques (des “formules”). Voici un exemple qui vous sera utile dans le cours d’analyse. C’est la définition de la convergence d’une suite vers 0.

**Exemple 1.1.2.** La convergence d’une suite numérique vers 0 :

$$\forall \varepsilon > 0, \exists N \in \mathbb{N} : n \geq N \Rightarrow |x_n| < \varepsilon.$$

Il est à noter qu’il y a un petit abus d’écriture dans cette notation, on devrait écrire

$$\forall \varepsilon > 0, \exists N \in \mathbb{N} : \forall n \in \mathbb{N}, (n \geq N \Rightarrow |x_n| < \varepsilon).$$

Mais en général, on écrira pas le quantificateur  $\forall n \in \mathbb{N}$ . On peut par exemple utiliser cette définition pour démontrer que la suite définie par  $x_n = \frac{1}{n}$  pour tout  $n \geq 1$  converge (tend) vers 0 quand  $n$  tend vers l’infini. On constatera que dans la démonstration, pour chaque  $\varepsilon$ , on peut trouver un  $N$  qui convient. Cela ne pourrait pas être le cas si on avait écrit une autre condition en inversant les quantificateurs, comme ceci :

$$\exists N \in \mathbb{N} : \forall \varepsilon > 0, \forall n \in \mathbb{N}, (n \geq N \Rightarrow |x_n| < \varepsilon).$$

Vous aurez le temps de vous familiariser avec ces expressions, mais j’insiste déjà sur le fait qu’il ne faut pas se contenter de les traduire mot à mot, mais bien essayer de se représenter ce qu’elles veulent dire.

<sup>h</sup>. Les personnes choquées par cette dernière assertion pourront échanger les mots “fille” et “garçon”, et s’interroger sur ses implications morales.

On peut également se demander quelle est la négation de propositions contenant des quantificateurs. Sans entrer dans les détails, notons que la négation d'une proposition contenant  $\forall$  s'exprime avec un  $\exists$  et vice-versa.

### Exemple 1.1.3.

1. La négation de "Tous les profs de math sont petits" est "Il existe un prof de math qui n'est pas petit".
2. La négation de "Il existe un cheval de course bon marché" est "Tous les chevaux de course coûtent cher".

**Exercice 1.1.1.** 1. Nier l'assertion  $\forall x > 0, \exists y > 0 : y < x$ .

2. Ecrivez l'assertion indiquant qu'une suite numérique ne converge pas vers 0.

En utilisant ces opérations logiques, on peut construire des assertions de plus en plus compliquées. Il faut être prudent et utiliser les parenthèses de la manière habituelle pour signifier l'ordre dans lequel il faut interpréter les connecteurs logiques. Par exemple,  $P \wedge Q \vee R$  n'a pas de sens, car on pourrait l'interpréter de deux façons différentes :  $(P \wedge Q) \vee R$  et  $P \wedge (Q \vee R)$ . Ces deux assertions ne sont pas logiquement équivalentes, comme le prouve le tableau de vérité suivant<sup>i</sup>.

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

En effet, la cinquième et la dernière colonne ne sont pas égales. Le tableau donne aussi des cas précis où les assertions sont différentes. Par exemple, la deuxième ligne correspond à un cas où  $P$  et  $Q$  sont faux et  $R$  vrai, on voit que dans ce cas,  $(P \wedge Q) \vee R$  est vraie et  $P \wedge (Q \vee R)$  est fausse.

Par contre, les assertions  $P \Rightarrow (Q \vee R)$  et  $(P \Rightarrow Q) \vee R$  sont logiquement équivalentes, comme le prouve le tableau suivant.

| $P$ | $Q$ | $R$ | $Q \vee R$ | $P \Rightarrow (Q \vee R)$ | $P \Rightarrow Q$ | $(P \Rightarrow Q) \vee R$ |
|-----|-----|-----|------------|----------------------------|-------------------|----------------------------|
| 0   | 0   | 0   | 0          | 1                          | 1                 | 1                          |
| 0   | 0   | 1   | 1          | 1                          | 1                 | 1                          |
| 0   | 1   | 0   | 1          | 1                          | 1                 | 1                          |
| 0   | 1   | 1   | 1          | 1                          | 1                 | 1                          |
| 1   | 0   | 0   | 0          | 0                          | 0                 | 0                          |
| 1   | 0   | 1   | 1          | 1                          | 0                 | 1                          |
| 1   | 1   | 0   | 1          | 1                          | 1                 | 1                          |
| 1   | 1   | 1   | 1          | 1                          | 1                 | 1                          |

**Exercice 1.1.2.** Etudier la validité des équivalences logiques suivantes :

1.  $(P \vee Q) \vee R \equiv P \vee (Q \vee R)$
2.  $(P \wedge Q) \wedge R \equiv P \wedge (Q \wedge R)$
3.  $P \wedge (Q \vee R) \equiv (P \wedge Q) \vee (P \wedge R)$
4.  $P \vee (Q \wedge R) \equiv (P \vee Q) \wedge (P \vee R)$

i. Les trois premières colonnes du tableau donnent toutes les valeurs de vérité possibles pour le triplet  $(P, Q, R)$ . De plus on calcule facilement les autres colonnes, en considérant une colonne à la fois.

On a également des relations entre les différentes constructions possibles. Ces relations sont données par des équivalences logiques entre certaines assertions. En voici un exemple classique.

**Proposition 1.1.2.** *Si  $P$  et  $Q$  sont deux assertions, l'assertion  $P \Rightarrow Q$  est logiquement équivalente à l'assertion  $\neg Q \Rightarrow \neg P$ . En d'autres termes, on a,*

$$(P \Rightarrow Q) \equiv (\neg Q \Rightarrow \neg P).$$

En conséquence, démontrer l'assertion  $\neg Q \Rightarrow \neg P$  est équivalent à démontrer  $P \Rightarrow Q$ . Cette technique de preuve s'appelle la *contraposition* et l'assertion  $\neg Q \Rightarrow \neg P$  est la *contraposée* de l'assertion  $P \Rightarrow Q$ .

*Démonstration.* Démontrons cette proposition en utilisant des tables de vérité. On calcule successivement les valeurs de vérités des assertions en question en fonction de tous les cas possibles pour  $P$  et  $Q$ . On obtient le tableau suivant

| $P$ | $Q$ | $\neg P$ | $\neg Q$ | $P \Rightarrow Q$ | $\neg Q \Rightarrow \neg P$ |
|-----|-----|----------|----------|-------------------|-----------------------------|
| 0   | 0   | 1        | 1        | 1                 | 1                           |
| 0   | 1   | 1        | 0        | 1                 | 1                           |
| 1   | 0   | 0        | 1        | 0                 | 0                           |
| 1   | 1   | 0        | 0        | 1                 | 1                           |

On constate que les deux dernières colonnes sont égales. D'après notre définition de l'équivalence, cela veut dire que ces assertions sont logiquement équivalentes.  $\square$

On peut également se convaincre intuitivement que les assertions “Si on est vendredi, je porte un pull rouge” et “Si je ne porte pas de pull rouge, on n'est pas vendredi” veulent dire la même chose, c'est-à-dire, sont logiquement équivalentes.

Terminons cette introduction élémentaire à la logique par les *tautologies*. Voici un exemple. Si  $P$  et  $Q$  sont deux assertions, on peut calculer la table de vérité de l'assertion

$$((P \Rightarrow Q) \text{ et } P) \Rightarrow Q.$$

On a directement

| $P$ | $Q$ | $P \Rightarrow Q$ | $(P \Rightarrow Q) \text{ et } P$ | $((P \Rightarrow Q) \text{ et } P) \Rightarrow Q$ |
|-----|-----|-------------------|-----------------------------------|---------------------------------------------------|
| 0   | 0   | 1                 | 0                                 | 1                                                 |
| 0   | 1   | 1                 | 0                                 | 1                                                 |
| 1   | 0   | 0                 | 0                                 | 1                                                 |
| 1   | 1   | 1                 | 1                                 | 1                                                 |

L'assertion  $((P \Rightarrow Q) \text{ et } P) \Rightarrow Q$  est donc toujours vraie, dans tous les cas de figure pour  $P$  et  $Q$ . C'est une **tautologie**.

**Définition 1.1.8.** Une assertion composée qui est vraie quelles que soient les valeurs de vérités des assertions qui la composent est une *tautologie*.

La notion de tautologie permet de formaliser d'un point de vue logique des techniques de démonstration. Voici premier exemple : dire que  $P$  et  $Q$  sont logiquement équivalentes, c'est dire que l'assertion  $P \Leftrightarrow Q$  est une tautologie. Voici maintenant quelques raisonnements couramment utilisés.

Voici quelques exemples supplémentaires et techniques de démonstration basées sur des équivalences logiques.

### 1.1.1 La contraposition : quelques exemples

Je vous rappelle la contraposition (proposition 1.1.2).

**Proposition 1.1.3.** *Si  $P$  et  $Q$  sont deux assertions, l'assertion  $P \Rightarrow Q$  est logiquement équivalente à l'assertion  $\neg Q \Rightarrow \neg P$ . En d'autres termes, on a,*

$$(P \Rightarrow Q) \equiv (\neg Q \Rightarrow \neg P).$$

D'après cette proposition, démontrer que  $P$  implique  $Q$  est équivalent à démontrer que  $\neg Q$  implique  $\neg P$ .

Voici deux exemples.

**Proposition 1.1.4.** *Si  $a, b \in \mathbb{R}$  sont tels que  $a + b$  est irrationnel, alors  $a$  ou  $b$  est irrationnel.*

*Démonstration.* Si on note  $P$  l'assertion “ $a + b$  est irrationnel” et  $Q$  l'assertion “ $a$  ou  $b$  est irrationnel”, on cherche bien à démontrer que  $P$  implique  $Q$ . Il est plus facile de démontrer que  $\neg Q$  implique  $\neg P$ . L'assertion  $\neg Q$  est “ $a$  et  $b$  sont rationnels”, tandis que l'assertion  $\neg P$  est “ $a + b$  est rationnel”. Alors  $\neg Q \Rightarrow \neg P$  est une assertion vraie par définition des rationnels.  $\square$

Bien sûr, j'ai rédigé la preuve pour faire apparaître le lien avec la proposition précédente. On rédigera plutôt comme ceci.

*Démonstration.* Procédons par contraposition et supposons que  $a$  et  $b$  soient rationnels. Alors  $a + b$  est rationnel, par définition des rationnels.  $\square$

**Proposition 1.1.5.** *Si  $n$  est un nombre entier tel que  $n^2$  est pair, alors  $n$  est pair.*

*Démonstration.* On démontre la proposition contraposée. Supposons que  $n$  soit impair. Alors, il existe  $k \in \mathbb{Z}$  tel que  $n = 2k + 1$ . On calcule alors  $n^2 = (2k + 1)^2 = 4k^2 + 4k + 1$  et on constate que  $n^2$  est impair.  $\square$

Avec l'habitude, on n'indique même plus qu'il s'agit d'une démonstration par contraposée, et cela n'inquiète pas le lecteur expérimenté.

### 1.1.2 La démonstration par l'absurde

Cette méthode de démonstration est très utilisée. Elle consiste à supposer que ce que l'on veut démontrer est faux et à arriver à une absurdité (encore appelée une contradiction), c'est à dire la conjonction d'une assertion et de sa négation. Ce type de démonstration est également basé sur une équivalence logique.

**Proposition 1.1.6.** *Pour toutes assertions  $P$  et  $Q$ , l'assertion  $P$  est logiquement équivalente à  $(\neg P) \Rightarrow (Q \wedge (\neg Q))$ .*

*Démonstration.* La preuve est immédiate, à l'aide de tables de vérités.  $\square$

Voici un exemple classique, qui peut être énoncé facilement quand on connaît les nombres réels : “Le nombre  $\sqrt{2}$  est irrationnel”. J'énonce ce résultat en ne présupposant pas que  $\sqrt{2}$  existe, puisque nous ne l'avons pas défini.

**Proposition 1.1.7.** *Il n'existe pas de nombre rationnel  $z$  tel que  $z^2 = 2$ .*

*Démonstration.* Nous allons supposer que l'assertion à démontrer est fausse, et montrer que cela conduit à une contradiction. On suppose donc qu'il existe un nombre rationnel  $z$  satisfaisant  $z^2 = 2$ . Par définition des rationnels<sup>j</sup>, il existe des nombres entiers  $p$  et  $q$  ( $q$

j. Prenez la définition de l'enseignement secondaire, si vous voulez, elle est équivalente à celle que je vous donnerai.

non nul), tels que  $z = \frac{p}{q}$ . On peut supposer que  $q$  est le nombre positif et minimal<sup>k</sup> tel que  $z = \frac{p}{q}$ . On a alors  $2 = \frac{p^2}{q^2}$ , ou encore  $p^2 = 2q^2$ .

Alors  $p^2$  est pair, et par la proposition 1.1.5,  $p$  est pair. Il existe alors  $r \in \mathbb{Z}$  tel que  $p = 2r$ . On a alors  $4r^2 = p^2 = 2q^2$ , qui donne  $q^2 = 2r^2$ . Alors  $q$  est pair : il existe  $s \in \mathbb{N}$ , tel que  $q = 2s$ . Puisque  $q$  n'est pas nul, on a  $s < q$ , et visiblement  $z = \frac{p}{q} = \frac{2r}{2s} = \frac{r}{s}$ . Donc  $q$  n'est pas minimal, cela donne la contradiction.  $\square$

Il n'existe des dizaines de preuves de la proposition précédente. On peut également trouver d'autres contradictions ou d'autres façons d'obtenir la contradiction utilisée ici.

### 1.1.3 Contre-exemple et démonstration d'une alternative

La proposition suivante, que l'on démontre sans difficulté, permet de justifier l'emploi du contre-exemple.

**Proposition 1.1.8.** *On a l'équivalence logique  $\neg(P \Rightarrow Q) \equiv P \wedge (\neg Q)$ .*

Il arrive souvent que l'on doive démontrer une assertion qui s'exprime comme une disjonction (un "ou"). On a une technique simple qui permet d'avoir une hypothèse en plus à sa disposition. Cette technique est basée sur le résultat suivant.

**Proposition 1.1.9.** *On a les équivalences logiques suivantes :*

$$P \Rightarrow (Q \vee R) \equiv (P \wedge (\neg Q)) \Rightarrow R \quad \text{et} \quad P \Rightarrow (Q \vee R) \equiv (P \wedge (\neg R)) \Rightarrow Q.$$

*Démonstration.* Démontrons la première équivalence. La deuxième se démontre de manière analogue. On a

$$P \Rightarrow (Q \vee R) \equiv (\neg P) \vee (Q \vee R) \equiv ((\neg P) \vee Q) \vee R \equiv \neg(P \wedge (\neg Q)) \vee R,$$

et la dernière assertion est logiquement équivalente à  $(P \wedge (\neg Q)) \Rightarrow R$ .  $\square$

Voici un exemple simple, non mathématique : pour prouver l'assertion "si je vis sur mars, alors je suis bleu ou j'ai quatre bras", on peut prouver "si je vis sur mars et si je ne suis pas bleu alors j'ai quatre bras".

### 1.1.4 La disjonction des cas

Il s'agit encore ici de prouver une implication où l'une des deux assertions contient une alternative (une disjonction "ou"). Il ne faut pas confondre avec la section précédente : l'alternative est ici avant l'implication.

**Proposition 1.1.10.** *On a l'équivalence  $(P \text{ ou } Q) \Rightarrow R \equiv (P \Rightarrow R) \text{ et } (Q \Rightarrow R)$ .*

La démonstration peut être faite à l'aide de tables de vérité. Je la laisse comme exercice. Cette proposition montre que pour démontrer  $(P \text{ ou } Q) \Rightarrow R$ , il faut traiter tous les cas. Voici un exemple simple.

**Proposition 1.1.11.** *Si  $n$  est entier naturel, alors  $n(n+1)$  est pair.*

*Démonstration.* Si  $n$  est un nombre naturel, alors il est pair ou impair. Nous pouvons donc démontrer que si  $n$  est pair ou impair, alors  $n(n+1)$  est pair. Si  $n$  est pair, alors il existe  $k \in \mathbb{N}$  tel que  $n = 2k$ . On a alors  $n(n+1) = 2k(2k+1)$  et ce nombre est pair. Si  $n$  est impair, alors il existe  $k \in \mathbb{N}$  tel que  $n = 2k+1$ . On a alors  $n(n+1) = (2k+1)(2k+2) = 2(2k+1)(k+1)$  et ce nombre est pair également.  $\square$

---

k. Ce sera intéressant de le montrer à l'aide de la définition de  $\mathbb{Q}$  que je vous donnerai, et de la propriété du bon ordre de  $\mathbb{N}$ , que nous verrons sous peu également.

## 1.2 Théorie des ensembles

Passons maintenant à la description des ensembles, avec une première définition.

**Définition 1.2.1.** Un *ensemble* est une collection d'objets possédant une ou plusieurs propriétés communes<sup>a</sup>. Ces objets sont les *éléments* ou *points* de l'ensemble.

On notera généralement un ensemble par une lettre majuscule.

Les *éléments* peuvent par exemple être donnés

1. de manière explicite, par des symboles tels que  $1, 2, 3, a, b \dots$  ;
2. par un symbole générique affecté par un ou plusieurs indices,  $x_i$  où  $i$  est un élément quelconque d'un autre ensemble.

Un ensemble peut être donné

1. de manière explicite, en donnant tous ses éléments, (définition en *extension*) comme par exemple  $A = \{1, 2, 3, 4\}$  ou  $B = \{a, b, c, d, e\}$  ;
2. de manière explicite, mais sans donner tous les éléments, que l'on peut remplacer par des points de suspension, comme  $C = \{1, 2, \dots, 100\}$ , ou encore  $D = \{a, b, c, \dots, z\}$ .
3. en décrivant la propriété caractérisant ses éléments, (définition en *compréhension*) comme dans

$$\{n : n \text{ est entier, pair et compris entre 1 et 99}\}.$$

En général, si  $P$  est une assertion, on désigne par  $\{x : P\}$  ou par  $\{x|P\}$  l'ensemble des objets  $x$  pour lesquels la propriété  $P$  est vérifiée.

Passons maintenant aux propriétés et aux relations entre éléments et ensembles.

1. **Ensemble vide** : il existe un ensemble qui ne contient pas d'éléments, l'ensemble vide, noté  $\emptyset$ .
2. **Appartenance** : on écrit  $x \in A$  ( $x$  appartient à  $A$ ) pour signifier que  $x$  est un élément de l'ensemble  $A$ .
3. **Inclusion** : on écrit  $B \subset A$  ( $B$  est inclus dans  $A$ , ou  $B$  est un sous-ensemble de  $A$ ) quand tout élément de  $B$  est aussi un élément de  $A$ . Dans ce cas, on dit que  $B$  est un sous-ensemble de  $A$ . L'ensemble vide est un sous-ensemble de tout ensemble donné.

Par exemple, si  $B = \{2, 4, 6, 8\}$ ,  $A = \{n : n \text{ est un nombre entier pair}\}$ , on a  $B \subset A$ .

4. **Égalité** : on écrit  $A = B$  ( $A$  et  $B$  sont égaux) quand les ensembles  $A$  et  $B$  ont les mêmes éléments. Cela se traduit aussi par le fait que  $A \subset B$  et  $B \subset A$ .

Les inclusions et l'égalité s'expriment également en termes d'implications : on a  $A \subset B$  si l'implication  $x \in A \Rightarrow x \in B$  est vraie, quel que soit l'objet  $x$  considéré. De même, on a  $A = B$  si l'équivalence  $x \in A \Leftrightarrow x \in B$  est vraie, quel que soit l'objet  $x$  considéré.

Enfin, on peut nier ces implications et écrire par exemple  $x \notin A$ ,  $B \not\subset A$  ou  $A \neq B$ , et par un léger abus de langage, on pourra écrire et lire ces symboles dans l'autre sens :  $A \ni a$ ,  $A \supset B$  ou  $A \not\supset a$ .

Soient  $A$  et  $B$  deux ensembles donnés, on peut construire les ensembles suivants :

1. **Union** : l'ensemble  $A \cup B$  est formé par les éléments qui appartiennent à  $A$  ou à  $B$ . On a donc, d'un point de vue logique

$$(x \in A \cup B) \equiv ((x \in A) \text{ ou } (x \in B)).$$

2. **Intersection** : l'ensemble  $A \cap B$  est formé par les éléments qui appartiennent à  $A$  et  $B$ . On a donc, d'un point de vue logique

$$(x \in A \cap B) \equiv ((x \in A) \text{ et } (x \in B)).$$

---

a. Ce n'est pas une définition extrêmement rigoureuse puisque le terme "collection" n'a pas été défini.

3. **Différence** : l'ensemble  $A \setminus B$  (lisez  $A$  moins  $B$ ) est formé par les éléments qui appartiennent à  $A$  et pas à  $B$ . On a donc, d'un point de vue logique

$$(x \in A \setminus B) \equiv ((x \in A) \text{ et } \neg(x \in B)).$$

On peut représenter des ensembles à l'aide de diagrammes. Les plus utilisés sont sans doute les diagrammes de Venn<sup>b</sup>. Ils permettent de visualiser les opérations qui ont été définies plus haut de manière très simple. Voici comment on les construit.

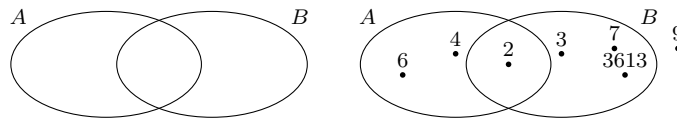
- On représente l'ensemble par une courbe fermée, généralement un cercle ou une ellipse (appelée parfois patate).
- Si on veut marquer qu'un objet est un élément de l'ensemble, on le place dans la région déterminée par la courbe<sup>c</sup>. On n'est pas obligé de représenter tous les éléments de l'ensemble en question, et c'est souvent impossible.
- On représente plusieurs ensembles (généralement 2, 3 ou 4) par plusieurs courbes fermées.

**Exemple 1.2.1.** On note  $A$  l'ensemble des nombres entiers pairs et strictement positifs. On le représente par le diagramme à gauche dans la figure suivante. Si on veut marquer que 2 et 4 sont des éléments de cet ensemble, on les y indique avec un point, comme dans le diagramme à droite dans la figure suivante. Notez que la position n'a pas d'importance, à l'intérieur de la région en forme de patate.



Prenons maintenant un exemple avec deux ensembles, donc avec un diagramme comportant deux patates.

**Exemple 1.2.2.** Soit  $A$  l'ensemble des nombres pairs strictement positifs et  $B$  l'ensemble des nombres premiers.<sup>d</sup> Voici à gauche la représentation générale des deux ensembles, et à droite quelques éléments des deux ensembles. On peut ajouter également des points “en dehors” des deux ensembles.



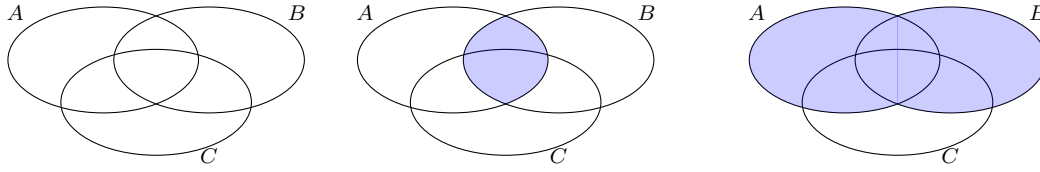
Cette représentation permet de visualiser aisément les unions et les intersections de deux ou plusieurs ensembles et d'avoir une intuition sur des égalités entre ensembles (sans toutefois constituer une démonstration rigoureuse). On peut en effet colorier ou hachurer les zones représentant les ensembles que l'on considère. Voici un exemple à trois ensembles<sup>e</sup>. A gauche, on a représenté la situation générale, au milieu, on a colorié la zone représentant  $A \cap B$ , à droite la zone représentant  $A \cup B$ .

b. John Venn (1834-1923) les formalisa en 1880.

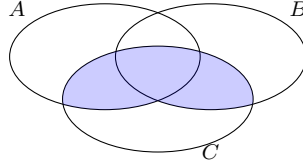
c. La plus petite des deux, évidemment.

d. Un nombre premier est un nombre entier naturel qui admet exactement deux diviseurs distincts.

e. Vous pouvez toujours adopter cette représentation pour trois ensembles



On peut aller encore plus loin et colorier par exemple la zone représentant  $(A \cup B) \cap C$  :



On peut alors constater sur cette représentation la relation

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

La proposition suivante présente un certain nombre de relations que nous avons déjà rencontrées en logique. Elles peuvent être démontrées en utilisant des tables de vérité, et j'en donne un exemple, ou en se ramenant à une propriété logique équivalente.

**Proposition 1.2.1.** *Si  $X$  est un ensemble et si  $A, B, C$  sont trois sous-ensembles de  $X$ , alors*

1.  $X \cup X = X, X \cap X = X$  ;
2.  $X \setminus X = \emptyset, X \setminus \emptyset = X, \emptyset \cup X = X, \emptyset \cap X = \emptyset$  ;
3.  $A \cup B = B \cup A, A \cap B = B \cap A$  ;
4.  $(A \cup B) \cup C = A \cup (B \cup C), (A \cap B) \cap C = A \cap (B \cap C)$  ;
5.  $A \cup (X \setminus A) = X, A \cap (X \setminus A) = \emptyset$  ;
6. *si  $A \subset B$ , alors  $A \cap C \subset B \cap C$  ;*
7. *si  $A \subset B$ , alors  $A \cup C \subset B \cup C$  ;*
8.  $A \cap B \subset A \subset A \cup B$  ;
9. *si  $C \subset A$  et  $C \subset B$ , alors  $C \subset A \cap B$  ;*
10. *si  $A \subset C$  et  $B \subset C$ , alors  $A \cup B \subset C$  ;*
11.  $A \cup (B \cap C) = (A \cup B) \cap (A \cup C)$  ;
12.  $A \cap (B \cup C) = (A \cap B) \cup (A \cap C)$  ;
13.  $(A \cup B) \setminus C = (A \setminus C) \cup (B \setminus C)$  ;
14.  $(A \cap B) \setminus C = (A \setminus C) \cap (B \setminus C)$  ;
15.  $X \setminus (A \cup B) = (X \setminus A) \cap (X \setminus B)$  ;
16.  $X \setminus (A \cap B) = (X \setminus A) \cup (X \setminus B)$ .

*Preuve de l'assertion 12.* Etant donné un objet  $x$ , on a les 3 assertions  $x \in A, x \in B$  et  $x \in C$ , qui peuvent toutes prendre les valeurs Vrai (1) ou faux (0). On calcule alors la table de vérité suivante.

| $x \in A$ | $x \in B$ | $x \in C$ | $x \in (B \cup C)$ | $x \in (A \cap B)$ | $x \in (A \cap C)$ | $x \in A \cap (B \cup C)$ | $x \in (A \cap B) \cup (A \cap C)$ |
|-----------|-----------|-----------|--------------------|--------------------|--------------------|---------------------------|------------------------------------|
| 0         | 0         | 0         | 0                  | 0                  | 0                  | 0                         | 0                                  |
| 0         | 0         | 1         | 1                  | 0                  | 0                  | 0                         | 0                                  |
| 0         | 1         | 0         | 1                  | 0                  | 0                  | 0                         | 0                                  |
| 0         | 1         | 1         | 1                  | 0                  | 0                  | 0                         | 0                                  |
| 1         | 0         | 0         | 0                  | 0                  | 0                  | 0                         | 0                                  |
| 1         | 0         | 1         | 1                  | 0                  | 1                  | 1                         | 1                                  |
| 1         | 1         | 0         | 1                  | 1                  | 0                  | 1                         | 1                                  |
| 1         | 1         | 1         | 1                  | 1                  | 1                  | 1                         | 1                                  |

On constate que les assertions  $x \in A \cap (B \cup C)$  et  $x \in (A \cap B) \cup (A \cap C)$  sont logiquement équivalentes.  $\square$

C'est un bon exercice de voir toutes ces assertions à l'aide de diagrammes de Venn. Il est également important d'essayer de les retenir. Les assertions 11 et 12 sont des règles de distributivité. Les assertions 15 et 16 correspondent à la négation logique et sont appelées lois de De Morgan.

Puisque les opérations d'union et d'intersection sur les ensembles sont associatives on peut définir l'union et l'intersection de plusieurs ensembles. On donnera facilement un sens aux expressions du type  $A_1 \cup \dots \cup A_n$  ou  $B_1 \cap \dots \cap B_p$ , où  $n$  et  $p$  sont des nombres naturels. Il arrivera aussi que l'on note ces ensembles

$$\cup_{k=1}^n A_k \quad \cap_{k=1}^p B_k.$$

Il s'agit d'un raccourci d'écritures fort utile, et que nous emploierons chaque fois que c'est possible, avec une opération associative et commutative. Vous les reverrez d'ici peu en algèbre avec les sommes. Il est important de noter que dans ces expressions, la lettre  $k$  est dite muette et peut être remplacée par n'importe quel autre symbole. C'est exactement le même principe que la lettre  $x$  dans  $\int_0^3 \sin(x)dx$ .

### 1.2.1 Un mot sur le paradoxe de Russell

J'ai présenté ci-dessus les rudiments de la théorie naïve des ensembles, et nous avons constaté que la définition même d'un ensemble n'en était pas une. Cela ne semble pas poser de problème puisque nous avons l'habitude dans la vie courante de travailler avec des collections d'objets et de pouvoir considérer intuitivement des unions, des intersections, des complémentaires...

Comme vous commencez à le percevoir, tout ce qui n'est pas parfaitement défini à partir d'axiomes ou de leurs conséquences peut comporter des risques en mathématiques. C'est également le cas de la théorie naïve des ensembles : elle est contradictoire. En effet, on peut, en utilisant les ensembles ainsi définis, trouver une proposition qui n'est ni vraie ni fausse. Cette proposition est un paradoxe, et montre que la théorie naïve des ensembles est incomplète. Rassurez-vous, on a depuis complété la théorie (en fait de plusieurs façons), et vous n'aurez pas à vous soucier du fait que la théorie des ensembles soit incomplète avant longtemps dans vos études.

Je vous livre cependant cette propriété paradoxale, publiée par B. Russell en 1903.

Considérons les ensembles suivants :  $A = \{1, 2, 3, a, b, 2\}$  et  $B = \{1, 2, 3, 4, B, a, u, v\}$ . On voit une différence entre les deux ensembles. En effet, on a  $B \in B$ . Puisque les ensembles sont des collections quelconques. Cela ne pose pas de problème.

Appelons donc ensembles extraordinaires ceux qui, comme  $B$ , se contiennent eux-mêmes. Appelons également ordinaires les ensembles qui ne se contiennent pas eux-mêmes, c'est à dire ceux qu'on a l'habitude de voir.

Considérons l'ensemble  $R$  des ensembles ordinaires. Cet ensemble conduit au paradoxe : on se demande si  $R$  est ordinaire ou extraordinaire.

S'il est ordinaire, alors il ne se contient pas comme élément, donc  $R \notin R$ , donc  $R$  est extraordinaire.

S'il est extraordinaire, alors  $R \in R$ , donc  $R$  est ordinaire.

## 1.3 Relations, applications, injections, surjections

Le but de cette section est de reconstruire la notion d'application. Elle généralise ce que vous avez appelé "fonction" dans l'enseignement secondaire. La plupart d'entre vous a eu une définition qui s'exprime comme ceci

"Une fonction est une loi qui à tout  $x$  associe un  $y$ , on note alors  $y = f(x)$ ."

On entend même dire (et aussi écrire) "la fonction  $f(x)$ ", ou "la fonction  $y = f(x)$ ", " $x$  et  $y$  sont des variables liées par  $f$ "...

Il y a plusieurs problèmes dans cette “définition” et dans les conceptions qui en découlent. Le problème le plus évident est qu’on n’a pas dit ce que pouvait être  $x$ . On peut donc préciser la définition comme ceci :

“Une fonction d’un ensemble  $A$  dans un ensemble  $B$  est une loi qui à tout  $x \in A$  associe un  $y \in B$ , on note alors  $y = f(x)$ .”

Cela a l’air plus scientifique, plus précis, mais il y a encore un problème : qu’est-ce qu’une “loi” ?

Dans la première partie de cette section, nous allons donner un sens à cette définition, en explicitant ce que l’on veut dire par “loi” dans cette définition un peu bancale. Cela se fait à l’aide du concept de relation, qui est à la fois plus simple et plus général que celui de fonction. Nous ne définirons pas ce que pourraient être les variables, car on ne parle pas de variables en mathématique. La notion d’ensemble remplit en effet le rôle que jouent les variables en sciences : quand un scientifique parle de la “variable température”, le mathématicien envisage l’ensemble de toutes les températures possibles, et c’est vrai pour toutes les “variables” scientifiques.

Pour définir ce qu’est une relation, nous aurons besoin de la notion de produit cartésien d’ensembles. Pour deux ensembles, cette notion est assez simple à définir.

**Définition 1.3.1.** Si  $A$  et  $B$  sont deux ensembles, alors le produit cartésien de  $A$  et  $B$  est l’ensemble

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

On peut bien entendu étendre cette définition de produit à plusieurs ensembles. L’opération “produit” n’est pas tout à fait associative, et on devrait en toute précision mettre des parenthèses, mais cela ne posera pas de problème de faire un léger abus et de les oublier.

**Définition 1.3.2.** Si  $A_1, \dots, A_n$  sont des ensembles ( $n \in \mathbb{N}$ ), alors le produit cartésien  $A_1 \times \dots \times A_n$  est l’ensemble

$$A_1 \times \dots \times A_n = \{(a_1, \dots, a_n) : a_1 \in A_1, \dots, a_n \in A_n\}.$$

Donnons directement la définition d’une relation d’un ensemble  $A$  vers un ensemble  $B$ .

**Définition 1.3.3.** Une relation  $\mathcal{R}$  de  $A$  dans  $B$  est une partie de  $A \times B$ . On appelle  $A$  l’ensemble de départ et  $B$  l’ensemble d’arrivée de  $\mathcal{R}$ .

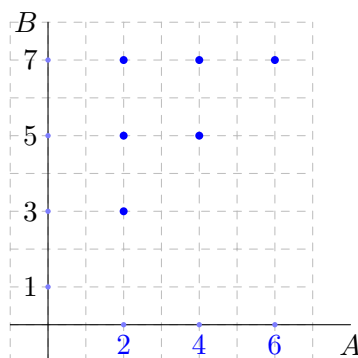
Si le couple  $(a, b)$  appartient  $\mathcal{R}$ , on note  $a\mathcal{R}b$  et on dit que  $a$  est en relation avec  $b$ . Le lien entre les deux notations est donc

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

**Exemple 1.3.1.** Posons  $A = \{2, 4, 6\}$  et  $B = \{1, 3, 5, 7\}$ . La relation “est plus petit que”, de  $A$  dans  $B$  est

$$\mathcal{R} = \{(2, 3), (2, 5), (2, 7), (4, 5), (4, 7), (6, 7)\}.$$

On peut bien sûr représenter le produit cartésien comme d’habitude par un graphique plan et on représente alors facilement la relation :



On peut définir une relation de  $A$  dans  $B$  par une assertion définissant l'appartenance à cette relation : par exemple, définissons  $\mathcal{R}'$  par  $x\mathcal{R}'y$  si et seulement si  $y = x + 3$ . On a alors  $\mathcal{R}' = \{(2, 5), (4, 7)\}$ . Je vous laisse faire la représentation graphique.

Voici encore trois exemples.

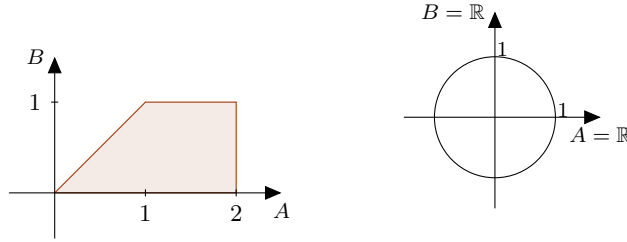
1. Considérons la relation  $\geq$  de  $A = [0, 2]$  dans  $B = [0, 1]$  donnée par

$$\mathcal{R}_1 = \{(x, y) \in [0, 2] \times [0, 1] : x \geq y\}.$$

2. Soit  $\mathcal{R}_2$  la relation de  $\mathbb{R}$  dans  $\mathbb{R}$  définie par

$$x\mathcal{R}_2y \quad \text{si, et seulement si} \quad x^2 + y^2 = 1.$$

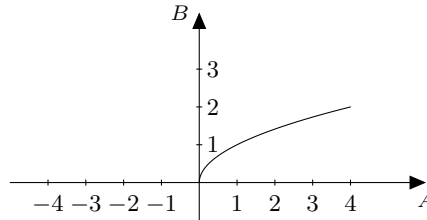
Elle se représentent visiblement par



3. Terminons par la relation  $\mathcal{R}_3$  définie de  $A = [-4, 4]$  dans  $B = [0, 3]$  par

$$x\mathcal{R}_3y \quad \text{si, et seulement si} \quad y^2 - x = 0.$$

Elle est représentée par



Remarquez que par convention, quand on peut représenter une relation par un graphique plan, on indique l'ensemble de départ sur un axe horizontal. Cette convention n'a rien de mathématique, puisque l'horizontalité n'est pas définie.

**Définition 1.3.4.** Soit  $\mathcal{R}$  une relation de  $A$  dans  $B$ . On appelle *domaine* de  $\mathcal{R}$  l'ensemble des points  $a$  de  $A$  qui sont en relation avec au moins un élément  $b$  de  $B$ . On le note  $\text{dom}_{\mathcal{R}}$  ou  $\text{dom}(\mathcal{R})$  ou encore  $D_{\mathcal{R}}$ . On a

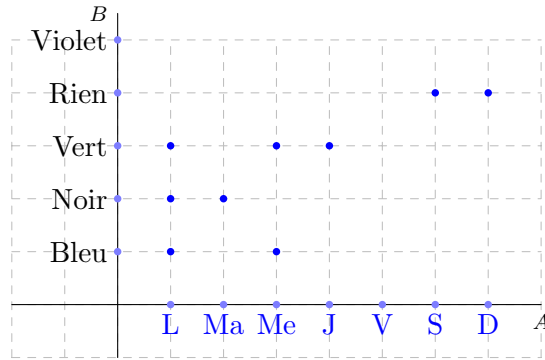
$$\text{dom}_{\mathcal{R}} = \text{dom}(\mathcal{R}) = D_{\mathcal{R}} = \{a \in A : \exists b \in B : a\mathcal{R}b\}.$$

On appelle *codomaine* ou *image* de  $\mathcal{R}$  l'ensemble  $\text{Im}(\mathcal{R})$  (ou  $\text{Im}_{\mathcal{R}}$ ) des points  $b$  de  $B$  tels qu'il existe au moins un élément  $a$  de  $A$  qui soit en relation avec  $b$ . On a

$$\text{Im}_{\mathcal{R}} = \text{Im}(\mathcal{R}) = \{b \in B : \exists a \in A : a\mathcal{R}b\}.$$

Voici quelques exemples :

1. si  $A$  est l'ensemble des jours de la semaine et  $B$  un ensemble de couleurs de chaussettes que je peux porter, on peut considérer la relation suivante, de l'ensemble  $A = \{\text{L, Ma, Me, J, V, S, D}\}$  dans l'ensemble  $B = \{\text{bleu, noir, vert, blanc, violet}\}$ .



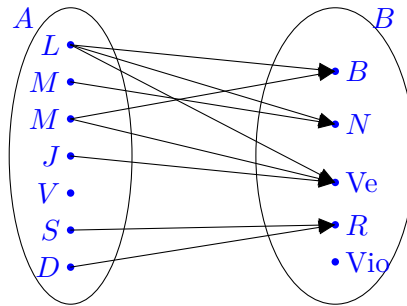
On constate qu'il n'y a pas de point de  $B$  qui soit en relation avec vendredi (V), donc  $\text{dom}_{\mathcal{R}} = A \setminus \{V\}$ . De même, aucun point de  $A$  n'est en relation avec violet, et on calcule donc que  $\text{Im}_{\mathcal{R}} = B \setminus \{\text{Violet}\}$ .

2. Soit la relation  $\mathcal{R}_2$  de  $\mathbb{N}$  dans  $\mathbb{N}$  définie par  $x\mathcal{R}_2y$  si, et seulement si  $x + y = 3$ . On constate que le domaine de  $\mathcal{R}_2$  est  $\{0, 1, 2, 3\}$ , tandis que son image est également  $\{0, 1, 2, 3\}$ .
3. Soit la relation  $\mathcal{R}_3$  de  $\mathbb{R}$  dans  $\mathbb{R}$  définie par  $x\mathcal{R}_3y$  si, et seulement si  $|x| + |y| = 3$ . Le domaine de  $\mathcal{R}_3$  est égal à  $[-3, 3]$ , et son image aussi. Il est intéressant de représenter cette relation dans le plan, identifié à  $\mathbb{R}^2$  au moyen d'un repère orthonormé : on obtient un carré. Vous verrez en analyse qu'il s'agit de la boule de centre  $(0, 0)$  et de rayon 3, pour la distance de Manhattan.

En ce qui concerne les représentations graphiques que l'on peut faire d'une relation, nous connaissons déjà une façon de représenter un produit, même de manière schématique, dans le plan. Nous avons utilisé cette représentation dans le premier exemple ci-dessus. Il existe une autre représentation graphique, en termes de diagramme de Venn. C'est la représentation *sagittale*.

**Définition 1.3.5.** La représentation sagittale d'une relation  $\mathcal{R}$  de  $A$  dans  $B$  est obtenue en représentant les ensembles par des diagrammes de Venn et en indiquant une flèche<sup>a</sup> de  $a \in A$  vers  $b \in B$  quand  $a\mathcal{R}b$ .

Dans notre exemple de chaussettes, cela donne ceci.



On constate facilement sur cette représentation que  $V$  n'est pas dans le domaine de  $\mathbb{R}$  car aucune flèche ne part de  $V$ . De même "Violet" n'est pas dans l'image car aucune flèche n'arrive à "Violet". Cette représentation sagittale a d'autres avantages : elle permet de visualiser facilement la définition de la composée de deux relations, ainsi que de la relation réciproque.

**Définition 1.3.6.** Si  $\mathcal{R} : A \rightarrow B$  et  $\mathcal{R}' : B \rightarrow C$  sont des relations, alors la relation composée  $\mathcal{R}' \circ \mathcal{R} : A \rightarrow C$  est définie par<sup>b</sup>

$$\mathcal{R}' \circ \mathcal{R} = \{(a, c) \in A \times C : \exists b \in B : a\mathcal{R}b \text{ et } b\mathcal{R}'c\}.$$

a. En latin, sagitta veut dire flèche, cela a également donné le mot sagittaire.

b. Attention à l'ordre dans lequel on écrit les relation  $\circ$  se lit "après".

En termes de représentation sagittale,  $a$  est donc en relation avec  $c$  pour  $\mathcal{R}' \circ \mathcal{R}$  si on peut connecter  $a$  à  $c$  par deux flèches successives (la première de  $\mathcal{R}$  et la seconde de  $\mathcal{R}'$ ), aboutissant et partant d'un point intermédiaire  $b \in B$ .

Cette définition nous permettra d'ici peu de composer des applications et d'obtenir par exemple la relation

$$\mathcal{R} = \{(x, y) \in \mathbb{R}^2 : y = \sin^2(x^3)\}$$

comme la composée des relations  $\mathcal{R}_1 = \{(x, y) \in \mathbb{R}^2 : y = x^3\}$ ,  $\mathcal{R}_2 = \{(y, z) \in \mathbb{R}^2 : z = \sin(y)\}$  et  $\mathcal{R}_3 = \{(z, a) \in \mathbb{R}^2 : a = z^2\}$ . On a alors  $\mathcal{R} = \mathcal{R}_3 \circ (\mathcal{R}_2 \circ \mathcal{R}_1)$ .

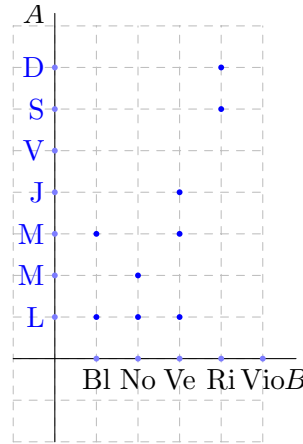
**Exercice 1.3.1.** Montrer que la composition des relations est associative : on a  $\mathcal{R}_3 \circ (\mathcal{R}_2 \circ \mathcal{R}_1) = (\mathcal{R}_3 \circ \mathcal{R}_2) \circ \mathcal{R}_1$ , quelles que soient les relations  $\mathcal{R}_1, \mathcal{R}_2, \mathcal{R}_3$ .

En ce qui concerne la relation réciproque, elle est obtenue sur le diagramme simplement en inversant le sens des flèches. Il s'agit donc d'une relation de  $B$  vers  $A$ , si  $\mathcal{R}$  est une relation de  $A$  dans  $B$ . Plus formellement, on a la définition suivante.

**Définition 1.3.7.** Si  $\mathcal{R} : A \rightarrow B$  est une relation, alors la relation inverse (ou réciproque) de  $\mathcal{R}$  est la relation  $\mathcal{R}^{-1} : B \rightarrow A$  définie par

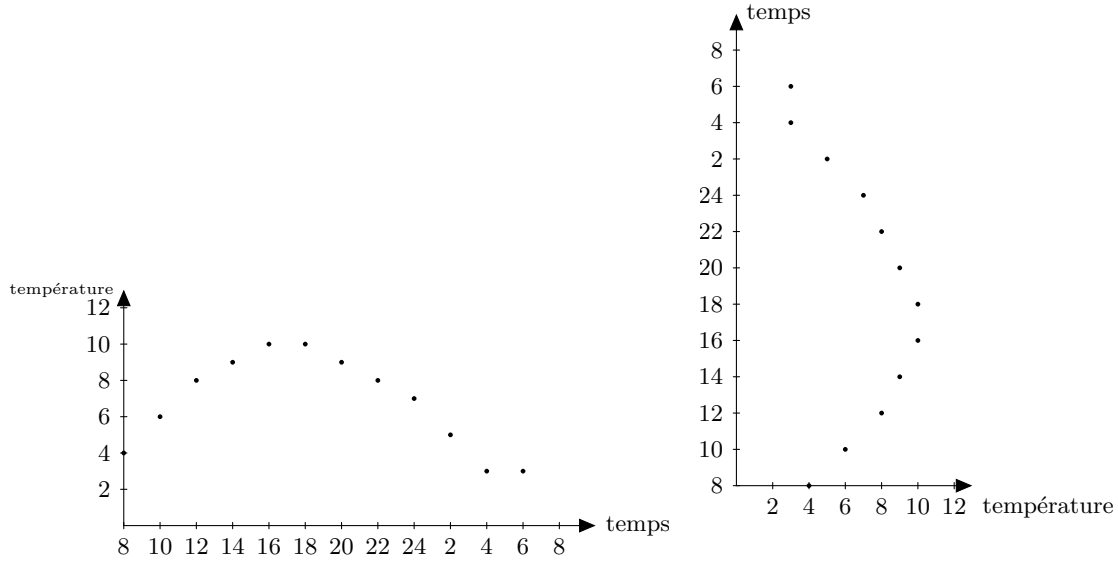
$$\mathcal{R}^{-1} = \{(b, a) \in B \times A : a\mathcal{R}b\}.$$

Nous avons déjà la représentation sagittale de  $\mathcal{R}^{-1}$ . Si on se concentre sur la représentation en produit cartésien, il suffit de lire le graphique dans l'autre sens : de  $B$  vers  $A$ . Cependant, comme on prend souvent la convention de représenter l'ensemble de départ horizontalement<sup>c</sup>, on place  $B$  sur l'axe horizontal et  $A$  sur l'axe vertical, et on considère les mêmes point qu'avant. Cela revient géométriquement à effectuer une symétrie orthogonale qui échange les axes. Pour l'exemple des chaussettes, on a la représentation suivante de  $\mathcal{R}^{-1}$  :



Voici encore un exemple d'une relation liant l'heure et la température observée, et de sa relation réciproque :

c. Comme je l'ai dit plus haut, cela n'a rien de mathématique, mais c'est très courant en sciences.



Il est également intéressant de calculer la composée d'une relation et de sa réciproque.

**Exercice 1.3.2.** Démontrer que  $\text{dom}_{\mathcal{R}^{-1}} = \text{Im}_{\mathcal{R}}$  et  $\text{Im}_{\mathcal{R}^{-1}} = \text{dom}_{\mathcal{R}}$ .

Ayant une relation à sa disposition, on peut en créer d'autres en restreignant l'ensemble de départ ou l'ensemble d'arrivée. C'est l'objet de la définition suivante.

**Proposition 1.3.1.** Si  $\mathcal{R} : A \rightarrow B$  est une relation, et si  $A'$  est un sous-ensemble de  $A$ , alors la restriction de  $\mathcal{R}$  à  $A'$  est la relation  $\mathcal{R}|_{A'} : A' \rightarrow B$  définie par

$$\mathcal{R}|_{A'} = \{(a, b) \in A' \times B : a\mathcal{R}b\}.$$

On restreint donc l'ensemble de départ. Cela a bien sûr une influence sur l'image de la relation en général. Par exemple, si  $\mathcal{R} = \{(x, y) \in \mathbb{R}^2 : y = x^2\}$ , alors  $\mathcal{R}|_{[0, +\infty[} = \{(x, y) \in [0, +\infty[ \times \mathbb{R} : y = x^2\}$  a la même image que  $\mathcal{R}$ , tandis que  $\mathcal{R}|_{[0, 2]}$  admet  $[0, 4]$  pour image.

On peut aussi restreindre l'ensemble d'arrivée et définir la restriction de  $\mathcal{R}$  à  $B' \subset B$  comme étant  $\{(a, b) \in A \times B' : a\mathcal{R}b\}$ . Je n'utiliserai pas directement cette restriction et je n'introduis donc pas de notation particulière.

## 1.4 Applications

Passons maintenant au premier type particulier de relation. Il s'agit des relations de type application. Je les définis en deux temps, car elles sont caractérisées par deux propriétés distinctes. La première de ces propriétés est le fait de pouvoir associer à tout élément de  $A$  *au plus* un élément de  $B$ . L'élément de  $B$  est alors déterminé *en fonction* de l'élément de  $A$  que l'on considère. Dans le cas de la relation  $\mathcal{R} : \text{temps} \rightarrow \text{température}$ , la température est exprimée en fonction de l'heure : à chaque heure, il correspond au plus une température. Si on regarde la relation réciproque, pour certaines températures, il y a plusieurs heures associées, donc l'heure ne s'exprime pas en fonction de la température. Passons maintenant à la définition formelle.

**Définition 1.4.1.** Une relation  $\mathcal{R}$  de  $A$  dans  $B$  est de type fonctionnel si tout point  $a$  de  $\text{dom}_{\mathcal{R}}$  est en relation avec exactement un élément de  $B$ .

On a comme d'habitude une formulation équivalente, puisqu'on analyse en fait l'unicité de l'élément de  $B$  qui est associé à chaque point de  $A$ . Une relation est donc de type fonctionnel si l'implication suivante est vraie

$$(a \in A, b_1, b_2 \in B, a\mathcal{R}b_1, a\mathcal{R}b_2) \Rightarrow b_1 = b_2.$$

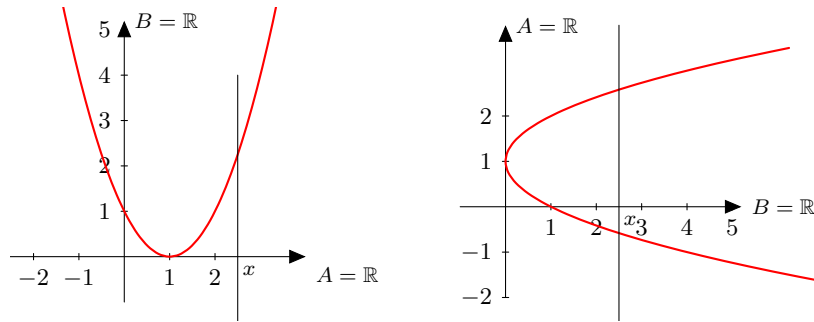


FIGURE 1.1 – Deux relations

Cela donne lieu sur la représentation graphique (quand elle est possible), à un test élémentaire.

La relation à gauche de  $A = \mathbb{R}$  dans  $B = \mathbb{R}$  est de type fonctionnel : à chaque point de  $A$  correspond au plus un point de  $B$ , comme on le voit en représentant l'ensemble des couples  $\{(x, b) : b \in B\}$  pour chaque  $x \in A^a$ . Pour information, cette relation s'écrit  $\{(x, (x - 1)^2) : x \in \mathbb{R}\}$ . A droite, il s'agit de la relation réciproque, de  $B = \mathbb{R}$  dans  $A = \mathbb{R}$ . On constate qu'elle n'est pas de type fonctionnel, puisqu'il existe au moins un point  $b$  de l'ensemble de départ (ici  $B$ ) qui est en relation plus d'un point de l'ensemble d'arrivée ( $A$ ).

Dans une relation de type fonctionnel  $\mathcal{R} : A \rightarrow B$ , le point  $b \in B$  associé à un point  $a$  de  $A$  est unique (quand il existe). Cela lui vaut un nom.

**Définition 1.4.2.** Si  $\mathcal{R} : A \rightarrow B$  est de type fonctionnel, alors si  $(a, b) \in \mathcal{R}$ , on dit que  $b$  est l'image de  $a$  par  $\mathcal{R}$ .

Parmi les relations de type fonctionnel, il en est qui sont particulières : celles pour lesquelles tout point de  $a$  admet une image. De manière équivalente, ce sont les relations  $\mathcal{R} : A \rightarrow B$  dont le domaine est  $A$ .

**Définition 1.4.3.** Une relation  $\mathcal{R} : A \rightarrow B$  est de type application si les deux conditions suivantes sont satisfaites :

1. La relation  $\mathcal{R}$  est de type fonctionnel ;
2. Le domaine de  $\mathcal{R}$  est égal à  $A$ .

On peut bien sûr résumer ces deux conditions : la première stipule que tout point  $a$  de  $A$  est en relation avec au plus un point  $b$  de  $B$ . La deuxième s'écrit également "tout point  $a$  de  $A$  est en relation avec au moins un point  $b$  de  $B$ ". La conjonction des deux conditions s'écrit donc de manière équivalente :

"Tout point  $a$  de  $A$  est en relation avec exactement un point de  $b$  de  $B$ ."

Bien entendu, si on dispose d'une relation de type fonctionnel  $\mathcal{R}$  de  $A$  dans  $B$ , on peut toujours en faire une relation de type application en restreignant son ensemble de départ à son domaine. Ce fait, bien qu'évident, vaut bien une proposition.

**Proposition 1.4.1.** Si  $\mathcal{R} : A \rightarrow B$  est de type fonctionnel, alors  $\mathcal{R}|_{\text{dom}(\mathcal{R})}$  est de type application.

*Remarque 1.2.* Il est important de remarquer que les définitions que nous venons de poser ne disent rien sur les points de  $B$  : il peut exister des points de  $B$  qui ne sont images d'aucun point de  $A$ , et aussi des points de  $B$  qui sont images de plusieurs points de  $A$ .

Voici quelques exemples.

a. C'est la droite verticale représentée sur le graphique.

b. Ici encore, on le voit en traçant des droites verticales sur les représentations graphiques.

**Exemple 1.4.1.** Les relations suivantes sont de type application :

1.  $\mathcal{R}_1 : \mathbb{R} \rightarrow \mathbb{R}$  définie par  $\mathcal{R}_1 = \{(x, x^2) : x \in \mathbb{R}\}$  ;
2.  $\mathcal{R}_2 : \mathbb{R} \rightarrow [0, +\infty[$  définie par  $\mathcal{R}_2 = \{(x, x^2) : x \in \mathbb{R}\}$  ;
3.  $\mathcal{R}_3 : [0, +\infty[ \rightarrow [0, +\infty[$  définie par  $\mathcal{R}_3 = \{(x, x^2) : x \in \mathbb{R}\}$ .

Par contre la relation  $\mathcal{R}_4 : \mathbb{R} \rightarrow \mathbb{R}$  définie par  $\mathcal{R}_4 = \{(x, \operatorname{tg}(x)) : x \in \mathbb{R}\}$  n'est pas une relation car  $\operatorname{tg}(x)$  n'est pas défini pour tout  $x \in \mathbb{R}$ .

Il est important de remarquer que l'ensemble d'arrivée peut dans une certaine mesure être modifié sans changer le caractère “application” d'une relation. Ainsi, la relation  $\mathcal{R}_2$  est de type application. En élargissant l'ensemble d'arrivée pour obtenir la relation  $\mathcal{R}_1$ , on conserve une relation de type application.

Je ne peux terminer cette section sur les applications sans faire le lien avec la notion que vous avez vue en secondaire, et qui est bien souvent la seule utilisée en sciences.

Elle n'est pas bien difficile à faire : une relation de type application  $\mathcal{R} : A \rightarrow B$  est la donnée, pour chaque point  $a$  de  $A$ , d'un unique point  $b \in B$ , qui est l'image de  $a$ . On peut donc voir une relation comme une “transformation” qui transforme chaque point de  $A$  en son image. On note alors l'application par une lettre (souvent  $f^c$ , et on précise la transformation en question : on a alors la notation complète suivante :

$$f : A \rightarrow B : a \mapsto f(a).$$

Elle indique que  $f$  est une application de  $A$  dans  $B$  qui à chaque point  $a$  de  $A$  associe son image  $f(a)$ . On obtient ainsi la notion d'application qui a été introduite dans l'enseignement secondaire.

**Définition 1.4.4** (Intuitive mais incomplète). Une application  $f$  de  $A$  dans  $B$  est une “loi de transformation” qui associe à tout point  $x$  de  $A$ , associe un point  $f(x)$  de  $B^d$ .

Bien entendu, partant de cette vision “loi de transformation” d'une application de  $A$  dans  $B$ , il n'est pas difficile de récupérer la relation de type application correspondante : si on considère la “loi de transformation” qui transforme tout  $x$  dans  $\mathbb{R}$  en  $\sin(x)$ , alors la relation de type application correspondante (au sens de notre définition officielle) est

$$\mathcal{R} = \{(x, \sin(x)) : x \in \mathbb{R}\}.$$

Cette relation s'écrit encore  $\{(x, y) \in \mathbb{R}^2 : y = \sin(x)\}$ . Nous avons donc fait le lien également avec une formulation utilisée dans l'enseignement secondaire où l'on parle parfois de la “fonction  $y = \sin(x)$ ”. L'avantage évident de notre approche est qu'il n'y a pas à se poser de question sur le statut de  $x$  et  $y$  (on parle dans l'enseignement secondaire de variables ou d'indéterminées). Ici, nous considérons la relation  $\{(x, y) \in \mathbb{R}^2 : y = \sin(x)\}$ , et  $x$  et  $y$  sont toujours des nombres réels.

L'avantage fondamental de notre approche de la notion d'application par rapport à celle qui a été généralement suivie dans l'enseignement secondaire est qu'elle ne fait appel à aucun concept non défini : elle est basée sur les notions de sous-ensemble et de produit cartésien et des notions d'existence et d'unicité. Si j'avais défini “Une application est une loi de transformation...”, il aurait été naturel que vous me demandiez de définir ce qu'est une loi de transformation. Je n'aurais pas eu de réponse à donner, à part une association particulière entre des points de deux ensembles... et on se serait ramené à la définition qui a été donnée plus haut.

Le point de vue intuitif “loi de transformation” que vous avez connu jusqu'à présent est maintenant intégré à une définition précise. Vous pouvez donc continuer à penser les applications comme des lois de transformation, mais si on vous demande de définir cette

c. On utilise souvent  $f$  parce que, lorsque  $B = \mathbb{R}$  ou  $B = \mathbb{C}$ , les applications sont appelées fonctions, mais tout autre symbole convient.

d. Donc  $f$  transforme  $x$  en  $f(x)$ .

notion avec toute la rigueur nécessaire (c'est-à-dire à partir des objets mathématiques définis en amont), vous savez maintenant que ces lois de transformations sont associées à des relations de type application. Résumons ces diverses constatations dans une proposition<sup>e</sup>.

**Proposition 1.4.2.** *Si  $\mathcal{R}$  est une relation de type application de  $A$  dans  $B$ , alors elle définit une “loi de transformation” de  $A$  dans  $B$  :  $a \in A$  est transformé en  $b \in B$  si, et seulement si  $(a, b) \in \mathcal{R}$ . Réciproquement, si  $f : A \rightarrow B : a \mapsto f(a)$  est une “loi de transformation”, elle définit une relation de type application, appelée le graphe de  $f$  et notée  $G_f$  :*

$$G_f = \{(a, f(a)) : a \in A\} = \{(a, b) \in A \times B : b = f(a)\}.$$

Dans la suite, nous utiliserons la notation  $f : A \rightarrow B : a \mapsto f(a)$ , et nous pourrions penser  $f$  comme une loi de transformation, tout en sachant que si nous devons utiliser la définition dans une preuve, il s'agit bien d'une relation.

Terminons cette section par la notion de composée d'applications. Nous avons déjà défini la composée de relations en général, il suffit donc de vérifier que cette définition s'applique aux relations de type application.

**Proposition 1.4.3.** *La composée de relations de type application est de type application. Plus précisément, si on a  $\mathcal{R} : A \rightarrow B$ ,  $\mathcal{R}' : B \rightarrow C$ ,  $\mathcal{R} = G_f$  et  $\mathcal{R}' = G_g$ , alors  $\mathcal{R}' \circ \mathcal{R} : A \rightarrow C$  s'écrit  $G_{g \circ f}$  où  $(g \circ f)(a) = g(f(a))$ , pour tout  $a \in A$ .*

*Démonstration.* Il suffit de vérifier que  $\mathcal{R}' \circ \mathcal{R}$  satisfait les deux conditions pour être de type application. Par définition, pour  $a \in A$  et  $c \in C$ , on a  $(a, c) \in \mathcal{R}' \circ \mathcal{R}$  si, et seulement si, il existe  $b \in B$  tel que  $a\mathcal{R}b$  et  $b\mathcal{R}'c$ . Mais puisque  $\mathcal{R}$  est de type application, pour tout  $a \in A$ , il existe un unique  $b$  tel que  $a\mathcal{R}b$ , c'est l'image de  $a$ , notée  $f(a)$ . De même, pour tout  $b \in B$  (et en particulier  $f(a)$ ), il existe un unique  $c$  tel que  $b\mathcal{R}'c$ , c'est  $g(b)$ . En conclusion, pour tout  $a \in A$ , il existe un unique  $c \in C$  tel que  $a\mathcal{R}' \circ \mathcal{R}c$ , c'est  $g(f(a))$ .  $\square$

## 1.5 Applications réciproques

Nous avons défini les relations réciproques en toute généralité, et la réciproque d'une relation existe toujours. Nous avons ensuite défini des relations particulières : les applications. Il est naturel de se demander si la réciproque d'une application est encore une application. Un bref coup d'oeil à la figure 1.1 montre que ce n'est pas toujours vrai. La question naturelle qui se pose est de déterminer des conditions sur  $\mathcal{R}$  pour que  $\mathcal{R}^{-1}$  soit de type application. La réponse à ce question est simple. La voici.

**Proposition 1.5.1.** *Soit  $\mathcal{R} : A \rightarrow B$  une relation de type application. Alors  $\mathcal{R}^{-1}$  est de type application si, et seulement si, pour tout  $b \in B$  il existe un unique  $a \in A$  tel que  $a\mathcal{R}b$ .*

*Démonstration.* Il suffit de traduire le fait que  $\mathcal{R}^{-1}$  soit de type application : c'est le cas si, et seulement si, pour tout  $b \in B$  il existe un unique  $a \in A$  tel que  $(b, a) \in \mathcal{R}^{-1}$ . Mais la condition  $(b, a) \in \mathcal{R}^{-1}$  est par définition équivalente à  $(a, b) \in \mathcal{R}$ , ou encore à  $a\mathcal{R}b$ .  $\square$

Cette proposition mène à la définition suivante.

**Définition 1.5.1.** Une application  $f : A \rightarrow B$  telle que pour tout  $b \in B$  il existe un unique  $a \in A$  tel que  $f(a) = b$  est une bijection.

Bien entendu, la condition pour que  $f$  soit une bijection est une conjonction : existence et unicité. Comme dans la définition des applications, il est utile de séparer ces deux conditions. Cela donne lieu à deux propriétés des applications, l'injectivité et la surjectivité.

Considérons les deux relations de la figure 1.1. La première est une application, et

e. Cette proposition est nécessairement bancal, puisqu'elle donne un lien entre un concept bien défini, celui de relation et un concept mal défini, celui de loi de transformation.

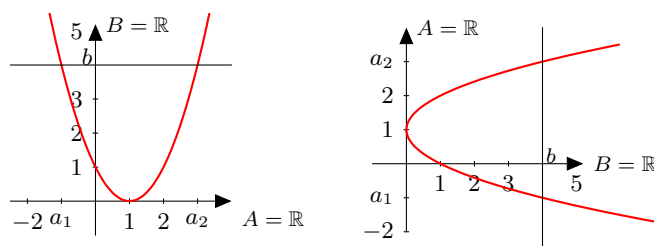


FIGURE 1.2 – Une application non injective et sa réciproque

sa réciproque ne l'est pas. Elle n'est en effet pas de type fonctionnel : il existe un point  $b \in B$  qui est en relation avec deux points  $a_1$  et  $a_2$  de  $A$ . Cela peut se voir directement sur l'application  $f$  et nous amène à la définition.

**Définition 1.5.2.** Une application  $f : A \rightarrow B$  est injective<sup>a</sup> si il n'existe pas  $a_1, a_2 \in A$  tels que  $a_1 \neq a_2$  et  $f(a_1) = f(a_2)$ .

Bien sûr cette définition n'est pas facile à mettre en oeuvre, puisqu'elle est basée sur l'inégalité  $a_1 \neq a_2$ . On passe donc naturellement à la contraposée qui fournit une définition équivalente pour l'injectivité.

**Proposition 1.5.2.** Une application  $f : A \rightarrow B$  est injective si, et seulement si, pour tous  $a_1, a_2 \in A$ , si  $f(a_1) = f(a_2)$ , alors  $a_1 = a_2$ .

*Démonstration.* Il suffit de contraposer la condition de la définition. □

*Remarque 1.3.* Il y a une erreur fréquemment commise avec cette définition, elle consiste à renverser le sens de l'implication et à écrire “si  $a_1 = a_2$  alors  $f(a_1) = f(a_2)$ .” Cette condition n'apporte évidemment aucune information supplémentaire au fait que  $f$  soit une application.

Voici quelques exemples et contre-exemples.

- Exemple 1.5.1.**
1. L'application  $f_1 : \mathbb{R} \rightarrow \mathbb{R} : x \mapsto x^2$  n'est pas injective. En effet, on constate que  $f_1(-2) = f_1(2) = 4$ .
  2. L'application  $f_2 : [0, +\infty[ \rightarrow \mathbb{R} : x \mapsto x^2$  est injective. En effet, soient  $x, y \in [0, +\infty[$  tels que  $f(x) = f(y)$ . On a alors  $x^2 = y^2$ , ou encore  $(x - y)(x + y) = 0$ . Si on fixe  $y \geq 0$  et que l'on cherche tous les  $x$  satisfaisant cette équation, on trouve  $x = y$  ou  $x = -y$ . La deuxième solution n'est possible pour  $x \geq 0$  que si  $x = y = 0$ . On a donc bien montré que  $f(x) = f(y)$  implique  $x = y$ , pour tous  $x, y \in [0, +\infty[$ .
  3. L'application  $f_3 : ]-\infty, 0] \rightarrow \mathbb{R} : x \mapsto x^2$  est injective, pour la même raison.
  4. Il en va de même pour l'application  $f_4 : A \rightarrow \mathbb{R} : x \mapsto x^2$ , pour tout sous-ensemble  $A$  de  $\mathbb{R}$  dont l'intersection avec la paire  $\{-x, x\}$  contient au plus un élément, pour tout  $x \in \mathbb{R}$ .
  5. L'application  $f_5 : \mathbb{R} \rightarrow \mathbb{R} : x \mapsto \sin(x)$ , appelée application sinus<sup>b</sup> n'est pas injective. En effet, on a  $\sin(\frac{\pi}{3}) = \sin(\frac{2\pi}{3})$ .
  6. L'application  $\sin : ]-\frac{\pi}{2}, \frac{\pi}{2}[ \rightarrow \mathbb{R}$  est injective. Cela peut se voir à partir du cercle trigonométrique, ou en résolvant l'équation  $\sin(x) = \sin(y)$  à partir des propriétés du sinus, pour tout nombre  $y$  fixé dans  $]-\frac{\pi}{2}, \frac{\pi}{2}[$ .

Passons maintenant à quelques résultats sur les applications injectives. Le premier n'est pas surprenant puisque c'est pour l'obtenir que l'on a posé la définition.

a. On dit aussi que  $f : A \rightarrow B$  est une injection.

b. Cette application sera redéfinie au cours d'analyse à partir de l'exponentielle des nombres complexes. Je me base ici sur la définition que vous en avez à partir du cercle trigonométrique.

**Proposition 1.5.3.** Si  $f : A \rightarrow B$  est une application injective, alors  $G_f^{-1}$  est une relation de type fonctionnel.

*Démonstration.* Par définition, puisque  $G_f^{-1}$  est une relation de  $B$  dans  $A$ , on doit montrer que si  $b \in B$ ,  $a_1, a_2 \in A$  sont tels que  $bG_f^{-1}a_1$  et  $bG_f^{-1}a_2$ , alors on a  $a_1 = a_2$ . Mais les deux conditions s'écrivent aussi  $a_1G_fb$  et  $a_2G_fb$ , ou encore  $b = f(a_1)$  et  $b = f(a_2)$ . Vu l'injectivité de  $f$ , on obtient  $a_1 = a_2$ .  $\square$

*Remarque 1.4.* La réciproque de cette proposition est vraie et est laissée comme exercice.

Nous pouvons également traduire l'injectivité en termes d'équations.

**Proposition 1.5.4.** Une application  $f : A \rightarrow B$  est injective si, et seulement si, pour tout  $b \in B$ , l'équation

$$f(x) = b, (x \in A) \quad (1.1)$$

admet au plus une solution.

*Démonstration.* Ici encore, on utilise une simple traduction de la définition.

Supposons que  $f$  est une application injective et montrons que (1.1) admet au plus une solution. Si tel n'est pas le cas, il existe  $x_1, x_2 \in A$  tels que  $x_1 \neq x_2$  et  $f(x_1) = b$  et  $f(x_2) = b$ . Mais alors l'injectivité de  $f$  implique  $x_1 = x_2$ , une absurdité.

Réciproquement, si l'équation (1.1) admet au plus une solution, alors l'application  $f$  est injective. Soient en effet  $a_1, a_2 \in A$  tels que  $f(a_1) = f(a_2)$ . On pose alors  $b = f(a_1)$  et on constate que  $a_1$  et  $a_2$  sont deux solutions de l'équation (1.1). On doit donc avoir  $a_1 = a_2$ .  $\square$

Passons maintenant à la surjectivité, qui correspondra à la deuxième condition pour que la réciproque d'une application soit une application. Considérons encore la figure 1.1. Une deuxième raison pour laquelle la réciproque n'est pas une application est que son domaine n'est pas égal à son ensemble de départ, à savoir  $B$  : On constate que le point

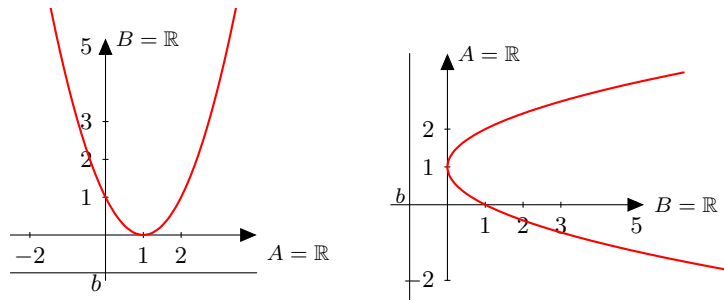


FIGURE 1.3 – Une application non surjective et sa réciproque

$b \in B$  n'est pas dans le domaine de la relation représentée à droite. Cela se voit sur l'application initiale représentée à gauche : il n'existe pas de point  $a \in A$  satisfaisant  $f(a) = b$ . En d'autres termes  $b$  n'est pas dans l'image de la relation  $G_f$ . Cela conduit à la définition.

**Définition 1.5.3.** Soit  $f : A \rightarrow B$  une application. L'image de  $f$ , notée  $Im(f)$  ou  $f(A)$  est égale à l'image de  $G_f$ . L'application  $f$  est surjective<sup>c</sup> si  $Im(f) = B$ .

Donnons tout de suite quelques exemples et contre-exemples.

**Exemple 1.5.2.** 1. L'application  $f_1 : \mathbb{R} \rightarrow \mathbb{R} : x \mapsto x^2$  n'est pas surjective. En effet, le nombre  $-1$  n'est pas dans l'image de  $f_1$ .

c. On dit aussi que  $f$  est une surjection.

2. L'application  $f_2 : \mathbb{R} \rightarrow [0, +\infty[ : x \mapsto x^2$  est surjective. Il s'agit d'une propriété des nombres réels, que vous avez admise, mais que vous serez à même de démontrer à l'issue du cours d'analyse.
3. L'application  $f_3 : \mathbb{N} \rightarrow \mathbb{N}_0 : n \mapsto n+1$  est surjective, puisque tout nombre naturel non nul est le successeur d'un nombre naturel. Nous reverrons sous peu cette propriété.
4. L'application  $f_4 : \mathbb{R} \rightarrow \mathbb{R} : x \mapsto \sin(x)$  n'est pas surjective puisque par exemple, 2 n'est le sinus d'aucun nombre réel.
5. L'application  $f_5 : \mathbb{R} \rightarrow [-1, 1] : x \mapsto \sin(x)$  est surjective. Ici encore, je vous renvoie au cours d'analyse pour une démonstration rigoureuse, mais vous pouvez vous convaincre de ce fait sur le cercle trigonométrique.

Nous avons bien sûr des résultats analogues à ceux concernant les applications injectives.

**Proposition 1.5.5.** *Une application  $f : A \rightarrow B$  est surjective si, et seulement si, on a  $\text{dom}(G_f^{-1}) = B$ .*

*Démonstration.* On sait que  $\text{dom}(G_f^{-1}) = \text{Im}(G_f) = \text{Im}(f)$ . Donc l'assertion de l'énoncé est équivalente à  $\text{Im}(f) = B$ , c'est-à-dire à la surjectivité de  $f$ .  $\square$

La traduction en termes d'équations est également simple.

**Proposition 1.5.6.** *Une application  $f : A \rightarrow B$  est surjective si, et seulement si, pour tout  $b \in B$ , l'équation*

$$f(x) = b, \quad (x \in A) \tag{1.2}$$

*admet au moins une solution.*

*Démonstration.* Il suffit de traduire la condition de l'énoncé. Le fait que l'équation (2.7) admette une solution est équivalent au fait que  $b$  soit dans l'image de  $f$ . La condition d'existence d'une solution pour tout  $b \in B$  s'écrit donc  $B = \text{Im}(f)$ .  $\square$

Enfin, il est utile de remarquer que l'on peut toujours rendre une application quelconque surjective par restriction de l'ensemble d'arrivée.

**Proposition 1.5.7.** *Pour toute application  $f : A \rightarrow B$ , l'application  $f : A \rightarrow \text{Im}(f) = f(A)$  est surjective.*

*Démonstration.* C'est évident.  $\square$

Passons maintenant à quelques propriétés des bijections. D'après ce que nous venons de voir, on a le résultat suivant.

**Proposition 1.5.8.** *Les assertions suivantes sont équivalentes :*

1. *L'application  $f : A \rightarrow B$  est une bijection ;*
2. *L'application  $f : A \rightarrow B$  est une injection et une surjection.*
3. *L'application  $f : A \rightarrow B$  est telle que  $G_f^{-1}$  est de type application.*

Bien entendu, nous avons déjà démontré toutes les équivalences. La dernière donne lieu à une autre définition.

**Définition 1.5.4.** Soit  $f : A \rightarrow B$  une bijection. La relation réciproque  $G_f^{-1}$  est de type application. On note cette application  $f^{-1} : B \rightarrow A$ , définie par  $G_{f^{-1}} = G_f^{-1}$  et on l'appelle application réciproque de  $f$ <sup>d</sup>.

Vous avez déjà rencontré des applications réciproques. En voici une petite liste.

---

d. On dit aussi application inverse de  $f$ .

- Exemple 1.5.3.**
1. L'application  $f : [0, +\infty[ \rightarrow [0, +\infty[ : x \mapsto x^2$  est une bijection. La réciproque est l'application racine carrée :  $\sqrt{\cdot} : [0, +\infty[ \rightarrow [0, +\infty[ : y \mapsto \sqrt{y}$ .
  2. L'application sinus  $\sin : [-\frac{\pi}{2}, \frac{\pi}{2}] \rightarrow [-1, 1]$  est une bijection<sup>e</sup>. L'application réciproque est l'application arc sinus  $\arcsin : [-1, 1] \rightarrow [-\frac{\pi}{2}, \frac{\pi}{2}]$ .
  3. L'application cosinus  $\cos : [0, \pi] \rightarrow [-1, 1]$  est une bijection. Sa réciproque est l'application arc cosinus  $\arccos : [-1, 1] \rightarrow [0, \pi]$ .

Il est utile de pouvoir caractériser l'application réciproque.

**Proposition 1.5.9.** *Si  $f : A \rightarrow B$  est une bijection, alors pour tous  $a \in A$ ,  $b \in B$ , on a  $f(a) = b$  si, et seulement si,  $a = f^{-1}(b)$ .*

*Démonstration.* Revenons à la définition des relations réciproques. On a alors les équivalences suivantes

$$b = f(a) \Leftrightarrow aG_fb \Leftrightarrow bG_f^{-1}a \Leftrightarrow bG_{f^{-1}}a \Leftrightarrow a = f^{-1}(b),$$

ce qui achève la preuve.  $\square$

Pour ce qui suit, nous avons besoin de définir l'égalité de deux applications.

**Définition 1.5.5.** Deux applications  $f$  et  $g$  sont égales si

1. Elles ont même domaine  $A$ ;
2. Pour tout  $a \in A$ , on a  $f(a) = g(a)$ .

Pour tout ensemble  $A$ , on définit l'application identique de  $A$ , notée  $\text{id}_A$  par  $\text{id}_A(a) = a$  pour tout  $a \in A$ .

On a alors le résultat suivant.

**Proposition 1.5.10.** *Si  $f : A \rightarrow B$  est une bijection, alors on a  $f^{-1} \circ f = \text{id}_A$ ,  $f \circ f^{-1} = \text{id}_B$  et  $(f^{-1})^{-1} = f$ .*

*Démonstration.* On sait que  $f^{-1} \circ f$  est une application définie sur  $A$ , tout comme  $\text{id}_A$ . Il reste maintenant à calculer  $(f^{-1} \circ f)(a)$  pour tout  $a \in A$  et à montrer que c'est  $a$ . Mais on a  $(f^{-1} \circ f)(a) = b$  si, et seulement si  $f^{-1}(f(a)) = b$ , qui est équivalent à  $f(a) = f(b)$ , ou encore à  $a = b$ , puisque  $f$  est injectif.

La deuxième égalité se démontre comme la première, ou en appliquant la première à  $f^{-1} : B \rightarrow A$ , une fois que l'on a démontré que  $(f^{-1})^{-1} = f$ . Pour montrer que  $f^{-1} : B \rightarrow A$  est une bijection et calculer son inverse, on peut procéder comme plus haut et montrer que pour tout  $a \in A$ , il existe un unique  $b \in B$  tel que  $f^{-1}(b) = a$ . Alors  $b$  sera égal à  $(f^{-1})^{-1}(a)$ . Mais encore une fois par la proposition 1.5.9, l'assertion  $f^{-1}(b) = a$  est équivalente à  $b = f(a)$ , ce qui suffit.  $\square$

**Proposition 1.5.11.** *La composée de deux bijections est une bijection. Plus précisément, si  $f : A \rightarrow B$  et  $g : B \rightarrow C$  sont des bijections, alors  $g \circ f : A \rightarrow C$  est une bijection. De plus, on a la relation  $(g \circ f)^{-1} = f^{-1} \circ g^{-1}$ .*

*Démonstration.* On sait déjà que la composée  $g \circ f : A \rightarrow C$  est une application. Il reste à montrer que pour tout  $c \in C$ , il existe un unique  $a \in A$  tel que  $(g \circ f)(a) = c$ . Si tel est le cas,  $g \circ f$  sera une bijection et l'élément  $a$  satisfaisant cette condition sera par définition  $(g \circ f)^{-1}(c)$ . Mais on a les équivalences suivantes :

$$(g \circ f)(a) = c \Leftrightarrow g(f(a)) = c \Leftrightarrow f(a) = g^{-1}(c) \Leftrightarrow a = f^{-1}(g^{-1}(c)).$$

La première équivalence vient de la caractérisation de la composée d'applications. La deuxième et la troisième découlent de la proposition 1.5.9.  $\square$

<sup>e</sup>. On a restreint l'ensemble d'arrivée de sinus pour le rendre injectif, ce n'est pas la seule façon de faire, comme nous le verrons bientôt. On a également restreint l'ensemble d'arrivée pour rendre cette application surjective.

## 1.6 Images et pré-images de sous-ensembles

Dans cette courte section, nous donnons quelques définitions très importantes dans tous les cours. Il s'agit des images et des pré-images (ou images inverses) de sous-ensembles.

**Définition 1.6.1.** Soit  $f : A \rightarrow B$  une application, soient  $X$  un sous-ensemble de  $A$  et  $Y$  un sous ensemble de  $B$ . Alors l'image de  $X$  par  $f$  est l'ensemble

$$f(X) = \{f(x) : x \in X\}.$$

La pré-image de  $Y$  par  $f$ , ou l'image inverse de  $Y$  par  $f$  est l'ensemble

$$f^{-1}(Y) = \{x \in A : f(x) \in Y\}.$$

L'ensemble  $f(X)$  est donc constitué des images par  $f$  de tous les points de  $X$ , tandis que  $f^{-1}(Y)$  est l'ensemble de tous les points dont l'image appartient à  $Y$ .

*Remarque 1.5.* 1. La notation  $f^{-1}(Y)$  est dangereuse elle pourrait vous faire penser qu'il faut que  $f$  soit une bijection pour que cet ensemble soit défini. Ce n'est pas le cas : cet ensemble existe toujours, et ce n'est en général pas l'image de  $Y$  par  $f^{-1}$ , qui n'existe que si  $f$  est bijectif.

2. Nous avons associé à une application  $f : A \rightarrow B$  deux nouvelles applications :  $f : \mathcal{P}(A) \rightarrow \mathcal{P}(B) : X \mapsto f(X)$  et  $f^{-1} : \mathcal{P}(B) \rightarrow \mathcal{P}(A) : Y \mapsto f^{-1}(Y)$ .

Voici quelques exemples.

**Exemple 1.6.1.** 1. Soit  $f : \mathbb{R} \rightarrow \mathbb{R} : x \mapsto \sin(x)$ , alors  $f([0, \frac{\pi}{2}]) = \{\sin(x) : x \in [0, \frac{\pi}{2}]\} = [0, 1]$ .

2. Pour la même application, on a

$$f^{-1}([0, 1]) = \{x \in \mathbb{R} : \sin(x) \in [0, 1]\} = [0, \pi] \cup [2\pi, 3\pi] \cup \dots = \cup_{k \in \mathbb{Z}} [2k\pi, (2k+1)\pi].$$

3. Pour la même application, on a

$$f^{-1}([\frac{1}{2}, \frac{3}{2}]) = \{x \in \mathbb{R} : \sin(x) \geq \frac{1}{2}\} = \cup_{k \in \mathbb{Z}} [\frac{\pi}{6} + 2k\pi, \frac{5\pi}{6} + 2k\pi].$$

4. Soit  $p_1 : \mathbb{R}^2 \rightarrow \mathbb{R} : (x, y) \mapsto x$ . On a

$$p_1^{-1}([\frac{1}{2}, 1]) = \{(x, y) \in \mathbb{R}^2 : x \in [\frac{1}{2}, 1]\} = [\frac{1}{2}, 1] \times \mathbb{R}.$$

On a le résultat suivant pour le comportement de l'image et de la pré-image vis-à-vis des unions et intersections. Vous pouvez retenir qu'il n'y a que l'image d'une intersection qui ne se comporte pas bien.

**Proposition 1.6.1.** Soit  $f : A \rightarrow B$  une application. On a alors

1.  $f^{-1}(Y \cup Z) = f^{-1}(Y) \cup f^{-1}(Z)$ , pour tous  $Y, Z \subset B$  ;
2.  $f^{-1}(Y \cap Z) = f^{-1}(Y) \cap f^{-1}(Z)$ , pour tous  $Y, Z \subset B$
3.  $f(Y \cup Z) = f(Y) \cup f(Z)$ , pour tous  $Y, Z \subset A$  ;
4.  $f(Y \cap Z) \subset f(Y) \cap f(Z)$ , pour tous  $Y, Z \subset A$ .

En général, la dernière inclusion est stricte.

*Démonstration.* On procède par double inclusion. Montrons par exemple la première égalité. Soit  $a \in f^{-1}(Y \cup Z)$ . Par définition, le point  $f(a)$  appartient à  $Y \cup Z$ . Si  $f(a) \in Y$ , alors  $a \in f^{-1}(Y) \subset f^{-1}(Y) \cup f^{-1}(Z)$ . Si  $f(a) \in Z$ , alors  $a \in f^{-1}(Z) \subset f^{-1}(Y) \cup f^{-1}(Z)$ .

On montre l'autre inclusion : soit  $a \in f^{-1}(Y) \cup f^{-1}(Z)$ . Si  $a \in f^{-1}(Y)$ , alors  $f(a) \in Y \subset Y \cup Z$  et donc  $a$  appartient à  $f^{-1}(Y \cup Z)$ . On fait de même si  $a \in f^{-1}(Z)$ .

Montrons la dernière inclusion. Si  $b$  appartient à  $f(Y \cap Z)$ , alors il existe  $a \in Y \cap Z$  tel que  $b = f(a)$ . Mais puisque  $a$  appartient à  $Y$ ,  $b$  appartient à  $f(Y)$ . De même, puisque  $a \in Z$ ,  $b$  appartient à  $f(Z)$ . Au total,  $b$  appartient à  $f(Y) \cap f(Z)$ .

Le deux autres égalités se montrent de manière analogue. □

On peut montrer sur un exemple simple que la dernière inclusion de la proposition précédente est stricte. Considérons la projection  $p_1 : \mathbb{R}^2 \rightarrow \mathbb{R}$ . Soient les sous ensembles  $X = [0, 1] \times [0, 1]$  et  $Y = [0, 1] \times [2, 3]$  dans  $\mathbb{R}^2$ . Ils sont disjoints : on a  $X \cap Y = \emptyset$ , donc  $p_1(X \cap Y) = p_1(\emptyset) = \emptyset$ . Par contre  $p_1(X) = p_1(Y) = [0, 1]$ , donc  $p_1(X) \cap p_1(Y) = [0, 1]$ .

Etant donné  $f : A \rightarrow B$ , on peut également se demander comment se composent les applications  $f$  et  $f^{-1}$  définies entre  $\mathcal{P}(A)$  et  $\mathcal{P}(B)$ . La notation suggère qu'elles sont inverses l'une de l'autre, mais ce n'est pas le cas. Voici un résultat cependant utile.

**Proposition 1.6.2.** *Soit  $f : A \rightarrow B$  une application. Pour tout  $X \subset A$  et tout  $Y \subset B$ ,*

1. *On a  $X \subset f^{-1}(f(X))$ , l'égalité ayant lieu notamment si  $f$  est injectif;*
2. *On a  $f(f^{-1}(Y)) \subset Y$ , l'égalité ayant lieu notamment si  $f$  est surjectif.*

*Démonstration.* On montre la première inclusion de manière classique : soit  $x \in X$  et montrons que  $x \in f^{-1}(f(X))$ . Par définition, cela est vrai si  $f(x) \in f(X)$ . Mais cette dernière assertion est vraie par définition de  $f(X)$ , puisque  $x$  appartient à  $X$ . Supposons maintenant  $f$  injectif et montrons l'autre inclusion. Soit  $a \in f^{-1}(f(X))$ . Par définition,  $f(a)$  appartient à  $f(X)$ . Il existe donc  $x \in X$  tel que  $f(a) = f(x)$ . Mais puisque  $f$  est injectif, cela implique  $a = x$ , donc  $a \in X$ .

Pour la deuxième inclusion, on considère  $z \in f(f^{-1}(Y))$ . Il existe  $x \in f^{-1}(Y)$  tel que  $z = f(x)$ . Mais puisque  $x$  appartient à  $f^{-1}(Y)$ , on doit avoir  $f(x) \in Y$ . Donc  $z$  appartient à  $Y$ . Supposons  $f$  surjectif et montrons l'autre inclusion. Soit  $y \in Y$ . Puisque  $f$  est surjectif, il existe  $a \in A$  tel que  $f(a) = y$ . Puisque  $f(a) \in Y$ , le point  $a$  appartient à  $f^{-1}(Y)$ , donc  $y$  est l'image d'un point de  $f^{-1}(Y)$  et appartient donc à  $f(f^{-1}(Y))$ .  $\square$

## Chapitre 2

# Nombres complexes

### 2.1 Introduction

Dans ce chapitre, nous commençons par une révision rapide des équations du second degré dans  $\mathbb{R}$  et des quelques résultats associés. Nous en profitons pour illustrer une technique, la complétion des carrés, qui est utile dans d'autres domaines. Nous constatons que certaines équations du second degré (et de degré supérieur) n'ont pas de solution dans  $\mathbb{R}$ , parce que les nombres négatifs ne sont pas des carrés. Nous faisons alors une extension de l'ensemble des nombres, et des opérations, pour obtenir les nombres complexes. Comme cela a été le cas dans le premier chapitre, nous construisons ces nombres à partir de ceux que nous connaissons<sup>a</sup>. Cette construction est importante car elle permet de prouver l'existence d'un ensemble de nombres ayant les propriétés voulues, mais il est encore plus important de pouvoir calculer avec les nombres complexes, et la construction classique de ces nombres ne permet pas de calculer facilement. Heureusement, nous trouvons deux façons de se donner des nombres complexes qui permettent de calculer facilement. Au passage, elles permettent, via l'exponentielle complexe, de récupérer facilement des résultats de trigonométrie, via par exemple, la formule de Moivre.

On pourrait croire que c'est encore un chapitre qui ne sert pas à grand chose et se demander si c'est bien raisonnable d'étendre encore les ensembles de nombres, après tout, si certaines équations n'ont pas de solution, elles n'ont pas de solution, un point c'est tout. Vous verrez dans la suite du cours d'algèbre que le travail avec les nombres complexes est paradoxalement plus simple qu'avec les nombres réels. Cela est dû au fait que non seulement toutes les équations du second degré auront deux solutions complexes (éventuellement confondues), mais en fait toutes les équations polynomiales de degré  $n \geq 1$  admettent  $n$  solutions complexes (comptées avec leurs multiplicités). Pour les mathématiciens, les nombres complexes sont donc ceux qui permettent d'obtenir les résultats les plus simples et d'apporter des solutions efficaces à de nombreux problèmes. Pour les physiciens, ils sont la pierre angulaire de bon nombre de théories. Citons par exemple cette équation qui a fait couler pas mal d'encre au 20e siècle :

$$i\hbar \frac{d}{dt} |\Psi(t)\rangle = \hat{H} |\Psi(t)\rangle.$$

Elle commence par l'unité complexe  $i$ . Vous la découvrirez bientôt, sous cette forme ou sous une autre. Plus généralement, la mécanique quantique est construite sur des espaces de Hilbert, fait appel à des opérateurs linéaires hermitiens, et tout cela ne peut être construit sans nombres complexes. Tout cela répond à la question que vous vous posiez certainement : *“Quand est-ce que j'en aurai fini avec les nombres complexes ?”*... Jamais. Alors autant s'y mettre tout de suite.

---

a. Nous nous sommes entendus sur le fait que je ne définirai pas les nombres réels. Cela fait partie des cours d'analyse, mais vous les connaissez raisonnablement.

## 2.2 Equations du second degré et complétion des carrés

Les équations du second degré sont des équations ayant la forme générale :

$$ax^2 + bx + c = 0, \quad a \in \mathbb{R}_0, b, c \in \mathbb{R}. \quad (2.1)$$

Ces équations sont très simples à résoudre quand  $b = 0$  : par exemple  $x^2 - 7 = 2$  a pour ensemble de solutions  $\{-3; 3\}$ , comme chacun sait. Mais attardons nous quelques secondes sur les étapes que l'on peut parcourir pour arriver à cette conclusion. Nous avons en fait procédé par équivalences :

$$x^2 - 7 = 2 \Leftrightarrow x^2 - 7 - 2 = 2 + (-2) \Leftrightarrow x^2 - 9 = 0 \Leftrightarrow (x+3)(x-3) = 0 \Leftrightarrow x+3 = 0 \text{ ou } x-3 = 0 \dots$$

Les équivalences (vraies) ont le même sens qu'au premier chapitre : elles indiquent que les deux assertions ont même valeur de vérité, c'est-à-dire que ces équations ont le même ensemble de solutions. Mais nous avons utilisé, pour cette simple équation, que 2 admet un *opposé* pour l'addition ; cet opposé est le nombre qui permet d'avoir le *neutre* pour l'addition, à savoir 0. Nous avons, sans l'écrire, déplacé les parenthèses, donc utilisé l'*associativité* de l'addition. Ensuite, nous avons utilisé l'identité remarquable  $(a+b)(a-b) = a^2 - b^2$ , qui résulte facilement du fait que la *multiplication distribue l'addition*, et du fait que la multiplication est *commutative*. Ensuite, nous avons utilisé la "règle du produit nul" ( $ab = 0$ ssi  $a = 0$  ou  $b = 0$ ), qui se déduit facilement du fait que *tout nombre non nul admet un inverse* (et de la distributivité, qui rend 0 absorbant). En bref, nous avons utilisé pratiquement toutes les propriétés bien connues des nombres réels.

Voici un deuxième exemple encore assez simple : on considère l'équation suivante (pour  $x \in \mathbb{R}$ )

$$3(x+5)^2 = 21. \quad (2.2)$$

On a la même résolution que ci-dessus :

$$3(x+5)^2 = 21 \Leftrightarrow 3(x+5)^2 - 21 = 0 \Leftrightarrow 3[(x+5)^2 - 7] = 0 \Leftrightarrow 3[(x+5) + \sqrt{7}][(x+5) - \sqrt{7}] = 0$$

Donc l'ensemble des solutions est  $S = \{-5 + \sqrt{7}, -5 - \sqrt{7}\}$ . Evidemment, vous me direz que je ne prends que des équations simples, qui mènent tout de suite à une factorisation car on a une différence de deux carrés. Je suis bien d'accord, mais en développant l'équation (2.2), nous venons en fait de résoudre

$$3x^2 + 30x + 54 = 0. \quad (2.3)$$

Tout le problème consiste à passer de l'équation (2.3), contenant un terme du premier degré en  $x$  à l'équation (2.2). C'est en fait assez simple : on met d'abord le facteur 3 en évidence et on se ramène à  $3(x^2 + 10x + 18) = 0$ . La seule façon de grouper  $x^2$  et  $10x$  pour former un carré est que  $10x$  soit un double produit (c'est  $2.5.x$ ), donc on souhaite avoir  $x^2 + 10x + 25$ , et on écrit donc

$$3(x^2 + 10x + 18) = 3(x^2 + 10x + 18 + 7 - 7) = 3[(x^2 + 2.5.x + 25) - 7] = 3[(x+5)^2 - 7].$$

On a donc complété  $x^2 + 10x$  pour obtenir un carré parfait et faire ainsi disparaître le terme du premier degré qui nous posait problème.

Dans le cas général, on procède de la même façon pour faire "disparaître" ce même terme dans le trinôme  $ax^2 + bx + c$ . On peut toujours le faire en **complétant le carré** :

$$\begin{aligned} ax^2 + bx + c &= a\left(x^2 + \frac{b}{a}x + \frac{c}{a}\right) = a\left(x^2 + 2\frac{b}{2a}x + \frac{c}{a}\right) = a\left(x^2 + 2\frac{b}{2a}x + \left(\frac{b}{2a}\right)^2 - \left(\frac{b}{2a}\right)^2 + \frac{c}{a}\right) \\ &= a\left(\left(x + \frac{b}{2a}\right)^2 - \frac{b^2 - 4ac}{4a^2}\right). \end{aligned} \quad (2.4)$$

Selon le signe de  $b^2 - 4ac$  on a une différence de deux carrés ou une somme de deux carrés. Ce nombre détermine donc l'existence des solutions. On l'appelle le *réalisant* ou *discriminant* de l'équation, et on le note  $\rho$  ou  $\Delta$  (je choisirai  $\Delta$ ).

La forme (2.4) du trinôme  $ax^2 + bx + c$  obtenue ci-dessus permet alors, via la factorisation et la "règle du produit nul" d'obtenir tous les résultats ci-dessous. Elle sera également utile pour étudier les fonctions du second degré.

**Proposition 2.2.1.** *La structure des solutions de l'équation (2.1) dépend du signe de  $\Delta$  :*

1. si  $\Delta > 0$ , l'équation admet deux solutions distinctes : on a

$$S = \left\{ \frac{-b - \sqrt{\Delta}}{2a}, \frac{-b + \sqrt{\Delta}}{2a} \right\}.$$

2. si  $\Delta = 0$ , l'équation admet une solution unique :  $S = \left\{ \frac{-b}{2a} \right\}$ . La solution  $\frac{-b}{2a}$  est dite double.

3. si  $\Delta < 0$ , l'équation est incompatible :  $S = \emptyset$ .

*Démonstration.* Si  $\Delta$  est positif ou nul, alors c'est un carré, le carré de  $\sqrt{\Delta}$ , donc on a

$$ax^2 + bx + c = a \left[ \left( x + \frac{b}{2a} \right)^2 - \frac{(\sqrt{\Delta})^2}{4a^2} \right] = a \left[ \left( x + \frac{b}{2a} \right)^2 - \left( \frac{\sqrt{\Delta}}{2a} \right)^2 \right] = a \left( x + \frac{b}{2a} + \frac{\sqrt{\Delta}}{2a} \right) \left( x + \frac{b}{2a} - \frac{\sqrt{\Delta}}{2a} \right),$$

et on conclut via la règle du produit nul.

On constate que les deux solutions sont distinctes si  $\Delta > 0$  et qu'elles se confondent (d'où l'appellation solution double) si  $\Delta = 0$ .

Si  $\Delta$  est strictement négatif, alors on a quand-même

$$ax^2 + bx + c = a \left( \left( x + \frac{b}{2a} \right)^2 - \frac{\Delta}{4a^2} \right).$$

Le premier terme  $\left( x + \frac{b}{2a} \right)^2$  est positif ou nul, tandis que le deuxième,  $-\frac{\Delta}{4a^2}$ , est strictement positif, quel que soit  $x$ . Donc la somme est strictement positive, pour tout  $x$ , et n'est donc jamais nulle.  $\square$

Tant que nous y sommes, rappelons un théorème sur la somme et le produit des solutions. Nous avons déjà fait une partie de la preuve, pour trouver les solutions des équations du second degré.

**Proposition 2.2.2.** *Si  $\Delta \geq 0$ , alors l'équation du second degré  $ax^2 + bx + c = 0$  admet les solutions  $x_1$  et  $x_2$  (éventuellement égales). De plus, le trinôme correspondant se factorise :*

$$ax^2 + bx + c = a(x - x_1)(x - x_2) \quad \forall x \in \mathbb{R}.$$

Enfin, on a  $x_1 + x_2 = \frac{-b}{a}$  et  $x_1 x_2 = \frac{c}{a}$ .

Cette proposition permet par exemple de vérifier les solutions que l'on a trouvé à peu de frais : on peut calculer leur produit et comparer à  $\frac{c}{a}$ . Ce n'est pas une vérification exacte, mais elle permet de détecter des erreurs : si vous avez fait une faute de calcul, il y a peu de chances que le produit de vos solutions donne quand-même la bonne valeur. Ce résultat permet aussi de trouver la deuxième solution si on connaît la première, en ramenant l'équation du second degré à une équation du premier degré. Par exemple, sachant que l'équation  $x^2 - 5x + 6 = 0$  admet  $x_1 = 2$  comme solution, la deuxième solution  $x_2$  satisfait  $2x_2 = 6$ , donc c'est 3.

*Démonstration.* La factorisation est acquise dans la preuve de la proposition 2.2.1. Pour la somme et le produit des solutions, on peut procéder à partir de leur expression :  $\frac{-b \pm \sqrt{\Delta}}{2a}$ . Dans la somme, les racines se simplifient, et on a donc  $\frac{-2b}{2a}$ . Pour le produit, on a

$$x_1 x_2 = \frac{1}{4a^2} (-b - \sqrt{\Delta})(-b + \sqrt{\Delta}) = \frac{1}{4a^2} (b^2 - \Delta) = \frac{c}{a},$$

et la preuve est terminée.  $\square$

Ce n'est cependant pas la preuve la plus instructive, alors j'en donne une deuxième.

*Démonstration.* Puisque le trinôme se factorise, on a

$$ax^2 + bx + c = a(x - x_1)(x - x_2) = ax^2 - a(x_1 + x_2)x + ax_1x_2,$$

quel que soit  $x$ . En identifiant les coefficients de chaque degré, on obtient directement  $b = -a(x_1 + x_2)$  et  $c = ax_1x_2$ .  $\square$

Cette preuve est efficace, mais elle soulève une question : comment sait-on que l'on peut identifier les coefficients, quand deux fonctions polynomiales sont égales ? Autrement dit, peut-on affirmer que  $ax^2 + bx + c = a'x^2 + b'x + c'$  (pour tout  $x \in \mathbb{R}$ ) implique  $a = a'$ ,  $b = b'$  et  $c = c'$  ? En soustrayant le membre de droite, on voit que cette question est équivalente à la suivante : l'assertion " $(a - a')x^2 + (b - b')x + (c - c') = 0$ , pour tout  $x \in \mathbb{R}$ " implique-t-elle  $a - a' = b - b' = c - c' = 0$  ? On peut encore écrire cette implication sous la forme

$$a''x^2 + b''x + c'' = 0 \quad \forall x \in \mathbb{R} \Rightarrow a'' = b'' = c'' = 0.$$

C'est un problème d'indépendance linéaire, comme vous le verrez dans la suite du cours d'algèbre. Dans notre cas, il y a plusieurs façons de se convaincre que l'implication est vraie. Par exemple, en exprimant que l'équation est satisfaite pour trois valeurs de  $x$  (0, 1 et  $-1$ ), on obtient un système d'équations en  $a''$ ,  $b''$  et  $c''$  qui n'admet que la solution nulle. On peut aussi se souvenir qu'une vraie équation du second degré admet au plus 2 solutions. Donc celle-ci doit être une "fausse" équation du second degré : on doit avoir  $a'' = 0$ , mais alors on a  $b''x + c'' = 0$  pour tout  $x \in \mathbb{R}$ , et on trouve de même  $b'' = 0$ , puis  $c'' = 0$ . On peut aussi utiliser les dérivées : si une fonction est nulle sur  $\mathbb{R}$ , elle est nulle en 0, et sa dérivée aussi sur  $\mathbb{R}$ , donc en 0, ... Ces arguments se généralisent à des fonctions polynomiales de degré supérieur à 2.

Cette proposition admet une réciproque, que je cite pour être complet.

**Proposition 2.2.3.** *Si  $n_1$  et  $n_2$  sont deux nombres dont la somme est  $s$  et le produit  $p$ , alors ces nombres sont solutions de l'équation*

$$x^2 - sx + p = 0.$$

*Démonstration.* Il suffit d'exprimer les conditions sur  $n_1$  et  $n_2$  ( $n_1 + n_2 = s$  et  $n_1n_2 = p$ ) et de résoudre ce système d'équations.  $\square$

Pour terminer cette section, voici quelques utilisations possibles de la complétion des carrés :

1. Dans un repère orthonormé, l'équation du cercle de centre  $C : (c_1, c_2)$  et de rayon  $r$  (pour la distance euclidienne) est donnée par

$$(x - c_1)^2 + (y - c_2)^2 = r^2,$$

ou un multiple non nul. On peut ramener un certain nombre d'équations du second degré à cette forme et ainsi trouver le centre et le rayon du cercle d'équation

$$4x^2 + 4y^2 + 6x - 12y - 25 = 0.$$

De manière générale, on peut appliquer cette méthode à l'équation

$$ax^2 + ay^2 + bx + cy + d = 0 \quad (a \neq 0).$$

2. Si ne demande pas d'avoir un cercle, on peut appliquer la même méthode à

$$4x^2 + 2y^2 + 6x - 12y - 25 = 0.$$

On obtient l'équation canonique d'une ellipse dont les axes sont parallèles à ceux du repère.

3. On peut même avoir des doubles produits en  $xy$ , c'est la même méthode ( $y$  est un nombre, après tout) :

$$4x^2 + 4xy + 2y^2 + 6x - 12y - 25 = 0.$$

4. La complétion des carrés permet également de factoriser dans  $\mathbb{R}$  la fonction polynomiale  $P(x) = x^4 + 4$ , bien que celle-ci n'admette pas de zéros dans  $\mathbb{R}$ . On doit ici compléter le carré en ajoutant et en soustrayant le double produit adéquat.
5. Enfin, on peut aussi démontrer que l'expression  $Q(x_1, x_2) = 7x_1^2 + 2x_2^2 + 4x_1x_2$  est positive ou nulle, quels que soient  $x_1, x_2 \in \mathbb{R}$ , en écrivant cette expression comme une somme de deux carrés.

## 2.3 Nombres complexes, introduction et définition

Nous venons de revoir les équations du second degré dans  $\mathbb{R}$ . Nous savons donc que l'équation

$$x^2 + 4x + 8 = 0 \tag{2.5}$$

n'admet pas de solution. Cela peut se voir de deux façons : d'une part en calculant le réalisant du trinôme du second degré en question :  $\Delta = 16 - 32 = -16 < 0$ , d'autre part en notant que l'équation s'écrit encore

$$(x + 2)^2 + 4 = 0 \quad \text{ou encore} \quad (x + 2)^2 - (-4) = 0.$$

Si nous voulons résoudre cette équation comme pour les équations réelles, cela revient, d'après la dernière formule à trouver un (ou plusieurs) nombre(s)  $z$  tel que  $z^2 = -4$ . On sait que ce n'est pas possible dans les nombres réels, puisque le carré de tout nombre réel est positif ou nul. La solution consiste à *étendre* l'ensemble des réels en ajoutant un nombre supplémentaire le nombre  $i$  tel que

$$i^2 = -1.$$

Cela vous paraît abstrait ? Absurde ? Pourtant, cela ne vous a pas tracassé que l'on ajoute un nombre aux rationnels, que l'on a appelé  $\sqrt{2}$  pour résoudre l'équation  $(x + 2)^2 - 2 = 0$ . On aurait pu appeler ce nombre  $j$ , et demander que  $j^2 = 2$ <sup>a</sup>. Cela vous semble naturel, mais cela ne l'était pas pour Pythagore, par exemple. On pourrait alors obtenir

$$(2i)^2 = (2i)(2i) = -4,$$

moyennant le fait qu'on puisse définir le nombre  $2i$  et faire le produit de deux tels nombres, avec les règles habituelles. On aurait alors également

$$(-2i)^2 = -4$$

et donc deux solutions à l'équation (2.5), à savoir :

$$x = -2 + 2i \quad \text{et} \quad x = -2 - 2i.$$

Pour que ces nombres soient solutions de l'équation (2.5), encore faut-il que l'on puisse calculer  $(-2 + 2i)^2$ , comme d'habitude (avec un produit bien défini, commutatif et associatif,...). Cela permettra de calculer  $(-2 + 2i)^2 + 4(-2 + 2i) + 8$  et de montrer que c'est 0. Nous avons déjà remarqué toutes les propriétés nécessaires à la résolution des équations

---

a. La construction que je vais détailler ci-dessous s'adapte pour définir le champ  $\mathbb{Q}(\sqrt{2})$ , qui est l'ensemble des nombres de la forme  $a + b\sqrt{2}$ , où  $a$  et  $b$  sont rationnels.

du premier et du second degré dans  $\mathbb{R}$ . Il serait de bon ton que ces propriétés soient également satisfaites pour les nombres de la forme  $a + ib$ . Si ces propriétés étaient satisfaites, on aurait alors

$$(a+ib)+(a'+ib') = (a+a')+i(b+b'), (a+ib)(a'+ib') = aa'+bb'i^2+(ab'+a'b)i = (aa'-bb')+(ab'+a'b)i, \quad (2.6)$$

par associativité de  $+$ , distributivité et commutativité de l'addition et de la multiplication.

On pourrait alors donner comme définition des nombres complexes

$$\mathbb{C} = \{a + ib : a, b \in \mathbb{R}\},$$

où les opérations sont définies par (2.6), et où  $i^2 = -1$ . Mais si on procède de la sorte, comme sait-on qu'un tel nombre  $i$  existe, quel statut donner au  $+$  entre  $a$  et  $ib$ , comment  $ib$  est-il d'ailleurs défini ? L'écriture  $a + ib$  est-elle unique ? Il y a plusieurs façon de présenter une définition qui répond à toutes ces questions. J'en choisis une qui me paraît intuitive à ce stade : dans  $a + ib$ , je ne sais ni ce qu'est ce  $+$ , ni ce qu'est ce  $i$ , tout ce que je sais, c'est qu'il y a  $a$  et  $b$ , donc j'écris plutôt  $(a, b)$ , en pensant très fort que je veux que ce soit  $a + ib$ . Je vais ensuite définir la somme et le produit pour obtenir ce qui est écrit à l'équation (2.6), mais en utilisant les couples au lieu d'expressions  $a + ib$  pas encore définies. J'arrive donc à la définition.

**Définition 2.3.1.** On définit l'ensemble des nombres complexes par

$$\mathbb{C} = \{(a, b) : a, b \in \mathbb{R}\}.$$

En tant qu'ensemble, on a donc  $\mathbb{C} = \mathbb{R}^2$ . On passe maintenant aux opérations en gardant en tête l'équation (2.6).

**Définition 2.3.2.** L'addition des nombres complexes est définie par

$$+ : \mathbb{C} \times \mathbb{C} \rightarrow \mathbb{C} : ((a, b), (a', b')) \mapsto (a + a', b + b'),$$

et la multiplication est donnée par

$$\cdot : \mathbb{C} \times \mathbb{C} \rightarrow \mathbb{C} : ((a, b), (a', b')) \mapsto (aa' - bb', ab' + a'b),$$

pour tous  $a, a', b, b' \in \mathbb{R}$ .

Les nombres complexes sont donc simplement des couples de nombres réels. l'addition est bien connue : c'est celle des composantes de vecteurs du plan. Seule la multiplication est nouvelle. Nous vérifions maintenant que l'addition et la multiplication munissent  $\mathbb{C}$  d'une structure de champ. Nous découpons ces propriétés en celles relatives à l'addition, celles de la multiplication et la distributivité.

**Proposition 2.3.1.** L'addition des nombres complexes a les propriétés suivantes :

1. elle est associative : on a  $(z + z') + z'' = z + (z' + z'') \quad \forall z, z', z'' \in \mathbb{C}$  ;
2. elle admet un neutre  $e = (0, 0)$  : on a  $z + e = e + z = z \quad \forall z \in \mathbb{C}$  ;
3. tout élément admet un opposé :  $\forall z \in \mathbb{C}, \exists z' \in \mathbb{C} : z + z' = z' + z = e$ .
4. elle est commutative : on a  $z + z' = z' + z, \quad \forall z, z' \in \mathbb{C}$ .

*Démonstration.* C'est une simple vérification. On a par exemple

$$((x, y) + (x', y')) + (x'', y'') = (x + x', y + y') + (x'', y'') = ((x + x') + x'', (y + y') + y'')$$

et

$$(x, y) + ((x', y') + (x'', y'')) = (x, y) + (x' + x'', y' + y'') = (x + (x' + x''), y + (y' + y'')).$$

On est donc ramené à l'associativité de la somme des nombres réels. De la même façon, on a

$$(x, y) + (0, 0) = (x + 0, y + 0) = (x, y) = (0, 0) + (x, y),$$

car 0 est neutre pour l'addition dans  $\mathbb{R}$ . On vérifie ensuite que l'opposé du nombre complexe  $(x, y)$  est  $(-x, -y)$ , et on montre la commutativité comme plus haut. Dans chaque cas, on est ramené à la propriété correspondante dans  $\mathbb{R}$ .  $\square$

**Remarque 2.1.** a) L'élément  $e = (0, 0)$  est plus simplement noté 0, l'opposé d'un élément  $z \in \mathbb{C}$  est noté  $-z$ .

b) Cette proposition signifie que la structure mathématique  $(\mathbb{C}, +, 0)$  est un *groupe* (les propriétés (1), (2) et (3)) *commutatif* (la propriété (4)).

On a des propriétés analogues pour la multiplication, qui montrent qu'elle se comporte comme la multiplication des nombres réels.

**Proposition 2.3.2.** *La multiplication des nombres complexes a les propriétés suivantes :*

1. elle est associative : on a

$$(z.z').z'' = z.(z'.z'') \quad \forall z, z', z'' \in \mathbb{C};$$

2. elle admet un neutre  $1 = (1, 0)$  : on a

$$1.z = z.1 = z \quad \forall z \in \mathbb{C};$$

3. tout élément non nul admet un inverse :

$$\forall z \in \mathbb{C}^* = \mathbb{C} \setminus \{0\}, \exists z' \in \mathbb{C} : z.z' = z'.z = 1.$$

4. elle est commutative : on a

$$z.z' = z'.z, \quad \forall z, z' \in \mathbb{C}.$$

*Démonstration.* Les points 1., 2., et 4. se démontrent comme plus haut et sont laissés à titre d'exercice. Seule l'existence d'un inverse demande une justification. Soit  $z = (x, y) \neq (0, 0)$ . Un inverse de  $z$  est alors  $z' = (x', y')$  satisfaisant

$$\begin{cases} xx' - yy' = 1 \\ yx' + xy' = 0 \end{cases}$$

On résout le système, par exemple par la méthode du pivot (il faut discuter si  $x = 0$  ou  $y = 0$ ), et on trouve que l'inverse de  $z = (x, y)$  est  $z' = (\frac{x}{x^2+y^2}, -\frac{y}{x^2+y^2})$ .  $\square$

**Remarque 2.2.** a) L'inverse d'un élément non nul  $z$  est généralement noté  $z^{-1}$  ou  $\frac{1}{z}$ . Cela permet de définir la division des nombres complexes : si  $z' \neq 0$ ,  $\frac{z}{z'}$  est le nombre par lequel il faut multiplier  $z'$  pour obtenir  $z$ . C'est donc  $zz'^{-1}$ . Cela permet aussi de vérifier qu'on peut "multiplier en haut et en bas :  $\frac{z}{z'} = \frac{wz}{wz'}$ , pour tout  $w \neq 0$ .

Une dernière propriété lie l'addition et la multiplication des nombres complexes. La démonstration est ici encore une simple vérification et est laissée comme exercice.

**Proposition 2.3.3.** *La multiplication des nombres complexes distribue l'addition : on a*

$$z.(z' + z'') = z.z' + z.z'' = (z' + z'').z,$$

pour tous  $z, z', z'' \in \mathbb{C}$ .

**Remarque 2.3.** Les neuf propriétés que nous venons de lister dans les trois propositions ci-dessus sont résumées en une phrase, "La structure  $(\mathbb{C}, +, 0, \cdot, 1)$  est un *champ*". Elles signifient qu'on peut additionner et multiplier les nombres complexes **comme on a l'habitude de le faire avec les nombres réels**.

En particulier les "produits remarquables" que nous avons vus jusqu'à présent, ou encore la théorie des systèmes linéaires, sont exactement identiques, que l'on travaille avec des nombres réels ou des nombres complexes.

## 2.4 Plongement de $\mathbb{R}$ , nombre $i$ , représentation, forme algébrique et nombres associés

Lors des extensions précédentes, nous avons écrit des inclusions  $\mathbb{N} \subset \mathbb{Z} \subset \mathbb{Q} \subset \mathbb{R}$ . En fait, ces inclusions ne sont obtenues qu'à partir d'identifications : les éléments de  $\mathbb{Z}$  s'écrivent  $+n$  ou  $-n$ , pour  $n \in \mathbb{N}$ . On identifie alors  $n$  à  $+n$ . De même, les rationnels sont définis comme des fractions (à équivalence près)  $\frac{a}{b}$ . Les nombres entiers sont des fractions :  $3 = \frac{3}{1}$ . On identifie alors ces deux nombres et on plonge  $\mathbb{Z}$  dans  $\mathbb{Q}$ . Pour obtenir l'inclusion  $\mathbb{R} \subset \mathbb{C}$ , il faut faire de même. On connaît l'identification parce que c'est celle utilisée en géométrie analytique : quand on a deux axes dans le plan, on porte des nombres sur les axes, et pas des couples de nombres (par exemple, on note 1 pour l'unité sur l'axe des abscisses au lieu de  $(1, 0)$ ). Pourtant ces points sont sur les axes, mais aussi dans le plan.

**Définition 2.4.1.** Le plongement de  $\mathbb{R}$  dans  $\mathbb{C}$  est défini par

$$j : \mathbb{R} \rightarrow \mathbb{C} : x \mapsto (x, 0).$$

Dans la suite, on ne notera plus cette identification, et on écrira  $x$  pour  $(x, 0)$ , exactement comme dans  $\mathbb{Q}$ , on écrit 3 au lieu de  $\frac{3}{1}$ . On peut alors écrire  $\mathbb{R} \subset \mathbb{C}$ . Il est cependant important de vérifier que cette identification a les propriétés adéquates : des nombres réels différents ne peuvent pas coïncider une fois qu'on les considère dans  $\mathbb{C}$ , et la multiplication et l'addition de deux nombres réels ne doit pas dépendre du fait que l'on considère ces nombres comme appartenant à  $\mathbb{R}$  ou à  $\mathbb{C}$ . C'est l'objet du résultat suivante.

**Proposition 2.4.1.** *L'application  $j$  est injective. Elle définit donc une bijection de  $\mathbb{R}$  sur son image. De plus,*

1.  $j(x + x') = j(x) + j(x')$  pour tous  $x, x' \in \mathbb{R}$  ;
2.  $j(xx') = j(x).j(x')$  pour tous  $x, x' \in \mathbb{R}$  ;

*Démonstration.* Montrons que  $j$  est injectif. On a  $j(x) = j(x') \Leftrightarrow (x, 0) = (x', 0)$ , et cette condition implique visiblement  $x = x'$ . De même,

$$j(x+x') = (x+x', 0) = (x, 0) + (x', 0) = j(x) + j(x'), \quad \text{et} \quad j(xx') = (xx', 0) = (x, 0) \cdot (x', 0) = j(x).j(x'),$$

pour tous  $x, x' \in \mathbb{R}$ . □

En utilisant cette identification, nous pouvons écrire des expressions du type  $a.z$ , où  $a \in \mathbb{R}$  et  $z \in \mathbb{C}$ . Si  $z = (x, y)$ , on a

$$a.z = (a, 0).(x, y) = (ax - 0y, ay + 0x) = (ax, ay).$$

La première égalité est obtenue en identifiant  $a$  à  $j(a)$ , la seconde en appliquant la définition du produit. La multiplication d'un nombre complexe par un nombre réel correspond donc à la multiplication usuelle des éléments de  $\mathbb{R}^2$  par les nombres réels.

*Remarque 2.4.* Dans la suite nous ne noterons plus le point pour la multiplication.

Il est maintenant grand temps de revenir aux préoccupations que nous avons dans l'introduction de cette section. Nous nous rappelons que nous voulions définir  $a + ib$ , mais que nous avons dû nous rabattre sur le couple  $(a, b)$ . Naturellement, le couple  $(0, 1)$  devrait correspondre au nombre  $0 + i1 = i$ . Cela amène la définition suivante.

**Définition 2.4.2.** Nous notons  $i$  le nombre complexe  $(0, 1)$ . On l'appelle unité imaginaire.<sup>a</sup>

Cette simple définition nous permet d'arriver au but recherché.

<sup>a</sup>. Cette notation remonte à Leonhard Euler (1707-1783).

**Proposition 2.4.2.** *Tout nombre complexe  $z$  s'écrit de manière unique  $x + iy$  ( $x, y \in \mathbb{R}$ ). On a  $i^2 = -1$ .*

*Démonstration.* Par définition, tout nombre complexe  $z$  s'écrit de manière unique  $z = (x, y)$ . De plus,

$$(x, y) = (x, 0) + (0, y) = x(1, 0) + y(0, 1) = x1 + yi = x + iy.$$

On a utilisé la multiplication dans  $\mathbb{C}$  et l'identification de  $\mathbb{R}$  à un sous-ensemble de  $\mathbb{C}$ . L'unicité est évidente : si  $x + iy = x' + iy'$ , alors  $(x, y) = (x', y')$ , donc  $x = x'$  et  $y = y'$ . De plus,  $i^2 = (0, 1)(0, 1) = (-1, 0) = -1$ .  $\square$

**Définition 2.4.3.** Si  $x, y$  sont réels, l'écriture  $x + iy$  est l'écriture sous *forme algébrique* du nombre complexe  $(x, y)$ .

Nous pouvons donc maintenant utiliser cette écriture unique pour calculer avec les nombres complexes. Il suffit de retenir que l'addition et la multiplication munissent  $\mathbb{C}$  d'une structure de champ, et que  $i^2 = -1$ . La proposition suivante insiste bien sur ce mode de calcul, que nous avons espéré dans la relation (2.6).

**Proposition 2.4.3.** *On a*

$$(x + iy) + (x' + iy') = x + x' + i(y + y') \quad \text{et} \quad (x + iy)(x' + iy') = xx' - yy' + i(xy' + x'y)$$

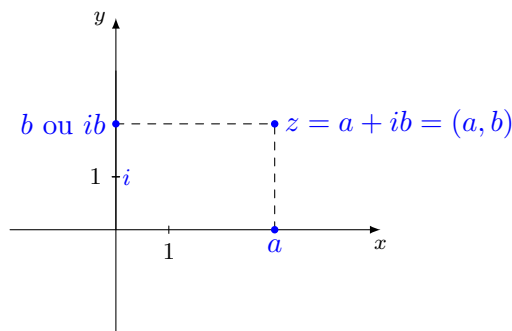
*quels que soient  $x, x', y, y'$  dans  $\mathbb{R}$*

*Démonstration.* L'addition est commutative et associative, de plus la multiplication distribue l'addition, et  $i^2 = -1$ .  $\square$

*Cette proposition est très importante, car elle vous permet d'additionner et multiplier les nombres complexes, comme si vous le faisiez avec les nombres réels, à l'exception du fait que le nombre  $i$  a un statut spécial, et qu'il faut tenir compte de la seule relation  $i^2 = -1$ .*

**Exemple 2.4.1.** Calculer  $5.(3 + 2i)$ ,  $(2 + i)^2$ ,  $(1 + 3i)^2$ ,  $(-2 + 2i)^2 + 4(-2 + 2i) + 8$ .<sup>b</sup>

Enfin, puisque les nombres complexes sont des couples de nombres réels, on peut les représenter dans le plan, alors appelé plan complexe (on parle aussi de plan d'Argand, ou de diagramme d'Argand), comme on a l'habitude de le faire en géométrie analytique :



Remarquons que le nombre  $b$  sur l'axe des *ordonnées* indique que l'on reporte le nombre réel  $b$  sur cet axe gradué, mais le point en question représente le nombre complexe  $ib = (0, b)$ . Cela ne créera pas de confusion pour la suite de noter  $b$  ou  $ib$  sur le deuxième axe. Les nombres correspondants sont les multiples réels de  $i$ . Ils sont appelés imaginaires purs.

Nous définissons ici des nombres naturellement associés à tout nombre complexe.

<sup>b</sup>. Solutions :  $15 + 10i$ ,  $3 + 4i$ ,  $-8 + 6i$ ,  $0$ .

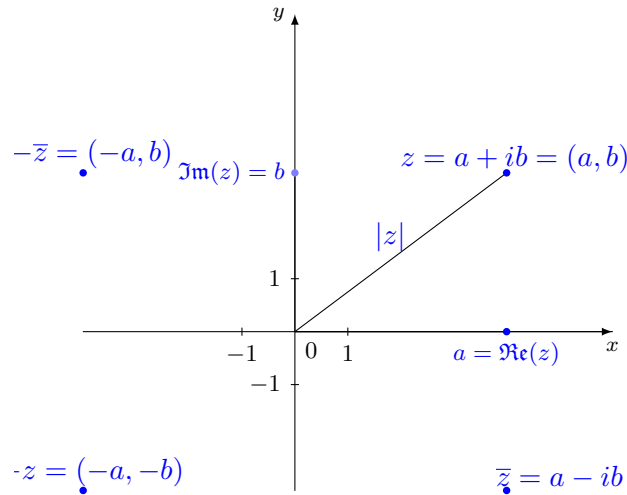
**Définition 2.4.4.** Pour tout nombre complexe  $z = a + ib$  ( $a, b \in \mathbb{R}$ ), on définit

1. la partie réelle de  $z$  par  $\Re(z) = a$ ;
2. la partie imaginaire de  $z$  par  $\Im(z) = b$ ;
3. le nombre conjugué de  $z$  par  $\bar{z} = a - ib$ ;
4. le module de  $z$  par  $|z| = \sqrt{a^2 + b^2} = \sqrt{z \cdot \bar{z}}$ ;

*Remarque 2.5.* Attention, la partie imaginaire de  $z$  est bien  **$b$** , et pas  **$ib$** . Ici on a fait un choix.

Le module n'est pas défini comme la valeur absolue des nombres réels. En effet, il n'y a pas d'ordre naturel sur  $\mathbb{C}$  et la tentative de définition  $\max\{-z, z\}$  n'aurait simplement pas de sens. Cependant, nous savons que la valeur absolue du nombre réel  $x$  mesure la distance de  $x$  à 0. C'est cette propriété qui est généralisée. Cela permet de montrer que le module du nombre réel  $x$  tel que défini ici correspond à sa valeur absolue. Cela découle aussi de la définition : Si  $x$  est réel, on a  $|(x, 0)| = \sqrt{x^2} = |x|$ .

On a donc pour tout nombre complexe  $z$ ,  $z = \Re(z) + i\Im(z)$ . Un nombre est dit *réel* si sa partie imaginaire est nulle. Si sa partie réelle est nulle, il est dit *imaginaire pur*. Tous ces nombres se représentent facilement dans le plan complexe :



**Exemple 2.4.2.** Calculer la partie réelle, la partie imaginaire, le module et le conjugué de

1.  $z_0 = i$ ,
2.  $z_1 = -2$ ,
3.  $z_2 = 1 + \sqrt{3}i$ ,
4.  $z_3 = 1 + i$ ,
5.  $z_4 = 3 + 2i$ .

On applique simplement les définitions. On peut bien sûr aussi représenter ces nombres pour “voir” les solutions, que voici :

- |                    |                    |                               |                       |                        |
|--------------------|--------------------|-------------------------------|-----------------------|------------------------|
| • $\Re(z_0) = 0$   | • $\Re(z_1) = -2$  | • $\Re(z_2) = 1$              | • $\Re(z_3) = 1$      | • $\Re(z_4) = 3$       |
| • $\Im(z_0) = 1$   | • $\Im(z_1) = 0$   | • $\Im(z_2) = \sqrt{3}$       | • $\Im(z_3) = 1$      | • $\Im(z_4) = 2$       |
| • $ z_0  = 1$      | • $ z_1  = 2$      | • $ z_2  = 2$                 | • $ z_3  = \sqrt{2}$  | • $ z_4  = \sqrt{13}$  |
| • $\bar{z}_0 = -i$ | • $\bar{z}_1 = -2$ | • $\bar{z}_2 = 1 - \sqrt{3}i$ | • $\bar{z}_3 = 1 - i$ | • $\bar{z}_4 = 3 - 2i$ |

Il est important de pouvoir travailler avec ces nombres associés, c'est-à-dire de connaître les propriétés de cette association. C'est l'objet des propriétés suivantes.

**Proposition 2.4.4** (Égalité). Si  $z_1, z_2$  sont des nombres complexes, on a

$$z_1 = z_2 \Leftrightarrow \begin{cases} \Re(z_1) = \Re(z_2) \\ \Im(z_1) = \Im(z_2) \end{cases}$$

*Démonstration.* Cela provient directement de l'unicité de la décomposition  $z = \Re(z) + i\Im(z)$ , valable pour tout nombre complexe.  $\square$

**Proposition 2.4.5** (Conjugué). *Pour tous nombres complexes  $z_1$  et  $z_2$ , on a*

1.  $\overline{z_1 + z_2} = \overline{z_1} + \overline{z_2}$ ,  $\overline{z_1 \cdot z_2} = \overline{z_1} \cdot \overline{z_2}$  et  $\overline{\overline{z_1}} = z_1$  ;
2. Si  $z_2 \neq 0$ , alors  $\overline{\left(\frac{z_1}{z_2}\right)} = \frac{\overline{z_1}}{\overline{z_2}}$  ;
3. Si  $z_1 \neq 0$ , on a  $\frac{1}{z_1} = \frac{\overline{z_1}}{|z_1|^2}$ .

*Remarque 2.6.* Les premières formules sont faciles à retenir : “conjugué d’une somme, d’un produit ou d’un quotient égale somme, produit, quotient des conjugués.” Il reste à savoir que conjuguer deux fois revient à ne rien faire, puis la formule la plus importante : le produit d’un nombre par son conjugué vaut son module au carré.

*Démonstration.* Pour la première assertion, c’est un simple calcul : on écrit  $z_1 = a_1 + ib_1$  et  $z_2 = a_2 + ib_2$ , pour  $a_1, a_2, b_1, b_2$  dans  $\mathbb{R}$  et on calcule.

Pour la deuxième assertion, on peut calculer de la même manière. Mais on peut aussi remarquer que  $\overline{\left(\frac{z_1}{z_2}\right)} = \frac{\overline{z_1}}{\overline{z_2}}$  est équivalente à  $\overline{z_2} \overline{\left(\frac{z_1}{z_2}\right)} = \overline{z_1}$ . Cette dernière égalité est vraie puisque le produit des conjugués vaut le conjugué du produit (par 1.). La troisième assertion est équivalente à  $z_1 \overline{z_1} = |z_1|^2$ . Avec les notations que nous avons choisies,  $z_1 \overline{z_1} = (a_1 + ib_1)(a_1 - ib_1) = a_1^2 + b_1^2 = |z_1|^2$ .  $\square$

Le nombre complexe conjugué permet de mettre sous forme algébrique le quotient de deux nombres complexes dont le dénominateur n’est pas nul, en utilisant le dernier point de la proposition précédente, ce qui revient à multiplier haut et bas par le conjugué du dénominateur :  $\frac{z_1}{z_2} = \frac{z_1 \overline{z_2}}{z_2 \overline{z_2}} = \frac{z_1 \overline{z_2}}{|z_2|^2}$ .

**Exemple 2.4.3.** Mettre sous forme algébrique les quotients

1.  $\frac{3+i}{4-i}$ ,                      2.  $\frac{i+5}{i-5}$ ,                      3.  $\frac{1}{i}$ ,                      4.  $\frac{5+3i}{4i+3}$ .

On applique la formule, ou on multiplie haut et bas par le conjugué (ou un multiple réel du conjugué) et on obtient

$$\frac{3+i}{4-i} = \frac{(3+i)(4+i)}{(4-i)(4+i)} = \frac{11+7i}{17}, \quad \frac{i+5}{i-5} = \frac{(i+5)(i+5)}{(i-5)(i+5)} = -\frac{24+10i}{26}.$$

De même, on trouve  $\frac{1}{i} = -i$  et  $\frac{5+3i}{4i+3} = \frac{1}{25}(27-11i)$ .

**Proposition 2.4.6** (Parties réelles et imaginaires). *Soit  $z$  un nombre. Alors,*

1. On a  $\Re(z) = \frac{z+\overline{z}}{2}$  et  $\Im(z) = \frac{z-\overline{z}}{2i}$  ;
2. Le nombre  $z$  est réel ssi  $\Im(z) = 0$  ssi  $z = \overline{z}$  ;
3. Le nombre  $z$  est imaginaire pur ssi  $\Re(z) = 0$  ssi  $z = -\overline{z}$ .

*Démonstration.* On considère  $z = a + ib$ , où  $a, b$  sont réels, et on calcule  $z + \overline{z} = 2a$ ,  $z - \overline{z} = 2ib$ . Pour la deuxième assertion, , par définition  $z$  est réel si  $b = 0$ , c’est-à-dire si  $\Im(z) = 0$ . On peut encore écrire cette condition  $\frac{z-\overline{z}}{2i} = 0$ , ou  $z = \overline{z}$ . On prouve la troisième assertion de la même façon.  $\square$

Passons maintenant aux propriétés du module.

**Proposition 2.4.7** (Module). *Pour tous nombres complexes  $z_1, z_2$*

1.  $|z_1 z_2| = |z_1| |z_2|$  et si  $z_2 \neq 0$ ,  $\left|\frac{z_1}{z_2}\right| = \frac{|z_1|}{|z_2|}$  et  $|\overline{z_1}| = |z_1|$  ;
2.  $|z_1|^2 = z_1 \overline{z_1}$ ,
3.  $|z_1 + z_2|^2 = |z_1|^2 + |z_2|^2 + 2\Re(z_1 \overline{z_2})$

On a donc de bonnes propriétés pour le module d'un produit, d'un quotient, et du conjugué. La deuxième propriété est déjà connue, mais c'est également une propriété du module. Enfin, puisque le module d'un nombre complexe correspond à la norme d'un élément de  $\mathbb{R}^2$ . On ne peut pas s'attendre à ce que le module d'une somme soit la somme des modules. Mais on a un théorème qui ressemble assez bien au théorème d'Al Kashi, ou aux propriétés du produit scalaire et de la norme. Ce n'est pas étonnant.

*Démonstration.* Pour la première propriété, il suffit de calculer, en prenant  $z_1 = a_1 + ib_1$  et  $z_2 = a_2 + ib_2$ . On peut démontrer l'égalité entre les carrés. On a

$$|z_1 z_2|^2 = (a_1 a_2 - b_1 b_2)^2 + (a_1 b_2 + a_2 b_1)^2 \quad \text{et} \quad |z_1|^2 |z_2|^2 = (a_1^2 + b_1^2)(a_2^2 + b_2^2).$$

On développe les deux expressions et on constate qu'elles sont égales.

On peut calculer aussi pour la propriété sur les quotients, mais on peut aussi constater qu'elle est équivalente à  $|z_2| \left| \frac{z_1}{z_2} \right| = |z_1|$ . Cette dernière propriété est vraie, puisque le produit des modules est égal au module du produit. Enfin,  $|a_1 - ib_1| = \sqrt{a_1^2 + (-b_1)^2} = |a_1 + ib_1|$ .

Nous avons déjà démontré la deuxième assertion.

Pour la troisième, on peut encore calculer les deux membres en utilisant les mêmes notations. Je vous laisse le faire. Mais on peut aussi calculer comme ceci :

$$|z_1 + z_2|^2 = (z_1 + z_2)(\overline{z_1 + z_2}) = (z_1 + z_2)(\overline{z_1} + \overline{z_2}) = z_1 \overline{z_1} + z_2 \overline{z_2} + (z_1 \overline{z_2} + z_2 \overline{z_1}).$$

Mais on a la relation  $z_2 \overline{z_1} = \overline{z_1 \overline{z_2}}$ , donc la dernière parenthèse vaut  $z_1 \overline{z_2} + \overline{z_1 \overline{z_2}}$ , ce qui donne le résultat annoncé vu la proposition 2.4.6.  $\square$

*Remarque 2.7.* Il est quand-même intéressant de calculer  $\Re(z_1 \overline{z_2})$ , avec les notations habituelles : on a alors

$$z_1 \overline{z_2} = (a_1 + ib_1)(a_2 - ib_2) = a_1 a_2 + b_1 b_2 + i(a_2 b_1 - a_1 b_2)$$

Donc  $\Re(z_1 \overline{z_2}) = a_1 a_2 + b_1 b_2$  n'est rien d'autre que le produit scalaire standard des éléments  $z_1$  et  $z_2$  de  $\mathbb{R}^2$ .

Enfin, nous terminons par quelques inégalités qui seront utiles par la suite. Il est également intéressant de réfléchir à l'interprétation géométrique de ces inégalités. Pour les premières, on a un triangle rectangle et pour la seconde, un triangle quelconque.

**Proposition 2.4.8** (Inégalités). *Soient des nombres complexes  $z, z_1, z_2$ .*

1. On a  $|\Re(z)| \leq |z|$ ,  $|\Im(z)| \leq |z|$  ;
2. On a l'inégalité triangulaire  $|z_1 + z_2| \leq |z_1| + |z_2|$ .
3. De plus  $|z_1 - z_2| \geq ||z_1| - |z_2||$

*Démonstration.* On écrit comme d'habitude  $z = a + ib$  où  $a, b$  sont réels. La première inégalité s'écrit alors  $|a| \leq \sqrt{a^2 + b^2}$ . Elle est équivalente, puisqu'il s'agit de nombres positifs, à  $a^2 \leq a^2 + b^2$ . L'inégalité concernant la partie imaginaire se démontre de la même façon.

Pour la deuxième assertion, il est équivalent de démontrer  $|z_1 + z_2|^2 \leq (|z_1| + |z_2|)^2$ . Mais le membre de gauche vaut  $|z_1|^2 + |z_2|^2 + 2\Re(z_1 \overline{z_2})$  et celui de droite  $|z_1|^2 + |z_2|^2 + 2|z_1||z_2|$ . L'inégalité à démontrer se réduit donc à  $\Re(z_1 \overline{z_2}) \leq |z_1||z_2| = |z_1||\overline{z_2}|$ , ou encore à  $\Re(z_1 \overline{z_2}) \leq |z_1 \overline{z_2}|$ . Cette dernière inégalité résulte du point 1.

Pour la troisième assertion, il y a un truc :  $z_1 = (z_1 - z_2) + z_2$ , donc  $|z_1| \leq |z_1 - z_2| + |z_2|$ , par le point 2. On a donc  $|z_1 - z_2| \geq |z_1| - |z_2|$ . Par symétrie, on obtient aussi  $|z_1 - z_2| \geq |z_2| - |z_1|$ .  $\square$

*Remarque 2.8.* La deuxième inégalité porte aussi le nom d'inégalité de Minkowski.

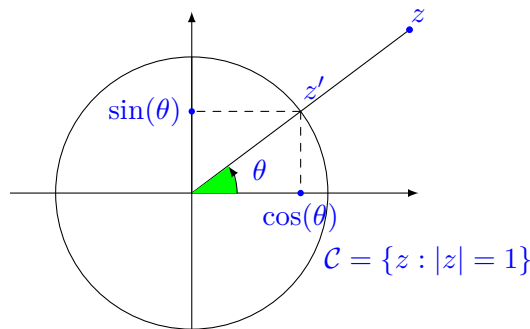
## 2.5 Forme trigonométrique, formule de Moivre

Jusqu'à présent, nous avons défini les nombres complexes comme des couples de nombres réels, et obtenu la représentation algébrique. Nous les avons représentés dans le plan muni d'un repère orthonormé, au moyen de *coordonnées cartésiennes*. Il existe cependant d'autres systèmes de coordonnées du plan. Les plus connues sont les coordonnées polaires. Elles correspondent à une autre écriture des nombres complexes, appelée forme trigonométrique.

Tout d'abord, rappelons que l'ensemble des nombres complexes de module égal à 1 est exactement le cercle trigonométrique. Tout nombre complexe  $z$  tel que  $|z| = 1$  s'écrit donc de manière unique

$$z = \cos(\theta) + i \sin(\theta),$$

pour  $\theta \in [0, 2\pi[$ . On peut de la même manière repérer tout nombre complexe  $z$  non nul, au moyen de l'angle entre le premier vecteur associé au repère et  $z$ , et du module de  $z$ . Voici l'idée :



La droite déterminée par 0 et  $z$  coupe le cercle en  $z'$ , qui a pour coordonnées  $(\cos(\theta), \sin(\theta))$ . Donc  $z' = \cos(\theta) + i \sin(\theta)$ . Pour obtenir  $z$ , on doit multiplier  $z'$  par un nombre positif  $\rho$ . Si  $z = \rho z'$ , alors  $|z| = \rho |z'| = \rho$ . L'idée géométrique est claire, mais comment décrire cette transformation sous forme algébrique ? C'est l'objet de la proposition suivante.

**Proposition 2.5.1.** *Si  $z = a + ib$  est un nombre complexe non nul, alors il existe un unique  $\theta \in [0, 2\pi[$  et un unique  $\rho \in ]0, +\infty[$  tels que  $z = \rho(\cos(\theta) + i \sin(\theta))$ . De plus  $\theta$  et  $\rho$  sont donnés par les équations suivantes :*

$$\begin{cases} \rho &= |z| = \sqrt{a^2 + b^2} \\ \cos(\theta) &= \frac{a}{\sqrt{a^2 + b^2}} \\ \sin(\theta) &= \frac{b}{\sqrt{a^2 + b^2}} \end{cases}$$

*Remarque 2.9.* Cette écriture du nombre complexe  $z$  est appelée forme trigonométrique de  $z$ .

*Démonstration.* Il suffit d'écrire les conditions demandées : on a  $a + ib = \rho(\cos(\theta) + i \sin(\theta))$  si, et seulement si, les parties réelles et imaginaires de ces deux nombres sont égales. On peut aussi ajouter la condition d'égalité des modules de ces nombres, qui est une conséquence des deux autres. On a donc les conditions équivalentes

$$\begin{cases} a &= \rho \cos(\theta) \\ b &= \rho \sin(\theta) \\ \sqrt{a^2 + b^2} &= \rho |\cos(\theta) + i \sin(\theta)| = \rho. \end{cases}$$

On obtient directement l'existence d'un unique  $\theta$ , puisque  $(\frac{a}{\sqrt{a^2 + b^2}}, \frac{b}{\sqrt{a^2 + b^2}})$  a un module égal à 1.  $\square$

- Définition 2.5.1.** 1. L'angle  $\theta$  de la proposition précédente est appelé *l'argument principal* de  $z$ . Par extension, on appellera aussi argument de  $z$  tout angle  $\theta'$  tel que  $z = \rho(\cos(\theta') + i \sin(\theta'))$ , c'est-à-dire tel que  $\theta' - \theta = 2k\pi$ ,  $k \in \mathbb{Z}$ .<sup>a</sup>
2. Un nombre complexe  $z$  écrit sous la forme  $z = \rho(\cos \theta + i \sin \theta)$  est mis sous *forme trigonométrique*.
3. On note  $e^{i\theta}$  le nombre  $\cos(\theta) + i \sin(\theta)$ . La forme trigonométrique est alors  $z = \rho e^{i\theta}$ .<sup>b</sup>

**Exemple 2.5.1.** Trouver les formes trigonométriques de

1.  $z_0 = 1 + i$ ,
2.  $z_1 = -1$ ,
3.  $z_2 = 2 + 3i$ ,
4.  $z_3 = -2 - 3i$ ,
5.  $z_4 = 1 + \sqrt{3}i$ .

Détaillons la solution pour  $z_0 = 1 + i$  (on peut le représenter pour voir l'angle). On a  $\rho = |z_0| = \sqrt{2}$ . On cherche alors  $\theta$  tel que

$$\begin{cases} \rho \cos(\theta) &= 1 \\ \rho \sin(\theta) &= 1, \end{cases}$$

ou encore

$$\begin{cases} \cos(\theta) &= \frac{\sqrt{2}}{2} \\ \sin(\theta) &= \frac{\sqrt{2}}{2} \end{cases}$$

Dès lors, on a  $z_0 = \sqrt{2}e^{i\frac{\pi}{4}}$ .

On procède de la même façon pour obtenir  $z_1 = e^{i\pi}$ ,  $z_2 = \sqrt{13}e^{i \arccos(\frac{2}{\sqrt{13}})}$ ,  $z_3 = \sqrt{13}e^{i(2\pi - \arccos(\frac{2}{\sqrt{13}}))}$  et  $z_4 = 2e^{i\frac{\pi}{3}}$ .

Ces formes alternatives permettent de **calculer plus facilement les produits et inverses** des nombres complexes, comme le montre la proposition suivante.

**Proposition 2.5.2.** Si  $z_1 = \rho_1 e^{i\theta_1}$  et  $z_2 = \rho_2 e^{i\theta_2}$  alors on a

$$z_1 \cdot z_2 = \rho_1 \rho_2 e^{i(\theta_1 + \theta_2)}$$

Si  $z_2 \neq 0$ , alors on a

$$\frac{z_1}{z_2} = \frac{\rho_1}{\rho_2} e^{i(\theta_1 - \theta_2)}.$$

Enfin, on a

$$z_1^n = \rho_1^n e^{in\theta_1} \quad \text{et} \quad \overline{z_1} = \rho_1 e^{-i\theta_1}.$$

*Démonstration.* On prouve la première formule en utilisant les définitions de  $e^{i\theta_1}$  et  $e^{i\theta_2}$ , et les formules d'addition pour le sinus et le cosinus, en effet, on a

$$z_1 \cdot z_2 = \rho_1 \rho_2 e^{i\theta_1} e^{i\theta_2}.$$

De plus

$$\begin{aligned} e^{i\theta_1} e^{i\theta_2} &= (\cos(\theta_1) + i \sin(\theta_1))(\cos(\theta_2) + i \sin(\theta_2)) \\ &= (\cos(\theta_1) \cos(\theta_2) - \sin(\theta_1) \sin(\theta_2)) + i(\cos(\theta_1) \sin(\theta_2) + \sin(\theta_1) \cos(\theta_2)) \\ &= \cos(\theta_1 + \theta_2) + i \sin(\theta_1 + \theta_2) \\ &= e^{i(\theta_1 + \theta_2)}. \end{aligned}$$

a. Le choix de considérer l'angle  $\theta$  dans  $[0, 2\pi[$  est arbitraire, on aurait pu le choisir dans  $] - \pi, \pi]$  par exemple. La correspondance se fait en ajoutant des multiples de  $2\pi$  si nécessaire.

b. Cette notation sera justifiée par les propriétés de ces nombres vis-à-vis de la multiplication, des puissances et du passage à l'inverse. Notez que dans certains exposés, on prend le problème dans l'autre sens : on définit l'exponentielle de tout nombre complexe, puis seulement le sinus et le cosinus.

Toutes les autres formules s'en déduisent : la deuxième assertion est équivalente à

$$z_1 = \frac{\rho_1}{\rho_2} e^{i(\theta_1 - \theta_2)} \rho_2 e^{i\theta_2},$$

qui est vraie par le point 1. L'expression de  $z_1^n$  est obtenue par récurrence directe, et  $\bar{z}_1 = \rho_1(\cos(\theta_1) - i \sin(\theta_1))$ .  $\square$

Le résultat est donc obtenu à partir des formules de trigonométrie. Il peut vous permettre de les retrouver, simplement en vous souvenant que  $e^{i\theta}$  a les propriétés bien connues de l'exponentielle.<sup>c</sup> Comme cas particulier des formules précédentes, on obtient le résultat suivant.

**Corollaire 2.5.1** (Formule de Moivre (Abraham de Moivre (1667-1754))). *On a*

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

*pour tout  $\theta \in \mathbb{R}$  et tout  $n$  entier.*

Cette formule permet, en développant le membre de gauche, d'exprimer  $\cos(n\theta)$  et  $\sin(n\theta)$  au moyen d'expressions polynomiales en  $\sin(\theta)$  et  $\cos(\theta)$ . Par exemple, on a

$$\begin{aligned} \cos(3\theta) &= \Re(\cos \theta + i \sin \theta)^3 = \Re(\cos^3(\theta) + 3i \cos^2(\theta) \sin(\theta) - 3 \cos(\theta) \sin^2(\theta) - i \sin^3(\theta)) \\ &= \cos^3(\theta) - 3 \cos(\theta) \sin^2(\theta). \end{aligned}$$

De même,  $\sin(3\theta) = 3 \cos^2(\theta) \sin(\theta) - \sin^3(\theta)$ .

Si on veut obtenir des formules qui permettent de transformer des produits, ou des puissances de sinus et cosinus, en sinus ou cosinus d'angles multiples, on peut utiliser la proposition suivante.

**Proposition 2.5.3.** *On a les formules*

$$\cos(x) = \Re(e^{ix}) = \frac{e^{ix} + e^{-ix}}{2}, \quad \text{et} \quad \sin(x) = \Im(e^{ix}) = \frac{e^{ix} - e^{-ix}}{2i}$$

*pour tout  $x \in \mathbb{R}$ .*

A l'aide de ces formules, on peut :

1. Récupérer les formules de Carnot (développer  $\cos^2(x) = (\frac{e^{ix} + e^{-ix}}{2})^2$ )
2. Exprimer  $\cos^3(x)$  en fonction de  $\cos(3x)$  et  $\cos(x)$  (développer  $\cos^3(x) = (\frac{e^{ix} + e^{-ix}}{2})^3$ ).
3. Faire de même avec  $\sin^3(x)$  (développer  $\sin^3(x) = (\frac{e^{ix} - e^{-ix}}{2i})^3$ ).
4. Se souvenir des formules d'addition :  $\cos(x+y) = \Re(e^{i(x+y)})$ , de même avec le sinus.

## 2.6 Un mot sur la forme exponentielle

Nous avons jusqu'à présent défini l'exponentielle d'un nombre complexe imaginaire pur :  $e^{ib} = \cos(b) + i \sin(b)$ . Nous connaissons aussi l'exponentielle des nombres réels (c'est l'exponentielle classique du nombre réel  $a$ , notée  $e^a$ ). Si on veut définir l'exponentielle sur les nombres complexes, et qu'elle garde la propriété de transformer les sommes en produits, on est amené naturellement à la définition suivante.

**Définition 2.6.1.** Si  $z = a + ib$ , où  $a, b \in \mathbb{R}$ , on définit

$$e^z = e^a(\cos(b) + i \sin(b)).$$

<sup>c</sup>. Le comportement en question est  $e^a e^b = e^{a+b}$ , quels que soient  $a$  et  $b$ . Dans la plupart des textes mathématiques, les exponentielles complexes sont définies avant les sinus et cosinus, et les propriétés de ces derniers se déduisent de cette propriété fondamentale des exponentielles.

Cette définition permet d'obtenir un résultat similaire à celui sur les formes trigonométriques.

**Proposition 2.6.1.** *Tout nombre complexe non nul  $z$  s'écrit  $e^w$  pour un  $w \in \mathbb{C}$ . Ce nombre complexe  $w$  n'est pas unique.*

*Démonstration.* On sait que  $z$  admet une forme trigonométrique  $z = \rho e^{i\theta} = \rho(\cos(\theta) + i \sin(\theta))$ . Puisque  $\rho$  est strictement positif, on a  $\rho = e^{\ln(\rho)}$ , donc  $z = e^{\ln(\rho) + i\theta}$ . Cette forme n'est pas unique puisqu'on peut ajouter à  $\theta$  tout multiple de  $2\pi$  sans changer la valeur de l'exponentielle.  $\square$

La notation exponentielle est justifiée par les propriétés suivantes, que l'on démontrera sans difficulté.

**Proposition 2.6.2.** *Pour tous nombres complexes  $z, z_1, z_2$ , on a*

1.  $e^{z_1} e^{z_2} = e^{z_1 + z_2}$  ;
2.  $\frac{1}{e^z} = e^{-z}$  ;
3.  $(e^z)^n = e^{nz}, \forall z \in \mathbb{C}$ .
4. Pour tout  $x \in \mathbb{R}$ , on a  $|e^{ix}| = 1$ .

## 2.7 Interprétation géométrique de l'addition et de la multiplication

Nous avons défini les opérations d'addition et de multiplication dans  $\mathbb{C}$ . Nous avons constaté que l'addition était celle des composantes de vecteurs. Elle a donc une interprétation géométrique simple. L'addition d'un élément fixe de  $\mathbb{C}$  est une transformation du plan bien connue : c'est une translation. En passant à la forme trigonométrique, nous pouvons aussi interpréter la multiplication par un élément fixe de  $\mathbb{C}$  comme une transformation de  $\mathbb{C}$ . Voici les définitions et résultats formels

**Définition 2.7.1.** Pour tout  $z_0 \in \mathbb{C}$  définissons les applications

1.  $s_{z_0} : \mathbb{C} \rightarrow \mathbb{C} : z \mapsto z + z_0$
2.  $m_{z_0} : \mathbb{C} \rightarrow \mathbb{C} : z \mapsto zz_0$

Ainsi définie,  $s_{z_0}$  est la "somme avec  $z_0$ ", tandis que  $m_{z_0}$  est la "multiplication par  $z_0$ ". Ces deux applications peuvent être inversées (si  $z_0 \neq 0$  dans le second cas), comme l'indique la proposition suivante.

**Proposition 2.7.1.** *Pour tout  $z_0 \in \mathbb{C}$ ,  $s_{z_0}$  est une bijection de  $\mathbb{C}$  dans  $\mathbb{C}$ . C'est aussi le cas pour  $m_{z_0}$  si  $z_0 \neq 0$ .*

*Démonstration.* On constate que l'inverse de  $s_{z_0}$  est  $s_{-z_0}$  et que l'inverse de  $m_{z_0}$  est  $m_{z_0}^{-1}$ .  $\square$

*Remarque 2.10.* On n'a pas du tout utilisé les structures spécifiques de  $\mathbb{C}$  dans cette preuve. On a juste utilisé la structure de groupe de  $(\mathbb{C}, +, 0)$  et de  $(\mathbb{C}_0, \cdot, 1)$ . Ce résultat se généralise donc à n'importe quel groupe.

Puisque  $\mathbb{C}$  a été identifié au plan, les applications  $s_{z_0}$  et  $m_{z_0}$  peuvent être vues comme des transformations du plan. Il reste à identifier ces transformations.

**Proposition 2.7.2.** *Si  $z_0 = a_0 + ib_0 = \rho_0 e^{i\theta_0}$ , alors la bijection  $s_{z_0}$  est une translation de vecteur  $z_0$  (ou  $\overrightarrow{Oz_0}$ ), et  $m_{z_0}$  est une similitude, c'est la composée d'une rotation centrée à l'origine du repère et d'angle  $\theta_0$  et d'une homothétie de rapport  $\rho_0$ , également centrée à l'origine.*

*Démonstration.* On calcule  $s_{z_0}(z)$  en considérant  $z = a + ib$ . On a  $s_{z_0}(z) = a + a_0 + i(b + b_0)$ . En termes géométriques,  $s_{z_0}$  transforme le point de coordonnées  $(a, b)$  en le point de coordonnées  $(a + a_0, b + b_0)$ . On a donc bien une translation.

On calcule  $m_{z_0}(z)$  en considérant  $z = \rho e^{i\theta}$ . On a  $m_{z_0}(z) = \rho_0 e^{i(\theta + \theta_0)}$ . On a donc ajouté  $\theta_0$  à l'argument de  $z$ , ce qui correspond bien à une rotation, puis multiplié par  $\rho_0 \in \mathbb{R}$ , ce qui correspond à une homothétie.  $\square$

Cette interprétation permet de “voir” certains résultats algébriques. Par exemple, le nombre  $e^{i\theta}$  correspond à une rotation d'angle  $\theta$ . Alors son carré correspond à appliquer deux fois cette rotation, donc à une rotation d'angle  $2\theta$ , sa puissance  $n$ -ème à une rotation d'angle  $n\theta$ . On voit donc la formule de Moivre. De même, multiplier par  $-1$  consiste à appliquer une rotation d'angle  $\pi$ . On n'est donc pas surpris d'avoir  $-1 = e^{i\pi}$ ...

## 2.8 Racines carrées et équations du second degré

Il est maintenant grand temps de revenir à nos moutons. Nous avons quand même introduit les nombres complexes pour obtenir des racines carrées.<sup>a</sup>. Nous allons montrer dans cette section que toute équation du second degré à coefficients complexes admet des solutions dans l'ensemble des nombres complexes. Passons toute de suite à la définition des racines carrées.

**Définition 2.8.1.** Si  $z$  est un nombre complexe, alors on appelle une racine carrée de  $z$  tout nombre complexe  $z'$  tel que  $z'^2 = z$ .

Nous avons construit les nombres complexes pour donner des racines aux nombres négatifs. On a en fait obtenu bien plus, comme l'indique le résultat suivant.

**Proposition 2.8.1.** *Tout nombre complexe  $z$  non nul admet exactement deux racines carrées opposées. Le nombre complexe  $z = 0$  admet une seule racine 0. Cette racine est dite double.*

*Démonstration.* Si  $z \neq 0$ , nous avons constaté que la forme trigonométrique était plus adaptée à la multiplication. On considère donc  $z = \rho e^{i\theta}$  et on cherche  $z' = \rho' e^{i\theta'}$  tel que  $z'^2 = z$ . Cette condition s'écrit encore  $\rho'^2 e^{i2\theta'} = \rho e^{i\theta}$ . On obtient donc les conditions équivalentes

$$\begin{cases} \rho'^2 &= \rho \\ \cos(2\theta') &= \cos(2\theta) \\ \sin(2\theta') &= \sin(2\theta) \end{cases}$$

Cela donne directement  $\rho' = \sqrt{\rho}$  et  $2\theta' = \theta + 2k\pi$ ,  $k \in \mathbb{Z}$ . Si on remarque que  $e^{i2\pi} = 1$ , on peut se limiter à  $k \in \{0, 1\}$  et on trouve les solutions  $\sqrt{\rho} e^{i\frac{\theta}{2}}$  et  $\sqrt{\rho} e^{i\frac{\theta}{2} + i\pi} = -\sqrt{\rho} e^{i\frac{\theta}{2}}$ . Le cas de  $z = 0$  est évident.  $\square$

*Remarque 2.11.* 1. Dans le plan complexe, il devient délicat de privilégier naturellement une des deux racines, pour définir une expression du type  $\sqrt{z}$ . On peut le faire en prenant les précautions, mais on ne peut plus garantir que cette racine ait les mêmes propriétés que dans  $\mathbb{R}$ , comme par exemple  $\sqrt{a}\sqrt{b} = \sqrt{ab}$  quand ces expressions sont définies. On n'utilisera donc pas le symbole  $\sqrt{z}$ , quand  $z$  est complexe. On ne l'utilisera que quand le radicand est réel et positif. C'est déjà le cas dans les lignes qui suivent.

2. Il est important d'insister sur le fait qu'il s'agit d'une preuve constructive. Les racines de  $z = \rho e^{i\theta}$  sont

$$\pm \sqrt{\rho} e^{i\frac{\theta}{2}}.$$

a. En fait, les nombres complexes n'ont pas été introduits à cet effet. Ils sont apparus comme un artifice de calcul pour résoudre certaines équations du troisième de degré.

3. Dans la preuve, il n'était pas nécessaire de passer aux équations trigonométriques : on a  $e^{i\theta} = e^{i\theta'}$  si et seulement si  $\theta - \theta' = 2k\pi$ , pour un  $k \in \mathbb{Z}$ .

**Exemple 2.8.1.** Déterminer les racines carrées de  $1 + i$  et de  $3 + 4i$

Solution : on calcule la forme trigonométrique de  $1 + i$ , et on obtient  $\sqrt{2}e^{i\frac{\pi}{4}}$ . Les racines carrées sont donc  $\pm\sqrt[4]{2}e^{i\frac{\pi}{8}}$ . Si on le souhaite, on peut développer l'exponentielle pour revenir à une forme algébrique.

On fait de même pour  $3 + 4i$ . Il s'écrit  $5e^{i\arccos(\frac{3}{5})}$ . Donc les racines sont  $\pm\sqrt{5}e^{i\arccos(\frac{3}{5})/2}$ . Il reste alors à calculer la forme algébrique en développant  $\cos(\frac{1}{2}\arccos(\frac{3}{5}))$  et  $\sin(\frac{1}{2}\arccos(\frac{3}{5}))$ . Les formules de Carnot donnent une solution : on a  $1 + \cos(2a) = 2\cos^2(a)$  et  $1 - \cos(2a) = 2\sin^2(a)$ , donc  $\cos^2(\frac{a}{2}) = \frac{1}{2}(1 + \cos(a))$  et  $\sin^2(\frac{a}{2}) = \frac{1}{2}(1 - \cos(a))$ . Donc finalement

$$\cos(\arccos(\frac{3}{5})/2) = \sqrt{\frac{1}{2}(1 + \frac{3}{5})} = \frac{2}{\sqrt{5}} \quad \text{et} \quad \sin(\arccos(\frac{3}{5})/2) = \frac{1}{\sqrt{5}}.$$

Donc les racines de  $3 + 4i$  sont  $\pm(2 + i)$ .

L'exemple précédent montre qu'il serait utile de pouvoir calculer directement les racines carrées à partir de la forme algébrique.

**Proposition 2.8.2.** Les racines carrées du nombre complexe  $z = a + ib$  sont les nombres  $x + iy$  où  $x, y$  sont solutions du système d'équations

$$\begin{cases} x^2 - y^2 &= a \\ 2xy &= b \\ x^2 + y^2 &= \sqrt{a^2 + b^2}. \end{cases}$$

*Démonstration.* Un nombre complexe  $z' = x + iy$  est une racine de  $z$  ssi on a  $z'^2 = z$ . On exprime que ces deux nombres sont égaux, c'est-à-dire qu'ils ont même partie réelle et même partie imaginaire. Puisque  $z'^2 = (x^2 - y^2) + 2ixy$ , ces conditions sont équivalentes à

$$\begin{cases} x^2 - y^2 &= a \\ 2xy &= b. \end{cases}$$

On peut ajouter une équation : les nombres  $z'^2$  et  $z$  sont égaux ssi, ils ont des même partie réelle, même partie imaginaire et même module. La condition sur le module est une conséquence des deux autres, mais elle sera bien utile pour résoudre le système d'équations. Puisque  $|z'^2| = |z'|^2$ , on a l'équation supplémentaire  $x^2 + y^2 = \sqrt{a^2 + b^2}$  et on obtient le résultat annoncé.  $\square$

Ce théorème est facile à retenir : on cherche  $z' = x + iy$  tel que  $z'^2 = z$ . On exprime que ces nombres ont même partie réelle, même partie imaginaire et même module.

**Exemple 2.8.2.** Revenons à  $z = 3 + 4i$ . Les racines carrées sont les nombres  $x + iy$  satisfaisant

$$\begin{cases} x^2 - y^2 &= 3 \\ 2xy &= 4 \\ x^2 + y^2 &= 5. \end{cases}$$

La première et la dernière équation nous permettent, en additionnant et en soustrayant membre à membre, d'obtenir le système équivalent

$$\begin{cases} x^2 &= 4 \\ y^2 &= 1 \\ 2xy &= 4. \end{cases}$$

Les solutions des deux premières équations sont  $x = \pm 2$  et  $y = \pm 1$ . Cela fait quatre possibilités. Mais la dernière équation indique que  $x$  et  $y$  ont même signe. On trouve donc  $x = 2$  et  $y = 1$  ou  $x = -2$  et  $y = -1$ . C'est quand-même plus facile.

Il est très simple de retenir les trois équations ci-dessus. On peut systématiser la résolution du système d'équations et obtenir une formule pour les racines. C'est l'objet du résultat suivant, que je donne à titre indicatif.

**Proposition 2.8.3.** *Les racines du nombre complexe  $z = a + ib$  sont données par*

1. Si  $b \neq 0$ ,  $\pm(\sqrt{\frac{\sqrt{a^2+b^2}+a}{2}} + \operatorname{sgn}(b)\sqrt{\frac{\sqrt{a^2+b^2}-a}{2}}i)$ , où  $\operatorname{sgn}(b)$  est le signe de  $b$ .
2. Si  $b = 0$  et  $a \geq 0$ ,  $\pm\sqrt{a}$ ;
3. Si  $b = 0$  et  $a \leq 0$ ,  $\pm i\sqrt{-a}$ .

*Démonstration.* On reprend le système d'équations obtenu à la proposition précédente et on additionne et soustrait, membre à membre, la première et la troisième équation. On obtient alors le système équivalent

$$\begin{cases} 2x^2 &= \sqrt{a^2+b^2}+a \\ 2y^2 &= \sqrt{a^2+b^2}-a \\ 2xy &= b. \end{cases}$$

Les membres de droite des deux premières équations sont positifs car  $|a| \leq \sqrt{a^2+b^2}$  pour tous  $a, b$  réels. Cela permet de déterminer  $x$  et  $y$ , au signe près (on a en général deux solutions pour  $x$  et deux pour  $y$ ). Si  $b \neq 0$ , la dernière équation permet d'éliminer deux possibilités sur les quatre et d'obtenir le résultat annoncé. Si  $b = 0$ , les équations se simplifient et donnent directement le résultat annoncé.  $\square$

Voici encore un exemple.

**Exemple 2.8.3.** Si on cherche les racines complexes de  $5 + 12i$ , on les cherche sous la forme  $x + iy$ . On exprime l'équation  $(x + iy)^2 = 5 + 12i$ , que l'on développe pour obtenir

$$\begin{cases} x^2 - y^2 &= 5 \\ 2xy &= 12. \end{cases}$$

On n'oublie pas la troisième équation, qui est  $x^2 + y^2 = |5 + 12i| = \sqrt{169} = 13$ . On résout le système

$$\begin{cases} x^2 - y^2 &= 5 \\ 2xy &= 12 \\ x^2 + y^2 &= 13. \end{cases}$$

En additionnant et en soustrayant les première et troisième équation, on obtient  $x^2 = 9$  et  $y^2 = 4$ . On a donc  $x = \pm 3$  et  $y = \pm 2$ . Mais puisque le produit  $xy$  doit être positif, soit on choisit les deux solutions positives, soit les deux négatives. Les racines carrées sont donc  $-3 - 2i$  et  $3 + 2i$ .

Maintenant que nous connaissons l'existence des racines carrées et que nous pouvons les calculer, nous pouvons également résoudre les équations du second degré à coefficients complexes.

**Définition 2.8.2.** Une équation du type

$$az^2 + bz + c = 0, \quad a \in \mathbb{C} \setminus \{0\}, b, c \in \mathbb{C} \quad (2.7)$$

est appelée équation du second degré à coefficients complexes.

La fonction  $P : \mathbb{C} \rightarrow \mathbb{C} : z \mapsto az^2 + bz + c$  est appelée *fonction trinôme du second degré*.

A ce trinôme on associe le nombre complexe  $\Delta = b^2 - 4ac$ .

**Proposition 2.8.4.** *Toute équation du second degré à coefficients complexes admet deux solutions complexes distinctes (si  $\Delta \neq 0$ ) ou une seule solution complexe, dite double (si  $\Delta = 0$ ). Ces solutions sont données par*

$$\frac{-b \pm \delta}{2a} \quad \text{où} \quad \delta^2 = \Delta.$$

*Remarque 2.12.* On peut simplifier le résultat en indiquant qu'il y a deux solutions, comptées avec leurs multiplicités. Le concept de multiplicité sera approfondi dans la suite du cours d'algèbre. On notera aussi que la formule est identique à celle que l'on connaît dans les nombres réels. On ne note pas  $\sqrt{\Delta}$ , et cela explique la notation  $\delta$ .

*Démonstration.* On procède comme dans le cas des équation à coefficients réels. On a encore

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

Puisque  $\Delta$  est toujours un carré, on peut écrire  $\Delta = \delta^2$  et continuer la factorisation.

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

L'équation admet donc les solutions de l'énoncé.  $\square$

Les résultats sur la factorisation du trinôme, et la somme et le produit des solutions se généralisent.

**Proposition 2.8.5.** *L'équation du second degré  $az^2 + bz + c = 0$  admet des solutions  $z_1$  et  $z_2$  (éventuellement égales). Le trinôme correspondant se factorise :*

$$az^2 + bz + c = a(z - z_1)(z - z_2) \quad \forall z \in \mathbb{C}.$$

De plus, on a  $z_1 + z_2 = -\frac{b}{a}$  et  $z_1 z_2 = \frac{c}{a}$ .

*Démonstration.* Nous avons déjà obtenu la factorisation. La preuve pour la somme et le produit est identique à celle du cas réel.  $\square$

## 2.9 Généralisations : puissances $n$ -èmes et racines $n$ -èmes

Pour calculer la puissance  $n$ -ème d'un nombre complexe, on peut utiliser une des formes que l'on a à disposition : algébrique ou trigonométrique. La forme trigonométrique est évidemment la plus simple à utiliser. En effet, calculer une puissance  $n$ -ème d'une forme algébrique  $a + ib$  donne lieu à l'expression  $(a + ib)^n$ . C'est quelque peu pénible, mais cela peut être mené à bien grâce à la formule du binôme de Newton, que je rappelle ici. Tout d'abord, voici la définition des coefficients binomiaux.

**Définition 2.9.1.** Pour tous  $k, n \in \mathbb{N}$  tels que  $0 \leq k \leq n$ , on définit le coefficient binomial

$$C_n^k = \frac{n!}{k!(n-k)!}.$$

*Remarque 2.13.* 1. Par convention, on pose  $0! = 1$ , donc  $C_n^0 = C_n^n = 1$ .

2. Le nombre  $C_n^k$  est le nombre de façons de choisir  $k$  objets parmi  $n$  objets distincts, sans remise, et l'ordre n'ayant pas d'importance.

Ces nombres satisfont la règle de construction connue sous le nom de triangle de Pascal.

**Proposition 2.9.1.** *Pour tout  $n \geq 1$  et tout  $k$  tel que  $0 \leq k \leq n-1$ ,*

$$C_n^k + C_n^{k+1} = C_{n+1}^{k+1}$$

*Démonstration.* Il suffit de calculer le membre de gauche, en mettant en évidence tout ce que l'on peut.  $\square$

*Remarque 2.14.* Il y a une preuve moins calculatoire : considérons  $n+1$  objets distincts et marquons le premier, disons  $o_1$ . Pour choisir les  $k+1$  objets, soit on sélectionne  $o_1$ , soit on ne le sélectionne pas. Dans le premier cas, il reste à choisir  $k$  objets parmi les  $n$  objets restants ( $C_n^k$  possibilités). Dans le deuxième cas, on doit choisir  $k+1$  objets parmi les  $n$  objets restants ( $C_n^{k+1}$  possibilités).

Voici une version du triangle de Pascal : sur la ligne  $n$ , on porte  $C_n^k$  pour  $k \in \{0, \dots, n\}$ . Les valeurs extrêmes sont égales à 1. Les autres sont obtenues en utilisant la formule de la proposition précédente : on additionne les valeurs adéquates de la ligne précédente :

$$\begin{array}{ccccccc}
 n = 0 : & & & & & & 1 \\
 n = 1 : & & & & 1 & & 1 \\
 n = 2 : & & & 1 & & 2 & & 1 \\
 n = 3 : & & 1 & & 3 & & 3 & & 1 \\
 n = 4 : & 1 & & 4 & & 6 & & 4 & & 1 \\
 n = 4 : & 1 & & 5 & & 10 & & 10 & & 5 & & 1
 \end{array}$$

**Proposition 2.9.2** (Formule du binôme de Newton). *Pour tous nombres complexes  $w$  et  $z$ , et tout  $n \in \mathbb{N}$ , on a*

$$(w + z)^n = \sum_{k=0}^n C_n^k w^k z^{n-k}.$$

*Démonstration.* La preuve la plus systématique se fait par récurrence. Pour  $n = 0, 1, 2$  le résultat est facile à obtenir. Supposons que le résultat soit vrai pour  $n \in \mathbb{N}$  et montrons qu'il est vrai pour  $n + 1$ . On a

$$(w + z)^{n+1} = (w + z)(w + z)^n = (w + z) \sum_{k=0}^n C_n^k w^k z^{n-k} = \sum_{k=0}^n C_n^k w^{k+1} z^{n-k} + \sum_{k=0}^n C_n^k w^k z^{n-k+1}.$$

Si on garde en tête la formule à prouver pour  $n + 1$ , la deuxième somme contient les bons monômes  $w^k z^{n-k+1}$ . On fait un changement d'indice ( $k' = k + 1$ , donc  $k = k' - 1$ ) dans la première somme, pour obtenir

$$\sum_{k=0}^n C_n^k w^{k+1} z^{n-k} = \sum_{k'=1}^{n+1} C_n^{k'-1} w^{k'} z^{n-k'+1} = \sum_{k=1}^{n+1} C_n^{k-1} w^k z^{n-k+1}.$$

On a donc

$$\begin{aligned}
 (w + z)^{n+1} &= \sum_{k=1}^{n+1} C_n^{k-1} w^k z^{n-k+1} + \sum_{k=0}^n C_n^k w^k z^{n-k+1} \\
 &= C_n^n w^{n+1} z^0 + C_0^0 w^0 z^{n+1} + \sum_{k=1}^n (C_n^{k-1} + C_n^k) w^k z^{n-k+1}.
 \end{aligned}$$

On conclut en utilisant le triangle de Pascal et en notant que  $C_n^n = 1 = C_{n+1}^{n+1}$ .  $\square$

*Remarque 2.15.* Un seul cas de base suffit, évidemment. Mais regarder ce qui se passe pour  $n = 2$  ou  $n = 3$  permet au lecteur débutant d'anticiper les manipulations des sommes qui apparaissent dans la preuve.

Il existe d'autres démonstrations plus rapides, mais moins systématiques. Par exemple, on peut se dire que  $(w + z)^n$  est le produit de  $n$  facteurs suivants :  $(w + z) \dots (w + z)$ . Pour développer ce produit, on distribue chaque terme. A chaque distribution, on est donc amené à choisir entre  $z$  et  $w$ . Dans le développement, on a donc des monômes de la forme  $w^k z^l$ , mais avec la condition  $k + l = n$ , ou  $l = n - k$ , car il y a  $n$  facteurs. Reste à déterminer le coefficient de ce monôme. Il correspond à toutes les façons distinctes de choisir les  $k$  facteurs où on développe sur  $w$  et non  $z$ . C'est évidemment  $C_n^k$ .

Il existe bien sûr d'innombrables généralisations de cette formule. On peut aussi noter que la preuve n'a pas fait appel aux définitions spécifiques des opérations dans  $\mathbb{C}$ . On a distribué, permuté des sommes, utilisé la commutativité... La formule est donc certainement vraie dans tout anneau commutatif.

Mais bon, revenons à nos moutons. On peut définir les racines  $n$ -èmes d'un nombre complexe :

**Définition 2.9.2.** Une racine  $n$ -ème d'un nombre complexe  $z$  est un nombre complexe  $w$  tel que  $w^n = z$ .

En utilisant la forme trigonométrique des nombres complexes, on obtient directement le résultat suivant. La preuve est un copier coller de celle qui concerne les racines carrées.

**Proposition 2.9.3.** *Tout nombre complexe non nul  $z$  admet  $n$  racines  $n$ -èmes distinctes.*

*Démonstration.* On considère donc  $z = \rho e^{i\theta}$  et on cherche  $z' = \rho' e^{i\theta'}$  tel que  $z'^n = z$ . Cette condition s'écrit encore  $\rho'^n e^{in\theta'} = \rho e^{i\theta}$ . On obtient donc les conditions équivalentes

$$\begin{cases} \rho'^n &= \rho \\ n\theta' &= \theta + 2k\pi, \quad k \in \mathbb{Z} \end{cases}$$

Cela donne directement  $\rho' = \sqrt[n]{\rho}$ . De plus, si on remarque que  $e^{i2\pi} = 1$ , on peut se limiter à  $k \in \{0, \dots, n-1\}$  et on trouve les solutions  $\sqrt[n]{\rho} e^{i(\frac{\theta}{n} + \frac{2k\pi}{n})}$ ,  $k \in \{0, \dots, n-1\}$ .  $\square$

Un cas particulier intéressant est le suivant.

**Définition 2.9.3.** On appelle racine  $n$ -ème de l'unité tout nombre complexe  $w$  tel que  $w^n = 1$ .

Le résultat suivant est en grande partie un corollaire de la proposition 2.9.3, plus précisément de la forme des racines  $n$ -èmes trouvée dans la preuve de cette proposition.

**Proposition 2.9.4.** *Les racines  $n$ -èmes de l'unité sont les nombres*

$$e^{i\frac{2k\pi}{n}}, \quad k \in \{0, \dots, n-1\}.$$

On en déduit directement les propriétés suivantes.

**Proposition 2.9.5.** — *Les racines  $n$ -èmes de l'unité sont les puissances successives de  $\omega_n = e^{i\frac{2\pi}{n}}$ .*

- *Si  $n \geq 2$ , la somme des racines  $n$ -èmes de l'unité est nulle.*
- *Les racines  $n$ -èmes de l'unité ont un module égal à 1.*
- *Les racines  $n$ -èmes de l'unité sont les sommets d'un polygone régulier à  $n$  côtés, inscrit dans le cercle trigonométrique.*
- *L'ensemble des racines  $n$ -èmes de l'unité est un sous-groupe de  $(\mathbb{C}_0, \cdot, 1)$ , noté  $U_n$ .*
- *L'ensemble des racines  $n$ -èmes de  $w$  s'écrit  $\{w_0 \omega_n^k : k \in \{0, \dots, n-1\}\}$ , si  $w_0$  est une racine  $n$ -ème de  $w$ .*

Terminons cette section par une petite discussion sur les racines  $n$ -èmes primitives de l'unité.

**Définition 2.9.4.** Une racine  $n$ -ème primitive de l'unité est une racine de l'unité  $\omega$  telle que  $n$  soit le plus petit  $t \in \mathbb{N}_0$  tel que  $\omega^t = 1$ .

**Exemple 2.9.1.** Le nombre  $-1$  n'est pas une racine *quatrième primitive* de l'unité. Parce que bien sûr  $(-1)^4 = 1$ , mais on a aussi  $(-1)^2 = 1$ , donc 4 n'est pas la plus petite puissance de  $-1$  qui soit égale à 1. Mais  $i$  et  $-i$  sont bien des racines quatrièmes primitives.

**Proposition 2.9.6.** *Si  $\omega$  est une racine  $n$ -ème primitive de l'unité, alors*

$$U_n = \{\omega, \omega^2, \dots, \omega^n = 1\}.$$

*Démonstration.* Il est facile de démontrer que les nombres  $\omega, \omega^2, \dots, \omega^n$  sont des racines  $n$ -èmes de l'unité :  $(\omega^k)^n = (\omega^n)^k = 1^k = 1$ . Pour terminer il suffit de montrer que ces nombres sont tous distincts. Mais si il existe  $k, l$  distincts compris entre 1 et  $n$  tels que  $\omega^k = \omega^l$ , si  $l > k$ , on a alors  $\omega^{l-k} = 1$ , et  $k - l < n$ , ce qui est contraire à l'hypothèse.  $\square$

On a aussi assez directement ce dernier résultat, qui permet de déterminer les racines  $n$ -èmes primitives de l'unité.

**Proposition 2.9.7.** *Si  $w = e^{i\frac{2k\pi}{n}}$  est une racine  $n$ -ème primitive, alors  $k$  et  $n$  n'ont pas de facteur commun.*

*Démonstration.* Si  $k$  et  $n$  ont un facteur commun  $q$ , on a  $k = k'q$  et  $n = n'q$ , où  $n < n'$ . Alors  $e^{i\frac{2k\pi}{n}} = e^{i\frac{2k'\pi}{n'}}$ , ce qui montre que  $w$  est une racine  $n'$ -ème de l'unité et n'est donc pas une racine  $n$ -ème primitive.  $\square$

*Remarque 2.16.* La réciproque de cette proposition est vraie, mais elle est plus délicate à établir.

# Chapitre 3

## Nombres naturels

Dans ce chapitre, nous allons étudier la définition axiomatique des nombres naturels due à Peano. Cela nous permettra d'utiliser intensivement la méthode de démonstration par récurrence (ou induction), qui découle de la définition. Ainsi, nous verrons comment l'induction permet de définir les opérations sur les nombres naturels et démontrer leurs propriétés. Bien sûr, ces propriétés ne sont pas inconnues : vous savez comment on compte et ce que l'on peut attendre comme propriétés des opérations. Les propriétés ne doivent pas vous surprendre et c'est dans les techniques de démonstration que se cache la difficulté de ce chapitre.

### 3.1 La définition de $\mathbb{N}$ et les récurrences

L'idée de la définition est que l'on part d'un nombre particulier, noté 0, et que l'on définit les propriétés de l'opération qui à tout nombre associe un successeur. L'axiomatique initiale de Peano définissait également les propriétés de l'égalité entre deux nombres, en précisant qu'il s'agit d'une relation d'équivalence (que nous définirons dans la suite). Nous allons donner cette définition axiomatique de  $\mathbb{N}$ . Il faudrait normalement montrer que les axiomes sont indépendants les uns des autres (aucun sous-ensemble d'axiomes n'implique un des autres axiomes) et qu'ils ne sont pas contradictoires, en montrant par exemple comment construire un ensemble ayant les propriétés demandées. Nous admettrons ces faits.

**Définition 3.1.1.** L'ensemble  $\mathbb{N}$  est défini par les conditions suivantes :

1. Il existe un nombre, noté 0, appartenant à  $\mathbb{N}$  ;
2. Tout nombre  $n \in \mathbb{N}$  admet un successeur unique, noté  $s(n)$ . Deux nombres distincts ont des successeurs distincts<sup>a</sup> ;
3. Le nombre 0 n'est le successeur d'aucun nombre ;
4. Si  $K$  est un ensemble tel que
  - on a  $0 \in K$
  - pour tout  $n \in \mathbb{N}$ , si  $n \in K$ , alors  $s(n) \in K$ ,alors  $K$  contient  $\mathbb{N}$ .

Les trois premiers axiomes sont assez naturels quand on pense à notre façon de compter : on ajoute des nombres petit à petit comme on ajoute des éléments à un ensemble. On définira l'addition de façon telle que le successeur d'un élément  $n$  soit  $n + 1$ . Le fait que 0 ne soit pas un successeur implique que  $s(0)$  est différent de 0. Puisque deux nombres différents ont des successeurs différents,  $s(s(0))$  et  $s(0)$  sont différents, et ainsi de suite. Le dernier axiome permet de montrer que  $\mathbb{N}$  est bien l'ensemble des successeurs que l'on construit de la sorte. Il permet cependant d'introduire une nouvelle technique de démonstration : la démonstration par récurrence. Cette technique permet de démontrer qu'une propriété dépendant d'un paramètre naturel est vraie pour toute valeur de ce paramètre.

---

a. En termes d'application :  $s : \mathbb{N} \rightarrow \mathbb{N}$  est une application injective telle que  $0 \notin \text{im}(s)$ .

**Proposition 3.1.1.** *Si pour tout  $n \in \mathbb{N}$ , on se donne une assertion  $P(n)$ , et si les conditions suivantes sont satisfaites :*

1. *L'assertion  $P(0)$  est vraie ;*
2. *Pour tout  $n$ , si  $P(n)$  est vrai alors  $P(s(n))$  est vrai ;*

*alors l'assertion  $P(n)$  est vraie pour tout  $n \in \mathbb{N}$ .*

*Démonstration.* Il suffit d'appliquer le dernier axiome de Peano à l'ensemble  $K = \{n \in \mathbb{N} : P(n) \text{ est vrai}\}$ .  $\square$

Cette méthode de démonstration est aussi appelée démonstration par induction. On constate qu'elle se décompose en deux parties. La première est appelée *cas de base* : il s'agit de montrer que  $P(0)$  est vraie. La deuxième est appelée *récurrence* ou *induction*. Elle se présente comme une implication "si  $P(n)$  est vraie, alors  $P(s(n))$  est vraie". La première assertion de cette implication " $P(n)$  est vraie" est appelée "hypothèse de récurrence" ou "hypothèse d'induction".

*Remarque 3.1.* Il ne faut pas confondre la partie d'induction, qui se présente comme une implication, vraie pour tout  $n \in \mathbb{N}$ , et la conséquence "l'assertion  $P(n)$  est vraie pour tout  $n \in \mathbb{N}$ ". Dans un premier temps, pour les reconnaître facilement, on peut utiliser des lettres différentes et par exemple écrire l'induction par "Pour tout  $k$ , si  $P(k)$  est vraie, alors  $P(s(k))$  est vraie". On formule alors la conclusion avec  $n$ .

On n'est pas obligé de commencer à 0. Si on veut démontrer qu'une propriété  $P$  est vraie pour tout  $n \geq n_0$ <sup>b</sup>, on peut prendre  $P(n_0)$  comme cas de base, puis démontrer l'induction suivante "Pour tout  $n \geq n_0$ , si  $P(n)$  est vraie, alors  $P(s(n))$  est vraie".

Comme exemple simple, on peut démontrer (en supposant que l'on sait additionner les nombres naturels) que pour tout  $n$  naturel, on a  $1 + \dots + n = \frac{n(n+1)}{2}$ . Voici un exercice que l'on peut faire, dès qu'on a vu les sommes.

**Exercice 3.1.1.** Démontrer que l'on a pour tout  $n \geq 1$  :

$$\sum_{j=1}^n (-1)^j j^2 = \frac{(-1)^n n(n+1)}{2}$$

et

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

Ce principe de démonstration par récurrence est très efficace. Il peut également être généralisé en un principe légèrement différent : la récurrence dite forte. Dans la récurrence classique, l'induction s'écrit "Si  $P(k)$  est vrai, alors  $P(s(k))$  est vraie". Dans la récurrence forte, l'induction s'écrit "Si  $P(0), \dots, P(k)$  sont vraies, alors  $P(k+1)$  est vrai".

**Proposition 3.1.2.** *Si pour tout  $n \in \mathbb{N}$ , on se donne une assertion  $P(n)$ , et si les conditions suivantes sont satisfaites :*

1.  *$P(0)$  est vrai ;*
2. *Pour tout  $k$ , si  $P(0), \dots, P(k)$  sont vrais, alors  $P(k+1)$  est vrai ;*

*alors la propriété  $P(n)$  est vraie pour tout  $n \in \mathbb{N}$ .*

*Démonstration.* Il suffit d'appliquer la récurrence classique à la propriété " $P'(n) : P(k)$  est vraie pour tout  $k \leq n$ ."  $\square$

---

b. Je définirai précisément l'ordre d'ici peu.

On peut montrer que la récurrence dite forte est équivalente à la récurrence classique. Elle permet également de démontrer une propriété  $P(n)$  pour tout  $n \geq n_0$ , comme la récurrence classique.

A titre d'exemple, considérons la décomposition de tout nombre naturel en un produit de nombres premiers. Nous utilisons des propriétés que nous allons revoir dans la suite de manière précise, mais que vous connaissez bien.

**Définition 3.1.2.** Un nombre naturel  $p$  est premier s'il admet exactement deux diviseurs, à savoir 1 et  $p$ .

En particulier, 0 n'est pas premier, 1 non plus car dans ce cas, il n'y a qu'un diviseur, et 2 est le seul nombre premier pair.

**Exemple 3.1.1.** Nous allons démontrer que tout nombre naturel supérieur ou égal à 2 se décompose en un produit de nombre premiers (on admet les produits à un seul facteur). En d'autres termes pour tout nombre  $n$  supérieur ou égal à 2, il existe des nombres premiers  $p_1, \dots, p_l$ , ( $l \geq 1$ ) (éventuellement égaux), tels que  $n = p_1 \cdots p_l$ .

**Cas de base :** Le nombre  $n = 2$  est premier, donc la condition précédente est satisfaite avec  $p_1 = 2$ .

**Induction forte :** Supposons que tout nombre naturel  $n \leq k$  se décompose en un produit de nombres premiers, et montrons que cette propriété est également vraie pour  $k + 1$ . Deux cas peuvent se produire. Soit  $k + 1$  est un nombre premier, et la propriété est vraie pour  $k + 1$ , avec  $p_1 = k + 1$ . Soit  $k + 1$  n'est pas un nombre premier, il admet alors un diviseur  $a$  différent de 1 et  $k + 1$ . On a donc  $k + 1 = ab$  pour deux nombres naturels  $a$  et  $b$  compris entre 2 et  $k$ . D'après l'hypothèse de récurrence (forte), les nombres  $a$  et  $b$  sont des produits de nombres premiers, et  $k + 1$  l'est donc aussi.

Le principe de démonstration par récurrence (forte) permet alors de conclure que la propriété demandée est vraie pour tout nombre supérieur ou égal à 2.

## 3.2 Addition et multiplication

Nous allons maintenant passer en revue les définitions des opérations sur les naturels en commençant par l'addition et la multiplication. Ces opérations sont définies de manière récursive. Le principe de démonstration par récurrence permet alors de démontrer que ces opérations sont bien définies sur  $\mathbb{N}$ . Nous donnerons alors les propriétés importantes des opérations, que vous connaissez bien depuis toujours. Ces propriétés se démontrent, bien entendu, par récurrence.

**Définition 3.2.1.** L'addition est définie récursivement comme suit. Pour tout  $a \in \mathbb{N}$ ,

1. on pose  $a + 0 = a$ .
2. pour tout  $b \in \mathbb{N}$ ,  $a + s(b) = s(a + b)$ <sup>a</sup>.

Cette définition fait appel à la récurrence : pour tout  $a \in \mathbb{N}$ , on considère l'ensemble  $K_a = \{b : a + b \text{ est défini}\}$ . Cet ensemble contient 0 par le premier point de la définition. S'il contient un nombre  $b$ , alors il contient  $s(b)$ , par le deuxième point de la définition. L'ensemble  $K_a$  est donc égal à  $\mathbb{N}$  quel que soit  $a$ , et on voit que l'addition  $a + b$  est définie quels que soient  $a$  et  $b$  dans  $\mathbb{N}$ .

**Définition 3.2.2.** On note 1 le successeur de 0.

On voit alors qu'en appliquant la définition pour  $b = 0$ , on a  $a + 1 = s(a)$ , quel que soit  $a$ . La définition plus haut a donc consisté à définir l'addition de  $b + 1$  à partir de l'addition de  $b$ . On a posé la condition naturelle  $a + (b + 1) = (a + b) + 1$ , what else ?

Passons maintenant aux propriétés de l'addition.

<sup>a</sup>. Cette condition définit l'addition de  $s(b)$  à  $a$  à partir de l'addition de  $b$  à  $a$ , on ne tourne donc pas en rond.

**Proposition 3.2.1.** *L'addition a les propriétés suivantes :*

1. Elle est associative : on a  $(a + b) + c = a + (b + c)$ , pour tous  $a, b, c \in \mathbb{N}$  ;
2. Elle admet 0 pour neutre : on a  $a + 0 = 0 + a = a$  pour tout  $a \in \mathbb{N}$  ;
3. Elle est commutative : on a  $a + b = b + a$  pour tous  $a, b \in \mathbb{N}$ .

Avant de passer à la démonstration, introduisons un peu de vocabulaire : les deux premières conditions font de la structure  $(\mathbb{N}, +, 0)$  un monoïde. On résume les trois conditions en indiquant que  $(\mathbb{N}, +, 0)$  un monoïde commutatif.

*Démonstration.* Démontrons tout d'abord la deuxième assertion. On sait que  $a + 0 = a$ , quel que soit  $a$ , par définition de l'addition. Démontrons donc que  $0 + a = a$  pour tout  $a$ . Cette démonstration se fait par récurrence sur  $a$ .

**Cas de base :** Pour  $a = 0$ , l'assertion à démontrer s'écrit  $0 + 0 = 0$ . Elle est par définition de l'addition.

**Induction :** Soit  $k \in \mathbb{N}$ . Supposons que  $0 + k = k$ , et montrons que  $0 + s(k) = s(k)$ . On calcule à l'aide de la définition :

$$0 + s(k) = s(0 + k) = s(k),$$

ce qu'il fallait démontrer. Bien sûr, il faut justifier : la première égalité vient de la définition de l'addition, tandis que la deuxième vient de l'hypothèse d'induction.

Passons maintenant à la démonstration de la troisième assertion de la proposition. Fixons  $a$  arbitraire dans  $\mathbb{N}$  et procédons par récurrence sur  $b$ .

**Cas de base :** Pour  $b = 0$ , la propriété à démontrer s'écrit  $a + 0 = 0 + a$ . C'est le point 2. de la proposition, que nous venons de démontrer.

**Induction :** Supposons que la propriété est vraie pour un nombre  $b$  et montrons la pour  $s(b)$ . On calcule

$$a + s(b) = s(a + b) = s(b + a).$$

Il reste à montrer que  $s(b + a) = s(b) + a$ , quels que soient  $a$  et  $b$ . Cette égalité se montre par récurrence sur  $a$ . Pour  $a = 0$ , elle s'écrit  $s(b + 0) = s(b) + 0$  et elle est vraie puisque 0 est neutre. On suppose qu'elle est vraie pour  $a$  et on montre qu'elle est vraie pour  $s(a)$  : on a

$$s(b + s(a)) = s(s(b + a)) = s(s(b) + a) = s(b) + s(a).$$

En effet, la première égalité vient de la définition de l'addition, la deuxième, de l'hypothèse de récurrence, et la troisième encore de la définition de l'addition. Nous avons donc montré que cette propriété est vraie pour tout  $a$ , pour  $b$  quelconque, c'est-à-dire pour tous  $a, b \in \mathbb{N}$ .

Finalement, la commutativité de l'addition est donc aussi démontrée par induction.

Passons maintenant à l'associativité. On fixe  $a$  et  $b$  quelconques dans  $\mathbb{N}$  et on procède par récurrence sur  $c$ .

**Cas de base :** Pour  $c = 0$ , l'assertion à démontrer s'écrit  $(a + b) + 0 = a + (b + 0)$ . Elle est vraie car les deux membres sont égaux à  $a + b$ , puisque 0 est neutre.

**Induction :** Supposons que la propriété est vraie pour  $k$  et montrons qu'elle est vraie pour  $s(k)$ . On calcule

$$(a + b) + s(k) = s((a + b) + k) = s(a + (b + k)) = a + s(b + k) = a + (b + s(k)).$$

Ici encore, la première égalité vient de la définition de l'addition, la deuxième de l'hypothèse de récurrence, la troisième et la quatrième, de la définition de l'addition. La propriété est donc vraie pour  $s(k)$  et on conclut par récurrence.  $\square$

Passons maintenant à la multiplication. On définit encore la multiplication (à droite) par 0. Sans surprise, le résultat est nul, puis on définit la multiplication par  $b + 1 = s(b)$ , à partir de la multiplication par  $b$ . Sans surprise non plus on voudrait avoir  $a.(b+1) = a.b + a$ . La définition n'est donc pas mystérieuse.

**Définition 3.2.3.** La multiplication est définie récursivement comme suit. Pour tout  $a \in \mathbb{N}$ ,

1. on pose  $a.0 = 0$ .
2. pour tout  $b \in \mathbb{N}$ , on pose  $a.s(b) = a.b + a$ .

Ici encore, il s'agit d'une définition récursive. La multiplication a les propriétés habituelles vis-à-vis d'elle-même et vis-à-vis de l'addition.

**Proposition 3.2.2.** *La multiplication distribue l'addition : on a  $a.(b + c) = a.b + a.c$  et  $(b + c).a = b.a + c.a$  pour tous  $a, b, c \in \mathbb{N}$ . De plus la structure  $(\mathbb{N}, ., 1)$  est un monoïde commutatif.*

*Démonstration.* La preuve se fait par récurrence, de manière analogue à la preuve des propriétés de l'addition. Il y a sans doute un ordre plus simple pour démontrer les propriétés. Je laisse cette démonstration comme exercice.  $\square$

### 3.3 L'ordre usuel sur $\mathbb{N}$

Nous avons déjà utilisé la notation  $n \leq n'$ , mais nous n'avons pas donné de définition de l'ordre en question, à partir de la définition de  $\mathbb{N}$  et des opérations que nous avons introduites dans la section précédente. Avant de traiter ce cas particulier, nous donnons la définition d'une relation en général et d'une relation d'ordre sur un ensemble en particulier.

#### 3.3.1 Relations et ordres

Passons à la définition des relations d'ordre. Ce sont des relations particulières.

**Définition 3.3.1.** Une relation  $\mathcal{R}$  de  $A$  dans  $A$  est une relation d'ordre si elle satisfait les conditions suivantes :

1. Elle est réflexive : on a  $a\mathcal{R}a$  pour tout  $a \in A$  ;
2. Elle est antisymétrique : pour tous  $a, b \in A$ , si  $a\mathcal{R}b$  et  $b\mathcal{R}a$ , alors on a  $a = b$  ;
3. Elle est transitive : pour tous  $a, b, c \in A$ , si  $a\mathcal{R}b$  et  $b\mathcal{R}c$ , alors on a  $a\mathcal{R}c$ .

Voici deux exemples.

**Exemple 3.3.1.** 1. Comme premier exemple simple, on peut citer l'égalité. Définissons dans un ensemble  $A$  quelconque  $a\mathcal{R}b$  si, et seulement si  $a = b$ . On peut vérifier que cette relation vérifie les trois conditions énoncées dans la définition.

2. Soit  $X$  un ensemble. On considère l'ensemble  $A = \mathcal{P}(X)$  des parties de  $X$ <sup>a</sup>. On a  $\mathcal{P}(X) = \{Y : Y \subset X\}$ . Alors l'inclusion définit un ordre : on pose  $Y_1\mathcal{R}Y_2$  si, et seulement si,  $Y_1 \subset Y_2$ . Cette relation satisfait également les trois conditions de la définition.

Une relation d'ordre est généralement notée  $\leq$ . Elle donne lieu à une autre relation d'ordre, dite duale, notée  $\geq$ . Elle est définie par  $a \geq b$  si, et seulement si  $b \leq a$ . On peut bien sûr définir des relations associées (qui ne sont pas des relations d'ordre)  $<$  et  $>$  par  $a < b$  si, et seulement si " $a \leq b$  et  $a \neq b$ ".

Passons maintenant à la définition de l'ordre usuel sur  $\mathbb{N}$ .

**Définition 3.3.2.** L'ordre usuel sur  $\mathbb{N}$  est la relation  $\leq$  définie par

$$a \leq b \quad \text{si, et seulement si} \quad \exists c \in \mathbb{N} : b = a + c.$$

---

a. Cet ensemble est également noté  $2^X$ .

Nous allons bien sûr montrer qu'il s'agit d'une relation d'ordre. Pour cela nous aurons besoin d'un résultat intermédiaire, qui a son intérêt propre. Un tel résultat, préliminaire à une proposition plus importante, est appelé *lemme*.

**Lemme 3.3.1.** *Si  $a, b, c \in \mathbb{N}$  sont tels que  $a + b = a + c$ , alors  $b = c$ . Si  $a, b \in \mathbb{N}$  sont tels que  $a + b = 0$ , alors  $a = b = 0$ .*

*Démonstration.* Pour le premier point, on procède par récurrence sur  $a$ . Pour le cas de base, on considère  $a = 0$ . La propriété à démontrer est alors vraie puisque 0 est neutre pour l'addition. Passons à l'induction. On suppose que la propriété est vraie pour  $a$  et on la démontre pour  $s(a) = a + 1$ . On doit donc démontrer que si  $(a + 1) + b = (a + 1) + c$ , alors  $b = c$ . Mais puisque l'addition est associative, on a  $(a + 1) + b = a + (1 + b)$  et  $(a + 1) + c = a + (1 + c)$ . Donc l'égalité  $(a + 1) + b = (a + 1) + c$  implique  $a + (1 + b) = a + (1 + c)$ . Par hypothèse de récurrence, cette dernière égalité implique  $1 + b = 1 + c$ , ou encore  $b + 1 = c + 1$ . On obtient alors  $b = c$ , puisque  $b$  et  $c$  ont le même successeur. La propriété est donc démontrée par récurrence.

Passons à la deuxième partie. Démontrons tout d'abord que tout nombre non nul est le successeur d'un nombre naturel. On considère pour cela l'ensemble  $K = \{n \in \mathbb{N} : n = 0 \text{ ou } \exists m \in \mathbb{N} : n = s(m)\}$ . Cet ensemble contient 0, et si  $n \in K$ , alors  $s(n) \in K$ , par définition. Donc  $K = \mathbb{N}$  et tout nombre est soit nul, soit le successeur d'un autre nombre.

Ensuite, on procède par contraposée : on suppose que soit  $a \neq 0$ , soit  $b \neq 0$  et on montre que  $a + b \neq 0$ . On peut supposer que  $b \neq 0$ . Il existe alors  $b' \in \mathbb{N}$  tel que  $b = s(b')$ . On a alors

$$a + b = a + s(b') = s(a + b') \neq 0.$$

La première égalité vient de la définition de  $b'$ , la deuxième de la définition de l'addition, et la conclusion du fait que 0 n'est le successeur d'aucun nombre.  $\square$

Nous pouvons maintenant démontrer que l'ordre usuel est bien un ordre.

**Proposition 3.3.1.** *La relation  $\leq$  introduite à la définition 3.3.2 est un ordre.*

*Démonstration.* La relation est réflexive. En effet, on a  $a \leq a$  pour tout  $a \in \mathbb{N}$ , car  $a = a + 0$ . La relation est antisymétrique. En effet, si on a  $a \leq b$  et  $b \leq a$ , alors il existe des nombres naturels  $c$  et  $c'$  tels que

$$\begin{cases} b &= a + c \\ a &= b + c'. \end{cases}$$

En substituant la valeur de  $b$  de la première équation dans la deuxième, on obtient

$$a = (a + c) + c' = a + (c + c').$$

D'après le lemme précédent, on obtient successivement  $c + c' = 0$ , puis  $c = c' = 0$ , donc  $b = a$ , puisque 0 est neutre pour l'addition. La relation est transitive. Supposons en effet avoir des naturels  $a, b, c$  satisfaisant  $a \leq b$  et  $b \leq c$ . Il existe alors des nombres naturels  $d$  et  $e$  tels que

$$\begin{cases} b &= a + d \\ c &= b + e. \end{cases}$$

On obtient alors  $c = (a + d) + e = a + (d + e)$  et on a  $a \leq c$ , par définition.  $\square$

Montrons maintenant que l'ordre est compatible avec les opérations.

**Proposition 3.3.2.** *Soient  $a, b, c \in \mathbb{N}$ . Si  $a \leq b$ , alors  $a + c \leq b + c$  et  $a \cdot c \leq b \cdot c$ .*

*Remarque 3.2.* On pourra démontrer facilement dans le cas de l'addition que la réciproque est vraie : si  $a + c \leq b + c$ , alors  $a \leq b$ . C'est un exercice. Le même constat vaut pour la multiplication pour autant que  $c$  soit non nul, mais il nous faudra en savoir un peu plus sur l'ordre pour le démontrer.

*Démonstration.* On traduit l'hypothèse : si  $a \leq b$ , il existe  $d \in \mathbb{N}$  tel que  $b = a + d$ . On obtient alors (en utilisant la commutativité et l'associativité)  $b + c = (a + c) + d$ , donc  $a + c \leq b + c$ . De même, on obtient  $b.c = a.c + d.c$ , donc  $a.c \leq b.c$ .  $\square$

Passons maintenant à une autre propriété importante de l'ordre usuel : c'est un ordre total.

**Définition 3.3.3.** Un ordre  $\mathcal{R}$  sur  $A$  est total si pour tous  $a, b \in A$ , on a  $a\mathcal{R}b$  ou  $b\mathcal{R}a$ .

L'habitude de traiter des inéquations avec des nombres peut faire penser que les ordres doivent être totaux. Mais voici deux contre-exemples.

**Exemple 3.3.2.** Soit  $A = \{1, 2, 3, 4\}$ , la relation d'ordre  $\mathcal{R}$  donnée par l'égalité n'est pas totale, puisque 2 n'est pas en relation avec 3 et 3 n'est pas en relation avec 2. Un exemple un peu moins trivial est donné par  $A = \mathcal{P}(X)$  où  $X = \{1, 2, 3, 4\}$  et où l'ordre est donné par inclusion. Si  $Y_1 = \{1, 2\} \in A$  et  $Y_2 = \{2, 3\} \in A$ , alors  $Y_1$  n'est pas en relation avec  $Y_2$  et  $Y_2$  n'est pas en relation avec  $Y_1$ .

En ce qui concerne l'ordre usuel sur  $\mathbb{N}$ , on a le résultat suivant.

**Proposition 3.3.3.** *L'ensemble ordonné  $(\mathbb{N}, \leq)$  est totalement ordonné.*

*Démonstration.* Fixons  $m \in \mathbb{N}$ . On doit prouver que pour tout  $n \in \mathbb{N}$ , on a  $m \leq n$  ou  $n \leq m$ . Pour ce faire, définissons

$$K = \{n \in \mathbb{N} : n \leq m \text{ ou } m \leq n\},$$

et montrons que  $K = \mathbb{N}$ . On utilise encore l'axiome de récurrence, sous sa forme initiale.

On montre que 0 appartient à  $K$ . C'est clair : puisque  $m = 0 + m$ , on a  $0 \leq m$ . Ensuite, on suppose que  $n$  appartient à  $K$  et on montre que  $s(n) = n + 1$  appartient à  $K$ . Deux cas peuvent se produire.

- Si  $m \leq n$  alors  $m \leq n + 1$ , car s'il existe  $a \in \mathbb{N}$  tel que  $n = m + a$ , alors par associativité, on a  $n + 1 = m + (a + 1)$ .
- Si  $n \leq m$ , alors il existe  $k \in \mathbb{N}$  tel que  $m = n + k$ . On a ici aussi deux possibilités :
  - a) Si  $k = 0$ , alors  $m = n$ , et  $m = n \leq n + 1$ , donc  $n + 1 \in K$ ;
  - b) Si  $k \neq 0$ , alors on peut écrire  $k = k' + 1$ . On a alors  $m = n + k' + 1 = (n + 1) + k'$ , donc  $n + 1 \leq m$ , et  $n + 1 \in K$ .

On a donc  $K = \mathbb{N}$ , ce qu'il fallait démontrer.  $\square$

Une propriété également importante de l'ensemble ordonné  $(\mathbb{N}, \leq)$  est qu'il est bien ordonné.

**Définition 3.3.4.** Soit  $(E, \leq)$  un ensemble ordonné et  $A \subset E$ . Un élément  $m$  est un minimum de  $A$  si

1. On a  $m \in A$
2. Pour tout  $a \in A$ , on a  $m \leq a$ .

Il est important de ne pas confondre élément minimum et élément minimal, dont la définition suit. Ces deux notions coïncident dans les ensembles ordonnés que vous avez rencontrés jusqu'à présent, parce qu'ils sont totalement ordonnés.

**Définition 3.3.5.** Soit  $(E, \leq)$  un ensemble ordonné et  $A \subset E$ . Un élément  $m$  est minimal dans  $A$  si

1. On a  $m \in A$
2. Pour tout  $a \in A$ ,  $a \leq m$  implique  $a = m$ .

On constate assez facilement qu'un minimum de  $A$  est toujours minimal dans  $A$ . Par contre la réciproque n'est pas vraie. Prenons par exemple  $E = \mathcal{P}(\{1, 2, 3\})$  et  $A = E \setminus \{\emptyset\}$ , où l'ordre est donné par l'inclusion. Alors  $\{1\}$  est minimal, car il n'y a pas d'élément strictement inférieur dans l'ensemble. Il ne s'agit cependant pas d'un minimum, puisque  $\{1\}$  n'est pas inférieur (inclus) à  $\{2\}$ .

**Définition 3.3.6.** Un ensemble ordonné  $(E, \leq)$  est *bien ordonné* si tout sous-ensemble non vide  $A$  de  $E$  admet un minimum.

Nous avons alors le résultat suivant.

**Proposition 3.3.4.** *L'ensemble ordonné  $(\mathbb{N}, \leq)$  est bien ordonné<sup>b</sup>.*

Avant de passer à la preuve, nous avons besoin d'un lemme technique.

**Lemme 3.3.2.** *Pour tout  $k \in \mathbb{N}$ , l'ensemble  $A_k = \{n \in \mathbb{N} : k < n < k + 1\}$  est vide.*

*Démonstration.* On procède par l'absurde. Supposons qu'il existe  $k \in \mathbb{N}$  tel que  $A_k \neq \emptyset$ . Soit alors  $n \in A_k$ . Puisque  $n > k$ , il existe  $c \in \mathbb{N}_0 = \mathbb{N} \setminus \{0\}$  tel que  $n = k + c$ . De même, puisque  $n < k + 1$ , il existe  $c' \in \mathbb{N}_0$  tel que  $k + 1 = n + c'$ . On substitue la valeur de  $n$  dans la deuxième équation, et on obtient la condition

$$k + 1 = k + c + c'$$

qui donne  $1 = c + c'$ . Puisque  $c' \neq 0$ , il existe  $c'' \in \mathbb{N}$  tel que  $c' = c'' + 1$ . On substitue dans l'équation précédente pour obtenir  $c + c'' = 0$ , qui donne  $c = 0$ , une contradiction.  $\square$

Passons maintenant à la preuve de la proposition 3.3.4.

On procède par contraposée. On considère un ensemble  $E \subset \mathbb{N}$  tel que  $E$  n'admet pas de minimum. On montre alors que  $E$  est vide. On utilise la propriété " $P(n)$  : pour tout  $i \leq n$ ,  $i \notin E$ ", et on montre par récurrence que  $P(n)$  est vraie pour tout  $n \in \mathbb{N}$ .

**Cas de base :** Pour  $n = 0$ , la propriété est vraie, car pour tout  $i \leq n$ , on a  $i \leq 0$ , mais on a toujours  $i \geq 0$ , donc on a  $i = 0$ . On doit donc seulement montrer que  $0 \notin E$ . Mais si 0 appartenait à  $E$ , il serait un élément minimum de  $E$ , car 0 est inférieur à tout nombre naturel.

**Induction :** Supposons que la propriété  $P(n)$  est vraie et montrons, par l'absurde, que  $P(n + 1)$  est vraie. Si  $P(n + 1)$  n'est pas vraie, alors il existe  $i \leq n + 1$  tel que  $i \in E$ . Montrons que  $i = n + 1$  : puisque la propriété  $P(n)$  est vraie, aucun nombre inférieur ou égal à  $n$  n'est dans  $E$ . Donc  $i > n$ . De plus  $i \leq n + 1$ . D'après le lemme 3.3.2, on a  $i = n + 1 \in E$ , et pour tout  $k < n + 1$ ,  $k \notin E$ . Puisque  $\mathbb{N}$  est totalement ordonné, tous les éléments de  $E$  sont alors supérieurs ou égaux à  $n + 1$ , qui en est donc un minimum. C'est contraire à l'hypothèse sur  $E$ .

## 3.4 Soustraction et division

Nous étudions dans cette section quelques propriétés des opérations secondaires, la soustraction et la division, exacte, puis avec reste. Nous appliquons ces propriétés pour obtenir le lemme d'Euclide et en déduire l'unicité de la décomposition de tout nombre en produit de facteurs premiers (à l'ordre des facteurs près).

**Définition 3.4.1.** Si  $a, b \in \mathbb{N}$  sont tels qu'il existe  $c \in \mathbb{N}$  satisfaisant  $b = a + c$ , alors le nombre  $c$  est appelé la différence de  $b$  et  $a$  et on note  $c = b - a$ .

<sup>b</sup>. On dit aussi que l'ordre usuel est un bon ordre sur  $\mathbb{N}$ .

<sup>c</sup>. Cette formulation permet d'éviter la récurrence forte, que l'on aurait pu utiliser avec l'assertion " $P'(n) : n \notin E$ ".

On constate que par définition de l'ordre dans  $\mathbb{N}$ , le nombre  $b - a$  est défini si, et seulement si,  $a \leq b$ . Cette définition n'est cependant pas complète : il faut montrer que le nombre  $c$  intervenant dans la définition est unique, sinon le nombre  $b - a$  n'est pas défini sans ambiguïté : mon voisin Raoul<sup>a</sup> pourrait obtenir  $b = a + c'$  et moi  $b = a + c$ . Je définirais alors  $b - a = c$  et lui  $b - a = c'$ . Cela ne se peut.

**Proposition 3.4.1.** *Si  $a, b, c, c' \in \mathbb{N}$  sont tels que  $b = a + c$  et  $b = a + c'$ , alors  $c = c'$ .*

*Démonstration.* Par transitivité de l'égalité, on obtient  $a + c = a + c'$ , et donc  $c = c'$ , par le lemme 3.3.1.  $\square$

Passons maintenant à la division.

**Définition 3.4.2.** Soient  $a, b \in \mathbb{N}$ , si il existe  $c \in \mathbb{N}$  tel que  $b = a.c$ , on dit que  $a$  divise  $b$ , ou que  $b$  est multiple de  $a$ . On note  $a|b$ .

Bien entendu, quand  $b = a.c$ , on veut définir la division de  $b$  par  $a$  comme étant le nombre  $c$  tel que  $a.c = b$ . Mais comme pour la soustraction, on a le problème de l'unicité, qui ne se laisse pas faire aussi facilement.

**Proposition 3.4.2.** *Tout nombre divise 0. Le nombre 0 ne divise que 0, mais pas de manière unique. Si  $a \neq 0$ , pour tout  $b$  il existe au plus un nombre  $c$  tel que  $b = a.c$ .*

*Démonstration.* Pour le premier point, on note qu'on a  $0 = a.0$ , quel que soit  $a$ , donc tout nombre  $a$  divise 0. Pour le deuxième, si  $a = 0$  et si  $b = a.c$ , alors  $b = 0$ , par définition de la multiplication. Dans ce cas, on a  $0 = 0.c$ , quel que soit  $c \in \mathbb{N}$ . Passons maintenant au cas où  $a$  est non nul, supposons qu'on a  $b = a.c$  et  $b = a.c'$ , et montrons qu'alors  $c = c'$ . On montre que  $a.c = a.c'$  implique  $c = c'$  par récurrence sur  $a \in \mathbb{N}_0$ .

**Cas de base :** pour  $a = 1$ , on a  $a.c = c$  et  $a.c' = c'$ , donc la propriété est évidente.

**Induction :** Supposons que la propriété est vraie pour  $a$  et montrons qu'elle est vraie pour  $s(a) = a + 1$ . On suppose donc qu'on a

$$(a + 1).c = (a + 1).c'$$

et on montre que cela implique  $c = c'$ . En développant la relation ci-dessus, on obtient

$$a.c + c = a.c' + c'.$$

Si  $c \neq c'$ , alors on a  $c > c'$  ou  $c' > c$  (puisque  $\mathbb{N}$  est totalement ordonné). Sans perte de généralité, on peut supposer que  $c' > c$ . Il existe alors  $k \neq 0$  tel que  $c' = c + k$ . L'équation ci-dessus devient alors

$$a.c + c = a(c + k) + c + k,$$

ou encore

$$a.k + k = 0.$$

On obtient alors  $k = 0$ , une contradiction.  $\square$

Au vu de cette proposition, on peut définir la division par un nombre non nul.

**Définition 3.4.3.** Soient  $a, b \in \mathbb{N}$  tels que  $a \neq 0$ . S'il existe  $c \in \mathbb{N}$  tel que  $b = a.c$ , on écrit  $c = b : a$  ou  $c = \frac{b}{a}$ . Le nombre  $c$  est appelé quotient de la division de  $b$  par  $a$ .

On obtient également comme corollaire le résultat suivant.

**Proposition 3.4.3.** *Si  $a, b \in \mathbb{N}$  satisfont  $a.b = 0$ , alors on a  $b = 0$  ou  $a = 0$ .*

*Démonstration.* Si  $a.b = 0$  et  $a \neq 0$ , alors on a  $a.b = a.0$ , et on conclut que  $b = 0$ , vu l'unicité démontrée dans la proposition précédente.  $\square$

---

<sup>a</sup>. Mon voisin Raoul est un imbécile, comme chacun sait. Il fait le malin parce que c'est le beau-frère de quelqu'un de connu, et dès qu'il peut, il ne fait pas comme moi.

Passons maintenant à une généralisation de la division que nous venons d'introduire, à savoir la division euclidienne.

**Proposition 3.4.4.** *Soient  $a \in \mathbb{N}$  et  $d \in \mathbb{N}_0$ . Il existe des nombres  $q$  et  $r$  satisfaisant les conditions  $a = qd + r$  et  $0 \leq r < d$ . De plus, le couple  $(q, r)$  est unique.*

Dans la proposition ci-dessus,  $a = qd + r$  s'appelle la division euclidienne de  $a$  par  $d$ . Le nombre  $d$  s'appelle le diviseur, le nombre  $q$  est le quotient et le nombre  $r$  est le reste de la division.

*Démonstration.* Démontrons l'existence par récurrence sur  $a$ ,  $d$  étant fixé. Pour  $a = 0$ , l'existence de  $q$  et  $r$  est claire : on a  $0 = 0.d + 0$ . C'est à dire  $q = 0$  et  $r = 0 < d$ . Supposons maintenant que la division existe pour  $a$  et montrons qu'elle existe pour  $a + 1$ . Si  $a = q.d + r$ , où  $r < d$ , alors on a  $a + 1 = q.d + r + 1$ . Deux cas peuvent se produire :

- Si  $r + 1 < d$ , alors on a une division de  $a + 1$  : le quotient est  $q$  et le reste  $r + 1$ .
- Si  $r + 1 \geq d$ , sachant que  $r < d$ , on trouve directement que  $r + 1 \leq d$ . On a donc  $r + 1 = d$ , et  $a + 1 = q.d + d = (q + 1).d$ . On a donc un reste nul, et la division euclidienne de  $a + 1$  par  $d$  existe.

Nous avons donc démontré l'existence par induction. Passons maintenant à l'unicité. Comme d'habitude, on suppose qu'on a deux couples  $(q, r)$  et  $(q', r')$  satisfaisant les conditions de la division euclidienne d'un nombre  $a$  et on montre qu'ils sont égaux. Supposons donc que l'on a

$$a = q.d + r, \quad r < d, \quad \text{et} \quad a = q'.d + r', \quad r' < d. \quad (3.1)$$

Montrons tout d'abord  $q = q'$ . Si tel n'est pas le cas, on a  $q > q'$  ou  $q' > q$ . On suppose sans perte de généralité que  $q' > q$ . Il existe alors  $k \in \mathbb{N}_0$  tel que  $q' = q + k$ . Les équations ci-dessus impliquent alors

$$q.d + r = (q + k).d + r'$$

ou encore  $r = k.d + r'$ . On a donc  $r \geq r'$  et  $r - r' = k.d$ . Mais on a  $r - r' \leq r < d$  et  $k.d \geq d$  car  $k \geq 1$ . L'égalité est donc impossible, c'est absurde. On a donc  $q = q'$  et les conditions dans (3.1) impliquent alors  $r = r'$ .  $\square$

Remarquons que le reste de la division de  $a$  par  $d$  permet de donner une condition nécessaire et suffisante pour que  $d$  divise  $a$  : il faut et il suffit que le reste de la division euclidienne de  $a$  par  $d$  soit nul.

Nous allons terminer ce chapitre sur les nombres naturels par le théorème fondamental de l'arithmétique.

**Théorème 3.4.1.** *Tout nombre naturel supérieur ou égal à 2 se décompose en un produit de facteur premiers (éventuellement réduit à un seul facteur). La décomposition est unique à l'ordre des facteurs près.*

La démonstration de ce théorème repose sur le lemme d'Euclide concernant la division d'un produit par un nombre premier.

**Lemme 3.4.1** (Euclide). *Si un nombre premier  $p$  divise le produit deux nombres  $a$  et  $b$  alors il divise  $a$  ou il divise  $b$ .*

Notons que l'hypothèse que  $p$  soit premier est fondamentale : en effet 4 divise  $12 = 2.6$ , mais 4 ne divise ni 2 ni 6.

*Démonstration.* • On procède par l'absurde : on suppose qu'il existe  $p$  premier, et  $a, b \in \mathbb{N}$  satisfaisant les conditions  $p|a.b$ ,  $p \nmid a$  et  $p \nmid b$ .

- On fixe un tel couple  $(p, a)$  et on définit l'ensemble  $E = \{b \in \mathbb{N} : p|a.b \text{ et } p \nmid b\}$ . Cet ensemble est non vide. Il ne contient pas 0 car  $p \nmid 0$ , et il ne contient pas 1 car  $p \nmid a$ . Il est non vide et il contient donc un élément minimum, que nous notons  $b_0$  (parce que  $\mathbb{N}$  est bien ordonné).

- On montre que  $1 < b_0 < p$ . D'une part, on sait que 0 et 1 n'appartiennent pas à  $E$ , donc  $b_0 > 1$ . Ensuite, on procède par l'absurde. Si  $b_0 \geq p$ , alors on peut effectuer la division euclidienne  $b_0 = q.p + r_0$ . On a alors  $r_0 > 0$ , car  $p \nmid b_0$  et  $r_0 < p \leq b_0$ . On montre que  $r_0 \in E$ , et cela contredit le fait que  $b_0$  est un minimum.
- On montre que  $p$  divise  $a.r_0$ . On a en effet  $a.b_0 = a.q.p + a.r_0$ , mais puisque  $p$  divise  $a.b_0$ , on a  $a.b_0 = q_0.p$  et on obtient  $a.r_0 = q_0.p - a.q.p = (q_0 - a.q).p$ , ce qui montre que  $p$  divise  $a.r_0$ .
- On sait que  $p \nmid r_0$  car  $0 < r_0 < p$ , donc la division euclidienne de  $r_0$  par  $p$  s'écrit  $r_0 = 0.p + r_0$ , et le reste est non nul.
- On trouve un élément  $b_1$  dans  $E$  qui est inférieur strictement à  $b_0$ . C'est une contradiction. On effectue la division euclidienne de  $p$  par  $b_0$ . On a

$$p = q_1.b_0 + b_1, \quad b_1 < b_0.$$

On a  $b_1 > 0$  car sinon  $b_0$  serait un diviseur de  $p$ , distinct de 1 et  $p$ . C'est impossible car  $p$  est premier. On montre que  $b_1$  appartient à  $E$  :

- On a  $p \nmid b_1$ , car  $0 < b_1 < b_0 < p$ .
  - On obtient que  $p \mid a.b_1$ , car  $a.b_1 = a.p - a.q_1.b_0 = a.p - q_1.q_0.p = p.(a - q_1.q_0)$ .
- Cela termine la preuve puisqu'on a obtenu une contradiction.  $\square$

Passons maintenant à la preuve du théorème fondamental de l'arithmétique.

*Preuve du Théorème fondamental.* Nous avons déjà démontré l'existence d'une factorisation en un produit de nombres premiers. Il nous reste à démontrer l'unicité. On procède par l'absurde et on suppose qu'il existe un nombre admettant deux décompositions en nombres premiers (au moins) distinctes. L'ensemble des nombres pour lesquels la décomposition n'est pas unique étant non vide, il admet un plus petit élément (puisque  $\mathbb{N}$  est bien ordonné). Notons  $n$  ce plus petit élément, et considérons deux décompositions de  $n$  :

$$n = p_1 \dots p_r = q_1 \dots q_s,$$

où  $p_1, \dots, p_r, q_1, \dots, q_s$  sont premiers. D'une part, aucun des nombres  $p_1, \dots, p_r$  n'est parmi  $q_1, \dots, q_s$ , sinon en divisant  $n$  par ce nombre, on obtient un nombre strictement inférieur à  $n$  qui admet deux décompositions distinctes. D'autre part, le nombre  $p_1$  divise  $n$ , donc il divise  $q_1 \dots q_s$ . Par le lemme d'Euclide (généralisé à un produit fini), on obtient que  $p_1$  divise un des facteurs de ce produit, disons  $q_j$ . Puisque  $q_j$  est premier,  $p_1$  est égal à 1 ou à  $q_j$ . Mais puisque  $p_1$  est premier,  $p_1 \neq 1$ , donc  $p_1 = q_j$ . Les deux affirmations que nous venons de faire sont contradictoires.  $\square$

Je termine ce chapitre sur les nombres par un résultat, classique également et que vous ne pouvez ignorer.

**Proposition 3.4.5.** *L'ensemble des nombres premiers est infini.*

*Démonstration.* Procédons par l'absurde et supposons qu'il y a un nombre fini de nombre premiers, que nous notons  $p_1, \dots, p_r$ . Le nombre  $n = p_1 \dots p_r + 1$  n'est alors divisible par aucun nombre premier (le reste de la division euclidienne par ces nombres est toujours 1). Il n'admet donc pas de décomposition comme un produit de facteurs premiers. C'est contraire au théorème fondamental.  $\square$

## Chapitre 4

# Bijections classiques, cardinal, relations d'équivalence

Nous savons ce qu'est une relation entre deux ensembles et nous avons déjà rencontré deux types particuliers de relation, à savoir les applications, parmi lesquelles on trouve les injections, les surjections et les bijections et les relations d'ordre. Dans ce chapitre, nous allons quelque peu approfondir les premières, en étudiant quelques bijections classiques, qui mènent aux notions de cardinal (d'un ensemble fini) et de dénombrabilité. Ensuite nous étudierons un troisième grand type de relation, *les relations d'équivalence*. Ce sont ces dernières qui permettent d'identifier (de coller ensemble) des objets a priori distincts. Elles jouent un rôle fondamental dans bon nombre de constructions mathématiques usuelles, notamment celles des nombres entiers, des rationnels, des champs modulaires ou encore, pour faire un peu de géométrie classique, dans celle des vecteurs libres.

### 4.1 Quelques bijections classiques, et un mot sur le cardinal

Voici quelques bijections classiques, qui sont parfois quelque peu surprenantes. Par exemple, si on vous demande s'il y a autant de nombres pairs que de nombres, vous aurez peut-être tendance à dire que non, puisqu'il manque les impairs. Cependant, il y a une correspondance parfaite (une bijection) entre l'ensemble des nombres pairs, et l'ensemble des nombres naturels.

**Définition 4.1.1.** Un nombre naturel est pair s'il est divisible par deux. Il est impair dans le cas contraire.

Nous avons vu la division euclidienne dans  $\mathbb{N}$ , et nous avons vu que  $n$  est pair si le reste de la division de  $n$  par 2 est nul, et qu'il est impair dans le cas contraire, c'est à dire si le reste de la division de  $n$  par 2 est 1. L'ensemble des nombres pairs est donc égal à  $2\mathbb{N} = \{2q : q \in \mathbb{N}\}$  tandis que l'ensemble des nombres impairs est  $I = \{2q + 1 : q \in \mathbb{N}\}$ . On a alors le résultat suivant.

**Proposition 4.1.1.** *Les ensembles  $\mathbb{N}$ ,  $2\mathbb{N}$  et  $I$  sont en bijection.*

*Démonstration.* On considère l'application  $f : \mathbb{N} \rightarrow \mathbb{N} : q \mapsto 2q$ . L'image de  $f$  est exactement  $2\mathbb{N}$ , donc  $f : \mathbb{N} \rightarrow 2\mathbb{N} : q \mapsto 2q$  est surjective. De plus  $f$  est injective : si  $f(q) = f(q')$ , alors  $2q = 2q'$  et donc  $q = q'$  (voir la section sur la multiplication des nombres entiers). On peut montrer que  $2\mathbb{N}$  et  $I$  sont en bijection. L'application

$$g : 2\mathbb{N} \rightarrow I : n \mapsto n + 1$$

est une bijection de  $2\mathbb{N}$  dans  $I$ . Tout d'abord, elle est bien définie : l'image d'un élément  $n \in 2\mathbb{N}$  est un élément impair : si  $n \in 2\mathbb{N}$ , alors il existe  $q \in \mathbb{N}$  tel que  $n = 2q$ . Alors  $g(n) = n + 1 = 2q + 1$  est un élément de  $I$ . Ensuite c'est une bijection : pour tout  $n \in I$  il

existe un unique  $n' \in \mathbb{N}$  tel que  $n = n' + 1$  (car  $n \neq 0$ ). Si  $n = 2q + 1$  ( $q \in \mathbb{N}$ ), alors on a  $n' = 2q$  et  $n'$  appartient à  $2\mathbb{N}$ .

En composant ces deux bijections, on obtient clairement une bijection entre  $\mathbb{N}$  et  $I$ . Pour être complet, la composition s'écrit

$$g \circ f : \mathbb{N} \rightarrow I : q \mapsto g(f(q)) = g(2q) = 2q + 1.$$

□

Voici un autre exemple, qui peut être présenté de manière imagée au moyen de l'hôtel de Hilbert<sup>a</sup>.

**Proposition 4.1.2.** *Les ensembles  $\mathbb{N} \times \{0, 1\}$  et  $\mathbb{N}$  sont en bijection.*

*Démonstration.* L'application

$$f : \mathbb{N} \times \{0, 1\} \rightarrow \mathbb{N} : (a, b) \mapsto 2a + b$$

définit une bijection. Bien sûr, à chaque couple  $(a, b)$  on associe bien un nombre naturel  $2a + b$ . Pour montrer que  $f$  est une bijection, il faut prouver que pour tout  $n \in \mathbb{N}$ , il existe un unique couple  $(a, b) \in \mathbb{N} \times \{0, 1\}$  tel que  $f((a, b)) = n$ . Cette dernière condition s'écrit  $n = 2a + b$ , et l'existence et l'unicité de  $a$  et  $b$  découlent de la division euclidienne par 2. □

On pourrait croire que tout fonctionne parce qu'on a considéré  $\mathbb{N} \times \{0, 1\}$ , c'est à dire deux copies de  $\mathbb{N}$ . Mais on peut en fait considérer un nombre infini de copies.

**Proposition 4.1.3.** *Les ensembles  $\mathbb{N} \times \mathbb{N}$  et  $\mathbb{N}$  sont en bijection.*

*Démonstration.* Il existe une bijection célèbre entre ces deux ensembles. Elle est due à Georg Cantor<sup>b</sup>. Il s'agit de numérotter les couples de points à coordonnées naturelles du plan. On commence par découper le plan en tranches  $T_n = \{(a, b) \in \mathbb{N} \times \mathbb{N} : a + b = n\}$ ,  $n \in \mathbb{N}$ . Dans chaque tranche, on numérote les éléments selon les ordonnées croissantes. Cela donne la numérotation, qui est une bijection de  $\mathbb{N} \times \mathbb{N}$  dans  $\mathbb{N}$ . Sachant que chaque tranche  $T_n$  contient  $n + 1$  éléments, le numéro du couple  $(i, j)$  est donné par

$$f(i, j) = \frac{(i + j)(i + j + 1)}{2} + j.$$

□

Passons maintenant à la définition précise du cardinal d'un ensemble fini, et des ensembles dénombrables.

**Définition 4.1.2.** Un ensemble  $A$  est fini s'il est vide ou s'il existe  $n \in \mathbb{N}_0$ , tel que  $A$  soit en bijection avec  $\{0, \dots, n - 1\}$ . Si  $A = \emptyset$ , le cardinal de  $A$ , noté  $|A|$  ou  $\sharp A$  est nul. Si  $A$  est en bijection avec  $\{0, \dots, n - 1\}$ , alors le cardinal de  $A$  est  $n$ .

Il faut bien sûr vérifier que la définition a un sens, c'est-à-dire qu'un ensemble  $A$  ne peut être en bijection avec deux ensembles  $\{0, \dots, n - 1\}$  et  $\{0, \dots, m - 1\}$ , pour  $m \neq n$ , et que l'ensemble vide n'est en bijection avec aucun de ces ensembles.

**Lemme 4.1.1.** *Si  $f : A \rightarrow B$  est une bijection, pour tout  $A' \subset A$ ,  $f|_{A'} : A' \rightarrow f(A')$  est une bijection.*

---

a. Proposé par David Hilbert, (Mathématicien Allemand 1862-1943). Il s'agit d'un hôtel qui a un nombre de places numérotées comme les entiers naturels. Elles sont toutes occupées. Un car arrive avec également un nombre de clients infini, mais ayant chacun un numéro d'ordre (un entier naturel) pour être logés. La réception indique alors qu'il n'y a aucun problème.

b. Mathématicien Allemand, (1845-1918).

*Démonstration.* La restriction n'altère pas le caractère injectif de  $f$ . Bien sûr, puisque l'ensemble d'arrivée de l'application  $f|_{A'}$  est son image, elle est surjective.  $\square$

**Lemme 4.1.2.** *L'ensemble vide n'est en bijection avec aucun ensemble non vide.*

*Démonstration.* Soit  $A$  un ensemble non vide. La seule relation  $\mathcal{R} : \emptyset \rightarrow A$  est la relation vide. C'est une relation de type application, mais elle ne peut pas être surjective.  $\square$

**Lemme 4.1.3.** *Pour  $n \geq 2$ , et pour tout  $i \in \{0, \dots, n-1\}$ , l'ensemble  $\{0, \dots, n-1\} \setminus \{i\}$  est en bijection avec  $\{0, \dots, n-2\}$ .*

*Démonstration.* Il suffit de construire une bijection : on définit

$$f : \{0, \dots, n-1\} \setminus \{i\} \rightarrow \{0, \dots, n-2\} : j \mapsto \begin{cases} j & \text{si } j < i \\ j-1 & \text{si } j > i \end{cases}$$

On vérifie que  $f$  est à valeurs dans  $\{0, \dots, n-2\}$  et que c'est une bijection.  $\square$

**Proposition 4.1.4.** *Si  $m, n \in \mathbb{N}_0$  sont tels que  $\{0, \dots, n-1\}$  et  $\{0, \dots, m-1\}$  soient en bijection, alors on a  $m = n$ . En particulier le cardinal d'un ensemble fini est bien défini.*

*Démonstration.* On peut procéder par récurrence sur  $n$ . Pour  $n = 1$ , si on a une bijection  $f$  de  $\{0\}$  dans  $\{0, \dots, m-1\}$ , alors on doit avoir  $m = 1$ . En effet, dans le cas contraire,  $\{0, \dots, m-1\}$  contient au moins deux éléments distincts, disons  $x_0$  et  $x_1$  qui doivent être dans l'image de  $f$ . On a donc  $x_0 = f(0)$  et  $x_1 = f(0)$ , et donc  $x_0 = x_1$ , ce qui est une contradiction.

Supposons le résultat acquis pour  $n$  et montrons-le pour  $n+1$ . Si on a une bijection de  $\{0, \dots, (n+1)-1\}$  dans  $\{0, \dots, m-1\}$ , alors  $f(n)$  appartient à  $\{0, \dots, m-1\}$ , et  $f$  induit une bijection de  $\{0, \dots, n\} \setminus \{n\} = \{0, \dots, n-1\}$  dans  $\{0, \dots, m-1\} \setminus \{f(n)\}$ . Ce dernier ensemble est en bijection avec  $\{0, \dots, m-2\}$ . On a donc  $n = m-1$ , donc  $n+1 = m$ .  $\square$

Passons maintenant aux ensembles infinis. Un type important d'ensemble est formé de ceux dont on peut compter les éléments, ils sont infinis, mais dénombrables.

**Définition 4.1.3.** Un ensemble  $A$  est infini dénombrable s'il est en bijection avec  $\mathbb{N}$ . Par extension, un ensemble  $A$  est dénombrable s'il existe une injection  $i : A \rightarrow \mathbb{N}$ ; c'est-à-dire si il existe une bijection entre  $A$  et une partie de  $\mathbb{N}$ .

Nous avons démontré plus haut que l'ensemble des nombres pairs, ou des nombres impairs sont infinis dénombrables. Voici maintenant un exemple célèbre d'ensemble infini non dénombrable. Il est encore dû à G. Cantor et a été publié en 1891.

**Théorème 4.1.1** (Cantor). *L'ensemble  $E$  des suites à termes dans  $\{0, 1\}$  n'est pas en bijection avec  $\mathbb{N}$ .*

*Démonstration.* Procédons par l'absurde et on suppose que  $E$  est en bijection avec  $\mathbb{N}$ . On peut donc numéroter toutes les suites et on a  $E = \{s_i : i \in \mathbb{N}\}$ . On construit alors une suite à termes dans  $\{0, 1\}$  qui n'est égale à aucun des éléments de  $E$ . C'est une contradiction. Définissons la suite  $\mathbf{a} = (a_n)_{n \in \mathbb{N}}$  par  $a_n = 1 - (s_n)_n^c$ . Donc  $a_n = 0$  si le terme  $n$  de  $s_n$  est égal à 1 et vice-versa. Pour tout  $j \in \mathbb{N}$ , la suite  $\mathbf{a}$  est différente de  $s_j$ . En effet, leur terme numéro  $j$  est différent.  $\square$

Pour généraliser, on peut remarquer que l'ensemble  $E$  du théorème précédent est en bijection avec l'ensemble des sous-ensembles de  $\mathbb{N}$ .

**Proposition 4.1.5.** *Avec les notations du théorème précédent, l'application  $f : E \rightarrow \mathcal{P}(\mathbb{N}) : s \mapsto \{n \in \mathbb{N} : s_n = 1\}$  est une bijection.*

On peut alors généraliser le théorème précédent : pour tout ensemble  $S$ ,  $S$  et  $\mathcal{P}(S)$  ne sont pas en bijection.

---

c. On parle de l'argument diagonal, car si on écrit les suites l'une en dessous de l'autre, on forme un tableau carré infini, dont les termes  $(s_n)_n$  forment la diagonale.

## 4.2 Relations d'équivalence et quotients

Nous avons vu comment rendre l'application  $f : \mathbb{R} \rightarrow \mathbb{R} : x \rightarrow x^2$  bijective. Nous avons restreint l'ensemble d'arrivée à  $[0, +\infty[$  pour la rendre surjective, et nous avons restreint l'ensemble de départ à  $[0, +\infty[$  pour la rendre injective, tout en maintenant son caractère surjectif. Ce n'est pas la façon la plus utilisée : l'idée qui est suivie généralement est que si l'application n'est pas injective, c'est parce qu'il existe des éléments distincts de l'ensemble de départ qui ont la même image par  $f$ . L'idée est de définir un nouvel ensemble où l'on décide que ces éléments sont égaux. Ce nouvel ensemble est appelé quotient.

Vous avez déjà utilisé ce processus qui consiste à identifier des objets différents. Par exemple, la fraction  $\frac{1}{2}$  correspond à couper un gâteau en deux et à en prendre une partie, tandis que pour obtenir  $\frac{2}{4}$ , on coupe le gâteau en quatre parties égales et on en prend deux morceaux. Ces fractions sont différentes parce que le gâteau n'est pas coupé de la même façon, mais vous avez bien vite considéré qu'il s'agissait de la même quantité de gâteau et identifié les fractions *équivalentes*  $\frac{1}{2}$  et  $\frac{2}{4}$ .

De la même façon, on peut définir un vecteur, dans l'enseignement secondaire, comme étant un couple  $(A, B)$  formé d'une origine  $A$  et d'une extrémité  $B$ . On dit que c'est un vecteur lié en  $A$ . Bien entendu, on décide assez vite, poussé par des motivations physiques, de déclarer que deux vecteurs  $(A, B)$  et  $(C, D)$  sont équivalents (dans ce cas, on dit équipollents) si  $ABDC$  est un parallélogramme. Assez rapidement, on déclare que des vecteurs équipollents *représentent* le même vecteur (libre).

Enfin, vous avez sans doute construit des cylindres en collant les bords d'une feuille de papier. Le fait de coller consiste à identifier (ne plus distinguer) des points a priori distincts. Je vous invite à refaire l'expérience en collant les bords d'une feuille de papier après avoir tordu la feuille pour retourner l'un des bords. Vous obtiendrez ainsi un ruban de Moebius, une surface qui n'a qu'une face.

Mais revenons à nos moutons, on ne peut évidemment pas déclarer égaux des points a priori distincts sans aucune précaution : si je déclare que  $x_1$  est égal à  $x_2$ , alors  $x_2$  doit aussi être déclaré égal à  $x_1$  : la relation d'égalité est *symétrique*. Si je déclare que  $x_1$  devient égal à  $x_2$  et que  $x_2$  devient égal à  $x_3$ , alors je dois déclarer que  $x_1$  est égal à  $x_3$  : la relation d'égalité doit être transitive. Et bien entendu, avant que je ne décide quoi que ce soit, un point  $x_1$  est toujours égal à lui-même.

Cela conduit à la définition d'une relation d'équivalence.

**Définition 4.2.1.** Soit  $A$  un ensemble. Une relation  $\mathcal{R} : A \rightarrow A$  est une relation d'équivalence si les trois conditions suivantes sont satisfaites :

1. Elle est *réflexive* : on a  $a\mathcal{R}a$  pour tout  $a \in A$  ;
2. Elle est *symétrique* : pour tous  $a, b \in A$ , si  $a\mathcal{R}b$ , alors  $b\mathcal{R}a$  ;
3. Elle est *transitive* : pour tous  $a, b, c \in A$ , si  $a\mathcal{R}b$  et  $b\mathcal{R}c$ , alors  $a\mathcal{R}c$ .

Voici quelques exemples.

**Exemple 4.2.1.** 1. La relation d'égalité sur n'importe quel ensemble est une relation d'équivalence.

2. Considérons les interrupteurs à poussoir simples : une impulsion donne un changement d'état : si la lampe était allumée, elle est éteinte, et vice-versa. On suppose par exemple que l'on commence avec une lampe éteinte. On voit qu'en poussant 2 fois, 4 fois etc..., on arrive au même résultat : la lampe reste éteinte. Par contre, une impulsion, 3, 5... donnent le même résultat : la lampe est allumée. On peut définir cette relation d'un point de vue mathématique : tous les nombres pairs sont équivalents entre eux. Tous les nombres impairs sont équivalents entre eux. Aucun nombre pair n'est équivalent à un nombre impair. C'est un peu fastidieux. On peut résumer comme ceci :  $a\mathcal{R}b$  si, et seulement si,  $a$  et  $b$  ont la même parité. De manière équivalente, puisque tout se ramène à la divisibilité par 2 :  $a\mathcal{R}b$  si, et seulement si,

les reste de la division de  $a$  par 2 et de  $b$  par 2 sont égaux (on dit que  $a$  et  $b$  sont égaux modulo 2).

3. On peut généraliser cet exemple avec des interrupteurs à impulsions à plusieurs positions 0 : éteint, 1 : faible, 2 : moyen, 3 : tamisé, 4 : fort et 5 : à nouveau éteint. Deux nombres d'impulsions  $n$  et  $n'$  donneront visiblement le même résultat s'ils diffèrent par un multiple de 5 (on dit qu'ils sont égaux modulo 5). On peut définir la relation de la manière suivante :  $n\mathcal{R}n'$  si, et seulement si,

$$n = n' \text{ ou } (n > n', n = n' + 5k, k \in \mathbb{N}) \text{ ou } (n' > n, n' = n + 5k, k \in \mathbb{N}).$$

Si on dispose d'une relation d'équivalence sur un ensemble  $A$ , le quotient consiste à déclarer égaux entre eux tous les points qui sont équivalents. C'est l'objet de la définition suivante.

**Définition 4.2.2.** Soit  $\mathcal{R}$  une relation d'équivalence sur un ensemble  $A$ . Pour tout  $a \in A$ , la classe d'équivalence de  $a$ , notée  $[a]_{\mathcal{R}}$  ou  $[a]$  est l'ensemble des éléments équivalents à  $a$ . On a donc  $[a]_{\mathcal{R}} = \{b \in A : b\mathcal{R}a\}$ . Le quotient de  $A$  par  $\mathcal{R}$ , noté  $A/\mathcal{R}$  est l'ensemble formé par les classes d'équivalence :

$$A/\mathcal{R} = \{[a] : a \in A\}.$$

On peut donc voir une classe d'équivalence de deux façons différentes : c'est soit un sous-ensemble de  $A$  formé d'éléments équivalents entre eux pour  $\mathcal{R}$ , soit un point de  $A/\mathcal{R}$ . On a donc bien réalisé le programme demandé : les point équivalents entre eux dans  $A$  deviennent un seul point dans le quotient.

Cela permet de rendre une application injective. Considérons encore l'application  $f : \mathbb{R} \rightarrow \mathbb{R} : x \mapsto x^2$ . Elle n'est pas injective, par exemple parce que 2 et  $-2$  ont le même image. Définissons la relation  $\mathcal{R} : \mathbb{R} \rightarrow \mathbb{R}$  par  $x\mathcal{R}y$  si  $f(x) = f(y)$ , i.e.  $x^2 = y^2$ . On vérifie facilement qu'il s'agit bien d'une relation d'équivalence. La classe de  $x$ , vue comme sous-ensemble de  $\mathbb{R}$ , est  $[x] = \{-x, x\}$ . On obtient une application injective en définissant  $\tilde{f} : \mathbb{R}/\mathcal{R} \rightarrow \mathbb{R} : [x] \mapsto x^2$ .

Il existe une manière simple de se donner une relation d'équivalence : il suffit de se donner les classes. Ici encore, il y a cependant des conditions à satisfaire pour obtenir une relation d'équivalence. Rappelons tout d'abord ce qu'est une partition.

**Définition 4.2.3.** Soit  $A$  un ensemble. Une partition de  $A$  est une famille de sous-ensembles de  $A$  deux à deux disjoints et dont l'union est égale à  $A$ .

- Exemple 4.2.2.**
1. Soit  $A = \{1, 2, 3, 4, 5\}$ . Alors  $\{\{1, 2\}, \{3, 4\}, \{5\}\}$  est une partition de  $A$ .
  2. Soit  $A = \mathbb{N}$ . Posons  $A_1 = 2\mathbb{N}$ , l'ensemble des nombres pairs et  $A_2 = I$ , l'ensemble des nombres impairs. Alors  $\{A_1, A_2\}$  est une partition de  $\mathbb{N}$ .
  3. Soit  $A = \mathbb{R}^2$ . Pour tout  $x \in \mathbb{R}$ , définissons  $E_x = \{(x, y) : y \in \mathbb{R}\}$ . Alors  $\{E_x : x \in \mathbb{R}\}$  est une partition de  $\mathbb{R}^2$ .

Remarquons que l'on peut ajouter l'ensemble vide aux ensembles définissant une partition, cela donne une autre partition. S'il est présent, on peut aussi l'enlever, et on conserve une partition.

Avant d'exposer le lien entre classes d'équivalence et partitions, j'ai besoin d'un lemme.

**Lemme 4.2.1.** Soit  $A$  un ensemble et  $\mathcal{R}$  une relation d'équivalence sur  $A$ . Alors pour tous  $a, a' \in A$ , on a  $[a] = [a']$  si, et seulement si  $a\mathcal{R}a'$ .

*Démonstration.* Si  $[a] = [a']$ , puisque  $a \in [a]$  (la relation est réflexive), on a  $a \in [a']$  donc  $a\mathcal{R}a'$ . Pour la réciproque, on montre deux inclusions. Supposons  $a\mathcal{R}a'$  et soit  $x \in [a]$ . On a donc  $x\mathcal{R}a$  et par transitivité  $x\mathcal{R}a'$ , donc  $x$  appartient à  $[a']$ . On a donc montré  $[a] \subset [a']$ . On montre de même  $[a'] \subset [a]$ , ou on utilise la symétrie de l'hypothèse.  $\square$

Voici maintenant un premier résultat qui montre qu'une relation d'équivalence induit une partition.

**Proposition 4.2.1.** *Soit  $A$  un ensemble et  $\mathcal{R}$  une relation d'équivalence sur  $A$ . Les classes d'équivalence de  $\mathcal{R}$ , vues comme sous-ensembles de  $A$ , forment une partition de  $A$ .*

*Démonstration.* On démontre d'abord par double inclusion que l'union des classes est égal à  $A$ . Toute classe d'équivalence est un sous-ensemble de  $A$ . Il en va donc de même pour leur union.

D'autre part, si  $a \in A$ , alors  $a \in [a]$ , ce qui montre que  $A$  est inclus dans l'union des classes d'équivalence.

Montrons maintenant que des classes distinctes sont disjointes. Il est plus facile de montrer que des classes non disjointes sont nécessairement égales. Supposons donc que pour  $a, a' \in A$ , on a  $[a] \cap [a'] \neq \emptyset$  et montrons que  $[a] = [a']$ . Il existe  $c \in [a] \cap [a']$ . Puisque  $c \in [a]$ , on a  $c\mathcal{R}a$ , et puisque  $c \in [a']$  on a  $c\mathcal{R}a'$ . Par transitivité, on obtient  $a\mathcal{R}a'$  donc  $[a] = [a']$  par le lemme 4.2.1.  $\square$

Passons maintenant à la réciproque : une partition détermine une relation d'équivalence.

**Proposition 4.2.2.** *Soit  $A$  un ensemble et  $\{A_i : i \in I\}$  une partition<sup>a</sup> de  $A$  par des sous-ensembles non vides. Il existe une unique relation d'équivalence  $\mathcal{R}$  dont les classes sont les sous-ensembles  $A_i$ .*

Remarquons que l'on aurait pu considérer une partition quelconque et imposer que les classes de  $\mathcal{R}$  soient les sous-ensembles non vides constituant la partition.

*Démonstration.* L'unicité est claire puisqu'une relation d'équivalence est définie par ses classes et puisque les classes sont données : pour tout  $x \in A$ , il existe un unique  $i \in I$  tel que  $x \in A_i$ . Alors  $A_{i_0}$  est une classe qui contient  $x$ . C'est donc  $[x]$ . On doit donc poser  $y\mathcal{R}x$  si, et seulement si  $y \in A_{i_0}$ . En général, on doit donc poser  $z\mathcal{R}t$  si, et seulement si, il existe  $i \in I$  tel que  $z \in A_i$  et  $t \in A_i$ .

Il reste à démontrer l'existence. Pour ce faire, il faut démontrer que la proposition faite ci-dessus répond bien à la question, c'est-à-dire qu'il s'agit bien d'une relation d'équivalence et que les classes associées à cette relation sont les sous-ensembles  $A_i$ .

C'est une relation d'équivalence :

1. Elle est transitive. Pour tout  $x \in A$ , il existe  $i \in I$  tel que  $x \in A_i$ . On a alors  $x\mathcal{R}x$  ;
2. Elle est symétrique : si  $x, y \in A$  sont tels que  $x\mathcal{R}y$ , alors il existe  $i \in I$  tel que  $x \in A_i$  et  $y \in A_i$ . Alors on a également  $y\mathcal{R}x$ .
3. Elle est transitive, si  $x, y, z \in A$  sont tels que  $x\mathcal{R}y$  et  $y\mathcal{R}z$ , il existe  $i_0 \in I$  tel que  $x \in A_{i_0}$  et  $y \in A_{i_0}$ , et il existe  $j \in I$  tel que  $y \in A_j$  et  $z \in A_j$ . Puisque  $y$  appartient à  $A_{i_0} \cap A_j$ , on a  $i_0 = j$ , et par suite  $x\mathcal{R}z$ .

Soit  $x \in A$ , il existe  $i \in I$  tel que  $x \in A_i$ . On a alors  $y\mathcal{R}x$  si et seulement si  $y \in A_i$ , donc  $[x] = A_i$ . Donc toutes les classes s'écrivent  $A_i$  pour un certain  $i$ . Réciproquement, tout sous-ensemble  $A_i$  est une classe : puisque  $A_i \neq \emptyset$ , il existe  $x \in A_i$  et on a donc  $A_i = [x]$ .  $\square$

Pour compléter notre programme de rendre des applications injectives en identifiant les points qui empêchent l'injectivité, c'est-à-dire en considérant un quotient, il faut encore pouvoir définir une application sur un ensemble quotient. L'idée est de définir l'image d'une telle application sur une classe  $[x]$  par son comportement sur  $x$ , mais il y a un problème potentiel : une classe peut être décrite par plusieurs *représentants*. Voici des exemples.

a. Ici  $I$  représente un ensemble d'indices qui peut être fini comme dans les exemples 1 et 2 ci-dessus ou infini comme dans l'exemple 3.

**Exemple 4.2.3.** 1. Reprenons l'exemple des interrupteurs à 5 états. On sait que deux nombres d'impulsions sont équivalents s'ils diffèrent pas un multiple de 5. On a donc un quotient de  $\mathbb{N}$  par cette relation d'équivalence, que nous notons  $\mathcal{R}$ . Dans le quotient, on a  $[2] = [7]$ . On dit que 2 et 7 sont des représentants de la même classe. Supposons que l'on veuille définir

$$f : \mathbb{N}/\mathcal{R} \rightarrow \mathbb{N} : [x] \mapsto x^3.$$

Cette définition n'est pas correcte, car le résultat va dépendre du représentant choisi dans la classe, et pas de classe elle-même. En effet, si je considère le représentant 2 de la classe  $[2]$  alors j'écris  $f([2]) = 2^3 = 8$ . Mais je peux considérer le représentant 7, puisque  $[2] = [7]$ . Alors j'écris  $f([2]) = f([7]) = 7^3 = 343$ . Je n'ai donc pas défini une application sur l'ensemble des classes d'équivalence.

2. Avec les mêmes notations, on peut définir

$$f : \mathbb{N}/\mathcal{R} \rightarrow \mathbb{N}/\mathcal{R} : [x] \mapsto [x^3].$$

Le problème évoqué ci-dessus n'existe plus. En effet, on a  $f([2]) = [2^3] = [8] = [3]$ , et  $f([7]) = [343] = [3]$ . De manière générale, si  $n$  et  $n'$  représentent la même classe, alors  $n$  et  $n'$  diffèrent pas un multiple de 5, et c'est aussi le cas de leurs cubes. On n'a donc pas de problème de définition, car l'image d'une classe, bien que définie à l'aide d'un représentant, est indépendante de celui-ci.

3. Toujours avec les mêmes notations, on peut être tenté de définir

$$g : \mathbb{N}/\mathcal{R} \rightarrow \mathbb{N}/\mathcal{R} : [x] \mapsto [\lceil \sin(\frac{\pi}{2}x) \rceil]$$

On constate que cette définition n'est pas correcte, car l'application qui à  $x \in \mathbb{N}$  associe  $\lceil \sin(\frac{\pi}{2}x) \rceil$  n'associe pas les mêmes valeurs à 2 et à 7.

On voit qu'une façon raisonnable de définir une application sur un quotient d'un ensemble  $A$  par une relation d'équivalence  $\mathcal{R}$  est de la définir sur  $A$ , de façon telle qu'elle donne la même image à des points équivalents. Je ne le démontrerai pas, pour ne pas allonger ce chapitre déjà copieux, mais il n'y a en fait pas d'autre façon de définir une application sur un quotient.

**Définition 4.2.4.** Soient  $A, B$  deux ensembles et  $\mathcal{R}$  une relation d'équivalence sur  $A$ .

1. On dit qu'une application  $f : A \rightarrow B$  passe au quotient  $A/\mathcal{R}$  s'il existe une application  $\tilde{f} : A/\mathcal{R} \rightarrow B$  telle que  $\tilde{f}([a]) = f(a)$  pour tout  $a \in A$ .
2. On dit qu'une application  $f : A \rightarrow B$  est constante sur les classes de  $\mathcal{R}$  si la condition

$$a\mathcal{R}a' \Rightarrow f(a) = f(a')$$

est satisfaite pour tous  $a, a' \in \mathcal{R}$ .

Remarquons que quand une application  $f$  passe au quotient, l'application  $\tilde{f}$  qu'elle induit est visiblement unique (elle est définie par  $f$ ). Sans surprise, on a la proposition suivante. La démonstration consiste en de longues vérifications.

**Proposition 4.2.3.** Soient  $A, B$  deux ensembles et  $\mathcal{R}$  une relation d'équivalence sur  $A$ . Une application  $f : A \rightarrow B$  passe au quotient si, et seulement si, elle est constante sur les classes de  $\mathcal{R}$ .

Terminons cette section par un théorème sur le passage au quotient. Il permet toujours de rendre une application bijective. On utilise la restriction sur l'ensemble d'arrivée pour rendre l'application surjective, et le quotient sur l'ensemble de départ pour la rendre injective.

**Théorème 4.2.1.** *Soient  $A$  et  $B$  deux ensembles et  $f : A \rightarrow B$  une application. On définit la relation  $\mathcal{R}$  sur  $A$  par  $a\mathcal{R}a'$  si, et seulement si  $f(a) = f(a')$ . Alors  $f$  passe au quotient en une application bijective  $\tilde{f} : A/\mathcal{R} \rightarrow f(A)$ .*

*Démonstration.* Dans les conditions de l'énoncé, on montre les trois assertions suivantes :

1. l'application  $f$  passe au quotient ;
2. L'application  $\tilde{f}$  est injective ;
3. L'application  $\tilde{f}$  est surjective.

Pour la première, il suffit de démontrer que si  $a\mathcal{R}a'$  alors  $f(a) = f(a')$ . C'est évident vu la définition de  $\mathcal{R}$ .

Pour la deuxième, considérons  $[a], [a'] \in A/\mathcal{R}$  tels que  $\tilde{f}([a]) = \tilde{f}([a'])$ . On a alors  $f(a) = f(a')$  par définition de  $\tilde{f}$ . Cela implique  $a\mathcal{R}a'$  par définition de  $\mathcal{R}$  et donc  $[a] = [a']$ , ce qu'il fallait démontrer.

Pour la troisième, soit  $y \in f(A)$ . Il existe  $a \in A$  tel que  $y = f(a)$ . Alors on a aussi  $y = \tilde{f}([a])$  et  $y$  appartient à  $Im(\tilde{f})$ .  $\square$

Terminons par une application simple de ce théorème. Considérons l'application

$$f : [0, 1] \rightarrow \mathbb{R}^2 : x \mapsto \begin{pmatrix} \cos(2\pi x) \\ \sin(2\pi x) \end{pmatrix}$$

Cette application n'est pas surjective, mais vous connaissez bien son image. C'est le cercle de  $\mathbb{R}^2$  centré en  $\begin{pmatrix} 0 \\ 0 \end{pmatrix}$  et de rayon 1, que l'on note  $S^1$  (la sphère de dimension 1). Cette application n'est pas injective. On a en effet  $f(x) = f(y)$  si, et seulement si,

$$\begin{pmatrix} \cos(2\pi x) \\ \sin(2\pi x) \end{pmatrix} = \begin{pmatrix} \cos(2\pi y) \\ \sin(2\pi y) \end{pmatrix}.$$

Vu les propriétés des applications  $\sin$  et  $\cos$ , cette condition donne  $x = y$  ou  $x, y \in \{0, 1\}$ . Cela définit une relation d'équivalence  $\mathcal{R}$  où toutes les classes sont réduites à un singleton, sauf une, qui est formée par les deux bornes du segment  $[0, 1]$ . Le théorème nous enseigne qu'en identifiant les bornes du segment  $[0, 1]$ , on obtient un ensemble en bijection avec le cercle  $S^1$ .

# Chapitre 5

## Nombres et structures algébriques

L'objet de ce chapitre est de construire les nombres entiers (relatifs), c'est-à-dire l'ensemble  $\mathbb{Z}$ . Ce n'est évidemment pas suffisant : il faut munir  $\mathbb{Z}$  des opérations d'addition et de soustraction que l'on connaît bien. La construction utilisée ici est assez naturelle. On considère d'abord des couples de nombres naturels, que l'on peut penser comme un "gain" et une "perte". Bien sûr, si on gagne trois unités et puis qu'on en perd une, ou si on gagne 5 unités et puis on en perd trois, le résultat final est le même : ces couples sont équivalents. On construit donc  $\mathbb{Z}$  comme un quotient de  $\mathbb{N} \times \mathbb{N}$ .

On définit ensuite les opérations. Pour l'addition, elle est tellement évidente que je la présente directement. Puis on vérifie qu'elle a les bonnes propriétés pour faire de  $\mathbb{Z}$  un groupe commutatif. Pour la multiplication, on la définit à partir d'un "cahier des charges", appelé en mathématiques une *analyse du problème* : on demande que l'on puisse considérer les nombres naturels comme des éléments de  $\mathbb{Z}$ , et que les opérations pour ces nombres soient celles que l'on connaît déjà : on plonge donc  $\mathbb{N}$  dans  $\mathbb{Z}$ . On veut bien sûr que la multiplication distribue l'addition. On démontre qu'il n'y a qu'un seul choix possible. Il reste à définir la multiplication selon le seul choix possible et à démontrer (à partir des propriétés des opérations sur les nombres naturels) qu'elle satisfait toutes les propriétés pour faire de  $(\mathbb{Z}, +, 0, \cdot, 1)$  un anneau commutatif.

Nous pourrions faire le même programme pour le corps des rationnels, et montrer que c'est la seule façon de définir l'addition et la multiplication des rationnels, mais nous nous contenterons de donner les définitions et de vérifier qu'il s'agit d'un champ.

### 5.1 Le groupe additif $(\mathbb{Z}, +, 0)$

**Définition 5.1.1.** On note  $\mathcal{R}$  la relation sur  $\mathbb{N} \times \mathbb{N}$  définie par  $(a, b)\mathcal{R}(a', b')$  si, et seulement si,  $a + b' = a' + b$ .

L'idée de cette équivalence est bien entendu que le couple  $(a, b)$  devrait correspondre à la différence  $a - b$ , et que l'équivalence est alors  $a - b = a' - b'$ , (avec les notations de la définition). Mais bien sûr, cette relation n'est pas définie sur  $\mathbb{N} \times \mathbb{N}$ . On l'exprime donc à l'aide d'une relation partout définie, et qui la prolonge sur  $\mathbb{N} \times \mathbb{N}$ . On a le résultat suivant sur cette relation.

**Proposition 5.1.1.** *La relation  $\mathcal{R}$  définie ci-dessus est une relation d'équivalence. De plus, on a  $(a, b)\mathcal{R}(a', b')$  si, et seulement si, il existe  $k \in \mathbb{N}$  tel que*

$$\begin{cases} a' &= a + k \\ b' &= b + k \end{cases} \quad \text{ou} \quad \begin{cases} a &= a' + k \\ b &= b' + k \end{cases}$$

*Démonstration.* On montre les trois conditions définissant une relation d'équivalence :

1. Elle est réflexive : pour tous  $a, b \in \mathbb{N}$ , la condition  $(a, b)\mathcal{R}(a, b)$  s'écrit par définition  $a + b = a + b$  ; c'est vrai.

2. Elle est symétrique : pour tous  $a, b, a', b' \in \mathbb{N}$ , si  $(a, b)\mathcal{R}(a', b')$ , alors on a  $a+b' = a'+b$ . La condition  $(a', b')\mathcal{R}(a, b)$  s'écrit  $a'+b = a+b'$  et on constate que ces deux conditions sont équivalentes, par symétrie de l'égalité.
3. Elle est transitive : supposons que pour  $x, y, z, t, r, s \in \mathbb{N}$ , on a  $(x, y)\mathcal{R}(z, t)$  et  $(z, t)\mathcal{R}(r, s)$ . Cela se traduit par les égalités  $x+t = y+z$  et  $z+s = t+r$ . On ajoute  $s$  aux deux membres de la première égalité et on obtient, en utilisant l'associativité de l'addition dans  $\mathbb{N}$

$$x+t+s = y+z+s = y+t+r.$$

On obtient alors également  $x+s = y+r$ , en simplifiant par  $t$  (voir les propriétés de l'addition dans  $\mathbb{N}$ ).

Nous avons donc montré que  $\mathcal{R}$  est une relation d'équivalence. Passons maintenant à la deuxième partie. Supposons que  $(a, b)\mathcal{R}(a', b')$ . Si  $a \geq a'$ , il existe  $k \in \mathbb{N}$  tel que  $a = a' + k$ . Par définition de  $\mathcal{R}$ , on a aussi  $a+b' = a'+b$ . Ces deux égalités fournissent  $a'+k+b' = a'+b$ . En simplifiant par  $a'$ , on obtient  $b' = b+k$ . On procède de la même façon si  $a \leq a'$ .

Réciproquement, supposons qu'il existe  $k \in \mathbb{N}$  tel que  $a' = a+k$  et  $b' = b+k$ . On a alors  $a+b' = a+b+k$  et  $a'+b = a+k+b$ . Ces deux nombres sont donc égaux et on a  $(a, b)\mathcal{R}(a', b')$ . On traite l'autre cas de la même façon.  $\square$

On peut maintenant passer à la définition de l'ensemble  $\mathbb{Z}$ , et de l'opération d'addition sur  $\mathbb{Z}$ . Cette dernière consiste à additionner les gains d'un côté et les pertes de l'autre, et à passer au quotient.

**Définition 5.1.2.** L'ensemble  $\mathbb{Z}$  est le quotient  $\mathbb{N} \times \mathbb{N} / \mathcal{R}$ . L'addition est l'application

$$+ : \mathbb{Z} \times \mathbb{Z} \rightarrow \mathbb{Z} : ([a, b], [c, d]) \mapsto [a+c, b+d].$$

Une première tâche consiste à montrer que l'addition est bien définie. Il s'agit en effet d'une opération définie sur un quotient, et il faut démontrer que la définition posée est indépendante des représentants choisis pour l'exprimer. C'est l'objet de la proposition suivante.

**Proposition 5.1.2.** Pour tout  $a, b, a', b', c, d, c', d' \in \mathbb{N}$ , si  $(a, b)\mathcal{R}(a', b')$  et  $(c, d)\mathcal{R}(c', d')$ , alors  $(a+c, b+d)\mathcal{R}(a'+c', b'+d')$ .

*Démonstration.* On commence par traduire les conditions de l'énoncé en utilisant la définition de  $\mathcal{R}$ . La première s'écrit  $a+b' = a'+b$ . La deuxième est  $c+d' = c'+d$ . Il faut en déduire  $(a+c)+(b'+d') = (b+d)+(a'+c')$ . La troisième est donc visiblement obtenue en sommant membre à membre les deux premières, et en utilisant les propriétés de l'addition dans  $\mathbb{N}$  pour grouper les termes.  $\square$

Pour exprimer les propriétés de l'addition, nous introduisons la définition de la structure de groupe. Elle aura une portée bien plus large dans votre parcours que la seule structure de  $\mathbb{Z}$ .

**Définition 5.1.3.** Un groupe est un triplet  $(G, \circ, e)$  où  $G$  est un ensemble (non vide),  $e$  un élément de  $G$  et  $\circ$  une application de  $G \times G$  dans  $G$  satisfaisant les propriétés suivantes.

1. L'application  $\circ$  est associative : on a  $a \circ (b \circ c) = (a \circ b) \circ c$  pour tous  $a, b, c \in G$ ;
2. L'élément  $e$  est neutre : on a  $e \circ a = a \circ e = a$  pour tout  $a \in G$ ;
3. Pour tout  $a \in G$ , il existe  $a' \in G$  tel que  $a \circ a' = a' \circ a = e$ .

Voici deux exemples.

**Exemple 5.1.1.** 1. Si  $A$  est un ensemble (non vide), alors l'ensemble des bijections de  $A$  dans  $A$ , muni de la composition des bijections, et avec le neutre  $\text{id}_A$ , est un groupe.

2. L'ensemble  $G = \{0, 1\}$ , muni de l'opération  $+$  définie par  $0+0=0$ ,  $0+1=1+0=1$  et  $1+1=0$  et du neutre  $e=0$  est un groupe.

Voici deux propriétés des groupes, utiles pour fixer le vocabulaire.

**Proposition 5.1.3.** *Dans tout groupe, il n'y a qu'un seul élément neutre. Pour tout  $a \in G$ , il existe un unique  $a' \in G$  tel que  $a \circ a' = a' \circ a = e$ .*

*Démonstration.* On suppose qu'il existe deux éléments neutres  $e$  et  $e'$ . On a alors  $e \circ e' = e'$ , puisque  $e$  est neutre, et  $e \circ e' = e$ , puisque  $e'$  est neutre.

De même, on suppose qu'un élément  $a$  admet deux inverses,  $a'$  et  $a''$ . On calcule alors  $a' \circ a \circ a''$  de deux façons différentes en utilisant l'associativité et on obtient  $a' = a''$ .  $\square$

On dit donc que dans un groupe  $(G, \circ, e)$ , l'élément  $e$  est le neutre de  $G$ . Pour tout  $a$ , l'élément  $a'$  de la définition est appelé l'inverse ou l'opposé de  $a$ , il est noté en général  $a^{-1}$ . Cependant, quand on utilise la notation  $+$  pour l'opération de groupe on dit que  $a'$  est l'opposé de  $a$  et on le note  $-a$ . On a également le résultat suivant sur les inverses.

**Proposition 5.1.4.** *Soit  $(G, \circ, e)$  un groupe. Pour tous  $a, b \in G$ , on a  $(a \circ b)^{-1} = b^{-1} \circ a^{-1}$ .*

*Démonstration.* Il suffit de montrer que  $b^{-1} \circ a^{-1}$  satisfait les conditions pour être l'opposé de  $a \circ b$ . Le résultat suivra vu l'unicité de ce dernier. On a

$$(b^{-1} \circ a^{-1}) \circ (a \circ b) = b^{-1} \circ (a^{-1} \circ a) \circ b = b^{-1} \circ e \circ b = b^{-1} \circ (e \circ b) = b^{-1} \circ b = e,$$

et de la même façon  $(a \circ b) \circ (b^{-1} \circ a^{-1}) = e$ .  $\square$

**Définition 5.1.4.** Un groupe commutatif est un groupe  $(G, \circ, e)$  tel que  $a \circ b = b \circ a$  pour tous  $a, b \in G$ .

Ces structures de groupe que vous allez souvent rencontrer dans la suite permettent d'énoncer facilement les propriétés de l'addition.

**Proposition 5.1.5.** *Le triplet  $(\mathbb{Z}, +, 0)$ , où  $0 = [(0, 0)]$  est un groupe commutatif.*

*Démonstration.* Il s'agit d'une simple vérification des quatre propriétés définissant un groupe commutatif. On se ramène dans chaque cas aux propriétés adéquates de l'addition dans  $\mathbb{N}$ .

Pour l'associativité, on calcule

$$([(a_1, b_1)] + [(a_2, b_2)]) + [(a_3, b_3)] = [(a_1 + a_2, b_1 + b_2)] + [(a_3, b_3)] = [((a_1 + a_2) + a_3, (b_1 + b_2) + b_3)],$$

par définition de l'addition dans  $\mathbb{Z}$ . D'autre part, on a

$$[(a_1, b_1)] + ([[(a_2, b_2)] + [(a_3, b_3)])] = [(a_1, b_1)] + [(a_2 + a_3, b_2 + b_3)] = [(a_1 + (a_2 + a_3), b_1 + (b_2 + b_3))],$$

toujours par définition. On remarque que ces deux nombres sont égaux, vu l'associativité de l'addition dans  $\mathbb{N}$ .

Montrons que  $0 = [(0, 0)]$  est neutre. Si  $x = [(a, b)]$  appartient à  $\mathbb{Z}$ , alors on a

$$0 + x = [(0, 0)] + [(a, b)] = [(0 + a, 0 + b)] = [(a, b)],$$

puisque  $0$  est neutre pour l'addition dans  $\mathbb{N}$ . On montre de même que  $x + 0 = x$ .

En ce qui concerne l'opposé, on vérifie directement<sup>a</sup> que l'opposé de  $[(a, b)]$  est donné par  $[(b, a)]$ . En effet, on a

$$[(a, b)] + [(b, a)] = [(a + b, a + b)] = 0,$$

puisque  $(a + b, a + b) \mathcal{R} (0, 0)$ , par définition.

Pour la commutativité, on se ramène également directement à la commutativité de l'addition dans  $\mathbb{N}$ .  $\square$

a. On peut avoir une intuition de la valeur de l'opposé en se souvenant que  $(a, b)$  représente intuitivement un gain de  $a$  et une perte de  $b$ . L'opposé est alors naturellement un gain de  $b$  et une perte de  $a$ . Mais on peut bien sûr se contenter de résoudre l'équation  $[(a, b)] + [(x, y)] = [(0, 0)]$ .

Pour terminer cette section, démontrons que  $\mathbb{N}$  peut être vu comme un sous-ensemble de  $\mathbb{Z}$ . A proprement parler, c'est impossible : les éléments de  $\mathbb{Z}$  sont des classes d'équivalences de couples de naturels. Il faut donc *identifier*  $\mathbb{N}$  à un sous-ensemble de  $\mathbb{Z}$ . Cela se fait naturellement en définissant une application injective  $\varphi$  de  $\mathbb{N}$  dans  $\mathbb{Z}$ . L'ensemble  $\mathbb{N}$  est alors en bijection avec son image  $\varphi(\mathbb{N})$ . Mais l'identification doit aller plus loin qu'une simple identification d'ensembles : si vous considérez deux nombres naturels, disons 3 et 5, vous pouvez les additionner dans  $\mathbb{N}$ . Mais vous pouvez aussi considérer que ce sont des éléments de  $\mathbb{Z}$  (via  $\varphi$ ), et les additionner. Les résultats devraient "être les mêmes", si ce n'est que le premier est dans  $\mathbb{N}$ , et le second dans  $\mathbb{Z}$ . Ils devraient donc se correspondre via  $\varphi$ .

Nous disposons d'une telle application, comme le montre la proposition suivante.

**Proposition 5.1.6.** *L'application<sup>b</sup>*

$$\varphi : \mathbb{N} \rightarrow \mathbb{Z} : n \mapsto [(n, 0)]$$

*est injective. De plus elle satisfait*

$$\varphi(n + n') = \varphi(n) + \varphi(n') \quad \text{et} \quad \varphi(0) = 0.$$

*Démonstration.* Supposons que  $\varphi(n) = \varphi(n')$ , pour  $n, n' \in \mathbb{N}$ . On a alors  $[(n, 0)] = [(n', 0)]$ , ce qui donne  $(n, 0)\mathcal{R}(n', 0)$ , ou encore  $n + 0 = n' + 0$ . Puisque 0 est neutre dans  $\mathbb{N}$ , on obtient  $n = n'$  et l'application  $\varphi$  est injective. Le deuxième point à démontrer découle directement de la définition de l'addition dans  $\mathbb{Z}$  et de son neutre.  $\square$

Dans la suite, nous identifierons souvent  $\mathbb{N}$  à  $\varphi(\mathbb{N})$ , et nous écrirons  $\mathbb{N} \subset \mathbb{Z}$ . Le nombre  $n \in \mathbb{N}$  est donc identifié à  $[(n, 0)]$ . L'opposé de  $[(n, 0)]$  est  $[(0, n)]$ . On doit le noter  $-[(n, 0)]$ , mais on peut également le noter<sup>c</sup>  $-n$ . On rejoint ainsi la définition que vous avez peut-être eue de  $\mathbb{Z}$  : on ajoute leurs opposés aux nombres naturels. Ce fait est précisé dans le résultat qui suit, où on note  $-\mathbb{N}$  le sous-ensemble de  $\mathbb{Z}$  défini par  $\{-n : n \in \mathbb{N}\}$ .

**Proposition 5.1.7.** *On a  $\mathbb{Z} = \mathbb{N} \cup -\mathbb{N}$  et  $\mathbb{N} \cap -\mathbb{N} = \{0\}$ .*

L'idée intuitive est ici encore que les éléments de  $\mathbb{Z}$  sont des classes de couples. Le couple  $(a, b)$  correspond à un élément de  $\mathbb{N}$  si le gain est supérieur à la perte et à un élément de  $-\mathbb{N}$  dans le cas contraire.

*Démonstration.* Pour montrer l'égalité des ensembles, on montre deux inclusions. On a par définition  $\mathbb{N} \cup -\mathbb{N} \subset \mathbb{Z}$ . Montrons maintenant l'autre inclusion. Soit  $x = [(a, b)] \in \mathbb{Z}$ . Deux cas peuvent se produire.

— Si  $a \geq b$ , alors  $a - b$  est un nombre naturel<sup>d</sup>. On a alors  $x = [(a - b, 0)] \in \mathbb{N}$ .

— Si  $a \leq b$ , alors  $b - a$  est défini dans  $\mathbb{N}$ . On a alors  $x = [(0, b - a)] = -[(b - a, 0)] \in -\mathbb{N}$ .

Passons maintenant à l'intersection. Ici encore, on montre deux inclusions. D'une part, le neutre de  $\mathbb{Z}$  s'écrit  $[(0, 0)] = \varphi(0)$ . Il appartient donc à  $\mathbb{N}$ . D'autre part, c'est également  $-[(0, 0)]$ , et il appartient donc à  $-\mathbb{N}$ . On a donc montré l'inclusion  $\{0\} \subset \mathbb{N} \cap -\mathbb{N}$ . Passons à l'autre inclusion. Soit  $x \in \mathbb{N} \cap -\mathbb{N}$ . Par définition, il existe  $n, m \in \mathbb{N}$  tel que  $x = [(n, 0)]$  et  $x = -[(m, 0)] = [(0, m)]$ . On a donc  $(n, 0)\mathcal{R}(0, m)$ , ce qui donne  $m + n = 0$ , et par suite  $m = n = 0$ . On a donc  $x = 0$ .  $\square$

Le fait que tout nombre admette un opposé permet de définir la soustraction dans  $\mathbb{Z}$ . La définition est la même que dans l'ensemble  $\mathbb{N}$ , mais la soustraction est une application de  $\mathbb{Z} \times \mathbb{Z}$  dans  $\mathbb{Z}$ .

**Proposition 5.1.8.** *Pour tous  $x, y \in \mathbb{Z}$ , il existe un unique  $z \in \mathbb{Z}$  tel que  $x + z = y$ .*

b. Ici encore, l'intuition est claire : on associe au nombre naturel un "gain" de  $n$ , et aucune perte.

c. Voir la notation de l'opposé quand l'opération est l'addition.

d. La soustraction est définie dans  $\mathbb{N}$ .

*Démonstration.* Si  $z$  est tel que  $x + z = y$ , alors on a  $(-x) + (x + z) = (-x) + y$  et donc  $z = (-x) + y$ . Cela prouve l'unicité. Pour l'existence, on voit que la solution  $z$  proposée convient : on a bien  $x + ((-x) + y) = (x + (-x)) + y = 0 + y = y$ .  $\square$

**Définition 5.1.5.** L'unique nombre  $z$  tel que  $x + z = y$  est noté  $y - x$ . C'est la différence entre  $y$  et  $x$ , ou la soustraction de  $x$  à  $y$ .

La proposition précédente permet également de calculer explicitement  $y - x$ . Les propriétés de l'opposé permettent également de démontrer le résultat suivant, que je vous laisse comme exercice.

**Proposition 5.1.9.** On a  $-(x + y) = (-x) + (-y) = -x - y$  pour tous  $x, y \in \mathbb{Z}$ .

Enfin, on peut définir à partir de la soustraction une relation d'ordre sur  $\mathbb{Z}$ , qui prolonge celle que nous avons utilisée sur  $\mathbb{N}$ .

**Définition 5.1.6.** La relation d'ordre usuelle sur  $\mathbb{Z}$  est définie par

$$x \leq y \Leftrightarrow \exists k \in \mathbb{N} : y = x + k.$$

Bien entendu, par définition, on a  $x \leq y \Leftrightarrow y - x \in \mathbb{N}$ , si on voit  $\mathbb{N}$  comme un sous-ensemble de  $\mathbb{Z}$ . Je vous laisse le soin de vérifier que c'est bien une relation d'ordre. On pourra également montrer que cet ordre est compatible avec l'addition, et étudier son comportement vis-à-vis de la multiplication que nous allons définir.

## 5.2 L'anneau $(\mathbb{Z}, +, 0, \cdot, 1)$

Nous avons défini une multiplication des nombres naturels et étudié ses propriétés. Mais vous avez déjà rencontré une multiplication des entiers, qui possède la propriété tant attendue : moins par moins donne plus. Cette propriété semble être un choix étrange pour le grand public.

Nous allons voir qu'il n'y a pas d'autre façon de faire si on souhaite que la multiplication ait des propriétés raisonnables, dont voici la liste.

1. Les nombres naturels se multiplient dans  $\mathbb{Z}$  comme dans  $\mathbb{N}$ . Autrement dit, le plongement  $\varphi$  de  $\mathbb{N}$  dans  $\mathbb{Z}$  doit satisfaire

$$\varphi(n \cdot n') = \varphi(n) \cdot \varphi(n'), \quad n, n' \in \mathbb{N}.$$

2. Il existe un neutre pour la multiplication, c'est-à-dire un élément  $e \in \mathbb{Z}$  tel que  $e \cdot x = x \cdot e = x$  pour tout  $x \in \mathbb{Z}$ .
3. La multiplication distribue l'addition : on a  $x \cdot (y + z) = x \cdot y + x \cdot z$  et  $(y + z) \cdot x = y \cdot x + z \cdot x$  pour tous  $x, y, z \in \mathbb{Z}$ .

Montrons que ces propriétés définissent complètement la multiplication. Nous avons besoin de résultats intermédiaires.

**Proposition 5.2.1.** Si la multiplication distribue l'addition et admet un neutre, alors 0 est absorbant : on a  $0 \cdot x = x \cdot 0 = 0$  pour tout  $x \in \mathbb{Z}$ .

*Démonstration.* On utilise successivement les propriétés adéquates pour obtenir

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

pour tout  $x \in \mathbb{Z}$ . En additionnant  $-x$  aux deux membres de cette égalité, on obtient bien  $0 = 0 \cdot x$ . On procède de façon analogue pour obtenir  $x \cdot 0 = 0$ .  $\square$

Passons à "moins par moins donne plus".

**Proposition 5.2.2.** *Si la multiplication distribue l'addition et admet un neutre, alors*

$$x \cdot (-y) = (-x) \cdot y = -(x \cdot y) \quad \text{et} \quad (-x) \cdot (-y) = x \cdot y$$

pour tous  $x, y \in \mathbb{Z}$ .

*Démonstration.* Pour les premières égalités, il suffit de montrer que  $x \cdot (-y)$  et  $(-x) \cdot y$  satisfont les propriétés définissant l'opposé de  $x \cdot y$ . On calcule donc

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

en utilisant la distributivité et le fait que 0 est absorbant. On montre de même que  $x \cdot y + (-x) \cdot y = 0$ .

Pour la dernière égalité, on utilise par exemple les deux premières et on obtient

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

ce qui achève la preuve. □

Déterminons maintenant le nombre  $e$ .

**Proposition 5.2.3.** *Si les trois conditions énoncées ci-dessus sont satisfaites, alors on a  $e = 1$ .*

*Démonstration.* On sait que si l'unité existe pour la multiplication, elle est unique. Montrons que 1 (plus précisément  $\varphi(1)$ ) est neutre pour la multiplication. Pour tout  $x \in \mathbb{Z}$ , on a  $x \in \mathbb{N}$  ou  $x \in -\mathbb{N}$ ; Si  $x \in \mathbb{N}$ , alors  $x = \varphi(n)$ ,  $n \in \mathbb{N}$ . On a alors

$$\varphi(1) \cdot x = \varphi(1) \cdot \varphi(n) = \varphi(1 \cdot n) = \varphi(n) = x.$$

Si  $x \in -\mathbb{N}$ , alors  $x = -\varphi(n)$  pour un  $n \in \mathbb{N}$ , et on a

$$\varphi(1) \cdot x = \varphi(1) \cdot (-\varphi(n)) = -(\varphi(1) \cdot \varphi(n)) = -\varphi(n) = x.$$

On montre de même que  $\varphi(1)$  est neutre à droite et cela achève la preuve. □

Nous pouvons maintenant passer à la détermination de la multiplication dans  $\mathbb{Z}$ . Nous procédons en deux étapes : en supposant que les trois propriétés données ci-dessus soient vraies, on détermine l'unique multiplication qui peut les satisfaire. Ensuite, ayant une formule explicite pour la multiplication, on démontre toutes les propriétés voulues, par simple vérification.

**Proposition 5.2.4.** *Si une multiplication  $\cdot : \mathbb{Z} \times \mathbb{Z} \rightarrow \mathbb{Z}$  est telle que  $\varphi(n \cdot n') = \varphi(n) \cdot \varphi(n')$  pour tous  $n, n' \in \mathbb{N}$ , si elle admet un neutre et distribue l'addition, alors on a nécessairement*

$$[(a, b)] \cdot [(c, d)] = [(ac + bd, ad + bc)],$$

pour tous  $a, b, c, d$  dans  $\mathbb{N}$ .

Remarquons que cette forme de la multiplication n'est pas étonnante si on pense  $[(a, b)]$  comme la différence  $a - b$ . Il est aussi important de remarquer que l'on ne tourne pas en rond : les produits dans le membre de droite sont des produits dans  $\mathbb{N}$ .

*Démonstration.* Utilisons les propriétés citées dans l'énoncé pour déterminer la multiplication. Par distributivité et définition de la somme, on a

$$\begin{aligned} [(a, b)] \cdot [(c, d)] &= ([(a, 0)] + [(0, b)]) \cdot ([(c, 0)] + [(0, d)]) \\ &= [(a, 0)] \cdot [(c, 0)] + [(a, 0)] \cdot [(0, d)] + [(0, b)] \cdot [(c, 0)] + [(0, b)] \cdot [(0, d)]. \end{aligned}$$

Il reste à calculer les quatre termes séparément. Le premier découle de la propriété demandée pour  $\varphi$ . Il s'écrit en effet  $\varphi(a) \cdot \varphi(c)$ . C'est donc  $\varphi(ac)$ , soit  $[(ac, 0)]$ . Le deuxième est

$$[(a, 0)] \cdot (-(d, 0)) = -([(a, 0)] \cdot [(d, 0)]) = -[(ad, 0)] = [(0, ad)].$$

Le troisième vaut

$$(-(b, 0)) \cdot [(c, 0)] = -([(b, 0)] \cdot [(c, 0)]) = -[(bc, 0)] = [(0, bc)],$$

et enfin le dernier vaut

$$(-(b, 0)) \cdot (-(d, 0)) = [(b, 0)] \cdot [(d, 0)] = [(bd, 0)].$$

Le résultat suit en additionnant les quatre termes obtenus suivant la définition de l'addition dans  $\mathbb{Z}$ .  $\square$

Puisque nous y sommes forcés, nous définissons la multiplication dans  $\mathbb{Z}$  comme suit.

**Définition 5.2.1.** La multiplication des nombres entiers est l'application

$$\cdot : \mathbb{Z} \times \mathbb{Z} \rightarrow \mathbb{Z} : [(a, b)], [(c, d)] \mapsto [(ac + bd, ad + bc)].$$

Nous allons maintenant vérifier toutes les propriétés de cette application, en commençant par la première, qui stipule qu'elle est bien définie.

**Proposition 5.2.5.** *La multiplication donnée par la définition 5.2.1 est indépendante du choix des représentants.*

*Démonstration.* On calcule la multiplication pour d'autres représentants et on montre que l'on obtient le même résultat. Supposons donc  $(a, b)\mathcal{R}(a', b')$  et  $(c, d)\mathcal{R}(c', d')$  et montrons que  $(ac + bd, ad + bc)\mathcal{R}(a'c' + b'd', a'd' + b'c')$ . Cela revient à montrer l'égalité

$$ac + bd + a'd' + b'c' = ad + bc + a'c' + b'd' \quad (5.1)$$

On peut supposer sans perte de généralité (vu la symétrie de  $\mathcal{R}$ ) que  $a' \geq a$  et  $c' \geq c$ . On écrit alors  $a' = a + k$  et  $b' = b + k$ , pour un  $k \in \mathbb{N}$ , on fait de même pour trouver  $c' = c + l$  et  $d' = d + l$ , pour un  $l \in \mathbb{N}$ . On développe la relation à démontrer et on constate qu'elle est vraie.  $\square$

Les propriétés de l'addition et de la multiplication dans  $\mathbb{Z}$  sont générales. Elles sont caractéristiques de la structure d'anneau, que nous définissons maintenant.

**Définition 5.2.2.** Un anneau est une structure  $(A, +, 0, \cdot, 1)$  où  $A$  est un ensemble non vide et où 0 et 1 sont des éléments de  $A$ , satisfaisant les propriétés suivantes :

1.  $(A, +, 0)$  est un groupe commutatif ;
2. La multiplication  $\cdot : A \times A \rightarrow A$  est associative : on a  $(a \cdot b) \cdot c = a \cdot (b \cdot c)$  pour tous  $a, b, c \in A$  ;
3. La multiplication distribue l'addition : on a  $a \cdot (b + c) = a \cdot b + a \cdot c$  et  $(b + c) \cdot a = b \cdot a + c \cdot a$  pour tous  $a, b, c \in A$  ;
4. L'élément 1 est neutre pour la multiplication<sup>a</sup> : on a  $a \cdot 1 = 1 \cdot a = a$  pour tout  $a \in A$ .

Voici encore deux structures qui ont plus de propriétés, on dit qu'elles sont plus riches.

**Définition 5.2.3.** Un anneau  $(A, +, 0, \cdot, 1)$  est dit commutatif si la multiplication est commutative, c'est-à-dire si on a  $a \cdot b = b \cdot a$  pour tous  $a, b \in A$ .

a. Dans la littérature, il existe plusieurs conventions à ce sujet : certains auteurs ne demandent pas qu'un anneau soit muni d'une unité multiplicative. Ils parlent alors d'anneau avec unité si cette propriété est satisfaite.

**Définition 5.2.4.** Un corps est un anneau  $(A, +, 0, \cdot, 1)$  tel que  $A \neq \{0\}$  et tel que pour tout  $a \in A \setminus \{0\}$  il existe  $a' \in A$  tel que  $a \cdot a' = a' \cdot a = 1$ .

**Définition 5.2.5.** Un champ est un corps commutatif.

Remarquons que si dans un anneau  $1 = 0$ , on a alors  $a = 1 \cdot a = 0 \cdot a = 0$ , puisque 0 est absorbant dans tout anneau. Donc on a  $1 = 0$  si, et seulement si  $A = \{0\}$ . Si  $1 \neq 0$ , il est impossible de trouver un élément  $a \in A$  tel que  $a \cdot 0 = 1$ . Cela explique la condition  $a \neq 0$  dans la définition d'un corps.

On peut bien sûr démontrer que dans tout anneau, il n'y a qu'un élément neutre pour la multiplication. On peut également montrer que l'inverse d'un élément pour la multiplication, quand il existe, est unique.

Enfin, la notation  $\frac{a}{b}$  pour noter le produit de  $a$  par l'inverse de  $b$  n'a de sens que si  $a$  et  $b^{-1}$  commutent. C'est toujours le cas dans un champ, mais pas dans un anneau, ni dans un corps en général.

Enfin, on ne note en général pas la multiplication par un  $\cdot$ , sauf si la clarté de l'exposé l'exige.

Nous verrons des exemples de corps et de champs, avec les nombres rationnels et les nombres complexes. Les nombres réels, que je ne définirai pas, forment également un champ. Passons maintenant aux propriétés de la multiplication dans  $\mathbb{Z}$ . Nous verrons aussi des exemples moins classiques, donnés par les anneaux modulaires.

**Proposition 5.2.6.** La structure  $(\mathbb{Z}, +, 0, \cdot, 1)$  est un anneau commutatif.

*Démonstration.* Nous savons déjà que  $(\mathbb{Z}, +, 0)$  est un groupe commutatif. Les propriétés qui restent à démontrer se vérifient par simple calcul. Dans la suite de cette démonstration, les lettres  $a, b, c, d, e, f$  désignent des nombres naturels quelconques.

Montrons que la multiplication est associative. D'une part, on a

$$\begin{aligned} ((a, b) \cdot ((c, d)) \cdot ((e, f)) &= [(ac + bd, ad + bc)] \cdot ((e, f)) \\ &= [((ac + bd)e + (ad + bc)f, (ac + bd)f + (ad + bc)e)]. \end{aligned}$$

D'autre part, on a

$$\begin{aligned} ((a, b) \cdot ((c, d)) \cdot ((e, f))) &= [(a, b) \cdot ((ce + df, cf + de))] \\ &= [((a(ce + df) + b(cf + de), a(cf + de) + b(ce + df))]. \end{aligned}$$

On constate que ces nombres sont égaux en utilisant la distributivité de la multiplication et la commutativité de l'addition dans  $\mathbb{N}$ .

Vérifions que la multiplication distribue l'addition. D'une part, on calcule

$$\begin{aligned} ((a, b) \cdot ((c, d) + ((e, f))) &= [(a, b) \cdot ((c + e, d + f))] \\ &= [(a(c + e) + b(d + f), a(d + f) + b(c + e))]. \end{aligned}$$

D'autre part, on a

$$\begin{aligned} ((a, b) \cdot ((c, d)) + ((a, b) \cdot ((e, f))) &= [(ac + bd, ad + bc)] + [(ae + bf, af + be)] \\ &= [(ac + ae + bd + bf, ad + af + bc + be)]. \end{aligned}$$

On constate encore une fois que ces nombres sont égaux. On procède de manière analogue pour l'autre relation de distributivité, ou on la déduit de la commutativité de la multiplication, que nous allons démontrer.

D'après notre analyse préparatoire, nous savons que le neutre de la multiplication de  $\mathbb{Z}$  ne peut être que le neutre de  $\mathbb{N}$ , vu comme un élément de  $\mathbb{Z}$ . Montrons donc que  $[(1, 0)]$  est neutre pour la multiplication. On calcule

$$[(1, 0)] \cdot [(a, b)] = [(1a + 0b, 1b + 0a)] = [(a, b)] \text{ et } [(a, b)] \cdot [(1, 0)] = [(a1 + b0, a0 + b1)] = [(a, b)],$$

vu les propriétés de 0 et 1 dans  $\mathbb{N}$ .

Enfin, terminons par la commutativité. On a

$$[(a, b)] \cdot [(c, d)] = [(ac + bd, ad + bc)] \quad \text{et} \quad [(c, d)] \cdot [(a, b)] = [(ca + db, cb + da)].$$

On constate encore une fois que ces nombres sont égaux, vu la commutativité de la multiplication et de l'addition dans  $\mathbb{N}$ .  $\square$

Pour compléter notre programme pour la multiplication, il nous faut montrer que le plongement que nous avons défini de  $\mathbb{N}$  dans  $\mathbb{Z}$  se comporte bien vis-à-vis de la multiplication.

**Proposition 5.2.7.** *L'application  $\varphi : \mathbb{N} \rightarrow \mathbb{Z} : n \mapsto [(n, 0)]$  satisfait la condition*

$$\varphi(n \cdot n') = \varphi(n) \cdot \varphi(n'), \quad \forall n, n' \in \mathbb{N}.$$

*Démonstration.* Ici encore, c'est une simple vérification.  $\square$

Pour terminer cette section sur les nombres entiers, maintenant que nous avons défini la multiplication, il est naturel de lui associer comme dans  $\mathbb{N}$  une opération secondaire, à savoir la division. La définition est identique à celle de la division dans  $\mathbb{N}$ .

**Définition 5.2.6.** Soient  $a \in \mathbb{Z}$  et  $b \in \mathbb{Z}_0^b$ . On dit que  $b$  divise  $a$  si il existe  $c \in \mathbb{Z}$  tel que  $a = b \cdot c$ . On écrit alors  $b|a$  et  $c = a : b$  ou  $c = \frac{a}{b}$ .

Bien sûr, pour que la notation  $a : b$  ait un sens, il faut que le nombre  $c$  intervenant dans la définition soit unique. Cela découle de l'intégrité de l'anneau  $\mathbb{Z}$ .

**Définition 5.2.7.** Un anneau  $(A, +, 0, \cdot, 1)$  est intègre si pour tous  $x, y \in A$ , si  $x \cdot y = 0$  alors  $x = 0$  ou  $y = 0$ .

A titre d'exemple, l'anneau des matrices carrées à éléments réels, que vous verrez d'ici peu en algèbre, n'est pas intègre : il existe des matrices non nulles dont le produit est nul. Nous verrons dans le chapitre suivant des exemples d'anneau non intègres.

**Proposition 5.2.8.** *L'anneau  $(\mathbb{Z}, +, 0, \cdot, 1)$  est intègre.*

*Démonstration.* Supposons que  $x, y \in \mathbb{Z}$  satisfont  $x \cdot y = 0$ . Par définition  $x$  s'écrit  $[(a, b)]$  et  $y$  s'écrit  $[(c, d)]$ , pour des nombres naturels  $a, b, c, d$ . On a alors

$$[(0, 0)] = [(a, b)] \cdot [(c, d)] = [(ac + bd, ad + bc)],$$

c'est-à-dire

$$ac + bd = ad + bc. \quad (5.2)$$

Supposons maintenant  $x \neq 0$ , c'est-à-dire  $a \neq b$ , et montrons que  $y = 0$ , i.e.  $c = d$ . Puisque  $a \neq b$ , on a  $a > b$  ou  $b > a$ . Traitons par exemple le deuxième cas, le premier étant analogue. Dans ce cas, il existe  $k \in \mathbb{N}_0$  tel que  $b = a + k$ . On utilise cette condition dans l'équation (5.2) et on obtient

$$ac + ad + kd = ad + ac + kc.$$

En simplifiant, on a  $kd = kc$ . Mais dans  $\mathbb{N}$ , nous avons vu que cela implique  $d = c$ , puisque  $k \neq 0$ .  $\square$

**Corollaire 5.2.1.** *Soient  $a \in \mathbb{Z}$ ,  $b \in \mathbb{Z}_0$ . Si  $a = b \cdot c$  et  $a = b \cdot c'$ , pour  $c, c' \in \mathbb{Z}$ , alors  $c = c'$ .*

---

b. Cela pourrait paraître arbitraire d'exclure 0, mais comme dans  $\mathbb{N}$ , on voit facilement que 0 ne peut diviser que 0, et il ne le fait pas de manière unique.

*Démonstration.* Les conditions de l'énoncé impliquent  $b \cdot c = b \cdot c'$ , ou encore  $b \cdot c - b \cdot c' = 0$ . Cette condition est équivalente à  $b \cdot (c - c') = 0$ , et puisque  $b \neq 0$ , elle implique  $c = c'$ .  $\square$

Enfin, voici quelques propriétés bien connues de la division, qui découlent directement de la définition et des propriétés correspondantes de la multiplication.

**Proposition 5.2.9.** *Soient  $a \in \mathbb{Z}$ ,  $b \in \mathbb{Z}_0$ , tels que  $b|a$ . Alors  $-b|a$ ,  $b|-a$  et  $-b|-a$ . On a de plus  $\frac{a}{-b} = \frac{-a}{b} = -\frac{a}{b}$  et  $\frac{-a}{-b} = \frac{a}{b}$ .*

### 5.3 Le champ $(\mathbb{Q}, +, 0, \cdot, 1)$

Nous avons défini dans les sections précédentes l'anneau des nombres entiers  $\mathbb{Z}$ . Il n'est pas difficile de voir que ce n'est pas un champ. Nous allons plonger  $\mathbb{Z}$  dans un champ, à savoir le champ des rationnels. L'idée consiste à ajouter les inverses des nombres entiers non nuls. Vous connaissez la construction, et il s'agit donc ici de reconstruire avec précision l'ensemble des rationnels : on considère d'abord les fractions, qui sont des objets du type  $\frac{a}{b}$ , c'est-à-dire des couples de nombres entiers, dont le second est non nul. Ensuite, on définit une équivalence dans l'ensemble des fractions. Le corps des rationnels sera alors le quotient de l'ensemble des couples par la relation d'équivalence.

Les opérations seront naturellement induites par celles des fractions. Ici encore, on pourrait démontrer que c'est dans une certaine mesure l'unique choix raisonnable pour définir un champ contenant les nombres entiers.

**Définition 5.3.1.** On note  $\mathcal{E}$  la relation d'équivalence définie sur  $\mathbb{Z} \times \mathbb{Z}_0$  par  $(a, b)\mathcal{E}(c, d)$  si, et seulement si  $ad = bc$ . L'ensemble  $\mathbb{Q}$  est le quotient  $(\mathbb{Z} \times \mathbb{Z}_0)/\mathcal{E}$ .

Pour que la définition soit licite, il faut bien sûr prouver ce que l'on a avancé.

**Proposition 5.3.1.** *La relation  $\mathcal{E}$  est une relation d'équivalence. De plus, on a  $(a, b)\mathcal{E}(a', b')$  si, et seulement si, il existe  $r, r' \in \mathbb{Z}_0$  tels que  $ar = a'r'$  et  $br = b'r'$ . En particulier  $(a, b)\mathcal{E}(r \cdot a, r \cdot b)$ , pour tout  $r \in \mathbb{Z}_0$ .*

*Démonstration.* On montre les trois conditions définissant les relations d'équivalence :

1. La relation est réflexive : on a  $(a, b)\mathcal{E}(a, b)$  car  $ab = ba$ , pour tout  $(a, b) \in \mathbb{Z} \times \mathbb{Z}_0$  ;
2. La relation est symétrique : soient  $(a, b)$  et  $(c, d)$  dans  $\mathbb{Z} \times \mathbb{Z}_0$  tels que  $(a, b)\mathcal{E}(c, d)$ . On a alors  $ad = bc$  et donc aussi  $(c, d)\mathcal{E}(a, b)$  car cette condition est équivalente à  $cb = da$  ;
3. La relation est transitive : soient  $(a, b)$ ,  $(c, d)$  et  $(e, f)$  dans  $\mathbb{Z} \times \mathbb{Z}_0$  tels que  $(a, b)\mathcal{E}(c, d)$  et  $(c, d)\mathcal{E}(e, f)$ . On a alors  $ad = bc$  et  $cf = de$ . En multipliant les deux membres de la première égalité par  $f$  et les deux membres de la seconde par  $b$ , on obtient  $adf = bcf$  et  $bcf = bde$ . On en déduit  $adf = bde$ , et puisque  $d \neq 0$ , on a  $af = be$ , donc finalement  $(a, b)\mathcal{E}(e, f)$ , ce qu'il fallait démontrer.

Pour le dernier point, on note que si il existe  $r, r'$  non nuls satisfaisant les conditions de l'énoncé, alors on a  $ab'rr' = a'r'br$ , donc  $ab' = a'b$ , puisque  $\mathbb{Z}$  est intègre. Réciproquement, si  $(a, b)\mathcal{E}(a', b')$ , alors  $r = b'$  et  $r' = b$  satisfont les conditions de l'énoncé. Le cas particulier est évident.  $\square$

Les opérations sur les fractions sont bien connues. Exprimons-les sur les classes d'équivalence. Nous noterons encore  $[(a, b)]$  la classe d'équivalence du couple  $(a, b) \in \mathbb{Z} \times \mathbb{Z}_0$ . Il faut garder à l'esprit que c'est une classe d'équivalence de la relation  $\mathcal{E}$  qui définit  $\mathbb{Q}$ .

**Définition 5.3.2.** L'addition de  $\mathbb{Q}$  est l'application définie par

$$+ : \mathbb{Q} \times \mathbb{Q} : ([a, b], [c, d]) \mapsto [(a, b)] + [(c, d)] = [(ad + bc, bd)].$$

La multiplication de  $\mathbb{Q}$  est l'application définie par

$$\cdot : \mathbb{Q} \times \mathbb{Q} : ([a, b], [c, d]) \mapsto [(a, b)] \cdot [(c, d)] = [(ac, bd)].$$

Pour compléter la définition, il faut montrer que les applications sont bien définies, c'est-à-dire que le résultat est bien une classe d'équivalence d'éléments de  $\mathbb{Z} \times \mathbb{Z}_0$  et que le résultat est indépendant du représentant de la classe choisi pour le calculer.

**Proposition 5.3.2.** *Les applications données dans la définition 5.3.2 sont bien définies.*

*Démonstration.* Tout d'abord, dans les deux cas, le résultat est une classe de couples de  $\mathbb{Z} \times \mathbb{Z}_0$ . En effet, puisque  $b$  et  $d$  sont non nuls, c'est aussi le cas de  $bd$ .

Dans les deux cas, on choisit des représentants différents et on montre que le résultat est le même. Supposons donc que  $(a, b)\mathcal{E}(a', b')$  et  $(c, d)\mathcal{E}(c', d')$ . On a donc  $ab' = ba'$  et  $cd' = dc'$ .

Dans le cas de l'addition, il faut montrer que l'on a

$$(ad + bc, bd)\mathcal{E}(a'd' + b'c', b'd'),$$

c'est-à-dire

$$(ad + bc)b'd' = bd(a'd' + b'c').$$

Travaillons sur le membre de gauche. En développant, on voit qu'il vaut  $ab'dd' + bb'cd'$ . Vu les égalités ci-dessus, il vaut aussi  $ba'dd' + bb'dc'$ . C'est exactement le développement du membre de droite.

Dans le cas de la multiplication, il faut montrer

$$[(ac, bd)] = [(a'c', b'd')],$$

c'est-à-dire

$$acb'd' = bda'c'.$$

Le membre de gauche vaut  $(ab')(cd')$ , c'est-à-dire  $(a'b)(c'd)$ , vu les égalités ci-dessus. C'est le membre de droite.  $\square$

**Proposition 5.3.3.** *L'ensemble  $\mathbb{Q}$ , muni des opérations données par la définition 5.3.2, est un champ. Le neutre pour l'addition est  $0 = [(0, 1)]$  et celui de la multiplication est  $1 = [(1, 1)]$ . Si  $(a, b) \in \mathbb{Z} \times \mathbb{Z}_0$  est tel que  $[(a, b)] \neq 0$ , alors l'inverse de  $[(a, b)]$  est  $[(b, a)]$ .*

*Démonstration.* Il s'agit de simples vérifications. Dans cette preuve, les couples  $(a, b)$ ,  $(c, d)$  et  $(e, f)$  sont des éléments quelconques de  $\mathbb{Z} \times \mathbb{Z}_0$ . Montrons d'abord que  $(\mathbb{Q}, +, 0)$  est un groupe commutatif.

1. L'addition est associative : d'une part, on a

$$\begin{aligned} [(a, b)] + [(c, d)] + [(e, f)] &= [(ad + bc, bd)] + [(e, f)] \\ &= [((ad + bc)f + bde, bdf)]. \end{aligned}$$

D'autre part, on a

$$\begin{aligned} [(a, b)] + [(c, d)] + [(e, f)] &= [(a, b)] + [(cf + de, df)] \\ &= [(adf + b(cf + de), bdf)]. \end{aligned}$$

On constate que ces nombres sont égaux.

2. L'élément  $0 = [(0, 1)]$  est neutre pour l'addition : on a

$$[(0, 1)] + [(a, b)] = [(0b + 1a, 1b)] = [(a, b)] \text{ et } [(a, b)] + [(0, 1)] = [(a1 + b0, b1)] = [(a, b)].$$

3. Tout élément admet un opposé. Le nombre  $[(-a, b)]$  est en effet opposé à  $[(a, b)]$  : on a en effet

$$[(a, b)] + [(-a, b)] = [(ab + b(-a), b^2)] = [(0, b)] = [(0, 1)] = 0.$$

On montre de la même façon que  $[(-a, b)] + [(a, b)] = 0$ .

4. L'addition est commutative : on calcule en effet

$$[(a, b)] + [(c, d)] = [(ad + bc, bd)] \quad \text{et} \quad [(c, d)] + [(a, b)] = [(cb + da, db)].$$

Ces deux nombres sont égaux vu la commutativité de l'addition et de la multiplication dans  $\mathbb{Z}$ .

Passons maintenant aux propriétés de la multiplication.

1. La multiplication est associative : on calcule

$$\begin{aligned} [(a, b)] \cdot [(c, d)] \cdot [(e, f)] &= [(ac, bd)] \cdot [(e, f)] \\ &= [((ac)e, (bd)f)]. \end{aligned}$$

D'autre part, on a

$$\begin{aligned} [(a, b)] \cdot [(c, d)] \cdot [(e, f)] &= [(a, b)] \cdot [(ce, df)] \\ &= [(a(ce), b(df))]. \end{aligned}$$

On constate que ces nombres sont égaux, vu l'associativité de la multiplication dans l'anneau  $\mathbb{Z}$ .

2. L'élément  $1 = [(1, 1)]$  est neutre pour la multiplication. On a en effet

$$[(a, b)] \cdot [(1, 1)] = [(a1, b1)] = [(a, b)],$$

et on montre l'égalité  $[(1, 1)] \cdot [(a, b)] = [(a, b)]$  de la même façon.

3. La multiplication est commutative : on a en effet

$$[(a, b)] \cdot [(c, d)] = [(ac, bd)], \quad \text{et} \quad [(c, d)] \cdot [(a, b)] = [(ca, db)],$$

et ces deux nombres sont égaux, vu la commutativité de la multiplication dans  $\mathbb{Z}$ .

Montrons maintenant que la multiplication distribue l'addition. D'une part, on calcule

$$\begin{aligned} [(a, b)] \cdot [(c, d)] + [(a, b)] \cdot [(e, f)] &= [(a, b)] \cdot [(cf + de, df)] \\ &= [(a(cf + de), bdf)]. \end{aligned}$$

D'autre part, on a en appliquant les définitions puis en simplifiant (par équivalences)

$$\begin{aligned} [(a, b)] \cdot [(c, d)] + [(a, b)] \cdot [(e, f)] &= [(ac, bd)] + [(ae, bf)] \\ &= [(acf + bdae, b^2df)] = [(acf + dae, bdf)]. \end{aligned}$$

On constate encore une fois que ces nombres sont égaux. On procède de manière analogue pour l'autre relation de distributivité, ou on la déduit de la commutativité.

Enfin, montrons que tout nombre non nul admet un inverse. Si  $[(a, b)] \in \mathbb{Q} \setminus \{0\}$ , alors  $a \neq 0$  et  $[(b, a)]$  est un inverse de  $[(a, b)]$ , puisqu'on a  $[(b, a)] \cdot [(a, b)] = [(ba, ab)] = [(1, 1)] = 1$  et de même  $[(a, b)] \cdot [(b, a)] = 1$ , puisque la multiplication est commutative.  $\square$

Il est commun d'affirmer que  $\mathbb{Z}$  est un sous-ensemble de  $\mathbb{Q}$ . Visiblement, la définition de  $\mathbb{Q}$  comme un quotient de  $\mathbb{Z} \times \mathbb{Z}_0$  ne le permet pas. On doit donc passer par un plongement.

**Proposition 5.3.4.** *L'application  $\psi : \mathbb{Z} \rightarrow \mathbb{Q} : x \mapsto [(x, 1)]$  est injective. De plus elle satisfait*

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

pour tous  $x, y \in \mathbb{Z}$ . On a de plus  $\psi(0) = 0$  et  $\psi(1) = 1$ .

*Démonstration.* Supposons que pour  $x, y \in \mathbb{Z}$ , on a  $\psi(x) = \psi(y)$ . On a alors  $(x, 1)\mathcal{E}(y, 1)$ , ce qui donne  $x \cdot 1 = 1 \cdot y$ , et donc  $x = y$ .

On calcule ensuite  $\psi(x + y) = [(x + y, 1)]$ , par définition de  $\psi$  et  $\psi(x) + \psi(y) = [(x, 1)] + [(y, 1)] = [(x + y, 1)]$ , par définition de l'addition dans  $\mathbb{Q}$ . La première égalité est donc prouvée. On procède de la même façon pour la deuxième. On a également  $\psi(0) = [(0, 1)]$ , et c'est le neutre de l'addition dans  $\mathbb{Q}$ . De même, on a  $\psi(1) = [(1, 1)]$ .  $\square$

## Chapitre 6

# Arithmétique modulaire

Ce chapitre est l'occasion de rencontrer de nouveaux exemples d'anneaux et de corps, qui ne contiennent qu'un nombre fini d'éléments. Ce sont les anneaux modulaires. Ce sera aussi l'occasion de voir quelques théorèmes d'arithmétique, incluant une étude élémentaire du pgcd, le théorème de Bezout, et le célèbre algorithme d'Euclide, que vous reverrez dans d'autres circonstances.

Ce chapitre est largement inspiré, en ce qui concerne l'ordre des définitions et des propositions, du passage correspondant du cours d'algèbre du professeur M. Rigo. Ce cours est disponible sur son site web [www.discmath.ulg.ac.be](http://www.discmath.ulg.ac.be).

### 6.1 Définitions et exemples

Les exemples d'arithmétique modulaire sont présents dans la vie de tous les jours, et nous les utilisons sans nous en rendre compte. En voici quelques-uns.

Le plus classique est sans doute la lecture de l'heure sur les horloges rondes marquant 12 heures. Nous admettons sans difficulté qu'elles ne marquent que douze heures et que le chiffre trois représente aussi bien trois heures que quinze heures. Si l'horloge marque 7 heures, et qu'on se demande combien elle marquera dans 8 heures, le calcul est simple : on additionne les heures pour obtenir 15, puis on retranche 12, pour obtenir 3. Si, à partir de minuit, on compte 5 périodes de 8 heures et qu'on se demande quelle heure sera indiquée après ce laps de temps, on multiplie, pour obtenir  $5 \times 8 = 40$ , c'est à dire 3 tours d'horloge et 4 heures : l'horloge marquera quatre heures.

Les interrupteurs (à poussoir) fournissent également un bon exemple. Dans un cas simple, on peut imaginer qu'une impulsion change l'état de la lampe : elle passe d'éteint à allumé, et d'allumé à éteint. Si on débute avec une lampe éteinte, quel est le résultat si on pousse 17 fois sur l'interrupteur ? La réponse est simple : puisque deux pressions sur l'interrupteur ramènent la lampe à l'état initial, 16 impulsions également, donc 17 impulsions ont le même effet qu'une seule. La lampe est donc allumée.

On peut bien sûr imaginer des exemples plus compliqués <sup>a</sup> : un interrupteur à poussoir pour une lampe à intensité variable, qui aurait cinq états possibles. On aurait les intensités 0 : éteint, 1 : faible, 2 : moyenne, 3 : fort, 4 : insoutenable. Bien entendu, sous un éclairage avec une intensité insoutenable, une pression sur l'interrupteur éteindrait la lampe. On peut se poser le même type de question : si je débute avec une lampe éteinte et que je pousse 17 fois sur l'interrupteur, comme 5 impulsions correspondent à ne rien faire, cela revient à pousser deux fois. On aura donc une intensité moyenne. Si quand la lampe est éteinte, je pousse 4 fois et que mon frère (ou mon voisin Raoul) pousse 4 fois, cela revient à pousser 8 fois, c'est à dire 3 fois, etc... Deux nombres d'impulsions sont donc équivalents s'ils diffèrent par un multiple de 5 impulsions.

Il ressort de ces exemples une structure commune qui est définie comme un quotient. Nous avons utilisé jusqu'à présent des quotients de l'ensemble des nombres naturels, et

---

a. Et un peu plus farfelus.

ce serait suffisant, mais il est plus naturel d'un point de vue algébrique de définir des quotients de l'anneau  $\mathbb{Z}$ .

**Définition 6.1.1.** Soit  $m$  un entier supérieur ou égal à 2. On définit la relation d'égalité modulo  $m$  dans  $\mathbb{Z}$  par  $x \equiv_m y$ , si et seulement si, il existe  $k \in \mathbb{Z}$  tel que  $y = x + km$ . On dit alors que  $y$  est égal (ou congru) à  $x$  modulo  $m$  et on note aussi  $x = y \pmod{m}$ .<sup>b</sup>

Evidemment, on ne peut échapper à la proposition qui rend cette définition utile pour définir des quotients.

**Proposition 6.1.1.** Pour tout entier  $m$  supérieur ou égal à 2, la relation d'égalité modulo  $m$  est une relation d'équivalence.

*Démonstration.* La relation est réflexive : pour tout  $x \in \mathbb{Z}$ , on a  $x = x + 0m$  ;

La relation est symétrique : si  $y = x + km$ , alors  $x = y + (-k)m$  ;

Enfin, elle est transitive : si  $x \equiv_m y$  et  $y \equiv_m z$ , alors il existe  $k, k' \in \mathbb{Z}$  tels que  $y = x + km$  et  $z = y + k'm$ . On a alors  $z = x + (k + k')m$ , et donc  $x \equiv_m z$ .  $\square$

Remarquons que la classe de 0 pour  $\equiv_m$  est l'ensemble des multiples de  $m$  :

$$m\mathbb{Z} = \{mz : z \in \mathbb{Z}\}.$$

On peut exprimer la relation  $\equiv_m$  à partir de cet ensemble : on a  $x \equiv_m y$  si, et seulement si,  $y - x \in m\mathbb{Z}$ .

**Définition 6.1.2.** On appelle  $\mathbb{Z}_m$  le quotient  $\mathbb{Z}/\equiv_m = \mathbb{Z}/m\mathbb{Z}$ .

L'ensemble  $\mathbb{Z}_m$  est fini et de cardinal  $m$ . Pour le voir, il faut étendre la division euclidienne à  $\mathbb{Z}$ . On rappelle que la valeur absolue d'un nombre entier est définie par

$$|z| = \begin{cases} z & \text{si } z \geq 0 \\ -z & \text{si } z \leq 0 \end{cases}$$

On peut alors prolonger la division euclidienne définie dans  $\mathbb{N}$  à  $\mathbb{Z}$ .

**Proposition 6.1.2.** Pour tout  $n \in \mathbb{Z}$  et  $d \in \mathbb{Z}_0$ , il existe un unique couple  $(q, r) \in \mathbb{Z}^2$  satisfaisant les conditions

1.  $n = qd + r$  ;
2.  $0 \leq r < |d|$ .

Bien sûr, comme dans  $\mathbb{N}$ , les nombres  $q$  et  $r$  sont appelés quotient et reste de la division euclidienne de  $n$  par  $d$ .

*Démonstration.* Prouvons d'abord l'existence, en nous ramenant à l'existence du quotient et du reste dans  $\mathbb{N}$ .

Si  $n$  et  $d$  sont positifs ou nuls, ils appartiennent à  $\mathbb{N}$ , et on est ramené directement à l'existence dans  $\mathbb{N}$ .

Si  $n < 0$  et  $d > 0$ , alors il existe  $q', r' \in \mathbb{N}$  tels que  $-n = q'd + r'$  et  $0 \leq r' < d$ . On a alors  $n = (-q')d + (-r')$ , où  $-d < -r' \leq 0$ . Si  $r' = 0$ , alors  $-r' < d$ , et  $(q, r) = (-q', -r')$  convient. Si  $-r' < 0$ , alors on peut écrire  $n = (-q' - 1)d + (-r' + d)$ , et on a  $0 \leq -r' + d < d$ . Le couple  $(q, r) = (-q' - 1, -r' + d)$  convient alors.

Nous avons donc démontré l'existence dans le cas où  $d$  est strictement positif. S'il est strictement négatif, on peut écrire  $n = q'(-d) + r' = (-q')d + r'$ , où  $0 \leq r' < -d$ , par le cas précédent, et le couple  $(q, r) = (-q', r')$  convient.

b. Pour  $m = 0$ , la relation que d'égalité modulo  $m$  est simplement l'égalité, tandis que pour  $m = 1$ , tous les nombres sont égaux modulo  $m$ . Enfin, l'égalité modulo  $-m$  est la même relation que l'égalité modulo  $m$ , il n'y a donc pas de restriction à considérer  $m \geq 2$ .

Passons maintenant à l'unicité. Supposons qu'il existe deux couples  $(q_1, r_1)$  et  $(q_2, r_2)$  satisfaisant les conditions de l'énoncé. On a donc  $n = q_1d + r_1 = q_2d + r_2$ ,  $0 \leq r_1 < |d|$ , et  $0 \leq r_2 < |d|$ . Montrons tout d'abord que  $r_1 = r_2$ . Si tel n'est pas le cas, on peut supposer sans perte de généralité que  $r_1 > r_2$ . On a alors  $r_1 - r_2 \in \mathbb{N}$  et  $0 < r_1 - r_2 \leq r_1 < |d|$ . De plus, on obtient directement  $r_1 - r_2 = (q_2 - q_1)d$ . Mais le membre de droite de cette égalité est soit nul (si  $q_2 = q_1$ ), soit négatif, soit supérieur à  $|d|$ . C'est donc absurde.

Sachant que  $r_1 = r_2$ , on obtient l'égalité  $q_1d = q_2d$ , ou encore  $(q_1 - q_2)d = 0$ . Puisque  $d$  n'est pas nul et puisque  $\mathbb{Z}$  est intègre, on obtient  $q_1 = q_2$ .  $\square$

La division permet de prouver ce qui est évident sur les exemples, à savoir que les ensembles  $\mathbb{Z}_m$  sont finis.

**Proposition 6.1.3.** *Pour tout entier  $m \geq 2$ , l'application  $f$  qui à  $[x]$ , ( $x \in \mathbb{Z}$ ) associe le reste de la division de  $x$  par  $m$  est une bijection entre  $\mathbb{Z}_m$  et  $\{0, \dots, m-1\}$ . En particulier, le cardinal de  $\mathbb{Z}_m$  est  $m$ .*

*Démonstration.* Il faut d'abord démontrer que l'application est bien définie. Si la division euclidienne de  $x \in \mathbb{Z}$  s'écrit  $x = km + r$ , alors on a  $f([x]) = r$ . Si  $x \equiv_m y$ , alors il existe  $k' \in \mathbb{Z}$  tel que  $y - x = k'm$ . On a donc  $y = (k + k')m + r$ , et donc  $f([y]) = r$ . L'image  $r$  est donc indépendante du représentant.

Montrons maintenant que l'application  $f$  est une bijection. Elle est injective. En effet, si  $f([x]) = f([y]) = r$ , pour  $x, y \in \mathbb{Z}$ , alors par définition de  $f$ , il existe  $k, k' \in \mathbb{Z}$  tels que  $x = km + r$  et  $y = k'm + r$ . On a donc  $y - x = (k' - k)m$ , qui donne  $y \equiv_m x$ , ou encore  $[x] = [y]$ .

L'application  $f$  est surjective : pour tout  $n \in \{0, \dots, m-1\}$ , la division euclidienne de  $n$  par  $m$  s'écrit  $n = 0m + n$ . On a donc  $n = f([n])$ .  $\square$

Comme dans les exemples, on peut munir  $\mathbb{Z}_m$  d'une addition, induite par celle de  $\mathbb{Z}$ .

**Définition 6.1.3.** L'addition de  $\mathbb{Z}_m$  est l'application

$$+ : \mathbb{Z}_m \times \mathbb{Z}_m \rightarrow \mathbb{Z}_m : ([x], [y]) \mapsto [x] + [y] = [x + y].$$

Bien sûr, il s'agit d'une application définie sur un quotient. Il faut dès lors montrer, pour que la définition soit légitime, que le résultat de l'addition est indépendant du choix des représentants utilisés pour la définir.

**Proposition 6.1.4.** *L'addition dans  $\mathbb{Z}_m$  est bien définie. De plus, elle munit  $\mathbb{Z}_m$  d'une structure de groupe commutatif.*

*Démonstration.* Il s'agit ici encore de simples vérifications. Montrons tout d'abord que l'application proposée est indépendante du choix des représentants. Soient  $x, y, x', y' \in \mathbb{Z}$  tels que  $x \equiv_m x'$  et  $y \equiv_m y'$ , et montrons que  $[x + y] = [x' + y']$ . Par définition, il existe  $k, k' \in \mathbb{Z}$  tels que  $x = x' + km$  et  $y = y' + k'm$ . On a alors  $x + y = x' + km + y' + k'm = (x' + y') + (k + k')m$ . On a donc  $x + y \equiv_m x' + y'$ , ou encore  $[x + y] = [x' + y']$ .

Montrons maintenant que l'addition définit une structure de groupe commutatif. Pour chaque propriété à démontrer, on déduit le résultat de la propriété correspondante dans  $\mathbb{Z}$ .

1. L'addition est associative : pour tous  $x, y, z \in \mathbb{Z}$ , on a

$$[x] + ([y] + [z]) = [x] + [y + z] = [x + (y + z)]$$

et

$$([x] + [y]) + [z] = [x + y] + [z] = [(x + y) + z].$$

Ces deux expressions sont égales vu l'associativité de l'addition dans  $\mathbb{Z}$ .

2. L'addition admet un neutre  $[0]$  : on a pour tout  $x \in \mathbb{Z}$

$$[x] + [0] = [x + 0] = [x] \quad \text{et} \quad [0] + [x] = [0 + x] = [x],$$

puisque 0 est neutre pour l'addition dans  $\mathbb{Z}$ .

3. Pour tout  $x \in \mathbb{Z}$ ,  $[-x]$  est l'opposé de  $[x]$  : on a en effet

$$[x] + [-x] = [x + (-x)] = [0] = 0, \quad \text{et} \quad [-x] + [x] = [(-x) + x] = [0] = 0,$$

puisque  $x$  admet pour opposé  $-x$  dans  $\mathbb{Z}^c$ .

4. L'addition est commutative : pour tous  $x, y \in \mathbb{Z}$ , on a

$$[x] + [y] = [x + y] = [y + x] = [y] + [x],$$

puisque l'addition dans  $\mathbb{Z}$  est commutative.

□

Passons maintenant à la multiplication.

**Définition 6.1.4.** La multiplication de  $\mathbb{Z}_m$  est l'application

$$\cdot : \mathbb{Z}_m \times \mathbb{Z}_m \rightarrow \mathbb{Z}_m : ([x], [y]) \mapsto [x] \cdot [y] = [x \cdot y].$$

On a bien sûr un résultat analogue à celui concernant la somme.

**Proposition 6.1.5.** La multiplication dans  $\mathbb{Z}_m$  est bien définie. De plus, elle permet de munir  $\mathbb{Z}_m$  d'une structure d'anneau commutatif.

*Démonstration.* Il s'agit de vérifications élémentaires, analogues à celles qui ont été menées concernant la somme dans  $\mathbb{Z}_m$ . Montrons juste que la multiplication est bien définie. Si  $[x'] = [x]$  et  $[y'] = [y]$ , alors il existe  $k, k' \in \mathbb{Z}$  tels que  $x' = x + km$  et  $y' = y + k'm$ . On a alors  $x'y' = (x + km)(y + k'm) = xy + kym + xk'm + kmk'm$ , donc  $[xy] = [x'y']$ . □

Voici un exemple élémentaire de l'utilité de ces calculs modulaires.

**Exemple 6.1.1.** Dans le système de numération décimal usuel, un nombre est divisible par 3 si, et seulement si, la somme de ses chiffres est divisible par 3. Pour le voir, on considère l'équivalence modulo 3. Un nombre  $n$  est divisible par 3 si, et seulement si  $[n] = [0]$  dans  $\mathbb{Z}_3$ . Prenons un exemple concret : 2148 est divisible par 3 si, et seulement si  $[2148] = [0]$ . Mais en utilisant la somme et la multiplication dans  $\mathbb{Z}_3$ , on calcule

$$[2148] = [2 \cdot 10^3 + 1 \cdot 10^2 + 4 \cdot 10 + 8] = [2] \cdot [10]^3 + [1] \cdot [10]^2 + [4][10] + [8],$$

et puisque visiblement  $[10] = [1]$ , on a

$$[2148] = [2] + [1] = [4] + [8] = [2 + 4 + 1 + 8].$$

La classe du nombre 2148 modulo 3 est donc égale à la classe modulo 3 de la somme de ses chiffres (dans le système décimal). Cette constatation peut être étendue à tout nombre  $n$ . Un nombre s'écrit  $n = a_l \dots a_0$  si, et seulement si,  $n = \sum_{k=0}^l a_k 10^k$ . On a alors

$$[n] = \sum_{k=0}^l [a_k][10]^k = \sum_{k=0}^l [a_k] = [\sum_{k=0}^l a_k],$$

et on conclut de la même façon.

Le même procédé permet de calculer le reste de la division par 3 de n'importe quel nombre  $n$ . C'est le même que le reste de la division par 3 de la somme des chiffres de  $n$ .

Bien entendu, les mêmes considérations peuvent s'appliquer pour d'autres diviseurs (vous pouvez par exemple prouver facilement les critères de divisibilité par 9, par 2, par 4, ou même par 11, etc...) mais aussi pour d'autres systèmes de numération (dans d'autres bases entières par exemple).

---

c. C'est ici qu'il est plus naturel d'avoir un quotient de  $\mathbb{Z}$  plutôt qu'un quotient de  $\mathbb{N}$  pour définir un groupe. Cela dit, pour tout  $x \in \{0, \dots, m-1\}$ ,  $[m-x]$  est également l'opposé de  $[x]$ , et on est resté dans  $\mathbb{N}$ .

## 6.2 Les champs finis $\mathbb{Z}_p$

Considérons les tables d'addition des anneaux  $\mathbb{Z}_3$  et  $\mathbb{Z}_4$ . Pour alléger les notations, on ne note pas les crochets, dès que cela ne risque pas de mener à des confusions. On voit que

| + | 0 | 1 | 2 |
|---|---|---|---|
| 0 | 0 | 1 | 2 |
| 1 | 1 | 2 | 0 |
| 2 | 2 | 0 | 1 |

| + | 0 | 1 | 2 | 3 |
|---|---|---|---|---|
| 0 | 0 | 1 | 2 | 3 |
| 1 | 1 | 2 | 3 | 0 |
| 2 | 2 | 3 | 0 | 1 |
| 3 | 3 | 0 | 1 | 2 |

FIGURE 6.1 – Les tables d'addition de  $\mathbb{Z}_3$  et  $\mathbb{Z}_4$

dans  $\mathbb{Z}_3$  et dans  $\mathbb{Z}_4$ , l'addition d'un élément fixe  $[x]$  quelconque définit une permutation de l'anneau dans lui-même (une bijection). L'inverse de cette bijection est bien sûr donné par l'addition de l'opposé  $-[x]$ . Cette constatation s'étend bien sûr à tous les anneaux commutatifs  $\mathbb{Z}_m$  (et en fait à tous les groupes), comme l'indique la proposition suivante, dont la preuve est immédiate.

**Proposition 6.2.1.** *Pour tout  $m \geq 2$ , et tout  $x \in \mathbb{Z}$ , l'application*

$$a_{[x]} : \mathbb{Z}_m \rightarrow \mathbb{Z}_m : [y] \mapsto [x] + [y]$$

*est une bijection dont l'inverse est donné par  $a_{-[x]}$ .*

Considérons maintenant les tables de multiplication des anneaux  $\mathbb{Z}_3$  et  $\mathbb{Z}_4$ .

| · | 0 | 1 | 2 |
|---|---|---|---|
| 0 | 0 | 0 | 0 |
| 1 | 0 | 1 | 2 |
| 2 | 0 | 2 | 1 |

| · | 0 | 1 | 2 | 3 |
|---|---|---|---|---|
| 0 | 0 | 0 | 0 | 0 |
| 1 | 0 | 1 | 2 | 3 |
| 2 | 0 | 2 | 0 | 2 |
| 3 | 0 | 3 | 2 | 1 |

FIGURE 6.2 – Les tables de multiplication de  $\mathbb{Z}_3$  et  $\mathbb{Z}_4$

On constate que dans  $\mathbb{Z}_3$ , la multiplication par un élément quelconque non nul est une permutation, c'est-à-dire une bijection de  $\mathbb{Z}_3$  dans lui-même. L'inverse de cette bijection est aussi une multiplication par un élément de  $\mathbb{Z}_3$ . Ce n'est pas le cas dans  $\mathbb{Z}_4$ , où la multiplication par 2 n'est pas injective (ni surjective d'ailleurs). Cette différence a des répercussions sur la recherche des solutions des équations à inconnues et coefficients dans  $\mathbb{Z}_m$ . Ce type d'équations est notamment utile pour la cryptographie élémentaire.

On peut voir que  $\mathbb{Z}_3$  est un champ, puisque l'inverse de 1 est 1 et l'inverse de 2 est 2. Par contre, dans  $\mathbb{Z}_4$ , l'élément 2 n'a pas d'inverse, puisqu'il n'existe pas d'élément  $x$  tel que  $2x = 1$ . On voit également que  $\mathbb{Z}_4$  n'est pas intègre, puisque  $2 \cdot 2 = 0$ , et  $2 \neq 0$ . Ces deux constatations sont liées par la proposition suivante.

**Proposition 6.2.2.** *Tout corps est intègre.*

*Démonstration.* Soit  $A$  un corps et soient  $a$  et  $b$  des éléments de  $A$  tels que  $a \cdot b = 0$ . Supposons  $a \neq 0$  et montrons que  $b = 0$ . Puisque  $A$  est un corps, l'élément non nul  $a$  admet un inverse  $a^{-1}$ . On a donc  $a^{-1} \cdot (a \cdot b) = a^{-1} \cdot 0 = 0$ . Mais on a aussi  $a^{-1} \cdot (a \cdot b) = (a^{-1} \cdot a) \cdot b = 1 \cdot b = b$ , donc  $b$  est nul.  $\square$

**Corollaire 6.2.1.** *L'anneau  $(\mathbb{Z}_4, +, 0, \cdot, 1)$  n'est pas un corps.*

Bien entendu, ces considérations s'étendent à tous les anneaux  $\mathbb{Z}_m$ .

**Proposition 6.2.3.** *Pour tout  $m \geq 2$ , si  $m$  n'est pas premier, alors  $\mathbb{Z}_m$  n'est pas un corps.*

*Démonstration.* Si  $m$  n'est pas premier, alors il existe  $a, b \in \mathbb{N}$  tels que  $1 < a, b < m$  et  $m = a \cdot b$ . Alors dans  $\mathbb{Z}_m$ , on a  $[a] \neq [0]$  et  $[b] \neq [0]$ . Pourtant on a  $[a] \cdot [b] = [a \cdot b] = [m] = [0]$ . Donc  $\mathbb{Z}_m$  n'est pas intègre et ce n'est donc pas un corps.  $\square$

Par contre, on constate par inspection directe que  $\mathbb{Z}_5, \mathbb{Z}_7$ , etc... sont des corps. Cela permet de résoudre des équations. Par exemple, dans  $\mathbb{Z}_5$ , l'équation  $3x + 2 = 0$  est facile à résoudre, comme nous allons le voir. La résolution<sup>a</sup> de telles équations est basée sur la notion d'équivalence.

**Définition 6.2.1.** Deux équations sont dites équivalentes si elles ont le même ensemble de solutions. Si les équations  $E_1$  et  $E_2$  sont équivalentes, on note  $E_1 \Leftrightarrow E_2$ .

Dans  $\mathbb{Z}_5$ , l'équation  $3x + 2 = 0$  est équivalente à l'équation  $3x = -2$ , ou encore à l'équation  $3x = 3$ . Cette équation est équivalente à l'équation  $2 \cdot (3x) = 2 \cdot 3$ , ou encore à l'équation  $x = 1$ . La dernière équation n'a évidemment qu'une solution, à savoir 1, et c'est donc aussi le cas de la première. De manière générale, on a le résultat suivant, qui se démontre par double inclusion et que je vous laisse comme exercice.

**Proposition 6.2.4.** *Soit  $(A, +, 0, \cdot, 1)$  un anneau. Une équation  $E$  est équivalente à toute équation obtenue en ajoutant un même élément de  $A$  aux deux membres de  $E$ . Une équation  $E$  est équivalente à toute équation obtenue en multipliant (à droite ou à gauche) les deux membres de  $E$  par un élément  $a$  inversible de  $A$ .*

**Corollaire 6.2.2.** *Si  $(A, +, 0, \cdot, 1)$  est un corps, alors l'équation  $ax + b = 0$  admet une solution unique, quel que soient  $a \in A \setminus \{0\}$  et  $b \in A$ .*

*Démonstration.* D'après la proposition précédente, on a

$$ax + b = 0 \Leftrightarrow ax = -b \Leftrightarrow x = a^{-1}(-b) \Leftrightarrow x = -a^{-1}b.$$

La dernière équation n'a qu'une solution. C'est donc aussi le cas de la première.  $\square$

La situation est plus délicate lorsque l'on considère une telle équation sur un anneau  $A$  qui n'est pas un corps.

**Proposition 6.2.5.** *Dans un anneau  $(A, +, 0, \cdot, 1)$  on considère l'équation  $ax + b = 0$ . Les cas suivants peuvent se produire.*

1. *Si  $a$  est inversible, alors cette équation admet la solution unique  $-a^{-1}b$ ;*
2. *En général, cette équation n'admet une solution que si  $b$  est dans l'image de l'application de multiplication à gauche par  $a$  définie par  $G_a : A \rightarrow A : x \mapsto ax$ ;*
3. *En général, si cette équation admet une solution  $x_0$ , alors l'ensemble de ses solutions est donné par*

$$S = \{x_0 + y : ay = 0\}.$$

*Démonstration.* Pour la première assertion, la preuve est identique à celle qui a été donnée lorsque  $a$  est un corps. Pour la deuxième assertion, on sait que l'équation  $ax + b = 0$  est équivalente à l'équation  $ax = -b$ . Or  $-b$  est multiple de  $a$  si, et seulement si  $b$  l'est.

Pour la troisième assertion, on montre par double inclusion que l'ensemble des solutions de  $ax + b = 0$  est bien l'ensemble proposé. Le fait que  $S$  soit inclus dans l'ensemble des solutions est clair : si  $s$  appartient à  $S$ , alors  $as = a(x_0 + y) = ax_0 + ay = b + 0 = b$ , donc  $s$  est une solution. Soit maintenant  $z$  une solution de  $ax + b = 0$ . On a alors directement  $z = x_0 + (z - x_0)$ . En posant  $y = z - x_0$ , on trouve  $ay = az - ax_0 = b - b = 0$ , donc  $z$  appartient à  $S$ .  $\square$

a. Je ne ferai pas ici une théorie des équations, que je n'ai d'ailleurs pas définies. Vous aurez l'occasion de la voir dans le cours d'algèbre.

*Remarque 6.1.* Lorsqu'une équation admet une solution au moins, elle est dite compatible. Par exemple, dans  $\mathbb{Z}_4$ , l'équation  $2x = 3$  n'est pas compatible. Par contre dans  $\mathbb{Z}_4$  toujours, l'équation  $2x = 2$  admet pour ensemble de solutions  $\{1, 3\}$ . Elle est donc compatible.

La troisième assertion de la proposition précédente ramène l'étude des solutions de l'équation  $ax = b$  à la recherche d'une solution particulière et à la recherche des solutions d'une équation plus simple, à savoir l'équation  $ax = 0$ . C'est l'équation homogène associée à l'équation  $ax + b = 0$ . Ces considérations seront bien sûr généralisées dans le cours d'algèbre, pour des systèmes d'équations.

Les résultats ci-dessus amènent trois questions naturelles :

1. Quand l'anneau  $\mathbb{Z}_m$  est-il un champ ?
2. Si  $\mathbb{Z}_m$  n'est pas un champ, alors quand un élément de  $\mathbb{Z}_m$  est-il inversible ?
3. Dans les deux cas, comment calculer l'inverse d'un élément  $a \in \mathbb{Z}_m$ , lorsqu'il existe ?

Pour répondre à ces trois questions, nous aurons besoin de quelques résultats d'arithmétique, que je vous propose de découvrir.

**Définition 6.2.2.** Si  $a, b \in \mathbb{Z}$  sont non nuls, alors le plus grand commun diviseur (PGCD) de  $a$  et  $b$  est le nombre entier  $d$  strictement positif satisfaisant :

1.  $d|a$  et  $d|b$ ;
2. Si  $c|a$  et  $c|b$ , alors  $c \leq d$ .

Ce nombre est noté  $\text{pgcd}(a, b)$ .

Bien que la notion de  $\text{pgcd}$  soit bien connue et assez intuitive, il n'est pas inutile d'examiner son existence et son unicité. Pour ce faire, on considère l'ensemble des diviseurs positifs communs de deux nombres.

**Définition 6.2.3.** Pour tous  $a, b \in \mathbb{Z}_0$ , on définit l'ensemble  $a$  et  $b$  :

$$E_{a,b} = \{c \in \mathbb{Z} : c > 0, c|a \text{ et } c|b\}$$

On peut alors prouver l'existence et l'unicité du  $\text{pgcd}$ .

**Proposition 6.2.6.** Si  $a, b \in \mathbb{Z}$  sont non nuls, alors le plus grand commun diviseur de  $a$  et  $b$  existe et est unique.

*Démonstration.* Pour tous  $a, b \in \mathbb{Z}_0$ ,  $E_{a,b}$  est non vide car il contient 1. De plus, si un nombre  $c$  divise  $a$ , on montre facilement que  $|c|$  divise  $|a|$ , et que cela implique  $|c| \leq |a|$  (c'est une propriété de la division dans les nombres naturels) et donc  $-|a| \leq c \leq |a|$ . On a donc  $E_{a,b} \subset \{x \in \mathbb{Z} : -|a| \leq x \leq |a|\}$ , et  $E_{a,b}$  admet un plus grand élément.

Prouvons maintenant l'unicité. Si  $d_1$  et  $d_2$  satisfont les conditions définissant le plus grand diviseur, alors  $d_1$  est le plus grand diviseur commun et  $d_2$  est un diviseur commun. Cela donne  $d_2 \leq d_1$ . En échangeant les rôles de  $d_1$  et  $d_2$ , on a aussi  $d_1 \leq d_2$ , ce qui achève la preuve.  $\square$

*Remarque 6.2.* L'argument qui vient d'être utilisé peut être étendu pour un ensemble  $E_{a,b}$  où l'un des deux nombres est nul, mais pas l'autre : si  $a \neq 0$ ,  $E_{a,0}$  est l'ensemble des diviseurs de  $a$ , dont le plus grand élément est  $|a|$ .

Le calcul du  $\text{pgcd}$  en général peut se ramener au calcul du  $\text{pgcd}$  de deux nombres positifs, en vertu de la proposition suivante.

**Proposition 6.2.7.** Pour tous  $a, b \in \mathbb{Z}_0$ , on a

$$\text{pgcd}(a, b) = \text{pgcd}(-a, b) = \text{pgcd}(a, -b) = \text{pgcd}(-a, -b).$$

*Démonstration.* On démontre par double inclusion que les ensembles  $E_{a,b}$ ,  $E_{a,-b}$ ,  $E_{-a,b}$  et  $E_{-a,-b}$  sont égaux. Bien entendu, cela vient du fait que si un nombre  $c$  divise un nombre  $a$ , il divise aussi  $-a$ . Ces ensembles ont donc tous le même plus grand élément.  $\square$

Afin d'obtenir une méthode pour calculer le pgcd, nous avons besoin d'une propriété supplémentaire.

**Proposition 6.2.8.** *Pour tous  $a, b \in \mathbb{Z}_0$  et tout  $m \in \mathbb{Z}$ , on a*

$$\text{pgcd}(a, b) = \text{pgcd}(a, b + ma).$$

Par exemple, le pgcd de 6 et 33 est identique au pgcd de 6 et 45.

*Démonstration.* On prouve ici encore que les ensembles  $E_{a,b}$  et  $E_{a,b+ma}$  coïncident, par double inclusion. Soit  $c \in E_{a,b}$ . Par définition, il existe  $k, k' \in \mathbb{Z}$  tels que  $a = kc$  et  $b = k'c$ . Alors  $b + ma = k'c + mkc = (k' + mk)c$  et  $c$  divise  $a$  et  $b + ma$ . On a donc démontré  $E_{a,b} \subset E_{a,b+ma}$ , quel que soit  $m \in \mathbb{Z}$ .

Pour l'autre inclusion, si  $c$  appartient à  $E_{a,b+ma}$ , il existe  $k, k' \in \mathbb{Z}$  tels que  $a = kc$  et  $b + ma = k'c$ . On a alors  $b = (k' - mk)c$  et  $c$  appartient à  $E_{a,b}$ .  $\square$

**Proposition 6.2.9** (Algorithme d'Euclide). *Soient deux nombres entiers  $a, b$  tels que  $b \geq a > 0$ . On pose  $a = r_0$  et on écrit la suite de divisions*

$$\begin{aligned} b &= aq_1 + r_1 \\ a &= r_1q_2 + r_2 \\ r_1 &= r_2q_3 + r_3 \\ &\vdots \\ r_{j-1} &= r_jq_{j+1} + r_{j+1} \\ &\vdots \\ r_{J-1} &= r_Jq_{J+1} + 0. \end{aligned}$$

Alors le dernier reste non nul  $r_J$  est le pgcd de  $a$  et  $b$ .

*Démonstration.* Tout d'abord, la suite des restes est formée de nombres entiers positifs et est strictement décroissante. Elle atteint donc 0 et il y a toujours un  $J \geq 0$  tel que  $r_{J+1} = 0$ . La procédure de calcul est donc bien définie. Montrons qu'elle conduit au bon résultat. Par la proposition précédente, puisque  $r_1 = b - aq_1$ , on a  $\text{pgcd}(a, b) = \text{pgcd}(a, r_1)$ . On montre de la même façon que  $\text{pgcd}(r_{j-1}, r_j) = \text{pgcd}(r_{j+1}, r_j)$  pour  $j \in \{1, \dots, J-1\}$ , puisque  $r_{j+1} = r_{j-1} - r_jq_{j+1}$ . Cette suite d'égalités montre que le pgcd de  $a$  et  $b$  est le pgcd de  $r_{J-1}$  et  $r_J$ , qui est  $r_J$ , puisque  $r_J$  divise  $r_{J-1}$ .  $\square$

**Théorème 6.2.1** (Bezout). *Si  $a, b \in \mathbb{Z}$  sont non nuls et si  $d = \text{pgcd}(a, b)$ , alors il existe  $x_0, y_0 \in \mathbb{Z}$  tels que  $d = ax_0 + by_0$ .*

*Démonstration.* On considère l'ensemble

$$S = \{ax + by : x, y \in \mathbb{Z} \text{ et } ax + by > 0\}.$$

Cet ensemble est formé de nombres entiers strictement positifs, il est donc inclus dans  $\mathbb{N}$ . Cet ensemble est non vide, car il contient  $a$  si  $a > 0$ , ou  $-a = (-1)a$  si  $a < 0$ . Puisque  $\mathbb{N}$  est bien ordonné, l'ensemble  $S$  admet un élément minimal, que nous notons  $e$ . Montrons que  $e$  est le pgcd de  $a$  et  $b$ . Par définition, il existe  $x_0, y_0 \in \mathbb{Z}$  tels que  $e = ax_0 + by_0$ .

On montre que  $e|a$  par l'absurde. Supposons que  $e$  ne divise pas  $a$ , la division de  $a$  par  $e$  s'écrit

$$a = qe + r, \quad 0 < r < e.$$

Alors on a  $r < e$  et  $r$  appartient à  $S$ , ce qui est absurde. On a en effet  $r > 0$  et  $r = a - qe = a - q(ax_0 + by_0) = a(1 - qx_0) + b(-qy_0)$ .

On montre de la même manière que  $e$  divise  $b$ .

On montre ensuite que  $e$  est le plus grand diviseur commun de  $a$  et  $b$ . En effet, si  $c$  divise  $a$  et  $b$ , soit  $c \leq 0$ , soit  $c > 0$ . Dans le premier cas, on a  $c < e$ . Dans le second,  $c$  divise  $e = ax_0 + by_0$ , donc  $c \leq e$ .  $\square$

**Définition 6.2.4.** Deux nombres entiers non nuls sont premiers entre eux si leur pgcd vaut 1.

**Proposition 6.2.10.** Deux nombres entiers non nuls  $a$  et  $b$  sont premiers entre eux si, et seulement si, il existe  $x_0, y_0 \in \mathbb{Z}$  tels que  $ax_0 + by_0 = 1$ .

Cette proposition implique que si deux nombres sont premiers entre eux, on peut obtenir tout nombre entier comme combinaison à coefficients entiers de ces deux nombres. Ce n'est clairement pas le cas si les deux nombres ne sont pas premiers entre eux. Par exemple, toutes les combinaisons entières des nombres 2 et 8 s'écrivent  $2k + 8k'$ , où  $k, k'$  sont entiers. Ces nombres sont clairement multiples de 2.

*Démonstration.* La nécessité de la condition est une conséquence du théorème de Bezout. Pour la suffisance, supposons qu'il existe  $x_0, y_0 \in \mathbb{Z}$  tels que  $ax_0 + by_0 = 1$  et notons  $d$  le pgcd de  $a$  et  $b$ . Le nombre  $d$  est strictement positif. De plus il divise  $a$  et  $b$ , donc il divise  $ax_0 + by_0$ . On a donc un nombre  $d$  strictement positif qui divise 1. Donc  $d = 1$ , ce qu'il fallait démontrer.  $\square$

Nous avons maintenant tous les outils en main pour caractériser les éléments inversibles dans les anneaux modulaires.

**Proposition 6.2.11.** L'élément  $x = [a] \in \mathbb{Z}_m \setminus \{0\}$  est inversible si, et seulement si,  $a$  est premier avec  $m$ .<sup>b</sup>

*Démonstration.* Il s'agit d'une simple traduction : l'élément  $x = [a]$  est inversible si, et seulement si, il existe  $y \in \mathbb{Z}_m$  tel que  $xy = 1$  dans  $\mathbb{Z}_m$ . Par définition, cette condition est équivalente à l'existence de  $b \in \mathbb{Z}$  tel que  $y = [b]$  et  $[a][b] = 1$  dans  $\mathbb{Z}_m$ . Cette condition se traduit par l'existence de  $b, k \in \mathbb{Z}$ , tels que  $ab = 1 + km$ , ou encore  $ab - km = 1$ . Par la proposition précédente, cette dernière condition est équivalente au fait que  $a$  et  $m$  soient premiers entre eux.  $\square$

Comme corollaire, nous avons le résultat suivant.

**Proposition 6.2.12.** L'anneau  $\mathbb{Z}_p$  est un champ si, et seulement si  $p$  est premier.

*Démonstration.* Nous avons déjà démontré que si  $p$  n'est pas premier, alors  $\mathbb{Z}_p$  n'est pas un champ. Supposons maintenant  $p$  premier et considérons  $x = [a] \in \mathbb{Z}_p \setminus \{0\}$ . Alors  $a$  n'est pas multiple de  $p$  et  $a$  et  $p$  sont donc premiers entre eux. Par la proposition précédente,  $x = [a]$  est inversible dans  $\mathbb{Z}_p$ .  $\square$

De manière générale, on peut calculer l'inverse de  $[a]$  dans  $\mathbb{Z}_p$  quand  $a$  et  $p$  sont premiers entre eux : on sait qu'il existe  $x_0, y_0 \in \mathbb{Z}$  tels que

$$ax_0 + py_0 = 1.$$

On a donc, dans  $\mathbb{Z}_p$ ,  $[a][x_0] + [p][y_0] = [1] = 1$ , qui donne  $[a][x_0] = 1$ . L'inverse de  $[a]$  est donc donné par  $[x_0]$ . De manière générale, si  $a, b \in \mathbb{Z}_0$  sont tels que  $\text{pgcd}(a, b) = d$ , l'existence de  $x_0, y_0$  tels que  $ax_0 + by_0 = d$  est donnée par le théorème de Bezout. Mais ce théorème ne permet pas de les calculer. On peut le faire en complétant l'algorithme d'Euclide.

<sup>b</sup>. Le résultat est encore vrai pour  $[a] = 0$  : dans ce cas,  $[a] = 0$  n'est pas inversible, et  $a$  n'est pas premier avec  $m$ , puisque  $m > 1$ .

**Algorithme 6.2.1.** Si  $a, b \in \mathbb{Z}_0$  sont tels que  $\text{pgcd}(a, b) = d$ , on complète l'algorithme d'Euclide :

$$\begin{array}{ll}
 b &= aq_1 + r_1 & r_1 &= b - aq_1 \\
 a &= r_1q_2 + r_2 & r_2 &= a - r_1q_2 \\
 r_1 &= r_2q_3 + r_3 & r_3 &= r_1 - r_2q_3 \\
 &\vdots & &\vdots \\
 r_{J-2} &= r_{J-1}q_J + r_J & r_J &= r_{J-2} - r_{J-1}q_J \\
 r_{J-1} &= r_Jq_{J+1} + 0.
 \end{array}$$

On sait que  $d = r_J$ . On remplace successivement les restes par leur valeur en fonction des restes précédents. Cela permet d'exprimer  $r_J$  en fonction de  $a$  et  $b$ .

Voici un exemple, où on inverse un élément dans un anneau commutatif, qui n'est pas un champ.

**Exemple 6.2.1.** Calcul de l'inverse de 11 modulo 26 (i.e. dans  $\mathbb{Z}_{26}$ ). On sait que 11 est inversible car il est premier avec 26. On peut le vérifier à l'aide de l'algorithme d'Euclide, et on en profite pour exprimer les restes successifs, en étendant l'algorithme.

$$\begin{array}{ll}
 26 &= 11 \cdot 2 + 4 & 4 &= 26 - 11 \cdot 2 \\
 11 &= 4 \cdot 2 + 3 & 3 &= 11 - 4 \cdot 2 \\
 4 &= 3 \cdot 1 + 1 & 1 &= 4 - 3 \cdot 1 \\
 3 &= 1 \cdot 3 + 0
 \end{array}$$

On utilise alors les égalités à droite pour exprimer un reste en fonction des deux précédents : on a

$$1 = 4 - 3 \cdot 1 = 4 - (11 - 4 \cdot 2) \cdot 1 = 4 \cdot 3 - 11 \cdot 1 = (26 - 11 \cdot 2) \cdot 3 - 11 \cdot 1 = 26 \cdot 3 - 11 \cdot 7.$$

L'inverse de  $[11]$  dans  $\mathbb{Z}_{26}$  est donc  $[-7] = [19]$ . On vérifie effectivement que  $11 \cdot 19 = 209 = 8 \cdot 26 + 1$ .

# Chapitre 7

## Addendum

### 7.1 Un mot d'analyse combinatoire

Dans cette section, je vais voir (ou revoir) les notions de permutations, de factorielles, d'arrangements et de combinaisons. Elles seront utiles en théorie des probabilités. Il s'agit en fait souvent de déterminer le cardinal d'un ensemble donné (essentiellement un ensemble de bijections). Dans la mesure du possible, je donnerai une preuve intuitive de chaque résultat, puis je la formaliserai en me basant sur les notions vues précédemment.

**Définition 7.1.1.** La factorielle d'un nombre naturel est définie récursivement par  $0! = 1$  et  $(n + 1)! = (n + 1)n!$ .

La factorielle intervient dans de nombreux comptages, notamment dans celui des permutations d'ensembles finis.

**Définition 7.1.2.** Soit  $A$  un ensemble. Une permutation de  $A$  est une bijection de  $A$  dans  $A$ . On note  $\text{Bij}(A, B)$  l'ensemble des bijections de  $A$  dans  $B$ . On note  $\text{Bij}(A)$  ou  $\mathfrak{S}(A)$  l'ensemble des bijections de  $A$  dans  $A$ .

On a le résultat suivant qui permet de simplifier le comptage des bijections d'un ensemble fini.

**Proposition 7.1.1.** Si  $f$  est une bijection de  $B$  dans  $B'$ , alors l'application

$$\varphi : \text{Bij}(A, B) \rightarrow \text{Bij}(A, B') : g \mapsto f \circ g$$

est une bijection.

*Démonstration.* On sait que la composée de bijections est une bijection. Donc si  $g$  est une bijection de  $A$  dans  $B$ ,  $f \circ g$  est une bijection de  $A$  dans  $B'$ . On montre que  $\varphi$  est une bijection soit en montrant qu'elle est injective et surjective, soit en donnant son inverse, qui n'est autre que

$$\psi : \text{Bij}(A, B') \rightarrow \text{Bij}(A, B) : h \mapsto f^{-1} \circ h.$$

Ceci achève la preuve. □

On montre de la même manière que si  $A$  est en bijection avec  $A'$ , alors les ensembles  $\text{Bij}(A, B)$  et  $\text{Bij}(A', B)$  sont en bijection, quel que soit  $B$ . En combinant ces deux constatations, on déduit que si  $A$  et  $B$  sont deux ensembles en bijection, alors  $\text{Bij}(A)$  et  $\text{Bij}(B)$  sont également en bijection.

Passons maintenant à une définition classique.

**Définition 7.1.3.** Soit  $A$  un ensemble de cardinal  $n \geq 1$ . Une façon d'ordonner les éléments de  $A^a$  est une bijection de  $\{1, \dots, n\}$  dans  $A$ .

---

a. C'est en fait un ordre total sur  $A$ .

L'idée est que l'on place les éléments dans  $A$  aux places  $1, 2, \dots, n$ , ou simplement que l'on numérote les éléments de  $A$ . On sait qu'une telle bijection existe, parce que  $A$  est en bijection avec  $\{0, \dots, n-1\}$ , lui-même en bijection avec  $\{1, \dots, n\}$ . Mais combien y a-t-il de façons d'ordonner les éléments de  $A$ ?

**Proposition 7.1.2.** *Si  $A$  est de cardinal  $n$ , alors le nombre de façons d'ordonner les éléments de  $A$  est égal à  $|\text{Bij}(A)|$  et à  $|\text{Bij}(\{1, \dots, n\})|$ . Ce nombre vaut  $n!$ .*

*Démonstration.* Puisque  $A$  et  $\{1, \dots, n\}$  sont en bijection, on sait que les ensembles  $\text{Bij}(A)$ ,  $\text{Bij}(\{1, \dots, n\})$ , et  $\text{Bij}(\{1, \dots, n\}, A)$  le sont aussi. Ils ont donc même cardinal. On montre par récurrence sur  $n$  que ce nombre vaut  $n!$ . Pour  $|A| = 1$ , il n'y a qu'une bijection (une seule place, et un seul objet). Supposons que le résultat soit vrai pour  $k$  objets, c'est-à-dire pour  $|A| = k$ , ( $k \geq 1$ ) et montrons qu'il est vrai pour  $k+1$  objets. Pour ordonner  $k+1$  objets, on doit choisir lequel sera en première place, et il y a  $k+1$  façons de le faire. Pour chaque choix de premier objet, il reste  $k$  places et  $k$  objets, que l'on doit ordonner de toutes les façons possibles. Par hypothèse de récurrence, il y a  $k!$  façon de le faire. Au total, il y a donc bien  $(k+1)k!$  façons d'ordonner  $k+1$  objets.  $\square$

J'ai rédigé ici la version intuitive, qui permet de se souvenir facilement du résultat. Le raisonnement utilisé peut être formalisé plus précisément en utilisant les bijections : on a en fait partitionné l'ensemble des bijections de  $\{1, \dots, k+1\}$  dans  $\{1, \dots, k+1\}$  selon l'image qu'elles donnent à l'élément 1. On a

$$\text{Bij}(\{1, \dots, k+1\}) = \bigcup_{i=1}^{k+1} \{f \in \text{Bij}(\{1, \dots, k+1\}) : f(1) = i\},$$

et l'union est disjointe. De plus, les ensembles  $A_i = \{f \in \text{Bij}(\{1, \dots, k+1\}) : f(1) = i\}$  et  $A_j = \{f \in \text{Bij}(\{1, \dots, k+1\}) : f(1) = j\}$  sont en bijection et ont donc le même cardinal. En effet, si on note  $\sigma_{i,j}$  la bijection de  $\{1, \dots, k+1\}$  qui échange  $i$  et  $j$  et fixe les autres éléments (la transposition  $(i, j)$ ), alors si  $f \in A_i$ , on a  $\sigma_{i,j} \circ f \in A_j$  et cela définit une bijection entre ces deux ensembles. Enfin,  $A_1$  est en bijection avec  $\text{Bij}(\{2, \dots, k+1\})$  de manière évidente<sup>b</sup>, ou encore avec  $\text{Bij}(\{1, \dots, k\})$ .

Passons maintenant aux arrangements de  $p$  objets parmi  $n$ .

**Définition 7.1.4.** Soit  $A$  un ensemble de cardinal  $n \in \mathbb{N}_0$  et  $p \in \{1, \dots, n\}$ . Un arrangement de  $p$  objets de  $A$  est un  $p$ -uplet d'éléments distincts de  $A$ . On dit aussi un arrangement de  $p$  objets parmi  $n$ .

Un arrangement de  $p$  objets parmi  $n$  est donné de manière équivalente par un sous-ensemble ordonné de cardinal  $p$  de  $A$ . Par exemple, si  $A = \{a, b, c, d, e\}$ , alors un arrangement de 3 objets dans  $A$  est un mot de trois lettres distinctes. On peut compter facilement qu'il y en a 60 : il y a 5 possibilités pour la première lettre. Ensuite, étant donné qu'on ne peut répéter les lettres, pour chaque choix de première lettre, il y a quatre possibilités pour la deuxième, donc au total, 20 possibilités pour un couple de première et deuxième lettre distinctes. Pour chaque choix d'un tel couple, il reste 3 possibilités pour la troisième lettre, donc 60 en tout.

Le nombre d'arrangements de  $p$  objets de  $A$  ne dépend de  $A$  qu'à travers son cardinal. On peut donc prendre  $A = \{1, \dots, n\}$ . On peut directement généraliser le constat que nous venons de faire pour les mots de trois lettres distinctes.

**Proposition 7.1.3.** *Soit  $n \in \mathbb{N}_0$  et  $p \in \{1, \dots, n\}$ . Le nombre d'arrangements de  $p$  objets parmi  $n$  est donné par*

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

b. Si  $f \in A_1$ , alors  $f|_{\{2, \dots, k+1\}}$  est une bijection de  $\{2, \dots, k+1\}$  dans lui-même.

Ici encore, on peut donner plusieurs preuves. En voici une qui est basée sur une décomposition de l'ensemble des permutations de  $\{1, \dots, n\}$ .

*Démonstration.* Si  $p = n$ , on est ramené au comptage des bijections de  $\{1, \dots, n\}$  dans lui-même. Si  $p < n$ , l'ensemble des permutations de  $\{1, \dots, n\}$  se partitionne en

$$\text{Bij}(\{1, \dots, n\}) = \bigcup_{(i_1, \dots, i_p)} A_{i_1, \dots, i_p}$$

où  $A_{i_1, \dots, i_p} = \{f \in \text{Bij}(\{1, \dots, n\}) : f(1) = i_1, \dots, f(p) = i_p\}$ . L'union est prise sur tous les choix de  $p$ -uplets d'éléments distincts parmi  $\{1, \dots, n\}$ , et les sous-ensembles dont on prend l'union sont disjoints. Tous ces sous-ensembles sont en bijection entre eux et ont donc le même cardinal. Pour le voir, on procède comme dans la démonstration précédente : l'ensemble  $A_{1, \dots, p}$  est en bijection avec  $A_{i_1, \dots, i_p}$  : si on se donne une bijection  $f_0$  de  $\{1, \dots, n\}$  dans lui-même, qui applique 1 sur  $i_1, \dots, p$  sur  $i_p$ , alors pour tout  $f \in A_{1, \dots, p}$ , alors  $f_0 \circ f$  appartient à  $A_{i_1, \dots, i_p}$ . Cela définit une bijection.

De plus,  $A_{1, \dots, p}$  est en bijection avec  $\text{Bij}(\{p+1, \dots, n\})$ , ou  $\text{Bij}(\{1, \dots, n-p\})$ . La décomposition ci-dessus donne donc directement  $n! = A_n^p(n-p)!$ , et le résultat suit.  $\square$

Bien sûr, la démonstration intuitive que voici est plus simple. Mais elle manque peut-être de justifications précises. Je la donne quand même, car elle permet de se souvenir aisément du résultat.

*Démonstration.* Pour choisir un  $p$ -uplet d'éléments distincts parmi  $n$  éléments, il faut choisir le premier. On a  $n$  possibilités. Pour chaque choix possible, on choisit le deuxième éléments parmi  $n-1$  éléments restants, et ainsi de suite jusqu'au choix du  $p$ -ème élément parmi les  $n-p+1$  restants.  $\square$

On peut bien sûr remplacer le “et ainsi de suite” de la preuve par une récurrence sur  $p$ , en commençant par l'évident  $A_n^1 = n$  et en démontrant comme ci-dessus la relation  $A_n^p = nA_{n-1}^{p-1}$ .

Terminons par les combinaisons de  $p$  objets parmi  $n$  objets distinguables.

**Définition 7.1.5.** Une combinaison de  $p$  objets parmi  $n$  objets distinguables est un choix d'un sous ensemble de cardinal  $p$  dans un ensemble de cardinal  $n$ . Le nombre de choix possibles est également appelé combinaison et est noté  $C_n^p$  (ou  $\binom{n}{p}$ ). Par extension, on pose  $C_n^0 = 1$ .

On calcule le nombre  $C_n^p$  comme nous l'avons fait plus haut pour les arrangements.

**Proposition 7.1.4.** Pour tout  $n \in \mathbb{N}_0$  et tout  $p \leq n$ , on a  $C_n^p = \frac{A_n^p}{p!} = \frac{n!}{(n-p)!p!}$ .

*Démonstration.* On prouve la relation  $A_n^p = C_n^p p!$  de la manière suivante. On décrit tous les arrangements de  $p$  éléments parmi  $n$  en choisissant d'abord un sous-ensemble de  $p$  éléments parmi  $n$  de toutes les façons possibles et en ordonnant ses éléments de toutes les façons possibles. Le nombre de choix possibles de sous-ensembles est par définition  $C_n^p$ . Pour chaque choix de sous-ensemble, le nombre de permutations est  $p!$ .  $\square$

On a les relations bien connues sur les combinaisons.

**Proposition 7.1.5.** On a  $C_n^p = C_n^{n-p}$  pour tout  $n \in \mathbb{N}_0$  et  $0 \leq p \leq n$ . De plus  $C_n^p = C_{n-1}^{p-1} + C_{n-1}^p$ , pour tous  $n, p$  tels que  $1 \leq p \leq n-1$ .

La deuxième relation donne lieu au fameux triangle de Pascal. Ces relations permettent de démontrer la formule du binôme de Newton.

**Proposition 7.1.6.** Soit  $A$  un anneau commutatif. Pour tout  $a, b \in A$  et tout  $n \in \mathbb{N}$ , on a

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

*Démonstration.* On démontre la première formule et on obtient la seconde en échangeant les rôles de  $a$  et  $b$  dans la première. On démontre la première par récurrence sur  $n$ . On peut prendre  $n = 0$  comme cas de base. La formule à démontrer se réduit alors à  $1 = 1$ . Elle est donc vraie.

Supposons la formule vraie pour  $n$  et montrons-la pour  $n + 1$ . On a directement, en utilisant l'hypothèse de récurrence :

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

On change d'indice dans la première somme en posant  $k' = k + 1$ , puis en rebaptisant  $k'$  : on a

$$\sum_{k=0}^n C_n^k a^{k+1} b^{n-k} = \sum_{k'=1}^{n+1} C_n^{k'-1} a^{k'} b^{n-k'+1} = \sum_{k=1}^{n+1} C_n^{k-1} a^k b^{n-k+1}.$$

On obtient donc

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

où  $D_{n+1}^0 = C_n^0 = C_{n+1}^0$ ,  $D_{n+1}^{n+1} = C_n^n = C_{n+1}^{n+1}$  et  $D_{n+1}^k = C_n^{k-1} + C_n^k = C_{n+1}^k$  pour tout  $k$  tel que  $1 \leq k \leq n$ . La formule est donc valable pour  $n + 1$ .  $\square$

Remarquons que nous n'avons utilisé dans la preuve que des propriétés des sommes qui sont vraies dans tout anneau, et le fait de pouvoir commuter  $a$  et  $b$ . La formule est donc valable dans un anneau non commutatif, pour tous  $a, b$  tels que  $ab = ba$ .

La formule de binôme de Newton peut être généralisée (par exemple par récurrence sur le nombre de termes) pour calculer une puissance d'une somme. C'est la formule multinomiale, qui permet de calculer des expressions de la forme  $(a + b + c)^n$ , ou plus généralement  $(a_1 + \dots + a_r)^n$ .

# Table des matières

|          |                                                                                                          |           |
|----------|----------------------------------------------------------------------------------------------------------|-----------|
| <b>1</b> | <b>Logique et ensembles</b>                                                                              | <b>2</b>  |
| 1.1      | Logique . . . . .                                                                                        | 2         |
| 1.1.1    | La contraposition : quelques exemples . . . . .                                                          | 9         |
| 1.1.2    | La démonstration par l'absurde . . . . .                                                                 | 9         |
| 1.1.3    | Contre-exemple et démonstration d'une alternative . . . . .                                              | 10        |
| 1.1.4    | La disjonction des cas . . . . .                                                                         | 10        |
| 1.2      | Théorie des ensembles . . . . .                                                                          | 11        |
| 1.2.1    | Un mot sur le paradoxe de Russell . . . . .                                                              | 14        |
| 1.3      | Relations, applications, injections, surjections . . . . .                                               | 14        |
| 1.4      | Applications . . . . .                                                                                   | 19        |
| 1.5      | Applications réciproques . . . . .                                                                       | 22        |
| 1.6      | Images et pré-images de sous-ensembles . . . . .                                                         | 27        |
| <b>2</b> | <b>Nombres complexes</b>                                                                                 | <b>29</b> |
| 2.1      | Introduction . . . . .                                                                                   | 29        |
| 2.2      | Equations du second degré et complétion des carrés . . . . .                                             | 30        |
| 2.3      | Nombres complexes, introduction et définition . . . . .                                                  | 33        |
| 2.4      | Plongement de $\mathbb{R}$ , nombre $i$ , représentation, forme algébrique et nombres associés . . . . . | 36        |
| 2.5      | Forme trigonométrique, formule de Moivre . . . . .                                                       | 41        |
| 2.6      | Un mot sur la forme exponentielle . . . . .                                                              | 43        |
| 2.7      | Interprétation géométrique de l'addition et de la multiplication . . . . .                               | 44        |
| 2.8      | Racines carrées et équations du second degré . . . . .                                                   | 45        |
| 2.9      | Généralisations : puissances $n$ -èmes et racines $n$ -èmes . . . . .                                    | 48        |
| <b>3</b> | <b>Nombres naturels</b>                                                                                  | <b>52</b> |
| 3.1      | La définition de $\mathbb{N}$ et les récurrences . . . . .                                               | 52        |
| 3.2      | Addition et multiplication . . . . .                                                                     | 54        |
| 3.3      | L'ordre usuel sur $\mathbb{N}$ . . . . .                                                                 | 56        |
| 3.3.1    | Relations et ordres . . . . .                                                                            | 56        |
| 3.4      | Soustraction et division . . . . .                                                                       | 59        |
| <b>4</b> | <b>Bijections classiques, cardinal, relations d'équivalence</b>                                          | <b>63</b> |
| 4.1      | Quelques bijections classiques, et un mot sur le cardinal . . . . .                                      | 63        |
| 4.2      | Relations d'équivalence et quotients . . . . .                                                           | 66        |
| <b>5</b> | <b>Nombres et structures algébriques</b>                                                                 | <b>71</b> |
| 5.1      | Le groupe additif $(\mathbb{Z}, +, 0)$ . . . . .                                                         | 71        |
| 5.2      | L'anneau $(\mathbb{Z}, +, 0, \cdot, 1)$ . . . . .                                                        | 75        |
| 5.3      | Le champ $(\mathbb{Q}, +, 0, \cdot, 1)$ . . . . .                                                        | 80        |
| <b>6</b> | <b>Arithmétique modulaire</b>                                                                            | <b>83</b> |
| 6.1      | Définitions et exemples . . . . .                                                                        | 83        |
| 6.2      | Les champs finis $\mathbb{Z}_p$ . . . . .                                                                | 87        |
| <b>7</b> | <b>Addendum</b>                                                                                          | <b>93</b> |
| 7.1      | Un mot d'analyse combinatoire . . . . .                                                                  | 93        |