

Ensembles et logique

Table des matières

1	Fondations	1
1.1	Ensembles et éléments	1
1.2	Quantificateurs	3
1.3	Parties d'un ensemble	4
1.4	Opérateurs sur les ensembles	4
1.5	L'ensemble $\mathbb{N}$ des entiers naturels	5
1.6	Produit cartésien	7
2	Formules propositionnelles	9
2.1	Syntaxe	9
2.2	Sémantique	10
2.3	Négation d'une proposition	14
3	Relations binaires	14
3.1	Définitions	14
3.2	Relations d'ordre	16
3.3	L'ordre naturel et la soustraction	18
3.4	Multiplication dans $\mathbb{N}$ et relation de divisibilité	20
3.5	Maximum et minimum dans $\mathbb{N}$	21
3.6	Relations d'équivalence	24
4	La logique mathématique	26
5	L'art de la démonstration	29
5.1	Conjonction et disjonction	29
5.2	Démonstration par disjonction de cas	29
5.3	Résoudre une équation	30
5.4	Implication	31
5.5	Quantificateurs	32
5.6	Existence et unicité	32
5.7	Démonstration par analyse-synthèse	33
5.8	Inclusion entre ensembles	34
5.9	Démonstrations par récurrence	34

6	Résoudre et rédiger un problème	36
6.1	Les préalables	36
6.2	Règles typographiques	36
6.3	Une rédaction claire	37
6.4	Un peu de stratégie	38

1 Fondations

1.1 Ensembles et éléments

Définition. Sur le plan intuitif, un ensemble E est une “collection” d’objets que l’on appelle les éléments de E , telle que l’on peut “décider” si un élément donné appartient ou non à E .

On écrit “ $x \in E$ ” si et seulement si x est un élément de E . Sinon, on écrit “ $x \notin E$ ”.

Axiome d’extensionnalité : Si E et F sont deux ensembles, alors $E = F$ si et seulement si pour tout $x \in E$, $x \in F$ et pour tout $x \in F$, $x \in E$.

Remarque. Cet axiome décrit une technique très utilisée pour montrer l’égalité entre deux ensembles E et F :

On écrit “soit $x \in E$ ” et l’on montre qu’alors $x \in F$, puis on écrit “soit $x \in F$ ” et l’on montre qu’alors $x \in E$.

Exemple. Montrer (de manière élémentaire) que $\mathbb{R}_+ = \{x \in \mathbb{R} / \forall \varepsilon > 0, x > -\varepsilon\}$.

Si $x \in \mathbb{R}_+$, pour tout $\varepsilon > 0$, on a $x \geq 0 > -\varepsilon$, donc $x > -\varepsilon$.

Réciproquement, soit $x \in \mathbb{R}$ tel que, pour tout $\varepsilon > 0$, $x > -\varepsilon$. Supposons que $x \notin \mathbb{R}_+$. Alors $-\frac{x}{2} > 0$, donc avec $\varepsilon = -\frac{x}{2}$, on obtient $x > -\varepsilon = \frac{x}{2}$, puis $1 < \frac{1}{2}$ car $x < 0$, ce qui est faux. Ainsi $x \in \mathbb{R}_+$.

Propriété. Il existe un unique ensemble ne contenant aucun élément, il est noté $\emptyset$ ¹.

Démonstration.

L’ensemble des x tels que $x \neq x$ ne possède aucun élément, d’où l’existence.

Soit E et F deux ensembles ne contenant aucun élément.

La propriété “pour tout $x \in E$, ...” est toujours vraie, donc l’après l’axiome d’extensionnalité, $E = F$. $\square$

Propriété. Si a est un objet, l’unique ensemble dont a est le seul élément est noté $\{a\}$. C’est un singleton.

De même, lorsque $a \neq b$, $\{a, b\}$ est appelé une paire.

1. L’invention de l’ensemble vide, c’est-à-dire du zéro a constitué un pas décisif dans l’histoire des mathématiques. C’est un concept fondamental de la théorie des ensembles et, en tant qu’élément neutre, de l’algèbre.

Le zéro a été inventé beaucoup plus tard que les entiers naturels supérieurs à un. En particulier, les mathématiciens grecs (Thalès, Pythagore, Euclide, etc.) n’en disposaient pas.

Le zéro apparaît chez les Babyloniens au III^e siècle avant notre ère, mais seulement sous la forme d’un symbole (en forme de double chevron) pour désigner l’absence d’unité d’un certain rang dans l’écriture d’un nombre (en base 60). C’était alors un chiffre, mais pas un nombre que l’on pouvait utiliser seul.

Un tel usage intrinsèque du zéro n’est apparu que vers le V^e siècle, en Inde : ce nombre est défini par le mathématicien indien Brahmagupta comme la soustraction d’un nombre par lui-même ($x - x = 0$). Il était représenté par un cercle et était appelé “sunya”, le mot qui signifie “vide” en sanskrit. La notion et la notation indiennes du zéro ont été empruntées à l’Inde par les mathématiciens arabes au IX^e siècle. Le mot “sunya” a été traduit en arabe par sifr, lequel est d’ailleurs la racine des mots “chiffre” et “zéro” en français.

Définition. Sur le plan intuitif, si E et F sont deux ensembles, une application f de E dans F associe à tout élément x de E un unique élément $f(x)$ de F .

Définition. “vrai” et “faux” sont appelés les deux valeurs booléennes, en mémoire de George Boole.² On les notera également V et F .

Définition. Un prédicat P sur un ensemble E est une application de E dans $\{V, F\}$. Dans ce cas, pour tout $x \in E$, $P(x)$ est vrai ou faux.

Exemple. L’assertion $P(x) : [x^2 > 4]$ est un prédicat sur $\mathbb{R}$. $P(\pi)$ est vrai, $P(0)$ est fausse.

Il existe plusieurs manières de définir un ensemble :

- **Définition par énumération** : on donne tout simplement la liste de ses éléments. Par exemple $E = \{-1, 0, 1\}$,
ou bien $F = \{\emptyset, \{0\}, \{-1, 1\}, \text{truc}\}$; on voit qu’un élément d’un ensemble peut lui-même être un ensemble.
On peut répéter plusieurs fois le même élément, et l’ordre d’apparition des éléments n’a pas d’importance : $\{-1, 0, 1, -1, 0, 0\} = \{0, 1, -1\}$.
- **Définition par construction** : L’ensemble est donné à partir d’ensembles déjà définis et d’opérateurs sur les ensembles, que nous étudierons plus loin. Par exemple $E \cap F$ et $E \cup F$.
- **Définition en compréhension** : Si l’on dispose d’un ensemble E déjà défini et d’un *prédicat* P sur E , alors $\{x \in E / P(x)\}$ représente l’ensemble des éléments x de E tels que $P(x)$ est vrai.
Par exemple $\mathbb{R}_+ = \{x \in \mathbb{R} / x \geq 0\}$.
Il est important de bien maîtriser cette syntaxe.
Si $F = \{x \in E / P(x)\}$, alors pour tout $x \in E$,
on a l’équivalence : $x \in F$ si et seulement si $P(x)$.
- **Définition axiomatique** : on admet l’existence d’un ensemble satisfaisant certains axiomes. Par exemple l’axiome de l’infini affirme qu’il existe un ensemble A tel que, $\emptyset \in A$ et pour tout $y \in A$, $y \cup \{y\} \in A$.³
On peut s’en servir pour construire $\mathbb{N}$, en posant $0 = \emptyset$, $1 = \emptyset \cup \{\emptyset\} = \{\emptyset\}$,
 $2 = 1 \cup \{1\} = \{\emptyset, 1\}$, $\dots$, $n + 1 = n \cup \{n\} = \{\emptyset, 1, \dots, n\} \dots$
- **Définition par induction structurelle** : informellement, pour définir un ensemble E par induction, on fournit un procédé de construction de nouveaux éléments de E à partir d’éléments de E déjà obtenus. On part de la donnée de quelques éléments initiaux de E , et on applique le procédé de construction étape par étape pour obtenir de plus en plus d’éléments. On obtient finalement E .

2. George Boole (1815-1864) est un logicien, mathématicien et philosophe britannique. Il est notamment à l’origine de la logique moderne. Autodidacte, il publia ses premiers travaux tout en exerçant son métier d’instituteur et de directeur d’école dans la région de Lincoln.

3. En théorie formelle des ensembles, tout objet mathématique est un ensemble

Par exemple, nous définirons par induction l'ensemble des formules propositionnelles au chapitre 2.1.

Le paradoxe de Russell (1901)⁴ :

Notons A l'ensemble de tous les ensembles et posons $B = \{x \in A / x \notin x\}$. Alors $B \in B$ si et seulement si $B \notin B$, ce qui est impossible. Cela signifie que A n'est pas un ensemble ! Ainsi, notre définition d'un ensemble est seulement intuitive, non mathématique, mais elle fut utilisée telle quelle par les mathématiciens jusqu'au début du vingtième siècle. A cette époque, la logique est devenue une branche à part entière des mathématiques. Elle permet notamment de donner une définition rigoureuse de la notion d'ensemble qui évite ce type de paradoxe. Il s'agit des axiomes de Zermelo⁵-Fraenkel⁶, qui sortent du programme de MPSI.

1.2 Quantificateurs

Définition du quantificateur universel :

Soit E un ensemble et P un prédicat sur E . La propriété " $\forall x \in E, P(x)$ " signifie que pour tous les éléments x de E , $P(x)$ est vraie, c'est-à-dire que $\{x \in E / P(x)\}$ est égal à E .

Définition du quantificateur existentiel :

Avec les mêmes notations, la propriété " $\exists x \in E, P(x)$ " signifie qu'il existe au moins un $x \in E$ tel que $P(x)$ est vraie, c'est-à-dire que $\{x \in E / P(x)\} \neq \emptyset$.

Existence et unicité : La propriété " $\exists! x \in E, P(x)$ " signifie qu'il existe un unique $x \in E$ tel que $P(x)$ est vraie, c'est-à-dire que $\{x \in E / P(x)\}$ est un singleton.

Remarque. Par exemple, l'axiome d'extensionnalité s'écrit :

$E = F$ si et seulement si $[\forall x \in E, x \in F]$ et $[\forall x \in F, x \in E]$.

Remarque. L'emploi des quantificateurs en guise d'abréviations est exclu : l'usage d'un " $\forall x$ " est toujours suivi d'un " $\in E, P(x)$ " (ou plus rarement d'un " $, P(x)$ "), où P est un prédicat sur E .

Exemple. Si l'on reprend la preuve de $\mathbb{R}_+ = \{x \in \mathbb{R} / \forall \varepsilon > 0, x > -\varepsilon\}$,

il serait incorrect de commencer par :

"Si $x \in \mathbb{R}_+, \forall \varepsilon > 0$ on a $x \geq 0 > -\varepsilon$, donc $x > -\varepsilon$ ".

Cependant, on pourrait continuer par :

"Réciproquement, soit $x \in \mathbb{R}$ tel que $\forall \varepsilon > 0, x > -\varepsilon$ ".

Remarque. Soit P un prédicat sur un ensemble E . Alors dans les phrases

" $\forall x \in E, P(x)$ " et " $\exists x \in E, P(x)$ ", on peut remplacer la variable x par y , ou n'importe quel autre symbole. On dit que, dans les phrases " $\forall x \in E, P(x)$ " et

4. Bertrand Russell, 1872-1970, 3e comte Russell, est un mathématicien, logicien, philosophe, épistémologue, homme politique et moraliste britannique.

5. Ernst Zermelo, 1871-1953, est un mathématicien allemand.

6. Abraham Fraenkel, 1891-1965, est un mathématicien d'abord allemand puis israélien.

“ $\exists x \in E, P(x)$ ”, x est une variable muette. On dit aussi que x est une variable liée. Dans la propriété “ $\exists y \in \mathbb{R}, x = y^2$ ”, y est une variable liée, et par opposition, on dit que x est une variable libre.

1.3 Parties d'un ensemble

Définition. Soit E et F deux ensembles.

On dit que F est inclus dans E et l'on note $F \subset E$ si et seulement si tout élément de F est un élément de E , c'est-à-dire si et seulement si $\forall x \in F, x \in E$.

Lorsque $F \subset E$, on dit que F est une partie de E ou encore un sous-ensemble de E .

Remarque. L'axiome d'extensionnalité peut donc s'écrire :

$E = F$ si et seulement si $E \subset F$ et $F \subset E$.

Remarque. Il faut éviter de confondre les symboles $\in$ et $\subset$.

Ainsi, $1 \in \{1, 2\}$ mais $1 \not\subset \{1, 2\}$, et $\{0\} \subset \{0, 2\}$, mais $\{0\} \not\in \{0, 2\}$.

Informellement, lorsqu'on écrit $A \subset B$, A et B sont des ensembles de même niveau.

Au contraire, lorsqu'on écrit $a \in B$, a et B peuvent être des ensembles, mais ils ne sont pas du même type : a est hiérarchiquement inférieur à B .

Transitivité de l'inclusion : Si $A \subset B$ et $B \subset C$, alors $A \subset C$.

Définition. Si E est un ensemble, on note $\mathcal{P}(E)$ l'ensemble de ses parties.

Exemple. $\mathcal{P}(\{0, 1\}) = \{\emptyset, \{0\}, \{1\}, \{0, 1\}\}$.

$\mathcal{P}(\emptyset) = \{\emptyset\}$.

Remarque. $A \subset B \iff A \in \mathcal{P}(B)$.

1.4 Opérateurs sur les ensembles

Définition. Soit E et F deux ensembles :

- **Intersection** : $x \in E \cap F$ si et seulement si $(x \in E \text{ et } x \in F)$.
- **Réunion** : $x \in E \cup F$ si et seulement si $(x \in E \text{ ou } x \in F)$.
- **Différence ensembliste** : $E \setminus F = \{x \in E / x \notin F\}$.
- **Différence symétrique** : $E \Delta F = (E \setminus F) \cup (F \setminus E)$.
- **Complémentaire de F dans E** : Si F est une partie de E , le complémentaire de F dans E est $\overline{F} = E \setminus F$, que l'on note plus rarement $\mathcal{C}_E^F$.

Propriété. Si F et G sont deux parties d'un ensemble E , alors $F \setminus G = F \cap \overline{G}$.

Remarque. $E \Delta F = F \Delta E = (E \cup F) \setminus (E \cap F)$.

On peut le démontrer par double inclusion en passant aux éléments.

Ainsi, les éléments de $E \Delta F$ sont les éléments qui sont exclusivement ou bien dans E , ou bien dans F , mais pas dans E et F en même temps.

Propriété. Commutativité de l'intersection et de la réunion :

Soit A et B deux ensembles. Alors, $A \cap B = B \cap A$ et $A \cup B = B \cup A$.

Propriété. Associativité de l'intersection et de la réunion : Soit A, B, C trois ensembles. Alors, $A \cap (B \cap C) = (A \cap B) \cap C$ et $A \cup (B \cup C) = (A \cup B) \cup C$.

Définition. Plus généralement, si I est un ensemble et si $(E_i)_{i \in I}$ est une famille d'ensembles (c'est-à-dire que pour chaque $i \in I$, "on se donne" un ensemble E_i), alors on peut définir $\bigcup_{i \in I} E_i$ et $\bigcap_{i \in I} E_i$ par :

$$x \in \bigcup_{i \in I} E_i \iff (\exists i \in I, x \in E_i) \text{ et}$$

$$x \in \bigcap_{i \in I} E_i \iff (\forall i \in I, x \in E_i).$$

Cette dernière définition n'est pas correcte lorsque $I = \emptyset$, car tout x serait élément de l'intersection vide, donc d'après le paradoxe de Russell, l'intersection vide n'est pas un ensemble.

Exemple. On a vu que $\mathbb{R}_+ = \{x \in \mathbb{R} / \forall \varepsilon > 0, x > -\varepsilon\}$, donc $\mathbb{R}_+ = \bigcap_{\varepsilon > 0}]-\varepsilon, +\infty[$.

De même, on peut montrer que $\mathbb{R}_+^* = \bigcup_{\varepsilon > 0}]\varepsilon, +\infty[$.

Propriété. A, B et C sont trois ensembles. $(E_i)_{i \in I}$ est une famille d'ensembles.

Distributivité de l'intersection par rapport à la réunion :

$$A \cap (B \cup C) = (A \cap B) \cup (A \cap C). \quad A \cap \bigcup_{i \in I} B_i = \bigcup_{i \in I} (A \cap B_i).$$

Distributivité de la réunion par rapport à l'intersection :

$$A \cup (B \cap C) = (A \cup B) \cap (A \cup C). \quad A \cup \bigcap_{i \in I} B_i = \bigcap_{i \in I} (A \cup B_i) \text{ (avec } I \neq \emptyset \text{)}.$$

Notation. Soit $(E_i)_{i \in I}$ une famille d'ensembles deux à deux disjoints, c'est-à-dire telle que, pour tout $i, j \in I$ avec $i \neq j$, $E_i \cap E_j = \emptyset$. Alors $\bigcup_{i \in I} E_i$ est appelée une réunion

disjointe et elle est notée $\bigsqcup_{i \in I} E_i$.

Remarque. Si A, B, C sont trois ensembles, la quantité $A \sqcup B \sqcup C$ désigne $A \cup B \cup C$, mais le fait d'utiliser le symbole $\sqcup$ au lieu du symbole $\cup$ indique que A, B et C sont deux à deux disjoints.

1.5 L'ensemble $\mathbb{N}$ des entiers naturels

On admet qu'il existe un ensemble, noté $\mathbb{N}$, satisfaisant les axiomes de Peano⁷ suivants :

- $\mathbb{N}$ est muni d'un élément particulier noté 0 et d'une application "successeur", notée s de $\mathbb{N}$ dans $\mathbb{N}$.

7. Giuseppe Peano, 1858-1932, mathématicien et linguiste italien, inventeur d'une langue auxiliaire internationale, le Latino sine flexione (LsF). Ses qualités d'enseignant ont parfois souffert d'un excès de détails concernant les notations utilisées et les concepts de base, au détriment de l'ensemble du programme qu'il devait traiter.

- 0 n'est le successeur d'aucun entier : $\forall n \in \mathbb{N}, s(n) \neq 0$.
- s est une application injective : pour tout $n, m \in \mathbb{N}$, si $s(n) = s(m)$, alors $n = m$.
- Pour toute partie F de $\mathbb{N}$, si $0 \in F$ et si pour tout $n \in F$, $s(n) \in F$, alors $F = \mathbb{N}$.

Remarques :

◇ Le dernier axiome commence par “pour toute partie de $\mathbb{N}$ ”. L'utilisation d'une telle quantification est caractéristique de la logique du second ordre. Pour se limiter à une logique du premier ordre, où les quantificateurs ne portent que sur des éléments de $\mathbb{N}$, il est nécessaire de compliquer les énoncés des axiomes de Peano, ce qui sort du programme.

◇ Informellement, d'après le dernier axiome, $\mathbb{N} = \{0, s(0), s(s(0)), s(s(s(0))), \dots\}$. On convient de poser $1 = s(0)$, $2 = s(1)$, $3 = s(2)$, etc.

Principe de récurrence : Soit $R(n)$ un prédicat sur $\mathbb{N}$.

Si $R(0)$ est vraie et si pour tout $n \in \mathbb{N}$, $R(n)$ implique $R(s(n))$, alors pour tout $n \in \mathbb{N}$, $R(n)$ est vraie.

Démonstration.

On applique le dernier axiome avec $F = \{n \in \mathbb{N} / R(n)\}$. $\square$

Addition entre entiers : Pour tout $m \in \mathbb{N}$, on pose

$$0 + m = m \text{ et}$$

$$\forall n \in \mathbb{N}, s(n) + m = s(n + m).$$

Ces conditions définissent l'addition entre entiers.

Démonstration.

Fixons $m \in \mathbb{N}$. Notons F_m l'ensemble des entiers n pour lesquels $n + m$ est défini.

$0 \in F_m$ et, pour tout $n \in \mathbb{N}$, si $n \in F_m$, alors $s(n) \in F_m$, donc $F_m = \mathbb{N}$. $\square$

Exemple. $1 + n = s(0) + n = s(0 + n) = s(n)$, $2 + n = s(1) + n = s(1 + n) = s(s(n))$.

Propriétés de l'addition :

- 0 est neutre : $\forall m \in \mathbb{N}, m + 0 = 0 + m = m$.
- Associativité : $\forall n, m, k \in \mathbb{N}, (n + m) + k = n + (m + k)$.
- Commutativité : $\forall n, m \in \mathbb{N}, n + m = m + n$.

Démonstration.

◇ Élément neutre : Soit $m \in \mathbb{N}$. Par définition, $0 + m = m$.

Montrons par récurrence $R(m) : m + 0 = m$.

On a bien $0 + 0 = 0$, d'où $R(0)$.

Pour $m \in \mathbb{N}$, supposons $R(m)$.

Alors $s(m) + 0 = s(m + 0) = s(m)$, d'où $R(s(m))$.

D'après le principe de récurrence, on a bien montré que $\forall m \in \mathbb{N}, m + 0 = m$.

◇ Associativité : Fixons $m, k \in \mathbb{N}$.

Pour tout $n \in \mathbb{N}$, notons $R(n) : (n + m) + k = n + (m + k)$.

Lorsque $n = 0$, $(0 + m) + k = m + k = 0 + (m + k)$, d'où $R(0)$.

Pour $n \in \mathbb{N}$, supposons $R(n)$. Alors,

$(s(n) + m) + k = s(n + m) + k = s((n + m) + k)$ et $s(n) + (m + k) = s(n + (m + k))$, donc d'après $R(n)$, $R(s(n))$ est aussi vraie.

◇ Commutativité : On sait déjà que $n + 0 = 0 + n$ pour tout $n \in \mathbb{N}$.

Montrons que pour tout $n \in \mathbb{N}$, $n + 1 = 1 + n$.

En effet, c'est vrai pour $n = 0$ et si $n + 1 = 1 + n$, alors

$$s(n) + 1 = s(n + 1) = s(1 + n) = s(s(0) + n) = s(s(0 + n)) = s(s(n)) = 1 + s(n).$$

Fixons enfin $m \in \mathbb{N}$ et montrons, toujours par récurrence, que $n + m = m + n$.

C'est vrai pour $n = 0$ et si $n + m = m + n$, alors

$$s(n) + m = s(n + m) = 1 + (n + m) = (m + n) + 1, \text{ puis d'après l'associativité,}$$

$$s(n) + m = m + (n + 1) = m + (1 + n) = m + s(n). \quad \square$$

1.6 Produit cartésien

Définition. Si a et b sont deux objets, posons $(a, b) = \{\{a\}, \{a, b\}\}$.

(a, b) est appelé le couple de composantes a et b .

Propriété. $(a, b) = (c, d)$ si et seulement si $a = c$ et $b = d$.

Remarque. En particulier, lorsque $a \neq b$, $(a, b) \neq (b, a)$, alors que $\{a, b\} = \{b, a\}$.

Démonstration.

On suppose que $\{\{a\}, \{a, b\}\} = \{\{c\}, \{c, d\}\}$.

Supposons que $a \neq c$. Alors $\{a\} \neq \{c\}$, donc $\{a\} = \{c, d\}$, puis $a = c = d$, ce qui est contradictoire.

Ainsi, $a = c$ et $\{\{a\}, \{a, b\}\} = \{\{a\}, \{a, d\}\}$.

Premier cas : Si $\{a, b\} = \{a\}$, alors $a = b = d$.

Second cas : Si $\{a, b\} \neq \{a\}$, alors $\{a, b\} = \{a, d\}$. Alors $b = a$ ou $b = d$, mais si $b = a$, alors $\{a, b\} = \{a\}$ ce qui est faux.

Ainsi dans tous les cas, $b = d$.

La réciproque est évidente. $\square$

Définition. Si A et B sont deux ensembles, on pose $A \times B = \{(a, b) / a \in A \text{ et } b \in B\}$. $A \times B$ s'appelle le produit cartésien de A et B .

Remarque. Selon les axiomes de Zermelo-Fraenkel (cf page 27), $A \times B$ est bien un ensemble, car $A \times B = \{z \in \mathcal{P}(\mathcal{P}(A \cup B)) / \exists a \in A, \exists b \in B, z = \{\{a\}, \{a, b\}\}\}$.

Exemple. $\mathbb{R}^2$ est l'ensemble des couples de nombres réels. On peut l'assimiler à un plan muni d'un repère cartésien en confondant chaque couple (x, y) avec le point de coordonnées (x, y) .

Définition.

- Si a, b, c sont trois objets, le triplet de composantes a, b et c est défini par $(a, b, c) = ((a, b), c)$.
- Si a, b, c, d sont quatre objets, le quadruplet de composantes a, b, c et d est défini par $(a, b, c, d) = ((a, b, c), d)$.

- Soit $n \geq 2$. Supposons définie la notion de n -uplet $(a_1, \dots, a_n)$. On pose alors $(a_1, \dots, a_{n+1}) = ((a_1, \dots, a_n), a_{n+1})$, quels que soient les $n+1$ objets $a_1, \dots, a_{n+1}$.

Propriété. Soit $n \geq 3$. Soit $a_1, \dots, a_n$ et $b_1, \dots, b_n$ $2n$ objets.

Alors $(a_1, \dots, a_n) = (b_1, \dots, b_n)$ si et seulement si $\forall i \in \{1, \dots, n\}, a_i = b_i$.

Démonstration.

Pour tout $n \in \mathbb{N}$, on note $R(n)$ l’assertion suivante : si $a_1, \dots, a_{n+2}$ sont $n+2$ objets, si $b_1, \dots, b_{n+2}$ sont aussi $n+2$ objets, alors $(a_1, \dots, a_{n+2}) = (b_1, \dots, b_{n+2})$ si et seulement si $\forall i \in \{1, \dots, n+2\}, a_i = b_i$.

On a déjà démontré $R(0)$ et, pour tout $n \in \mathbb{N}$, il est simple de montrer que $R(n)$ implique $R(n+1)$. $\square$

Notation. $\mathbb{N}^*$ désigne $\mathbb{N} \setminus \{0\}$.

Définition. Soit $n \in \mathbb{N}^*$. Si $A_1, \dots, A_n$ sont n ensembles, on pose

$$A_1 \times \dots \times A_n = \{(a_1, \dots, a_n) / \forall i \in \{1, \dots, n\}, a_i \in A_i\}.$$

Si E est un ensemble, on note $E^n = \underbrace{E \times \dots \times E}_{n \text{ fois}}$.

Les éléments de E^n s’appellent des n -uplets ou encore des n -listes de E .

Remarque. Avec notre construction, lorsque $n \geq 2$,

$$A_1 \times \dots \times A_n = (A_1 \times \dots \times A_{n-1}) \times A_n.$$

Remarque. Convention, lorsque $n = 1$, le “1-uplet” (a) est égal à a .

Avec cette convention, $E^1 = E$.

Commutativité de deux quantificateurs universels :

Soit E et F deux ensembles. Notons $P(x, y)$ un prédicat défini sur $E \times F$. L’affirmation “ $\forall (x, y) \in E \times F, P(x, y)$ ” est équivalente⁸ à l’affirmation “ $\forall x \in E, \forall y \in F, P(x, y)$ ”, car elles signifient toutes les deux que “pour tout x dans E et y dans F , $P(x, y)$ est vrai”, ou bien encore que “pour tout y dans F et x dans E , $P(x, y)$ est vrai”.

Ainsi, l’affirmation “ $\forall x \in E, \forall y \in F, P(x, y)$ ” est équivalente à “ $\forall y \in F, \forall x \in E, P(x, y)$ ”.

Commutativité de deux quantificateurs existentiels :

On peut de même se convaincre que les deux affirmations

“ $\exists x \in E, \exists y \in F, P(x, y)$ ” et “ $\exists y \in F, \exists x \in E, P(x, y)$ ” sont équivalentes.

ATTENTION :

Un quantificateur universel ne commute pas avec un quantificateur existentiel.

En effet l’affirmation “ $\forall x \in E, \exists y \in F, P(x, y)$ ” signifie que pour tout $x \in E$, il existe un y , qui dépend a priori de x tel que $P(x, y)$, c’est-à-dire qu’il existe une application $x \mapsto y(x)$ de E dans F telle que, pour tout $x \in E, P(x, y(x))$.

Quant à l’affirmation “ $\exists y \in F, \forall x \in E, P(x, y)$ ”, elle signifie qu’il existe $y \in F$ tel que, pour tout $x \in E, P(x, y)$. Ici, y ne dépend pas de x .

En résumé, “ $\forall x \in E, \exists y \in F, P(x, y)$ ” si et seulement si il existe une application

8. Exercice

$x \mapsto y(x)$ de E dans F tel que, pour tout $x \in E$, $P(x, y(x))$,

et “ $\exists y \in F, \forall x \in E, P(x, y)$ ” si et seulement si il existe une application **constante** $x \mapsto y_0$ de E dans F , telle que pour tout $x \in E$, $P(x, y_0)$.

On voit qu’en général, la seconde affirmation implique la première mais que la réciproque est fausse.

Exemple. $\forall x \in \mathbb{Q}, \exists d \in \mathbb{N}^*, dx \in \mathbb{Z}$. En effet, tout rationnel x possède au moins une écriture fractionnelle avec un dénominateur $d \in \mathbb{N}^*$ (on le verra lors de la construction de $\mathbb{Q}$). Cependant, si on intervertit les deux quantificateurs, la phrase $\exists d \in \mathbb{N}^*, \forall x \in \mathbb{Q}, dx \in \mathbb{Z}$ est fausse.

2 Formules propositionnelles

On souhaite construire des propositions logiques structurées à partir de propositions plus simples. Outre les quantificateurs, nous allons définir dans ce but des connecteurs logiques puis définir leur sens : $\vee$ pour “ou non exclusif”, $\wedge$ pour “et”, $\implies$ pour l’implication, $\iff$ pour l’équivalence, $\neg$ pour la négation.

2.1 Syntaxe

Définition par induction des formules propositionnelles : on part d’un ensemble $\mathcal{V}$ dont les éléments sont appelés des variables propositionnelles. On utilise également les “connecteurs logiques” suivants : $\wedge, \vee, \implies, \iff, \neg$.

On suppose que $\mathcal{V} \cap \{\wedge, \vee, \implies, \iff, \neg, \}\emptyset$ est vide.

L’ensemble F des formules propositionnelles est défini par induction structurale :

- Les variables propositionnelles sont des formules propositionnelles.
- si $P, Q \in F$, alors $(P \wedge Q), (P \vee Q), (P \implies Q), (P \iff Q)$ et $\neg P$ sont aussi des formules propositionnelles.

Plus précisément, si l’on note $F_0 = \mathcal{V}$, et pour tout $n \in \mathbb{N}$,

$$F_{n+1} = F_n \cup \{\neg P / P \in F_n\} \cup \{(P \alpha Q) / P, Q \in F_n \text{ et } \alpha \in \{\wedge, \vee, \implies, \iff\}\},$$

$$\text{alors } F = \bigcup_{n \in \mathbb{N}} F_n.$$

Remarque. L’écriture “ $F = \bigcup_{n \in \mathbb{N}} F_n$ ” définit bien F car on peut montrer par récurrence que F_n est défini pour tout $n \in \mathbb{N}$; en effet, si l’on note $R(n)$ la propriété “ F_n est défini”, alors on a $R(0)$ et pour tout $n \in \mathbb{N}$, $R(n)$ implique $R(n+1)$.

Remarque. Une formule propositionnelle s’appelle aussi une proposition, une assertion, une formule, un énoncé, une expression booléenne, etc.

Exemple. Si A, B et C sont des variables propositionnelles, alors

$(A \implies (B \iff A))$ et $((A \wedge (\neg B \implies \neg A)) \wedge (\neg B \vee \neg C)) \implies (C \implies \neg A)$ sont des formules propositionnelles.

$A \implies (B \vee C)$ n'est pas une formule car elle n'est pas correctement parenthésée, mais on relâche en général les règles syntaxiques pour accepter ce type de formule.

$A \implies BC \neg$ n'est pas une formule, en aucun cas.

Exemple. On obtient des “formules” plus habituelles en mathématiques par substitution : par exemple, si l'on part de la formule propositionnelle $(A \wedge B) \implies C$, en substituant A , B et C respectivement par $x \geq 0$, $y \geq 0$ et $x + y \geq 0$, on obtient la formule mathématique suivante : $((x \geq 0) \wedge (y \geq 0)) \implies (x + y \geq 0)$.

Remarque. Autre exemple, la formule d'extensionnalité s'écrit :

$$(E = F) \iff \{[\forall x \in E, x \in F] \wedge [\forall x \in F, x \in E]\}.$$

Exemple. L'axiome de l'infini s'écrit : il existe un ensemble A tel que $(\emptyset \in A) \wedge (\forall y (y \in A \implies y \cup \{y\} \in A))$.

Définition. Si P et Q sont deux formules propositionnelles, $P \wedge Q$ (prononcer “ P et Q ”) s'appelle la conjonction de P et de Q ,

$P \vee Q$ (prononcer “ P ou Q ”) s'appelle la disjonction de P et de Q ,

$P \implies Q$ s'appelle une implication,

$P \iff Q$ est une équivalence,

et $\neg P$ est la négation de la proposition P .

2.2 Sémantique

Définition. Une distribution de valeurs de vérité sur l'ensemble $\mathcal{V}$ des variables propositionnelles est une application de $\mathcal{V}$ dans l'ensemble $\{0, 1\}$. La donnée d'une telle distribution attribue donc à chaque variable propositionnelle l'une des deux valeurs booléennes 0 (qui signifie “faux”) ou 1 (qui signifie “vrai”).

Définition. Soit v une distribution de valeurs de vérité sur l'ensemble $\mathcal{V}$. On prolonge v sur l'ensemble des formules propositionnelles construites à partir de $\mathcal{V}$ de la manière suivante : pour toutes formules propositionnelles P et Q ,

- $v(P \wedge Q) = 1$ si et seulement si $v(P) = v(Q) = 1$ (on dira aussi que $P \wedge Q$ est vraie si et seulement si P et Q sont toutes deux vraies).
- $v(P \vee Q) = 1$ si et seulement si $v(P) = 1$ ou $v(Q) = 1$.

Il convient de noter qu'en mathématiques, le “ou” est par défaut un “ou non exclusif”.

- $v(P \implies Q) = 0$ si et seulement si $v(P) = 1$ et $v(Q) = 0$.

Cela signifie que l'implication $P \implies Q$ est fautive si et seulement si P est vraie alors que Q est fautive.

Ainsi l'implication $P \implies Q$ est vraie si et seulement si P est fautive ou bien Q est vraie.

C'est assez peu intuitif, mais “faux” $\implies$ “n'importe quoi”.

- $v(P \iff Q) = 1$ si et seulement si $v(P) = v(Q)$.
- $v(\neg P) = 1$ si et seulement si $v(P) = 0$.

Définition. La définition précédente est équivalente à la donnée des “tables de vérité” des connecteurs logiques $\wedge$, $\vee$, $\implies$, $\iff$ et $\neg$:

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

Remarque. Les symboles $\implies$ et $\iff$ ne doivent pas être utilisés comme des abréviations, tout au moins en mathématiques. Notamment, écrire “ P , donc Q ” signifie que P est vraie, ce qui permet d’en déduire que Q est vraie. Au contraire, écrire “ $P \implies Q$ ” ne préjuge pas de la valeur booléenne de P .

Le symbole $\iff$ est à réserver pour la résolution d’équations ou d’inéquations :

Exemple. Pour déterminer les couples d’entiers naturels non nuls dont le produit est égal à la somme, fixons $n, m \in \mathbb{N}^*$.

(C) : $nm = n + m \iff n(m-1) - m = 0 \iff n(m-1) - (m-1) - 1 = 0$: astucieux ! Ensuite, (C) $\iff (n-1)(m-1) = 1 \iff n-1 = m-1 = 1$, car $n-1$ et $m-1$ sont des entiers naturels (cf page 20). On peut conclure : (C) $\iff n = m = 2$.

L’astuce utilisée, assez fréquente, consiste à ajouter et à retrancher la même quantité, ou bien à multiplier et diviser par la même quantité non nulle.

Mais pourquoi remplacer ici m par $m-1+1$?

C’est afin de faire apparaître la différence de deux termes semblables :

$n(m-1)$ et $1.(m-1)$.

Lorsqu’on retranche (resp : divise) deux termes, il est souvent intéressant de les mettre sous la même forme.

Définition. Supposons que $P \implies Q$.

- On dit alors que P est une *condition suffisante* pour Q , car il *suffit* que P soit vérifiée pour que Q le soit.
- On dit aussi que Q est une *condition nécessaire* pour P , car lorsque P est vérifiée, alors Q est *nécessairement* vérifiée.

Lorsque $P \iff Q$, on dit que P est une *condition nécessaire et suffisante* pour Q .

Remarque. On rencontre souvent des énoncés de la forme “donnez une condition C pour que P ”. Par exemple, donner une condition portant sur deux ensembles A et B pour que $A \cap B = A \cup B$.

De tels énoncés sont ambigus car ils ne précisent pas si l’on doit chercher une condition suffisante (i.e : trouver C telle que $C \implies P$) ou bien une condition nécessaire (i.e : trouver C telle que $P \implies C$). Le mieux dans cette situation est de déterminer une condition nécessaire et suffisante (CNS).

Exemple. Etudions l’exemple précédent.

Analyse : Supposons que $A \cap B = A \cup B$. Alors $A \subset A \cup B = A \cap B \subset B$ et de même, $B \subset A$, donc $A = B$.

Ainsi, une condition nécessaire est : $A = B$.

Synthèse : Réciproquement, si $A = B$, alors $A \cap B = A = A \cup B$, donc la condition $A = B$ est une CNS pour que $A \cap B = A \cup B$.

Définition. Une tautologie est une formule propositionnelle qui est toujours vraie, quelle que soit la distribution de valeurs de vérité des variables propositionnelles qui interviennent dans la formule.

Exemple. Voici quelques exemples fondamentaux de tautologies, où A, B, C sont des formules propositionnelles quelconques :

1. $A \vee \neg A$: principe du tiers-exclus,
2. $(A \vee A) \implies A$: idempotence de $\vee$ (on a aussi l'idempotence de $\wedge$),
3. $(A \wedge B) \iff (B \wedge A)$: commutativité de $\wedge$ ($\vee$ est également commutatif),
4. $(A \vee (B \vee C)) \iff ((A \vee B) \vee C)$: associativité de $\vee$ ($\wedge$ est aussi associatif),
5. $(A \wedge (B \vee C)) \iff ((A \wedge B) \vee (A \wedge C))$: distributivité de $\wedge$ par rapport à $\vee$,
6. $(A \vee (B \wedge C)) \iff ((A \vee B) \wedge (A \vee C))$: distributivité de $\vee$ par rapport à $\wedge$,
7. $(A \wedge (A \vee B)) \iff A$: première loi d'absorption,
8. $((A \vee (A \wedge B)) \iff A$: seconde loi d'absorption,
9. $(\neg(A \vee B)) \iff (\neg A \wedge \neg B)$: loi de Morgan⁹,
10. $(\neg(A \wedge B)) \iff (\neg A \vee \neg B)$: loi de Morgan,
11. $(A \implies B) \iff (\neg B \implies \neg A)$: contraposition.
12. $(A \implies B) \iff (\neg A) \vee B$ (une définition de l'implication),
13. $\neg(A \implies B) \iff A \wedge (\neg B)$,

Remarque. Les 4 premières tautologies ne sont pas à apprendre car on les utilise en pratique sans même y réfléchir.

Démonstration.

Une méthode systématique de démonstration de telles formules consiste à utiliser les tables de vérité.

À titre d'exemple, démontrons la seconde loi de Morgan :

A	B	$A \vee B$	$\neg(A \vee B)$	$\neg A \wedge \neg B$
V	V	V	F	F
V	F	V	F	F
F	V	V	F	F
F	F	F	V	V

On peut aussi construire des démonstrations qui font appel à ... la logique.

Démontrons ainsi la première loi d'absorption :

Si A est vraie, alors $A \vee B$ puis $A \wedge (A \vee B)$ sont aussi vraies.

Si A est fausse, a fortiori, $A \wedge (A \vee B)$ est fausse. $\square$

9. Auguste De Morgan (1806-1871), mathématicien et logicien britannique.

Autres exemples de tautologies : (inutile de les apprendre)

- $((A \implies B) \wedge A) \implies B$ (règle du modus ponens),
- $((A \implies B) \wedge \neg B) \implies \neg A$ (règle du modus tollens),
- $(\neg A \implies A) \implies A$,
- $(A \implies B) \vee (C \implies A)$,
- $((A \implies B) \wedge (B \implies C)) \implies (A \implies C)$ (transitivité de l'implication).

Démonstration.

◇ Montrons la transitivité de l'implication.

Première méthode : Supposons que $(A \implies B) \wedge (B \implies C)$.

On veut montrer que $A \implies C$. Pour cela, supposons A et montrons C .

A et $A \implies B$ étant vraies, B est vraie.

De même, B et $B \implies C$ étant vraies, C est vraie.

Avec les tables de vérité :

A	B	C	$A \implies B$	$B \implies C$	$(A \implies B) \wedge (B \implies C)$	$A \implies C$	$((A \implies B) \wedge (B \implies C)) \implies (A \implies C)$
V	V	V	V	V	V	V	V
V	V	F	V	F	F	F	V
V	F	V	F	V	F	V	V
V	F	F	F	V	F	F	V
F	V	V	V	V	V	V	V
F	V	F	V	F	F	V	V
F	F	V	V	V	V	V	V
F	F	F	V	V	V	V	V

◇ $(A \implies B) \vee (C \implies A)$ est une tautologie, car lorsque A est fausse, $A \implies B$ est vraie, et lorsque A est vraie, $C \implies A$ est vraie. □

Définition. On dit que deux propositions P et Q sont logiquement équivalentes si et seulement si la proposition $P \iff Q$ est une tautologie. On note alors $P \equiv Q$.

Ainsi, lorsque l'on ne s'intéresse qu'à la valeur booléenne des propositions, on peut remplacer toute proposition par une proposition qui lui est logiquement équivalente.

Exemple. La distributivité de $\wedge$ par rapport à $\vee$ signifie que les formules $(A \wedge (B \vee C))$ et $((A \wedge B) \vee (A \wedge C))$ sont logiquement équivalentes.

Définition. La contraposée de l'implication $A \implies B$ est égale à $\neg B \implies \neg A$.

Toute implication est logiquement équivalente à sa contraposée.

Démonstration.

C'est la tautologie 11, ou bien :

d'après la tautologie 12, $(A \implies B) \equiv (\neg A) \vee B$ et

$(\neg B \implies \neg A) \equiv B \vee (\neg A)$. □

Remarque. Les tautologies 3 à 13 sont des $\iff$, donc elles énoncent que deux propositions sont logiquement équivalentes. Elles permettent d'écrire une succession d'expressions connectées par le symbole $\equiv$, c'est-à-dire de faire du *calcul booléen*.

2.3 Négation d'une proposition

- ◇ Les lois de Morgan permettent de nier une conjonction ou une disjonction :
 $\neg(A \vee B)$ est logiquement équivalente à $(\neg A) \wedge (\neg B)$,
 $\neg(A \wedge B)$ est logiquement équivalente à $(\neg A) \vee (\neg B)$.
- ◇ La négation d'une négation redonne la propriété initiale :
 $\neg(\neg A)$ est logiquement équivalente à A .
- ◇ La négation d'une implication est plus délicate. Informellement, $A \implies B$ est fausse si et seulement si A est vraie alors que B est fausse. Ainsi, lorsque A est fausse, il n'y a pas de problème, $A \implies B$ est vraie.
Formellement, $\neg(A \implies B)$ est logiquement équivalente à $A \wedge (\neg B)$.
- ◇ Une équivalence est la conjonction de deux implications, donc
 $\neg(A \iff B)$ est logiquement équivalente à $[\neg(A \implies B)] \vee [\neg(B \implies A)]$.

Propriété. Soit P un prédicat sur un ensemble E .

L'assertion " $\forall x \in E, P(x)$ " est fausse si et seulement si il existe $x \in E$ tel que $\neg P(x)$.
Ainsi, $\neg[\forall x \in E, P(x)] \iff [\exists x \in E, \neg P(x)]$.

De même, $\neg[\exists x \in E, P(x)] \iff [\forall x \in E, \neg P(x)]$.

Exemple. On dispose maintenant de toutes les propriétés nécessaires pour nier n'importe quelle formule mathématique. Par exemple :

- Un ensemble A est inclus dans un ensemble B si et seulement si : $\forall x \in A, x \in B$.
On en déduit que A n'est pas inclus dans B si et seulement si : $\exists x \in A, x \notin B$.
- Une fonction f de $\mathbb{R}$ dans $\mathbb{R}$ est croissante si et seulement si :
 $\forall x, y \in \mathbb{R}, x \leq y \implies f(x) \leq f(y)$.
Ainsi, f n'est pas croissante si et seulement si :
 $\exists x, y \in \mathbb{R}, (x \leq y) \wedge (f(x) > f(y))$: peu de rapport avec la décroissance de f .
- Une suite $(x_n)_{n \in \mathbb{N}}$ de réels converge vers 0 si et seulement si :
 $\forall \varepsilon > 0, \exists N \in \mathbb{N}, \forall n \in \mathbb{N}, [n \geq N \implies |x_n| \leq \varepsilon]$. Ainsi, une suite $(x_n)_{n \in \mathbb{N}}$ de réels ne converge pas vers 0 si et seulement si :
 $\exists \varepsilon > 0, \forall N \in \mathbb{N}, \exists n \in \mathbb{N}, (n \geq N) \wedge (|x_n| > \varepsilon)$.

Propriété. Soit A et B deux ensembles de E .

Soit $(E_i)_{i \in I}$ une famille de parties de E , avec $I \neq \emptyset$. Alors,

- $\overline{\overline{A}} = A, \quad \overline{A \cup B} = \overline{A} \cap \overline{B}, \quad \overline{A \cap B} = \overline{A} \cup \overline{B},$
- $A \subset B \iff \overline{B} \subset \overline{A},$
- $\bigcap_{i \in I} E_i = \overline{\bigcup_{i \in I} \overline{E_i}}, \quad \bigcup_{i \in I} E_i = \overline{\bigcap_{i \in I} \overline{E_i}}.$

3 Relations binaires

3.1 Définitions

Définition. Soit E et F deux ensembles.

Une relation binaire R sur $E \times F$ est une partie de $E \times F$.

Lorsque la partie R est appelée une relation binaire, on convient, pour tout $(x, y) \in E \times F$, de noter “ xRy ” au lieu de “ $(x, y) \in R$ ”.

Le graphe de la relation binaire R est $\{(x, y) \in E \times F / xRy\}$, donc le graphe de R est ... égal à R .

Remarque. Il s’agit donc juste d’un changement de notations et de vocabulaire. Sans cela, la définition donnée ci-dessous de la transitivité d’une relation binaire serait peu intuitive.

Le fait de changer de langage n’est pas anodin en mathématiques. Une même propriété se traduit différemment selon le langage utilisé. Elle peut être plus facile à démontrer dans un langage plutôt qu’un autre.

Exemple. Prenons pour E l’ensemble des êtres humains, et pour F l’ensemble de tous les prénoms. Lorsque $(x, y) \in E \times F$, on convient que xRy si et seulement si le prénom y est l’un des prénoms de l’être humain x .

R est une relation binaire sur $E \times F$.

Remarque. Lorsque E et F sont finis, il existe plusieurs manières de représenter une relation R sur $E \times F$:

- Par un tableau : on dispose les éléments de E sur la première colonne, ceux de F sur la première ligne. On place une croix dans les cases correspondant à la ligne $x \in E$ et à la colonne $y \in F$ si et seulement si xRy .
- Avec des patates : on représente les ensembles E et F par des formes patatoïdales. On dessine une flèche d’un élément x de E vers un élément y de F si et seulement si xRy .
- Lorsque $E = F$, la représentation précédente devient un graphe orienté dont les sommets sont les éléments de E .

Définition. Lorsque $E = F$, on dit que R est une relation binaire sur E (plutôt que sur E^2). Dans ce cas,

- R est réflexive si et seulement si $\forall x \in E, xRx$,
- R est symétrique si et seulement si $\forall x, y \in E, (xRy) \implies (yRx)$,
- R est antisymétrique si et seulement si $\forall x, y \in E, [(xRy) \wedge (yRx) \implies x = y]$,
- et R est transitive si et seulement si $\forall x, y, z \in E, [(xRy) \wedge (yRz) \implies (xRz)]$.

Exemple. L’égalité est une relation binaire sur tout ensemble E . Elle est réflexive, symétrique, antisymétrique et transitive.

Exemple. La relation $\neq$ n’est pas transitive. Aussi est-il préférable d’éviter d’écrire “ $x \neq y \neq z$ ”.

Exemple. Si E est l’ensemble des droites de l’espace usuel (de dimension 3), la relation d’orthogonalité est symétrique, mais elle n’est ni réflexive, ni antisymétrique, ni transitive.

3.2 Relations d'ordre

Définition. Une relation binaire R sur un ensemble E est appelée une relation d'ordre si et seulement si R est réflexive, antisymétrique et transitive.

Exemple. Si A est un ensemble, la relation d'inclusion est une relation d'ordre sur $\mathcal{P}(A)$.

Définition. Une relation d'ordre R sur un ensemble E est totale si et seulement si pour tout couple (x, y) de E^2 , x et y sont comparables, c'est-à-dire $(xRy) \vee (yRx)$. Sinon, on dit que R est un ordre partiel.

Exemple. La relation d'inclusion sur $\mathcal{P}(A)$ n'est pas totale dès que A possède plus de deux éléments.

En effet, si a et b sont deux éléments distincts de A , $\{a\}$ et $\{b\}$ ne sont pas comparables.

Exemple. Prenons $E = \mathcal{F}([0, 1], \mathbb{R})$. Si $f, g \in E$, on convient que $f \leq g$ si et seulement si pour tout $t \in [0, 1]$, $f(t) \leq g(t)$. On définit ainsi une relation d'ordre sur E . C'est un ordre partiel, car en posant par exemple $f(x) = x$ et $g(x) = 1 - x$, on a $\neg(f \leq g)$ et $\neg(g \leq f)$.

Notation. Pour la suite de ce paragraphe, on fixe une relation d'ordre sur un ensemble E , que l'on note " $\preceq$ ".

Pour tout $x, y \in E$, on convient de noter $x \prec y$ si et seulement si $x \preceq y$ et $x \neq y$ et on convient que $x \succeq y$ si et seulement si $y \preceq x$.

La relation $\succeq$ est aussi une relation d'ordre (appelée l'ordre réciproque de $\preceq$), mais la relation d'ordre *strict* $\prec$ n'est pas réflexive, donc ce n'est pas une relation d'ordre.

Attention : lorsque l'ordre n'est pas total, $\neg(x \preceq y)$ n'est pas équivalent à $x \succ y$.

Remarque. Pour la relation d'inclusion, l'usage est d'utiliser le symbole $\subset$ pour désigner la relation d'ordre d'inclusion, et non $\subseteq$.

Exemple. On définit sur E^2 la relation binaire $\mathcal{L}$ par :

$$\forall (x, y), (a, b) \in E^2, (x, y)\mathcal{L}(a, b) \iff [(x \prec a) \vee ((x = a) \wedge (y \preceq b))].$$

Vérifier que $\mathcal{L}$ est une relation d'ordre.

Montrer que $\mathcal{L}$ est totale si et seulement si $\preceq$ est totale.

Lorsque E est l'ensemble des lettres de l'alphabet et que $\preceq$ est l'ordre usuel $a, b, c, \dots, z$, on vient de définir l'ordre lexicographique entre mots de deux lettres.

Généraliser aux mots de n lettres. On définit ainsi une relation binaire $\mathcal{L}_n$ sur E^n . Montrer que c'est un ordre et qu'il est total si et seulement si $\preceq$ est total.

Définition. Soit F une partie de E et $m \in E$. On dit que m est un majorant de F si et seulement si pour tout $a \in F$, $a \preceq m$. On définit de même la notion de minorant d'une partie de E .

On dit qu'une partie est majorée si et seulement si elle possède au moins un majorant.

On dit qu'une partie est minorée si et seulement si elle possède au moins un minorant.

On dit qu'une partie est bornée si et seulement si elle est majorée et minorée.

Exemple. Posons $\mathcal{P}(A) = E$, muni de la relation d'inclusion.

Toute partie F de E est minorée par $\emptyset$ et majorée par A .

Si F est une partie de E et si $B \in E$, B est un majorant de F si et seulement si $B \supset \bigcup_{D \in F} D$, et B est un minorant de F si et seulement si $B \subset \bigcap_{D \in F} D$.

Dans ce contexte, on peut convenir que l'intersection vide est égale à A .

Définition. Si F est une partie de E et $m \in E$, on dit que m est le maximum de F si et seulement si m majore F et $m \in F$. On définit de même le minimum de F .

Ainsi que cette définition le sous-entend, si une partie possède un maximum (resp : un minimum), il est unique. Il est noté $\max(E)$ (resp : $\min(E)$).

Exemple. Reprenons l'exemple précédent. Alors $\emptyset = \min(E)$ et $A = \max(E)$.

Si $F = \{\{a\}, \{b\}\}$ avec $a \neq b$, F ne possède ni maximum, ni minimum.

Si F est une partie quelconque de E , l'ensemble des majorants de F possède un minimum, égal à $\bigcup_{D \in F} D$. De même, l'ensemble des minorants de F possède un maximum, égal à $\bigcap_{D \in F} D$.

Définition. Si F est une partie de E , un élément m de F est dit maximal si et seulement si $\forall x \in F (x \succeq m \implies x = m)$, i.e si et seulement si $\forall x \in F \neg(m \prec x)$.

Il est minimal si et seulement si $\forall x \in F (x \preceq m \implies x = m)$, i.e si et seulement si $\forall x \in F \neg(m \succ x)$.

Démonstration.

Soit $x \in F : (x \succeq m \implies x = m) \equiv \neg(x \succeq m) \vee (x = m)$ et $\neg(m \prec x) \equiv \neg(m \preceq x \wedge m \neq x) \equiv \neg(x \succeq m) \vee (x = m)$. $\square$

Exemple. Toujours avec le même exemple, Si $F = \mathcal{P}(A) \setminus \{\emptyset\}$, les éléments minimaux de F sont exactement les singletons.

Propriété. Lorsque la relation d'ordre est totale, toute partie F de E possède au plus un élément maximal et dans ce cas, c'est le maximum de F . Idem avec minimal et minimum.

Remarque. Ainsi, pour une relation d'ordre totale, les notions d'éléments minimaux et maximaux sont inutiles.

Remarque. La réciproque de la propriété précédente est vraie. En effet, supposons que E est muni d'une relation d'ordre partiel. Alors il existe deux éléments a et b de E qui ne sont pas comparables. L'ensemble $\{a, b\}$ admet a et b comme éléments maximaux.

Exercice. Si E est un ensemble fini et non vide, pour tout ordre défini sur E , montrer que E possède au moins un élément minimal.

Solution : Considérons un ordre $\leq$ sur E et supposons que E ne possède aucun élément minimal pour cet ordre.

E étant non vide, il existe $a_0 \in E$.

a_0 n'est pas un élément minimal de E , donc il existe $a_1 \in E$ avec $a_1 < a_0$.

a_1 n'est pas un élément minimal de E , donc il existe $a_2 \in E$ avec $a_2 < a_1$.
 Soit $n \in \mathbb{N}$. Supposons construits $a_0, \dots, a_n \in E$ tels que,
 pour tout $i \in \{0, \dots, n-1\}$, $a_{i+1} < a_i$.
 a_n n'est pas un élément minimal de E , donc il existe $a_{n+1} \in E$ avec $a_{n+1} < a_n$.
 On a ainsi fourni un procédé de construction par récurrence d'une suite (a_k)
 strictement décroissante d'éléments de E .
 Pour tout $i < j$, $a_j < a_i$, donc $\{a_i \mid i \in \mathbb{N}\}$ est infini, ce qu'il fallait démontrer.

3.3 L'ordre naturel et la soustraction

L'ordre naturel : Pour tout $n, m \in \mathbb{N}$,
 on convient que $n \leq m$ si et seulement si $\exists k \in \mathbb{N}$, $m = n + k$.
 Dans ce cas, k est unique. On le note $k = m - n$.

Démonstration.

Pour l'unicité, il faut montrer que, si $n + k = n + h$, alors $k = h$.
 Fixons $h, k \in \mathbb{N}$. Notons $R(n)$: si $n + k = n + h$, alors $k = h$.
 $R(0)$ est vraie. Pour $n \in \mathbb{N}$, supposons $R(n)$.
 Supposons que $s(n) + k = s(n) + h$. Alors $s(n+k) = s(n+h)$, donc d'après le troisième
 axiome de Peano, $n + k = n + h$. D'après $R(n)$, $h = k$. $\square$

Définition. On vient de montrer que, si n est un entier naturel,
 pour tout $h, k \in \mathbb{N}$, $n + h = n + k$ implique $h = k$.
 On dit que n est régulier.

Remarque. Par définition de l'opérateur “-”, on a :
 pour tout $m, n \in \mathbb{N}$ tel que $n \leq m$, $n + (m - n) = m$.
 Pour tout $n \in \mathbb{N}$, $n \leq n$ et $n - n = 0$.
 Pour tout $n \in \mathbb{N}$, $0 \leq n$ et $n - 0 = n$.

Propriété. Sous condition d'existence des différences, on a, pour tout $n, m, k \in \mathbb{N}$:
 $(n - m) - k = n - (m + k)$ et $n - (m - k) = (n - m) + k = (n + k) - m$.
 En particulier l'opérateur “-” n'est pas associatif.

Démonstration.

◇ On suppose que $m \leq n$: il existe $a \in \mathbb{N}$ tel que $n = m + a$. Alors $a = n - m$.
 On suppose de plus que $k \leq a$: il existe $b \in \mathbb{N}$ tel que $a = k + b$.
 Alors $(n - m) - k = a - k = b$. De plus $n = m + a = m + k + b$, donc $m + k \leq n$ et
 $n - (m + k) = b$. On a bien montré que $(n - m) - k = n - (m + k)$.
 ◇ On suppose que $k \leq m$: il existe $a \in \mathbb{N}$ tel que $m = k + a$. Alors $m - k = a$.
 On suppose de plus que $a \leq n$: il existe $b \in \mathbb{N}$ tel que $n = a + b$.
 Alors $n - (m - k) = n - a = b$. Or $n + k = (k + a) + b = m + b$,
 donc $m \leq n + k$ et $(n + k) - m = b$.
 On suppose de plus que $m \leq n$: il existe $c \in \mathbb{N}$ tel que $n = m + c$.
 $n = a + b = m + c = k + a + c$, donc par régularité de a , $b = k + c$.
 Ainsi $(n - m) + k = c + k = b$. $\square$

Lemme : 0 est l'unique élément de $\mathbb{N}$ qui n'est pas le successeur d'un autre entier. Il est équivalent de dire que, pour tout $n \in \mathbb{N}$, $n \neq 0$ si et seulement si il existe $k \in \mathbb{N}$ tel que $n = s(k)$.

Démonstration.

Notons $F = \{0\} \cup \{s(k)/k \in \mathbb{N}\}$. D'après le dernier axiome de Peano, $F = \mathbb{N}$. Ainsi chaque élément non nul de $\mathbb{N}$ est le successeur d'un entier, mais d'après le second axiome de Peano, ce n'est pas le cas de 0. $\square$

Propriétés de l'ordre naturel :

- Réflexivité : $\forall n \in \mathbb{N}, n \leq n$.
- Antisymétrie : Pour tout $n, m \in \mathbb{N}$, si $n \leq m$ et $m \leq n$, alors $n = m$.
- Transitivité : Pour tout $n, m, p \in \mathbb{N}$, si $n \leq m$ et $m \leq p$, alors $n \leq p$.

Ainsi, l'ordre naturel est bien une relation d'ordre sur $\mathbb{N}$.

Démonstration.

- $\diamond$ $n = n + 0$, donc $n \leq n$.
- $\diamond$ Supposons que $n \leq m$ et $m \leq n$. Il existe $a, b \in \mathbb{N}$ tels que $m = n + a$ et $n = m + b$. Alors $m = m + b + a$, or m est régulier, donc $b + a = 0$. Si $a \neq 0$, d'après le lemme, il existe $c \in \mathbb{N}$ tel que $a = s(c)$, donc $0 = s(c) + b = s(c + b)$, ce qui est impossible. Ainsi $a = 0$ puis $m = n$.
- $\diamond$ Supposons que $n \leq m$ et $m \leq p$. Il existe $a, b \in \mathbb{N}$ tels que $m = n + a$ et $p = m + b$. Alors $p = n + (a + b)$, donc $n \leq p$. $\square$

Propriété. Soit $m, n \in \mathbb{N}$. On note $m < n$ lorsque $m \leq n$ et $m \neq n$.

Si $m < n$, alors $m + 1 \leq n$.

Démonstration.

Supposons que $m < n$. Il existe $a \in \mathbb{N}$ tel que $n = m + a$ et $a \neq 0$.

D'après le lemme, il existe $b \in \mathbb{N}$ tel que $a = s(b)$, donc $n = m + b + 1$ et $n \geq m + 1$. $\square$

Propriété. L'ordre naturel est un ordre total sur $\mathbb{N}$.

Démonstration.

Soit $N \in \mathbb{N}$. Notons $R(N)$: pour tout $n, m \in \mathbb{N}$ tels que $n \leq N$ et $m \leq N$, n et m sont comparables.

$R(0)$ est vraie.

Supposons $R(N)$. Soit $n, m \in \mathbb{N}$ tels que $n \leq N + 1$ et $m \leq N + 1$.

Si $n \leq N$ et $m \leq N$, n et m sont comparables d'après $R(N)$.

Sinon, on a par exemple $n > N$, donc $n \geq N + 1$, puis $n = N + 1$. Alors $m \leq n$.

On a prouvé $R(N + 1)$. $\square$

La relation d'ordre est compatible avec l'addition :

Pour tout $a, b, c, d \in \mathbb{N}$, si $a \leq b$ et $c \leq d$, alors $a + c \leq b + d$.

Démonstration.

Supposons que $a \leq b$ et $c \leq d$. Il existe $e, f \in \mathbb{N}$ tels que $b = a + e$ et $d = c + f$. Alors $b + d = (a + c) + (e + f)$, donc $a + c \leq b + d$. $\square$

3.4 Multiplication dans $\mathbb{N}$ et relation de divisibilité

Multiplication entre entiers :

Pour tout $m \in \mathbb{N}$, on pose

$$0 \times m = 0 \text{ et}$$

$$\forall n \in \mathbb{N}, s(n) \times m = n \times m + m.$$

Ces conditions définissent la multiplication entre entiers.

Démonstration.

Exercice. $\square$

Remarque. On note souvent nm au lieu de $n \times m$.

Exemple. $3m = s(s(s(0))) \times m = s(s(0)) \times m + m = s(0) \times m + m + m = m + m + m$.

Ainsi, nm est défini, moins formellement, par la propriété $nm = \underbrace{m + \dots + m}_{n \text{ fois}}$.

Propriétés de la multiplication :

- 0 est absorbant : $\forall m \in \mathbb{N}, m \times 0 = 0 \times m = 0$.
- 1 est neutre : $\forall m \in \mathbb{N}, m \times 1 = 1 \times m = m$.
- Distributivité de la multiplication par rapport à l'addition :
 $\forall n, m, p \in \mathbb{N}, n(m+p) = (nm) + (np) = nm + np$: les dernières parenthèses sont inutiles si l'on convient que la multiplication est prioritaire devant l'addition.
- Associativité : $\forall n, m, k \in \mathbb{N}, (n \times m) \times k = n \times (m \times k)$.
- Commutativité : $\forall n, m \in \mathbb{N}, n \times m = m \times n$.

Démonstration.

Exercice. $\square$

Propriété. Pour tout $a, b, c \in \mathbb{N}$, si $a \leq b$, alors $c(b - a) = cb - ca$.

Démonstration.

Posons $r = b - a$. On sait que $b = a + r$, donc $cb = ca + cr$. Ainsi, $ca \leq cb$ et $cb - ca = cr = c(b - a)$. $\square$

La relation d'ordre est compatible avec la multiplication :

Pour tout $a, b, c, d \in \mathbb{N}$, si $a \leq b$ et $c \leq d$, alors $ac \leq bd$.

Démonstration.

Supposons que $a \leq b$ et $c \leq d$. Il existe $e \in \mathbb{N}$ tels que $b = a + e$. Alors $bc = ac + ec \geq ac$. De même, on montre que $db \geq cb$ et on conclut par transitivité. $\square$

Propriété. Soit $n, k \in \mathbb{N}$.

Si $nk = 0$, alors $n = 0$ ou $k = 0$.

Si $nk = 1$, alors $n = k = 1$.

Démonstration.

◇ Raisonnons par contraposition : supposons que $n \neq 0$ et $k \neq 0$.

Alors $n > 0$, donc $n \geq 1$. De même $k \geq 1$, puis $kn \geq 1 \times 1 = 1$ et $kn \neq 0$.

◇ Supposons que $nk = 1$.

$n \neq 0$ et $k \neq 0$, donc $n \geq 1$ et $k \geq 1$.

Si $n \neq 1$, alors $n \geq 2$ puis $nk \geq 2 > 1$. C'est impossible. $\square$

Définition. Soit $n, m \in \mathbb{N}$. On dit que n divise m , que n est un diviseur de m , ou encore que m est un multiple de n si et seulement si il existe $k \in \mathbb{N}$ tel que $m = kn$. On note $n|m$.

Remarque. Tout entier divise 0 mais 0 ne divise que lui-même.

Définition. On appelle nombre premier tout entier n supérieur à 2 dont les seuls diviseurs sont 1 et n .

Propriété. La relation de divisibilité est une relation d'ordre partielle sur $\mathbb{N}$.

Démonstration.

Exercice. $\square$

3.5 Maximum et minimum dans $\mathbb{N}$

Propriété. Toute partie non vide et majorée de $\mathbb{N}$ possède un maximum.

Démonstration.

Notons $R(n)$ la propriété suivante : toute partie non vide de $\mathbb{N}$ majorée par n possède un plus grand élément.

◇ Pour $n = 0$, l'unique partie non vide de $\mathbb{N}$ majorée par 0 est $\{0\}$. Elle possède bien un plus grand élément.

◇ Soit $n \geq 1$. Supposons $R(n - 1)$.

Soit A une partie non vide de $\mathbb{N}$ majorée par n .

Premier cas : Si $n \notin A$, alors A est majorée par $n - 1$ et on utilise $R(n - 1)$.

Deuxième cas : Si $n \in A$, n est le plus grand élément de A . $\square$

Remarques :

— La réciproque est vraie.

— $\mathbb{N}$ n'est pas majoré.

Propriété. Soit $a, b \in \mathbb{N}$ avec $b \neq 0$. Il existe un unique couple $(q, r) \in \mathbb{N}^2$ tel que $a = bq + r$ et $0 \leq r < b$. On dit que q et r sont le quotient et le reste de la division euclidienne de a par b .

Exemple. $27 = 6 \times 4 + 3 = 6 \times 3 + 9 = 6 \times 5 - 3$, mais seul 3 est le reste de la division euclidienne de 27 par 6.

$31 = 10 \times 3 + 1$ et $-31 = (-11) \times 3 + 2$, donc, une fois que nous aurons prolongé la division euclidienne sur $\mathbb{Z}$, le reste de la division euclidienne de -31 par 3 n'est pas l'opposé du reste de 31 par 3.

Démonstration.

◇ Posons $A = \{k \in \mathbb{N} / bk \leq a\}$.

$b \geq 1$, donc pour tout $k \in A$, $k \leq bk \leq a$. Ainsi, A est majorée, or elle est non vide car $0 \in A$, donc A possède un maximum noté q . Par construction, $bq \leq a < b(q + 1)$.

Ainsi, il existe $r \in \mathbb{N}$ tel que $a = bq + r$ et $r = a - bq < b$. Ceci prouve l'existence.

◇ Pour montrer l'unicité, supposons de plus que $a = bq' + r'$ avec $q', r' \in \mathbb{N}$ et $0 \leq r' < b$. Supposons que $r \geq r'$ (la démonstration s'adapte lorsque $r' \geq r$). Il existe $r'' \in \mathbb{N}$ tel que $r = r' + r''$. $r = r' + r''$, donc $bq' + r' = a = bq + r = bq + r' + r''$. Par régularité de r' , $bq' = bq + r''$, donc $bq' \geq bq$. $b \geq 1$, donc $q' \geq q$. Ainsi, $r'' = bq' - bq = b(q' - q)$. Supposons que $q \neq q'$. Alors $q' - q \geq 1$, donc $r'' \geq b$ puis $r \geq b + r' \geq b$ ce qui est faux. Ainsi, $q = q'$, puis $r = r'$. □

Exercice. Diviser 3378 par 53.

Solution : $53 \times 6 = 318$, donc $337 = 53 \times 6 + 19$. Ainsi, $3378 = 53 \times 60 + 198$. Mais $198 = 3 \times 53 + 39$, donc $3378 = 53 \times 63 + 39$.

On présente en général ce calcul sous la forme suivante :

$$\begin{array}{r} 3 \ 3 \ 7 \ 8 \mid 5 \ 3 \\ - \ 3 \ 1 \ 8 \mid 6 \\ \hline 1 \ 9 \mid \\ \\ 3 \ 3 \ 7 \ 8 \mid 5 \ 3 \\ - \ 3 \ 1 \ 8 \mid 6 \ 3 \\ \hline 1 \ 9 \ 8 \mid \\ - \ 1 \ 5 \ 9 \mid \\ \hline 3 \ 9 \mid \end{array}$$

Propriété. Toute partie non vide de $\mathbb{N}$ possède un minimum.

Démonstration.

Soit A une partie non vide de $\mathbb{N}$.

Notons M l'ensemble des minorants de A .

$0 \in M$, donc M est non vide. Elle est majorée par n'importe quel élément de A .

D'après la propriété précédente, M possède un plus grand élément, que l'on notera m . Si $m \notin A$, pour tout $n \in A$, $n > m$, donc $n \geq m + 1$. Alors $m + 1$ est encore un minorant de A , donc il est plus petit que m , ce qui est faux. Ainsi $m \in A$. C'est le plus petit élément de A . □

Remarque. Un ensemble ordonné dont toute partie non vide possède un plus petit élément est appelé un ensemble bien ordonné.

Paradoxe de Berry¹⁰ : Notons n le plus petit entier qui n'est pas définissable en moins de 100 mots.

L'ensemble des entiers définissables en moins de 100 mots est une partie finie de $\mathbb{N}$, car le nombre de mots de la langue française est fini, donc son complémentaire dans $\mathbb{N}$ est une partie non vide de $\mathbb{N}$, qui possède bien un plus petit élément. Ainsi, la définition de n est correcte.

Cependant c'est une définition de n utilisant strictement moins de 100 mots, alors que par définition, n n'est pas définissable en moins de 100 mots !

Les logiciens du début du vingtième siècle ont compris que la solution de ce paradoxe nécessite de distinguer le langage utilisé au sein d'une théorie du métalangage utilisé

10. Personnage fictif inventé par Russell.

pour parler de cette théorie. Ainsi, dans la première phrase, “définissable” signifie “que l’on peut définir dans le cadre du langage de l’arithmétique”. Cette même phrase définit bien l’entier n , mais cette définition utilise un métalangage décrivant des phrases du langage arithmétique. Le paradoxe est ainsi levé; la phrase “notons n le plus petit entier qui n’est pas définissable en moins de 100 mots selon le langage de la théorie arithmétique” est une définition de n selon le métalangage de cette même théorie.

Remarque. Le fait que $\mathbb{N}$ est bien ordonné se démontre essentiellement à partir du principe de récurrence, c’est-à-dire à partir du dernier axiome de Peano. Mais on peut également démontrer la réciproque : remplaçons le dernier axiome de Peano par le fait que $\mathbb{N}$ est bien ordonné et démontrons le principe de récurrence.

Soit F une partie de $\mathbb{N}$ contenant 0 et telle que pour tout $n \in F$, $s(n) \in F$. Il faut montrer que $F = \mathbb{N}$.

Sinon, $\mathbb{N} \setminus F$ est une partie non vide de $\mathbb{N}$, donc elle possède un plus petit élément que l’on notera m . $m \neq 0$ car $0 \in F$.

Il faut de plus remplacer le second axiome par l’énoncé plus fort du lemme de la page 19. Ainsi il existe $n \in \mathbb{N}$ tel que $m = s(n)$.

$n < m$ donc par définition de m , $n \in F$. Mais alors $m = s(n) \in F$, ce qui est faux.

Ainsi, le principe de récurrence est équivalent au fait que $\mathbb{N}$ est bien ordonné.

Il importe de retenir de cette remarque que lorsqu’une démonstration par récurrence semble délicate, il peut être plus simple d’utiliser directement que $\mathbb{N}$ est bien ordonné.

Exercice. En adaptant cette démonstration, montrer le principe de récurrence finie.

Principe de la descente infinie : plus précisément, pour montrer que “ $\forall n \in \mathbb{N}, R(n)$ ”, une alternative à la récurrence est de raisonner par l’absurde en supposant qu’il existe $n \in \mathbb{N}$ tel que $\neg[R(n)]$ et d’en déduire l’existence d’un entier naturel m tel que $m < n$ et tel que $\neg[R(m)]$. Cela permet de construire une suite strictement décroissante d’entiers naturels n_k vérifiant $\neg[R(n_k)]$, ce qui est absurde. Mais pourquoi au fait ?

Il est préférable de remplacer cet argument de descente infinie par le suivant :

On note $F = \{n \in \mathbb{N} / \neg R(n)\}$. Il est non vide par hypothèse, donc il possède un minimum n_0 , mais il existe alors $m < n_0$ tel que $m \in F$, ce qui est impossible.

Exemple. C’est ainsi que Gauss¹¹ démontre le lemme d’Euclide¹² (dans le cours

11. Johann Carl Friedrich Gauss, 1777-1855, est un mathématicien, astronome et physicien allemand. Il a apporté de très importantes contributions à ces trois domaines. Surnommé “le prince des mathématiciens”, il est considéré comme l’un des plus grands mathématiciens de tous les temps.

Refusant de publier un travail qu’il ne considérait pas au-dessus de toute critique, son journal montre qu’il avait fait plusieurs importantes découvertes mathématiques des années, voire des décennies, avant qu’elles ne soient publiées par ses contemporains.

Il rechignait à présenter l’intuition derrière ses très élégantes démonstrations. Il préférait qu’elles apparaissent comme sorties de nulle part et effaçait toute trace du processus de sa découverte, “de même qu’un architecte ne laisse pas l’échafaudage une fois l’édifice achevé”.

12. Euclide est un mathématicien de la Grèce antique. Il est vraisemblable qu’il ait vécu vers 300 avant notre ère. Son ouvrage principal, “les Éléments” aborde la géométrie (euclidienne) et l’arithmétique selon une démarche axiomatique et déductive.

d'arithmétique, on démontrera autrement un théorème plus général, appelé le lemme de ... Gauss) :

Soit $a, b, p \in \mathbb{N}$ tels que p est premier et $p|ab$. Il s'agit de montrer que $[p|a] \vee [p|b]$.

Pour cela, fixons p premier et $a \in \mathbb{N}$ tel que $p \nmid a$.

Il s'agit alors de montrer que pour tout $b \in \mathbb{N}$, si $p|ab$, alors $p|b$.

Pour cela, raisonnons par l'absurde en supposant que l'ensemble

$B = \{b \in \mathbb{N} / p|ab \text{ et } p \nmid b\}$ est non vide. Alors il possède un minimum noté b_0 .

Si $b_0 \geq p$, alors $b_0 - p \in B$ ce qui contredit la minimalité de b_0 , donc $0 < b_0 < p$.

On peut écrire la division euclidienne de p par b_0 : il existe $q, b_1 \in \mathbb{N}$ tels que $p = qb_0 + b_1$ avec $0 \leq b_1 < b_0$.

Si $b_1 = 0$, alors $p = qb_0$, mais p est premier, donc $b_0 = 1$, or $p|ab_0$, donc $p|a$ ce qui est faux. Ainsi, $b_1 \neq 0$, donc $1 \leq b_1 < p$ puis $p \nmid b_1$.

$ab_1 = ap - aqb_0$, donc $p|ab_1$. Ainsi $b_1 \in B$ et $b_1 < b_0$, ce qui est faux.

Remarque. Pour démontrer une propriété de la forme " $\forall n \in \mathbb{N}, R(n)$ ", on ne raisonne cependant pas a priori par récurrence. C'est seulement lorsqu'on peut relier $R(n+1)$ à $R(n)$ que c'est une méthode envisageable.

3.6 Relations d'équivalence

Définition. Une relation binaire sur un ensemble E est une relation d'équivalence si et seulement si R est réflexive, symétrique et transitive.

Exemple. Les entiers divisibles par 2 sont dits pairs et les autres impairs.

On note P la relation binaire sur $\mathbb{N}$ définie par : nPm si et seulement si n et m ont la même parité. P est une relation d'équivalence.

Définition. Soit R une relation d'équivalence sur E .

Si $x \in E$, on note $\bar{x}$ l'ensemble des $y \in E$ tels que xRy .

$\bar{x}$ s'appelle la classe d'équivalence de x .

On désigne par E/R l'ensemble des classes d'équivalence : $E/R = \{\bar{x} / x \in E\}$.

E/R s'appelle l'ensemble quotient de E par R .

Exemple fondamental : Soit f une application de E dans F . En convenant, pour tout $x, y \in E$, que $x R y \iff f(x) = f(y)$, on définit sur E une relation d'équivalence.

Exemples :

- La relation d'égalité est l'unique relation d'équivalence dont les classes d'équivalence sont toutes des singletons. On a $(E / =) = \{\{x\} / x \in E\}$.
- La relation de parité P possède deux classes d'équivalence : $\bar{0}$ est l'ensemble des entiers pairs et $\bar{1}$ est l'ensemble des entiers impairs. Ainsi, $\mathbb{N} / P = \{\bar{0}, \bar{1}\}$.
- Sur l'ensemble des formules propositionnelles, la relation "être logiquement équivalente à" est une relation d'équivalence.
- Si E est l'ensemble des droites du plan, la relation de parallélisme est une relation d'équivalence dont les classes d'équivalence sont les directions du plan.

Propriété. Avec les hypothèses et notations précédentes, pour tout $x, y \in E$, $xRy \iff \bar{x} = \bar{y}$.

Démonstration.

Supposons que xRy . Soit $z \in \bar{x}$. On a xRz , donc yRz puis $z \in \bar{y}$.

Ainsi, $\bar{x} \subset \bar{y}$. Par symétrie, $\bar{y} \subset \bar{x}$, donc $\bar{x} = \bar{y}$.

Réciproquement, supposons que $\bar{x} = \bar{y}$. xRx , donc $x \in \bar{x}$. Ainsi $x \in \bar{y}$ et xRy . $\square$

Définition. Une partition $\mathcal{P}$ de E est une partie de $\mathcal{P}(E)$ telle que :

- pour tout $A, B \in \mathcal{P}$, $A \neq B \implies A \cap B = \emptyset$,
- pour tout $A \in \mathcal{P}$, $A \neq \emptyset$,
- et $\bigcup_{A \in \mathcal{P}} A = E$.

Exemple. Si $E = \{1, \dots, 10\}$, un exemple de partition de E est $\mathcal{P} = \{\{1\}, \{2, 3, 7, 10\}, \{4, 9\}, \{5, 6, 8\}\}$.

Théorème. Si R est une relation d'équivalence sur E , son ensemble quotient E/R est une partition de E .

Réciproquement, si $\mathcal{P}$ est une partition de E , il existe une unique relation d'équivalence R sur E telle que $\mathcal{P} = E/R$.

Elle est définie par : $\forall x, y \in E$, $[xRy \iff (\exists C \in \mathcal{P}, x, y \in C)]$.

En résumé, la donnée d'une relation d'équivalence sur E est équivalente à la donnée d'une partition de E .

Remarque. On peut formaliser un peu plus : si l'on note $\mathcal{R}$ l'ensemble des relations d'équivalence sur E et $\mathbb{P}$ l'ensemble des partitions de E , alors le théorème précédent énonce que l'application

$$\begin{array}{ccc} \mathcal{R} & \longrightarrow & \mathbb{P} \\ R & \longmapsto & E/R \end{array}$$

est une bijection.

Démonstration.

- Supposons que R est une relation d'équivalence.
 - Pour tout $x \in E$, $x \in \bar{x}$, donc $\bar{x} \neq \emptyset$.
 - Soit $x, y \in E$ tels que $\bar{x} \cap \bar{y} \neq \emptyset$. Il existe $z \in E$ tel que $z \in \bar{x} \cap \bar{y}$. On a xRz et zRy , donc xRy puis $\bar{x} = \bar{y}$.
 - Pour tout $x \in E$, $x \in \bar{x} \subset \bigcup_{x \in E} \bar{x}$, donc $E \subset \bigcup_{x \in E} \bar{x}$. L'inclusion réciproque est claire car, pour tout $x \in E$, $\bar{x} \subset E$.

Ainsi E/R est bien une partition de E .

- Réciproquement, soit $\mathcal{P}$ une partition de E . Raisonnons par analyse-synthèse.

Analyse : Supposons qu'il existe une relation d'équivalence R telle que $E/R = \mathcal{P}$.

Soit $x, y \in E$. Si xRy , alors $x, y \in \bar{x}$, donc il existe $C \in \mathcal{P}$ tel que $x, y \in C$.

Réciproquement, supposons qu'il existe $C \in \mathcal{P}$ tel que $x, y \in C$. Sachant que $E/R = \mathcal{P}$, il existe $z \in E$ tel que $C = \bar{z}$.

Alors $\bar{x} = \bar{z} = \bar{y}$, donc xRy . On a donc montré que

(1) : $\forall x, y \in E$, $[xRy \iff (\exists C \in \mathcal{P}, x, y \in C)]$.

Ainsi, sous condition d'existence, il existe une unique relation d'équivalence R telle que $E/R = \mathcal{P}$. Elle est définie par (1).

Synthèse : Notons R la relation binaire définie par (1). Elle est clairement symétrique. Soit $x \in E$. $E = \bigcup_{C \in \mathcal{P}} C$, donc il existe $C \in \mathcal{P}$ tel que $x \in C$. Ainsi, xRx . R est donc réflexive.

Soit $x, y, z \in E$ tels que xRy et yRz . Il existe $C, D \in \mathcal{P}$ tels que $x, y \in C$ et $y, z \in D$. En particulier $y \in C \cap D$, donc $C \cap D \neq \emptyset$, donc $C = D$. Alors $x, z \in C$, donc xRz . On a montré que R est transitive.

Ainsi R est une relation d'équivalence. Pour terminer, il reste à montrer que $E/R = \mathcal{P}$.

◇ Commençons par montrer que, si $x \in C \in \mathcal{P}$, alors $\bar{x} = C$:

Si $y \in \bar{x}$, il existe $D \in \mathcal{P}$ tel que $x, y \in D$. Alors $x \in C \cap D \neq \emptyset$, donc $C = D$ et $y \in C$. Ainsi $\bar{x} \subset C$. Réciproquement, si $y \in C$, alors $x, y \in C$, donc xRy puis $y \in \bar{x}$. On a bien montré que $\bar{x} = C$.

◇ Soit $x \in E$. Il existe $C \in \mathcal{P}$ tel que $x \in C$. On a $x \in C \in \mathcal{P}$, donc $\bar{x} = C \in \mathcal{P}$. DONC $E/R \subset \mathcal{P}$.

◇ Réciproquement, soit $C \in \mathcal{P}$. $C \neq \emptyset$, donc il existe $x \in E$ tel que $x \in C$.

On a encore $x \in C \in \mathcal{P}$, donc $C = \bar{x} \in E/R$. DONC $\mathcal{P} \subset E/R$.

Finalement $E/R = \mathcal{P}$, ce qui clôt la démonstration. □

Remarque. Informellement, les éléments de E/R sont les éléments de E , mais en acceptant d'identifier deux éléments x et y de E lorsque xRy , c'est-à-dire en faisant abstraction¹³ de certaines caractéristiques des éléments de E .

Nous allons voir que c'est un procédé puissant pour élaborer des constructions mathématiques.

4 La logique mathématique

Ce chapitre est informel, il est seulement destiné à votre culture personnelle¹⁴.

Définition. Un langage (du premier ordre) est un ensemble de mots (appelés des formules), se conformant à des règles syntaxiques non précisées ici, dont les lettres appartiennent à l'alphabet $\{ (,), \neg, \wedge, \vee, \implies, \iff, \forall, \exists \} \cup C \cup V \cup F \cup R$, où

- C est un ensemble de constantes.
- V est l'ensemble des variables : x, y, z etc.
- F est un ensemble de fonctions : si $f \in F$ et si f possède n arguments (on dit que f est d'arité n), on peut par exemple considérer $f(x_1, \dots, x_n)$ où les x_i sont dans V , mais aussi $f(h_1(x_1), h_2(x_2, \dots, x_p), \dots, h_n(x_1, \dots, x_k))$ où $h_1, \dots, h_n$ sont d'autres fonctions d'arités convenables.
- R est un ensemble de relations : si $r \in R$, r est un prédicat dépendant d'un nombre variable d'arguments.

13. Voici la définition du mot "abstraction" proposée par le Larousse : Opération intellectuelle qui consiste à isoler par la pensée l'un des caractères de quelque chose et à le considérer indépendamment des autres caractères de l'objet.

14. Pour plus de détails, on pourra consulter *Logique mathématique*, Tome 1, de René Cori et Daniel Lascar

Exemple. Le langage de l'arithmétique est donné (par exemple) par $C = \{0\}$, $F = \{s, +, \times\}$, $R = \{=, \leq\}$.

Voici un élément de ce langage, en posant pour simplifier $1 = s(0)$ et $2 = s(1)$:

$$(2 \leq p) \wedge (\forall k, \forall h((\neg(h = 1) \wedge \neg(k = 1)) \implies \neg(h \times k = p))).$$

Cette formule exprime que p est un nombre premier. Dans cette formule, p est une variable libre alors que k et h sont des variables liées.

Définition. Une formule est close lorsqu'elle ne contient aucune variable libre.

- Une théorie T d'un langage L est un ensemble de formules closes de L .
- Les axiomes logiques de L sont des formules de L qui sont “logiquement vraies”, comme toutes les tautologies (par exemple en arithmétique $(0 \leq k) \iff (0 \leq k)$), la loi de Morgan (pour toute formule F , $\exists v F \iff \neg \forall v \neg F$), etc. (il s'agit d'une présentation informelle).
- On utilise deux règles de déduction :
 - Le modus ponens : à partir des deux formules F et $F \implies G$, on peut déduire la formule G .
 - La règle de généralisation : à partir de la formule F et d'une variable $v \in V$, on peut déduire la formule $\forall v F$.
- Une démonstration formelle d'une formule F de L dans une théorie T est une suite finie de formules $F_0, \dots, F_n$ telle que $F_n = F$, et pour tout $i \in \{0, \dots, n\}$,
 - $F_i \in T$, ou bien
 - F_i est un axiome logique, ou bien
 - F_i se déduit d'une ou de deux formules F_k avec $k < i$ par l'une des deux règles de déduction.

Exemple de la théorie des ensembles¹⁵ :

Le langage correspondant est défini par : $C = \{\emptyset\}$, $F = \emptyset$, $R = \{=, \in, \subset\}$.

La théorie ZF de Zermelo-Fraenkel est constituée des axiomes suivants

- L'axiome d'extensionnalité : $\forall x \forall y [\forall z (z \in x \iff z \in y) \implies x = y]$.
- L'axiome de la réunion, affirmant que pour tout ensemble x , il existe un ensemble y dont les éléments sont les éléments des éléments de x : $\forall x \exists y \forall z [z \in y \iff \exists t (t \in x \wedge z \in t)]$.
- L'axiome de l'ensemble des parties, qui affirme que pour tout ensemble a , il existe un ensemble b dont les éléments sont les parties de a .
- L'axiome de l'infini, cf page 2.
- L'axiome de substitution, qui donne un procédé de construction d'ensembles, similaire mais plus précis que la “définition par compréhension”, vue page 2. En gros, cet axiome dit que, pour toute formule $P(x)$ construite à partir des lettres “ $(,), \neg, \wedge, \vee, \forall, \exists, x \in \mathcal{V}, \emptyset, =, \in, \subset$ ”, si A est un ensemble alors $\{x \in A / P(x)\}$ est aussi un ensemble.

Une telle précaution évite l'apparition du paradoxe de Russell et de ses variantes : la collection de tous les ensembles n'est pas un ensemble.

15. Pour plus de détails, on pourra consulter *Théorie des ensembles*, de Jean-Louis Krivine.

Définition. Une théorie T d'un langage L est non contradictoire (on dit aussi consistante ou cohérente) si et seulement si il n'existe aucune formule F telle que F et $\neg F$ soient toutes deux démontrables à partir de T .

Une théorie T d'un langage L est complète si et seulement si pour toute formule F de L , F ou $\neg F$ est démontrable à partir de T .

Théorème d'incomplétude de Gödel¹⁶ : Soit T une théorie d'un langage L . On suppose que T "contient ZF".

Si T est non contradictoire, alors la propriété " T est non contradictoire" peut être écrite comme une formule du langage L , et cette formule, qui est vraie, n'est pas démontrable à partir de T .

En résumé, une théorie ne peut pas démontrer sa propre consistance.

En particulier, toute théorie non contradictoire contenant ZF n'est pas complète.

Remarque. Le fait de supposer que T est non contradictoire est indispensable. En effet, une théorie contradictoire serait tout à fait capable de démontrer qu'elle ne l'est pas, car une théorie contradictoire permet de montrer $F \wedge (\neg F)$ pour une certaine formule, mais comme $Faux \implies G$ pour toute formule G , une théorie contradictoire est capable de démontrer que toute formule de son langage est vraie (et fausse).

Axiome du choix : en voici deux énoncés équivalents.

- Pour tout ensemble I , pour toute famille $(E_i)_{i \in I}$ d'ensembles tous non vides, il existe une famille $(x_i)_{i \in I}$ telle que, pour tout $i \in I$, $x_i \in E_i$.
- Pour tout ensemble E , pour toute relation d'équivalence sur E , il existe un ensemble R tel que l'intersection de R avec chaque classe d'équivalence est un singleton. Cela signifie qu'on peut choisir dans chaque classe d'équivalence un représentant et considérer l'ensemble de ces représentants.

Théorème d'indépendance de l'axiome du choix :

Notons AC l'axiome du choix.

Si ZF est consistante, alors ZF+AC est aussi consistante (Gödel, 1938).

Si ZF est consistante, alors ZF+non(AC) est aussi consistante (Cohen¹⁷, 1963).

On a donc le *choix* d'accepter ou de réfuter l'axiome du *choix*.

Par défaut, en mathématiques modernes, on travaille avec ZF+AC, même si cela conduit à quelques résultats étranges :

Paradoxe de Banach-Tarski (1924) :

Notons S la boule unité de l'espace usuel : $S = \{(x, y, z) \in \mathbb{R}^3 / x^2 + y^2 + z^2 \leq 1\}$.

Il existe une partition de S en un nombre fini de parties $D_1, \dots, D_n$ telles que, en soumettant chaque partie à un déplacement (composition d'une rotation et d'une translation) approprié dans l'espace, leur réunion après déplacements est égale à la réunion disjointe de deux sphères de rayon 1.

16. Kurt Gödel, 1906-1978, logicien et mathématicien autrichien naturalisé américain. Enfant, sa famille l'avait surnommé "monsieur Pourquoi".

17. Paul Cohen, 1934-2007, est un mathématicien américain, médaille Fields 1966

Les parties D_i ne sont pas “mesurables”, ce qui résout le paradoxe.

Propriété. (admise) : Sous l'hypothèse ZF , l'axiome du choix est équivalent à l'axiome de Zermelo ainsi qu'à celui de Zorn, où :

- Axiome de Zermelo : Tout ensemble E peut-être muni d'un bon ordre.
- Axiome de Zorn : Un ensemble ordonné $(E, \leq)$ possède un élément maximal dès que toutes ses chaînes sont majorées, en convenant qu'une chaîne de E désigne une partie de E totalement ordonnée par $\leq$.

5 L'art de la démonstration

La structure d'une démonstration se construit avant tout en fonction de la structure de la propriété à démontrer. En conséquence, on regarde d'abord la cible à atteindre et seulement lorsque c'est nécessaire les hypothèses dont on dispose pour y parvenir. On ne sait pas a priori sous quelles formes ces hypothèses seront utilisées.

Lorsqu'une telle démarche n'aboutit pas, on essaie plus généralement de décomposer le résultat à atteindre et de traduire les hypothèses pour les rapprocher, par exemple en essayant de les écrire dans un langage commun.

Dans ce chapitre, on notera R le résultat à démontrer.

5.1 Conjonction et disjonction

◇ Lorsque $R = P \wedge Q$, bien entendu, on montre le plus souvent P et Q .

Dans ce cas, il importe de commencer par la propriété qui semble la plus simple, afin de se familiariser en douceur avec le problème à résoudre.

◇ Lorsque $R = P \vee Q$, on peut supposer que P est fausse et démontrer Q , ou bien supposer que Q est fausse et montrer P .

On ajoute ainsi une hypothèse. Il importe de bien choisir parmi ces deux possibilités.

5.2 Démonstration par disjonction de cas

Pour démontrer une propriété dépendant de certains paramètres, on peut être amené à étudier plusieurs cas selon les valeurs de ces paramètres. Il importe que la réunion des différents cas étudiés recouvre tout l'ensemble des valeurs des paramètres.

Exemple. Soit $a, b \in \mathbb{R}$. Résoudre le système $(S) : \begin{cases} x + (4 - a)y = 0 \\ (1 - a)x - 2y = b \end{cases}$.

Ici, la question est ouverte en ce sens que le résultat à obtenir n'est pas explicitement donné. On recherche l'ensemble des couples (x, y) de $\mathbb{R}^2$ (a priori) vérifiant les deux égalités de (S) .

Solution :

$$(S) \iff \begin{cases} x = (a - 4)y \\ (1 - a)(a - 4)y - 2y = b \end{cases} \iff \begin{cases} x = (a - 4)y \\ (-a^2 + 5a - 6)y = b \end{cases}$$

Le discriminant de l'équation $t^2 - 5t + 6 = 0$ vaut $\Delta = 25 - 24 = 1$, donc les racines sont $\frac{5 \pm 1}{2}$, soit 2 et 3.

◇ *Premier cas* : on suppose que $a \notin \{2, 3\}$.

$(S) \iff \begin{cases} y = \frac{b}{-a^2 + 5a - 6} \\ x = \frac{b(a-4)}{-a^2 + 5a - 6} \end{cases}$. Ce système possède alors une unique solution dans $\mathbb{R}^2$.

Le fait d'avoir raisonné par équivalence nous affranchit de vérifier que le couple obtenu est bien solution.

◇ *Second cas* : on suppose que $a \in \{2, 3\}$.

Alors $(S) \iff \begin{cases} x = (a-4)y \\ b = 0 \end{cases}$.

- *Premier sous-cas* : si $b \neq 0$, alors (S) ne possède aucune solution.
- *Second sous-cas* : on suppose que $b = 0$. Alors $(S) \iff x = (a-4)y$. L'ensemble des solutions est une droite de $\mathbb{R}^2$, d'équation $x+2y = 0$ lorsque $a = 2$ et $x+y = 0$ lorsque $a = 3$.

5.3 Résoudre une équation

On vient de résoudre un système d'équations. C'est un cas particulier d'équation, dont voici la définition la plus générale possible :

Définition. Si P est un prédicat sur un ensemble E , “résoudre l'équation $P(x)$, en l'inconnue $x \in E$ ”, c'est calculer $\{x \in E/P(x)\}$ qu'on appelle alors l'ensemble des solutions de l'équation.

“calculer” signifie “donner l'ensemble des solutions sous la forme la plus simple possible”.

Remarque. La plupart des équations sont de la forme “ $f(x) = g(x)$ ”, où f et g sont deux applications de E dans un autre ensemble F .

Lorsque $F = \mathbb{R}$, on rencontre parfois des équations de la forme “ $f(x) \leq g(x)$ ”, ou “ $f(x) < g(x)$ ”. Dans ce cas, on parle plutôt d'*inéquations*.

Méthode :

- Précisez d'abord pour quelles valeurs $x \in E$ l'équation a bien un sens. Par exemple, pour une équation de la forme “ $f(x) = g(x)$ ”, il faudra d'abord rechercher les domaines de définition de f et de g .
- Autant que possible, raisonnez par équivalence comme dans l'exemple précédent. Cependant le fait de raisonner par équivalence impose parfois trop de lourdeur à la rédaction. Lorsqu'on choisit de raisonner par implications, après avoir montré que $P(x) \implies x \in S$, pour une certaine partie S de E , il restera à rechercher quels sont les éléments de S qui sont effectivement solutions.

Exemple. Résoudre dans $\mathbb{R}$ l'équation du second degré générale $(E) : x^2 + ax + b = 0$.

Solution : Soit $x \in \mathbb{R}$.

$$(E) \iff \left(x + \frac{a}{2}\right)^2 + b - \frac{a^2}{4} = 0 \iff \left(x + \frac{a}{2}\right)^2 = \frac{a^2 - 4b}{4}.$$

Posons $\Delta = a^2 - 4b$: c'est le discriminant.

Premier cas : Si $\Delta < 0$, l'ensemble des solutions est vide.

Second cas : On suppose que $\Delta \geq 0$. La fonction racine carrée est ici supposée connue.

On pose $\delta = \sqrt{\Delta}$. Alors $(E) \iff \exists \varepsilon \in \{-1, 1\}, x + \frac{a}{2} = \varepsilon \frac{\delta}{2}$, donc les solutions de (E)

sont exactement $\frac{-a \pm \sqrt{\Delta}}{2}$.

5.4 Implication

◇ Lorsque $R = [P \implies Q]$, on suppose que P est vraie (hypothèse supplémentaire) et on démontre Q .

Raisonnement par contraposition : l'implication $P \implies Q$ est logiquement équivalente à $(\neg Q) \implies (\neg P)$, qui est appelée sa contraposée. Ainsi,

pour démontrer $P \implies Q$, on peut raisonner par contraposition, c'est-à-dire démontrer $(\neg Q) \implies (\neg P)$: on suppose que Q est fausse et on démontre que P est fausse.

Exemple. Montrer que pour tout entier $n \in \mathbb{N}$, si n^2 est pair, alors n est aussi pair.

Solution : Soit $n \in \mathbb{N}$. Raisonnons par contraposition en supposant que n est impair. Il existe $k \in \mathbb{N}$ tel que $n = 2k + 1$. Alors $n^2 = 4k^2 + 4k + 1$ est impair.

Le raisonnement par l'absurde : cela consiste à supposer que R est fausse et à aboutir à une contradiction, souvent de la forme $S \wedge (\neg S)$.

Le raisonnement par l'absurde repose sur le postulat que la théorie dans laquelle on travaille est consistante. Il est donc impossible d'aboutir à une contradiction ce qui prouve que R est vraie.

Remarque. Les raisonnements par contraposition et par l'absurde sont très voisins.

D'une part, si l'on parvient à montrer R par l'absurde, on a montré que

$\neg R \implies S \wedge (\neg S)$, dont la contraposée est $S \vee (\neg S) \implies R$, mais $S \vee (\neg S)$ est vraie, donc R est prouvée. Ainsi, démontrer R par l'absurde revient à démontrer $S \vee (\neg S) \implies R$ par contraposition.

D'autre part, en notant H l'ensemble des hypothèses, on souhaite démontrer $H \implies R$.

Le raisonnement par contraposition consiste à montrer $\neg R \implies \neg H$, mais on peut le maquiller en un raisonnement par l'absurde : on suppose H et on veut démontrer R . Pour cela on raisonne par l'absurde en supposant $\neg R$ et on en déduit $\neg H$. On a donc $H \wedge (\neg H)$ ce qui est contradictoire, donc R est vraie.

Un tel maquillage est maladroit, il produit une démonstration inutilement alambiquée.

Il y a cependant une réelle différence entre ces deux raisonnements ; lorsqu'on raisonne vraiment par l'absurde, la propriété S est quelconque, on en cherche seulement une qui convient. Lorsqu'on raisonne par contraposition, la propriété H est donnée.

◇ Lorsque $R = [P \iff Q]$, on regarde souvent R comme la conjonction

$[P \implies Q] \wedge [Q \implies P]$ à laquelle on applique les méthodes précédentes.

◇ Lorsque R est “montrer que les propriétés $P_1, \dots, P_k$ sont équivalentes”, on peut se contenter de montrer le cycle d’implications $P_1 \implies P_2 \implies \dots \implies P_k \implies P_1$. Mais la liste $P_1, \dots, P_k$ n’est pas toujours donnée dans l’ordre idéal. Il convient donc parfois de la réordonner.

5.5 Quantificateurs

◇ Lorsque $R = [\forall x \in E, P(x)]$, le plus souvent, on prend x quelconque dans E , en écrivant “soit $x \in E$ ”, puis on démontre $P(x)$.

La portion de phrase “soit $x \in E$ ” ne doit pas être sous-entendue. Il ne faut pas démontrer d’emblée $P(x)$ en considérant que tout le monde comprendra que x est un élément quelconque de E . D’une manière générale, toutes les variables que vous utilisez au cours d’une démonstration doivent préalablement avoir été définies sans ambiguïté.

Exemple. Soit $f : \mathbb{R} \longrightarrow \mathbb{R}$ une fonction telle que, pour tout $x, y \in \mathbb{R}$, $f(x + y) = f(x)f(y)$. Montrer que $\forall x \in \mathbb{R}, f(x) \geq 0$.

Solution : Soit $x \in \mathbb{R}$. $f(x) = f(\frac{x}{2} + \frac{x}{2}) = f(\frac{x}{2})^2 \geq 0$.

◇ Lorsque $R = [\exists x \in E, P(x)]$, la méthode directe consiste à construire un élément x de E satisfaisant $P(x)$. Cela nécessite de l’imagination et de la créativité. C’est souvent un passage délicat dans une démonstration, mais cela participe à la “beauté” des mathématiques. Dans ce contexte, on procède souvent par analyse-synthèse (cf ci-dessous).

On peut aussi raisonner par l’absurde, en supposant $\forall x \in E, \neg(P(x))$ et en recherchant une contradiction. Il faut cependant que cette nouvelle hypothèse se marie bien avec les autres hypothèses.

Exemple. Soit $x_1, x_3 \in \mathbb{Q}$ avec $x_1 < x_3$.

Montrer qu’il existe $x_2 \in \mathbb{Q}$ tel que $x_1 < x_2 < x_3$.

Solution : $x_2 = \frac{x_1 + x_3}{2}$ convient.

◇ Lorsque $R = \neg(\forall x \in E, P(x))$, on a vu que R est logiquement équivalente à $[\exists x \in E, \neg(P(x))]$, donc pour montrer R , on peut rechercher un x dans E tel que $P(x)$ est fausse. Dans ce contexte, x est appelé un contre-exemple du prédicat $P(x)$.

Par exemple pour montrer que $\mathbb{R} \neq \mathbb{Q}$, c’est-à-dire $\neg(\forall x \in \mathbb{R}, x \in \mathbb{Q})$, il suffit de contruire un réel non rationnel. On verra que $\sqrt{2}$ est un contre-exemple.

5.6 Existence et unicité

On suppose que R est de la forme $R = [\exists! x \in E, P(x)]$.

Dans de nombreux exercices et problèmes, l’énoncé d’une telle propriété se présente sous la forme : “montrer qu’il existe $x \in E$ tel que $P(x)$, puis montrer que x est unique”. Sur le plan ontologique, tout objet mathématique est unique, mais ce n’est pas du tout ce qui est demandé par l’énoncé. La propriété “ x est unique” dépend de P .

En mathématiques, l'unicité est toujours prononcée relativement à un prédicat. Par exemple, 2 est l'unique entier premier et pair, mais 2 n'est pas l'unique entier pair inférieur à 10.

◇ Pour montrer qu'il existe un unique $x \in E$ tel que $P(x)$, il est souvent préférable de séparer l'existence et l'unicité. Pour l'existence, on en a déjà parlé, pour l'unicité, il faut montrer que $\{x \in E/P(x)\}$ ne possède pas deux éléments distincts, par exemple en supposant qu'il existe $x, y \in E$ vérifiant $P(x)$ et $P(y)$ et en prouvant que $x = y$.

Mais il y a d'autres méthodes :

- On peut montrer que $\{x \in E/P(x)\}$ est un singleton.
- On peut résoudre l'équation " $P(x)$ " en l'inconnue x pour montrer qu'elle admet une seule solution.
- On peut raisonner par analyse-synthèse :

5.7 Démonstration par analyse-synthèse

On a déjà rencontré un tel raisonnement page 25, pour montrer que si $\mathcal{P}$ est une partition d'un ensemble E , alors il existe une unique relation d'équivalence R sur E telle que $\mathcal{P} = E/R$.

Ce mode de raisonnement est envisageable lorsque R est de la forme $[\exists x \in E, P(x)]$.

Il se décompose en deux parties :

◇ **L'analyse** : on suppose qu'il existe $x \in E$ tel que $P(x)$.

C'est a priori très étrange, car on suppose justement ce qu'il faut démontrer !

A partir du fait que x vérifie $x \in E$ et $P(x)$, on cherche à préciser x , en montrant que x est nécessairement de la forme $y(t)$, où t est un paramètre et où pour tout t , $y(t)$ est parfaitement définie.

Il est fréquent que l'analyse conduise à une seule valeur possible pour x .

◇ **La synthèse** : On cherche pour quel(s) t la quantité $y(t)$ vérifie $P(y(t))$.

Remarque. Le raisonnement par analyse-synthèse est bien adapté pour démontrer une propriété d'existence et d'unicité de la forme $[\exists! x \in E, P(x)]$: la partie analyse doit mener à une unique solution y possible, ce qui prouve l'unicité *sous condition d'existence* puis la partie synthèse prouve que y est bien solution, donc elle établit l'existence.

Exemple. Montrer que toute fonction de $\mathbb{R}$ dans $\mathbb{R}$ se décompose de manière unique en la somme d'une fonction paire et d'une fonction impaire.

Solution :

Soit f une fonction de $\mathbb{R}$ dans $\mathbb{R}$.

Analyse : Supposons qu'il existe deux applications p et i de $\mathbb{R}$ dans $\mathbb{R}$, respectivement paire et impaire, telles que $f = p + i$.

Pour tout $x \in \mathbb{R}$, $f(x) = p(x) + i(x)$ et $f(-x) = p(x) - i(x)$, en sommant et en retranchant ces deux égalités, $p(x) = \frac{1}{2}(f(x) + f(-x))$ et $i(x) = \frac{1}{2}(f(x) - f(-x))$.

Ainsi, sous condition d'existence, il y a unicité de l'écriture de f comme somme d'une fonction paire et d'une fonction impaire.

Synthèse : Posons pour tout $x \in \mathbb{R}$, $p(x) = \frac{1}{2}(f(x) + f(-x))$ et $i(x) = \frac{1}{2}(f(x) - f(-x))$. p et i sont des applications de $\mathbb{R}$ dans $\mathbb{R}$. On vérifie qu'elles sont respectivement paire et impaire et que, pour tout $x \in \mathbb{R}$, $p(x) + i(x) = f(x)$.

5.8 Inclusion entre ensembles

◇ Pour montrer que $A \subset B$, on peut “passer aux éléments”, c'est-à-dire montrer que $\forall x \in A, x \in B$. On commence donc par écrire “soit $x \in A$ ”, puis on cherche à démontrer que $x \in B$.

Exemple. On pose $A = \{\frac{k(k+1)}{2} / k \in \mathbb{N}\}$. Montrer que $A \subset \mathbb{N}$.

◇ Lorsque $R \iff [A = B]$, où A et B sont deux ensembles, on peut regarder R comme la conjonction de $A \subset B$ et $B \subset A$.

5.9 Démonstrations par récurrence

On a vu précédemment que :

Principe de récurrence 1 : Soit $R(n)$ un prédicat sur $\mathbb{N}$.

Si $R(0)$ est vraie et si pour tout $n \in \mathbb{N}$, $R(n)$ implique $R(n+1)$, alors pour tout $n \in \mathbb{N}$, $R(n)$ est vraie.

Remarque. On montre ainsi que : $\forall n \in \mathbb{N}$, $R(n)$. Mais bien entendu, l'assertion $R(n)$ ne commence pas par $\forall n$.

Il arrive que le prédicat $R(n)$ ne soit défini que pour $n \geq n_0$, où $n_0 \in \mathbb{N}^*$, ou bien qu'il soit faux pour des valeurs inférieures. On peut généraliser le principe de récurrence à cette situation :

Principe de récurrence 2 :

Soit $n_0 \in \mathbb{N}^*$. Soit $R(n)$ un prédicat défini pour tout entier $n \geq n_0$.

Si $R(n_0)$ est vraie et si pour tout $n \geq n_0$, $R(n)$ implique $R(n+1)$, alors pour tout $n \in \mathbb{N}$ tel que $n \geq n_0$, $R(n)$ est vraie.

Démonstration.

On pose $S(n) = R(n+n_0)$ et on applique le principe de récurrence 1 au prédicat $S(n)$.
□

Exemple. Pour $n \in \mathbb{N}$, comparer n^2 et 2^n .

On peut aussi énoncer des principes de récurrence sur $\llbracket n, m \rrbracket$:

Principe de récurrence ascendante finie : Soit $n, m \in \mathbb{N}$ avec $n \leq m$.

Soit $R(k)$ un prédicat défini pour $k \in \llbracket n, m \rrbracket$.

Si $R(n)$ est vraie et si pour tout $k \in \llbracket n, m-1 \rrbracket$, $R(k)$ implique $R(k+1)$, alors $R(k)$ est vraie pour tout $k \in \llbracket n, m \rrbracket$.

Démonstration.

Raisonnons par l'absurde en supposant que $\{k \in \llbracket n, m \rrbracket / \neg(R(k))\}$ est non vide. Alors cet ensemble d'entiers possède un plus petit élément, noté k_0 .

$k_0 > n$, car $R(n)$ est vrai, donc $k_0 - 1 \in \llbracket n, m - 1 \rrbracket$ et $R(k_0 - 1)$ est vraie. Alors d'après les hypothèses, $R(k_0)$ est également vraie, ce qui est faux. $\square$

Principe de récurrence descendante finie : Soit $n, m \in \mathbb{N}$ avec $n \leq m$.

Soit $R(k)$ un prédicat défini pour $k \in \llbracket n, m \rrbracket$.

Si $R(m)$ est vraie et si pour tout $k \in \llbracket n + 1, m \rrbracket$, $R(k)$ implique $R(k - 1)$, alors $R(k)$ est vraie pour tout $k \in \llbracket n, m \rrbracket$.

Démonstration.

Adapter la démonstration précédente. $\square$

Principe de récurrence forte :

Soit $n_0 \in \mathbb{N}$. Soit $R(n)$ un prédicat défini pour tout entier $n \geq n_0$.

Si $R(n_0)$ est vraie et si pour tout $n \geq n_0$, $[\forall k \in \{n_0, \dots, n\}, R(k)]$ implique $R(n + 1)$, alors pour tout $n \in \mathbb{N}$ tel que $n \geq n_0$, $R(n)$ est vraie.

Démonstration.

On pose $S(n) = [\forall k \in \{n_0, \dots, n\}, R(k)]$ et on utilise le principe de récurrence 2. $\square$

Exemple. Tout entier naturel $n \geq 2$ se décompose sous la forme d'un produit de nombres premiers.

En effet, notons $R(n)$ cette propriété et montrons-la par récurrence forte.

$R(2)$ est évidente car 2 est premier.

Soit $n \geq 2$. On suppose $R(k)$ pour tout $k \in \{2, \dots, n\}$.

Si $n + 1$ est premier, c'est bien un produit de nombres premiers.

Sinon, il existe $d \notin \{1, n + 1\}$ tel que $d | n + 1$. Il existe donc $d' \in \mathbb{N}$ tel que $n + 1 = dd'$. $d \neq 0$ et $d' \neq 0$. Alors $d > 1$, donc $d' < dd' = n + 1$.

De plus, $d \leq dd' = n + 1$ et $d \neq n + 1$, donc $d \leq n$. On peut donc utiliser $R(d)$ et $R(d')$ ce qui permet de conclure.

On démontrera dans le cours d'arithmétique que cette décomposition est unique.

Exemple. Le raisonnement suivant comporte heureusement une erreur. Laquelle ?

Pour $n \in \mathbb{N}^*$, notons $R(n)$ la propriété suivante : n droites quelconques du plan sont toujours deux à deux parallèles.

Initialisation : Pour $n = 1$, on a bien $R(1)$.

Hérédité : Soit $n \geq 1$ tel que $R(n)$.

Considérons $n + 1$ droites du plan, notées $D_1, \dots, D_{n+1}$.

D'après $R(n)$, les droites $D_1, \dots, D_n$ sont deux à deux parallèles.

Toujours d'après $R(n)$, les droites $D_2, \dots, D_{n+1}$ sont deux à deux parallèles.

On en déduit $R(n + 1)$.

D'après le principe de récurrence, pour tout $n \in \mathbb{N}^*$, n droites du plan sont toujours 2 à 2 parallèles.

Solution : $R(1)$ est vraie et le raisonnement utilisé pour montrer que

$R(n) \implies R(n+1)$ est vrai, mais seulement lorsque $n \geq 2$, par transitivité de la relation de parallélisme.

Cependant, on n'a jamais démontré $R(2)$, qui bien sûr est fausse. C'est donc l'initialisation de la récurrence qui est incorrecte.

Ceci montre combien l'initialisation est importante. Pour éviter ce type d'erreur, il est bienvenu de vérifier qu'à partir de l'initialisation, on peut effectivement faire fonctionner l'itération pour montrer la propriété sur les premières valeurs de l'entier n .

Principe de récurrence double :

Soit $n_0 \in \mathbb{N}$. Soit $R(n)$ un prédicat défini pour tout entier $n \geq n_0$.

Si $R(n_0)$ et $R(n_0 + 1)$ sont vraies et si

pour tout $n \geq n_0$, $[R(n) \wedge R(n+1)]$ implique $R(n+2)$,

alors pour tout $n \in \mathbb{N}$ tel que $n \geq n_0$, $R(n)$ est vraie.

Démonstration.

On pose $S(n) = [R(n) \wedge R(n+1)]$ et on utilise le principe de récurrence 2. $\square$

Exemple. Considérons la suite $(u_n)_{n \in \mathbb{N}}$ définie par : $u_0 = 2$, $u_1 = 5$ et, pour tout $n \in \mathbb{N}$, $u_{n+2} = 5u_{n+1} - 6u_n$.

Montrer que, pour tout $n \in \mathbb{N}$, $u_n = 2^n + 3^n$.

Solution : Pour tout $n \in \mathbb{N}$, posons $R(n) : u_n = 2^n + 3^n$.

Initialisation : On vérifie $R(0)$ et $R(1)$.

Hérédité : Soit $n \in \mathbb{N}$. On suppose $R(n)$ et $R(n+1)$.

Ainsi, $u_n = 2^n + 3^n$ et $u_{n+1} = 2^{n+1} + 3^{n+1}$.

Alors $u_{n+2} = 5(2^{n+1} + 3^{n+1}) - 6(2^n + 3^n) = 2^n(10 - 6) + 3^n(15 - 6)$, d'où $R(n+2)$.

Principe de la descente infinie : cf page 23.

6 Résoudre et rédiger un problème

Bien que les correcteurs de concours soient relativement indulgents, il y a un minimum de règles à respecter et quelques conseils à suivre.

6.1 Les préalables

- Pour appliquer correctement le cours, il est indispensable de connaître précisément les énoncés des résultats du cours, et d'en maîtriser le vocabulaire.
- Il faut bien sûr dormir suffisamment **avant** les épreuves.
- Prenez de quoi vous alimenter pendant les épreuves (du sucre notamment).
- Une montre est indispensable pendant les devoirs surveillés (l'usage des smartphones est interdit).

6.2 Règles typographiques

- Écrivez à l'encre noire ou bleue.

- Laissez une marge à gauche d'au moins 4 cm.
- N'utilisez pas l'effaceur ou le correcteur blanc sur de trop grande surface.
- Faites figurer clairement les références des questions traitées.
- Sautiez une ligne entre deux questions.
- Sautiez une page (voire changez de feuille) entre deux exercices, entre deux parties d'un problème.
- Respectez les notations de l'énoncé même si ce ne sont pas celles auxquelles vous êtes habitué. N'hésitez pas, en revanche, à introduire des notations supplémentaires, en les définissant clairement.

6.3 Une rédaction claire

- Ecrivez lisiblement. Encadrez vos résultats.
- Une rédaction mathématique est d'abord un texte en français; chaque phrase doit posséder un sujet (clairement identifié), un verbe etc. Il faut faire des phrases en bon français, simples et claires.
- N'hésitez pas à faire figurer de nombreux dessins, même s'ils ne sont pas demandés, à condition toutefois qu'ils soient en rapport avec la question posée.
- Au cours des calculs, différenciez nettement les différents symboles utilisés; par exemple une écriture hâtive confond souvent " n " et " x ".
Lorsque le calcul est terminé, simplifiez au maximum le résultat obtenu, puis encadrez-le.
Ne trichez pas dans un calcul; il est préférable de donner son résultat, même s'il n'est pas correct et d'indiquer ce que *devrait* donner le calcul.
- Dans vos démonstrations, on doit pouvoir distinguer du premier coup d'oeil :
 - un résultat connu auquel vous faites référence, indiqué par "d'après le cours", "d'après le théorème <nom>", "on sait que" etc.
 - l'utilisation du résultat d'une question précédente, référencée par "d'après la question <numéro>", etc.
 - un résultat déduit de vos raisonnements, qui doit être précédé d'un "donc", d'un "d'où", d'un "par conséquent", d'un "il s'ensuit que", etc;
 - Une hypothèse intermédiaire, repérée par "supposons que" etc.
 - un résultat que vous allez démontrer, qui doit être annoncé par un "Démontrons que", un "Il faut établir que", etc.
 - La conclusion d'une question, ou bien une conclusion intermédiaire, repérée par "j'ai prouvé que", "nous venons de montrer que" etc.
- Il faut justifier chaque étape d'un raisonnement :
 - en utilisant une hypothèse de l'énoncé;
 - en citant un résultat obtenu dans les questions précédentes (même si l'on n'est pas parvenu à le démontrer);
 - en faisant explicitement référence à un résultat du cours dont il faut alors vérifier soigneusement toutes les hypothèses.
- Lorsque vous venez de trouver une solution, mais qu'elle vous paraît très com-

pliquée, perdez encore une ou deux minutes à essayer de la simplifier.

6.4 Un peu de stratégie

- Il est fréquent qu'un problème soit trop long dans le temps imparti. Il ne faut donc pas chercher à traiter *toutes* les questions le plus vite possible. L'objectif est de traiter *des* questions avec précision, clarté et rigueur. Ensuite bien sûr, ce niveau de qualité étant compris, il faut en faire le plus possible.
- Commencer par lire l'énoncé en diagonale : quelques minutes, pour prendre connaissance des objets mathématiques manipulés par l'énoncé.
Si vous détectez une question que vous connaissez, il faudra aussi lui accorder du temps.
Notez au brouillon les contraintes horaires que ceci impose.
- Avant d'aborder une nouvelle partie, refaites une lecture en diagonale de cette partie.
Avant d'aborder le 2.a.1, il faut lire en détail l'énoncé de la question 2 et tenter de comprendre ce qu'elle démontre.
- Pour une bonne utilisation des feuilles de brouillon, éviter bien sûr de tout écrire directement au propre et éviter également d'écrire tous les détails sur une feuille de brouillon puis de les recopier. Le bon usage du brouillon se situe entre les deux attitudes précédentes.
- Distinguer les questions fermées des questions ouvertes.
Pour les premières, comme la question posée donne précisément le résultat à atteindre, c'est la justification de la réponse qui est importante et qui vous apportera des points. Votre rédaction ne peut se limiter à "je suis d'accord avec l'énoncé", aussi triviale que soit la question.
Pour les questions ouvertes, il faut bien sûr conserver une rédaction précise et claire, mais avec un peu moins d'exigence, car dans ce cas, les points seront surtout décernés en fonction de la justesse de la réponse.
- Lorsque vous remarquez que vous êtes en train de faire une erreur de calcul ou de raisonnement, corrigez bien sûr, mais ayez le réflexe de rechercher dans les questions déjà traitées si vous n'avez pas commis la même erreur : il est en effet très fréquent de répéter plusieurs fois la même erreur.
- Cherchez chaque question en tenant compte des questions précédentes.
Lorsque la question posée vous résiste, consultez votre montre et accordez-vous 10 minutes de recherche. En cas d'échec, passez aux questions suivantes : les 10 minutes ne sont pas une perte pure, car elles vous auront permis de mieux vous imprégner du sujet.
Si vous parvenez à résoudre les 2 questions suivantes, pensez à revenir quelques secondes sur la question qui vous a résisté : vous aurez peut-être alors la solution, soit parce que les méthodes que vous avez utilisées pour les questions suivantes s'adaptent, soit parce que les questions suivantes donnent des informations supplémentaires, soit enfin parce que votre cerveau a inconsciemment

continué à travailler sur la question résistante.

Pendant les 10 minutes de recherche, consultez l'énoncé des questions suivantes.

Cela peut vous donner des informations très pertinentes pour avancer.

En résumé, les questions sont à chercher dans l'ordre, mais devant une question résistante, il ne faut pas hésiter à regarder plus loin, et/ou à revenir dessus plus tard.

- Chaque partie d'un problème est souvent largement indépendante des précédentes. Les questions faciles se trouvent en général au début de chaque partie et les questions difficiles à la fin. Si vous séchez sur les dernières questions d'une partie, passez donc à la partie suivante, après avoir pris connaissance des résultats démontrés dans les questions que vous ne traitez pas.