

THÉORIE DES ENSEMBLES

1. Un *ensemble* est une collection d'objets appelés les *éléments* de l'ensemble. Si a est un élément d'un ensemble E , on écrit $a \in E$; dans le cas contraire on écrit $a \notin E$.

Parmi les ensembles de nombres, on notera surtout les suivants :

$\mathbb{N}$ désigne l'ensemble des nombres entiers naturels $0, 1, 2, 3, \dots$

$\mathbb{Z}$ désigne l'ensemble des nombres entiers relatifs (positifs et négatifs) $\dots, -3, -2, -1, 0, 1, 2, 3, \dots$

$\mathbb{Q}$ désigne l'ensemble des nombres rationnels, c'est-à-dire des fractions $\frac{m}{n}$ avec $m \in \mathbb{Z}$, $n \in \mathbb{Z}$, $n \neq 0$.

$\mathbb{R}$ désigne l'ensemble des nombres réels (comprenant non seulement les nombres rationnels, mais aussi les nombres ayant un développement décimal arbitraire comme $\sqrt{2} = 1,41421\dots$, $\sqrt{30} = 5,47723\dots$, $\pi = 3,141592\dots$, $e = 2,71828\dots$, etc.).

$\mathbb{C}$ désigne l'ensemble des nombres complexes (qui seront introduits dans ce cours).

Un ensemble E peut se définir

- en énumérant ses éléments : $E = \{x_1, x_2, \dots, x_n\}$; par exemple $\mathbb{N}_n = \{1, 2, \dots, n\}$;
- à l'aide d'un autre ensemble F et d'une propriété $\mathbf{P}$; on écrit $E = \{x \in F \mid \mathbf{P}(x)\}$, ce qui signifie que E est constitué de tous les éléments x de F pour lesquels la propriété $\mathbf{P}(x)$ est vérifiée; par exemple

$$\mathbb{N}_n = \{m \in \mathbb{N} \mid 1 \leq m \leq n\} \quad (\text{nombres entiers entre 1 et } n),$$

$$\mathbb{R}^* = \{x \in \mathbb{R} \mid x \neq 0\} \quad (\text{nombres réels non nuls}),$$

$$\mathbb{R}_+ = \{x \in \mathbb{R} \mid x \geq 0\} \quad (\text{nombres réels positifs}),$$

$$\mathbb{Q}_+^* = \{x \in \mathbb{Q} \mid x > 0\} \quad (\text{nombres rationnels strictement positifs}).$$

La notation entre crochets $\{\dots\}$ est la notation universellement adoptée par les mathématiciens pour désigner des ensembles. Lorsqu'on énumère les éléments d'un ensemble, leur ordre ne joue aucun rôle; ainsi, par exemple, on a $\mathbb{N}_3 = \{1, 2, 3\} = \{3, 1, 2\}$. L'ensemble qui ne possède aucun élément est appelé *ensemble vide* et se note $\emptyset$.

2. Étant donné un ensemble E , on dit qu'un ensemble A est une *partie* de E (ou bien un *sous-ensemble* de E) si tout élément de A est aussi un élément de E . Dans ce cas on écrit $A \subset E$. On dit aussi que A est *inclus* dans E et on parle de *l'inclusion* de A dans E . Par exemple $\mathbb{N}_n \subset \mathbb{N}$, $\mathbb{Z} \subset \mathbb{Q}$, $\mathbb{Q} \subset \mathbb{R}$, etc. L'égalité de deux ensembles E et F est caractérisée par :

$$E = F \quad \Longleftrightarrow \quad E \subset F \quad \text{et} \quad F \subset E.$$

Ainsi, pour démontrer que deux ensembles E et F sont égaux, on doit généralement démontrer deux inclusions, c'est-à-dire montrer que tout élément de E est dans F et que tout élément de F est dans E . L'ensemble de toutes les parties de E se note $\mathcal{P}(E)$. On a toujours $\emptyset \in \mathcal{P}(E)$ et $E \in \mathcal{P}(E)$. Par exemple

$$\mathcal{P}(\mathbb{N}_3) = \{ \emptyset, \{1\}, \{2\}, \{3\}, \{1, 2\}, \{1, 3\}, \{2, 3\}, \{1, 2, 3\} \} \quad (\text{ensemble ayant 8 éléments}).$$

Si A et B sont deux parties d'un ensemble E , on définit :

$$\text{la réunion de } A \text{ et } B : \quad A \cup B = \{x \in E \mid x \in A \text{ ou } x \in B\},$$

$$\text{l'intersection de } A \text{ et } B : \quad A \cap B = \{x \in E \mid x \in A \text{ et } x \in B\},$$

$$\text{la différence de } A \text{ et } B : \quad A - B = \{x \in A \mid x \notin B\}, \quad \text{souvent notée aussi } A \setminus B,$$

$$\text{la différence symétrique de } A \text{ et } B : \quad A \triangle B = (A - B) \cup (B - A).$$

Par exemple $\mathbb{Q}_+ = \mathbb{Q} \cap \mathbb{R}_+$, $\mathbb{Q}^* = \mathbb{R}^* \cap \mathbb{Q}$, $\mathbb{Q}_+^* = \mathbb{Q}^* \cap \mathbb{Q}_+$. Notons qu'on a aussi $A \triangle B = (A \cup B) - (A \cap B)$. Deux ensembles A et B sont appelés *disjoints* si $A \cap B = \emptyset$.

Si $A_1, A_2, A_3, \dots, A_n$ sont des parties d'un ensemble E , on définit de même :

$$A_1 \cup A_2 \cup A_3 \cup \dots \cup A_n = \{x \in E \mid \text{il existe } i \in \mathbb{N}_n \text{ tel que } x \in A_i\},$$

$$\text{noté plus simplement } \bigcup_{i=1}^n A_i = A_1 \cup \dots \cup A_n,$$

$$A_1 \cap A_2 \cap A_3 \cap \dots \cap A_n = \{x \in E \mid x \in A_i \text{ pour tout } i \in \mathbb{N}_n\},$$

$$\text{noté plus simplement } \bigcap_{i=1}^n A_i = A_1 \cap \dots \cap A_n.$$

3. Un *couple* est un système ordonné de deux éléments x_1 et x_2 , non nécessairement distincts. On le note (x_1, x_2) . Il faut faire la distinction entre l'ensemble $\{x_1, x_2\}$ (non ordonné) et le couple (x_1, x_2) où l'ordre est déterminé. Ainsi on a par exemple $\{1, 2\} = \{2, 1\}$ (égalité d'ensembles), mais les couples $(1, 2)$ et $(2, 1)$ ne sont pas égaux. Le *produit cartésien* des ensembles E et F est, par définition, l'ensemble des couples

$$E \times F = \{(x, y) \mid x \in E, y \in F\}.$$

Plus généralement, un n -uple $(x_1, x_2, \dots, x_n)$ est un système ordonné de n objets $x_1, x_2, \dots, x_n$, non nécessairement distincts. Un 2-uple est donc un *couple*, un 3-uple s'appelle un *triple*, etc. Le *produit cartésien* des ensembles $E_1, E_2, \dots, E_n$ est l'ensemble

$$E_1 \times E_2 \times \dots \times E_n = \{(x_1, x_2, \dots, x_n) \mid x_j \in E_j \text{ pour tout } j \text{ avec } 1 \leq j \leq n\}.$$

On note aussi $F^n = \overbrace{F \times \dots \times F}^{n \text{ fois}}$. Par exemple $\mathbb{R}^2$ est l'ensemble de tous les couples de nombres réels (qui représentent des points du plan), $\mathbb{R}^3$ est l'ensemble de tous les triples de nombres réels (qui représentent des points de l'espace), etc.

4. Une *application* $f : E \longrightarrow F$ d'un ensemble E dans un ensemble F est une correspondance qui associe à tout élément x de E un élément $f(x)$ de F . L'élément $f(x)$ est donc déterminé de manière unique par x et s'appelle l'*image* de x par f . On écrit $x \longmapsto f(x)$ pour spécifier l'image d'un élément x . L'ensemble E s'appelle l'*ensemble de départ* (ou *source*) et F l'*ensemble d'arrivée* (ou *but*) de l'application f . Le *graphe* de f est l'ensemble

$$\Gamma(f) = \{(x, f(x)) \mid x \in E\}, \quad \text{qui est un sous-ensemble de } E \times F.$$

Une application n'est bien déterminée que si l'on connaît son ensemble de départ E , son ensemble d'arrivée F et son graphe (ou en d'autres termes l'image de chaque élément de E). Dans le cas particulier où l'ensemble d'arrivée est l'ensemble $\mathbb{R}$ des nombres réels ou l'ensemble $\mathbb{C}$ des nombres complexes, une application s'appelle aussi une *fonction*.

Une *famille* $\{y_x\}_{x \in E}$ d'éléments de F indexée par E est définie comme étant une application de E dans F telle que $x \longmapsto y_x$; en d'autres termes, pour chaque élément $x \in E$, on spécifie un élément de F , noté y_x .

On dit qu'une application $f : E \longrightarrow F$ est

- *surjective* (ou une *surjection*) si tout élément de F est l'image d'au moins un élément de E ;
- *injective* (ou une *injection*) si deux éléments distincts de E ont toujours des images distinctes dans F , (en d'autres termes : si $f(x) = f(x')$, alors $x = x'$);
- *bijective* (ou une *bijection*) si f est injective et surjective.

Dans ce dernier cas, et dans ce cas seulement, il existe une application $f^{-1} : F \longrightarrow E$, appelée *application inverse* de f , telle que $f(f^{-1}(y)) = y$ pour tout $y \in F$ et $f^{-1}(f(x)) = x$ pour tout $x \in E$.

Soit $f : E \longrightarrow F$ une application, $A \subset E$ et $B \subset F$. On définit

$$\begin{aligned} \text{l'image (directe) de } A \text{ par } f : & \quad f(A) = \{f(x) \mid x \in A\}, \\ \text{l'image réciproque de } B \text{ par } f : & \quad f^{-1}(B) = \{x \in E \mid f(x) \in B\}. \end{aligned}$$

La notation $f^{-1}(B)$ ne signifie pas qu'on a affaire à une application f^{-1} de F dans E , car une telle application inverse n'existe généralement pas (elle n'existe que si f est bijective). L'écriture $f^{-1}(B)$ n'est donc qu'une notation commode.

Dans le cas particulier où A est l'ensemble E tout entier, l'ensemble $f(E)$ s'appelle l'*image* de f et se note $\text{Im}(f)$. Par exemple, dire que $\text{Im}(f) = F$ est équivalent à dire que f est surjective.

Si $A \subset E$, la *restriction* de f à A est l'application $f|_A : A \longrightarrow F$ définie par $x \longmapsto f(x)$ pour tout $x \in A$. C'est donc "la même" application que f , sauf que l'ensemble de départ a été restreint à A .

Si $f : E \longrightarrow F$ et $g : F \longrightarrow G$ sont des applications, l'*application composée* de f et g , notée $g \circ f$, est l'application $g \circ f : E \longrightarrow G$ définie par $(g \circ f)(x) = g(f(x))$ pour tout $x \in E$. La composée n'est définie que si l'ensemble d'arrivée de f coïncide avec l'ensemble de départ de g .

L'application $\text{id}_E : E \longrightarrow E$, définie par $\text{id}_E(x) = x$ pour tout $x \in E$, s'appelle l'*identité* de E . Pour toute application $f : E \longrightarrow F$, on a $\text{id}_F \circ f = f$ et $f \circ \text{id}_E = f$. Si $f : E \longrightarrow F$ est bijective, alors l'application inverse $f^{-1} : F \longrightarrow E$ existe et on a les relations $f^{-1} \circ f = \text{id}_E$ et $f \circ f^{-1} = \text{id}_F$.

5. Un ensemble E est dit *fini* s'il existe $n \in \mathbb{N}$ et une bijection $f : E \longrightarrow \mathbb{N}_n$. Dans ce cas n est unique et est appelé le *cardinal* de E (ou simplement le nombre d'éléments de E). On note $n = \text{Card}(E)$. Pour donner un sens à cette définition lorsque $n = 0$, on convient aussi que $\text{Card}(\emptyset) = 0$.

Si E et F sont des ensembles finis, alors

$$\text{Card}(E \times F) = \text{Card}(E) \text{Card}(F), \quad \text{Card}(E \cup F) = \text{Card}(E) + \text{Card}(F) - \text{Card}(E \cap F).$$

Théorème. Soit $f : E \longrightarrow F$ une application entre deux ensembles finis de même cardinal.

- (a) Si f est injective, alors f est aussi surjective (et donc f est une bijection).
- (b) Si f est surjective, alors f est aussi injective (et donc f est une bijection).

Preuve. (a) Comme f est injective, si on restreint l'ensemble d'arrivée à $\text{Im}(f)$, on obtient une bijection $E \rightarrow \text{Im}(f)$, $x \mapsto f(x)$. Donc $\text{Im}(f)$ a le même cardinal que E . Mais comme $\text{Im}(f)$ est un sous-ensemble de F qui a lui aussi le même cardinal, on doit avoir $\text{Im}(f) = F$, ce qui montre que f est surjective.

(b) Soit n le cardinal de F et écrivons $F = \{x_1, \dots, x_n\}$. Les images réciproques $f^{-1}(\{x_i\})$ de chacun des éléments x_i de F sont deux à deux disjointes et leur réunion donne E tout entier : $E = \bigcup_{i=1}^n f^{-1}(\{x_i\})$.

On a donc

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

Chacun des nombres $\text{Card}(f^{-1}(\{x_i\}))$ est ≥ 1 par surjectivité de f . Comme la somme de ces n nombres vaut n , chacun doit être égal à 1. Dire que $\text{Card}(f^{-1}(\{x_i\})) = 1$ pour chaque x_i revient à dire que f est injective. $\square$

6. Une *relation binaire* sur un ensemble E est une partie R de $E \times E$; on note également xRy à la place de $(x, y) \in R$ et on dit que x est en relation avec y . Une relation binaire R est appelée une *relation d'équivalence* si elle est à la fois
- *réflexive* : xRx pour tout $x \in E$;
 - *symétrique* : xRy implique yRx ;
 - *transitive* : xRy et yRz impliquent xRz .

Dans ce cas, l'ensemble $[x] = \{y \in E \mid yRx\}$ est appelé la *classe d'équivalence* de x pour la relation d'équivalence R .

Théorème. Soit R une relation d'équivalence sur un ensemble E . Deux éléments x et x' de E sont en relation si et seulement si leurs classes d'équivalence $[x]$ et $[x']$ sont égales.

Preuve. Supposons qu'on a xRx' . Alors pour tout $y \in [x]$, on a yRx , et comme on a aussi xRx' , on en déduit que yRx' par transitivité de R , c'est-à-dire $y \in [x']$. De même si $y \in [x']$, alors on a yRx' et $x'Rx$ (car xRx' et R est symétrique), donc yRx par transitivité, c'est-à-dire $y \in [x]$. Ainsi on a montré que $[x] = [x']$.

Supposons maintenant l'égalité $[x] = [x']$. Comme on a xRx par réflexivité de R , on a $x \in [x]$, donc $x \in [x']$, c'est-à-dire xRx' , ce qui achève la preuve. $\square$

UN PEU DE LOGIQUE

1. On doit souvent considérer une propriété $P(x)$ qui dépend d'un élément x variant dans un ensemble E . Elle peut être vraie pour certains éléments x et fausse pour d'autres. Par exemple le fait d'être divisible par 3 définit une propriété $P(n)$ dépendant de chaque entier $n \in \mathbb{N}$; elle est fausse si $n = 1$, $n = 2$, $n = 4$, etc., et elle est vraie si $n = 0$, $n = 3$, $n = 6$, etc.
La négation d'une propriété $P(x)$ est la propriété $\text{NON } P(x)$, qui est vraie lorsque $P(x)$ est fausse et qui est fausse lorsque $P(x)$ est vraie.

2. Si une propriété $P(x)$ est vraie pour tout élément $x \in E$, on écrit

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

le signe $\forall$ se lisant "pour tout". Pour démontrer cela dans le cas concret d'une propriété $P(x)$ donnée, il faut faire un raisonnement qui prouve la véracité de $P(x)$ pour un élément x arbitraire dans E .

Si une propriété $P(x)$ est vraie pour au moins un élément $x \in E$, on écrit

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

le signe $\exists$ se lisant "il existe". Pour démontrer cela dans le cas concret d'une propriété $P(x)$ donnée, il suffit de trouver un $x \in E$ pour lequel la propriété $P(x)$ est vraie (et on n'a alors pas besoin de se préoccuper des autres éléments de E).

La négation de chacune de ces deux dernières propriétés fait intervenir logiquement l'autre, de la manière suivante :

$$\text{NON } (\forall x \in E, P(x)) \quad \text{est équivalent à} \quad \exists x \in E, \text{NON } P(x),$$

$$\text{NON } (\exists x \in E, P(x)) \quad \text{est équivalent à} \quad \forall x \in E, \text{NON } P(x),$$

On fait un large usage de ce type de raisonnement logique.

Exemple. L'assertion que tout entier est un carré parfait (qui est évidemment fausse) se traduit plus précisément par la phrase : pour tout entier $n \in \mathbb{N}$, il existe un entier $m \in \mathbb{N}$ tel que $n = m^2$, ou encore :

$$\forall n \in \mathbb{N}, \exists m \in \mathbb{N}, n = m^2 \quad (\text{faux}).$$

C'est la négation de cette phrase qui est vraie. En vertu des règles ci-dessus, elle s'écrit :

$$\exists n \in \mathbb{N}, \forall m \in \mathbb{N}, n \neq m^2 \quad (\text{vrai}),$$

ce qui donne en français : il existe un entier $n \in \mathbb{N}$ tel que, pour tout entier $m \in \mathbb{N}$, n n'est pas égal au carré de m , ou encore : il existe un entier n qui n'est égal à aucun carré parfait. La preuve de cette assertion consiste simplement à trouver explicitement un tel entier n et à montrer qu'il n'est pas un carré, ce qui est évidemment très facile; par exemple on peut prendre $n = 2$ (ou bien $n = 3$, ou bien $n = 21$).

3. On dit qu'une propriété $P(x)$ *implique* une autre propriété $Q(x)$ si, pour chaque élément x pour lequel $P(x)$ est vraie, alors $Q(x)$ l'est aussi. L'implication logique se note à l'aide du signe $\Rightarrow$:

$$P(x) \Rightarrow Q(x)$$

Il faut remarquer ici que, par définition, on n'exige rien du tout dans le cas d'un élément x pour lequel $P(x)$ est fausse : $Q(x)$ peut être vraie ou fausse dans ce cas.

On dit qu'une propriété $P(x)$ est *équivalente* à une autre propriété $Q(x)$ si on a à la fois $P(x) \Rightarrow Q(x)$ et $Q(x) \Rightarrow P(x)$. On écrit dans ce cas $P(x) \iff Q(x)$.

L'implication $P(x) \Rightarrow Q(x)$ est équivalente à l'implication $\text{NON } Q(x) \Rightarrow \text{NON } P(x)$ (et non pas à l'implication $\text{NON } P(x) \Rightarrow \text{NON } Q(x)$, faute commune de raisonnement logique !).

Par exemple, le fait que tous les hommes sont mortels peut s'écrire :

x est un homme $\Rightarrow x$ est mortel, ou encore : x est immortel $\Rightarrow x$ n'est pas un homme (mais surtout pas : x n'est pas un homme $\Rightarrow x$ est immortel, car alors tous les chats seraient immortels !).

Lorsque x parcourt les éléments d'un ensemble E , l'implication $P(x) \Rightarrow Q(x)$ est équivalente à l'inclusion d'ensembles

$$\{x \in E \mid P(x)\} \subset \{x \in E \mid Q(x)\}.$$

Elle se démontre en prouvant que si un élément x est dans le premier ensemble alors il est dans le second, ou encore en prouvant que si un élément x n'est pas dans le second ensemble alors il n'est pas dans le premier.

4. Le *raisonnement par récurrence* joue un rôle très important en mathématiques. Il intervient lorsqu'on veut démontrer qu'une propriété $P(n)$, qui dépend d'un entier naturel $n \geq 0$, est vraie pour toute valeur de $n \geq m$, pour un nombre naturel m fixé. Le raisonnement par récurrence consiste à démontrer deux choses :

- (a) la propriété $P(m)$ est vraie;
- (b) l'implication $P(n) \Rightarrow P(n+1)$ est vraie.

Si ces deux assertions sont démontrées, alors la propriété $P(n)$ est vraie pour tout $n \geq m$. En effet, elle est vraie pour $n = m$ en vertu de (a); vu qu'elle est vraie pour $n = m$ elle l'est pour $n = m + 1$ en vertu de (b); vu qu'elle est vraie pour $n = m + 1$ elle l'est pour $n = m + 2$ en vertu de (b); vu qu'elle est vraie pour $n = m + 2$ elle l'est pour $n = m + 3$ en vertu de (b); etc.

Pour démontrer (b), il faut faire l'hypothèse que $P(n)$ est vraie et, à l'aide d'un raisonnement adapté au problème considéré, prouver que $P(n+1)$ est alors aussi vraie. L'hypothèse que $P(n)$ est vraie s'appelle *l'hypothèse de récurrence*.

Exemple. On démontre par récurrence l'égalité

$$1 + 2 + 3 + \dots + n = \frac{n(n+1)}{2} \quad (n \geq 1).$$

Pour $n = 1$, l'égalité est vraie car $1 = \frac{1(1+1)}{2}$. En faisant l'hypothèse de récurrence que l'égalité est vraie pour n , on a donc $1 + 2 + 3 + \dots + n = \frac{n(n+1)}{2}$. Alors, en ajoutant $n + 1$, on obtient :

$$1 + 2 + 3 + \dots + n + (n + 1) = \frac{n(n+1)}{2} + (n + 1) = \frac{n^2 + n + 2n + 2}{2} = \frac{(n+1)(n+2)}{2},$$

ce qui démontre l'égalité pour $n + 1$. Ainsi l'égalité est vraie pour tout entier ≥ 1 .

LOIS DE COMPOSITION

1. Une *loi de composition* (ou simplement *loi*) sur un ensemble E est une application de $E \times E$ dans E . On attribue à une loi de composition un symbole, par exemple $*$, ou bien $\diamond$, l'image de (x, y) par la loi étant alors notée $x * y$ (respectivement $x \diamond y$). Lorsque la loi est l'addition, respectivement la multiplication, on utilise évidemment
 - la notation additive : $(x, y) \mapsto x + y$,
 - la notation multiplicative : $(x, y) \mapsto x \cdot y$ ou xy .

Exemples. 1) L'addition est une loi de composition sur chacun des ensembles $\mathbb{R}$, $\mathbb{Q}$, $\mathbb{Z}$ et $\mathbb{N}$.

2) La multiplication est une loi de composition sur chacun des ensembles $\mathbb{R}$, $\mathbb{Q}$, $\mathbb{Z}$ et $\mathbb{N}$.

3) La réunion $(A, B) \mapsto A \cup B$ et l'intersection $(A, B) \mapsto A \cap B$ sont des lois de composition sur l'ensemble $\mathcal{P}(X)$ de toutes les parties de X .

3) $(f, g) \mapsto f \circ g$ est une loi de composition sur l'ensemble $\mathcal{F}(X, X)$ de toutes les applications de X dans X .

4) Sur $\mathbb{R}^3$, on a la loi $\wedge$ définie par $(x_1, x_2, x_3) \wedge (y_1, y_2, y_3) = (x_2y_3 - x_3y_2, x_3y_1 - x_1y_3, x_1y_2 - x_2y_1)$, appelée produit vectoriel.

2. Soit $*$ une loi sur un ensemble E . On dit qu'une partie $A \subset E$ est *stable* par la loi $*$ si, pour tout choix de $x, y \in A$, on a $x * y \in A$. Lorsque A est stable par la loi $*$, cette loi définit, par restriction, une loi sur A , qu'on dit *induite* par celle de E et qu'on désigne par le même symbole. Ainsi, l'addition usuelle sur $\mathbb{N}$, $\mathbb{Z}$, $\mathbb{Q}$ est induite par l'addition usuelle sur $\mathbb{R}$. Il en est de même de la multiplication.
3. Une loi $*$ sur un ensemble E est dite *associative* si $(x * y) * z = x * (y * z)$ pour tout choix de $x, y, z \in E$. Pour composer une suite finie $x_1, \dots, x_n$ d'éléments de E dans l'ordre donné, on a, a priori, plusieurs possibilités, chacune d'elles spécifiée par un système particulier de parenthèses. Ainsi, lorsque $n = 4$, on a cinq possibilités :

$$\begin{aligned} &((x_1 * x_2) * x_3) * x_4, \quad (x_1 * (x_2 * x_3)) * x_4, \quad (x_1 * x_2) * (x_3 * x_4), \\ &x_1 * ((x_2 * x_3) * x_4), \quad x_1 * (x_2 * (x_3 * x_4)). \end{aligned}$$

Lorsque la loi $*$ est associative, le résultat est indépendant du système de parenthèses choisi et se note simplement $x_1 * x_2 * \dots * x_n$. En notation additive, on utilise la notation

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

et en notation multiplicative

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

Lorsque $x_1 = x_2 = \dots = x_n = x$, la notation $\sum_{i=1}^n x_i$ s'abrège en nx et $\prod_{i=1}^n x_i$ s'abrège en x^n .

4. Soit $*$ une loi de composition sur un ensemble E . On dit que $x, y \in E$ *commutent* si $x * y = y * x$. On dit que la loi $*$ est *commutative* si $x * y = y * x$ pour tout choix de $x, y \in E$. Dans toutes les circonstances où elle a un sens, la loi d'addition $+$ est commutative. En revanche, bien que la loi de multiplication soit aussi commutative lorsqu'on multiplie des nombres, on verra des exemples (les matrices) où la multiplication n'est pas commutative.

Si $*$ est une loi associative sur E et si $x_1, \dots, x_n \in E$ est une suite finie d'éléments qui commutent deux à deux, alors non seulement la composition de $x_1, \dots, x_n$ est indépendante du système de parenthèses mais encore indépendante de l'ordre des éléments. Il s'ensuit que, si la loi $*$ est associative et commutative, on peut définir sans ambiguïté la composition des éléments d'une partie finie non vide A de E . Dans ce cas, on notera simplement

$$\sum_{x \in A} x \quad \text{en notation additive} \quad \text{et} \quad \prod_{x \in A} x \quad \text{en notation multiplicative.}$$

Plus généralement, si $\{x_a\}_{a \in A}$ est une famille d'éléments de E indexée par A , on définit de manière analogue les expressions

$$\sum_{a \in A} x_a \quad \text{et} \quad \prod_{a \in A} x_a.$$

5. Soit $*$ une loi de composition sur un ensemble E . On dit qu'un élément $e \in E$ est *neutre* si $x * e = e * x = x$ pour tout $x \in E$. Une loi $*$ admet au plus un élément neutre, car si $e, e' \in E$ sont neutres, on a $e = e * e' = e'$.

En notation additive, l'élément neutre se note 0 et s'appelle l'élément nul. Ainsi $x + 0 = 0 + x = x$ pour tout x .

En notation multiplicative, l'élément neutre se note le plus souvent 1 et s'appelle l'élément unité. Ainsi $x \cdot 1 = 1 \cdot x = x$ pour tout x .

6. Soit $*$ une loi de composition sur un ensemble E , admettant l'élément neutre e . On dit que $y \in E$ est un *inverse* de $x \in E$ si $x * y = y * x = e$. Un élément qui admet un inverse est dit *inversible*.

En notation multiplicative, l'inverse de x , s'il existe, se note x^{-1} . Ainsi $x \cdot x^{-1} = x^{-1} \cdot x = 1$.

En notation additive, l'inverse de x se note $-x$ et s'appelle l'*opposé* de x . Ainsi $x + (-x) = (-x) + x = 0$.

Théorème. Soit $*$ une loi de composition associative sur E , admettant l'élément neutre e .

(a) Chaque élément inversible n'a qu'un seul inverse.

(b) Si $x, y \in E$ sont inversibles, alors $x * y$ l'est aussi. Plus précisément, si x' est l'inverse de x et si y' est l'inverse de y , alors $y' * x'$ est l'inverse de $x * y$.

En notation multiplicative, $(x \cdot y)^{-1} = y^{-1} \cdot x^{-1}$ (dans l'ordre inverse!).

Preuve. (a) Soient $a', a'' \in E$ des inverses de $a \in E$. Alors $a' = a' * e = a' * (a * a'') = (a' * a) * a'' = e * a'' = a''$ et donc $a' = a''$.

(c) On a $(x * y) * (y' * x') = (x * (y * y')) * x' = (x * e) * x' = x * x' = e$ et de même $(y' * x') * (x * y) = e$. $\square$

Plus généralement, sous les hypothèses du théorème, si $x_1, \dots, x_n \in E$ sont inversibles et si $x'_1, \dots, x'_n \in E$ sont leurs inverses respectifs, alors $x'_n * x'_{n-1} * \dots * x'_1$ est l'inverse de $x_1 * \dots * x_n$.

7. On a déjà défini nx en notation additive et x^n en notation multiplicative, pour tout entier $n \geq 1$. Lorsque $n = 0$, on définit encore $0x = 0$ en notation additive et $x^0 = 1$ en notation multiplicative.

Si de plus x est inversible et n est un entier négatif, alors $n = -m$ (avec m positif) et on définit :

en notation additive : $(-m)x = m(-x)$ où $-x$ est l'opposé de x .

en notation multiplicative : $x^{-m} = (x^{-1})^m$ où x^{-1} est l'inverse de x .

On a les règles familières suivantes :

en notation additive

$$nx + kx = (n + k)x$$

$$n(kx) = (nk)x$$

$$n(x + y) = nx + ny$$

en notation multiplicative

$$x^n \cdot x^k = x^{n+k}$$

$$(x^n)^k = x^{nk}$$

$$(xy)^n = x^n y^n \quad (\text{valable uniquement si } x \text{ et } y \text{ commutent})$$

Ces règles sont valables en général pour $n \geq 0$ et $k \geq 0$. Si x et y sont inversibles, elles sont valables pour tout choix de $n, k \in \mathbb{Z}$.

GROUPES

1. Un *groupe* est un ensemble muni d'une loi de composition associative, admettant un élément neutre et telle que tout élément est inversible. Un groupe est dit *abélien* (ou *commutatif*) si la loi de composition est commutative. Si le groupe est noté G et la loi de composition $*$, on écrit souvent $(G, *)$ pour spécifier le groupe avec sa loi de composition.

Exemples. 1) $\mathbb{Z}$, $\mathbb{Q}$, $\mathbb{R}$ sont des groupes abéliens pour l'addition usuelle.

2) $\{+1, -1\}$, $\mathbb{Q}^*$, $\mathbb{R}^*$, $\mathbb{R}_+^*$ sont des groupes abéliens pour la multiplication usuelle.

3) L'ensemble S_n de toutes les permutations de $\mathbb{N}_n$ (c'est-à-dire les bijections de $\mathbb{N}_n$ sur lui-même) est un groupe pour la composition usuelle des applications. L'élément neutre est l'application identité et l'inverse d'une permutation σ est l'application inverse σ^{-1} . Le groupe S_n n'est pas abélien si $n \geq 3$.

Dans la suite de ces notes, on utilisera les notations suivantes. La loi de composition d'un groupe G sera notée $*$. L'élément neutre sera noté e . L'inverse de x sera noté x' (mais il convient de se souvenir que la notation est x^{-1} si la loi est la multiplication et $-x$ si la loi est l'addition).

2. Soit G un groupe. On appelle *sous-groupe* de G toute partie $H \subset G$ vérifiant les trois conditions suivantes :
 - (i) H est non vide.
 - (ii) H est stable par la loi de composition de G (c'est-à-dire, quels que soient $x, y \in H$, on a $x * y \in H$).
 - (iii) La loi sur H induite par celle de G est une loi de groupe sur H .

Critère. Pour qu'une partie H d'un groupe G soit un sous-groupe de G , il faut et il suffit qu'elle satisfasse aux trois conditions suivantes :

- (a) H est non vide.
- (b) H est stable par la loi de composition de G .
- (c) Pour tout $x \in H$, l'inverse x' (qui existe dans G) appartient à H .

Preuve. La nécessité est évidente. Pour montrer la suffisance, il s'agit de montrer que les conditions (a), (b) et (c) impliquent que la loi sur H induite par celle de G est une loi de groupe. L'associativité est évidente car elle est vérifiée dans G tout entier. On a $e \in H$, car il suffit de choisir $x \in H$, grâce à (a), d'où on tire $x' \in H$, grâce à (c), et donc $e = x * x' \in H$, grâce à (b). L'existence d'inverses est garantie par (c). $\square$

Exemples. 1) Pour l'addition usuelle, $\mathbb{Z}$ et $\mathbb{Q}$ sont des sous-groupes de $\mathbb{R}$.

2) Pour la multiplication usuelle, $\{+1, -1\}$, $\mathbb{Q}^*$, $\mathbb{R}_+^*$ sont des sous-groupes de $\mathbb{R}^*$.

3) Pour tout groupe G , le sous-ensemble $\{e\}$ est un sous-groupe de G . De même G tout entier est un sous-groupe de G .

4) Soit $x \in \mathbb{R}^*$. L'ensemble $\{x^m \mid m \in \mathbb{Z}\}$ de toutes les puissances positives et négatives de x est un sous-groupe de $\mathbb{R}^*$. C'est le plus petit sous-groupe de $\mathbb{R}^*$ contenant l'élément x .

5) Si $n \in \mathbb{Z}$, on note $n\mathbb{Z}$ l'ensemble $\{nk \mid k \in \mathbb{Z}\}$ des multiples de n . Clairement, $n\mathbb{Z}$ est un sous-groupe de $(\mathbb{Z}, +)$. En fait, on peut montrer que tout sous-groupe de $\mathbb{Z}$ est de cette forme.

3. Soient $(G, *)$ et $(H, \diamond)$ des groupes. On appelle *homomorphisme de groupes* de G dans H toute application $f : G \rightarrow H$ telle que $f(x * y) = f(x) \diamond f(y)$ pour tout choix de $x, y \in G$.
 En notation additive, cela revient à demander que $f(x + y) = f(x) + f(y)$ pour tout choix de $x, y \in G$.
 En notation multiplicative, cela revient à demander que $f(xy) = f(x)f(y)$ pour tout choix de $x, y \in G$.

Exemples. 1) Soit $a \in \mathbb{R}^*$. L'application $\mathbb{Z} \rightarrow \mathbb{R}^*$, définie par $m \mapsto a^m$ pour tout $m \in \mathbb{Z}$, est un homomorphisme du groupe $(\mathbb{Z}, +)$ dans $(\mathbb{R}^*, \cdot)$. C'est en fait l'unique homomorphisme de $(\mathbb{Z}, +)$ dans $(\mathbb{R}^*, \cdot)$ qui envoie 1 sur a .

2) Soit $n \in \mathbb{Z}$. L'application $\mathbb{R}^* \rightarrow \mathbb{R}^*$, définie par $x \mapsto x^n$ pour tout $x \in \mathbb{R}^*$, est un homomorphisme de groupes de $(\mathbb{R}^*, \cdot)$ dans $(\mathbb{R}^*, \cdot)$. On a de même un homomorphisme $(\mathbb{R}, +) \rightarrow (\mathbb{R}, +)$, $x \mapsto nx$.

3) L'application $\mathbb{R}_+^* \rightarrow \mathbb{R}$, définie par $x \mapsto \log x$ pour tout $x \in \mathbb{R}_+^*$, est un homomorphisme du groupe $(\mathbb{R}_+^*, \cdot)$ dans le groupe $(\mathbb{R}, +)$. En effet, la propriété fondamentale du logarithme affirme précisément que $\log(xy) = \log x + \log y$ pour tout choix de $x, y \in \mathbb{R}_+^*$.

Théorème. Tout homomorphisme de groupes $f : (G_1, *) \rightarrow (G_2, \diamond)$ possède les propriétés suivantes :

(a) $f(e_1) = e_2$ où e_1 désigne l'élément neutre de G_1 , et e_2 celui de G_2 .

En notation additive, $f(0) = 0$. En notation multiplicative, $f(1) = 1$.

(b) $f(x') = f(x)'$ pour tout $x \in G$, où $'$ désigne l'inverse.

En notation additive, $f(-x) = -f(x)$. En notation multiplicative, $f(x^{-1}) = f(x)^{-1}$.

Preuve. (a) Comme $e_1 * e_1 = e_1$, on a $f(e_1) \diamond f(e_1) = f(e_1 * e_1) = f(e_1)$. De plus $f(e_1) = e_2 \diamond f(e_1)$, si bien que $f(e_1) \diamond f(e_1) = e_2 \diamond f(e_1)$. Après simplification, $f(e_1) = e_2$.

(b) On a $f(e_1) = e_2$ par a), et donc $e_2 = f(e_1) = f(x * x') = f(x) \diamond f(x')$. De manière analogue, on a $e_2 = f(x') \diamond f(x)$. Il en résulte que $f(x')$ est l'inverse de $f(x)$, c'est-à-dire $f(x') = f(x)'$. $\square$

4. Si $f : G_1 \rightarrow G_2$ est un homomorphisme de groupes et si e_2 désigne l'élément neutre de G_2 , l'ensemble

$$f^{-1}(\{e_2\}) = \{x \in G_1 \mid f(x) = e_2\}$$

est appelé le *noyau* de f et noté $\text{Ker}(f)$. On vérifie aisément que $\text{Ker}(f)$ est un sous-groupe de G_1 . En notation additive (la plus importante pour la suite de ce cours),

$$\text{Ker}(f) = \{x \in G_1 \mid f(x) = 0\}.$$

5. Le groupe $\mathbb{Z}/n\mathbb{Z}$.

Fixons un entier $n \geq 1$. On définit sur $\mathbb{Z}$ une relation binaire, dite *congruence modulo n* , comme suit. Si $x, y \in \mathbb{Z}$,

$$x \equiv y \pmod{n} \iff x - y \in n\mathbb{Z}.$$

La relation $x \equiv y \pmod{n}$ se lit " x est congru à y modulo n " et signifie donc que l'entier $x - y$ est multiple de n . On vérifie sans peine que c'est une relation d'équivalence sur $\mathbb{Z}$. La classe d'équivalence de $k \in \mathbb{Z}$ se note souvent $\bar{k}$ et s'appelle la *classe de k modulo n* .

La classe de 0 est constituée de $\dots, -2n, -n, 0, n, 2n, 3n, \dots$.

La classe de 1 est constituée de $\dots, -2n+1, -n+1, 1, n+1, 2n+1, 3n+1, \dots$, etc.

Dans chaque classe, il existe un unique $r \in \mathbb{Z}$ tel que $0 \leq r < n$. Il s'ensuit qu'il y a exactement n classes, à savoir $\bar{0}, \bar{1}, \dots, \bar{n-1}$.

L'ensemble des classes modulo n se note $\mathbb{Z}/n\mathbb{Z}$, et donc $\text{Card}(\mathbb{Z}/n\mathbb{Z}) = n$.

La relation de congruence modulo n possède la propriété remarquable suivante :

Si $k, k', \ell, \ell' \in \mathbb{Z}$ sont tels que $k \equiv k' \pmod{n}$ et $\ell \equiv \ell' \pmod{n}$, on a aussi $k + \ell \equiv k' + \ell' \pmod{n}$.

Cela permet de définir une addition sur $\mathbb{Z}/n\mathbb{Z}$ par la règle

$$\bar{k} + \bar{\ell} = \overline{k + \ell}.$$

Il est facile de vérifier que $\mathbb{Z}/n\mathbb{Z}$, muni de cette addition, est un groupe abélien. L'application naturelle $\mathbb{Z} \rightarrow \mathbb{Z}/n\mathbb{Z}$, définie par $k \mapsto \bar{k}$ pour tout $k \in \mathbb{Z}$, est un homomorphisme surjectif de groupes; son noyau est $n\mathbb{Z}$.

Exemple. On considère le cas $n = 2$. Le groupe $\mathbb{Z}/2\mathbb{Z}$ est constitué de deux éléments, à savoir les classes $\bar{0}$ et $\bar{1}$, soumises aux règles de calcul suivantes :

$$\bar{0} + \bar{0} = \bar{0}, \quad \bar{0} + \bar{1} = \bar{1}, \quad \bar{1} + \bar{0} = \bar{1}, \quad \bar{1} + \bar{1} = \bar{0}.$$

C'est le calcul modulo 2, qui est celui utilisé par les ordinateurs.

Bibliographie

Les livres d'algèbre linéaire sont très nombreux. Le cours ne suit aucun de ces ouvrages de manière systématique, mais bien des matières du cours se trouvent dans ces livres. Voici un choix (non définitif !) :

- R. Cairoli, *Algèbre linéaire*, Presses Polytechniques Universitaires Romandes, 2^e édition 1999.
- S. Friedberg, A. Insel, L. Spence, *Linear Algebra*, any edition.
- J. Grifone, *Algèbre linéaire*, Cepadues-Editions, 1990.
- K. Hoffman, R. Kunze, *Linear Algebra*, Prentice-Hall, second edition, 1971.

De plus, pour trouver des résumés et de nombreux exercices, on pourra consulter :

- R. Dalang, A. Chabouni, *Algèbre linéaire*, Presses Polytechniques Universitaires Romandes, 2^e édition, 2004.
- S. Lipschutz, *Algèbre linéaire*, Mc Graw-Hill, Série Schaum, 1973.

Lettres grecques

Les lettres grecques sont des symboles très pratiques et très utilisés. En voici la liste. On y trouvera aussi les majuscules les plus courantes. Les minuscules peu utilisées sont entre parenthèses.

α	alpha			ν	nu		
β	bêta			ξ	xi		
γ	gamma	Γ	Gamma	(o	omicron)		
δ	delta	Δ	Delta	π	pi	Π	Pi
ε	epsilon			ρ	rho		
ζ	zêta			σ	sigma	Σ	Sigma
η	êta			τ	tau		
θ	thêta	Θ	Thêta	(v	upsilon)		
(ι	iota)			φ	phi	Φ	Phi
(κ	kappa)			χ	chi		
λ	lambda	Λ	Lambda	ψ	psi	Ψ	Psi
μ	mu			ω	omega	Ω	Omega