

UE 111 : Fondements Mathématiques

Le document suivant sera actualisé au fur et à mesure de la progression de l'enseignement, et la version actuelle sera toujours accessible via Arche, <https://arche.univ-lorraine.fr/course/view.php?id=28468>. Cette version correspond à l'avancement du cours dans le groupe EI 3 ; dans les autres groupes, les contenus et l'avancement du cours peuvent être légèrement différents.

Ces notes sont conçues pour accompagner, mais non pour remplacer le cours de l'UE 111 Fondements Mathématiques. D'une part, elles contiennent plus que le cours : nous cherchons à mettre le sujet des "fondements" dans leur contexte (historique, philosophique, épistémologique, et proprement mathématique) ; ces questions dépassent un cours de première année, mais le lecteur curieux est encouragé à suivre les hyperliens donnés en gris dans le texte ; pour la plupart, ces liens vont vers des pages wikipedia, qui à leur tour pointent vers d'autres sources disponibles sur internet. D'autre part, ces notes contiennent moins que le cours : comme l'UE 111 est un "enseignement intégré", il n'y a pas de frontière bien distincte entre cours et TD ; et lire ce texte ne peut pas remplacer le travail sur les exercices en parallèle avec le cours.

Fondements mathématiques – Introduction

Par “fondements mathématiques” on peut entendre

- (1) d’un point de vue *historique*, les origines des maths dans l’histoire de l’humanité ;
- (2) d’un point de vue *philosophique*, la description de la place et du rôle des mathématiques dans l’ensemble du savoir humain, et en particulier de comprendre ses relations spécifiques avec les *sciences naturelles* ;
- (3) d’un point de vue *conceptuel, intérieur aux mathématiques*, décrire les structures et concepts les plus basiques et fondamentaux des mathématiques.

C’est le dernier aspect qui fera l’objet de ce cours. Mais il n’est pas isolé des deux autres. Commençons par en dire quelques mots.

§1 Fondements de maths – repères historiques. Il faut remonter très loin dans l’histoire de l’humanité pour trouver les origines des mathématiques. Mais, non sans raison, on cite souvent les *mathématiques grecques anciennes* comme étant la véritable origine des mathématiques que nous connaissons : surtout, aux “*Eléments*” d’Euclide on doit la *méthode rigoureuse*, “*axiomatique*” qu’on emploie encore aujourd’hui :

- formuler clairement les propriétés et hypothèses fondamentales que nous acceptons comme point de départ : les *axiomes*, et en déduire tout le reste par un *raisonnement logique* : les *preuves*.

Autrement dit, nous mettons toutes les cartes sur la table au début du jeu ; ensuite, tout se déroule selon des règles strictes et logiques. Euclide a appliqué cette méthode surtout dans deux domaines :

1. la *géométrie*,
2. les *nombres* (l’arithmétique).

Cependant, après l’ “explosion des mathématiques” qui débuta au XVIIe siècle avec l’invention du *calcul infinitésimal*, par Newton et Leibniz, les mathématiciens se rendaient compte du fait que les fondements posés par Euclide étaient insuffisants : ces “fondements” ne permettaient pas de dire ce qu’est un “nombre réel”, ou de développer une théorie rigoureuse du calcul “infinitésimal”, ni même d’éclaircir vraiment les fondements de la géométrie dite “euclidienne” !

Le XIXe et le XXe siècle donnent la “solution” à ces problèmes et mènent ainsi à une nouvelle ouverture, mais aussi à de nouveaux problèmes. :

1. *Georg Cantor* crée la *Théorie des ensembles* qui donne un fondement rigoureux de toutes les mathématiques,
2. *Augustin Cauchy* et *Richard Dedekind* proposent des définitions rigoureuses des nombres réels,
3. *David Hilbert* analyse et repose en 1899 dans les “*Fondements de la géométrie*” les axiomes de la géométrie euclidienne,
4. le groupe Bourbaki essaie au XXe siècle de rassembler en plusieurs tomes les *Éléments de mathématique*, comme sorte de réponse moderne à Euclide,
5. la *Théorie des catégories* s’apprête au XXe siècle à remplacer la théorie des ensembles comme fondement universel des maths...

De tout ces développements, pour les études de maths en Licence, sont les plus importants : l'*analyse*, comportant comme outil de base les *nombres réels*, et l'*algèbre linéaire*, qui est la version moderne de constructions de modèles de la géométrie, les *espaces vectoriels*. Les deux théories utilisent le langage de la théorie des ensembles.

§2 Fondements de maths – point de vue philosophique. Les mathématiques se distinguent profondément des sciences naturelles par le fait que ses “objets” *ne se situent ni dans l’espace, ni dans le temps* : un nombre, ou un théorème sont “éternels”, ils existent “toujours et partout”. Pourtant, les mathématiques s’appliquent dans la plupart des sciences naturelles, surtout en physique : dans son célèbre article *The unreasonable effectiveness of mathematics in the natural sciences* (lien) le physicien Eugene Wigner concède que c’est un grand mystère. Nous n’en dirons pas plus ici.

Il existe aussi quelques domaines scientifiques et philosophiques qui partagent ces traits caractéristiques avec les maths :

1. La *logique*, i.e., l’étude des règles formelles que doit respecter toute argumentation correcte. Elle s’applique aussi bien à l’extérieur des mathématiques –dans notre discours sur le “monde réel” –, qu’à l’intérieur des mathématiques, comme une sorte de “règlement interne de la maison des maths”; on parle aussi de *logique mathématique*. Ce dernier aspect nous occupera dans ce cours. Soulignons, cependant, que le rôle de la logique “dans le monde réel” (par exemple, dans les sciences expérimentales) peut être très différent du rôle de la logique en maths.
2. L’*informatique* – le concept d’*information* est difficile à cerner ; comme les objets mathématiques, l’information n’est pas vraiment localisée dans l’espace et le temps. L’*informatique théorique* est parfois considérée comme un sous-domaine des maths. Souvent, on assimile le “digital” avec les “maths discrètes”, opposé au “continu”, “analogue”.

§3 Les concepts fondateurs mathématiques, et l’infini. Comme dit ci-dessus, au début du XXe siècle, la théorie des ensembles a accédé au rôle de langage commun des fondements des maths. Une grande partie de ce cours sera consacrée à “apprendre ce langage”. Dans ce contexte, une distinction est d’importance capitale : il faut distinguer le *fini* de l’*infini*. Au niveau des *ensembles finis*, on n’a pas besoin d’une théorie très élaborée : une “théorie naïve” convient bien, voire mieux. Par contre, *la théorie des ensembles prend sa vraie forme et vraie puissance seulement quand elle parle des ensembles infinis*. Ces *ensembles infinis* posent le problème suivant : on ne peut pas les écrire, en un temps fini et sur un bout de papier (ou ordinateur) fini, sous forme d’une “liste complète” ou “fichier complet” ! Or, les “ensembles de nombres usuels”,

1. les nombres naturels (entiers naturels) $\mathbb{N} = \{0, 1, 2, 3, 4, 5, 6, \dots\}$,
2. les entiers relatifs $\mathbb{Z} = \{\dots, -2, -1, 0, 1, 2, \dots\}$,
3. les nombres rationnels $\mathbb{Q}$ (fractions $r = \frac{m}{n}$ avec m, n des entiers relatifs et $n \neq 0$),
4. les nombres réels $\mathbb{R}$ (“tous les nombres de la droite réelle”),
5. les nombres complexes $\mathbb{C}$ (“tous les nombres du plan d’Argand”)

sont bien infinis, et là, une théorie “naïve” n’est plus possible - il faut préciser : que signifie le symbole “ $\dots$ ” – par exemple, peut-on écrire $\mathbb{R} = \{0, e, \pi, \sqrt{2}, \dots\}$? Comment peut-on savoir que de tels ensembles “existent vraiment” si l’on ne peut pas les écrire sous forme d’une liste explicite ? Le mathématicien Hermann Weyl a caractérisé les maths ainsi : *les mathématiques sont la science de l’infini* ; elles commencent vraiment dès qu’on fait face au problème de l’infini. – Voici une piste de lecture recommandée pour le lecteur qui souhaite poursuivre ces réflexions : “Le commencement de l’infini” par D. Deutsch.

Première partie : les ensembles finis

Dans cette partie du cours, nous suivons une “approche naïve” des mathématiques : la notion d'*ensemble fini* fait abstraction de notre expérience de regrouper un nombre fini d'objets ; de les “mettre dans un sac” ; de compter des moutons, ou autre chose, en gravant un trait dans un morceau de bois pour chaque mouton, etc. En somme, nous suivons la démarche par laquelle nous commençons à compter et à raisonner sur des nombres dans l'enseignement primaire et secondaire. Néanmoins, dans ce cadre, on peut déjà démontrer quelques “théorèmes” : des énoncés mathématiques “non triviaux”, qui demandent une *preuve*. Pour l'instant, par “preuve” nous entendons “justification” ou “explication”, et il suffira de la rédiger en l'ange française usuelle, en respectant bien les règles d'une bonne rédaction en français. Cela veut dire, entre autres, qu'il faudra écrire des phrases grammaticalement correctes, avec une ponctuation correcte, et surtout, *qui ont un sens* : si vous utilisez des symboles, la phrase doit être compréhensible pour quelqu'un à qui vous la lisez à voix haute, mais qui ne voit pas la phrase écrite. Par ailleurs, ces règles de bon sens garderont leur pertinence pour toute votre carrière en maths, pendant et après les études !

Chapitre 1 : Ensembles et parties

§1 Ensembles. Voici comment Georg Cantor a défini sa notion d’“ensemble” :

Définition 1. Par ensemble, nous entendons toute collection d’objets de notre intuition ou de notre pensée, définis et distincts, ces objets étant appelés les éléments de l’ensemble. On utilise la notation $m \in M$ pour dire que m est un élément de l’ensemble M .

Définition 2. Un ensemble fini est donné par une liste finie de ses éléments, écrite entre accolades :

$$A = \{a_1, a_2, \dots, a_n\}.$$

Si les éléments $a_1, \dots, a_n$ sont deux à deux distincts (ce qui veut dire : $a_i \neq a_j$ lorsque $i \neq j$, pour $i, j = 1, \dots, n$), alors on dit que A est un ensemble de cardinal fini n , et on écrit

$$n = \text{card}(A) = |A|.$$

Notations et explications. Ainsi un ensemble fini est donné par la “liste” de ses éléments. Deux ensembles sont les mêmes s’ils ont les mêmes éléments. L’ordre dans lequel on présente cette liste ne joue aucun rôle : par exemple, si $M = \{a_1, a_2, \dots, a_n\}$, on aurait pu écrire aussi $M = \{a_n, a_{n-1}, \dots, a_1\}$. Aussi est-il possible d’écrire un élément plusieurs fois sans changer l’ensemble : ainsi $\{a, a\} = \{a\}$, ou $\{a, b, a\} = \{a, b\} = \{b, a, a, a\}$. Quand on écrit $M = \{a, b\}$, il faut bien préciser si les lettres a et b désignent le même élément ($a = b$; cardinal 1), ou non ($a \neq b$; cardinal 2). Attention à la typographie et à l’écriture : si un élément est noté a , il faut utiliser dans toute la suite le même symbole (reconnaissable), et non le changer en (par exemple) $A, \alpha, a, \mathbf{a}, \mathfrak{a}, \dots$

Exemple. L’ensemble “standard” de cardinal n sera noté

$$[[1, \dots, n]] := \{1, 2, \dots, n\}$$

(ou parfois juste $[[1, n]]$). Un ensemble de cardinal $2n + 1$ est, par exemple,

$$[[-n, n]] = \{-n, -n + 1, \dots, 0, \dots, n\}.$$

Définition 3. L’ensemble vide est l’ensemble qui ne contient aucun élément. Il est noté

$$\emptyset = \{\}.$$

§2 Parties d’un ensemble.

Définition 1. Soit M un ensemble. Une partie de M , ou sous-ensemble de M , est un ensemble A tel que tout élément de A est un élément de M . On écrit alors $A \subset M$.

Description d’une partie. Si M est fini, il y a deux façons de décrire ses parties :

(1) par une liste explicite. Par exemple, soit $M = \{1, 2, 3, 4, 5, 6, 7\}$, alors $A = \{2, 4, 6\}$ est une partie de M ;

(2) par une propriété : par exemple, l’ensemble A ci-dessus est l’ensemble des éléments de M qui sont pairs. On écrit :

$$A = \{x \in M \mid x \text{ est pair}\}.$$

De manière générale, soit $P(x)$ une certaine “propriété” que x peut avoir (“vrai”) ou non (“faux”). La partie des éléments de M vérifiant la propriété P (i.e., pour lesquels $P(x)$ est vraie) s’écrit

$$C = \{x \in M \mid P(x)\}$$

Dans ce contexte, on dit aussi que $P(x)$ est un prédicat (une proposition logique dépendant d’une “variable” x), et on utilise souvent les symboles

$\forall$ pour abrégé : *quel que soit* ou *pour tout*,

$\exists$ pour abrégé : *il existe*.

Par exemple, la partie A du point (1) s’écrit aussi

$$A = \{k \in M \mid \exists \ell \in \{1, 2, 3\} : k = 2\ell\}$$

On reparlera plus tard des règles de bonne utilisation de ces *quantificateurs* (“calcul des prédicats”).

Remarque 1. Deux ensembles sont égaux (notation : $A = B$), si, et seulement si, on a la double inclusion $A \subset B$ et $B \subset A$. Une inclusion $A \subset B$ est dite stricte s’il existe un élément $b \in B$ qui *n’est pas* un élément de A .

Remarque 2. L’ensemble vide fait partie de n’importe quel autre ensemble : $\emptyset \subset M$. Aussi, tout ensemble A est une partie de lui-même : $A \subset A$.

Exemple. L’ensemble $M = \{1\}$ a deux parties : $\emptyset$ et M .

Voici une liste de *toutes les parties* de l’ensemble $M = \{1, 2\}$: il y en a 4,

$$\emptyset, \{1\}, \{2\}, M.$$

Voici une liste de *toutes les parties* de l’ensemble $M = \{1, 2, 3\}$: il y en a 8,

$$\emptyset, \{1\}, \{2\}, \{3\}, \{1, 2\}, \{2, 3\}, \{1, 3\}, M.$$

TD : faire la même chose pour $M = \{1, 2, 3, 4\}$.

Définition 2. Soit M un ensemble (fini). On note :

$\mathcal{P}(M)$ l’ensemble de toutes les parties de M (ensemble des parties),

$\mathcal{P}_k(M)$ l’ensemble de toutes les parties $A \subset M$ telles que $\text{card}(A) = k$.

Théorème 1. Soit M un ensemble fini de cardinal n . Alors $\mathcal{P}(M)$ est un ensemble fini de cardinal 2^n :

$$\text{card}(\mathcal{P}(M)) = 2^{\text{card}(M)}$$

Preuve. Soit $M = \{a_1, a_2, \dots, a_n\}$, et $A \subset M$. Combien de possibilités de choisir A y a-t-il ? Pour chaque élément a_i de M , il existe deux possibilités :

- soit, $a_i \in A$,
- soit, $a_i \notin A$ (la notation $x \notin M$ signifie : x n’appartient pas à M).

Comme M contient n éléments, il y a n choix (chacun binaire : deux issues possibles) à effectuer ; cela donne $2 \times \dots \times 2$ (n fois) issues possibles, soit 2^n issues. Graphiquement, on peut représenter la séquence des choix par un arbre : la première étage (choix concernant a_1) a deux branches, la deuxième (choix concernant a_2) a 4 branches, et ainsi de suite.

Définition 3. On note $\binom{n}{k}$ le cardinal de l’ensemble $\mathcal{P}_k([1, \dots, n])$; autrement dit, c’est le nombre de parties de cardinal k dans un ensemble de cardinal n ; on appelle ce nombre “ k parmi n ”.

Théorème 2.

1. Si $k > n$, alors $\binom{n}{k} = 0$;
2. $\binom{n}{0} = 1$ et $\binom{n}{n} = 1$;
3. si $k \in [[1, \dots, n]]$,

$$\binom{n+1}{k} = \binom{n}{k} + \binom{n}{k-1}$$

4. pour tout $k \in [[0, \dots, n]]$,

$$\binom{n}{k} = \binom{n}{n-k}$$

Preuve. 1. Intuitivement, il est clair qu'aucune partie de M ne peut avoir plus d'éléments que M . Pour donner une preuve plus formelle de ce fait, il faudrait anticiper quelques notions sur les *applications* (voir plus tard).

2. La seule partie de M ayant autant d'éléments que M est M lui-même, et la seule partie ayant 0 éléments est l'ensemble vide.

3. Soit $A \subset [[1, \dots, n+1]]$ une partie de cardinal k . Alors il existe deux possibilités :

- (a) soit, l'élément $n+1$ appartient à A . Dans ce cas on peut choisir les $k-1$ autres éléments librement dans $[[1, \dots, n]]$; il y a donc $\binom{n}{k-1}$ possibilités dans ce cas ; ou bien,
- (b) tous les éléments de A appartiennent à $[[1, \dots, n]]$: cela fait $\binom{n}{k}$ possibilités.

Comme les cas (a) et (b) sont exhaustifs et s'excluent mutuellement, le nombre total de possibilités est la somme des deux, soit $\binom{n}{k} + \binom{n}{k-1}$.

4. Choisir k éléments parmi n (on peut les "colorier en blanc") revient exactement à choisir les $n-k$ éléments restants (on peut les "colorier en noir").

Le triangle de Pascal. La propriété 3 du théorème 2 permet de calculer les coefficients $\binom{n}{k}$ de proche en proche : on les retrouve pour $k = 0, \dots, n$ dans la $(n+1)$ -ième ligne du triangle de Pascal suivant. Par exemple, on y trouve que $\binom{5}{3} = 10$:

$$\begin{array}{ccccccc} & & & & 1 & & & \\ & & & & 1 & & 1 & \\ & & & 1 & & 2 & & 1 \\ & & 1 & & 3 & & 3 & & 1 \\ & 1 & & 4 & & 6 & & 4 & & 1 \\ 1 & & 5 & & 10 & & 10 & & 5 & & 1 \\ 1 & & 6 & & 15 & & 20 & & 15 & & 6 & & \dots \end{array}$$

§3 Des formules remarquables. Tout le monde connaît la formule remarquable

$$(a+b)^2 = a^2 + 2ab + b^2.$$

Pour la prouver, on utilise les propriétés suivantes :

Propriétés de la somme et du produit usuels (de nombres réels ou complexes) :

- (1) distributivité : $a(b+c) = ab+ac$
- (2) commutativité : $a+b = b+a$ et $ab = ba$,
- (3) associativité : $(a+b)+c = a+(b+c)$ et $(ab)c = a(bc)$.

Dans le calcul suivant, pour chaque égalité, notez à la marge laquelle de ces propriétés est utilisée pour la justifier :

$$\begin{aligned}
 (a+b)^2 &= (a+b)(a+b) \\
 &= a(a+b) + b(a+b) \\
 &= aa + ab + ba + bb \\
 &= a^2 + 2ab + b^2.
 \end{aligned}$$

Procédez de la même façon pour

$$\begin{aligned}
 (a+b)^3 &= (a+b)(a+b)(a+b) \\
 &= a(a+b)(a+b) + b(a+b)(a+b) \\
 &= a(aa + ab + ba + bb) + b(aa + ab + ba + bb) \\
 &= aaa + aab + aba + baa + abb + bab + bba + bbb \\
 &= a^3 + 3a^2b + 3ab^2 + b^3.
 \end{aligned}$$

Théorème 3 (Formule du binôme de Newton). *Soit a, b deux nombres réels ou complexes et n un nombre naturel. Alors*

$$(a+b)^n = a^n + na^{n-1}b + \binom{n}{2}a^{n-2}b^2 + \dots + \binom{n}{k}a^{n-k}b^k + \dots + nab^{n-1} + b^n.$$

Preuve. En utilisant la distributivité et l'associativité, on développe le produit $(a+b)^n = (a+b) \cdots (a+b)$. On obtient 2^n termes. Plus précisément, pour chaque partie $A \subset [[1, n]]$, il y a un terme de la forme

$$c_1 \cdot c_2 \cdot \dots \cdot c_n,$$

où $c_i = b$ si $i \in A$, et $c_i = a$ si $i \notin A$. On utilisant la commutativité et l'associativité, si $\text{card}(A) = k$, ce terme vaut $b^k a^{n-k}$. Il y a exactement $\binom{n}{k}$ termes de ce type ; ainsi ils contribuent $\binom{n}{k} b^k a^{n-k}$ au résultat final. On peut faire cela pour tout $k = 0, \dots, n$, et la somme de tout ces termes vaut donc $(a+b)^n$. (Si vous avez des difficultés à suivre l'argument dans le cas de n général, détaillez-le encore pour le cas $n = 4$.)

Le signe somme. Si $a_1, \dots, a_n$ sont des nombres réels ou complexes, on abrège

$$\sum_{i=1}^n a_i := a_1 + a_2 + \dots + a_n.$$

L'index i est "muet" : on peut utiliser n'importe quelle autre lettre, hormis n qui est déjà pris. Avec cette convention, la formule du binôme de Newton s'écrit :

$$(a+b)^n = \sum_{k=0}^n \binom{n}{k} a^{n-k} b^k.$$

Chapitre 2 : Opérations sur les ensembles

Etant donné des parties $A, B, C, \dots$ de M , on peut en fabriquer d'autres par les opérations intersection, union, différence, complémentaire. Aussi, à partir de deux ensembles M et N , on peut former un autre, leur produit cartésien, noté $M \times N$. Ces opérations vérifient certaines règles que nous allons étudier.

§1 Opérations sur les parties d'un ensemble M .

Définition 1. Soit A et B deux parties d'un ensemble M . Leur intersection $A \cap B$ est définie par

$$A \cap B := \{x \in M \mid x \in A \text{ et } x \in B\}.$$

Cela veut dire que $x \in A \cap B$ est vrai si, et seulement si, à la fois $x \in A$ et $x \in B$. On définit l'union $A \cup B$ par

$$A \cup B := \{x \in M \mid x \in A \text{ ou } x \in B\}.$$

Cela veut dire que $x \in A \cup B$ est vrai si, et seulement si, x appartient à A , ou à B , ou aux deux à la fois (il s'agit du "ou logique"; on parlera plus tard de la logique en général). Puis, la différence ensembliste de B et de A , ou : B privé de A , $B \setminus A$, est défini par

$$B \setminus A := \{x \in M \mid x \in B \text{ et } x \notin A\}.$$

Ainsi $x \in B \setminus A$ est vrai si, et seulement si, x appartient à B , mais non à A . Enfin, le complémentaire de A dans M est l'ensemble des éléments qui n'appartiennent pas à A :

$$A^c := \{x \in M \mid x \notin A\} = M \setminus A$$

On dit que A et B sont disjointes si elles n'ont aucun élément en commun : $A \cap B = \emptyset$.

Représentation logique. On peut résumer ces définitions par la table de vérité suivante : un symbole 0 dans la colonne marquée $x \in A$ veut dire que cette propriété est fausse (donc, $x \notin A$ est vraie), et un symbole 1 qu'elle est vraie (donc, $x \in A$ est vraie), etc. Par exemple, $x \in A \cap B$ est vraie dans le seul cas où $x \in A$ et $x \in B$, et fausse dans les autres 3 cas :

$x \in A$	$x \in B$	$x \in A \cap B$	$x \in A \cup B$	$x \in B \setminus A$
0	0	0	0	0
1	1	1	1	0
0	1	0	1	1
1	0	0	1	0

Représentation graphique. En négligeant des propriétés "géométriques" ou autres, on représente des ensembles souvent par des "diagrammes en patates" (*diagrammes de Venn*), en coloriant ou en hachurant la partie $A \cap B$, etc. Cette représentation est souvent utile pour "fixer les idées" (arguments heuristiques), mais ne constitue jamais une preuve mathématique !

Théorème 1. Les opérations ensemblistes vérifient les propriétés suivantes :

1. *associativité* : $(A \cap B) \cap C = A \cap (B \cap C)$, $(A \cup B) \cup C = A \cup (B \cup C)$,
2. *commutativité* : $A \cap B = B \cap A$, $A \cup B = B \cup A$,
3. *distributivité* : $(A \cap B) \cup C = (A \cup C) \cap (B \cup C)$, $(A \cup B) \cap C = (A \cap C) \cup (B \cap C)$,
4. *lois de complément* : $(A^c)^c = A$, $A \cap A^c = \emptyset$, $A \cup A^c = M$,
5. *lois de de Morgan* : $(A \cup B)^c = A^c \cap B^c$, $(A \cap B)^c = A^c \cup B^c$.

Preuve. Il est conseillé d'illustrer chacune des propriétés par un diagramme de Venn. Pour donner une preuve formelle, il faut utiliser les propriétés logiques : on distingue tous les cas possibles – s'il y a deux ensembles en jeu, il y a 4 cas, s'il y en a trois, cela fait 8 cas. Par exemple, pour la première loi de de Morgan, dressons la "table de vérité" avec les 4 cas possibles :

$x \in A$	$x \in B$	$x \in A \cup B$	$x \in (A \cup B)^c$	$x \in A^c$	$x \in B^c$	$x \in A^c \cap B^c$
0	0	0	1	1	1	1
1	1	1	0	0	0	0
0	1	1	0	1	0	0
1	0	1	0	0	1	0

Les valeurs des colonnes 4 et 7 coïncident, ce qui prouve que les conditions décrivant ces ensembles sont les mêmes ; donc ces ensembles sont égaux. Les autres points sont démontrés de la même façon.

Définition 2. Soit $A, B \subset M$. Leur différence symétrique est définie par

$$A \Delta B := \{x \in M \mid x \in A \text{ ou bien } x \in B\},$$

où "ou bien" est le "ou exclusif" : $x \in A \Delta B$ si et seulement si x appartient à $A \setminus B$, ou à $B \setminus A$ (mais non à A et B à la fois). Exercice de TD : représenter ceci par un diagramme de Venn ; utilisant la table de vérité, montrer que $A \Delta B = (A \cup B) \setminus (A \cap B)$, et établir certaines propriétés de cette construction.

Théorème 2. Soit M un ensemble fini et $A, B \subset M$.

1. Si A et B sont disjointes, alors $\text{card}(A \cup B) = \text{card}(A) + \text{card}(B)$,
2. Si $A_1, \dots, A_k$ sont des parties de M , deux à deux disjointes, alors

$$\text{card}(A_1 \cup A_2 \cup \dots \cup A_k) = \sum_{i=1}^k \text{card}(A_i).$$

3. En général, $\text{card}(A \cup B) = \text{card}(A) + \text{card}(B) - \text{card}(A \cap B)$.

Preuve. 1. Si $A \cap B = \emptyset$, tout élément de $A \cup B$ appartient, soit à A , soit à B (et non au deux en même temps). Ainsi le nombre d'éléments de $A \cup B$ est bien la somme des nombres d'éléments de A et de B .

2. Tout élément de $E = A_1 \cup A_2 \cup \dots \cup A_k$ appartient à un seul des ensembles A_i , $i = 1, \dots, k$. Ainsi le nombre d'éléments de E est bien la somme des nombres d'éléments de A_1, A_2 , jusqu'à A_k .

3. Remarquons d'abord que

$$A = (A \setminus B) \cup (A \cap B)$$

est une réunion disjointe ; donc d'après 1., $\text{card} A = \text{card}(A \setminus B) + \text{card}(A \cap B)$. Ensuite, remarquons que

$$A \cup B = (A \setminus B) \cup (A \cap B) \cup (B \setminus A)$$

est une réunion disjointe ; d'où d'après 2.,

$$\begin{aligned} \text{card}(A \cup B) &= (\text{card}(A) - \text{card}(A \cap B)) + \text{card}(A \cap B) + (\text{card}(B) - \text{card}(A \cap B)) \\ &= \text{card}(A) + \text{card}(B) - \text{card}(A \cap B). \end{aligned}$$

(Exercice, TD 1 : généraliser au cas d'une réunion de trois ensembles, ou plus.)

Définition 3. On dit que des parties $A_1, \dots, A_k$ de M forment une partition de M si

- les A_i sont tous non-vides ;
- elles sont deux à deux disjointes ($\forall i \neq j : A_i \cap A_j = \emptyset$) ;
- leur réunion est $M : A_1 \cup \dots \cup A_k = M$.

D'après le théorème, on a alors $\text{card}(M) = \sum_{i=1}^k \text{card}(A_i)$.

Idée : les ensembles A_i sont des “casiers”, ou des “ tiroirs”, et chaque élément de M est rangé dans un unique casier. – Exemple : les groupes de TD de la promo de L1 forment une partition de cette promo. L'effectif total est donc la somme des effectifs des groupes de TD.

Le “principe des tiroirs”. Soit $\{A_1, \dots, A_k\}$ une partition de M , et $\text{card}(M) = n$. Alors :

1. il ne peut pas y être plus de casiers que d'éléments : il n'est pas possible que $k > n$;
2. si $k = n$, alors les A_i sont tous des singletons (i.e., elles sont de cardinal 1 : chaque casier contient un unique élément) ;
3. et si $k < n$, il existe au moins un index i tel que $\text{card}(A_i) > 1$ (s'il y a moins de casiers que d'objets, alors au moins un casier contient plus qu'un objet).

§2 Produit cartésien.

Définition 1. Soit M_1 et M_2 deux ensembles. Le produit cartésien $M_1 \times M_2$ est l'ensemble des couples (x, y) avec $x \in M_1$ et $y \in M_2$. La règle de base opératoire concernant les couples est de savoir que

$\text{deux couples } (x, y) \text{ et } (x', y') \text{ sont égaux si et seulement si : } x = x' \text{ et } y = y'.$

Cela s'écrit aussi : $(x_1, x_2) = (x'_1, x'_2)$ ssi $(x_1 = x'_1 \text{ et } x_2 = x'_2)$.

Retenir : Contrairement aux ensembles $\{x, y\}$, pour les couples (x, y) l'ordre est important : en général, $(x, y) \neq (y, x)$ (le seul cas où il y a égalité est $x = y$).

Dans le cas $M_1 = M_2 = M$, on note $M^2 = M \times M$.

Remarque. Nous admettons que le produit cartésien de deux ensembles existe toujours. Intuitivement, cela semble bien crédible. En “théorie axiomatique des ensembles”, on étudie de façon plus approfondie la question comment justifier ce genre d'hypothèses.

Théorème 1. Si M_1 est de cardinal fini n et M_2 de cardinal fini m , alors $M_1 \times M_2$ est de cardinal fini $n \cdot m$:

$$\text{card}(M_1 \times M_2) = \text{card}(M_1) \cdot \text{card}(M_2).$$

Preuve. Pour choisir un élément $(x, y) \in M_1 \times M_2$, on a *d'abord* le choix entre les n éléments $x_1, \dots, x_n$ de M_1 ; ce choix fixé, on a *ensuite* le choix entre les m éléments $y_1, \dots, y_m$ de M_2 . On pourra représenter ces choix par un *arbre* : la première étage a n branches; chaque branche se ramifie en m branches au deuxième étage; cela fait donc $n \cdot m$ branches au total.

Représentation graphique. En théorie des probabilités, on utilise parfois une représentation par des arbres; par contre, dans ce cours, il sera plus approprié de représenter $M_1 \times M_2$ par un schéma rectangulaire – une sorte de tableau à double entrée avec n lignes et m colonnes, et l'élément (x, y) au croisement de la “colonne sur x ” avec “ligne au niveau y ”.

Exemple. La partie

$$\Delta := \{(x, y) \in M \times M \mid x = y\} = \{(x, x) \mid x \in M\}$$

est alors représentée par la *diagonale* du schéma (carré, car ici $M_1 = M_2$). Pour d'autres exemples, voir TD 2, exercice 1.

Définition 2. Soit $M_1, \dots, M_n$ des ensembles. Le produit cartésien $M_1 \times \dots \times M_n$, ou $\times_{i=1}^n M_i$, est l'ensemble des n -uplets $(x_1, \dots, x_n)$ avec $\forall i = 1, \dots, n : x_i \in M_i$. Par *définition*, deux n -uplets $(x_1, \dots, x_n)$ et $(x'_1, \dots, x'_n)$ sont égaux si et seulement si $\forall i = 1, \dots, n : x_i = x'_i$. Dans le cas où tous les M_i sont le même ensemble M , on note $M^n = M \times \dots \times M$.

Surtout dans le cas où $M = \mathbb{R}$, on note souvent les n -uplets sous forme de colonne $\begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix}$.

Théorème 2. Les ensembles $M_1 \times M_2 \times M_3$ et $(M_1 \times M_2) \times M_3$ et $M_1 \times (M_2 \times M_3)$ sont “les mêmes”, c'est-à-dire, en théorie des ensembles, ils sont “indistinguishables” : nous pouvons alors écrire

$$M_1 \times (M_2 \times M_3) = M_1 \times M_2 \times M_3 = (M_1 \times M_2) \times M_3.$$

Preuve. Les paires $((x, y), z)$ se comportent exactement comme les triplets : on a $((x, y), z) = ((x', y'), z')$ ssi $[(x, y) = (x', y') \text{ et } z = z']$, ssi $[x = x' \text{ et } y = y' \text{ et } z = z']$.

L'égalité $(M_1 \times M_2) \times M_3 = M_1 \times (M_2 \times M_3)$ se traduit en disant que le produit cartésien est “associatif”. (En revanche, on prendra soin de distinguer $M_1 \times M_2$ et $M_2 \times M_1$.)

Théorème 3. Si, pour $i = 1, \dots, m$, l'ensemble M_i est de cardinal fini n_i , alors

$$\text{card}(M_1 \times \dots \times M_m) = n_1 \cdots n_m.$$

En particulier, M^m est de cardinal $(\text{card}(M))^m$.

Chapitre 3 : Applications, fonctions

On peut dire sans exagérer que les *fonctions* et *applications* sont l'objet d'étude le plus important des maths : dans le théâtre des maths, la théorie des ensembles met en place la scène, mais la pièce de théâtre elle-même est joué par les fonctions.

§1 Définitions fondamentales.

Définition 1. Une application $f : A \rightarrow B, a \mapsto f(a)$ entre deux ensembles A et B est la donnée, pour tout élément $a \in A$, d'un élément dit image de a par f et noté $f(a) \in B$. Souvent, on utilise aussi le mot fonction, en particulier si $B = \mathbb{R}$ ("fonction réelle"); et si $A = B$, on parle aussi d'une transformation (de A). Le symbole désignant la variable est "muet" : on peut utiliser n'importe quelle lettre – par exemple, $f : A \rightarrow B, x \mapsto f(x)$ est la même application que celle écrite ci-dessus.

Noter bien : chaque $a \in A$ a une unique image $f(a)$ dans B ; mais il n'est pas toujours vrai que chaque $b \in B$ soit image d'un élément a de A ; et s'il l'est, il peut y être plusieurs éléments a avec cette propriété :

Définition 2. Soit $f : A \rightarrow B$ une application, et $b \in B$. Tout élément $a \in A$ tel que $b = f(a)$ est dit un antécédent de b .

Définition 3 (injectif, surjectif, bijectif). On dit qu'une application $f : A \rightarrow B$ est

1. surjective si tout élément $b \in B$ a (au moins) un antécédent :
$$\forall b \in B : \exists a \in A : b = f(a);$$
2. injective si tout élément de B a au plus un antécédent :
$$\forall a, a' \in A : [\text{si } a \neq a', \text{ alors } f(a) \neq f(a')] ; \text{ ce qui équivaut à :}$$
$$\forall a, a' \in A : [\text{si } f(a) = f(a'), \text{ alors } a = a'] ;$$
3. bijective si elle est surjective et injective : tout $b \in B$ a exactement un antécédent. En utilisant le quantificateur $\exists!$ ("il existe un unique") cela s'écrit,
$$\forall b \in B : \exists! a \in A : b = f(a).$$

Exemple 1. L'application $f : \mathbb{R} \rightarrow \mathbb{R}, x \mapsto x^2$ n'est ni injective (car $b = 1$ a deux antécédents : $a = 1$ et $a' = -1$) ni surjective (car -2 n'a pas d'antécédent). L'application $f : \mathbb{R} \rightarrow \mathbb{R}^+, x \mapsto x^2$ est surjective (car tout nombre réel positif admet une racine carrée, son antécédent), mais non injective. L'application $f : \mathbb{R}^+ \rightarrow \mathbb{R}^+, x \mapsto x^2$ est bijective (l'unique antécédent de $r \in \mathbb{R}^+$ est sa racine carrée $\sqrt{r} \in \mathbb{R}^+$). On dira que l'application $g : \mathbb{R}^+ \rightarrow \mathbb{R}^+, r \mapsto \sqrt{r}$ est l'*application réciproque* ou *inverse* de f :

Définition 4. Si $f : A \rightarrow B$ est bijective, on définit l'application réciproque

$$f^{-1} : B \rightarrow A, b \mapsto f^{-1}(b)$$

en associant à $b \in B$ son unique antécédent dans A , donc $f^{-1}(b) = a$ où $f(a) = b$. Il s'ensuit que $f^{-1}(f(a)) = a$ et $f(f^{-1}(b)) = b$.

Exemple 2. Si A et B sont des ensembles finis, une application peut être décrite par la liste des images $f(a)$ quand a parcourt A . Par exemple, si $A = B = \{1, 2\}$, alors il existe 4 applications $f : A \rightarrow A$ possibles, à savoir, $f_i, i = 1, 2, 3, 4$, données par :

	$f_i(1)$	$f_i(2)$
f_1	1	1
f_2	2	2
f_3	1	2
f_4	2	1

Parmi ces applications, f_3 et f_4 sont bijectives, et f_1 et f_2 ne sont ni injectives, ni surjectives. En fait, elles sont *constantes* :

Définition 5. Une application $f : A \rightarrow B$ est dite constante s'il existe $c \in B$ tel que $f(x) = c$ pour tout $x \in A$.

Généralisant cet exemple, si A et B sont finis, on peut toujours énumérer toutes les applications possibles de A vers B :

Théorème 1. Soit A un ensemble fini de cardinal n et B de cardinal m . Alors il existe exactement m^n applications $f : A \rightarrow B$.

Preuve. Soit $A = \{a_1, \dots, a_n\}$ et $B = \{b_1, \dots, b_m\}$. Nous avons m possibilités de choisir $f(a_1)$ dans B , puis encore m possibilités de choisir $f(a_2)$, etc. On peut de nouveau représenter ces choix par un arbre, cette fois-ci avec n étages, et m ramifications à chaque étage. Au total, cela fait $m \cdots m = m^n$ choix possibles.

Théorème 2. Soit $n = \text{card}(A)$ et $m = \text{card}(B)$.

1. Si $n > m$, alors il n'existe aucune application injective $f : A \rightarrow B$;
2. si $n < m$, alors il n'existe aucune application surjective $f : A \rightarrow B$;
3. si $n \leq m$, alors il y a exactement

$$m(m-1) \cdots (m-n+1)$$

applications injectives $f : A \rightarrow B$;

4. si $n = m$, il y a exactement

$$n! := n(n-1) \cdots 2 \cdot 1$$

applications bijectives $f : A \rightarrow B$.

Preuve. 1. On utilise le “principe des tiroirs” : chaque élément de B correspond à une case qu'on remplit par ses antécédents ; il y a plus d'éléments dans A que de cases ; donc au moins une case contient au moins deux éléments : donc f ne peut pas être injective.

2. Argument analogue : il y a moins d'éléments dans A que de cases ; donc au moins une case doit être vide : cet élément n'a pas d'antécédent : donc f ne peut pas être surjective.

3. Soit $A = \{a_1, \dots, a_n\}$. Pour définir une permutation $f : A \rightarrow B$, nous avons d'abord m possibilités pour définir $f(a_1)$, disons $f(a_1) = b$; une fois ce choix fixé, nous pouvons choisir $f(a_2)$ librement parmi les $m-1$ éléments différents de b ; ce choix fixé, il reste $m-2$ possibilités pour choisir $f(a_3)$, et ainsi de suite – pour la dernière valeur $f(a_n)$ il reste $m-n+1$ choix. Ensemble, cela donne un arbre de décision à n étages, avec $m(m-1) \cdots (m-n+1)$ branches au niveau final.

4. D'après 3., on sait déjà qu'il y a exactement $n!$ applications injectives $f : A \rightarrow B$. Or, toujours d'après le principe des tiroirs, chacune de ces $n!$ applications sera alors aussi surjective : le nombre de tiroirs est n ; chaque tiroir contient au plus un élément ; comme $m = n$, chaque tiroir contient alors exactement un élément, qui est l'unique antécédent.

Remarque. Si $n > m$, il est plus difficile de compter les applications surjectives. Cf. TD 2 pour les cas (relativement faciles) $m = 2$ et $n = m + 1$.

§2 Composée d'applications.

Définition 1. Soient $f : A \rightarrow B$, $a \mapsto f(a)$ et $g : B \rightarrow C$, $b \mapsto g(b)$ deux applications. On définit une application $g \circ f$ (lire “ g rond f ” ou “ g après f ”), la composée de f et de g

$$g \circ f : A \rightarrow C, \quad a \mapsto (g \circ f)(a) = g(f(a))$$

Théorème 1. La composition d'applications est associative : si $f : A \rightarrow B$ et $g : B \rightarrow C$ et $h : C \rightarrow D$, alors $\boxed{h \circ (g \circ f) = (h \circ g) \circ f}$.

Preuve. Pour tout $a \in A$, $(h \circ (g \circ f))(a) = h(g \circ f(a)) = h(g(f(a))) = ((h \circ g) \circ f)(a)$.

Exemple 1. Si $A = B$ et $f, g : A \rightarrow A$, alors on peut composer : $g \circ f$ est encore une application $A \rightarrow A$. Par exemple, si $A = B = \{1, 2\}$, en utilisant les notations de l'exemple ci-dessus,

$$\begin{aligned} f_1 \circ f_4(1) &= f_1(2) = 1, \quad f_1 \circ f_4(2) = f_1(1) = 1, \text{ donc } f_1 \circ f_4 = f_1; \\ f_4 \circ f_1(1) &= f_4(1) = 2, \quad f_4 \circ f_1(2) = f_4(1) = 2, \text{ donc } f_4 \circ f_1 = f_2. \end{aligned}$$

Conclusion : la composée **n'est pas commutative** – en général, $g \circ f$ est différent de $f \circ g$! Exercice (cf TD 3) : compléter la table

$\circ$	f_1	f_2	f_3	f_4
f_1				f_2
f_2				
f_3				
f_4	f_1			

On constatera que $f_3 \circ f_i = f_i = f_i \circ f_3$ pour tout $i = 1, 2, 3, 4$:

Définition 2. Pour tout ensemble M , l'application

$$\text{id}_M : M \rightarrow M, \quad x \mapsto x = \text{id}_M(x)$$

s'appelle l'(application) identité de M .

Proposition 1. Pour toute application $f : A \rightarrow B$, on a $\boxed{f \circ \text{id}_A = f = \text{id}_B \circ f}$.

Théorème 2. Pour une application $f : A \rightarrow B$, les propriétés suivantes sont équivalentes :

- (1) f est bijective ;
- (2) il existe une application $g : B \rightarrow A$ telle que $g \circ f = \text{id}_A$ et $f \circ g = \text{id}_B$.

Dans ce cas, g est égale à l'application réciproque f^{-1} définie ci-dessus.

Preuve. Nous avons vu ci-dessus que (1) implique (2) (en posant $g = f^{-1}$).

Montrons que (2) implique (1). Supposons donc que $g \circ f = \text{id}_A$ et $f \circ g = \text{id}_B$.

Montrons que f est injective : si $f(a) = f(a')$, alors $g(f(a)) = g(f(a'))$, car g est une application, et donc $a = a'$ car $g \circ f = \text{id}$.

Montrons que f est surjective : soit $b \in B$; posons $a := g(b)$, alors $f(a) = f(g(b)) = b$. Donc tout $b \in B$ admet un antécédent.

Proposition 2. Si $h : B \rightarrow C$ et $f : A \rightarrow B$ sont bijectives, alors $h \circ f$ est bijective, et

$$(h \circ f)^{-1} = f^{-1} \circ h^{-1}.$$

En effet : $(h \circ f) \circ (f^{-1} \circ h^{-1}) = \dots = \text{id}_C$ et $(f^{-1} \circ h^{-1}) \circ (h \circ f) = \dots = \text{id}_A$, donc d'après le théorème, $h \circ f$ est bijective, et $f^{-1} \circ h^{-1}$ est son inverse.

§3 Permutations.

Définition 1. Une permutation d'un ensemble M est une application bijective $f : M \rightarrow M$. L'ensemble $\mathfrak{S}(M)$ des permutations de M est appelé le groupe symétrique de M .

Théorème 1. Les permutations d'un ensemble M forment un groupe $G := \mathfrak{S}(M)$, avec la composée $\circ$ comme loi de groupe et $e = \text{id}_M$ comme élément neutre.

Définition 2. Un groupe est un ensemble G muni d'une "loi de groupe" qui associe à un couple (g, h) un troisième élément noté $g \cdot h$ (ou juste gh , ou $g * h$, ou encore autrement...), tel que :

- la loi est associative : $\forall f, g, h \in G : (f \cdot g) \cdot h = f \cdot (g \cdot h)$,
- il existe un élément neutre $e \in G$ tel que : $\forall g \in G : e \cdot g = g = g \cdot e$,
- chaque $g \in G$ admet un élément inverse g^{-1} vérifiant $g \cdot g^{-1} = e = g^{-1} \cdot g$.

Notation. Si $M = \{1, \dots, n\}$, le groupe $\mathfrak{S}(M)$ est noté $\mathfrak{S}_n$.

Attention : quelle que soit la notation du produit d'un groupe, il n'est pas toujours *commutatif* : en général, $g \cdot h \neq h \cdot g$. La théorie des groupes, et en particulier des groupes de permutations $\mathfrak{S}_n$, est un objet important des études de maths en Licence. Voir TD 3 pour une première approche.

Exemple. Rappelons que le cardinal de $\mathfrak{S}_n$ est $n!$. Si $n = 1$, la seule permutation est $f = \text{id}$, donc $\mathfrak{S}_1 = \{e\}$ est le "groupe trivial".

Si $n = 2$, on peut écrire $\mathfrak{S}_2 = \{\text{id}, f\}$ avec $f : 1 \mapsto 2, 2 \mapsto 1$ ("transposition"). Sa table de groupe est très simple, et ce groupe est encore commutatif.

Si $n = 3$, on fera en TD une liste des 6 permutations de $\{1, 2, 3\}$, et on calculera la *table de groupe* de $\mathfrak{S}_3$. Ce groupe n'est pas commutatif (et c'est en fait *le plus petit groupe non-commutatif possible*).

Le groupe $\mathfrak{S}_4$ a 24 éléments, et $\mathfrak{S}_5$ en a 120. Ces groupes seront étudiés dans les cours d'algèbre de L2 et L3.

Théorème 2. Pour tout $k = 0, \dots, n$, le coefficient $\binom{n}{k}$ est donné par les formules

$$\binom{n}{k} = \frac{n(n-1)(n-2) \cdots (n-k+1)}{k!} = \frac{n!}{k!(n-k)!}.$$

Preuve. Choisir un ensemble $\{i_1, \dots, i_k\}$ de cardinal k dans $\{1, \dots, n\}$ se fait en deux étapes : d'abord, choisir le k -uplet $(i_1, \dots, i_k)$ (d'après le §1, théorème 2, le nombre de tels choix est de $n(n-1) \cdots (n-k+1)$) ; puis, comme chacune des $k!$ permutations de $\{1, \dots, k\}$ donne le même ensemble, il faut diviser cette quantité par $k!$ pour obtenir le nombre de parties de cardinal k dans $\{1, \dots, n\}$.

Exemple. Loto (ancienne version) : choisir 6 parmi 49 chiffres donne $\binom{49}{6} = \frac{49 \cdot 48 \cdot 47 \cdot 46 \cdot 45 \cdot 44}{2 \cdot 3 \cdot 4 \cdot 5 \cdot 6} = 49 \cdot 2 \cdot 47 \cdot 46 \cdot 3 \cdot 22 = 13983816$ possibilités. (Attention : toujours simplifier numérateur et dénominateur avant d'utiliser votre calculette!)

§4 D'autres éléments de langage : image directe et image réciproque d'une partie. Soit $f : A \rightarrow B$ une application.

Définition 1. Soit $U \subset A$ une partie. L'image directe $f(U)$ est la partie de B définie par

$$f(U) := \{b \in B \mid \exists a \in U : b = f(a)\} = \{f(a) \mid a \in U\}.$$

L'image directe $f(A)$ est l'image de f , notée aussi $\text{im}(f)$.

Noter : f est surjective si et seulement si $\text{im}(f) = B$.

TD : on a toujours $f(U \cup V) = f(U) \cup f(V)$, mais $f(U \cap V)$ est en général différent de $f(U) \cap f(V)$. Aussi, on a toujours $(g \circ f)(U) = g(f(U))$.

Définition 2. Soit $W \subset B$ une partie. L'image réciproque $f^{-1}(W)$ est la partie de A

$$f^{-1}(W) := \{a \in A \mid f(a) \in W\}.$$

Attention : ne pas confondre ici le symbole f^{-1} avec une “application réciproque” – cette notation est définie même si f n'est pas bijective ; mais si f est bijective, il s'agit en effet de l'image directe de W par f^{-1} .

TD : on a toujours $f^{-1}(W \cup V) = f^{-1}(W) \cup f^{-1}(V)$, et $f^{-1}(W \cap V) = f^{-1}(W) \cap f^{-1}(V)$.

Définition 3. Si $W = \{b\}$ est un singleton (partie contenant un seul élément), alors on appelle $f^{-1}(\{b\})$ (= ensemble des antécédents de b) aussi la b -fibre de f .

Noter : b appartient à $\text{im}(f)$ si et seulement si la b -fibre de f est non-vide.

Et f est injective si et seulement si chaque fibre contient au plus un élément.

Théorème 1. Soit $f : A \rightarrow B$ une application, et A et B des ensembles finis. Alors les ensembles $f^{-1}(b)$, lorsque b parcourt $\text{im}(f)$, forment une partition de A . Il s'ensuit que

$$\text{card}(A) = \sum_{b \in \text{im}(f)} \text{card}(f^{-1}(\{b\})).$$

Preuve. Les fibres sont deux à deux disjointes : soit $a \in f^{-1}(\{b\}) \cap f^{-1}(\{b'\})$; alors $b = f(a) = b'$, ainsi, si deux fibres ont une intersection non-vide, elles sont égales. La réunion des fibres est A , car si $a \in A$, alors a appartient à la fibre $f^{-1}(\{f(a)\})$. Ainsi les fibres forment une partition. Si A et B sont finis, l'affirmation sur les cardinalités s'ensuit par le théorème 2 du §1 chap. 1.

Corollaire. Soit A et B des ensembles finis et $f : A \rightarrow B$ une application.

1. Si f est injective, alors $\text{card}(A) = \text{card}(\text{im } f) \leq \text{card}(B)$.
2. Si $\text{card}(A) > \text{card}(B)$, alors f ne peut pas être injective.
3. Si f est surjective, alors $\text{card}(A) = \sum_{b \in B} \text{card}(f^{-1}(\{b\})) \geq \text{card}(B)$.
4. Si $\text{card}(A) < \text{card}(B)$, alors f ne peut pas être surjective.
5. Supposons $\text{card}(A) = \text{card}(B)$. Alors f est injective si, et seulement si, f est surjective.

Remarque : nous avons déjà vu une partie de ces affirmations (§1, Thm. 2 ; “principe des tiroirs”).

Chapitre 4 : Relations

Les *relations* généralisent les applications : toute application est une relation, mais la réciproque est fausse.

§1 Relations et graphes.

Définition 1. Si $f : M_1 \rightarrow M_2$ est une application, son graphe est la partie de $M_1 \times M_2$ définie par

$$R_f = \{(x, y) \in M_1 \times M_2 \mid y = f(x)\}.$$

Exemple 1. Si $M_1 = M_2 = \mathbb{R}$, le graphe R_f est une partie du plan $\mathbb{R}^2$. Par exemple, si $f(x) = x^2$, R_f est une *parabole* ; si $f(x) = ax + b$, R_f est une *droite*, de *coefficient directeur* a , et coupant l'axe des ordonnées au point $(0, b)$. On observera que *toute droite qui n'est pas verticale* est graphe d'une fonction ; mais les droites verticales ne peuvent pas être réalisées de cette façon ! Par contre, ce sont des *relations* :

Définition 2. Une relation entre deux ensembles M_1 et M_2 est une partie

$$R \subset (M_1 \times M_2).$$

Si $(x, y) \in R$, nous écrirons aussi xRy , et nous dirons que y est en relation R avec x .

Si $M_1 = M_2 = M$, on parle aussi d'une relation "sur M ".

Si R est de la forme $R = R_f$, où $f : M_1 \rightarrow M_2$ est une application, on dit que R est une relation fonctionnelle.

Exemple 2. Sur M , on a la relation "égalité" : xRy ssi $x = y$, donc $R = \Delta_M = \{(x, x) \mid x \in M\}$ est la *diagonale*. C'est une relation fonctionnelle : graphe de la fonction id_M .

Exemple 3. Si $R = \emptyset$, aucun élément de M_2 n'est en relation R avec un élément de M_1 . Si $R = M_1 \times M_2$, alors tout élément de M_2 est en relation R avec tout élément de M_1 (relation totale). Ces relations ne sont pas fonctionnelles.

Types de relations. Supposons $M_1 = M_2 = M$.

Définition 1. Une relation R sur un ensemble M est dite

- (1) symétrique si elle vérifie : pour tout $x, y \in M$, xRy si et seulement si yRx ;
- (2) antisymétrique si, pour tout $x, y \in M$: si $[xRy \text{ et } yRx]$, alors $x = y$;
- (3) réflexive si elle vérifie : pour tout $x \in M$, xRx ;
- (4) transitive si elle vérifie : pour tout $x, y, z \in M$, si xRy et yRz , alors xRz .

§2 Relations d'ordre.

Ce sont des relations qui formalisent l'idée d'une relation xRy signifiant que " x est inférieur ou égal à y " :

Définition 1. Une relation R sur M est dite relation d'ordre (partielle) si elle est

- antisymétrique,
- réflexive, et
- transitive.

Dans ce cas, on écrit souvent $x \leq y$ au lieu de xRy . On écrit aussi $a \geq b$ ssi $b \leq a$.

On parle d'une relation d'ordre total si, de plus :

- $\forall (x, y) \in M^2 : x \leq y \text{ ou } y \leq x$.

Exemple 1. Soit $M_1 = M_2 = \mathbb{N}$ et

$$R = \{(i, j) \in \mathbb{N}^2 \mid i \leq j\}$$

(inégalité au sens usuel). C'est une relation d'ordre totale. Idem pour $M = \mathbb{Z}$ ou $M = \mathbb{R}$.

Exemple 2. Soit A un ensemble et $M = \mathcal{P}(A)$ son ensemble de parties. Alors

$$bRc \quad \text{ssi} \quad b \subset c$$

définit une relation sur M (relation d'inclusion), qui n'est pas totale. Par exemple, si $A = \{1, 2, 3\}$, et $a = \{1, 2\}$ et $b = \{2, 3\}$, on n'a ni $a \leq b$ ni $b \leq a$.

Définition 2. Si $\leq$ est une relation d'ordre, on définit la relation $<$ par :
 $a < b$ ssi ($a \leq b$, et $a \neq b$).

Définition 3. Si $\leq$ est une relation d'ordre, et $(a, b) \in M^2$, on définit les intervalles

$$[a, b] := \{x \in M \mid a \leq x \text{ et } x \leq b\}, \quad]a, b[:= \{x \in M \mid a < x \text{ et } x < b\}.$$

Remarque. Cette définition semble évidente, mais elle donne lieu à des questions profondes, qui sont importantes en analyse : par exemple, a priori, l'ensemble $I = \{x \in \mathbb{R} \mid x^2 < 2\}$ n'est pas défini comme un intervalle. Et pourtant, on aurait envie de dire que c'en est un, à savoir $] -\sqrt{2}, \sqrt{2}[$

§3 Relations d'équivalence.

C'est un autre type de relation très important, xRy formalisant l'idée que “ x et y ont une certaine propriété en commun” :

Définition 1. Une relation R sur M est dite relation d'équivalence sur M si elle est

1. symétrique,
2. réflexive,
3. transitive.

Dans ce cas, on utilise souvent la notation $x \sim y$ au lieu de xRy .

Exemple 1. Soit M l'ensemble des étudiants de la promo L1. On définit une relation par : aRb si les étudiants a et b sont dans la même classe (i.e., appartiennent au même groupe de TD). Elle est clairement symétrique, transitive et réflexive.

Proposition 1. Soit $M = \sqcup_{i=1}^n A_i$ une partition de M . Alors la relation R définie par :

$$xRy \quad \text{ssi} \quad \exists i = 1, \dots, n : x \in A_i \text{ et } y \in A_i$$

est une relation d'équivalence sur M .

Preuve : arguments identitiques à ceux de l'exemple 1 :

1. si x est dans la même classe que y , alors y est dans la même classe que x ;
2. tout x est dans la même classe que x ;
3. si x est dans la même classe que y , et y dans la même que z , alors x est dans la même que z .

Théorème 1. *Toute relation d'équivalence sur M est donnée par la construction précédente : si R est une relation d'équivalence, il existe une partition de M donnant lieu à R comme dans la proposition 1. Ainsi “partitions et relations d'équivalence sur M sont la même chose”.*

Preuve (supposons que M soit un ensemble fini ; mais tout reste valable aussi si M est infini) :

Définition 2. *Si R est une relation d'équivalence sur M , on note, pour tout $x \in M$, par $[x] = \{y \in M \mid yRx\}$, la classe d'équivalence de x (= l'ensemble des éléments qui sont en relation avec x).*

La clé de la preuve est de montrer : Soit $x, z \in M$. Alors des deux cas l'un :

$$\boxed{\text{soit } [x] = [z], \text{ soit } [x] \cap [z] = \emptyset}.$$

En effet, si $y \in [x] \cap [z]$, alors montrons que $[x] \subset [z]$. Soit $u \in [x]$, alors : uRx, xRy, yRz , donc uRz , donc $u \in [z]$. Par symétrie, on aura de même $[z] \subset [x]$, donc $[x] = [z]$.

Ensuite, fixons dans chaque classe $[x]$ le choix d'un élément $i \in [x]$, qu'on appellera un représentant de la classe (de sorte que $[i] = [x]$; penser à i comme une sorte de “délégué de classe”), et soit I l'ensemble de tout ces représentants. Alors on a $M = \cup_{i \in I} A_i$ avec $A_i = [i]$ (réunion disjointe, la partition cherchée).

Exemple 2. La relation égalité est une relation d'équivalence. Ses classes d'équivalence sont les singletons (parties ayant un seul élément). C'est la partition “la plus fine possible”.

Exemple 3. La relation totale est une relation d'équivalence. Il n'y a qu'une seule classe : l'ensemble M lui-même. C'est la partition “la plus grossière possible”.

Exemple 4. On peut partitionner les entiers en deux parties : les entiers *paires* (les doubles, $2\mathbb{Z}$) et les entiers *impaires* ($2\mathbb{Z} + 1$).

Idem pour les multiples de 3 : on a une partition en trois parties, $A_0 = 3\mathbb{Z}$, $A_1 = 3\mathbb{Z} + 1$, $A_2 = 3\mathbb{Z} + 2$ (cf. TD 3).

Remarque. Ces résultats restent valables au cas des ensembles infinis. On utilise alors souvent les relations d'équivalence pour *définir de nouveaux ensembles*, qu'on appelle “ensemble quotient”. Nous en parlerons plus tard.

Seconde Partie du cours : Les ensembles infinis

La définition, due à Cantor, d'un ensemble quelconque (chapitre 1) ne suppose en aucune manière que l'ensemble A soit fini. Or, si on ne peut pas écrire A sous forme d'une "liste" ou d'un "fichier fini", la question se pose : *comment "définir" ou "décrire" un ensemble infini ?* Dans le cours, nous allons étudier ce problème pour les ensembles infinis "les plus élémentaires" : les ensembles $\mathbb{N}$ et $\mathbb{Z}$ des nombres naturels, respectivement entiers. Une fois ce problème résolu, dans le cours d'analyse on enchaînera pour définir les nombres réels, $\mathbb{R}$, et complexes, $\mathbb{C}$.

L'ensemble $\mathbb{N}$ est absolument fondamental pour toutes les mathématiques. Nous allons *poser comme axiome qu'un ensemble $\mathbb{N}$ avec certaines propriétés existe*. Ces propriétés peuvent être formulées de diverses manières – la plus courante est les *axiomes de Peano*. Cependant, Cantor, et ses successeurs (Zermelo, Frankel), ont cherché à jeter les bases à un niveau plus profond encore : il faudra non seulement énoncer les "axiomes de $\mathbb{N}$ " de façon très précise, mais aussi *toute la théorie des ensembles*. Entrer dans les détails dépasserait le cadre de ce cours ; mais nous expliquerons les grandes lignes.

Les ensembles infinis ont des propriétés surprenantes. Par exemple, il existe "autant de nombres naturels que de nombres naturels pairs" (à n on fait correspondre $2n$), et pourtant, il devait y en avoir "moins". (Pour la petite histoire : l'Hôtel de Hilbert.) D'autres propriétés sont franchement paradoxales : la plus célèbre est le *paradoxe de Russell* : soit M l'ensemble dont les éléments sont tous les ensembles n'appartenant pas à eux-mêmes :

$$M = \{x \mid x \notin x\}.$$

Alors, est-ce que $M \in M$? si oui, par définition M n'est pas dans M ; donc c'est impossible. Mais alors si non, par définition on a $M \in M$; donc c'est impossible aussi. En fait, dans le cadre de la logique ce paradoxe est connu depuis l'antiquité : le *paradoxe du menteur* (Un homme dit : "Je mens maintenant." Vrai ou faux ? Si vrai, c'est qu'il ne ment pas, donc l'affirmation est fausse. Donc c'est impossible. Alors l'affirmation est fausse : cela veut dire qu'il dit la vérité, i.e., il ment - impossible également.) Il ne faut pas voir en ces paradoxes un simple tour de passe-passe : ce sont de vrais problèmes qui apparaissent quand on manipule de façon trop imprudente des univers de discours *infinis*, qui offrent la possibilité d'*autoréférence* (une phrase fait référence à elle-même). En effet, on constate que les problèmes de la théorie des ensembles infinis et de la logique sont étroitement liés entre eux. Pour cette raison, nous allons évoquer dans cette partie du cours aussi quelques notions fondamentales de la *logique mathématique*.

Chapitre 5 : Les nombres naturels $\mathbb{N}$

Le mathématicien Leopold Kronecker a écrit, en 1891, “*les nombres naturels sont la création de Dieu, tout le reste est l’œuvre de l’homme*”. Qu’on soit croyant ou non, il faut fixer, quelque part, un point de départ, considéré comme accepté par tous, pour ce qui concerne les théories mathématiques. Ces fondements s’appellent, en mathématiques, les axiomes. Voici les axiomes dits *axiomes de Peano* qui caractérisent les nombres naturels

§1 Les axiomes de Peano, et le principe de récurrence.

Définition 1. Les entiers naturels sont la donnée d’un ensemble $\mathbb{N}$, d’un élément particulier noté $0 \in \mathbb{N}$, et d’une application $s : \mathbb{N} \rightarrow \mathbb{N}$, $n \mapsto s(n)$; l’élément $s(n)$ s’appelle le successeur de n ; et ces données vérifient les propriétés suivantes :

- (P1) 0 n’est pas un successeur : $s(n)$ est différent de 0 pour tout $n \in \mathbb{N}$;
- (P2) l’application s est injective : si $s(n) = s(m)$ pour $n, m \in \mathbb{N}$, alors $n = m$;
- (P3) soit $M \subset \mathbb{N}$ une partie qui contient 0 et telle que, si $m \in M$, alors $s(m) \in M$; alors on a $M = \mathbb{N}$.

Remarque 1. L’application $s : \mathbb{N} \rightarrow \mathbb{N}$ est donc *injective* (d’après (P2)) mais pas *surjective* (d’après (P1)). Cela confirme que $\mathbb{N}$ ne peut pas être un ensemble fini car nous avons vu (“principe des tiroirs”) qu’une application injective $f : M \rightarrow M$ d’un ensemble fini M dans lui-même est toujours surjective.

Notation. Nous utilisons les symboles usuels : $1 := s(0)$, $2 := s(1)$, $3 := s(2)$, $4 := s(3)$, et nous utilisons aussi la notation : $n + 1 := s(n)$.

L’axiome (P3) permet de justifier le *principe de preuve par récurrence* : supposons que nous conjecturons qu’une certaine proposition (énoncé mathématique) dépendant de n soit vraie, quel que soit $n \in \mathbb{N}$. Pour la prouver, il suffit

- de la démontrer “à la main” pour $n = 0$;
- de vérifier que, si elle vraie pour un $n \in \mathbb{N}$, alors elle l’est aussi pour $s(n)$:

Théorème 1 (Principe de récurrence). Supposons donné, pour tout $n \in \mathbb{N}$, une proposition $P(n)$, telle que

1. (*initialisation*) : la proposition $P(0)$ est vraie ;
2. (*hérédité*) : pour tout $n \in \mathbb{N}$: [si $P(n)$ est vraie, alors $P(s(n))$ est vraie aussi].

Alors $P(n)$ est vraie pour tout $n \in \mathbb{N}$.

Preuve. On considère l’ensemble $M := \{n \in \mathbb{N} \mid P(n) \text{ est vraie}\}$. D’après (1) et (2), M vérifie les hypothèses de l’axiome (P3) ; donc, d’après cet axiome, on a $M = \mathbb{N}$. Donc $P(n)$ est vraie pour tout $n \in \mathbb{N}$.

Exemple 1. Soit $P(n)$ la proposition : pour tout nombre réel ou complexe a tel que $a \neq 1$,

$$\sum_{k=0}^n a^k = \frac{1 - a^{s(n)}}{1 - a} .$$

Initialisation : $P(0)$ revient à $a^0 = \frac{1-a^1}{1-a}$; c’est évidemment vrai (car $a^0 = 1$).

Hérédité : supposons que $P(n)$ soit vraie pour un $n \in \mathbb{N}$ (hypothèse de récurrence).

Montrons que $P(s(n))$ est vraie : en effet

$$\begin{aligned}
 \sum_{k=0}^{s(n)} a^k &= \sum_{k=0}^n a^k + a^{s(n)} \\
 &= \frac{1 - a^{s(n)}}{1 - a} + a^{s(n)} \quad (\text{par hypothèse de récurrence}) \\
 &= \frac{1 - a^{s(n)} + (1 - a)a^{s(n)}}{1 - a} \\
 &= \frac{1 - a^{s(n)}a}{1 - a} = \frac{1 - a^{s(n)+1}}{1 - a} = \frac{1 - a^{s(s(n))}}{1 - a}
 \end{aligned}$$

donc $P(s(n))$ est vraie. Par le principe de récurrence, $P(n)$ est donc vraie pour tout $n \in \mathbb{N}$.

Sous-entendue dans ce raisonnement est la *définition des puissances* : on définit $a^0 := 1$, puis $a^{s(n)} := a^n \cdot a$; ainsi la puissance a^n est définie pour tout $n \in \mathbb{N}$ (par récurrence).

On écrira désormais plutôt $n + 1$ au lieu de $s(n)$; la formule précédente, si $a \neq 1$,

$$\boxed{\sum_{k=0}^n a^k = \frac{1 - a^{n+1}}{1 - a}}$$

s'appelle la *formule de la somme géométrique*. Voir TD pour d'autres exemples !

Exemple 2. Tout entier naturel non-nul x admet un prédécesseur, i.e., $\exists y \in \mathbb{N} : x = s(y)$.

En effet, soit $S = \{x \in \mathbb{N} \mid x = 0 \text{ ou } \exists y \in \mathbb{N} : x = s(y)\}$. Clairement, $0 \in S$; et si $n \in S$, alors $x = s(n) \in S$ (en prenant $y = n$), donc d'après (P3), on a $S = \mathbb{N}$.

Remarque 2. Si l'initialisation se fait au rang 1 (ou 2 ou...), et l'hérédité est vérifiée pour tout n sauf 0 (sauf 0 et 1...), alors $P(n)$ est vrai pour tout n à partir du rang 1, etc.

Remarque 3. Comme tous les mathématiciens, nous supposons qu'un système $(\mathbb{N}, 0, s)$ vérifiant les axiomes de Peano *existe* – nous ne cherchons pas à *démontrer* ce fait. En revanche, on peut prouver que, si un tel ensemble existe, alors il est *unique*, dans un sens à préciser : il ne peut pas y être “deux espèces différentes” de nombres naturels.

§2 Addition et multiplication dans $\mathbb{N}$.

Nous avons défini $n + 1 := s(n)$ et $2 = s(1) = 1 + 1$. On définit maintenant

$$n + 2 := s(s(n)) = (n + 1) + 1.$$

Ainsi $n + (1 + 1) = (n + 1) + 1$ pour tout $n \in \mathbb{N}$. Par récurrence, nous définissons : soit $k \in \mathbb{N}$; on pose $k + 0 := k$, et si $k + n$ pour un $n \in \mathbb{N}$ est déjà défini, on pose

$$k + s(n) := s(k + n).$$

Autrement dit, $k + (n + 1) := (k + n) + 1$. Par (P3), on a ainsi défini la *somme* $k + n$ pour tout $n \in \mathbb{N}$.

Théorème 1 (associativité). Pour tous $k, m, n \in \mathbb{N}$,

$$\boxed{k + (m + n) = (k + m) + n}.$$

Preuve par récurrence : soit $P(n)$ la proposition “pour tout k et m dans $\mathbb{N}$, l'égalité $k + (m + n) = (k + m) + n$ est vraie”.

$P(0)$ est vrai car $k + (m + 0) = k + m = (k + m) + 0$, pour tout k et m . Par ailleurs, $P(1)$ est vraie comme vue ci-dessus.

Supposons que $P(n)$ est vraie. Alors, $\forall k, m \in \mathbb{N}$:

$$\begin{aligned} k + (m + (n + 1)) &= k + ((m + n) + 1) && \text{(par le cas } n = 1 \text{ déjà prouvé)} \\ &= (k + (m + n)) + 1 && \text{(par le cas } n = 1 \text{ déjà prouvé)} \\ &= ((k + m) + n) + 1 && \text{(par hypothèse de récurrence)} \\ &= (k + m) + (n + 1) && \text{(par le cas } n = 1 \text{ déjà prouvé).} \end{aligned}$$

Donc $P(n + 1)$ est vraie aussi.

Par le principe de récurrence, $P(n)$ est donc vraie pour tout $n \in \mathbb{N}$.

Théorème 2 (commutativité). *Pour tous $m, n \in \mathbb{N}$,*

$$\boxed{n + m = m + n}.$$

Par exemple, $2 + 1 = (1 + 1) + 1 = 1 + (1 + 1) = 1 + 2$. Montrer d'abord par récurrence que $n + 1 = 1 + n$ pour tout $n \in \mathbb{N}$, puis par une autre récurrence que $k + n = n + k$ (cf. TD). – Ensuite, nous définissons par récurrence le *produit*

$$0 \cdot m = 0, \quad 1 \cdot m = m, \quad 2 \cdot m := m + m, \quad 3 \cdot m := 2 \cdot m + m, \quad 4 \cdot m := 3 \cdot m + m,$$

et si $n \cdot m$ est déjà défini, on pose

$$(n + 1) \cdot m := n \cdot m + m.$$

Ainsi $n \cdot m$ est défini pour tout $(n, m) \in \mathbb{N}^2$. On écrit souvent mn au lieu de $n \cdot m$.

Théorème 3 (règles de calcul). *Pour tous $k, m, n, p, q \in \mathbb{N}$,*

(D) *distributivité* : $k(m + n) = km + kn$,

(A) *associativité* : $k(mn) = (km)n$

(C) *commutativité* : $mn = nm$,

(S) *règle de simplification additive* : $p + m = q + m \Rightarrow p = q$,

(S') *règle de simplification multiplicative* : si $m \neq 0$, alors : $pm = qm \Rightarrow p = q$.

Preuve cf. TD. Utilisant l'associativité, on peut omettre des parenthèses dans les sommes

$$\sum_{i=1}^n a_i = a_1 + a_2 + \dots + a_n$$

avec $a_i \in \mathbb{N}$, et dans les produits

$$\prod_{i=1}^n a_i = a_1 \cdot a_2 \cdots a_n.$$

Plus généralement, si I est un ensemble *fini* d'indices (quelconque : pas forcément d'entiers), et $a_i \in \mathbb{N}$ est donné pour tout $i \in I$, on peut définir les sommes, resp. produits, de

tout ces éléments, car le résultat ne dépend ni de l'ordre ni de la façon de regrouper les termes ; on note ces sommes, resp. produits, par

$$\sum_{i \in I} a_i, \quad \prod_{i \in I} a_i.$$

Proposition 1. *La seule solution de l'équation $p + q = 0$ dans $\mathbb{N}$ est $p = q = 0$.*

Preuve. Il s'agit d'une *preuve par l'absurde* : supposons que $p + q = 0$ et $q \neq 0$. D'après §1, exemple 2, q admet alors un prédécesseur : $\exists m \in \mathbb{N}, q = s(m) = m + 1$. Donc $0 = p + q = (p + m) + 1$ admet un prédécesseur – mais ceci est en contradiction avec l'axiome (P1). Ainsi l'hypothèse que $p + q = 0$ avec $q \neq 0$ mène à une contradiction logique ; il faut donc qu'elle est fausse.

Remarque : cf. TD pour une autre preuve célèbre qui procède par l'absurde – celle d'Euclide, prouvant qu'il existe une infinité de nombres premiers.

§3 La relation d'ordre sur $\mathbb{N}$.

Définition 1. *Soit $(m, n) \in \mathbb{N}^2$. On écrira $m \leq n$ s'il existe $p \in \mathbb{N}$ tel que $n = m + p$.*

Lemme (règles de calcul). *La relation $\leq$ sur $\mathbb{N}$ vérifie :*

- (1) $\forall n \in \mathbb{N} : 0 \leq n$,
- (2) *si $p \leq q$ et $n \in \mathbb{N}$, alors $p + n \leq q + n$ et $np \leq nq$.*

La preuve est facile à partir des règles de calcul du théorème 3 ci-dessus.

Théorème 1. *La relation $\leq$ est une relation d'ordre total sur $\mathbb{N}$.*

Preuve. Il faut vérifier les 4 propriétés qui définissent une relation d'ordre total :

- Réflexivité : $n \leq n$ clair (prendre $p = 0$)
- Transitivité : si $k \leq m$ et $m \leq n$, alors $m = k + p$, $n = m + q$, donc $n = (k + p) + q = k + (p + q)$, donc $k \leq n$.
- Antisymétrie : soit $m \leq n$ et $n \leq m$, donc $n = m + p$ et $m = n + q$, donc $n = (n + q) + p = n + (p + q)$. D'après la règle de simplification, il s'ensuit que $p + q = 0$. D'après la proposition 1 ci-dessus, il s'ensuit que $p = q = 0$, donc $m = n$.
- Totalité : pour $n \in \mathbb{N}$, posons

$$S_n := \{m \in \mathbb{N} \mid m < n\} \cup \{n\} \cup \{m \in \mathbb{N} \mid n < m\}$$

et montrons par récurrence que $S_n = \mathbb{N}$. L'initialisation vient du fait que $0 \leq m$ pour tout $m \in \mathbb{N}$, et l'hérédité se démontre à l'aide de la règle de calcul $p < q \Rightarrow p + 1 < q + 1$.

Une fois ces règles de calcul (“bien connues”) posées, on peut procéder à des choses plus intéressantes : définir les nombres premiers, prouver que tout entier naturel se factorise de façon unique en nombres premiers, et étudier ces propriétés, i.e., faire de l'*arithmétique*. Ce sera l'objet d'une option de maths en S2. Cf. aussi TD.

Chapitre 6 : Nombres

Suivant la phrase de Kronecker, une fois posé les axiomes de $\mathbb{N}$, c'est "l'œuvre de l'homme" de construire les "ensembles des nombres usuels", de plus en plus grands :

1. les entiers relatifs $\mathbb{Z} = \{\dots, -2, -1, 0, 1, 2, \dots\}$,
2. les nombres rationnels $\mathbb{Q}$ (fractions $r = \frac{m}{n}$ avec m, n des entiers relatifs et $n \neq 0$),
3. les nombres réels $\mathbb{R}$ ("tous les nombres de la droite réelle"),
4. les nombres complexes $\mathbb{C}$ ("tous les nombres du plan d'Argand")

En général, on entend par "nombres" en mathématiques un ensemble $\mathbb{K}$, dont les éléments sont appelés des "nombres", muni d'opérations, "addition" $+$ et "multiplication" $\times$, ou $\cdot$, vérifiant certaines propriétés. Pour prouver que ces nombres "existent",

- soit qu'on pose ces propriétés comme "axiomes" (approche axiomatique),
- soit qu'on base la théorie sur celle de $\mathbb{N}$ (approche par construction) :

en supposant que l'existence de $\mathbb{N}$ est acquise, on *construit les nombres* $\mathbb{K}$ à partir de $\mathbb{N}$, via certaines constructions venant de la théorie des ensembles (produit cartésien ; relations d'équivalence,...). Pour $\mathbb{Z}, \mathbb{Q}, \mathbb{R}, \mathbb{C}$, c'est tout à fait possible – l'étape la plus difficile est de *construire* $\mathbb{R}$ à partir de $\mathbb{Q}$ (la *construction des nombres réels* est un sujet intéressant, qui n'est malheureusement pas au programme des études universitaires, car jugé trop long et technique – pour abrégé, on introduit, au cours d'analyse 1 en S2, les *nombres réels* de façon axiomatique). En revanche, il n'est pas trop difficile de construire $\mathbb{Z}$ à partir de $\mathbb{N}$, puis $\mathbb{Q}$ à partir de $\mathbb{Z}$. Vous allez voir ces constructions aux cours d'algèbre de Licence ; elles ne font pas objet de ce cours, mais il est utile d'en avoir un aperçu dès maintenant.

§1 Construction de $\mathbb{Z}$: idée. Nous avons vu que l'équation $a + x = b$ n'admet pas toujours de solution dans $\mathbb{N}$: elle en admet une si $a \leq b$, mais non si $b < a$. Or, on aimerait bien pouvoir dire qu'il existe *toujours* une solution – pour cela, il faut agrandir l'ensemble $\mathbb{N}$ en l'ensemble $\mathbb{Z}$ des *entiers relatifs*. On notera alors cette solution (qui sera unique), $x = b - a$. Ainsi, si cet ensemble de nombres $\mathbb{Z}$ existe, on aura une application

$$f : \mathbb{N} \times \mathbb{N} \rightarrow \mathbb{Z}, \quad (a, b) \mapsto b - a.$$

Cette application sera surjective (car on aura bien $m = m - 0 = f((0, m))$), mais non injective : pour $m \in \mathbb{Z}$, la *fibre* $f^{-1}(m)$ est l'ensemble $[m] := \{(a, b) \in \mathbb{N}^2 \mid b - a = m\}$. On pourra dire que le couple $(a, b) \in \mathbb{N}^2$ "représente" l'entier m ; et deux couples $(a, b), (a', b')$ représentent le même entier si $a - b = a' - b'$, autrement dit, si $a + b' = a' + b$.

Dit encore autrement, la relation R sur $M = \mathbb{N}^2$ définie par :

$$(a, b)R(a', b') \quad \text{ssi} \quad a + b' = a' + b$$

devait être une *relation d'équivalence*, et $\mathbb{Z}$ devait s'identifier à l'ensemble des classes d'équivalence de cette relation. Et en effet, ce programme permet de *construire* $\mathbb{Z}$ via une *relation d'équivalence* R sur $\mathbb{N}^2$: voir TD, et Annexe 1 du cours pour quelques détails de plus. Par cette construction, on obtient un ensemble ayant les propriétés suivantes :

Théorème 1. *Les entiers relatifs $\mathbb{Z}$ forment un ensemble, muni d'un élément particulier noté 0, et de deux opérations $+$ et $\cdot$ associant à deux nombres un troisième, vérifient les propriétés d'un anneau commutatif, i.e.,*

1. *addition et multiplication sont associatives et commutatives,*
2. *ces opérations sont distributives entre elles ;*
3. *0 est neutre pour l'addition : $\forall m \in \mathbb{Z}, 0 + m = m$,*
4. *pour tout $(m, n) \in \mathbb{Z}^2$, l'équation $m + x = n$ admet une unique solution (qu'on note $x = n - m \in \mathbb{Z}$) ;*

de plus cet anneau est ordonné : il existe une relation d'ordre total $\leq$ sur $\mathbb{Z}$, telle que

1. $\forall a, b, c \in \mathbb{Z} : a \leq b \Rightarrow a + c \leq b + c$,
2. $\forall a, b, c \in \mathbb{Z} : (a \leq b \text{ et } 0 \leq c) \Rightarrow ac \leq bc$.

Finalement, on a une bijection "canonique" entre $\mathbb{Z}^+ = \{m \in \mathbb{Z} \mid 0 \leq m\}$ et $\mathbb{N}$.

Toutes ces propriétés ensemble caractérisent $\mathbb{Z}$ de façon unique, et peuvent donc être utilisées pour définir $\mathbb{Z}$ "axiomatiquement".

§2 Construction de $\mathbb{Q}$: idée. Dans $\mathbb{Z}$, on ne peut pas toujours résoudre des équations du type $ax = b$: une solution $x = \frac{b}{a}$ serait une *fraction d'entiers* ; si $a = 1$ ou $a = -1$, c'est un entier, mais pas dans les autres cas (et si $a = 0$ elle n'est pas définie, car $0x = 0$ pour tout x). On veut donc agrandir le domaine des nombres encore une fois, pour qu'on puisse diviser par tout nombre non-nul. Ce genre de "domaine de nombres" joue un rôle très important en maths :

Définition 1. *Un corps est un ensemble $\mathbb{K}$, tel que pour tout $(a, b) \in \mathbb{K}^2$ soient définis*

- *la somme $a + b \in \mathbb{K}$,*
- *le produit $ab \in \mathbb{K}$,*
- *la différence $a - b \in \mathbb{K}$,*
- *si $b \neq 0$, le quotient $\frac{a}{b} \in \mathbb{K}$,*

et tel que toutes les lois usuelles (associativité, commutativité, distributivité,...) du calcul littéral des 4 opérations soient vérifiées. Il existe un élément neutre 0 pour l'addition, et un élément neutre 1 pour la multiplication.

Les corps les plus importants sont $\mathbb{Q}, \mathbb{R}, \mathbb{C}$; mais il y en a beaucoup d'autres (cours d'algèbre de L2-L3....). On peut prouver que *si $\mathbb{Z}$ existe, alors $\mathbb{Q}$ existe aussi*, en construisant le corps $\mathbb{Q}$ selon la même méthode qu'au paragraphe précédent : si $\mathbb{Q}$ existe, alors on a une application (surjective)

$$f : \mathbb{Z} \times (\mathbb{Z} \setminus \{0\}) \rightarrow \mathbb{Q}, \quad (a, b) \mapsto \frac{a}{b}.$$

On a $f((a, b)) = f((a', b'))$ ssi $\frac{a}{b} = \frac{a'}{b'}$, ssi $ab' = a'b$. Ceci définit une relation d'équivalence, et on pourra définir $\mathbb{Q}$ comme l'ensemble de classes d'équivalences. De plus, par les règles du calcul littéral

$$\frac{a}{b} + \frac{c}{d} = \frac{ad + bc}{bd}, \quad \frac{a}{b} - \frac{c}{d} = \frac{ad - bc}{bd}, \quad \frac{a}{b} \cdot \frac{c}{d} = \frac{ac}{bd}, \quad \frac{a}{b} : \frac{c}{d} = \frac{ad}{bc},$$

on saura comment il faut définir les 4 opérations sur de telles classes. Voir Annexe 2 pour plus de détails. Voici le résultat :

Théorème 1. *Les nombres rationnelles $\mathbb{Q}$ forment un corps qui est ordonnée, i.e., muni d'un ordre total $\leq$ tel que :*

1. *si $a \leq b$, alors pour tout $c \in \mathbb{Q}$, on a $a + c \leq b + c$,*
2. *si $a \leq b$ et $c \geq 0$ dans $\mathbb{Q}$, alors $ac \leq bc$.*

De plus, $\mathbb{Q}$ contient $\mathbb{Z}$ comme partie ; et tout élément de $\mathbb{Q}$ est de la forme $q = \frac{m}{n}$ avec $m \in \mathbb{Z}$ et $n \in \mathbb{N}^*$.

Là aussi, ces propriétés caractérisent $\mathbb{Q}$ de façon unique, et pourrait servir à le définir axiomatiquement.

§3 Calcul dans $\mathbb{Z}$ et $\mathbb{Q}$. Tout ce qui relève des “règles du calcul littéral des 4 opérations” est valable dans tout corps (dont $\mathbb{Q}, \mathbb{R}$ et $\mathbb{C}$). En particulier :

Définition 1. Dans tout corps ou anneau commutatif $\mathbb{K}$, sommes et produits de plusieurs éléments sont définis comme nous l’avons fait dans $\mathbb{N}$; et de même, les signes somme $\sum_{i \in I}$ et signe produit $\prod_{i \in I}$, pour tout ensemble fini d’indices, sont définies comme nous l’avons fait dans $\mathbb{N}$.

En revanche, il existe des propriétés qui sont spécifiques à $\mathbb{Q}$: on a déjà vu une propriété, à savoir que $\mathbb{Q}$ est un corps ordonné.

Théorème 1. Dans un corps ordonné, les carrés sont positifs : $\forall a \in \mathbb{Q}, a^2 \geq 0$.

Preuve par disjonction de cas : soit, $a \geq 0$, alors $a^2 \geq 0$ par l’item 2 du théorème 1 précédent. Soit $a < 0$. Alors $0 < -a$ par l’item 1, et alors $0 < (-a) \cdot (-a)$ par l’item 2. Or $(-a)(-a) = a^2$ (calcul littéral), donc encore $a^2 \geq 0$.

Corollaire. Le corps $\mathbb{C}$ n’est pas un corps ordonné.

Preuve. Si $\mathbb{C}$ était ordonné, on aurait, d’une part, $-1 = i^2 > 0$; d’autre part, $-1 < 0$: contradiction !

Bien sûr, $\mathbb{R}$ sera un corps ordonné, comme $\mathbb{Q}$. Alors, que distingue $\mathbb{R}$ de $\mathbb{Q}$?

Théorème 2. Il n’existe aucun nombre $c \in \mathbb{Q}$ tel que $c^2 = 2$.

Preuve. cf. TD !

On “sait” qu’il existe $c = \sqrt{2}$ dans $\mathbb{R}$; ainsi on a trouvé une propriété qui distingue $\mathbb{Q}$ de $\mathbb{R}$. Cependant, $\mathbb{R}$ n’est pas le seul corps qui a cette propriété (cf. TD). On peut aller plus loin et dire que $\mathbb{R}$ contient toutes les racines de tous les entiers positifs. Mais là encore, il existe d’autres corps qui ont cette propriété (cours d’algèbre de L2-L3). Alors, quelle est “la” propriété qui distingue $\mathbb{R}$ de $\mathbb{Q}$? En fait, il s’agit d’une propriété d’analyse, et non d’algèbre, et elle sera précisée en cours d’analyse de S2 : le corps $\mathbb{R}$ est un corps ordonné complet ; le mot “complet” signifie que $\mathbb{R}$ réalise la droite numérique “continue”, “sans trou” (tandis que $\mathbb{Q}$ n’est pas complet : au lieu où devait se trouver le nombre $\sqrt{2}$, il n’y a rien, il y a un “trou”). Pour bien comprendre cette propriété, il faudra mobiliser à peu près tous les fondements qu’on vient de poser dans ce cours... et alors cela permettra de définir de “nouveaux nombres”, comme e ou π , qui n’existent pas dans $\mathbb{Q}$, et qui sont définis par des propriétés d’analyse.

Chapitre 7 : Théorie des ensembles

Avec la définition de $\mathbb{N}$, $\mathbb{Z}$ et $\mathbb{Q}$ nous sommes entrés dans le “royaume des ensembles infinis”. C’est le moment d’évoquer quelques questions de fondements des maths : la *théorie des ensembles* est précisément la “théorie mathématique de l’infini”. La découverte révolutionnaire de Cantor était qu’il existent plusieurs “infinis”, et qu’il est possible d’en dire des choses précises en langage mathématique.

§1 Équipotence ; ensembles dénombrables.

Pour comparer la “taille” de deux ensembles, on définit :

Définition 1. Deux ensembles A et B sont dits équipotents, ou : de même cardinalité, s’il existe une application bijective $f : A \rightarrow B$. On écrit alors $\text{card}(A) = \text{card}(B)$.

Remarque. Nous avons les propriétés d’une relation d’équivalence :

$\text{card}(A) = \text{card}(A)$ (prendre $f = \text{id}_A$),

$\text{card}(A) = \text{card}(B)$ ssi $\text{card}(B) = \text{card}(A)$ (prendre $f^{-1} = B \rightarrow A$),

si $\text{card}(A) = \text{card}(B)$ et $\text{card}(B) = \text{card}(C)$, alors $\text{card}(A) = \text{card}(C)$ (prendre $h = g \circ f$).

Définition 2. Un ensemble A est fini, de cardinal $n \in \mathbb{N}$ si A est équipotent à $[1, \dots, n]$. Un ensemble A est dit dénombrable s’il est équipotent à $\mathbb{N}$.

Exemples (cf TD 4) :

l’ensemble $2\mathbb{N}$ des nombres naturels *pairs* est dénombrable,

L’ensemble $A = \{n^2 \mid n \in \mathbb{N}\}$ est dénombrable.

$\mathbb{Z}$ est dénombrable.

Exemple : $\{1, 2\}$ et $\{1, 2, 3\}$ ne sont pas équipotents. En effet, d’après le chapitre 3 (§1, thm. 2) il ne peut pas exister de bijection entre ces ensembles !

Conclusion : un ensemble *fini* M n’est jamais équipotent à une partie $A \subset M$, $A \neq M$. Pour un ensemble *infini* M , c’est tout à fait possible.

Définition 3. On écrit $\text{card}(A) \leq \text{card}(B)$ s’il existe une application injective $f : A \rightarrow B$.

La preuve du théorème suivant est difficile dans le cas de cardinalité infinie ; dans le cas de cardinalité finie c’est un exercice (cf. TD) :

Théorème 1 (de Cantor-Bernstein). Si $\text{card}(A) \leq \text{card}(B)$ et $\text{card}(B) \leq \text{card}(A)$, alors $\text{card}(A) = \text{card}(B)$. Autrement dit, s’il existe $f : A \rightarrow B$, et $g : B \rightarrow A$, les deux injectives, alors il existe aussi $h : A \rightarrow B$, bijective.

§2 L’argument diagonal de Cantor. Existe-t-il des ensembles A “plus grands que $\mathbb{N}$ ”, dans le sens que A n’est pas fini et $\text{card}(A) \neq \text{card}(\mathbb{N})$? La réponse, due à Cantor, est “oui” : par exemple, $A = \mathbb{R}$ est d’un infini plus grand que celui de $\mathbb{N}$, voir ci-dessous. Dans un premier temps, on pourrait penser que l’ensemble $M := \mathbb{N} \times \mathbb{N}$ est déjà “bien plus grand que $\mathbb{N}$ ”. En fait, cette impression trompe :

Théorème 1. L’ensemble $\mathbb{N} \times \mathbb{N}$ est dénombrable.

Preuve. On peut “numéroter” les couples d’entiers : il existe une fonction $f : \mathbb{N} \times \mathbb{N} \rightarrow \mathbb{N}$ telle que chaque couple (i, j) correspond à un numéro unique $f((i, j))$. Voici le schéma :

$$f((0, 0)) = 0, \quad f((1, 0)) = 1, \quad f((0, 1)) = 2, \quad f((2, 0)) = 3, \quad f((1, 1)) = 4, \quad f((0, 2)) = 5$$

Faire un graphique pour comprendre le principe : on “numérote les couples en procédant selon les diagonales” ; en TD on donnera une “formule explicite” pour calculer $f((x, y))$.

Corollaire. *Les nombres rationnels $\mathbb{Q}$ sont dénombrables.* (Si on peut numérotter les couples (a, b) , on peut aussi numérotter les fractions $\frac{a}{b}$ avec $a \in \mathbb{Z}, b \in \mathbb{Z}^*$.)

Alors, jusque-là nous n'avons pas trouvé d'ensemble qui soit "plus grand que $\mathbb{N}$ ". Y en-a-t-il? La réponse est "oui", et la méthode pour le démontrer est le célèbre *argument de la diagonale de Cantor* :

Théorème 2. *L'ensemble $\mathcal{P}(\mathbb{N})$ de toutes les parties de $\mathbb{N}$ n'est pas dénombrable.*

Preuve. Il s'agit d'une *preuve par l'absurde* : supposons que $\mathcal{P}(\mathbb{N})$ soit dénombrable – i.e., qu'il existe une bijection $f : \mathbb{N} \rightarrow \mathcal{P}(\mathbb{N})$. Autrement dit, on peut faire une liste $A_1 = f(1)$, $A_2 = f(2), \dots$ de toutes les parties de $\mathbb{N}$. Alors Cantor considère la partie

$$M := \{m \in \mathbb{N} \mid m \notin A_m\}$$

et il montre que M est différent de tous les A_n . Pour le montrer, il suppose qu'il existe $n \in \mathbb{N}$ tel que $M = A_n$. Alors, soit $n \in A_n$, soit $n \notin A_n$. Or, dans le premier cas, il s'ensuit que $n \notin A_n$, donc c'est impossible. Mais dans le deuxième cas, il s'ensuit que $n \in M$, donc $n \in A_n$, donc c'est également impossible. C'est contradictoire, et ainsi l'hypothèse qu'il existe n tel que $M = A_n$ ne peut pas être vraie : ainsi f ne peut pas être surjective. Pour que la théorie soit non-contradictoire, il faut donc que $\mathcal{P}(\mathbb{N})$ soit non-dénombrable.

Dans la suite, Cantor appliqua le même genre de raisonnement pour montrer :

Théorème 3. *l'ensemble $\mathbb{R}$ des nombres réels n'est pas dénombrable.* (Cf. TD.)

Ainsi, il existe bien des ensembles M tels que $\text{card}(M) > \text{card}(\mathbb{N})$, et ainsi de suite : $\text{card}(\mathbb{N}) < \text{card}(\mathcal{P}(\mathbb{N})) < \text{card}(\mathcal{P}(\mathcal{P}(\mathbb{N}))) < \dots$: une chaîne infinie de cardinaux infinis toujours strictement plus grands! Cette découverte marqua le véritable début de la théorie des ensembles.

§3 Axiomes de la théorie des ensembles. L'argument de la diagonale de Cantor ressemble dans une certaine mesure à celui du "paradoxe de Russell" connu en logique et dont on parlera au chapitre suivant. Et, en effet, il faut être très prudent en utilisant ce genre d'arguments. Pour pouvoir les utiliser correctement, les mathématiciens (dont Zermelo et Frankel dans la suite de Cantor) ont posé une *base axiomatique de la théorie des ensembles* (axiomes *ZF*, ou *ZFC*). Il n'est pas question dans ce cours d'expliquer ces axiomes en détail. Pour plus d'information, cf. par exemple la [page wikipedia](#).

Chapitre 8 : Logique et maths

La logique est une branche de la philosophie qui traite de la cohérence de notre discours sur le monde en général, et la *logique mathématique* traite de notre discours concernant les maths en particulier. Ainsi la logique ne fait pas partie des maths, mais elle pose le cadre nécessaire pour commencer à les développer.

§1 Logique mathématique : propositions. Les mathématiques ont pour objectif d'étudier des *objets mathématiques*, et de *prouver des théorèmes* portant sur ces objets, c'est-à-dire, d'établir la *vérité de certaines propositions* concernant ces objets. En mathématiques, nous appelons proposition un énoncé portant sur des objets mathématiques, et qui a un sens. Cela veut dire que, au moins théoriquement, il devra être possible d'attribuer exactement l'une des deux *valeurs de vérité*, vraie ou fausse, à cet énoncé. Par exemple, "3 est plus grand" n'est pas une proposition (la phrase est syntaxiquement incomplète), et "il pleut" non plus (même si la phrase est syntaxiquement correcte, elle peut être ni vraie ni fausse : une goutte par $(km)^2$, c'est de la pluie ou pas ? On n'accepte pas ce genre de situation en maths...) ; mais " $3^2 + 4^2 = 5^2$ " et " $4^2 + 5^2 = 6^2$ " le sont.

§2 Calcul des propositions. Le calcul des propositions porte sur les *connecteurs logiques* : pour une proposition α , on notera la valeur de vérité $v(\alpha) = 0$ si α est fausse et $v(\alpha) = 1$ si α est vraie.¹ À partir de deux propositions α et β , on en définit d'autres. D'abord, $\neg\alpha$, la négation logique, est fausse si α est vraie, et vraie si α est fausse : $v(\neg\alpha) = 1$ si $v(\alpha) = 0$, et réciproquement. Ensuite, on définit d'autres propositions

$\alpha \wedge \beta$ (conjonction logique : " α et β ")

$\alpha \vee \beta$ (disjonction logique : " α ou β ")

$\alpha \oplus \beta$ (disjonction exclusive : " α ou bien β ")

$\alpha \Rightarrow \beta$ (implication logique : " α implique β ", "si α , alors β ")

$\alpha \Leftrightarrow \beta$ (équivalence logique : " α est équivalent à β ", "ssi"),

par la table de vérité suivante (qui, non par hasard, ressemble celle vue en chapitre 2 sur les opérations sur les ensembles) :

$v(\alpha)$	$v(\beta)$	$v(\alpha \wedge \beta)$	$v(\alpha \vee \beta)$	$v(\alpha \oplus \beta)$	$v(\alpha \Rightarrow \beta)$	$v(\alpha \Leftrightarrow \beta)$
0	0	0	0	0	1	1
1	1	1	1	0	1	1
0	1	0	1	1	1	0
1	0	0	1	1	0	0

Par exemple, la troisième colonne dit que $\alpha \wedge \beta$ est vraie si α et β le sont, et faux dans les autres 3 cas. L'avant-dernière colonne dit que l'implication logique $\alpha \Rightarrow \beta$ est fausse si [α est vraie et β est fausse], et elle est vraie dans les autres 3 cas. La dernière colonne dit que α et β sont équivalentes si elles ont mêmes valeurs de vérité.

Théorème 1. *Quelles que soient les propositions α , β , les propositions suivantes ont même valeur de vérité et sont donc équivalentes : " $\alpha \Rightarrow \beta$ " et " $\neg(\alpha \wedge \neg\beta)$ ".*

1. On pourrait utiliser deux autres symboles, comme v et f ; mais le choix des symboles 1 et 0 est judicieux dans beaucoup de domaines, comme par exemple dans l'électronique numérique : 1 = courant passe, 0 = pas de courant, etc.

Preuve. En utilisant la table, on constate que $\neg(\alpha \wedge \neg\beta)$ est fausse si α est vraie et β est fausse, et vraie dans les autres 3 cas ; ainsi les valeurs de vérité sont les mêmes, ce qui signifie que ces propositions sont logiquement équivalentes.

Commentaire. Cette définition de l'implication logique peut paraître étrange : par exemple, la proposition “si $1 = 0$, alors $4^2 + 5^2 = 6^2$ ” est vraie, toute comme “si $1 = 0$, alors $3^2 + 4^2 = 5^2$ ”. Mais on peut la motiver en remarquant que la *seule* façon de retorque une affirmation de type “ α implique β ” est en montrant que α est vraie, mais β est fausse. Noter que ceci n'est pas forcément l'interprétation qu'on donne au mot “si” dans la vie réelle, où on suppose le plus souvent qu'il doit y être un “lien de causalité” quand on dit “si α alors β ”. Or, la notion de “causalité” n'a de sens qu'en présence de celle du temps (la cause “précède” l'effet) ; elle a sa place dans les sciences naturelles, mais non en mathématiques. En effet, on peut douter que les règles du calcul propositionnel ci-dessus s'appliquent à la “logique quotidienne du monde réel” : pourquoi seulement deux valeurs de vérité, 0 et 1 ; comment décider si une proposition portant sur le “monde réel” est vraie ou fausse, etc. ? Les sciences expérimentales doivent tenir compte de cette situation complexe peu claire ; ainsi leurs règles de déduction et de raisonnement sont souvent différentes de celles utilisées en maths.

A partir de la table, on peut prouver les *lois de la logique formelle* les plus importantes. Le modus ponens, ou *règle de détachement*, est la règle de base du raisonnement logique.

Théorème 2 (modus ponens). *Si α est vraie, et $\alpha \Rightarrow \beta$ est vraie, alors β est vraie aussi.*

Preuve. Soit $v(\alpha) = 1$ et $v(\alpha \Rightarrow \beta) = 1$. Ce cas de figure correspond uniquement à la deuxième ligne du tableau ; et alors le tableau nous donne que $v(\beta) = 1$.

Théorème 3. *Quelle que soit la proposition α , on a :*

$v(\alpha \vee \neg\alpha) = 1$ (“ α ou non α ” est toujours vrai : tertium non datur = tiers exclu),

$v(\alpha \wedge \neg\alpha) = 0$ (“ α et non α ” est impossible : principe de non-contradiction).

Preuve. Si $v(\alpha) = 0$, alors $v(\neg\alpha) = 1$, et donc $v(\alpha \vee \neg\alpha) = 1$. Si $v(\alpha) = 1$, alors $v(\neg\alpha) = 0$, et donc $v(\alpha \vee \neg\alpha) = 1$. Dans tous les cas, $v(\alpha \vee \neg\alpha) = 1$. De la même façon, on prouve la deuxième affirmation.

Théorème 4. *Quelles que soient les propositions α, β , pour chaque item suivant, les propositions ont même valeur de vérité et sont donc équivalentes :*

- (1) $\neg(\neg\alpha)$ et α (double négation)
- (2) $\neg(\alpha \vee \beta)$ et $(\neg\alpha) \wedge (\neg\beta)$;
 $\neg(\alpha \wedge \beta)$ et $(\neg\alpha) \vee (\neg\beta)$ (lois de de Morgan) ;
- (3) $\alpha \Rightarrow \beta$ et $(\neg\beta) \Rightarrow (\neg\alpha)$ (contraposition)
- (4) $(\alpha \vee \beta) \vee \gamma$ et $\alpha \vee (\beta \vee \gamma)$ (associativité) ; idem pour $\wedge$;
 $\alpha \vee \beta$ et $\beta \vee \alpha$ (commutativité) ; idem pour $\wedge$;
- (5) $\alpha \vee (\beta \wedge \gamma)$ et $(\alpha \vee \beta) \wedge (\alpha \vee \gamma)$ (distributivité) ;
- (6) $\alpha \Leftrightarrow \beta$ et $\neg(\alpha \oplus \beta)$, et $(\alpha \Rightarrow \beta) \wedge (\beta \Rightarrow \alpha)$ (double implication),

Preuve. Les preuves sont presque identiques à celles du théorème 1 du §1 Chap.2, où nous avons explicité le cas d'une des lois de de Morgan : on dresse les tables de vérité, et on constate que les valeurs de vérité sont identiques pour les propositions en question.

Fin cours du 18/12/19 en EI 3. Voir TD 6, exercices 1, 2, 3, à préparer. Lire la suite du cours :

Exemple. Pour prouver une proposition du type $\alpha \Leftrightarrow \beta$, on procède souvent en démontrant séparément la direction $\beta \Rightarrow \alpha$ (on dit souvent : α est “condition nécessaire” pour β), puis $\alpha \Rightarrow \beta$ (on dit : α est “condition suffisante” pour β).

§3 Calcul des prédicats. Un *prédicat* $P(x)$ est une proposition P qui dépend d’une “variable” x : pour quelques x , la proposition $P(x)$ peut être vraie, et pour d’autres x , elle peut être fausse. En *calcul des prédicats*, ou : *logique de premier ordre*, on étudie les règles d’opération sur ces expressions. Dans le contexte d’une formule mathématique, la variable x appartient toujours à un ensemble (qu’il convient de préciser). Ainsi, en mathématiques, la logique de premier ordre est une combinaison de la théorie des ensembles avec le calcul des propositions. Par exemple, la loi de Morgan se généralise par la “règle de négation” suivante : “non quel que soit” devient “il existe non” (voir TD pour des exemples),

$\neg(\forall x : P(x))$ équivaut à $\exists x : \neg P(x)$

$\neg(\exists x : P(x))$ équivaut à $\forall x : \neg P(x)$

§4 Théorèmes et preuves. Un *théorème* est une affirmation (mathématique ou logique) qui peut être démontrée, c’est-à-dire une assertion qui peut être établie comme vraie au travers d’un raisonnement logique. Un théorème peut être *démontré* de plusieurs façons – par exemple,

- par disjonction de cas,
- par contre-exemple,
- par récurrence,
- par analyse-synthèse,
- par l’absurde.

Il n’y a pas de “méthode générale” : c’est pendant les études de maths qu’on apprend quand et comment on applique ces types de raisonnement – vous en avez déjà vu quelques-uns. Souvent, en cherchant une preuve, on commence par “analyser” la situation (étude de cas particuliers, d’exemples numériques, faire des diagrammes et des figures...) ; la structure logique de la preuve en ressort lentement ; une fois qu’on a compris cette structure, la rédaction de la preuve “part à l’inverse” (i.e., le point de départ de la rédaction sont des principes généraux, et non l’analyse de cas particuliers).

Quant à la structure de l’énoncé du théorème, on distingue :

- les hypothèses, et
- la conclusion.

La forme générale d’un théorème est donc “si α [hypothèses], alors β [conclusion]”, ou : “Supposons α [hypothèses]. Alors β [conclusion]”. Souvent on se dispense d’énumérer toutes les hypothèses ; mais il est important de pouvoir le faire, si besoin est !

Annexe 1 : Sur la construction de $\mathbb{Z}$

§1 Construction de l'ensemble $\mathbb{Z}$.

Lemme 1. Nous définissons une relation R sur $M = \mathbb{N}^2$ par :

$$(a, b)R(a', b') \quad \text{ssi} \quad a + b' = a' + b.$$

Alors R est une relation d'équivalence.

Preuve : vérification directe des 3 propriétés (cf. TD)

Définition 1. Rappelons la définition de la classe d'équivalence $[x] = \{y \in M \mid xRy\}$, pour tout $x \in M$. Alors on définit $\mathbb{Z}$ comme étant l'ensemble des classes d'équivalence de la relation R définie dans le lemme 1 :

$$\mathbb{Z} = \{[(a, b)] \mid (a, b) \in \mathbb{N}^2\}.$$

D'après la définition, $\mathbb{Z}$ est une certaine partie de $\mathcal{P}(M)$, et on a $[(a, b)] = [(a', b')] \text{ ssi } a + b' = a' + b$.

Définition 2. L'application canonique de $\mathbb{N}^2$ dans $\mathbb{Z}$ est l'application

$$\pi : \mathbb{N}^2 \rightarrow \mathbb{Z}, \quad (a, b) \mapsto \pi(a, b) := [(a, b)].$$

On pose ces définitions pour n'importe quelle relation d'équivalence sur n'importe quel ensemble. Alors l'application canonique est *surjective*, car (a, b) est un antécédent de $[(a, b)]$, et les fibres de π sont précisément les classes d'équivalence (cf. §1).

§2 L'addition dans $\mathbb{Z}$. Expliquons comment additionner deux classes $[(a, b)]$ et $[(c, d)]$:

Lemme 2. Soit $m = [(a, b)] = [(a', b')]$, $n = [(c, d)] = [(c', d')] \in \mathbb{Z}$. Alors

$$[(a + c, b + d)] = [(a' + c', b' + d')],$$

et ainsi l'addition $m + n := [(a + c, b + d)]$ est bien définie, i.e., elle ne dépend pas des représentants des classes m et n qu'on a choisis.

Preuve. $(a + c) + (b' + d') = (a + b') + (c + d') = (b + a') + (c' + d) = (b + d) + (a' + c')$

Théorème 1. L'addition dans $\mathbb{Z}$ a les propriétés suivantes :

1. elle est associative,
2. elle est commutative,
3. elle a un élément neutre $0 = [(0, 0)]$,
4. l'équation $m + x = n$ admet une unique solution, pour tout $m, n \in \mathbb{Z}$. On la note $x = n - m$.

Preuve. 1., 2., 3. : direct ; 4. l'unique solution est $[(a, b)] - [(c, d)] = [(a + d, b + c)]$

Remarque : le théorème dit que $(\mathbb{Z}, +)$ est un *groupe commutatif*.

§3 Relation d'ordre sur $\mathbb{Z}$.

Lemme 3. L'application

$$\kappa : \mathbb{N} \rightarrow \mathbb{Z}, \quad n \mapsto [(n, 0)]$$

est injective, et elle vérifie $\kappa(a + b) = \kappa(a) + \kappa(b)$.

Définition 3. Si $n \in \mathbb{N}$, nous allons identifier l'entier $[(n, 0)]$ avec n , et ainsi considérer $\mathbb{N}$ comme une partie de $\mathbb{Z}$. Alors, pour $(m, n) \in \mathbb{Z}^2$, on écrira

$$m \leq n \quad \text{ssi} \quad n - m \in \mathbb{N}.$$

Théorème 2. La définition précédente défine une relation d'ordre total sur $\mathbb{Z}$, qui coïncide avec celle de $\mathbb{N}$ sur la partie $\mathbb{N} \subset \mathbb{Z}$.

§4 Produit dans $\mathbb{Z}$. Le calcul littéral usuel donne

$$(a - b)(c - d) = (ac + bd) - (ad + bc).$$

Motivé par cela, on pose :

Lemme 4. Soit $m = [(a, b)] = [(a', b')]$, $n = [(c, d)] = [(c', d')] \in \mathbb{Z}$. Alors

$$[(ac + bd, ad + bc)] = [(a'c' + b'd', a'd' + b'c')],$$

et ainsi le produit $m \cdot n := [(ac + bd, ad + bc)]$ est bien défini.

Théorème 3. Addition et multiplication dans $\mathbb{Z}$ ont les propriétés suivantes :

1. $(\mathbb{Z}, +)$ est un groupe commutatif,
2. la multiplication est associative et commutative,
3. la multiplication a un élément neutre $1 = [(1, 0)]$,
4. on a la loi de distributivité : $(k + m)n = kn + mn$.

Preuve. Vérification directe.

Définition 4. On résume les propriétés du théorème précédent en disant que $(\mathbb{Z}, +, \cdot)$ est un anneau commutatif.

Remarque. Nous venons de démontrer le “théorème” suivant : Si $\mathbb{N}$ existe, alors $\mathbb{Z}$ existe aussi. Il est important de connaître ce fait ; mais dans la suite, nous allons “oublier” la construction de $\mathbb{Z}$ par classes d'équivalence : il faut retenir les *propriétés* de $\mathbb{Z}$, et non une construction particulière (il en existe d'autres !). Cependant, le principe général de la construction est très important en maths : souvent, on construit de nouveaux objets en utilisant des relations et classes d'équivalence – la construction des nombres rationnels $\mathbb{Q}$, et celle des nombres réels $\mathbb{R}$ en sont d'autres exemples.

Annexe 2 : Sur la construction de $\mathbb{Q}$

§1 Construction de l'ensemble $\mathbb{Q}$.

Lemme 1. Soit $M = \mathbb{Z} \times (\mathbb{Z} \setminus \{0\})$. Alors la relation suivante est une relation d'équivalence sur M :

$$(a, b)R(a', b') \quad \text{ssi} \quad ab' = a'b.$$

Définition 1. Soit $\mathbb{Q}$ l'ensemble des classes d'équivalence de la relation R du lemme 1, et $\pi : M \rightarrow \mathbb{Q}$ l'application canonique.

Théorème 1. Pour $[(a, b)], [(c, d)] \in \mathbb{Q}$, les éléments suivants sont bien définis :

$$\begin{aligned} [(a, b)] + [(c, d)] &= [(ad + bc, bd)], \\ [(a, b)] - [(c, d)] &= [(ad - bc, bd)], \\ [(a, b)] \cdot [(c, d)] &= [(ac, bd)], \\ [(a, b)] : [(c, d)] &= [(ad, bc)] \quad (\text{ si } c \neq 0), \end{aligned}$$

et $\mathbb{Q}$ muni de ces 4 opérations est un corps. De plus, l'application

$$f : \mathbb{Z} \rightarrow \mathbb{Q}, \quad m \mapsto [(m, 1)]$$

est injective et vérifie $f(ab) = f(a)f(b)$ et $f(a+b) = f(a)+f(b)$. On pourra donc identifier $\mathbb{Z}$ avec la partie $\text{im}(f) \subset \mathbb{Q}$, et la classe $[(a, b)]$ avec la fraction $a : b = \frac{a}{b}$.

Preuve. Les vérifications sont similaires à celles du l'annexe précédent, et nous n'allons pas les détailler ici (en cours d'algèbre de L2-L3 vous allez les revoir ; mot-clef : “corps des fractions”). Signalons juste que, pour montrer que les opérations sont bien définies, nous avons besoin d'une propriété particulière du produit dans $\mathbb{Z}$: $\mathbb{Z}$ est *intègre* (cf. TD). De même, nous admettons la preuve du théorème suivant :

Théorème 2. Le corps $\mathbb{Q}$ admet un ordre total $\leq$ tel que :

1. si $a, b \in \mathbb{Z}$, alors $a \leq b$ dans $\mathbb{Q}$ ssi $a \leq b$ dans $\mathbb{Z}$;
2. si $a \leq b$, alors pour tout $c \in \mathbb{Q}$, on a : $a + c \leq b + c$,
3. si $a \leq b$ et $c \geq 0$ dans $\mathbb{Q}$, alors $ac \leq bc$.

Pour résumer, on dit que $(\mathbb{Q}, \leq)$ est un corps ordonné.

Pour plus d'informations, sur la construction des entiers et des nombres rationnels et des questions liées, voici quelques pistes de lecture :

https://fr.wikipedia.org/wiki/Nombre_rationnel,

<https://fr.wikipedia.org/wiki/Nombre>,

https://fr.wikipedia.org/wiki/Entier_naturel,

https://fr.wikipedia.org/wiki/Construction_du_nombre_chez_l%27enfant