

Cours de Licence 3: Géométrie et Topologie

Cours de: Laurent Bessières

Rédigé par Hugo Clouet*

Année universitaire 2023 - 2024

Version: 9 novembre 2025

Table des matières

1	Topologie générale	4
1.1	Espaces topologiques	4
1.1.1	Topologie	4
1.1.2	Notions fondamentales : fermés, voisinages, intérieur, adhérence	5
1.1.3	Bases d'une topologie	6
1.1.4	Continuité	7
1.2	Opérations sur les espaces topologiques	9
1.2.1	Produit	9
1.2.2	Union disjointe	10
1.2.3	Topologie quotient	11
1.2.4	Séparation	12
1.2.5	Factorisation d'applications	12
1.2.6	Exemples : cône, suspension, recollement, somme pointée	13
1.3	Connexité, connexité par arc	14
1.3.1	Connexité	14
1.3.2	Connexité par arc	16
1.4	Compacité	16
1.4.1	Propriété de Borel-Lebesgue	16
1.4.2	Compacité dans les espaces métriques	17
2	Variétés	19
2.1	Variétés topologiques de dimension m	19
2.2	Sphères et projectifs	20
2.2.1	Sphères	20
2.2.2	Espaces projectifs	21
2.3	Opérations sur les variétés topologiques	22
2.3.1	Quotients de variétés	22
2.3.2	Exemples de surfaces obtenues comme quotient	24
2.3.3	Somme connexe	25
3	Groupe fondamental	27
3.1	Groupe fondamental	27
3.2	Espace simplement connexe	29
3.3	Effet d'applications continues	30
3.4	Homotopie et revêtement	32
3.5	Applications du calcul de $\pi_1(\mathbf{S}^1)$	34
3.6	Théorème de Van Kampen	34
3.6.1	Produit libre de groupes et groupe libre	34
3.6.2	Théorème de Van Kampen	37
3.6.3	Applications	38

*email: hugo.clouet@etu.u-bordeaux.fr

4	Classification des surfaces compactes	43
4.1	Quotients de polygones	43
4.2	Classification : opérations sur les symboles	44
4.2.1	Symbole et somme connexe	44
4.2.2	Simplification de symboles	45
4.3	Toute surface connexe est quotient de P	47
4.3.1	Triangulations	48
4.3.2	Preuve du théorème	49

Introduction

Ce cours est une initiation à la topologie algébrique, branche des mathématiques qui utilise des outils d'algèbre pour étudier des espaces topologiques.

De quoi s'agit-il ? Une question fondamentale pour les topologues est de classer les espaces topologiques. On veut, par exemple, faire la liste de toutes les surfaces, à homéomorphisme près, pour savoir si deux surfaces données sont homéomorphes.

L'idée centrale pour attaquer ce genre de question est d'associer à l'objet topologique un objet algébrique (entier, groupe, ...) de manière qu'à deux objets topologiques homéomorphes soient associées deux objets algébriques isomorphes.

Un des premiers succès de la topologie algébrique (aux alentours des années 1900) a été la classification des surfaces, qui sera démontrée dans ce cours.

Nous commencerons par introduire les principales notions de topologie générale, qui sera plus efficace pour construire surfaces et variétés (l'analogue n -dimensionnel d'une surface) que les simples espaces métriques. En particulier on verra comment classer les surfaces en les considérant comme des quotients de polygones où les côtés sont recollés selon certaines règles.

On abordera également la notion de groupe fondamental, groupe associé à un espace topologique.

Mentionnons que le problème de topologie le plus célèbre du 20^e siècle, la *conjecture de Poincaré*, est une question de classification. On peut l'énoncer comme suit : « une variété de dimension 3 compacte, simplement connexe est homéomorphe à $\mathbf{S}^3$. » (Poincaré, 1904).

Cette conjecture à 1 million de dollars fait partie des 7 problèmes du millénaire de l'Institut Clay. Elle a été résolue en 2002 par *Grigori Perelman*, par des méthodes d'analyse et de géométrie.

1 Topologie générale

Une partie $O \subset E$ d'un espace métrique est *ouverte* si

$$(\forall x \in O) (\exists \varepsilon > 0) \quad B(x, \varepsilon) \subset O.$$

L'ensemble des ouverts de E , qu'on appelle sa *topologie*, est stable par union quelconque et intersection finie. Elle contient $\emptyset$ et E . Sur un ensemble X quelconque on prend ces propriétés comme axiomes pour définir une topologie. On peut alors définir les notions de limite, continuité, compacité, connexité, etc... en terme d'ouverts, sans requérir de distance.

1.1 Espaces topologiques

1.1.1 Topologie

Définition 1.1.1.1 (TOPOLOGIE). Une *topologie* sur un ensemble X est un ensemble $\mathcal{T}$ de parties de X vérifiant :

- (i) $\emptyset$ et X sont dans $\mathcal{T}$;
- (ii) une union quelconque d'éléments de $\mathcal{T}$ appartient à $\mathcal{T}$;
- (iii) une intersection finie d'éléments de $\mathcal{T}$ appartient à $\mathcal{T}$;

Le couple $(X, \mathcal{T})$ est un *espace topologique* et les éléments de $\mathcal{T}$ sont des *ouverts*.

L'axiome (iii) peut être remplacé par : l'intersection de deux éléments de $\mathcal{T}$ appartient à $\mathcal{T}$. L'ensemble des ouverts d'un espace métrique forme une topologie. Tout ensemble X admet des topologies : la *topologie grossière* $\mathcal{T}_g = \{\emptyset, X\}$; la *topologie discrète* $\mathcal{T}_d = \mathcal{P}(X)$ formée des parties de X . On dit qu'une topologie $\mathcal{T}'$ est plus *fine* que $\mathcal{T}$ si $\mathcal{T} \subset \mathcal{T}'$ (elle a plus d'ouverts). On pourra noter un espace topologique simplement X s'il n'y a pas d'ambiguïté.

Intersection de topologies : Soit $(\mathcal{T}_i)_{i \in I}$ une famille de topologies sur un ensemble X . On définit l'intersection de topologies par :

$$\bigcap_{i \in I} \mathcal{T}_i = \{U \subset X; U \in \mathcal{T}_i (\forall i \in I)\}$$

Proposition 1.1.1.2 (INTERSECTION DE TOPOLOGIES). Si $(\mathcal{T}_\alpha)_{\alpha \in I}$ est une famille de topologies sur X , alors $\bigcap_{\alpha \in I} \mathcal{T}_\alpha$ est une topologie sur X .

Démonstration. Soit $(\mathcal{T}_\alpha)_{\alpha \in I}$ une famille de topologies sur X . Vérifions simplement les axiomes :

- Comme $X, \emptyset \in \mathcal{T}_\alpha$ pour tout $\alpha \in I$, alors $X, \emptyset \in \bigcap_{\alpha \in I} \mathcal{T}_\alpha$.
- Soit $(O_j)_{j \in J}$ une famille d'ouverts de $\bigcap_{\alpha \in I} \mathcal{T}_\alpha$ i.e. $(\forall j \in J) (\forall \alpha \in I) O_j \in \mathcal{T}_\alpha$. On a donc directement : $(\forall \alpha \in I) \bigcup_{j \in J} O_j \in \mathcal{T}_\alpha$ puisque les $\mathcal{T}_\alpha$ sont des topologies sur X , d'où $\bigcup_{j \in J} O_j \in \bigcap_{\alpha \in I} \mathcal{T}_\alpha$.
- Soit $V_1, \dots, V_n$ une famille finies d'ouverts de $\bigcap_{\alpha \in I} \mathcal{T}_\alpha$ i.e. $(\forall k \in \llbracket 1, n \rrbracket) (\forall \alpha \in I) V_k \in \mathcal{T}_\alpha$. On a donc : $(\forall \alpha \in I) \bigcap_{k=1}^n V_k \in \mathcal{T}_\alpha$, d'où $\bigcap_{k=1}^n V_k \in \bigcap_{\alpha \in I} \mathcal{T}_\alpha$.

On a bien montré que $\bigcap_{\alpha \in I} \mathcal{T}_\alpha$ est une topologie sur X . □

Corollaire 1.1.1.3 (TOPOLOGIE ENGENDRÉE). Soit X un ensemble et soit $\mathcal{A} \subset \mathcal{P}(X)$. On appelle *topologie engendrée* par $\mathcal{A}$ l'intersection $\mathcal{T}_{\mathcal{A}}$ des topologies contenant $\mathcal{A}$. C'est la plus petite (pour l'inclusion) topologie sur X contenant $\mathcal{A}$.

Ses ouverts non vides sont les unions quelconques d'intersection finies d'éléments de $\mathcal{A}$.

Définition 1.1.1.4. Soit $(X, \mathcal{T})$ un espace topologique et soit $Y \subset X$. L'ensemble

$$\mathcal{T}_Y := \{U \cap Y; U \in \mathcal{T}\}$$

est une topologie sur Y , appelée la *topologie induite*.

C'est une vérification immédiate. Lorsqu'on omettra les topologies, on dira simplement que U est ouvert dans Y s'il est ouvert pour la topologie induite et ouvert dans X s'il appartient à la topologie de X .

Remarque 1.1.1.5. Attention ! L'ensemble $Y = [0, 1] \subset \mathbf{R} = X$ est ouvert dans Y mais pas dans X . Aussi, $[0, 1/2[$ est ouvert dans Y mais pas dans X . Cependant, quand Y est ouvert dans X , ses ouverts sont ouverts dans X .

Définition 1.1.1.6 (TOPOLOGIE MÉTRISABLE, SÉPARÉE). Une topologie $\mathcal{T}$ sur X est *métrisable* s'il existe une distance d sur X dont les ouverts sont ceux de $\mathcal{T}$. Elle est *séparée* si pour tout $x \neq y$ dans X , il existe deux ouverts U_x, U_y tels que

$$(x \in U_x) \ (y \in U_y) \quad U_x \cap U_y = \emptyset.$$

Une topologie métrisable est séparée car si $x \neq y$, alors les boules ouvertes $U_x = B(x, d(x, y)/2)$ et $U_y = B(y, d(x, y)/2)$ sont disjointes. Une topologie non séparée n'est donc pas métrisable. La topologie discrète est métrisable (en posant $d(x, y) = 1$ si $x \neq y$ et $d(x, x) = 0$) et la topologie grossière ne l'est pas si X a au moins deux éléments car elle n'est pas séparée.

1.1.2 Notions fondamentales : fermés, voisinages, intérieur, adhérence

Définition 1.1.2.1. Une partie A d'un espace topologique X est *fermée* si $X \setminus A$ est ouverte.

Immédiatement, $\emptyset$ et X sont fermés et il résulte de

$$X \setminus \bigcap_{\alpha} U_{\alpha} = \bigcup_{\alpha} (X \setminus U_{\alpha}) \quad \text{et} \quad X \setminus \bigcup_{\alpha} U_{\alpha} = \bigcap_{\alpha} (X \setminus U_{\alpha})$$

que l'ensemble des fermés est stable par union finie et intersection quelconque.

Exercice 1.1.2.2. Soit X un espace topologique et $Y \subset X$. Les fermés de Y sont les intersections de Y avec les fermés de X .

Conséquence immédiate : si Y est fermé, ses fermés sont fermés dans X .

Définition 1.1.2.3 (VOISINAGES). Une partie $V \subset X$ d'un espace topologique X est un *voisinage* d'un point $x \in X$ s'il existe un ouvert U tel que $x \in U \subset V$. On note $\mathcal{V}_x$ l'ensemble des voisinages de x .

Une partie de X est ouverte si et seulement si elle est voisinage de chacun de ses points. Il est clair qu'une partie V' contenant $V \in \mathcal{V}_x$ est aussi un voisinage de x et qu'une intersection finie (ou union quelconque) de voisinages de x est un voisinage de x .

Définition 1.1.2.4 (BASE DE VOISINAGES D'UN POINT). Soit X un espace topologique et soit $x \in X$. Une famille $\mathcal{W} \subset \mathcal{V}_x$ de voisinages de x est une *base de voisinages* ou *système fondamental de voisinages* en x si pour tout voisinage $V \in \mathcal{V}_x$, il existe un voisinage $W \in \mathcal{W}$ tel que $x \in W \subset V$.

Dans un espace métrique, $\mathcal{W} = \{B(x, \varepsilon); \varepsilon > 0\}$ est une base de voisinages (ouverts) en x .

Définition 1.1.2.5 (INTÉRIEUR, ADHÉRENCE, FRONTIÈRE). Soit X un espace topologique et soit $A \subset X$.

- (i) Un point $a \in A$ est *intérieur* à A si A est voisinage de a i.e. $(\exists O \in \mathcal{T}) \ a \in O \subset A$. L'*intérieur* de A est l'ensemble des points intérieurs à A , noté $\text{Int}(A)$.
- (ii) Un point $x \in X$ est *adhérent* à A si tout voisinage de x rencontre A i.e. $(\forall V \in \mathcal{V}_x) \ V \cap A \neq \emptyset$. L'*adhérence* de A est l'ensemble des points adhérents à A , noté $\overline{A}$.
- (iii) La *frontière* de A est $\partial A = \overline{A} \setminus \text{Int}(A)$.

Proposition 1.1.2.6 (PROPRIÉTÉS DE L'INTÉRIEUR ET DE L'ADHÉRENCE). Soit X un espace topologique et A une partie de X .

- (i) Pour U ouvert, $\text{Int}(A) = \bigcup_{U \subset A} U$ est un ouvert de X , c'est le plus grand contenu dans A .
- (ii) Pour F fermé, $\overline{A} = \bigcap_{A \subset F} F$ est fermé dans X , c'est le plus petit contenant A .
- (iii) On a $\overline{X \setminus A} = X \setminus \text{Int}(A)$ et $\text{Int}(X \setminus A) = X \setminus \overline{A}$.
- (iv) Si $A \subset B$, alors $\text{Int}(A) \subset \text{Int}(B)$ et $\overline{A} \subset \overline{B}$.

Démonstration. (i) Si $x \in \text{Int}(A)$, alors A est voisinage de x et par définition il existe U un ouvert tel que $x \in U \subset A$, donc $x \in \bigcup_{U \subset A} U$, d'où l'inclusion $\text{Int}(A) \subset \bigcup_{U \subset A} U$. L'autre inclusion est évidente. Il s'ensuit que $\text{Int}(A)$ est une union d'ouverts donc est ouvert et que tout ouvert $O \subset A$ vérifie $O \subset \text{Int}(A)$.

(ii) Soit $x \in \overline{A}$ et soit F fermé contenant A . Si $x \notin F$, alors $x \in X \setminus F$ est ouvert, donc il existe $V \in \mathcal{V}_x$ contenu dans $X \setminus F$ i.e. disjoint de F et *a fortiori* de A , contredisant $x \in \overline{A}$, d'où l'inclusion $\overline{A} \subset F$ puis $\overline{A} \subset \bigcap_{A \subset F} F$.

Réciproquement, supposons que $x \notin \overline{A}$, alors il existe $V \in \mathcal{V}_x$ disjoint de A . On peut supposer V ouvert, alors $F = X \setminus V$ est un fermé contenant A et $x \notin F$, donc $x \notin \bigcap_{A \subset F} F$, d'où l'égalité voulue. Il s'ensuit que $\overline{A}$ est fermé

comme intersection de fermés et clairement $A \subset \overline{A}$.

(iii) On a les équivalences :

$$\begin{aligned} x \in \overline{X \setminus A} &\Leftrightarrow \text{tout voisinage de } x \text{ rencontre } X \setminus A \\ &\Leftrightarrow \text{il n'existe pas de voisinage de } x \text{ contenu dans } A \\ &\Leftrightarrow x \notin \text{Int}(A) \\ &\Leftrightarrow x \in X \setminus \text{Int}(A). \end{aligned}$$

En appliquant cette égalité à $B = X \setminus A$ et en prenant le complémentaire, on obtient $X \setminus \overline{B} = X \setminus \text{Int}(B)$ la deuxième égalité.

(iv) Soit $A \subset B$. Si $x \in \text{Int}(A)$, alors il existe U_x ouvert tel que $x \in U_x \subset A$, donc $x \in U_x \subset B$ et $x \in \text{Int}(B)$. La partie $\overline{B}$ est un fermé et $A \subset B \subset \overline{B}$, or $\overline{A}$ est le plus petit fermé contenant A , donc $\overline{A} \subset \overline{B}$. $\square$

Rappel. Une *suite* dans X est une application $x: \mathbf{N} \rightarrow X$, notée $(x_n)_{n \in \mathbf{N}}$ ou simplement (x_n) où $x_n = x(n)$.

Définition 1.1.2.7. Soit X un espace topologique et soit $(x_n)_{n \in \mathbf{N}}$ une suite de X . On dit que $(x_n)_{n \in \mathbf{N}}$ *converge* vers $x \in X$ si pour tout voisinage V de x , il existe $n_0 \in \mathbf{N}$ tel que pour tout $n \geq n_0$, on a $x_n \in V$.

Remarque 1.1.2.8. « La » limite d'une suite n'est pas forcément unique : pour la topologie grossière, une suite quelconque converge vers tout les points ! Si la topologie est séparée, la limite est unique (exercice).

Lemme 1.1.2.9 (CRITÈRE SÉQUENTIEL). Soit A une partie d'un espace topologique X et soit $x \in X$. S'il existe une suite $(x_n)_{n \in \mathbf{N}}$ de A convergeant vers x , alors $x \in \overline{A}$. La réciproque est vraie si on suppose que x admet une base de voisinages dénombrable.

Démonstration. • Supposons que $(x_n)_{n \in \mathbf{N}}$ est une suite dans A qui converge vers x dans X . Soit V un voisinage de x dans X . Il existe $N \in \mathbf{N}$ tel que pour tout $n \geq N$, on a $x_n \in V$. En particulier, $x_N \in V \cap A$, donc $V \cap A \neq \emptyset$.
• Réciproquement, supposons $x \in \overline{A}$. On utilise l'existence d'une base dénombrable de voisinages. Ici, $\mathcal{B}_x = \{B(x, \frac{1}{n}); n \in \mathbf{N}_{>0}\}$ est une base de voisinages de x . Soit $n \in \mathbf{N}_{>0}$, on a $B(x, \frac{1}{n}) \cap A \neq \emptyset$ par hypothèse. On choisit $x_n \in B(x, \frac{1}{n}) \cap A$. On a alors construit $(x_n)_{n \in \mathbf{N}_{>0}}$ dans A . Vérifions que $\lim_{n \rightarrow +\infty} x_n = x$. Soit V un voisinage de x . On sait que $\mathcal{B}_x$ est une base, donc il existe $N \in \mathbf{N}_{>0}$ tel que $B(x, \frac{1}{N}) \subset V$. Pour $n \geq N$, on a alors $x_n \in B(x, \frac{1}{n}) \subset B(x, \frac{1}{N}) \subset V$. $\square$

Dans un espace métrique, $\mathcal{W} = \{B(x, \frac{1}{n}); n \in \mathbf{N}_{>0}\}$ est une base de voisinages dénombrable en x , donc le critère séquentiel s'applique.

1.1.3 Bases d'une topologie

Définition 1.1.3.1 (BASE DE TOPOLOGIE). Soit $\mathcal{T}$ une topologie sur X et soit $\mathcal{B} \subset \mathcal{T}$. On dit que $\mathcal{B}$ est une *base* de $\mathcal{T}$ si tout élément non vide de $\mathcal{T}$ est réunion d'éléments de $\mathcal{B}$.

De manière équivalente : $(\forall U \in \mathcal{T}) (\forall x \in U) (\exists B \in \mathcal{B}) x \in B \subset U$. L'ensemble des boules ouvertes d'un espace métrique est une base de sa topologie.

Proposition 1.1.3.2 (CARACTÉRISATION DES BASES). Soient X un ensemble et $\mathcal{B} \subset \mathcal{P}(X)$, alors $\mathcal{B}$ est la base d'une topologie sur X si et seulement si :

- (i) pour chaque $x \in X$, il existe $B \in \mathcal{B}$ tel que $x \in B$;
- (ii) pour tout $B_1, B_2 \in \mathcal{B}$ et $x \in B_1 \cap B_2$, il existe $B_3 \in \mathcal{B}$ tel que $x \in B_3 \subset B_1 \cap B_2$.

La topologie dont $\mathcal{B}$ est la base est unique. Elle est formée de l'ensemble vide et des réunions d'éléments de $\mathcal{B}$. C'est la topologie $\mathcal{T}_{\mathcal{B}}$ engendrée par $\mathcal{B}$.

On dit que $\mathcal{B}$ est une *base de topologie*. L'axiome (i) équivaut à ce que X soit réunion d'éléments de $\mathcal{B}$ et (ii) à ce qu'une intersection finie d'éléments de $\mathcal{B}$ soit réunion d'éléments de $\mathcal{B}$.

Démonstration. • Soit $\mathcal{T}$ une topologie dont $\mathcal{B}$ est la base. On a $X \in \mathcal{T}$, donc X est réunion d'éléments de $\mathcal{B}$ par définition d'une base et (i) est vérifié. Soit $B_1, B_2 \in \mathcal{B} \subset \mathcal{T}$, alors $B_1 \cap B_2$ est ouvert et l'axiome (ii) suit par définition d'une base.

• Soit $\mathcal{B} \subset \mathcal{P}(X)$ vérifiant (i) et (ii). Considérons $\mathcal{T} \subset \mathcal{P}(X)$ formé de l'ensemble vide et des réunions d'éléments de $\mathcal{B}$. Montrons que c'est une topologie, nécessairement la plus petite contenant $\mathcal{B}$. On a déjà que $\emptyset \in \mathcal{T}$ par définition et $X \in \mathcal{T}$ par (i). Soit $(U_\alpha)_\alpha$ une collection d'éléments de $\mathcal{T}$, on doit montrer que $\bigcup_\alpha U_\alpha \in \mathcal{T}$ i.e. (en supposant U_α non vide) qu'il est réunion d'éléments de $\mathcal{B}$. Par définition de $\mathcal{T}$, chaque U_α est réunion d'éléments de $\mathcal{B}$, donc $\bigcup_\alpha U_\alpha \in \mathcal{T}$ également. Soient U_1, U_2 deux éléments de $\mathcal{T}$, on doit montrer que $U_1 \cap U_2 \in \mathcal{T}$. Soit

$x \in U_1 \cap U_2$, alors $x \in U_1 \in \mathcal{T}$, donc il existe $B_1 \in \mathcal{B}$ tel que $x \in B_1 \subset U_1$. De même, il existe $B_2 \in \mathcal{B}$ tel que $x \in B_2 \subset U_2$. En particulier, $x \in B_1 \cap B_2$. Par (ii), il existe $B_x \in \mathcal{B}$ tel que $x \in B_x \subset B_1 \cap B_2 \subset U_1 \cap U_2 = \bigcup_{x \in U_1 \cap U_2} B_x \in \mathcal{T}$. Puisque $\mathcal{B} \subset \mathcal{T}$, on a $\mathcal{T}_{\mathcal{B}} \subset \mathcal{T}$ par définition de la topologie engendrée. Réciproquement, tout élément $U \in \mathcal{T}$ non vide est réunion d'éléments de $\mathcal{B} \subset \mathcal{T}_{\mathcal{B}}$, donc appartient à $\mathcal{T}_{\mathcal{B}}$. On a ainsi l'égalité $\mathcal{T} = \mathcal{T}_{\mathcal{B}}$. $\square$

Exercice 1.1.3.3 (TD). Soit X un espace topologique et soit $\mathcal{B} \subset \mathcal{P}(X)$. Montrer que $\mathcal{B}$ est une base de la topologie de X si et seulement si $\mathcal{B} \cap \mathcal{V}_x$ est une base de voisinages de x pour tout $x \in X$.

Solution : Supposons que $\mathcal{B}$ est une base de $\mathcal{T}$. Soit $V \in \mathcal{V}_x$. Par définition, il existe $U \in \mathcal{T}$ tel que $x \in U \subset V$. Comme $\mathcal{B}$ est une base, il existe $B \in \mathcal{B}$ tel que $x \in B \subset U \subset V$, ce qui montre le sens direct. Réciproquement, supposons que $\mathcal{B} \cap \mathcal{V}_x$ est une base de voisinages de x pour tout $x \in X$. Soit $U \in \mathcal{T}$. Comme U est ouvert, il est voisinage de chacun de ses points *i.e.* $(\forall x \in U) U \in \mathcal{V}_x$. Par hypothèse, il existe $V_x \in \mathcal{B} \cap \mathcal{V}_x$ tel que $x \in V_x \subset U$. En particulier, $V_x \in \mathcal{B}$, ce qui conclut.

Proposition 1.1.3.4 (DESCRIPTION DE LA TOPOLOGIE ENGENDRÉE). Soient X un ensemble et $\mathcal{A} \subset \mathcal{P}(X)$. Soit $\mathcal{B}$ l'ensemble des intersections finies d'éléments de $\mathcal{A}$ (y compris l'intersection vide qui vaut X). L'ensemble $\mathcal{B}$ est alors une base de la topologie $\mathcal{T}_{\mathcal{A}}$ engendrée par $\mathcal{A}$.

Démonstration. Puisque $X \in \mathcal{B}$, l'axiome (i) de la proposition 1.1.3.2 est vérifiée. L'axiome (ii) est également vérifié puisque $\mathcal{B}$ est stable par intersection finie : si $B_1 = A_1 \cap \dots \cap A_k$, et $B_2 = A'_1 \cap \dots \cap A'_\ell$ sont deux éléments de $\mathcal{B}$ (les A_i et A'_j sont des éléments de $\mathcal{A}$), alors clairement $B_1 \cap B_2$ est une intersection finie d'éléments de $\mathcal{A}$, donc $B_1 \cap B_2 \in \mathcal{B}$. C'est alors une base d'une topologie, la topologie engendrée $\mathcal{T}_{\mathcal{B}}$. Toute topologie contenant $\mathcal{A}$ contient les réunions d'intersections finies d'éléments de $\mathcal{A}$, donc contient $\mathcal{T}_{\mathcal{B}}$. Il s'ensuit que $\mathcal{T}_{\mathcal{A}} = \mathcal{T}_{\mathcal{B}}$. $\square$

En résumé, $\mathcal{T}_{\mathcal{A}}$ est formée de l'ensemble vide et des réunions d'intersections finies d'éléments de $\mathcal{A}$.

1.1.4 Continuité

Sur les espaces métriques, la continuité d'une application se caractérise par « l'image réciproque de tout ouvert et un ouvert ». Sur les espaces topologiques, c'est la définition suivante.

Définition 1.1.4.1. Une application $f: X \rightarrow Y$ entre deux espaces topologiques est *continue* si, pour tout ouvert U de Y , l'ensemble $f^{-1}(U)$ est ouvert dans X .

Remarque 1.1.4.2. Si la topologie de Y est engendrée par $\mathcal{A} \subset \mathcal{P}(Y)$, il suffit de vérifier que $f^{-1}(A)$ est ouvert pour tout $A \in \mathcal{A}$. En effet, tout ouvert non vide U de Y est réunion d'intersections finies d'éléments de $\mathcal{A}$. L'union et l'intersection se comportent bien par image réciproque : pour toute famille $(A_i)_{i \in I}$ de partie de Y ,

$$f^{-1}\left(\bigcup_{i \in I} A_i\right) = \bigcup_{i \in I} f^{-1}(A_i) \quad \text{et} \quad f^{-1}\left(\bigcap_{i \in I} A_i\right) = \bigcap_{i \in I} f^{-1}(A_i).$$

L'ensemble $f^{-1}(U)$ est alors réunion d'intersections finies d'éléments de $\{f^{-1}(A); A \in \mathcal{A}\}$, qui sont des ouverts, donc c'est un ouvert.

Théorème 1.1.4.3 (CARACTÉRISATION DE LA CONTINUITÉ). Soient X et Y deux espaces topologiques et $f: X \rightarrow Y$ une application. Sont équivalents :

- (i) L'application f est continue.
- (ii) Pour toute partie $A \subset X$, on a $f(\overline{A}) \subset \overline{f(A)}$.
- (iii) Pour tout fermé B de Y , l'ensemble $f^{-1}(B)$ est fermé dans X .
- (iv) Pour tout $x \in X$ et pour tout voisinage V de $f(x)$, il existe un voisinage U de x tel que $f(U) \subset V$.

Démonstration. Exercice. $\square$

Définition 1.1.4.4 (CONTINUITÉ EN UN POINT). On dit que f est *continue au point* $x \in X$ si

$$(\forall V \in \mathcal{V}_{f(x)}) \quad f^{-1}(V) \in \mathcal{V}_x.$$

Lemme 1.1.4.5 (CRITÈRE SÉQUENTIEL). Soit X et Y deux espaces topologiques, $x \in X$ un point et $f: X \rightarrow Y$ une application. Si f est continue en x , alors pour toute suite $(x_n)_{n \in \mathbb{N}}$ convergeant vers x , la suite $(f(x_n))_{n \in \mathbb{N}}$ converge vers $f(x)$. La réciproque est vraie s'il existe une base dénombrable de voisinages en x .

Démonstration. • Supposons que f est continue en x . Soit $(x_n)_{n \in \mathbb{N}} \in X$ tel que x_n converge vers x . Soit $W \in \mathcal{V}_{f(x)}$. Par hypothèse, $f^{-1}(W) \in \mathcal{V}_x$ et comme x est une limite de $(x_n)_{n \in \mathbb{N}}$, il existe $N \in \mathbb{N}$ tel que pour tout $n \geq N$, on a $x_n \in f^{-1}(W)$ i.e. $f(x_n) \in W$.

• Réciproquement, supposons que pour toute suite $(x_n)_{n \in \mathbb{N}}$ convergeant vers x , la suite $(f(x_n))_{n \in \mathbb{N}}$ converge vers $f(x)$. Par l'absurde, supposons que f n'est pas continue en x : il existe $V \in \mathcal{V}_x$ tel que $f^{-1}(V) \notin \mathcal{V}_x$. Comme il existe une base dénombrable $(W_k)_{k \in \mathbb{N}}$ de voisinage, cela revient à dire que pour tout $k \in \mathbb{N}$, $W_k \not\subset f^{-1}(V)$ et donc qu'il existe $x_k \in W_k$ tel que $x_k \in f^{-1}(V)$ i.e. $f(x_k) \notin V$. Ainsi, $x_k \in W_k$ mais $f(x_k) \notin V$. On a donc montré que x_k converge vers x mais que $f(x_k)$ ne converge pas vers $f(x)$. $\square$

Théorème 1.1.4.6 (RÈGLES DE CONSTRUCTION D'APPLICATIONS CONTINUES). Soit X, Y et Z des espaces topologiques.

- (i) (*fonction constante*). Si l'application $f: X \rightarrow Y; x \mapsto y_0$ avec y_0 fixé est constante, alors f est continue.
- (ii) (*inclusion*). Si $A \subset X$, alors l'inclusion $j: A \rightarrow X; x \mapsto x$ est continue.
- (iii) (*composition*). Si $f: X \rightarrow Y$ et $g: Y \rightarrow Z$ sont continues, alors $g \circ f: X \rightarrow Z$ est continue.
- (iv) (*restriction*). Si $A \subset X$ et si $f: X \rightarrow Y$ est continue, alors la restriction $f|_A: A \rightarrow Y$ est continue.
- (v) (*corestriction*). Si $f: X \rightarrow Y$ est continue et $f(X) \subset B \subset Y$, alors la corestriction $X \rightarrow B; x \mapsto f(x)$ est continue.
- (vi) (*formulation locale*). Si $(U_\alpha)_\alpha$ est une famille d'ouverts recouvrant X , alors $f: X \rightarrow Y$ est continue si et seulement si les restrictions $f|_{U_\alpha}$ sont continues.
- (vii) (*recollement de fermés*). Si $A, B \subset X$ sont deux fermés tel que $X = A \cup B$, alors $f: X \rightarrow Y$ est continue si et seulement si $f|_A$ et $f|_B$ sont continues.

Démonstration. Les points (i) à (v) sont évidents.

(vi) Il suffit de voir que pour un ouvert $V \subset Y$, on a $f|_{U_\alpha}^{-1}(V) = f^{-1}(V) \cap U_\alpha$ qui est ouvert dans X car ouvert dans U_α par hypothèse et U_α est ouvert dans X . Ensuite, puisque $X = \bigcup_\alpha U_\alpha$, alors

$$f^{-1}(V) = f^{-1}(V) \cap \bigcup_\alpha U_\alpha = \bigcup_\alpha (f^{-1}(V) \cap U_\alpha)$$

est une union d'ouverts de X donc ouverte.

(vii) Soit F un fermé de Y . On a $f^{-1}(F) = (f^{-1}(F) \cap A) \cup (f^{-1}(F) \cap B) = (f|_A)^{-1}(F) \cup (f|_B)^{-1}(F)$ qui est l'union de deux fermés de X , donc c'est un fermé. $\square$

Exercice 1.1.4.7. La propriété (vii) se généralise à un nombre fini de fermés, mais pas une infinité. Trouver un contre-exemple.

Définition 1.1.4.8 (HOMÉOMORPHISME). Soit X et Y deux espaces topologiques. On dit qu'une application $f: X \rightarrow Y$ est *ouverte* si l'image par f de tout ouvert est un ouvert et qu'elle est *fermée* si l'image par f de tout fermé est fermé. On dit que f est un *homéomorphisme* si f est bijective, continue et ouverte.

Remarque 1.1.4.9. Si f est bijective alors : « f est ouverte $\Leftrightarrow f$ est fermée $\Leftrightarrow f^{-1}$ est continue ». Une bijection f est un homéomorphisme si f et f^{-1} sont continues.

Définition 1.1.4.10 (PLONGEMENT). Soit $f: X \rightarrow Y$ une application entre espaces topologiques. On dit que f est un *plongement* si sa corestriction $f: X \rightarrow f(X)$ est un homéomorphisme, où $f(X)$ est munie de la topologie induite.

Exercice 1.1.4.11. (TD). Soit $f:]-3, \pi[\rightarrow \mathbf{R}^2$ définie par $f(t) = (t+1, 0)$ pour $-3 < t \leq 0$ et $f(t) = (\cos(t), \sin(t))$ pour $0 \leq t < \pi$. Montrer que f est continue injective, que sa restriction à $] -3, \pi - \varepsilon[$ est un plongement pour tout $\varepsilon > 0$ mais que f n'est pas un plongement. Qu'en est-il de la restriction de f à $] -2, \pi[$?

Solution : • On va utiliser le point (vii) du théorème 1.1.4.6. On remarque que $f|_{]-3, 0]}$ et $f|_{[0, \pi]}$ sont continue et $] -3, 0] \cup [0, \pi[$ est un recouvrement fini de fermé, donc f est continue. L'application f est injective puisque ses restrictions sont injectives et $f(]-3, 0]) \cap f([0, \pi]) = \emptyset$.

• Notons $\bar{f}$ la restriction de f à $] -3, \pi - \varepsilon[$. L'application est évidemment continue injective et $[-3, \pi - \varepsilon]$ est compact. On conclut que $\bar{f}$ est un plongement et $\bar{f}|_{]-3, \pi - \varepsilon[}$ est un plongement. En particulier, $\bar{f}|_{]-3, \pi - \varepsilon[} = f|_{]-3, \pi - \varepsilon[}$ est aussi un plongement.

• Montrons que f n'est pas un plongement. Pour tout voisinage V de -2 , $f(V)$ est un voisinage de $(-1, 0)$ dans $f(]-3, \pi])$. Si on prend $V =] -\frac{5}{2}, -\frac{3}{2}[$, alors $f(V) =] -\frac{3}{2}, -\frac{1}{2}[\times \{0\}$, donc $f(V)$ n'est pas un voisinage de $(-1, 0)$ dans $f(]-3, \pi])$.

Pour conclure, on a un petit résultat bien utile :

Lemme 1.1.4.12 (INJECTION CONTINUE DANS UN SÉPARÉ). Si $f: X \rightarrow Y$ est continue injective et Y est séparé, alors X est séparé.

Démonstration. Soit $x \neq y \in X$, alors $f(x) \neq f(y) \in Y$ (car f injective) peuvent être séparés par $U, V \subset Y$ ouverts disjoints. Par continuité de f , on a $x \in f^{-1}(U)$ et $y \in f^{-1}(V)$. Les ouverts $f^{-1}(U)$ et $f^{-1}(V)$ (f est continue) sont disjoints et séparent x et y . $\square$

1.2 Opérations sur les espaces topologiques

1.2.1 Produit

Soit $(X_i)_{i \in I}$ une famille d'ensembles (I est un ensemble quelconque). On définit l'ensemble produit $X = \prod_{i \in I} X_i$ comme l'ensemble des applications $x: I \rightarrow \bigcup_{i \in I} X_i$ telles que $x(i) \in X_i$. Pour $x \in X$ on écrit $x = (x_i)$ où $x_i = x(i)$ ¹. On note $p_i: X \rightarrow X_i$ tel que $x \mapsto x_i$ les projections naturelles.

Définition 1.2.1.1 (TOPOLOGIE PRODUIT). Soit $(X_i)_{i \in I}$ une famille d'espaces topologiques, on appelle *topologie produit* sur $\prod_{i \in I} X_i$ la topologie engendrée par les $p_i^{-1}(U_i)$, avec $i \in I$ et U_i ouvert de X_i .

On appelle *cylindre ouvert* $p_i^{-1}(U_i) = \{x \in X; x_i \in U_i\} = U_i \times \prod_{i \neq j} X_j$ où $U_i \subset X_i$ est ouvert. En dimension 2, on a $p_1^{-1}(U_1) = U_1 \times X_2$ et $p_2^{-1}(U_2) = X_1 \times U_2$ (cf. dessin ci-dessous).

On appelle *pavé ouvert* une intersection finie de cylindres ouverts, donc égale à $\bigcap_{j \in J} p_j^{-1}(U_j) = \prod_{j \in J} U_j \times \prod_{i \in I \setminus J} X_i$ où $J \subset I$ est fini.

Proposition 1.2.1.2 (DESCRIPTION DE LA TOPOLOGIE PRODUIT). Soit $\mathcal{T}$ la topologie produit sur un produit $X = \prod_{i \in I} X_i$ d'espaces topologiques.

- (i) $\mathcal{T}$ est la topologie la moins fine (la plus petite) telle que les p_i soient continues.
- (ii) L'ensemble des pavés ouverts est une base de $\mathcal{T}$.

Démonstration. (i) Les p_i sont continues par définition de $\mathcal{T}$ et de la continuité. Toute topologie rendant continues les p_i contient les $p_i^{-1}(U_i)$ donc contient $\mathcal{T}$ par définition de la topologie engendrée.

(ii) Découle de la proposition 1.1.3.4 : si on pose

$$\mathcal{A} = \{p_i^{-1}(U_i); i \in I, U_i \subset X_i \text{ ouvert}\}$$

et on a vu que la topologie $\mathcal{T}_{\mathcal{A}}$ engendrée par $\mathcal{A}$ a pour base les intersections finis d'éléments de $\mathcal{A}$ i.e. les pavés ouverts. Ainsi, les ouverts sont les unions de pavés ouverts. $\square$

Remarque 1.2.1.3. Attention ! Un produit $\prod_{i \in I} U_i \subset X$ d'ouverts U_i n'est ouvert que si c'est un pavé, i.e. si les indices i où $U_i \neq X_i$ sont en nombre fini (si ce nombre est infini le produit ne contient pas de pavé donc n'est pas ouvert).

Lorsque $|I| = n$ est fini, les pavés ouverts sont les $U = U_1 \times \cdots \times U_n$ où U_i est ouvert dans X_i .

Proposition 1.2.1.4 (PROPRIÉTÉS FONDAMENTALES). Soit $(X_i)_{i \in I}$ une famille d'espaces topologiques et soit $X = \prod_{i \in I} X_i$ leur produit.

- (i) Les projections $p_i: X \rightarrow X_i$ sont ouvertes.
- (ii) Si $f = (f_i)_{i \in I}: Y \rightarrow X$ est une application, alors f est continue si et seulement si les composantes $f_i = p_i \circ f: Y \rightarrow X_i$ sont continues (cette propriété caractérise la topologie produit).
- (iii) Soit $j \in I$ et soit $a \in \prod_{i \neq j} X_i$. L'application $i_a: X_j \rightarrow X; x_j \mapsto x$ où $x(j) = x_j$ et $x(i) = a_i$ pour $i \neq j$, est un plongement.
- (iv) Le produit $X = \prod_{i \in I} X_i$ est séparé si et seulement si les X_i sont séparés.

Démonstration. (i) Soit $U \subset X$ un ouvert et soit $x_i = p_i(x) \in p_i(U)$. Soit $V = \prod_{i \in I} V_i$ un pavé ouvert tel que $x \in V \subset U$. On a alors $p_i(U) \supset p_i(V) = V_i \ni x_i$ où V_i est ouvert, donc $p_i(U)$ est ouvert.

(ii) Si f est continue, ses composantes $f_i = p_i \circ f$ sont continues par composition. Réciproquement, supposons

1. Une application x est une fonction qui choisit un élément $x_i \in X_i$ pour chaque i . L'existence d'une telle fonction si et seulement si les X_i sont non vides est exactement l'axiome du choix. L'axiome du choix est donc équivalent au fait qu'un produit d'ensembles non vides est un ensemble non vide.

les f_i continues et montrons que f est continue. Puisque les $p_i^{-1}(U_i)$ engendrent la topologie produit, il suffit de montrer que les $f^{-1}(p_i^{-1}(U_i))$ sont ouverts lorsque U_i est ouvert. Ainsi, $f^{-1}(p_i^{-1}(U_i)) = (p_i \circ f)^{-1}(U_i) = f_i^{-1}(U_i)$ est ouvert puisque f_i est continue.

(iii) L'application est clairement injective donc bijective sur son image $X_i \times \{a\} \subset X$. Les composantes $p_i \circ i_a$ de i_a sont soit constantes $X_j \rightarrow X_i$ si $i \neq j$, soit l'identité $X_j \rightarrow X_j$, donc toutes continues. Ainsi, i_a est continue par (ii). Soit $U_j \subset X_j$ un ouvert, alors le pavé $U_j \times \prod_{i \neq j} X_i$ est ouvert dans X et $i_a(U_j) = X_i \times \{a\} \cap U_j \times \prod_{i \neq j} X_i$ est ouvert dans $X_i \times \{a\}$, prouvant que i_a est ouverte vers son image.

(iv) Supposons X séparé. D'après (iii) chaque X_i s'injecte continûment dans X séparé, donc est séparé par le lemme 1.1.4.12. Réciproquement, supposons les X_i séparés et soit $x \neq y \in X$. Il existe alors $i \in I$ tel que $x_i \neq y_i$. Puisque X_i est séparé, il existe deux ouverts disjoints U et V de X_i tel que $x_i \in U$ et $y_i \in V$, donc $p_i^{-1}(U)$ et $p_i^{-1}(V)$ sont deux ouverts disjoints de X contenant x et y respectivement. $\square$

Remarque 1.2.1.5. Si on avait défini la topologie produit sur X pour que tous les produits d'ouverts $U = \prod_{i \in I} U_i$ soient ouverts, alors

$$f^{-1}(U) = \{x \in X; f_i(x) \in U_i \ (\forall i \in I)\} = \bigcap_{i \in I} f_i^{-1}(U_i)$$

aurait peu de chances d'être ouvert lorsque I est infini, même si les $f_i^{-1}(U_i)$ sont ouverts. La propriété (ii) serait fausse.

Exemple 1.2.1.6. (1) Si X est un espace topologique, on appelle *cylindre* sur X le produit cartésien $X \times [0, 1]$.

(2) On appelle *tore de dimension n* le produit $\mathbf{T}^n = \mathbf{S}^1 \times \dots \times \mathbf{S}^1$ de n copies du cercle unité $\mathbf{S}^1 \subset \mathbf{R}^2$.

(3) Soit Y un espace topologique, I un ensemble d'indices et posons $X_i = Y$ pour tout $i \in I$. Un élément $x \in \prod_{i \in I} X_i$ est exactement une application $I \rightarrow Y$ et on note également $\prod_{i \in I} X_i = Y^I$.

Exercice 1.2.1.7. On peut montrer que la topologie produit sur Y^I est la topologie de la convergence simple : une suite $(f_n)_{n \in \mathbf{N}}$ de Y^I converge vers $f \in Y^I$ si et seulement si les fonctions $f_n : I \rightarrow Y$ convergent simplement vers $f : I \rightarrow Y$. Si I n'est pas dénombrable, cette topologie n'est pas métrisable (bien que séparée).

1.2.2 Union disjointe

Définition 1.2.2.1 (UNION DISJOINTE). Soit $(X_i)_{i \in I}$ une famille d'ensembles. Leur *union disjointe* est l'ensemble

$$\coprod_{i \in I} X_i = \{(x, i); x \in X_i, i \in I\}.$$

Exemple 1.2.2.2. Pour $X_1 = X_2 = [0, 1]$, on a

$$[0, 1] \coprod [0, 1] = \{(x, i); x \in [0, 1], i = 1, 2\} = [0, 1] \times \{1\} \cup [0, 1] \times \{2\}.$$

Chaque X_i s'injecte dans $\coprod_{i \in I} X_i$ par l'application

$$\begin{aligned} \phi_i : X_i &\rightarrow \coprod_{i \in I} X_i \\ x_i &\mapsto (x_i, i) \end{aligned}$$

et $\phi_i(X_i) \cap \phi_j(X_j) = \emptyset$ si $i \neq j$. Ainsi, $X_1 \coprod X_2 = X_1 \times \{1\} \cup X_2 \times \{2\}$ et quand $X_1 = X_2 = X$, pour tout $x \in X$, on a $(x, 1) \neq (x, 2)$, donc $X \coprod X$ est constitué de deux copies disjointes de X .

Définition 1.2.2.3 (Topologie de l'union disjointe). Soit $(X_i)_{i \in I}$ une famille d'espaces topologiques. La *topologie de l'union disjointe* sur $\coprod_{i \in I} X_i$ est

$$\mathcal{T}_u = \left\{ \coprod_{i \in I} U_i; U_i \text{ ouvert de } X_i \right\}.$$

On vérifie que c'est une topologie grâce aux égalités $\coprod_{i \in I} U_i \cup \coprod_{i \in I} V_i = \coprod_{i \in I} (U_i \cup V_i)$ et $\coprod_{i \in I} U_i \cap \coprod_{i \in I} V_i = \coprod_{i \in I} (U_i \cap V_i)$. C'est la topologie la plus fine pour laquelle les ϕ_i sont continues. En effet, soit $U = \coprod_{i \in I} U_i \in \mathcal{T}_u$, alors $\phi_i^{-1}(U) = U_i$ est ouvert, donc ϕ_i est continue. Réciproquement, soit $\mathcal{T}$ une topologie sur $\coprod_{i \in I} X_i$ rendant les ϕ_i continues. Soit $U \in \mathcal{T}$, alors $U_i := \phi_i^{-1}(U)$ est ouvert dans X_i pour tout $i \in I$, donc $U = \coprod_{i \in I} U_i \in \mathcal{T}_u$ puis $\mathcal{T} \subset \mathcal{T}_u$.

Exercice 1.2.2.4. Les ϕ_i sont des plongements.

1.2.3 Topologie quotient

Une *relation d'équivalence* sur un ensemble X est une partie $\mathcal{R} \subset X \times X$ telle que :

- (1) $(\forall x \in X) (x, x) \in \mathcal{R}$;
- (2) $(\forall (x, y) \in X \times X) (x, y) \in \mathcal{R} \Leftrightarrow (y, x) \in \mathcal{R}$;
- (3) $(\forall (x, y, z) \in X^3) ((x, y) \in \mathcal{R}, (y, z) \in \mathcal{R}) \Rightarrow (x, z) \in \mathcal{R}$.

On note $x \mathcal{R} y$ ou $x \sim y$ si $(x, y) \in \mathcal{R}$. La *classe* de x est $\{y \in X; x \mathcal{R} y\}$. L'ensemble des classes, noté $X/\mathcal{R}$, forme une partition de X et on note $\pi: X \rightarrow X/\mathcal{R}$ la *projection canonique*. On pourra noter $\pi(x) = \bar{x} = [x]$ la classe de x .

Exercice 1.2.3.1. Soit X un ensemble et $(\mathcal{R}_i)_{i \in I}$ une famille de relations d'équivalence sur X . L'ensemble $\bigcap_{i \in I} \mathcal{R}_i$ est alors une relation d'équivalence sur X .

La plus petite (pour l'inclusion) relation d'équivalence est l'égalité, définie par la diagonale $\mathcal{R} = \{(x, x); x \in X\}$, clairement contenue dans toute relation d'équivalence. À l'inverse, $\mathcal{R} = X \times X$ les contient toutes. Étant donné une partie $\mathcal{A} \subset X \times X$, on appelle *relation d'équivalence engendrée* par $\mathcal{A}$ l'intersection des relations d'équivalence contenant $\mathcal{A}$. C'est la plus petite contenant $\mathcal{A}$.

Proposition-Définition 1.2.3.2 (TOPOLOGIE QUOTIENT). Soit $(X, \mathcal{T})$ un espace topologique, $\mathcal{R}$ une relation d'équivalence sur X et $\pi: X \rightarrow X/\mathcal{R}$ la projection canonique. L'ensemble

$$\mathcal{T}_{\mathcal{R}} = \{U \subset X/\mathcal{R}; \pi^{-1}(U) \in \mathcal{T}\}$$

est alors une topologie sur $X/\mathcal{R}$, appelée *topologie quotient*, la plus fine rendant π continue.

Démonstration. C'est une topologie car l'union et l'intersection se comportent bien par image réciproque : étant donné $U_i \subset X/\mathcal{R}$ des ouverts pour $i \in I$ i.e. tels que $\pi^{-1}(U_i) \in \mathcal{T}$ et $J \subset I$ fini, on a

$$\pi^{-1}\left(\bigcup_{i \in I} U_i\right) = \bigcup_{i \in I} \pi^{-1}(U_i) \in \mathcal{T} \quad \text{et} \quad \pi^{-1}\left(\bigcap_{j \in J} U_j\right) = \bigcap_{j \in J} \pi^{-1}(U_j) \in \mathcal{T}$$

donc $\bigcup_{i \in I} U_i \in \mathcal{T}_{\mathcal{R}}$ et $\bigcap_{j \in J} U_j \in \mathcal{T}_{\mathcal{R}}$. De plus, $\pi^{-1}(\emptyset) = \emptyset$ et $\pi^{-1}(X/\mathcal{R}) = X$ par surjectivité de π . Par définition, π est continue pour cette topologie. Soit $\mathcal{S}$ une topologie sur $X/\mathcal{R}$ rendant π continue. Si $U \in \mathcal{S}$, alors $\pi^{-1}(U)$ est ouvert, donc $U \in \mathcal{T}_{\mathcal{R}}$. Ainsi, $\mathcal{S} \subset \mathcal{T}_{\mathcal{R}}$. $\square$

En pratique, le quotient $X/\mathcal{R}$ d'un espace topologique est toujours muni de la topologie quotient.

Remarque 1.2.3.3. Attention ! L'image par π d'un ouvert de X n'est pas nécessairement un ouvert de $X/\mathcal{R}$ i.e. π n'est pas nécessairement ouverte (car l'intersection se comporte mal par image directe).

Exemple 1.2.3.4. Soit $A =]0, 1[\cup [1, 2] =]0, 2] \subset \mathbf{R}$ muni de $\mathcal{R}$ engendrée par $\mathcal{A} = A \times A$ i.e. vérifiant $x \mathcal{R} y$ pour tout $x, y \in A$ ou si $x = y \in \mathbf{R}$. L'ensemble $\pi(]0, 1[) = \pi(A) \subset \mathbf{R}/\mathcal{R}$ a pour image réciproque A qui n'est pas ouvert, donc $\pi(]0, 1[)$ n'est pas ouvert.

Définition 1.2.3.5 (SATURÉ). Soient X un ensemble et soit $\mathcal{R}$ une relation d'équivalence sur X . Le *saturé* par $\mathcal{R}$ d'une partie et $A \subset X$ est l'ensemble $\pi^{-1}(\pi(A)) \subset X$. On dit que A est saturé s'il est égal à son saturé.

Dans l'exemple précédent, $]0, 1[$ n'est pas saturé, son saturé étant $]0, 1[\cup [1, 2]$. Il suit immédiatement de la définition 1.2.3.5 que l'image directe d'un ouvert saturé est un ouvert : si $U \subset X$ est ouvert, alors $\pi^{-1}(\pi(U)) = U$ est ouvert, donc $\pi(U) \subset X/\mathcal{R}$ est ouvert.

Définition 1.2.3.6 (RELATION OUVERTE). On dit qu'une relation d'équivalence $\mathcal{R}$ sur un espace topologique est *ouverte* si la projection canonique $\pi: X \rightarrow X/\mathcal{R}$ est ouverte (ou de façon équivalente, si le saturé par $\mathcal{R}$ de tout ouvert de X est ouvert).

Lemme 1.2.3.7 (ACTION PAR HOMÉOMORPHISMES). Soit X un espace topologique et Γ un sous-groupe de $\text{Homeo}(X)$. La relation associée à l'action de Γ sur X ($x \mathcal{R} y$ s'il existe $g \in \Gamma$ tel que $y = g \cdot x$) est alors ouverte.

Démonstration. Soit $U \subset X$ un ouvert. On a alors

$$\begin{aligned} \pi^{-1}(\pi(U)) &= \{y \in X; \pi(y) \in \pi(U)\} = \{y \in X; (\exists x \in U) \pi(y) = \pi(x)\} \\ &= \{y \in X; (\exists g \in \Gamma) (\exists x \in U) y = g \cdot x\} = \bigcup_{g \in \Gamma} g(U) \end{aligned}$$

qui est une réunion d'ouverts, donc c'est ouvert dans X , prouvant que $\pi(U)$ est ouvert dans $X/\mathcal{R}$. $\square$

Exemple 1.2.3.8. L'ensemble $\mathbf{Z}$ agissant sur $\mathbf{R}$ par translation $k \cdot x = x + k$ induit sur $\mathcal{R}$ la relation $x \mathcal{R} y \Leftrightarrow x - y \in \mathbf{Z}$, qui est donc ouverte.

1.2.4 Séparation

Remarque 1.2.4.1. Le quotient $X/\mathcal{R}$ d'un espace topologique séparé X n'est pas nécessairement séparé (ce qui est prévisible vu qu'on identifie des points...).

Exemple 1.2.4.2. (1) Soit $X = \mathbf{R}$ muni de $\mathcal{R}$ définie par $x \sim y \Leftrightarrow x - y \in \mathbf{Q}$. L'ensemble $\mathbf{R}/\mathbf{Q}$ a alors la topologie grossière et plus de deux éléments, donc n'est pas séparé. En effet, soit $U \subset \mathbf{R}/\mathbf{Q}$ un voisinage ouvert de $[x]$ dans $\mathbf{R}/\mathbf{Q}$. On voit alors que $\pi^{-1}(U)$ est un voisinage ouvert de $x + \mathbf{Q}$ dans $\mathbf{R}$. Il existe $\varepsilon > 0$ tel que $]x - \varepsilon, x + \varepsilon[\subset \pi^{-1}(U)$. On a alors $\pi^{-1}(U) \supset]x - \varepsilon, x + \varepsilon[= \mathbf{R}$, donc $U = \pi(\pi^{-1}(U)) = \mathbf{R}/\mathbf{Q}$. De plus, $\mathbf{R}/\mathbf{Q}$ a au moins deux éléments car $[\sqrt{2}] = \sqrt{2} + \mathbf{Q} \neq \mathbf{Q} = [0]$.

(2) Soit $\mathcal{R}$ engendrée par $A \times A$ avec $A =]0, 1[\cup]2, 3]$. On a $[\frac{1}{2}] =]0, 1[\cup]2, 3]$. L'ensemble $X/\mathcal{R}$ n'est pas séparé car $[0]$ et $[\frac{1}{2}]$ ne peuvent être séparés par des ouverts de $X/\mathcal{R}$. En effet, soit $U, V \subset X/\mathcal{R}$ ouverts tel que $[0] \in U$ et $[\frac{1}{2}] \in V$, alors $\pi^{-1}(U)$ est un ouvert de $\mathbf{R}$ contenant 0 et $\pi^{-1}(V)$ est un ouvert de $\mathbf{R}$ contenant $\pi^{-1}([\frac{1}{2}]) =]0, 1[\cup]2, 3]$. On a donc $\pi^{-1}(U) \cap \pi^{-1}(V) \neq \emptyset$, d'où $U \cap V \neq \emptyset$.

Proposition 1.2.4.3 (CRITÈRE DE SÉPARATION). Soit X un espace topologique et soit $\mathcal{R}$ une relation d'équivalence sur X .

(i) L'espace $X/\mathcal{R}$ est séparé si et seulement si tout $x \not\sim y \in X$ peuvent être séparés par des ouverts saturés.

(ii) Si $\mathcal{R}$ est ouverte, alors on a l'équivalence : « $X/\mathcal{R}$ est séparé $\Leftrightarrow \mathcal{R} \subset X \times X$ est fermé ».

Démonstration. (i) • Supposons $X/\mathcal{R}$ séparé. Soit $x, y \in X$ tel que $\pi(x) \neq \pi(y)$. Soit $U, V \subset X/\mathcal{R}$ des ouverts séparant $[x]$ de $[y]$, alors $\pi^{-1}(U)$ et $\pi^{-1}(V)$ sont des ouverts saturés disjoints séparant x et y .

• Réciproquement, supposons la condition vérifiée. Soit $[x] \neq [y] \in X/\mathcal{R}$, soit U et V des ouverts saturés disjoints contenant x et y respectivement. Les ensembles $\pi(U) \ni [x]$ et $\pi(V) \ni [y]$ sont alors ouverts, disjoints puisqu'aucun élément de U n'est équivalent à un élément de V .

(ii) • Soit $(x, y) \in X \times X \setminus \mathcal{R}$. Par hypothèse, on a $x \not\sim y$, donc (i) donne des ouverts saturés disjoints U et V séparant x et y . On a alors $(x, y) \in U \times V \subset X \times X \setminus \mathcal{R}$, prouvant que $X \times X \setminus \mathcal{R}$ est ouvert.

• Réciproquement, si $[x] \neq [y] \in X/\mathcal{R}$, alors $(x, y) \in X \times X \setminus \mathcal{R}$ qui par hypothèse est ouvert. Il existe donc un cube ouvert $U \times V$ dans $X \times X$ tel que $(x, y) \in U \times V \subset X \times X \setminus \mathcal{R}$. Par conséquent, U et V sont des ouverts (pas nécessairement saturés) de X séparant x et y tels que $\pi(U) \cap \pi(V) = \emptyset$. Comme π est ouverte, ce sont des ouverts séparant $[x]$ et $[y]$. $\square$

1.2.5 Factorisation d'applications

Proposition 1.2.5.1 (PROPRIÉTÉ UNIVERSELLE). Si X et Y sont deux espaces topologiques, $\mathcal{R}$ une relation d'équivalence sur X et $\phi: X/\mathcal{R} \rightarrow Y$ une application, alors ϕ est continue si et seulement si $\phi \circ \pi$ est continue.

$$\begin{array}{ccc} X & \xrightarrow{\phi \circ \pi} & Y \\ \pi \searrow & & \nearrow \phi \\ & X/\mathcal{R} & \end{array}$$

Démonstration. Si ϕ est continue, alors $\phi \circ \pi$ est continue par composition. Supposons $\phi \circ \pi$ continue et soit $U \subset Y$ un ouvert. L'ensemble $\phi^{-1}(U)$ est alors ouvert puisque $\pi^{-1}(\phi^{-1}(U)) = (\phi \circ \pi)^{-1}(U)$ est ouvert dans X , prouvant que $\phi^{-1}(U)$ est ouvert dans $X/\mathcal{R}$. Ainsi, ϕ est continue. $\square$

Remarque 1.2.5.2. On peut montrer que cette propriété caractérise la topologie quotient (cf. exercice de TD sur la topologie finale).

Soit $f: X \rightarrow Y$ une application entre deux ensembles et soit $\mathcal{R}$ une relation d'équivalence sur X . Supposons que $(\forall x, y \in X) x \mathcal{R} y \Rightarrow f(x) = f(y)$. L'application $\bar{f}: X/\mathcal{R} \rightarrow Y$ définie par $\bar{f}(\pi(x)) = f(x)$ pour tout $x \in X$ est alors bien définie puisque si $\pi(x) = \pi(y)$, alors on a $f(x) = f(y)$.

$$\begin{array}{ccc} X & \xrightarrow{f} & Y \\ \pi \searrow & & \nearrow \bar{f} \\ & X/\mathcal{R} & \end{array}$$

On dit que f se « factorise » au quotient. Notons $\mathcal{R}_f$ la relation d'équivalence sur X définie par $x \mathcal{R}_f y \Leftrightarrow f(x) = f(y)$, alors $x \mathcal{R} y \Rightarrow f(x) = f(y)$ équivaut à $\mathcal{R} \subset \mathcal{R}_f$ en tant que partie de $X \times X$.

Théorème 1.2.5.3 (FACTORISATION CANONIQUE D'UNE APPLICATION CONTINUE). Soient X, Y deux espaces topologiques et $\mathcal{R}$ une relation d'équivalence sur X . Soit $f: X \rightarrow Y$ continue telle que $\mathcal{R} \subset \mathcal{R}_f$. On a alors :

(i) f se factorise en une application continue $\bar{f}$ telle que $\bar{f} \circ \pi = f$.

(ii) $\bar{f}$ est injective si et seulement si $\mathcal{R} = \mathcal{R}_f$.

(iii) Si f est ouverte, alors $\bar{f}$ est ouverte et si f est fermée, alors $\bar{f}$ est fermée.

Démonstration. (i) Évident par la propriété universelle 1.2.5.1.

(ii) On a déjà $\mathcal{R} \subset \mathcal{R}_f$ et

$$\begin{aligned}\bar{f} \text{ injective} &\Leftrightarrow (\forall x, y \in X) \bar{f}(\pi(x)) = \bar{f}(\pi(y)) \Rightarrow \pi(x) = \pi(y) \\ &\Leftrightarrow (\forall x, y \in X) f(x) = f(y) \Rightarrow \pi(x) = \pi(y) \\ &\Leftrightarrow (\forall x, y \in X) x \mathcal{R}_f y \Rightarrow x \mathcal{R} y \\ &\Leftrightarrow \mathcal{R}_f \subset \mathcal{R}.\end{aligned}$$

(iii) Soit U un ouvert de $X/\mathcal{R}$, alors $\pi^{-1}(U)$ est ouvert de X et $\bar{f}(U) = f(\pi^{-1}(U))$ est ouvert si f est ouverte. Même argument si f est fermée. $\square$

Corollaire 1.2.5.4 (FACTORISATION). Soit X, Y deux espaces topologiques et $f: X \rightarrow Y$ une application continue.

(i) f se factorise en une bijection continue $\bar{f}: X/\mathcal{R}_f \rightarrow f(X) \subset Y$.

(ii) Si Y est séparé, alors $X/\mathcal{R}_f$ est séparé.

(iii) Si f est ouverte ou fermée, alors $\bar{f}$ est un plongement.

On obtient le diagramme commutatif suivant :

$$\begin{array}{ccc} X & \xrightarrow{f} & f(X) \subset Y \\ \pi \searrow & & \nearrow \bar{f} \\ & X/\mathcal{R}_f & \end{array}$$

Démonstration. (i) et (iii) découlent du théorème 1.2.5.3 et (ii) du lemme 1.1.4.12 puisque $X/\mathcal{R}_f$ s'injecte continûment dans Y qui est séparé. $\square$

Exemple 1.2.5.5. Sur $X = [0, 1]$, soit $\sim$ la relation d'équivalence engendrée par $0 \sim 1$. On montre que $[0, 1]/\sim$ est homéomorphe à $\mathbf{S}^1$ en factorisant l'application $t \in [0, 1] \mapsto f(t) = (\cos(2\pi t), \sin(2\pi t)) \in \mathbf{R}^2$.

$$\begin{array}{ccc} [0, 1] & \xrightarrow{f} & \mathbf{S}^1 \subset \mathbf{R}^2 \\ \pi \searrow & & \nearrow \bar{f} \\ & [0, 1]/\sim & \end{array}$$

En effet, f est continue, d'image $\mathbf{S}^1$ et vérifie $f(x) = f(y) \Leftrightarrow x \sim y$. Elle se factorise donc en une bijection continue $\bar{f}: [0, 1]/\sim \rightarrow f([0, 1]) = \mathbf{S}^1$. De plus, f est fermée : si $U \subset [0, 1]$ est fermé, alors U est compact dans $[0, 1]$ aussi compact, d'où $f(U)$ est compact donc fermé dans $\mathbf{R}^2$ et $\mathbf{S}^1$. D'après le corollaire 1.2.5.4, $\bar{f}$ est un plongement *i.e.* un homéomorphisme sur son image $\mathbf{S}^1$.

Exercice 1.2.5.6. Sur $[0, 1] \times [0, 1]$, soit $\sim$ engendrée par $(s, 0) \sim (s, 1)$ et $(0, t) \sim (1, t)$ pour tous $s, t \in [0, 1]$. Montrer que $([0, 1] \times [0, 1])/\sim$ est homéomorphe au tore $\mathbf{S}^1 \times \mathbf{S}^1$.

Exercice 1.2.5.7 (en TD). Démontrer que $\mathbf{R}/\mathbf{Z}$ est homéomorphe à $\mathbf{S}^1$.

Exercice 1.2.5.8 (en TD). Soit $\mathbf{D}^2 = \{z \in \mathbf{C}; |z| \leq 1\}$ et $\sim$ la relation d'équivalence sur $\mathbf{D}^2$ engendrée par $x \sim y$ pour tous $x, y \in \mathbf{S}^1 = \partial \mathbf{D}^2$. Montrer que $\mathbf{D}^2/\sim$ est homéomorphe à $\mathbf{S}^2$.

1.2.6 Exemples : cône, suspension, recollement, somme pointée

La topologie quotient permet de nombreuses constructions d'espaces topologiques.

Définition 1.2.6.1 (CÔNE). Soit X un espace topologique. On appelle *cône* sur X le quotient

$$C(X) = (X \times [0, 1])/\sim$$

où $\sim$ est engendrée par $(x, 1) \sim (x', 1)$ pour tous $x, x' \in X$. L'image de $X \times \{1\}$ est le *sommet* du cône.

Exercice 1.2.6.2. L'application $X \rightarrow C(X); x \mapsto [(x, 0)]$ est un plongement (permettant d'identifier X avec la base du cône). Si X et Y sont homéomorphes, alors $C(X)$ et $C(Y)$ sont homéomorphes.

Définition 1.2.6.3 (SUSPENSION). Soit X un espace topologique. On appelle *suspension* de X l'espace topologique quotient

$$S(X) = (X \times [-1, 1])/\sim$$

où $\sim$ est engendrée par $(x, 1) \sim (x', 1)$ et $(x, -1) \sim (x', -1)$ pour tous $x, x' \in X$.

Exercice 1.2.6.4. L'application $X \rightarrow S(X); x \mapsto [(x, 0)]$ est un plongement. Si X et Y sont homéomorphes, alors $S(X)$ et $S(Y)$ sont homéomorphes.

Exercice 1.2.6.5 (en TD). Le cône $C(\mathbf{S}^n)$ est homéomorphe à $\mathbf{B}^{n+1}$ et $S(\mathbf{S}^n)$ est homéomorphe à $\mathbf{S}^{n+1}$.

Définition 1.2.6.6 (RECOLLEMENT). Soient X, Y deux espaces topologiques, A une partie de X , B une partie de Y et $f: A \rightarrow B$ continue. Le *recollement* de X sur Y par f est l'espace topologique quotient

$$X \coprod_f Y = (X \coprod Y) / \sim$$

où $\sim$ est la relation d'équivalence engendrée par $a \sim f(a)$ pour tout $a \in A$.

Exercice 1.2.6.7. Soit $X = Y = \mathbf{D}^2 = \{z \in \mathbf{C}; |z| \leq 1\}$, $A = \mathbf{S}^1 = \partial \mathbf{D}^2$ et $f = \text{Id}_{\mathbf{S}^1}$. Montrer que $X \coprod_f Y$ est homéomorphe à $\mathbf{S}^2$.

Définition 1.2.6.8 (SOMME POINTÉE). Soit $(X_i, x_i)_{i \in I}$ une famille d'espaces topologiques pointés ($x_i \in X_i$). La *somme pointée* de cette famille est l'espace topologique quotient

$$\bigvee_{i \in I} (X_i, x_i) = \left(\coprod_{i \in I} X_i \right) / \sim$$

où $\sim$ est la relation d'équivalence engendrée par $(x_i, i) \sim (x_j, j)$ pour tous $i, j \in I$.

Il est clair que l'inclusion $X_i \rightarrow \coprod_{i \in I} X_i$ induit un plongement de X_i dans $\bigvee_{i \in I} (X_i, x_i)$.

Car particulier : si $X_1 = X_2 = \dots = X_n = \mathbf{S}^1$ et $x_i = 1$ pour tout $i \in \llbracket 1, n \rrbracket$, alors $\bigvee_{i \in \llbracket 1, n \rrbracket} (X_i, x_i)$ s'appelle un *bouquet de n cercles*.

1.3 Connexité, connexité par arc

1.3.1 Connexité

Définition 1.3.1.1 (CONNEXITÉ). Un espace topologique X est *connexe* s'il n'existe pas de partition de X en deux ouverts non vides. Une partie $A \subset X$ est connexe si elle est connexe pour la topologie induite.

Cela équivaut à dire, par passage au complémentaire, que X est connexe s'il n'existe pas de partition de X en deux fermés non vides. Également que les seules parties ouvertes et fermées de X sont $\emptyset$ et X .

Proposition 1.3.1.2 (PROPRIÉTÉS FONDAMENTALES). Soit X un espace topologique.

- (i) X est connexe si et seulement si toute application continue $X \rightarrow \{0, 1\}$ discret est constante.
- (ii) Si X est connexe et $f: X \rightarrow Y$ est continue, alors $f(X)$ est une partie connexe de Y .
- (iii) Si $A, B \subset X$ telles que $A \subset B \subset \bar{A}$ et A est connexe, alors B est connexe. En particulier, $\bar{A}$ est connexe.
- (iv) Soient $(A_i)_{i \in I}$ une famille de parties connexes de X telle que $\bigcap_{i \in I} A_i \neq \emptyset$. L'ensemble $\bigcup_{i \in I} A_i$ est alors connexe.
- (v) Un produit fini d'espaces connexes est connexe.

Démonstration. (i) • Soit $f: X \rightarrow \{0, 1\}$ continue et soit $y \in f(X)$. Puisque $\{0, 1\}$ est discret, le singleton $\{y\}$ est à la fois ouvert et fermé. Par continuité de f , l'application $f^{-1}(y)$ est ouvert et fermé dans X , donc égal à X par connexité de X . Ainsi, $f(X) = \{y\}$.

• Montrons la contraposée. Supposons X non connexe et soit $\{U, X \setminus U\}$ une partition de X en deux ouverts non vides. Définissons $f: X \rightarrow \{0, 1\}$ par $f(U) = 0$ et $f(X \setminus U) = 1$. Les restrictions de f à U et à $X \setminus U$ sont continues car constantes. Puisque U et $X \setminus U$ sont ouverts, alors f est continue.

(ii) Soit X connexe et soit $f: X \rightarrow Y$ continue. Soit $g: f(X) \rightarrow \{0, 1\}$ continue. L'application composée $g \circ f: X \rightarrow \{0, 1\}$ est alors continue, à valeurs dans $\{0, 1\}$ donc constante par (i) puisque X est connexe. Il s'ensuit que g est constante, d'où $f(X)$ est connexe par (i).

(iii) Soit $f: B \rightarrow \{0, 1\}$ continue. La restriction de f à A est continue à valeurs dans $\{0, 1\}$, donc constante puisque A est connexe. Soit $x \in B$. Puisque $f(x)$ est un ouvert de $\{0, 1\}$, par continuité, $U = f^{-1}(f(x))$ est un voisinage (ouvert) de x . Néanmoins, on a $x \in \bar{A}$, donc tout voisinage de x rencontre A . En particulier, on a $U \cap A \neq \emptyset$, d'où $f(x) = f(A)$. Il s'ensuit que f est constante sur B . On conclut avec (i).

(iv) On utilise encore (i). Soit $f: \bigcup_{i \in I} A_i \rightarrow \{0, 1\}$ continue. Pour chaque $i \in I$, la restriction de f à A_i est continue, à valeurs dans $\{0, 1\}$, donc constante puisque A_i est connexe. Fixons un indice $i_0 \in I$, alors pour tout $i \in I$, puisque $A_{i_0} \cap A_i \neq \emptyset$, on a $f(A_{i_0}) = f(A_i)$. Ainsi, f est constante et $\bigcup_{i \in I} A_i$ est connexe.

(v) Exercice. □

Le prototype d'ensemble connexe est l'intervalle :

Théorème 1.3.1.3. Les parties connexes de $\mathbf{R}$ sont les intervalles.

Démonstration. Rappelons qu'une partie $A \subset \mathbf{R}$ est un intervalle si elle vérifie la propriété :

$$(\forall a < b \in A) (\forall z \in \mathbf{R}) \quad a < z < b \Rightarrow z \in A$$

• Montrons d'abord qu'une partie connexe est un intervalle. Soit $A \subset \mathbf{R}$ connexe, soit $a, b \in A$ tel que $a < b$ et soit $a < z < b$. Si $z \notin A$ alors

$$A = (A \cap]-\infty, z[) \cup (A \cap]z, +\infty[)$$

est réunion de deux ouverts de A , disjoints, contenant a et b respectivement donc non vides. On a alors une contradiction, donc $z \in A$ et A est un intervalle.

• Montrons qu'un intervalle est connexe. Soit $A \subset \mathbf{R}$ un intervalle (non vide) et $f: A \rightarrow \{0, 1\}$ continue. Nous allons montrer que f est constante, ce qui sera conclusif. Supposons qu'il existe $a < b \in A$ tel que $f(a) \neq f(b)$; quitte à remplacer f par $1 - f$, supposons que $f(a) = 0$. Soit

$$J = \{t \in [a, b]; f(s) = 0 (\forall s \in [a, t])\}.$$

Cet ensemble est non vide car contient a , majoré par b , donc admet une borne supérieure $z \in [a, b] \subset A$. Par définition de la borne sup, il existe $(z_n)_{n \in \mathbf{N}}$ dans J convergeant vers z . Par continuité de f , on a donc $f(z) = \lim_{n \rightarrow +\infty} f(z_n) = 0$, donc $z < b$. Toujours par continuité de f , on a que $f^{-1}(\{0\})$ est ouvert dans A . Comme $z < b$, on en déduit qu'il existe $\varepsilon > 0$ tel que $f = 0$ sur $[z, z + \varepsilon] \subset [a, b]$, ce qui contredit la définition de z . En conclusion, f est constante. $\square$

Obstruction à l'homéomorphie : La connexité est une propriété invariante par homéomorphisme, ce qui donne donc une obstruction très simple à l'homéomorphie entre deux espaces topologiques. Ainsi, $\mathbf{S}^1$ n'est pas homéomorphe à un intervalle car $\mathbf{S}^1$ privé d'un point reste connexe alors que l'intervalle privé d'un point (qu'on peut choisir intérieur) ne l'est pas. De même, $\mathbf{R}^2$ n'est pas homéomorphe à $\mathbf{R}$ car $\mathbf{R}^2 \setminus \{\text{point}\}$ est connexe alors que $\mathbf{R} \setminus \{\text{point}\}$ ne l'est pas. Plus généralement $\mathbf{R}^p$ et $\mathbf{R}^q$ ne sont homéomorphes que si $p = q$ mais la démonstration demande des outils plus sophistiqués.

Soit X un espace topologique. On définit une relation $\sim$ sur X par $x \sim y$ s'il existe une partie connexe $C \subset X$ contenant x et y .

Proposition 1.3.1.4. La relation $\sim$ définie ci-dessus est une relation d'équivalence sur X . Ses classes sont des parties connexes et fermées de X , appelées *composantes connexes* de X . La classe de x est la réunion des parties connexes contenant x .

Démonstration. La relation est clairement réflexive et symétrique. Soit $x, y, z \in X$ tel que $x \sim y$ et $y \sim z$. Par définition, il existe des parties connexes C et C' de X telles que $x, y \in C$ et $y, z \in C'$. Puisque $y \in C \cap C'$, on a $C \cup C'$ est connexe donc $x \sim z$. Montrons que la classe de x est égale à C_x la réunion des parties connexes de X contenant x . Soit $y \sim x$ et $C \subset X$ contenant x et y , alors $C \subset C_x$ et en particulier, $y \in C_x$. Réciproquement, tout $y \in C_x$ est clairement équivalent à x , donc la classe de x est égale à C_x . L'ensemble C_x est connexe par la proposition 1.3.1.2 (iv). Par ailleurs, $C_x \subset \overline{C_x}$ qui est connexe par la proposition 1.3.1.2 (iii), donc $C_x = \overline{C_x}$ est fermée. $\square$

Une composante connexe C_x avale tous les connexes qu'elle touche : si C_x intersecte un connexe C , alors $C_x \supset C$. L'espace X est connexe si et seulement si $C_x = X$. On peut détecter une composante connexe avec le lemme suivant.

Lemme 1.3.1.5. Si $A \subset X$ une partie non vide, ouverte, fermée et connexe, alors A est une composante connexe de X .

Démonstration. Soit $x \in A$, alors $A \subset C_x$ car A est connexe. Réciproquement, $C_x \cap A$ est non vide, ouvert et fermé dans C_x connexe : il est donc égal à C_x . On a donc $C_x \subset A$. $\square$

Une composante connexe n'est pas nécessairement ouverte : les composantes connexes de $\mathbf{Q}$ sont les singletons (exercice) qui ne sont pas ouverts.

1.3.2 Connexité par arc

Définition 1.3.2.1 (CONNEXITÉ PAR ARC). Soit X un espace topologique et $x, y \in X$ deux points. Un *arc* de x à y est une application continue $\gamma: [0, 1] \rightarrow X$ telle que $\gamma(0) = x$ et $\gamma(1) = y$. On dit que X est *connexe par arc* si pour tout $x, y \in X$, il existe un arc de x à y . Une partie $A \subset X$ est connexe par arc si elle est connexe par arc pour la topologie induite.

Proposition 1.3.2.2. Soit X un espace topologique.

- (i) Si X est connexe par arc, alors il est connexe.
- (ii) Si X est connexe par arc et $f: X \rightarrow Y$ est continue, alors $f(X)$ est connexe par arc.
- (iii) Si $(A_i)_{i \in I}$ est une famille de parties connexes par arc de X telle que $\bigcap_{i \in I} A_i \neq \emptyset$, alors $\bigcup_{i \in I} A_i$ est connexe par arc.
- (iv) Un produit fini d'espaces connexes par arc est connexe par arc.

Démonstration. (i) Supposons X connexe par arc et utilisons la caractérisation (proposition 1.3.1.2 (i)) de la connexité. Soit $f: X \rightarrow \{0, 1\}$ continue et soit $x, y \in X$. Puisque X est connexe par arc, il existe un arc γ de x à y . L'application composée $f \circ \gamma: [0, 1] \rightarrow \{0, 1\}$ est continue, à valeurs dans $\{0, 1\}$, donc constante puisque $[0, 1]$ est connexe. Ainsi, $f(x) = f \circ \gamma(0) = f \circ \gamma(1) = f(y)$.

(ii) et (iii) sont laissés en exercice. $\square$

Remarque 1.3.2.3. Attention! La réciproque à (i) est fausse et l'adhérence d'une partie connexe par arc n'est pas nécessairement connexe par arc. Le contre-exemple classique est fourni par l'adhérence $\bar{A}$ de $A = \{(x, \sin(1/x)); 0 < x < 1\} \subset \mathbf{R}^2$. L'ensemble A est connexe par arc, donc connexe comme image continue de $]0, 1[$ qui est connexe par arc. Son adhérence $\bar{A}$ est donc connexe par la proposition 1.3.1.2 (iii). Par contre, $\bar{A}$ n'est pas connexe par arc (exercice).

Exercice 1.3.2.4. Soit X un espace topologique. Le cône $C(X)$ et la suspension $S(X)$ sont alors connexes par arc.

Composantes connexes par arc : De même qu'avec la connexité, on définit une relation d'équivalence $\sim$ sur X par $x \sim y$ s'il existe une partie $C \subset X$ connexe par arc contenant x et y . Ses classes sont les *composantes connexes par arc* de X . Elles sont connexes par arc mais pas nécessairement fermées. Dans le contre-exemple ci-dessus, A est une composante connexe par arc de $\bar{A}$.

1.4 Compacité

1.4.1 Propriété de Borel-Lebesgue

Si X est un ensemble, on dit qu'une famille $(A_i)_{i \in I}$ de parties de X est un *recouvrement* de X (ou recouvre X) si $\bigcup_{i \in I} A_i = X$. Un *sous-recouvrement* est une sous-famille $(A_i)_{i \in J}$ (avec $J \subset I$) qui recouvre encore X . Si X est un espace topologique, un recouvrement $(A_i)_{i \in I}$ de X est dit ouvert ou fermé si les A_i le sont.

Définition 1.4.1.1 (ESPACE COMPACT). On dit qu'un espace topologique X est *compact* si les conditions suivantes sont satisfaites :

- (i) X est séparé.
 - (ii) Tout recouvrement ouvert de X admet un sous-recouvrement fini.
- Une partie $A \subset X$ sera compacte si A muni de la topologie induite est compacte.

L'hypothèse (ii) est appelée *axiome de Borel-Lebesgue*. Par passage au complémentaire, (ii) équivaut à : si $(F_i)_{i \in I}$ est une famille de fermés telle que $\bigcap_{i \in I} F_i = \emptyset$, alors il existe $J \subset I$ fini tel que $\bigcap_{i \in J} F_i = \emptyset$. Si la famille $(F_n)_{n \in \mathbf{N}}$ est décroissante ($F_{n+1} \subset F_n$), alors (ii) implique qu'il existe n tel que $F_n = \emptyset$.

Par contraposée, (ii) implique :

« Pour toute suite décroissante $(F_n)_{n \in \mathbf{N}}$ de fermés non vides de X , $\bigcap_{n \in \mathbf{N}} F_n$ est non vide. »

Si $A \subset X$ est compacte, l'assertion (ii) se traduit par : pour toute famille $(U_i)_{i \in I}$ d'ouverts de X tel que $A \subset \bigcup_{i \in I} U_i$, il existe $J \subset I$ fini tel que $A \subset \bigcup_{i \in J} U_i$. Par abus de langage, on pourra appeler $\bigcup_{i \in I} U_i$ un recouvrement ouvert de A (c'est les $U_i \cap A$ qui forment un recouvrement ouvert de A).

Proposition 1.4.1.2 (PROPRIÉTÉS FONDAMENTALES). (i) Si X est un espace compact, Y un espace séparé et $f: X \rightarrow Y$ continue, alors $f(X)$ est compact.

(ii) Si X est séparé et $K \subset X$ compact, alors K est fermé.

(iii) Si X est compact et $K \subset X$, alors K est compact si et seulement si K est fermé.

Démonstration. (i) L'image $f(X)$ est séparé comme partie d'un espace séparé. Soit $(U_i)_{i \in I}$ une famille d'ouverts de Y telle que $f(X) \subset \bigcup_{i \in I} U_i$. Puisque f est continue, chaque $f^{-1}(U_i)$ est ouvert et $(f^{-1}(U_i))_{i \in I}$ est un recouvrement ouvert de X . On extrait un sous-recouvrement fini $(f^{-1}(U_i))_{i \in J}$, $J \subset I$. Ainsi, $(U_i)_{i \in J}$ recouvre encore $f(X)$.

(ii) Soit $x \in X \setminus K$. Puisque X est séparé, pour chaque $a \in K$, il existe un ouvert U_a de X contenant x et un ouvert U_{xa} de X contenant a tel que $U_a \cap U_{xa} = \emptyset$. La famille $(U_a)_{a \in K}$ est un recouvrement ouvert de K . Par compacité de K , il existe $a_1, \dots, a_n \in K$ tel que $(U_{a_i})_{i=1, \dots, n}$ recouvre encore K . Posons $U_x = \bigcap_{i=1, \dots, n} U_{a_i}$, c'est un ouvert de X car l'intersection est finie, contenant x . De plus, $U_x \cap U_{a_i} = \emptyset$ pour chaque $i = 1, \dots, n$. On pose alors $U_K = \bigcup_{i=1, \dots, n} U_{a_i}$, c'est un ouvert de X contenant K et $U_x \cap U_K = \emptyset$. En particulier, $U_x \subset X \setminus K$ et $X \setminus K = \bigcup_{x \in X \setminus K} U_x$ est un ouvert de X , d'où K est fermé dans X .

(iii) Soit X compact. Il reste à montrer que si $K \subset X$ est fermé, alors K est compact. On sait que K est séparé comme partie de X séparé. Soit $(U_i)_{i \in I}$ une famille d'ouverts de X tel que $K \subset \bigcup_{i \in I} U_i$. En ajoutant l'ouvert $X \setminus K$ à la famille (U_i) , on obtient un recouvrement ouvert de X , donc on extrait un sous-recouvrement fini, de la forme $(U_i)_{i \in J}$, avec $J \subset I$ fini, plus éventuellement $X \setminus K$. Ainsi, $(U_i)_{i \in J}$ est un sous-recouvrement fini de K . $\square$

Seront très utiles :

Corollaire 1.4.1.3. Si X est un espace compact, Y un espace séparé et $f: X \rightarrow Y$ continue, bijective, alors f est un homéomorphisme.

Démonstration. Il suffit simplement de montrer que f est fermée. Soit F un fermé de X , alors F est compact, donc $f(F)$ est compact, puis fermé. $\square$

Corollaire 1.4.1.4. Soit $f: X \rightarrow Y$ continue où Y est séparé. Si $X/\mathcal{R}_f$ est compact, alors f se factorise en un homéomorphisme $\bar{f}: X/\mathcal{R}_f \rightarrow f(X)$.

Démonstration. L'application f se factorise en une bijection continue $\bar{f}: X/\mathcal{R}_f \rightarrow f(X)$ où $f(X)$ est séparé, on applique le corollaire 1.4.1.3 pour conclure. $\square$

Exemple 1.4.1.5. Soit $f: \mathbf{R} \rightarrow \mathbf{S}^1$ où $f(t) = (\cos(2\pi t), \sin(2\pi t))$. L'application f se factorise en une bijection continue $\bar{f}: \mathbf{R}/\mathbf{Z} \rightarrow \mathbf{S}^1$. L'espace $\mathbf{R}/\mathbf{Z}$ est séparé (car s'injecte continûment dans un séparé). En notant $\pi: \mathbf{R} \rightarrow \mathbf{R}/\mathbf{Z}$ la projection canonique, on a alors que $\mathbf{R}/\mathbf{Z} = \pi(\mathbf{R}) = \pi([0, 1])$ est compact comme image continue d'un compact dans un séparé. On conclut avec le corollaire 1.4.1.4 que $\bar{f}$ est un homéomorphisme.

Très surprenant :

Théorème 1.4.1.6 (TYCHONOFF). Un produit $\prod_{i \in I} X_i$ d'espaces topologiques compacts est compact.

Exercice 1.4.1.7. Soit X et Y deux espaces topologiques. On a l'équivalence :

$$X \times Y \text{ est compact} \Leftrightarrow X \text{ et } Y \text{ sont compacts.}$$

1.4.2 Compacité dans les espaces métriques

On rappelle qu'étant donnée une suite $(x_n)_{n \in \mathbf{N}}$ d'un espace topologique X , on dit que $y \in X$ est *valeur d'adhérence* de la suite $(x_n)_{n \in \mathbf{N}}$ si pour tout voisinage V de y , l'ensemble $\{n \in \mathbf{N}; x_n \in V\}$ est infini.

Exercice 1.4.2.1. Si X est métrisable ou plus généralement si y admet une base dénombrable de voisinages, alors $y \in X$ est valeur d'adhérence de $(x_n)_{n \in \mathbf{N}}$ si et seulement si il existe une sous-suite de $(x_n)_{n \in \mathbf{N}}$ convergeant vers y .

Théorème 1.4.2.2 (BOLZANO-WEIERSTRASS). Soit (E, d) un espace métrique et A une partie de E . L'ensemble A est compact si et seulement si toute suite de E admet une valeur d'adhérence.

Cela équivaut donc à dire que toute suite a une sous-suite convergente.

Corollaire 1.4.2.3. Soit (E, d) un espace métrique. Si $A \subset E$ est compact, alors A est fermée et bornée.

Démonstration. Soit $A \subset E$ compacte. Un espace métrique est séparé, donc A est fermé. Fixons $x \in E$. Si A n'est pas borné, pour tout $n \in \mathbf{N}_{>0}$, $A \not\subset B(x, n)$. On peut alors construire une suite $(x_n)_{n \in \mathbf{N}}$ de A telle que $d(x, x_n)$ tend vers $+\infty$ quand $n \rightarrow \infty$. Une telle suite n'a pas de sous-suite convergente car une suite convergente est bornée. $\square$

Théorème 1.4.2.4. Soit $(E, \|\cdot\|)$ un espace vectoriel normé de dimension finie. Les compacts de E sont alors les fermés bornés.

Remarque 1.4.2.5. En fait, la compacité des fermés bornés, ou de manière équivalente de la boule unité fermée, caractérise les e.v.n de dimension finie : c'est le *théorème de compacité de Riesz*.

2 Variétés

2.1 Variétés topologiques de dimension m

Définition 2.1.0.1 (VARIÉTÉ TOPOLOGIQUE). Une *variété topologique* de dimension $m \in \mathbf{N}$ est un espace topologique M localement homéomorphe à $\mathbf{R}^m$ (i.e. tout point de M admet un voisinage homéomorphe à $\mathbf{R}^m$).

Puisque tout ouvert de $\mathbf{R}^m$ est localement homéomorphe à $\mathbf{R}^m$, (en effet, si $f: x \in U \xrightarrow{\sim} V$ ouvert de $\mathbf{R}^m$, alors $x \in U \simeq B(a, \varepsilon) \subset V$, mais $B(a, \varepsilon) \simeq \mathbf{R}^m$) cela revient à dire que M est localement homéomorphe à un ouvert de $\mathbf{R}^m$. On pourra appeler M une m -variété ou simplement une variété.

Remarque 2.1.0.2. (1) Tout ouvert de $\mathbf{R}^m$ est une m -variété.

(2) Toute sous-variété différentiable de dimension m de $\mathbf{R}^n$ est une m -variété. Si $\phi: U \subset \mathbf{R}^n \rightarrow \mathbf{R}^n$ est un redressement local de M , alors sa restriction et corestriction $M \cap U \rightarrow \mathbf{R}^m \times \{0\}$ est un homéomorphisme vers un ouvert de $\mathbf{R}^m$.

(3) L'espace $M =]-1, 1[\times \{0\} \cup \{0\} \times]0, 1[\subset \mathbf{R}^2$ muni de la topologie induite n'est pas une 1-variété. Sinon il existerait $U \subset M$ un voisinage ouvert de 0 homéomorphe à $\mathbf{R}$, mais alors $U \setminus \{0\}$ serait homéomorphe à $\mathbf{R} \setminus \{\text{point}\}$, ce qui est impossible car $U \setminus \{0\}$ a trois composantes connexes et $\mathbf{R} \setminus \{\text{point}\}$ en a deux.

On a immédiatement :

Proposition 2.1.0.3. Si M et N sont des variétés topologiques de dimension m et n respectivement, alors l'espace produit $M \times N$ est une variété de dimension $m + n$.

La dimension m d'une variété topologique est constante par définition. Si on autorise la dimension $m(x) \in \mathbf{N}$ à dépendre du point $x \in M$, chaque composante connexe de M aurait néanmoins une dimension fixe car la fonction $x \mapsto m(x)$ est localement constante. Cela semble évident mais c'est en fait non trivial à prouver rigoureusement : cela requiert le théorème suivant, dont il n'y a pas de preuve simple (sans introduire de nouveaux outils).

Théorème 2.1.0.4 (INVARIANCE DU DOMAINE DE BROUWER). Soit U un ouvert de $\mathbf{R}^n$ et $f: U \rightarrow \mathbf{R}^n$ continue et injective, alors $V = f(U)$ est ouvert et f est un homéomorphisme entre U et V . En particulier, $\mathbf{R}^n$ est homéomorphe à $\mathbf{R}^m$ si et seulement si $n = m$.

La dernière assertion du théorème vient du fait que l'inclusion $i: \mathbf{R}^n \rightarrow \mathbf{R}^n \times \{0\} \subset \mathbf{R}^m$ est continue et injective si $n \leq m$, mais d'image $i(\mathbf{R}^n) = \mathbf{R}^n \times \{0\}$ clairement pas ouvert si $n < m$, donc s'il existait un homéomorphisme $f: \mathbf{R}^m \rightarrow \mathbf{R}^n$ avec $n < m$, alors $i \circ f: \mathbf{R}^m \rightarrow \mathbf{R}^m$ serait continue injective, ce qui contredit la première assertion.

Il suit du théorème que la dimension d'une variété topologique est un invariant d'homéomorphie : deux variétés topologiques homéomorphes ont même dimension. Une variété de dimension 2 est une *surface*.

Définition 2.1.0.5 (CARTE, ATLAS). Soit M une m -variété.

- (i) Une *carte* de M est un couple (U, ϕ) où U est un ouvert de M , l'application $\phi: U \rightarrow \phi(U)$ un homéomorphisme et $\phi(U)$ un ouvert de $\mathbf{R}^m$.
- (ii) L'ouvert U est le *domaine* de la carte et l'application réciproque $\phi^{-1}: \phi(U) \rightarrow U$ est une *coordonnée locale*.
- (iii) Étant donné deux cartes (U, ϕ) et (V, ψ) de M , on appelle l'application de changement de coordonnées (qui est un homéomorphisme entre ouverts de $\mathbf{R}^m$)

$$\psi \circ \phi^{-1}: \phi(U \cap V) \rightarrow \psi(U \cap V)$$

un *changement de cartes* (avec un abus de langage).

- (iv) Un *atlas* de M est une famille de cartes $\{(U_i, \phi_i); i \in I\}$ de M telle que $\bigcup_{i \in I} U_i = M$.

Toute variété topologique admet un atlas.

Remarque 2.1.0.6. (1) Une variété topologique M peut-être enrichie d'une structure supplémentaire en demandant qu'un atlas $\mathcal{A} = \{(U_i, \phi_i); i \in I\}$ de M ait des changements de cartes préservant la structure.

(2) M est une variété différentiable de classe $\mathcal{C}^k$ si elle admet un atlas dont les changements de cartes sont des $\mathcal{C}^k$ -difféomorphismes. Exemples : les sous-variétés différentiables $\mathcal{C}^k$.

(3) M est une variété affine si les changements de cartes sont des restrictions d'applications affines.

Proposition 2.1.0.7 (PROPRIÉTÉS DES VARIÉTÉS TOPOLOGIQUES). Soit M une variété topologique.

- (i) M est localement connexe par arc (tout point admet une base de voisinages connexes par arc).
- (ii) M est localement séparée (tout point admet un voisinage séparé).
- (iii) Tout point de M admet une base de voisinages compacts. Si M est séparée, alors M est localement compacte.

Démonstration. Soit $a \in M$ et soit (U, ϕ) une carte autour de a telle que $\phi(U) = \mathbf{R}^m$. Tout d'abord, U est un voisinage séparé de a car homéomorphe à $\mathbf{R}^m$ qui est séparé. Posons pour $n \in \mathbf{N}_{>0}$, $U_n := \phi^{-1}(B(0, 1/n))$. On a alors que U_n est connexe par arcs et $\overline{U_n} = \phi^{-1}(\overline{B(0, 1/n)})$ est compact dans U (donc dans M) comme image continue d'un compact dans un séparé. Ainsi, $\{U_n\}_{n \in \mathbf{N}}$ est une base de voisinages connexes par arc de a et $\{\overline{U_n}\}_{n \in \mathbf{N}}$ est une base de voisinages compacts de a . $\square$

Remarque 2.1.0.8. Une variété topologique n'est pas nécessairement séparée!

Exercice 2.1.0.9 (LA DROITE À DEUX ORIGINES - EN TD). Soit $M = \mathbf{R} \times \{1, 2\} / \sim$ où la relation d'équivalence est engendrée par $(x, 1) \sim (x, 2)$ pour tout $x \neq 0$. Montrer que M est une 1-variété connexe avec un atlas à 2 cartes, mais non séparée.

La même construction en remplaçant $\mathbf{R}$ par $\mathbf{S}^1$ donne une 1-variété qui a la propriété de Borel-Lebesgue (axiome (ii) de la compacité) mais n'est pas compacte car pas séparée.

Exercice 2.1.0.10. Montrer que si M est une variété compacte, elle est à base dénombrable (sa topologie admet une base dénombrable).

Pour la culture générale on mentionne le théorème suivant :

Théorème 2.1.0.11 (PLONGEMENT DE WITNEY | ~ 1935). Soit M une m -variété topologique séparée à base dénombrable. Il existe alors $\phi: M \rightarrow \mathbf{R}^{2m+1}$ un plongement. En particulier, M est métrisable.

Ce sera le cas des variétés compactes qu'on verra par la suite. Par contre, il est fréquent de rencontrer des m -variétés très raisonnables qui ne se plongent pas dans $\mathbf{R}^{m+1}$. Ainsi, on verra que parmi les surfaces compactes, la moitié exactement, ne se plonge pas dans $\mathbf{R}^3$. On n'utilisera pas le théorème de Whitney, on admet sa preuve.

MAINTENANT, TOUTES LES VARIÉTÉS SERONT SUPPOSÉES SÉPARÉES À BASE DÉNOMBRABLE.

On ne sera donc pas embêté avec les pathologies sur les critères séquentiels, entre autres. Vérifions que nos variétés favorites en font partie.

2.2 Sphères et projectifs

2.2.1 Sphères

On note $\mathbf{S}^n$ la sphère unité de $\mathbf{R}^{n+1}$

$$\mathbf{S}^n = \left\{ x = (x_1, \dots, x_{n+1}) \in \mathbf{R}^{n+1}; |x|^2 = \sum_{i=1}^{n+1} x_i^2 = 1 \right\}$$

munie de la topologie induite.

Proposition 2.2.1.1. La sphère $\mathbf{S}^n$ est une n -variété compacte admettant un atlas à deux cartes. Elle est connexe si $n \geq 1$.

Démonstration. • Si $n = 0$, la sphère $\mathbf{S}^0 = \{-1, 1\}$ est une variété de dimension 0 compacte à deux composantes connexes.

• Si $n \geq 1$, les cartes sont les deux projections stéréographiques associées au pôle $N = (0, 1) \in \mathbf{R}^n \times \mathbf{R}$ et $S = (0, -1)$ qu'on rappelle. Pour $x \in \mathbf{S}^n \setminus \{N\}$, on définit $\phi_N(x)$ comme l'intersection de la droite (Nx) avec $\mathbf{R}^n \times \{0\}$ et pour $x \in \mathbf{S}^n \setminus \{S\}$, $\phi_S(x)$ est l'intersection de la droite (Sx) avec $\mathbf{R}^n \times \{0\}$. On peut calculer (exercice, penser à Thalès!) que

$$\phi_N(u, t) = \frac{u}{1-t}, ((u, t) \in \mathbf{S}^n \setminus \{N\}) \quad \text{et} \quad \phi_S(u, t) = \frac{u}{1+t}, ((u, t) \in \mathbf{S}^n \setminus \{S\}).$$

Ce sont des homéomorphismes vers $\mathbf{R}^n$ de réciproques

$$\phi_N^{-1}(y) = \frac{(2y, \|y\|^2 - 1)}{\|y\|^2 + 1} \quad \text{et} \quad \phi_S^{-1}(y) = \frac{(2y, -\|y\|^2 + 1)}{\|y\|^2 + 1}.$$

Comme $\phi_S(N) = 0 = \phi_N(S)$, le changement des cartes $\phi_S \circ \phi_N^{-1}$ est défini sur $\mathbf{R}^n \times \{0\} \rightarrow \mathbf{R}^n \times \{0\}$ et on calcule que c'est l'inversion

$$y \mapsto \frac{y}{\|y\|^2}.$$

L'espace topologique $\mathbf{S}^n$ est fermé borné, donc compact. Les deux domaines de cartes sont homéomorphes à $\mathbf{R}^n$ donc connexes et d'intersection non vide et ainsi, $\mathbf{S}^n$ est connexe. $\square$

Exercice 2.2.1.2 (en TD). Montrer que la restriction à $\mathbf{S}^1 \cap \{y > 0\}$ de la projection $(x, y) \mapsto x$ définit une carte de $\mathbf{S}^1$. Déterminer un atlas de $\mathbf{S}^1$ à l'aide de projections. Montrer que les changements de carte sont $\mathcal{C}^\infty$.

Exercice 2.2.1.3 (en TD). Montrer que l'application $]0, \infty[\times \mathbf{S}^n \rightarrow \mathbf{R}^{n+1} \setminus \{0\}; (r, x) \mapsto rx$ est un homéomorphisme.

2.2.2 Espaces projectifs

Définition 2.2.2.1 (ESPACE PROJECTIF). Soit $\mathbf{k} = \mathbf{R}$ ou $\mathbf{C}$. L'espace projectif $\mathbf{P}_{\mathbf{k}}^n$ (aussé noté $\mathbf{P}^n(\mathbf{k})$) est l'ensemble des droites vectorielles de $\mathbf{k}^{n+1}$. Il est topologisé en l'identifiant au quotient

$$(\mathbf{k}^{n+1} \setminus \{0\}) / \sim$$

où $x \sim y$ si et seulement si x et y sont colinéaires, i.e. $x = \lambda y$ pour $\lambda \in \mathbf{k} \setminus \{0\}$.

Proposition 2.2.2.2. L'espace projectif réel $\mathbf{P}_{\mathbf{R}}^n$ (resp. complexe $\mathbf{P}_{\mathbf{C}}^n$) est une variété topologique de dimension n (resp. $2n$) connexe et compacte. Il admet un atlas à $n + 1$ cartes.

Démonstration. • On note $\pi: \mathbf{k}^{n+1} \setminus \{0\} \rightarrow \mathbf{P}_{\mathbf{k}}^n$ la projection et si $x \in \mathbf{k}^{n+1} \setminus \{0\}$, on note $[x] = \pi(x)$ sa classe. On munit $\mathbf{P}_{\mathbf{k}}^n$ de l'atlas suivant : pour $i \in \{1, \dots, n+1\}$, soit $U_i = \{[x_1, \dots, x_{n+1}] \in \mathbf{P}_{\mathbf{k}}^n; x_i \neq 0\}$. C'est un ouvert puisque $\pi^{-1}(U_i) = \{x \in \mathbf{k}^{n+1} \setminus \{0\}; x_i \neq 0\}$ est ouvert. On définit $\phi_i: U_i \rightarrow \mathbf{k}^n$ ($\mathbf{k}^n = \mathbf{R}^n$ ou $\mathbf{k}^n = \mathbf{C}^n = \mathbf{R}^{2n}$) par

$$\phi_i([x_1, \dots, x_{n+1}]) = \left(\frac{x_1}{x_i}, \dots, \frac{\hat{x}_i}{x_i}, \dots, \frac{x_{n+1}}{x_i} \right) \in \mathbf{k}^n$$

(le symbole $\hat{x}_i$ sur x_i veut dire que l'on retire la i -ème coordonnée). On obtient le diagramme commutatif suivant :

$$\begin{array}{ccc} \pi^{-1}(U_i) & \xrightarrow{\phi_i \circ \pi} & \mathbf{k}^n \\ \pi \searrow & & \nearrow \phi_i \\ & U_i \subset \mathbf{P}_{\mathbf{k}}^n & \end{array}$$

L'application est bien définie et continue puisque c'est la factorisation de

$$x \in \pi^{-1}(U_i) \mapsto \left(\frac{x_1}{x_i}, \dots, \frac{\hat{x}_i}{x_i}, \dots, \frac{x_{n+1}}{x_i} \right) \in \mathbf{k}^n$$

qui est constante sur chaque classe et continue (théorème 1.2.5.3).

- Montrons qu'elle est injective : soit $[x], [x'] \in U_i$ tel que $\phi_i([x]) = \phi_i([x'])$. Puisque $\phi_i([x]) = \phi_i([\lambda x])$ pour tout $\lambda \in \mathbf{k} \setminus \{0\}$, on peut supposer que $x_i = 1 = x'_i$. On a donc par hypothèse l'égalité $(x_1, \dots, \hat{1}, \dots, x_{n+1}) = (x'_1, \dots, \hat{1}, \dots, x'_{n+1})$, d'où $[x] = [x']$.
- L'application ϕ_i est clairement surjective si pour tout $y \in \mathbf{k}^n$, on prend

$$(y_1, \dots, y_n) = \phi_i([y_1, \dots, 1, \dots, y_n]) \in U_i$$

(avec 1 sur la i -ème coordonnée). L'application est donc bijective de réciproque la composée

$$\phi_i^{-1}(y_1, \dots, y_n) = [y_1, \dots, 1, \dots, y_n]$$

qui est continue. Ainsi, ϕ_i est un homéomorphisme et (U_i, ϕ_i) une carte de $\mathbf{P}_{\mathbf{k}}^n$.

- Clairement, les U_i recouvrent $\mathbf{P}_{\mathbf{k}}^n$, donc la famille $\{(U_i, \phi_i); 1 \leq i \leq n+1\}$ est un atlas de $\mathbf{P}_{\mathbf{k}}^n$.
- Si $n = 0$, alors $\mathbf{P}_{\mathbf{k}}^0 = \{[1]\}$ qui est bien compact et connexe. Supposons $n \geq 1$, alors $\mathbf{k}^{n+1} \setminus \{0\}$ est connexe, donc $\mathbf{P}_{\mathbf{k}}^n$ est connexe comme image continue d'un connexe.
- Montrons que $\mathbf{P}_{\mathbf{k}}^n$ est compact. On commence par la séparation. Soit $[x] \neq [y] \in \mathbf{P}_{\mathbf{k}}^n$, i.e. x et y ne sont pas colinéaires. On se donne deux formes linéaires $a, b: \mathbf{k}^{n+1} \rightarrow \mathbf{k}$ telles que $a(x) = 0, b(x) = 1$ et $a(y) = 1, b(y) = 0$ puis on définit

$$C_x := \left\{ z \in \mathbf{k}^{n+1} \setminus \{0\}; \frac{|a(z)|}{|b(z)|} < 1 \right\} \quad C_y := \left\{ z \in \mathbf{k}^{n+1} \setminus \{0\}; \frac{|b(z)|}{|a(z)|} < 1 \right\}.$$

Ces ensembles sont ouverts, disjoints et invariants par multiplication scalaire (i.e. $z \in C$ implique que $\lambda z \in C$ pour $\lambda \in \mathbf{k} \setminus \{0\}$), donc saturés. Évidemment, $x \in C_x$ et $y \in C_y$. Il s'ensuit que $\pi(C_x)$ et $\pi(C_y)$ sont des ouverts disjoints séparant $[x]$ et $[y]$, prouvant que $\mathbf{P}_{\mathbf{k}}^n$ est séparé. On montre maintenant que $\mathbf{P}_{\mathbf{k}}^n$ est compact comme image continue d'un compact dans un séparé. Notons $\mathbf{S}$ la sphère unité dans $\mathbf{R}^{n+1} \setminus \{0\}$ si $\mathbf{k} = \mathbf{R}$ ou dans $\mathbf{R}^{2n+2} \setminus \{0\} = \mathbf{C}^{n+1} \setminus \{0\}$ lorsque $\mathbf{k} = \mathbf{C}$. On a alors $\mathbf{P}_{\mathbf{k}}^n = \pi(\mathbf{S})$. En effet, soit $[x] \in \mathbf{P}_{\mathbf{k}}^n$. On a donc $\frac{x}{|x|} \in \mathbf{S}$ et $[x] = [\frac{x}{|x|}] = \pi(\frac{x}{|x|}) \in \pi(\mathbf{S})$. $\square$

Proposition 2.2.2.3. Soit $n \in \mathbf{N}$.

- (i) L'espace projectif réel $\mathbf{P}_{\mathbf{R}}^n$ est homéomorphe à $\mathbf{S}^n / \sim$ où $\sim$ est engendrée par $x \sim -x$ pour tout $x \in \mathbf{S}^n$.
(ii) L'espace projectif complexe $\mathbf{P}_{\mathbf{C}}^n$ est homéomorphe à $\mathbf{S}^{2n+1} / \sim$ où $\sim$ est engendrée par $z \sim e^{i\theta} z$ pour tout $z \in \mathbf{S}^{2n+1} \subset \mathbf{C}^{n+1}$ et $\theta \in \mathbf{R}$.

Démonstration. (i) On considère l'application $f: \mathbf{S}^n \rightarrow \mathbf{P}_{\mathbf{R}}^n$ définie par

$$\begin{array}{ccccc} f: & \mathbf{S}^n & \xrightarrow{i} & \mathbf{R}^{n+1} \setminus \{0\} & \xrightarrow{\pi} & \mathbf{P}_{\mathbf{R}}^n \\ & x & \mapsto & x & \mapsto & [x] \end{array}$$

Clairement, f est surjective continue et $f(x) = f(y) \Leftrightarrow x = \lambda y$ où $\lambda \in \mathbf{R} \setminus \{0\} \Leftrightarrow x = \pm y$ (car $\|x\| = |\lambda| \cdot \|y\|$ et $\|x\| = \|y\| = 1$) $\Leftrightarrow x \sim y$. L'application f se factorise donc en une bijection continue $\bar{f}: \mathbf{S}^n / \sim \rightarrow \mathbf{P}_{\mathbf{R}}^n$:

$$\begin{array}{ccc} \mathbf{S}^n & \xrightarrow{f} & \mathbf{P}_{\mathbf{R}}^n \\ \pi \searrow & & \nearrow \bar{f} \\ & \mathbf{S}^n / \sim & \end{array}$$

L'espace $\mathbf{S}^n / \sim$ est séparé car s'injecte continûment dans $\mathbf{P}_{\mathbf{R}}^n$ séparé. Ainsi, $\mathbf{P}_{\mathbf{R}}^n$ est compact comme image continue de $\mathbf{S}^n$ compact. On conclut que la bijection continue $\bar{f}$ est un homéomorphisme à l'aide du corollaire 1.4.1.3.

(ii) Même argument (exercice). □

Exercice 2.2.2.4 (en TD). (1) Montrer que $\mathbf{P}_{\mathbf{R}}^1$ est homéomorphe à $\mathbf{S}^1$.

(2) Montrer que $\mathbf{P}_{\mathbf{C}}^1$ est homéomorphe à $\mathbf{S}^2$.

Solution : On cherche une application F telle que la diagramme ci-dessous est commutatif :

$$\begin{array}{ccc} \mathbf{S}^1 & \xrightarrow{F} & \mathbf{P}_{\mathbf{R}}^1 \\ \pi \searrow & & \nearrow f \\ & \mathbf{S}^1 / \{\pm 1\} & \end{array}$$

Prenons par exemple $F(z) = z^2$ et notons $[z]$ la classe d'équivalence de z dans $\mathbf{S}^1 / \{\pm 1\}$ i.e. $[z] = \{-z, z\}$. Supposons que $f([z]) = f([z'])$, donc $z^2 = z'^2$, d'où $z = \pm z'$ et on en déduit que $[z] = [z']$. Comme F est surjective continue, alors f est surjective continue, donc f est bijective, $\mathbf{S}^1 / \{\pm 1\}$ compact car image continue d'un compact et $\mathbf{P}_{\mathbf{R}}^1$ est compact, donc séparé. En outre, f est un homéomorphisme.

On verra que $\mathbf{P}_{\mathbf{R}}^2$ n'est pas homéomorphe à $\mathbf{S}^2$ (en fait il n'est pas plongeable dans $\mathbf{R}^3$).

Exercice 2.2.2.5 (DROITE AFFINES - EN TD). Montrer que l'ensemble des droites affines $\{ax + by + c = 0\}$ de $\mathbf{R}^2$ admet une structure de 2-variété topologique.

2.3 Opérations sur les variétés topologiques

2.3.1 Quotients de variétés

On investigate les conditions pour que M/G soit une variété (séparée) lorsqu'un groupe G agit par homéomorphisme sur une variété M . On rappelle que M/G désigne l'espace topologique quotient $M / \sim$ où $\sim$ est engendrée par $x \sim gx$ pour tout $x \in M$ et $g \in G$. Deux exemples déjà vus :

- $\mathbf{S}^2 / (\mathbf{Z} / 2 \mathbf{Z}) = \mathbf{P}_{\mathbf{R}}^2$ (où $(\mathbf{Z} / 2 \mathbf{Z}, +)$ agit par antipodie : $[1] \cdot x = -x$) ;
- $\mathbf{R} / \mathbf{Z} = \mathbf{S}^1$ (où $(\mathbf{Z}, +)$ agit par translation : $k \cdot x = x + k$).

Introduisons auparavant une notion reliée, celle de revêtement.

Définition 2.3.1.1 (REVÊTEMENT). Une application $p: X \rightarrow Y$ entre deux espaces topologiques est un *revêtement* si :

- (i) p est continue et surjective ;
(ii) pour tout $y \in Y$, il existe un voisinage ouvert V de y tel que $p^{-1}(V) = \bigcup_{i \in I} U_i$ où les U_i sont des ouverts disjoints dans X tel que $p: U_i \rightarrow V$ soit un homéomorphisme pour chaque $i \in I$.

En particulier, p est un homéomorphisme local mais un homéomorphisme local n'est pas nécessairement un revêtement. Il est clair que $y \mapsto \text{Card}(p^{-1}(\{y\})) \in \mathbf{N} \cup \{+\infty\}$ est localement constant, donc constant si Y est connexe. Quand ce nombre est fini, on l'appelle le *degré* du revêtement. L'ensemble $p^{-1}(\{y\})$ est appelé la *fibres* au dessus de y .

Exemple 2.3.1.2. (1) L'application quotient $p: \mathbf{S}^2 \rightarrow \mathbf{P}_{\mathbf{R}}^2$ est un revêtement de degré 2 : posons H^+ l'hémisphère nord de $\mathbf{S}^2$ et H^- l'hémisphère sud. On a $V := p(H^+) = p(H^-)$ et $p^{-1}(V) = H^+ \cup H^-$. Ainsi, on a bien un homéomorphisme $p: H^\pm \xrightarrow{\sim} V$.
 (2) L'application $f: \mathbf{R} \rightarrow \mathbf{R}/\mathbf{Z} \simeq \mathbf{S}^1; t \mapsto (\cos(t), \sin(t))$ définit un revêtement de $\mathbf{R}$ sur $\mathbf{S}^1$ (de degré infini).
 (3) Sa restriction $] - 2\pi, 2\pi[\rightarrow \mathbf{S}^1$ est encore continue, surjective et un homéomorphisme local mais n'est pas un revêtement.

On rappelle qu'une action d'un groupe G sur un ensemble X est un morphisme $\rho: G \rightarrow \text{Bij}(X) = \mathfrak{S}_X$. On doit donc avoir $\rho(gh)(x) = \rho(g)(\rho(h)(x))$ pour tout $g, h \in G$ et $x \in X$. En général, on note simplement gx l'élément $\rho(g)(x)$. Ainsi, $(gh)x = g(hx) = ghx$. Nécessairement, le neutre $e \in G$ agit comme Id_X puisque $e(ex) = (ee)x = ex$ pour tout $x \in X$, donc e fixe tous les ex de X i.e. tous les points. L'orbite d'un élément x est l'ensemble $Gx = \{gx; g \in G\}$. Selon la structure que porte X , on peut demander que $\rho(G)$ arrive dans une partie sympathique de $\text{Bij}(X)$, comme $\text{Homeo}(X)$ (si X est un espace topologique), $\text{Diffeo}(X)$ (si X a une structure différentiable), $\text{Isom}(X)$ (si X est un espace métrique ou un groupe), $\text{GL}(X)$ (si X est un espace vectoriel), etc...

Exercice 2.3.1.3 (EN TD). (1) Déterminer les actions par isométries de $(\mathbf{Z}/2\mathbf{Z}, +)$ sur $\mathbf{R}$.
 (2) Même question sur $\mathbf{R}^2$.

Définition 2.3.1.4 (ACTION SÉPARANTE, TOTALEMENT DISCONTINUE). Une action de G sur un espace topologique X est :

- (i) *séparante* si pour tous $x, y \in X$ tel que $y \notin Gx$, il existe des voisinages ouverts U de x et V de y tels que $GU \cap V = \emptyset$.
- (ii) *totalement discontinue* si pour tout $x \in X$, il existe un voisinage ouvert U de x tel que pour tout $g \in G \setminus \{e\}$, on a $gU \cap U = \emptyset$.

Remarque 2.3.1.5. Dans (i), on peut renforcer la conclusion en $GU \cap GV = \emptyset$. En effet, pour $g, h \in G$ on a

$$gU \cap hV = \emptyset \Leftrightarrow h^{-1}gU \cap V = \emptyset.$$

Si l'action de G est séparante, alors X/G est séparé. En effet, si $[x] \neq [y] \in X/G$, alors $y \notin Gx$. On prend des voisinages ouverts U et V de x et y respectivement tels que $GU \cap GV = \emptyset$. Les ouverts GU et GV sont saturés disjoints, leur quotient dans X/G sont des ouverts disjoints séparant $[x]$ et $[y]$.

Exercice 2.3.1.6. L'action par translation de $G = (\mathbf{Z}^n, +)$ sur $\mathbf{R}^n$, définie par $gx = x + g$, où $g = (g_1, \dots, g_n) \in \mathbf{Z}^n$, est séparante et totalement discontinue. Pour (i), soit $x, y \in \mathbf{R}^n$ avec $y \notin Gx = x + \mathbf{Z}^n$. Clairement, $B(y, 2d(x, y)) \cap Gx$ est fini, donc $\varepsilon := d(y, Gx) > 0$. On prend alors $U = B(x, \varepsilon/2)$ et $V = B(y, \varepsilon/2)$. L'ensemble GU est une réunion de boule de rayon $\varepsilon/2$ centrées sur les éléments de Gx , donc $GU \cap V = \emptyset$. Pour (ii), il suffit de prendre $U = B(x, 1/2)$.

Théorème 2.3.1.7. Soit G un groupe agissant de manière totalement discontinue et séparante sur une m -variété M . L'espace topologique quotient M/G est une m -variété séparée et la projection $p: M \rightarrow M/G$ est un revêtement.

Démonstration. L'espace M/G est muni de la topologie quotient. Il est séparé par la remarque ci-dessus. Soit $[x] = Gx \in M/G$ et U un voisinage ouvert de x comme dans la définition précédente. Puisque gU est disjoint de U pour tout $g \in G$ différent de e , la projection p est injective $U \rightarrow p(U) \subset M/G$. Elle est continue par définition de la topologie quotient et ouverte puisque G agit par homéomorphisme. Il s'ensuit que $p(U)$ est un ouvert et que $p: U \rightarrow p(U)$ est un homéomorphisme. Cela montre que M/G est localement homéomorphe à $\mathbf{R}^m$, puisque U l'est. Ensuite, $gU \cap hU = \emptyset$ pour tout $g \neq h \in G$, donc $\bigcup_{g \in G} gU$ est une union d'ouverts disjoints. Comme

au-dessus, $p: gU \rightarrow p(U)$ est un homéomorphisme. Il s'ensuit que $p: GU \rightarrow p(U)$ est un revêtement et donc que $p: M \rightarrow M/G$ est un revêtement. $\square$

Définition 2.3.1.8 (ACTION SANS POINT FIXE). Un groupe G agit *sans point fixe* sur un espace X si pour tout $g \in G \setminus \{e\}$ et tout $x \in X$, on a $gx \neq x$.

L'action par antipodie de $\mathbf{Z}/2\mathbf{Z}$ sur $\mathbf{S}^n$ et par translations de $\mathbf{Z}^n$ sur $\mathbf{R}^n$ sont sans point fixes. Une action totalement discontinue est sans point fixe.

Théorème 2.3.1.9. Si G est un groupe fini agissant sans point fixe sur une m -variété M séparée, alors M/G est une m -variété séparée.

Démonstration. • Montrons d'abord que l'action est totalement discontinue. Notons $G = \{e = g_1, g_2, \dots, g_k\}$. Soit $x \in M$ et soit $x_i = g_i x$ pour $i \in \{1, \dots, k\}$. Puisque G agit sans point fixe, on a $x_i \neq x_j$ si $i \neq j$. Puisque $\{x_1, \dots, x_k\}$ est fini et M est séparé, on peut trouver des voisinages ouverts U_i de x_i tels que $U_i \cap U_j = \emptyset$

si $i \neq j$ (en prenant l'intersection de voisinages séparant les points deux à deux). On pose alors $V_1 = U_1 \cap g_2^{-1}(U_2) \cap \dots \cap g_k^{-1}(U_k)$, puis $V_i = g_i(V_1)$. Ce sont des voisinages de x_i contenus dans U_i , donc $V_i \cap V_j = \emptyset$ si $i \neq j$. En particulier, pour $g \in G \setminus \{e\}$, on a $g(V_1) \cap V_1 = \emptyset$.

• On montre maintenant que l'action est séparante. Soit $x, x' \in M$ tel que $x' \notin Gx$. On a alors $Gx = \{x_1, \dots, x_k\}$ et $Gx' = \{x'_1, \dots, x'_k\}$ où tous les points sont distincts. On procède comme au dessus pour construire des voisinages U_i des x_i et U'_i des x'_i deux à deux disjoints puis on pose $V_1 = \bigcap_i g_i^{-1}(U_i)$ et $V'_1 = \bigcap_i g_i^{-1}(U'_i)$. Ainsi, $GV_1 \cap GV'_1 = \emptyset$. On conclut en appliquant le théorème 2.3.1.7. $\square$

2.3.2 Exemples de surfaces obtenues comme quotient

Exemple 2.3.2.1 (CYLINDRE REVISITÉ). Le cylindre $\mathbf{S}^1 \times]-1, 1[$ est homéomorphe à $(\mathbf{R} \times]-1, 1[)/\mathbf{Z}$ où $(\mathbf{Z}, +)$ agit par translation sur le premier facteur, i.e. est engendrée par $1 \cdot (x, y) = (x + 1, y)$. On obtient l'homéomorphisme par factorisation :

$$\begin{array}{ccc} \mathbf{R} \times]-1, 1[& \xrightarrow{f(x,t)=(e^{2i\pi x}, t)} & \mathbf{S}^1 \times]-1, 1[\\ \pi \searrow & & \nearrow \bar{f} \\ & (\mathbf{R} \times]-1, 1[)/\mathbf{Z} & \end{array}$$

Exemple 2.3.2.2 (TORE REVISITÉ). Le tore $\mathbf{T}^n = \mathbf{S}^1 \times \dots \times \mathbf{S}^1$ est homéomorphe à $\mathbf{R}^n / \mathbf{Z}^n$ où $(\mathbf{Z}^n, +)$ agit par translation. On a vu que l'action est totalement discontinue et séparante, donc $\mathbf{R}^n / \mathbf{Z}^n$ est une n -variété. On peut vérifier que $f: \mathbf{R}^n \rightarrow \mathbf{T}^n$ tel que $(x_1, \dots, x_n) \mapsto (e^{2i\pi x_1}, \dots, e^{2i\pi x_n})$ se factorise en une bijection continue $\bar{f}: \mathbf{R}^n / \mathbf{Z}^n \rightarrow \mathbf{T}^n$.

$$\begin{array}{ccc} \mathbf{R}^n & \xrightarrow{f} & \mathbf{T}^n \\ \pi \searrow & & \nearrow \bar{f} \\ & \mathbf{R}^n / \mathbf{Z}^n & \end{array}$$

On en déduit que $\mathbf{R}^n / \mathbf{Z}^n$ est compact : séparé car s'injecte continûment dans $\mathbf{T}^n$ séparé puis image continue de $[0, 1]^n$ compact. On conclut que $\bar{f}$ est un homéomorphisme.

Exemple 2.3.2.3 (RUBAN DE MOEBIUS). Soit $M = \mathbf{S}^1 \times]-1, 1[$ (le cylindre). On définit une action par homéomorphisme de $G = (\mathbf{Z}/2\mathbf{Z}, +)$ sur M par $[1] \cdot (z, t) = (-z, -t)$. C'est bien une action car

$$[1]([1](z, t)) = [1](-z, -t) = (z, t) = [0](z, t) = ([1] + [1])(z, t).$$

L'action est sans point fixe, donc d'après le théorème précédent, M/G est une surface, qu'on appelle *ruban de Moebius* (ainsi que toute surface qui lui est homéomorphe).

Exercice 2.3.2.4 (RUBAN DE MOEBIUS - EN TD). (1) Montrer que le ruban de Moebius est homéomorphe à $M = \mathbf{R} \times]-1, 1[/\mathbf{Z}$ où l'action de $\mathbf{Z}$ est engendrée par $1 \cdot (x, t) = (x + 1, t)$. En déduire que le ruban de Moebius s'obtient d'un rectangle en identifiant deux côtés opposés.

(2) Montrer que le projectif réel $\mathbf{P}_{\mathbf{R}}^2$ est réunion d'un ruban de Moebius et d'un disque fermé.

Exemple 2.3.2.5 (BOUEILLE DE KLEIN). Soit $\mathbf{T}^2 = \mathbf{S}^1 \times \mathbf{S}^1 \subset \mathbf{C}^2$ (le tore). On fait agir $\mathbf{Z}/2\mathbf{Z}$ sur $\mathbf{T}^2$ par $[1](z, z') = (-z, -z')$. C'est clairement une action par homéomorphisme, sans point fixe, donc $\mathbf{T}^2/(\mathbf{Z}/2\mathbf{Z})$ est une surface compacte connexe, qu'on appelle *bouteille de Klein* notée $\mathbf{K}^2$.

Exercice 2.3.2.6 (BOUEILLE DE KLEIN - EN TD). (1) Montrer que la bouteille de Klein $\mathbf{K}^2$ est homéomorphe à $(\mathbf{R}^2 / \mathbf{Z}^2)/(\mathbf{Z}/2\mathbf{Z})$ où

$$[1] \cdot [(x, y)] = [(x + 1/2, -y)].$$

En déduire que $\mathbf{K}^2$ peut s'obtenir d'un rectangle en identifiant les côtés par paires.

On veut montrer que $\mathbf{K}^2$ est homéomorphe à $\mathbf{R}^2/G$ où $G = (\mathbf{Z} \times \mathbf{Z}, \bullet)$ agit par isométries sur $\mathbf{R}^2$ et la loi $\bullet$ du groupe est à déterminer. On note $a = (1, 0)$ et $b = (0, 1)$ et on pose que

$$a(x, y) = (x + 1/2, -y) \quad b(x, y) = (x, y + 1).$$

(2) Calculer $a(b(x, y))$ et $b(a(x, y))$. Qu'en conclure ?

(3) Déterminer une formule pour $(k, \ell)(x, y)$ qui étende a et b à tout $(k, \ell) \in \mathbf{Z} \times \mathbf{Z}$.

(4) Déterminer une formule pour $(k, \ell) \bullet (k', \ell')$ compatible avec l'action formulée précédemment.

(5) Vérifier que $\bullet$ est bien une loi de groupe sur $\mathbf{Z} \times \mathbf{Z}$.

Bravo ! Vous avez découvert le produit semi-direct $\mathbf{Z} \rtimes \mathbf{Z}$ et montré que $\mathbf{K}^2 = \mathbf{R}^2/(\mathbf{Z} \rtimes \mathbf{Z})$.

2.3.3 Somme connexe

Heuristique : la somme connexe de deux m -variétés M et N est la variété obtenue en enlevant une boule ouverte à chacune et en recollant les sphères de bord par un homéomorphisme. Soit donc (U, ϕ) une carte de M et (V, ψ) une carte de N telles que $\phi(U) = \psi(V) = \mathbf{R}^m$. Soit $U_1 \subset U$ et $V_1 \subset V$ les domaines correspondants à $\mathbf{B}(0, 1) \subset \mathbf{R}^m$. Soit $f: \partial U_1 \rightarrow \partial V_1$ définie par $f(x) = \psi^{-1} \circ \phi(x)$. Notons $M' = M \setminus U_1$, $N' = N \setminus V_1$ puis observons que $\partial U_1 \subset M'$ et $\partial V_1 \subset N'$. On définit un espace topologique W comme le recollement de M' sur N' par f ,

$$W = M' \coprod_f N'$$

(i.e. on identifie $x \in \partial U_1$ à $f(x) \in \partial V_1$). On admettra la proposition suivante :

Proposition 2.3.3.1. (i) W est une m -variété topologique.

(ii) Si M et N sont compactes, alors W est compacte.

(iii) Si M et N sont connexes, alors W est connexe.

(iv) Si M et N sont des surfaces connexes, alors la classe d'homéomorphie de W ne dépend pas des cartes (U, ϕ) et (V, ψ) .

La propriété (iv) reste vraie pour des m -variétés orientées où une carte préserve l'orientation et l'autre carte la renverse. L'orientabilité est facile à définir pour des variétés différentiables (il existe un atlas où les changements de carte ont tous une jacobienne de déterminant positif) mais plus délicate pour des variétés topologiques. On ne le fera pas en toute généralité. Dans le cas des surfaces, on peut prendre comme définition : une surface est orientable si et seulement si elle ne contient pas de ruban de Moebius.

Définition 2.3.3.2 (SOMME CONNEXE). On appelle *somme connexe* de surfaces connexes M et N et on notera $M \# N$ toute surface homéomorphe à W .

Proposition 2.3.3.3 (PROPRIÉTÉS DE LA SOMME CONNEXE). Soit M_0, M_1, M_2 des surfaces connexes. On note $\cong$ la relation « être homéomorphe ».

(i) $M_0 \# M_1 \cong M_1 \# M_0$;

(ii) $(M_0 \# M_1) \# M_2 \cong M_0 \# (M_1 \# M_2)$;

(iii) $M \# \mathbf{S}^2 \cong M$.

Démonstration. (i) Soit (U, ϕ) et (V, ψ) des cartes comme ci-dessus pour M_0 et M_1 respectivement. Soit $f: \partial U_1 \rightarrow \partial V_1$ l'homéomorphisme associé. On réalise $M_1 \# M_0$ à l'aide des cartes (V, ψ) , (U, ϕ) et de l'homéomorphisme f^{-1} . On peut vérifier que pour $i = 0, 1$, l'application $(x, i) \mapsto [(x, 1 - i)]$

$$M'_0 \times \{0\} \cup M'_1 \times \{1\} \rightarrow M'_1 \times \{0\} \cup M'_0 \times \{1\} \rightarrow M'_1 \coprod_{f^{-1}} M'_0$$

se factorise en un homéomorphisme $[(x, i)] \mapsto [(x, 1 - i)]$ entre $M'_0 \coprod_f M'_1$ et $M'_1 \coprod_{f^{-1}} M'_0$.

(ii) On choisit des cartes (U, ϕ) et (V, ψ) de M_0 et M_1 définissant $M_0 \# M_1$, puis (U', ϕ') et (W, φ) de M_1 et M_2 définissant $M_1 \# M_2$ telle que $V \cap V' = \emptyset$. L'application adéquate se factorise en un homéomorphisme.

(iii) Soit (U, ϕ) une carte de M telle que $\phi(U) = \mathbf{R}^2$. Soit (V, ψ) une carte de $\mathbf{S}^2$ telle que $\psi(V) = \mathbf{R}^2$ et V_1 soit égale à l'hémisphère inférieur ouvert. On prend par exemple la projection stéréographique du pôle nord. Ainsi, $\mathbf{S}^{2'} = \mathbf{S}^2 \setminus V_1$ est l'hémisphère supérieur fermé. On note que ψ est l'identité en restriction à $\partial V_1 = \mathbf{B}(0, 1) \subset \mathbf{R}^m$. On définit une application $h: M' \coprod \mathbf{S}^{2'} \rightarrow M = M' \cup \overline{U_1}$, qui va se factoriser en un homéomorphisme $M \# \mathbf{S}^2 \rightarrow M$, comme suit. Sur M' , on pose $h(x) = x$. Sur $\mathbf{S}^{2'} = \{(u, t) \in \mathbf{S}^2 \subset \mathbf{R}^2 \times \mathbf{R}; t \geq 0\}$, on pose $h(u, t) = \phi^{-1}(u)$ de sorte que $h(\mathbf{S}^{2'}) = \overline{U_1}$ (car u décrit la boule fermée $\mathbf{B}(0, 1)$ quand (u, t) décrit $\mathbf{S}^{2'}$). On note que $h(x) = h(y)$ pour $x \in M'$ et $y \in \mathbf{S}^{2'}$ si et seulement si $h(x) \in \partial U_1$. Dans ce cas, $y \in \partial \mathbf{B}(0, 1)$ et $x = \phi^{-1}(y) = \phi^{-1} \circ \psi(y) = f^{-1}(y)$ i.e. $x \sim y$. On peut vérifier ensuite que la factorisation de h convient. $\square$

Exercice 2.3.3.4. Soit M, M', N, N' des surfaces connexes telles que $M \cong M'$ et $N \cong N'$. Montrer que $M \# N \cong M' \# N'$.

La propriété (iii) dit que $\mathbf{S}^2$ est le neutre de la somme connexe. On peut maintenant construire une infinité de surfaces compactes connexes en faisant des sommes connexes des briques $\mathbf{T}^2, \mathbf{P}^2 = \mathbf{P}_{\mathbf{R}}^2$ et $\mathbf{K}^2$ (et $\mathbf{S}^2$ mais celle là n'ajoute rien). Y-a-t-il d'autres surfaces que celles produites comme cela ? Le théorème de classification des surfaces dit que non ! Et qu'en plus, deux briques suffisent :

Théorème 2.3.3.5 (CLASSIFICATION DES SURFACES COMPACTES). Toute surface compacte et connexe est homéomorphe à l'une des surfaces suivantes :

(i) $\Sigma_g = \mathbf{T}^2 \# \cdots \# \mathbf{T}^2$ (g termes, $g \geq 0, \Sigma_0 = \mathbf{S}^2$);

(ii) $N_h = \mathbf{P}^2 \# \cdots \# \mathbf{P}^2$ (h termes, $h \geq 1$).

Ce théorème est démontré plus tard. Il dit que $\mathbf{K}^2$ n'est pas une brique fondamentale, puisqu'on peut la décomposer. Question : en quoi ?

Exercice 2.3.3.6 (EN TD). $\mathbf{K}^2 = \mathbf{P}^2 \# \mathbf{P}^2$.

On peut se demander ce qui se passe quand on mélange des $\mathbf{T}^2$ et des $\mathbf{P}^2$. Question : où est $\mathbf{T}^2 \# \mathbf{P}^2$? On verra la relation fondamentale $\mathbf{T}^2 \# \mathbf{P}^2 \cong \mathbf{P}^2 \# \mathbf{P}^2 \# \mathbf{P}^2$. On comprend alors que toute somme connexe qui comporte un $\mathbf{P}^2$ s'écrit comme somme connexe de $\mathbf{P}^2$. Question : les Σ_g et les N_h sont-elles toutes différentes ? Question plus basique : comment prouve-t-on que $\mathbf{S}^2, \mathbf{P}^2, \mathbf{T}^2$ et $\mathbf{K}^2$ ne sont pas homéomorphes ? À suivre...

3 Groupe fondamental

Dans ce chapitre, on apprend que :

- À X connexe par arc, on associe un groupe $\pi_1(X)$ appelé son groupe fondamental, qui d'une certaine manière, encode la structure de X . Si $\pi_1(X) = 0$, on dit que X est simplement connexe.
- Un homéomorphisme $X \rightarrow Y$ induit un isomorphisme $\pi_1(X) \rightarrow \pi_1(Y)$. La notion plus faible d'équivalence d'homotopie $X \rightarrow Y$ induira également un isomorphisme $\pi_1(X) \rightarrow \pi_1(Y)$.
- Si un groupe G agit de manière séparante et totalement discontinue sur X simplement connexe, alors $\pi_1(X/G) \simeq G$. Exemple fondamental : $\pi_1(\mathbf{S}^1) = \pi_1(\mathbf{R}/\mathbf{Z}) \simeq \mathbf{Z}$.
- Le théorème de Van Kampen exprime $\pi_1(U_1 \cup U_2)$ en termes de $\pi_1(U_1)$, $\pi_1(U_2)$ et $\pi_1(U_1 \cap U_2)$.

3.1 Groupe fondamental

On appelle *chemin* dans un espace topologique X une application $\gamma : [0, 1] \rightarrow X$ continue. Son origine est $\gamma(0)$ et son extrémité est $\gamma(1)$.

Définition 3.1.0.1 (HOMOTOPIES DE CHEMINS). Deux chemins γ_0, γ_1 dans X sont *équivalents*, ou *homotopes à extrémités fixes*, noté $\gamma_0 \sim \gamma_1$ s'il existe $H : [0, 1] \times [0, 1] \rightarrow X$ continue tel que :

$$\begin{aligned} H(t, 0) &= \gamma_0(t) & (\forall t \in [0, 1]) \\ H(t, 1) &= \gamma_1(t) & (\forall t \in [0, 1]) \\ H(0, s) &= \gamma_0(0) = \gamma_1(0) & (\forall s \in [0, 1]) \\ H(1, s) &= \gamma_0(1) = \gamma_1(1) & (\forall s \in [0, 1]) \end{aligned}$$

L'application H est une *homotopie* (à extrémités fixes) entre γ_0 et γ_1 .

Des chemins équivalents ont même origine et ont même extrémité. Si on note $\gamma_s(t) = H(t, s)$, on pense à $s \mapsto \gamma_s$ comme à une déformation continue de γ_0 en γ_1 .

Ainsi, tout chemin $\gamma_0 : [0, 1] \rightarrow \mathbf{R}^n$ joignant x à y est homotope au segment $\gamma_1(t) = x + t(y - x)$ via $\gamma_s(t) = (1 - s)\gamma_0(t) + s\gamma_1(t)$. Pour tout espace X et tout chemin $\gamma : [0, 1] \rightarrow X$, pour tout « reparamétrage » positif $\phi : [0, 1] \rightarrow [0, 1]$ continue telle que $\phi(0) = 0$ et $\phi(1) = 1$, on a $\gamma \circ \phi \sim \gamma$ via $H(t, s) = \gamma((1 - s)\phi(t) + st)$.

Lemme 3.1.0.2. L'homotopie à extrémités fixes définit une relation d'équivalence.

Démonstration. Soit F une homotopie entre deux chemins α et β alors $G(\cdot, s) = F(\cdot, 1 - s)$ est une homotopie de β à α . Soient α, β, γ des chemins tels que $\alpha \sim \beta$ et $\beta \sim \gamma$. Soit F une homotopie entre α et β et soit G une homotopie entre β et γ . On voit alors que

$$H(\cdot, s) = \begin{cases} F(\cdot, 2s) & (\forall s \in [0, 1/2]) \\ G(\cdot, 2(s - 1/2)) & (\forall s \in [1/2, 1]) \end{cases}$$

est bien définie car $F(\cdot, 1) = \beta(\cdot) = G(\cdot, 0)$ et est continue par recollement d'applications continues sur les fermés $[0, 1] \times [0, 1/2]$ et $[0, 1] \times [1/2, 1]$ (cf. théorème 1.1.4.6). On a $H(t, 0) = F(t, 0) = \alpha(t)$ et $H(t, 1) = G(t, 1) = \gamma(t)$. On a $H(0, s) = \alpha(0) = \beta(0) = \gamma(0)$ et $H(1, s) = \alpha(1) = \beta(1) = \gamma(1)$ pour tout $s \in [0, 1]$. Ainsi, H est une homotopie de α à γ . $\square$

Définition 3.1.0.3 (CONCATÉNATION DE CHEMINS). Soit α, β des chemins dans X tels que $\alpha(1) = \beta(0)$. On définit le *chemin concaténé* $\alpha\beta$: $[0, 1] \rightarrow X$, par

$$\begin{cases} \alpha(2t) & (t \in [0, 1/2]) \\ \beta(2(t - 1/2)) & (t \in [1/2, 1]) \end{cases}$$

Le chemin $\alpha\beta$ est continu comme recollement de deux applications continues sur les fermés $[0, 1/2]$ et $[1/2, 1]$.

Lemme 3.1.0.4 (LA CONCATÉNATION DE CHEMINS EST COMPATIBLE AVEC LA RELATION D'ÉQUIVALENCE). Soit α_0, β_0 des chemins tels que $\alpha_0\beta_0$ soit défini. Si $\alpha_0 \sim \alpha_1$ et $\beta_0 \sim \beta_1$, alors $\alpha_0\beta_0 \sim \alpha_1\beta_1$.

Démonstration. Soit F une homotopie entre α_0 et α_1 et G une homotopie entre β_0 et β_1 . En particulier, $\alpha_1(1) = \alpha_0(1) = \beta_0(0) = \beta_1(0)$, donc $\alpha_1\beta_1$ est défini. On voit alors que

$$\begin{cases} F(2t, s) & (t \in [0, 1/2]) \\ G(2(t - 1/2), s) & (t \in [1/2, 1]) \end{cases}$$

définit une homotopie de $\alpha_0\beta_0$ à $\alpha_1\beta_1$. $\square$

Lemme 3.1.0.5. La concaténation est associative pour les classes d'équivalence :

$$(\alpha\beta)\gamma \sim \alpha(\beta\gamma)$$

pour des chemins concaténables.

Démonstration. On peut vérifier que

$$(\alpha\beta)\gamma \sim (\alpha\beta)\gamma \circ \phi = \alpha(\beta\gamma)$$

où $\phi: [0, 1] \rightarrow [0, 1]$ affine par morceaux envoie $[0, 1/2]$ sur $[0, 1/4]$, $[1/2, 3/4]$ sur $[1/4, 1/2]$ et $[3/4, 1]$ sur $[1/2, 1]$. $\square$

Exercice 3.1.0.6. Étant donné γ un chemin dans X et $[a, b] \subset [0, 1]$ avec $a \leq b$, notons $\gamma_{ab} = \gamma \circ \phi_{ab}$ où $\phi_{ab}([0, 1]) = [a, b]$ et ϕ_{ab} est affine. Soit $a \in [0, 1]$. Montrer que $\gamma \sim \gamma_{0a}\gamma_{a1}$ (chercher ϕ tel que $\gamma \circ \phi = \gamma_{0a}\gamma_{a1}$).

Plus généralement, si $0 = a_0 \leq a_1 \leq \dots \leq a_N = 1$, on a

$$\gamma \sim \gamma_{a_0 a_1} \gamma_{a_1 a_2} \dots \gamma_{a_{N-1} a_N}.$$

Neutre et inverse : Pour tout $x \in X$ on note $e_x: [0, 1] \rightarrow X$ le chemin *constant* $e_x(t) = x$ pour tout $t \in [0, 1]$. Soit α un chemin, son chemin *inverse* α^{-1} est défini par $\alpha^{-1}(t) = \alpha(1 - t)$. On a $(\alpha\beta)^{-1} = \beta^{-1}\alpha^{-1}$.

Lemme 3.1.0.7. Pour tout chemin α de x à y , on a : (i) $\alpha e_y \sim \alpha \sim e_x \alpha$.

(ii) $\alpha\alpha^{-1} \sim e_x$.

(iii) $\alpha^{-1}\alpha \sim e_y$.

Démonstration. Exercice. $\square$

On a donc presque une loi de groupe. On appelle *lacet* basé en x un chemin de x à x . Sur l'ensemble $\mathcal{L}(X, x)$ des lacets basés en x , la concaténation de chemins est toujours possible.

Définition 3.1.0.8 (GROUPE FONDAMENTAL). On appelle *groupe fondamental* de X basé en $x \in X$ l'ensemble quotient $\mathcal{L}(X, x)/\sim$ munie de la loi de groupe $[\alpha][\beta] := [\alpha\beta]$. On le note $\pi_1(X, x)$.

Les considérations précédentes montrent que c'est bien un groupe. Le lacet constant $[e_x]$ est le neutre. Ce groupe ne dépend pas vraiment de x (si X est connexe par arc) :

Lemme 3.1.0.9. Si X est connexe par arc, pour tous $x, y \in X$, $\pi_1(X, x)$ est isomorphe à $\pi_1(X, y)$.

Démonstration. Puisque X est connexe par arc, il existe un chemin γ de x à y . On définit $\gamma_{\#}: \pi_1(X, x) \rightarrow \pi_1(X, y)$ par $[\alpha] \mapsto [\gamma^{-1}\alpha\gamma]$. C'est un morphisme de groupe :

$$\gamma_{\#}[\alpha\beta] = [\gamma^{-1}\alpha\beta\gamma] = [\gamma^{-1}\alpha\gamma\gamma^{-1}\beta\gamma] = [\gamma^{-1}\alpha\gamma][\gamma^{-1}\beta\gamma] = \gamma_{\#}[\alpha]\gamma_{\#}[\beta].$$

Sa réciproque est $(\gamma^{-1})_{\#}: \pi_1(X, y) \rightarrow \pi_1(X, x)$ puisque

$$(\gamma^{-1})_{\#}(\gamma_{\#}[\alpha]) = [\gamma\gamma^{-1}\alpha\gamma\gamma^{-1}] = [\alpha].$$

$\square$

L'isomorphisme $\gamma_{\#}$ dépend en général de la classe d'homotopie de γ .

Exemple 3.1.0.10. On a $\pi_1(\mathbf{R}^n, 0) = 0$. En effet, tout lacet α basé en 0 est homotope au lacet constant via l'homotopie $H(t, s) = s\alpha(t)$. De même, $\pi_1(U, x) = 0$ pour toute partie convexe (ou même étoilée) U d'un espace vectoriel normé.

On rappelle que le produit direct $G \times H$ de deux groupes est l'ensemble $\{(g, h); (g \in G) (h \in H)\}$ muni de la loi $(g, h) \cdot (g', h') = (gg', hh')$.

Proposition 3.1.0.11. On a $\pi_1(X \times Y, (x, y)) = \pi_1(X, x) \times \pi_1(Y, y)$.

Démonstration. Exercice. $\square$

3.2 Espace simplement connexe

Définition 3.2.0.1 (SIMPLE CONNEXITÉ). Soit X connexe par arc. On dit que X est *simplement connexe* si $\pi_1(X, x) = 0$.

Autrement dit, si tout lacet basé en x est homotopiquement trivial, cela ne dépend pas de x . Intuitivement, $\mathbf{S}^2$ est simplement connexe alors que $\mathbf{S}^1$ ne l'est pas et $\mathbf{T}^2$ encore moins (en fait, on verra que $\pi_1(\mathbf{S}^1) = \mathbf{Z}$ et $\pi_1(\mathbf{T}^2) = \mathbf{Z}^2$). Vérification à venir...

Exercice 3.2.0.2. Montrer que sont équivalents :

- (1) X est simplement connexe.
- (2) Pour tous chemins α, β de X tel que $\alpha(0) = \beta(0)$ et $\alpha(1) = \beta(1)$ alors $\alpha \sim \beta$.
- (3) Toute application continue $f: \mathbf{S}^1 \rightarrow X$ prolonge continûment $\mathbf{D}^2 \rightarrow X$.

Solution : (1) $\Rightarrow$ (2). Soit $x = \alpha(0)$ et $y = \alpha(1)$. Le chemin $\alpha\beta^{-1}$ est un lacet en x , donc $e_x \sim \alpha\beta^{-1}$. Ceci implique

$$\beta \sim e_x \sim \alpha\beta^{-1}\beta \sim \alpha e_y \sim \alpha.$$

(2) $\Rightarrow$ (1). Il suffit de prendre $\alpha(0) = \alpha(1)$ et β le lacet constant.

(1) $\Rightarrow$ (3). Étant donné $f: \mathbf{S}^1 \rightarrow X$ continue on définit un lacet $\alpha: [0, 1] \rightarrow X$ par $\alpha(t) = f(e^{2i\pi t})$. Par hypothèse, il existe $H: [0, 1] \times [0, 1] \rightarrow X$ une homotopie de α au lacet constant $x = \alpha(0) = \alpha(1)$. On veut définir $F: \mathbf{D}^2 \rightarrow X$ par $F(se^{2i\pi t}) = H(t, 1-s)$. C'est bien défini en $se^{2i\pi 0} = se^{2i\pi}$ car $H(0, 1-s) = x = H(1, 1-s)$ pour tout $s \in [0, 1]$. C'est bien défini en $0 = 0 \times e^{2i\pi t} \in \mathbf{D}^2$ car $H(t, 1) = x$ pour tout $t \in [0, 1]$. On vérifie la continuité de F grâce au diagramme

$$\begin{array}{ccc} [0, 1] \times [0, 1] & \xrightarrow{H} & X \\ \rho \searrow & & \nearrow F \\ & se^{2i\pi t} \in \mathbf{D}^2 & \end{array}$$

commutatif suivant où $\rho(t, s) := se^{2i\pi t}$. Montrons d'abord que $U \subset \mathbf{D}^2$ est fermé si et seulement si $\rho^{-1}(U)$ est fermé. Si $U \subset \mathbf{D}^2$ est fermé, alors $\rho^{-1}(U)$ est fermé par continuité de ρ . Réciproquement, si $\rho^{-1}(U)$ est fermé, alors $\rho^{-1}(U)$ est compact dans $[0, 1] \times [0, 1]$ compact dans son image continue $\rho(\rho^{-1}(U))$ est compact dans $\mathbf{D}^2$ donc fermé. Mais ρ est surjective, donc $\rho(\rho^{-1}(U)) = U$, d'où la conclusion. On sait alors que F est continue si et seulement si $F \circ \rho = H$ est continue, d'où la conclusion. (Soit $V \subset X$ un fermé alors $F^{-1}(V)$ est fermé si $\rho^{-1}(F^{-1}(V))$ est fermé, or $\rho^{-1}(F^{-1}(V)) = H^{-1}(V)$ est fermé par continuité de H)

(3) $\Rightarrow$ (1). On pose $H(t, s) = F((s-1)e^{2i\pi t} + s)$.

Théorème 3.2.0.3 (PETIT THÉORÈME DE VAN KAMPEN). Soit $X = U_1 \cup U_2$ connexe par arc tel que U_1 et U_2 sont ouverts, simplement connexes et $U_1 \cap U_2$ est connexe par arc. L'espace X est alors simplement connexe.

C'est un cas particulier du théorème de Van Kampen, qui calcule le groupe fondamental de $U_1 \cup U_2$, en fonction des groupes fondamentaux de U_1, U_2 et $U_1 \cap U_2$. Le petit théorème de Van Kampen découle immédiatement du lemme suivant.

Lemme 3.2.0.4. Soit $X = U_1 \cup U_2$ connexe par arc tel que U_1 et U_2 sont ouverts, connexes par arc et $U_1 \cap U_2$ est connexe par arc. Si $x \in U_1 \cap U_2$, alors tout lacet γ basé en x est homotope à une concaténation de lacets d'origine x , chacun contenu dans U_1 ou dans U_2 .

Démonstration. Si $\gamma \subset U_1$ ou $\gamma \subset U_2$, la conclusion est vérifiée. Supposons $\gamma \not\subset U_i$ pour tout $i = 1, 2$. On construit d'abord une subdivision $0 = t_0 < t_1 < \dots < t_N = 1$ de $[0, 1]$ tel que chaque $[t_i, t_{i+1}]$ est envoyé par γ dans un des U_j . En effet, par continuité de γ chaque $t \in [0, 1]$ a un voisinage V_t tel que $\gamma(V_t)$ est contenu dans un des U_j . Par compacité, $[0, 1]$ est recouvert par un nombre fini de tels V_t et les extrémités des V_t donnent la subdivision adéquate. Quitte à supprimer des t_k on peut supposer que si $\gamma([t_i, t_{i+1}]) \subset U_j$, alors $\gamma([t_{i+1}, t_{i+2}]) \not\subset U_j$. Ainsi, $\gamma(t_i) \in U_1 \cap U_2$ pour tout i . Pour $i = 0, \dots, N-1$, notons γ_i le chemin $\gamma_{t_i t_{i+1}}$ (obtenu en reparamétrisant sur $[0, 1]$ la restriction $\gamma|_{[t_i, t_{i+1}]}$). Soit c_i un chemin dans $U_1 \cap U_2$ de $\gamma(t_{i+1})$ à x . On a alors

$$\gamma \sim (\gamma_0 c_0)(c_0^{-1} \gamma_1 c_1)(c_1^{-1} \gamma_2 c_2) \cdots (c_{N-2}^{-1} \gamma_{N-1} c_{N-1})(c_{N-1}^{-1} \gamma_N)$$

une concaténation de lacets basés en x , chacun contenu dans U_1 ou dans U_2 . □

Ce résultat peut être généralisé à $X = \bigcup U_i$ si les U_i sont simplement connexes, $U_i \cap U_j$ est connexe par arc pour tout i, j et il existe x appartenant à tous les U_i .

Corollaire 3.2.0.5. Pour tout entier $n \geq 2$, la sphère $\mathbf{S}^n$ est simplement connexe.

Démonstration. Soit S, N les pôles sud et nord de $\mathbf{S}^n$. On pose $U_1 = \mathbf{S}^n \setminus \{S\}$ et $U_2 = \mathbf{S}^n \setminus \{N\}$. Ainsi, U_1 et U_2 sont homéomorphes à $\mathbf{R}^n$ donc simplement connexes. De plus, $U_1 \cap U_2$ est connexe par arc si $n \geq 2$ (mais pas pour $n = 1$). On applique le théorème 3.2.0.3. □

3.3 Effet d'applications continues

Définition 3.3.0.1 (MORPHISME INDUIT). Soit $f: X \rightarrow Y$ continue entre deux espaces connexes par arc. Pour $x_0 \in X$, on définit un morphisme de groupes dit *morphisme induit* par f , noté f_* : pour tout $[\alpha] \in \pi_1(X, x_0)$,

$$\begin{aligned} f_*: \pi_1(X, x_0) &\rightarrow \pi_1(Y, f(x_0)) \\ [\alpha] &\mapsto [f \circ \alpha] \end{aligned}$$

L'application est bien définie car $\alpha \sim \beta \Rightarrow f \circ \alpha \sim f \circ \beta$. C'est bien un morphisme de groupes car

$$\begin{aligned} f_*([\alpha][\beta]) &= f_*([\alpha\beta]) = [f \circ (\alpha\beta)] \\ &= [(f \circ \alpha)(f \circ \beta)] = [f \circ \alpha][f \circ \beta] \\ &= f_*([\alpha])f_*([\beta]) \end{aligned}$$

On l'appelle morphisme induit par f . On voit immédiatement que $(\text{Id}_X)_* = \text{Id}_{\pi_1(X, x)}$.

Remarque 3.3.0.2. Attention ! Cependant, f injective n'implique pas f_* injective (considérer $\mathbf{S}^1 \rightarrow \mathbf{R}^2$) et f surjective n'implique pas f_* surjective (considérer $\mathbf{R} \rightarrow \mathbf{S}^1$).

Définition 3.3.0.3. Deux applications continues $f_0, f_1: X \rightarrow Y$ sont *homotopes* ($f_0 \sim f_1$) s'il existe $H: X \times [0, 1] \rightarrow Y$ continue telle que $H(\cdot, 0) = f_0(\cdot)$ et $H(\cdot, 1) = f_1(\cdot)$. On appelle H une *homotopie* de f_0 à f_1

Si on note $f_s(t) = H(t, s)$, on pense à $s \mapsto f_s$ comme à une déformation continue de f_0 en f_1 . Précisément :

Exercice 3.3.0.4. On munit l'espace $\mathcal{C}(X, Y)$ des fonctions continues $X \rightarrow Y$ de la topologie compacte ouverte, engendrée par les

$$V(K, U) = \{f \in \mathcal{C}(X, Y); f(K) \subset U\}$$

où $K \subset X$ parcourt les compacts et $U \subset Y$ les ouverts. On suppose X localement compact. Montrer qu'une homotopie H est continue si et seulement si $s \mapsto f_s \in \mathcal{C}(X, Y)$ est continue.

On pourra noter simplement $f_s: X \rightarrow Y$ une homotopie où $f_s = H(\cdot, s)$.

Proposition 3.3.0.5 (PROPRIÉTÉS DU MORPHISME INDUIT). Soit X, Y, Z connexes par arc et soit $x_0 \in X$.

(i) Si $f: X \rightarrow Y$ et $g: Y \rightarrow Z$ continue alors

$$(g \circ f)_* = g_* \circ f_*.$$

(ii) Soit $f_s: X \rightarrow Y$ une homotopie et γ le chemin $s \mapsto f_s(x_0)$ parcourue par l'image d'un point base $x_0 \in X$, alors $(f_1)_* = \gamma_{\#} \circ (f_0)_*$

$$\begin{array}{ccc} & \pi_1(Y, f_0(x_0)) & \\ (f_0)_* \nearrow & \downarrow \gamma_{\#} \simeq & \searrow (f_1)_* \\ \pi_1(X, x_0) & & \pi_1(Y, f_1(x_0)) \end{array}$$

Remarque 3.3.0.6. (1) Le point (i) peut s'appliquer en particulier à un homéomorphisme $f: X \rightarrow Y$ et $f^{-1}: Y \rightarrow X$ tel que

$$f_* \circ (f^{-1})_* = (f \circ f^{-1})_* = (\text{Id})_* = \text{Id} = (f^{-1})_* \circ f_*$$

donc que $f_*: \pi_1(X, x) \rightarrow \pi_1(Y, f(x))$ est un isomorphisme de réciproque $(f^{-1})_*$. Deux espaces homéomorphes ont des groupes fondamentaux isomorphes.

(2) Le point (ii) dit que deux applications homotopes induisent le même morphisme modulo un isomorphisme prenant en compte le changement éventuel de l'image du point base d'une application à l'autre. Si $f_0(x_0) = f_1(x_0)$ et l'homotopie H fixe cette image (i.e. $H(x_0, s) = f_0(x_0) = f_s(x_0)$ pour tout $s \in [0, 1]$) on dit que f_s est une homotopie relative à $\{x_0\}$, alors $(f_0)_* = (f_1)_*$. Deux applications homotopes relativement au point base induisent le même morphisme.

Démonstration. (i) Soit $[\alpha] \in \pi_1(X, x_0)$, on a alors

$$(g \circ f)_*([\alpha]) = [g \circ f \circ \alpha] = g_*([f \circ \alpha]) = g_*(f_*([\alpha])).$$

(ii) Soit $[\alpha] \in \pi_1(X, x_0)$, il s'agit de montrer que

$$(f_1)_*([\alpha]) = [f_1 \circ \alpha] = \gamma_{\#} \circ (f_0) \circ ([\alpha]) = [\gamma^{-1}(f_0 \circ \alpha)\gamma]$$

autrement dit que

$$f_1 \circ \alpha \sim \gamma^{-1}(f_0 \circ \alpha)\gamma$$

soit aussi

$$\gamma(f_1 \circ \alpha)\gamma^{-1} \sim f_0 \circ \alpha. \quad (3.1)$$

On construit une homotopie entre ces lacets fixant le point base $f_0(x_0)$ comme suit. Notons que $f_s \circ \alpha$ est un lacet basé en $f_s(x_0) = \gamma(s)$. Soit γ_s le chemin de $f_0(x_0)$ à $f_s(x_0)$ obtenu en reparamétrisant $\gamma|_{[0,s]}$ sur $[0,1]$: explicitement, $\gamma_s(t) = \gamma(st)$. En particulier, γ_0 est constant égal à $e_{f_0(x_0)}$ et $\gamma_1 = \gamma$. Ainsi, les concaténations

$$s \mapsto \gamma_s(f_s \circ \alpha)\gamma_s^{-1}$$

définissent une homotopie de lacets (voir continuité ci-dessous). En $s = 0$, on a

$$e_{f_0(x_0)}(f_0 \circ \alpha)e_{f_0(x_0)}^{-1} \sim f_0 \circ \alpha$$

et en $s = 1$, on a exactement $\gamma(f_1 \circ \alpha)\gamma^{-1}$. $\square$

Lemme 3.3.0.7 (CONTINUITÉ DE LA CONCATÉNATION À PARAMÈTRES). Soit X un espace topologique et soit $(t, s) \in [0, 1]^2 \mapsto \alpha_s(t), \beta_s(t) \in X$ continues tel que $\alpha_s(1) = \beta_s(0)$ pour tout $s \in [0, 1]$, alors

$$(t, s) \mapsto (\alpha_s \beta_s)(t)$$

est continue.

Démonstration. On a

$$(\alpha_s \beta_s)(t) = \begin{cases} \alpha_s(2t) & (t \in [0, 1/2]) \\ \beta_s(2(t - 1/2)) & (t \in [1/2, 1]) \end{cases}$$

qui est continue par recollement d'applications continues sur les fermés $[0, 1/2] \times [0, 1]$ et $[1/2, 1] \times [0, 1]$ (sur l'intersection $\{1/2\} \times [0, 1]$, on a $\alpha_s \beta_s(1/2) = \alpha_s(1) = \beta_s(0)$). $\square$

Dans beaucoup de situations, il n'est pas nécessaire d'avoir un homéomorphisme $X \rightarrow Y$ pour conclure que $\pi_1(X) \simeq \pi_1(Y)$. La notion plus malléable d'équivalence d'homotopie suffit.

Définition 3.3.0.8 (ÉQUIVALENCE D'HOMOTOPIE). On dit que $f: X \rightarrow Y$ continue est une *équivalence d'homotopie* s'il existe $g: Y \rightarrow X$ continue telle que $f \circ g \sim \text{Id}_Y$ et $g \circ f \sim \text{Id}_X$. On dit alors que X et Y ont même type d'homotopie.

En quelque sorte g est une réciproque de f à homotopie près. Cela ne demande pas que f soit injective ou surjective.

Exemple 3.3.0.9. L'inclusion $f: x \in \mathbf{S}^1 \mapsto x \in \mathbf{R}^2 \setminus \{0\}$ est une équivalence d'homotopie, sa réciproque à homotopie près étant $g(x) = \frac{x}{\|x\|}$. En effet, on a trivialement $g \circ f = \text{Id}_{\mathbf{S}^1}$. Réciproquement, $f \circ g(x) = \frac{x}{\|x\|}$ est homotope à l'identité $\mathbf{R}^2 \setminus \{0\} \rightarrow \mathbf{R}^2 \setminus \{0\}$ via $H(x, s) = sx + (1-s)\frac{x}{\|x\|}$.

Proposition 3.3.0.10 (MÊME TYPE D'HOMOTOPIE $\Rightarrow \pi_1$ ISOMORPHES). Si $f: X \rightarrow Y$ est une équivalence d'homotopie, alors $f_*: \pi_1(X, x) \rightarrow \pi_1(Y, f(x))$ est un isomorphisme. Par conséquent, des espaces ayant même type d'homotopie ont des groupes fondamentaux isomorphes.

Ainsi, $\pi_1(\mathbf{S}^1) \simeq \pi_1(\mathbf{R}^2 \setminus \{0\})$.

Démonstration. Soit $g: Y \rightarrow X$ tel que $f \circ g \sim \text{Id}_Y$ et $g \circ f \sim \text{Id}_X$. On déduit de la proposition 3.3.0.5 que $g_* \circ f_* = (g \circ f)_* = \gamma_{\#} \circ (\text{Id}_X)_*$ (où γ est un chemin de x à $g(f(x))$) qui est un isomorphisme, donc f_* est injective. Le même raisonnement avec $f \circ g$ prouve que f_* est surjective. $\square$

Exemple 3.3.0.11. (1) La formule dans l'exemple 3.3.0.9 montre que $\mathbf{S}^n$ et $\mathbf{R}^{n+1} \setminus \{0\}$ ont même type d'homotopie. On déduit de la simple connexité de $\mathbf{S}^n$ pour $n \geq 2$ que $\pi_1(\mathbf{R}^{n+1} \setminus \{0\}) = 0$ si $n \geq 2$.

(2) L'inclusion $f: x \in \mathbf{S}^1 \mapsto (x, 0) \in \mathbf{S}^1 \times \mathbf{R}$ est une équivalence d'homotopie. Ainsi, le cylindre et le cercle ont même type d'homotopie.

(3) Le ruban de Moebius a le type d'homotopie d'un cercle (exercice).

(4) Un espace X est *contractile* s'il a le type d'homotopie d'un point. C'est en particulier le cas d'un espace convexe. Si X est contractile, alors $\pi_1(X) = 0$.

Exercice 3.3.0.12. Montrer qu'une couronne

$$C = \{z \in \mathbf{C}; r_1 < |z| < r_2\}$$

a le type d'homotopie d'un cercle.

Exercice 3.3.0.13. Montrer que le demi-plan fermé $\mathbf{R} \times [0, \infty[$ n'est pas homéomorphe à $\mathbf{R}^2$ (en utilisant $\pi_1(\mathbf{S}^1) \neq 0$).

3.4 Homotopie et revêtement

On rappelle que $p: \tilde{X} \rightarrow X$ continue surjective est un revêtement si pour tout $x \in X$, il existe un voisinage U de x tel que $p^{-1}(U)$ soit une union disjointe d'ouverts U_i (avec $i \in I$) et chaque $p: U_i \rightarrow U$ est un homéomorphisme. On appelle un tel U un voisinage élémentaire et un U_i un *relevé* de U .

Définition 3.4.0.1 (RELEVÉ D'UNE APPLICATION). Si $f: Y \rightarrow X$ est une application continue, on appelle *relevé* de f une application $\tilde{f}: Y \rightarrow \tilde{X}$ telle que $p \circ \tilde{f} = f$.

$$\begin{array}{ccc} & \tilde{X} & \\ \tilde{f} \nearrow & & \downarrow p \\ Y & \xrightarrow{f} & X \end{array}$$

Il est clair que si deux applications $\tilde{f}, \tilde{g}: Y \rightarrow \tilde{X}$ sont homotopes, leur composée $p \circ \tilde{f}, p \circ \tilde{g}: Y \rightarrow X$ sont homotopes (composer p avec l'homotopie). Une propriété fondamentale des revêtements est qu'on peut souvent faire l'opération inverse : relever deux applications f, g homotopes en des relevés $\tilde{f}, \tilde{g}$ homotopes. C'est le cas en particulier pour des chemins et leurs homotopies.

Théorème 3.4.0.2 (RELÈVEMENT DES CHEMINS ET HOMOTOPIES). Soit $p: \tilde{X} \rightarrow X$ un revêtement, $x_0 \in X$ et $\tilde{x}_0 \in p^{-1}(x_0)$.

(i) Pour tout chemin $\gamma: [0, 1] \rightarrow X$ d'origine x_0 , il existe un unique chemin relevé $\tilde{\gamma}: [0, 1] \rightarrow \tilde{X}$ d'origine $\tilde{x}_0$.

$$\begin{array}{ccc} & \tilde{X} & \\ \tilde{\gamma} \nearrow & & \downarrow p \\ [0, 1] & \xrightarrow{\gamma} & X \end{array}$$

(ii) Pour toute homotopie $H: [0, 1]^2 \rightarrow X$ de chemins d'origine x_0 , il existe une unique homotopie relevée $\tilde{H}: [0, 1]^2 \rightarrow \tilde{X}$ de chemins d'origine $\tilde{x}_0$.

$$\begin{array}{ccc} & \tilde{X} & \\ \tilde{H} \nearrow & & \downarrow p \\ [0, 1]^2 & \xrightarrow{H} & X \end{array}$$

Cela dit qu'étant donnés deux chemins homotopes α, β d'origine x_0 , leur unique relevés $\tilde{\alpha}, \tilde{\beta}$ d'origine $\tilde{x}_0$ sont homotopes. Le théorème 3.4.0.2 découle du résultat plus général suivant :

Théorème 3.4.0.3. Soit $p: \tilde{X} \rightarrow X$ un revêtement. Étant donné $F: Y \times [0, 1] \rightarrow X$ continue et $\tilde{F}_0: Y \times \{0\} \rightarrow \tilde{X}$ un relevé de $F|_{Y \times \{0\}}$, il existe un unique relevé continue $\tilde{F}: Y \times [0, 1] \rightarrow \tilde{X}$ de F égal à $\tilde{F}_0$ en restriction à $Y \times \{0\}$.

$$\begin{array}{ccc} & \tilde{X} & \\ \tilde{F}_0 \nearrow & & \downarrow p \\ Y \times \{0\} & \xrightarrow{F|_{Y \times \{0\}}} & X \end{array} \quad \begin{array}{ccc} & \tilde{X} & \\ \tilde{F} \nearrow & & \downarrow p \\ Y \times [0, 1] & \xrightarrow{F} & X \end{array}$$

La partie (i) du théorème 3.4.0.2 suit de 3.4.0.3 appliqué au point $Y = \{x_0\}$, $F(y_0, t) = \gamma(t)$ et $\tilde{F}_0(x_0, 0) = \tilde{x}_0$. Pour la partie (ii), on pose $\gamma(t) = H(t, 0)$. Soit $\tilde{\gamma}$ le chemin relevé obtenu en (i), alors on prend $F = H$ et $\tilde{F}_0 = \tilde{\gamma}$. On obtient un unique relevé $\tilde{H} := \tilde{F}$ de $F = H$. Les restrictions de $\tilde{F}$ à $\{0\} \times [0, 1]$ et $\{1\} \times [0, 1]$ sont des relevés de chemins constants $s \mapsto H(0, s)$ et $s \mapsto H(1, s)$ donc sont constants par unicité du (i). Ainsi, $\tilde{H}$ est bien une homotopie de chemins ($\tilde{H}$ fixe origines et extrémités).

Avant de prouver le théorème 3.4.0.3, on définit la notion d'ouvert élémentaire et on démontre une propriété d'unicité des relèvements.

Définition 3.4.0.4 (OUVERT ÉLÉMENTAIRE). Soit X un espace topologique. Soit $U \subset X$ un ouvert. On dit que U est un *ouvert élémentaire* si $p^{-1}(U) = \bigcup_{i \in I} U_i$ avec les U_i ouverts disjoints et les $p: U_i \rightarrow U$ sont des homéomorphismes.

Proposition 3.4.0.5 (PROPRIÉTÉ DU RELÈVEMENT UNIQUE). Soit $p: \tilde{X} \rightarrow X$ un revêtement, $f: Y \rightarrow X$ continue et deux relevés $\tilde{f}_0, \tilde{f}_1: Y \rightarrow \tilde{X}$ de f coïncidant en un point $y_0 \in Y$. Si Y est connexe, alors $\tilde{f}_0 = \tilde{f}_1$.

Démonstration. Montrons que $A := \{y \in Y; \tilde{f}_0(y) = \tilde{f}_1(y)\}$ est ouvert et fermé. Par hypothèse, A est non vide. Soit $y \in A$, $x = f(y) \in X$ et $\tilde{x} = \tilde{f}_0(y) = \tilde{f}_1(y)$. Soit U un voisinage élémentaire contenant x , $\tilde{U}$ un relevé de U

contenant $\tilde{x}$. Par continuité des $\tilde{f}_i$, il existe un voisinage V de y tel que $f_i(V) \subset \tilde{U}$ pour $i = 0, 1$. Comme $p|_{\tilde{U}}$ est injective, $p \circ f_0 = f = p \circ f_1$ sur V implique que $\tilde{f}_0 = \tilde{f}_1$ sur V . En outre, $V \subset A$ et A est ouvert. Soit $y \notin A$, alors $\tilde{f}_0(y) \neq \tilde{f}_1(y)$. Soit U un voisinage élémentaire de $f(y)$, alors il existe des relevés $\tilde{U}_i$ de U contenant les $\tilde{f}_i(y)$. Nécessairement, $\tilde{U}_0 \neq \tilde{U}_1$, donc $\tilde{U}_0 \cap \tilde{U}_1 = \emptyset$. Par continuité des $\tilde{f}_i$, il existe un voisinage V de y tel que $\tilde{f}_i(V) \subset \tilde{U}_i$. En particulier, $\tilde{f}_0(z) \neq \tilde{f}_1(z)$ pour tout $z \in V$, prouvant que $V \subset Y \setminus A$ et que A est fermé. $\square$

Démonstration du théorème 3.4.0.3. On commence par construire, pour tout $y \in Y$, un relevé $\tilde{F}: V_y \times [0, 1] \rightarrow X$ pour un certain voisinage $V_y \subset Y$ de y . Fixons y . Par continuité de F , chaque (y, s) a un voisinage $V_s \times (a_s, b_s)$ tel que $F(V_s \times (a_s, b_s))$ est contenu dans un voisinage élémentaire de $F(y, s)$. Par compacité de $\{y\} \times [0, 1]$, on peut recouvrir $\{y\} \times [0, 1]$ par un nombre fini de tels voisinages. On peut donc trouver un voisinage V de y et $0 = s_0 < s_1 < \dots < s_N = 1$ tel que chaque $F(V \times [s_i, s_{i+1}])$ soit contenu dans un voisinage élémentaire. Supposons $\tilde{F}$ construit sur $V \times [0, s_i]$, coïncidant avec $\tilde{F}_0$ sur $V \times \{0\}$. Soit U un voisinage élémentaire contenant $F(V \times [s_i, s_{i+1}])$. Il existe un unique relevé $\tilde{U}$ de U , homéomorphe à U via p , contenant $\tilde{F}(y, s_i)$. Quitte à rapetisser V , on peut supposer que $\tilde{F}(V \times \{s_i\}) \subset \tilde{U}$. On définit $\tilde{F} = p|_{\tilde{U}}^{-1} \circ F$ sur $V \times [s_i, s_{i+1}]$. Après un nombre fini d'étapes, on obtient un relevé $\tilde{F}: V_y \times [0, 1] \rightarrow \tilde{X}$. La proposition 3.4.0.5 implique l'unicité du relevé $\tilde{F}$ sur $\{z\} \times [0, 1]$ pour $z \in V_y$. Deux relevés ainsi construits sur $V_y \times [0, 1]$ et $V_{y'} \times [0, 1]$ coïncident sur $\{z\} \times [0, 1]$ pour tout $y \in V \cap V'$, donc sur $(V \cap V') \times [0, 1]$. Il s'ensuit qu'il existe un relevé $\tilde{F}: Y \rightarrow \tilde{X}$, continu puisque continu sur chaque ouvert $V_y \times [0, 1]$ et unique puisque unique sur chaque $\{y\} \times [0, 1]$. $\square$

Corollaire 3.4.0.6. Soit $p: \tilde{X} \rightarrow X$ un revêtement, $x \in X$ et $\tilde{x}_0 \in p^{-1}(x_0)$. Le morphisme induit $p_*: \pi_1(\tilde{X}, \tilde{x}_0) \rightarrow \pi_1(X, x_0)$ est injectif.

Démonstration. Soit $\tilde{\alpha}$ un lacet de $\tilde{X}$ d'origine $\tilde{x}_0$ tel que $p_*([\tilde{\alpha}]) = 0$. Ceci signifie que $p \circ \tilde{\alpha} = \alpha$ est un lacet de X homotope à un lacet constant. En relevant une telle homotopie, on obtient une homotopie de $\tilde{\alpha}$ vers un lacet constant, i.e. $[\tilde{\alpha}] = 0$. $\square$

Corollaire 3.4.0.7. Soit $p: \tilde{X} \rightarrow X$ un revêtement, $x_0 \in X$ et $\tilde{x}_0 \in p^{-1}(x_0)$. On suppose $\tilde{X}$ simplement connexe. L'application

$$d: \pi_1(X, x_0) \rightarrow p^{-1}(x_0) \\ [\alpha] \mapsto \tilde{\alpha}(1)$$

où $\tilde{\alpha}$ est l'unique relevé de α d'origine $\tilde{x}_0$, est alors une bijection.

Démonstration. L'application d est bien définie car $p \circ \tilde{\alpha}(1) = \alpha(1) = x_0$ et si $\beta \sim \alpha$, alors une homotopie de β à α se relève à partir de $\tilde{x}_0$ en une homotopie entre $\tilde{\beta}$ et $\tilde{\alpha}$. En particulier, $\tilde{\beta}(1) = \tilde{\alpha}(1)$. L'application est surjective : étant donné $\tilde{x} \in p^{-1}(x_0)$, on se donne un chemin γ dans $\tilde{X}$ de $\tilde{x}_0$ à $\tilde{x}$. On a alors que $\alpha = p \circ \gamma$ est un lacet en x_0 et son relevé $\tilde{\alpha}$ à partir de $\tilde{x}_0$ est égal à γ par unicité du relevé. Ainsi, $d([\alpha]) = \tilde{\alpha}(1) = \gamma(1) = \tilde{x}$. Observons que $[\alpha]$ ne dépend pas du chemin γ choisi. En effet, soit β un autre chemin de $\tilde{x}_0$ à $\tilde{x}$. Par simple connexité de $\tilde{X}$, on a $\gamma \sim \beta$, donc $p \circ \gamma \sim p \circ \beta$ i.e. $[\alpha] = [p \circ \gamma] = [p \circ \beta]$. L'application $\tilde{x} \in p^{-1}(x_0) \mapsto \pi_1(X, x_0)$ ainsi définie est la réciproque de d , prouvant que d est injective. $\square$

Théorème 3.4.0.8 ($\pi_1(\tilde{X}/G, x_0) \simeq G$). Soit $\tilde{X}$ simplement connexe et G un groupe agissant par homéomorphisme sur $\tilde{X}$ de manière totalement discontinue et séparante. Soit $x_0 \in X := \tilde{X}/G$, alors $\pi_1(X, x_0)$ est isomorphe à G .

Démonstration. Notons $p: \tilde{X} \rightarrow \tilde{X}/G = X$ la projection, c'est un revêtement. Fixons $\tilde{x}_0 \in p^{-1}(x_0)$. On construit un morphisme $\phi: \pi_1(X, x_0) \rightarrow G$. D'après 3.4.0.7, on a une bijection $d: \pi_1(X, x_0) \rightarrow p^{-1}(x_0); [\alpha] \mapsto \tilde{\alpha}(1)$ où $\tilde{\alpha}$ est l'unique relevé de α d'origine $\tilde{x}_0$. Puisque G agit de manière totalement discontinue, il agit sans point fixe. À $\tilde{x} \in X$ fixé, l'application $g \in G \mapsto g\tilde{x} \in G\tilde{x}$ est une bijection. Appliquant cela à $\tilde{x}_0$ et $G\tilde{x}_0 = p^{-1}(x_0)$, on associe à $\tilde{\alpha}(1)$ l'unique $g \in G$ tel que $\tilde{\alpha}(1) = g\tilde{x}_0$. Ainsi, $\phi([\alpha])$ est l'unique $g \in G$ tel que $\tilde{\alpha}(1) = g\tilde{x}_0$. C'est une bijection, reste à voir que c'est un morphisme de groupe. Soit α_1, α_2 des lacets en x_0 , $\tilde{\alpha}_1$ et $\tilde{\alpha}_2$ leur relevés à partir de $\tilde{x}_0$. Soit $g_1 = \phi([\alpha_1])$ i.e. $g_1\tilde{x}_0 = \tilde{\alpha}_1(1)$ et $g_2 = \phi([\alpha_2])$ i.e. $g_2\tilde{x}_0 = \tilde{\alpha}_2(1)$. Observons que $g_1 \circ \tilde{\alpha}_2$ est un chemin d'origine $g_1\tilde{x}_0 = \tilde{\alpha}_1(1)$, donc le chemin concaténé $\tilde{\alpha}_1(g_1 \circ \tilde{\alpha}_2)$ est bien défini. C'est un relevé du lacet $\alpha_1\alpha_2$, l'unique d'origine $\tilde{x}_0$. Son extrémité est $g_1(\tilde{\alpha}_2(1)) = g_1g_2\tilde{x}_0$. On voit donc que

$$\phi([\alpha_1\alpha_2])\tilde{x}_0 = (\tilde{\alpha}_1(g_1 \circ \tilde{\alpha}_2))(1) = g_1(\tilde{\alpha}_2(1)) = g_1g_2\tilde{x}_0 = \phi([\alpha_1])\phi([\alpha_2])\tilde{x}_0.$$

Comme l'action est sans point fixe, on conclut $\phi([\alpha_1, \alpha_2]) = \phi([\alpha_1])\phi([\alpha_2])$. $\square$

Corollaire 3.4.0.9. On déduit les groupes fondamentaux suivants :

$$\pi_1(\mathbf{S}^1) = \mathbf{Z} \quad \pi_1(\mathbf{P}^2) = \mathbf{Z}/2\mathbf{Z} \quad \pi_1(\mathbf{T}^2) = \mathbf{Z}^2 \quad \pi_1(\mathbf{K}^2) = \mathbf{Z} \rtimes \mathbf{Z}.$$

En rappelant que $\pi_1(\mathbf{S}^2) = 0$, on voit que les espaces $\mathbf{S}^2$, $\mathbf{P}^2$, $\mathbf{T}^2$ et $\mathbf{K}^2$ sont homéomorphiquement distincts.

Remarque 3.4.0.10 (REVÊTEMENT UNIVERSEL). La situation du théorème 3.4.0.8 semble très particulière, elle est en fait très générale. Un résultat important, qu'on n'abordera pas, dit que tout espace X connexe par arc, localement connexe par arc, semi-localement connexe par arc (*i.e.* tout point $x \in X$ a un voisinage U tel que l'inclusion $\pi_1(U, x) \rightarrow \pi_1(X, x)$ est triviale), a un revêtement $p: \tilde{X} \rightarrow X$ par un espace simplement connexe $\tilde{X}$ (appelé revêtement universel de X) muni d'une action par homéomorphisme de $\pi_1(X, x_0)$ tel que $X \simeq \tilde{X}/\pi_1(X, x_0)$. Le revêtement universel de $\mathbf{S}^2$ et $\mathbf{P}^2$ est $\mathbf{S}^2$, pour toutes les autres surfaces, c'est $\mathbf{R}^2$.

3.5 Applications du calcul de $\pi_1(\mathbf{S}^1)$

Théorème 3.5.0.1 (FONDAMENTAL DE L'ALGÈBRE). Tout polynôme $P \in \mathbf{C}[X]$ non constant a au moins une racine dans $\mathbf{C}$.

Démonstration. Exercice de TD. □

Théorème 3.5.0.2 (POINT FIXE DE BROUWER). Tout application continue $h: \mathbf{D}^2 \rightarrow \mathbf{D}^2$ a au moins un point fixe *i.e.* un point $x \in \mathbf{D}^2$ tel que $h(x) = x$.

Démonstration. Par l'absurde, supposons h sans point fixe. On définit une application continue $f: \mathbf{D}^2 \rightarrow \mathbf{S}^1$, tel que $f(x) = x$ sur $\mathbf{S}^1 = \partial \mathbf{D}^2$, comme suit : $f(x)$ est l'intersection avec $\mathbf{S}^1$ de la demi-droite $[h(x), x]$ issue de $h(x)$ contenant x . Notons $i: \mathbf{S}^1 \rightarrow \mathbf{D}^2$ l'inclusion. On a $f \circ i = \text{Id}$ sur $\mathbf{S}^1$, donc $f_* \circ i_* = \text{Id}$ sur $\pi_1(\mathbf{S}^1, 1)$ par 3.3.0.5. Ceci implique que $f_*: \pi_1(\mathbf{D}^2, 1) \rightarrow \pi_1(\mathbf{S}^1, 1)$ est surjective, ce qui est absurde puisque $\pi_1(\mathbf{D}^2, 1) = 0$ alors que $\pi_1(\mathbf{S}^1, 1) = \mathbf{Z}$. □

Exercice 3.5.0.3 (en TD). Montrer que sont équivalents :

- (1) $\mathbf{S}^n$ n'est pas contractile.
- (2) (Brouwer) Toute application continue $\mathbf{B}^{n+1} \rightarrow \mathbf{B}^{n+1}$ a un point fixe.
- (3) (Borsuk) Il n'existe pas d'application continue $r: \mathbf{B}^{n+1} \rightarrow \mathbf{S}^n$ qui fixe tout $x \in \mathbf{S}^n$.

Théorème 3.5.0.4 (BORSUK-ULAM EN DIMENSION 2). Toute application continue $f: \mathbf{S}^2 \rightarrow \mathbf{R}^2$ admet une paire de point antipodaux x et $-x$ tel que $f(x) = f(-x)$. En particulier, il n'existe pas d'application continue injective $\mathbf{S}^2 \rightarrow \mathbf{R}^2$.

Démonstration. Exercice de TD. □

3.6 Théorème de Van Kampen

Heuristique : On veut calculer par exemple la somme connexe $\pi_1(\mathbf{T}^2 \# \mathbf{T}^2, x)$. On peut voir $\mathbf{T}^2 \# \mathbf{T}^2$ comme une réunion $V_1 \cup V_2$ obtenue en recollant deux copies de $\mathbf{T}^2 \setminus D$ (avec D un disque plongé) sur leur bord, de sorte que $V_1 \cap V_2 = \partial V_1 = \partial V_2 \sim \mathbf{S}^1$ est l'image du bord. Prenons $x \in V_1 \cap V_2$. Un lacet α basé en x est homotope à un produit de lacets $\alpha_1 \cdots \alpha_N$ chacun contenu dans V_1 ou V_2 (*cf.* preuve du théorème 3.2.0.3). Ainsi, $\pi_1(V_1 \cup V_2, x)$ doit être en relation avec une sorte de « produit » des groupes $\pi_1(V_1, x)$ et $\pi_1(V_2, x)$. Si α est dans $V_1 \cap V_2$, on a le choix de l'envoyer dans $\pi_1(V_1, x)$ ou dans $\pi_1(V_2, x)$. Plutôt que de choisir de manière arbitraire (et pas canonique) un des groupes, on peut l'envoyer dans les deux, quitte à identifier ensuite les images dans le « produit » des deux groupes, en faisant le quotient adéquat. C'est ce que dira le théorème de Van Kampen. La première étape est d'introduire la notion de produit adéquate, celle de produit libre de groupes.

3.6.1 Produit libre de groupes et groupe libre

Soit $(G_\alpha)_\alpha$ une famille de groupes. Leur *produit libre* $*_\alpha G_\alpha$ est le groupe défini comme suit.

Ensemblistement :

$$*_\alpha G_\alpha = \left\{ g_1 g_2 \cdots g_m; g_i \in \bigcup_\alpha (G_\alpha \setminus e_\alpha) \ (g_i \in G_\alpha \Rightarrow g_{i+1} \notin G_\alpha) \ (\forall m \geq 0) \right\}$$

i.e. l'ensemble des mots $g_1 g_2 \cdots g_m$ de longueur arbitraire $m \geq 0$ finie, où chaque lettre g_i appartient à un des G_α , n'est pas neutre et où deux lettres consécutives g_i et g_{i+1} appartiennent à des groupes différents. De tels mots sont dits *réduits*. Le mot vide est autorisé, ce sera l'identité.

Exemple 3.6.1.1. Si $G_1 = \mathbf{Z}$ est engendrée par a (avec la notation multiplicative $\mathbf{Z} = \{a^k; k \in \mathbf{Z}\}$) et $G_2 = \mathbf{Z}$ est engendré par b alors $*_2 \mathbf{Z}$ (qu'on note plutôt $\mathbf{Z} * \mathbf{Z}$) le mot $a^2 b^{-3} a b$ est réduit mais $a^2 b^{-3} b b$ ne l'est pas.

Loi de groupe : Elle est simplement définie par juxtaposition et réductions itérées. On pose que

$$(g_1 \cdots g_m)(h_1 \cdots h_n) = g_1 \cdots g_m h_1 \cdots h_n$$

si g_m et h_1 n'appartiennent pas au même G_α . Si $g_m, h_1 \in G_\alpha$ on considère l'élément $(g_m h_1) \in G_\alpha$ et si $(g_m h_1) \neq e_\alpha$, on pose

$$(g_1 \cdots g_m)(h_1 \cdots h_n) = g_1 \cdots g_{m-1}(g_m h_1)h_2 \cdots h_n \in *_\alpha G_\alpha$$

qui est réduit. Si $(g_m h_1) = e_\alpha$, alors on le supprime et on réitère l'opération ci-dessus avec $(g_1 \cdots g_{m-1})(h_2 \cdots h_n)$. Après un nombre fini d'itérations on obtient un mot réduit, possiblement vide si on part de $(g_1 \cdots g_m)(g_m^{-1} \cdots g_1^{-1})$. Ainsi dans $\mathbf{Z} * \mathbf{Z}$,

$$(a^2 b^{-3})(b^2) = a^2 b^{-1}.$$

Pour vérifier qu'on a un groupe, la seule chose délicate est l'associativité.

Lemme 3.6.1.2. La loi précédemment définie sur $*_\alpha G_\alpha$ est associative.

Démonstration. Soit $\mathcal{M} = *_\alpha G_\alpha$ l'ensemble des mots réduits. À chaque $g \in G_\alpha$, on associe

$$L_g: \mathcal{M} \rightarrow \mathcal{M}, \quad g_1 g_2 \cdots g_m \mapsto \begin{cases} (gg_1)g_2 \cdots g_m & \text{si } g_1 \in G_\alpha \\ gg_1 g_2 \cdots g_m & \text{si } g_1 \notin G_\alpha \end{cases}$$

Cette application vérifie $L_{gg'} = L_g \circ L_{g'}$ pour tout $g, g' \in G_\alpha$ puisque $(gg')g_1 = g(g'g_1)$ dans G_α . Il s'ensuit que L_g est bijective d'inverse $L_{g^{-1}}$. Tout ceci montre que $g \mapsto L_g$ définit un morphisme de groupes $G_\alpha \rightarrow \text{Bij}(\mathcal{M})$. On définit alors $L: \mathcal{M} \rightarrow \text{Bij}(\mathcal{M})$ par $L(g_1 \cdots g_m) = L_{g_1} \circ \cdots \circ L_{g_m}$. L'application est injective car $L(g_1, \dots, g_m)(\emptyset) = g_1 \cdots g_m$. La formule $L_{gg'} = L_g \circ L_{g'}$ implique que la loi définie sur $\mathcal{M}$ correspond via L à la composition dans $\text{Bij}(\mathcal{M})$. Comme la composition est associative dans $\text{Bij}(\mathcal{M})$, on conclut que la loi de $\mathcal{M}$ est associative. $\square$

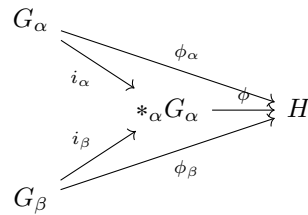
Exercice 3.6.1.3. (1) Montrer que dans $*_\alpha G_\alpha$

$$g_1 \cdots g_m = h_1 \cdots h_n \Leftrightarrow m = n \text{ et } g_i = h_i \ (\forall i = 1, \dots, m).$$

(2) Soit G, H des groupes non triviaux. Montrer que le centre de $G * H$ est trivial et que les éléments de $G * H$ d'ordre fini sont les conjugués d'éléments d'ordre fini de G ou de H .

Notons $i_\alpha: G_\alpha \rightarrow *_\alpha G_\alpha$ les morphismes injectifs envoyant e_α sur le mot vide et $g \in G_\alpha$ sur le mot formé de la lettre g . Ils permettent d'identifier naturellement les G_α avec les sous-groupes $i_\alpha(G_\alpha)$. La réunion des $i_\alpha(G_\alpha)$ engendre $*_\alpha G_\alpha$ et chaque intersection $i_\alpha(G_\alpha) \cap i_\beta(G_\beta)$ est réduite au mot vide (le neutre) quand $\alpha \neq \beta$. Le produit libre de groupe est caractérisé par la propriété fondamentale d'extension suivante :

Proposition 3.6.1.4 (PROPRIÉTÉ UNIVERSELLE D'UNIQUE EXTENSION). Étant donné un groupe H et des morphismes $\phi_\alpha: G_\alpha \rightarrow H$, il existe un unique morphisme $\phi: *_\alpha G_\alpha \rightarrow H$ tels que $\phi \circ i_\alpha = \phi_\alpha$ pour tout α .



Cette propriété caractérise $*_\alpha G_\alpha$ à isomorphisme près : si G est un groupe muni de morphismes $j_\alpha: G_\alpha \rightarrow G$ tels que la réunion des $j_\alpha(G_\alpha)$ engendre G et G a la propriété d'unique extension (remplaçant ci-dessus les i_α par j_α et $*_\alpha G_\alpha$ par G) alors les j_α sont injectifs et G est isomorphe à $*_\alpha G_\alpha$.

Démonstration. D'abord l'unicité. Si $\phi: *_\alpha G_\alpha \rightarrow H$ est un morphisme vérifiant $\phi \circ i_\alpha = \phi_\alpha$ pour tout α , alors pour tout mot (réduit) $g_1 \cdots g_m \in *_\alpha G_\alpha$, on a

$$\phi(g_1 \cdots g_m) = \phi_{\alpha_1}(g_1) \cdots \phi_{\alpha_m}(g_m)$$

où $g_i \in G_{\alpha_i}$, ce qui caractérise ϕ uniquement. On définit ϕ par cette formule, il reste à vérifier que c'est un morphisme. Considérons le produit $(g_1 \cdots g_m)(g_{m+1} \cdots g_{m+n})$ de mots réduits. Si $\alpha_m \neq \alpha_{m+1}$, on a

$$\begin{aligned} \phi((g_1 \cdots g_m)(g_{m+1} \cdots g_{m+n})) &= \phi(g_1 \cdots g_m g_{m+1} \cdots g_{m+n}) \\ &= \phi_{\alpha_1}(g_1) \cdots \phi_{\alpha_m}(g_m) \phi_{\alpha_{m+1}}(g_{m+1}) \cdots \phi_{\alpha_{m+n}}(g_{m+n}) \\ &= \phi(g_1 \cdots g_m) \phi(g_{m+1} \cdots g_{m+n}). \end{aligned}$$

Sinon, $\alpha_m = \alpha_{m+1}$, alors $g_m g_{m+1} \in G_{\alpha_m}$. En utilisant que ϕ_{α_m} est un morphisme, si $g_m g_{m+1} \neq e$, on écrit

$$\begin{aligned} \phi((g_1 \cdots g_m)(g_{m+1} \cdots g_{m+n})) &= \phi(g_1 \cdots g_{m-1}(g_m g_{m+1})g_{m+2} \cdots g_{m+n}) \\ &= \phi_{\alpha_1}(g_1) \cdots \phi_{\alpha_{m-1}}(g_{m-1})\phi_{\alpha_m}(g_m g_{m+1})\phi_{\alpha_{m+2}}(g_{m+2}) \\ &\quad \cdots \phi_{\alpha_{m+n}}(g_{m+n}) \\ &= \phi_{\alpha_1}(g_1) \cdots \phi_{\alpha_{m-1}}(g_{m-1})\phi_{\alpha_m}(g_m)\phi_{\alpha_{m+1}}(g_{m+1})\phi_{\alpha_{m+2}}(g_{m+2}) \\ &\quad \cdots \phi_{\alpha_{m+n}}(g_{m+n}) \\ &= \phi(g_1 \cdots g_m)\phi(g_{m+1} \cdots g_{m+n}). \end{aligned}$$

Si $g_m g_{m+1} = e$, alors $\phi_{\alpha_m}(g_m)\phi_{\alpha_{m+1}}(g_{m+1}) = e_H$. Si $\alpha_{m-1} \neq \alpha_{m+2}$, alors

$$\begin{aligned} \phi((g_1 \cdots g_m)(g_{m+1} \cdots g_{m+n})) &= \phi(g_1 \cdots g_{m-1}g_{m+2} \cdots g_{m+n}) \\ &= \phi_{\alpha_1}(g_1) \cdots \phi_{\alpha_{m-1}}(g_{m-1})\phi_{\alpha_{m+2}}(g_{m+2}) \cdots \phi_{\alpha_{m+n}}(g_{m+n}) \\ &= \phi_{\alpha_1}(g_1) \cdots \phi_{\alpha_{m-1}}(g_{m-1})\phi_{\alpha_m}(g_m)\phi_{\alpha_{m+1}}(g_{m+1})\phi_{\alpha_{m+2}}(g_{m+2}) \\ &\quad \cdots \phi_{\alpha_{m+n}}(g_{m+n}) \\ &= \phi(g_1 \cdots g_m)\phi(g_{m+1} \cdots g_{m+n}). \end{aligned}$$

Après itérations éventuelles, on conclut que ϕ est un morphisme. On montre que la propriété d'extension caractérise le produit libre en considérant la diagramme suivant

$$\begin{array}{ccccc} G_\alpha & & & & \\ i_\alpha \downarrow & \searrow j_\alpha & & i_\alpha \searrow & \\ *_\alpha G_\alpha & \xrightarrow{j} & G & \xrightarrow{i} & *_\alpha G_\alpha \\ i_\beta \uparrow & \nearrow j_\beta & & i_\beta \nearrow & \\ G_\beta & & & & \end{array}$$

où i et j sont les morphismes universels étendant les i_α et j_α respectivement. D'abord, j_α est injectif car $i \circ j_\alpha = i_\alpha$ est injectif. Montrons que i est surjectif et pour cela que $i \circ j = \text{Id}$ (ce qui prouve aussi que j est injective). On a pour tout $g_1 \cdots g_m \in *_\alpha G_\alpha$,

$$\begin{aligned} i \circ j(g_1 \cdots g_m) &= i \circ j(i_{\alpha_1}(g_1) \cdots i_{\alpha_m}(g_m)) \\ &= i \circ j(i_{\alpha_1}(g_1)) \cdots i \circ j(i_{\alpha_m}(g_m)) \\ &= i_{\alpha_1}(g_1) \cdots i_{\alpha_m}(g_m) \\ &= g_1 \cdots g_m \end{aligned}$$

donc $i \circ j = \text{Id}$. Cet argument n'utilise pas que $g_1 \cdots g_m$ est réduit, uniquement que $*_\alpha G_\alpha$ est engendré par les $i_\alpha(G_\alpha)$. Inversant le rôle de G et de $*_\alpha G_\alpha$, on obtient que $j \circ i = \text{Id}$, donc que i est injective. $\square$

Définition 3.6.1.5 (GROUPE LIBRE). On appelle *groupe libre* un produit libre $*_\alpha \mathbf{Z}$ de copies de $\mathbf{Z}$ (en nombre fini ou infini).

Corollaire 3.6.1.6. Tout groupe est le quotient d'un groupe libre.

Démonstration. Soit H un groupe quelconque. On se donne une famille $\{a_\alpha\}_{\alpha \in I}$ de générateurs de H (en prenant éventuellement tous les éléments de H). Pour chaque $\alpha \in I$, on pose que $G_\alpha = \mathbf{Z}$ où on note $a_\alpha \in G_\alpha$ l'élément générateur ($(G_\alpha, \cdot) \simeq (\mathbf{Z}, +)$; $a_\alpha \mapsto 1 \in \mathbf{Z}$). Soit $L = *_\alpha G_\alpha$ leur produit libre. On définit $\phi_\alpha: G_\alpha \rightarrow H$ par $\phi_\alpha(a_\alpha) = a_\alpha$, de sorte que $\phi_\alpha(G_\alpha)$ est le sous-groupe de H engendré par a_α . Leur réunion engendre H , donc le morphisme universel $\phi: L \rightarrow H$ vérifiant $\phi \circ i_\alpha = \phi_\alpha$ est surjectif. Il s'ensuit par factorisation que H est isomorphe à $L/\mathcal{N}$ où $\mathcal{N} = \text{Ker}(\phi) \subset L$:

$$\begin{array}{ccc} L & \xrightarrow{\phi} & H \\ & \searrow & \nearrow \bar{\phi} \\ & L/\mathcal{N} & \end{array}$$

$\square$

Description courte d'un groupe - présentation par générateurs et relations : L'écriture $L/\mathcal{N}$ de tout groupe H réduit sa description à celle de L et $\mathcal{N}$. Évidemment, on voudrait des descriptions les plus courtes possibles. Par définition, le groupe libre $L = L(a_i, i \in I)$ est déterminé par un choix de générateur a_i de H (ils deviennent les générateurs de L). Reste à décrire $\mathcal{N} \subset L$. Comme c'est un sous-groupe distingué, plutôt que de

donner ses générateurs, on procède comme suit. Soit $(r_j)_{j \in J}$ une famille d'éléments de L . Les r_j sont des mots en les a_i qu'on appelle les *relations*. On note $\mathcal{N}(r_j, j \in J)$ le plus petit sous-groupe distingué de L contenant les r_j . Cela raccourcit beaucoup la description : dans $\mathbf{Z} * \mathbf{Z} = L(a, b)$, les générateurs de $\mathcal{N}(b)$ sont

$$b, aba^{-1}, a^2ba^{-2}, \dots$$

Exercice 3.6.1.7. Si S est une partie d'un groupe G , alors $\mathcal{N}(S)$ (le plus petit sous-groupe distingué de G contenant S) est engendré par les conjugués des éléments de S .

On appellera $\mathcal{N}(S)$ le sous-groupe distingué engendré par S (mais cela ne veut pas dire que S est une partie génératrice).

Définition 3.6.1.8. Soit H un groupe. On dit que $\langle a_i, i \in I; r_j, j \in J \rangle$ est une *présentation de H par générateurs et relations* si les a_i engendrent H et si le morphisme universel $\phi: L(a_i, i \in I) \rightarrow H$ induit un isomorphisme $L(a_i, i \in I)/\mathcal{N}(r_j, j \in J) \simeq H$.

On pense au r_j comme aux relations par lesquelles il faut quotienter L pour obtenir H . Le corollaire 3.6.1.6 dit que tout groupe a une présentation par générateurs et relations.

La présentation n'est pas unique : $\mathbf{Z} = \langle a; \rangle = \langle a, b; b \rangle = \langle a, b, c; b, c \rangle = \dots$

Exercice 3.6.1.9 (en TD). (1) $\mathbf{Z}/n\mathbf{Z} = \langle a; a^n \rangle$;

(2) $\mathbf{Z}^2 = \langle a, b; aba^{-1}b^{-1} \rangle$;

(3) $\mathbf{Z}^n = \langle a_1, \dots, a_n; a_i a_j a_i^{-1} a_j^{-1}, \forall i, j \rangle$.

3.6.2 Théorème de Van Kampen

Théorème 3.6.2.1 (VAN KAMPEN). Soit $X = U_1 \cup U_2$ où U_1, U_2 et $U_1 \cap U_2$ sont des ouverts connexes par arc et soit $x \in U_1 \cap U_2$. Soit $i_j: \pi_1(U_1 \cap U_2, x) \rightarrow \pi_1(U_j, x)$ avec $j = 1, 2$ les morphismes induits par l'inclusion $U_1 \cap U_2 \subset U_j$. On a alors

$$\pi_1(X, x) \simeq (\pi_1(U_1, x) * \pi_1(U_2, x)) / \mathcal{N}$$

où $\mathcal{N} = \mathcal{N}(i_1(\gamma)i_2(\gamma)^{-1}; \gamma \in \pi_1(U_1 \cap U_2, x))$, i.e. $\mathcal{N}$ est le sous-groupe distingué engendré par les $i_1(\gamma)i_2(\gamma)^{-1}$ pour $\gamma \in \pi_1(U_1 \cap U_2, x)$.

La relation $i_1(\gamma)i_2(\gamma)^{-1}$ signifie que les lacets $i_1(\gamma)$ et $i_2(\gamma)$, qui sont une image différente dans le produit libre $\pi_1(U_1) * \pi_1(U_2)$, sont identifiés dans $\pi_1(X)$, ce qui est normal vu que c'est l'image d'un même lacet $\gamma \subset U_1 \cap U_2 \subset X$.

Démonstration. La preuve est assez longue, nous n'en ferons qu'une partie. Considérons le diagramme

$$\begin{array}{ccccc} & & \pi_1(U_1, x) & & \\ & \nearrow i_1 & \downarrow & \searrow \phi_1 & \\ \pi_1(U_1 \cap U_2, x) & & \pi_1(U_1, x) * \pi_1(U_2, x) & \xrightarrow{\phi} & \pi_1(X, x) \\ & \searrow i_2 & \uparrow & \nearrow \phi_2 & \\ & & \pi_1(U_2, x) & & \end{array}$$

où les ϕ_j sont les morphismes induits par les inclusions $U_j \subset X$ et ϕ est le morphisme universel donné par 1.2.5.1. D'après le lemme 3.2.0.4, les images des ϕ_j engendrent $\pi_1(X, x)$, donc le morphisme universel ϕ est surjectif. Il s'ensuit que $\pi_1(X, x) \simeq (\pi_1(U_1, x) * \pi_1(U_2, x)) / \text{Ker}(\phi)$. D'après la proposition 3.3.0.5, le morphisme $\phi_1 \circ i_1$, induit par les inclusions $U_1 \cap U_2 \subset U_1 \subset X$, est égal au morphisme induit par l'inclusion directe $U_1 \cap U_2 \subset X$. Il en est de même pour $\phi_2 \circ i_2$, donc $\phi_1 \circ i_1 = \phi_2 \circ i_2$. D'après le diagramme, on a donc $\phi \circ i_1 = \phi \circ i_2$ i.e. $\phi(i_1(\gamma)i_2(\gamma)^{-1}) = 0$ pour tout $\gamma \in \pi_1(U_1 \cap U_2, x)$. Ceci montre que $\mathcal{N} \subset \text{Ker}(\phi)$. Ceci prouve par factorisation l'existence d'un morphisme surjectif

$$(\pi_1(U_1, x) * \pi_1(U_2, x)) / \mathcal{N} \rightarrow \pi_1(X, x).$$

La partie la plus délicate de la preuve, que nous ne ferons pas, consiste à montrer l'injectivité de ce morphisme, i.e. que $\mathcal{N} = \text{Ker}(\phi)$. $\square$

En particulier :

Corollaire 3.6.2.2. Sous les hypothèses du théorème de Van Kampen :

(i) Si U_1 est simplement connexe, alors

$$\pi_1(X, x) \simeq \pi_1(U_2, x) / \mathcal{N}$$

où $\mathcal{N}$ est le sous-groupe distingué engendré par l'image $i_2(\pi_1(U_1 \cap U_2, x))$.

(ii) Si $U_1 \cap U_2$ est simplement connexe, alors

$$\pi_1(X, x) \simeq \pi_1(U_1, x) * \pi_1(U_2, x).$$

Démonstration. (i) Puisque $i_1(\gamma) = 1 \in \pi_1(U_1, x) = \{1\}$, le sous-groupe distingué $\mathcal{N}$ engendré par les $i_1(\gamma)i_1(\gamma)^{-1}$ est engendré par l'image de i_2 , donc

$$\pi_1(X, x) \simeq (\{1\} * \pi_1(U_2, x)) / \mathcal{N} \simeq \pi_1(U_2, x) / \mathcal{N}.$$

(ii) Puisque $\pi_1(U_1 \cap U_2, x) = \{1\}$, on a $\mathcal{N} = \{1\}$ □

Deux exemples très simples pour tester la machinerie :

Le disque : On sait que $X = \mathbf{D}^2$ est simplement connexe, donc que $\pi_1(\mathbf{D}^2) = \{1\}$. Écrivons $X = U_1 \cup U_2$ où $U_1 = \mathbf{B}(0, 2r)$ pour $0 < r < 1/2$ et $U_2 = \mathbf{D} \setminus \overline{\mathbf{B}(0, r)}$ de sorte que $U_1 \cap U_2 = \{z \in \mathbf{D}^2; r < |z| < 2r\}$. Soit $x \in U_1 \cap U_2$. Notons a le générateur de $\pi_1(U_1 \cap U_2, x) = \mathbf{Z}$ et b le générateur de $\pi_1(U_2, x) = \mathbf{Z}$. Clairement, $i_1(a) = 0 \in \pi_1(U_1, x) = 0$ et $i_2(a) = \pm b \in \pi_1(U_2, x)$. Ainsi, $\mathcal{N} = \pi_1(U_2, x)$, d'où on a bien

$$\pi_1(X, x) \simeq \pi_1(U_2, x) / \pi_1(U_2, x) = \{1\}.$$

Bouquet de cercles : Un bouquet de cercles est la somme pointée $(\mathbf{S}^1 \vee \mathbf{S}^1, x)$ de deux cercles, elle est homéomorphe à l'union de deux cercles dans $\mathbf{R}^2$ tangents en un point.

Proposition 3.6.2.3. On a

$$\pi_1(\mathbf{S}^1 \vee \mathbf{S}^1, x) = \mathbf{Z} * \mathbf{Z}.$$

Démonstration. Notons A l'un des cercles et B l'autre et soit $x = A \cap B$. On peut se donner un voisinage U_1 de A et U_2 de B tel que U_1 a le type d'homotopie de A , U_2 a le type d'homotopie de B et $U_1 \cap U_2$ est simplement connexe. On a alors $\pi_1(U_i, x) = \pi_1(\mathbf{S}^1) = \mathbf{Z}$ pour $i = 1, 2$ et l'assertion découle du corollaire 3.6.2.2 (ii). □

3.6.3 Applications

On peut recalculer le groupe fondamental du projectif. Détaillons cette nouvelle preuve, qui sera utile pour calculer le groupe fondamental des surfaces $N_h = \mathbf{P}^2 \# \cdots \# \mathbf{P}^2$.

Proposition 3.6.3.1. Si $x \in \mathbf{P}^2$, alors

$$\pi_1(\mathbf{P}^2, x) \simeq \mathbf{Z}/2\mathbf{Z}$$

Démonstration. Heuristique : $\mathbf{P}^2$ est la réunion d'un ruban de Moebius compact $\mathbf{M}$ (quotient de $\mathbf{S}^1 \times [-1, 1]$ par antipodie) et d'un disque fermé $\mathbf{D}$ le long de leur bord commun $\partial \mathbf{M} = \partial \mathbf{D} =: S \simeq \mathbf{S}^1$. Notons C le cercle central du ruban (le quotient de $\mathbf{S}^1 \times \{0\}$). On a

$$\begin{aligned} \pi_1(\mathbf{D}) &= 0 \\ \pi_1(S) &= \mathbf{Z} && \text{engendré par } \alpha \text{ parcourant une fois } S \\ \pi_1(\mathbf{M}) &= \mathbf{Z} && \text{engendré par } \beta \text{ parcourant une fois } C. \end{aligned}$$

Dans $\mathbf{M}$, le lacet α est librement homotope (*i.e.* sans fixer les extrémités) à un lacet parcourant deux fois le cercle central C , donc $\alpha \sim \beta^2$. Si $i: \pi_1(S) \rightarrow \pi_1(\mathbf{M})$ est le morphisme induit par l'inclusion, on a donc $i(\mathbf{Z}) = 2\mathbf{Z} \subset \mathbf{Z}$ et comme $\mathcal{N}(i(\mathbf{Z})) = i(\mathbf{Z})$ puisque $\mathbf{Z}$ est abélien, on obtient du corollaire 3.2.0.5 que $\pi_1(\mathbf{P}^2) = \pi_1(\mathbf{M})/i(\pi_1(S)) = \mathbf{Z}/2\mathbf{Z}$.

• On nettoie maintenant cette idée, en travaillant avec des ouverts et en tenant compte des points base, ce qu'on a allègrement négligé ci-dessus. On écrit $\mathbf{P}^2$ comme réunion d'un ruban de Moebius M (ouvert : quotient de $\mathbf{S}^1 \times]-1, 1[$ par antipodie) et d'un disque ouvert D , l'intersection $M \cap D$ étant un cylindre isomorphe à $\mathbf{S}^1 \times I$. En effet, $\mathbf{P}^2 \simeq \mathbf{S}^2 / \sim$ où $x \sim -x$ pour tout $x \in \mathbf{S}^2$. En notant $p: \mathbf{S}^2 \rightarrow \mathbf{P}^2$ la projection (c'est un revêtement de degré 2) et $\mathbf{S}^2 = \{(z, t) \in \mathbf{C} \times \mathbf{R}; |z|^2 + t^2 = 1\}$, on définit

$$\begin{aligned} \widetilde{M} &= \mathbf{S}^2 \cap \{-3/4 < t < 3/4\} & M &= p(\widetilde{M}) \\ \widetilde{D} &= \mathbf{S}^2 \cap \{1/4 < t\} & D &= p(\widetilde{D}) \end{aligned}$$

avec M un ruban de Moebius et D un disque (homéomorphe à $\tilde{D}$ via p). L'intersection $M \cap D$ est homéomorphe à $\tilde{M} \cap \tilde{D} = \mathbf{S}^2 \cap \{1/4 < t < 3/4\} =$ une couronne (homéomorphe à $\mathbf{S}^1 \times]-1, 1[$).

• Notons $\tilde{S}$ la parallèle $\mathbf{S}^2 \times \{t = 1/2\}$ et $S = p(\tilde{S})$ son image dans $M \cap D$. Fixons $\tilde{x} \in \tilde{S}$ et soit $x = p(\tilde{x}) \in S$. Soit $y = p(1, 0)$ et soit $C = p(\{(e^{it}, 0); t \in [0, \pi]\})$ le cercle central du ruban M , c'est l'image d'un demi-cercle de $\mathbf{S}^2 \times \{0\}$. On a

$$\begin{aligned}\pi_1(D, x) &= 0 \\ \pi_1(M \cap D, x) &= \mathbf{Z} \quad \text{engendré par } \alpha \text{ parcourant une fois } S \\ \pi_1(M, y) &= \mathbf{Z} \quad \text{engendré par } \beta \text{ parcourant une fois } C.\end{aligned}$$

Écrivons $\alpha = p \circ \tilde{\alpha}$ où $\tilde{\alpha}$ parcourt une fois le méridien $\tilde{S}$. Écrivons $\beta = p \circ \tilde{\beta}$ où $\tilde{\beta}(t) = (e^{i\pi t}, 0)$ avec $t \in [0, 1]$, parcourant une moitié de l'équateur. Soit $\tilde{\gamma}$ un chemin de $\tilde{M}$ joignant $\tilde{x}$ à $\tilde{y}$ et soit $\gamma = p \circ \tilde{\gamma}$. Clairement, $\tilde{\gamma}^{-1} \tilde{\alpha} \tilde{\gamma}$ est homotope dans $\tilde{M}$ au cercle complet $t \mapsto (e^{2i\pi t}, 0)$ avec $t \in [0, 1]$, parcourant l'équateur i.e. au lacet $\tilde{\beta}(-\tilde{\beta})$. Composant avec p , on en déduit que $\gamma^{-1} \alpha \gamma$ est homotope dans M à β^2 (β parcouru 2 fois). Puisque α engendre $\pi_1(M \cap D, x)$ et β engendre $\pi_1(M, y)$, cela veut dire que

$$\gamma_{\#}(i_*(1)) = 2 \in \pi_1(M, y)$$

et comme $\gamma_{\#}: \pi_1(M, x) \rightarrow \pi_1(M, y)$ est un isomorphisme, on a

$$i_*(1) = \pm 2 \in \pi_1(M, x)$$

d'où $\mathcal{N} = 2\mathbf{Z}$. On conclut via le corollaire 3.6.2.2 que

$$\pi_1(\mathbf{P}^2, x) \simeq \pi_1(M, x)/\mathcal{N} \simeq \mathbf{Z}/2\mathbf{Z}.$$

□

Proposition 3.6.3.2 (π_1 D'UN TORE PERCÉ). Soit $D \subset \mathbf{T}^2 = \mathbf{R}^2/\mathbf{Z}^2$ l'intérieur d'un disque fermé plongé dans $\mathbf{T}^2$ et soit $x \in \partial D$. (i) On a

$$\pi_1(\mathbf{T}^2 \setminus D, x) = \mathbf{Z} * \mathbf{Z} = L(a, b)$$

et le groupe à deux générateurs. À changement de point base près, a et b sont les générateurs $([0, 1] \times \{0\})/\mathbf{Z}$ et de $(\{1\} \times [0, 1])/\mathbf{Z}$ respectivement.

(ii) Si α est un lacet générateur sur ∂D et $i: \partial D \rightarrow \mathbf{T}^2 \setminus D$ l'inclusion, alors

$$i_*([\alpha]) = \pm aba^{-1}b^{-1} = \pm[a, b] \in \pi_1(\mathbf{T}^2 \setminus D, x).$$

Démonstration. (i) Notons $p: \mathbf{R}^2 \rightarrow \mathbf{R}^2/\mathbf{Z}^2 = \mathbf{T}^2$ la projection (c'est un revêtement). On a alors que $A = p([0, 1] \times \{0\})$ est un cercle dans $\mathbf{T}^2$ ainsi que $B = p(\{0\} \times [0, 1])$. Ces deux cercles s'intersectent uniquement en $y := p(0, 0)$, donc $A \cup B$ est homéomorphe à un bouquet de cercle $\mathbf{S}^1 \vee \mathbf{S}^1$. On va montrer que $\mathbf{T}^2 \setminus D$ a le type d'homotopie de $A \cup B$, dont le groupe fondamental est $\mathbf{Z} * \mathbf{Z}$ d'après la proposition 3.6.2.3. Notons que $\mathbf{T}^2 = p([0, 1]^2)$. Sans perte de généralité, on peut supposer que D est l'image par p d'un petit disque $\tilde{D} \subset [0, 1]^2$ centré sur le milieu $m = (1/2, 1/2)$ du carré. Clairement, $[0, 1]^2 \setminus \tilde{D}$ a le type d'homotopie du bord $\partial[0, 1]^2$: on définit une équivalence d'homotopie

$$[0, 1]^2 \setminus \tilde{D} \xrightarrow{\tilde{r}} \partial[0, 1]^2 \xrightarrow{i} [0, 1]^2 \setminus \tilde{D}$$

en posant que $\tilde{r}(z)$ est l'intersection de la demi-droite issue de m contenant z avec le bord $\partial[0, 1]^2$. Sur le bord, $r = \text{Id}$, donc $r \circ i = \text{Id}$ trivialement et sur $([0, 1]^2 \setminus \tilde{D}) \times [0, 1]$, l'application $(z, s) \mapsto H(z, s) = sz + (1-s)\tilde{r}(z)$ réalise une homotopie entre $i \circ \tilde{r}$ et l'identité de $[0, 1]^2 \setminus \tilde{D}$. Un quotient et une factorisation induisent une équivalence d'homotopie $\mathbf{T}^2 \setminus D \rightarrow A \cup B$:

$$\begin{array}{ccccc} [0, 1]^2 \setminus \tilde{D} & \xrightarrow{\tilde{r}} & \partial[0, 1]^2 & \xrightarrow{i} & [0, 1]^2 \setminus \tilde{D} \\ p \downarrow & \searrow & \downarrow p & & \downarrow p \\ \mathbf{T}^2 \setminus D & \xrightarrow{r} & A \cup B & \xrightarrow{i} & \mathbf{T}^2 \setminus D \end{array}$$

Sans perte de généralité, on suppose que $x \in \partial D$ vérifie $r(x) = y$ (en prenant $x = p(\tilde{x})$ où $\tilde{x} \in \partial \tilde{D}$ est l'intersection de $\partial \tilde{D}$ avec la demi-droite passant par $(0, 0)$ issue de m). Ainsi, on a l'isomorphisme

$$r_*: \pi_1(\mathbf{T}^2 \setminus D, x) \rightarrow \pi_1(A \cup B, y) \simeq \pi_1(\mathbf{S}^1 \vee \mathbf{S}^1) = L(a, b)$$

où a correspond au générateur de $\pi_1(A, y) = \mathbf{Z}$ et b à celui de $\pi_1(B, x) = \mathbf{Z}$. Si γ est le chemin joignant x à y tel que $r(\gamma) = y$, il est clair que $r(\gamma a \gamma^{-1}) = a$ et $r(\gamma b \gamma^{-1}) = b$. Il s'ensuit que $\pi_1(\mathbf{T}^2 \setminus D, x)$ est engendré par

$\gamma a \gamma^{-1}$ et $\gamma b \gamma^{-1}$. Si on pose $c = \gamma a \gamma^{-1}$ et $d = \gamma b \gamma^{-1}$, on a exactement $\pi_1(\mathbf{T}^2 \setminus D, x) = L(c, d)$.

(ii) Notons $\tilde{\alpha}$ le lacet base en $\tilde{x}$ faisant une fois le tour de $\partial \tilde{D}$, obtenu en relevant α . On peut supposer que $\tilde{\alpha}$ tourne dans le sens direct. Il est clair que $\tilde{\alpha}$ est homotope dans $[0, 1]^2 \setminus \tilde{D}$, en tant qu'application (sans fixer les extrémités, ça s'appelle une homotopie libre) à un lacet parcourant le bord $\partial[0, 1]^2$ dans le sens direct. Composant avec p , cela montre que α est librement homotope dans $\mathbf{T}^2 \setminus D$ à $aba^{-1}b^{-1}$. Si pour être plus précis on veut fixer le point base x , on voit que α est homotope à

$$\gamma aba^{-1}b^{-1}\gamma^{-1} \sim \gamma a \gamma^{-1} \gamma b \gamma^{-1} \gamma a^{-1} \gamma^{-1} \gamma b^{-1} \gamma^{-1} \sim cdc^{-1}d^{-1}$$

en posant $c = \gamma a \gamma^{-1}$ et $d = \gamma b \gamma^{-1}$, i.e. on a $[\alpha] = cdc^{-1}d^{-1} = [c, d]$. $\square$

Remarque 3.6.3.3. On peut recalculer $\pi_1(\mathbf{T}^2, x)$ comme on l'a fait pour $\pi_1(\mathbf{P}^2)$ en appliquant le théorème de Van Kampen à la décomposition $\mathbf{T}^2 = \mathbf{T}^2 \setminus D \cup \bar{D}$. Comme $\bar{D}$ est simplement connexe et qu'on a calculé que l'image de $\pi_1(\partial D, x) \simeq \mathbf{Z}$ dans $\pi_1(\mathbf{T}^2 \setminus D, x) = L(c, d)$ est engendré par le commutateur $[c, d]$, on obtient donc

$$\pi_1(\mathbf{T}^2, x) = L(c, d) / \mathcal{N}(cdc^{-1}d^{-1}) = \langle c, d; cdc^{-1}d^{-1} \rangle$$

qui est le groupe libre abélien à deux générateurs, i.e. $\mathbf{Z}^2$.

On peut maintenant calculer le groupe fondamental de toutes les surfaces $\Sigma_g = \mathbf{T}^2 \# \dots \# \mathbf{T}^2$, somme connexe de g tores. On va considérer dans les arguments ci-dessus des disques ouverts $D \subset \Sigma_g$ qui seront l'intérieur de disques fermés $\mathbf{D}^2$ plongés dans Σ_g . En particulier, $\partial D \subset \Sigma_g$ est homéomorphe à $\mathbf{S}^1$.

Proposition 3.6.3.4 (GROUPE FONDAMENTAL DE Σ_g). Pour tout entier $g \geq 1$ on a les propriétés suivantes. Soit $D \subset \Sigma_g$ l'intérieur d'un disque fermé plongé dans Σ_g et soit $x \in \partial D$.

(i) On a que

$$\pi_1(\Sigma_g \setminus D, x) = L(a_1, b_1, \dots, a_g, b_g)$$

est un groupe libre à $2g$ générateurs.

(ii) Si α est un lacet générateur du bord ∂D et $i: \partial D \rightarrow \Sigma_g \setminus D$ l'inclusion, alors

$$i_*([\alpha]) = [a_1, b_1] \cdots [a_g, b_g] \in \pi_1(\Sigma_g \setminus D, x).$$

(iii) On a enfin

$$\pi_1(\Sigma_g) = \langle a_1, b_1, \dots, a_g, b_g; [a_1, b_1] \cdots [a_g, b_g] \rangle.$$

Démonstration. On montre (i) et (ii) par récurrence sur g . On a déjà traité le cas $g = 1$. Supposons (i) et (ii) vrais au rang g et considérons une surface Σ_{g+1} . Puisque $\Sigma_{g+1} = \Sigma_g \# \mathbf{T}^2$, elle s'obtient en recollant $\Sigma_g \setminus D_1$ et $\mathbf{T}^2 \setminus D_2$ sur leur bord $\partial D_1 \simeq \partial D_2$. On peut donc l'écrire

$$\Sigma_{g+1} = V_1 \cap V_2$$

où

$$V_1 \simeq \Sigma_g \setminus D_1 \quad V_2 \simeq \mathbf{T}^2 \setminus D_2 \quad V_1 \cap V_2 =: S \simeq \mathbf{S}^1.$$

Soit $D \subset \Sigma_{g+1}$ un disque ouvert centré sur S (son « centre » appartient à S). On écrit les deux cercles ∂D et S comme une union d'intervalles :

$$\partial D = I_1 \cup I_2 \quad I_j \subset V_j \quad S = I \cup J \quad I = \bar{D} \cap S \quad J \cap D = \emptyset \quad x = I \cap I_1 \cap I_2.$$

On a alors

$$\begin{aligned} \Sigma_{g+1} \setminus D &= (V_1 \setminus D) \cup (V_2 \setminus D) \\ &= V'_1 \cup V'_2 \end{aligned}$$

où

$$V'_1 \simeq V_1 \quad \partial V'_1 = I_1 \cup J, \quad V'_2 \simeq V_2 \quad \partial V'_2 = I_2 \cup J \quad V'_1 \cap V'_2 = J.$$

Comme l'intervalle J est simplement connexe, l'hypothèse de récurrence et le théorème de Van Kampen implique (en prenant des voisinages ouverts adéquats de V'_1, V'_2 et $V'_1 \cap V'_2$) que

$$\begin{aligned} (\Sigma_{g+1} \setminus D, x) &\simeq \pi_1(V'_1, x) * \pi_1(V'_2, x) \\ &= L(a_1, b_1, \dots, a_g, b_g) * L(a, b) \\ &= L(a_1, b_1, \dots, a_{g+1}, b_{g+1}) \end{aligned}$$

est un groupe libre à $2g + 2$ générateurs, prouvant (i). Maintenant, considérons α un lacet générateur de ∂D . On paramètre I_1, I_2 et J comme des chemins de sorte que $I_1 I_2 \sim \alpha$ et $I_1(0) = J(0)$. Dans $\Sigma_{g+1} \setminus D$, on a alors

$$I_1 I_2 \sim (I_1 J^{-1})(J I_2)$$

deux lacets basés en x , générateurs respectivement de $\partial V'_1 \simeq \partial V_1$ et $\partial V'_2 \simeq \partial V_2$. L'hypothèse de récurrence appliquée à V'_1 et V'_2 implique

$$i_*([\alpha]) = [I_1 J^{-1}][J I_2] = [a_1, b_1] \cdots [a_g, b_g][a, b] = [a_1, b_1] \cdots [a_{g+1}, b_{g+1}]$$

prouvant (ii) au rang $g + 1$.

On en déduit (iii) en appliquant le théorème de Van Kampen 3.2.0.4 (i) (en prenant des voisinages ouverts adéquats) à la réunion

$$\Sigma_{g+1} = (\Sigma_{g+1} \setminus D) \cup \overline{D}$$

où $\overline{D}$ est simplement connexe et l'intersection des deux domaines est ∂D . Le groupe $\pi_1(\partial D, x)$ est engendré par $[\alpha]$ dont on a calculé en (ii) l'image dans $\pi_1(\Sigma_{g+1} \setminus D)$, donc

$$\begin{aligned} \pi_1(\Sigma_{g+1}, x) &\simeq \pi_1(\Sigma_{g+1} \setminus D, x) / \mathcal{N}(i_*([\alpha])) \\ &= L(a_1, b_1, \dots, a_{g+1}, b_{g+1}) / \mathcal{N}\left(\prod_{i=1}^{g+1} [a_i, b_i]\right) \\ &= \left\langle a_1, b_1, \dots, a_{g+1}, b_{g+1}; \prod_{i=1}^{g+1} [a_i, b_i] \right\rangle. \end{aligned}$$

□

On en vient au calcul du groupe fondamental de $N_h = \mathbf{P}^2 \# \cdots \# \mathbf{P}^2$. La preuve de la proposition 3.6.3.1 établit que

$$\pi_1(\mathbf{P}^2 \setminus D) \simeq \pi_1(\mathbf{S}^1) \simeq \mathbf{Z} = \langle a \rangle$$

et que l'image dans $\pi_1(\mathbf{P}^2 \setminus D)$ du générateur de $\pi_1(\partial D)$ est a^2 . Une preuve par récurrence similaire à celle de la proposition 3.6.3.4 montre que :

Proposition 3.6.3.5 (π_1 DES SURFACES N_h). Pour tout $h \geq 1$, on a la propriété suivante. Soit $D \subset N_h$ l'intérieur d'un disque fermé plongé dans N_h et soit $x \in \partial D$.

(i) On a que

$$\pi_1(N_h \setminus D, x) = L(a_1, a_2, \dots, a_h)$$

est un groupe libre à h générateurs.

(ii) Si α est un lacet générateur du bord ∂D et $i: \partial D \rightarrow N_h \setminus D$ l'inclusion, alors

$$i_*([\alpha]) = a_1^2 a_2^2 \cdots a_h^2 \in \pi_1(N_h \setminus D, x).$$

(iii) On a enfin

$$\pi_1(N_h, x) = \langle a_1, a_2, \dots, a_h; a_1^2 a_2^2 \cdots a_h^2 \rangle.$$

Corollaire 3.6.3.6. (i) Σ_g et N_h ne sont pas homéomorphes.

(ii) $\Sigma_g \simeq \Sigma_{g'} \Leftrightarrow g = g'$.

(iii) $N_h \simeq N_{h'} \Leftrightarrow h = h'$.

Démonstration. Rappelons que des espaces homéomorphes ont des groupes fondamentaux isomorphes. Si G est un groupe, notons $G^{ab} = G/[G, G]$ l'abélianisé de G , qui est le groupe abélien obtenu en quotientant G par $[G, G]$, le sous-groupe de G engendré par ses commutateurs. Si G est isomorphe à H , alors G^{ab} est isomorphe à H^{ab} . On a

$$\begin{aligned} \pi_1(\Sigma_g)^{ab} &\simeq \mathbf{Z} \times \cdots \times \mathbf{Z} = \mathbf{Z}^{2g} \\ \pi_1(N_h)^{ab} &\simeq \mathbf{Z}^{h-1} \times \mathbf{Z}/2\mathbf{Z}. \end{aligned}$$

La deuxième égalité résulte de

$$\begin{aligned} \pi_1(N_h)^{ab} &\simeq \langle a_1, a_2, \dots, a_h; a_1^2 a_2^2 \cdots a_h^2, [a_i, a_j] \forall i, j \rangle \\ &\simeq \langle a_1, a_2, \dots, a_h; (a_1 a_2 \cdots a_h)^2, [a_i, a_j] \forall i, j \rangle \\ &\simeq \langle b_1, b_2, \dots, b_h; b_h^2, [b_i, b_j] \forall i, j \rangle. \end{aligned}$$

On voit que $\pi_1(\Sigma_g)^{ab}$ n'a pas d'élément d'ordre fini alors que $\pi_1(N_h)^{ab}$ si, donc ils ne sont pas isomorphes, ce qui implique (i). Par ailleurs, $\mathbf{Z}^n$ n'est isomorphe à $\mathbf{Z}^m$ que si $n = m$: en quotientant $\mathbf{Z}^n$ par le sous-groupe $2(\mathbf{Z}^n) = (2\mathbf{Z})^n$, on obtient $(\mathbf{Z}/2\mathbf{Z})^n$, de cardinal 2^n . On en déduit (ii) et (iii). □

4 Classification des surfaces compactes

Dans ce chapitre on démontre le théorème 2.3.3.5. Les deux étapes clés sont :

Théorème 4.0.0.1. Toute surface compacte connexe est le quotient d'un polygone de $\mathbf{R}^2$ par identification des côtés par paires.

Théorème 4.0.0.2. Tout quotient d'un polygone de $\mathbf{R}^2$ par identification des côtés par paires est une surface, homéomorphe à $\mathbf{S}^2$ ou à une somme connexe de tores et de projectifs.

4.1 Quotients de polygones

Définition 4.1.0.1 (POLYGONE). Soit $n \geq 3$. Soit P_n l'enveloppe convexe dans $\mathbf{R}^2$ des n points $\{z_k = e^{2i\pi k/n}, k = 1, \dots, n\}$ de $\mathbf{S}^1$. Les z_k sont les *sommets* de P_n et les segments $e_k = [z_k, z_{k+1}]$ (orientés) sont ses *côtés*. On dit que $P \subset \mathbf{R}^2$ est un *polygone* à n côtés (n -gone) s'il existe un homéomorphisme $f: P_n \rightarrow P$ affine sur chaque côté de P_n . Les sommets et côtés de P sont images des sommets et côtés de P_n .

Étant donné un polygone P , on fixe un homéomorphisme $f: P_n \rightarrow P$ et la numérotation $p_1, \dots, p_n$ des sommets de P ainsi induite ($p_k = f(z_k)$).

Définition 4.1.0.2 (SYMBOLES). Soit $\mathcal{A} = \{a_1, a_2, \dots, a_p\}$ un ensemble de lettres distinctes. On appelle *symbole* de longueur n l'expression

$$\sigma = a_{i_1}^{\varepsilon_1} a_{i_2}^{\varepsilon_2} \cdots a_{i_n}^{\varepsilon_n}$$

où $a_{i_j} \in \mathcal{A}$ et $\varepsilon_j \in \{-1, 1\}$ pour $j = 1, \dots, n$. On dit que σ est *pair* si chaque lettre a_i apparaît deux fois. En particulier, σ est de longueur paire.

En général, on omet l'exposant $+1$. Par exemple $aba^{-1}b^{-1}$ est un symbole pair sur $\mathcal{A} = \{a, b\}$ et $aba^{-1}b^{-1}c$ est un symbole (non pair) sur $\{a, b, c\}$.

Association polygone-symbole : Étant donné un polygone P à n côtés et un symbole $\sigma = a_{i_1}^{\varepsilon_1} a_{i_2}^{\varepsilon_2} \cdots a_{i_n}^{\varepsilon_n}$ de longueur n , on lie σ et P en associant $a_{i_k}^{\varepsilon_k}$ au côté $[p_k, p_{k+1}]$. On définit alors une relation d'équivalence $\sim$ sur P comme suit. Un point de $\text{Int}(P)$ n'est identifié qu'à lui-même. On identifie deux côtés $e_j \sim e_k$ lorsqu'il sont associés à une même lettre, par l'isomorphisme affine :

- qui envoie $[p_j, p_{j+1}]$ sur $[p_k, p_{k+1}]$ si $\varepsilon_j = \varepsilon_k$;
- qui envoie $[p_j, p_{j+1}]$ sur $[p_{k+1}, p_k]$ si $\varepsilon_j = -\varepsilon_k$.

L'identification des côtés se fait donc en respectant le sens horaire lorsque les exposants sont les mêmes et en l'inversant sinon. On note $X(P, \sigma) = P / \sim$ l'espace topologique quotient obtenu (P étant muni de la topologie induite par $\mathbf{R}^2$). Par exemple, $X([0, 1]^2, aba^{-1}b^{-1})$ est homéomorphe à $\mathbf{T}^2$. On vérifie sans peine que

$$X(P, \sigma) \approx X(P', \sigma)$$

s'il existe un homéomorphisme $P \rightarrow P'$ qui envoie p_i sur p'_i et qui est affine sur les côtés (l'homéomorphisme passe au quotient). La classe d'homéomorphie de $X(P, \sigma)$ ne dépendant que de σ , on la notera donc $X(\sigma)$. Sans hypothèse supplémentaire, l'espace quotient $X(\sigma)$ n'est pas nécessairement une surface.

Lemme 4.1.0.3. Si σ est pair, $X(\sigma)$ est une surface compacte et connexe.

Démonstration. On donne les idées de la construction. Montrons d'abord que $X(\sigma)$ est localement homéomorphe à $\mathbf{R}^2$. Soit $P = P_{2n}$ tel que $X(P, \sigma) = X(\sigma)$. Notons $\pi: P \rightarrow P / \sim$ la projection et pour tout $x \in P$, notons $[x]$ sa classe d'équivalence. Il y a 3 cas à considérer : (i) x est intérieur à P , (ii) x est intérieur à un côté, (iii) x est un sommet.

(i) Si $x \in \text{Int}(P)$, il existe un voisinage ouvert U de x dans $\mathbf{R}^2$ tel que $U \subset P$. Comme $U \cap \partial P = \emptyset$, l'application $\pi: U \rightarrow \pi(U)$ est un homéomorphisme, d'où la conclusion.

(ii) Supposons x intérieur à un côté e de P . Il existe un unique autre côté $e' \sim e$, donc $[x] = \{x_1, x_2\}$ où $x = x_1$ et $x_2 \in e'$. En posant $D_i = \mathbf{B}(x_i, \varepsilon) \cap P$ pour $\varepsilon > 0$, alors $\pi(D_1 \cup D_2)$ est un voisinage de $[x]$. Comme $e \sim e'$, il existe clairement des homéomorphismes affines

$$\begin{aligned} h_1: D_1 &\rightarrow \mathbf{B}(0, \varepsilon) \cap (\mathbf{R}_+ \times \mathbf{R}) \\ h_2: D_2 &\rightarrow \mathbf{B}(0, \varepsilon) \cap (\mathbf{R}_- \times \mathbf{R}) \end{aligned}$$

tel que $h_1(z) = h_2(z)$ si et seulement si $z \sim z'$. L'application $h: D_1 \cup D_2 \rightarrow \mathbf{B}(0, \varepsilon)$ égale à h_i sur D_i est surjective continue ouverte et se factorise au quotient en un homéomorphisme d'un voisinage de $[x]$ sur $\mathbf{B}(0, \varepsilon)$.

(iii) Si x est un sommet, alors $[x] = \{x_1, \dots, x_k\}$, des sommets.

• Supposons d'abord que $[x] = \{x_1\}$. Cela n'arrive que si les 2 côtés de part et d'autre de x_1 sont labélisés aa^{-1} ou $a^{-1}a$. Dans ce cas, on peut trouver un voisinage U de x dans P homéomorphe à $V = [0, \infty) \times \mathbf{R}$ et où $(0, y) \sim (0, -y)$ pour tout $y \in \mathbf{R}$ (le point x correspond à $(0, 0)$). Le quotient de V est homéomorphe à $\mathbf{R}^2$, d'où la conclusion.

• Supposons maintenant $[x] = \{x_1, \dots, x_k\}$ où $k \geq 2$. Aucun des x_i n'appartient alors à des côtés labélisés aa^{-1} ou $a^{-1}a$. Pour $\varepsilon > 0$ petit, posons $A_i = B(x_i, \varepsilon) \cap P$. C'est une portion de secteur angulaire délimité par les deux côtés contenant x_i . Le quotient $\pi(\bigcup_{i=1}^k A_i)$ est un voisinage de $[x]$. Soient $B_1, B_2, \dots, B_k$ des secteurs angulaires de $B(0, \varepsilon)$ d'angles $2\pi/k$, numérotés dans le sens horaire, dont la réunion est $B(0, \varepsilon)$. Après renumérotation des x_i , on va définir des homéomorphismes $h_i: A_i \rightarrow B_i$ compatibles avec $\sim$, ce qui induira par factorisation un homéomorphisme de $\pi(\bigcup_i A_i) \rightarrow B(0, \varepsilon)$.

On part de $x_1 = x = p_i$ et on appelle $c_1 = [p_{i-1}, p_i]$ et $c'_0 = [p_i, p_{i+1}]$ les deux côtés bordant A_i . On se donne un homéomorphisme affine $h_1: A_1 \rightarrow B_1$ envoyant $c_1 \cap A_1$ sur $B_1 \cap B_2$. Comme le symbole est pair, il existe un unique $j \pmod n$ tel que c_1 est équivalent à $[p_j, p_{j+1}]$ ou $[p_{j+1}, p_j]$. On appelle c'_1 ce côté $[p_j, p_{j+1}]$ ou $[p_{j+1}, p_j]$ et x_2 l'extrémité équivalente à x_1 . Comme on suppose $k \geq 2$, on a $x_2 \neq x_1$ (sinon $[x] = \{x_1\}$). Le secteur A_2 contenant x_2 est bordé par c'_1 et par un autre côté qu'on appelle $B_2 \cap B_3$. Par construction, $h_1(z) = h_2(z')$ si $c_1 \ni z \sim z' \in c'_1$. On itère la construction. Ayant défini x_i et deux côtés c'_{i-1}, c_i de A_i , on définit c'_i comme côté équivalent à c_i , $x_{i+1} \in c'_i$ est l'extrémité équivalente à x_i , c_{i+1} est l'autre côté de A_{i+1} . On a $x_{i+1} \neq x_1, \dots, x_i$ car le symbole est pair, donc c'_i ne peut pas être un des côtés déjà appariés $c_1, c'_1, c_2, c'_2, \dots, c_{i-1}, c'_{i-1}$. La construction se termine après j itérations avec (c'_{j-1}, x_j, c_j) lorsque $c_j \sim c'_0$. Nécessairement, $j = k$ et on a parcouru tous les points de $[x]$: si un point $x' \in [x]$ était manqué, comme la construction a produit des paires de côtés appariés et que le symbole est pair, il ne serait plus possible de passer de x' à x via des identifications de côté. À la fin, on définit $h_k(A_k) = B_k$, $h_k(z) = h_1(z')$ si $c_k \ni z \sim z' \in c'_0$, ce qui permet de boucler la boucle. On montre que le quotient est séparé, étant donné $[x] \neq [x']$ et $\varepsilon > 0$ assez petit pour que

$$\left(\bigcup_{x_i \in [x]} B(x_i, \varepsilon) \right) \cap \left(\bigcup_{x'_i \in [x']} B(x'_i, \varepsilon) \right) = \emptyset.$$

C'est facile car les classes d'équivalence sont finies. Pour $\varepsilon > 0$ assez petit, l'intersection avec P de chaque union est saturée, donc l'image est un voisinage ouvert de $[x]$, d'où la conclusion. Il s'ensuit que $X(\sigma)$ est compacte comme image continue d'un compact dans un séparé. C'est donc une surface compacte. Elle est connexe comme image continue d'un connexe. $\square$

Cela prouve la première partie du théorème 4.0.0.2. Quelques exemples :

$$X(aba^{-1}b^{-1}) = \mathbf{T}^2 \quad X(abab^{-1}) = \mathbf{K}^2 \quad X(aa^{-1}bb^{-1}) = \mathbf{S}^2 \quad X(aabb^{-1}) = \mathbf{P}^2.$$

Dans le cas de $\mathbf{S}^2$ et $\mathbf{P}^2$, on a envie de supprimer le bb^{-1} qui ne semble pas fondamental. On peut le faire en étendant la définition 4.1.0.1 pour disposer d'un polygone à deux côtés : on décide que $P_2 = \mathbf{D}^2 \subset \mathbf{C}$ (le disque unité fermé), ses deux « sommets » sont $p_1 = 1$ et $p_2 = i$ et donc ses deux côtés sont les arcs de cercle supérieur $e_1 = (p_1 p_2) = \{z \in \mathbf{S}^1; \operatorname{Im}(z) \geq 0\}$ et inférieur $e_2 = (p_2 p_1) = \{z \in \mathbf{S}^1; \operatorname{Im}(z) \leq 0\}$. Si le symbole est aa^{-1} , on identifie tout $z \in e_1$ avec son conjugué $\bar{z}$ i.e. on quotiente $\partial \mathbf{D}^2$ par symétrie d'axe les réels. On obtient bien $X(aa^{-1}) = \mathbf{S}^2$. Si le symbole est aa , on identifie tout $z \in e_1$ avec $-z \in e_2$ i.e. on quotiente $\partial \mathbf{D}^2$ par l'antipodie. On obtient bien $X(aa) = \mathbf{P}^2$.

4.2 Classification : opérations sur les symboles

On travaille sur les surfaces $X(\sigma)$. On va simplifier le symbole par certaines opérations pour se ramener à des formes simples.

4.2.1 Symbole et somme connexe

On commence par une formule magique de simplicité :

Proposition 4.2.1.1. Si σ_1, σ_2 sont deux symboles pairs sur des alphabets disjoints, alors

$$X(\sigma_1 \sigma_2) = X(\sigma_1) \# X(\sigma_2).$$

Démonstration. Soit $n_i \in \mathbf{N}$ avec $i = 1, 2$ la longueur de σ_i . Soit P_i un polygone à n_i côtés tel que $X(P_i, \sigma_i) \in X(\sigma_i)$. En ajoutant un côté supplémentaire à P_1 entre p_{n_1} et p_1 , on obtient un polygone P'_1 à $(n_1 + 1)$ côtés auquel on associe $\sigma'_1 = \sigma_1 c$ où c n'appartient pas aux alphabets de σ_1, σ_2 . Ainsi, $X(P'_1, \sigma'_1)$ est homéomorphe à $X(\sigma_1) \setminus D_1$ où $D_1 \subset X(\sigma_1)$ est l'intérieur d'un disque plongé. En effet, le bord ∂D_1 est l'image du côté

supplémentaire $p_{n_1+1}p_1$ labélisé c et la preuve du lemme 4.1.0.3 montre que $p_{n_1+1} \sim p_1$. De même, soit P'_2 un polygone à $(n_2 + 1)$ côtés auquel on associe $\sigma'_2 = c^{-1}\sigma_2$. En outre, $X(P'_2, \sigma'_2)$ est homéomorphe à $X(\sigma_2) \setminus D_2$ où $D_2 \subset X(\sigma_1)$ est l'intérieur d'un disque plongé et ∂D_2 est l'image du côté supplémentaire labélisé c^{-1} . L'espace quotient obtenu en identifiant les bords $\partial D_1 \sim \partial D_2$ est $X(\sigma_1) \# X(\sigma_2)$. Soit P' le polygone obtenu en recollant P'_1 et P'_2 le long de cc^{-1} , il est muni naturellement du symbole $\sigma_1\sigma_2$. L'application continue adéquate se factorise en un homéomorphisme $X(\sigma_1) \# X(\sigma_2) \rightarrow X(\sigma_1\sigma_2)$.

$$\begin{array}{ccc}
 P'_1 \coprod P'_2 & \xrightarrow{c \sim c^{-1}} & P' \xrightarrow{\pi} P' / \sim \in X(\sigma_1\sigma_2) \\
 \pi \downarrow & & \nearrow \simeq \\
 (X(\sigma_1) \setminus D_1) \coprod (X(\sigma_2) \setminus D_2) & & \\
 \partial D_1 \sim \partial D_2 \downarrow & & \\
 X(\sigma_1) \# X(\sigma_2) & &
 \end{array}$$

□

4.2.2 Simplification de symboles

Définition 4.2.2.1. On dit que deux symboles σ, σ' sont équivalents si $X(\sigma) = X(\sigma')$.

C'est évidemment une relation d'équivalence. On note $\sigma \sim \sigma'$ la relation. On va opérer sur les symboles un certain nombre de manipulations/simplifications qui respectent $\sim$. Pour certaines c'est trivial à vérifier :

- *Échange de lettre* : (si les σ_i ne contiennent ni a , ni b)

$$\sigma_1 a^\varepsilon \sigma_2 a^{\varepsilon'} \sigma_3 \sim \sigma_1 b^\varepsilon \sigma_2 b^{\varepsilon'} \sigma_3. \quad (4.1)$$

- *Permutation d'orientation* :

$$\sigma_1 a^\varepsilon \sigma_2 a^{\varepsilon'} \sigma_3 \sim \sigma_1 a^{-\varepsilon} \sigma_2 a^{-\varepsilon'} \sigma_3. \quad (4.2)$$

- *Permutation circulaire* :

$$a_{i_1}^{\varepsilon_1} a_{i_2}^{\varepsilon_2} \cdots a_{i_k}^{\varepsilon_k} \sim a_{i_2}^{\varepsilon_2} \cdots a_{i_k}^{\varepsilon_k} a_{i_1}^{\varepsilon_1}. \quad (4.3)$$

- *Inversion* :

$$a_{i_1}^{\varepsilon_1} a_{i_2}^{\varepsilon_2} \cdots a_{i_k}^{\varepsilon_k} \sim a_{i_k}^{-\varepsilon_k} \cdots a_{i_2}^{-\varepsilon_2} a_{i_1}^{-\varepsilon_1}. \quad (4.4)$$

Les simplifications suivantes sont moins triviales :

Proposition 4.2.2.2. Soient A, B deux symboles tel que AB est pair, non vide (A ou B peut être vide) et ne contient pas la lettre a , alors on a :

- (i) (Simplification des aa^{-1})

$$Aaa^{-1}B \sim AB. \quad (4.5)$$

- (ii) (Regroupement des aa)

$$aAaB \sim aaAB^{-1}. \quad (4.6)$$

- (iii) (Inversion des $bbcc$)

$$aabbcc \sim aabcb^{-1}c^{-1}. \quad (4.7)$$

- (iv) (Regroupement des $aba^{-1}b^{-1}$) Soit A, B, C, D des symboles tels que $ABCD$ est pair, non vide (A, B, C ou D peut être vide) et ne contient pas a, b , alors

$$aAbBa^{-1}Cb^{-1}D \sim aba^{-1}b^{-1}DCBA. \quad (4.8)$$

Remarque 4.2.2.3. (1) Un cas particulier de (ii) est $abab^{-1} \sim aabb$. En effet, comme $X(aabb) = X(aa) \# X(bb)$, cela équivaut à $\mathbf{K}^2 \simeq \mathbf{P}^2 \# \mathbf{P}^2$. Dans le cas général, cela dit que $X(aAaB) = \mathbf{P}^2 \# X(AB^{-1})$.

(2) La formule (4.7) équivaut à $\mathbf{P}^2 \# \mathbf{P}^2 \# \mathbf{P}^2 \simeq \mathbf{P}^2 \# \mathbf{T}^2$.

(3) La formule (4.8) dit que $X(aAbBa^{-1}Cb^{-1}D) = \mathbf{T}^2 \# X(DCBA)$.

Démonstration. (i) Par permutation circulaire et comme $\mathbf{S}^2$ est le neutre de la somme connexe (cf. proposition 2.3.3.3), on a

$$X(Aaa^{-1}B) = X(aa^{-1}BA) = X(aa^{-1}) \# X(BA) = \mathbf{S}^2 \# X(BA) = X(BA) = X(AB)$$

donc $Aaa^{-1}B \sim AB$.

Voici une autre preuve. Soit P un polygone associé à $Aaa^{-1}B$ et P' un polygone associé à AB . On se donne $f: P \rightarrow P'$ continue, affine sur les côtés (compatible avec les relations d'équivalence) qui envoie a et a^{-1} sur un même segment $f(a)$ dans P' et est un homéomorphisme de $P \setminus \{a, a^{-1}\}$ sur $P' \setminus f(a)$. Ainsi, $p \circ f$ passe au

quotient en un homéomorphisme $P/\sim \rightarrow P'/\sim$.

(ii) Si A ou B est vide, c'est trivial, supposons qu'ils ne le sont pas. Soit P un polygone associé à $aAaB$. Soit $p_i p_{i+1}$ et $p_j p_{j+1}$ les côtés (disjoints) labélisés a . En découpant P le long d'un segment c joignant p_{i+1} à p_{j+1} , on obtient un polygone P_1 associé à acB (P_1 contient p_i) et un polygone P_2 associé à Aac^{-1} . Soit P' le polygone obtenu en recollant le long de a le polygone P_1 et le symétrique de P_2 par rapport à $p_j p_{j+1}$. On a alors que P' est associé à $ccBA^{-1}$. En considérant les applications adéquates de $P_1 \coprod P_2$ vers P et vers P' , prolongées vers $X(P, aAaB)$ et $X(P', ccBA^{-1})$, on obtient par factorisation un homéomorphisme $X(P, aAaB) \simeq X(P', ccBA^{-1})$, donc $aAaB \sim ccBA^{-1}$. Ensuite

$$X(ccBA^{-1}) = X(cc) \# X(BA^{-1}) = X(aa) \# X(AB^{-1}) = X(aaAB^{-1})$$

en utilisant (4.1) et (4.4), d'où la conclusion.

(iii) Montrons

$$aabc b^{-1} c^{-1} \sim aabc b^{-1} c \quad (4.9)$$

soit $\mathbf{P}^2 \# \mathbf{T}^2 \simeq \mathbf{P}^2 \# \mathbf{K}^2$. Le résultat en découlera car $\mathbf{K}^2 \simeq \mathbf{P}^2 \# \mathbf{P}^2$. On a

$$\begin{aligned} aabc b^{-1} c^{-1} &= aa(bc)(cb)^{-1} \\ &\sim a(bc)a(cb) && \text{par (4.6)} \\ &\sim cacbab && \text{par permutation circulaire (4.3)} \\ &\sim cca(bab)^{-1} && \text{par (4.6)} \\ &\sim ccab^{-1} a^{-1} b^{-1} \sim aabc b^{-1} c && \text{par échange de lettre (4.1) et d'orientation (4.2).} \end{aligned}$$

(iv) Soit P un polygone labélisé par $aAbBa^{-1}Cb^{-1}D$. On découpe P le long d'un segment c joignant deux sommets initiaux de b et on recolle les deux morceaux le long de a en un polygone P' . On découpe à nouveau P' le long d'un segment d joignant deux extrémités opposés de c et on recolle les deux morceaux le long de b pour obtenir un polygone P'' équivalent à P . On obtient

$$\sigma \sim d^{-1}cdc^{-1}DCBA \sim aba^{-1}b^{-1}DCBA.$$

□

On peut maintenant classifier. On dit qu'un symbole pair σ contient un projectif s'il contient une lettre apparaissant deux fois avec le même exposant. À permutation circulaire et changement d'orientation près, il est de la forme $aAaB$ où AB est pair et ne contient pas a .

Théorème 4.2.2.4 (CLASSIFICATION). Si σ est un symbole pair, alors σ est équivalent à aa^{-1} ou à un des symboles suivants :

- (i) $a_1 a_1 a_2 a_2 \cdots a_h a_h$ si σ contient un projectif;
- (ii) $a_1 b_1 a_1^{-1} b_1^{-1} a_2 b_2 a_2^{-1} b_2^{-1} \cdots a_g b_g a_g^{-1} b_g^{-1}$ sinon.

Par conséquent, $X(\sigma)$ est $\mathbf{S}^2$, N_h (somme connexe de h projectifs) ou Σ_g (somme connexe de g tores).

Tout d'abord, énonçons et démontrons un lemme.

Lemme 4.2.2.5 (RÉDUCTION DES CLASSES DE SOMMETS). (). Si σ est un symbole pair, alors il existe $\sigma' \sim \sigma$ vérifiant :

- (i) $|\sigma'| \leq |\sigma|$;
- (ii) σ' n'a qu'une classe d'équivalence de sommets;
- (iii) σ' contient un projectif si et seulement si σ contient un projectif.

Démonstration. Par application de (4.5), on peut supposer que σ ne contient pas de aa^{-1} . Supposons que le polygone P associé à σ admette deux classes d'équivalence de sommets au moins $[p]$ et $[q]$. Elles sont de cardinal ≥ 2 car la seule configuration où un sommet n'est équivalent qu'à lui-même est le sommet milieu de aa^{-1} , ce qu'on a supprimé. On peut supposer que $p = p_i$ et $q = p_{i+1}$ sont sur un même côté (quitte à changer de classe d'équivalence), labélisé par a . Soit b la lettre sur $rp := p_{i-1}p$. Ainsi, $b \neq a$ sinon $p \sim q$ et $b \neq a^{-1}$ puisqu'on a supprimé les aa^{-1} . En outre, $p_{i-1}p$ est équivalent à un côté $p_j p_{j+1}$ différent de pq . On découpe P le long du segment $c = rq$ et on recolle rq sur $p_j p_{j+1}$, directement si $p_j p_{j+1}$ est labélisé b^{-1} , après symétrie le long de b si $p_j p_{j+1}$ est labélisé b . Ainsi, partant de $baAb^{-1}B$, on obtient $\sigma' = cAac^{-1}B$; partant de $baAbB$ on obtient $\sigma' = cAca^{-1}B$. On a $\sigma' \sim \sigma$, la classe d'équivalence $[p]$ diminue d'une unité et $[q]$ augmente d'une unité. On réitère la procédure ci-dessus (incluant des simplifications éventuelles des dd^{-1} par (4.5)) de diminution de $[p]$ d'une unité et augmentation d'une autre classe d'une unité jusqu'à la disparition de $[p]$. On réitère le tout tant qu'il reste plusieurs classes d'équivalence. □

Démonstration du théorème 4.2.2.4. On procède par récurrence sur la longueur de σ . Soit σ un symbole pair. Si $|\sigma| = 2$ ou 4 , la conclusion est vraie. Supposons la conclusion vraie pour les symboles de longueur $\ell \geq 4$. Soit σ de longueur $\ell + 2$.

1er cas : σ contient un projectif i.e. $\sigma \sim aAaB$. D'après (4.6),

$$\sigma \sim aaAB^{-1} = aa\sigma' \quad |\sigma'| = |\sigma| - 2$$

On utilise l'hypothèse de récurrence.

- Si $\sigma' \sim bb^{-1}$, alors $\sigma \sim aa$.
- Si $\sigma' \sim a_1a_1 \cdots a_ka_k$, alors $\sigma \sim aaa_1a_1 \cdots a_ka_k$, donc (i) est satisfait.
- Sinon, $\sigma' \sim T_1T_2 \cdots T_k$ où $T_i = a_ib_ia_i^{-1}b_i^{-1}$.

Par application de (4.7), on peut remplacer un projectif et un tore par trois projectifs. En notant $A_j = u_ju_j$, on obtient par itération

$$\begin{aligned} \sigma \sim aa\sigma' &= aaT_1T_2 \cdots T_k \\ &\sim A_1A_2A_3T_2 \cdots T_k \\ &\sim A_1A_2A_3A_4A_5T_3 \cdots T_k \\ &\sim \dots \\ &\sim A_1A_2 \cdots A_{2k+1} \end{aligned}$$

vérifiant (i) et terminant la récurrence.

2e cas : σ ne contient pas de projectif. Il est de la forme $aAa^{-1}B$ où AB est pair, non vide et ne contient pas de projectif. Montrons que

$$\sigma \sim aba^{-1}b^{-1}\sigma' \quad |\sigma'| = |\sigma| - 4 \quad (4.10)$$

où σ' ne contient pas de projectif et la conclusion en découlera par récurrence. Avant de procéder, on simplifie d'abord le symbole pour que le polygone associé n'ait qu'une classe de sommets. Il peut en effet en avoir plusieurs. Par exemple

$$\sigma = aa^{-1}bb^{-1}e^{-1}gcc^{-1}g^{-1}dd^{-1}e$$

en a 7.

On en vient à (4.10). Soit $\sigma = aAa^{-1}B$. Après application du lemme 4.2.2.5, on suppose qu'il n'y a qu'une classe d'équivalence de sommets. Montrons qu'il existe $b \in A$ et $b^{-1} \in B$ i.e.

$$\sigma = a \cdots b \cdots a^{-1} \cdots b^{-1} \cdots$$

et on conclut alors grâce à (4.8). Si les côtés de P labélisés par A sont stables pour la relation d'équivalence, ceux marqués par B aussi et il y a plusieurs classes d'équivalence de sommets, contrairement à l'hypothèse. Il s'ensuit qu'un côté de P labélisé dans A est équivalent à un côté de P labélisé dans B au moins. Comme σ ne contient pas de projectif, il existe donc $b \in A$ tel que $b^{-1} \in B$. Ceci termine la preuve du théorème 4.2.2.4 et conclut celle du théorème 4.0.0.2. $\square$

Corollaire 4.2.2.6 (CALCUL DU π_1). Si σ est un symbole de longueur 2ℓ ne comportant qu'une classe de sommets, alors

$$\pi_1(X(\sigma)) = \langle a_1, \dots, a_\ell; \sigma \rangle.$$

Démonstration. Soit P un polygone associé à σ et $D \subset P$ l'intérieur d'un disque fermé plongé dans l'intérieur de P , alors D se plonge dans $S = X(P, \sigma)$. On applique le théorème de Van Kampen à la décomposition $S = S \setminus D \cup \overline{D}$. Notons $p: P \rightarrow S$ la projection. L'équivalence d'homotopie $P \setminus D \rightarrow \partial P$ passe au quotient et se factorise en une équivalence d'homotopie $S \setminus D \rightarrow p(\partial D) = \bigvee_{\ell} (\mathbf{S}^1)$ un bouquet de ℓ cercles. Ainsi, $\pi_1(S \setminus D) =$

$L(a_1, \dots, a_\ell)$ le groupe libre à ℓ générateurs. Le générateur a_i correspond au lacet $p(a_i) \subset S$ image du côté labélisé a_i . Comme D est simplement connexe, on sait que

$$\pi_1(S) = \pi_1(S \setminus D) / \mathcal{N}(i_*([\alpha]))$$

où $[\alpha]$ est un générateur de $\pi_1(\partial D)$. Dans $P \setminus D$, on voit que α est librement homotope au générateur de ∂P , i.e. $\pm\sigma$ (vu comme chemin) ce qui veut dire en projetant que le générateur de ∂D est librement homotope dans $S \setminus D$ à σ (vu comme produit de lacets), d'où $\pi_1(S) = \langle a_1, \dots, a_\ell; \sigma \rangle$. $\square$

4.3 Toute surface connexe est quotient de P

Dans cette section on prouve le théorème 4.0.0.1 qui affirme que toute surface compacte connexe est le quotient d'un polygone par identification des côtés par paires. Cela repose sur un résultat non trivial de Radó que nous ne feront que citer.

4.3.1 Triangulations

Définition 4.3.1.1 (TRIANGULATION). Une *triangulation* (finie) d'un espace topologique X est la donnée d'une famille $\mathcal{T} = \{T_1, \dots, T_N\}$ de fermés recouvrant X et d'homéomorphismes $\phi_i: P_i \rightarrow T_i$, où chaque $P_i \subset \mathbf{R}^2$ est un polygone à 3 côtés (*i.e.* un triangle non plat). Les T_i sont appelés *triangles*, les images des sommets et côtés des P_i sont les sommets et côtés des T_i . On demande de plus que pour tout $i \neq j$, l'un des points suivant soit vérifié :

- $T_i \cap T_j$ est vide;
- $T_i \cap T_j$ est un sommet;
- $T_i \cap T_j$ est un côté de T_i et de T_j .

On peut supposer de plus que sur chaque côté commun $e = T_i \cap T_j$, la composée $\phi_j \phi_i^{-1}$ est affine entre les côtés de P_i et P_j correspondants. On utilise souvent le terme *arêtes* pour désigner les côtés de la triangulation. On admettra :

Théorème 4.3.1.2 (RADÓ, 1925). Toute surface compacte admet une triangulation.

Une triangulation de surface a la propriété suivante :

Proposition 4.3.1.3. Soit $\mathcal{T}$ une triangulation d'une surface. Chaque arête de $\mathcal{T}$ est alors le côté de 2 triangles exactement.

C'est clairement faux pour une triangulation d'une réunion de plans $(\mathbf{R}^2 \times \{0\}) \cup (\{0\} \times \mathbf{R}^2)$ mais celui-ci n'est pas localement homéomorphe à $\mathbf{R}^2$, qui est la propriété clé.

Avant de démontrer la proposition précédente, énonçons et démontrons deux lemmes.

Lemme 4.3.1.4 (UN DEMI-PLAN FERMÉ N'EST PAS LOCALEMENT HOMÉOMORPHE À $\mathbf{R}^2$). Précisément, soit $A = \{z \in \mathbf{C}; z \geq 0\}$, alors aucun point de ∂A n'admet de voisinage homéomorphe à $\mathbf{R}^2$.

Démonstration. On procède par contradiction. Supposons que $a \in \partial A$ admette un voisinage $U \subset A$ homéomorphe à $\mathbf{R}^2$. Soit $f: U \rightarrow \mathbf{R}^2$ un homéomorphisme tel que $f(a) = 0$. Pour $\varepsilon > 0$ assez petit, $V := B(a, \varepsilon) \cap A \subset U$. Clairement, $V \setminus \{a\}$ est étoilé (depuis $a + i\varepsilon/2$) donc simplement connexe. Soit $W = f(V) \subset \mathbf{R}^2$, c'est un voisinage de 0 et $W \setminus \{0\} = f(V \setminus \{a\})$ est simplement connexe. Ceci est impossible. Pour $r > 0$ assez petit, $B := B(0, r) \subset V'$ est également un voisinage de 0 mais $B \setminus \{0\}$ n'est pas simplement connexe : il a le type d'homotopie de $\mathbf{S}^1$. L'inclusion et l'application $r: W \setminus \{0\} \rightarrow B \setminus \{0\}$ définie par $r(z) = zr^{-1}$ pour $|z| \geq r$ et $r(z) = z$ si $|z| \leq r$

$$B \setminus \{0\} \xrightarrow{i} W \setminus \{0\} \xrightarrow{r} B \setminus \{0\}$$

induisent les morphismes

$$\begin{array}{ccccc} \pi_1(B \setminus \{0\}, x) & \xrightarrow{i_*} & \pi_1(W \setminus \{0\}, x) & \xrightarrow{r_*} & \pi_1(B \setminus \{0\}, x) \\ \mathbf{Z} & \xrightarrow{i_*} & 0 & \xrightarrow{r_*} & \mathbf{Z} \end{array}$$

où i_* est injectif (car $r_* \circ i_* = (r \circ i)_* = (\text{Id})_* = \text{Id}$) ce qui est absurde. $\square$

Lemme 4.3.1.5. La réunion de plans $A = (\mathbf{R}^2 \times \{0\}) \cup (\{0\} \times \mathbf{R}^2)$ n'est pas localement homéomorphe à $\mathbf{R}^2$. Précisément, si $a \in (\mathbf{R}^2 \times \{0\}) \cap (\{0\} \times \mathbf{R}^2)$, alors aucun voisinage de a n'est homéomorphe à $\mathbf{R}^2$.

Démonstration. On procède par contradiction. Soit $U \subset A$ un voisinage de a homéomorphe à $\mathbf{R}^2$. Sans perte de généralité, on suppose $a = 0$. Notons de suite que $U \setminus \{0\}$ a le type d'homotopie de $\mathbf{R}^2$ moins un point *i.e.* de $\mathbf{S}^1$, donc que $\pi_1(U \setminus \{0\}, b) = \mathbf{Z}$ pour tout point base b . Soit $\varepsilon > 0$ assez petit pour que $B := \overline{B}(0, \varepsilon) \cap A \subset U$. On affirme que $B \setminus \{0\}$ a le type d'homotopie de $\mathbf{S}^1 \cup J$ où J est un intervalle attaché à $\mathbf{S}^1$ en deux points. En effet, $\mathbf{S}^1 \cup J = S(0, \varepsilon) \cap A$ où l'équivalence d'homotopie est donnée par $h: B \setminus \{0\} \rightarrow S(0, \varepsilon) \cap A$ définie par $h(x) = \varepsilon|x|^{-1}$ (la réciproque est l'inclusion). Il s'ensuit que $\pi_1(B \setminus \{0\}, b) = \mathbf{Z} * \mathbf{Z}$ (ce qu'on voit par Van Kampen appliqué à $\mathbf{S}^1 \cup J$ vu comme réunion de 2 cercles s'intersectant sur l'intervalle J). Considérant $r: U \setminus \{0\} \rightarrow B \setminus \{0\}$ définie par $r(x) = x$ si $|x| \leq \varepsilon$ et $r(x) = \varepsilon|x|^{-1}$ sinon,

$$B \setminus \{0\} \xrightarrow{i} U \setminus \{0\} \xrightarrow{r} B \setminus \{0\}$$

on obtient les morphismes

$$\begin{array}{ccccc} \pi_1(B \setminus \{0\}, b) & \xrightarrow{i_*} & \pi_1(U \setminus \{0\}, b) & \xrightarrow{r_*} & \pi_1(B \setminus \{0\}, b) \\ \mathbf{Z} * \mathbf{Z} & \xrightarrow{i_*} & \mathbf{Z} & \xrightarrow{r_*} & \mathbf{Z} * \mathbf{Z} \end{array}$$

où i_* est injectif (car $r_* \circ i_* = (r \circ i)_* = (\text{Id})_* = \text{Id}$). Comme $\mathbf{Z}$ est abélien et $\mathbf{Z} * \mathbf{Z}$ ne l'est pas c'est absurde. $\square$

Démonstration de la proposition 4.3.1.3. Soit $\mathcal{T}$ une triangulation d'une surface S et soit e une arête de $\mathcal{T}$, par définition e est le côté d'un $T_i \in \mathcal{T}$.

(i) Montrons d'abord qu'il existe $T_j \neq T_i$ tel que $e = T_i \cap T_j$. Le point essentiel est le lemme 4.3.1.4. Étant donné $P \subset \mathbf{R}^2$ un triangle et $a \in \partial P$ n'étant pas un sommet, tout voisinage U (assez petit) de a dans P est homéomorphe à un voisinage d'un point de ∂A et par conséquent n'est pas homéomorphe à $\mathbf{R}^2$. Soit $f: P \rightarrow T_i$ un homéomorphisme et $b \in \partial T$ qui n'est pas un sommet. Montrons que pour tout voisinage U de b dans S , U rencontre $S \setminus T$. Clairement, on peut supposer U arbitrairement petit. Comme S est une surface, on peut également supposer que U est homéomorphe à $\mathbf{R}^2$. Néanmoins, U n'est pas inclus dans T , sinon $f^{-1}(U)$ serait un voisinage de a dans P homéomorphe à $\mathbf{R}^2$, ce que le lemme interdit. Puisque $\mathcal{T}$ recouvre S , une suite de points de $S \setminus T_i$, convergeant vers b , doit être contenu dans un triangle T_j et $T_i \cap T_j = e$.

(ii) Montrons maintenant que T_j est unique. Le point clé est le lemme 4.3.1.5. Comme dans la preuve du lemme, on montrerait de même qu'une union de k plans $A_1, A_2, \dots, A_k$ s'intersectant sur une droite commune Δ n'est pas localement homéomorphe à $\mathbf{R}^2$, en chaque point de Δ . L'ensemble $B \setminus \{0\}$ dans la preuve du lemme aurait le type d'homotopie d'un bouquet de k cercles.

Maintenant, si des triangles $T_1, T_2, \dots, T_k$ de S , avec $k \geq 3$, avaient un côté commun δ , on aurait un voisinage dans S d'un point de δ à la fois homéomorphe à $\mathbf{R}^2$ et à un voisinage dans $\bigcup_i A_i$ d'un point de Δ ci-dessus et on vient de voir que c'est impossible. Ceci conclut la preuve de la proposition. $\square$

4.3.2 Preuve du théorème

On montre le théorème 4.0.0.1.

Démonstration. Soit S une surface compacte connexe. On se donne une triangulation $T_1, \dots, T_N$ de S , fournie par le théorème de Radó 4.3.1.2. Avec ces derniers, viennent les homéomorphismes $\phi_i: P_i \rightarrow T_i$ où les P_i sont des triangles de $\mathbf{R}^2$. On peut supposer les P_i disjoints dans $\mathbf{R}^2$ et ainsi définir $\phi: \bigcup_i P_i \rightarrow S$ continue par $\phi = \phi_i$ sur P_i . La relation d'équivalence $\sim$ sur $\bigcup_i P_i$ induite par ϕ ($x \sim y$ si $\phi(x) = \phi(y)$) correspond à identifier un côté de P_i à un de côté de P_j si leur image par ϕ est un côté commun $T_i \cap T_j$. L'identification des côtés se fait donc par paires d'après la proposition 4.3.1.3. On construit maintenant un polygone P à partir des P_i . Via un isomorphisme affine, on déplace P_1 dans la boule unité, positionnant ses 3 sommets sur $\mathbf{S}^1$. On appelle encore P_1 le nouveau triangle. On choisit un côté e de P_1 , il est équivalent au côté e' d'un certain P_i . On déplace P_i par un isomorphisme affine pour que $e = e' = P_1 \cap P_i$ et que le dernier sommet de P_i soit sur $\mathbf{S}^1$. On obtient un polygone avec ses sommets sur $\mathbf{S}^1$. On itère : on choisit un côté du polygone et on y recolle le triangle correspondant par un isomorphisme affine, obtenant un nouveau polygone. Comme la surface est connexe, on obtient après N étapes un polygone P qui contient tous les P_i . Identifiant les côtés de P par paires selon $\sim$, on obtient S . Ceci conclut la preuve du théorème 4.0.0.2 et du théorème de classifications des surfaces compactes 2.3.3.5. $\square$