

SÉMINAIRE N. BOURBAKI

VALENTIN POÉNARU

Extension des immersions en codimension 1

Séminaire N. Bourbaki, 1968, exp. n° 342, p. 473-505

http://www.numdam.org/item?id=SB_1966-1968__10__473_0

© Association des collaborateurs de Nicolas Bourbaki, 1968, tous droits réservés.

L'accès aux archives du séminaire Bourbaki (<http://www.bourbaki.ens.fr/>) implique l'accord avec les conditions générales d'utilisation (<http://www.numdam.org/conditions>). Toute utilisation commerciale ou impression systématique est constitutive d'une infraction pénale. Toute copie ou impression de ce fichier doit contenir la présente mention de copyright.

NUMDAM

Article numérisé dans le cadre du programme
Numérisation de documents anciens mathématiques

<http://www.numdam.org/>

EXTENSION DES IMMERSIONS EN CODIMENSION 1

(d'après Samuel BLANK)

par Valentin POÉNARU

1. Position du problème.

On va considérer des variétés différentiables et des applications différentiables $f : M \rightarrow N$. La source M sera toujours compacte (à bord non nécessairement vide) et le but N sans bord. On va désigner par $T(M)$ l'espace tangent de M et par $T(M)_x$ sa fibre au point $x \in M$. Je rappelle que f est une IMMERSION si l'application tangente $df_x : T(M)_x \rightarrow T(N)_{f(x)}$ est toujours injective.

On considère maintenant une sous-variété de M , désignée par P ; l'inclusion canonique sera désignée par $i : P \rightarrow M$; p, m, n seront les dimensions de P, M, N , respectivement. On suppose que $m \leq n$.

Le problème de l'extension des immersions est le suivant : une immersion $f : P \rightarrow N$ étant donnée quand peut-on la prolonger en une immersion $F : M \rightarrow N$? On veut donc trouver une immersion $F : M \rightarrow N$ qui rende le diagramme suivant commutatif :

$$(1) \quad \begin{array}{ccc} & P & \\ i \swarrow & & \searrow f \\ M & \xrightarrow{\quad F \quad} & N \end{array}$$

Le problème se présente sous deux formes complètement différentes, suivant que $m < n$ ou $m = n$. La situation vraiment intéressante, comme on va le voir

dans ce qui suit, est le PROBLÈME d'EXTENSION EN CODIMENSION 1, où $m = n$,
 $p = m - 1$.

Je vais dire d'abord quelques mots sur le cas $m < n$. De la théorie de Smale-Hirsch [7], [3], [9], [4], [5], [10], on peut extraire le théorème suivant :

THÉOREME 1 (théorème d'extension).- Soient M, P, N, i, f comme avant, et supposons que $m < n$. Alors, une immersion $F : M \rightarrow N$ telle que (1) soit commutatif, existe si et seulement si il existe une application d'espaces fibrés vectoriels $\Phi : T(M) \rightarrow T(N)$, monomorphisme sur chaque fibre tel que le diagramme suivant :

$$(2) \quad \begin{array}{ccc} & T(P) & \\ di \swarrow & & \searrow df \\ T(M) & \xrightarrow{\Phi} & T(N) \end{array}$$

soit commutatif, à une homotopie (fibrée) près.

Le problème d'extension des immersions, pour $m < n$, se trouve ainsi complètement résolu, en principe tout au moins, puisqu'il se réduit à une question d'homotopie, donc, de notre point de vue à un "problème qu'on sait toujours résoudre".

Le théorème 1 est une conséquence de deux autres théorèmes, dus essentiellement à Smale et Hirsch, et qu'on va rappeler :

THÉOREME 2 (théorème de relèvement des homotopies).- Soient, comme avant, P, M, N trois variétés différentiables, P étant sous-variété de M , telles que $\dim M = m < n = \dim N$. On considère les espaces $\text{Imm}(P, N)$, $\text{Imm}(M, N)$ d'immersions de P dans N et de M dans N , munis de la topologie C^r , $r > 2$. Soit

$$\pi : \text{Imm}(M, N) \rightarrow \text{Imm}(P, N)$$

la projection canonique, obtenue par restriction des immersions de M à P.

Dans ces conditions,

$$\pi : \text{Imm}(M, N) \rightarrow \text{Image } \pi$$

est un fibré de Serre.

Remarque.— Le théorème 2 reste vrai si la condition $m < n$ est remplacée par :

$$\dim P = \dim M = m \leq n = \dim N$$

mais M est obtenue à partir de P en ajoutant des anses d'indices inférieurs à n.

Le théorème 2 implique facilement le suivant :

THÉOREME 3 (théorème d'équivalence d'homotopie faible).— Soient M, N deux variétés différentiables avec $\dim M = m < n = \dim N$. On désigne par $R(TM, TN)$ l'espace des applications d'espaces fibrés vectoriels $TM \rightarrow TN$, qui sont des monomorphismes, sur chaque fibre. On donne à cet espace sa topologie naturelle ce qui fait de l'application tangente un morphisme continu :

$$d : \text{Imm}(M, N) \rightarrow R(TM, TN) ;$$

d est une équivalence d'homotopie faible.

(Je rappelle la "bonne définition" des équivalences d'homotopies faibles : si K, X sont deux espaces topologiques, $[K, X]$ désigne l'ensemble des classes d'homotopie d'applications continues $K \rightarrow X$. Tout $g : X \rightarrow T$, application continue, induit, d'une manière naturelle, une application

$$g_* : [K, X] \rightarrow [K, Y] ;$$

g est, par définition, une équivalence d'homotopie faible si, pour tout K , complexe simplicial fini, g_* est une bijection.)

Deux immersions $f_0, f_1 \in \text{Imm}(M, N)$ sont dites régulièrement homotopes si l'une des deux conditions suivantes, équivalentes, est satisfaite :

- (i) f_0 et f_1 sont dans la même composante connexes par arcs, de $\text{Imm}(M, N)$.
- (ii) Il existe

$$F \in \text{Imm}(M \times I, N \times I)$$

compatible avec les projections $M \times I \rightarrow I$ et $N \times I \rightarrow I$ telle que

$$F|_{M \times 1} = f_1 \quad \text{et} \quad F|_{M \times 0} = f_0 .$$

Le théorème 3 implique que f_0 et f_1 sont régulièrement homotopes si et seulement si les applications tangentes

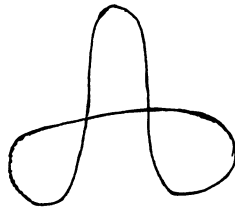
$$df_0, df_1 : TM \rightarrow TN$$

sont homotopes (par une homotopie fibrée).

(Evidemment pour tout $m \leq n$ l'homotopie régulière implique l'homotopie (fibrée) des applications tangentes, mais le fait non trivial est que pour $m < n$ la réciproque est vraie aussi.)

Les théorèmes 1, 2 et 3 sont en général faux si $m = n$. La figure 1 représente une immersion f de S_1 dans R_2 qui est régulièrement homotope à un plongement (donc à une immersion qui peut s'étendre à D_2), telle que f elle-même ne puisse pas s'étendre à D_2 . Ceci contredit les théorèmes 1 et 2.

Figure 1 :



Ceci nous montre tout de suite que le problème "en codimension 1" est beaucoup plus compliqué. Il est aussi très important puisque, par exemple, la conjecture suivante est équivalente à la conjecture de Poincaré, en dimension 3 :

Conjecture.— Soit M_3 une variété différentiable de dimension 3 sans bord, et soit $g : S_2 \rightarrow M_3$ un plongement différentiable qui est régulièrement homotope à une immersion $f : S_2 \rightarrow M_3$ qui peut s'étendre à D_3 . Alors g peut s'étendre à une immersion $D_3 \rightarrow M_3$.

(La conjecture analogue, en dimensions $n > 3$, implique la conjecture de Poincaré en dimension n .)

Le cas le plus simple du problème d'extension en codimension 1 est celui où

$$P = S_1, \quad M = D_2, \quad N = R_2$$

(et $i : S_1 \rightarrow D_2$ est l'inclusion canonique $S_1 = \partial D_2 \subset D$).

Ce problème avait déjà été soulevé par H. Hopf et R. Thom, [8], il y a longtemps, et il vient d'être complètement résolu par S. Blank [1]. C'est ce résultat qui fera l'objet de notre exposé.

2. L'énoncé du théorème de Blank.

On part donc d'une immersion

$$f : S_1 \rightarrow R_2.$$

A f on associe une application continue $f' : S_1 \rightarrow S_1$ définie par

$$f'(x) = \frac{(\text{grad } f)_x}{\|(\text{grad } f)_x\|}.$$

Notre S_1 initial de même que R_2 seront donnés avec des orientations, fixées une fois pour toutes. Le S_1 , but de f' , est le cercle unité de R_2 , lui-même

orienté d'une manière conforme avec celle du disque unité. On peut donc parler du degré de f , qu'on va désigner par $R(f) \in \mathbb{Z}$. $R(f)$ ne dépend que de la classe d'homotopie régulière de f et la détermine même complètement. En fait, $R(f)$ établit une correspondance biunivoque entre les classes d'homotopie régulière de S_1 dans R_2 et $\mathbb{Z}$. Ceci est très élémentaire.

On va considérer seulement des immersions

$$f : S_1 \rightarrow R_2$$

qui sont génériques, c'est-à-dire telles que f admette des points doubles, au plus, avec intersection transversale. On considère l'ouvert $R_2 - f(S_1) \subset R_2$ et ses composantes connexes bornées : $P_1, P_2, \dots, P_n$. Dans chaque P_i on choisit un point $p_i \in P_i$. L'ouvert $A = R_2 - \bigcup_{i=1}^n P_i$ admet comme groupe fondamental le groupe libre avec n générateurs. On va choisir une base de $\pi_1(A)$ de la manière suivante :

On commence par prendre un point de base $*$ $\in A$ (qui sera supposé "très loin" de $f(S_1)$, dans la suite). De plus pour chaque p_i on prend une courbe simple fermée, orientée, de A , passant par $*$, contenant exactement p_i dans son intérieur (et pas les autres p_j , $j \neq i$) et dont l'orientation soit compatible avec celle de son intérieur, orienté comme R_2 . Cette courbe simple, orientée, sera désignée par α_i . On va prendre les α_i sans autre point commun (deux à deux) que $*$ $\in A$. Les classes d'homotopie des α_i que l'on va désigner par $a_i \in \pi_1(A)$ forment une base de $\pi_1(A)$. Dorénavant $\pi_1(A)$ sera muni de cette base.

Sur S_1 on choisit aussi un point de base. Puisque S_1 est orienté, on a un isomorphisme bien déterminé $\pi_1(S_1) = \mathbb{Z}$. Soit $u \in \pi_1(S_1)$ l'élément qui correspond à $1 \in \mathbb{Z}$. L'application $f : S_1 \rightarrow A$ ne respecte pas les points de base

donc l'homomorphisme

$$f_* : \pi_1(S_1) \rightarrow \pi_1(A)$$

n'est pas univoquement déterminé. Il est tout de même déterminé à un automorphisme intérieur de $\pi_1(A)$, près. On va désigner par $m(f) \in \pi_1(A)$ l'image $f_*(u)$. ("Le mot de $\pi_1(A)$ est déterminé par la courbe $f(S_1)$ elle-même.") D'après ce que l'on vient de dire $m(f)$ est déterminé seulement à une conjugaison près. (Il peut être remplacé par $am(f)a^{-1}$ avec $a \in \pi_1(A)$, arbitraire.) La condition nécessaire et suffisante pour que f puisse s'étendre à une immersion $D_2 \rightarrow R_2$ s'exprime en termes de $m(f) \in \pi_1(A)$. Pour pouvoir l'exprimer, on a besoin d'une définition.

Soit $\pi_1(A)$ groupe libre engendré par les "lettres" $a_1, \dots, a_n$, comme auparavant. Soit

$$x = a_{i_1}^{\epsilon_{i_1}} a_{i_2}^{\epsilon_{i_2}} \dots a_{i_k}^{\epsilon_{i_k}}$$

avec $1 \leq i_j \leq n$, $\epsilon_{i_j} = \pm 1$ un "mot" écrit avec les lettres

$a_1, \dots, a_n, a_1^{-1}, \dots, a_n^{-1}$; ($x \in \pi_1(A)$). Si $1 \leq \ell \leq m \leq k$, on dit que

$$y = a_{i_\ell}^{\epsilon_{i_\ell}} \dots a_{i_m}^{\epsilon_{i_m}}$$

est un "sous-mot" de x . Supposons que x contienne un "sous-mot" de la forme

$a_i p a_i^{-1}$ (ou $a_i^{-1} p a_i$) où p est écrit seulement avec des lettres positives (sans employer $a_1^{-1}, \dots, a_n^{-1}$). Donc :

$$x = x_1 a_i p a_i^{-1} x_2 \quad (\text{ou } x = x_1 a_i^{-1} p a_i x_2).$$

L'opération d'effacer $a_i p a_i^{-1}$ (ou $a_i^{-1} p a_i$), donc de remplacer x par $x' = x_1 x_2$ s'appelle un groupement (du sous-mot respectif de x).

Les lettres a_i et a_i^{-1} (ou a_i^{-1} et a_i) sont couplées par le groupement. Elles forment un couple attaché au groupement. Deux groupements sont différents, par définition, si leurs couples sont différents.

Enfin, le mot x est dit réduit si les deux conditions suivantes sont satisfaites :

- (i) x ne contient pas de sous-mot de la forme $a_i a_i^{-1}$ ou $a_i^{-1} a_i$.
- (ii) x n'est pas de la forme zyz^{-1} .

On va voir que sans perdre la généralité, on peut toujours supposer $m(f)$ réduit. Maintenant, on peut énoncer le théorème principal de Blank.

Je rappelle qu'on veut caractériser les immersions $f : S_1 \rightarrow R_2$ qui s'étendent à des immersions $F : D_2 \rightarrow R_2$. Si f se prolonge à F , f est régulièrement homotope à un plongement $S_1 \rightarrow R_2$ donc, puisque $R(f)$ est invariant de l'homotopie régulière $R(f) = \pm 1$. En choisissant convenablement les orientations, on peut prendre $R(f) = +1$, ce que l'on va faire dorénavant.

THÉOREME 4 (théorème d'extension).— Soit $f : S_1 \rightarrow R_2$ une immersion générique, avec $R(f) = +1$. Comme avant on définit

$$A = R_2 - \bigcup p_i$$

où $p_i \in P_i$ et P_i est une composante connexe bornée de $R_2 - f(S_1)$. Pour

$\pi_1(A)$ on prend comme générateurs les a_i décrits plus haut. Soit

$$m(f) = a_{i_1}^{\epsilon_{i_1}} \dots a_{i_k}^{\epsilon_{i_k}} \quad (1 \leq i_j \leq n, \quad \epsilon_{i_j} = \pm 1) \quad \text{un "mot" représentant la classe}$$

d'homotopie de f dans le groupe libre $\pi_1(A)$.

f peut s'étendre à une immersion $F : D_2 \rightarrow R_2$ si et seulement si, la con-

dition suivante est satisfaite :

(C) Par un nombre fini de groupements, $m(f)$ peut se réduire à un mot formé seulement de lettres positives (qui peut être vide).

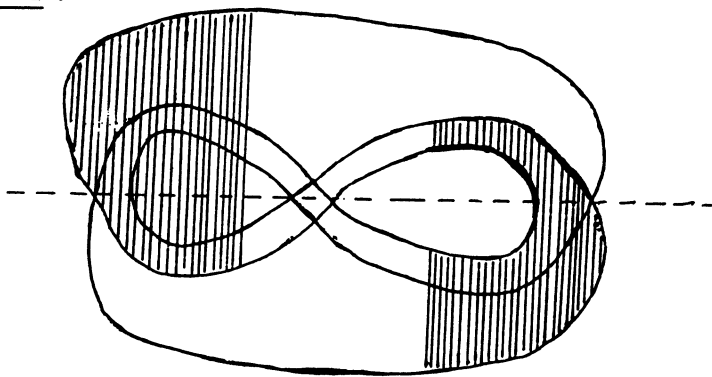
Remarque : $m(f)$ n'est déterminé que modulo l'application (en nombre fini) d'opérations "élémentaires" du type suivant :

- (i) conjugaison par a_i (ou a_i^{-1}) ;
- (ii) effacement de sous-mots du type $a_i a_i^{-1}$ (ou $a_i^{-1} a_i$) ;
- (iii) l'inverse de l'opération (ii) ; mais la condition (C) est clairement invariante par rapport à ces opérations élémentaires.

Supposons maintenant que $f : S_1 \rightarrow R_2$ puisse s'étendre à une immersion $F : D_2 \rightarrow R_2$. Deux extensions $F, F' : D_2 \rightarrow R_2$ seront considérées comme équivalentes si elles diffèrent seulement par un difféomorphisme de D_2 . On pourrait croire alors que si F existe il est unique, à une équivalence près. Ceci est faux, comme le montre l'exemple suivant, dû à Milnor :

On considère l'immersion générique $f : S_1 \rightarrow R_2$ dont l'image est donnée par la figure 2.

Figure 2 :



Il y a deux manières différentes de "remplir" $f(S_1)$ avec un $F(D_2)$, "étalé" sur R_2 . Sur la figure 2 on a commencé les deux manières différentes de "remplir". (Par rotation on augmente la dimension de l'exemple.)

En regardant de près la démonstration du théorème 4, Blank a pu déterminer explicitement le nombre d'extensions possibles (à une équivalence près).

THÉORÈME 5 (théorème de non-unicité).— Soit $f : S_1 \rightarrow R_2$ une immersion générique avec $R(f) = +1$. Soit $m(f)$ un mot RÉDUIT représentant la classe d'homotopie de f dans $\pi_1(A)$.

Les deux nombres suivants sont égaux :

- (i) Le nombre des classes d'équivalence d'extensions de f à des immersions $F : D_2 \rightarrow R_2$.
- (ii) Le nombre de réductions différentes de $m(f)$ à un mot positif par des groupements successifs ; on regarde comme étant les mêmes, deux réductions telles que les deux ensembles de paires (*) (a, a^{-1}) (ou (a^{-1}, a)) correspondant aux groupements respectifs soient égaux (constitués par les mêmes éléments).

Ce théorème implique, entre autre, que pour l'exemple de Milnor, il y a exactement deux extensions possibles. En général, on a :

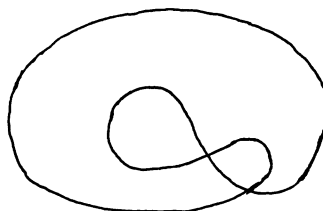
THÉORÈME 6.— Pour tout $n \geq 1$, il existe une immersion $f_n : S_1 \rightarrow R_2$ qui peut s'étendre exactement de n façons à une immersion $D_2 \rightarrow R_2$.

Il suffit de regarder les exemples de la figure 3 et de leur appliquer le théorème 5.

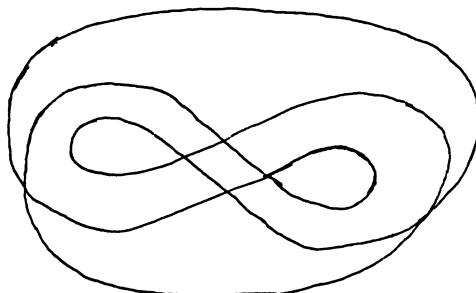
(*) considérées comme paires de lettres du mot $m(f)$.

Figure 3 :

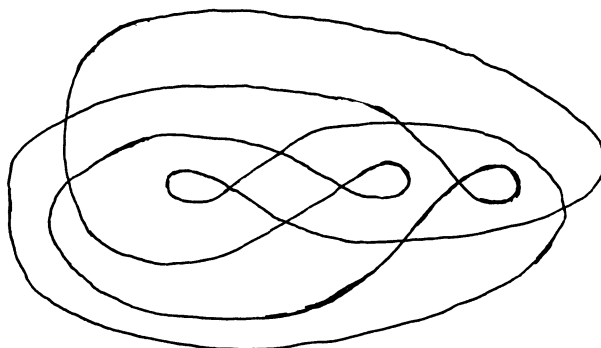
$f_1(s_1)$



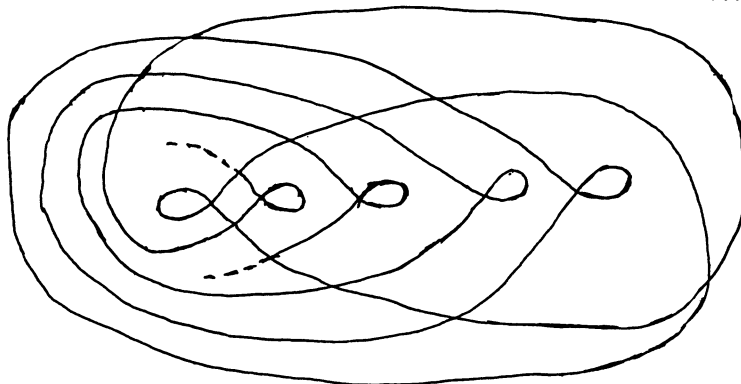
$f_2(s_1)$ = exemple de Milnor



$f_3(s_1)$



$f_n(s_1)$

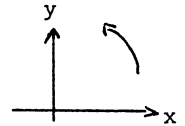
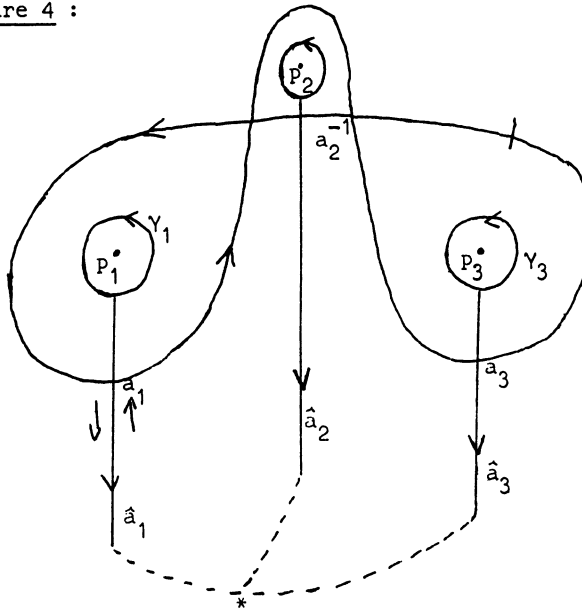


On va donner maintenant la démonstration du théorème 4. Pour ne pas trop allonger l'exposé, un certain nombre de points ont été laissés au lecteur.

3. Préliminaires pour la démonstration du théorème 4.

On va commencer par donner un procédé pour le "calcul" de $m(f)$. Tout d'abord les générateurs $a_i \in \pi_1(A)$ seront représentés par des courbes $S_1 \rightarrow A$ plus explicitement définies.

Figure 4 :



$$m(f) = a_2^{-1} a_1 a_3$$

On commence par choisir, autour de chaque p_i , dans P_i , un petit cercle γ_i , orienté dans le sens direct (figure 4). Pour chaque γ_i on considère un rayon $\hat{a}_i$, à l'extérieur de γ_i , allant vers l'infini. On suppose les $\hat{a}_i$ parallèles, et puis, très loin, près de ∞ , on les courbe un peu, pour rejoindre

le point base $*$ $\in A$. On suppose que les $\hat{\alpha}_i$ coupent $f(S_1)$ transversalement.

Les courbes $\alpha_i : S_1 \rightarrow A$, dont les classes d'homotopie sont nos générateurs $a_i \in \pi_1(A)$ seront décrites de la manière suivante :

on part de $*$ sur $\hat{\alpha}_i$ vers p_i , puis on tourne (dans le sens direct) autour de p_i , sur γ_i , puis on retourne à $*$ sur $\hat{\alpha}_i$.

D'autre part, sur chaque $\hat{\alpha}_i$ on choisit une orientation, une fois pour toutes, allant de γ_i vers $*$.

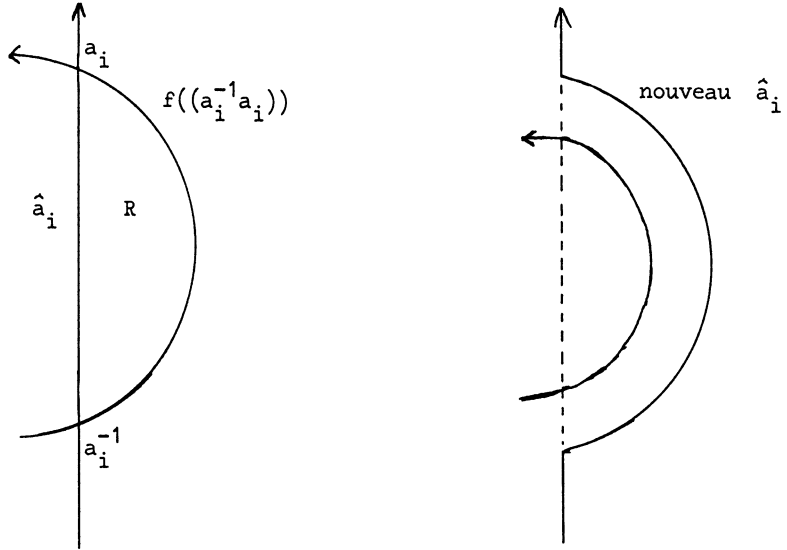
On rappelle que S_1 (donc $f(S_1)$) est orienté (et de telle manière que $R(f) = +1$). Chaque point d'intersection de $\hat{\alpha}_i$ avec $f(S_1)$ est baptisé a_i ou a_i^{-1} suivant que l'orientation de R_2 , autour du point d'intersection, donnée par $\hat{\alpha}_i$ (orienté comme plus haut) et $f(S_1)$ (avec son orientation) est la bonne ou la mauvaise orientation de R_2 (voir figure 4). Sur S_1 on considère les images inverses des intersections de $f(S_1)$ avec les $\hat{\alpha}_i$, avec leurs noms donnés comme avant. Ceux qui ont des noms positifs seront appelés "intersections positives", ceux qui ont des noms négatifs, "intersections négatives" ; (quelquefois : points positifs ou points négatifs). A partir d'un point quelconque de S_1 (le point de base par exemple) on écrit ces noms, de gauche à droite, dans l'ordre avec lequel les points respectifs apparaissent sur S_1 . Le mot qu'on obtient de cette sorte sera désigné par $m(f)$ et c'est facile à voir qu'il représente, effectivement, dans $\pi_1(A)$ la classe d'homotopie de f . Pour le reste de la démonstration, d'ailleurs, toute considération de groupes fondamentaux est inutile et on raisonne tout simplement sur le $m(f)$ qu'on vient de définir.

En changeant, au besoin, les $\hat{\alpha}_i$ et le point de base sur S_1 (le point qui

nous sert de repère pour écrire les images inverses des intersections ...), $m(f)$ devient réduit. Démontrons-le :

Tout d'abord par changement du point de base (de S_1) le remplacement de $m(f) = a_i^{\pm 1} z a_i^{\mp 1}$ pour z se réduit au remplacement de $m(f) = x_1 a_i^{\pm 1} a_i^{\mp 1} x_2$ par $x_1 x_2$.

Figure 5 :



Cette dernière opération peut toujours se réaliser par un changement convenable des $\hat{a}_i$.

Supposons, pour fixer les idées, qu'il s'agit du cas $m(f) = x_1 a_i^{-1} a_i x_2$ (l'autre cas se traite de la même façon). Sur S_1 il y a un arc $(a_i^{-1}, a_i) \subset S_1$, d'extrémités a_i^{-1} et a_i , déterminé univoquement par la condition suivante :

Le sens de parcours de $(a_i^{-1}a_i)$ en allant de a_i^{-1} vers a_i , est le bon sens sur S_1 ; $f|(a_i^{-1}a_i)$ est clairement un plongement, car autrement $f((a_i^{-1}a_i))$ donnerait naissance à des P_j' s, donc à des $\hat{a}_j'$ s qui devraient le couper. Il y aurait donc des lettres dans $m(f)$, entre a_i^{-1} et a_i , ce qui n'est pas le cas. De même $\hat{a}_i$ ne coupe pas $f((a_i^{-1}a_i))$ (pour la même raison). Donc $\hat{a}_i$ et $f((a_i^{-1}a_i))$ déterminent un disque R de R_2 (figure 5). Si

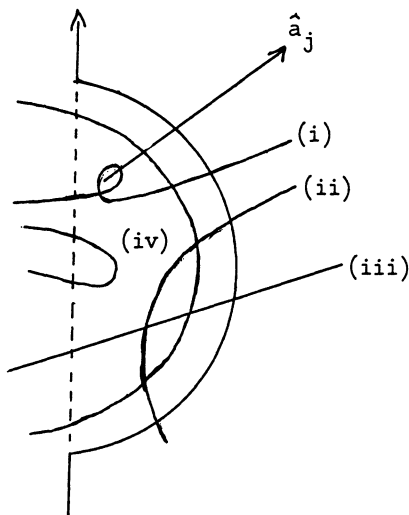
$R \cap f(S_1) = f((a_i^{-1}a_i))$, on change $\hat{a}_i$ comme la figure 5 l'indique, et $a_i^{-1}a_i$ disparaît. Si

$$X = R \cap f(S_1) - f((a_i^{-1}a_i)) \neq \emptyset$$

sa fermeture se décompose en "branches", images d'arcs contenus dans S_1 , qui sont de trois sortes :

- (i) branches à points multiples ou branches simples qui :
- (ii) ne coupent que $f((a_i^{-1}a_i))$
- (iii) coupent $f((a_i^{-1}a_i))$ et $\hat{a}_i$
- (iv) ne coupent que $\hat{a}_i$ (voir figure 6).

Figure 6 :



Les cas (i), (ii) ne peuvent pas se produire, parce que, comme tout à l'heure, on aurait des $\hat{a}_j$ coupant $f((a_i^{-1}, a_i))$.

Dans le cas (iii), il n'y a aucune différence en ce qui concerne le changement de $m(f)$, par rapport à la figure 5. Dans le cas (iv) une autre paire $a_j^{\pm 1} a_j^{\mp 1}$ est tuée en même temps, en supplément gratuit.

Dorénavant, on suppose que $m(f)$ est réduit. (Ce qui sera utile pour démontrer la proposition B, du paragraphe 7.)

4. La nécessité de la condition (C).

Supposons que $f : S_1 \rightarrow R_2$ admette une extension $F : D_2 \rightarrow R_2$ qui est une immersion. On va obtenir une réduction de $m(f)$ à un mot positif, par groupements successifs, en regardant les images inverses $F^{-1}(\hat{a}_j)$.

Chaque composante connexe de $F^{-1}(U \hat{a}_j)$ est un intervalle fermé. Il y a exactement deux types de composantes connexes possibles :

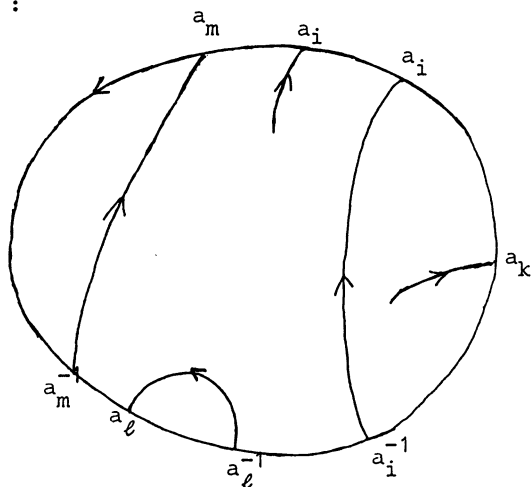
- (i) des intervalles avec les deux extrémités sur S_1 (et l'intérieur dans D_2) ;
- (ii) des intervalles avec une extrémité sur S_1 et le reste à l'intérieur de D_2 .

Les $\hat{a}_i$ sont orientés donc par transport de structure, les composantes connexes de $F^{-1}(U \hat{a}_i)$. Donc à chaque intersection d'un intervalle (de $F^{-1}(U \hat{a}_i)$) correspond, sur cet intervalle, un sens qui "entre", ou qui "sort" de D_2 .

S_1 étant orienté, induit une orientation de D_2 et le fait que $R(f) = +1$ implique que l'immersion $F : D_2 \rightarrow R_2$ conserve les orientations. On voit alors, facilement, que les points d'entrée correspondent aux intersections négatives, et les points de sortie aux intersections positives. Un $F^{-1}(\hat{a}_i)$ qui entre dans D_2

doit bien sortir, puisque son sens de parcours va très loin, vers l'infini. Donc les intervalles du type (ii) sortent au point où ils rencontrent S_1 donc leur unique extrémité sur S_1 est positive. Les intervalles du type (ii) entrent par un bout et sortent par l'autre. Si l'intervalle est contenu dans $F^{-1}(\hat{a}_i)$ les deux extrémités sont baptisées respectivement a_i et a_i^{-1} . La situation se présente donc comme sur la figure 7.

Figure 7 :



Les intervalles (i) vont déterminer les groupements respectifs qui réduisent $m(f)$ à un mot positif. On laisse au lecteur le soin de s'en convaincre.

5. La condition (C) est suffisante.

Ceci est la partie difficile de la démonstration.

Traitons d'abord le cas trivial, où il n'y a pas de lettres négatives du tout.

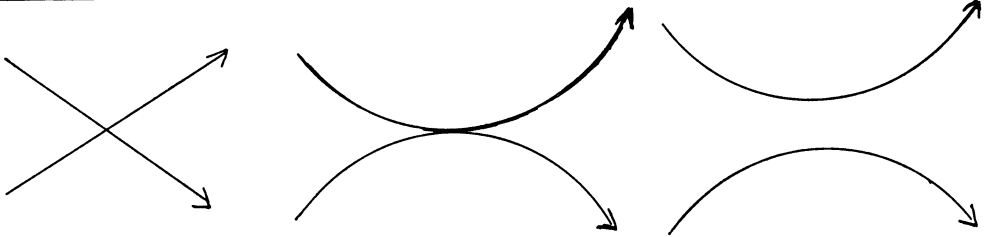
LEMME 1.- Soit $f : S_1 \rightarrow R_2$ une immersion générique telle que :

- (i) $R(f) = +1$;
- (ii) $m(f)$ est un mot constitué seulement par des lettres positives.

Alors f est un plongement ; en particulier il s'étend à une immersion
(même à un plongement) $D_2 \rightarrow R_2$.

Démonstration. On va commencer en donnant une construction qui sera très utile. Chaque immersion générique $f : S_1 \rightarrow R_2$ peut être décomposée en plongements, de la manière suivante : pour chaque point d'auto-intersection de f , par une petite homotopie régulière, on fait les deux arcs qui traversent ce point, tangents. En changeant la paramétrisation on obtient deux immersions "ayant la même image que f " (figure 8). On remarque que cette construction dépend de l'orientation de la source.

Figure 8 :



En répétant le même procédé pour chaque point d'auto-intersection, on remplace $f : S_1 \rightarrow R_2$, pour un nombre fini de plongements (disjoints) $f_i : S_1 \rightarrow R_2$ ($i = 1, \dots, \ell$). Les sources S_1 des f_i sont canoniquement orientées par la source de f et on voit très facilement que :

$$R(f) = \sum_{i=1}^{\ell} R(f_i) .$$

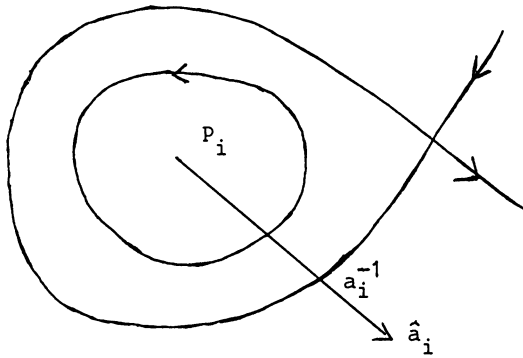
Chaque $R(f_i)$ est ± 1 , suivant que $f_i(S_1)$ va dans le sens direct ou pas.

Appliquons cette construction à l'immersion f de notre énoncé (lemme 1). Je dis qu'on ne trouvera pas de f_i avec $R(f_i) = -1$. Ceci est clair, car les f_i sont en correspondance biunivoque canonique avec une partie de nos régions P_i , et si $R(f_{i_0}) = -1$, vu que $p_i \in \text{intérieur } f_{i_0}$, $\hat{a}_i$ devrait traverser $f_{i_0}(S_1)$. Ceci donnerait une intersection négative donc a_i^{-1} devrait apparaître dans $m(f)$, ce qui contredirait nos hypothèses (voir figure 9). Donc :

$$1 = R(f) = \sum_{i=1}^{\ell} R(f_i) = \ell.$$

Donc $\ell = 1$, donc il n'y a pas d'auto-intersections du tout. c.q.f. à d.

Figure 9 :

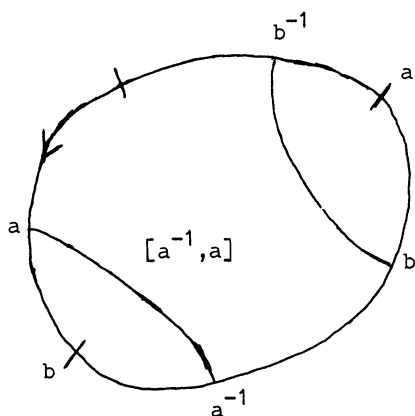


Supposons maintenant que $m(f)$, après un nombre fini de groupements correspondant aux couples $(a_{j_1}, a_{j_1}^{-1})$, $(a_{j_2}, a_{j_2}^{-1})$, ..., $(a_{j_N}, a_{j_N}^{-1})$, se réduit à un mot positif. (Donc $a_{j_1}^{-1}$, $a_{j_2}^{-1}$, ..., $a_{j_N}^{-1}$ sont toutes les lettres négatives de $m(f)$.) Chaque couple a_{j_1} , $a_{j_1}^{-1}$ correspond à une paire de points de S_1 , qu'on

désigne, par abus de langage, par les mêmes lettres a_{j_1} , $a_{j_1}^{-1}$. On remarque que, sur S_1 , un couple $(a_{j_1}, a_{j_1}^{-1})$ ne sépare jamais les points d'un autre. Donc, si S_1 est considéré comme bord de D_2 (le D_2 sur lequel on veut étendre f), chaque paire de points a_{j_1} , $a_{j_1}^{-1}$ peut être reliée par un segment, qu'on va désigner par $[a_{j_1}^{-1}, a_{j_1}]$ de telle manière que ces segments soient deux à deux disjoints (figure 10).

Nos a_{j_i} , $a_{j_i}^{-1}$ correspondent aussi à des points sur $\hat{a}_{j_i}$ que l'on va, par un autre abus de langage, désigner aussi a_{j_i} , $a_{j_i}^{-1}$. Dorénavant, par un abus de langage supplémentaire, j_i va devenir i , donc $a_{j_i}^{\pm 1}$ va devenir a_i et $\hat{a}_{j_i}$ va devenir $\hat{a}_i$.

Figure 10 :



$$m(f) = aba^{-1}bab^{-1}$$

On rappelle que, sur $\hat{a}_i$, on a un ordre choisi (allant vers ∞) que l'on désigne par $<$.

DÉFINITION.- Si $a_i > a_i^{-1}$ (sur $\hat{a}_i$) le couple (a_i, a_i^{-1}) est dit positif ; (de même le segment respectif).

Si $a_i < a_i^{-1}$ (sur $\hat{a}_i$) le couple (a_i, a_i^{-1}) est dit négatif (de même le segment respectif).

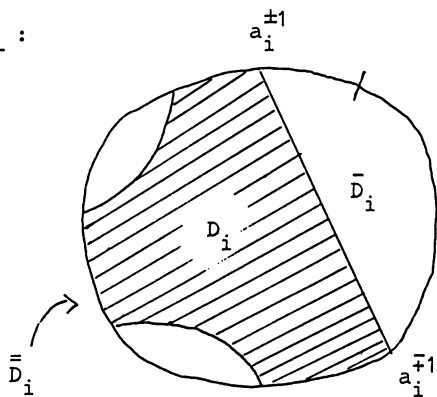
Soit maintenant $a_i^{\pm 1}, a_i^{\mp 1}$ l'ordre dans lequel les lettres a_i, a_i^{-1} apparaissent dans S_1 . Au paragraphe 3, on a défini sur S_1 , un arc (a_i^{-1}, a_i) qui est déterminé par les deux conditions suivantes :

- (i) Il a a_i et a_i^{-1} comme extrémités.
- (ii) En le parcourant dans le bon sens, on va de $a_i^{\pm 1}$ vers $a_i^{\mp 1}$.

Soit $\bar{D}_i \subset D_2$ le disque de bord $(a_i^{-1}, a_i) + [a_i^{-1}, a_i]$. L'autre disque sera désigné par $\bar{D}_i$.

L'ensemble des $[a_i^{-1}, a_i]$ décompose D_2 en une série de disques. Celui de ces disques qui a $[a_i^{-1}, a_i]$ sur son bord, et est contenu dans $\bar{D}_i$ sera désigné par D_i (voir figure 11).

Figure 11 :



(Donc les D_i dépendent du choix du point base de S_1 .)

On définit

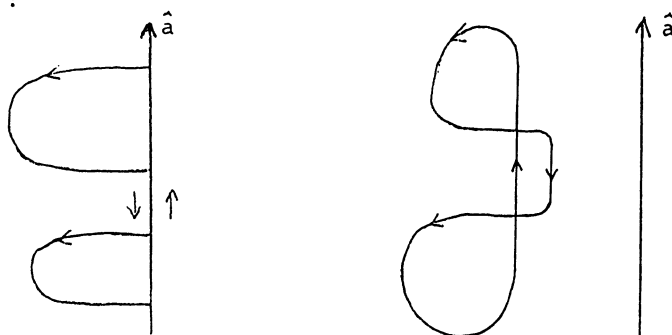
$$f^* : (S_1 \cup (\cup [a_i^{-1}, a_i])) \rightarrow R_2$$

de la manière suivante :

- (i) $f^*|_{S_1} = f$;
 (ii) $f|_{[a_i^{-1}, a_i]}$ est un difféomorphisme de $[a_i^{-1}, a_i]$ avec le segment de $\hat{a}_i$ d'extrémités a_i^{-1} et a_i (donc f^* applique le $a_i^{\pm 1}$ de S_1 sur le $a_i^{\pm 1}$ de $\hat{a}_i$).

On va considérer maintenant les $f^*|\partial D_i$, qu'on rend différentiables par un simple arrondissement d'angles. Ce sont alors des immersions, pas nécessairement génériques (voir figure 12).

Figure 12 :

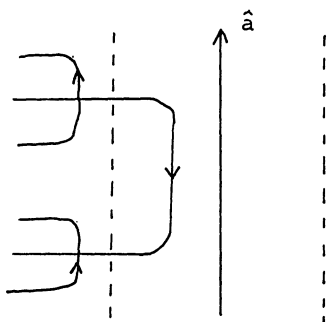


On va rendre les $f^*|\partial D_j$ génériques, par un procédé canonique, et les résultats seront désignés par $f_j : S_1 \rightarrow R_2$.

On remarque que $f^*(\partial D_i)$ se compose d'arcs provenant de $f(S_1)$ et d'intervalles (segments) contenus dans les $\hat{a}_i$. Les intervalles sont les images des $[a_j^{-1}, a_j]$ et correspondent donc à des couples positifs ou négatifs. Le procédé canonique consiste à retirer les intervalles des $\hat{a}_i$, sans introduire des nouvelles intersections. On retire les intervalles positifs d'abord, de telle façon, qu'à la fin un voisinage tubulaire des $\hat{a}$ contient tous les intervalles négatifs,

et seulement les intervalles négatifs (figure 13).

Figure 13 :



On remarque que l'image de f_i est constituée par des arcs provenant de $f(S_1)$ et des intervalles (segments) provenant des $[a_i^{-1}, a_i] \subset \hat{a}_i$.

LEMME 2.- Pour chaque D_j , f_j est un plongement.

Ce lemme constitue la partie la plus importante de la démonstration, que l'on va reléguer au paragraphe suivant.

Montrons seulement, pour le moment, comment le lemme 2 nous permet de finir la démonstration.

Tout d'abord, vu que le rôle des $\bar{D}_j$ et $\bar{\bar{D}}_j$ peut être interchangé, par un changement du point de base, le lemme 2 implique que f^* , restreint à chacun des bords des disques, deux à deux disjoints, dans lesquels $\cup [a_i^{-1}, a_i]$ décompose D_2 , est un plongement. En fait, vue notre hypothèse (C) il y a exactement $N + 1$ tels disques, nos D_j ($j = 1, \dots, N$), et en plus, un dernier disque, correspondant à ce qui reste de $m(f)$ à la fin de la réduction. On va le désigner par D_{N+1} . (Je rappelle que N est le nombre des groupements et des couples.) On peut prolonger f^* à $F : D_2 \rightarrow R_2$ en demandant que $F(D_i)$ soit l'intérieur de $f^*(\partial D_i)$. Je dis qu'on obtient de cette manière une immersion (le problème se pose, évidemment

au voisinage de $U [a_i^{-1}, a_i]$, seulement).

Le fait que F est une immersion, résulte facilement du lemme suivant :

LEMME 3.- Soient $[a_i^{-1}, a_i] \subset \partial D_j$, $j = 1, \dots, N + 1$, et b_i^{-1} , b_i deux petits segments sur S_1 , ayant une extrémité dans a_i^{-1} et a_i , respectivement, pas contenus dans ∂D_j . Alors, dans R_2 , $f(b_i^{-1})$ et $f(b_i)$ ne sont pas à l'intérieur de $f^*(\partial D_j)$.

Démonstration. Si l'un des segments était à l'intérieur et l'autre à l'extérieur on contredirait le fait que chaque $f^*|\partial D_k$ est un plongement. Si les deux étaient à l'intérieur, on aurait sur $\partial D_j - U [a_i^{-1}, a_i]$ des points négatifs (figure 14).

Figure 14 :



Reste donc à prouver le lemme 2.

6. La démonstration du lemme 2.

On a défini au paragraphe précédent les couples positifs et négatifs. Cette définition va jouer maintenant un rôle fondamental.

Esquisse de la démonstration.1ère étape :

Par un dévissage facile on prouve les formules suivantes :

1. Si (a_i^{-1}, a_i) est positif :

$$R(f) = R(\bar{f}_i) + R(\bar{\bar{f}}_i) - 1 ,$$

donc $R(\bar{f}_i) + R(\bar{\bar{f}}_i) = 2$

où $\bar{f}_i = f^*|_{\partial \bar{D}_i}$ et $\bar{\bar{f}}_i = f^*|_{\partial \bar{\bar{D}}_i}$.

De même, si (a_i^{-1}, a_i) est négatif, $R(f) = R(\bar{f}_i) + R(\bar{\bar{f}}_i) + 1$, donc

$$R(\bar{f}_i) + R(\bar{\bar{f}}_i) = 0 .$$

$$2. \quad R(\bar{\bar{f}}_i) = R(f_i) + \sum_{\oplus} [R(\bar{\bar{f}}_j) - 1] + \sum_{\ominus} [R(\bar{\bar{f}}_j) + 1]$$

où $\sum_{\oplus}$ est la somme de tous les intervalles $[a_j^{-1}, a_j]$ correspondant aux couples positifs dans ∂D_i , et $\sum_{\ominus}$ est la somme de tous les intervalles correspondant aux couples négatifs de ∂D_i à l'exception de (a_i^{-1}, a_i) (si c'est le cas). On désigne le nombre de ces intervalles (négatifs, sans (a_i^{-1}, a_i)) par L , et le nombre de tous les intervalles négatifs, par Q .

Donc, si (a_i^{-1}, a_i) est positif $Q = L$, sinon $Q - 1 = L$.

2e étape :

On démontre la proposition suivante :

PROPOSITION A.- Si l'on applique à $f_i : D_i \rightarrow R_2$ le procédé de décomposition en plongements (lemme 1, paragraphe 5), on trouve P cercles C_j avec

$$R(C_j) = +1 \text{ et } N \text{ cercles } C_j \text{ avec } R(C_j) = -1 :$$

$$R(f_i) = \sum R(C_j) + \sum R(C'_j) = P - N .$$

De plus : $N \leq Q$.

3e étape :

On démontre la proposition suivante :

PROPOSITION B.- $P + (Q - N) > 0$ (donc : il est impossible que $P = 0$
et $Q = N$ en même temps).

4e étape :

On démontre la proposition suivante :

PROPOSITION C.- Si (a_i^{-1}, a_i) est positif

$$R(\bar{f}_i) = 1 .$$

Si (a_i^{-1}, a_i) est négatif

$$R(\bar{f}_i) = 0 .$$

Démonstration. Il suffit de prouver

$$R(\bar{f}_i) \geq 1 \quad (\text{dans le cas positif})$$

$$\text{et} \quad R(\bar{f}_i) \geq 0 \quad (\text{dans le cas négatif}).$$

En effet, par changement de point de base la même chose sera vraie pour $\bar{f}_i$ et
alors les formules 1 de l'étape 1, impliquent la proposition C.

On remarque que, dans la formule 2 : $j < i$. Supposons, par induction, nos
inégalités valables pour $j < i$.

Les formules 2 donnent :

$$R(\bar{f}_i) \geq P - N + L .$$

Ceci et la proposition B impliquent la proposition C.

La formule 2 donne alors :

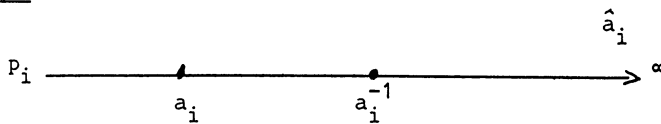
$$R(f_i) = 1 - Q.$$

5e étape : (surprise !)

On démontre le fait suivant :

LEMME 4.- Si $m(f)$ peut être réduit à un mot positif par une série de groupements, tous les couples du groupement sont positifs. (C'est-à-dire la situation de la figure 15 n'arrive jamais.)

Figure 15 :



Ce lemme admis, on a $Q = 0$, donc $R(f_i) = P = 1$, donc en décomposant f_i , on trouve un seul plongement, donc f_i est un plongement.

Ceci finit l'esquisse de la démonstration. On va prouver maintenant les propositions A et B, et le lemme 4.

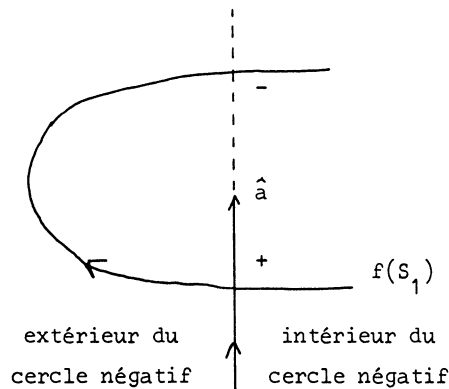
7. Fin de la démonstration.

Démonstration de la proposition A. On applique à f_i le procédé de décomposition en plongements (du lemme 1). On doit montrer qu'on a au plus un cercle négatif (avec $R = -1$), pour chaque couple négatif. On peut énoncer le lemme suivant, qui se démontre de la même manière que le lemme 1 (vu que, d'après nos hypothèses, tous les points négatifs sont couplés) :

LEMME 5.- Si l'on applique le procédé de décomposition en plongements à f_i , il n'y aura aucun point $p_i \in P_i$ contenu à l'intérieur d'un cercle négatif. En particulier, il n'y aura pas de cercle négatif formé avec des arcs de $(\partial D_i) \cap S_1$ seulement (les cercles négatifs qui apparaissent vont utiliser les $[a_i^{-1}, a_i]$).

Donc tout cercle négatif va contenir des sous-segments de $[a_i^{-1}, a_i]$. Supposons qu'un cercle négatif contienne une intersection positive. Alors, il doit contenir une intersection négative aussi (figure 16).

Figure 16 :



(car on voit facilement que les $\hat{a}_i$, au point d'intersection positive, pénètrent à l'intérieur du cercle négatif. Au point de sortie, il y aura une intersection négative).

Donc chaque cercle négatif contient au moins une intersection négative.

Nous allons montrer qu'à chaque intersection négative (ou plutôt à chaque intervalle négatif) il correspond un cercle négatif au plus. Ceci se démontre en

faisant appel au procédé "canonique" qui nous a permis de définir f_i . On rappelle qu'il y a un voisinage tubulaire T des $\hat{a}_i$ qui contient, exactement, tous les intervalles négatifs.

Pour chaque intervalle négatif, il y a deux petits arcs de $f_i(S_i)$ qui entrent et sortent de T en parcourant l'intervalle respectif. Si ces arcs ne traversent pas $\hat{a}$, ils ne contiennent pas d'autre intervalle négatif, donc ils ne peuvent engendrer plus de cercles que d'intervalles négatifs. S'ils traversent $\hat{a}$ (auquel cas ils pourront contenir encore un intervalle négatif, de l'autre côté de $\hat{a}$), ils ne peuvent pas provenir d'un cercle négatif (autrement $\hat{a}$ entrerait à l'intérieur du cercle négatif, donc à la sortie, donnerait un point négatif ; mais tout point d'intersection de f_i avec $\hat{a}$ est positif). La proposition A est démontrée.

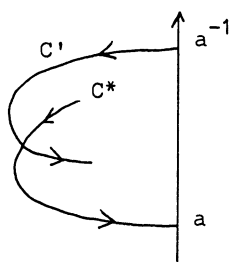
Démonstration de la proposition B. On va prouver une assertion plus forte :

Soit a^{-1} un point négatif de ∂D_i , appartenant à un couple négatif. Soit C le cercle de la décomposition de f_i qui le contient et soit a le premier point positif de $\hat{a}$ sur C , après a^{-1} (en supposant qu'un tel point existe). Soit ℓ l'arc sur ∂D_i , allant de a^{-1} à a . Supposons que l'on applique à ℓ le processus de décomposition en plongements. Il est impossible que la situation suivante arrive : il n'y a pas de cercle positif, et chaque cercle négatif contient exactement un segment négatif.

Démonstration. ℓ n'est pas un plongement car autrement le cercle négatif $C = \ell + [a, a^{-1}]$ serait $f_i(S_i)$. Alors $m(f)$ ne serait pas réduit (si d'autres rayons coupaient C , on aurait sûrement des points négatifs ...).

Soit $C' = C \mid \overset{\curvearrowright}{a^{-1}a}$. Puisque ℓ n'est pas un plongement, il y a un autre cercle de la décomposition de ∂D_i , C^* , tangent à C' . Puisque $R(C^*) = -1$, les intérieurs de C et C^* ont des points communs dans le voisinage de leur point de tangence (figure 17).

Figure 17 :



Si C^* ne coupait pas $\hat{a}$, on trouverait un cercle négatif formé par des arcs de $f(S_1)$, ce qui est impossible d'après le lemme 5. Donc C^* touche $\hat{a}$, on peut répéter notre raisonnement avec un C' plus court, e.a.d.s.

Démonstration du lemme 5. En utilisant les techniques développées dans les étapes précédentes (de la démonstration du lemme 2), on peut montrer que :

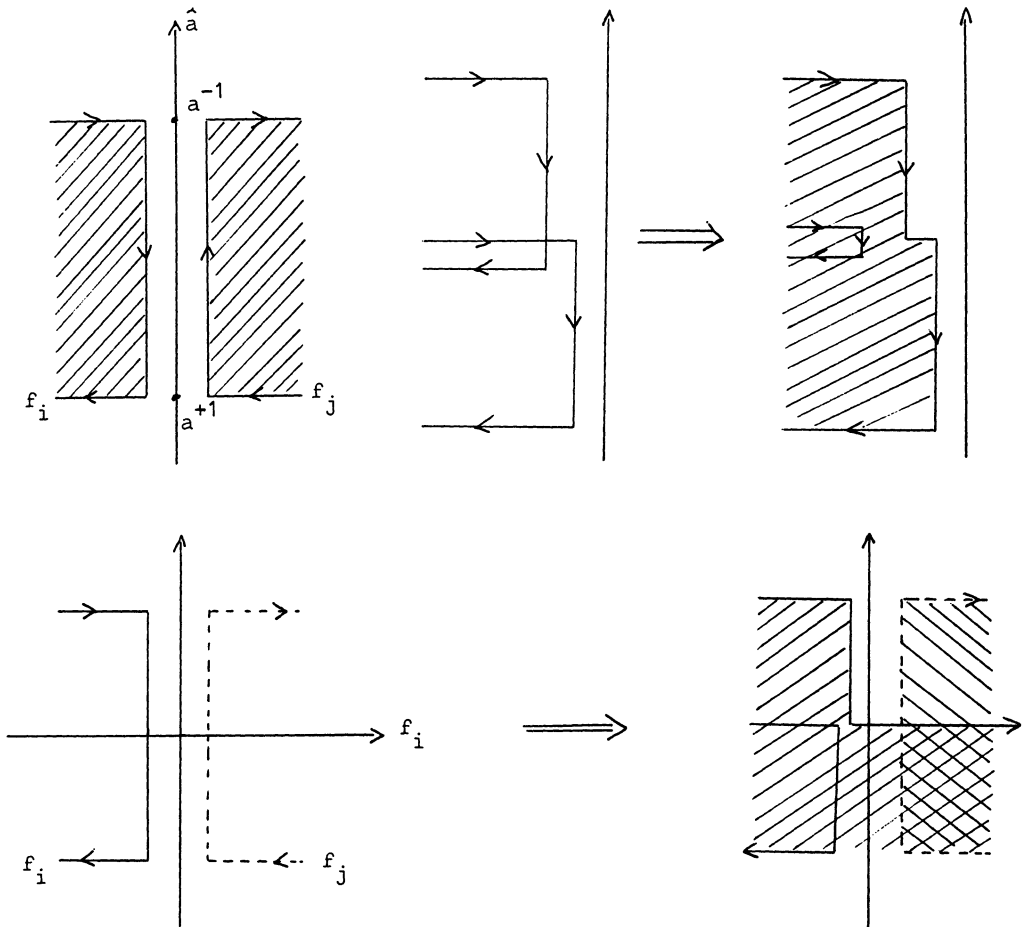
Un cercle négatif obtenu dans la décomposition de ∂D_i ne contient aucun segment positif et, réciproquement, un cercle positif obtenu dans la décomposition de ∂D_i ne contient aucun segment négatif.

On laisse au lecteur le soin de faire cette démonstration.

En utilisant ce fait, on peut voir la chose suivante : soient $[a^{-1}, a]$ un segment négatif et D_i, D_j les deux disques de la décomposition de D_2 (par les $[a_i^{-1}, a_i]$) qui le contiennent, sur leur bord.

On décompose f_i et f_j en cercles. Parce que les cercles qui touchent $[a^{-1}, a]$ sont tous négatifs, la réunion de leurs intérieurs va recouvrir un petit voisinage tubulaire de $[a^{-1}, a]$. (Voir figure 18, qui montre le genre de situations qui peuvent arriver. En la regardant de près, le lecteur pourra donner, sans beaucoup de difficulté, une démonstration générale.)

Figure 18 :



Je dis que $[a^{-1}, a] \subset \hat{a}$ ne touche pas à la composante infinie de $R_2 - f(S_1)$. En effet, un voisinage de $[a^{-1}, a]$ est, comme on vient de le voir contenu dans une réunion d'intérieurs de cercles négatifs. Partout où les bords de ces cercles sont constitués par des segments, les segments sont négatifs aussi, donc leurs voisinages sont recouverts aussi par nos cercles négatifs, et d'autres encore, e.a.d.s.

Soit donc P_j la composante bornée de $R_2 - f(S_1)$, contenant a^{-1} sur son bord. Elle contient aussi un morceau de $[a^{-1}, a]$.

Par le même raisonnement que tout à l'heure, P_j est recouverte par des intérieurs de cercles négatifs, donc $p_j \in P_j$ se trouve à l'intérieur d'un cercle négatif, ce qui contredit le lemme 5.

8. Quelques résultats apparentés.

En codimension 0, les singularités les plus simples qui puissent intervenir sont celles du type :

$$X_1 = x_1^2, \quad X_i = x_i \quad (i = 2, \dots, n).$$

Si $f : M_n \rightarrow N_n$ est une application différentiable ayant seulement des singularités du type décrit et si les composantes connexes de l'ensemble singulier $\Sigma \subset M_n$ sont des sous-variétés parallèles à ∂M_n (supposé connexe et non vide), on dit que f est une pseudo-immersion. Si l'espace des pseudo-immersions $M_n \rightarrow N_n$ est muni d'une topologie convenable, sa projection naturelle sur $\text{Imm}(\partial M_n, N_n)$ est une fibration de Serre [6].

On en déduit [6] l'existence d'une pseudo-immersion $S_3 \rightarrow R_3$, dont l'ensem-

ble singulier est constitué par un nombre pair, $2k$, de sphères de dimension 2, parallèles les unes aux autres.

Par d'autres techniques, on prouve, pour tout disque d'homotopie M_3 , l'existence d'une immersion $F : M_3 \rightarrow R_3$, telle que $F|_{\partial M_3}$ puisse s'étendre à une immersion $C : D_3 \rightarrow R_3$, [2].

Enfin, en revenant au théorème 4, on peut se demander ce qui se passe si R_2 est remplacé par une variété de dimension 2, M_2 , fermée. Vu que le problème peut toujours s'attaquer par les revêtements universels, les seuls cas qui restent à étudier sont $M_2 = S_2$ et $M_2 = P_2(R)$.

BIBLIOGRAPHIE

- [1] S. BLANK - Thèse, (Brandeis 1967).
- [2] S. BLANK et V. POÉNARU - Papiers secrets.
- [3] M. HIRSCH - Immersions of manifolds, Trans. Ann. Math. Soc. 93 (1959), 242-276.
- [4] V. POÉNARU - Sur la théorie des immersions, Topology 1 (1962), 81-100.
- [5] V. POÉNARU - Regular homotopy ..., (miméographié) Harvard 1964.
- [6] V. POÉNARU - Regular homotopy in codimension 1, Ann. of Math. (1966), 257-265.
- [7] S. SMALE - The classification of immersions ..., Ann. of Math. 69 (1959), 327-344.
- [8] R. THOM - Remarques sur les ... inégalités différentielles globales, Bull. Soc. Math. France (1959), 455-461.
- [9] R. THOM - La classification des immersions, Séminaire Bourbaki, 1957-1958, exposé 157.
- [10] A. HAEFLIGER et V. POÉNARU - La classification des immersions combinatives, Publications mathématiques de l'I. H. E. S., 1965.