

Sur les théorèmes de de Rham

Par ANDRÉ WEIL (Chicago)

La démonstration actuellement la plus satisfaisante des célèbres théorèmes de de Rham est celle qui résulte de la théorie de l'homologie de H. Cartan, qui les renferme, ainsi que le théorème de dualité de Poincaré, comme cas particuliers. Mais cette théorie n'a fait l'objet que de publications partielles sous forme de notes de cours miméographiées¹). A son origine se trouvent d'ailleurs, d'une part un mémoire de Leray, et d'autre part justement une démonstration des théorèmes de de Rham que je communiquai à Cartan en 1947. A défaut d'autre utilité, celle-ci peut encore servir d'introduction aux méthodes de Cartan; et c'est avant tout à ce titre que je la présente ici, avec des améliorations dont je dois quelques-unes à G. de Rham et à N. Hamilton; j'y joins une démonstration (datant aussi de 1947) du fait que tout espace possédant un recouvrement d'un certain type (dit «topologiquement simple») a même type d'homotopie que le nerf de ce recouvrement.

§ 1. Construction d'un recouvrement simple

Soit $(X_i)_{i \in I}$ une famille de parties d'un espace E , à ensemble d'indices I quelconque; on dit, comme on sait, que cette famille est *localement finie* si tout point de E a un voisinage qui ne rencontre qu'un nombre fini des X_i ; si E est localement compact, il revient au même de dire que toute partie compacte de E ne rencontre qu'un nombre fini des X_i . Nous conviendrons une fois pour toutes, si $(X_i)_{i \in I}$ est une famille localement finie et si $J \subset I$, de poser $X_J = \bigcap_{i \in J} X_i$; l'ensemble N des parties non vides J de I telles que X_J ne soit pas vide s'appelle le nerf de la famille (X_i) ; si $J \in N$, J est finie.

L'objet de notre étude sera une variété différentiable V de dimension n , «paracompacte» c'est-à-dire dont toute composante connexe est dénombrable à l'infini; il revient au même de dire que V admet un recouvrement localement fini par des «cartes», c'est-à-dire par des parties ouvertes munies chacune d'un isomorphisme différentiable sur une partie

¹) Cours de Harvard, 1948; Séminaire de l'E. N. S., Paris 1948—1949 et 1950—1951.

ouverte de $\mathbf{R}^n$. Le mot «différentiable» sera toujours pris au sens «indéfiniment différentiable» (ou «de classe C^∞ »). Cela n'est pas vraiment une restriction si on tient compte du théorème de Whitney d'après lequel toute variété de classe C^n , pour $n \geq 1$, admet un homéomorphisme de classe C^n sur une variété de classe C^∞ ; d'ailleurs la méthode qui va être exposée s'applique aussi aux variétés de classe C^n pour $n \geq 2$.

Notre outil principal sera un recouvrement $\mathcal{U} = (U_i)_{i \in I}$ localement fini de V par des ensembles ouverts U_i relativement compacts, qui devra avoir de plus la propriété suivante: chaque ensemble non vide $U_j = \bigcap_{i \in J} U_i$ possède une «rétraction différentiable» c'est-à-dire une application différentiable φ_j de $U_j \times \mathbf{R}$ dans U_j telle que $\varphi_j(x, t) = x$ chaque fois que $x \in U_j$ et que $t \geq 1$, et que φ_j soit constante sur $U_j \times]-\infty, 0]$. Un tel recouvrement, muni de la donnée des rétractions φ_j , sera dit *différentiablement simple*.

Pour construire un tel recouvrement, on peut, comme le fait de Rham²⁾, se servir d'un ds^2 , mais il est peut-être plus élémentaire de procéder comme suit. Partons d'un recouvrement localement fini de V par des cartes ouvertes relativement compactes V_i ; à V_i sera donc attaché un isomorphisme différentiable de V_i sur une partie ouverte de $\mathbf{R}^n$ au moyen de «coordonnées locales» $t_1^{(i)}, \dots, t_n^{(i)}$. On peut alors, pour chaque i , définir des ouverts W_i, W'_i et une fonction f_i différentiable sur V de manière que les W_i forment encore un recouvrement de V , que l'on ait $\overline{W_i} \subset W'_i$ et $\overline{W'_i} \subset V_i$, et que f_i ait la valeur 1 sur $\overline{W_i}$ et 0 en dehors de W'_i . Posons $f_{i0} = f_i$, et désignons par f_{ij} la fonction égale à $f_i t_j^{(i)}$ dans V_i et à 0 en dehors de V_i ; l'ensemble des fonctions f_{ij} pour $0 \leq j \leq n$, et pour toutes les valeurs de i , détermine une application de V dans l'espace $\mathbf{R}^{(A)}$, où A est l'ensemble des couples (i, j) ; on sait qu'on désigne ainsi l'espace vectoriel des applications de A dans $\mathbf{R}$ qui prennent la valeur 0 partout sauf en un nombre fini d'éléments de A . De plus, l'application (f_{ij}) de V dans $\mathbf{R}^{(A)}$ détermine sur toute partie ouverte relativement compacte Z de V un isomorphisme différentiable de Z sur une sous-variété d'un sous-espace vectoriel de dimension finie de $\mathbf{R}^{(A)}$. On pourra donc simplifier le langage en identifiant V avec son image dans $\mathbf{R}^{(A)}$. Sur $\mathbf{R}^{(A)}$, nous mettrons une structure d'espace métrique («pré-hilbertien») au moyen de la distance $d(x, y) = [\sum_{i,j} (x_{ij} - y_{ij})^2]^{\frac{1}{2}}$; elle

²⁾ Cf. G. de Rham, Complexes à automorphismes et homéomorphie différentiable, Ann. Gren. 2 (1950) p. 51. Ce dernier exposé, comme ma démonstration de 1947, reste limité au cas compact; mais c'est de Rham qui m'a indiqué la possibilité d'étendre l'une et l'autre méthode aux variétés non compactes.

fait de tout sous-espace de dimension finie de $\mathbf{R}^{(4)}$ un espace euclidien. D'après ce qui précède, la distance de $\bar{W}_i$ à $V - W'_i$ est ≥ 1 , puisque la coordonnée x_{i0} a la valeur 1 sur le premier ensemble et 0 sur le second.

Pour tout $x \in V$, désignons par T_x la variété linéaire tangente à V en x , et par P_x la projection orthogonale de $\mathbf{R}^{(4)}$ sur T_x , considérée comme application linéaire de $\mathbf{R}^{(4)}$ sur T_x ; et désignons par $U(x, r)$ l'intersection de V avec la boule ouverte de centre x et de rayon r ; si $x \in \bar{W}_i$ et $r < 1$, on aura $U(x, r) \subset W'_i$; donc $U(x, r)$ est relativement compact pourvu que $r < 1$.

Soit $x \in \bar{W}_i$; soit E un espace vectoriel de dimension finie contenant $\bar{W}'_i$. En prenant dans E des coordonnées orthogonales d'origine x , les n premiers vecteurs coordonnés étant choisis dans T_x , on voit que x possède un voisinage ouvert U contenu dans W'_i et ayant les propriétés suivantes: (a) quel que soit $y \in U$, P_y induit sur U un isomorphisme différentiable (c'est-à-dire une application biunivoque, partout de rang n) de U sur son image $U_y = P_y(U)$ dans T_y ; (b) quels que soient y, z_1, z_2 dans $\bar{U}$, on a $d(z_1, z_2) < 2d(P_y(z_1), P_y(z_2))$; (c) quel que soit $z_0 \in U$, $d(z_0, z)^2$ est une fonction convexe de $P_y(z)$ dans U_y . En effet, cette dernière condition signifie que la matrice des dérivées secondes de $d(z_0, z)^2$ par rapport aux coordonnées de $P_y(z)$ dans T_y est la matrice d'une forme quadratique définie positive; or, dès que U est assez petit, cette matrice est aussi voisine qu'on veut de sa valeur pour $y = z_0 = z = x$, valeur qui n'est autre que la matrice unité. Soit alors K une partie compacte de V ; recouvrons K par des ensembles U_α en nombre fini ayant les propriétés (a), (b), (c); et prenons $r(K) > 0$ et < 1 tel que $U(x, r(K))$ soit contenu dans l'un des U_α quel que soit $x \in K$; ainsi $U(x, r(K))$ aura les propriétés (a), (b), (c) quel que soit $x \in K$. De plus, pour $x \in K$ et $r = r(K)$, la projection $P_x[U(x, r)]$ de $U(x, r)$ sur T_x contiendra tous les points de T_x à distance $< r/2$ de x ; en effet, si z' est un point frontière de cette projection, z' sera point limite de points $z'_\nu = P_x(z_\nu)$ avec $z_\nu \in U(x, r)$; $U(x, r)$ étant relativement compact sur V , on pourra remplacer les z_ν par une suite partielle ayant une limite z sur V . Comme P_x est un isomorphisme différentiable de $U(x, r)$ sur son image, tout point intérieur de $U(x, r)$ se projette sur un point intérieur de $P_x[U(x, r)]$; donc z est un point frontière de $U(x, r)$, et on a $d(x, z) = r$, d'où $d(x, z') > r/2$ en vertu de (b). Montrons maintenant que, si $x \in K$, $0 < r \leq r(K)/4$, et $y \in U(x, r)$, P_x induit sur $U(y, r)$ un isomorphisme différentiable de $U(y, r)$ sur une partie convexe de T_x . Comme on a $U(y, r) \subset U(x, 2r)$, le seul point à démontrer est la convexité de $P_x[U(y, r)]$. Or c'est là l'ensemble des points $z' = P_x(z)$

pour $z \in U(x, r(K))$ et $d(y, z)^2 < r^2$. Considérons deux tels points $z'_1 = P_x(z_1)$, $z'_2 = P_x(z_2)$; on a, pour $h = 1$ et $h = 2$, $d(z_h, x) < 2r$, donc $d(z'_h, x) < 2r$, donc aussi $d(z', x) < 2r \leq r(K)/2$ quel que soit z' sur le segment de droite qui joint z'_1 et z'_2 dans T_x ; ce segment est donc contenu dans $P_x[U(x, r(K))]$; comme $d(y, z)^2$ est une fonction convexe de $z' = P_x(z)$ dans ce dernier ensemble, c'est une fonction convexe de z' sur le segment qui joint z'_1 et z'_2 ; la valeur de cette fonction étant $< r^2$ aux extrémités du segment, elle l'est aussi sur tout le segment; celui-ci est donc bien contenu dans $P_x[U(y, r)]$.

Cela étant, choisissons pour chaque i des points $x_{i\lambda}$ de $\overline{W}_i$, en nombre fini, tels que les ensembles $U_{i\lambda} = U(x_{i\lambda}, r(\overline{W}'_i)/4)$ forment un recouvrement de $\overline{W}_i$; je dis que les $U_{i\lambda}$ forment un recouvrement différentiablement simple de V . Comme on a $x_{i\lambda} \in \overline{W}_i$ et $r(\overline{W}'_i)/4 < 1$, on a $U_{i\lambda} \subset W'_i$, donc les $U_{i\lambda}$ sont relativement compacts et forment un recouvrement localement fini de V . Soit x un point commun à des ensembles $U_{i\lambda}, U_{j\mu}, U_{k\nu}, \dots$, en nombre nécessairement fini; soit r le plus grand des nombres $r(\overline{W}'_i), r(\overline{W}'_j), r(\overline{W}'_k), \dots$; supposons par exemple qu'on ait $r = r(\overline{W}'_i)$. Alors chacun des ensembles $U_{i\lambda}, U_{j\mu}, \dots$, est de la forme $U(y, r')$, avec $y \in U(x, r')$ et $r' \leq r(\overline{W}'_i)/4$; ils sont tous contenus dans $U(x, r(\overline{W}'_i))$, et, comme $x \in \overline{W}'_i$, P_x induit sur $U(x, r(\overline{W}'_i))$ un isomorphisme différentiable dans lequel chacun des $U_{i\lambda}, U_{j\mu}, \dots$ a pour image une partie ouverte convexe de T_x d'après ce qu'on a démontré plus haut; P_x induit donc aussi sur leur intersection un isomorphisme différentiable sur une partie ouverte convexe U' de T_x . Celle-ci admet la rétraction $(z', t) \rightarrow x + \lambda(t)(z' - x)$, où $\lambda(t)$ est une fonction différentiable sur $\mathbf{R}$, égale à 0 pour $t \leq 0$ et à 1 pour $t \geq 1$; en vertu de l'isomorphisme induit par P_x , cette rétraction se transporte à l'intersection des $U_{i\lambda}, U_{j\mu}, \dots$, ce qui achève la démonstration.

Nous n'avons fait usage en réalité que du fait que, lorsque V est plongée dans $\mathbf{R}^{(d)}$, toute partie compacte de V est de courbure bornée, ou encore que tout point de V a un voisinage qui peut se représenter paramétriquement au moyen de fonctions de classe C^1 dont les dérivées d'ordre 1 ont leurs nombres dérivés bornés. Déjà pour une variété V de classe C^1 , il ne semble pas aisé de construire un recouvrement simple sans définir d'abord sur V une structure de classe C^2 au moyen du théorème de Whitney déjà cité; et le problème de l'existence d'un recouvrement simple reste ouvert en ce qui concerne les variétés de classe C^0 ; bien entendu, pour une telle variété, on n'imposerait plus aux rétractions φ_J que d'être continues. En revanche, tout complexe simplicial localement fini admet trivialement un tel recouvrement, formé des

étoiles ouvertes de ses sommets ; en vue de ce qui va suivre, rappelons brièvement quelques définitions relatives à ces complexes. Par un complexe simplicial abstrait, on entend un ensemble N de parties finies non vides d'un ensemble quelconque I , tel que, si $J \in N$, toute partie non vide de J appartienne aussi à N ; N est dit localement fini (ou star-fini) si tout $i \in I$ n'appartient au plus qu'à un nombre fini d'éléments de N . Nous conviendrons d'identifier le complexe abstrait N avec sa «réalisation géométrique», c'est-à-dire avec l'ensemble des points $x = (x_i)_{i \in I}$ de l'espace $\mathbf{R}^{(I)}$ tels que $\sum_{i \in I} x_i = 1$, $x_i \geq 0$ pour tout i , et que l'en-

semble des $i \in I$ tels que $x_i \neq 0$ appartienne à N . Sans restreindre la généralité, on peut supposer que I est la réunion des ensembles de N (sinon on remplacerait I par cette réunion) ; pour chaque i , soit e_i le point de $\mathbf{R}^{(I)}$ dont la coordonnée d'indice i est 1 et les autres sont nulles ; les éléments i de I , ou aussi les points e_i qui leur correspondent, seront appelés les sommets de N . A tout $J \in N$, on fera correspondre, d'une part le simplexe Σ_J , ensemble des points $x = (x_i)$ de N tels que $x_i = 0$ pour i n'appartenant pas à J , d'autre part l'étoile ouverte St_J , ensemble des points $x = (x_i)$ de N tels que $x_i > 0$ pour $i \in J$; si $J = \{i\}$, Σ_J se réduit au sommet e_i de N , et St_J , qu'on écrira St_i , est dite l'étoile ouverte de e_i ; pour $J \in N$, on a $St_J = \bigcap_{i \in J} St_i$. Si J a m éléments, donc si Σ_J est de dimension $m - 1$, le centre de gravité (ou barycentre) de Σ_J sera le point $e_J = (x_i)$, avec $x_i = 1/m$ pour $i \in J$, $x_j = 0$ pour j n'appartenant pas à J . Si la fonction $\lambda(t)$ est définie comme plus haut, $(x, t) \rightarrow e_J + \lambda(t)(x - e_J)$ est une rétraction de St_J ; les St_i forment donc bien un recouvrement simple de N .

§ 2. Les formes différentielles

Par une forme différentielle, on entendra toujours une telle forme dont les coefficients, lorsqu'on exprime localement la forme au moyen de coordonnées locales, soient des fonctions de classe C^∞ de ces coordonnées. Une forme ω est dite fermée si $d\omega = 0$; elle est dite homologue à 0, sur la variété où elle est définie, s'il existe sur cette variété une forme η telle que $\omega = d\eta$.

Soit U une partie ouverte d'une variété différentiable V , munie d'une rétraction φ ; soit ω une forme de degré m sur U ; considérons sur $U \times \mathbf{R}$ la forme $\omega[\varphi(x, t)]$, image réciproque de ω par φ . Si, au voisinage d'un point de U , $x_1, \dots, x_n$ sont des coordonnées locales, on pourra écrire :

$\omega[\varphi(x, t)] = \sum_{(i)} f_{(i)}(x, t) dx_{i_1} \wedge \dots \wedge dx_{i_m} + \sum_{(j)} g_{(j)}(x, t) dt \wedge dx_{j_1} \wedge \dots \wedge dx_{j_{m-1}}$,
où $\wedge$ désigne le produit extérieur. Dans le même voisinage, considérons la forme $I\omega$ de degré $m - 1$ définie par

$$I\omega = \sum_{(j)} \left(\int_0^1 g_{(j)}(x, t) dt \right) dx_{j_1} \wedge \dots \wedge dx_{j_{m-1}}.$$

On vérifie immédiatement que cet opérateur est compatible avec les changements de coordonnées locales et peut donc être considéré comme défini globalement dans U ; si $m = 0$, on a $I\omega = 0$. Au moyen de l'expression locale de I , on vérifie aussitôt que l'on a $\omega = Id\omega + dI\omega$ si $m > 0$; si $m = 0$, on a $\omega = Id\omega + \omega(a)$ si a est la valeur constante de $\varphi(x, t)$ pour $t \leq 0$. Il s'ensuit que, si $m > 0$, $d\omega = 0$ entraîne $\omega = dI\omega$.

Supposons maintenant donné, une fois pour toutes, un recouvrement différentiablement simple $\mathfrak{U} = (U_i)_{i \in I}$ de V ; soit N le nerf de $\mathfrak{U}$. Si $H = (i_0, i_1, \dots, i_p)$ est une suite quelconque d'éléments (distincts ou non) de I , on désignera par $|H|$ l'ensemble des i_ν distincts. Par un *coélément différentiel de bidegré (m, p)* , on entendra un système $\Omega = (\omega_H) = (\omega_{i_0 i_1 \dots i_p})$ de formes de degré m , respectivement attachées aux suites $H = (i_0 i_1 \dots i_p)$ de $p + 1$ éléments de I telles que $|H| \in N$, ω_H étant pour tout H une forme définie dans $U_{|H|} = \bigcap_{0 \leq \nu \leq p} U_{i_\nu}$. Le coélément Ω sera dit *fini* s'il ne comprend qu'un nombre fini de formes $\omega_H \neq 0$; il sera dit *alterné* si $\omega_H = \omega_{i_0 \dots i_p}$ est une fonction alternée des indices $i_0, \dots, i_p$, ce qui implique que cette forme est nulle si les i_ν ne sont pas tous distincts.

Si $\Omega = (\omega_H)$ est un coélément de bidegré (m, p) , $d\Omega = (d\omega_H)$ est un coélément de bidegré $(m + 1, p)$. Comme par hypothèse on s'est donné une rétraction φ_J de U_J pour tout $J \in N$, on peut définir comme ci-dessus, pour tout $J \in N$, un opérateur I_J tel que $\omega = dI_J\omega$ pour toute forme fermée ω de degré $m > 0$, définie dans U_J . Alors, si $\Omega = (\omega_H)$ est un coélément de bidegré (m, p) , $I\Omega = (I_{|H|}\omega_H)$ est un coélément de bidegré $(m - 1, p)$; si $m > 0$, on a $\Omega = Id\Omega + dI\Omega$, et par suite $d\Omega = 0$ entraîne $\Omega = dI\Omega$. Si $m = 0$, $\Omega = (f_H)$ est un système de fonctions; comme les U_H sont rétractiles et par suite connexes, les f_H seront des constantes si $d\Omega = 0$; donc, en ce cas, Ω n'est pas autre chose qu'un système (ξ_H) de nombres réels respectivement attachés aux suites H de $p + 1$ éléments de I telles que $|H| \in N$; c'est là ce qu'on appelle, comme on sait, une *cochaîne* de N (à coefficients réels), finie ou alternée si Ω est fini ou est alterné. Il est clair que les

opérateurs d , I transforment les coéléments finis en coéléments finis, les coéléments alternés en coéléments alternés.

Soit encore $\Omega = (\omega_H) = (\omega_{i_0 \dots i_p})$ un coélément de bidegré (m, p) ; on appellera *cobord* de Ω , et on désignera par $\delta\Omega$, le coélément $\delta\Omega = (\eta_{i_0 \dots i_{p+1}})$ de bidegré $(m, p+1)$ défini par

$$\eta_{i_0 \dots i_{p+1}} = \sum_{\nu=0}^{p+1} (-1)^\nu \omega_{i_0 \dots i_{\nu-1} i_{\nu+1} \dots i_{p+1}},$$

où il doit être entendu que chacun des termes du second membre est à remplacer par la forme qu'il induit sur $U_{|i_0 \dots i_{p+1}|}$, ce qui a un sens puisque ce dernier ensemble est l'intersection des ensembles où ces termes sont définis. De même, si ω est une forme de degré m définie sur V , et si ω induit sur U_i la forme ω_i , on posera $\delta\omega = (\omega_i)$; $\delta\omega$ est donc un coélément alterné de bidegré $(m, 0)$, fini si ω est à support compact et dans ce cas seulement. Il est clair que δ est permutable avec d et transforme tout coélément fini en un coélément fini et tout coélément alterné en un coélément alterné; et on vérifie immédiatement que $\delta^2 = 0$.

Pour définir le dernier opérateur dont nous avons besoin, donnons-nous une fois pour toutes une partition différentiable de l'unité subordonnée au recouvrement $\mathfrak{U}$; on entend par là, comme on sait, une famille $(f_i)_{i \in I}$ de fonctions différentiables et ≥ 0 sur V , telles que $\sum_{i \in I} f_i = 1$ et que le support de f_i (c'est-à-dire l'adhérence de l'ensemble où $f_i > 0$) soit contenu dans U_i pour tout $i \in I$. Cela posé, soient $J \in N$, $i \in I$ et $J' = J \cup \{i\}$; si ω est une forme définie dans $U_{J'}$, on conviendra de désigner par $f_i \omega$ la forme définie dans U_J qui est égale à $f_i \omega$ dans $U_{J'}$ et à 0 dans $U_J \cap \mathbb{C}(U_{J'})$; il est immédiat en effet que c'est bien là une forme (à coefficients différentiables) dans U_J ; si J' n'appartient pas à N , c'est-à-dire si $U_{J'} = \emptyset$, cette définition entraîne que $f_i \omega = 0$. Avec cette convention, si $\Omega = (\omega_H)$ est un coélément de bidegré (m, p) avec $p > 0$, nous poserons³⁾ $K\Omega = (\zeta_{i_0 \dots i_{p-1}})$, avec

$$\zeta_{i_0 \dots i_{p-1}} = \sum_{k \in I} f_k \omega_{k i_0 \dots i_{p-1}},$$

où les termes du second membre doivent être entendus comme il vient d'être dit. De même, si $\Omega = (\omega_i)$ est un coélément de bidegré $(m, 0)$, on désignera par $K\Omega$ la forme $\omega = \sum_{k \in I} f_k \omega_k$, où on doit entendre par $f_k \omega_k$ la forme définie sur V , égale à $f_k \omega_k$ dans U_k et à 0 en dehors de

³⁾ Je dois l'opérateur K à N. Hamilton. Ma démonstration primitive se servait, au lieu de K , du théorème de prolongement de Whitney.

U_k ; ω est donc une forme définie sur V . Si Ω est de bidegré (m, p) et est fini, $K\Omega$ est fini si $p > 0$, et est une forme à support compact si $p = 0$; si Ω est alterné et $p > 0$, $K\Omega$ est alterné. On vérifie immédiatement qu'on a $\Omega = K\delta\Omega + \delta K\Omega$, donc que $\delta\Omega = 0$ entraîne $\Omega = \delta K\Omega$, pour $p \geq 0$; si ω est une forme, on a $\omega = K\delta\omega$, et $\delta\omega = 0$ entraîne donc $\omega = 0$.

Dans ces conditions, considérons toutes les suites $(\omega, \Omega_0, \Omega_1, \dots, \Omega_{m-1}, \mathcal{E})$ où ω est une forme de degré $m > 0$ sur V , Ω_h un coélément de bidegré $(m - h - 1, h)$ pour $0 \leq h \leq m - 1$, et $\mathcal{E}$ un coélément de bidegré $(0, m)$, satisfaisant aux relations

$$\delta\omega = d\Omega_0; \quad \delta\Omega_h = d\Omega_{h+1} \quad (0 \leq h \leq m - 2); \quad \delta\Omega_{m-1} = \mathcal{E}. \quad (\text{I})$$

S'il en est ainsi, on a $d\delta\Omega_h = 0$ pour $0 \leq h \leq m - 1$, $d\mathcal{E} = 0$ et $\delta\mathcal{E} = 0$, et $\delta d\omega = 0$ d'où $d\omega = K\delta d\omega = 0$. Donc ω appartient à l'espace vectoriel $\mathfrak{F}_m$ (sur $\mathbf{R}$) des formes fermées sur V ; Ω_h appartient à l'espace vectoriel $\mathfrak{F}_{m,h}$ des coéléments de bidegré $(m - h - 1, h)$ qui satisfont à $d\delta\Omega = 0$; quant à $\mathcal{E}$, puisqu'on a $d\mathcal{E} = 0$, on peut, comme on a vu, le considérer comme une cochaîne de N ; comme $\delta\mathcal{E} = 0$, c'est un cocycle; donc $\mathcal{E}$ appartient à l'espace vectoriel des cocycles de dimension m de N (à coefficients réels). Supposons Ω_h donné dans $\mathfrak{F}_{m,h}$, et $h < m - 1$; alors la relation $\delta\Omega_h = d\Omega_{h+1}$ est satisfaite pour $\Omega_{h+1} = I\delta\Omega_h$. Supposons que Ω_h soit dans la somme $\mathfrak{H}_{m,h}$ des sous-espaces de $\mathfrak{F}_{m,h}$ respectivement déterminés par les conditions $d\Omega = 0$ et $\delta\Omega = 0$; on aura donc $\Omega_h = X + Y$, $dX = 0$, $\delta Y = 0$; comme X est de bidegré $(m - h - 1, h)$, et qu'on a $m - h - 1 > 0$, $dX = 0$ entraîne $X = dIX$; on aura donc $\delta\Omega_h = \delta d(IX)$, d'où $dZ = 0$ en posant $Z = \Omega_{h+1} - \delta IX$; comme on a $\Omega_{h+1} = \delta(IX) + Z$, $dZ = 0$, Ω_{h+1} est dans $\mathfrak{H}_{m,h+1}$. Exactement de même, on voit que, si Ω_{h+1} est donné dans $\mathfrak{F}_{m,h+1}$, la relation $\delta\Omega_h = d\Omega_{h+1}$ est satisfaite par $\Omega_h = Kd\Omega_{h+1}$, puis que $\Omega_{h+1} \in \mathfrak{H}_{m,h+1}$ entraîne $\Omega_h \in \mathfrak{H}_{m,h}$. Il s'ensuit que la relation $\delta\Omega_h = d\Omega_{h+1}$ détermine un isomorphisme entre les espaces vectoriels $\mathfrak{F}_{m,h}/\mathfrak{H}_{m,h}$ et $\mathfrak{F}_{m,h+1}/\mathfrak{H}_{m,h+1}$.

De même, si Ω_0 est donné dans $\mathfrak{F}_{m,0}$, on satisfera à $\delta\omega = d\Omega_0$ en prenant $\omega = Kd\Omega_0$; si Ω_0 est dans $\mathfrak{H}_{m,0}$, on aura $\Omega_0 = X + Y$, $dX = 0$, $\delta Y = 0$, d'où $Y = \delta(KY)$, et $\delta\omega = dY = \delta d(KY)$, d'où, en posant $\eta = KY$, $\delta(\omega - d\eta) = 0$, donc $\omega = d\eta$. Réciproquement, si ω est donnée dans $\mathfrak{F}_m$, on satisfera à $\delta\omega = d\Omega_0$ en prenant $\Omega_0 = I\delta\omega$; si $\omega = d\eta$, on aura $\delta d\eta = d\Omega_0$, donc, en posant $X = \Omega_0 - \delta\eta$, $\Omega_0 = X + \delta\eta$, $dX = 0$, donc $\Omega_0 \in \mathfrak{H}_{m,0}$. En désignant par $\mathfrak{H}_m$ l'espace vectoriel des formes de degré m homologues à 0 sur V , on voit donc que

la relation $\delta\omega = d\Omega_0$ détermine un isomorphisme entre le «groupe de de Rham» $\mathfrak{F}_m/\mathfrak{S}_m$ et $\mathfrak{F}_{m,0}/\mathfrak{S}_{m,0}$. Enfin, si $\Omega_{m-1} = X + Y$, $dX = 0$, $\delta Y = 0$, on a $\mathcal{E} = \delta X$, et X est une cochaîne de N , donc $\mathcal{E}$ est un cobord de N ; réciproquement, si $\mathcal{E}$ est donné et $\delta\mathcal{E} = 0$, on satisfait à $\delta\Omega_{m-1} = \mathcal{E}$ en prenant $\Omega_{m-1} = K\mathcal{E}$; si $\mathcal{E} = \delta X$, où X est une cochaîne c'est-à-dire un coélément satisfaisant à $dX = 0$, on aura $\Omega_{m-1} = X + Y$, $dX = 0$, $\delta Y = 0$. Donc la relation $\delta\Omega_{m-1} = \mathcal{E}$ détermine un isomorphisme entre $\mathfrak{F}_{m,m-1}/\mathfrak{S}_{m,m-1}$ et le groupe de cohomologie $H^m(N)$ de dimension m de N à coefficients réels. En définitive, (I) établit donc un isomorphisme entre le groupe de de Rham $\mathfrak{F}_m/\mathfrak{S}_m$ de V , et le groupe $H^m(N)$; et cet isomorphisme est canoniquement déterminé par la seule donnée du recouvrement simple $\mathfrak{U}$.

On voit de plus que, si on se donne la forme fermée ω , on peut prendre $\Omega_h = (I\delta)^{h+1}\omega$, $\mathcal{E} = \delta(I\delta)^m\omega$; réciproquement, si on se donne le cocycle $\mathcal{E} = (\xi_{i_0\dots i_m})$, on pourra prendre $\Omega_h = K(dK)^{m-h-1}\mathcal{E}$, $\omega = K(dK)^m\mathcal{E}$, c'est-à-dire :

$$\omega = (-1)^{\frac{m(m-1)}{2}} \sum_{i_0, i_1, \dots, i_m} \xi_{i_0\dots i_m} f_{i_m} df_{i_0} \wedge \dots \wedge df_{i_{m-1}}.$$

Pour $m = 0$, on substituera aux relations (I) l'unique relation $\delta\omega = \mathcal{E}$, d'où on déduit trivialement les mêmes résultats.

Il n'y a rien à changer à ce qui précède si on désire considérer exclusivement des coéléments et cochaînes alternés. Il n'y a rien à y changer si, au lieu des formes, on désire considérer les «courants» (ce sont les formes dont les coefficients, quand on les exprime au moyen de coordonnées locales, sont des distributions au lieu d'être des fonctions différentiables). Enfin, il n'y a rien à y changer non plus si on désire considérer exclusivement les coéléments et cochaînes finis et les formes à support compact; en ce cas, bien entendu, on n'aboutit pas en général aux mêmes groupes que précédemment, mais on obtient un isomorphisme entre les groupes de de Rham à support compact et les groupes de cohomologie de N relatifs aux cochaînes finies.

Enfin, supposons qu'on se soit donné deux formes fermées ω , ω' de degrés respectifs m , r , et qu'on ait formé deux suites $(\omega, \Omega_0, \dots, \Omega_{m-1}, \mathcal{E})$ et $(\omega', \Omega'_0, \dots, \Omega'_{r-1}, \mathcal{E}')$ satisfaisant à (I). On peut alors former, sans nouvelle intégration, une suite $(\omega'', \Omega''_0, \dots, \Omega''_{m+r-1}, \mathcal{E}'')$ satisfaisant à (I) et commençant par le produit extérieur $\omega'' = \omega \wedge \omega'$. Posons en effet $\Omega_h = (\omega_{i_0\dots i_h}^h)$, $\mathcal{E} = (\xi_{i_0\dots i_m})$, et de même pour Ω'_k , $\mathcal{E}'$; on pourra alors prendre :

$$\Omega_h'' = (\omega_{i_0 \dots i_h}^h \wedge \omega') \quad (0 \leq h \leq m-1),$$

$$\Omega_{m+k}'' = (\xi_{i_0 \dots i_m} \omega_{i_m \dots i_{m+k}}^k) \quad (0 \leq k \leq r-1),$$

$$\mathcal{E}'' = (\xi_{i_0 \dots i_m} \xi'_{i_m \dots i_{m+r}}),$$

c'est-à-dire $\mathcal{E}'' = \mathcal{E} \cup \mathcal{E}'$. Si on s'était servi exclusivement de coéléments alternés, les formules ci-dessus seraient à modifier; la manière la plus simple de le faire est d'ordonner une fois pour toutes les $i \in I$ et de convenir que les Ω_{m+k}'' et $\mathcal{E}''$ sont alternés et ont leurs composantes données par les formules ci-dessus pour $i_0 < i_1 < \dots < i_{m+r}$; cela donne le «cup-product» de Whitney.

§ 3. Les cycles singuliers

Les hypothèses et notations restant les mêmes qu'au § 2, nous passons maintenant à l'étude des cycles singuliers différentiables.

Dans un espace affine, considérons $m+1$ points $a_0, \dots, a_m$; soit K le plus petit ensemble convexe contenant les a_μ , et soit L la variété linéaire qui porte K ; L est l'ensemble des points $\sum_{\mu=0}^m x_\mu a_\mu$ pour $\sum_{\mu} x_\mu = 1$, et K est l'ensemble des points de cette forme pour lesquels $\sum_{\mu} x_\mu = 1$ et $x_\mu \geq 0$ pour tout μ . Si L est de dimension m , K est un «simplexe euclidien» de dimension m , de sommets $a_0, \dots, a_m$. En particulier, si e_μ est le vecteur dans $\mathbf{R}^{m+1}$ dont la μ -ième composante est 1 et les autres sont nulles, on notera Σ^m le simplexe de sommets $e_0, \dots, e_m$, c'est-à-dire l'ensemble des $x = (x_\mu)$ de $\mathbf{R}^{m+1}$ tels que $\sum_{\mu} x_\mu = 1$ et $x_\mu \geq 0$ pour tout μ .

Par un simplexe singulier différentiable de dimension m dans V , on entendra, suivant S. Eilenberg ⁴⁾, la restriction à Σ^m d'une application différentiable f dans V d'un voisinage de Σ^m ; $f(\Sigma^m)$ sera dit le support de ce simplexe. Si de plus K et L sont définis comme ci-dessus à partir de points $a_0, \dots, a_m$, et que f soit une application différentiable dans V d'un voisinage de K (dans l'espace ambiant ou seulement dans L), l'application $(x_0, \dots, x_m) \rightarrow f(\sum_{\mu} x_\mu a_\mu)$, restreinte à Σ^m , est un simplexe singulier différentiable qui sera noté $[f; a_0 \dots a_m]$; il est dégénéré si L est de dimension $< m$.

Le mot «différentiable» sera en général sous-entendu dans ce qui suit. Par une *chaîne* (ou plus explicitement une chaîne singulière différen-

⁴⁾ S. Eilenberg, Singular homology in differentiable manifolds, Ann. Math. 48 (1947) p. 670.

tiable) de dimension m dans V , à coefficients dans un groupe abélien G , on entendra toute expression de la forme $t = \sum_e c_e s_e$, où les c_e sont dans G et les s_e sont des simplexes singuliers de dimension m dans V dont les supports forment une famille localement finie; une telle expression sera dite réduite si tous les s_e sont distincts et tous les c_e sont $\neq 0$. Toute chaîne possède une expression réduite et une seule; le support $|t|$ d'une chaîne t sera la réunion des supports des simplexes figurant dans l'expression réduite de t ; on dira que t est contenue dans une partie U de V si $|t| \subset U$. Une chaîne est dite finie si son expression réduite est une somme finie, ou, ce qui revient au même, si son support est compact. Si $t = \sum_e c_e s_e$ est une chaîne finie, on posera $\deg(t) = \sum_e c_e$.

Si $s = [f; a_0 \dots a_m]$ et qu'on pose $s_\mu = [f; a_0 \dots a_{\mu-1} a_{\mu+1} \dots a_m]$, la chaîne finie $bs = \sum_{\mu=0}^m (-1)^\mu s_\mu$ s'appelle le bord de s ; cet opérateur s'étend aux chaînes par linéarité; une chaîne de bord nul s'appelle un cycle; on a $b^2 = 0$, ce qui permet de définir des groupes d'homologie de V au moyen de b et du groupe des chaînes (ou encore du groupe des chaînes finies) à coefficients dans G . Si t est de dimension 0, on a $bt = 0$, mais on posera $b_0 t = \deg(t)$ si t est fini; on a $b_0 bt = 0$ si t est de dimension 1. Plus généralement, on a $\deg(bt) = \deg(t)$ si t est de dimension m paire > 0 , et $\deg(bt) = 0$ en tout autre cas.

Soit s un simplexe singulier défini par une application différentiable f dans V d'un voisinage W de Σ^m ; si ω est une forme de degré m dans V , son image réciproque $\omega[f(x)]$ par f est une forme de degré m dans W , dont l'intégrale sur Σ^m est par définition l'intégrale $\int_s \omega$ de ω sur s ; cette définition s'étend par linéarité aux chaînes finies à coefficients réels, et même à toutes les chaînes à coefficients réels si ω est à support compact. On a la formule de Stokes $\int_t d\omega = \int_{bt} \omega$, valable chaque fois que t est une chaîne finie ou que ω est à support compact. Au moyen de $\int_t \omega$, qui est une forme bilinéaire en t et ω , les chaînes finies sont mises en dualité avec les formes, et les chaînes avec les formes à support compact, ce qui permet de transposer aux chaînes, par dualité, les opérations et les résultats du § 2; mais nous allons en donner un exposé indépendant, de manière à ne pas avoir à supposer $G = \mathbf{R}$.

Soit d'abord U une partie ouverte de V munie d'une rétraction différentiable φ ; soit p la valeur constante de $\varphi(x, t)$ pour $t \leq 0$. On désignera par $\bar{s}_m$ le simplexe dégénéré $[f; aa \dots a]$ de dimension m , où

$f(a) = p$, ou, ce qui revient au même, le simplexe défini par la restriction à Σ^m de l'application constante de R^{m+1} sur p ; on a $b\bar{s}_m = \bar{s}_{m-1}$ si m est pair et > 0 , $b\bar{s}_m = 0$ si m est impair ou 0 . Considérons un simplexe singulier $s = [f; a_0 \dots a_m]$ dans U , les a_μ étant des points d'un espace affine E ; désignons par a_μ^0, a_μ^1 les points $(a_\mu, 0)$ et $(a_\mu, 1)$ de $E \times \mathbf{R}$. Par définition, f est une application différentiable dans U d'un voisinage du plus petit ensemble convexe K contenant les a_μ ; alors, si on pose $f'(x, t) = \varphi[f(x), t]$, f' est une application différentiable dans U d'un voisinage de $K \times \mathbf{R}$. Posons, dans ces conditions :

$$Ps = \sum_{\mu=0}^m (-1)^\mu [f'; a_0^0 \dots a_\mu^0 a_\mu^1 \dots a_m^1] + \bar{s}_{m+1},$$

et étendons cet opérateur par linéarité aux chaînes finies dans U . Un calcul facile donne $bPs + Pbs = s$ pour $m > 0$ et $bPs + Pbs = s - \bar{s}_0$ pour $m = 0$, donc, pour toute chaîne finie de dimension m , $t = bPt + Pbt$ si $m > 0$ et $t = bPt + (b_0 t) \bar{s}_0$ si $m = 0$. Donc $bt = 0$ entraîne $t = bPt$ si $m > 0$, et $b_0 t = 0$ entraîne $t = bPt$ si $m = 0$.

Par un $\mathfrak{U}$ -simplexe, on entendra un simplexe singulier contenu dans l'un au moins des ensembles U_i du recouvrement $\mathfrak{U}$; par une $\mathfrak{U}$ -chaîne, on entendra une chaîne dont tous les simplexes sont des $\mathfrak{U}$ -simplexes. L'application de notre méthode exige qu'on se restreigne aux $\mathfrak{U}$ -chaînes; d'après un théorème de S. Eilenberg ⁵⁾, cela ne change rien aux groupes d'homologie; rappelons les points principaux de sa démonstration. Soit $s = [f; a_0 \dots a_m]$ un simplexe singulier. Posons $I_\mu = \{0, 1, \dots, \mu\}$ pour $0 \leq \mu \leq m$; et, si $I = \{\mu_1, \dots, \mu_k\}$ est une partie quelconque de I_m , posons $a_I = \sum_{h=1}^k (1/k) a_{\mu_h}$. Alors on appelle subdivision barycentrique de s la chaîne finie

$$\sigma s = \sum_{\pi} \varepsilon_{\pi} [f; a_{\pi(I_0)} \dots a_{\pi(I_m)}],$$

où la somme est étendue à toutes les permutations π de I_m , et où $\varepsilon_{\pi} = \pm 1$ suivant que π est paire ou impaire; on étend l'opérateur aux chaînes par linéarité; on vérifie qu'on a $b\sigma = \sigma b$. D'autre part, Eilenberg (loc. cit., note 5, p. 429) définit un autre opérateur ϱ , analogue mais dont l'expression explicite serait plus compliquée, tel que $b\varrho + \varrho b = \sigma - 1$; ϱs est une chaîne finie de dimension $m + 1$, somme de termes de la forme $\pm [f; b_0 \dots b_{m+1}]$, où chacun des b_μ est l'un des a_I . Cela posé, si s est un simplexe singulier, on peut trouver un entier ν assez grand pour que $\sigma^\nu s$ soit une $\mathfrak{U}$ -chaîne; soit $\nu(s)$ le plus petit entier ayant cette propriété; soit τ l'opérateur défini sur les simplexes singuliers par

⁵⁾ S. Eilenberg, Singular homology theory, Ann. Math. 45 (1944) p. 407.

$$\tau s = \varrho(1 + \sigma + \dots + \sigma^{\nu(s)-1}) s ,$$

et étendu aux chaînes par linéarité. Alors, si on pose comme plus haut $bs = \sum_{\mu} (-1)^{\mu} s_{\mu}$, on vérifie immédiatement qu'on a

$$(b\tau + \tau b)s = (\sigma^{\nu(s)} - 1)s - \sum_{\mu} (-1)^{\mu} \sum_{j=\nu(s_{\mu})}^{\nu(s)-1} \varrho \sigma^j s_{\mu} ,$$

ce qui montre que $(1 + b\tau + \tau b)s$ est une $\mathcal{U}$ -chaîne finie, de support contenu dans celui de s . Donc, si t est une chaîne, $\bar{t} = (1 + b\tau + \tau b)t$ est une $\mathcal{U}$ -chaîne, finie si t est finie. Si t est un cycle, on a $\bar{t} = t + b\tau t$, donc $\bar{t}$ est un cycle et t est homologue à $\bar{t}$. De plus, la formule ci-dessus montre que $(b\tau + \tau b)s = 0$ si $\nu(s) = 0$, c'est-à-dire si s est un $\mathcal{U}$ -simplexe, donc $\bar{t} = t$ si t est une $\mathcal{U}$ -chaîne. Supposons qu'une $\mathcal{U}$ -chaîne t' soit le bord d'une chaîne t ; on aura $t' = bt$, et $b\bar{t} = bt + b\tau bt = \bar{t}' = t'$, donc t' est aussi le bord d'une $\mathcal{U}$ -chaîne. Il s'ensuit bien que la restriction aux $\mathcal{U}$ -chaînes ne change rien aux groupes d'homologie; désormais nous ne considérerons que celles-là, et pour abrégé nous dirons «chaîne» au lieu de « $\mathcal{U}$ -chaîne».

Par un *élément singulier de bidegré* (m, p) , on entendra un système $T = (t_H) = (t_{i_0 \dots i_p})$ de chaînes finies t_H de dimension m respectivement attachées aux suites $H = (i_0 \dots i_p)$ de $p + 1$ éléments de I telles que $|H| \in N$, t_H étant contenue dans $U_{|H|}$ pour tout H . L'élément T sera dit *fini* si les t_H sont tous nuls à l'exception d'un nombre fini d'entre eux, *alterné* si $t_{i_0 \dots i_p}$ est une fonction alternée de ses indices.

Si $T = (t_H)$ est un élément de bidegré (m, p) , $bT = (bT_H)$ est un élément de bidegré $(m - 1, p)$, fini si T est fini, alterné si T est alterné; on a $b^2 = 0$. Si de plus $m = 0$, $b_0 T = (b_0 t_H)$ fait correspondre à tout H un élément $b_0 t_H$ du groupe de coefficients G ; c'est là ce qu'on appelle une chaîne de N à coefficients dans G . Si $m = 1$, on a $b_0 bT = 0$.

Le recouvrement $\mathcal{U}$ étant simple, on peut, au moyen des rétractions φ_J attachées à tout $J \in N$, définir dans les U_J des opérateurs P_J ayant les propriétés décrites plus haut, et tels en particulier que, si t est une chaîne finie de dimension $m > 0$ dans U_J , $bt = 0$ entraîne $t = bPt$. Alors, si $T = (t_H)$ est un élément de bidegré (m, p) , on posera $PT = (P_{|H|} t_H)$; c'est un élément de bidegré $(m + 1, p)$; si $m > 0$, $bT = 0$ entraîne $T = bPT$; si $m = 0$, $b_0 T = 0$ entraîne $T = bPT$; en général, on a $T = bPT + PbT$ si $m > 0$; PT est fini si T est fini, alterné si T est alterné.

D'autre part, si $T = (t_{i_0 \dots i_p})$ est un élément de bidegré (m, p) , et si $p > 0$, nous définirons un élément $\partial T = (u_{i_0 \dots i_{p-1}})$ de bidegré $(m, p - 1)$ au moyen de la formule

$$u_{i_0 \dots i_{p-1}} = \sum_{\mu, k} (-1)^\mu t_{i_0 \dots i_{\mu-1} k i_\mu \dots i_{p-1}},$$

où la sommation doit être étendue aux valeurs de μ, k pour lesquelles $|i_0 \dots i_{\mu-1} k i_\mu \dots i_{p-1}| \in N$; ces valeurs sont en nombre fini, et tous les termes du second membre sont des chaînes finies dans $U_{|i_0 \dots i_{p-1}|}$, donc ces formules définissent bien un élément ∂T , qui est fini si T est fini. De même, si $T = (t_i)$ est un élément de bidegré $(m, 0)$, on posera $\partial T = \sum_k t_k$; ∂T est alors une chaîne, finie si T est fini. On a $\partial^2 = 0$, et ∂ est permutable avec b . D'autre part, on peut aussi, dans la formule qui définit ∂T , interpréter T comme une chaîne de N , les t_H étant alors des éléments de G ; cette formule, où la sommation est étendue aux mêmes valeurs de μ, k que tout à l'heure, définit alors ∂T comme chaîne de N ; les groupes d'homologie de N sont ceux qui sont définis au moyen des chaînes de N et de l'opérateur ∂ , ou encore au moyen des chaînes finies de N et de ∂ . Dans ces conditions, ∂ est permutable avec b_0 .

On va définir un opérateur L tel que $\partial T = 0$ entraîne $T = \partial LT$. Pour cela, convenons de choisir une fois pour toutes, pour tout $\mathcal{U}$ -simplexe s , l'un des U_i dans lesquels il est contenu; soit $U_{f(s)}$ cet ensemble. Soit $T = (t_H)$ un élément de bidegré (m, p) ; soit $t_H = \sum_q c_H^q s_q$ l'expression réduite de t_H . Si $H = (i_0 \dots i_p)$, on posera $iH = (i_{i_0} \dots i_{i_p})$. Alors on définira un élément $LT = (v_H)$ de bidegré $(m, p+1)$ en posant $v_{iH} = \sum_{f(s_q)=i} c_H^q s_q$ chaque fois que $|iH| \in N$; cela veut dire que la somme est étendue à toutes les valeurs de q telles que $f(s_q) = i$. Puisque t_H est une somme finie, v_{iH} en est une aussi; et chaque simplexe s_q figurant dans v_{iH} est contenu dans $U_{|H|}$ parce qu'il figure dans t_H , et dans U_i parce que $i = f(s_q)$, donc aussi dans $U_{|iH|}$; LT est donc bien un élément, fini si T est fini. De même, si $t = \sum_q c_q s_q$ est l'expression réduite d'une $\mathcal{U}$ -chaîne de dimension m , on définit, au moyen de $v_i = \sum_{f(s_q)=i} c_q s_q$, un élément $Lt = (v_i)$ de bidegré $(m, 0)$, fini si t est finie. On a $T = \partial LT + L\partial T$ si T est un élément, et $t = \partial Lt$ si t est une chaîne. Si donc T est un élément tel que $\partial T = 0$, on a $T = \partial LT$.

Il n'est pas vrai que LT soit alterné chaque fois que T est alterné. Si on veut se servir exclusivement d'éléments alternés, il faut substituer à ∂, L les opérateurs ∂', L' qui, avec les mêmes notations que ci-dessus, sont définis par les formules

$$\partial' T = \left(\sum_k t_{k i_0 \dots i_{p-1}} \right),$$

où la sommation est étendue aux k tels que $|ki_0 \dots i_{p-1}| \in N$, et

$$L'T = \left(\sum_{\mu=0}^{p+1} \sum_{f(s_\varrho)=i_\mu} (-1)^\mu c_{i_0 \dots i_{\mu-1} i_{\mu+1} \dots i_{p+1}}^{\varrho} s_\varrho \right).$$

On vérifie facilement qu'ils possèdent des propriétés semblables à celles de ∂ et L lorsqu'on les applique à des éléments alternés et qu'ils transforment ceux-ci en éléments alternés.

Considérons maintenant toutes les suites $(t, T_0, \dots, T_m, Z)$, où t est une chaîne de dimension $m > 0$ de V , T_h un élément de bidegré $(m-h, h)$ pour $0 \leq h \leq m$, et Z une chaîne de dimension m de N , satisfaisant aux relations

$$t = \partial T_0; \quad bT_h = \partial T_{h+1} \quad (0 \leq h \leq m-1); \quad b_0 T_m = Z. \quad (\text{II})$$

S'il en est ainsi, on a $b\partial T_h = 0$ ($0 \leq h \leq m-1$), $b_0 \partial T_m = 0$, $bt = 0$ et $\partial Z = 0$. Donc t appartient au groupe $\mathfrak{C}_m$ des cycles singuliers différentiables à coefficients dans G sur V , et Z au groupe des cycles de N à coefficients dans G ; T_h appartient au groupe $\mathfrak{C}_{m,h}$ des éléments de bidegré $(m-h, h)$ qui satisfont à $b\partial T = 0$ pour $h < m$ et à $b_0 \partial T = 0$ pour $h = m$. Soit $\mathfrak{B}_m$ le groupe des bords dans V , c'est-à-dire le groupe des éléments de $\mathfrak{C}_m$ de la forme bt' ; soit $\mathfrak{B}_{m,h}$, pour $0 \leq h \leq m$, le groupe des éléments de $\mathfrak{C}_{m,h}$ de la forme $bX + \partial Y$, où X, Y sont des éléments de bidegrés respectifs $(m-h+1, h)$ et $(m-h, h+1)$. On satisfera à la relation $bT_h = \partial T_{h+1}$ en prenant $T_h = P\partial T_{h+1}$ si T_{h+1} est donné dans $\mathfrak{C}_{m,h+1}$, et $T_{h+1} = LbT_h$ si T_h est donné dans $\mathfrak{C}_{m,h}$; on satisfera à $t = \partial T_0$ en prenant $T_0 = Lt$ si t est donné dans $\mathfrak{C}_m$; enfin il est clair qu'on peut former T_m satisfaisant à $b_0 T_m = Z$ si Z est donné. Si $T_h \in \mathfrak{B}_{m,h}$, donc si $T_h = bX + \partial Y$, on aura, en posant $U = T_{h+1} - bY$, $\partial U = 0$, d'où $U = \partial(LU)$ et $T_{h+1} = bY + \partial(LU) \in \mathfrak{B}_{m,h+1}$; de même, si $T_{h+1} = bY + \partial V$, on aura $bW = 0$ en posant $W = T_h - \partial Y$, d'où $W = bPW$ puisque W est de bidegré $(m-h, h)$ et que $m-h > 0$; on a donc $T_h = b(PW) + \partial Y \in \mathfrak{B}_{m,h}$. Donc la relation $bT_h = \partial T_{h+1}$ détermine un isomorphisme entre $\mathfrak{C}_{m,h}/\mathfrak{B}_{m,h}$ et $\mathfrak{C}_{m,h+1}/\mathfrak{B}_{m,h+1}$. De même, si $t = \partial T_0$ et $T_0 = bX + \partial Y$, on a $t = b(\partial X) \in \mathfrak{B}_m$; si $t = bt'$, et qu'on pose $U = T_0 - b(Lt')$, on a $\partial U = 0$, donc $U = \partial(LU)$, et $T_0 = b(Lt') + \partial(LU) \in \mathfrak{B}_{m,0}$. Si $b_0 T_m = Z$ et $T_m = bX + \partial Y$, on a $Z = \partial(b_0 Y)$, donc Z est homologue à 0; et, si $Z = \partial Z'$ et $b_0 T' = Z'$, on aura, en posant $X = T_m - \partial T'$, $b_0 X = 0$, donc $X = bPX$ et $T_m = b(PX) + \partial T' \in \mathfrak{B}_{m,m}$. En définitive, on voit que les relations (II) établissent un isomorphisme entre le groupe d'homologie singulière diffé-

rentiable $\mathfrak{C}_m/\mathfrak{B}_m$ de V et le groupe d'homologie des chaînes de N , pour la dimension m et le groupe de coefficients G ; et cet isomorphisme est canoniquement déterminé par la donnée du recouvrement simple $\mathfrak{U}$.

Pour $m = 0$, on partira des relations $t = \partial T_0$, $b_0 T_0 = Z$, où T_0 est un élément de bidegré $(0, 0)$, et on arrive au même résultat par des raisonnements analogues mais plus simples.

Il n'y a rien à changer à ce qui précède si on veut considérer exclusivement les éléments et chaînes finis; on obtient ainsi un isomorphisme entre les groupes d'homologie de V et de N obtenus au moyen de chaînes finies. Il n'y a rien à y changer si on veut se servir de chaînes de classe C^k , c'est-à-dire dont les simplexes sont définis par des applications k fois continument différentiables, k étant un entier quelconque; en ce cas, il suffit que les rétractions φ_J soient elles-mêmes de classe C^k ; pour $k = 0$, on voit qu'on obtient les mêmes résultats au moyen de chaînes singulières continues, les φ_J étant alors seulement assujetties à être continues; ce résultat s'applique en particulier au recouvrement simple d'un complexe simplicial localement fini par les étoiles ouvertes des sommets (voir § 1), et contient donc une démonstration de l'invariance topologique des groupes d'homologie combinatoires d'un tel complexe, qui d'ailleurs ne diffère qu'en apparence de la démonstration classique. Il n'y a rien à changer non plus à ce qui précède si l'on veut se servir exclusivement d'éléments alternés, et de chaînes alternées de N , sauf qu'il faut substituer ∂' , L' à ∂ , L .

Si on prend $G = \mathbf{R}$, les opérateurs qu'on a défini sur les éléments singuliers sont en dualité avec ceux qu'on a défini sur les coéléments différentiels. Soient en effet $\Omega = (\omega_H)$ et $T = (t_H)$ un coélément différentiel et un élément singulier, tous deux de bidegré (m, p) , dont l'un soit fini; on posera alors

$$(T, \Omega) = \sum_H \int_H \omega_H,$$

et, s'ils sont tous deux alternés :

$$(T, \Omega)' = \frac{1}{(p+1)!} (T, \Omega) = \sum'_H \int_H \omega_H$$

où Σ' indique qu'on prend une fois seulement chaque combinaison $i_0, \dots, i_p$ de $p+1$ éléments de I , rangés dans un ordre quelconque; c'est de $(T, \Omega)'$ qu'il faut se servir dans la théorie alternée. La formule de Stokes donne $(bT, \Omega) = (T, d\Omega)$; et on vérifie facilement qu'on a $(\partial T, \Omega) = (T, \delta\Omega)$, et de même $(\partial' T, \Omega)' = (T, \delta\Omega)'$ si T, Ω sont alternés; de même, si ω est une forme de degré m sur V et T un élément de

bidegré $(m, 0)$, et que ω soit à support compact ou T fini, on a $(T, \delta\omega) = \int_{\partial T} \omega$. Enfin, si T est un élément de bidegré $(0, p)$, et $\mathcal{E}$ un coélément

de bidegré $(0, p)$ satisfaisant à $d\mathcal{E} = 0$ ou autrement dit une cochaîne de N , et si T ou $\mathcal{E}$ est fini, on a $(T, \mathcal{E}) = (b_0 T, \mathcal{E})$, où dans le second membre figure le produit scalaire des chaînes et cochaînes de N défini par $(Z, \mathcal{E}) = \sum_H z_H \xi_H$ pour $Z = (z_H)$, $\mathcal{E} = (\xi_H)$; ces formules sont à modifier d'une manière évidente dans la théorie alternée.

Considérons alors deux suites $(\omega, \Omega_0, \dots, \Omega_{m-1}, \mathcal{E})$ et $(t, T_0, \dots, T_m, Z)$, satisfaisant respectivement aux relations (I) du § 2 et aux relations (II) ci-dessus; supposons ω à support compact et les Ω_h et $\mathcal{E}$ finis, ou t , les T_h et Z finis. Au moyen des formules ci-dessus, on obtient immédiatement :

$$\int_t \omega = (T_0, d\Omega_0) = \dots = (T_{m-1}, d\Omega_{m-1}) = (T_m, \delta\Omega_{m-1}) = (Z, \mathcal{E}) \quad (\text{III})$$

Il s'ensuit que les groupes de de Rham et les groupes d'homologie singulière à coefficients réels de V ont entre eux les mêmes relations de dualité que les groupes de cohomologie et d'homologie de N . En particulier, il existe toujours une forme fermée ω sur V telle que $\int_t \omega$ soit une fonction

linéaire arbitrairement donnée sur le groupe d'homologie singulière finie de V , ou autrement dit soit égale à une fonction linéaire $L(t)$ donnée sur l'espace vectoriel des cycles finis de V , nulle sur les bords de chaînes finies. D'autre part, si une forme fermée ω à support compact sur V est telle que $\int_t \omega = 0$ pour tout cycle t , fini ou non, de V , elle est de la

forme $\omega = d\eta$, où η est à support compact; de même, si une forme fermée ω est telle que $\int_t \omega = 0$ pour tout cycle fini t , elle est de la

forme $\omega = d\eta$. En effet, d'après ce qui précède, il suffit, pour obtenir ces résultats, de vérifier les résultats analogues pour N , ce qui est immédiat.

Les espaces vectoriels dont il s'agit ici sont en général de dimension infinie si V n'est pas compacte; on ne peut donc espérer établir entre eux de relations de dualité tout à fait satisfaisantes à moins d'y introduire des topologies convenables; c'est là un terrain sur lequel nous ne nous engagerons pas. En revanche, si V est compacte, le recouvrement $\mathcal{U}$ est fini; ce qui précède montre donc que tous les groupes d'homologie de V sont alors de type fini, et s'annulent au-dessus d'une certaine dimension; sur $\mathbf{R}$, en particulier, tous ces groupes sont des espaces vectoriels de dimension finie. On conclut alors de ce qui précède que la fonc-

tion bilinéaire $\int \omega$ met en dualité le groupe de de Rham de degré m et le groupe d'homologie différentiable de dimension m à coefficients réels.

On peut compléter ces résultats au moyen des remarques suivantes, que nous bornerons au cas compact, où toute chaîne est finie. Il est immédiat que toute chaîne t à coefficients réels peut se mettre sous la forme $t = \sum_i \xi_i t_i$, où les t_i sont des chaînes à coefficients entiers et les ξ_i sont des nombres réels linéairement indépendants sur le corps $\mathbf{Q}$ des rationnels ; alors $bt = 0$ entraîne $bt_i = 0$ pour tout i , donc tout cycle réel est combinaison linéaire de cycles entiers ; et, si un cycle entier t' est le bord bt d'une chaîne réelle t , alors, en mettant t sous la forme ci-dessus, on voit que l'un des ξ_i , par exemple ξ_1 , doit être rationnel et qu'alors on a $t' = b(\xi_1 t_1)$, donc qu'un multiple entier de t' est le bord d'un cycle entier. Le groupe d'homologie entière de dimension m étant de type fini, il est somme directe d'un groupe fini et d'un groupe abélien libre engendré par des classes d'homologie entière en nombre fini ; soient $t_1, \dots, t_r$ des cycles entiers appartenant respectivement à ces classes ; d'après ce qui précède, les classes d'homologie réelle de $t_1, \dots, t_r$ forment alors une base du groupe d'homologie réelle de dimension m considéré comme espace vectoriel sur $\mathbf{R}$; et on peut identifier les formes linéaires sur ce dernier groupe avec les homomorphismes dans $\mathbf{R}$ du groupe d'homologie entière, une telle forme ou un tel homomorphisme étant complètement déterminé par ses valeurs sur les classes des cycles t_i .

Par une *période* d'une forme ω , on entend son intégrale $\int \omega$ sur un cycle entier t ; pour un choix déterminé des cycles $t_1, \dots, t_r$, on appelle souvent «*périodes fondamentales*» de ω les intégrales de ω sur les t_i . On voit donc qu'il revient au même de se donner, soit la forme linéaire $\int \omega$ sur le groupe d'homologie réelle de V , soit l'homomorphisme $\int \omega$ du groupe d'homologie entière de V dans $\mathbf{R}$, soit les périodes fondamentales de ω . On a donc retrouvé les «*théorèmes de de Rham*» sous leur forme classique :

Sur une variété différentiable compacte V , il existe des formes fermées dont les périodes fondamentales sont arbitrairement données ; toute forme fermée dont les périodes fondamentales sont nulles est homologue à 0 sur V .

Quant au «*troisième théorème de de Rham*», une partie en est contenue dans le résultat de la fin du § 2, d'après lequel le «*cup-product*» des cocycles de N correspond au produit extérieur des formes sur V . Pour passer de là à l'énoncé classique du même théorème, il faut se servir de la dualité de Poincaré établie par le nombre d'intersection

entre les cycles réels de dimensions m et $n - m$, ou encore (ce qui au fond revient au même) passer au produit de la variété par elle-même, puis à la diagonale dans ce produit. Je n'insisterai pas sur ces questions déjà classiques ; mais il ne sera pas superflu de faire apparaître une conséquence importante de nos résultats, qui d'habitude se déduit du troisième théorème de de Rham. Bornons-nous toujours au cas compact ; considérons une forme ω sur V dont toutes les périodes sont entières ; soit $\mathcal{E}$ un cocycle de N correspondant à ω , cocycle qui est bien déterminé à un cobord arbitraire près. Alors $(Z, \mathcal{E})$ est entier pour tout cycle entier Z . Mais le groupe des cycles entiers de N est le sous-groupe du groupe des chaînes entières déterminé par les conditions $\partial Z = 0$, donc toute chaîne entière dont un multiple est un cycle est elle-même un cycle ; d'après la théorie des diviseurs élémentaires, le groupe des chaînes entières est donc somme directe du groupe des cycles entiers et d'un autre groupe, de sorte qu'on peut étendre au groupe des chaînes tout homomorphisme donné sur le groupe des cycles. Comme tout homomorphisme du groupe des chaînes entières dans le groupe additif des entiers peut s'écrire sous la forme $Z \rightarrow (Z, \mathcal{E}_0)$, où $\mathcal{E}_0$ est une cochaîne entière, on voit qu'il existe une cochaîne entière $\mathcal{E}_0$ telle que $(Z, \mathcal{E}_0) = (Z, \mathcal{E})$ pour tout cycle entier Z , donc aussi pour tout cycle réel Z . Il s'ensuit que $\mathcal{E}_0 - \mathcal{E}$ est le cobord d'une cochaîne réelle, donc que $\mathcal{E}_0$ est, aussi bien que $\mathcal{E}$, un cocycle correspondant à ω . Par suite, *pour qu'une forme ω corresponde à un cocycle $\mathcal{E}$ à coefficients entiers, il faut et il suffit que toutes ses périodes soient des entiers*. De là et du résultat final du § 2, on conclut que, *si ω et ω' sont à périodes entières, il en est de même de leur produit extérieur $\omega \wedge \omega'$* . Bien entendu, on peut obtenir aussi ce même résultat en passant au produit de V par elle-même et en se servant du théorème de Künneth.

§ 4. La dualité de Poincaré

Tout ce que nous avons fait jusqu'ici repose en réalité sur une seule propriété du recouvrement $\mathcal{U}$: c'est que les U_j sont homologiquement triviaux, c'est-à-dire ont l'homologie d'un espace réduit à un point. Nous nous sommes servis, il est vrai, des rétractions φ_j , mais seulement pour obtenir un exposé à la fois plus élémentaire et plus élégant grâce à la possibilité de définir explicitement les opérateurs I et P . L'exposé ci-dessus renferme donc, du moins pour l'homologie singulière, une démonstration du théorème de Leray d'après lequel, si un recouvrement $\mathcal{U}$ d'un espace X est tel que les U_j soient homologiquement triviaux, l'homologie de X est la même que celle du nerf N de $\mathcal{U}$.

En revanche, puisque tout complexe simplicial admet un recouvrement simple, il est évident que l'existence d'un tel recouvrement n'entraîne pas le théorème de dualité de Poincaré. Pour obtenir ce théorème sur une variété au moyen du recouvrement $\mathcal{U}$, il faut mettre en œuvre une propriété des U_j qui n'est pas encore intervenue, à savoir que leur homologie modulo leur frontière est triviale dans toutes les dimensions sauf la dimension n de V . Ce n'est pas là une propriété «élémentaire» sauf en ce qui concerne les formes différentielles ; aussi nous bornerons-nous à celles-ci, et par conséquent à la dualité de Poincaré à coefficients réels. Pour les formes, la propriété en question des U_j n'est autre que le résultat suivant, qui est bien connu et facile à démontrer élémentairement :

Soit ω une forme différentielle à support compact contenu dans une partie ouverte convexe U de $\mathbf{R}^n$. Alors, pour que ω soit la différentielle $d\eta$ d'une forme η à support compact contenu dans U , il faut et il suffit qu'on ait $d\omega = 0$ si ω est de degré $< n$, et qu'on ait $\int_U \omega = 0$ si ω est de degré n .

Comme les ensembles U_j formés au moyen de notre recouvrement simple $\mathcal{U}$ de V sont différentiablement isomorphes à des parties ouvertes convexes de $\mathbf{R}^n$, le résultat ci-dessus leur est applicable.

On supposera V orientable ; dans le cas contraire, il faudrait se servir de formes «de deuxième espèce» au sens de de Rham, c'est-à-dire à «coefficients locaux» qui sont les «réels tordus» ; cela ne fait aucune difficulté mais entraîne quelques complications de langage qu'il vaut mieux éviter ici puisqu'il ne s'agit que de résultats bien connus par ailleurs. On supposera donc tous les U_j orientés d'une manière cohérente au moyen d'une orientation de V choisie une fois pour toutes ; c'est sur les U_j ainsi orientés qu'on intégrera les formes différentielles de degré n à supports contenus dans ces ensembles.

Par un *élément différentiel de bidegré* (m, p) on entendra un système $\Theta = (\theta_H)$ de formes de degré m respectivement attachées aux suites H de $p + 1$ éléments de I telles que $|H| \in N$, θ_H étant pour tout H une forme à support compact contenu dans $U_{|H|}$; l'élément Θ sera dit *fini* si les θ_H sont nuls sauf un nombre fini d'entre eux. On pose $d\Theta = (d\theta_H)$; c'est là un élément de bidegré $(m + 1, p)$. Si $\Theta = (\theta_H)$ est un élément de bidegré (n, p) , on désignera par $\int \Theta$ la chaîne $Z = (z_H)$ de N définie par $z_H = \int_V \theta_H = \int_{U_{|H|}} \theta_H$. On a $d^2 = 0$, et $\int d\Theta = 0$ si Θ est de bidegré $(n - 1, p)$. Pour que l'élément Θ de bidegré (m, p) soit de la forme $d\Theta'$, où Θ' est de bidegré $(m - 1, p)$, il faut et il suffit que $d\Theta = 0$ si $m < n$, et que $\int \Theta = 0$ si $m = n$.

Si $\Theta = (\theta_H) = (\theta_{i_0 \dots i_p})$ est un élément de bidegré (m, p) , on définira pour $p > 0$ un élément $\partial\Theta = (\eta_{i_0 \dots i_{p-1}})$ de bidegré $(m, p-1)$ au moyen de la formule

$$\eta_{i_0 \dots i_{p-1}} = \sum_{\mu, k} (-1)^\mu \theta_{i_0 \dots i_{\mu-1} k i_\mu \dots i_{p-1}}$$

où la sommation est étendue aux valeurs de μ, k pour lesquelles on a $|i_0 \dots i_{\mu-1} k i_\mu \dots i_{p-1}| \in N$; ces valeurs sont en nombre fini, et chaque terme du second membre est une forme à support compact contenu dans $U_{|i_0 \dots i_{p-1}|}$, donc cette formule définit bien un élément $\partial\Theta$, qui est fini si Θ est fini. De même, si $\Theta = (\theta_i)$ est de bidegré $(m, 0)$, on posera $\partial\Theta = \sum_k \theta_k$; $\partial\Theta$ est alors une forme sur V , à support compact si Θ est fini. On a $\partial^2 = 0$; et ∂ est permutable avec d et $\int$.

Si (f_i) désigne de nouveau une partition différentiable de l'unité subordonnée à $\mathcal{U}$, on désignera par L l'opérateur qui, à tout élément $\Theta = (\theta_H)$ de bidegré (m, p) , fait correspondre l'élément $L\Theta = (\zeta_H)$ de bidegré $(m, p+1)$ défini par $\zeta_{iH} = f_i \theta_H$; de même, si θ est une forme de degré m sur V , on désignera par $L\theta$ l'élément de bidegré $(m, 0)$ défini par $L\theta = (f_i \theta)$; on a alors $\theta = \partial L\theta$. Si Θ est un élément, on a $\Theta = \partial L\Theta + L\partial\Theta$; donc $\partial\Theta = 0$ entraîne $\Theta = \partial L\Theta$.

Cela posé, la théorie du § 3 s'applique sans aucun changement si on substitue les éléments différentiels de bidegré $(n-m, p)$ aux éléments singuliers de bidegré (m, p) , les formes de degré $n-m$ aux chaînes de dimension m , et les opérateurs $d, \int, \partial$ aux opérateurs b, b_0, ∂ . On partira donc des relations

$$\theta = \partial\Theta_0; \quad d\Theta_h = \partial\Theta_{h+1} \quad (0 \leq h \leq m-1); \quad \int \Theta_m = Z, \quad (\text{III})$$

où θ est une forme de degré $n-m$, Θ_h un élément différentiel de bidegré $(n-m+h, h)$ pour $0 \leq h \leq m$, et Z une chaîne de N de dimension m à coefficients réels; et on conclut, comme au § 3, que (III) établit un isomorphisme entre le groupe de de Rham de V de degré $n-m$ et le groupe d'homologie de N de dimension m à coefficients réels. Il n'y a rien à changer si on se borne aux formes à support compact sur V et aux éléments et chaînes finis. On pourrait aussi, naturellement, se servir d'éléments alternés en modifiant ∂ et L comme il a été dit au § 3.

Enfin, la dualité établie au § 3 entre coéléments différentiels et éléments singuliers se transporte ici aux coéléments et éléments différentiels. Si $\Omega = (\omega_H)$ est un coélément différentiel de bidegré (m, p) , et $\Theta = (\theta_H)$ un élément différentiel de bidegré $(n-m, p)$, et que l'un d'eux soit fini, on posera $(\Theta, \Omega) = \sum_H \int_{U_H} \theta_H \wedge \omega_H$; on a alors

$(\partial\Theta, \Omega) = (\Theta, \delta\Omega)$, mais la formule de Stokes donne cette fois $(d\Theta, \Omega) = (-1)^{n-m}(\Theta, d\Omega)$ si Θ est de bidegré $(n - m - 1, p)$ et Ω de bidegré (m, p) , de sorte qu'on a, pour deux suites satisfaisant respectivement à (III) et à (I), $\int_V \theta \wedge \omega = (-1)^{mn + \frac{m(m-1)}{2}}(Z, E)$. La conclusion est que les relations de dualité entre l'homologie et la cohomologie de N se transportent aux groupes de de Rham de dimensions complémentaires sur V . En particulier, si V est compacte, on voit que la forme bilinéaire $\int_V \theta \wedge \omega$ met en dualité les groupes de de Rham de degrés $n - m$ et m , respectivement.

A titre d'exemple, considérons les groupes relatifs aux dimensions 0 et n ; pour simplifier le langage, supposons V connexe, le cas général se déduisant trivialement de là par formation de sommes directes ou de produits, suivant qu'il s'agit des groupes à support compact ou non. Les groupes de dimension 0 se déterminent immédiatement; le groupe d'homologie finie de V de dimension 0 est libre et engendré par la classe d'un cycle réduit à un point; si V n'est pas compacte, le groupe d'homologie infinie de dimension 0 s'annule; le groupe de de Rham de degré 0 à support quelconque est engendré par la forme 1, et le même groupe à support compact s'annule si V n'est pas compacte. D'après les résultats du présent §, on en conclut que le groupe de de Rham de degré n à support compact est engendré par la classe d'une forme ω_0 de degré n telle que $\int_V \omega_0 = 1$, et que le groupe de de Rham de degré n à support quelconque s'annule si V n'est pas compacte; et, pour qu'une forme ω de degré n à support compact puisse s'écrire $\omega = d\eta$, avec η à support compact, il faut et il suffit que $\int_V \omega = 0$. Au moyen des résultats du § 3, on peut alors conclure qu'il existe un cycle singulier différentiable t_0 de dimension n tel que $\int_{t_0} \omega_0 = 1$; alors, si ω est à support compact, on a $\omega = c\omega_0 + d\eta$ avec $c = \int_V \omega$ et η à support compact, donc $\int_{t_0} \omega = c$, et par suite on a $\int_{t_0} \omega = \int_V \omega$ quel que soit ω à support compact, ce qui entraîne que le support de t_0 est V ; on peut conclure aussi que tout cycle fini t tel que $\int_t \omega_0 = 0$ est le bord d'une chaîne finie; donc le groupe d'homologie de dimension n à support compact, à coefficients réels, s'annule si V n'est pas compacte. Si on suppose V compacte, on peut conclure de plus que le groupe d'homologie de V de dimension n , à coefficients réels, est engendré par t_0 . Mais nous n'avons pas

prouvé qu'on puisse prendre pour t_0 un cycle entier, ni que ce cycle engendre le groupe d'homologie de V de dimension n à coefficients entiers ; pour cela il faudrait faire usage, soit d'une triangulation de V , soit d'une théorie du degré d'application pour les applications différentiables, soit de moyens topologiques plus puissants tels que ceux que fournit la théorie de Cartan, qui contient bien entendu les résultats en question.

§ 5. Le théorème d'homotopie

Comme on l'a remarqué, le fait que le nerf N de $\mathcal{U}$ a même homologie que V dépend seulement des propriétés homologiques des ensembles U_J . Si on tient compte du fait qu'ils sont homotopiquement triviaux, on obtient un résultat beaucoup plus précis ; c'est que N a même type d'homotopie que V . Il s'ensuit que N peut être substitué à V dans tout problème qui ne dépend que du type d'homotopie, et par exemple dans la plupart des questions concernant les espaces fibrés de base V ; dans de telles circonstances, le nerf d'un recouvrement simple de V peut donc souvent servir aux mêmes usages qu'une triangulation de V ; il semble qu'on ait là un outil élémentaire très maniable dans l'étude des variétés. C'est ce que montre aussi l'application qu'en a faite récemment G. de Rham à l'étude des invariants dits de torsion ⁶⁾ ; il est remarquable que ce ne sont pas là des invariants du type d'homotopie. Il se peut donc que les nerfs des recouvrements simples aient des propriétés encore plus précises que celle qui va être indiquée maintenant.

Le résultat qui va suivre est de nature purement topologique. Pour l'énoncer, rappelons qu'on dit qu'un espace B a la *propriété d'extension* si toute application continue dans B d'une partie fermée X d'un espace normal A peut être prolongée à une application continue de A dans B .

Soit alors $\mathcal{U} = (U_i)_{i \in I}$ un recouvrement localement fini d'un espace E par des ouverts U_i ; soit N son nerf. On dira que $\mathcal{U}$ est *topologiquement simple* si, pour tout $J \in N$, l'ensemble $U_J = \bigcap_{i \in J} U_i$ possède la propriété d'extension.

Notre théorème s'énonce alors comme suit ⁷⁾ : *si E est un espace tel que $E \times E \times [0, 1]$ soit normal, et si $\mathcal{U}$ est un recouvrement topologiquement simple de E , le nerf N de $\mathcal{U}$ a même type d'homotopie que E .*

La démonstration s'appuyera sur le lemme suivant :

⁶⁾ loc. cit., note 2.

⁷⁾ Dans le travail déjà cité (note 2), de Rham reproduit une partie de la démonstration qui suit, réduite à ce qui suffit au cas particulier qu'il a en vue. Un résultat apparenté au nôtre a été publié par K. Borsuk pour les espaces de dimension finie (On the imbedding of systems of compacta in simplicial complexes, Fund. Math. 35 (1948) p. 217); les démonstrations n'ont, semble-t-il, rien de commun.

Lemme. Soit E un espace tel que $E \times E \times [0, 1]$ soit normal. Soit $(X_i)_{i \in I}$ une famille localement finie de parties fermées de E ; soit N son nerf; pour $J \in N$, soit $X_J = \bigcap_{i \in J} X_i$. Soit $(U_J)_{J \in N}$ une famille de parties de E telle que, pour tout $J \in N$, U_J ait la propriété d'extension et contienne X_J , et qu'on ait $U_J \subset U_{J'}$, chaque fois que $J \supset J'$, $J \in N$, $J' \in N$. Alors il existe une application continue $F(x, y, t)$ de $\bigcup_{i \in I} (X_i \times X_i \times [0, 1])$ dans E telle que, pour tout $J \in N$, $x \in X_J$ et $y \in X_J$ entraîne $F(x, y, t) \in U_J$ et $F(x, x, t) = x$ quel que soit t , $F(x, y, 0) = x$, et $F(x, y, 1) = y$.

Pour toute partie N' de N , posons $Y(N') = \bigcup_{J \in N'} (X_J \times X_J \times [0, 1])$; considérons toutes les applications continues F' d'ensembles $Y(N')$ dans E qui satisfont, là où elles sont définies, à toutes les conditions du lemme; on les ordonnera en disant que $F' > F''$ si $Y(N') \supset Y(N'')$ et si F' coïncide avec F'' sur $Y(N'')$. En tenant compte du fait que (X_i) est localement finie, et que par suite tout $x \in E$ a un voisinage qui ne rencontre qu'un nombre fini des X_J , on voit immédiatement qu'on peut appliquer aux F' ainsi ordonnées le théorème de Zorn. Soit donc F' une telle application maximale c'est-à-dire non prolongeable, définie sur $Y(N')$. Supposons qu'il existe $J \in N$ tel que $X_J \times X_J \times [0, 1]$ ne soit pas contenu dans $Y(N')$; parmi les $J' \in N$ en nombre fini qui contiennent J , prenons-en un qui ait la même propriété et qui ait le plus grand nombre possible d'éléments; en remplaçant J par celui-ci, on voit qu'on peut supposer de plus que $X_{J'} \times X_{J'} \times [0, 1] \subset Y(N')$ pour tout $J' \neq J$ tel que $J' \supset J$. Comme $X_J \times X_J \times [0, 1]$ est une partie fermée de $E \times E \times [0, 1]$, c'est un espace normal; les points (x, y, t) de cet espace qui satisfont à $x = y$, à $t = 0$ et à $t = 1$ en forment des parties fermées; son intersection avec $Y(N')$ est fermée aussi en raison du caractère localement fini de la famille (X_i) ; il s'ensuit qu'il y a une application continue $G(x, y, t)$ de $X_J \times X_J \times [0, 1]$ dans U_J qui coïncide avec F' sur l'intersection de cet ensemble avec $Y(N')$ et qui satisfasse à $G(x, x, t) = x$, $G(x, y, 0) = x$, $G(x, y, 1) = y$. Montrons que la fonction qui coïncide avec F' sur $Y(N')$ et avec G sur $X_J \times X_J \times [0, 1]$ a toutes les propriétés énoncées dans le lemme, contrairement à l'hypothèse que F' n'est pas prolongeable. Le seul point à vérifier est que, si $J' \in N$ et si x et y sont dans $X_J \cap X_{J'}$, $G(x, y, t)$ est dans $U_{J'}$; c'est évident si $J' \subset J$, puisqu'alors $U_{J'} \supset U_J$; dans le cas contraire, posons $J'' = J \cup J'$; on aura $J'' \neq J$ et $J'' \in N$, donc, en vertu de l'hypothèse faite sur J , $X_{J''} \times X_{J''} \times [0, 1] \subset Y(N')$, donc $F'(x, y, t) \in U_{J''} \subset U_{J'}$, d'où la conclusion annoncée puisque G coïncide avec F' en (x, y, t) . Donc, quel que soit $J \in N$, on a

$X_J \times X_J \times [0, 1] \subset Y(N')$, et en particulier $X_i \times X_i \times [0, 1] \subset Y(N')$ pour tout $i \in I$; F' est donc la fonction F qu'il s'agissait de construire.

Corollaire. *Les hypothèses étant celles du lemme, soient f, f' deux applications continues d'un espace A dans E , telles que, quel que soit $u \in A$, il y ait un $i \in I$ pour lequel $f(u) \in X_i$ et $f'(u) \in X_i$. Alors f et f' sont homotopes.*

En effet, $F(f(u), f'(u), t)$ est une homotopie joignant f à f' .

Nous pouvons passer maintenant à la démonstration de notre théorème. Soit d'abord $\mathfrak{U} = (U_i)$ n'importe quel recouvrement localement fini d'un espace normal E par des ouverts U_i ; alors il y a une partition (f_i) de l'unité subordonnée à $\mathfrak{U}$; posons, pour $p \in E$, $f(p) = (f_i(p))$; f est une application continue de E dans le nerf N de $\mathfrak{U}$, réalisé géométriquement conformément aux définitions rappelées à la fin du § 1. Si $p \in E$, et si J est l'ensemble des $i \in I$ tels que $p \in U_i$, $f(p)$ est dans le simplexe Σ_J de N ; si donc (f'_i) est une autre partition de l'unité subordonnée à $\mathfrak{U}$, le segment de droite joignant $f(p)$ et $f'(p)$ est contenu dans Σ_J , donc dans N ; par suite, l'application $p \rightarrow (1-t)f(p) + tf'(p)$ est une homotopie joignant f à f' ; la classe d'homotopie de f est donc complètement déterminée par la donnée de $\mathfrak{U}$.

Supposons maintenant que les $U_J = \bigcap_{i \in J} U_i$, pour $J \in N$, aient tous leurs groupes d'homotopie nuls; autrement dit, toute application continue dans l'un des U_J de la frontière d'un simplexe de dimension m peut se prolonger à tout le simplexe; pour $m = 1$, cela veut dire que U_J est connexe par arcs. Pour tout $J \in N$, soit e_J le centre de gravité de Σ_J . Considérons toutes les suites croissantes $J_0 \subset J_1 \subset \dots \subset J_m$ d'éléments tous distincts de N ; pour une telle suite, soit $\Sigma'(J_0, \dots, J_m)$ le simplexe de sommets $e_{J_0}, \dots, e_{J_m}$; N est la réunion de tous ces simplexes, qui en forment la subdivision barycentrique. On va définir par récurrence une application continue g de N dans E telle que $g(\Sigma'(J_0, \dots, J_m)) \subset U_{J_0}$ pour toute suite $J_0, \dots, J_m$. On prendra $g(e_J)$ quelconque dans U_J pour tout $J \in N$. Supposons g définie sur les simplexes de la subdivision barycentrique de N de dimension $\leq m-1$; alors g est définie sur la frontière du simplexe $\Sigma'(J_0, \dots, J_m)$, qui est la réunion des simplexes $\Sigma'_\mu = \Sigma'(J_0, \dots, J_{\mu-1}, J_{\mu+1}, \dots, J_m)$ pour $0 \leq \mu \leq m$. D'après l'hypothèse de récurrence, on a $g(\Sigma'_0) \subset U_{J_1} \subset U_{J_0}$, et $g(\Sigma'_\mu) \subset U_{J_0}$ pour $1 \leq \mu \leq m$; donc on peut prolonger g à une application de $\Sigma'(J_0, \dots, J_m)$ dans U_{J_0} . Si d'ailleurs g' est une autre application de N dans E satisfaisant à la même condition, on peut, par

une récurrence tout à fait analogue, construire une homotopie joignant g à g' ; la classe d'homotopie de g est donc bien déterminée par la condition qu'on s'est imposée.

Montrons que, dans ces conditions, $f \circ g$ est une application de N dans N homotope à l'application identique. Soit en effet F_i la réunion des images par g de tous les simplexes $\Sigma'(J_0, \dots, J_m)$ pour lesquels $i \in J_0$; comme ces simplexes sont en nombre fini, F_i est une partie compacte de U_i . Pour chaque i , soit U'_i une partie ouverte de U_i contenant F_i , telle que $\overline{U'_i} \subset U_i$, et que les U'_i forment encore un recouvrement de E ; puisque le choix de la partition (f_i) subordonnée à $\mathcal{U}$ est sans influence sur la classe d'homotopie de f , on peut la supposer choisie de telle sorte que $f_i > 0$ sur F_i et $f_i = 0$ en dehors de U'_i , pour tout $i \in I$. Soit $x = (x_i)$ un point de N ; choisissons un i tel que $x_i = \max_{j \in I} (x_j)$; alors x est dans un simplexe de la subdivision barycentrique de N ayant un sommet en e_i , et on a donc $g(x) \in F_i$, d'où $f_i(g(x)) > 0$. Pour tout $i \in I$ et tout $x = (x_i) \in N$, posons $\varphi_i(x) = \min[x_i, f_i(g(x))]$; les φ_i sont des fonctions continues ≥ 0 sur N , et on a $\varphi_i(x) = 0$ si $x_i = 0$, c'est-à-dire si x n'appartient pas à St_i ; de plus, d'après ce qu'on vient de montrer, il y a pour tout $x \in N$ un i tel que $\varphi_i(x) > 0$. Il s'ensuit que $\varphi = \sum_i \varphi_i$ est une fonction continue partout > 0 sur N , et par suite que les $h_i = \varphi_i/\varphi$ forment sur N une partition de l'unité subordonnée au recouvrement (St_i) ; si donc on pose $h(x) = (h_i(x))$, h est une application de N dans N . Si, pour $x \in N$, J est l'ensemble des $i \in I$ tels que $h_i(x) > 0$, on aura, pour tout $i \in J$, $x_i > 0$ et $f_i(g(x)) > 0$; alors $h(x)$ est dans Σ_J , et x et $f(g(x))$ sont tous deux dans St_J , de sorte que les segments de droite qui joignent $h(x)$ à x d'une part et à $f(g(x))$ d'autre part sont contenus dans N ; comme tout à l'heure on conclut de là que h est homotope à l'application identique d'une part, et à $f \circ g$ d'autre part.

Enfin, soit $p \in E$, et soit J l'ensemble des $i \in I$ tels que $f_i(p) > 0$; on aura donc $p \in U'_i$ pour tout $i \in J$; on aura $f(p) \in \Sigma_J$, donc $f(p)$ appartiendra à un des simplexes de la subdivision barycentrique de Σ_J ; mais ce sont là, avec les notations employées plus haut, les simplexes $\Sigma'(J_0, \dots, J_m)$ avec $J_m \subset J$; si alors on prend $i \in J_0$, on aura $g(f(p)) \in F_i$; donc p et $g(f(p))$ sont tous deux dans U'_i . Posons $X_i = \overline{U'_i}$; soit N' le nerf de la famille (X_i) ; on aura $N' \subset N$. Si de plus on suppose maintenant que les U_J ont la propriété d'extension, c'est-à-dire que $\mathcal{U}$ est topologiquement simple, on voit que les familles $(X_i)_{i \in I}$ et $(U_J)_{J \in N'}$ satisfont à toutes les conditions du lemme de tout à l'heure; d'après le corollaire de ce lemme, on peut donc affirmer que $g \circ f$ est homotope à

l'application identique de E pourvu que $E \times E \times [0, 1]$ soit normal. Le théorème annoncé est donc complètement démontré.

Supposons en particulier que l'un des U_i soit recouvert par la réunion des autres, donc qu'on ait $U_i = \bigcup_{j \neq i} U_{ij}$, et que $U_i \times U_i \times [0, 1]$ soit normal. Alors les U_{ij} non vides forment un recouvrement topologiquement simple de U_i , dont le nerf a donc même type d'homotopie que U_i ; ce type est trivial, puisque U_i a la propriété d'extension et est donc contractile. Si on omet U_i dans le recouvrement $\mathfrak{U}$, ce qui reste est encore un recouvrement $\mathfrak{U}'$ de E en vertu de l'hypothèse; le nerf N' de $\mathfrak{U}'$ se déduit de N en retranchant St_i ; et la frontière de St_i n'est autre que le nerf du recouvrement $(U_{ij})_{j \neq i}$ de U_i , donc est un complexe fini homotopiquement trivial (c'est-à-dire contractile); comme on le voit facilement, cela équivaut à dire qu'il existe une rétraction de l'adhérence $\overline{St_i}$ de St_i sur sa frontière $\overline{St_i} \cap N'$, donc une rétraction de N sur N' , et même qu'il existe une telle rétraction dépendant continuellement d'un paramètre, c'est-à-dire une application continue $F(x, t)$ de $N \times [0, 1]$ dans N telle que $F(x, 0) = x$ et $F(x, 1) \in N'$ pour tout $x \in N$, $F(x, t) = x$ quel que soit t pour tout $x \in N'$, et $F(x, t) \in \overline{St_i}$ quel que soit t pour tout $x \in \overline{St_i}$. En particulier, de Rham a montré (loc. cit., note 6) que, si on se borne à considérer la famille des recouvrements simples qu'il appelle «convexes» d'une variété différentiable compacte, on peut toujours passer de l'un à l'autre de ces recouvrements par insertions et omissions successives d'ensembles superflus; le résultat que nous venons de démontrer indique, d'une manière un peu plus précise que ne le faisait de Rham, l'effet de ces opérations sur les nerfs des recouvrements correspondants.

(Reçu le 22 novembre 1951.)