

ANNALES DE L'I. H. P.

OTTON MARTIN NIKODYM

Un nouvel appareil mathématique pour la théorie des quanta

Annales de l'I. H. P., tome 11, n° 2 (1949), p. 49-112

[<http://www.numdam.org/item?id=AIHP_1949__11_2_49_0>](http://www.numdam.org/item?id=AIHP_1949__11_2_49_0)

© Gauthier-Villars, 1949, tous droits réservés.

L'accès aux archives de la revue « Annales de l'I. H. P. » 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/>

Un nouvel appareil mathématique pour la théorie des quanta

par

Otton Martin NIKODYM.

PRÉFACE.

Je remercie le Conseil de l'Institut Henri Poincaré qui m'a fait le grand honneur de m'inviter à donner quatre conférences sur un appareil mathématique susceptible d'être utilisé dans les théories quantiques ⁽¹⁾.

Il y a quinze ans que j'ai commencé à étudier de plus près le spectre continu des transformations hermitiennes selfadjointes, en me proposant d'en trouver une théorie libérée des procédés artificiels et à caractère non géométrique, et de mettre mieux en évidence les différents phénomènes relatifs à ce spectre, par exemple celui de sa multiplicité. Dans cet ordre d'idées, la voie était ouverte par M. J. v. Neumann, qui, dans son ouvrage, *Mathematische Grundlagen der Quantenmechanik* (Berlin, 1932) a remarqué certaines relations intéressantes entre la logique des propositions et les projecteurs. Si l'on envisage des opérations géométriques bien simples portant sur les sous-espaces fermés ⁽²⁾ de l'espace de Hilbert-Hermite et qu'on pense aux principes fondamentaux de l'algèbre de Boole, on s'aperçoit que ce sont précisément des corps de Boole composés de tels espaces qui permettent de traduire

⁽¹⁾ Les conférences ont été faites les 4, 6, 11 et 13 février 1947 à l'Institut Henri Poincaré à Paris. Le résumé se trouve dans : O. Nikodym, Sur les tribus de sous-espaces d'un espace de Hilbert-Hermite. *C. R.*, 224, p. 522-524 (1947); tribus et lieux attachés à une classe ordonnée de sous-espaces d'un espace de Hilbert-Hermite. *C. R.*, 224, p. 628-630 (1947); système général de coordonnées dans l'espace séparable de Hilbert-Hermite. *C. R.*, 224, p. 778-780 (1947).

⁽²⁾ M. H. Stone, *Linear transformations in Hilbert space*, New-York, 1932.

en langage géométrique la découverte de Neumann, et que ces corps de Boole doivent servir de point de départ à la théorie en question. Un autre principe se dégage de la théorie des fonctions d'ensemble abordée par Lebesgue et développée par MM. Ch. de la Vallée-Poussin, Hahn, Carathéodory et surtout par M. Fréchet qui, en admettant un point de vue abstrait et général, a créé une importante théorie de la mesure et de l'intégration généralisées.

Ces deux principes, un peu généralisés et modifiés, permettent de rendre plus géométriques qu'elles n'étaient jusqu'à présent, la théorie des opérateurs hermitiens et celle des opérateurs normaux, bien que je n'aie pas encore réussi en ce qui concerne la démonstration du théorème spectral.

La théorie de l'espace hilbertien a subi, dès sa naissance, des modifications considérables. Abordée par D. Hilbert, elle a imité les méthodes de la théorie des équations intégrales qui avaient été résolues par Fredholm d'une manière ingénieuse par l'introduction d'un paramètre arbitraire complexe. Hilbert, tout en suivant cette voie, a utilisé comme outil les systèmes infinis d'équations linéaires à une infinité d'inconnues. Les méthodes axiomatiques et abstraites n'ayant pas encore été développées, Hilbert s'est servi des formes linéaires et bilinéaires à un nombre infini de variables, ces éléments étant alors les meilleurs qu'il eût à sa disposition. La méthode d'approximation lui a permis de montrer, dans toute sa généralité, le phénomène du spectre continu dont les exemples avaient été à peine découverts par ses prédécesseurs et, de plus, l'a conduit à sa célèbre formule spectrale pour les matrices réelles symétriques bornées, formule s'exprimant par une intégrale de Stieltjes. La théorie des quanta a attiré l'attention des mathématiciens sur les matrices et formes complexes hermitiennes auxquelles s'appliquent une théorie analogue et les mêmes méthodes. Dès lors on n'étudie plus que le cas hermitien.

La méthode de Hilbert possède un caractère arithmétique et est analogue à la méthode de l'algèbre des matrices finies dans leurs rapports avec les systèmes d'équations linéaires. C'est E. Schmidt qui a appliqué, pour la première fois, aux équations intégrales, l'idée d'une géométrie analytique dans l'espace euclidien à une infinité de dimensions. La théorie est ainsi devenue plus simple et plus harmonieuse et les théorèmes fondamentaux ont obtenu une explication géométrique.

intuitive. Cette voie géométrique a été poursuivie par M. F. Riesz qui a étudié très profondément la théorie hilbertienne en se plaçant dans l'espace à une infinité de dimensions et en remplaçant les matrices par des transformations linéaires portant sur des vecteurs de cet espace. On doit un grand progrès à M. J. v. Neumann qui, suivant la même voie, a fondé la théorie de l'espace de Hilbert-Hermite de manière axiomatique. Il a pu se débarrasser ainsi des éléments non essentiels et, en même temps, il a ouvert la voie à diverses généralisations importantes. C'est aussi von Neumann qui a montré, dans son Mémoire de *Journ. f. d. r. u. ang. Mathem.*, 161, (1931), que les matrices infinies, employées surtout par les théoriciens de la physique moderne, représentent un instrument très mal commode et très peu adapté à leurs problèmes. Le point de vue géométrique des transformations linéaires dans un espace vectoriel a ainsi triomphé et la théorie des transformations symétriques s'est développée rapidement grâce aux travaux de MM. von Neumann et Stone.

On peut suivre cette évolution vers la géométrie en considérant les changements de terminologie au cours du temps. L'*Einzelform*, respectivement l'*Einzelmatrix* (dénominations empruntées à l'algèbre), s'est transformée en un projecteur, c'est-à-dire un opérateur effectuant la projection orthogonale de l'espace total sur un sous-espace fermé. Néanmoins jusqu'à présent règne l'expression algébrique, *résolution de l'identité* pour désigner l'échelle spectrale d'espaces invariants. Même le terme *selfadjoint* se rattachant plus à l'algèbre ou à l'analyse qu'à la géométrie montre que, jusqu'à présent, on ne s'est pas débarrassé complètement des suggestions arithmétiques. Dans les travaux récents, on préfère même quelquefois se laisser guider par l'algèbre moderne abstraite plutôt que par la géométrie, preuve de l'engouement pour l'algèbre moderne, qu'ont suscité ses méthodes puissantes.

La théorie des opérateurs symétriques, développée surtout par MM. F. Riesz, J. von Neumann et M. H. Stone qui en a fait un exposé très complet et très précis dans son ouvrage connu, semble être achevée ⁽³⁾ et l'on pourrait croire que, en ce qui concerne les principes,

⁽³⁾ M. G. Julia dans ses nombreux travaux insérés surtout dans les *C. R. de l'Académie des Sciences*, semble lui donner les derniers coups de pinceau, en découvrant des propriétés les plus cachées.

Voir aussi J. DIXMIER, *C. R. Acad. Sc.*, Paris, 1947, et les travaux de M. A. NEUMARK dans le *Bull. Ac. Sc. U.R.S.S.* dès 1943.

il n'y ait rien à ajouter. Néanmoins, de même que dans la géométrie élémentaire les vrais éléments sont le point, la droite, le plan et leurs relations mutuelles, il me semble qu'on doit fonder aussi la théorie de l'espace à une infinité de dimensions, en considérant à la base surtout les points, respectivement les vecteurs, les sous-espaces et leurs relations mutuelles comme par exemple l'orthogonalité; les transformations et les fonctionnelles paraissent être des éléments dérivés. De plus, il me semble plus naturel de fonder la théorie des opérations linéaires sur la considération de leurs variétés linéaires invariantes que sur l'utilisation d'un paramètre emprunté à la théorie des équations intégrales. En suivant systématiquement ce programme, nous jetterons de la lumière sur le phénomène du spectre et nous lui donnerons à peu près le même caractère intuitif que possédaient déjà le développement de Fourier, qui n'est autre que la décomposition d'un vecteur en ses composantes orthogonales deux à deux, et le théorème de Parseval-Riesz-Fischer qui n'est autre que le théorème généralisé de Pythagore. Les différentes démonstrations connues du théorème spectral n'ont pas jusqu'à présent un caractère purement géométrique : on obtient ce théorème, abstraction faite d'une approximation arithmétique, par l'intermédiaire soit d'une intégrale singulière de Stieltjes, soit par des artifices suggérés par l'analogie entre les polynômes et les formes hermitiennes quadratiques définies positives. Dans cet ordre d'idées on en possède même une démonstration très courte de Wecken ⁽⁴⁾. Pour ma part je préférerais une démonstration longue mais naturelle et sans artifices, et il me semble très intéressant de chercher une démonstration géométrique de ce théorème fondamental. Le théorème fondamental de Hellinger-Hahn considéré dans la théorie des invariants unitaires d'une transformation selfadjointe se démontre par le détour de l'*operational calculus* fondé par J. von Neumann et basé sur l'emploi des intégrales de Stieltjes; et cependant il semble qu'on devrait mettre en évidence d'une manière directe le phénomène qui s'y rattache, le phénomène de la multiplicité du spectre.

Les physiciens, qui sont toujours doués d'une intuition excellente et contrôlée par l'expérience utilisent la décomposition d'un vecteur en ses composantes d'une manière directe, même quand leur nombre

(4) F. J. WECKEN, *Zur Theorie linearer Operatoren* (Math. Ann., 110, 1935).

est non dénombrable, et la célèbre formule de M. Dirac a donné également de bons résultats, bien que ces procédés manquent jusqu'à présent d'une justification précise et bien adaptée à la pratique.

Toutes ces considérations m'ont guidé dans mes recherches. Les travaux de MM. von Neumann, Stone, Fréchet et Tarski ^(*) m'ont dirigé vers l'étude des tribus d'espaces et de la mesure sur ces tribus, en connexion avec les opérations spatiales étudiées dans l'ouvrage de M. Stone. La notion de tribu d'espaces (corps de Boole) se montrait vraiment fondamentale et elle sert de base aux considérations du présent travail. Notons, en passant, que l'importance de telles tribus fut aussi remarquée, pendant la guerre, par F. Wecken et par des mathématiciens japonais comme par exemple H. Nakano, bien que la considération algébrique d'un anneau semble quelquefois prédominer chez eux. Les problèmes de prolongement d'une tribu m'ont conduit à reconnaître dans la permutabilité de deux projecteurs une notion purement géométrique de *compatibilité de deux espaces*. J'introduis la notion de *champ de rayonnement d'un ensemble de vecteurs sur une tribu d'espaces* et celle de *lieu*, notions qui m'ont permis de trouver une démonstration géométrique simple de la partie essentielle du théorème de Hellinger-Hahn. La notion de *tribu saturée*, en connexion avec la notion de *vecteur générateur de l'espace par rapport à une tribu* m'a permis de m'orienter dans la multitude de toutes les mesures dénombrablement additives définies sur une tribu d'espaces. En désirant trouver une théorie précise de l'intégrale de M. Dirac je me suis aperçu qu'on ne peut y parvenir entièrement avec les seules notions simples de vecteur ou de fonction ordinaire. Il fallut créer une notion nouvelle qui, tout en conservant un caractère de vecteur, puisse servir d'axe dans un système général de coordonnées, même dans le cas où leur nombre n'est pas dénombrable. Mais cela ne peut se faire que dans le cas où l'on peut *linéariser* (ou *planifier*) une tribu d'espaces, c'est-à-dire où l'on peut trouver dans la tribu donnée une sous-classe d'espaces ordonnée d'une manière linéaire (respectivement plane), engendrant, par les opérations spatiales, la tribu donnée. Ces considérations m'ont conduit à la notion de lieu d'une classe ordonnée d'espaces et à la notion de *quasi-vecteur* et de *quasi-nombre*. Ces

(*) A. TARSKI, *Zur Grundlegung der Booleschen Algebra*, I (*Fund. Math.*, 24, 1935).

notions sont liées à celle de *système général de coordonnées* et, en particulier, à celle de système continu de coordonnées dans l'espace séparable de Hilbert-Hermite. Pour avoir la possibilité de décomposer un vecteur en ses composantes dans le cas où leur nombre est non dénombrable, il fallait développer une nouvelle théorie de l'intégration analogue à celle de Burkill.

La guerre et le régime barbare des rafles de la police hitlérienne ne m'ont pas permis de consulter, durant la guerre, la littérature scientifique et m'ont obligé de travailler dans un isolement complet. Ce n'est qu'après avoir trouvé déjà tous les principes et théorèmes fondamentaux de ma théorie que j'ai pu, une fois la guerre finie, et après la lecture des travaux récents, perfectionner et simplifier les détails et, de plus, démontrer quelques théorèmes nouveaux. Aussi ai-je retrouvé quelques-unes de mes idées chez différents auteurs ⁽⁶⁾.

La théorie que je vais exposer n'est pas simple. Les notions à introduire sont nombreuses, mais il me semble qu'elles ne sont pas artificielles, et qu'elles sont liées à la nature des choses. L'algorithme que je vais créer au moyen de considérations assez longues est, néanmoins, simple et maniable. Une situation analogue se présente dans la théorie classique de Cantor et Dedekind relative aux nombres réels : l'exposé est long et fourmille de lemmes; et cependant l'usage des nombres réels est plus simple que celui des nombres rationnels. La raison de la complication du fondement dans cette théorie est que ses sources sont assez profondes.

Je ne suis nullement d'avis que la théorie qui va être exposée est fondée de la manière la plus simple possible. On doit la regarder comme un essai de ce qui doit être fait. Les applications éventuelles seront la meilleure preuve de son utilité et de plus elles suggéreront de changements et des simplifications qui amèneront cet appareil mathématique à être toujours plus pratiqué et mieux adapté aux exigences de la Physique théorique moderne.

⁽⁶⁾ Surtout chez M. HIDEGORO NAKANO, *Unitäriinvariante hypermaximale normale Operatoren*, (*Ann. of Math.*, 42, 1941; *Unitäriinvarianten im allgemeinen Euklidischen Raum* (*Mathem. Annalen*, 43, 1941, p. 118) et chez M. G. JULIA, *Sur les projecteurs de l'espace hilbertien ou unitaire* (*C. R. Acad. Sc.*, 214, 1942, p. 456).

I.

PRÉLIMINAIRES.

1. Par l'espace de Hilbert-Hermite nous comprendrons un ensemble non vide d'êtres quelconques doué de la structure suivante. Il y a trois opérations portant sur les éléments $\vec{x}, \vec{y}, \dots$, de 1 appelés vecteurs : l'addition $\vec{x} + \vec{y}$, la multiplication $\lambda \vec{x}$ par un nombre complexe λ quelconque, ces deux opérations étant toujours possibles et fournissant un vecteur, et le produit scalaire $(\vec{x}, \vec{y})$ de deux vecteurs, donnant un nombre complexe. En outre, il y a une relation $\vec{x} = \vec{y}$ appelée égalité, cette relation jouissant des propriétés formelles de l'identité. Voici les axiomes auxquels doivent obéir la relation d'égalité et les opérations ci-dessus :

I. L'addition, la multiplication et le produit scalaire sont invariants par rapport à l'égalité ⁽¹⁾;

II. $\vec{x} + \vec{y} = \vec{y} + \vec{x}$; $\vec{x} + (\vec{y} + \vec{z}) = (\vec{x} + \vec{y}) + \vec{z}$; si $\vec{x} + \vec{y} = \vec{x} + \vec{y}'$, on a $\vec{y} = \vec{y}'$; $\lambda(\vec{x} + \vec{y}) = \lambda\vec{x} + \lambda\vec{y}$; $(\lambda + \mu)\vec{x} = \lambda\vec{x} + \mu\vec{x}$; $\lambda(\mu\vec{x}) = (\lambda\mu)\vec{x}$; $1\vec{x} = \vec{x}$ ⁽²⁾;

III. $(\vec{x}, \vec{y})$ est une forme bilinéaire hermitienne définie positive, c'est-à-dire :

$(\vec{x}, \vec{y}) = \overline{(\vec{y}, \vec{x})}$; $(\vec{x}, \vec{y} + \vec{z}) = (\vec{x}, \vec{y}) + (\vec{x}, \vec{z})$; $\lambda(\vec{x}, \vec{y}) = (\vec{x}, \lambda\vec{y})$, $(\vec{x}, \vec{x}) \geq 0$; la relation $(\vec{x}, \vec{x}) = 0$ équivaut à $\vec{x} = \vec{0}$.

⁽¹⁾ Nous ne considérerons que des ensembles propres $\mathcal{S}$ de vecteurs, ce qui exprime que, si $\vec{x} \in \mathcal{S}$, $\vec{x} = \vec{y}$, on a aussi $\vec{y} \in \mathcal{S}$.

⁽²⁾ Voir par exemple M. H. STONE, *Linear transformations in Hilbert space and their applications to analysis*. New-York, 1932. (Amer. Mat. Soc. Coll. Publ., Vol. XV).

De I et II il résulte l'existence et l'unicité d'un vecteur nul $\vec{0}$ pour lequel on a $\vec{x} + \vec{0} = \vec{x}$, $0\vec{x} = \vec{0}$ quel que soit $\vec{x}$.

On définit : $\underset{\text{df}}{\vec{x} - \vec{y}} = \vec{x} + (-1)\vec{y}$ ⁽¹⁾; $-\vec{x} \underset{\text{df}}{=} \overset{\rightarrow}{0} - \vec{x}$. On a $\lambda(\vec{x}, \vec{y}) = (\bar{\lambda}\vec{x}, \vec{y})$ où $\bar{\lambda}$ désigne le nombre conjugué de λ . Le produit scalaire induit la notion de *module* d'un vecteur $|\vec{x}| \underset{\text{df}}{=} \sqrt{(\vec{x}, \vec{x})}$. On a $|(\vec{x}, \vec{y})| \leq |\vec{x}| |\vec{y}|$ (inégalité de Cauchy-Schwarz) et $|\vec{x} + \vec{y}| \leq |\vec{x}| + |\vec{y}|$ (inégalité de Minkowski). La notion d'*écart de deux vecteurs* $|\vec{x}, \vec{y}| \underset{\text{df}}{=} |\vec{x} - \vec{y}|$ obéit aux lois : 1° $|\vec{x}, \vec{y}| = |\vec{y}, \vec{x}|$; 2° $|\vec{x}, \vec{y}| \leq |\vec{x}, \vec{z}| + |\vec{z}, \vec{y}|$; 3° $|\vec{x}, \vec{y}| \geq 0$; 4° $|\vec{x}, \vec{y}| = 0$ équivaut à $\vec{x} = \vec{y}$. Par conséquent cet écart organise l'espace $\mathfrak{r}$ en *espace métrique*, donc en *topologie* dans laquelle la limite $\vec{x}$ d'une suite infinie $\{\vec{x}_n\}$ de vecteurs se définit par la relation $|\vec{x}_n, \vec{x}| \rightarrow 0$. Dans le cas où cette métrique n'est pas complète elle peut toujours être complétée par le procédé connu de Cantor-Charles Meray au moyen des suites fondamentales de Cauchy.

Nous supposons dans ce qui va suivre que

IV. *La métrique engendrée par le produit scalaire est complète.* — La topologie en question n'est pas nécessairement séparable; l'espace $\mathfrak{r}$ peut avoir un nombre fini ou infini de dimensions. Nous supposons que $\mathfrak{r}$ ne se réduit pas au seul vecteur $\overset{\rightarrow}{0}$.

2. Les vecteurs $\vec{x}, \vec{y}$ sont dits *orthogonaux* ($\vec{x} \perp \vec{y}$), lorsque $(\vec{x}, \vec{y}) = 0$. Un ensemble non vide de vecteurs s'appelle *système orthonormal* lorsqu'il se compose de vecteurs de modules $= 1$ et orthogonaux deux à deux. Un tel système s'appelle *saturé* lorsque tout système orthonormal qui le contient coïncide avec lui.

Par *variété linéaire* nous entendons tout ensemble $\mathcal{E}$ non vide de vecteurs jouissant de la propriété suivante : si $\vec{x}, \vec{y} \in \mathcal{E}$ on a $\lambda\vec{x} + \mu\vec{y} \in \mathcal{E}$. Une variété linéaire s'appelle *sous-espace* (ou plus simplement, *espace*) lorsqu'elle est un ensemble fermé. Parmi les espaces mentionnons l'*espace nul* $\mathfrak{o}$ composé du seul vecteur $\overset{\rightarrow}{0}$ et l'*espace total* $\mathfrak{r}$ de tous les vecteurs. Si a, b sont des espaces, $a \subseteq b$ signifie que tout vecteur

(1) $\underset{\text{df}}{=}$ désigne : égal par définition (B. RUSSELL, *Principia Mathem.*).

appartenant à a appartient aussi à b . Si $a \subseteq b$ et si $a \neq b$, nous écrirons $a \subset b$. Le vecteur $\vec{x}$ est dit *orthogonal* à l'espace a ($\vec{x} \perp a$) si $\vec{x}$ est orthogonal à tout vecteur de a ; les espaces a, b sont dits *orthogonaux* ($a \perp b$), si chaque vecteur $\vec{x} \in a$ est orthogonal à chaque vecteur $\vec{y} \in b$. Si a est un espace, l'ensemble b de tous les vecteurs orthogonaux à a est un espace fermé; il sera désigné par $\text{co}a$ et nommé le *complémentaire* de a . On a $\text{co}(\text{co}a) = a$, $\text{co}0 = 1$. Étant donné un espace a et un vecteur $\vec{x}$, il existe une décomposition unique

$$\vec{x} = \vec{x}_a + \vec{x}_{\text{co}a},$$

telle que $\vec{x}_a \in a$, $\vec{x}_{\text{co}a} \in \text{co}a$. Le vecteur $\vec{x}_a$ s'appelle la *a-projection* de $\vec{x}$ et sera désigné aussi par $P_a \vec{x}$ ou $\text{Proj}_a \vec{x}$. Ce théorème est profond et sert de base à beaucoup de démonstrations dans ce qui va suivre ⁽¹⁾. Citons quelques propriétés bien connues de la projection : $P_a \vec{x} \in a$; $|P_a \vec{x}| \leq |\vec{x}|$; $(P_a \vec{x}, \vec{y}) = (\vec{x}, P_a \vec{y})$; $P_a(\lambda \vec{x} + \mu \vec{y}) = \lambda P_a \vec{x} + \mu P_a \vec{y}$; si $\vec{x}_n \rightarrow \vec{x}$, on a $P_a \vec{x}_n \rightarrow P_a \vec{x}$; si $a \subseteq b$, on a $P_a P_b \vec{x} = P_a \vec{x}$.

3. Il existe plusieurs réalisations de l'espace abstrait de Hilbert-Hermite. On en obtient une en appelant vecteur toute fonction $f(x)$ à valeurs complexes, définie presque partout sur le segment $\langle 0, 1 \rangle$ et au carré sommable ⁽²⁾. Par la somme $\vec{f} + \vec{g}$ on entend la fonction $f(x) + g(x)$, par $\lambda \vec{f}$, la fonction $\lambda f(x)$ et par le produit scalaire $(\vec{f}, \vec{g})$ le nombre $\int_0^1 \overline{f(x)} g(x) dx$ ⁽³⁾. Les vecteurs $\vec{f}, \vec{g}$ s'appellent égaux lorsque $f(x) = g(x)$ presque partout. Les axiomes I, II, III, IV se trouvent vérifiés et l'espace 1 de toutes les fonctions $f(x)$ aux carrés sommables est séparable et complet. Une autre interprétation analogue s'obtient en envisageant les fonctions aux carrés sommables définies presque partout dans $(-\infty, +\infty)$, ou bien dans $(0, +\infty)$. Les fonctions presque périodiques et continues dans $(-\infty, +\infty)$, si l'on définit

⁽¹⁾ Voir par exemple O. NIKODYM, *On Boolean fields of subspaces in an arbitrary Hilbert space* (Ann. de la Soc. polon. de Math., 17, 1938, p. 143).

⁽²⁾ $\langle 0, 1 \rangle$ désigne l'intervalle fermé $0 \leq x \leq 1$.

⁽³⁾ Voir par exemple G. JULIA, *Introduction mathématique aux théories quantiques*, Paris, Gauthier-Villars, 1939.

le produit scalaire d'une manière convenable, fournissent un exemple d'un espace non séparable. L'espace ordinaire complexe à 3 dimensions de vecteurs aux multiplicateurs complexes est un espace de Hilbert-Hermite complet et séparable, lorsqu'on définit le produit scalaire de la manière suivante $(\vec{x}, \vec{y}) = \bar{x}_1 y_1 + \bar{x}_2 y_2 + \bar{x}_3 y_3$.

4. Nous allons définir quelques opérations sur les espaces. Si a, b sont des espaces, appelons *somme* de a et b , ou $(a + b)$, le plus petit espace contenant a et b . Un tel espace existe toujours et est unique. Par ab , *produit* des espaces a et b , nous désignerons le plus grand espace contenu dans a et b . Un tel espace existe toujours et coïncide avec l'intersection des ensembles a et b de vecteurs. Il est unique. Définissons la *différence* $a - b$ comme $acob$. Deux espaces a, b s'appellent *disjoints* lorsque $ab = 0$ ⁽¹⁾.

L'addition et la multiplication spatiales sont des opérations commutatives et associatives; elles sont tautologiques $a + a = a$, $aa = a$. Elles obéissent vis-à-vis de l'inclusion, à des lois analogues à celles de la théorie générale des ensembles comme par exemple $a \subseteq a + b$; si $a \subseteq b$ et $a' \subseteq b'$, on a aussi $a + a' \subseteq b + b'$. On a, de plus, $a + 0 = a$; $a0 = 0$; $a1 = a$, $acoa = 0$, $a + coa = 1$ etc. Les lois de Morgan $co(a + b) = coacob$, $co(ab) = coa + cob$ sont valables sans restriction, mais la loi de distribution $(a + b)c = ac + bc$ n'est pas valable, de même que la loi $a - b = a - ab$. On en peut donner des exemples très simples dans l'espace à trois dimensions. Pour conserver toutes les lois formelles de la théorie des classes (ensembles), sauf, évidemment, celles qui portent sur la relation d'appartenance $x \in A$, il faut restreindre d'une manière convenable l'ensemble de tous les espaces possibles. C'est précisément ce que nous ferons un peu plus tard, en introduisant la notion de tribu d'espaces.

⁽¹⁾ Voir Stone, *loc. cit.*, où l'on emploie les symboles $\oplus$ et $\ominus$. Dans le livre cité, on trouve un exemple de deux espaces a, b où $a + b$ n'est pas identique à l'ensemble de tous les vecteurs $\vec{x} + \vec{y}$ où $\vec{x} \in a$ et $\vec{y} \in b$.

Les symboles $a + b$, $a - b$, ab , $a \subseteq b$ ressemblent à ceux qui sont admis dans la théorie générale des ensembles. Dans le dernier cas ces opérations seront appelées *opérations logiques*.

Les opérations spatiales $a + b$, $a - b$, coa diffèrent des opérations logiques analogues, tandis que ab , et la relation $a \subseteq b$ ont des significations respectivement identiques.

3. Pour le moment il est préférable d'introduire la notion de compatibilité de deux espaces. Les espaces a et b sont dits *compatibles* lorsque $a - ab \perp b - ab$ ⁽²⁾. Cette notion rappelle celle de perpendicularité dans l'espace euclidien à trois dimensions, mais on peut avoir aussi $a \subset b$, $b \subset a$ ou encore $a = b$. Avant de donner quelques conditions nécessaires et suffisantes de compatibilité de deux espaces, convenons de désigner par $P_a b$ la variété linéaire composée de tous les vecteurs de a qui sont des projections sur a de vecteurs de b ⁽¹⁾.

Les propriétés suivantes sont équivalentes deux à deux ⁽²⁾ :
 1° a, b sont compatibles; 2° $P_a b \subseteq b$ (ici il s'agit d'une inclusion logique d'ensembles de vecteurs); 3° $P_a b \subseteq ab$; 4° $P_a b = ab$, ce qui prouve que dans le cas de compatibilité $P_a b$ est un espace; 4.1° $P_a b = P_b a$; 5° $a - b = a - ab$; 6° $a - ab \subseteq a - b$; 7° $a = ab + a \cap b$; 8° $P_a P_b \vec{x} = P_b P_a \vec{x}$ pour tout $\vec{x} \in 1$; 9° $P_a P_b \vec{x} = P_{ab} \vec{x}$ pour tout $\vec{x} \in 1$; 10° il existe un espace c tel que $P_a P_b \vec{x} = P_c \vec{x}$ pour tout $\vec{x}$.

Pour démontrer cela on peut utiliser les formules $ac + bc \subseteq (a+b)c$; $a = ab + (a - ab)$ et $(b - ab)(a - ab) = 0$, valables pour n'importe

⁽¹⁾ On peut trouver des exemples de deux espaces a, b pour lesquels $P_a b$ n'est pas un espace. Cela résulte des considérations développées par M. G. Julia dans ses Notes des *C. R. de l'Académie des Sciences*, 1944. Le phénomène en question se rattache à ce que M. Julia appelle variétés linéaires fermées asymptotiques l'une à l'autre. Voir aussi les Notes de M. Dixmier insérées dans les *C. R. de l'Académie des Sciences*, 224, 1947. Je me permets de donner un exemple effectif, qui est en solidarité étroite avec ce qui précède, et qui m'a été communiqué par M. J. Ninot Nolla.

Soit $\{\vec{\varphi}_n\}$ un système orthonormal et saturé dans l'espace 1 séparable à une infinité de dimensions, $\{\alpha_n\}$ une suite de nombres $\neq 0$ et tels que $\sum_{n=1}^{\infty} |\alpha_n|^2 < +\infty$, $\{\theta_n\}$ une suite de nombres réels tels que $0 < \theta_n < \frac{\pi}{2}$ et que la série $\sum_{n=1}^{\infty} \sec^2 \theta_n |\alpha_n|^2$ diverge.

Soit a l'espace déterminé par $\{\vec{\varphi}_{2n-1}\}$ et b l'espace déterminé par $\{\cos \theta_n \vec{\varphi}_{2n-1} + \sin \theta_n \vec{\varphi}_{2n}\}$ ($n = 1, 2, \dots$). Posons $\mathcal{E} = P_a b$ et $\nu = a - p$ où p est le plus petit espace contenant

l'ensemble $\mathcal{E}$. Si $\vec{y} \in \nu$, on a $\vec{y} \in a, \vec{y} \perp b$, d'où $\vec{y} = 0$. Il en résulte que $p = a$. On démontre, par l'absurde, que le vecteur $\sum_{n=1}^{\infty} \alpha_n \vec{\varphi}_{2n-1} \in a$ n'appartient pas à $\mathcal{E}$. Par conséquent,

$P_a b$ n'est pas un espace. Cette construction a été suggérée par l'exemple cité de M. Stone concernant une somme spatiale siglière.

⁽²⁾ G. JULIA, *Sur les projecteurs de l'espace hilbertien ou unitaire* (*C. R. Acad. Sc.*, 9 mars 1942, p. 456).

quels espaces et utiliser le théorème sur la décomposition d'un vecteur : $\vec{x} = \vec{x}_a + \vec{x}_{\text{co } a}$ ainsi que les propriétés élémentaires de la projection, par exemple celle-ci : si $a \perp b$, on a $P_{a+b} \vec{x} = P_a \vec{x} + P_b \vec{x}$.

6. Nous aurons besoin d'opérations spatiales infinies. Le plus petit espace contenant tous les éléments d'une suite infinie $\{a_n\}$ d'espaces s'appelle leur *somme*, $\sum_{n=1}^{\infty} a_n$ et, le plus grand espace contenu dans chaque a_n s'appelle le *produit* $\prod_{n=1}^{\infty} a_n$. Nous admettons des définitions analogues pour la *somme* $\sum_{\alpha} a_{\alpha}$ et le *produit* $\prod_{\alpha} a_{\alpha}$ où α varie dans un ensemble donné quelconque non vide d'indices. La somme et le produit existent toujours. On a

$$a_{\beta} \subseteq \sum_{\alpha} a_{\alpha}, \quad \prod_{\alpha} a_{\alpha} \subseteq a_{\beta},$$

et pour les opérations dénombrables les lois de Morgan subsistent :

$$\text{co} \sum_{n=1}^{\infty} a_n = \prod_{n=1}^{\infty} \text{co } a_n, \quad \text{co} \prod_{n=1}^{\infty} a_n = \sum_{n=1}^{\infty} \text{co } a_n.$$

7. Par *tribu* ⁽¹⁾ d'espaces nous comprendrons tout ensemble (T) non vide d'espaces tel que : 1° si $a, b \in (T)$, on ait $a + b \in (T)$; 2° si $a \in (T)$, on ait $\text{co } a \in (T)$; 3° si $a, b \in (T)$ et si $ab = 0$, on ait $a \perp b$.

(1) L'expression « tribu », empruntée à M. RENÉ DE POSSEL (*Journ. d. Math.*, t. 101, 1936) semble être plus convenable que le terme « corps », étant donné que les méthodes de l'algèbre abstraite où celui-ci a une signification différente jouent un rôle de plus en plus grand dans l'analyse moderne.

Une tribu d'espaces est un cas particulier d'une *tribu abstraite de Boole* (connu sous le nom « corps de Boole »). D'une manière assez vague, une tribu de Boole peut être définie comme une sorte d'algèbre où les opérations $a + b$, $a.b$, $a - b$, $\text{co } a$ et la relation $a \subseteq b$ obéissent à toutes les lois formelles de la théorie générale des ensembles. Pour une axiomatisation précise voir A. TARSKI, *Zur Grundlegung d. Boole'schen Algebra I.* (*Fund. Math.*, t. 24, 1935, p. 160 et suiv.). Voici une définition directe d'une tribu de Boole. Soit (B) un anneau abstrait possédant une unité 1 et dont chaque élément a est idempotent : $a.a = a$. Si dans (B) (que nous voudrions appeler *anneau de Stone*), on remplace son addition « + » par l'opération suivante $a + b \stackrel{\text{df}}{=} a \dot{+} a.b \dot{+} b$ et qu'on définit $\text{co } a \stackrel{\text{df}}{=} 1 \dot{+} a$, $a - b \stackrel{\text{df}}{=} a.\text{co } b$, on obtient la tribu de Boole la plus générale. Voir M. H. STONE, *The theory of representation for*

Les espaces d'une tribu sont compatibles entre eux et l'on peut leur appliquer toutes les lois formelles finies de la théorie abstraite des ensembles. On peut démontrer le théorème suivant :

Si (T) est un ensemble non vide d'espaces satisfaisant aux conditions suivantes : 1° si $a, b \in (T)$, on a $a + b \in (T)$; 2° si $a \in (T)$, on a $coa \in (T)$, alors la condition nécessaire et suffisante pour que (T) soit une tribu est que les éléments de (T) soient compatibles entre eux.

8. Une tribu (T) s'appelle *dénombrablement* (resp. *complètement*) *additive* lorsque toute somme et tout produit dénombrables (resp. arbitraires) de ses éléments appartient à (T) .

THÉORÈME. — *Pour chaque tribu (T) il existe une tribu dénombrablement additive (T') qui contient (T) .*

THÉORÈME. — *Chaque tribu dénombrablement additive de sous-espaces d'un espace séparable est aussi complètement additive; chaque somme est égale à une somme au plus dénombrable de ses éléments et il en est de même pour un produit quelconque ⁽¹⁾.*

II.

PROLONGEMENT D'UNE TRIBU.

1. Si (T) est une tribu et p un espace il n'est pas vrai, en général, qu'il existe une tribu (T') comprenant (T) et p simultanément. *Pour que p soit capable d'être adjoint à (T) il faut et il suffit que p soit*

Boolean algebras (Trans. of the Amer. Math. Soc., t. 40, 1936). On a réciproquement

$$a + b = (a - b) + (b - a).$$

L'introduction des tribus d'espaces nous a été suggérée par le livre de J. V. Neumann (*Mathematische Grundlagen d. Quantenmechanik*. Berlin 1932). O. Nikodym, On Boolean fields etc. *loc. cit.* Les mathématiciens japonais par ex. HIDEGARÔ NAKANO (*Math. Ann.*, t. 118, 1941-3), ont aussi employé de telles tribus. Voir aussi F. WECKEN, *Math. Ann.*, t. 116, 1939, et *Math. Zeitschr.*, t. 45, 1939, où cet auteur utilise des tribus.

(¹) O. NIKODYM, *loc. cit.* En ce qui concerne le deuxième théorème, le cas général d'une tribu de Boole est établi par F. WECKEN. *Math. Zeitschr.*, t. 45, 1939.

compatible avec tout espace de la tribu. Dans ce cas nous disons que p est compatible avec la tribu (T) .

La condition étant évidemment nécessaire, il suffit d'en démontrer la suffisance. Voici l'esquisse d'une démonstration. On peut établir, d'une manière générale que, lorsque $a_1, a_2, \dots, a_n$ ($n \geq 2$) sont orthogonaux et que b est compatible avec tous les espaces, on a

$$(a_1 + \dots + a_n)b = a_1b + \dots + a_nb.$$

Ce lemme, dans le cas où $n = 2$ ou 3 permet de démontrer que, lorsque p est compatible avec (T) , les relations $a, b \in (T)$, $abp = 0$ entraînent $ap \perp bp$. Cela étant posé, on obtient le fait que l'ensemble de tous les espaces ap où $a \in (T)$, forme une tribu. En utilisant la formule générale

$$P_a b + P_{co a} b = b$$

valable pour tout couple a, b compatible d'espaces, on trouve que l'ensemble de tous les espaces $a \text{ cop}$ où $a \in (T)$ est aussi une tribu. Cela montre que la classe des espaces

$$ap + b \text{ cop},$$

où $a, b \in (T)$, est une tribu. Celle-ci est la plus petite tribu comprenant (T) et p .

Si (T) est dénombrablement additive et p compatible avec (T) , l'ensemble des espaces $ap + b \text{ cop}$ où $a, b \in (T)$ est, non seulement la plus petite tribu simplement additive comprenant (T) et p , mais aussi la plus petite tribu dénombrablement additive jouissant de cette propriété.

2. Une tribu (T) s'appelle *saturée* lorsque chaque tribu (T') contenant (T) coïncide avec elle. En utilisant le résultat précédent on peut démontrer par l'induction transfinie que

quelle que soit la tribu (T) il existe toujours une tribu saturée (T') contenant (T) .

Chaque tribu (T) saturée est dénombrablement additive, ce qui résulte du fait que toute tribu est contenue dans une tribu dénombrablement additive.

Si l'espace $\mathfrak{I}$ est séparable, chaque tribu saturée est complètement additive. C'est une conséquence de ce que dans ces conditions une tribu dénombrablement additive est aussi complètement additive.

3. **Champ de rayonnement.** — Soit (T) une tribu et $\mathcal{E} \neq 0$ un ensemble non vide de vecteurs. Appelons *champ de rayonnement de $\mathcal{E}$ sur (T)* l'ensemble $M_T(\mathcal{E})$ de tous les vecteurs limites de suites convergentes de combinaisons linéaires finies

$$\sum_i \lambda_i P_{a_i} \vec{\xi}_i,$$

où $a_i \in (T)$, $\vec{\xi}_i \in \mathcal{E}$, et où les λ_i sont des nombres complexes. On démontre sans peine que $M_T(\mathcal{E})$ est identique au plus petit espace contenant tous les vecteurs $P_a \vec{\xi}$ où a varie dans (T) et $\vec{\xi}$ dans $\mathcal{E}$ ⁽¹⁾.

On voit que $\mathcal{E} \subseteq M_T(\mathcal{E})$. Si $\mathcal{E} \subseteq \mathcal{F}$, on a $M_T(\mathcal{E}) \subseteq M_T(\mathcal{F})$. Si $\mathcal{F}$ est partout dense dans $\mathcal{E}$, on a $M_T(\mathcal{F}) = M_T(\mathcal{E})$. Si b est l'espace déterminé par $\mathcal{E}$, on a $M_T(\mathcal{E}) = M_T(b)$. Si $a \in (T)$, on a $M_T(a) = a$. En outre l'égalité suivante a lieu : $M_T(M_T(\mathcal{E})) = M_T(\mathcal{E})$. On peut aussi démontrer que

$$M_T(\mathcal{E}) = \sum_{\vec{\xi} \in \mathcal{E}} M_T\left(\left(\begin{smallmatrix} \vec{\xi} \\ \vec{\xi} \end{smallmatrix}\right)\right) \quad (2).$$

THÉORÈME. — *La condition nécessaire et suffisante pour qu'un espace p soit compatible avec tous les espaces d'une tribu (T) est qu'il existe un ensemble $\mathcal{E}$ de vecteurs tel qu'on ait*

$$p = M_T(\mathcal{E}).$$

Démonstration. — Si l'on pose $p \stackrel{\text{df}}{=} M_T(\mathcal{E})$ on trouve que $P_a p \subseteq p$ pour tout $a \in (T)$, ce qui exprime que p est compatible avec (T) . Réciproquement, si p est compatible avec (T) , on trouve que $M_T(p) = p$, donc pour $\mathcal{E} \stackrel{\text{df}}{=} p$ on a $p = M_T(\mathcal{E})$.

4. **Vecteur générateur.** — En général $M_T(\mathcal{E})$ n'est pas identique au champ de rayonnement d'un ensemble composé d'un seul vecteur. Il

⁽¹⁾ Le cas particulier où $\mathcal{E}$ se réduit à un seul vecteur est connu et a été considéré dans la théorie de la multiplicité du spectre continu et dans la théorie des invariants unitaires des transformations selfadjointes. Voir p. ex. le livre de M. STONE, *Linear transformations*, loc. cit.

⁽²⁾ $\left(\begin{smallmatrix} \vec{\xi} \\ \vec{\xi} \end{smallmatrix}\right)$ désigne l'ensemble composé du seul vecteur $\vec{\xi}$. La somme est spatiale.

n'est pas vrai non plus que pour deux vecteurs $\vec{x}_1, \vec{x}_2$ il existe toujours un troisième vecteur $\vec{x}_3$ tel que

$$M_T((\vec{x}_1) + (\vec{x}_2)) = M_T((\vec{x}_3)).$$

Mais, si la tribu (T) est saturée, cela est vrai.

Soient, en effet, $\vec{x}_1, \vec{x}_2$ deux vecteurs quelconques, et supposons que (T) soit saturée. Posons

$$p_1 \stackrel{\text{df}}{=} M_T((\vec{x}_1)), \quad p_2 \stackrel{\text{df}}{=} M_T((\vec{x}_2)).$$

p_1 ainsi que p_2 est compatible avec (T); donc, comme (T) est saturée, $p_1 \in (T)$ et $p_2 \in (T)$. Posons

$$\vec{v} \stackrel{\text{df}}{=} P_{p_1 - p_2} \vec{x}_2.$$

On démontre sans peine que le vecteur $\vec{x}_3 \stackrel{\text{df}}{=} \vec{x}_1 + \vec{v}$ possède la propriété désirée

$$M_T((\vec{x}_3)) = M_T((\vec{x}_1) + (\vec{x}_2)).$$

Remarquons que si (T) n'est pas saturée, le théorème est en défaut. Soit, en effet, (T) la tribu (dans l'espace 1 à trois dimensions), composée des espaces : 0, le plan xy , l'axe z et 1 et, soient $\vec{x}_1, \vec{x}_2$ deux vecteurs $\neq \vec{0}$ aux coordonnées différentes de zéro, mais non situés dans un même plan passant par l'axe z . On a $M_T((\vec{x}_1) + (\vec{x}_2)) = 1$, mais, quel que soit le vecteur $\vec{x}_3 \neq \vec{0}$, il y a seulement trois possibilités : $M_T((\vec{x}_3))$ est un plan passant par l'axe z ; $M_T((\vec{x}_3)) = (x, y)$; $M_T((\vec{x}_3))$ coïncide avec l'axe z .

Un vecteur $\vec{\omega}$ tel que $M_T((\vec{\omega})) = 1$ s'appelle *vecteur générateur* de l'espace 1 par rapport à la tribu (T).

Les deux propriétés suivantes sont équivalentes :

1° *Il existe un vecteur générateur de 1 par rapport à (T);*

2° Pour tout ensemble $\mathcal{E} \neq 0$ de vecteurs il existe un vecteur $\vec{\xi}$ tel que

$$M_T(\mathcal{E}) = M_T\left(\left(\vec{\xi}\right)\right).$$

Si $\vec{\omega}$ est un vecteur générateur, on a pour tout $b \in (T)$:

$$b = M_T\left(\left(P_b \vec{\omega}\right)\right).$$

En ce qui concerne l'existence d'un vecteur générateur, nous ne savons la démontrer que dans le cas où $\mathfrak{r}$ est séparable et (T) saturée :

THÉORÈME. — Si l'espace $\mathfrak{r}$ est séparable et la tribu (T) saturée, il existe au moins un vecteur générateur de $\mathfrak{r}$ par rapport à (T) .

Démonstration. — Soit $\mathcal{E} \stackrel{\text{df}}{=} (\vec{z}_1) + (\vec{z}_2) + \dots$ (somme logique), un ensemble dénombrable de vecteurs partout dense dans $\mathfrak{r}$. Posons

$$p_n \stackrel{\text{df}}{=} M_T\left[\sum_{i=1}^n (\vec{z}_i)\right] \quad (n = 1, 2, \dots).$$

On a $p_1 \subseteq p_2 \subseteq \dots \subseteq p_n \subseteq \dots$. Nous allons définir une suite au plus dénombrable $\{\vec{\xi}_n\}$ de vecteurs. Posons $\vec{\xi}_1 \stackrel{\text{df}}{=} \vec{z}_1$ et supposons qu'on ait déjà défini les vecteurs $\vec{\xi}_i$ pour $i < n$. Supposons, en outre, que $p_i = M_T\left((\vec{\xi}_{i-1}) + (\vec{z}_i)\right)$ pour $i = 2, 3, \dots, n-1$. Choisissons $\vec{\xi}_n$ de manière qu'on ait

$$M_T\left((\vec{\xi}_{n-1}) + (\vec{z}_n)\right) = M_T\left((\vec{\xi}_n)\right),$$

ce qui est possible en vertu d'un théorème précédent. On démontre alors que

$$p_n = M_T\left((\vec{\xi}_{n-1}) + (\vec{z}_n)\right).$$

Cela donne l'égalité

$$p_n = M_T\left((\vec{z}_n)\right) \quad \text{pour } n = 1, 2, \dots$$

L'ensemble $\mathcal{E}$ étant partout dense dans $\mathfrak{r}$, on a $\sum_{n=1}^{\infty} p_n = \mathfrak{r}$. Si les espaces p_n sont égaux à partir d'un certain indice, le théorème en résulte immédiatement. Sinon, choisissons une suite partielle $\{p_{v_n}\}$ telle que

$$p_{v_1} \subset p_{v_2} \subset \dots \subset p_{v_n} \subset \dots,$$

et posons

$$q_n \stackrel{\text{df}}{=} p_{\gamma_n}, \quad x_n \stackrel{\text{df}}{=} \xi_{\gamma_n}.$$

Posons, en outre :

$$\gamma_1 \stackrel{\text{df}}{=} x_1, \quad \gamma_n \stackrel{\text{df}}{=} p_{q_n - q_{n-1}}, x_n \quad (n = 2, 3, \dots).$$

On démontre que

$$q_1 = M_T((\gamma_1)), \quad q_n - q_{n-1} = M_T((\gamma_n)) \quad \text{pour } n \geq 2.$$

Les vecteurs γ_n sont différents de $\vec{o}$ et orthogonaux deux à deux. Fixons une suite $\{\sigma_n\}$ de nombres complexes où $\sigma_n \neq 0$ et où la série $\sum_{n=1}^{\infty} |\sigma_n|^2$ converge. On démontre que le vecteur

$$\omega \stackrel{\text{df}}{=} \sum_{n=1}^{\infty} \sigma_n \frac{\gamma_n}{|\gamma_n|}$$

est un vecteur générateur.

Un peu plus tard nous démontrerons que la réciproque de ce théorème est vraie aussi (dans le cas où $\mathfrak{I}$ est séparable).

III.

MESURE ET TOPOLOGIE.

1. La tribu (T) d'espaces est un cas particulier d'une tribu de Boole. On peut donc introduire une notion de mesure ⁽¹⁾ sur (T). En voici la définition :

(1) Nous avons introduit la notion de mesure sur une tribu de Boole dans le travail : O. NIKODYM, *Sur la mesure vectorielle parfaitement additive dans un corps abstrait de Boole* [Mém. de l'Acad. Roy. de Belgique (Cl. d. Sc., t. 17, 1938)].

En même temps : C. CARATHÉODORY, *Entwurf für eine Algebraisierung d. Integralbegriffs*. Münch. Sitzber, 1938. La notion de mesure sur une tribu générale s'est introduite d'elle-même aussi chez d'autres mathématiciens, p. ex. F. WECKEN, *Abstrakte Integrale u. fastperiodische Funktionen* (Math. Zeitschr., t. 45, 1939). C'est une généralisation naturelle de la mesure sur une tribu d'ensembles dont l'étude a été abordée par M. M. FRÉCHET, *Des familles et fonctions additives d'ensembles abstraits*. (Fund. Math., t. 4, 1923, Fund. Math., t. 5, 1924), et M. de la Vallée-Poussin. La mesure d'un ensemble abstrait joue un rôle fondamental dans le Calcul des Probabilités.

Par une *mesure* (simplement additive) sur une tribu (T) , nous comprendrons chaque fonction $\mu(a)$ qui fait correspondre un nombre fini, non négatif à tout $a \in (T)$, ce nombre satisfaisant à la condition suivante :

Si $a, b \in (T)$ et si $ab = 0$, on a $\mu(a+b) = \mu(a) + \mu(b)$. La mesure μ sera dite *dénombrablement additive* lorsque pour chaque suite infinie $\{a_n\}$ d'espaces disjoints de (T) on a

$$\mu\left(\sum_{n=1}^{\infty} a_n\right) = \sum_{n=1}^{\infty} \mu(a_n).$$

Mis à part le cas de la mesure triviale identique à zéro, on vérifie aisément que si $\xi \in \mathcal{I}$ et qu'on pose

$$\mu(a) = \left| P_{a\xi} \right|_{\text{df}}^2 \quad \text{pour } a \in (T),$$

on obtient une mesure dénombrablement additive. Cela se démontre au moyen du théorème généralisé de Pythagore.

Remarquons qu'il y a des cas où une mesure simplement additive sur (T) ne peut pas être prolongée de manière qu'elle devienne dénombrablement additive sur un prolongement (T') dénombrablement additif de (T) . En voici un exemple construit par un procédé connu : Soit $\{\vec{\varphi}_i\}$ un système orthonormal et saturé de vecteurs dans $\mathcal{I}$, espace séparable à un nombre infini de dimensions, et soit $\{\omega_i\}$ une suite composée de tous les nombres rationnels distincts de l'intervalle $(0, 1)$ semi-fermé $0 < x \leq 1$. A chaque somme logique finie $s = \sum_k (\alpha_k, \beta_k)$ de sous intervalles semi-fermés de $(0, 1)$, attachons l'espace a_s déterminé par les vecteurs $\vec{\varphi}_v$ pour lesquels $\omega_v \in s$. La classe composée de tous les a_s et de l'espace 0 est une tribu. Définissons $\mu(a_s)$ comme la mesure lebesguienne de l'ensemble s et posons $\mu(0) = 0$. Cette mesure ne peut pas être prolongée de la manière indiquée.

2. Dans ce qui va suivre nous nous bornons aux tribus dénombrablement additives dans un espace $\mathcal{I}$ séparable. La mesure μ sur une tribu (T) s'appelle *effective* lorsque la relation $\mu(a) = 0$ entraîne $a = 0$ ⁽¹⁾. Montrons qu'il existe toujours une telle mesure.

(1) La mesure effective sur une tribu abstraite de Boole a été étudiée par F. WECKEN. (*Math. Zeitschr.*, t. 45, 1939) et par M^{me} DOROTHY MAHARAM, *On homogeneous measure algebra...* (*Proceedings of the Nat. Acad. of Sc.*, vol. 28, 1942).

En effet soit $\{\vec{\xi}_n\}$, ($n = 1, 2, \dots$) une suite infinie de vecteurs non nuls formant un ensemble partout dense dans $\mathfrak{I}$. Posons

$$\mu(a) \stackrel{\text{df}}{=} \sum_{n=1}^{\infty} \frac{1}{n!} \cdot \frac{|P_a \vec{\xi}_n|^2}{|\vec{\xi}_n|^2} \quad \text{pour } a \in (T).$$

On démontre aisément que $\mu(a)$ est une mesure effective sur (T) et cela quelle que soit (T) .

THÉORÈME. — Si $\vec{\omega}$ est un vecteur générateur de $\mathfrak{I}$ par rapport à (T) , la mesure définie par l'équation

$$\mu(a) \stackrel{\text{df}}{=} |P_a \vec{\omega}|^2 \quad \text{pour } a \in (T)$$

est effective sur (T) .

La démonstration ne présente pas de difficultés.

3. Une mesure effective μ engendre dans (T) une topologie induite par la définition suivante d'écart de deux espaces $(^1)$:

$$\|a, b\|_{\mu} \stackrel{\text{df}}{=} \mu(a - b) + \mu(b - a) \quad (a, b \in (T)).$$

Cet écart jouit des propriétés usuelles :

1° $\|a, b\|_{\mu} = \|b, a\|_{\mu}$; 2° $\|a, b\|_{\mu} \leq \|a, c\|_{\mu} + \|c, b\|_{\mu}$; 3° $\|a, b\|_{\mu} \geq 0$; 4° l'équation $\|a, b\|_{\mu} = 0$ équivaut à $a = b$, et, par conséquent, elle organise la tribu en espace métrique. La topologie qui lui correspond est déterminée par la notion suivante de la limite a d'une suite infinie $\{a_n\}$ d'espaces. On dit que a_n converge vers a , ($a_n \rightarrow a$), lorsque $\|a_n, a\|_{\mu}$ tend vers zéro.

(¹) Voir O. NIKODYM, *Sur une généralisation des intégrales de M. J. Radon* (*Fund. Math.*, t. 14, 1929), où cette notion est étudiée dans le cas d'une mesure dénombrablement additive sur une tribu dénombrablement additive d'ensembles abstraits. Cette notion d'écart a été trouvée d'une manière indépendante par M. N. ARONSZAJN (dans un manuscrit inédit de 1929-30) où il considère le cas d'une mesure simplement additive et, d'une manière plus générale, une mesure « convexe » : $\mu(\mathcal{S} + \mathcal{T}) \leq \mu(\mathcal{S}) + \mu(\mathcal{T})$.

La notion considérée d'écart de deux ensembles n'est qu'un cas particulier de la notion plus générale d'écart de deux fonctions introduites par M. M. FRÉCHET, *Sur les divers modes de convergence d'une suite infinie de fonctions d'une variable* (*Bull. of the Calcutta Math. Soc.*, 1921).

On peut démontrer que la métrique définie ci-dessus est complète ⁽¹⁾.

Si μ et μ' sont deux mesures effectives sur (T) , les métriques qui y correspondent sont équivalentes, c'est-à-dire qu'elles engendrent la même topologie ⁽²⁾. La séparabilité de cette topologie semble être une propriété un peu plus cachée :

THÉORÈME. — La topologie induite dans une tribu (T) par une mesure effective μ est séparable (si $\mathfrak{I}$ est séparable).

Démonstration. — Soit μ une mesure effective sur (T) .

Nous allons remplacer la métrique μ par une autre équivalente mais plus maniable. Soit $\{\vec{\xi}_k\}$, $(k = 1, 2, \dots)$ un ensemble de vecteurs $\neq \vec{0}$, partout dense dans $\mathfrak{I}$. Fixons arbitrairement les nombres positifs λ_k tels que $\sum_{k=1}^{\infty} \lambda_k$ converge. Posons

$$\vec{x}_k = \frac{\vec{\xi}_k}{|\vec{\xi}_k|} \quad \text{et} \quad \nu(\alpha) = \sum_{k=1}^{\infty} \lambda_k |P_{\alpha} \vec{x}_k|^2 \quad \text{pour } \alpha \in (T).$$

C'est une mesure effective, donc la métrique engendrée par ν est équivalente à la métrique μ . Nous allons maintenant prolonger cette métrique qui n'est définie que sur (T) , par une autre, plus générale qui organise la classe (H) de tous les sous-espaces de $\mathfrak{I}$ en topologie. Dans ce but, introduisons dans (H) la notion auxiliaire suivante d'écart de deux espaces arbitraires

$$\|p, q\| = \sum_k \lambda_k |P_p \vec{x}_k - P_q \vec{x}_k|^2.$$

On voit que $\|p, q\| \geq 0$, $\|p, q\| = \|q, p\|$ et l'on démontre que $\|p, q\| \leq \|p, r\| + \|r, q\|$ et, que les équations $\|p, q\| = 0$ et $p = q$ sont équivalentes. La restriction à (T) de la métrique ci-dessus est équivalente à la métrique ν . Cela se démontre à l'aide de l'identité suivante :

$$|P_{a\gamma} \vec{x} - P_{b\gamma} \vec{x}|^2 = |P_{a-b\gamma} \vec{x}|^2 + |P_{b-a\gamma} \vec{x}|^2.$$

⁽¹⁾ Pour démontrer cela, il suffit d'appliquer le raisonnement que nous avons employé [dans le travail (1)] dans le cas d'une tribu d'ensembles abstraits. Un théorème un peu plus général se trouve dans le manuscrit cité de M. Aronszajn. Voir aussi F. WECKEN, *Math. Zeitschr.*, t. 45, 1939.

⁽²⁾ Démonstration comme chez F. Wecken.

valable pour tout $a, b \in (T)$ et pour tout $\vec{\gamma} \in \mathbf{1}$. En effet celle-ci permet de montrer que, si $a_n \in (T)$ et si $\nu(a_n - a) + \nu(a - a_n) \rightarrow 0$, on a $\|a_n, a\| \rightarrow 0$ et réciproquement, si $a_n, a \in (T)$ et si $\|a_n, a\| \rightarrow 0$, on a $\nu(a_n - a) + \nu(a - a_n) \rightarrow 0$. Cela étant donné, il suffit de montrer que la topologie définie sur (H) par l'intermédiaire de l'écart $\| \ \|$ est séparable. Cela se fait en deux étapes. On montre d'abord que si $p_1 \subseteq p_2 \subseteq \dots \subseteq p_n \subseteq \dots$, $\sum_{n=1}^{\infty} p_n = p$, on a $\|p_n, p\| \rightarrow 0$. Or chaque espace $p \in (H)$ peut être considéré comme somme spatiale infinie d'une suite non décroissante d'espaces dont chacun a un nombre fini de dimensions. Il en résulte que l'ensemble (H_1) de tous les espaces dont chacun a un nombre fini de dimensions est $\| \ \|$ — partout dense dans (H) . La seconde étape consiste dans le choix d'un ensemble dénombrable (H_2) partout dense dans (H_1) . Dans ce but reprenons les vecteurs $\vec{\xi}_n$ ($n = 1, 2, \dots$) et définissons (H_2) comme la classe de tous les espaces dont chacun est déterminé par un nombre fini de ces vecteurs. (H_2) est dénombrable et représente un sous-ensemble de (H_1) . Le fait que (H_2) est partout dense dans (H_1) se démontre par un procédé d'orthogonalisation accompagné d'un procédé d'approximation bien visible. La $\| \ \|$ — topologie dans (H) étant ainsi séparable, il en résulte qu'il en est de même de sa restriction sur (T) . Le théorème est ainsi démontré.

IV.

CLASSES ORDONNÉES.

1. Par une *classe ordonnée* nous allons entendre toute classe (L) d'espaces contenant $0, 1$ et jouissant de la propriété suivante : si $a, b \in (L)$, on a ou bien $a \subseteq b$ ou bien $b \subseteq a$. Chaque espace $a - b$ où $b \subset a$ et où $a, b \in (L)$ s'appelle *segment de la classe ordonnée* (L) . Un segment n'est jamais $= 0$, et, si $a - b = a' - b'$ est un segment, on a $a = a'$ et $b = b'$. L'espace a s'appelle *extrémité droite du segment* $a - b$ et b son *extrémité gauche*. Il existe une plus petite tribu (S) contenant tous les éléments de (L) . La tribu (S) c'est l'ensemble composé de 0 et de toutes les sommes spatiales finies de segments.

La classe ordonnée (L) s'appelle *dénombrablement fermée* lorsqu'elle contient toute somme spatiale $\sum_{n=1}^{\infty} a_n$ et tout produit $\prod_{n=1}^{\infty} a_n$ d'éléments de (L) . La classe ordonnée (L) s'appelle *complètement fermée* lorsque toute somme $\sum_{\alpha} a_{\alpha}$ et tout produit $\prod_{\alpha} a_{\alpha}$ de ses éléments lui appartiennent. Dans le cas considéré où $\mathbf{1}$ est séparable, ces deux notions coïncident. La plus petite classe ordonnée fermée $(\bar{L})$ contenant une classe ordonnée s'appelle la *fermeture* de (L) . Elle existe toujours et est unique. Il existe une plus petite tribu (T_L) , unique, dénombrablement additive contenant (L) . Cette tribu s'appelle *l'extension borélienne* de (L) . La tribu (T_L) se compose de tous les espaces $\prod_{i=1}^{\infty} \sum_{k=1}^{\infty} p_{ik}$ où les p_{ik} sont des segments de (L) . L'expression ci-dessus peut être remplacée par $\sum_{i=1}^{\infty} \prod_{k=1}^{\infty} q_{ik}$ où les q_{ik} sont des segments de (L) . On voit que $(T_L) = (T_{\bar{L}})$.

THÉORÈME. — Si la tribu (T) est dénombrablement additive (et $\mathbf{1}$ séparable) il existe une classe ordonnée (L) telle que son extension borélienne coïncide avec (T) .

Démonstration. — Considérons la topologie engendrée dans (T) par une mesure effective. Cette topologie est séparable. Soit $\alpha_1, \alpha_2, \dots, \alpha_n, \dots$ un ensemble d'espaces partout dense dans (T) par rapport à cette topologie. Définissons les classes ordonnées suivantes de plus en plus larges. Soit (L_1) la classe

$$0 \leq \alpha_1 \leq 1.$$

Supposons qu'on ait déjà défini les classes $(L_1), \dots, (L_n)$ ($n \geq 1$) et, supposons que : 1° $(L_1) \subseteq (L_2) \subseteq \dots \subseteq (L_n)$; 2° que la plus petite tribu (T_i) engendrée par (L_i) où $i \leq n$, contienne les espaces $\alpha_1, \dots, \alpha_n$.

Soit $0 \leq \beta_1 \leq \beta_2 \leq \dots \leq 1$ la classe (L_n) .

Définissons alors (L_{n+1}) comme

$$0 \leq \beta_1 \alpha_{n+1} \leq \beta_1 \leq \beta_1 + (\beta_2 - \beta_1) \alpha_{n+1} \leq \beta_2 \leq \dots \leq 1.$$

On démontre aisément que (L_n) est contenue dans (L_{n+1}) et que la plus petite tribu (T_{n+1}) contenant (L_{n+1}) contient les espaces $\alpha_1, \dots, \alpha_n$,

α_{n+1} . La somme logique de tous les (L_n) ($n = 1, 2, \dots$) est une classe linéaire cherchée. Remarquons que la démonstration ci-dessus s'appuie sur la séparabilité de la topologie (T) donc sur la séparabilité de l'espace 1 ⁽¹⁾.

V.

LIEUX.

1. Pour aller plus loin faisons quelques remarques d'un caractère intuitif. L'existence d'une classe linéaire engendrant la tribu donnée rapproche les tribus d'espaces des tribus de sous-ensembles du segment ordinaire. On peut mettre en évidence cette analogie, en faisant correspondre aux espaces $a \neq 0$ d'une classe ordonnée (L) les segments semi-ouverts $(0, \mu(a))$ ⁽²⁾ où μ est une mesure effective. Cette correspondance se prolonge aux segments $a - b$ de (L), il leur correspond les intervalles $(\mu(b), \mu(a))$. Même la plus petite tribu simplement additive (S) contenant (L) se traduit par la tribu (σ) composée de l'ensemble vide de points de $(0, \mu(1))$ et des sommes logiques finies de sous-intervalles semi-ouverts

$$\sum_i (\mu(b_i), \mu(a_i)).$$

Les tribus (S) et (σ) sont isomorphes, aux opérations spatiales correspondent les opérations logiques analogues. Les deux tribus peuvent être prolongées pour devenir des tribus énumérablement additives, mais après ce prolongement l'analogie cesse de se conserver dans le cas général où (L) se compose d'une infinité d'espaces. En effet, la tribu prolongée d'ensembles se compose, entre autres, de points qui sont des *atomes* de la tribu ⁽³⁾ tandis que pour les espaces cela n'existe plus, Il serait bien commode d'obtenir une analogie plus complète entre

⁽¹⁾ Il serait intéressant à l'examiner si le théorème est vrai pour toute tribu abstraite de Boole.

⁽²⁾ (α, β) désigne l'intervalle $\alpha < x \leq \beta$.

⁽³⁾ Par un *atome* d'une tribu de Boole on entend chaque élément p tel que : 1° $p \neq 0$; 2° si $q \subseteq p$, on a ou bien $q = 0$ ou bien $q = p$. Voir par exemple A. TARSKI, *loc. cit.* ou bien STONE, *Transactions*, loc. cit.

Pour la tribu composée de tous les sous-ensembles de $\langle 0, 1 \rangle$ les atomes sont les ensembles composés d'un seul point.

les tribus d'espaces et celles d'ensembles, on gagnerait beaucoup si l'on pouvait construire, même d'une manière artificielle, de tels atomes, disons « points ». Pour faire cela nous pouvons imiter la méthode suivante qui permet de construire la théorie des nombres irrationnels et, plus généralement, les nombres réels en partant des nombres rationnels. Ici on peut considérer des suites infinies décroissantes d'intervalles ouverts emboîtés l'un dans l'autre, leur longueur tendant vers 0 et définir les nombres réels par de telles suites. Les nombres réels seront considérés ainsi comme des intervalles « infiniment petits ». La même chose peut être faite avec les segments d'une classe ordonnée d'espaces, on peut introduire les points-atomes d'espaces comme des segments-espaces « infiniment petits ». Nous les appellerons *lieux*. C'est ce que nous voudrions faire d'une manière rigoureuse. Mais il s'avère qu'il est préférable, au point de vue pratique, de suivre une méthode voisine de celle de Dedekind plutôt que la voie indiquée ci-dessus. Le nombre réel se définit avec cette méthode comme une coupure, c'est-à-dire comme un couple de classes de nombres rationnels. Mais on peut modifier cette définition en ne considérant que des classes gauches par exemple. Quel que soit le mode de construction on obtient des êtres équivalents au point de vue des opérations arithmétiques.

Procédons maintenant à des considérations rigoureuses.

2. Soit (L) une classe ordonnée fermée et soit $a \in (L)$, où $a \neq 0$. Appelons *lieu A en a* l'ensemble de tous les segments de (L) admettant a pour extrémité droite, A désigne donc l'ensemble de tous les espaces $a - x$ où $x \in (L)$ et où $x \subset a$. Les espaces $a - x$ s'appelleront voisinages du lieu A . Chaque segment de (L) peut être considéré comme voisinage d'un lieu bien déterminé.

Nous allons considérer les ensembles de lieux et en déduire toute une théorie de la mesurabilité. Dans ce but introduisons quelques notions auxiliaires. Désignons par I l'ensemble de tous les lieux. Appelons *couverture* l'espace 0 et toute somme spatiale de segments. Une telle somme, même non dénombrable est toujours égale à une somme au plus dénombrable de segments. Convenons de dire qu'un lieu A est *agrégé à une couverture α* lorsqu'il existe un voisinage p de A tel que $p \subseteq \alpha$. Par *couverture d'un ensemble $\mathcal{E}$ de lieux* nous comprendrons toute couverture α pour laquelle si $A \in \mathcal{E}$, A est agrégé à α .

3. Soit $\mathcal{E}$ un ensemble de lieux et considérons toutes ses couvertures β , l'espace $\prod_{\beta} \beta$ s'appelle le *support extérieur* de $\mathcal{E}$, nous le désignerons par e^* . Considérons en même temps l'ensemble $\mathcal{F} = I - \mathcal{E}$ de lieux complémentaire de $\mathcal{E}$ par rapport à la totalité I des lieux. Si f^* est le support extérieur de $\mathcal{F}$, appelons l'espace $I - f^*$ le *support intérieur* de $\mathcal{E}$ et désignons-le par e_* . On a toujours $e_* \subseteq e^*$.

Cela se démontre en établissant que, lorsque α et β sont des couvertures respectives de $\mathcal{E}$ et $\mathcal{F} = I - \mathcal{E}$, on a $\alpha + \beta = I$. L'ensemble $\mathcal{E}$ de lieux sera dit *mesurable* lorsque ses deux supports sont égaux $e_* = e^*$. Dans ce cas l'espace $e = e_* = e^*$ s'appelle le *support* de $\mathcal{E}$.

L'ensemble vide de lieux est mesurable et son support est l'espace 0. L'ensemble des lieux agrégés à un segment p est mesurable et son support est p .

L'ensemble $\mathcal{E}$ de lieux et son complémentaire $\text{co } \mathcal{E} = I - \mathcal{E}$ sont toujours en même temps mesurables ou non mesurables. Le support de $\text{co } \mathcal{E}$ est $\text{co } e$. Fixons une mesure μ effective et dénombrablement additive sur la tribu (T) , extension borélienne de la classe ordonnée (L) .

Si $\mathcal{E}$ est un ensemble mesurable de lieux appelons la μ -*mesure* de $\mathcal{E}$, $\mu(\mathcal{E})$, la mesure $\mu(e)$ du support e de $\mathcal{E}$. Toute la théorie de la mesurabilité lebesguienne peut être appliquée aux ensembles mesurables de lieux. Même les démonstrations des théorèmes sont analogues bien que celles-là soit facilitées beaucoup par le fait que $\mu(a)$ est une mesure dénombrablement additive sur (T) .

Voici les lemmes à démontrer successivement pour développer toute théorie de la mesurabilité. Une somme spatiale quelconque de couvertures ainsi que le produit d'un nombre fini d'entre elles est aussi une couverture. Ceci reste vrai pour les couvertures d'un ensemble de lieux. Si $\mathcal{E}$ n'est pas vide, il existe une suite infinie $\{\alpha_n\}$ de couvertures telle

que $\alpha_1 \supseteq \alpha_2 \supseteq \dots \supseteq \alpha_n \supseteq \dots$, et $e^* = \prod_{n=1}^{\infty} \alpha_n$. Si $\mathcal{E}$ est mesurable et

si $e = \prod_{n=1}^{\infty} \alpha_n$, où $\alpha_{n+1} \subseteq \alpha_n$, on a $\mu(\mathcal{E}) = \lim \mu(\alpha_n)$. Soient α_n respectivement β_n des couvertures des $\mathcal{E}$ resp. $\mathcal{F} = I - \mathcal{E}$ telles que $\alpha_n \supseteq \alpha_{n+1}$,

$\beta_n \supseteq \beta_{n+1}$, $e^* = \prod_{n=1}^{\infty} \alpha_n$, $f^* = \prod_{n=1}^{\infty} \beta_n$. Dans ces conditions, si $\mathcal{E}$ est

mesurable, on a $\lim_{n \rightarrow \infty} \mu(\alpha_n \beta_n) = 0$. Réciproquement, si α_n, β_n sont des couvertures de $\mathcal{E}$, resp. $\mathcal{F}$, et, si $\lim_{n \rightarrow \infty} \mu(\alpha_n \beta_n) = 0$, l'ensemble $\mathcal{E}$ est mesurable et l'on a $e = \prod_n \alpha_n$. Ces lemmes permettent de démontrer

que si $\mathcal{E}', \mathcal{E}''$ sont mesurables, $\mathcal{E}' + \mathcal{E}''$ est aussi mesurable et le support de $\mathcal{E}' + \mathcal{E}''$ est $e' + e''$. Si $\mathcal{E}', \mathcal{E}''$ sont mesurables, il en est de même pour $\mathcal{E}' - \mathcal{E}''$ et $\mathcal{E}' \mathcal{E}''$; les supports respectifs sont $e' - e'', e' e''$. Mentionnons que pour les ensembles mesurables $\mathcal{E}' \subseteq \mathcal{E}''$ entraîne $e' \subseteq e''$; de plus on a $\mu(\mathcal{E}') \leq \mu(\mathcal{E}'')$. Si $\mathcal{E}', \mathcal{E}''$ sont mesurables et disjoints, on a $\mu(\mathcal{E}' + \mathcal{E}'') = \mu(\mathcal{E}') + \mu(\mathcal{E}'')$. Si $\mathcal{E}_1, \mathcal{E}_2, \dots, \mathcal{E}_n, \dots$ sont mesurables et disjoints et si l'on pose

$$\mathcal{E} = \sum_{n=1}^{\infty} \mathcal{E}_n,$$

l'ensemble $\mathcal{E}$ est aussi mesurable et l'on a

$$e = \sum_{n=1}^{\infty} e_n \quad \text{et} \quad \mu(\mathcal{E}) = \sum_{n=1}^{\infty} \mu(\mathcal{E}_n).$$

Si les ensembles $\mathcal{E}_n$ ($n = 1, 2, \dots$) sont mesurables et quelconques,

leur somme $\mathcal{E} = \sum_{n=1}^{\infty} \mathcal{E}_n$ est aussi mesurable, et l'on a pour les supports

la relation analogue $e = \sum_{n=1}^{\infty} e_n$. Cela subsiste pour un produit

dénombrable d'ensembles mesurables de lieux : le support de $\prod_{n=1}^{\infty} \mathcal{E}_n$

est $\prod_{n=1}^{\infty} e_n$.

4. L'ensemble $\mathcal{E}$ de lieux s'appelle de mesure nulle lorsque son support extérieur $e^* = 0$. Si $\mathcal{E}$ est de mesure nulle et $\mathcal{F} \subseteq \mathcal{E}$, $\mathcal{F}$ est aussi de mesure nulle. Évidemment un ensemble $\mathcal{E}$ de mesure nulle est mesurable, son support $e = 0$ et $\mu(\mathcal{E}) = 0$. Pour qu'un ensemble $\mathcal{E}$ soit de mesure nulle il faut et il suffit que quel que soit $\sigma > 0$ il existe une couverture α de $\mathcal{E}$ telle que $\mu(\alpha) \leq \sigma$. Si $\mathcal{E}_n$ ($n = 1, 2, \dots$) sont de mesure nulle, leur somme est aussi de mesure nulle. Si $\mathcal{E}$,

$\mathcal{E}'$ sont des ensembles mesurables de lieux, $\mathcal{E} + \mathcal{M} = \mathcal{E}' + \mathcal{M}'$, $\mu(\mathcal{M}) = \mu(\mathcal{M}') = 0$, on a $\mu(\mathcal{E}) = \mu(\mathcal{E}')$. Si $\mathcal{E}, \mathcal{E}'$ sont mesurables, les propriétés suivantes sont équivalentes deux à deux : 1° Il existe des ensembles $\mathcal{M}, \mathcal{M}'$ de mesure nulle tels que $\mathcal{E} + \mathcal{M} = \mathcal{E}' + \mathcal{M}'$; 2° Il existe des ensembles $\mathcal{M}, \mathcal{M}'$ de mesure nulle tels que $\mathcal{E} = \mathcal{E}' + \mathcal{M} - \mathcal{M}'$; 3° de même pour $\mathcal{E} = \mathcal{E}' - \mathcal{M} + \mathcal{M}'$; 4° de même pour $\mathcal{E} - \mathcal{M} = \mathcal{E}' - \mathcal{M}'$. 5° les supports de $\mathcal{E}$ et $\mathcal{E}'$ sont identiques; 6° $\mu(\mathcal{E} - \mathcal{E}') + \mu(\mathcal{E}' - \mathcal{E}) = 0$; Si l'une de ces conditions est réalisée, nous disons que $\mathcal{E}$ et $\mathcal{E}'$ sont *égaux à un ensemble de mesure nulle près*, qu'ils ne diffèrent que par un ensemble de mesure nulle ou qu'ils sont μ -égaux.

Remarquons que la mesurabilité d'un ensemble de lieux est une notion qui ne dépend pas du choix de la mesure μ effective.

5. Il y a, en principe, deux sortes de lieux : les *lieux ordinaires* pour lesquels le produit de tous les voisinages de A est $= 0$, et les *lieux extraordinaires* où ce produit est différent de 0.

L'ensemble composé d'un seul lieu ordinaire a une mesure nulle tandis que l'ensemble composé d'un seul lieu extraordinaire possède toujours une mesure positive. De l'additivité de la mesure il résulte que le nombre des lieux extraordinaires est au plus dénombrable. Pour qu'il existe deux ensembles $\mathcal{E}, \mathcal{E}'$ différents de lieux tels que $e = e'$, il faut et il suffit qu'il existe au moins un lieu ordinaire. Cela ne se produit jamais lorsque I a un nombre fini de dimensions. Dans ce cas I se compose d'un nombre fini de lieux extraordinaires.

6. Si par égalité de deux ensembles de lieux on entend l'identité logique $\mathbf{J}$, la classe de tous les ensembles mesurables de lieux forme une tribu dénombrablement additive $(\mathbf{M})$ d'ensembles. Si par égalité de deux ensembles on entend l'égalité presque μ -partout $\mathbf{J}'$, la classe ci-dessus forme aussi une tribu dénombrablement additive $(\mathbf{M}')$ d'ensembles. Il en est de même pour le quotient $\frac{(\mathbf{M})}{\mathbf{J}}$ qui est d'ailleurs une tribu dénombrablement isomorphe et isométrique à $\mathbf{J}'$.

La notion de distance $\|\mathcal{E}, \mathcal{F}\|_{\text{af}} = \mu(\mathcal{E} - \mathcal{F}) + \mu(\mathcal{F} - \mathcal{E})$ organise $(\mathbf{M}')$ en espace métrique dont la topologie ne dépend pas du choix de la mesure effective μ sur (T). Si $\mathbf{C}$ désigne l'application qui fait correspondre à chaque ensemble mesurable $\mathcal{E}$ de lieux son support e ,

cette correspondance est dénombrablement hémimorphe ⁽¹⁾ pour $(\mathbf{M})$ et dénombrablement isomorphe pour $(\mathbf{M}')$. L'application $\mathbf{C}$ est aussi une correspondance homéomorphe et isométrique entre les métriques $(\mathbf{M}')$ et $(\mathbf{T})$ parce que tout espace de $(\mathbf{T})$ est support d'un ensemble mesurable de lieux.

7. Nous avons indiqué qu'on peut représenter les espaces α d'une classe ordonnée $(\mathbf{L})$ par les segments ordinaires $(0, \mu(\alpha))$ semi-ouverts, cette correspondance pouvant être étendue à la plus petite tribu $(\mathbf{S})$ simplement additive comprenant $(\mathbf{L})$. Remarquons que cette représentation est isométrique si l'on considère la μ -mesure des segments de $(\mathbf{S})$ et la mesure lebesguienne ordinaire des images qui y correspondent. Maintenant, après avoir développé la théorie des lieux, on peut prolonger cette correspondance de la manière suivante.

Faisons correspondre à chaque lieu ordinaire A en a l'ensemble composé du point dont l'abscisse est $\mu(a)$. Soit B un lieu extraordinaire, en b . Le produit spatial de tous les voisinages p de B est un segment $b - b^*$ où $b^* \neq b$. Faisons correspondre à B le segment semi-ouvert $(\mu(b^*), \mu(b))$. On a $\mu(b) - \mu(b^*) = \text{mes. lebesg.}(\mu(b^*), \mu(b)) = \mu((B))$. Cela étant posé, considérons la correspondance inverse de celle-ci et désignons-la par $\mathbf{r}^0$. Elle fait correspondre des lieux à des *points* et à des intervalles semi-ouverts contenus dans $(0, \mu(1))$. Appelons *atomes* ces êtres qui se transforment en lieux par l'intermédiaire de $\mathbf{r}^0$. Nous avons donc une correspondance biunivoque entre les atomes et les lieux. $\mathbf{r}^0$ se prolonge aisément ⁽²⁾ en une correspondance $\mathbf{r}$ pour les ensembles d'atomes d'une part et les ensembles de lieux d'autre part. Pour qu'à un ensemble E d'atomes corresponde un ensemble mesurable $\mathcal{E}$ de lieux il faut et il suffit que la somme logique $\check{E}$ de tous les ensembles dont chacun se compose d'un seul atome de E soit mesurable suivant Lebesgue. Dans ce cas on a $\text{mes } \check{E} = \mu(\mathcal{E})$. On voit donc qu'au lieu d'opérer sur les lieux de $(\mathbf{L})$ on peut considérer les atomes-images

(1) Nous admettons la terminologie suivante : une *homomorphie* est une application (non nécessairement bi-univoque) d'une totalité sur une partie (ou totalité) de l'autre; une *hémimorphie* est une application (non nécessairement bi-univoque) d'une totalité sur une autre totalité; une *isomorphie* applique une totalité sur l'autre d'une manière biunivoque. Naturellement les opérations doivent correspondre à des opérations respectives. Une *homéomorphie* est une isomorphie par rapport à la notion de limite.

(2) En changeant un peu son type logique.

qui y correspondent et, par conséquent, remplacer des êtres abstraits par des nombres resp., des segments ordinaires.

8. Si l'on attache à chaque lieu A de I un nombre $f(A)$ on obtient une fonction du lieu variable A . On peut à ces fonctions appliquer la théorie, de M. M. Fréchet, concernant l'intégration sur un ensemble abstrait ⁽¹⁾. Cette théorie est analogue à celle de l'intégration lebesguienne. En voici les principes que nous voudrions rappeler brièvement, en renvoyant le lecteur, pour les détails, aux travaux cités correspondant. Considérons d'abord les fonctions $f(A)$ du lieu A et à valeurs réelles. Une telle fonction s'appelle *mesurable* ⁽²⁾ lorsque quel que soit le nombre réel λ , l'ensemble

$$\hat{A} \{f(A) \leq \lambda\},$$

de tous les A pour lesquels $f(A) \leq \lambda$, est mesurable. Si l'on passe aux $\mathbf{r}$ -images, à une telle fonction il correspond une fonction mesurable (d'une manière lebesguienne) si l'on convient d'attribuer à tout point d'un atome-segment la même valeur. Les fonctions- $\mathbf{r}$ -images sont donc des fonctions mesurables pour lesquelles chaque intervalle qui correspond à un lieu extraordinaire est un intervalle de constance.

Une fonction complexe $f(A)$ du lieu s'appelle mesurable lorsque sa partie réelle et sa partie imaginaire sont mesurables toutes les deux. L'intégrale Fréchetienne d'une fonction réelle $f(A)$ se définit d'une manière analogue à celle d'une fonction réelle ordinaire dans la théorie de Lebesgue. Une fonction $f(A)$ pour laquelle l'intégrale

$$\int_I f(A) d\mu_A$$

existe, s'appelle μ -sommable ou, tout simplement *sommable*. Aux fonctions μ -sommables correspondent les fonctions $\mathbf{r}$ -images sommables suivant Lebesgue et inversement. La fonction complexe $F(A)$ s'appelle sommable lorsque sa partie réelle et sa partie imaginaire sont sommables. L'intégrale

$$\int_{\mathcal{S}} f(A) d\mu_A,$$

⁽¹⁾ M. FRÉCHET, *Sur l'intégrale d'une fonctionnelle étendue à un ensemble abstrait* (Bull. Soc. Math. de France, 43, 1915), voir aussi NIKODYM, *Fund. Math.* 14, 1929.

⁽²⁾ Ou compatible avec la tribu des ensembles mesurables de lieux.

où $\mathcal{E}$ est un ensemble mesurable de lieux a un sens bien déterminé lorsque f est sommable sur I . Nous attirerons l'attention sur les *fonctions complexes au carré μ -sommable* : ce sont les fonctions $F(A)$ mesurables pour lesquels $\int_I |F(A)|^2 d\mu_A$ existe.

9. Nous avons construit une image numérique des lieux et des fonctions dans $(0, \mu(I))$. Remarquons qu'en partant de cette image nous pouvons aisément, par un changement convenable de l'échelle de la mesure, trouver une image analogue dans l'intervalle $(-\infty, +\infty)$, $(0, +\infty)$, etc., au moyen d'une fonction continue strictement croissante. Nous n'insistons plus là-dessus.

10. Soit maintenant (L) une classe ordonnée et fermée dont l'extension borélienne est une tribu saturée (T) . Nous savons qu'il existe alors un vecteur générateur de l'espace $\mathfrak{r}$ par rapport à (T) . Soit $\vec{\omega}$ l'un de ces vecteurs générateurs et envisageons la mesure effective sur (T)

$$\mu(a) = |\text{Proj}_a \vec{\omega}|^2 \quad (a \in T).$$

Reprenons la théorie des lieux de la classe (L) et de la μ -mesure des ensembles de lieux.

Nous allons trouver une correspondance entre les fonctions complexes $X(A)$ au carré sommable et les vecteurs $\vec{x}$ de l'espace $\mathfrak{r}$.

Convenons, d'une manière générale, à désigner par $X_{\mathcal{E}}(A)$ la fonction égale à $X(A)$ lorsque $A \in \mathcal{E}$ et à 0 lorsque A n'appartient pas à $\mathcal{E}$. Soit $\Omega(A)$ la fonction constante égale partout à 1. Appelons *fonction en escalier* toute fonction

$$(1) \quad \sum_{i=1}^n \lambda_i \Omega_{\mathcal{E}_i}(A),$$

où les $\mathcal{E}_i$ sont des ensembles mesurables de lieux, les λ_i des nombres complexes et où la somme est finie ($n \geq 1$). Faisons correspondre à chaque fonction (1) le vecteur

$$\sum_{i=1}^n \lambda_i \text{Proj}_{e_i} \vec{\omega},$$

où e_i est le support de $\mathcal{E}_i$ et $\vec{\omega}$ le vecteur générateur choisi ci-dessus.

On démontre sans peine que cette correspondance est indépendante de la manière de représenter la fonction en escalier par la formule (1). Cette correspondance $\mathbf{R}^0$ est hémimorphe et isométrique, car si $X(A) \sim \vec{x}$, $Y(A) \sim \vec{y}$, on a $\alpha X(A) + \beta Y(A) \sim \alpha \vec{x} + \beta \vec{y}$, et

$$\int_I \overline{X(A)} Y(A) d\mu_A = (\vec{x}, \vec{y}).$$

Pour qu'à deux fonctions en escalier corresponde le même vecteur, il faut et il suffit que les fonctions ne diffèrent que sur un ensemble de μ -mesure nulle.

On peut prolonger $\mathbf{R}^0$ de manière qu'elle devienne une correspondance $\mathbf{R}$ hémimorphe entre la classe $(\mathbf{F})$ de toutes les fonctions complexes définies presque μ -partout sur I et au carré sommable d'une part et les vecteurs de l'espace $\mathbf{1}$ d'autre part.

En effet, soit $F(A)$ une fonction complexe définie presque μ -partout et au carré μ -sommable. Il existe une suite infinie $X_n(A)$ de fonctions en escalier telle que

$$\lim_{n \rightarrow \infty} \int_I |F(A) - X_n(A)|^2 d\mu_A = 0.$$

Soient $\vec{x}_n$ les vecteurs $\mathbf{R}^0$ -correspondant aux $X_n(A)$.

Quel que soit $\varepsilon > 0$ il existe un indice n_0 tel que, si $n, m \geq n_0$, on ait

$$\int_I |X_n(A) - X_m(A)|^2 d\mu_A \leq \varepsilon.$$

En vertu de l'isométrie de $\mathbf{R}^0$ on a donc

$$|\vec{x}_n - \vec{x}_m| \leq \varepsilon \quad \text{pour } n, m \geq n_0.$$

La suite $\{\vec{x}_n\}$ étant convergente, il existe un vecteur $\vec{f}$ tel que

$$\vec{x}_n \rightarrow \vec{f}.$$

Faisons correspondre à $F(A)$ le vecteur $\vec{f}$. On démontre que 1° $\vec{f}$ ne dépend pas du choix de la suite $\{X_n(A)\}$ tendant en moyenne carrée vers la fonction donnée $F(A)$; 2° la correspondance $\mathbf{R}^0$ n'est qu'un cas particulier de la correspondance définie ci-dessus. Cette correspondance, que nous désignerons par $\mathbf{R}$, est homomorphe. Si $F(A) \underset{\mathbf{R}}{\sim} \vec{f}$,

$G(A) \underset{\mathbf{R}}{\sim} \vec{g}$, on a $\alpha F(A) + \beta G(A) \underset{\mathbf{R}}{\sim} \alpha \vec{f} + \beta \vec{g}$ et

$$\int_1 \overline{F(A)} G(A) d\mu_A = (\vec{f}, \vec{g}).$$

A deux fonctions $F_1(A)$, $F_2(A)$ correspond le même vecteur si et seulement si F_1 et F_2 ne diffèrent que sur un ensemble de μ -mesure nulle. $\mathbf{R}$ est homéomorphe, c'est-à-dire que, si

$$\int_1 |F_n(A) - F(A)|^2 \rightarrow 0,$$

on a pour les vecteurs correspondant

$$\vec{f}_n \rightarrow \vec{f}.$$

Mais $\mathbf{R}$ est hémimorphe, c'est-à-dire qu'elle applique la classe $(\mathbf{F})$ sur l'espace entier 1. En effet, $\overset{\rightarrow}{\omega}$ étant un vecteur générateur, tout vecteur $\vec{f}$ peut être approximé par des sommes finies

$$\sum_i \lambda_i \text{Proj}_{e_i} \overset{\rightarrow}{\omega}, \quad \text{où } e_i \in (T).$$

Si l'on convient de considérer deux fonctions comme égales lorsqu'elles ne diffèrent que sur un ensemble de mesure nulle, l'hémimorphisme devient isomorphisme, et l'on a affaire à une correspondance biunivoque entre les fonctions de $(\mathbf{F})$ et les vecteurs de 1.

Si l'on veut considérer les images numériques, on obtient aussi une correspondance biunivoque entre des fonctions ordinaires au carré sommable dans $(0, \mu(1))$ et les vecteurs de l'espace 1.

11. La correspondance $\mathbf{R}^{-1}$ transforme les vecteurs de 1 en des fonctions d'un lieu variable. Si $\alpha \in (T)$ où (T) est l'extension borélienne de la classe ordonnée envisagée, l'image de α est la classe de toutes les fonctions $F(A)$ au carré sommable telles que $F(A) = 0$ presque partout, lorsque A n'appartient pas à un ensemble $\mathcal{E}$ dont le support est α . Au moyen de la théorie des lieux on peut transformer toute tribu d'espaces en une tribu très simple, comme celle ci-dessus, de manière que la structure de la tribu devienne plus évidente. En effet, chaque tribu peut être prolongée de manière qu'elle devienne saturée et, dans chaque

tribu dénombrablement additive on peut choisir une classe ordonnée qui la détermine parfaitement.

Remarquons que la notion de lieu *dépend essentiellement du choix de la classe ordonnée*. Les lieux ne restent pas toujours les mêmes si l'on envisage dans la tribu donnée une autre classe ordonnée.

VI.

APPLICATIONS A LA THÉORIE DES TRIBUS.

1. Le lemme suivant se démontre sans difficulté :

Soit $X(A)$ une fonction mesurable du lieu A variable dans I . Soit $\mathcal{E}$ un ensemble mesurable de lieux, $0 < \lambda \leq |X(A)| < \mu$ dans $\mathcal{E}$ et $\sigma > 0$. Dans ces conditions il existe un nombre fini $\mathcal{E}_1, \dots, \mathcal{E}_n$ de sous-ensembles mesurables de $\mathcal{E}$ et des nombres complexes $\lambda_1, \dots, \lambda_n$ tels que

$$\left| \sum_{i=1}^n \lambda_i X_{\mathcal{E}_i}(A) - 1 \right| \leq \sigma.$$

Il en résulte le lemme plus général suivant :

Si $X(A)$ est sommable et différente de 0 sur un ensemble mesurable $\mathcal{E}$, alors quel que soit $\sigma > 0$, il existe des ensembles mesurables $\mathcal{E}_1, \dots, \mathcal{E}_n$ et des nombres $\lambda_1, \dots, \lambda_n$ tels que

$$\int_{\mathcal{E}} \left| \sum_{i=1}^n \lambda_i X_{\mathcal{E}_i}(A) - 1 \right|^2 d\mu_A \leq \sigma, \quad \sum_{i=1}^n \mathcal{E}_i = \mathcal{E}.$$

2. THÉORÈME. — Soit $\mathfrak{I}$ séparable, (T) une tribu énumérablement additive de sous-espaces de $\mathfrak{I}$. Supposons qu'il existe un vecteur générateur de $\mathfrak{I}$ par rapport à (T) . Dans ces conditions (T) est saturée.

Démonstration. — Soit (L) une sous-classe ordonnée et fermée de (T) et telle que (T) soit l'extension borélienne de (L) . Soit $\vec{\omega}$ un vecteur générateur et soit

$$\mu(a) = \left| P_a \vec{\omega} \right|^2, \quad (a \in (T))$$

la mesure effective correspondante. Envisageons les lieux de (L) et la théorie des fonctions μ -images des vecteurs de $\mathbf{1}$.

Supposons que (T) ne soit pas saturée et que p soit un espace n'appartenant pas à (T) , mais qui puisse être adjoint à (T) . L'espace p est donc compatible avec (T) . On a $M_T(p) = p$. Posons

$$a'' = \prod_{\substack{\text{df} \\ p \subseteq a}} a, \quad a' = \sum_{\substack{\text{df} \\ a \subseteq p}} a$$

pour les $a \in (T)$. On a $a' \subseteq p \subseteq a''$. Il suffit de démontrer que $a' = a''$. Supposons que $a' \neq a''$. On a

$$\text{Proj}_{a''-a'} p = (a'' - a') p,$$

d'où il résulte que $\text{Proj}_{a''-a'} p \neq 0$. Soit $\vec{\eta}$ un vecteur de p tel que $P_{a''-a'} \vec{\eta} \neq 0$. Envisageons une fonction-image $\mathcal{H}(A)$ du vecteur $\vec{\eta}$. Envisageons des ensembles $\mathcal{E}'$, $\mathcal{E}''$ dont les supports sont a' , resp a'' . L'ensemble $\mathcal{F}$ de tous les A où $\mathcal{H}_{\mathcal{E}''-\mathcal{E}'}(A) \neq 0$ est de μ -mesure positive. Si l'on y applique le lemme énoncé ci-dessus, on parvient au résultat

$$P_{f\omega} \vec{\eta} \in M_T\left(\left(\vec{\eta}\right)\right),$$

où f est le support de $\mathcal{F}$. Cela donne $f \subseteq p$, d'où l'on obtient la fausse égalité $f = 0$. L'identité $a' = a''$ démontrée ainsi implique que $p \in (T)$ contrairement à l'hypothèse.

Le théorème ci-dessus combiné avec un autre du paragraphe 2 n° 2 montre que, dans le cas de séparabilité de $\mathbf{1}$:

La condition nécessaire et suffisante pour qu'il existe un vecteur générateur de $\mathbf{1}$ par rapport à la tribu (T) dénombrablement additive est que (T) soit saturée.

THÉORÈME. — *Si $\mathbf{1}$ est séparable, (T) une tribu saturée et $\vec{\xi}$ un vecteur de $\mathbf{1}$, la condition nécessaire et suffisante pour que $\vec{\xi}$ soit un vecteur générateur de $\mathbf{1}$ par rapport à (T) est que la mesure sur (T) , définie par l'égalité*

$$\nu(a) = \left| P_a \vec{\xi} \right|^2$$

soit effective.

Pour démontrer cela on envisage un vecteur générateur $\vec{\omega}$ et l'on considère les fonctions μ -images où $\mu(a) = \left| P_a \vec{\omega} \right|^2$.

Par application du lemme mentionné on obtient que $\vec{\omega} \in M_T(\vec{\xi})$ d'où $M_T(\vec{\xi}) = 1$.

Voici une autre caractérisation du vecteur générateur :

THÉORÈME. — *Si $\mathfrak{I}$ est séparable, (T) une tribu saturée et $\vec{\xi}$ un vecteur, la condition nécessaire et suffisante pour que $\vec{\xi}$ soit un vecteur générateur de $\mathfrak{I}$ par rapport à (T) est que $\vec{\xi}$ n'appartienne à aucun espace de (T) différent de $\mathfrak{I}$.*

Ce théorème résulte du théorème précédent.

3. Nous savons que si l'on pose $\mu(a) = |\text{Proj}_a \vec{x}|^2$ pour tout a appartenant à une tribu, on obtient une mesure dénombrablement additive sur cette tribu. Il est remarquable que la réciproque est aussi vraie,

THÉORÈME. — *Si $\mathfrak{I}$ est séparable, (T) une tribu saturée et $\nu(a)$ une mesure dénombrablement additive sur (T) , il existe un vecteur $\vec{x}$ tel que pour tout $a \in (T)$ on ait*

$$\nu(a) = |\text{Proj}_a \vec{x}|^2.$$

Pour démontrer cela prenons un vecteur générateur $\vec{\omega}$ et choisissons dans (T) une classe ordonnée fermée (L) engendrant (T) . Envisageons les lieux de (L) et les fonctions μ -images des vecteurs de $\mathfrak{I}$ où l'on a posé

$$\mu(a) = |\text{Proj}_a \vec{\omega}|^2 \quad \text{pour tout } a \in (T).$$

D'une manière générale, si $\mathcal{E}$ est un ensemble mesurable de lieux, posons

$$f(\mathcal{E}) \stackrel{\text{df}}{=} \nu(e), \quad \mu(\mathcal{E}) \stackrel{\text{df}}{=} \mu(e),$$

où e est le support de $\mathcal{E}$. La fonction $f(\mathcal{E})$ est dénombrablement additive sur la tribu dénombrablement additive formée des ensembles mesurables de lieux. Comme si $\mu(\mathcal{E}) = 0$, on a aussi $f(\mathcal{E}) = 0$, il existe ⁽¹⁾ une fonction μ -sommable $S(A) \geq 0$ telle que

$$f(\mathcal{E}) = \int_{\mathcal{E}} S(A) d\mu_A.$$

⁽¹⁾ O. NIKODYM, *Sur une généralisation des intégrales de M. J. Radon* (Fund. Math., 14, 1929).

Soit $\Xi(A)$ une fonction complexe et mesurable telle que

$$|\Xi(A)|^2 = S(A),$$

et soit $\vec{\xi}$ le vecteur correspondant. On obtient alors $\nu(e) = |\text{Proj}_e \vec{\xi}|^2$, ce qu'il fallait démontrer.

Le théorème ci-dessus est aussi vrai pour la mesure dénombrablement additive sur une tribu quelconque. De plus, on peut démontrer que si l'on prolonge une tribu dénombrablement additive (T) en une tribu (T') ayant le même caractère, toute mesure $\mu(a)$ dénombrablement additive sur (T) peut être aussi prolongée sur (T') de manière qu'elle reste dénombrablement additive. Le même résultat subsiste lorsqu'on ne suppose que l'additivité simple de (T) , mais en conservant l'additivité dénombrable de la mesure μ .

VII.

APPLICATIONS AUX TRANSFORMATIONS SELFADJOINTES.

1. Rappelons brièvement les notions fondamentales concernant ces transformations ⁽¹⁾.

Soit $\vec{y} = \mathcal{A}(\vec{x})$ une transformation linéaire continue ou non dont le domaine d'existence $\mathcal{D}\mathcal{A}$ est partout dense dans $\mathbf{1}$. Considérons les vecteurs $\vec{\xi}$ pour lesquels il existe au moins un vecteur $\vec{\eta}$ satisfaisant à l'équation

$$(1) \quad (\mathcal{A}(\vec{x}), \vec{\xi}) = (\vec{x}, \vec{\eta}),$$

pour tous les $\vec{x} \in \mathcal{D}\mathcal{A}$. Un tel vecteur $\vec{\xi}$ existe toujours, par exemple $\vec{\xi} = \vec{o}$. Quel que soit le cas, étant donné un tel vecteur $\vec{\xi}$, il lui correspond un seul vecteur $\vec{\eta}$, cela à cause de l'hypothèse $\overline{\mathcal{D}\mathcal{A}} = \mathbf{1}$. L'application qui fait correspondre le vecteur $\vec{\eta}$ au vecteur $\vec{\xi}$ est une transformation linéaire, qui s'appelle *transformation adjointe* à $\mathcal{A}$.

⁽¹⁾ Voir par exemple M. H. STONE, *Linear transformations* (loc. cit.) ou BÉLA V. SZ. NAGY, *Spektraldarstellung linearer Transformationen des Hilbertschen Raumes* (Ergebnisse d. Math. u. ihrer Grenzgeb., V).

On désigne celle-ci par $\mathcal{A}^*$ et son domaine est l'ensemble de tous les $\vec{x}$ en question.

La transformation $\mathcal{A}$ s'appelle *selfadjointe* lorsque $\mathcal{A} = \mathcal{A}^*$. Une telle transformation est *symétrique-hermitienne*, ce qui signifie que

$$(\mathcal{A}\vec{x}, \vec{y}) = (\vec{x}, \mathcal{A}\vec{y})$$

pour tous les $\vec{x}$ et tous les $\vec{y}$ appartenant à $\mathcal{D}\mathcal{A}$. Elle est *maximale*, en ce sens qu'il n'existe aucune transformation symétrique-hermitienne qui serait un vrai prolongement de $\mathcal{A}$.

Le théorème fondamental sur les transformations self-adjointes est le suivant, appelé *théorème spectral* de Hilbert :

Si $\mathcal{H}$ est une transformation self-adjointe, il existe un système d'espaces $\{E_\lambda\}$ bien déterminé dépendant du paramètre réel λ variant dans $(-\infty, +\infty)$ et jouissant des propriétés suivantes :

1° E_λ est une fonction non décroissante du paramètre λ : si $\lambda' \leq \lambda''$, on a $E_{\lambda'} \subseteq E_{\lambda''}$;

2° $E_{-\infty} = 0$, $E_{+\infty} = 1$;

3° E_λ est continue du côté droit, c'est-à-dire que si

$$\lambda_1 > \lambda_2 > \dots > \lambda_n > \dots, \quad \lim \lambda_n = \lambda_0,$$

on a

$$E_{\lambda_0} = \bigcap_{n=1}^{\infty} E_{\lambda_n};$$

4° Si $\vec{x} \in \mathcal{D}\mathcal{H}$ et $\vec{y} \in 1$, on a

$$(1) \quad (\mathcal{H}\vec{x}, \vec{y}) = \int_{-\infty}^{+\infty} \lambda d(\text{Proj}_{E_\lambda} \vec{x}, \vec{y}),$$

$$(2) \quad |\mathcal{H}\vec{x}|^2 = \int_{-\infty}^{+\infty} \lambda^2 d|\text{Proj}_{E_\lambda} \vec{x}|^2,$$

où les intégrales sont celles de Stieljes ;

5° La condition nécessaire et suffisante pour que $\vec{x} \in \mathcal{D}\mathcal{H}$ est que l'intégrale

$$\int_{-\infty}^{+\infty} \lambda^2 d|\text{Proj}_{E_\lambda} \vec{x}|^2$$

existe.

On voit que $\{E_\lambda\}$ peut être considéré comme une classe ordonnée. Nous appellerons $\{E_\lambda\}$ *échelle spectrale* et l'extension borélienne de la classe ordonnée $\{E_\lambda\}$, *tribu spectrale* de $\mathcal{H}$.

2. Remarquons qu'au lieu de (1) on peut écrire

$$(3) \quad \mathcal{H} \vec{x} = \int_{-\infty}^{+\infty} \lambda \, d\text{Proj}_{E_\lambda} \vec{x},$$

où l'intégrale (3) se définit d'une manière analogue à celle qui est employée d'ordinaire pour les intégrales ordinaires de Stieltjes. Les passages à la limite peuvent être exécutés par la convergence forte ou faible ⁽¹⁾ parce que ces deux modes de convergence sont équivalents dans le cas considéré. De plus, on peut démontrer que si, pour un $\vec{x}$ donné, l'intégrale

$$\int_{-\infty}^{+\infty} \lambda \, d(\text{Proj}_{E_\lambda} \vec{x}, \vec{y})$$

existe pour tous les $\vec{y} \in 1$, il en résulte l'existence de l'intégrale

$$(1) \quad \int_{-\infty}^{+\infty} \lambda \, d\text{Proj}_{E_\lambda} \vec{x},$$

et réciproquement.

Si l'on introduit pour les segments semi-ouverts $(\lambda', \lambda'']$ une mesure vectorielle

$$\overset{\rightarrow}{\mu}(\lambda', \lambda''] \stackrel{\text{df}}{=} \text{Proj}_{E_{\lambda''}} \vec{x} - \text{Proj}_{E_{\lambda'}} \vec{x} = \text{Proj}_{E_{\lambda''} - E_{\lambda'}} \vec{x},$$

cette mesure se généralise aux ensembles ⁽²⁾ mesurables $\mathcal{E}$ selon Lebesgue et l'on peut introduire les intégrales du type de Radon-Fréchet

$$(2) \quad \int_{\mathcal{E}} \lambda \, d\text{Proj}_{E_\lambda} \vec{x}.$$

Ces intégrales se définissent de manière Lebesguienne et (2) considéré comme fonction d'ensemble variable $\mathcal{E}$ est une fonction dénombrablement additive.

⁽¹⁾ $\vec{x}_n$ converge *fortement* vers $\vec{x}$ lorsque $|\vec{x}_n, \vec{x}| \rightarrow 0$, la convergence est *faible* lorsque $\lim (\vec{x}_n, \vec{y}) = (\vec{x}, \vec{y})$ pour tout $\vec{y} \in 1$.

⁽²⁾ Comparer par exemple O. NIKODYM, *Remarques sur les intégrales de Stieltjes en connexion avec celles de MM. Radon et Fréchet* (Ann. Soc. Math. Polon., 18, 1945). Voir aussi B. v. SZ. NAGY, *loc. cit.* et K. FRIEDRICH, *Math. Ann.*, 110, 1935.

En particulier l'intégrale $\int_{\lambda'}^{\lambda''} \lambda d\text{Proj}_{E_\lambda} \vec{x}$ doit être considéré comme l'intégrale Fréchetienne étendue sur (λ', λ'') . Il est donc préférable d'écrire

$$\int_{(\lambda')}^{\lambda''} \lambda d\text{Proj}_{E_\lambda} \vec{x} \quad \text{ou} \quad \int_{(\lambda', \lambda'')} \lambda d\text{Proj}_{E_\lambda} \vec{x},$$

au lieu de $\int_{\lambda'}^{\lambda''} \lambda d\text{Proj}_{E_\lambda} \vec{x}$.

3. Outre le théorème spectral qui nous sert de base pour plusieurs démonstrations, nous aurons besoin de la notion d'espace invariant d'une transformation self-adjointe.

Soit $\mathcal{H}$ une transformation self-adjointe. Un espace p s'appelle *espace invariant* (propre) de $\mathcal{H}$ lorsqu'il jouit des propriétés suivantes :

- 1° Si $\vec{x} \in p$ et $\vec{x} \in \mathbb{Q}\mathcal{H}$, on a $\mathcal{H}(\vec{x}) \in p$;
- 2° Si $\vec{x} \in \mathbb{Q}\mathcal{H}$, le vecteur $\text{Proj}_p \vec{x}$ appartient aussi à $\mathbb{Q}\mathcal{H}$ ⁽¹⁾.

En particulier les espaces 0 et 1 sont des espaces invariants.

Un théorème connu énonce que la condition nécessaire et suffisante pour qu'un espace p soit invariant de $\mathcal{H}$ est que l'équation suivante soit satisfaite pour tous les $\vec{x} \in \mathbb{Q}\mathcal{H}$

$$\text{Proj}_p \mathcal{H}(\vec{x}) = \mathcal{H} \text{Proj}_p(\vec{x}).$$

Si p est un espace invariant, l'espace cop est aussi un espace invariant. Tout espace E_λ de l'échelle spectrale est un espace invariant.

On peut démontrer que si α appartient à la tribu spectrale, α est un espace invariant. La réciproque n'est pas toujours vraie.

Soit p un espace invariant et posons

$$\mathcal{H}_p \stackrel{\text{def}}{=} p \mathcal{H} p,$$

ce qui désigne la restriction de $\mathcal{H}$ sur $p \mathbb{Q}\mathcal{H}$. Il est connu que $\mathcal{H}_p$ est une transformation self-adjointe dans l'espace p .

(1) M. STONE dit « p reduces $\mathcal{H}$ » (*Linear Transf.*, loc. cit.).

Si l'on pose

$$\mathcal{H}'_p(\vec{x}) \stackrel{\text{df}}{=} \mathcal{H}(\vec{x}_p), \quad \mathcal{H}'_{\text{cop}}(\vec{x}) \stackrel{\text{df}}{=} \mathcal{H}(\vec{x}_{\text{cop}}), \quad \text{pour } \vec{x} \in \mathbb{Q}\mathcal{H},$$

on a

$$\mathcal{H}(\vec{x}) = \mathcal{H}'_p(\vec{x}) + \mathcal{H}'_{\text{cop}}(\vec{x}) \quad \text{pour } \vec{x} \in \mathbb{Q}\mathcal{H}.$$

On a

$$\mathcal{H}'_p(\vec{z}) = \vec{0} \quad \text{lorsque } \vec{z} \in \text{cop } \mathbb{Q}\mathcal{H}$$

et

$$\mathcal{H}'_{\text{cop}}(\vec{z}) = \vec{0} \quad \text{lorsque } \vec{z} \in p \mathbb{Q}\mathcal{H}.$$

On voit ainsi qu'un espace invariant décompose $\mathcal{H}$ en deux parties $\mathcal{H}_p$ et $\mathcal{H}_{\text{cop}}$ qui sont self-adjointes resp. dans les deux espaces complémentaires p et cop .

4. Les préparatifs terminés passons à l'application de notre théorie.

THÉORÈME. — *La condition nécessaire et suffisante pour qu'un espace p soit invariant pour $\mathcal{H}$ est que p soit compatible avec la tribu spectrale.*

Démonstration. — Soit p un espace invariant et soient $\{F_\lambda\}$ et $\{G_\lambda\}$ les échelles spectrales resp. de $\mathcal{H}_p$ et $\mathcal{H}_{\text{cop}}$. On a

$$\mathcal{H}_p \vec{x} = \int_{-\infty}^{\infty} \lambda d(\text{Proj}_{F_\lambda} \vec{x}) \quad \text{pour } \vec{x} \in p \mathbb{Q}\mathcal{H}$$

et

$$\mathcal{H}_{\text{cop}} \vec{y} = \int_{-\infty}^{\infty} \lambda d(\text{Proj}_{G_\lambda} \vec{y}) \quad \text{pour } \vec{y} \in \text{cop } \mathbb{Q}\mathcal{H}.$$

Donc, si $\vec{z} \in \mathbb{Q}\mathcal{H}$, on a

$$\mathcal{H} \vec{z} = \int_{-\infty}^{\infty} \lambda d(\text{Proj}_{F_\lambda + G_\lambda} \vec{z}),$$

car $F_{\lambda'} \perp G_{\lambda''}$ quels que soient λ' et λ'' . La représentation spectrale étant unique, on a $F_\lambda + G_\lambda = E_\lambda$. On en déduit que $F_\lambda = E_\lambda p$ et $G_\lambda = E_\lambda \text{cop}$, d'où

$$E_\lambda = E_\lambda p + E_\lambda \text{cop},$$

ce qui exprime que p est compatible avec tout E_λ . On en déduit que p est compatible avec la tribu spectrale.

Nous avons établi la nécessité de la condition. La démonstration de sa suffisance ne présente pas de difficulté. Elle s'obtient par l'approximation de l'intégrale spectrale par des sommes.

La tribu spectrale peut être caractérisée par la propriété connue suivante (1) : c'est la classe de tous les espaces dont chacun est compatible avec tous les espaces invariants.

5. Nous avons déjà mentionné que la tribu spectrale (T) n'est pas nécessairement saturée. La condition nécessaire et suffisante pour que (T) soit saturée est que $\mathcal{H}$ ait un spectre simple (2). Nous nous contentons de cette remarque sans même définir le terme « spectre simple » parce que tout cela s'éclaircira un peu plus loin.

Supposons que la tribu spectrale (T) ne soit pas saturée. Démontrons qu'il existe alors un vecteur $\vec{\varphi}$ tel que $M_T\left(\left(\begin{smallmatrix} \vec{\varphi} \\ \vec{\varphi} \end{smallmatrix}\right)\right)$ n'appartienne pas à (T). En effet, si (T) n'est pas saturée, il existe un espace p compatible avec (T) mais n'y appartenant pas. On a

$$p = \sum_{\vec{\varphi} \in p} M_T\left(\left(\begin{smallmatrix} \vec{\varphi} \\ \vec{\varphi} \end{smallmatrix}\right)\right)$$

où la somme est spatiale. S'il n'existait aucun $\vec{\varphi}$ jouissant de la propriété indiquée, tous les $M_T\left(\left(\begin{smallmatrix} \vec{\varphi} \\ \vec{\varphi} \end{smallmatrix}\right)\right)$ où $\vec{\varphi} \in p$ appartiendraient à (T) et, par conséquent, p aussi, ce qui est contradictoire.

Cela étant posé, on peut démontrer par l'induction transfinie et, en se servant des remarques concernant les espaces invariants, qu'il existe une suite au plus dénombrable de vecteurs

$$\vec{\varphi}_1, \vec{\varphi}_2, \dots, \vec{\varphi}_n, \dots$$

telle que si l'on pose

$$g_n = M_T\left(\left(\begin{smallmatrix} \vec{\varphi} \\ \vec{\varphi}_n \end{smallmatrix}\right)\right),$$

les espaces g_n sont orthogonaux deux à deux et $\sum_n g_n = 1$ (3). C'est la

(1) FRIEDRICHs.

(2) Ce théorème est connu et se trouve chez STONE, *Linear Transformations*, loc. cit.

(3) Ce théorème est un cas particulier du théorème se trouvant dans le travail de M. HIDEGORO NAKANO, *Unitäriinvariante hypermaximale normale Operatoren* (*Ann. of Math.* t. 42, 1941, et *Unitäriinvarianten im allgemeinen Euklidischen Raume* (*Mathem. Annalen*, t. 118, 1941-1943).

partie la plus importante du théorème de Hellinger et Hahn ⁽¹⁾ utilisé dans la théorie des invariants unitaires des opérations self-adjointes.

Chaque espace g_n est capable d'être adjoint à la tribu spectrale (T), mais on peut démontrer qu'ils peuvent y être adjoints simultanément. Cela se démontre par l'adjonction successive des g_n . Il existe donc une plus petite tribu (S) comprenant (T) et tous les espaces g_n ($n = 1, 2, \dots$). Soit (S*) la plus petite tribu dénombrablement additive contenant (S). On démontre sans difficulté que (S*) est identique à l'ensemble de tous les espaces

$$\sum a_i g_{v_i} \quad \text{où } a_i \in (T),$$

la somme spatiale étant au plus dénombrable. De plus, on démontre que la tribu (S*) est saturée.

Cela étant, on peut nommer la suite des vecteurs $\{\vec{\varphi}_n\}$ *suite saturante de vecteurs* et, la suite $\{g_n\}$, *suite saturante d'espaces*. Le fait que (S*) est saturée se démontre en établissant que le vecteur

$$\vec{\xi} = \sum_{n=1}^{\infty} \frac{1}{n!} \frac{\vec{\varphi}_n}{|\vec{\varphi}_n|}$$

est un vecteur générateur de 1 par rapport à (S*).

Nous appellerons (S*) la *tribu saturée de (T) au moyen des vecteurs saturants* $\vec{\varphi}_n$ (*espaces saturants* g_n).

6. La tribu (S*) étant saturée, il existe un vecteur générateur par exemple $\vec{\xi}$ du numéro précédent. Mais on peut démontrer un peu plus : *il existe un vecteur générateur $\vec{\omega}$ de 1 par rapport à (S*), ce vecteur appartenant au domaine $\mathcal{D}\mathcal{H}$ de la transformation self-adjointe donnée $\mathcal{H}$.*

Démonstration. — Soit $\{E_\lambda\}$ l'échelle spectrale de $\mathcal{H}$. Posons

$$e_N = E_{N+1} - E_N \quad \text{où } N = 0, \pm 1, \pm 2, \dots$$

Désignons par (S_N*) l'ensemble de tous les espaces $e_N a$ où $a \in (S^*)$. (S_N*) est une tribu de sous-espaces de l'espace e_N . La tribu (S*) étant saturée, il en est de même de (S_N*) par rapport à l'espace total e_N . On

(1) Voir par exemple F. J. WECKEN, *Math. Ann.*, t. 30, 1929.

démontre sans peine que $e_N \in \mathcal{H}$. Soit $\vec{\omega}_N$ un vecteur générateur de l'espace e_N par rapport à (S_N^*) . Choisissons des nombres complexes σ_N tels que $\sum_{N=-\infty}^{+\infty} |\sigma_N|^2$ converge et posons

$$\vec{\omega} \stackrel{\text{df}}{=} \sum_{N=-\infty}^{+\infty} \frac{\sigma_N}{|N|+1} \frac{\vec{\omega}_N}{|\vec{\omega}_N|}.$$

On démontre que $\vec{\omega} \in \mathcal{H}$ et que $M_{(S^*)} \left(\begin{pmatrix} \vec{\omega} \\ \vec{\omega} \end{pmatrix} \right) = 1$.

Le vecteur $\vec{\omega}$ est le vecteur cherché.

7. Étant donnée une transformation self-adjointe $\mathcal{H}$, ou bien sa tribu spectrale (T) est saturée ou bien on peut trouver une tribu (S^*) saturée contenant (T) , composée d'espaces invariants de $\mathcal{H}$ et déduite de (T) par l'adjonction d'un ensemble au plus dénombrable d'espaces saturants $p_1, p_2, \dots$. Nous allons construire une sous-classe ordonnée (L) de (S^*) appropriée à l'étude de la transformation $\mathcal{H}$. Posons $(S^*) \stackrel{\text{df}}{=} (T)$, lorsque la tribu spectrale (T) est saturée. Dans ce cas posons aussi $p_1 \stackrel{\text{df}}{=} 1$.

Considérons les classes ordonnées :

$$\begin{array}{ll} (C_1) & p_0 + E_\lambda p_1 \quad \text{où } p_0 \stackrel{\text{df}}{=} 0, \\ (C_2) & p_0 + p_1 + E_\lambda p_2. \\ & \dots\dots\dots \\ (C_n) & p_0 + p_1 + \dots + p_{n-1} + E_\lambda p_n, \\ & \dots\dots\dots \end{array}$$

Leur réunion est une classe ordonnée (L') . Désignons par (L) la fermeture de (L') . On démontre sans peine que (S^*) est l'extension borélienne de (L) . On obtient (L) par l'adjonction à (L') d'un nombre au plus dénombrable d'espaces.

Envisageons la classe ordonnée (L) et, soit $\vec{\omega} \in \mathcal{H}$ un vecteur générateur de 1 par rapport à (S^*) . Envisageons la mesure effective

$$\mu(a) = |\text{Proj}_a \vec{\omega}|^2$$

sur (S^*) . Envisageons les lieux A de (L) , la μ -mesure de leurs ensembles et les $\mathbf{R}^{-1}$ — fonctions — images des vecteurs de 1. Soit (F) l'ensemble de

toutes les fonctions complexes $F(A)$ et aux μ -carrés sommables. Soit $\Omega(A)$ la fonction constante partout égale à 1 et, définissons la fonction $\Omega_{\mathcal{E}}(A)$ comme étant égale à 1 lorsque $A \in \mathcal{E}$ et, égale à zéro dans le cas contraire. A la transformation self-adjointe $\mathcal{H}$ il correspond dans l'espace (F) une transformation self-adjointe bien déterminée H . Désignons son domaine par QH .

La correspondance R étant hémimorphe et isométrique la relation $\omega \in Q\mathcal{H}$ entraîne $\Omega(A) \in QH$.

Soit

$$\Phi(A) = \underset{\text{df}}{H}(\Omega(A)), \quad \Phi(A) \in (F).$$

Soit $\mathcal{E}$ un ensemble mesurable de lieux tel que son support e soit un espace invariant de $\mathcal{H}$.

On a

$$\text{Proj}_e \mathcal{H} \left(\overset{\rightarrow}{\omega} \right) = \mathcal{H} \text{Proj}_e \overset{\rightarrow}{\omega},$$

d'où

$$\Phi_{\mathcal{E}}(A) = H(\Omega_{\mathcal{E}}(A)).$$

Il en résulte que pour les fonctions en escalier

$$X(A) = \sum_i \lambda_i \Omega_{\mathcal{E}_i}(A) \quad (\text{somme finie})$$

on a

$$H(X(A)) = \sum_i \lambda_i \Phi_{\mathcal{E}_i}(A) = \sum_i \lambda_i \Omega_{\mathcal{E}_i}(A) \Phi(A),$$

donc

$$(1) \quad H(X(A)) = \Phi(A) X(A).$$

Cette formule s'étend à toutes les fonctions $F(A)$ appartenant à QH de manière qu'on obtient la formule

$$(2) \quad H(F(A)) = \Phi(A) F(A).$$

qui généralise (1).

Par la considération du produit scalaire on démontre que $\Phi(A)$ est une fonction presque partout réelle. De plus on trouve que QH coïncide avec l'ensemble de toutes les fonctions $F(A)$ de (F) pour lesquelles $\Phi(A)F(A)$ est au carré μ -sommable.

La formule (2) représente une forme canonique simple pour n'importe quelle transformation self-adjointe $\mathcal{H}$ dans un espace séparable 1. La

théorie reste valable même pour les espaces à un nombre fini de dimensions.

8. La fonction $\Phi(A)$ qui caractérise $\mathcal{H}$ par rapport à l'application R^{-1} n'est pas arbitraire. Nous allons montrer qu'elle se compose d'un nombre au plus dénombrable de fonctions monotones non décroissantes (l'abstraction faite d'un ensemble de μ -mesure nulle).

Reprenons la classe ordonnée (L') (du n° 7) dont la fermeture est (L) . Nous avons déjà mentionné que (L') ne diffère de (L) que par la présence d'un nombre au plus dénombrable d'éléments. Soit A un lieu ordinaire de (L) en α .

Faisons l'hypothèse que α n'est égal à aucune somme finie

$$p_0 + p_1 + \dots + p_i$$

et que $\alpha \in (L')$. Il existe alors une représentation unique

$$\alpha = p_0 + \dots + p_{n-1} + p_n E_\lambda,$$

où n et λ sont bien déterminés. Le point λ est un point de continuité de l'échelle spectrale $\{E_\lambda\}$. On peut désigner A par le symbole $A(n, \lambda)$. Considérons la fonction $\psi(A)$ définie par l'équation

$$\psi(A) = \lambda \quad \text{lorsque} \quad A = A(n, \lambda).$$

La fonction $\psi(A)$ est définie pour tous les lieux ordinaires, exception faite d'un nombre au plus dénombrable d'entre eux. Si A varie dans un des segments p_i la fonction $\psi(A)$ croît.

Soient maintenant A_0 un des lieux envisagés et n_0, λ_0 les nombres correspondants. Soit $\varepsilon > 0$ et considérons l'espace

$$v_\varepsilon \stackrel{\text{def}}{=} p_{n_0}(E_{\lambda_0} - E_{\lambda_0 - \varepsilon}).$$

On a

$$\mathcal{H}(\text{Proj}_{v_\varepsilon} \vec{\omega}) = \int_{(\lambda_0 - \varepsilon)}^{\lambda_0} \lambda d(P_{E_\lambda} \vec{\omega}_{v_\varepsilon}).$$

On obtient sans peine l'inégalité

$$|\mathcal{H} \vec{\omega}_{v_\varepsilon} - \lambda_0 \vec{\omega}_{v_\varepsilon}|^2 \leq \varepsilon^2 \mu(v_\varepsilon).$$

Cette inégalité se traduit pour les fonctions images par la suivante

$$|\mathbf{H} \Omega_{v_\varepsilon}(A) - \lambda_0 \Omega_{v_\varepsilon}(A)|^2 \leq \varepsilon^2 \mu(v_\varepsilon),$$

c'est-à-dire, par

$$(1) \quad \int_{\nu_\varepsilon} |\Phi(A) - \lambda_0|^2 d\mu_A \leq \varepsilon^2 \mu(\nu_\varepsilon).$$

Mais, on a, d'autre part,

$$(2) \quad \int_{\nu_\varepsilon} |\Psi(A) - \lambda_0|^2 d\mu_A \leq \varepsilon^2 \mu(\nu_\varepsilon),$$

parce que

$$\lambda_0 - \varepsilon \leq \varphi(A) \leq \lambda_0 \quad \text{pour } A \in \nu_\varepsilon.$$

L'inégalité de Minkowski $\sqrt{f|f+g|^2} \leq \sqrt{f|f|^2} + \sqrt{f|g|^2}$ donne, compte tenu de (1) et (2)

$$\int_{\nu_\varepsilon} |\Phi(A) - \Psi(A)|^2 d\mu \leq 4\varepsilon^2 \mu(\nu_\varepsilon),$$

c'est-à-dire

$$\frac{1}{\mu(\nu_\varepsilon)} \int_{\nu_\varepsilon} |\Phi(A) - \Psi(A)|^2 d\mu \leq 4\varepsilon^2.$$

En considérant les images numériques des lieux et des fonctions du lieu, on se trouve dans les conditions bien connues d'application d'un théorème de Lebesgue d'après lequel la dérivée d'une intégrale indéfinie d'une fonction sommable existe presque partout et est égale à la fonction à intégrer.

On parvient ainsi au résultat que

$$\Phi(A) = \Psi(A) \quad \text{presque partout.}$$

Jusqu'à présent nous n'avons considéré que des lieux ordinaires. Mais le cas des lieux extraordinaires ne présente pas de difficulté, étant donné qu'ils correspondent aux valeurs propres de la transformation $\mathcal{BC}$.

Nous avons ainsi démontré que $\Phi(A)$ est monotone non décroissante dans tout segment p_i . Si l'on considère les images numériques, on voit très nettement en quoi consiste le phénomène de multiplicité même pour les points du spectre continu : la valeur λ possède la multiplicité $0, 1, \dots, \omega$ lorsque la droite dont l'ordonnée est λ coupe $0, 1, \dots, \omega$ branches différentes de la courbe $y = \Phi(x)$. Il serait intéressant d'étudier les petites perturbations d'une opération self-adjointe, par la méthode développée ci-dessus : la tribu (S^*) et l'aspect de la courbe $y = \Phi(x)$ ne devraient pas changer trop, si la perturbation est suffisamment faible.

9. Jusqu'à présent, en construisant les fonctions $F(A)$ images des vecteurs, nous sommes partis du fait que nous avons appliqué le vecteur $\vec{\omega}$ générateur de 1 sur la fonction constante $\Omega(A) = 1$ qui représente un vecteur générateur de l'espace (F) . Mais on peut choisir, au lieu de $\Omega(A)$, une fonction arbitraire $\Xi(A)$ presque partout différente de zéro qui est donc aussi un vecteur générateur de (F) et faire correspondre $\vec{\omega}$ à $\Xi(A)$. Les nouvelles fonctions images donnent aussi une représentation canonique de $\mathcal{H}$, mais dans ce cas le facteur $\Phi(A)$ ne serait pas nécessairement aussi simple.

Remarquons que la théorie des opérateurs $f(\mathcal{H})$ appliqués aux transformations self-adjointes se développe avec une facilité remarquable. Nous avons suivi ainsi un chemin inverse de celui suivi par M. J. v. Neumann et ses successeurs qui ont obtenu la représentation canonique de $\mathcal{H}$ après avoir développé la théorie des opérations $f(\mathcal{H})$ au moyen des intégrales de Stieltjes; notre théorie peut être appliquée avec succès aussi aux transformations normales, c'est-à-dire aux transformations $\mathcal{N}$ telles que

$$\mathcal{N}\mathcal{N}^* = \mathcal{N}^*\mathcal{N}, \quad \overline{\mathcal{N}}\mathcal{N} = 1.$$

VIII.

REMARQUES PRÉLIMINAIRES SUR LE SYSTÈME DES COORDONNÉES.

1. Notre tâche finale consiste à introduire *un système général de coordonnées dans l'espace de Hilbert-Hermite* et, en particulier, un système continu des coordonnées. Pour expliquer l'idée directrice de notre construction et justifier quelques définitions qui pourraient paraître un peu bizarres, considérons le cas très simple d'un système orthonormal des coordonnées dans l'espace à trois dimensions.

Ce système se compose de trois vecteurs unitaires $\vec{e}_1, \vec{e}_2, \vec{e}_3$ orthogonaux deux à deux. Un vecteur arbitraire $\vec{x}$ s'exprime par la formule (développement de Fourier),

$$(1) \quad \vec{x} = x_1 \vec{e}_1 + x_2 \vec{e}_2 + x_3 \vec{e}_3,$$

où x_1, x_2, x_3 sont des nombres complexes.

Pour que cette décomposition du vecteur $\vec{x}$ en ses composantes $x_i \vec{e}_i$ puisse se prêter à une généralisation, on doit la modifier un peu. Pour cela, remarquons d'abord que $\vec{e}_1, \vec{e}_2, \vec{e}_3$ peuvent être considérés comme projections sur les axes e_1, e_2, e_3 d'un seul vecteur $\vec{e}_1 + \vec{e}_2 + \vec{e}_3$. Remplaçons ce vecteur par un vecteur arbitraire $\vec{\omega}$ qui ne soit pas situé dans aucun plan du système des coordonnées. Ses projections $\vec{\omega}_{e_1}, \vec{\omega}_{e_2}, \vec{\omega}_{e_3}$ sur les axes e_1, e_2, e_3 sont $\neq 0$. On a

$$\vec{x} = \sum_i (\vec{e}_i, \vec{x}) \vec{e}_i = \sum_i \left(\frac{\vec{\omega}_{e_i}}{|\vec{\omega}_{e_i}|}, \vec{x} \right) \frac{\vec{\omega}_{e_i}}{|\vec{\omega}_{e_i}|},$$

donc

$$(2) \quad \vec{x} = \sum_i x_{e_i} \vec{\omega}_{e_i},$$

où l'on a posé

$$(3) \quad x_{e_i} = \frac{(\vec{\omega}_{e_i}, \vec{x})}{|\vec{\omega}_{e_i}|^2}.$$

Nous verrons que ce sont les formules (2) et (3) qui sont bien adaptées à une généralisation dans l'espace à une infinité de dimensions.

2. Dans le système des coordonnées envisagé, il y a une tribu, à savoir la tribu (T) composée des espaces

$$(T) \quad 0, \quad e_1, \quad e_2, \quad e_3, \quad e_1 + e_2, \quad e_2 + e_3, \quad e_3 + e_1, \quad 1 = e_1 + e_2 + e_3.$$

C'est une tribu saturée et $\vec{\omega}$ en est un vecteur générateur. Si l'on introduit la mesure effective sur (T),

$$\mu(a) = |\text{Proj}_a \vec{\omega}|^2,$$

on obtient

$$(3) \quad x_{e_i} = \frac{(\vec{\omega}_{e_i}, \vec{x})}{\mu(e_i)}.$$

Le système généralisé de coordonnées diffère du système ordinaire (1) par un choix différent des vecteurs unitaires sur les axes e_1, e_2, e_3 : au lieu de prendre les vecteurs $\vec{e}_1, \vec{e}_2, \vec{e}_3$ on prend $\vec{e}_1 \sqrt{\mu(e_1)}, \vec{e}_2 \sqrt{\mu(e_2)}, \vec{e}_3 \sqrt{\mu(e_3)}$ et, par conséquent, au lieu des coefficients x_1, x_2, x_3 ,

on obtient les nombres

$$\frac{x_1}{\sqrt{\mu(e_1)}}, \quad \frac{x_2}{\sqrt{\mu(e_2)}}, \quad \frac{x_3}{\sqrt{\mu(e_3)}}.$$

La formule (3) où il y a une division par la mesure suggère le terme *coordonnée-densité* pour x_{e_1} , x_{e_2} , x_{e_3} , que nous voudrions employer au lieu du terme *coordonnée*, que nous préférons attacher aux nombres

$$x_{e_i} \stackrel{\text{df}}{=} \left(\overset{\rightarrow}{\omega_{e_i}}, \vec{x} \right).$$

Les axes sont des espaces à une dimension et ce sont aussi des segments les plus petits de la classe ordonnée (L) des espaces de (T)

$$(L) \quad 0 \subset e_1 \subset e_1 + e_2 \subset e_1 + e_2 + e_3 = I.$$

Il n'y a que des lieux extraordinaires A_1 , A_2 , A_3 attachés aux espaces $a_1 \stackrel{\text{df}}{=} e_1$, $a_2 \stackrel{\text{df}}{=} e_1 + e_2$, $a_3 \stackrel{\text{df}}{=} I$. Les intervalles de ces lieux sont précisément les axes e_1 , e_2 , e_3 du système des coordonnées. Cette circonstance nous donne la suggestion que, dans le cas général, ce sont précisément les lieux qui devraient constituer une sorte d'axes du système des coordonnées.

(Dans le cas simple considéré la définition admise plus haut d'un lieu en a , comme ensemble des segments $a - x$, est évidemment superflue et inutilement compliquée : il suffirait de le définir comme l'axe e_i , mais le cas général exige une telle complication.)

Or nous allons, en nous plaçant dans le cas général, considérer les lieux comme une sorte d'axes du système des coordonnées qui sera définie par une classe ordonnée extraite d'une tribu saturée et par un vecteur générateur.

Le but d'un tel système général des coordonnées consiste surtout en la possibilité de décomposer un vecteur en ses composantes, même dans le cas où il y en aurait une infinité non dénombrable. On s'aperçoit immédiatement que ces composantes ne peuvent pas être des vecteurs ordinaires, mais des êtres nouveaux qui conserveraient une partie du caractère d'un vecteur.

De plus, pour avoir la possibilité d'une sommation non dénombrable, il faut être en possession d'un instrument mathématique approprié, d'une sorte de sommation-intégration portant sur des êtres mentionnés ci-dessus.

IX.

QUASI-VECTEURS ET LEUR SOMMATION.

1. Avant de donner la définition précise d'un système général des coordonnées, faisons les préparatifs nécessaires. Commençons par l'introduction de deux notions nouvelles.

Soit (L) une classe ordonnée et fermée dont l'extension borélienne soit une tribu saturée (T) . Soit $\vec{\omega}$ un vecteur générateur de l'espace r (séparable) par rapport à (T) et $\mu(a) = |\text{Proj}_a \vec{\omega}|^2$ une mesure effective sur (T) . Envisageons les lieux respectifs.

Soit (V) un espace vectoriel normé de F. Riesz-Banach ⁽¹⁾. Dans ce qui va suivre (V) sera ou bien l'espace r ou bien l'espace de tous les nombres complexes. Dans le premier cas la norme $\|\vec{\varphi}\|$ d'un vecteur $\vec{\varphi}$ est son module, dans le second, c'est sa valeur absolue.

Étant donné un lieu A , appelons *quasi-vecteur en A* une fonction $\vec{\varphi}(p)$ définie pour tous les voisinages p de A , les valeurs de $\vec{\varphi}$ étant des vecteurs de l'espace (V) . En particulier, si (V) est l'espace des nombres complexes, nous emploierons la dénomination *quasi-nombre* et supprimerons la flèche.

Le quasi-vecteur $\frac{\vec{\varphi}(p)}{\mu(p)}$ respectivement le quasi-nombre $\frac{\varphi(p)}{\mu(p)}$, sera appelé *densité du quasi-vecteur* respectivement *quasi-nombre*. Nous désignerons les quasi-vecteurs par $\vec{\varphi}_A$, φ_A et leurs densités par $\vec{\varphi}_A^*$, φ_A^* .

Les quasi-vecteurs en A peuvent être ajoutés, multipliés par un scalaire, on peut en faire un produit scalaire, les définitions respectives étant aisées à deviner. Signalons à l'avance que ce seront des quasi-vecteurs qui seront les composantes d'un vecteur.

Supposons qu'à chaque lieu A soit attaché un quasi-vecteur $\vec{\varphi}_A$. Dans ce cas on a affaire à un *champ de quasi-vecteurs* ou si on le préfère, à une fonction définie dans I et dont les valeurs sont des quasi-vecteurs.

⁽¹⁾ S. BANACH, *Théorie des opérations linéaires*, Varsovie, 1932; FR. RIESZ, *Über lineare Funktionalgleichungen* (Acta Math., 41, 1918).

Un tel champ est équivalent à une fonction qui fait correspondre à chaque segment de (L) un vecteur de (V) .

2. Il faut créer une sorte de sommation des quasi-vecteurs d'un champ. Dans cet ordre d'idées introduisons quelques notions auxiliaires, quelques termes, et faisons quelques remarques concernant les lieux.

Appelons *agrégat* l'espace o et tout ensemble fini de segments disjoints de la classe ordonnée envisagée (L) . Comme nous aurons fréquemment besoin d'envisager la somme spatiale des éléments $p_1, p_2, \dots, p_m$ d'un agrégat P , alternativement avec l'agrégat lui-même, nous emploierons pour ces deux êtres différents le même symbole

$$\sum_i p_i,$$

où le point veut dire que les p_i sont disjoints. On évitera ainsi des complications de langage sans craindre des confusions. D'une manière analogue P désignera, selon le cas, l'agrégat et la somme spatiale des segments dont P est composé. Nous utiliserons aussi les symboles de l'écart

$$\|\mathcal{E}, P\|, \|P, Q\|,$$

où $\mathcal{E}$ est un ensemble mesurable de lieux et P, Q des agrégats, ces symboles désignant

$$\|e, \alpha\|, \|\alpha, \beta\|,$$

où α, β sont des sommes spatiales des éléments des agrégats P, Q , et e le support de $\mathcal{E}$. Nous allons aussi pour des raisons de brièveté identifier quelquefois $\mathcal{E}$ avec son support e , l'orthographe précise étant facile à rétablir.

Les agrégats $\sum_i p_i$ formant dans la μ -topologie de la tribu (T) un ensemble partout dense, chaque ensemble mesurable $\mathcal{E}$ de lieux peut être considéré comme limite d'une suite infinie d'agrégats. Nous écrirons $P_n \rightarrow \mathcal{E}$ et aussi $P_n \rightarrow e$ (où e est le support de $\mathcal{E}$), ce qui exprime que $\lim_{n \rightarrow \infty} \|P_n, \mathcal{E}\| = 0$.

3. Nous savons qu'il y a deux sortes de lieux : les lieux ordinaires et les lieux extraordinaires. Comme nous voulons introduire une sorte d'intégration basée sur le principe des segments « infiniment petits », la

présence des lieux extraordinaires nécessite l'introduction de la notion suivante : appelons *nombre caractéristique d'un segment* $b - a = q$ le nombre non négatif

$$Nq \stackrel{\text{df}}{=} \mu(b - a) - \mu(b - b^*),$$

où $b - b^* = \prod_{\text{df}}^p p$, et où le produit spatial est étendu à tous les voisins du lieu B placé en b .

On a évidemment $a \subseteq b^* \subseteq b$, $a \subset b$. Si B est un lieu ordinaire, on a $b^* = b$ et $Nq = \mu(q)$. Si B est extraordinaire, Nq est égal à $\mu(q)$ diminué de la mesure de l'intervalle $b - b^*$ du lieu B. On a $Nq > 0$, lorsque B est ordinaire, mais si B est extraordinaire, Nq peut être même égal à zéro.

Si $P = \sum_i p_i$ est un agrégat appelons *nombre caractéristique de P* le maximum NP des nombres $Np_1, Np_2, \dots$. Nous allons très souvent considérer des suites $\{P_n\}$ d'agrégats convergeant vers un ensemble $\mathcal{E}$ et telles que $NP_n \rightarrow 0$.

4. Les préliminaires établis, procédons à la définition de la somme d'un champ de quasi-vecteurs.

Soit $\vec{\varphi}_A$ un champ de quasi-vecteurs (quasi-nombres), $\mathcal{E}$ un ensemble mesurable de lieux et e son support. Considérons une suite infinie

$$\{P_n\} = \left\{ \sum_i p_{in} \right\} \text{ d'agrégats telle que}$$

$$(1) \quad P_n \rightarrow \mathcal{E}, \quad NP_n \rightarrow 0,$$

et faisons la somme

$$\vec{\Phi}_n = \sum_i \vec{\varphi}(p_{in}).$$

Lorsque le vecteur $\vec{\Phi}_n$ tend (suivant sa norme) vers un vecteur $\vec{\Phi}$ bien déterminé qui ne dépend pas du choix de la suite P_n satisfaisant à (1), le vecteur $\vec{\Phi}$ s'appellera la *somme du champ* $\vec{\varphi}_A$ et sera désigné par le symbole

$$\sum_{\mathcal{E}}^{\vec{\varphi}_A} \quad \text{ou} \quad \sum_e^{\vec{\varphi}_A} \quad (1).$$

(1) Cette notion ressemble aux intégrales de Weierstrass-Burkill. Pour une étude détaillée de telles intégrales, voir N. ARONSZAJN, *Quelques recherches sur l'intégrale de Weierstrass* (I, *Revue Scientifique*, 8, IX, 1939; II, *Revue Scientifique*, 3, III, 1940).

Nous dirons alors que $\sum_I \varphi_A$ est *sommable sur $\mathcal{E}$ resp. e.*

§. La théorie générale de la sommation introduite ci-dessus quoique ne présentant pas de difficultés, est assez pénible à cause de la présence éventuelle des lieux extraordinaires. Je me bornerai à citer quelques théorèmes généraux concernant le cas le plus général, en supprimant les démonstrations qui seront publiées dans un travail plus détaillé.

THÉORÈME. — Si $\sum_I \varphi$ existe, on peut trouver pour chaque $\varepsilon > 0$ un nombre $\eta > 0$ tel que les inégalités

$$\mu(P) \leq \eta, \quad NP \leq \eta,$$

pour l'agrégat P entraînent

$$\|\sum_I \varphi(P)\| < \varepsilon,$$

où $\sum_I \varphi(P)$ désigne $\sum_i \varphi(p_i)$ lorsque $P = \sum_i p_i$.

Ce théorème exprime que la somme est approximée d'une manière uniforme par les sommes $\sum_i \varphi(p_i)$.

THÉORÈME. — Si $\sum_I \varphi$ existe et si $\mathcal{E}$ est un ensemble mesurable de lieux, la somme

$$\sum_{\mathcal{E}} \varphi$$

existe aussi.

Ce théorème se démontre de proche en proche, en commençant par le cas où $\mathcal{E}$ correspond à un segment de (L) et en envisageant des ensembles de plus en plus compliqués.

THÉORÈME. — Si $\sum_I \varphi$ existe et que l'on pose

$$\tilde{J}(\mathcal{E}) = \sum_{\mathcal{E}} \varphi,$$

la fonction $\tilde{J}(\mathcal{E})$ d'ensemble $\mathcal{E}$ est une fonction dénombrablement

additive, c'est-à-dire que, si $\mathcal{E}_1, \mathcal{E}_2, \dots, \mathcal{E}_n, \dots$ sont disjoints, on a

$$\vec{J}\left(\sum_{i=1}^{\infty} \mathcal{E}_i\right) = \sum_{i=1}^{\infty} \vec{J}(\mathcal{E}_i),$$

où la sommation à gauche est logique et où celle de droite est une sommation de vecteurs suivie d'un passage à la limite suivant la norme dont l'espace (V) est doué.

THÉORÈME. — Si $\sum_I \vec{\varphi}$ existe, pour tout $\varepsilon > 0$ il existe un $\eta > 0$ tel que si $\mu(a) \leq \eta$, $a \in (T)$, on ait

$$\left\| \sum_a \vec{\varphi}_a \right\| \leq \varepsilon.$$

Ce théorème exprime la continuité uniforme de la fonction $\sum_{\mathcal{E}} \vec{\varphi}$ de l'ensemble $\mathcal{E}$, la continuité étant celle de la métrique induite, sur la tribu d'ensembles mesurables de lieux, par la mesure μ .

Étant donné deux champs de quasi-vecteurs $\vec{\varphi}_A, \vec{\psi}_A$ sommables dans I , convenons de dire que $\vec{\varphi}_A$ équivaut à $\vec{\psi}_A$, lorsque

$$\sum_{\mathcal{E}} \vec{\varphi}_A = \sum_{\mathcal{E}} \vec{\psi}_A,$$

pour tout ensemble $\mathcal{E}$ mesurable. On voit que cette notion satisfait aux lois formelles de l'identité. Nous écrirons $\vec{\varphi}_A \approx \vec{\psi}_A$.

6. Dans le cas d'un champ M_A de quasi-nombres on peut démontrer ce qui suit :

THÉORÈME. — Si M_A est sommable sur I , il existe une fonction $\Phi(A)$ du lieu A , bien déterminée à un ensemble de μ -mesure nulle près, telle que

$$\sum_{\mathcal{E}} M_A = \int_{\mathcal{E}} \Phi(A) d\mu_A.$$

Si l'on considère le champ J_A

$$\{J(q)\}_{\text{df}} = \left\{ \int_q \Phi(A) d\mu_A \right\},$$

de quasi-nombres, on a

$$J_A \approx M_A.$$

X.

CHAMPS RÉGULIERS DE QUASI-VECTEURS.

1. Dans le cas où (V) est l'espace $\mathbf{1}$, on obtient des résultats assez précis si l'on se borne à considérer des champs $\vec{\varphi}_\Lambda$ pour lesquels $\vec{\varphi}(q) \in q$ quel que soit le segment q de (L) .

Nous appellerons ces champs de quasi-vecteurs des *champs réguliers*.

THÉORÈME. — Si $\vec{\varphi}_\Lambda$ est un champ régulier et sommable et qu'on définit le champ F_Λ de quasi-nombres par

$$\{F(q)\} = \left\{ \left| \vec{\varphi}(q) \right|^2 \right\},$$

le champ F_Λ est aussi sommable et l'on a

$$\sum_a \left| \vec{\varphi}_\Lambda \right|^2 = \left| \sum_a \vec{\varphi}_\Lambda \right|^2 \quad \text{pour tout } a \in (T).$$

THÉORÈME. — Si $\vec{\varphi}_\Lambda$ est un champ régulier et sommable, on a pour tout $a \in (T)$:

$$\sum_a \vec{\varphi}_\Lambda \in a.$$

La démonstration est basée sur le lemme suivant :

Si $a_n, a \in (T)$, $\lim_{n \rightarrow \infty} \|a_n, a\|_u = 0$, $\vec{\xi}_n \in a_n$, et $\lim_{n \rightarrow \infty} \vec{\xi}_n = \vec{\xi}$, alors $\vec{\xi} \in a$.

THÉORÈME. — Si $\vec{\varphi}_\Lambda$ est un champ régulier et sommable, il existe un vecteur $\vec{\psi} \in \mathbf{1}$, tel que pour tout $a \in (T)$

$$\sum_a \vec{\varphi}_\Lambda = \text{Proj}_a \vec{\psi}.$$

Si l'on pose $\vec{\varphi}'(q) \stackrel{\text{df}}{=} \text{Proj}_q \vec{\psi}$ pour tout segment q le champ $\vec{\varphi}'_\Lambda$ est équivalent à $\vec{\varphi}_\Lambda$.

THÉORÈME. — Soient $\vec{\varphi}_A, \vec{\psi}_A$ des champs réguliers et sommables. Dans ces conditions, si $\vec{\varphi}_A \approx \vec{\psi}_A$, on a pour toute suite infinie $P_n = \sum_i p_{in}$ d'agrégats satisfaisant aux conditions

$$\sum_i p_{in} = 1, \quad NP_n \rightarrow 0,$$

la relation suivante :

$$(1) \quad \lim_{n \rightarrow \infty} \sum_i \left| \vec{\varphi}(p_{in}) - \vec{\psi}(p_{in}) \right|^2 = 0.$$

La condition ci-dessus exprimée par (1) sera appelée l'équivalence forte de $\vec{\varphi}_A, \vec{\psi}_A$ et nous l'écrirons : $\vec{\varphi}_A \approx \vec{\psi}_A$.

Pour les champs sommables et réguliers l'équivalence forte entraîne l'équivalence ordinaire.

THÉORÈME. — Si $\vec{\varphi}_A, \vec{\psi}_A$ sont des champs réguliers et sommables, la relation $\vec{\varphi}_A \approx \vec{\psi}_A$ entraîne $\left| \vec{\varphi}_A \right|^2 \approx \left| \vec{\psi}_A \right|^2$.

THÉORÈME. — Si $\vec{\varphi}_A, \vec{\psi}_A$ sont des champs réguliers et sommables, le champ de quasi-nombres $(\vec{\varphi}_A, \vec{\psi}_A)$ définie par

$$F(q) \stackrel{\text{df}}{=} (\vec{\varphi}(q), \vec{\psi}(q))$$

est aussi sommable et l'on a

$$\left| S_a(\vec{\varphi}_A, \vec{\psi}_A) \right| \leq \left| S_a \vec{\varphi}_A \right| \left| S_a \vec{\psi}_A \right| \quad \text{pour tout } a \in (T).$$

THÉORÈME. — Si $\vec{\varphi}_A \approx \vec{\varphi}'_A, \vec{\psi}_A \approx \vec{\psi}'_A$, on a dans le cas de sommabilité et régularité

$$(\vec{\varphi}_A, \vec{\psi}_A) \approx (\vec{\varphi}'_A, \vec{\psi}'_A).$$

THÉORÈME. — Si $\vec{\varphi}_A$ est régulier et sommable, $\vec{x} \in 1$ et $a \in (T)$, on a

$$S_a(\vec{\varphi}_A, \vec{x}) = \left(S_a \vec{\varphi}_A, \vec{x} \right).$$

2. Soit $f(x)$ une fonction complexe et μ -sommable du lieu variable x et, soit A un lieu. Le champ de quasi-nombres, défini par

$$\frac{1}{\mu(q)} \int_q f(x) d\mu_x$$

pour tous les voisinages q du lieu A , est un quasi-nombre en A . Nous le désignerons par $\text{val}_A f$ (valeur moyenne de $f(x)$ en A).

Si l'on fait varier A dans I le champ $\text{val}_A f$ de quasi-nombres n'est pas sommable en général.

En s'appuyant sur un théorème de M. Aronszajn concernant les intégrales de Hellinger ⁽¹⁾, on peut démontrer le théorème suivant :

THÉORÈME. — Si $f(x)$ est une fonction μ -sommable du lieu variable x et $\vec{\omega}$ le vecteur générateur (envisagé dans la théorie présente), pour que le champ de quasi-vecteurs

$$\text{val}_A f \vec{\omega}_A,$$

défini par

$$\left\{ \frac{1}{\mu(q)} \int_q f(x) d\mu_x \text{Proj}_q \vec{\omega} \right\}$$

soit sommable, il faut et il suffit que $f(x)$ soit une fonction au carré μ -sommable (donc μ -image d'un vecteur de I).

3. Parmi les champs de quasi-nombres le champ μ_A défini par

$$\mu(q) = |\text{Proj}_q \vec{\omega}|^2$$

joue un rôle important. Il est sommable et l'on a pour tout $\mathcal{E}$ mesurable :

$$\sum_{\mathcal{E}} \mu_A = \mu(\mathcal{E}).$$

Si $\vec{\varphi}_A$ est un champ sommable, il est toujours équivalent (et même identique) au champ $\vec{\varphi}_A \cdot \mu_A$ parce que celui-ci est le champ défini par

$$\frac{\vec{\varphi}_A(q)}{\mu(q)} \mu(q).$$

(1) N. ARONSZAJN, *Revue Scientifique*, VI, 3 mars 1940.

Par conséquent chaque somme $\sum_{\mathcal{E}}^{\rightarrow} \varphi_A$ peut être mise sous la forme $\sum_{\mathcal{E}}^{\rightarrow} \varphi_A \cdot \mu_A$ qui ressemble aux intégrales ordinaires, le quasi-nombre μ_A étant une sorte de différentielle employée sous le signe d'intégration. On pourrait convenir d'écrire

$$\int_{\mathcal{E}}^{\rightarrow} \varphi_A \cdot dA \quad \text{au lieu de} \quad \sum_{\mathcal{E}}^{\rightarrow} \varphi_A \cdot \mu_A.$$

Remarquons enfin que la somme $\sum_{\mathcal{E}}^{\rightarrow} \varphi$ et la notion de sommabilité ne dépendent pas du choix de la mesure effective μ , et par conséquent, du choix du vecteur générateur ω .

XI.

SYSTÈME GÉNÉRAL DE COORDONNÉES.

1. Nous avons maintenant les moyens nécessaires pour définir les systèmes généraux de coordonnées dans l'espace de Hilbert-Hermite. Nous nous bornerons à en donner les principes, en réservant les applications pour un Mémoire ultérieur.

Soit (L) une classe ordonnée et fermée de sous-espaces d'un espace $\mathfrak{r}$ séparable de Hilbert-Hermite. Soit (T) l'extension borélienne de (L), ω un vecteur générateur de $\mathfrak{r}$ par rapport à (T).

Nous appellerons le système composé de (L) et de ω *système général des coordonnées dans $\mathfrak{r}$* .

Envisageons la mesure $\mu(a) \stackrel{\text{df}}{=} \left| \text{Proj}_a \omega \right|^2$ effective sur (T) et envisageons les lieux de (L) et les fonctions $-\mu-$ images des vecteurs de $\mathfrak{r}$. Étant donné un vecteur $\vec{x} \in \mathfrak{r}$, le système $[(L), \omega]$ des coordonnées, et un lieu A quelconque, appelons *A-composante de $\vec{x}$* le quasi-vecteur $\vec{x}_A$ défini par $\text{Proj}_q \vec{x}$, pour tous les voisinages q du lieu A. Appelons *A-composante-densité de $\vec{x}$* , le quasi-vecteur $\vec{x}_A^*$, défini par

$$\frac{\text{Proj}_q \vec{x}}{\mu(q)},$$

A-coordonnée de $\vec{x}$, le quasi-nombre $\vec{x}_A$, défini par

$$(\text{Proj}_q \vec{\omega}, \vec{x})$$

et, A-coordonnée-densité de $\vec{x}$, le quasi-nombre $\vec{x}_A^\bullet$, défini par

$$\frac{(\text{Proj}_q \vec{\omega}, \vec{x})}{\mu(q)}.$$

Les grandeurs $\vec{x}_A$, $\vec{x}_A^\bullet$, x_A , $x_A^\bullet$ sont d'accord avec les remarques intuitives faite dans le paragraphe VIII.

On voit que $\vec{x}_A$ et $\vec{x}_A^\bullet$, lorsque A varie dans I, sont des champs réguliers et sommables de quasi-vecteurs.

2. On démontre aisément les formules suivantes

$$(1) \quad \vec{x} = \sum_I \vec{x}_A, \quad \text{Proj}_a \vec{x} = \sum_a \vec{x}_A$$

ou, ce qui revient au même,

$$\vec{x} = \sum_I \vec{x}_A^\bullet \mu_A, \quad \text{Proj}_a \vec{x} = \sum_a \vec{x}_A^\bullet \mu_A,$$

ces formules exprimant que le vecteur est égal à la somme de ses A-composantes.

On a aussi

$$|\vec{x}|^2 = \sum_I |\vec{x}_A|^2, \quad |\text{Proj}_a \vec{x}|^2 = \sum_a |\vec{x}_A|^2$$

et, pour deux vecteurs $\vec{x}, \vec{y}$ on a

$$(\vec{x}, \vec{y}) = \sum_I (\vec{x}_A, \vec{y}_A), \quad (\text{Proj}_a \vec{x}, \text{Proj}_a \vec{y}) = \sum_a (\vec{x}_A, \vec{y}_A),$$

c'est-à-dire

$$(\vec{x}, \vec{y}) = \sum_I (\vec{x}_A^\bullet, \vec{y}_A^\bullet) \mu_A = \sum_I (\vec{x}_A, \vec{y}_A) \mu_A.$$

3. Par un procédé d'approximation par des vecteurs dont les μ -images sont des fonctions à escalier, on démontre le théorème suivant :

THÉORÈME. — *Le champ des composantes $\vec{x}_A$ d'un vecteur $\vec{x}$ est équivalent au champ $\vec{\omega}_A x_A$, défini par*

$$\text{Proj}_q \vec{\omega} \frac{(\text{Proj}_q \vec{\omega}, \vec{x})}{\mu(q)}.$$

De plus, si $\Xi(x)$ est une fonction μ -image de $\vec{x}$, le champ $\vec{x}_A$ est aussi équivalent au champ

$$\text{val}_A \Xi \vec{\omega}_A$$

défini par

$$\frac{1}{\mu(q)} \int_q \Xi(x) d\mu_x \text{Proj}_q \vec{\omega}.$$

Ce théorème permet de démontrer le suivant :

THÉORÈME. — *Pour tout vecteur $\vec{x} \in \mathbf{I}$ a lieu la décomposition de Fourier :*

$$\begin{aligned} \vec{x} &= \sum_{\mathbf{I}} x_A \vec{\omega}_A, \\ \vec{x} &= \sum_{\mathbf{I}} \text{val}_A \Xi \vec{\omega}_A. \end{aligned}$$

Ici $\vec{\omega}_A$ généralise le système orthogonal et x_A les coefficients de Fourier.

Remarquons que c'est pour qu'on puisse obtenir ce théorème important qu'on a défini plus haut la sommation de façon si générale, en choisissant pour l'intégrale la forme de Weierstrass-Burkill.

Le théorème ci-dessus donne naissance à un cortège de formules

$$\begin{aligned} (\vec{x}, \vec{y}) &= \sum_{\mathbf{I}} \bar{x}_A \cdot y_A \cdot (\vec{\omega}_A, \vec{\omega}_A) = \sum_{\mathbf{I}} \bar{x}_A \cdot y_A \cdot \mu_A, \\ (\text{Proj}_a \vec{x}, \text{Proj}_a \vec{y}) &= \sum_a \bar{x}_A \cdot y_A \cdot \mu_A \end{aligned}$$

qui généralisent la formule de Parseval connue dans la théorie des séries orthogonales. On a aussi

$$(\vec{x}, \vec{y}) = \sum_I \overline{\text{val}_A \Xi} \text{val}_A H \mu_A,$$

formule qui peut être confrontée avec la suivante déjà connue

$$(\vec{x}, \vec{y}) = \int_I \overline{\Xi(A)} H(A) d\mu_A,$$

mais dont le caractère est complètement différent.

En connexion avec cela remarquons qu'on a pour μ — presque tous les lieux A

$$\Xi(A) = \lim_{N(q) \rightarrow 0} \frac{1}{\mu(q)} \int_q \Xi(A) d\mu_A$$

et

$$\Xi(A) = \lim_{N(q) \rightarrow 0} \frac{(\text{Proj}_q \overset{\rightarrow}{\omega}, \vec{x})}{\mu(q)},$$

ce qui est une conséquence d'un théorème bien connu de Lebesgue.

4. Remarquons que la théorie développée ci-dessus permet de donner un sens rigoureux à la célèbre formule intégrale de M. Dirac.

Prenons, en effet, pour espace $\mathbf{r}$ l'espace $(\mathbf{F})$ des fonctions μ -images des vecteurs et reprenons la formule de Fourier :

$$\overset{\rightarrow}{\xi} = \sum_I \xi_A \cdot \overset{\rightarrow}{\omega}_A = \sum_I \text{val}_A \Xi \overset{\rightarrow}{\omega}_A$$

qui peut être écrite sous la forme

$$(1) \quad \overset{\rightarrow}{\xi} = \sum_I \text{val}_A \Xi \overset{\rightarrow}{\omega}_A \cdot \mu_A.$$

Le quasi-vecteur $\overset{\rightarrow}{\omega}_A$ est défini par l'expression

$$(2) \quad \frac{\text{Proj}_q \overset{\rightarrow}{\omega}}{\mu(q)} \quad \text{pour tout voisinage } q \text{ de } A.$$

A (2) il correspond dans $\mathbf{r}$ l'expression

$$(3) \quad \frac{\Omega_q(x)}{\mu(q)},$$

donc une fonction du lieu variable x qui est partout égale à zéro sauf dans l'ensemble des lieux agrégés au segment q où elle possède la valeur constante $\frac{1}{\mu(q)}$ qui tend vers $+\infty$ lorsque A est un lieu ordinaire et $\mu(q) \rightarrow 0$.

Désignons la quasi-fonction (3) par $\Omega_A^*(x)$. La formule (1) devient alors

$$(4) \quad \Xi(x) = \sum_i \text{val}_A \Xi \Omega_A^*(x) \mu_A, \quad (\text{pour presque tout } x),$$

ce qui ressemble à la formule célèbre de Dirac

$$f(x) = \int \delta(x-y) f(y) dy.$$

Lorsque $\Xi(x)$ est continue, on peut remplacer $\text{val}_A \Xi$ par $\Xi(A)$ de manière qu'on obtient

$$\Xi(x) = \sum_i \Xi(A) \Omega_A^*(x) \mu_A.$$

Mais la formule (4) est plus générale et s'applique même aux espaces $\mathbf{1}'$ à un nombre fini de dimensions.

5. Les quasi-vecteurs et quasi-nombres sont des êtres abstraits et en apparence ils sont loin des êtres composés de nombres. Mais on peut les rendre plus intuitifs, si on les considère dans l'espace $(\mathbf{F})$ des fonctions μ -images. Les vecteurs $\vec{f}$ de $\mathbf{1}' = (\mathbf{F})$ sont des fonctions aux carrés sommables. Un espace e de la tribu $(\mathbf{T})$ se compose de toutes les fonctions de $(\mathbf{F})$ qui sont presque partout égales à zéro en dehors d'un ensemble mesurable $\mathcal{E}$ dont le support est identique à e . Le vecteur

$$\text{Proje} \vec{f}$$

c'est la fonction $f_{\mathcal{E}}(A)$ presque partout égale à $f(A)$ sur $\mathcal{E}$ et presque partout égale à zéro dans le complémentaire $I - \mathcal{E}$ de $\mathcal{E}$.

La A -composante de $f(x)$ c'est donc la *partie* de $f(x)$ *infiniment petite* en A , la A -coordonnée-densité de $f(x)$ c'est la valeur en A de cette partie infiniment petite, donc quelque sorte de valeur moyenne locale en A . La formule $\vec{f} = \sum_i \vec{f}_A$ exprime que la fonction totale $f(A)$

s'obtient par la superposition de ses parties infiniment petites. Le vecteur générateur $\vec{\omega}$ c'est la fonction constante $\omega(x) = 1$ presque partout et, $\vec{\omega}_A$ désigne sa partie *infiniment petite* en A. Ces parties ce sont des *vecteurs unitaires* posés sur les *axes* A du système des coordonnées, deux *vecteurs unitaires* $\vec{\omega}_A$ et $\vec{\omega}_B$ pour $A \neq B$ sont *orthogonaux*; en effet, les parties infiniment petites de $\omega(x)$ en A et en B sont disjointes et leur produit scalaire est $= 0$.

Remarquons que si l'on envisage les images numériques pour les lieux, on a affaire à des fonctions ordinaires, donc encore plus aisées à étudier.

6. On peut développer une théorie de l'intégration double (resp. multiple) basée sur la considération de couples de lieux. On peut ainsi obtenir la formule

$$(\vec{x}, \vec{y}) = \sum_I (\vec{x}_A, \vec{y}_B) = \sum_I \bar{x}_A \cdot y_B \cdot (\vec{\omega}_A, \vec{\omega}_B).$$

Cette théorie permet aussi de donner des formules simples pour un changement du système des coordonnées.

Les détails seront publiés dans un travail spécial traitant la totalité du sujet.