

SÉMINAIRE DE PROBABILITÉS (STRASBOURG)

PAUL-ANDRÉ MEYER

Éléments de probabilités quantiques (exposés VI à VIII)

Séminaire de probabilités (Strasbourg), tome 21 (1987), p. 34-78

<http://www.numdam.org/item?id=SPS_1987__21__34_0>

© Springer-Verlag, Berlin Heidelberg New York, 1987, tous droits réservés.

L'accès aux archives du séminaire de probabilités (Strasbourg) (<http://portail.mathdoc.fr/SemProba/>) 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/>

ELEMENTS DE PROBABILITES QUANTIQUES. VI
compléments aux exposés IV-V

Cet exposé reprend certains points de la théorie du calcul stochastique sur l'espace de Fock, en développant (parfois en corrigeant) ce qui a été dit dans les exposés du volume XX. Il est donc nécessairement un peu disparate. Bien que nous soyons forcés de renvoyer souvent au texte déjà rédigé, nous avons tenté d'écrire un exposé lisible, en rappelant des formules ou des motivations qui permettent de suivre les idées.

I. RAPPELS SUR LES DEVELOPPEMENTS EN CHAOS DE WIENER

1. L'idée fondamentale des exposés IV-V est la suivante : l'espace de Fock (symétrique ou antisymétrique) construit sur l'espace de Hilbert $L^2(\mathbb{R}_+)$ n'est rien d'autre que l'espace $L^2(P)$ associé à la mesure de Wiener, ou à l'une des mesures de Poisson. Pour fixer les idées, nous travaillerons dans l'interprétation brownienne.

Soit donc (X_t) un mouvement brownien réel issu de 0, réalisé canoniquement sur l'espace de Wiener $(\Omega, \mathcal{F}, P)$. Une v.a. appartenant à $\mathcal{H} = L^2_{\mathcal{F}}(P)$ admet un développement en chaos de Wiener, que l'on peut écrire soit en << notation courte >> , soit en << notation longue >> .

La seconde est plus familière : soit ρ_n le n-èdre croissant de $\mathbb{R}_+^n$ (points $(s_1, \dots, s_n)$ avec $s_1 < \dots < s_n$) muni de la mesure de Lebesgue $ds_1 \dots ds_n$; pour $n=0$, ρ_0 a un seul élément ($\emptyset$) et la mesure est la masse unité. On pose alors

$$(1) \quad f = \sum_n J_n(\hat{f}_n) = \sum_n \int_{\rho_n} \hat{f}_n(s_1, \dots, s_n) dX_{s_1} \dots dX_{s_n}$$

et l'on a

$$(2) \quad \|f\|^2 = \sum_n \int_{\rho_n} |\hat{f}_n(s_1, \dots, s_n)|^2 ds_1 \dots ds_n .$$

La notation courte consiste à identifier ρ_n à l'ensemble des parties à n éléments de $\mathbb{R}_+$, à noter ρ l'ensemble de toutes les parties finies de $\mathbb{R}_+$ (la réunion des ρ_n) et à écrire

$$(1') \quad f = \int_{\rho} \hat{f}(A) dX_A \quad , \quad \|f\|^2 = \int_{\rho} |\hat{f}(A)|^2 dA .$$

La notation $\hat{f}$ suggère une sorte d'analyse harmonique de la v.a. f (ceci est justifié par la théorie du << bébé Fock >> : exposé IV, p. IV.10,

où P est l'ensemble des parties de $\{1, \dots, N\}$ muni de sa structure naturelle de groupe compact). Il est parfois intéressant de remplacer $L^2(\Omega) = \mathbb{H}$ par $L^2_{\mathbb{H}}(\Omega) = \mathbb{H} \otimes \mathbb{H}$, où $\mathbb{H}$ est un espace de Hilbert arbitraire. Alors les fonctions $\hat{f}_n, \hat{f}$ sont à valeurs dans $\mathbb{H}$ au lieu d'être complexes, c'est la seule différence ($\mathbb{H}$ est l'espace de Hilbert initial au sens de Hudson-Parthasarathy, mentionné p. V.15).

Prenons par exemple $\mathbb{H} = L^2(\mathbb{R}_+)$. Une v.a. à valeurs dans $\mathbb{H}$ s'interprète comme une famille $F = (f_t)$ de v.a. complexes, telle que

$$(3) \quad \|F\|^2 = \int \|f_t\|^2 dt < \infty.$$

On peut considérer (f_t) comme un processus (non nécessairement adapté) de carré intégrable, à condition d'élargir un peu la définition usuelle : on identifie deux processus (f_t) et (g_t) tels que $f_t = g_t$ p.s. pour presque tout t (et non tout t comme d'habitude). Plus généralement un processus de carré intégrable à valeurs dans un Hilbert $\mathbb{H}$ s'identifie à une v.a. à valeurs dans $L^2(\mathbb{R}_+) \otimes \mathbb{H}$.

Du point de vue de la représentation en chaos, $f_t = \int_P \hat{f}_t(A) dX_A$, le processus est adapté si et seulement si les noyaux $\hat{f}_t$ peuvent être choisis de telle sorte que

$$(4) \quad \hat{f}_t(A) = 0 \text{ si } A \not\subset [0, t].$$

Nous avons indiqué dans l'exposé V (note au bas de la page V.25) comment se calcule l'intégrale stochastique $g = \int f_t dX_t$ d'un processus adapté de carré intégrable : on a

$$(5) \quad \hat{g}(A) = \hat{f}_{\vee A}(A \setminus \{\vee A\}) \text{ où } \vee A \text{ est le dernier élément de } A.$$

2. Gradient, divergence, intégrale de Skorokhod.

Le sujet présenté dans cette section n'est pas très important pour le calcul stochastique non commutatif, mais il était à la mode cette année ! La présentation rapide que nous en donnons ici résulte de discussions avec Yan Jia-An.

Rappelons que si f est une v.a., on désigne par $\nabla_u f$, pour $f \in L^2_{\mathbb{H}}(\mathbb{R}_+)$, la dérivée de f suivant la fonction de Cameron-Martin $u = \int_0^\cdot u_s ds$, et que l'on a $\nabla_u f = a_u^- f$, l'opérateur d'annihilation (cf. p. IV.18, formule (10)). Celui-ci est étendu aux fonctions u complexes, de manière anti-linéaire (mais la dérivée suivant une fonction complexe n'a pas de sens).

On dit que la v.a. f admet un gradient si

$$1) \nabla_u f \text{ existe pour tout } u \in L^2_{\mathbb{H}}(\mathbb{R}_+)$$

$$2) \nabla_u f \text{ peut s'écrire } \int u(s) g_s ds, \text{ où } (g_t) \text{ est un processus de carré intégrable } G, \text{ appelé le gradient de } f. \text{ Cela signifie encore que}$$

l'opérateur linéaire $u \mapsto \nabla_u f$ est donné par un noyau de carré intégrable

$$\nabla_u f(\omega) = \int u(t) G(t, \omega) dt$$

autrement dit, est du type de Hilbert-Schmidt.

Nous rappelons maintenant deux formules, qui donnent explicitement les opérateurs de création et d'annihilation (p. IV.12, formules (35))

$$(6) \quad (\hat{a}_u^- f)(A) = \int \bar{u}(t) \hat{f}(A \cup \{t\}) dt, \quad (\hat{a}_u^+ f)(A) = \sum_{t \in A} u(t) \hat{f}(A \setminus \{t\}).$$

La première formule montre que, si le gradient $(g_t) = G$ de f existe, il est donné par la formule

$$(7) \quad \hat{g}_t(A) = \hat{f}(A \cup \{t\}).$$

Pour transformer cette remarque en un raisonnement correct, on commence par le cas où $f = J_n(\hat{f}_n)$ appartient au n -ième chaos. Alors on vérifie sans peine que le gradient de f existe, qu'il est donné par (7) (tous les g_t appartenant au $n-1$ -ième chaos) et que l'on a $\|G\|^2 = n \|f\|^2$. Utilisant alors le fait que grad est un opérateur fermé, on peut montrer

- que le domaine du gradient est l'ensemble des $f = \sum_n J_n(\hat{f}_n)$ tels que $\sum_n n \|\hat{f}_n\|^2 < \infty$ (c'est aussi le domaine de l'opérateur $\sqrt{N}$, où N est l'opérateur nombre de particules) ;
- que $\|G\|^2 = \sum_n n \|\hat{f}_n\|^2$;
- que le gradient est donné par (7) sur tout son domaine.

Ceci s'étend aux v.a. à valeurs dans un Hilbert $\mathcal{H}$, le gradient étant un processus à valeurs dans $\mathcal{H}$ (ou une v.a. à valeurs dans $L^2(\mathbb{R}_+) \otimes \mathcal{H}$, ce qui permet d'itérer l'opération grad).

On appelle divergence l'opérateur adjoint de l'opérateur gradient. C'est donc un opérateur qui va des processus vers les v.a.. Au vu de la seconde formule (6), un bon candidat pour la divergence d'un processus $K = (k_t)$ est la v.a. j donnée par

$$(8) \quad \hat{j}(A) = \sum_{t \in A} \hat{k}_t(A \setminus \{t\}).$$

Et en effet, si l'on prend pour (k_t) un processus dont toutes les v.a. appartiennent au $n-1$ -ième chaos, on obtient une v.a. du n -ième chaos, et il est facile de vérifier que l'opérateur ainsi défini (les ordres des chaos restant fixés) est borné, et qu'il est bien l'adjoint de grad restreint au n -ième chaos. Il faut encore un peu de travail pour montrer que l'opérateur div défini par (7) (son domaine étant l'ensemble des processus K tels que (7) définisse une v.a. de L^2) est exactement l'adjoint de grad . J'avoue n'avoir jamais vérifié les détails !

Dans la formule (8), supposons que le processus (k_t) soit adapté. Alors la somme se réduit à son dernier terme $\hat{k}_{V_A}(A \setminus V_A)$, et d'après (5) j est l'intégrale stochastique du processus K . Autrement dit, (8) est

une extension de l'intégrale stochastique aux processus non adaptés, et sous cette forme elle a été introduite par Skorokhod. La surprenante remarque que cette << intégrale >> est en réalité un opérateur différentiel sur l'espace de Wiener est due à Gaveau et Trauber (J. Funct. Anal. 46, 1982) .

3. Formes différentielles sur l'espace de Wiener.

La considération de l'opérateur grad ci-dessus nous a montré que les processus non adaptés jouent le rôle des formes différentielles d'ordre 1 en géométrie différentielle classique. Les processus non adaptés sont les v.a. à valeurs dans $L^2(\mathbb{R}_+)$. De la même manière, les v.a. à valeurs dans la n-ième puissance extérieure de $L^2(\mathbb{R}_+)$ joueront le rôle des formes différentielles d'ordre n, et le rôle de l'algèbre de toutes les formes extérieures de tous les degrés est joué par l'espace des v.a. à valeurs dans $\bigoplus_n (L^2(\mathbb{R}_+))^{\wedge n}$, c'est à dire dans l'espace de Fock antisymétrique $\mathfrak{F}^\wedge$.

Or celui-ci est isomorphe à l'espace L^2 d'un second mouvement brownien (Y_t) . Autrement dit, une forme différentielle sur l'espace de Wiener (mélange de formes de tous les ordres) est une expression du type suivant

$$(9) \quad \Theta = \int \hat{\Theta}(A, B) dX_A dY_B \quad (\text{notation courte})$$

$$(9') \quad \Theta = \sum_n \int_{t_1 < \dots < t_n} (\int \hat{\Theta}(A, t_1, \dots, t_n) dX_A) dY_{t_1} \wedge \dots \wedge dY_{t_n}$$

$$(9'') \quad \Theta = \sum_{m, n} \frac{1}{m!n!} \int \hat{\Theta}(s_1, \dots, s_m, t_1, \dots, t_n) dX_{s_1} \dots dX_{s_m} dY_{t_1} \wedge \dots \wedge dY_{t_n}$$

(notations longues)

où les symboles dX_s commutent entre eux et avec les symboles dY_t , tandis que les dY_t anticommulent. Dans la troisième écriture, la fonction $\hat{\Theta}$ est symétrique en ses m premières variables, antisymétrique en les n dernières. Nous munirons l'espace de toutes les formes différentielles de la norme hilbertienne

$$(10) \quad \begin{aligned} \|\Theta\|^2 &= \int_{P \times P} |\hat{\Theta}(A, B)|^2 dA dB \\ &= \sum_{m, n} \frac{1}{m!n!} \int_{\mathbb{R}_+^m \times \mathbb{R}_+^n} |\hat{\Theta}(s_i, t_j)|^2 ds_1 \dots ds_m dt_1 \dots dt_n \end{aligned}$$

L'expression (9) a l'air d'être une v.a. sur un espace de Wiener bidimensionnel, et non une forme différentielle ! Pour l'interpréter comme forme différentielle, indiquons sa valeur sur un n-vecteur $\underline{u}_1 \wedge \dots \wedge \underline{u}_n$ de fonctions de Cameron-Martin : c'est la fonction sur l'espace de Wiener

$$(11) \quad \int_P dX_A \left(\int_P \hat{\Theta}(A, ds_1, \dots, ds_n) u_1(s_1) \dots u_n(s_n) ds_1 \dots ds_n \right)$$

Nous allons définir ensuite la différentielle d et la codifférentielle δ . Commençons par la première : si $\Theta=f$ est de degré 0 (i.e. est une fonction) de gradient (g_t) , $d\Theta$ est la forme de degré 1 $\int g_t dY_t$. D'autre part, le calcul usuel de $d(fdY_B)$ est $df \wedge dY_B = g_t dY_t \wedge dY_B$, ce qui fait apparaître un signe alternant lorsqu'on remet en ordre croissant l'ensemble $BU\{t\}$. Cela justifie la formule suivante : si $\Theta = \int \hat{\Theta}(A,B) dX_A dY_B$,

$$(12) \quad d\Theta = \int_{P \times P} (\sum_{t \in B} (-1)^{n(t,B)} \hat{\Theta}(AU\{t\}, B \setminus \{t\})) dX_A dY_B$$

où $n(t,B)$ est le nombre d'éléments de B strictement majorés par t . De même, la codifférentielle est donnée par

$$(13) \quad \delta\Theta = \int_{P \times P} (\sum_{t \in A} (-1)^{n(t,B)} \hat{\Theta}(A \setminus \{t\}, BU\{t\})) dX_A dY_B.$$

Nous allons illustrer le calcul, et le fait que ces deux opérateurs sont adjoints l'un de l'autre, par un exemple : prenons

$$\Theta = \int_{s_1 < s_2, t_1 < t_2} f(s_1, s_2; t_1, t_2) dX_{s_1} dX_{s_2} dY_{t_1} \wedge dY_{t_2}$$

$$\Omega = \int_{u, v_1 < v_2 < v_3} g(u; v_1, v_2, v_3) dX_u dY_{v_1} \wedge dY_{v_2} \wedge dY_{v_3}$$

les fonctions étant symétriques par rapports aux premiers arguments et antisymétriques par rapport aux derniers. Alors (on omet désormais les $\wedge$)

$$d\Theta = \int_{u, v_1 < v_2 < v_3} (f(u, v_1; v_2, v_3) - f(u, v_2; v_1, v_3) + f(u, v_3; v_1, v_2)) dX_u dY_{v_1} dY_{v_2} dY_{v_3}$$

$$\delta\Omega = \int_{s_1 < s_2, t_1 < t_2} (g(s_1; s_2, t_1, t_2) + g(s_2; s_1, t_1, t_2)) dX_{s_1} dX_{s_2} dY_{t_1} dY_{t_2}$$

(l'antisymétrie de g en ses derniers arguments a absorbé les facteurs $(-1)^{n(s_i, \{t_1, t_2\})}$). On a alors

$$\begin{aligned} \langle d\Theta, \delta\Omega \rangle &= \int_{s_1 < s_2, t_1 < t_2} f(s_1, s_2; t_1, t_2) (g(s_1; s_2, t_1, t_2) + g(s_2; s_1, t_1, t_2)) ds_1 ds_2 dt_1 dt_2 \\ &= \int_{t_1 < t_2} f(s_1, s_2; t_1, t_2) g(s_1; s_2, t_1, t_2) ds_1 ds_2 dt_1 dt_2 \\ &= \frac{1}{2!} \int_{\mathbb{R}_+^4} f(s_1, s_2; t_1, t_2) g(s_1; s_2, t_1, t_2) ds_1 ds_2 dt_1 dt_2. \\ \langle \delta\Theta, \Omega \rangle &= \int_{u, v_1 < v_2 < v_3} (f(u, v_1; v_2, v_3) - f(\dots) + f(\dots)) g(u; v_1, v_2, v_3) du dv_1 dv_2 dv_3 \\ &= \frac{1}{3!} \int_{\mathbb{R}_+^4} (f(u, v_1; v_2, v_3) - f(\dots) + f(\dots)) f(u; v_1, v_2, v_3) du dv_1 dv_2 dv_3 \\ &= \frac{1}{2} \int_{\mathbb{R}_+^4} f(u, v_1; v_2, v_3) g(u; v_1, v_2, v_3) du dv_1 dv_2 dv_3. \end{aligned}$$

et on a bien l'égalité. Le principe général de vérification est le même.

Il est immédiat de vérifier que $dd = \delta\delta = 0$.

Bien que l'idée de considérer des formes différentielles sur l'espace de Wiener ait attiré de nombreux auteurs, Shigekawa a été le premier, semble-t-il, à présenter une théorie complète. Nous allons retrouver très facilement, à partir des expressions explicites (12)-(13), le calcul du << laplacien de de Rham >> $d\delta + \delta d$ obtenu par Shigekawa (une démonstration simplifiée du résultat de Shigekawa, communiquée par P. Krée, nous a aidés à mettre au point la présentation ci-dessus).

Le coefficient de $dX_A dY_B$ dans $d\delta$ et δd a la valeur suivante (nous pouvons supposer A et B disjoints, l'éventualité contraire n'étant p.s. pas réalisée). Les notations de théorie des ensembles ont été allégées !

$$d\delta : \sum_{\substack{seB+t \\ teA}} (-1)^{n(t,B)+n(s,B+t)} \hat{\Theta}(A-t+s, B+t-s)$$

$$\delta d : \sum_{\substack{seB \\ teA+s}} (-1)^{n(s,B)+n(t,B-s)} \hat{\Theta}(A+s-t, B-s+t)$$

Il est facile de voir que les couples seB , teA ont une contribution nulle dans la somme, la seule contribution non nulle provenant de $s=teA$ dans la première somme, $t=seB$ dans la seconde. Reste donc $(|A|+|B|)\hat{\Theta}(A,B)$. Ou encore, si l'on a une forme de degré n

$$\Theta = \int_{t_1 < \dots < t_n} f(\omega, \dots) dY_{t_1} \wedge \dots \wedge dY_{t_n},$$

son laplacien de de Rham vaut

$$\square \Theta = \int_{t_1 < \dots < t_n} (N+nI) f(\omega, \dots) dY_{t_1} \wedge \dots \wedge dY_{t_n}$$

l'opérateur nombre de particules agissant sur la variable ω .

On peut aller beaucoup plus loin dans la direction indiquée ici, qui est celle d'un mélange d'espaces de Fock symétriques et antisymétriques (voir par ex. un article récent d'Y. LeJan, intitulé << Temps local et superchamp >>).

II. SUR LES NOYAUX DE MAASSEN

1. Dans ce paragraphe, nous nous proposons de redonner quelques définitions de base, et ensuite d'augmenter notre collection d'exemples de noyaux. Quelques erreurs seront corrigées, et quelques compléments apportés à la théorie.

Un opérateur K associé à un noyau est formellement du type

$$(1) \quad K = \int_{P \times P \times P} K(A,B,C) da_A^+ da_B da_C^- \quad (\text{notation courte})$$

a^+, a^0, a^- sont respectivement des opérateurs de création, de comptage, et d'annihilation, et $K(A,B,C)=0$ si A,B,C ne sont pas disjoints.

La notation longue est vraiment beaucoup plus longue :

$$\Sigma_{m,n,p} \int_{\substack{r_1 < \dots < r_m \\ s_1 < \dots < s_m \\ t_1 < \dots < t_p}} K(r_1, \dots, r_m, s_1, \dots, s_m, t_1, \dots, t_p) da_{r_1}^+ \dots da_{r_m}^+ da_{s_1}^0 \dots da_{s_m}^0 da_{t_1}^- \dots da_{t_p}^-$$

L'opérateur K transforme $f = \hat{f}(A) dX_A$ en $g = Kf = \hat{g}(A) dX_A$ donnée par

$$(2) \quad \hat{g}(A) = \int_{\mathcal{P}} \Sigma_{U+V+W=A} K(U, V, M) \hat{f}(V \cup W \cup M) dM \quad (1)$$

Pour tout cela, voir l'exposé IV, p. 33-34 et l'exposé V, § III. Nous utiliserons parfois des notations allégées de théorie des ensembles (par exemple, $f(A+t) = f(A \cup \{t\})$ si $t \notin A$, 0 si $t \in A$, $f(A-t) = f(A \setminus \{t\})$ si $t \in A$, 0 si $t \notin A$).

Maassen définit les fonctions-test par une majoration du type

$$(3) \quad |\hat{f}(A)| \leq \varrho^{|A|}, \quad \hat{f}(A) = 0 \text{ si } A \not\subset [0, t[\quad (t > 0, \varrho > 0)$$

et les noyaux réguliers par une majoration du type

$$(4) \quad |K(A, B, C)| \leq \varrho^{|A|+|B|+|C|}, \quad K(A, B, C) = 0 \text{ si } A \cup B \cup C \not\subset]0, t[$$

Il montre alors (cf. p. V.21) que les noyaux réguliers transforment les fonctions-test en fonctions-test, et forment une classe d'opérateurs stable par composition et passage à l'adjoint (exposé V, p. V.23-24).

Un noyau K est s-adapté si l'on a

$$(5) \quad K(A, B, C) = 0 \text{ pour } A \cup B \cup C \not\subset [0, s[$$

Cette définition n'est pas identique à celle de la page V.24, qui est fautive (cette erreur nous a été signalée par H. Maassen). Une famille (K_s) de noyaux est dite adaptée si K_s est s-adapté pour tout s .

En analogie complète avec la formule (5) du § I (ci-dessus), donnant l'intégrale stochastique d'un processus adapté de vecteurs (f_t) , Maassen définit les intégrales stochastiques de processus adaptés de noyaux (K_t) . Les noyaux des trois intégrales stochastiques $\int K_s da_s^+$, $\int K_s da_s^0$, $\int K_s da_s^-$, sont donnés respectivement par

$$(6) \quad I(A, B, C) = K_{VA}(A - VA, B, C), \quad K_{VB}(A, B - VB, C), \quad K_{VC}(A, B, C - VC)$$

($A - VA$ pour $A \setminus \{VA\}$, etc.). Cf. p. V.25.

2. Une liste d'opérateurs donnés par des noyaux.

a) Opérateurs élémentaires. Les opérateurs de création, d'annihilation et de comptage sont donnés par des noyaux très simples (V.22).

1. Dans le cas discret ($\ll$ bébé Fock $\gg$) la sommation ne porterait que sur les M disjoints de A ; ici, cela ne change rien.

$$\begin{aligned}
 (7) \quad & \widehat{(a_u^- f)}(A) = \int \bar{u}(t) \hat{f}(A+t) dt \quad ; \quad K(\emptyset, \emptyset, \{t\}) = \bar{u}(t) \quad (0 \text{ dans les autres cas}) \\
 & \widehat{(a_u^+ f)}(A) = \sum_{t \in A} u(t) \hat{f}(A-t) \quad ; \quad K(\{t\}, \emptyset, \emptyset) = u(t) \quad \text{''''''''} \\
 & \widehat{(a_u^0 f)}(A) = (\sum_{t \in A} u(t)) \hat{f}(A) \quad ; \quad K(\emptyset, \{t\}, \emptyset) = u(t) \quad \text{''''''''}
 \end{aligned}$$

b) Opérateurs divers de multiplication (Wick, Wiener)

Nous avons étudié dans l'exposé IV (p.IV.27) la formule de multiplication des intégrales stochastiques multiples (dans l'interprétation brownienne). Soit $h=fg$ le produit de Wiener ; on a

$$(8) \quad \hat{h}(A) = \int_{\rho} \sum_{U+V=A} \hat{f}(U) \hat{g}(V) dM$$

Cette formule est une réécriture de la formule qui donne le produit des "éléments différentiels" eux mêmes

$$(9) \quad dX_A dX_B = dX_{A \Delta B} d(A \cap B)$$

qui exprime simplement que dans un produit $dX_{s_1} \dots dX_{s_m} dX_{t_1} \dots dX_{t_n}$ (avec $s_1 < \dots < s_m, t_1 < \dots < t_n$) on remplace tout terme carré dX_s^2 (correspondant à l'égalité entre un s_i et un t_j) par ds . La formule (8) exprime aussi que le noyau de l'opérateur de multiplication de Wiener par f est donné par

$$(10) \quad K(A, \emptyset, C) = \hat{f}(A \cup C) \quad (0 \text{ si l'argument médian est } \neq \emptyset).$$

Au lieu d'adopter la règle $dX_s^2 = ds$, adoptons la règle $dX_s^2 = \theta ds$, où θ est un nombre (si $\theta = \sigma^2 > 0$, c'est la règle de multiplication correspondant au mouvement brownien de variance $\sigma^2 t$). Un cas intéressant est celui où $\theta = 0$: la multiplication est alors le produit de Wick.

Dans ce cas, la formule (9) est modifiée par l'apparition d'un facteur $\theta^{|A \cap B|}$ au second membre, et la formule (8) par l'apparition d'un facteur $\theta^{|M|}$; quant au noyau (10), il est remplacé par

$$(10') \quad K(A, \emptyset, C) = \hat{f}(A \cup C) \theta^{|C|}$$

En particulier, l'expression du produit de Wick $h=fg$ est

$$\hat{h}(A) = \sum_{U+V=A} \hat{f}(U) \hat{g}(V) \quad .$$

Parmi tous ces opérateurs, les opérateurs de multiplication de Wiener et de Wick ($\theta = 1$ ou 0) jouissent d'une propriété spéciale : la multiplication par un élément réel f définit un opérateur autoadjoint (ou du moins symétrique !) pour la structure hilbertienne donnée sur l'espace de Fock.

Exercice : En utilisant la formule (2), calculer le produit de deux vecteurs exponentiels $f=e(\lambda), g=e(\mu)$ (autrement dit

$$\hat{f}(A) = \prod_{s \in A} \lambda(s) \text{ , et de même pour } g \text{)}$$

et retrouver le résultat $e(\lambda)e(\mu) = e(\lambda+\mu) \exp(\theta \int (r) (r) dr)$ connu pour les produits de Wiener et de Wick .

c) Produit de Poisson.

L'espace de Fock admet une interprétation probabiliste (IV, p. 14) dans laquelle $X_t = (N_t - ct)/\sqrt{c}$, où N_t est un processus de Poisson d'intensité c . Dans ce cas, on a la règle suivante de multiplication

$$(11) \quad dX_s^2 = ds + \rho dX_s \quad \text{avec} \quad \rho = 1/\sqrt{c}$$

On retrouve le produit de Wiener lorsque $\rho \rightarrow 0$.

Surgailis a publié des formules de multiplication de Poisson, mais écrites sous une forme assez compliquée. En réalité, si l'on part de (11), que l'on calcule $dX_A dX_B$ à la manière de (9), puis que l'on interprète le résultat comme provenant d'un noyau, on obtient un résultat simple : le produit de Poisson $h = fg$ est donné par

$$(12) \quad \hat{h}(A) = \int \Sigma_{U+V+W=A} \hat{f}(UVUVM) \hat{g}(VUWUM) \rho^{|V|} dM$$

et le noyau de la multiplication de Poisson par f est

$$(13) \quad K(A, B, C) = \rho^{|B|} \hat{f}(A \cup B \cup C) .$$

Exercice. Calculer directement grâce à (13) le produit de Poisson $\mathcal{E}(\lambda)\mathcal{E}(\mu)$ de deux vecteurs exponentiels, qui vaut $\exp(\int \lambda(r)\mu(r)dr)\mathcal{E}(\lambda+\mu+\rho\lambda\mu)$ d'après la formule explicite de l'exponentielle de Poisson (p. IV.22).

Remarque. Si f est une fonction-test, il est clair que le noyau (13) est régulier ; donc le produit de Poisson de deux fonctions-test est une fonction-test, ce qui semble contredire la remarque page IV.35 (suivant le théorème relatif au produit de Wiener)... à vrai dire, il n'y a pas de contradiction, car le mot « fonction-test » est pris en deux sens différents aux deux endroits : le sens de la page IV.34-35 est moins intéressant que celui de Maassen.

(Incidentement, en haut de la page IV.35, lire $(p-1)^{n/2}$ au lieu de $(p-1)^n$, et ajouter que $p \geq 2$).

d) Opérateurs de Weyl

Si f est un vecteur de $L^2(\mathbb{R}_+)$, U un opérateur unitaire sur $L^2(\mathbb{R}_+)$, l'opérateur de Weyl $W = W_{f,U}$ est défini par

$$(14) \quad W\mathcal{E}(h) = \exp(-\|f\|^2/2 - \langle f, Uh \rangle) \mathcal{E}(f + Uh)$$

Prenons en particulier pour U l'opérateur de multiplication par $e^{i\lambda}$, où $\lambda(t)$ est une fonction réelle. Le noyau de W est donné par

$$(15) \quad K(A, B, C) = e^{-\|f\|^2/2} \prod_{s \in A} f(s) \prod_{s \in B} (e^{i\lambda(s)} - 1) \prod_{s \in C} (-e^{i\lambda(s)} \bar{f}(s))$$

Nous allons donner de ce fait une démonstration plus simple que celle de la page V.22, consistant tout juste à appliquer le noyau (15) à $\mathcal{E}(h)$, et à vérifier que l'on obtient (14). Ce calcul est très proche des deux "exercices" proposés plus haut, mais nous le traiterons en détail.

Dans l'expression $K\hat{\mathcal{E}}(h)(A) = \int \Sigma_{R+S+T=A} K(R,S,M) \prod_{u \in S \cup T \cup M} h(u)$, où nous remplaçons K par sa valeur (15), nous trouvons d'abord en tête le facteur $\exp(-\|f\|^2/2)$. Ensuite, un facteur qui ne contient pas M

$$\begin{aligned} \Sigma_{R+S+T=A} \prod_{r \in R} f(r) \prod_{s \in S} (e^{i\lambda(s)} - 1)h(s) \prod_{t \in T} h(t) \\ = \prod_{u \in A} (f(u) + (e^{i\lambda(u)} - 1)h(u) + h(u)) \end{aligned}$$

et enfin, l'intégrale en facteur

$$\int dM \prod_{u \in M} (-e^{i\lambda(u)} \bar{f}(u)h(u))$$

Or remarquons que pour toute fonction g , $\int \prod_{u \in M} g(u) dM$ vaut $(\int g du)^n/n!$, donc la somme sur tous les n vaut $\exp(\int g(u)du)$, et ici l'on trouve bien $\exp(-\langle f, e^{i\lambda} h \rangle)$. Il ne reste plus qu'à remarquer que le second facteur était le coefficient de A dans le développement en chaos de $\mathcal{E}(f + (e^{i\lambda} - 1)h + h) = \mathcal{E}(f + e^{i\lambda} h)$.

Exercice. On laisse fixes f et λ , et on désigne par W_t l'opérateur de Weyl relatif aux fonctions $fI_{[0,t]}$ et $\lambda I_{[0,t]}$. Connaissant le noyau de W_t (15), vérifier en utilisant (6) que les noyaux W_t satisfont à l'équation différentielle stochastique

$$(16) \quad dW_t = W_t(f(t)da_t^+ - \bar{f}(t)e^{i\lambda(t)}da_t^- + (e^{i\lambda(t)} - 1)da_t^0 - \frac{1}{2}\|fI_{[0,t]}\|^2 dt)$$

Manière de procéder : Il est plus simple de rechercher l'e.d.s. satisfait par $\exp(\frac{1}{2}\|fI_{[0,t]}\|^2)W_t$, dont le noyau K_t est de la forme

$$K_t(A,B,C) = \prod_{s \in A} u_t(s) \prod_{s \in B} v_t(s) \prod_{s \in C} w_t(s)$$

où u,v,w sont trois fonctions sur $\mathbb{H}_+$, et $u_t = uI_{[0,t]}$ par ex.. On calcule alors $\int_0^\infty u_s K_s da_s^+$ par la formule (15), et l'on trouve que ce noyau vaut $K_\infty(A,B,C)I_{\{v(BUC) < vA\}}$. En ajoutant les deux autres noyaux analogues, on trouve K_∞ .

e) Noyaux dépendant du second argument seulement. Considérons un noyau de la forme $K(\emptyset, B, \emptyset) = k(B)$, $K(A, B, C) = 0$ dans les autres cas. On a alors

$$(17) \quad (\hat{K}f)(A) = j(A)\hat{f}(A) \quad \text{avec} \quad j(A) = \Sigma_{B \subset A} k(B)$$

Si l'on interprète $\hat{f}(\cdot)$ comme une sorte de transformée de Fourier de f (ce qui est exactement le cas sur le << bébé Fock >>, ces noyaux apparaissent comme des sortes de convolutions. Partant du << multiplicateur >> j , on obtient le noyau k par la formule d'inversion de Moebius

$$(18) \quad k(A) = \Sigma_{B \subset A} (-1)^{|A-B|} j(B).$$

On peut remarquer qu'une inégalité $|j(A)| \leq M^{|A|}$ (fonctions-test de Maassen) entraîne $|k(B)| \leq (M+1)^{|B|}$ (noyau régulier), et de même en sens

inverse.

L'un des plus remarquables parmi ces opérateurs est l'opérateur $(-1)^{Nt} = J_t$ introduit par Hudson-Parthasarathy, et défini à la dernière page de l'exposé V. Son noyau est

$$(19) \quad J_t(\emptyset, B, \emptyset) = j_t(B) = (-2)^{|B|} \quad \text{si } B \subset [0, t[, 0 \text{ sinon .}$$

et l'on a

$$(20) \quad (\hat{J}_t f)(A) = (-1)^{|A \cap [0, t[|} \hat{f}(A)$$

Rappelons que les opérateurs $\beta_u^+ = \int u_s J_s da_s^+$, $\beta_u^- = \int \bar{u}_s J_s da_s^-$, sont des opérateurs de fermions. Leurs noyaux sont

$$(21) \quad \beta_u^+(\{s\}, B, \emptyset) = u(s)(-2)^{|B|}, \quad \beta_u^-(\emptyset, B, \{s\}) = \bar{u}(s)(-2)^{|B|}$$

et 0 dans les autres cas. Leur effet sur f est

$$(22) \quad \begin{aligned} (\hat{\beta}_u^+ f)(A) &= \sum_{t \in A} (-1)^{n(t, A)} u(t) \hat{f}(A - t) \\ (\hat{\beta}_u^- f)(A) &= \int (-1)^{n(s, A)} \hat{f}(A + s) \bar{u}(s) ds \end{aligned}$$

f) Opérateurs de multiplication de Clifford.

Le produit de Clifford $h = f * g$ de deux v.a. f et g a été défini p. IV.30. Il est donné par la formule

$$(23) \quad \hat{h}(A) = \int \sum_{U+V=A} \hat{f}(MUU) \hat{g}(NUV) (-1)^{n(MUU, MUV)} dM$$

Comme pour les produits de Wiener de paramètre Θ , on obtiendrait d'autres produits associatifs en insérant un facteur $\Theta^{|M|}$ devant dM . Pour $\Theta = 0$, on obtient le produit extérieur (produit de Grassmann) $f \wedge g$

$$(24) \quad \hat{h}(A) = \sum_{U+V=A} (-1)^{n(U, V)} \hat{f}(U) \hat{g}(V).$$

Nous avons affirmé à tort p. V.26 que ces opérateurs de multiplication par $f \ll$ ne semblaient pas $\gg$ être donnés par des noyaux. Nous allons ici déterminer leurs noyaux.

Nous commençons par définir, pour K, L fixés, une fonction d'ensemble $a_{K, L}(H)$ possédant la propriété suivante

$$(25) \quad \sum_{A \subset H} a_{K, L}(A) = (-1)^{n(K \cup L, H \cup L)}$$

Ceci est toujours possible, en vertu de la formule d'inversion de Moebius (en fait, la détermination explicite du noyau reviendra au calcul de cette fonction, dont nous reparlerons plus bas). Posons alors

$$(26) \quad K(A, B, C) = \hat{f}(A \cup C) \Theta^{|C|} a_{A, C}(B)$$

et calculons la fonction $h = Kg$:

$$\hat{h}(A) = \int \sum_{U+V+W=A} \hat{f}(U \cup M) a_{U, M}(V) \hat{g}(V \cup W \cup M) \Theta^{|M|} dM$$

Posons $V+W=Z$, et sommons en laissant U, M (et donc Z) fixes : V parcourt l'ensemble de toutes les parties de Z , et d'après (25) nous trouvons

$$\hat{h}(A) = \int \Sigma_{U+Z=A} \hat{f}(U, M) (-1)^{n(U, M, Z, M)} \hat{g}(Z, M) \Theta^{|M|} dM$$

ce qui est précisément (23) (ou (24) si $\Theta=0$; dans ce cas, il suffit de calculer $a_{K, \emptyset}(\cdot)$).

On voit donc que le produit de Clifford est effectivement donné par un noyau à trois arguments. Nous allons maintenant préciser le calcul de celui-ci.

Nous commençons par remarquer que dans la formule (25), nous nous intéressons uniquement au cas où H, K, L sont disjoints. Alors nous pouvons écrire $(-1)^{n(K \cup L, H \cup L)} = (-1)^{n(K \cup L, H)} (-1)^{n(K \cup L, L)}$, et le second facteur ne dépend pas de H . Autrement dit, nous pouvons poser

$$(27) \quad a_{K, L}(H) = (-1)^{n(K \cup L, L)} j_{K \cup L}(H)$$

où la fonction j est déterminée par

$$(28) \quad \Sigma_{A \subset H} j_B(A) = (-1)^{n(B, H)} \quad (B, H \text{ disjoints})$$

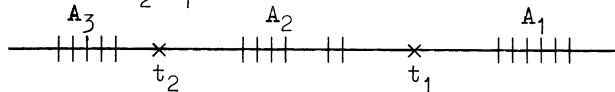
C'est la même fonction qui intervient dans les produits de Clifford et de Grassmann, celui-ci n'étant pas plus simple. Elle nous est donnée par la formule d'inversion de Moebius

$$(29) \quad j_B(A) = \Sigma_{H \subset A} (-1)^{|A-H|} (-1)^{n(B, H)} \quad (A, B \text{ disjoints})$$

Pour calculer cela, désignons par $t_n < t_{n-1} < \dots < t_1$ les points de B , et par $A_1, A_2, \dots, A_{n+1}$ les parties de A situées respectivement dans les intervalles $]t_1, \infty[$, $]t_2, t_1[$, $\dots$, $]t_n, t_{n-1}[$, $] -\infty, t_n[$ (attention, les indices croissent lorsqu'on se déplace de droite à gauche). Désignons par A^- la réunion $A_1 \cup A_3 \cup \dots$ des A_i impairs ; on a

$$(30) \quad j_B(A) = 0 \text{ si } A^- \neq \emptyset, \quad j_B(A) = (-2)^{|A|} \text{ si } A^- = \emptyset.$$

La raison de cette règle bizarre apparaîtra clairement sur le cas où B admet deux éléments $t_2 < t_1$:



Une partie arbitraire de A s'écrit $H = \cup H_i$ avec $H_i \subset A_i$. On a alors $|A-H| = \Sigma_i |A_i - H_i|$, et $n(B, H) = n(t_2, H) + n(t_1, H) = |H_3| + |H_3| + |H_2|$, de sorte que $|H_3|$ s'élimine de $(-1)^{n(B, H)}$. Reste donc

$$\Sigma_{H_1 \subset A_1, H_2 \subset A_2, H_3 \subset A_3} (-1)^{|A_1 - H_1|} (-1)^{|A_3 - H_3|} (-1)^{|A_2|}$$

Si $A_1 \neq \emptyset$ ou $A_3 \neq \emptyset$ on trouve 0, et si ces ensembles sont vides ($A=A_2$) on trouve simplement $2^{|A_2|} (-1)^{|A_2|}$.

3. Unicité du noyau associé à un opérateur.

H. Maassen nous a fait remarquer qu'il n'est nullement évident, sur la formule (2), que le noyau associé à un opérateur donné soit unique. Nous allons esquisser ici une démonstration de ce fait.

Nous désignons donc par K un noyau qui représente l'opérateur nul, et nous prouvons que $K(A,B,C)=0$ dAdBdC-p.p.. Afin de nous épargner toutes les difficultés d'intégrabilité, nous supposons par exemple que K est régulier, mais l'hypothèse n'est aucunement nécessaire.

a) Nous allons commencer par évaluer, en fonction du noyau K , la forme bilinéaire $\langle g, Kf \rangle$ (pour simplifier, nous nous plaçons dans le cas réel). D'après (2), cette forme est égale à

$$\int \sum_{A+B+C=H} \hat{g}(H) \hat{f}(B+C+M) K(A,B,M) dM$$

Nous transformons cette expression au moyen du lemme de l'exposé V, p. V.23 : ce lemme nous dit que

$$\int \sum_{A+B+C=H} u(A,B,C) dH = \int u(A,B,C) dAdBdC$$

d'où l'on tire

$$\begin{aligned} (31) \quad \langle g, Kf \rangle &= \int \hat{g}(A+B+C) \hat{f}(M+B+C) K(A,B,M) dAdBdCdM \\ &= \int \hat{g}(A+L) \hat{f}(M+L) \sum_{B+C=L} K(A,B,L) dAdBdCdM \\ &= \int \hat{g}(A+L) \hat{f}(M+L) \phi(A,L,M) dAdLdM \end{aligned}$$

où l'on a posé $\phi(A,L,M) = \sum_{B \subset L} K(A,B,M)$. D'après la formule d'inversion de Moebius, il est équivalent de montrer que $K=0$ p.p. ou que $\phi=0$ p.p..

b) Nous allons établir cela maintenant. Nous allons supposer g nulle

hors du chaos d'ordre n , f nulle hors du chaos d'ordre m , et identifier une fonction $j(s_1, \dots, s_k)$ sur $\mathcal{P}_k$ à la restriction d'une fonction symétrique sur $\mathbb{E}_+^k$, désignée par la même lettre. Plutôt que d'utiliser des notations générales lourdes, écrivons ce que signifie la nullité de (31) dans le cas où $n=3$, $m=2$:

$$\begin{aligned} &\frac{1}{3!2!} \int g(r,s,t) f(u,v) \phi(r,s,t;\emptyset;u,v) drdsdtdu dv \quad (|A|=3, |L|=0, |M|=2) \\ &+ \frac{1}{2!} \int g(r,s,t) f(t,u) \phi(r,s,t;u) drdsdtdu \quad (|A|=2, |L|=1, |M|=1) \\ &+ \frac{1}{2!} \int g(r,s,t) f(s,t) \phi(r,s,t;\emptyset) drdsdt \quad (|A|=1, |L|=2, |M|=0) \\ &= 0 \end{aligned}$$

quelles soient les fonctions g (symétrique), f (symétrique).

Désignons par α la mesure $\phi(r,s,t;\emptyset;u,v) drdsdtdu dv$; par β la mesure image de $\phi(r,s,t;u) drdsdtdu$ par $(r,s,t,u) \mapsto (r,s,t,t,u)$, et

par β' sa symétrisée relativement aux trois premières variables, et aux deux dernières variables ; par γ la mesure image de $\phi(r;s,t;\emptyset)drdsdt$ par $(r,s,t) \mapsto (r,s,t,s,t)$ et par γ' sa symétrisée comme ci-dessus. On a $\alpha/3! + \beta' + \gamma' = 0$; or ces trois mesures sont absolument continues sur des réunions finies de variétés affines de dimensions différentes. Donc elles sont nulles toutes trois. La mesure symétrique β' est de même une somme finie de mesures absolument continues sur des variétés affines de dimension 4 dans $\mathbb{R}^{3+2}$ (chacune desquelles s'obtient en égalant l'une des trois premières coordonnées à l'une des deux dernières) ; la relation $\beta'=0$ entraîne la nullité de la restriction de β' à chacune de ces variétés affines. Or la seule permutation des trois premières variables (r,s,t) et des deux dernières variables (u,v) qui préserve la variété $\{t=u\}$ est l'échange de r et s , et $\phi(r,s;t;u)$ est symétrique en (r,s) : donc $\beta=0$, et de même $\gamma=0$.

Il semble qu'il n'y ait aucune difficulté, autre que de notation, à étendre ce raisonnement à tous les couples (n,m) , et à en déduire la nullité p.p. de $\phi(A,L,M)$ pour $|A+L|=n$, $|M+L|=m$.

La formule (31) semble aussi présenter un certain intérêt propre.

III. CONSTRUCTION DE LOIS I.D.

1. Si l'on relit l'exposé III, p.III.12-15, à la lumière des exposés

IV-V, on peut l'interpréter de la manière suivante : sur l'espace de Fock le plus simple (construit sur un espace de Hilbert de dimension 1), on sait construire des opérateurs autoadjoints admettant (dans l'état vide) une loi normale ou une loi de Poisson. Comme toute loi indéfiniment divisible est une somme continue de lois du type précédent, il est naturel de chercher à construire, sur un espace de Fock convenable, un opérateur autoadjoint admettant une loi i.d. donnée, par un procédé qui généralise raisonnablement celui de l'exposé III. Il n'est pas inattendu que l'espace de Fock utilisé soit construit sur $L^2(\lambda)$, où λ est la mesure de Lévy de la loi i.d..

Ce paragraphe est entièrement emprunté à l'article écrit par Parthasarathy pour l'Encyclopédie Italienne.

2. Nous commençons par quelques généralités. Soit H un espace de

Hilbert complexe, et soit $\mathfrak{F}$ l'espace de Fock construit sur H . On se donne un groupe à un paramètre (θ_t) (fortement continu) d'opérateurs unitaires sur H . On se propose de lui associer un groupe à un paramètre $(\tilde{\theta}_t)$ d'opérateurs unitaires de $\mathfrak{F}$. Une solution triviale consiste à prendre $\tilde{\theta}_t = W(0, \theta_t)$ (où $W(f,U)$ désigne l'opérateur de

Weyl correspondant à la translation f et à la rotation (opérateur unitaire) U : page IV.8). Nous chercherons des solutions non triviales, de la forme

$$(1) \quad \tilde{\theta}_t = e^{ib_t} W(u_t, \theta_t)$$

où b_t est un scalaire réel, et où (u_t) est une hélice dans l'espace de Hilbert H (on dit aussi un cocycle)

$$(2) \quad u_0 = 0, \quad u_{t+s} = u_s + \theta_s u_t$$

Cette relation rappelle la théorie des fonctionnelles additives de Markov. Compte tenu des relations de Weyl et de (2), on voit que les opérateurs (1) forment un groupe unitaire dans $\mathfrak{F}$ dès lors que

$$(3) \quad b_{s+t} - b_s - b_t = -\text{Im} \langle u_s, \theta_s u_t \rangle.$$

Exemple : On peut aussitôt indiquer deux types élémentaires de cocycles.

a) $u_t = tu$, où u est un vecteur de H invariant par le groupe (θ_t)

b) $u_t = (\theta_t - I)u$, où u est un vecteur arbitraire de H .

Dans le premier cas, on prendra $b_t = ct$, et dans le second cas

$$(4) \quad b_t = \text{Im} \langle u, \theta_t - I \rangle + ct.$$

2. Soit λ une mesure de Lévy sur $\mathbb{T}$: une mesure positive, non nécessairement bornée au voisinage de 0, telle que la fonction $x^2 \lambda$ soit λ -intégrable. On cherche à construire une v.a. quantique de transformée de Fourier e^ϕ , où l'on a posé

$$(5) \quad \phi(y) = \text{imy} - \frac{1}{2} \lambda(0) y^2 + \int_{|x| > 1} (e^{iyx} - 1) \lambda(dx) + \int_{0 < |x| \leq 1} (e^{iyx} - 1 - iyx) \lambda(dx).$$

A cet effet, on prend comme espace de Hilbert H l'espace $L^2(\lambda)$, pour $\mathfrak{F}$ l'espace de Fock sur H comme plus haut. On pose $\theta_t u(x) = e^{itx} f(x)$ (ce sera le groupe unitaire). On choisit une partition E_n de la droite en ensembles λ -intégrables, tels que

$$(6) \quad E_0 = \{0\}, \quad E_1 = \{|x| > 1\} \quad (\text{les autres } E_i \text{ dans } [-1, 0[\cup]0, 1]).$$

Pour chaque n , soit h_n une fonction de module 1 sur E_n ; une telle fonction appartient à H . Nous posons

$$(7) \quad u_t(x) = t h_0(x) + \sum_{n \geq 0} (e^{itx} - 1) h_n(x)$$

Cela définit un cocycle (qui n'appartient pas en général aux types élémentaires décrits ci-dessus). Nous définissons d'autre part

$$(8) \quad b_t = mt + \text{Im} \langle h_1, (e^{itx} - 1) h_1 \rangle_\lambda + \sum_1^\infty \text{Im} \langle h_n, (e^{itx} - 1 - itx) h_n \rangle_\lambda$$

et avec cela la relation (3) est satisfaite. Désignant alors par $\mathbb{1}$ le vecteur vide dans l'espace de Fock, on a

$$\begin{aligned}\langle \mathbb{1}, \mathcal{G}_t \mathbb{1} \rangle &= e^{ib_t} \langle \mathbb{1}, W(u_t, \theta_t) \mathbb{1} \rangle = e^{ib_t} e^{-\|u_t\|^2/2} \langle \mathbb{1}, \mathcal{E}(u_t) \rangle \\ &= \exp(ib_t - \|u_t\|^2/2)\end{aligned}$$

et il est très facile de vérifier que $ib_t - \|u_t\|^2/2$ est exactement (5) (écrit avec t au lieu de y).

La méthode utilisée ici pour construire un groupe unitaire indexé par $\mathbb{R}$ s'applique à des représentations unitaires de groupes généraux.

En travaillant sur l'espace de Fock construit sur $L^2(\mathbb{R}_+) \otimes L^2(\lambda)$, on peut de manière analogue construire une famille d'opérateurs autoadjoints qui commutent, et qui dans l'état vide ont la loi d'un processus à accroissements indépendants et homogènes, de mesure de Lévy λ . Nous ne donnerons pas de détails.

Eléments de Probabilités Quantiques

VII. QUELQUES REPRESENTATIONS DES RELATIONS DE COMMUTATION, DE TYPE << NON-FOCK >>

Dans l'exposé III, nous avons indiqué à la fin du n°4 une construction simple de lois gaussiennes quantiques pour le couple canonique, qui ne sont pas d'incertitude minimale. Le but de cet exposé est de présenter une construction analogue en dimension infinie. On sait que l'espace de Fock est un objet très familier pour les probabilistes, l'espace L^2 de la mesure de Wiener. Les divers espaces << non-Fock >> que nous allons considérer ici sont toujours l'espace L^2 d'une mesure de Wiener - à deux dimensions cette fois - mais vu sous un angle assez nouveau.

Nous suivons divers articles dus à Hudson, Hudson-Lindsay, Lindsay-Maassen. Les résultats de ce dernier travail nous ont été aimablement communiqués par Hans Maassen, que nous remercions aussi pour de très utiles conversations.

I. RAPPELS SUR L'ESPACE DE FOCK ET GENERALITES SUR LES RCC

1. Soit H un espace de Hilbert complexe : il sera souvent commode de munir H d'une conjugaison $h \mapsto \bar{h}$. Un déplacement $\lambda = (u, U)$ ($u \in H$, U unitaire) opère sur $h \in H$ par $\lambda \cdot h = Uh + u$. Il sera commode aussi de noter $t\lambda$ ($t \in \mathbb{C}$) le déplacement (tu, U) . Le composé de $t\lambda$ et $t\mu$ n'est pas $t^2(\lambda\mu)$ mais $t(\lambda\mu)$! Le déplacement $\bar{\lambda} = (\bar{u}, \bar{U})$ est défini par

$$\bar{\lambda} \cdot h = \overline{\lambda \cdot \bar{h}}.$$

Considérons un second espace de Hilbert complexe Γ . On appelle représentation des RCC (ou plus souvent CCR en anglais : canonical commutation relations) sur H et dans Γ , la donnée d'une famille d'opérateurs de Weyl $(W_h)_{h \in H}$, unitaires sur Γ , possédant la propriété

$$(1) \quad W_0 = I, \quad W_h W_k = \exp(-i \operatorname{Im} \langle h, k \rangle) W_{h+k}$$

ainsi qu'une propriété de régularité minimale : la continuité forte de $W_\cdot$ restreinte à tout sous-espace de dimension finie de H , par exemple.

Il arrive fréquemment que l'on sache représenter dans Γ , non seulement le groupe des translations de H (i.e. H lui même), mais le groupe des déplacements. La forme correspondante de (1) est alors si $\lambda=(u,U)$, $\mu=(v,V)$

$$(2) \quad W_\lambda W_\mu = \exp(-i\text{Im}\langle u, Uv \rangle) W_{\lambda\mu}$$

(cf. exposé IV, n°4, formule (14)). Cependant, (2) est considérée comme moins essentielle que (1) par les physiciens, semble t'il.

La plus simple de toutes les représentations des CCR est la représentation de Fock, étudiée dans l'exposé IV. L'espace de Fock peut être caractérisé de la manière suivante : c'est un espace de Hilbert complexe $\mathfrak{F}$, muni d'une application exponentielle $h \mapsto \mathcal{E}(h)$, dont l'image engendre $\mathfrak{F}$, satisfaisant à

$$(3) \quad \langle \mathcal{E}(f), \mathcal{E}(g) \rangle = e^{\langle f, g \rangle}.$$

Les opérateurs de Weyl sont définis par

$$(4) \quad W_\lambda \mathcal{E}(f) = \exp(-A_\lambda(f)) \mathcal{E}(\lambda \cdot f), \quad A_\lambda(f) = \frac{1}{2}|u|^2 + \langle u, Uf \rangle.$$

Le vecteur-vide $\mathcal{E}(0)=1$ est vecteur cyclique pour la représentation de Weyl : les $W_h 1$ engendrent un sous-espace dense dans $\mathfrak{F}$ (d'une manière générale, les représentations que nous considérerons admettront toutes un tel << vecteur-vide >> cyclique, choisi une fois pour toutes). On a $\langle 1, W_h 1 \rangle = \exp(-\frac{1}{2}\|h\|^2)$ ce qui signifie que le vide définit un état gaussien d'incertitude minimale (le mot << état quasi-libre >> est utilisé un peu partout, mais je ne suis jamais arrivé à savoir ce que cela veut vraiment dire : l'adjectif quasi-libre s'applique aussi à l'état vide dans les représentations non-Fock, que nous verrons plus loin).

Lorsque H est de dimension 1, la représentation de Fock s'identifie à la représentation classique de Schrödinger vue dans l'exposé III. Lorsque H est de dimension finie, le théorème de Stone-von Neumann dit que toute représentation est une somme directe de copies de la représentation de Schrödinger. En dimension infinie, il n'existe aucun modèle irréductible universel jouant le rôle du couple canonique : en un sens, Fock sert tout de même de modèle, en raison de sa simplicité, mais pour construire la C^* -algèbre des relations de commutation sur H , i.e. la C^* -algèbre dans $\mathcal{B}(\mathfrak{F})$ engendrée par les opérateurs W_h : pour toute représentation (W'_h) des RCC dans un espace de Hilbert quelconque Γ , on peut montrer qu'il existe une représentation de la C^* -algèbre des RCC dans $\mathcal{B}(\Gamma)$, qui transforme W_h en W'_h . Autrement dit, la recherche des représentations des RCC revient à la recherche des représentations d'une certaine grosse C^* -algèbre abstraite, et les C^* -algébristes

se frottent les mains.

2. Voici quelques résultats sur les représentations des RCC, qui ont lieu en toute généralité et ne sont pas difficiles. On en trouvera d'autres dans le second volume du livre de Bratteli-Robinson.

Soit K un sous-espace de dimension finie de H , stable par la conjugaison. La restriction de la représentation W à K est (en vertu de l'hypothèse de régularité minimale) une somme directe de copies de la représentation de Schrödinger. Il est donc facile d'étendre les résultats établis pour celle-ci. Ecrivant $h = q + ip$ (p et q réels appartenant à K) le générateur du groupe unitaire (W_{th}) s'écrit

$$(5) \quad \frac{1}{i} \frac{d}{dt} W_{th} \Big|_{t=0} = Q_p - P_q, \quad Q_p = a_p^+ + a_p^-, \quad P_q = i(a_q^+ - a_q^-)$$

où les opérateurs a.a. P_q, Q_p (p, q parcourant le sous-espace réel de K) ont un domaine commun stable dense. Les opérateurs de création a_h^+ et d'annihilation a_h^- sont alors définis de manière à être respectivement linéaires et antilinéaires en h , ils sont fermés, mutuellement adjoints, et l'on a sur le domaine commun stable

$$(6) \quad [a_h^-, a_k^+] = \langle h, k \rangle I \quad ; \quad [P_q, Q_p] = \frac{2}{i} \langle q, p \rangle I$$

Il est commode, comme moyen mnémotechnique, de se rappeler que l'on a (formellement : l'exponentielle n'est pas bien définie)

$$W_h = \exp(a_h^+ - a_h^-)$$

soit $a_h^+ - a_h^- = \frac{d}{dt} W_{th} \Big|_{t=0}$.

En substance, tout cela figure dans l'exposé III. Nous n'en aurons guère besoin, car nous allons nous occuper de représentations concrètes.

II. DEFINITION DE NOUVELLES REPRESENTATIONS

1. Nous avons vu dans l'exposé IV comment on vérifie la relation de commutation (2) à partir de (4) : tout revient à montrer que

$$(7) \quad A_{\lambda\mu}(f) = A_\mu(f) + A_\lambda(\mu.f) - i \operatorname{Im} \langle u, v \rangle$$

Si l'on remplace A par son conjugué, on transforme le signe $-$ en signe $+$ dans le dernier terme. Par conséquent, on retombe sur la relation (7) en remplaçant $A_\lambda(f)$ par $A_{\alpha\lambda}(\alpha f) + \overline{A_{\beta\lambda}(\beta f)}$, à condition de supposer que $|\alpha|^2 - |\beta|^2 = 1$, ce que nous ferons dans toute la suite. Nous poserons $|\alpha|^2 = a$, $|\beta|^2 = b$, $a + b = c$ (> 1 si $b \neq 0$). En fait, nous n'avons pas d'intérêt à supposer α, β complexes, nous les prendrons réels et positifs.

Les << opérateurs de Weyl >>

$$(8) \quad W_\lambda \mathcal{E}(f) = \exp(-\frac{c}{2} \|u\|^2 - a\langle u, Uf \rangle - b\langle Uf, u \rangle) \mathcal{E}(\lambda.f)$$

satisfont donc à la relation de commutation (2), mais ils ne sont plus unitaires. Pour les rendre unitaires, il convient de modifier (1) en prenant

$$(9) \quad \langle \mathcal{E}(f), \mathcal{E}(g) \rangle = e^{a\langle f, g \rangle + b\langle g, f \rangle}.$$

Ainsi on est ramené à la construction d'un espace de Hilbert Ψ muni d'une application exponentielle $\mathcal{E} : H \rightarrow \Psi$ satisfaisant à (9). Nous verrons plus loin que c'est possible. Il est facile de voir, comme dans le cas de l'espace de Fock, que des vecteurs exponentiels distincts en nombre fini sont toujours linéairement indépendants (cf. exposé IV, §I, fin de la section 2). La formule (8) définit alors des opérateurs linéaires sur l'espace des combinaisons linéaires de vecteurs exponentiels, dont l'unitarité se vérifie comme dans le cas Fock.

L'existence peut se ramener au théorème classique suivant lequel le produit (ponctuel) de deux noyaux de type positif est encore de type positif, sachant que $e^{a\langle f, g \rangle}$ et $e^{b\langle g, f \rangle}$ sont de type positif. Nous verrons bien mieux dans un instant.

Enfin, il est clair que cet « espace exponentiel » au dessus de H est unique à isomorphisme près, et que l'espace de Fock correspond au cas $a=1, b=0$.

2. Voici une réalisation concrète de cet espace. Prenons deux copies

de l'espace de Fock que nous appellerons - pour utiliser une terminologie suggestive, mais a priori sans signification physique - l'espace des particules et celui des antiparticules, et que nous affecterons des marques + et - . Nous poserons

$$(10) \quad \Psi = \Phi_+ \otimes \Phi_- , \quad \mathcal{E}(f) = \mathcal{E}_+(\alpha f) \otimes \mathcal{E}_-(-\beta \bar{f})$$

Alors la relation (9) est évidemment satisfaite, et nous verrons un peu plus bas que les vecteurs $\mathcal{E}(f)$ engendrent Φ . Si A est un opérateur sur Φ , nous noterons A^+ ou A'^+ l'opérateur $A \otimes I_-$ sur Ψ , et de même pour $A^- = I_+ \otimes A^{(1)}$. Quant au signe - devant β , nous ne verrons que bien plus tard à quoi il sert. On pose $\mathcal{E}(0) = I = I_+ \otimes I_-$.

Les opérateurs de Weyl peuvent être définis directement par

$$(11) \quad W_\lambda = W_{\alpha\lambda}^+ \otimes W_{-\beta\bar{\lambda}}^- \quad \text{en particulier} \quad W_h = W_{\alpha h}^+ \otimes W_{-\beta\bar{h}}^-$$

Cette formule permet de mettre en évidence un fait important : les opérateurs

$$(12) \quad \tilde{W}_h = W_{-\beta\bar{h}}^+ \otimes W_{\alpha h}^-$$

constituent une autre représentation des RCC dans Ψ , qui commute aux

1. Ou A'^{-} (cette notation nous servira au n°4).

W_h . Si l'on désigne par G l'algèbre de v.N. engendrée par les W_h , son commutant G' contient au moins l'algèbre de v.N. $\tilde{G}$ engendrée par les $\tilde{W}_h$. Nous verrons dans un instant que $\mathbb{I}$ est cyclique pour G , donc pour $\tilde{G}$ par échange des $+$ et $-$, donc a fortiori pour G' . On montre (ce n'est pas difficile) qu'un vecteur cyclique pour le commutant de G est séparant pour G (Bratteli-Robinson, prop. 2.5.3, p. 85), autrement dit que deux opérateurs appartenant à l'algèbre G et qui coïncident sur le vecteur vide sont égaux. Cela peut s'étendre aux opérateurs qui nous intéressent vraiment, c'est à dire aux opérateurs non bornés " affiliés " à l'algèbre u , comme les $Q_h, P_h, a_h^\pm$ et leurs produits.

Remarques. a) Cette situation n'a évidemment pas lieu pour l'espace de Fock, puisque les a_h^- tuent le vecteur $\mathbb{I}$. Elle signifie que les opérateurs de G peuvent être identifiés à des vecteurs de Ψ - et pourront en particulier être développés suivant les chaos de Wiener.

b) On peut montrer que le commutant de G est exactement $\tilde{G}$: ce résultat est cité sans référence, ce qui signifie probablement qu'il est facile, mais je n'en ai pas de démonstration.

c) La présence d'un large commutant signifie, bien sûr, que la représentation (W_h) n'est pas irréductible. En revanche, elle l'est si l'on passe aux W_λ . Je ne connais pas non plus la démonstration de ce résultat.

3. Identifions chaque espace de Fock $\mathfrak{F}^\pm$ à l'espace L^2 d'un mouvement brownien $X^\pm$; il est clair qu'alors Ψ s'identifie à l'espace $L^2(\Omega)$, où Ω est l'espace canonique d'un mouvement brownien à deux dimensions (X^+, X^-) . On pourrait dire " un mouvement brownien complexe ", mais cela n'ajouterait pas grand chose. A tout élément f de $L^2_\phi(\mathbb{R}_+)$ associons la v.a.

$$(13) \quad \tilde{f} = \int_0^\infty (\alpha f_s dX_s^+ - \beta \tilde{f}_s dX_s^-) = \alpha \tilde{f}^+ - \beta \tilde{f}^-$$

Nous considérons la martingale correspondante, et son exponentielle stochastique, ce qui nous donne la v.a. $\mathcal{E}(\tilde{f}) = \mathcal{E}(\alpha \tilde{f}^+) \mathcal{E}(-\beta \tilde{f}^-)$, puisque le crochet des martingales X^+, X^- est nul. Compte tenu de l'identification des exponentielles stochastiques et des vecteurs exponentiels, faite sur l'espace de Fock, on voit sur (10) que $\mathcal{E}(f)$ est l'exponentielle stochastique $\mathcal{E}(\tilde{f})$.

Sachant cela, nous allons démontrer que l'espace vectoriel (que nous notons provisoirement G) engendré par les vecteurs $\mathcal{E}(f)$ est dense dans $L^2_\phi(\Omega)$. Désignons par G_r et G_i les espaces analogues engendrés (sur ϕ) par les $\mathcal{E}(f)$, où f est réelle, resp. imaginaire pure, et par U^r (U^i) les mouvements browniens $\alpha X^+ - \beta X^-$ ($\alpha X^+ + \beta X^-$),

de processus croissants $(a+b)t$, et non orthogonaux, engendrant les tribus respectives $\mathcal{F}_r$ ($\mathcal{F}_i$). L'espace G_r est dense dans $L^2_\phi(\mathcal{F}_r)$, donc il existe une suite d'éléments de G_r convergeant dans L^2_ϕ et p.p. vers $\exp(i \int_0^\infty h_s dU_s^r)$, où h est réelle. D'autre part $\exp(i \int_0^\infty k_s dU_s^i)$, où k est réelle, appartient à G_i , et est bornée. Remarquons que G est une algèbre, donc le produit d'un élément de G_r par un élément de G_i appartient à G . Par multiplication (et la multiplication par une v.a. bornée respecte la convergence L^2) on construit une suite d'éléments de G approchant $\exp(i \int_0^\infty (h_s dU_s^r + k_s dU_s^i))$. Or il est clair que les "polynômes trigonométriques" de ce type engendrent $L^2(\Omega)$.

L'identification de la définition algébrique et de la construction probabiliste est donc achevée. L'idée de considérer la v.a. $\tilde{f}$ de (13) semble vraiment bizarre pour un probabiliste. Compte tenu de l'orthogonalité de X^+ et X^- on a

$$\langle \tilde{f}, \tilde{g} \rangle = a \langle f, g \rangle + b \langle g, f \rangle$$

l'application $\sim$ étant seulement $\mathbb{R}$ -linéaire.

4. Revenons aux RCC. La représentation que nous avons construite satisfait à $\langle \mathbb{I}, W_{\mathbb{I}} \mathbb{I} \rangle = \exp(-\frac{c}{2} \|\mathbb{I}\|^2)$, avec $c=a+b>1$; sous l'état vide, on a donc bien construit l'analogue des états gaussiens du couple canonique qui ne sont pas d'incertitude minimale.

Nous passons au calcul des opérateurs de création et d'annihilation. Compte tenu du travail fait sur l'espace de Fock, nous disposons d'un bon domaine stable dense, constitué par les combinaisons linéaires finies de produits $u^+ \otimes v^-$, où u et v sont des fonctions-test sur les espaces de Fock correspondants (exposé V, §III). Différentiant la relation

$$W_{th} = W_{tch}^{+/+} \otimes W_{-t\beta h}^{-/-}$$

(nous mettons maintenant des $^{\pm}$ pour éviter les confusions) nous obtenons

$$a_h^+ - a_h^- = \alpha(a_h^{+/+} - a_h^{-/+}) - \beta(a_h^{+/-} - a_h^{-/-})$$

- ici, on a utilisé le fait que α et β sont réels, sans quoi on aurait $\bar{\alpha}, \bar{\beta}$ devant $a_h^{-/+}, a_h^{-/-}$. Séparant les parties linéaire et antilinéaire, on a

$$(14) \quad a_h^+ = \alpha a_h^{+/+} + \beta a_h^{-/-}, \quad a_h^- = \alpha a_h^{+/-} + \beta a_h^{-/+}$$

Si α, β sont positifs, on voit que l'on combine positivement la création d'une particule et l'annihilation d'une antiparticule pour définir la création, et de même de l'autre côté. Cela explique le signe mis plus haut devant β .

Les opérateurs d'annihilation ne tuent plus le vecteur vide. Plus

précisément, (14) nous donne (en notation différentielle)

$$(15) \quad da_t^+ \Pi = dY_t^+ = \alpha dX_t^+ , \quad da_t^- \Pi = dY_t^- = \beta dX_t^-$$

Les notations $Y_t^\pm$ ainsi définies seront constamment utilisées dans la suite. En utilisant les notations de l'exposé V, où $\mathcal{P}$ désigne l'ensemble de toutes les parties finies de $\mathbb{R}_+$, nous écrirons

$$dY_A^+ = dY_{s_1}^+ \dots dY_{s_n}^+ , \quad dA = ds_1 \dots ds_n , \quad |A|=n$$

si $A=\{s_1, \dots, s_n\}$, les instants $s_1, \dots, s_n$ étant tous distincts. Dans ces conditions, la décomposition d'un élément de Ψ selon les chaos de Wiener peut s'écrire

$$(16) \quad f = \int_{\mathcal{P} \times \mathcal{P}} \hat{f}(A,B) dY_A^+ dY_B^-$$

avec

$$(17) \quad \|f\|^2 = \int_{\mathcal{P} \times \mathcal{P}} |\hat{f}(A,B)|^2 a^{|A|} b^{|B|} dA dB .$$

La notation $\hat{f}$ suggère qu'il s'agit d'une sorte d'analyse harmonique du vecteur f : nous regrettons de ne pas l'avoir utilisée dès l'exposé IV ! Le calcul de l'effet d'un opérateur de création ou d'annihilation sur le vecteur f s'obtient alors par intégration à partir de la table suivante

$$(18) \quad \begin{aligned} da_t^+ (dY_A^+ dY_B^-) &= dY_{A+t}^+ dY_B^- + b dY_A^+ dY_{B-t}^- dt \\ da_t^- (dY_A^+ dY_B^-) &= a dY_{A-t}^+ dY_B^- dt - dY_A^+ dY_{B+t}^- \end{aligned}$$

où $dY_{H+t}^\pm$ désigne $dY_{H \cup \{t\}}^\pm$ si $t \notin H$, et 0 sinon, $dY_{H-t}^\pm$ désigne $dY_{U \setminus \{t\}}^\pm$ si $t \in H$, et 0 sinon.

Ces formules seront transformées plus tard en définitions d'opérateurs associés à des noyaux, à la manière des noyaux de Maassen de l'exposé V.

5. Pour l'instant, nous allons tenter de répondre à des questions élémentaires, suggérées par l'analogie avec l'espace de Fock.

a) Que sont les opérateurs de position $Q_h = a_h^+ + a_h^-$, pour h réelle ?

Il s'agit en effet d'opérateurs a.a. qui commutent, et ils doivent pouvoir s'interpréter comme des v.a. classiques. Or il est trivial que $Q_h = \alpha Q_h^+ + \beta Q_h^-$: c'est donc l'opérateur de multiplication par l'intégrale stochastique $\int_0^\infty h_s (\alpha dX_s^+ + \beta dX_s^-)$, relative au mouvement brownien $Y^+ + Y^-$ de processus croissant ct . Nous ne chercherons pas à décrire les opérateurs d'impulsion P_h .

b) L'analogue des opérateurs de jauge (opérateurs de comptage) a_h^0 sur l'espace de Fock est fourni par les opérateurs v_h définis,

pour h réelle, par la formule

$$e^{itv_h} \mathcal{E}(f) = \mathcal{E}(e^{ith} f) = \mathcal{E}^+(\alpha e^{ith} f) \otimes \mathcal{E}^-(\beta e^{-ith} f)$$

et par conséquent $v_h = a_h^{\circ/+} - a_h^{\circ/-}$. En particulier pour $h=I[0,t]$

$$(19) \quad v_t = a_t^{\circ/+} - a_t^{\circ/-}$$

Il s'agit encore d'un opérateur à valeurs entières, mais celles-ci ne sont plus positives. D'autre part, les opérateurs e^{itv_h} laissent fixe le vecteur-vide 1 , et donc n'appartiennent plus à l'algèbre de Weyl. Cependant, les résultats probabilistes établis dans l'espace de Fock s'étendent ici : par exemple, les opérateurs $v_t + \lambda Q_t$ sont a.a. et commutent, et peuvent donc être interprétés comme des v.a. classiques. Il est facile de voir, en utilisant les résultats de l'exposé IV, qu'ils constituent en fait des différences de deux processus de Poisson compensés indépendants, à sauts unité : le premier d'intensité $\lambda^2 a$, le second d'intensité $\lambda^2 b$.

c) On peut appeler m-ième chaos dans l'espace Ψ le sous-espace engendré par les vecteurs de la forme

$$J_m(f) = \frac{1}{m!} D_t^m \mathcal{E}(tf) |_{t=0}$$

$$\text{On a} \quad \langle J_m(f), J_n(g) \rangle = \delta_{mn} (a\langle f, g \rangle + b\langle g, f \rangle)^m.$$

Les chaos forment une décomposition orthogonale de l'espace Ψ , et l'on peut définir un opérateur a.a. N , qui admet le n -ième chaos comme sous-espace propre pour la valeur propre n (et qui tue donc le vecteur vide, ce qui exclut qu'il soit affilié à l'algèbre de Weyl). Il n'est pas difficile de voir que $N = N^{\circ/+} + N^{\circ/-}$.

IV. CALCUL STOCHASTIQUE SUR Ψ

1. Nous allons d'abord passer en revue, de manière rapide, une extension du calcul stochastique de Hudson-Parthasarathy de l'exposé V. Cette extension n'est pas très intéressante, mais elle est entièrement élémentaire et permet de comprendre un certain nombre de choses.

Nous désignons par Ψ_t (Ψ_t , $\Psi_{[s,t]}$) le sous-espace de Ψ engendré par les vecteurs exponentiels $\mathcal{E}(f)$, où f est nulle hors de $[0,t]$ ($[t, \infty[$, $[s,t]$). La notion de processus adapté d'opérateurs s'étend sans modification à l'espace Ψ : il s'agit toujours d'opérateurs définis sur les combinaisons linéaires finies de vecteurs exponentiels $\mathcal{E}(u)$, où u est bornée à support compact. Nous ne considérons que des processus adaptés (H_t) pour lesquels $H_t=0$ pour t grand, et l'intégrale $\int \|H_t \mathcal{E}(u)\|^2 dt$ est finie pour tout u .

Nous disposons maintenant de six martingales fondamentales $a_t^{\varepsilon/\pm}$, où ε prend les valeurs $-, \circ, +$, par rapport auxquelles nous pouvons intégrer le processus (H_t) , obtenant ainsi des intégrales $I^{\varepsilon/\pm}(H)$. L'extension est absolument évidente. Il serait possible de dresser un énorme formulaire, mais nous nous en abstenons : nous dresserons seulement la table d'Ito, dont les seuls termes non nuls sont les termes déjà connus (en particulier, tous les termes $da_t^{\varepsilon/+}da_t^{\eta/-}$ sont nuls).

$$da_t^{-/+}da_t^{\circ/\pm}=da_t^{-/\pm}, da_t^{-/\pm}da_t^{+/\pm}=dt, da_t^{\circ/\pm}da_t^{\circ/\pm}=da_t^{\circ/\pm}, da_t^{\circ/\pm}da_t^{+/\pm}=da_t^{+/\pm}$$

De la même manière, on peut étendre à la présente situation, sans changement autre que la lourdeur des notations, la résolution des équations différentielles stochastiques.

En particulier, d'après (14) on peut définir des i.s. et résoudre des é.d.s. par rapport aux processus da_t^+ et da_t^- : on obtient une sous-table d'Ito

$$(20) \quad da_t^+da_t^+=da_t^-da_t^-=0, \quad da_t^-da_t^+=adt, \quad da_t^+da_t^-=bdt.$$

Il est impossible d'adjoindre un << opérateur de comptage >> a_t° , combinaison linéaire de $a_t^{\circ/+}$ et $a_t^{\circ/-}$, de manière à engendrer une sous-table de la table d'Ito distincte de la table complète.

2. Nous allons adopter dans cette section (dont les résultats sont dus à Lindsay et Maassen) un point de vue différent, proche de celui des noyaux sur l'espace de Fock (exposé V). Nous allons commencer par des calculs formels, que nous rendrons rigoureux ensuite en introduisant une classe de fonctions-test.

Nous recopions d'abord la représentation (16)

$$(16) \quad f = \int \hat{f}(A,B) dY_A^+ dY_B^-$$

Nous dirons que $\hat{f}$ est le noyau du vecteur f . Si $\hat{f}(.,.)$ est nul p.p. pour la mesure $dAdB$, il représente le vecteur nul : en particulier, on peut supposer dans la représentation (16) que $\hat{f}(A,B)=0$ si $A \cap B = \emptyset$, puisque l'ensemble des diagonales (couples d'intersection non vide) est de mesure nulle. Nous dirons dans ce cas que le noyau est restreint.

Nous allons d'autre part, comme sur l'espace de Fock, considérer des opérateurs (désignés par des lettres majuscules) représentés par des noyaux

$$(21) \quad K = \int K(S,T) da_S^+ da_T^-$$

Ne sachant pas encore ce que signifie cette notation, nous ignorons a priori si la contribution des diagonales est négligeable. Nous dirons que le noyau est restreint si $K(S,T)=0$ pour $S \cap T \neq \emptyset$.

Pour donner un sens à (21), il est naturel de chercher à calculer d'abord la << table de multiplication >>, par itération des formules (18). On a

$$(22) \quad \begin{array}{c} (da_S^+ da_T^-)(dY_A^+ dY_B^-) = \sum_{\substack{S_1+S_2=S \\ T_1+T_2=T}} dY_{A-T_1+S_1}^+ dY_{B+T_2-S_2}^- a^{|T_1|} b^{|S_2|} dT_1 dS_2 \\ \text{opérateur.vecteur} \end{array}$$

Une notation comme dY_{A-T+S}^+ représente 0 si l'on n'a pas $T \subset S$, et $S \cap (A \setminus T) = \emptyset$; si ces conditions sont satisfaites, elle représente $dY_{(A \setminus T) \cup S}^+$

Ensuite, nous essayons de calculer l'action du noyau K sur le vecteur f : désignant le vecteur Kf par $h = \hat{h}(U, V) dY_U^+ dY_V^-$, la recherche de $h(U, V)$ nous amène à résoudre

$$A - T_1 + S_1 = U, \quad B + T_2 - S_2 = V$$

Nous poserons aussi

$$T_1 = N, \quad S_2 = M, \quad S_1 = U_1, \quad A - T_1 = U_2 \quad (\text{de sorte que } U = U_1 + U_2)$$

donc $A = U_2 + N$, $S = U_1 + M$. Supposant le noyau K restreint, nous pouvons nous borner dans (22) à considérer des couples S, T disjoints, et en particulier T_2, S_2 sont disjoints, donc $B + T_2 - S_2 = V$ entraîne $T_2 \subset V$; posons donc $T_2 = V_1$, $V - T_2 = V_2$. Nous avons $T = T_1 + T_2 = V_1 + N$, et $B = V_2 + M$. Ainsi le coefficient de $Y_U^+ Y_V^-$ dans $Kf = h$ est

$$(23) \quad \hat{h}(U, V) = \int \sum_{\substack{U_1+U_2=U \\ V_1+V_2=V}} K(U_1+M, V_1+N) \hat{f}(U_2+N, V_2+M) a^{|N|} b^{|M|} dM dN$$

En principe, chaque fois que les deux ensembles séparés par un $+$ ne sont pas disjoints, le terme correspondant est remplacé par 0 - mais en fait, M et N sont presque certainement ($dM dN$) disjoints de l'ensemble fini UV , donc on peut remplacer $+$ par U si on le désire.

Pour établir cette formule, nous avons supposé le noyau K restreint. Maintenant, prenant la formule telle quelle, mais supposant K quelconque, nous remarquons que si $U \cap V = \emptyset$, $U_1 + M$ et $V_1 + N$ sont presque certainement disjoints ($dM dN$), donc la contribution des termes diagonaux est nulle, et d'autre part, la contribution dans h des $\hat{h}(U, V)$ diagonaux est nulle. Donc on obtient le même vecteur en modifiant arbitrairement K sur les diagonales. De même, la contribution des termes diagonaux éventuels de $\hat{f}$ ne modifie pas $h = Kf$. En définitive, on peut considérer sans contradiction que les termes diagonaux sont négligeables dans la formule (21).

C'est la seule attitude cohérente, car on ne peut pas travailler sur les vecteurs et noyaux restreints, et ignorer entièrement les autres: dans la formule (23), même si f et K sont restreints, $\hat{h}$ ne l'est pas.

Un peu plus généralement, on voit que si l'on modifie le noyau K sur un ensemble de mesure nulle ($dS dT$), on obtient le même vecteur par la formule (23). D'autre part, si on calcule Kf par la formule (23), on

obtient simplement $\hat{h}(U,V)=K(U,V)$. Autrement dit, le noyau de l'opérateur K est simplement le noyau du vecteur KL.

Considérons ensuite deux noyaux L, K , que pour simplifier nous pouvons prendre restreints. Définissons un nouveau noyau $L \star K$ par la formule imitée de (23)

$$(24) \quad L \star K(W, Z) = \int \sum_{\substack{A_1 + A_2 = W \\ B_1 + B_2 = Z}} L(A_1 + H, B_1 + J) K(A_2 + J, B_2 + H) a^{|J|} b^{|H|} dH dJ$$

On a d'après (23) $(L \star K)l = Lk$, où k est le vecteur de noyau égal à $K(.,.)$. Or nous avons vu plus haut que $KL=k$. Donc $L \star K$ est un candidat à représenter le noyau composé LK . La vérification de la relation $L \star K(f) = L(Kf)$ - qui équivaut à celle de l'associativité de l'opération $\star$ - a été faite par Lindsay-Maassen de manière directe, et nous ne la reproduirons pas ici.

Comme vecteurs et opérateurs sont identifiés par leurs noyaux, (24) peut aussi être considérée comme une multiplication associative sur les vecteurs.

REMARQUE. On peut aussi calculer la composition des opérateurs au moyen de la table d'Ito. Cela amène aux formules

$$(25) \quad \begin{aligned} da_T^- da_A^+ da_B^+ &= a^{|T \cap A|} d(T \cap A) da_{A \setminus T}^+ da_{B \setminus (T \setminus A)}^- \\ da_S^+ da_A^+ da_B^+ &= b^{|S \cap B|} d(S \cap B) da_{A \setminus (S \setminus B)}^+ da_{B \setminus S}^- \end{aligned}$$

et si $S \cap T = \emptyset$

$$(26) \quad da_S^+ da_T^- da_A^+ da_B^- = a^{|T \cap A|} b^{|S \cap B|} d(T \cap A) d(S \cap B) da_{(A \setminus T) \cup (S \setminus B)}^+ da_{(B \setminus S) \cup (T \setminus A)}^-$$

Cette formule est plus simple que (22), car il n'y a au second membre qu'un seul terme au lieu d'une somme. Néanmoins, si l'on se sert de (26) pour calculer le composé de deux noyaux, on retombe sur (24). Cela explique (de manière non entièrement rigoureuse) l'associativité de (24).

3. Digression : recherche d'un modèle fini.

Dans le cas de l'espace de Fock, nous avons à notre disposition un modèle fini, le << bébé Fock >>, qui rendait de grands services. Nous allons tenter de faire de même ici, en partant de la table de multiplication discrète analogue à (22) : cette tentative se soldera par un échec.

Soit E un ensemble fini. Nous construisons un espace de Hilbert admettant pour base orthonormale des vecteurs $e_A^+ e_B^-$, où A, B sont des parties disjointes de E : en temps continu, cela correspond aux vecteurs restreints $dX_A^+ dX_B^- / \sqrt{dA dB}$. En chaque <<site>> i nous définissons les opérateurs $a_i^\pm$ par les formules

$$\begin{aligned} a_i^+(e_A^+ e_B^-) &= \beta e_{A-B-i}^+ e_B^- \text{ si } i \in B, \quad \alpha e_{A+i}^+ e_B^- \text{ si } i \notin B \\ a_i^-(e_A^+ e_B^-) &= \alpha e_{A-i}^+ e_B^- \text{ si } i \in A, \quad \beta e_{A+B+i}^+ e_B^- \text{ si } i \notin A \end{aligned}$$

Les opérateurs $a_{ST}^+ a_T^-$ sont définis par composition. Comme les divers sites sont indépendants, pour comprendre ce qui se passe, il suffit d'étudier le cas où $E=\{i\}$. Il y a alors trois vecteurs de base $e_\emptyset^+ e_\emptyset^-$, $e_\emptyset^+ e_i^-$ et $e_i^+ e_\emptyset^-$, et les matrices de a_i^+, a_i^- sont alors

$$a_i^+ = \begin{pmatrix} 0 & \beta & 0 \\ 0 & 0 & 0 \\ \alpha & 0 & 0 \end{pmatrix}, \quad a_i^- = \begin{pmatrix} 0 & 0 & \alpha \\ \beta & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix}$$

Il n'est pas difficile de vérifier que l'algèbre d'opérateurs engendrée par ces matrices contient tous les opérateurs. On a $a^+ = \alpha a^{+/+} + \beta a^{-/-}$, $a^- = \alpha a^{-/+} + \beta a^{+/-}$, avec

$$a^{+/+} = \begin{pmatrix} 0 & 0 & 0 \\ 0 & 0 & 0 \\ 1 & 0 & 0 \end{pmatrix}, \quad a^{-/+} = \begin{pmatrix} 0 & 0 & 1 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix}, \quad a^{+/-} = \begin{pmatrix} 0 & 0 & 0 \\ 1 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix}, \quad a^{-/-} = \begin{pmatrix} 0 & 1 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix}.$$

Cela ressemble beaucoup à la construction donnée plus haut des représentations non-Fock à partir du produit de deux Fock, mais nous avons économisé une dimension. L'espérance d'un opérateur dans l'état vide est le coefficient diagonal supérieur de sa matrice. Pour engendrer les opérateurs d'espérance nulle (les martingales), il faut donc huit opérateurs de base : les quatre ci-dessus, encore

$$a^{0/+} = \begin{pmatrix} 0 & 0 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 1 \end{pmatrix} \quad \text{et} \quad a^{0/-} = \begin{pmatrix} 0 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 0 \end{pmatrix},$$

et deux autres qui s'évanouissent en dimension infinie,

$$(a^{+/+})^2 = \begin{pmatrix} 0 & 0 & 0 \\ 0 & 0 & 0 \\ 0 & 1 & 0 \end{pmatrix}, \quad (a^{-/-})^2 = \begin{pmatrix} 0 & 0 & 0 \\ 0 & 0 & 1 \\ 0 & 0 & 0 \end{pmatrix}.$$

Ce << bébé non-Fock >> n'explique donc pas le miracle qui se produit en dimension infinie.

4. Résumé de quelques résultats supplémentaires.

Nous n'avons pas l'intention d'aller beaucoup plus loin dans l'étude des représentations des RCC du type << non Fock >> : nous renverrons le lecteur à l'article de Hudson et Lindsay Uses of non-Fock quantum brownian motion and a quantum martingale representation theorem, Quantum probability and applications II, Accardi et v. Waldenfelds éd., p. 276-305, Lect. Notes in M. 1136, Springer 1985. Nous nous bornerons à quelques remarques sur la comparaison entre le cas de l'espace de Fock, que nous avons étudié de manière assez approfondie, et le cas << non Fock >>.

Dans le cas Fock, l'algèbre de von Neumann engendrée par les opérateurs de Weyl comprenait tous les opérateurs bornés ; le vecteur vide $\mathbb{1}$ était cyclique, mais non séparant ; il existait des opérateurs bornés non donnés par un noyau de Maassen, et des martingales d'opérateurs non représentables en intégrales stochastiques par rapport aux trois processus a_t^ε , $\varepsilon = -, 0, +$.

Dans le cas $\langle\langle \text{non-Fock} \rangle\rangle$, la situation est en principe beaucoup plus simple, parce que le vecteur vide $\mathbb{1}$ est cyclique et séparant pour l'algèbre de von Neumann engendrée par les opérateurs de Weyl W_h : tout opérateur A appartenant à cette algèbre de von Neumann (ou plus généralement " affilié " à celle-ci, s'il s'agit d'un opérateur non borné : cela signifie que A commute à tous les opérateurs du commutant, en un sens convenable) admet alors un noyau, qui est tout simplement celui du vecteur $\mathbb{1}$. A partir de là, il est possible de montrer que toute martingale d'opérateurs, affiliée aussi à cette algèbre de v.N., est représentable en intégrale stochastique par rapport aux deux martingales de création et d'annihilation.

Bien entendu, ceci n'est que le principe de la démonstration : il reste à faire tout le travail précis - or dans cet exposé, nous n'avons fait que des calculs formels. Contrairement au cas de l'espace de Fock, nous n'avons même pas examiné les fonctions-test et les noyaux réguliers. Nous sommes donc bien loin d'avoir étudié les opérateurs dont le noyau est simplement dans L^2 !

Une partie du travail récemment accompli par Hudson et ses collaborateurs consiste à étendre aux représentations $\langle\langle \text{non-Fock} \rangle\rangle$ du type de cet exposé, sous une forme en général plus simple, les résultats déjà connus pour la représentation de Fock. Celle-ci apparaît alors comme une sorte de cas limite dégénéré des représentations $\langle\langle \text{non-Fock} \rangle\rangle$.

Eléments de Probabilités Quantiques

VIII . TEMPS D'ARRET SUR L'ESPACE DE FOCK (d'après Parthasarathy-Sinha)

Les temps d'arrêt sont un élément essentiel de la théorie générale des processus. Il ne manque pas d'extensions de la notion de temps d'arrêt à des situations non-commutatives abstraites (filtrations d'algèbres de von Neumann), mais elles manquent un peu de motivations, et s'accordent mal avec le caractère assez élémentaire de ces notes. Fait exception un article de Hudson (J. Funct. Anal. 34, 1979) qui traite une situation concrète, celle du mouvement brownien quantique, le couple (Q_t, P_t) . Cet article datant d'avant le calcul stochastique quantique de Hudson-Parthasarathy, il vient d'être repris et considérablement amplifié par un travail de Parthasarathy et Sinha, que nous présentons ici. Cet article est très riche en idées nouvelles pour les probabilités non-commutatives. Je remercie vivement les auteurs pour leur autorisation d'insérer ces notes dans le séminaire, bien que leur article ne soit pas encore paru.

I. TEMPS D'ARRET DISCRETS

1. Notations. $\mathfrak{F}$ est l'espace de Fock, identifié pour fixer les idées à l'espace $L^2(\Omega, \mathfrak{F}, P)$ engendré par le mouvement brownien (X_t) issu de 0. Pour toute fonction h sur $\mathbb{R}_+$, on pose $h_s = hI_{[0,s]}$, $h_s = hI_{[s,\infty]}$ et l'on désigne par $\mathfrak{F}_s$, $\mathfrak{F}_{[s}$ les espaces engendrés par les vecteurs exponentiels $\mathcal{E}(h_s)$, $\mathcal{E}(h_{[s})$ lorsque h varie. Le] du passé sera assez souvent omis, mais bien sûr jamais le [du futur : ainsi, on désignera par E_s le projecteur orthogonal sur $\mathfrak{F}_s$. Dans l'interprétation brownienne on a $\mathfrak{F}_s = L^2(\mathfrak{F}_s)$, et E_s est l'espérance conditionnelle $E[.|\mathfrak{F}_s]$.

Le vecteur vide $\mathcal{E}(0)$ est noté $\mathbb{1}$.

On note Θ_t l'isométrie de $\mathfrak{F}$ sur $\mathfrak{F}_{[t}$ définie par $\Theta_t \mathcal{E}(h) = \mathcal{E}(\Theta_t h)$, avec $\Theta_t h(s) = h(s-t)I_{\{s \geq t\}}$. On pose $\Theta_\infty = 0$.

Nous utiliserons rarement la multiplication des v.a., qui fait sortir de L^2 en général, et dépend de l'identification de l'espace de Fock à l'une de ses interprétations probabilistes. Toutefois, le produit uv de deux éléments de l'espace de Fock appartenant l'un à $\mathfrak{F}_t$, l'autre à $\mathfrak{F}_{[t}$ pour un $t \geq 0$ peut être défini sans recourir à une interprétation probabiliste, et sera très fréquemment utilisé.

Un opérateur borné H sera dit s-adapté si, pour $u \in \mathfrak{F}_s$, $v \in \mathfrak{F}_{[s}$ on a

$H(uv)=H(u)v$, $H(u)$ appartenant encore à Φ_s . Notons

LEMME 1. Un opérateur s-adapté H commute à E_s .

Démonstration. Il suffit de raisonner sur les vecteurs exponentiels $x = \mathcal{E}(g) = \mathcal{E}(g_s) \mathcal{E}(g_{[s]})$: on a $Hx = (H\mathcal{E}(g_s)) \mathcal{E}(g_{[s]})$, $E_s Hx = H\mathcal{E}(g_s)$, $E_s x = \mathcal{E}(g_s)$, $HE_s x = H\mathcal{E}(g_s)$.

2. Soit T une mesure spectrale étalée sur $[0, \infty]$. Nous noterons $I_{\{T \in A\}}$ le projecteur spectral associé au borélien A de $\mathbb{R}_+$. On dit que T est un temps d'arrêt si pour tout s le projecteur $I_{\{T < s\}}$ est s -adapté.

Le projecteur $I_{\{T=0\}}$ étant 0-adapté ne peut être que 1 ou 0. Le premier cas étant trivial, nous nous placerons toujours dans le second. Nous dirons que T est essentiellement fini (essentiellement infini) si $\langle I_{\{T=\infty\}}, 1, 1 \rangle = 0$ (resp. 1). On verra que l'état 1 joue un rôle bien particulier dans cette affaire, et on distinguera essentiellement fini de p.s. fini (qui suppose la donnée d'une loi de probabilité).

Comme en théorie des martingales ou des processus de Markov, on commence par étudier le cas des temps d'arrêt discrets, pour lesquels la mesure spectrale a un support fini (s_i) ($0 \leq s_1 \dots \leq s_N \leq +\infty$), puis on tente de passer du discret au continu en approchant T par une suite décroissante de t. d'a. discrets. La suite que l'on utilise est la même qu'en probabilités classiques

$$(1) \quad I_{\{T_n = \infty\}} = I_{\{T \geq 2^n\}} \quad , \quad I_{\{T_n = (k+1)2^{-n}\}} = I_{\{k2^{-n} \leq T < (k+1)2^{-n}\}}$$

pour $k=0, 1, \dots, 2^{2^n}-1$. Il est clair que T_n est un t. d'a. qui commute à T - les T_n commutant tous entre eux : ce sont des fonctions $f_n(T)$, où $f_n(x) \rightarrow x$ pour $n \rightarrow \infty$.

EXEMPLE. Soit $(W, \mathcal{G}, (\mathcal{G}_t))$ un espace mesurable avec une filtration. Soit Y une mesure spectrale sur Φ , étalée sur $(W, \mathcal{G})$, et telle que pour tout $A \in \mathcal{G}_s$ le projecteur spectral $I_{\{Y \in A\}}$ soit s -adapté. Soit S un t. d'a. au sens ordinaire sur W . Définissons une mesure spectrale T par

$$I_{\{T \in B\}} = I_{\{Y \in S^{-1}(B)\}} \quad (B \text{ borélien de } [0, \infty]).$$

Alors T est un t. d'a. sur l'espace de Fock. Par exemple, l'interprétation brownienne de l'espace de Fock définit une mesure spectrale étalée sur l'espace Ω des trajectoires continues, et tout t. d'a. du mouvement brownien (au sens usuel) devient un t. d'a. sur l'espace de Fock, commutant aux opérateurs Q_t .

Voici un exemple plus intéressant. Soit $(W, \mathcal{G}, P, (N_t))$ la réalisation canonique d'un processus de Poisson (non compensé), nul en 0, non homogène, d'intensité $h_t^2 dt$ où $h \in L^2(\mathbb{R}_+)$ est partout >0 . Soit $J(w)$

l'ensemble des sauts de la trajectoire $N_\bullet(w)$, p.s. fini. Soit $Q = e^{\|h\|^2}_P$.
 Un calcul simple d'exponentielle de Doléans montre que

$$E_P[e^{-\langle h, f \rangle} \prod_{s \in J} (1 + \frac{f(s)}{h(s)})] = 1 \quad \text{pour } f \in L^2(\mathbb{R}_+).$$

Posons alors $\varepsilon(f) = \prod_{s \in J} \frac{f(s)}{h(s)}$ (de sorte que $\varepsilon(h)=1$, $\varepsilon(0)=I_{\{J=\emptyset\}}$).

On vérifie que $\varepsilon(f)\varepsilon(g)=\varepsilon(fg/h)$, et on en déduit que $\langle \varepsilon(f), \varepsilon(g) \rangle_Q = \exp(\langle f, g \rangle)$, d'où sans peine un isomorphisme entre l'espace de Fock et $L^2(Q)$, dans lequel les $\varepsilon(f)$ correspondent aux vecteurs cohérents. Les opérateurs de multiplication par X_t correspondent aux opérateurs a_t° . Enfin, le premier saut T du processus (N_t) est un t. d'a. fort intéressant, qui est essentiellement infini.

3. Arrêt d'un processus de vecteurs.

Soit T un t. d'a. discret, la mesure spectrale étant concentrée sur un ensemble fini (r_i) .

Soit (y_t) une courbe dans l'espace de Fock, définie aussi pour $t=\infty$. Nous poserons

$$(2) \quad y_T = \sum_i I_{\{T=r_i\}} y_{r_i}.$$

Nous avons alors aussi

$$(3) \quad I_{\{T \in A\}} y_T = \sum_{i \in A} I_{\{T=r_i\}} y_{r_i}$$

Le cas des martingales fermées $x_t = E_t x$ est particulièrement important. On pose dans ce cas $x_T = E_T x$, et l'on a

$$(4) \quad I_{\{T \in A\}} x_T = \sum_{i \in A} I_{\{T=r_i\}} E_{r_i} x$$

LEMME 2. Les opérateurs $I_{\{T=r_i\}} E_{r_i}$ sont des projecteurs deux à deux orthogonaux.

Démonstration. D'après le lemme 1, $I_{\{T=r_i\}}$ et E_{r_i} commutent. Donc leur produit est un projecteur. Pour $i \neq j$, on a $I_{\{T=r_i\}} I_{\{T=r_j\}} = 0$, donc le produit est nul.

Donc l'opérateur $x \mapsto I_{\{T \in A\}} x_T$ est un projecteur (que nous désignerons plus bas par $I_{\{T \in A\}} E_T$). Comme le produit de deux projecteurs n'est un projecteur que s'ils commutent, on voit que E_T commute à la mesure spectrale T .

Nous supposons toujours que (x_t) est une martingale fermée $E_t x$, et nous représentons x sous la forme $x_{0+} + \int_0^\infty h_s dX_s$, où (X_t) est la courbe brownienne dans l'espace de Fock, et (h_t) est une courbe adaptée telle que $\int_0^\infty \|h_s\|^2 ds < \infty$. On a alors

LEMME 3. On a

$$(5) \quad E_T x = x_T = x_0 + \int (I_{\{T \geq s\}} h_s) dX_s$$

Démonstration. On peut supposer que $x_0 = 0$. Alors

$$\begin{aligned} x_T &= \sum_j I_{\{T=r_j\}} \int_0^{r_j} h_s dX_s = \sum_{j, i < j} I_{\{T=r_j\}} \int_{r_i}^{r_{i+1}} h_s dX_s = \\ &= \sum_i \int_{r_i}^{r_{i+1}} (I_{\{T > r_i\}} h_s) dX_s = \int (I_{\{T \geq s\}} h_s) dX_s. \end{aligned}$$

REMARQUE. Il est clair sur la formule (5) que si S et T sont deux t. d'a. qui commutent, E_S et E_T commutent, leur produit étant $E_{S \wedge T}$ (en revanche, il n'y a aucune raison pour que $I_{\{S \in A\}} E_S$ et $I_{\{T \in B\}} E_B$ commutent : cela n'a déjà pas lieu en théorie générale des processus).

NOTATION. L'image de Φ par le projecteur E_T (resp. $I_{\{T \in A\}} E_T$) sera notée Φ_T , ou $\Phi_T]$ si cette précision est nécessaire (resp. $I_{\{T \in A\}} \Phi_T$).

On peut remarquer que, si (y_t) est une courbe adaptée quelconque, on a

$$\begin{aligned} I_{\{T \in A\}} E_T (I_{\{T \in A\}} y_T) &= \sum_{r_i \in A} E_{r_i} I_{\{T=r_i\}} \sum_{r_j \in A} I_{\{T=r_j\}} y_{r_j} \\ &= \sum_{r_i \in A} I_{\{T=r_i\}} y_{r_i} = I_{\{T \in A\}} y_T \end{aligned}$$

Autrement dit, $I_{\{T \in A\}} y_T$ appartient à $I_{\{T \in A\}} \Phi_T$ (en particulier, on obtient un élément de Φ_T en arrêtant à T une courbe adaptée quelconque, pas nécessairement une martingale). Comme en théorie générale des processus, on peut caractériser les éléments de Φ_T par

$$(6) \quad (x \in \Phi_T) \Leftrightarrow (\forall s \ I_{\{T < s\}} x \in \Phi_s) .$$

Esquissons une démonstration qui s'applique aux t. d'a. non nécessairement discrets : Ecrivant x sous la forme $x_0 + \int h_r dX_r$, on a (cf. (5))

$$(x \in \Phi_T) \Leftrightarrow (\text{pour presque tout } r, h_r = I_{\{T \geq r\}} h_r)$$

$$(I_{\{T < s\}} x \in \Phi_s) \Leftrightarrow (I_{\{T < s\}} \int_s^\infty h_r dX_r \in \Phi_s)$$

Or cette dernière intégrale peut aussi s'écrire $\int_s^\infty I_{\{T < s\}} h_r dX_r$, et ne peut appartenir à Φ_r que si $I_{\{T < s\}} h_r = 0$ pour presque tout $r > s$, autrement dit $I_{\{T \geq s\}} h_r = h_r$ pour presque tout $r > s$. Par Fubini, cela entraîne que pour presque tout r , $h_r = I_{\{T \geq s\}} h_r$ pour presque tout $s < r$, et $h_r = I_{\{T \geq r\}} h_r$ par passage à la limite. Nous laissons au lecteur l'implication inverse.

4. Arrêt de processus d'opérateurs.

La théorie de l'arrêt de processus d'opérateurs est peu développée dans le travail de Parthasarathy-Sinha (elle l'est davantage dans un article à paraître de D. Applebaum).

On pourrait songer par exemple à une formule analogue à (5) pour définir l'arrêt d'un processus intégrale stochastique $\int_0^t H_s da_s^\varepsilon$ ($\varepsilon = -, 0, +$).

En fait, nous n'étudierons seulement l'arrêt de certains processus d'opérateurs bornés, non nécessairement adaptés. Si (H_t) est un tel processus, défini aussi pour $t=\infty$, nous poserons à la manière de (2), toujours dans le cas discret

$$(7) \quad H_T = \sum_i I_{\{T=r_i\}} H_{r_i}$$

Il se trouve que ceci est la forme d'arrêt la plus utile, mais dès que l'on veut passer à l'adjoint, on rencontre une autre forme, pour laquelle les indicatrices sont placées à droite :

$$(8) \quad H_{T*} = \sum_i H_{r_i} I_{\{T=r_i\}}.$$

Alors l'adjoint de H_T est H_{T*}^* , ce qui n'est pas trop ridicule. Pour tout $y \in \mathfrak{H}$, $H_T y$ est l'arrêté à T de la courbe $(H_t y)$.

Exemple 1. Si (E_t) est le processus - non adapté - des espérances conditionnelles, projecteurs sur les $\mathfrak{H}_t$, les opérateurs E_T et E_{T*} coïncident avec l'opérateur noté E_T avant (4).

Exemple 2. Considérons le processus - non adapté - (Θ_t) , où l'on convient de poser $\Theta_\infty = 0$. Nous avons

$$\langle \Theta_T x, \Theta_T y \rangle = \sum_{i,j} \langle I_{\{T=r_i\}} \Theta_{r_i} x, I_{\{T=r_j\}} \Theta_{r_j} y \rangle$$

Mais pour $r_i < \infty$ on a $\Theta_{r_i} x = \mathbb{1} \Theta_{r_i} x$, et l'adaptation de $I_{\{T=r_i\}}$ nous permet d'écrire

$$I_{\{T=r_i\}} \Theta_{r_i} x = (I_{\{T=r_i\}} \mathbb{1}) \Theta_{r_i} x.$$

D'autre part, la sommation ne porte que sur les couples $i=j$, et il nous reste

$$\begin{aligned} \langle \Theta_T x, \Theta_T x \rangle &= \sum_{r_i < \infty} \langle (I_{\{T=r_i\}} \mathbb{1}) \Theta_{r_i} x, (I_{\{T=r_i\}} \mathbb{1}) \Theta_{r_i} x \rangle \\ &= c \langle x, x \rangle \quad \text{où} \quad c = \sum_{r_i < \infty} \langle I_{\{T=r_i\}} \mathbb{1}, \mathbb{1} \rangle = \langle I_{\{T < \infty\}} \mathbb{1}, \mathbb{1} \rangle. \end{aligned}$$

On voit donc que si T est essentiellement fini ($c=1$), Θ_T est une isométrie. Si T n'est pas essentiellement infini, $\Theta_T / \sqrt{c}$ est une isométrie. Enfin, si T est essentiellement infini, Θ_T est l'opérateur nul.

Dans les trois cas, l'image de Θ_T est un sous-espace fermé de $\mathfrak{H}$, que l'on notera $\mathfrak{H}_T$.

Exemple 3 : opérateurs de Weyl. Fixons deux fonctions sur $\mathbb{R}_+$, $h \in L^2$ et ϕ réelle. Introduisons l'opérateur de Weyl

$$(9) \quad W(h, \phi) \mathcal{E}(f) = \exp(-\frac{1}{2} \|h\|^2 - \langle h, e^{i\phi} f \rangle) \mathcal{E}(e^{i\phi} f + h)$$

(c'est l'identité si $h=\phi=0$), et posons $W_t = W_t = W(h_t, \phi_t)$, et aussi

$W_{[t]} = W(h_{[t]}, \phi_{[t]})$ - cela forme deux processus d'opérateurs unitaires, le premier adapté, le second non ; leurs arrêts sont notés W_T et $W_{[T]}$. Il est commode de poser aussi $W_{st} = W(h_{[s,t]}, \phi_{[s,t]})$ pour $s < t$.

LEMME 4. W_T et $W_{[T]}$ sont unitaires.

Démonstration. Nous traiterons seulement le premier opérateur. On a

$$W_T(W_T)^* = \sum_{i,j} I_{\{T=r_j\}} W_{r_j} W_{r_i}^* I_{\{T=r_i\}} = \sum_{i < j} + \sum_{i=j} + \sum_{i > j}.$$

Si $i < j$, $W_{r_j} W_{r_i}^* = W_{r_i} r_j$ est dans le futur de r_i , donc commute à $I_{\{T=r_i\}}$, et le terme correspondant est nul ; de même si $i > j$. Pour $i=j$, il reste $I_{\{T=r_i\}}$, et la somme est l'identité.

Regardons ensuite $(W_T)^* W_T = \sum_i W_{r_i}^* I_{\{T=r_i\}} W_{r_i}$. Le i -ième terme peut aussi s'écrire $W_{\infty}^* I_{\{T=r_i\}} W_{\infty}$, car $W_{\infty} = W_{[r_i]}$, son adjoint se factorise de même, et $W_{[r_i]} W_{r_i}^*$ se rejoignent à travers $I_{\{T=r_i\}}$ pour se réduire à I . Il reste donc $W_{\infty}^* (\sum_i I_{\{T=r_i\}}) W_{\infty} = I$.

II. TEMPS D'ARRET QUELCONQUES

1. Commençons par une remarque de probabilités classiques : nous cherchons à arrêter à un instant T une courbe (z_t) de l'espace de Fock, c'est à dire en fait une famille de classes de v.a.. Même lorsque T est un t. d'a. ordinaire, ceci n'est pas une opération bien définie : il faut choisir une version du processus (z_t) à ensemble évanescent près. Cela n'a pas de sens sur l'espace de Fock, mais Parthasarathy et Sinha tournent la difficulté en se limitant à des courbes qui admettent une représentation

$$(10) \quad (x_t) \text{ est une martingale fermée}$$

$$z_t = x_t y_t \text{ où } (y_t) \text{ est une courbe adaptée au futur et bornée en norme.}$$

Le vecteur obtenu par arrêt se note correctement $\int I_{\{T \leq s\}} x_s y_s$, mais on peut le noter z_T ou $x_T y_T$, à condition de ne pas être trop naïf : z_T peut a priori dépendre de la représentation choisie, tandis que $x_T y_T$ peut dépendre des processus tout entiers, et non seulement de x_T et y_T .

Plus généralement, nous chercherons à définir certaines intégrales d'arrêt du type $\int W_s^I \{T \leq s\} x_s y_s$, ou $\int W_{[s]}^I \{T \leq s\} x_s y_s$, les processus d'opérateurs de Weyl étant ceux de l'exemple 3 (cf. (9)). La notation suggérée par le cas discret pour une telle intégrale est $W_{T*} z_T$, ou $W_{[T*} z_T$, avec une mise en garde de plus : il serait naïf de croire sans justification que ceci représente l'opérateur W_{T*} appliqué à z_T .

Commentaire. Pour éviter au lecteur familier avec les processus de Markov une erreur commise par le rédacteur, rappelons que le futur $\mathcal{F}_t$ est engendré par les accroissements du processus (X_s) après t , et ne contient pas X_t . Les processus adaptés au futur utilisés en théorie des processus de Markov (par ex. Sém. Prob. VIII p. 310-315) ne sont pas adaptés au futur au sens du travail de Parthasarathy-Sinha.

2. Nous arrivons maintenant au résultat fondamental de passage en temps continu. En première lecture, il peut être bon de prendre les opérateurs de Weyl égaux à l'identité. Nous traitons le cas des opérateurs de Weyl W_t plutôt que W_t , parce que ces derniers sont traités explicitement par Parthasarathy-Sinha : nous démontrons donc quelque chose d'un peu nouveau. Toutefois, le cas des W_t est un peu plus simple (formellement, $I_{\{\text{Teds}\}}$ commute avec W_s , mais non avec W_s), de sorte que $W_{T \times T} = W_{T \times T}$: nous avons volontairement renoncé à utiliser ce petit avantage dans la preuve.

Première étape. Nous traitons le cas d'une martingale $x_t = \mathcal{E}(g_t)$, et d'un processus (y_t) adapté au futur, uniformément continu en norme sur $[0, \infty]$, ou seulement borné sur $[0, \infty]$ et continu sur $[0, \infty[$: cette petite extension est utile dans les applications. Nous allons établir la convergence forte des approximations discrètes $W_{T_n \times T_n}$.

Soient m, n deux entiers avec $m < n$; nous posons $T_m = R$, $T_n = S$ (cf. (1)). Les points $i2^{-n}$ ($i \leq 2^{2n}$ ou $i = +\infty$) sont notés s_i ; pour tout i , l'intervalle $[s_i, s_{i+1}[$ est contenu dans un intervalle de la subdivision moins fine, que nous notons $[r_i, r_{i+1}[$. Alors

$$\begin{aligned} W_{[S \times S]} &= \sum_{s_i < \infty} W_{[s_{i+1}]} I_{\{s_i \leq T < s_{i+1}\}}^{z_{s_{i+1}}} + W_{\infty} I_{\{2^{2n} \leq T\}}^{z_{\infty}} \\ W_{[R \times R]} &= \sum_{s_i < \infty} W_{[r_{i+1}]} I_{\{s_i \leq T < s_{i+1}\}}^{z_{r_{i+1}}} + W_{\infty} I_{\{2^{2n} \leq T\}}^{z_{\infty}} \end{aligned}$$

Notons D_{mn} la différence : il s'agit de montrer que pour m, n assez grands, $\langle D_{mn}, D_{mn} \rangle = \langle \sum_i \dots, \sum_j \dots \rangle$ est petit. Nous commençons par remarquer que les termes pour lesquels $i \neq j$ ont une contribution nulle. Pour le voir, considérons par exemple un terme

$$\langle W_{[s_{i+1}]} I_{\{s_i \leq T < s_{i+1}\}}^a, W_{[t_{j+1}]} I_{\{s_j \leq T < s_{j+1}\}}^b \rangle \text{ avec } i < j$$

où t_{j+1} représente r_{j+1} ou s_{j+1} , mais en tout cas est $\geq s_{i+1}$. Faisons passer tous les opérateurs du côté gauche, et remarquons que $W_{[t_{j+1}]} W_{[s_{i+1}]} = W_{[s_{i+1}]} W_{[t_{j+1}]}$ est adapté à $\mathcal{F}_{[s_{i+1}]}$, donc commute à $I_{\{s_j \leq T < s_{i+1}\}}$. Alors les deux projecteurs se tuent en duel, et le terme est nul.

Ne gardant que les termes avec $i=j$, et notant I_j le projecteur, le carré de $\|D_{mn}\|$ vaut

$$\Sigma_j \parallel W_{[s_{j+1}]}^{I_j(\mathcal{E}(g_{s_{j+1}}))y_{s_{j+1}}} - W_{[r_{j+1}]}^{I_j(\mathcal{E}(g_{r_{j+1}}))y_{r_{j+1}}} \parallel^2$$

Nous écrivons $W_{[s_{j+1}]}$ comme $W_{[r_{j+1}]} W_{s_{j+1}r_{j+1}}$: le premier unitaire figure dans les deux termes de la différence et n'affecte donc pas la norme. Le second unitaire étant adapté à $\mathcal{E}_{[s_{j+1}]}$ n'affecte que le terme $y_{s_{j+1}}$, tandis que I_j adapté à $\mathcal{E}_{s_{j+1}}$ n'opère que sur $\mathcal{E}(g_{s_{j+1}})$. A droite, nous écrivons $\mathcal{E}(g_{r_{j+1}}) = \mathcal{E}(g_{s_{j+1}}) \mathcal{E}(g_{s_{j+1}r_{j+1}})$, et I_j n'opère que sur le premier facteur. Reste donc

$$\Sigma_j \parallel I_j \mathcal{E}(g_{s_{j+1}}) \parallel^2 \parallel W_{s_{j+1}r_{j+1}} y_{s_{j+1}} - \mathcal{E}(g_{s_{j+1}r_{j+1}}) y_{r_{j+1}} \parallel^2 = \Sigma_j a_j b_j$$

Nous allons montrer que $\Sigma_j a_j$ est bornée indépendamment de m, n , tandis que $\sup_j b_j$ est petit pour m, n grands, à peu de chose près.

Introduisons la mesure bornée de fonction de répartition

$$G(t) - G(s) = \parallel I_{\{s < T \leq t\}} \mathcal{E}(g) \parallel^2$$

Alors a_j vaut $(G(s_{j+1}) - G(s_j)) / \parallel \mathcal{E}(g_{[s_{j+1}]}) \parallel^2$, et le dénominateur est borné inférieurement par $e^{-\parallel g \parallel^2}$. Cela règle le sort de $\Sigma_i a_i$.

Pour étudier b_j , on passe $W_{s_{j+1}r_{j+1}}$ de l'autre côté, en remarquant que $W_{s_{j+1}r_{j+1}}^*$ n'affecte que $\mathcal{E}(g_{s_{j+1}r_{j+1}})$. Posant $W_{s_{j+1}r_{j+1}}^* \mathcal{E}(g_{s_{j+1}r_{j+1}}) = w_j$, un calcul direct sur les transformations de Weyl montre que w_j est uniformément voisin de $\mathbb{I}$ en norme pour m et n grands. On écrit alors

$$b_j = \parallel y_{s_{j+1}} - w_j y_{r_{j+1}} \parallel^2 \leq 2(\parallel y_{s_{j+1}} - y_{r_{j+1}} \parallel^2 + \parallel (1 - w_j) y_{r_{j+1}} \parallel^2)$$

Le second terme s'écrit $\parallel (1 - w_j) \parallel^2 \parallel y_{r_{j+1}} \parallel^2$: il est uniformément petit.

Le premier terme est aussi uniformément petit si (y_t) est uniformément continue pour la structure uniforme naturelle de $[0, \infty]$. Si (y_t) est seulement continue sur $[0, \infty[$, il faut faire un peu plus attention : il faut partager les s_i entre les $s_i \leq A$ et les $s_i > A$, où A est choisi assez grand pour que $\Sigma_{s_i > A} a_i$ soit petite. Les détails sont laissés au lecteur.

REMARQUES sur la première étape. La convergence forte établie lorsque (x_t) est une exponentielle s'étend bien sûr aux combinaisons linéaires finies d'exponentielles.

La limite sera notée $W_{[T]^*} z_T$, $W_{[T]^*} x_T y_T$ ou $W_{[S]^I \{Teds\}} x_S y_S$.

Lorsque les opérateurs de Weyl sont égaux à $\mathbb{I}$ et $y_t = \mathbb{I}$, de sorte que $z_t = x_t$ a une représentation $x_0 + \int_0^t h_s dX_s$, on a $z_T = x_0 + \int_{\{T \geq s\}} h_s dX_s$ comme dans le cas discret.

Seconde étape. Dans cette étape, nous nous affranchissons de l'hypothèse de continuité sur (y_t) , nous passons des combinaisons linéaires finies de vecteurs cohérents à des martingales fermées quelconques, enfin nous établissons l'importante formule d'isométrie (11).

Rappelons que pour toute mesure spectrale H sur $[0, \infty]$ et tout couple (u, v) de vecteurs, il existe une mesure $\mu_H(u, v, ds)$ sur $[0, \infty]$ de masse au plus $\|u\| \|v\|$, définie par $\mu_H(u, v, A) = \langle I_{\{H \in A\}} u, v \rangle$.

Considérons deux courbes du type précédent, $z_t = x_t y_t$, $z'_t = x'_t y'_t$; les deux martingales sont donc pour l'instant des c.l. finies de vecteurs exponentiels, et s'écrivent $x_t = x_0 + \int_0^t h_s dX_s$, $x'_t = \dots$. On se propose de calculer $\langle W_{[T]_*} z_T, W_{[T]_*} z'_T \rangle$. Le résultat de convergence forte établi plus haut permet de commencer par le cas discret, puis de passer à la limite.

Si T est discret, prend les valeurs s_j , on pose $I_j = I_{\{T = s_j\}}$ pour simplifier. On a

$$\langle W_{[T]_*} z_T, W_{[T]_*} z'_T \rangle = \langle \sum_k W_{[s_k]_*} I_k z_{s_k}, \sum_j W_{[s_j]_*} I_j z'_{s_j} \rangle.$$

Les termes avec $j \neq k$ sont nuls, par le même raisonnement que plus haut. Dans les termes avec $j = k$, les $W_{[s_j]_*}$ disparaissent (unitarité), et comme $z_{s_j} = x_{s_j} y_{s_j}$ on a une factorisation

$$\sum_k \langle I_j x_{s_j}, I_j x'_{s_j} \rangle \langle y_{s_j}, y'_{s_j} \rangle.$$

On a d'autre part (en écrivant x au lieu de x_∞)

$$I_j x_{s_j} = I_j (x - \int_{s_j}^\infty h_s dX_s) = I_j x - \int_{s_j}^\infty I_j h_s dX_s$$

On a ensuite $\langle I_j x, \int_{s_j}^\infty I_j h'_s dX_s \rangle = \langle x, \int_{s_j}^\infty I_j h'_s dX_s \rangle = \int_{s_j}^\infty \langle h_s, I_j h'_s \rangle ds$.

D'où

$$\langle I_j x_{s_j}, I_j x'_{s_j} \rangle = \langle I_j x, x' \rangle - \int_{s_j}^\infty \langle I_j h_s, h'_s \rangle ds$$

Regroupant le tout, nous obtenons une expression qui ne présente plus une allure discrète :

$$(11) \quad \langle W_{[T]_*} z_T, W_{[T]_*} z'_T \rangle = \int_{[0, \infty]} \mu_T(x, x', ds) \langle y_s, y'_s \rangle - \int_0^{\infty} ds \int_0^s \mu_T(h_s, h'_s, dr) \langle y_r, y'_r \rangle$$

Remarquons que la masse totale de $\mu_T(h_s, h'_s, \cdot)$ est au plus $\|h_s\| \|h'_s\|$, et que $\int \|h_s\| \|h'_s\| ds \leq \|x\| \|x'\|$ d'après l'inégalité de Schwarz. Si les deux processus $(y_t), (y'_t)$ sont continus sur $[0, \infty]$, il n'y a aucune difficulté (par convergence étroite, puis convergence dominée) à vérifier que cette formule passe à la limite par l'approximation discrète.

Nous écrirons (11) sous la forme abrégée

$$(11') \quad \langle W_{[T]_*} z_T, W_{[T]_*} z'_T \rangle = \int \langle y_r, y'_r \rangle v_T(x, x', dr) \\ v_T(x, x', dr) = \mu_T(x, x', dr) - \int_0^s \mu_T(h_s, h'_s, dr).$$

Prenons $x=x'$, $y_t=y'_t=f(t)\mathbb{1}$ où f est une fonction complexe continue sur $[0, \infty]$: il vient que $v_T(x, x, \cdot)$ est une mesure positive de masse totale au plus $\|x\|^2$. Le procédé habituel de polarisation montre alors que la masse totale de $v_T(x, x', \cdot)$ est au plus $\|x\| \|x'\|$.

Prenons $x_t=\mathbb{1}$, $y_t=f(t)\mathbb{1}$, $W_t=I$; alors $W_{[T]*}^{x_T y_T}$ vaut simplement $f(T)\mathbb{1}$, et $v_T(\mathbb{1}, \mathbb{1}, \cdot)=\mu_T(\mathbb{1}, \mathbb{1}, \cdot)$ est simplement la loi de T dans l'état vide.

Première conséquence. Si $\|y_t\| \leq M$, on a $\|W_{[T]*}^{x_T y_T}\|^2 \leq M^2 \|x\|^2$, indépendamment du temps d'arrêt T . Soit x un vecteur quelconque, que nous approchons par une suite x^n de c.l. finies de vecteurs exponentiels. Il résulte de cette inégalité que $W_{[T]*}^{x_T y_T}$ converge uniformément en T . Nous désignons par $W_{[T]*}^{x_T y_T}$ la limite, et le résultat de convergence forte des approximations discrètes s'étend par convergence uniforme aux martingales fermées quelconques, ainsi que la formule (11)-(11').

Seconde conséquence. Laissant fixe x , et faisant varier (y_t) , on peut prolonger $W_{[T]*}^{x_T y_T}$ à tous les processus (y_t) mesurables bornés, ou simplement tels que $\|y_t\|^2$ soit $v_T(x, x, \cdot)$ -intégrable. La formule (11)-(11') est préservée dans cette extension.

Tout ce qu'on vient de dire pour $W_{[T]*}^{x_T y_T}$ s'applique de même à $W_{[T]*}^{x_T y_T} - y$ compris la formule (11)-(11') avec la même mesure v_T .

3. Dans cette section, nous allons nous occuper des ambiguïtés de notation mentionnées au début.

Nous commençons par prendre $W_t=I$. Est ce que $x_T y_T$ ne dépend que de x_T et y_T ? Est ce que z_T ne dépend que du processus (z_t) ?

a) Si le processus (z_t) est de la forme $(x_t y_t)$, avec (y_t) continu sur $[0, \infty[$, l'approximation discrète z_T^n converge en norme vers z_T , et ne dépend que du processus (z_t) lui-même. Donc (z_T) peut être défini indépendamment du choix d'une représentation particulière.

b) La relation $x_T=0$ entraîne $x_T y_T=0$ quel que soit (y_t) . En effet, $x_T=0$ entraîne $v_T(x, x, \cdot)=0$, et la mesure $v_T(x, x, \cdot)$ est positive. On applique alors (11').

c) La relation $y_T=0$ équivaut à $\|y_t\|=0$ $v_T(\mathbb{1}, \mathbb{1}, \cdot)$ -p.p.. Donc la propriété $(y_T=0 \Rightarrow x_T y_T=0$ pour tout x) a lieu si et seulement si toutes les mesures $v_T(x, x, \cdot)$ sont absolument continues par rapport à $v_T(\mathbb{1}, \mathbb{1}, \cdot)$. Ceci est une propriété du t. d'a. T , qui peut être ou ne pas être satisfaite.

- Cas où elle est satisfaite. On a $v_T(x, x, \cdot) \leq \mu_T(x, x, \cdot)$, donc il suffit que la loi de T dans tout état normalisé x soit absolument continue par rapport à la loi de T dans l'état vide $\mathbb{1}$. Par exemple, si T est un

t.d'a. de la filtration du brownien (Q_t) , ou des processus de Poisson $Q_t + ca_t^\circ$, ou plus généralement d'une interprétation de l'espace de Fock comme un espace $L^2(P)$, dans laquelle le vecteur vide $\mathbf{1}$ correspond à la fonction 1, on a

$$\langle I_{\{T \in A\}} x, x \rangle = \int_{T \in A} x^2 dP, \quad \langle I_{\{T \in A\}} \mathbf{1}, \mathbf{1} \rangle = P(A)$$

et le résultat est clair.

- Cas où elle n'est pas satisfaite. Au § I, 2, nous avons vu l'exemple du premier saut T du processus (a_t°) . Dans l'état vide, T est p.s. infini, et la loi $v_T(\mathbb{I}, \mathbb{I})$ est une masse unité à l'infini. D'autre part, dans tout état cohérent normalisé $\varepsilon(g)/e^{\langle g, g \rangle / 2}$, le processus (a_t°) est un processus de Poisson non homogène d'intensité $|g(t)|^2 dt$. Prenant g réelle pour simplifier, il est facile de calculer la mesure

$$\mu_T(\varepsilon(g), \varepsilon(g), dr) = \varepsilon_\omega(dr) + I_{\{r < \infty\}} g^2(r) \exp\left(\int_r^\infty g^2(u) du\right) dr$$

On calcule alors la mesure $v(\varepsilon(g), \varepsilon(g), dr) = \varepsilon_\omega(dr) + g^2(r) I_{\{r < \infty\}} dr$ et l'on peut constater qu'elle n'est pas absolument continue par rapport à $v(\mathbf{1}, \mathbf{1}, \cdot)$.

Cet exemple peut sembler peu convaincant, du fait que T est essentiellement infini, mais cette difficulté se lève en remplaçant T par $T \wedge t$, t fini.

d) Nous nous attaquons maintenant à une autre ambiguïté de notation : définir l'opérateur $W_{[T]*}$ ou W_{T*} , et examiner si $W_{T*} z_T$ est le résultat de W_{T*} appliqué à z_T . Nous préférons traiter $W_{[T]*} = W_{T*}$ pour gagner un signe !

Soit $\varepsilon(g)$ un vecteur cohérent. Le processus constant $\varepsilon(g) = z_t$ admet la décomposition $x_t y_t$ avec $x_t = \varepsilon(g_t)$, $y_t = \varepsilon(g_t)$. Si T est un t. d'a. discret prenant les valeurs s_i , on a

$$z_T = \sum_i I_{\{T=s_i\}} z_{s_i} = \left(\sum_i I_{\{T=s_i\}}\right) \varepsilon(g) = \varepsilon(g).$$

Par le résultat de convergence forte établi plus haut, on a encore $z_T = \varepsilon(g)$ pour un t. d'a. quelconque. Nous définissons alors $W_{T*} \varepsilon(g)$ comme l'intégrale d'arrêt $\int W_s I_{\{T \leq s\}} \varepsilon(g_s) \varepsilon(g_s)$. L'approximation discrète montre que W_{T*} est isométrique sur les c.l. finies de vecteurs exponentiels, donc W_{T*} se prolonge en une isométrie sur l'espace de Fock, limite forte des W_{T_n*} de l'approximation discrète.

Soit alors $z_t = x_t y_t$, où (x_t) est une martingale, et (y_t) est continu sur $[0, \infty]$, adapté au futur. Nous écrivons l'égalité (immédiate)

$$W_{T_n*} z_{T_n} = (W_{T_n*})(z_{T_n})$$

Puis nous passons à la limite : z_{T_n} converge en norme vers z_T , W_{T_n*}

converge fortement vers W_{T*} , donc l'égalité passe à la limite :
 $W_{T*}z_T = (W_{T*})z_T$. Enfin, on fait varier (y_t) en laissant fixes T et (x_t) pour prolonger l'égalité au cas où (y_t) est seulement mesurable et borné.

e) Nous terminerons cette section en montrant que W_{T*} et $W_{[T*}$ sont unitaires. Ce résultat est connu pour les t. d'a. discrets. D'autre part les W_{T_n*} convergent fortement vers W_{T*} ; il suffit donc de vérifier la convergence forte des adjoints $W_{T_n*}^*$: la limite étant nécessairement l'adjoint de W_{T*} , les relations $W_{T_n*}W_{T_n*}^* = W_{T_n*}^*W_{T_n*} = I$ passeront à la limite.

Comme toutes les normes d'opérateurs sont bornées par 1, il suffit de montrer que $W_{T_n*}^* \varepsilon(g)$ converge en norme pour tout vecteur cohérent $\varepsilon(g)$. Mais ceci est l'intégrale d'arrêt $\int I_{\{T_n \leq s\}} W_s^* \varepsilon(g_s) \varepsilon(g_{[s]}$, et le résultat fondamental de convergence en norme nous ramène à écrire le processus $z_t = W_t^* \varepsilon(g_t) \varepsilon(g_{[t]}$ sous la forme d'un produit $x_t y_t$ d'une martingale et d'un processus adapté au futur, borné et continu sur $[0, \infty]$. Or $W_t^* \varepsilon(g_t)$ est de la forme $c_t \varepsilon((e^{iq} g + p)_t]$, p appartenant à $L^2(\mathbb{R}_+)$, q étant réelle, et c_t étant une fonction continue de t . Il suffit donc de poser

$$x_t = \varepsilon((e^{iq} g + p)_t]) , \quad y_t = c_t \varepsilon(g_{[t]}).$$

4. Factorisation de l'espace de Fock à un t. d'a. T.

Nous allons maintenant récolter sans difficulté les principaux résultats de Parthasarathy-Sinha, étendant ainsi aux t. d'a. (presque) quelconques les résultats du §I, et obtenant aussi une << propriété de Markov forte >> sur l'espace de Fock.

a) Opérateurs E_T et $I_{\{T \in A\}} E_T$.

Soit $x_t = E_t x$ une martingale fermée, qui s'écrit aussi $x_0 + \int_0^t h_s dX_s$. Soit (T_n) l'approximation discrète de T (cf. (1)). Une première conséquence du résultat fondamental de la section précédente est la convergence forte de $E_{T_n} x$ vers une limite que nous notons $E_T x$. Ce résultat est en fait immédiat à partir de la formule (5), et l'on a en toute généralité

$$(12) \quad E_T x = x_0 + \int_0^T I_{\{T \geq s\}} h_s dX_s.$$

L'opérateur E_T est un projecteur. Si S et T commutent on a $E_S E_T = E_{S \wedge T}$. D'autre part, E_{T_n} commute à tout opérateur de la forme $f(T_n)$, où f est borélienne bornée. Comme pour $m < n$ T_m est une fonction de T_n , E_{T_n} commute à $f(T_m)$. Faisant tendre n vers l'infini, E_T commute à $f(T_m)$. Prenant f continue et faisant tendre m vers l'infini, E_T commute à $f(T)$.

Cela permet de définir les projecteurs $I_{\{T \in A\}} E_T$. Appliquant le résultat principal au processus $z_t = x_t y_t$ avec $y_t = f(t) \mathbb{1}$, f continue bornée, on voit sans peine que

$$(13) \quad f(T) E_T x = \int f(s) I_{\{T \leq s\}} x_s$$

pour f continue, puis f borélienne bornée.

On introduit alors les notations Φ_T ($\Phi_T]$) et $I_{\{T \in A\}} \Phi_T$, comme pour les t . d'a. discrets au paragraphe I.

b) Opérateurs Θ_T . Soit $y \in \Phi$, et soit $y_t = \Theta_t y$, processus adapté au futur, et uniformément continu sur $[0, \infty[$ ($\Theta_\varepsilon y$ converge en norme vers y lorsque $\varepsilon \downarrow 0$). Il y a deux conventions raisonnables à l'infini, mais aucune ne fait de Θ_∞ une isométrie ! Celle que nous avons prise au §I consiste à poser $\Theta_\infty y = 0$; celle de Parthasarathy-Sinha consiste à poser $\Theta_\infty y = \langle \mathbb{1}, y \rangle \mathbb{1}$ (justifiée par la convergence en moyenne de Cesaro, et par le fait que cet opérateur est la seconde quantification $\Phi(\Theta_\infty) = \Phi(0)$). Nous conservons ici la convention simple $\Theta_\infty y = 0$.

Quoi qu'il en soit, le résultat principal du n° précédent nous permet de définir $\Theta_T y = \int I_{\{T \leq s\}} \Theta_s y$, et la formule (11') nous donne

$$(14) \quad \langle \Theta_T y, \Theta_T y' \rangle = \langle y, y' \rangle v_T(\mathbb{1}, \mathbb{1}, [0, \infty[) = c \langle y, y' \rangle$$

La constante c vaut aussi $\langle I_{\{T < \infty\}} \mathbb{1}, \mathbb{1} \rangle$. Ainsi, comme pour les t . d'a. discrets

- Si T est essentiellement fini, Θ_T est une isométrie,
- Si T n'est pas essentiellement infini, $\Theta_T / \sqrt{c}$ est une isométrie
- Si T est essentiellement infini, $\Theta_T = 0$.

Dans les trois cas, nous noterons Φ_T l'image de Θ_T : c'est un sous-espace fermé de Φ .

Dans les deux premiers cas, si l'on pose $y_t = \Theta_t y$, la relation $y_T = 0$ entraîne $y = 0$; on a donc $x_T y_T = 0$ pour toute martingale (x_t) . La difficulté du c) de la section précédente ne se présente donc pas pour les processus homogènes.

c) Multiplication d'éléments de Φ_T et Φ_T .

Nous supposons ici que T n'est pas essentiellement infini, mais cette hypothèse ne sera pas suffisante pour obtenir un résultat parfait (cf. la remarque à la fin de la démonstration).

Soient $u \in \Phi_T$ et $v \in \Phi_T$. Par définition, il existe un $x \in \Phi$ tel que $u = E_T x$, et un $y \in \Phi$ unique tel que $v = \Theta_T y$. Nous poserons alors

$$(15) \quad uv = x_T y_T \quad \text{où } (x_t) \text{ est la martingale } (E_t x) \text{ et} \\ \text{où } (y_t) = (\Theta_t y)$$

Cela ne dépend que de x_T et y_T , autrement dit de u et v . En fait on

peut toujours choisir $x=u$, de sorte que l'approximation discrète nous donne

$$(16) \quad uv = u\theta_T y = \lim_n \sum_i (I_{\{s_i \leq T < s_{i+1}\}} E_{s_{i+1}} u) \theta_{s_{i+1}} y$$

la sommation étant étendue aux $s_i = i2^{-n}$ avec $i < 2^{2n}$. Notre première remarque sera que, si u appartient à $I_{\{T=\infty\}}^{\Phi_T}$, uv est toujours nul (car $E_{s_{i+1}}$ commute à $I_{\{\cdot\}}$, et $I_{\{\cdot\}} u = 0$). Donc le produit n'est intéressant que si $u \in I_{\{T<\infty\}}^{\Phi_T}$.

Soient ensuite $u, u' \in I_{\{T<\infty\}}^{\Phi_T}$, et $y, y' \in \Phi$. Nous avons d'après (11)

$$\langle u\theta_T y, u'\theta_T y' \rangle = \langle y, y' \rangle v_T(u, u', [0, \infty[)$$

Remplaçons ensuite y, y' par $\mathbb{1}$. L'approximation discrète (16) montre aisément que $u\theta_T \mathbb{1} = I_{\{T<\infty\}} u = u$. Par conséquent $v_T(u, u', [0, \infty[) = \langle u, u' \rangle$ et l'on obtient la formule

$$(17) \quad \langle u\theta_T y, u'\theta_T y' \rangle = \langle u, u' \rangle \langle y, y' \rangle \quad \text{pour } u, u' \in I_{\{T<\infty\}}^{\Phi_T}.$$

Nous allons montrer maintenant que l'espace fermé engendré par les $u\theta_T y$ ($u \in I_{\{T<\infty\}}^{\Phi}$, $y \in \Phi$) est exactement $I_{\{T<\infty\}}^{\Phi}$.

Tout d'abord, posons $h_i = (I_{\{s_i \leq T < s_{i+1}\}} E_{s_{i+1}} u) \theta_{s_{i+1}} y$ (cf. (16)); il est clair que $I_{\{T < s_{i+1}\}} h_i = h_i$, donc $I_{\{T < \infty\}} h_i = h_i$. Sommant sur i et passant à la limite en (16), on voit que $I_{\{T < \infty\}}(uv) = uv$.

Inversement, il nous faut montrer que les vecteurs de la forme $u\theta_T y$ sont denses dans $I_{\{T < \infty\}}^{\Phi}$ - et pour cela, qu'ils permettent d'approcher $I_{\{T < \infty\}}^{\Phi} \mathcal{E}(g)$ pour tout $g \in L^2(\mathbb{R}_+)$.

Nous remarquons d'abord que Φ_T est stable par les opérateurs de la forme $f(T)$, puisque ceux-ci commutent à E_T . Il en est de même de

$I_{\{T < \infty\}}^{\Phi_T}$. Nous écrivons alors

$$\mathcal{E}(g) = u_s v_s, \quad u_s = \mathcal{E}(g_s), \quad v_s = \mathcal{E}(g[s]) = \theta_{s_{i+1}} y_s$$

avec $y_s = \mathcal{E}(g^s)$, $g^s(t) = g(s+t)$. Alors, posant $s_i = i2^{-n}$ ($i \leq 2^{2n}$), $s_{2^{2n}+1} = \infty$

$$\begin{aligned} I_{\{T < \infty\}} \mathcal{E}(g) &= \int I_{\{T \in ds\}} u_s \theta_{s_{i+1}} y_s = \sum_i \int_{s_i}^{s_{i+1}} I_{\{T \in ds\}} u_s \theta_{s_{i+1}} y_s \\ &= \lim_n \sum_i \int_{s_i}^{s_{i+1}} I_{\{T \in ds\}} u_s \theta_{s_{i+1}} (y_{s_i}) \end{aligned}$$

Chacun des termes de cette somme est de la forme $I_{\{T \in A\}} u_T \theta_T y$, $A \subset [0, \infty[$; le résultat est prouvé.

REMARQUE. Pour que $I_{\{T < \infty\}}^{\Phi}$ soit égal à Φ , il ne suffit pas que T soit essentiellement fini, il faut que T soit fini, i.e. que le projecteur $I_{\{T=\infty\}}$ soit nul.

d) $\mathfrak{F}_T$ comme nouvel espace de Fock $\tilde{\mathfrak{F}}$.

Nous supposons ici que T n'est pas essentiellement infini. Alors $\Theta_T/\sqrt{c}$ est une isométrie de $\mathfrak{F}$ sur $\mathfrak{F}_T$, c étant la constante $\langle I_{\{T<\infty\}} \mathbf{1}, \mathbf{1} \rangle$. Nous préférons que Θ_T lui-même soit une isométrie, en modifiant le produit scalaire sur $\mathfrak{F}_T$: si u, v sont deux éléments de $\mathfrak{F}_T$, nous poserons

$$(18) \quad \langle u, v \rangle_{\sim} = c \langle u, v \rangle.$$

Nous poserons $\tilde{\mathcal{E}}(f) = \Theta_T \mathcal{E}(f)$: comme on a $\langle \tilde{\mathcal{E}}(f), \tilde{\mathcal{E}}(g) \rangle_{\sim} = e^{\langle f, g \rangle}$, ce choix de vecteurs exponentiels fait de $\mathfrak{F}_T$ un espace de Fock. On se propose d'expliciter les opérateurs fondamentaux de ce nouvel espace.

Remarquons que

$$(19) \quad \tilde{\mathcal{E}}(f) = \int_0^\infty I_{\{T \leq s\}} \mathcal{E}(\theta_s f), \quad \theta_s f(t) = f(t-s) I_{\{t > s\}}$$

Cela suggère que les nouveaux opérateurs de Weyl $\tilde{W}(h, \phi)$ pour le nouvel espace de Fock $\tilde{\mathfrak{F}} = \mathfrak{F}_T$ peuvent s'écrire formellement

$$\tilde{W}(h, \phi) = \int \tilde{W}_s I_{\{T \leq s\}} = \int I_{\{T \leq s\}} \tilde{W}_s, \quad \tilde{W}_s = W(\theta_s h, \theta_s \phi)$$

(dans la notation du §I, ceci vaut $\tilde{W}_T$ et $\tilde{W}_{T*}$). Pour vérifier cela, on commence par le cas discret : si T prend les valeurs s_i

$$(20) \quad \begin{aligned} \tilde{W}_T &= \sum_i I_{\{T=s_i\}} W(\theta_{s_i} h, \theta_{s_i} \phi) = \sum_i W(\theta_{s_i} h, \theta_{s_i} \phi) I_{\{T=s_i\}} = \tilde{W}_{T*} \\ \tilde{W}_T \tilde{\mathcal{E}}(f) &= \sum_i I_{\{T=s_i\}} W(\theta_{s_i} h, \theta_{s_i} \phi) \mathcal{E}(\theta_{s_i} f) \quad (\text{somme sur } s_i < \infty) \\ &= \sum_i I_{\{T=s_i\}} \exp\left(-\frac{1}{2} \|\theta_{s_i} h\|^2 - \dots\right) \mathcal{E}(\theta_{s_i} (h + e^{i\phi} f)) \end{aligned}$$

le coefficient $\exp(\dots)$ ne dépend pas de i , et vaut $\exp(-\frac{1}{2} \|h\|^2 - \langle h, e^{i\phi} f \rangle)$. La somme restante vaut $\tilde{\mathcal{E}}(h + e^{i\phi} f)$: c'est le résultat cherché.

Pour montrer la convergence forte de ces approximations discrètes $\tilde{W}_{T_n}$, il suffit de

$$\begin{aligned} \tilde{W}_s \mathcal{E}(g) &= \exp\left(-\frac{1}{2} \|h\|^2 - \langle \theta_s h, e^{i\theta_s \phi} g \rangle\right) \mathcal{E}(g e^{i\theta_s \phi} + \theta_s h) \\ &= x_s y_s \quad \text{où } x_s = \mathcal{E}(g_s) \\ &\quad y_s = \exp(\dots) \mathcal{E}(g_s e^{i\theta_s \phi} + \theta_s h) \end{aligned}$$

D'autre part, $\tilde{W}_{T_n} \mathcal{E}(g) = \int_0^\infty I_{\{T_n \leq s\}} \tilde{W}_s \mathcal{E}(g)$: d'après le résultat principal de convergence, les intégrales étendues à $[0, \infty]$ convergent vers l'intégrale analogue relative à T . Pour avoir le même résultat pour les intégrales sur $[0, \infty[$, il suffit de regarder la convergence des intégrales à l'infini, que nous laissons au lecteur.

Ayant établi la convergence forte, soit $u \in \mathfrak{F}_T$: on a aussi $u \in \mathfrak{F}_{T_n}$ ($E_{T_n} u = E_T u = E_{T_n} u = u$). Or dans le cas d'un t . d'a. discret, la formule (20) sous la seconde forme montre que $\tilde{W}_T u = I_{\{T < \infty\}} u$ si $u \in \mathfrak{F}_T$. Donc ici

$\tilde{W}_{T_n} u = I_{\{T_n < \infty\}} u$. Par passage à la limite, on a $\tilde{W}_T u = I_{\{T < \infty\}} u$.

En utilisant la propriété de factorisation, on voit que $\tilde{W}_T$, en tant qu'opérateur sur $I_{\{T < \infty\}}^{\Phi}$, est unitaire et adapté à $\Phi|_T$. En tant qu'opérateur sur Φ , il tue le sous-espace $I_{\{T = \infty\}}^{\Phi}$, et il n'est donc unitaire que si T est fini.