

Fluctuations du champ électromagnétique du vide ou réaction de rayonnement ?

Sybil G. de Clark

Résumé

Le fait que divers effets physiques généralement attribués aux fluctuations du vide puissent également s'expliquer par le champ de réaction au rayonnement suggère qu'il existe peut-être peu de preuves des fluctuations du champ du vide. Mais est-ce vraiment le cas ? D'où vient cette sous-détermination entre le champ du vide et le champ de réaction au rayonnement, et que peut-on proposer pour la lever ? Quelles autres preuves avons-nous que le champ du vide est au moins en partie responsable de ces effets ? En outre, les fluctuations du vide ont été attribuées aux relations d'incertitude (RI). Dans quelle mesure ces affirmations sont-elles justifiées, et quelle interprétation des RI les fluctuations du vide suggèrent-elles ?

Sommaire

1	Introduction	3
1.1	Les questions	3
1.2	Fluctuations du vide, particules virtuelles et énergies du point zéro : qu'est-ce que c'est ?	4
1.2.1	Théorie quantique des champs (TQP) vs théorie classique des champs et la mécanique quantique non relativiste (NRQM)	4
1.2.2	Énergies du point zéro et état du vide	5
1.3	Effets physiques considérés comme des preuves des fluctuations du vide	7
1.4	Champ source et réaction de rayonnement : de quoi s'agit-il ?	8
2	Ordre des opérateurs et ambiguïté dans l'interprétation physique	9
2.1	Ordre des opérateurs comme source de sous-détermination	9
2.2	Réactions à l'ambiguïté	10
3	Le théorème de fluctuation-dissipation (FDT)	11
3.1	Description	11
3.2	Exemples	12
3.2.1	Relation de Nyquist	12
3.2.2	Mouvement brownien	13
3.2.3	Charge rayonnante	13
3.3	Interprétations possibles	14
3.4	Interprétations du théorème de dissipation des fluctuations dans la littérature physique	18
3.4.1	Le FDT implique l'existence de fluctuations du vide	18
3.4.2	Le FDT n'implique pas l'existence de fluctuations du vide	22
4	Cohérence de la théorie quantique du rayonnement : nécessité du du champ du vide pour que les relations de commutation soient valables.	25
5	Discussion	26
5.1	Ordre des opérateurs	27
5.2	S'agit-il vraiment d'un champ de vide ?	29
5.3	Implications pour les relations d'incertitude	30
5.3.1	Préliminaire : interprétations des relations d'incertitude dans le contexte de la NRQM	30
5.3.2	Relations d'incertitude, fluctuations du vide et particules virtuelles	33
6	Conclusions	38
7	Annexe : Sous-détermination et ordonnancement des opérateurs RÉFÉRENCES	40

1 Introduction

1.1 Les questions relatives à l'

La question de savoir si les fluctuations du point zéro du champ électromagnétique sont physiquement réelles a fait l'objet d'un débat parmi les physiciens au cours du XXe siècle. Bien que l'opinion soit aujourd'hui largement en leur faveur, la force des preuves sur lesquelles repose cette opinion est encore parfois remise en question ⁽¹⁾ et il n'est pas clair a priori si certains des arguments qui s'y opposent ont encore une importance. En outre, les questions soulevées dans ces débats approfondissent notre compréhension du champ électromagnétique du vide.

Les preuves de l'existence des fluctuations du vide ont naturellement été recherchées dans des phénomènes physiques observables supposés les impliquer. Cependant, il a été démontré que la plupart d'entre eux, y compris le célèbre effet Casimir, peuvent également s'expliquer sans faire appel aux fluctuations du vide, mais plutôt au champ rayonné par la charge elle-même, généralement appelé « champ de réaction de rayonnement » ou « champ source » dans ce contexte. Il est très intéressant de noter que ces deux alternatives (c'est-à-dire le vide et le champ source) ne constituent pas des modèles totalement indépendants. Il est clair que l'on peut passer d'une image physique à l'autre par un changement spécifique de formalisme : selon la manière dont on choisit d'ordonner l'opérateur relatif au système et l'opérateur relatif au champ électromagnétique, on constate que le champ du vide et le champ de réaction de rayonnement ont des contributions différentes à l'effet.

Dans le même temps, les interprétations alternatives en termes de vide ou de champ source ont souvent été considérées comme étroitement liées au théorème de dissipation des fluctuations (FDT).

Les alternatives sont-elles le résultat du choix d'ordonner une instance de sous-détermination théorique, ou pouvons-nous trouver un fait concret permettant de déterminer laquelle est valable ? Quels arguments et critères ont été utilisés pour soutenir ces alternatives ? Quel rôle a joué la FDT dans ces débats ? Comment a-t-elle été interprétée et que peut-elle réellement nous apprendre sur les fluctuations du vide ? Après avoir passé en revue ces débats, où en sommes-nous aujourd'hui ? Quelle semble être actuellement la meilleure preuve en faveur des fluctuations du vide ?

Afin d'étudier ces questions, cet article est organisé comme suit. Le reste de cette section introductive est consacré à des clarifications préliminaires : nous discuterons d'abord de ce que l'on entend par champ du vide, énergies du point zéro et fluctuations du vide, puis nous passerons en revue les effets physiques qui ont été considérés comme des preuves des fluctuations du vide. Comme ceux-ci ont également été interprétés en termes de champ de réaction de rayonnement, nous examinerons brièvement ce dernier.

La deuxième section examine comment l'ordre des opérateurs se référant au champ électromagnétique et aux variables atomiques conduit à une sous-détermination dans l'interprétation des effets physiques précédemment étudiés : sont-ils dus aux fluctuations du vide, au champ source, ou aux deux ? Nous présenterons les résultats obtenus par les physiciens qui ont cherché à trancher la question en arguant que, malgré

¹ Voir notamment [1]

menant aux mêmes prédictions expérimentales, un ordre est supérieur aux autres.

La troisième section est consacrée au théorème de dissipation des fluctuations (FDT). Si sa signification a toujours semblé claire dans la limite classique à haute température, nous verrons qu'il n'en va pas de même dans la limite $T = 0$ qui implique des énergies de point zéro. La manière dont il a été interprété dans ce contexte, ainsi que les implications correspondantes pour la question qui nous occupe, seront ensuite examinées dans une revue de la littérature physique pertinente. La plupart des physiciens ont fait valoir qu'il fournit la preuve des fluctuations du vide, mais certains en ont douté, et Edwin Jaynes, en particulier, a souligné à un moment donné un point de vue tout à fait différent.

La quatrième section présentera un argument en faveur des effets du champ du vide. Elle montrera que la théorie quantique du rayonnement serait incohérente s'il n'y avait pas le champ du vide et s'il n'était pas lié d'une manière spécifique à la force de réaction du rayonnement. Cet argument a été associé à la FDT, mais il repose en fait sur la relation de commutation entre la position et la quantité de mouvement d'une particule.

Enfin, l'article se terminera par une discussion des preuves fournies dans les sections précédentes en faveur ou contre l'existence des fluctuations du vide. Il examinera également comment ces dernières sont liées aux relations d'incertitude (RI) et quelle interprétation de ces dernières elles suggèrent.

1.2 Fluctuations du vide, particules virtuelles et énergies du point zéro : de quoi s'agit-il ?

Le concept de fluctuations du vide apparaît dans la théorie quantique des champs (TQP), en relation avec l'énergie du point zéro du champ électromagnétique. La notion de particules virtuelles y est souvent associée, et les deux sont souvent présentées comme résultant des relations d'incertitude (RI).

1.2.1 Théorie quantique des champs (TQFT) vs théorie classique des champs et mécanique quantique non relativiste (NRQM)

Le postulat de base de la QFT est que la matière et le rayonnement peuvent être décrits en termes de champs quantiques.

La différence entre les champs quantiques de la QFT et leurs équivalents classiques réside dans le fait que les premiers sont « quantifiés », d'une manière analogue à la façon dont l'énergie d'une particule est quantifiée dans la mécanique quantique non relativiste (NRQM). La principale différence entre la QFT et la NRQM réside dans le fait que les entités fondamentales qui nous intéressent dans la QFT sont les champs, tandis que les particules sur lesquelles se concentre directement la NRQM sont des concepts dérivés de la QFT, dans le sens suivant. Un champ est, par définition, une entité qui s'étend dans tout l'espace et dont la valeur diffère généralement d'un point à l'autre ⁽²⁾. En QFT, un champ est défini pour chaque type de particule que l'on souhaite traiter (il existe un « champ photonique », un « champ fermionique » pour les électrons, un pour les positons, etc.). Ces entités étant les entités fondamentales qui nous intéressent, les états en QFT

² Et généralement entre les événements, c'est-à-dire les points dans l'espace-temps.

sont des états *d'un champ*. La notion de particule apparaît dans ce contexte comme une propriété du champ : l'état d'un champ est un « état multiparticulaire », qui spécifie combien de particules ont, par exemple, une impulsion d'une valeur donnée lorsque le champ est dans cet état.³

Dans cette description, les particules sont des quanta d'énergie localisés, que l'on peut considérer comme des « excitations locales » du champ correspondant. Par exemple, un photon de fréquence ω est une excitation du champ photonique d'énergie $\hbar\omega$. Maintenant, parler des particules comme d'excitation du champ semble certainement impliquer que le champ est considéré comme existant de toute façon, même si aucune excitation ne se produit. Et en effet, une telle situation est considérée comme un état du champ : son « état de vide ». Cet état est l'état fondamental du champ au même titre qu'une particule NRQM a un état fondamental : c'est son état d'énergie la plus basse.⁽⁴⁾

1.2.2 Énergies du point zéro et état du vide

Rappelons que dans le NRQM, une particule dans un puits d'énergie potentielle « oscillateur harmonique » ne peut avoir que des énergies de $(n + \frac{1}{2})\hbar\omega$.⁵ Cela implique notamment que son énergie minimale

n'est pas nulle, mais $(\frac{1}{2})\hbar\omega$, qui est appelée son *énergie du point zéro*. Celle-ci est souvent décrite comme une conséquence des relations d'incertitude (ci-après « RI ») entre la position et la quantité de mouvement, qui impliquent des incertitudes correspondantes pour les énergies potentielle et cinétique, respectivement. Ce qui motive ce point de vue sera discuté ci-dessous, dans la section 5, ainsi que d'autres questions relatives aux RI.

Tout comme il existe des règles de commutation de la position et de la quantité de mouvement dans le NRQM, il existe dans la QFT des relations de commutation entre le champ quantique et la quantité de mouvement canonique⁶. Ces dernières entraînent à leur tour des relations de commutation entre le champ

³ D'autres propriétés intéressantes sont l'énergie, le spin ou la polarisation.

⁴ Cela ignore de nombreuses questions sérieuses liées à l'interprétation particulière des QFT. En réalité, la validité d'une description particulière dans les QFT a fait l'objet de vifs débats dans la littérature philosophique. Il existe un certain nombre de préoccupations liées aux différentes exigences auxquelles les particules devraient satisfaire, notamment la dénombrabilité ou la localisabilité. Lorsque l'on exige que les particules soient dénombrables et obéissent aux conditions d'énergie relativiste, nous avons des raisons de penser qu'une interprétation particulière de la représentation de Fock des champs *libres* est possible. Les vecteurs propres de l'opérateur du nombre total ont des propriétés appropriées pour les états contenant un nombre défini de particules. L'état de vide, sans particule, étant invariant sous le groupe de Poincaré, il semble identique à tous les observateurs inertiels ou $e\sqrt{}$ et surtout, les valeurs propres de l'énergie d'un et les états multi-particules sont $n^2 + m^2$, en accord avec la relativité. Cependant, cela n'est plus le cas obtient pour les champs en interaction. Il existe de sérieux problèmes, tant lorsqu'on tente de représenter le champ en interaction à l'aide de la représentation de Fock du champ libre (aucun état ne peut être interprété comme contenant zéro quantum, pour des raisons liées au théorème de Haag), que lorsqu'on cherche plutôt une représentation différente de l'espace de Hilbert (les tentatives impliquent des expressions qui ne sont pas covariantes relativistes ou qui ne garantissent plus que les états à une particule ont l'énergie relativiste souhaitée). De plus, même sans tenir compte des interactions, l'effet Unruh, selon lequel un observateur accélérant dans le vide détecterait des quanta de Rindler, pose également un problème de dénombrabilité [2], [3], [4], [5]. La localisabilité pose également des problèmes qui lui sont propres. Le fait que, dans un ensemble de conditions raisonnables, la probabilité de trouver une particule dans un ensemble spatial quelconque soit nulle constitue un théorème d'impossibilité, et cette conclusion a été généralisée aux espaces-temps génériques, ainsi qu'à la localisation non précise. [6], [7]

⁵ C'est-à-dire dont l'énergie potentielle est proportionnelle au carré de la coordonnée de position.

⁶ Tout comme il existe une correspondance entre les crochets de Poisson des théories classiques *des particules* et les commutateurs de la NRQM, il existe également une correspondance entre les crochets de Poisson

opérateurs de création et d'annihilation de Poisson.⁷ Comme dans le NRQM, le caractère non commutatif de ces opérateurs conduit à nouveau à une valeur propre non nulle pour l'état fondamental,

C'est-à-dire ce qu'on appelle « l'état de vide » — ou simplement « le vide ». Et là encore, on retrouve une énergie de $\frac{1}{2} k\omega$. Cependant, alors que dans le NRQM, cette valeur représentait l'énergie d'une particule, dans la QFT, elle représente désormais une énergie à chaque point de l'espace. Une autre différence par rapport à la NRQM est que cette « énergie du point zéro du champ » n'implique pas simplement une valeur de ω , mais est une somme sur une gamme de valeurs de ω , qui a priori pourrait prendre des valeurs arbitrairement grandes. C'est ce qui conduit à l'énergie infinie de l'état de vide.⁽⁸⁾ Ces énergies de $\frac{1}{2} k\omega$ sont

considérés comme des demi-quanta, car dans la QFT comme dans le NRQM, les quanta d'énergie sont $k\omega$.⁹

Pour qu'une

particule soit présente, il faut une excitation du champ d'un quantum complet, de sorte que l'état de vide représente un état où aucune particule n'est présente, justifiant ainsi le terme « vide ».

Cependant, on affirme souvent qu'il fourmille de « particules virtuelles », associées aux demi-quanta. L'expression « particules virtuelles » est souvent utilisée de manière interchangeable avec celle de « fluctuations du vide ». Dans le contexte du vide

impliquant des variables de la théorie classique *des champs* et les relations de commutation canoniques de la QFT.

⁷ Comme nous trouvons également des relations de commutation entre les opérateurs d'élévation et d'abaissement dans le NRQM

⁸ L'hamiltonien pour un champ scalaire libre s'avère être donné par :

$$H = \sum_k \frac{1}{2} \omega (a(\vec{k})a^\dagger(\vec{k}) + a^\dagger(\vec{k})a(\vec{k})) \quad (1)$$

pour les particules, avec des termes similaires pour les antiparticules. Maintenant, si $a(\vec{k})$ et $a^\dagger(\vec{k})$ commutent, par définition, cela signifierait que les deux termes $a(\vec{k})a^\dagger(\vec{k})$ et $a^\dagger(\vec{k})a(\vec{k})$ étaient égaux. On peut montrer que $a^\dagger(\vec{k})a(\vec{k})$ est un opérateur de nombre, $N(\vec{k})$, c'est-à-dire que lorsqu'il agit sur un état, il renvoie le nombre de particules de momentum ($\vec{k}$) présentes dans cet état (de manière analogue à l'opérateur de nombre NRQM qui renvoie le nombre de quanta d'énergie d'une particule). On obtiendrait alors :

$$a(\vec{k})a^\dagger(\vec{k}) = a^\dagger(\vec{k})a(\vec{k}) + N(\vec{k}) \quad (2)$$

et lorsque H agit sur le champ vide, on obtient :

$$H|0\rangle = \sum_k \omega_k N(\vec{k})|0\rangle = \sum_k \omega_k \cdot 0|0\rangle = 0, \quad (3)$$

c'est-à-dire que l'état de vide aurait une valeur propre d'énergie égale à zéro.

Ce n'est pas le cas car $a(\vec{k})$ et $a^\dagger(\vec{k})$ ne commutent pas, et à la place $[a(\vec{k}), a^\dagger(\vec{k})] = 1$. On obtient alors :

$$a(\vec{k})a^\dagger(\vec{k}) = a^\dagger(\vec{k})a(\vec{k}) + 1 = N(\vec{k}) + 1 \quad (4)$$

$$a(\vec{k})a^\dagger(\vec{k}) + a^\dagger(\vec{k})a(\vec{k}) = 2a^\dagger(\vec{k})a(\vec{k}) + 1 = 2N(\vec{k}) + 1, \quad (5)$$

et lorsque H agit sur le champ du vide, nous trouvons :

$$H|0\rangle = \sum_k \omega_k N(\vec{k})|0\rangle + \frac{1}{2} \sum_k \omega_k |0\rangle = \frac{1}{2} \sum_k \omega_k |0\rangle. \quad (6)$$

La somme portant sur un nombre infini de valeurs de k, cette somme est infinie, et nous trouvons que l'énergie de l'état de vide est infinie.

Ce problème est souvent résolu en imposant que les opérateurs soient « ordonnés normalement », c'est-à-dire ordonnés de telle sorte que l'opérateur d'annihilation se trouve toujours à droite de l'opérateur de création. En imposant cela à l'équation (1), on obtient l'équation (3). Cependant, ce que nous faisons alors, c'est imposer arbitrairement que $a(\vec{k})$ et $a^\dagger(\vec{k})$ commutent.

⁹ L'hamiltonien implique toujours $\frac{1}{2} \omega N$, et non la moitié de celui-ci.

Dans l'état de « quasi-stationnaire », les « particules virtuelles » renvoient à l'idée que l'UR entre l'énergie et le temps, $\Delta E \Delta t \geq \frac{\hbar}{2}$, permet au champ d'avoir une excitation d'un quantum complet d'énergie pendant un temps très court.¹⁰ Ce temps est considéré comme trop court *en principe* pour que les particules puissent être observées, en raison de l'UR. Le rôle exact des demi-quanta dans cette explication n'est souvent pas explicité, mais on pense qu'ils apparaissent grâce à l'UR par analogie avec l'énergie de l'état fondamental de l'oscillateur harmonique NRQM.⁽¹¹⁾

La question de savoir si ce schéma a un sens et quelles implications il aurait pour notre compréhension de l'UR est abordée ci-dessous dans la section 5.

Enfin, mentionnons déjà une définition opérationnelle du champ électromagnétique du vide que nous rencontrerons plus loin. Dans certaines recherches sur le sujet, on utilise plutôt le concept de « champ libre », qui désigne « la solution de l'équation du champ homogène (sans terme source atomique), [qui] coïncide avec le « champ du vide » lorsqu'aucun photon [réel] n'est initialement présent ».¹²

De plus, la manière dont le vide est identifié comme tel dans les dérivations passe souvent par sa densité d'énergie spectrale $\rho_0(\omega)$, c'est-à-dire sa densité d'énergie non seulement par unité de volume, mais aussi par intervalle de fréquence $d\omega$. Cette quantité est le produit de l'énergie d'un mode — qui, dans le cas du vide, est la moitié d'un quantum, $\frac{1}{2} \hbar \omega$, et de la densité des modes par unité de volume.¹³ L'expression qui en résulte est la suivante :

2

$$\rho_0(\omega) = \frac{k\omega^3}{2\pi^2 c^3} \quad (7)$$

1.3 Effets physiques considérés comme des preuves de fluctuations du vide

Les effets physiques attribués aux fluctuations du vide comprennent le décalage de Lamb, l'émission spontanée et la force de Casimir. Cependant, ils peuvent également s'expliquer par la réaction de radiation.

Le décalage de Lamb consiste en un écart par rapport à la prédiction de la mécanique quantique non relativiste concernant le spectre de l'atome d'hydrogène, selon laquelle les états $2s_{1/2}$ et $2p_{1/2}$ ¹⁵ devraient avoir la même énergie. Le décalage de Lamb est généralement interprété comme résultant de fluctuations de la position de l'électron dans le potentiel atomique. Ces fluctuations elles-mêmes sont attribuées à des interactions entre l'électron et les fluctuations du vide du champ électromagnétique⁽¹⁶⁾. Mais elles peuvent également être considérées comme résultant du champ électromagnétique associé à la réaction de rayonnement.

L'émission spontanée est généralement interprétée comme une forme d'émission *stimulée*.

¹⁰ D'où la notion de « fluctuations », c'est-à-dire de variations dans le temps.

¹¹ Particule dans un puits de potentiel d'oscillateur harmonique.

¹² [8], p. 1618. Il s'agissait d'une utilisation courante, que l'on retrouve notamment dans [9].

¹³ Également parfois appelée densité d'états.

¹⁴ La densité des modes est $\frac{\omega^2 d\omega}{\pi^2 c^3}$

¹⁵ le nombre quantique principal $n = 2$, le nombre quantique orbital $l = 0$ et 1 respectivement.

¹⁶ [10], pp.635-636 ; [11], pp.61-62.

sion, avec les fluctuations du vide comme origine de la perturbation.¹⁷ Avant cette explication, cependant, la source de la perturbation était attribuée au champ de réaction de rayonnement de l'électron ([12], pp.142-143).

La force de Casimir est une force d'attraction entre des plaques conductrices neutres. Elle est traditionnellement attribuée à la densité d'énergie des fluctuations du vide entre les plaques, qui est plus faible qu'à l'extérieur de celles-ci : cette différence de densité d'énergie diminue lorsque la distance entre les plaques est réduite. Il ne semble pas y avoir beaucoup de place pour la réaction de radiation dans ce tableau, mais la force de Casimir peut également être dérivée en considérant l'énergie de polarisation des atomes dans les plaques et la façon dont cette énergie change lorsque la distance entre les plaques varie. Cette polarisation peut être attribuée au champ de vide, mais aussi au champ de réaction de radiation.

Comme nous le verrons bientôt, il est possible de montrer que le champ responsable dépend de la manière dont on ordonne les opérateurs pour le champ et l'atome.

1.4 Champ source et réaction de rayonnement : que sont-ils ?

Contrairement au champ électromagnétique du vide, la réaction au rayonnement est un concept classique. La force de réaction au rayonnement, également connue sous le nom de force d'Abraham-Lorentz, est la force de recul exercée sur une charge rayonnante en raison de son propre rayonnement, tout comme la poussée est la force de recul exercée sur une fusée en raison des gaz d'échappement qu'elle émet, ou le recul est la force exercée sur un pistolet lorsqu'il tire une balle. Dans toutes ces situations, la conservation de l'énergie et de la quantité de mouvement ne s'applique pas au système plus massif qui nous intéresse (charge, fusée, arme à feu), mais au système plus large qui comprend ce qui est éjecté (rayonnement électromagnétique, gaz d'échappement ou balle, respectivement). Comme ces derniers emportent une partie de la quantité de mouvement du système total, la quantité de mouvement du système qui nous intéresse change également, de sorte que la quantité de mouvement totale reste la même. Cela implique qu'une force est exercée sur le système étudié par la partie éjectée. Ainsi, pour une charge rayonnante, la force de recul, ou « réaction au rayonnement », est exercée sur la charge par le rayonnement qu'elle émet. C'est pourquoi le champ concerné est appelé « champ source » : il est émis par une source.

c'est-à-dire la charge, par opposition au champ de vide.

La force de réaction au rayonnement est proportionnelle à $\ddot{x}$, la dérivée temporelle de la accélération de la charge :

$$F = - \frac{2}{3} \frac{e^2}{c^3} \ddot{x} \quad (8)$$

où e est la charge, c la vitesse de la lumière et x la position de la charge.¹⁸

¹⁷ L'émission spontanée est l'émission d'un photon par un atome, au cours de laquelle l'atome passe d'un état excité supérieur à un état excité inférieur, sans que ce processus soit dû à une perturbation visible de l'atome.

¹⁸[13], chapitre 28, pp. 6-7. Il existe des questions intéressantes concernant la force de réaction au rayonnement, sur lesquelles nous ne nous attarderons pas ici, mais qui méritent au moins d'être mentionnées. Elle rend compte de l'énergie qu'une charge perd lorsqu'elle rayonne, mais ce qu'elle conserve est la perte *moyenne*, sur toute une période d'oscillation d'un dipôle émetteur par exemple, et non la perte instantanée.

2 Ordre des opérateurs et ambiguïté dans l'interprétation physique de l'

Comme nous l'avons vu plus haut, les effets physiques attribués au champ du vide (décalage de Lamb, émission spontanée, force de Casimir) peuvent également être attribués à la réaction de radiation, c'est-à-dire à un champ source. Il a été démontré que la responsabilité de l'un ou l'autre de ces deux champs dépend de la manière dont on ordonne les opérateurs pour le champ et l'atome.

2.1 L'ordre des opérateurs comme source d'indétermination

Dans le contexte actuel, l'ambiguïté de l'image physique que dépeint le formalisme provient du fait que les quantités qui nous intéressent impliquent le produit de deux opérateurs, et qu'a priori, nous sommes libres de les ordonner comme bon nous semble. Pour comprendre cela, représentons par Q l'opérateur correspondant à la quantité que nous souhaitons trouver. Dans le cas le plus simple, Q est proportionnel au produit d'un opérateur pour le champ électromagnétique total, E , par un opérateur qui se réfère à une variable atomique, que nous appellerons P . C'est-à-dire :

$$Q \sim PE \quad (9)$$

E et P commutent à des instants égaux, nous sommes donc libres de les ordonner comme nous le souhaitons : les dérivations qui suivent donnent les mêmes résultats.¹⁹ Cependant, des choix d'ordre différents conduisent à des images physiques différentes, car E est la somme d'opérateurs qui *ne* commutent *pas* avec P — le champ de vide au point 0 E_0 et le champ source E_S :

$$E = E_0 + E_S, \quad (10)$$

La quantité d'intérêt, Q , est due aux deux contributions correspondantes, Q_0 et Q_S .²⁰ Si nous choisissons l'ordre $P E$, nous obtenons alors :

$$Q = PE = P(E_0 + E_S) = PE_0 + PE_S = Q_0 + Q_S \quad (11)$$

Et pour $E P$:

$$Q = EP = (E_0 + E_S)P = E_0P + E_SP = Q_0 + Q_S \quad (12)$$

¹⁹ Ceci est discuté par Dalibard *et al.* avant qu'ils ne traitent spécifiquement du décalage de Lamb et de l'anomalie de spin de l'électron.[8] En raison de l'application qu'ils envisagent, ils considèrent l'évolution temporelle d'un opérateur G , $\frac{dG(t)}{dt}$, plutôt qu'un opérateur arbitraire Q comme je l'ai fait ici.

²⁰ [8] désigne ces quantités par « gradient de $\frac{dG(t)}{dt}$ » et « $\frac{dG(t)}{dt}$ », respectivement, où le sous-potential $\frac{dG(t)}{dt}$ signifie « champ de vide » et « auto-réaction ».

Tout semble correct, mais bien que leur somme soit correcte, E_0 et E_S ne commutent pas avec P :

$$E_0 P \neq P E_0 \quad E_S P \neq P E_S. \quad ^{21} \quad (13)$$

Ainsi, Q_0 n'est pas la même chose dans l'équation (11) et dans l'équation (12) — et il en va de même pour Q_S bien sûr. Par conséquent, selon la façon dont nous choisissons d'ordonner P et E , l'effet total représenté par Q sera dû aux contributions de Q_0 et Q_S dans des proportions différentes — même si l'on obtient le même résultat pour Q lui-même, quel que soit le choix effectué.

Nous voyons que la sous-détermination apparaît parce que, dans la mesure où certains opérateurs commutent, nous obtenons des résultats finaux similaires, dans la mesure où les opérateurs qui les composent ne commutent pas, nous obtenons des descriptions différentes du champ responsable des effets.

2.2 Réactions à l'ambiguïté de l'

On trouve dans la littérature différentes attitudes face à cette sous-détermination. Il existe un consensus quant à l'équivalence formelle du choix de l'ordre. Cependant, certains chercheurs ont fait valoir que l'interprétation physique elle-même devrait servir de fondement pour considérer un ordre comme étant le bon. ^{Plus} précisément, ils estiment que chaque contribution distincte (telle que $Q_{(0)}$ et $Q_{(S)}$ ci-dessus) doit avoir une signification physique claire et que, pour cela, elle doit être décrite formellement par un opérateur hermitien. Seul l'ordre symétrique correspond à ce cas. ⁽²²⁾

Par conséquent, à ce stade, il semblerait que le fait d'imposer la condition que Q_0 et Q_S doit être hermitien implique un ordre symétrique.

Ce n'est toutefois pas vrai dans tous les cas, notamment lorsqu'il s'agit d'un modèle pertinent pour notre propos : l'atome à deux niveaux. ⁽²³⁾

En effet, dans ce cas, l'opérateur d'intérêt Q n'est pas donné par PE , mais par une somme de termes de la forme $E^- P^- + E^+ P^+$. ⁽²⁴⁾ Dans ce cas les choix d'ordre qui conduisent à des contributions relatives différentes de E_0 et E_S correspondent néanmoins

²¹ Une explication donnée pour cela est la suivante :

La liberté de modifier l'ordre des opérations en cours de calcul est limitée, car différents opérateurs peuvent acquérir des arguments temporels différents et parce que l'interaction du champ avec l'atome peut rendre difficile la séparation des différents degrés de liberté à un moment autre que l'instant initial. [9], p.157.

Il peut sembler étrange que des opérateurs se référant à un champ et à un atome, donc à des systèmes différents, puissent ne pas commuter. Cependant, si l'on considère par exemple la polarisation d'un atome, le fait qu'elle soit due au champ conduit à exprimer la polarisation en termes d'opérateurs de création et d'annihilation du champ.

²² [8]. Voir l'annexe pour la dérivation de ce résultat. L'ordre normal consiste à placer les opérateurs d'annihilation à droite des opérateurs de création, l'ordre anti-normal les place à gauche, et l'ordre symétrique consiste en une combinaison linéaire des deux.

²³ L'atome à deux niveaux est souvent utilisé pour discuter de l'émission spontanée et du décalage de Lamb, d'où son importance.

²⁴ $L^- E$ est la composante de fréquence négative du champ qui implique le seul opérateur de création seul, tandis que la composante de fréquence positive E^+ implique l'opérateur d'annihilation. Il en va de même pour P^- et P^+ .

à Q_0 et Q_S étant hermitiennes. Il ne suffit donc pas d'imposer que Q_0 et Q_S soient hermitiennes pour supprimer la liberté d'ordre et la sous-détermination qui y est associée. Il faut en outre exiger que Q soit écrit sous une forme qui implique uniquement des produits de E avec P ,²⁵ et non des produits de E^+ avec P^+ (ou $E^- P^-$) comme ci-dessus. Autrement dit, les opérateurs pour le système et pour les champs doivent être des opérateurs hermitien tout au long de la dérivation. L'ordre qui correspond

à cette exigence est à nouveau l'ordre symétrique. Par

conséquent, l'ordre symétrique est le seul ordre où :

- La contribution du champ de vide et la contribution du champ source à l'effet sont représentées par des opérateurs hermitiques.

- Les opérateurs pour le système et pour les champs sont des opérateurs hermitien tout au long de la dérivation.

Pour ces raisons, il a été avancé que la sous-détermination n'est que superficielle. Le raisonnement est que l'ordre symétrique devrait être préféré, car les observables sont représentés par des opérateurs hermitiques, de sorte que préférer l'ordre symétrique revient à imposer que les opérateurs aient une signification physique tout au long de la dérivation. Et l'ordre symétrique implique le champ du vide ainsi que le champ source — en fait, des contributions égales des deux ([8], [14]). Cependant, comme nous le verrons ci-dessous, ce raisonnement ne fait pas l'unanimité.

3 Le théorème de fluctuation-dissipation (, FDT)

3.1 Description

Outre l'hermiticité de l'opérateur, un autre élément a joué un rôle important dans la manière dont les physiciens ont envisagé les rôles relatifs des champs de vide et des champs sources : le théorème de fluctuation-dissipation (FDT). La relation entre les deux champs a été considérée comme un cas particulier de ce théorème ([15]).

Le FDT a été proposé en 1951 par Herbert B. Callen et Theodore A. Welton. Il généralise un résultat obtenu par H. Nyquist en 1928, qui relie les fluctuations de tension à la résistance dans les circuits électriques ([16]). Ce qui a fasciné Callen et Welton dans la relation de Nyquist et les a incités à la généraliser, c'est qu'elle ne se contente pas de relier des grandeurs physiques pour une situation physique donnée, mais deux processus différents : elle « relie une propriété d'un système en *équilibre* (c'est-à-dire les fluctuations de tension) à un paramètre qui caractérise un processus irréversible (c'est-à-dire la résistance électrique) » ([15], p. 34).

Callen et Welton ont considéré un système dissipatif arbitraire, c'est-à-dire un système capable « d'absorber de l'énergie lorsqu'il est soumis à une perturbation périodique dans le temps » ([15], p. 34). Leur FDT stipule alors que, pour un tel système, la force fluctuante qui s'exerce sur lui, F_f obéit à :

$$\langle F_f^2 \rangle = \frac{2}{\pi} \int_0^\infty R(\omega) E(\omega, T) d\omega, \quad (14)$$

²⁵ ou les produits des combinaisons appropriées de $E^+ - E^-$ et $P^+ - P^-$. [8], p.1625.

où :

$$E(\omega, T) = \frac{1}{2} \frac{k\omega + k\omega [e^{\frac{k\omega}{kT}} - 1]^{-1}}{2}, \quad (15)$$

est l'énergie du système, qui correspond à l'expression de l'énergie moyenne d'un oscillateur de fréquence naturelle ω , à la température T , et la « résistance généralisée » $R(\omega)$ caractérise la *dissipation*, donnée par ([15], p.37) :

$$R(\omega) = Re \frac{F_d}{\dot{Q}}, \quad (16)$$

où $\dot{Q}$ est la réponse du système et F_d la force dissipative.²⁶

Notez les deux termes de l'équation (15) : le premier terme correspond aux énergies du point zéro, et dans le cas $T = 0$, $E(\omega, T)$ se réduit à ce seul terme. En revanche, dans la limite des températures élevées, $E(\omega, T)$ peut être approximé par le deuxième terme qui est alors égal à kT .²⁷

L'application du théorème nécessite d'identifier la résistance généralisée, donc généralement la force dissipative et la réponse dans l'équation (16).

3.2 Exemples

3.2.1 Relation de Nyquist

En dérivant la relation qui a conduit à la généralisation de Callen et Welton sous la forme de la FDT, Nyquist voulait expliquer l'observation d'une force électromotrice (*f.e.m.*) dans les conducteurs en l'absence de potentiel externe, due uniquement à leur agitation thermique.

Lorsque nous identifions :

- la force fluctuante F_f avec la force électromotrice (*f.e.m.*)

- la résistance généralisée $R(\omega)$ avec la résistance électrique habituelle

et que nous considérons la limite de température élevée où $E(\omega, T) \sim kT$, la FDT (c'est-à-dire l'équation 14) devient ([15], p.37) :

$$\langle e.m.f.^2 \rangle = \frac{2}{\pi} kT \int_0^\infty R(\omega) d\omega, \quad (19)$$

en accord avec le résultat de Nyquist.

²⁶ [15], p.35. Plus précisément, ils ont considéré une force perturbatrice périodique de la forme :

$$F_f = F_{f0} \sin(\omega t). \quad (17)$$

Ils ont constaté que la puissance dissipée par le système était proportionnelle au carré de cette perturbation, et ont relié la constante de proportionnalité à des notions généralisées d'impédance $Z(\omega)$ et de résistance $R(\omega)$:

$$Puissance = F_f^2 \frac{R}{Z}. \quad (18)$$

Callen et Welton n'ont en fait pas fait de distinction entre F_f et F_d comme je l'ai fait ici : ils ont utilisé le même symbole pour les deux, suggérant que la force qui fluctue est la même que celle responsable de la dissipation. Cependant, dans les exemples qu'ils ont discutés comme applications de leur théorème, la distinction est intuitivement utile, comme nous le verrons ci-dessous.

²⁷ Cela peut être vu en utilisant le développement binomial du deuxième terme avec $kT \ll L\omega$, qui sert de condition pour la limite de température élevée.

3.2.2 Mouvement brownien

En appliquant le FDT au mouvement brownien (les mouvements aléatoires des petites particules dans un fluide), on retrouve la forme connue de la force fluctuante sur les particules. Dans ce cas, les identifications suivantes doivent être faites :

- F_f : force(s) due(s) au fluide lorsque ses molécules heurtent le système,
 - F_d : force visqueuse, responsable de la dissipation, et égale à $-\eta v$, où v est la vitesse de la particule et $\eta = 6\pi\alpha$ de viscosité, rayon de la particule,
 - réponse Q : vitesse v des particules,
- par conséquent, d'après l'équation (16) :
- la résistance généralisée $R(\omega)$ est égale à $-\eta$.

Ensuite, dans la limite de température élevée $E(\omega, T) \sim kT$, la FDT (c'est-à-dire l'équation (14)) devient : ²⁸

$$\langle F^2 \rangle = 2 \int_0^\infty \pi k T \eta \, d\omega, \quad (20)$$

en accord avec la force connue responsable du mouvement brownien.

3.2.3 Charge rayonnante

Les deux exemples ci-dessus donnent une idée intuitive du FDT, notamment parce qu'ils sont classiques, puisque nous avons pris dans les deux cas la limite de température élevée. Cependant, ce qui nous intéresse davantage dans le cadre de notre réflexion actuelle sur les champs électromagnétiques, c'est que Callen et Welton ont également appliqué leur théorème au rayonnement émis par une charge en accélération. Plus précisément, ils ont étudié un oscillateur dipolaire physique ([15], pp. 37-38).

Ils ont cherché à trouver la force responsable des fluctuations que leur théorème implique, sachant qu'un effet dissipatif se produit en raison de la force de réaction du rayonnement, ²⁹ responsable de la perte d'énergie de la charge. Pour ce faire, ils ont dû à nouveau déterminer la valeur de la résistance généralisée $R(\omega)$ dans la situation considérée.

Dans ce cas, les identifications pertinentes sont les suivantes :

- F_f : la force fluctuante exercée sur la charge x , où $x = \frac{p_0}{e} \sin(\omega t)$ avec P le ₀
 - F_d : la force de réaction au rayonnement, ³⁰
 - Q : la vitesse de la charge $\dot{x}$, par
- conséquent, d'après l'équation (16)
- la résistance généralisée $R(\omega)$ est égale à $-2e^2 \omega^2 \cdot \frac{1}{3c^3}$

Ce que la FDT implique, c'est qu'il doit exister sur la charge une force $F_f = eE_x$, de forme aléatoire et fluctuante, de la forme :

$$\langle F^2 \rangle_f = \langle e^2 E^2 \rangle_x = \frac{2}{\pi_0} \int_0^\infty E(\omega, T) \frac{2e^2 \omega^2}{3c^3} d\omega \quad (21)$$

Cette fois-ci, Callen et Welton n'ont pas utilisé la limite de température élevée pour obtenir un résultat classique, mais ont conservé la forme exacte de $E(\omega, T)$

(équation (15)), y compris ²⁸ $\langle F^2 \rangle = 2 \int_0^\infty R(\omega) E(\omega, T) d\omega$

terme d'énergie du point zéro $\frac{1}{2} \hbar \omega$ qui se réduirait à $T=0$. Ils ont noté que la densité d'énergie d'un champ de rayonnement isotrope est directement liée au champ électrique fluctuant par :

$$\text{densité d'énergie} = \frac{\langle E^2 \rangle}{4\pi} = \frac{1}{4\pi} \langle E_x^2 \rangle \quad (22)$$

et donc la FDT conduit directement à :

$$= \text{de densité d'énergie} = \frac{3}{4\pi e^2} \int_0^\infty \frac{1}{\pi^2 c^3} \frac{1}{2} \hbar \omega + \frac{\hbar \omega}{2} \exp\left(-\frac{\hbar \omega}{kT}\right) \omega^2 d\omega. \quad (23)$$

Cette densité d'énergie implique clairement deux contributions :

- la loi de Planck, donnée par le deuxième terme :

$$\frac{1}{\pi^2 c^3} \int_0^\infty \frac{\hbar \omega^3}{\exp\left(\frac{\hbar \omega}{kT}\right) - 1} d\omega, \quad (24)$$

- le terme d'énergie du point zéro, auquel la densité d'énergie se réduit dans la limite de $T \rightarrow 0$ (mais qui est également présent à des températures non nulles) :

$$\int_0^\infty \frac{\hbar \omega^3}{2\pi^2 c^3} d\omega. \quad (25)$$

Rappelons que la densité d'énergie spectrale du vide $\rho_0(\omega)$ a la forme suivante :

$$\rho_0 = \frac{\hbar \omega^3}{2\pi^2 c^3}. \quad (26)$$

Ainsi, la FDT stipule que la force fluctuante due au terme d'énergie du point zéro implique une expression de même forme fonctionnelle que la densité d'énergie spectrale du vide $\rho_0(\omega)$.

3.3 Interprétations possibles de l'

En regardant les équations (23) et (25), on pourrait penser que les implications de la FDT pour l'existence de fluctuations du vide du champ électromagnétique sont assez simples. Les fluctuations dans ce champ se produisent sous forme d'effets thermiques à une température non nulle, mais ne disparaissent pas à $T=0$ lorsque le champ est dans son état fondamental. Bien sûr, compte tenu de sa dérivation, l'utilisation du résultat de la FDT pour une charge rayonnante comme preuve de l'existence de fluctuations du vide ne semble certainement pas exempte de circularité : nous avons utilisé une contribution du point zéro dans l'équation (15) pour $E(\omega, T)$, il n'est donc guère surprenant que nous en obtenions une dans notre résultat final. Nous en obtiendrions également une pour le circuit de Nyquist et le mouvement brownien si nous n'avions pas choisi de la négliger en prenant la limite de température élevée. Cependant, lorsque nous examinons la littérature physique relative aux débats sur cette question, nous constatons effectivement des divergences quant aux implications de la FDT pour l'existence des fluctuations du vide. À cela s'ajoutent des différences dans la manière dont la FDT a été interprétée. Cela n'est d'ailleurs pas toujours clair dès le départ.

si le FDT (quelle que soit son interprétation) sert de prémisse aux discussions. Souvent, il semble plutôt que les considérations à la lumière desquelles il est interprété soient elles-mêmes avancées comme arguments pour ou contre l'idée des fluctuations du vide. C'est à ces interprétations que nous allons maintenant nous intéresser.

Les applications classiques du FDT à haute température évoquées ci-dessus nous permettent de comprendre intuitivement la signification de ce théorème. Peter Milonni a décrit ce dernier comme suit :

Si un système est couplé à un « bain » qui peut lui retirer de l'énergie de manière effectivement irréversible, alors le bain doit également provoquer des fluctuations. Les fluctuations et la dissipation vont de pair, l'une ne peut exister sans l'autre. ([12], p. 54)

Dans le cas du circuit électrique étudié par Nyquist, la force dissipative et la force fluctuante sont toutes deux exercées sur les électrons par le réseau ionique du métal, de sorte que le système peut être décrit comme étant constitué des électrons et du bain (le réseau).²⁹

Pour le mouvement brownien, l'effet dissipatif est dû à la force visqueuse exercée sur la particule par le fluide, et les forces fluctuantes sont également exercées par les molécules du fluide lorsqu'elles heurtent la particule. Dans ce cas, le système est donc la particule et le bain est constitué des molécules du fluide environnant.

Ceci est résumé dans le tableau 1, et nous notons que dans ces cas, les forces fluctuantes et dissipatives sont toutes deux finalement exercées par la même entité (c'est-à-dire le réseau métallique, les molécules du fluide).³⁰

La question qui nous intéresse est la suivante : si nous essayons maintenant de donner une description analogue de la FDT pour une charge rayonnante, que peut-elle nous apprendre, le cas échéant, sur le champ de réaction au rayonnement, les fluctuations du champ du vide et la relation entre eux ?

Ici, l'effet dissipatif est dû à la réaction de rayonnement, de sorte que la force dissipative est exercée par le champ de réaction de rayonnement. Il s'agit en réalité d'une simple question de définition, puisque le champ de réaction de rayonnement est un concept qui a été introduit afin de rendre compte de la perte d'énergie (moyenne) lorsqu'une charge rayonne.

Comme nous n'avons pas pris la limite de température élevée cette fois-ci, nous obtenons deux contributions aux forces fluctuantes : la force électromagnétique due au champ de rayonnement de Planck et la force électromagnétique F_{f0} . La question est de savoir si cette dernière peut être

²⁹En fait, Nyquist lui-même a pris en compte deux conducteurs. Dans son schéma, les fluctuations trouvent leur origine dans un conducteur et la dissipation dans l'autre : l'effet dissipatif est la chaleur générée dans ce dernier, que Nyquist a identifiée comme étant due aux fluctuations *de la force électromotrice* dans le circuit, qui sont finalement attribuées à l'agitation thermique dans le premier conducteur. Mais naturellement, la distinction entre ces deux conducteurs est artificielle, destinée à associer différentes parties d'un seul et même système, à savoir le circuit, aux deux processus distincts de fluctuation et de dissipation, afin de clarifier les concepts lors de l'examen de ces processus. En réalité, la fluctuation et la dissipation se produisent toutes deux dans l'ensemble du circuit résistif.

³⁰ Cette source commune des forces est vraisemblablement la raison pour laquelle Callen et Welton n'ont pas fait de distinction entre F_f et F_d dans leur dérivation, les représentant tous deux par F . La raison pour laquelle j'ai fait la distinction entre F_f et F_d est qu'ils jouent des rôles différents, apparaissant dans le contexte de processus qualitativement différents : F_d , qui est dissipative, augmente directement l'entropie, alors que F_f ne le fait pas — elle se rapporte à une situation d'équilibre dans la mesure où sa valeur attendue est nulle.

	Système	Bain	Résistance généralisée	Force dissipative F_d	Force fluctuante F_f
Circuit étudié par Nyquist	électrons	réseau ionique du métal	résistance électrique	• origine : due à agitation thermique • exercée par : le métal réseau	• origine : fluctuations de la force électromotrice dues à agitation thermique • exercée par : le métal réseau
mouvement brownien	particules	molécules du fluide	viscosité	• origine : force visqueuse due à l'agitation thermique • exercée par : le fluide	• origine : impacts provenant de l'agitation thermique des molécules • exercées par : le fluide

Tableau 1 : Application de la FDT au circuit de Nyquist et au mouvement brownien

attribué au champ électromagnétique du vide est ce qui est en jeu.

Maintenant, pour la contribution à température non nulle, la force dissipative et la force fluctuante sont exercées par la même entité : le champ qui exerce une réaction sur la charge lorsque celle-ci émet est le champ qu'elle émet (voir tableau 2). Le bain est donc le champ rayonné par la charge, le champ source.

Pour la contribution $T=0$, les choses deviennent toutefois plus intéressantes. La force dissipative pertinente est la même que pour la contribution à température non nulle, car Callen et Welton, pour obtenir leur résultat, l'ont remplacée uniquement par la force de réaction au rayonnement et aucune autre contribution. Il s'agit donc à nouveau du champ rayonné par la charge. Qu'est-ce que cela implique pour F_{f0} ? Si nous considérons que la FDT implique que la force fluctuante est exercée par la même entité que la force dissipative, nous sommes alors amenés à conclure que F_{f0} est due au champ de réaction au rayonnement, c'est-à-dire au champ source rayonné. On pourrait soutenir qu'il s'agit là de l'interprétation naturelle du FDT, dans la mesure où, si l'on examine les équations (23) et (25) et notre tableau, il semble clair que les effets thermiques et les effets du point zéro ont la même relation avec la force dissipative, et donc avec le champ de réaction au rayonnement. Ainsi, si la force fluctuante associée aux effets thermiques est due au champ rayonné, pourquoi celle associée aux effets du point zéro ne le serait-elle pas ? De plus, on considère généralement que la FDT implique un système et un seul bain qui exerce à la fois F_d et F_f .

Mais là encore, la plupart des applications de la FDT concernent la physique classique, d'où la limite de température élevée. Il ne faut donc peut-être pas tenir pour acquis que l'identité de F_d et F_f s'applique nécessairement aux contributions de l'énergie du point zéro. Et si l'on part du principe que la FDT n'exige pas les deux types de

	système	Bain	Résistance généralisée	Force dissipative F_d	Force fluctuante F_f
Charge rayonnante	la charge	?	$-\frac{2e^2\omega}{3c^3}$	• origine : réaction de radiation • exercée par : le champ de réaction de rayonnement $-\frac{2e^2}{3c^3}\ddot{x}$ c'est-à-dire le champ électromagnétique rayonné par la charge	• origine : - contribution de la température non nulle : rayonnement de Planck - = de $T=0$ contribution : effets du point zéro • exercée par : - contribution de la température non nulle : le rayonnement de Planck électromagnétique c'est-à-dire le champ électromagnétique rayonné par la charge - $T=0$ contribution F_{f0} : le champ électromagnétique au point zéro c'est-à-dire le champ électromagnétique à l'état fondamental, « vide » ?? Ou champ de réaction ?? Dans les deux cas : $\int_0^\infty \frac{\omega^3}{2\pi^2 c^3} d\omega$

Tableau 2 : La FDT appliquée à une charge rayonnante

forces à exercer par la même entité, on est alors tenté d'identifier F_{j0} comme étant due au champ du vide.

Il est intéressant de comparer ces deux alternatives aux exemples à haute température que nous avons évoqués, car étant classiques, elles sont intuitives. Il existe une différence intéressante entre le bain pour le circuit de Nyquist et le mouvement brownien d'une part, et pour le rayonnement de Planck d'autre part. Dans les deux premiers cas, le bain est un milieu qui existe indépendamment du système — et par conséquent, les forces dissipatives exercées sur le système sont clairement des forces externes, dont l'apparition n'a rien à voir avec le système. En revanche, dans le cas du rayonnement de Planck, le bain a été émis par le système lui-même, d'où le caractère de force de réaction de la force dissipative. Attribuer F_{j0} au champ du vide ou à la réaction de rayonnement implique une différence similaire, qui est sans doute la principale différence entre les deux champs : on considère que le champ du vide existe indépendamment du fait qu'une charge y accélère ou ait jamais accéléré.

3.4 Interprétations du théorème d' -dissipation des fluctuations dans la littérature physique

Lorsque l'on examine la littérature physique sur le sujet, les désaccords portent principalement sur la question de savoir si F_f et F_d peuvent être considérés comme exercés par la même entité, pour la contribution $T=0$, et si les forces fluctuantes F_f sont dues ou non aux fluctuations du champ du vide.³¹

De manière peut-être surprenante, la discussion ne s'est pas principalement concentrée sur la question de savoir si le FDT implique nécessairement que F_f et F_d sont tous deux exercés par la même entité, afin de tirer ensuite des conclusions concernant le cas spécifique d'une charge en accélération. En fait, il est souvent difficile de savoir si le FDT joue réellement le rôle de prémisse dans les discussions, et il semble être interprété à la lumière d'autres considérations, elles-mêmes avancées comme arguments pour ou contre l'idée que le champ du vide est responsable de F_{j0} . Nous examinons maintenant ces arguments et les critères qui ont été implicitement utilisés pour évaluer l'origine de F_{j0} .

3.4.1 La FDT implique l'existence de fluctuations du vide

3.4.1.1 Callen et Welton La plupart des physiciens qui ont travaillé sur cette question semblent avoir considéré que le théorème de fluctuation-dissipation impliquait l'existence d'un champ de vide, compris comme une entité distincte de la réaction de rayonnement. Cela semble certainement avoir été le cas de Callen et Welton eux-mêmes. En effet, après avoir dérivé l'équation (23), ils ont commenté leur résultat :

Le premier terme de cette équation donne la contribution infinie familière du « point zéro », et le second terme donne la loi de Planck ([15], p. 38).

Ils ont donc identifié F_f comme étant dû en partie au champ du vide, et le critère qu'ils ont utilisé comme base pour cette identification était l'expression fonctionnelle

³¹ Plus précisément, qu'ils soient en partie dus aux fluctuations du vide dans le cas général, et entièrement dus à celles-ci à $T=0$.

de leur densité d'énergie spectrale : l'expression de la densité d'énergie spectrale du champ fluctuant prédite par la FDT (c'est-à-dire l'intégrande de l'équation (25)) est identique à la forme connue de la densité d'énergie spectrale du champ du vide.

3.4.1.2 Peter Milonni Peter Milonni est un physicien qui a consacré une grande partie de ses recherches au rôle de la réaction de rayonnement et du champ du vide dans les divers effets physiques censés les impliquer, et notamment à la question de l'ordonnancement des opérateurs abordée dans la section précédente. Il s'efforce de rendre explicite son interprétation de la FDT. Le texte complet de l'extrait cité plus haut est le suivant :

Ce que nous avons ici ³² est un exemple de « relation fluctuation-dissipation ». D'une manière générale, si un système est couplé à un « bain » qui peut lui retirer de l'énergie de manière effectivement irréversible, alors le bain doit également provoquer des fluctuations. Les fluctuations et la dissipation vont de pair ; l'une ne peut exister sans l'autre. Dans l'exemple présent, le couplage d'un oscillateur dipolaire au champ électromagnétique comporte une composante dissipative, sous la forme d'une réaction de rayonnement, et une *composante fluctuante, sous la forme du champ du point zéro (vide)* ; étant donné l'existence de la réaction de rayonnement, le champ du vide doit également exister afin de préserver la règle de commutation canonique et tout ce qu'elle implique (³³).

Milonni considère donc que la FDT implique l'existence du champ du vide, en tant qu'entité physique distincte du champ de réaction de rayonnement, et attribue $\mathbf{F}_f$ au champ du vide.

Cependant, il ne fonde pas ces opinions uniquement sur la FDT. Sa mention de la « règle canonique de commutation » renvoie à une dérivation qui précède les remarques que nous venons de citer. Nous en discuterons en détail dans la section 4, mais il convient déjà de préciser qu'elle ne fait pas appel à la FDT. En outre, outre cette dérivation impliquant la règle de commutation, Milonni doit également avoir à l'esprit la question de l'ordre des opérateurs, qui a fait l'objet d'une grande partie de ses recherches. Contrairement à la FDT *en soi*, il s'agit d'une approche dans laquelle le champ de réaction de rayonnement et le champ de vide peuvent facilement se manifester comme des alternatives compatibles autant que mutuellement exclusives, puisque l'ordre peut conduire à des contributions des deux champs aussi facilement que de l'un d'entre eux.

3.4.1.3 Dennis Sciama Dans le chapitre qu'il a contribué à *The philosophy of vacuum* ([14]), Sciama a discuté d'un certain nombre d'effets impliquant « l'énergie du point zéro » et « le bruit du point zéro », y compris ceux du champ électromagnétique. Le plus intéressant ici est son traitement de la transition d'un atome à deux niveaux de l'état excité à l'état fondamental, accompagnée de l'émission d'un photon.³⁴ Comme

³² Ces remarques font suite à une discussion sur la relation de commutation pour un oscillateur dipolaire.

³³ [12], p.54. C'est moi qui souligne.

³⁴[14], pp. 148-150. C'est ce qu'on appelle généralement l'émission spontanée, mais Sciama préfère éviter cette expression, car la question de savoir si elle est appropriée est précisément celle qui se pose.

Dans les travaux de Callen et Welton, nous avons affaire à une charge rayonnante, mais cette fois-ci, il s'agit clairement d'une charge liée, c'est-à-dire un électron dans un atome.

Sciama soutient que c'est dans ce contexte spécifique que le désaccord entre les physiciens sur le champ du vide est apparu.³⁵ Il formule la question de manière légèrement différente : au lieu de se demander quel *champ* (rayonnement ou champ du vide) est responsable de la transition en jouant le rôle de F_0 , il se demande quel est *le processus* : émission stimulée ou émission spontanée. En fait, les deux approches sont similaires : « l'émission stimulée » est « une réaction à un rayonnement spontané émis par l'atome », tandis que ce qu'on appelle « l'émission spontanée » est considérée comme stimulée par le champ du vide.

Sciama soutient que les contributions relatives de l'émission stimulée et spontanée au taux de transition sont « exactement égales » et que « chaque contribution est physiquement réelle ».

Il ajoute que ces processus se produisent même pour un atome à l'état fondamental. Dans ce cas, il parle d'émission par opposition à *absorption* plutôt que d'émission stimulée par opposition à émission *spontanée* ⁽³⁶⁾. La manière dont il visualise la situation dans ce contexte suggère qu'il a à l'esprit la FDT ou des considérations connexes : il parle de l'atome « rayonnant vers et absorbant l'énergie des fluctuations du point zéro du champ électromagnétique du vide », et ajoute : le taux de rayonnement est déterminé par la puissance du bruit dans les fluctuations du point zéro du moment dipolaire atomique, et le taux d'absorption est déterminé par la puissance du bruit dans les fluctuations du point zéro du champ électromagnétique ([14], p. 149).⁽³⁷⁾

Ce qui nous intéresse, ce sont les raisons qui l'amènent à penser que les processus stimulés et spontanés (donc à la fois les champs de vide et les champs de réaction de rayonnement) contribuent « exactement à parts égales » à l'effet. Bien qu'il envisage la situation physique en termes de système dans un bain et qu'il ait certainement à l'esprit la FDT, son argument ne repose pas sur la FDT, mais sur des considérations liées à l'ordre des opérateurs, à savoir que seul l'ordre symétrique implique uniquement des opérateurs hermitien ([14], p. 148).

3.4.1.4 J. Dalibard, J. Dupont-Roc et C. Cohen-Tannoudji La plupart des conclusions de Sciama s'inspirent des travaux de J. Dalibard, J. Dupont-Roc et C. Cohen-Tannoudji, dans lesquels ces derniers affirment que l'ordre symétrique est plus correct physiquement que les autres ordres en raison du caractère hermitien qu'il implique pour les opérateurs. Outre cette considération, ils ont toutefois d'autres raisons de préférer l'ordonnement symétrique, qui sont liées à la FDT. En effet, ils affirment que lorsque l'on choisit l'ordonnement symétrique, on obtient des résultats physiquement intuitifs concernant les propriétés du système S et du réservoir R.

³⁵ Cela s'explique probablement par le fait que c'est précisément cet effet qui a été traité par Jaynes pour critiquer l'idée selon laquelle les fluctuations électromagnétiques du vide sont réelles, comme nous le verrons bientôt.

³⁶ Le changement de langage correspond naturellement au fait que les processus doivent désormais rendre compte non plus d'une émission, et donc d'une perte nette d'énergie par l'atome, mais de l'absence de celle-ci, et donc d'un équilibre dynamique entre perte et gain.

³⁷ La formulation de Sciama suggère qu'il considère que le bain est constitué uniquement du champ de vide, mais cela est sans doute dû au fait qu'il parle maintenant d'un atome dans son état fondamental, qui n'émet pas de rayonnement.

— c'est-à-dire le bain.³⁸

Leur traitement est censé s'appliquer de manière encore plus générale que la FDT dans la mesure où il est valable pour « un état arbitraire du réservoir », plutôt que d'exiger qu'il soit en équilibre thermique ([8], p.13). Parallèlement, la discussion est menée en tenant compte du décalage de Lamb. Ils dérivent des expressions pour le décalage énergétique du système dû au réservoir, (δE_{ar}) , et dû à l'auto-réaction, (δE_{asr}) .³⁹ Ils trouvent respectivement :

$$(\delta E_{ar}) = -\frac{1}{2} \sum_{ij} \int_{-\infty}^{+\infty} d\tau C_{ij}^{(R)}(\tau) \chi_{ij}^{(S)}(\tau) \quad (27)$$

$$(\delta E_{asr}) = -\frac{1}{2} \sum_{ij} \int_{-\infty}^{+\infty} d\tau \chi_{ij}^{(R)}(\tau) C_{ij}^{(S)}(\tau) \quad (28)$$

où :

- les C_{ij} font référence aux fonctions de corrélation, qui décrivent la dynamique des fluctuations dans le réservoir $C_{ij}^{(R)}$ ou dans le système $C_{ij}^{(S)}$
- les χ_{ij} font référence aux susceptibilités linéaires et représentent la réponse du système $\chi_{ij}^{(S)}$ ou réservoir $\chi_{ij}^{(R)}$ à une perturbation.

Ils soulignent que ces résultats appellent une interprétation physique simple :

On peut considérer que les fluctuations de R [le réservoir], caractérisées par $C_{ij}^{(R)}(\tau)$, polarisent S [le système] qui répond à cette perturbation caractérisée par $\chi_{ij}^{(S)}(\tau)$. L'interaction entre les fluctuations de R et la polarisation qu'elles induisent dans S a une valeur non nulle en raison des corrélations qui existent entre les fluctuations de R et la polarisation induite dans S . [...] Les mêmes commentaires peuvent être faits [pour l'équation (28)] que pour [l'équation (27)], les rôles [du système] et [du réservoir] étant intervertis. ([8], p.12.)

Ils résument ce modèle physique à l'aide du diagramme suivant :

³⁸ [8], p.1626. Ils posent la question suivante : « Est-il possible de comprendre l'évolution de S [le système] comme étant due à l'effet des fluctuations du réservoir agissant sur S , ou devons-nous invoquer une sorte d'auto-réaction, S perturbant R qui réagit à son tour sur S ? » et notent que « Pour l'émission spontanée, [...] le champ du vide, avec son nombre infini de modes, joue le rôle de R . »

³⁹ Plus généralement, ils considèrent le taux moyen de la variable du système G , où l'indice R indique que la moyenne est calculée sur l'ensemble des états du réservoir. Ils identifient les contributions du champ de vide et de l'autoréaction à ce taux respectivement représentées par $\frac{dG(t)}{dt} \big|_{\text{vide}}$ et $\frac{dG(t)}{dt} \big|_{\text{sr}}$.⁴⁰

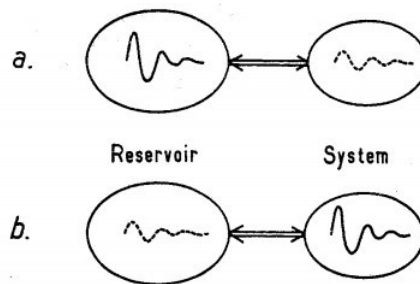


Fig. 2. — Physical pictures for the effect of reservoir fluctuations and self reaction. *a)* Reservoir fluctuations : the reservoir fluctuates and interacts with the polarization induced in the small system. *b)* Self reaction : the small system fluctuates and polarizes the reservoir which reacts back on the small system.

Fig. 1 : Interaction entre le système et le réservoir ([8], p.12.)

3.4.2 La FDT n'implique pas l'existence de fluctuations du vide.

3.4.2.1 Edwin Jaynes Edwin T. Jaynes a soutenu que des effets tels que le décalage de Lamb et l'émission spontanée pouvaient être dérivés de ce qu'il appelait la « théorie néoclassique », qui signifiait que la matière était effectivement quantifiée, mais pas le champ électromagnétique, en contraste évident avec la QED.⁴¹ Il interprétait bien le champ de réaction de rayonnement comme un opérateur, mais comme « un opérateur non pas sur l'espace de Hilbert de Maxwell d'un champ quantifié, mais sur l'espace de Hilbert de Dirac des électrons », et il ne voyait pas la nécessité de l'espace de Hilbert de Maxwell pour rendre compte des effets qui l'intéressaient ([17], pp. 10, 18).

Contrairement aux physiciens dont nous venons de parler, Jaynes a insisté sur le fait que la FDT n'implique pas l'existence de fluctuations du vide ([17]). Dans une conférence provocatrice donnée lors de la Conférence sur la cohérence et l'optique quantique en 1977, il a déclaré :

Cette indépendance de l'ordre initial n'est donc qu'un théorème de fluctuation-dissipation très simple, général et élégant ; mais permettez-moi de proposer une interprétation physique différente de celle qui est habituellement donnée. Cette interchangeabilité totale des effets de champ source et des effets de fluctuation du vide ne prouve pas que les fluctuations du vide sont « réelles ». Elle montre que les effets de champ source sont les mêmes que si des fluctuations du vide étaient présentes.

Depuis de nombreuses années, à partir de la relation établie par Einstein entre le coefficient de diffusion et la mobilité, les théoriciens ont découvert un flux constant de liens mathématiques étroits entre les problèmes stochastiques et les problèmes dynamiques. Il nous a fallu beaucoup de temps pour reconnaître

⁴¹ Cela en faisait une théorie semi-classique. Ce qui est particulièrement intéressant dans la théorie néoclassique de Jaynes, c'est ce qui l'a motivée. Peter Milonni l'a décrite comme « fondée sur la reconnaissance du fait que très peu de phénomènes dans la théorie non relativiste du rayonnement nécessitent réellement une quantification du champ pour être expliqués, et son objectif était d'explorer jusqu'où l'on pouvait aller sans quantification du champ et, éventuellement, d'indiquer la voie vers des alternatives à la QED ». [12], p. 148.

Je reconnais que la QED n'était qu'un autre exemple de ce phénomène.
 Mais dans un autre sens, je dois admettre que les fluctuations du vide
 sont, après tout, « des choses très réelles » ([17], p. 12).

Jaynes n'a clarifié ces déclarations contradictoires qu'en discutant d'un exemple
 — précisément, le cas d'un atome émettant spontanément. Il en a finalement conclu :

L'atome rayonnant interagit effectivement avec un champ
 électromagnétique dont l'intensité est prédite par l'énergie du point zéro ;
 mais il s'agit simplement du champ de réaction de rayonnement propre à
 l'atome ([17], p. 13).

Jaynes a essentiellement fait valoir que la présence dans les résultats de Callen et Welton de l'expression
 l'expression $\propto \omega^3$ ne devait pas être interprétée comme une preuve de l'existence de
 fluctuations du vide (équation (25)). Au cœur de son argumentation se trouve l'idée
 que ce qui est physiquement significatif n'est pas la densité *spectrale* d'énergie $\rho(\omega)$,
 mais $\rho(\omega)$ intégrée sur les fréquences appropriées, c'est-à-dire la densité d'énergie, W_f .
 Il convient de noter que dans les résultats de Callen et Welton, $\rho(\omega)$ apparaît dans
 une intégrale sur toutes les fréquences dans l'expression de la densité d'énergie, qui
 provient directement de $F(\rho)$. Le raisonnement de Jaynes consiste à dire que toutes
 ces fréquences ne sont pas physiquement présentes. Il a plutôt considéré $\rho(\omega)$ intégrée
 uniquement sur la petite bande passante de fréquences $\Delta\omega$, qu'il a jugée « efficace
 pour provoquer le rayonnement de l'atome » ([17], p. 12).

L'argument de Jaynes consiste à dériver et à comparer deux expressions :

- W_{0eff} , la partie « effective » de la densité d'énergie dont l'énergie spectrale est
 $\rho_o(\omega)$.
- W_{RReff} , la densité d'énergie du champ de réaction de
 rayonnement. Il a découvert qu'elles étaient données par la même
 expression :

$$\frac{(1)}{18\pi} \mu^{(2)}_c \frac{\omega}{c} \quad (6) , \quad (29)$$

où μ représente le moment dipolaire électrique de la transition atomique concernée et ω la
 fréquence naturelle de cette transition.⁴²

La dérivation de Jaynes montre que si l'on ne tient compte que de certains modes et
 composantes du champ (jugés physiquement significatifs), la densité d'énergie due à « ρ_o
 (ω) », W_{0eff} , est donnée par la même expression que la densité d'énergie du champ de
 rayonnement à la position de l'atome, W_{RR} . Il en a conclu que ce qui avait été
 interprété comme une contribution du champ du vide est en fait dû à la réaction de
 rayonnement.⁽⁴³⁾

À ce stade, il peut être utile de se demander ce que l'on entend par champ du vide dans le
 contexte des atomes rayonnants et des charges en accélération. Rappelons que, strictement
 parlant, le champ du vide est le champ électromagnétique dans son état fondamental, libre

⁴²La notation de Jaynes n'est pas tout à fait cohérente, car il utilise d'abord ω_0 pour représenter la
 fréquence naturelle de la ligne, puis simplement ω dans ses expressions finales pour les deux densités
 d'énergie. Cependant, son raisonnement et sa dérivation indiquent qu'il entend par là la fréquence naturelle
 de la ligne.

⁴³ Il s'agit très certainement de l'« autre sens » dans lequel il entendait que les fluctuations du vide sont «
 des choses très réelles », c'est-à-dire qu'elles sont en fait le champ de réaction au rayonnement.

de particules (réelles). Mais en réalité, dès qu'une charge s'accélère, elle est baignée dans son propre champ de rayonnement, les photons réels qu'elle émet. Des expressions telles que « fluctuations du vide » et « champ du point zéro (vide) » ne peuvent plus se référer *en soi* au vide entourant la particule dans ce contexte, puisque le champ quantique électromagnétique n'est plus dans son état de vide. Comment alors donner un sens à l'expression « champ du vide » dans ce contexte ?

Il semble juste d'interpréter ces expressions comme un raccourci pour « fluctuations au-dessus et en dessous de l'énergie (qui n'est plus nulle) du champ électromagnétique ». Ce qui peut justifier cette utilisation, c'est que ces fluctuations peuvent être considérées comme « les mêmes » que celles qui existeraient même s'il n'y avait pas de photons réels, dans le sens où elles résultent elles aussi de l'UR appliquée au champ.⁽⁴⁴⁾ Ainsi, lorsqu'on parle de la contribution du « champ du vide » dans une telle situation, on fait référence aux contributions au champ électromagnétique qui sont dues à l'UR. Et c'est précisément ce qui rend ce sujet si fascinant : il implique des implications possibles de l'UR. Étant le résultat de l'UR, les fluctuations du vide seraient présentes qu'il y ait ou non un rayonnement, c'est-à-dire qu'une charge ait déjà été en accélération ou non. Essentiellement, la question que Jaynes pourrait se poser est la suivante : si nous imaginons d'abord l'univers complètement vide de particules réelles (en particulier de photons et de charges en accélération), et que nous y plaçons un atome dans un état excité, cet atome va-t-il émettre un photon ? Une « émission spontanée » aurait-elle lieu, même dans ce cas ? C'est ce qu'il veut dire lorsqu'il demande si c'est le champ du vide qui est responsable de l'effet.⁽⁴⁵⁾

De toute évidence, Jaynes pense que la réponse est « non, un atome excité placé dans un univers dépourvu de photons réels n'émettrait pas de photon ». Il semble clair que la raison pour laquelle il pense ainsi est que, puisque les modes responsables de l'émission ont les mêmes fréquences que le rayonnement émis lorsque l'atome émet, il est plus économique de supposer que ces modes appartiennent en fait au champ dont nous savons qu'il existe de toute façon. Cependant, pour Jaynes, cette économie va au-delà d'une simple préoccupation pour le rasoir d'Occam. Elle élimine une difficulté cruciale, à savoir les infinis qui affligent la QED : pas de champ de vide, pas d'infinis :

Les chiffres fantastiques mentionnés précédemment disparaissent dès que l'on se rend compte que, pour expliquer l'émission spontanée, il n'est pas nécessaire que cette densité d'énergie soit présente dans tout l'espace, à tout moment, dans toutes les bandes de fréquences. Elle est produite automatiquement par l'atome rayonnant, mais de manière plus économique : seule la composante du champ qui est nécessaire, là où elle est nécessaire, quand elle est nécessaire et dans la bande de fréquences nécessaire ([17], p. 14).⁽⁴⁶⁾

⁴⁴ Plus précisément, ils résultent des relations de commutation entre les opérateurs de champ.

⁴⁵ Cela explique également pourquoi la question de savoir quel champ est responsable a été décrite en termes de spontanéité ou de stimulation du processus d'émission. Si les modes impliqués sont uniquement ceux rayonnés par l'atome, ils ne sont qu'une conséquence du processus, et non une cause, et celui-ci est alors véritablement spontané. En revanche, si l'émission est causée par le champ du vide, elle est bien sûr stimulée par ce dernier.

⁴⁶ Un aspect intéressant de la démonstration de Jaynes mérite d'être souligné : $\mathcal{L}$ s'annule dans sa dérivation de W_{eff} , comme le montre le résultat final. Cela se produit parce que son expression de la largeur de bande $\Delta\omega$ est proportionnelle à L^{-1} . Cela s'explique par le fait qu'il trouve $\Delta\omega$

Jaynes a donc interprété la FDT à partir des physiciens évoqués précédemment. Il pensait que le bain ne se composait que d'un seul champ, le champ de rayonnement, seul responsable de $\mathbf{F}_r$. Son point de vue n'était pas motivé par une compréhension indépendante de la FDT, mais par les autres considérations qui viennent d'être évoquées, qui l'ont ensuite conduit à son interprétation de la FDT.

4 Cohérence de la théorie quantique du rayonnement : nécessité du champ du vide pour que les relations de commutation soient valables.

Outre les arguments en faveur d'un ordre préférentiel et les considérations relatives à la FDT, d'autres preuves relatives au champ du vide impliquent la relation de commutation position-impulsion. En effet, une démonstration très intéressante consiste à montrer que pour $[X, P] = i\hbar$ pour une charge en accélération — qui subit bien sûr un amortissement par réaction de rayonnement au cours du processus —, la charge doit être entraînée par le champ du vide. Et non seulement les champs du vide et de rayonnement sont nécessaires, mais ils doivent également être liés l'un à l'autre d'une manière spécifique.⁴⁷

Peter Milonni l'a notamment démontré en considérant un électron non relativiste dans l'espace libre. L'équation du mouvement de Heisenberg pour l'électron est la suivante :

$$-\frac{d^2 X}{dt^2} - \frac{2e^2}{3mc^3} \frac{d^3 X}{dt^3} = -\frac{e}{m} E_0, \quad (30)$$

où nous reconnaissons le deuxième terme comme étant la force de réaction au rayonnement.

Si l'on résout cette équation pour X , puis que l'on utilise cette valeur pour trouver également $P = m\dot{X}$, il s'avère que le commutateur de ces opérateurs prend la forme suivante :

$$[X, P] = \frac{8\pi^2 i}{3m} \int_0^\infty d\omega \frac{\rho_0(\omega)}{[\omega^3(1 + \gamma^2 \omega^2)]}, \quad (31)$$

où $\gamma = \frac{2e^2}{3mc^3}$ et $\rho_0(\omega)$ est la densité spectrale du champ du vide. Lorsque nous puis substituons $\rho_0(\omega)$ par $\frac{\hbar \omega^3}{2\pi^2 c^3}$ dans cette expression, nous trouvons le commutateur égal à

être proportionnel au coefficient A d'Einstein (c'est-à-dire le coefficient d'émission spontanée, A_{21} étant la probabilité par unité de temps qu'un atome dans l'état d'énergie 2 émette spontanément un photon et passe à l'état d'énergie 1) et ce dernier étant proportionnel à $\hbar^{-1}$. En effet, Jaynes raisonne que bien que la densité spectrale $I(\omega)$ du rayonnement émis spontanément par l'atome ait un profil lorentzien, et donc aucune largeur bien définie, $I(\omega_0) \Delta\omega$ doit être égal à l'énergie totale rayonnée par l'atome, $\int I(\omega) d\omega$, et que ce dernier implique naturellement le coefficient A . Ce raisonnement suppose que la largeur de bande $\Delta\omega$, qui se réfère aux fréquences du champ fluctuant provoquant le rayonnement de l'atome, peut être identifiée à la largeur de la raie d'émission naturelle de l'atome. Jaynes ne défend pas fermement cette position, se contentant d'affirmer que c'est « vraisemblablement » le cas. Peut-être estimait-il que son résultat final, à savoir l'identité des expressions de $W_{(0) \text{ (eff)}}$ et $W_{(RR)}$, était suffisamment remarquable pour donner du crédit aux hypothèses qu'il avait dû formuler pour y parvenir.

⁴⁷ [18], p. 106 ; [19], p. 1322, 1323 ; [20].

⁴⁸ Notez que E_0 ici est la solution homogène de l'équation de Maxwell (Heisenberg) pour le champ électrique.

à ik comme il se doit.

Le point crucial est le suivant : si nous avons négligé le champ du vide dans l'équation (30) en fixant le côté droit à zéro, $[X, P]$ ne serait pas égal à ik . Ce ne serait pas non plus le cas si le champ responsable de la force motrice avait une densité d'énergie spectrale différente

de $\frac{L\omega^3}{2\pi c^3}$. Il est à noter que le spectre d'énergie doit être proportionnel au cube de la fréquence ω car la force de réaction au rayonnement implique la dérivée *troisième* de la position.⁴⁹

On peut peut-être se demander s'il n'y a pas une certaine circularité dans cet argument. Après tout, la cohérence avec la mécanique quantique qui est en jeu ici est la relation de commutation qui exprime formellement l'UR, et c'est encore l'UR qui a conduit à la prédiction des fluctuations du vide en premier lieu. Pourquoi devrions-nous maintenant considérer un autre argument *formel* basé sur le *même* « principe »⁽⁵³⁾ comme une preuve supplémentaire et (comme le suggère Milonni) encore plus forte des fluctuations du vide ? Cependant, la relation de commutation en jeu ici n'est pas celle qui donne naissance à l'énergie du point zéro du champ dans la QFT : il s'agit de la relation de commutation NRQM pour la charge. On peut maintenant dire qu'on utilise toujours le même « principe », même s'il se réfère maintenant à des variables différentes : position et impulsion, par opposition à champ et impulsion canonique. Cependant, comme on le verra ci-dessous, il y a une différence importante entre ces deux ensembles de variables, qui peut justifier de considérer l'UR NRQM comme plus solide que son équivalent QFT : le champ quantique et l'impulsion canonique ne sont pas observables, dans la mesure où leur valeur attendue est nulle.

5 Discussion

Les arguments avancés par Dalibard et al. et Peter Milonni invalident-ils le point de vue de Jayne ? Il y a au moins deux questions distinctes en jeu : premièrement, peut-on encore dire que l'ordre des opérateurs conduit à une sous-détermination en ce qui concerne

⁴⁹ L'inverse est également vrai : si nous avons négligé l'effet de réaction de rayonnement dans l'équation (30) en fixant le deuxième terme à zéro, $[X, P]$ ne serait pas non plus égal à ik .

Milonni donne également ailleurs un traitement analogue pour le cas d'un oscillateur dipolaire.⁵⁰

Il utilise ici l'approximation de faible amortissement $x'' \cong -\omega_0^2 x(t)$ dans l'équation pour le opérateur de position dans le cadre de Heisenberg $x(t)$:

$$x'' + \omega_0^2 x = \frac{e}{m} E_0(t) \quad (32)$$

ce qui donne :

$$x'' + \tau\omega_0^2 x' + \omega_0^2 x \cong \frac{e}{m} E_0(t).^{51} \quad (33)$$

Comme il le note :

Sans le champ libre $E_0(t)$ dans cette équation, l'opérateur $x(t)$ serait amorti de manière exponentielle, et les commutateurs tels que $[z(t), p_z(t)]$ tendraient vers zéro pour $t \gg (\tau\omega_0^2)^{-1}$. Cependant, avec le champ vide inclus, le commutateur est ik à tout moment.⁵²

En d'autres termes, la présence de $E_0(t)$ est nécessaire pour que les variables qui décrivent le dipôle prennent leur caractère quantique : sans ce champ, leur commutateur n'est plus proportionnel à ik ; il disparaît alors, de sorte que ces variables commutent comme elles le font classiquement.

⁵³ Dans la mesure où les UR sont souvent décrits comme un principe.

D'une part, il faut déterminer quel champ est responsable des effets qui nous intéressent, et d'autre part, lorsque nous interprétons une expression comme étant une signature du champ du vide, sommes-nous sûrs qu'il s'agit bien du champ du vide qui est en jeu ?

5.1 Ordre des opérateurs d'

Comme expliqué ci-dessus, l'ambiguïté quant à savoir si le champ du vide est définitivement impliqué dans les divers effets pertinents peut être considérée comme découlant de questions relatives aux propriétés de commutation des opérateurs, d'où notre liberté ou notre absence de liberté dans leur ordonnancement. Dans la mesure où les opérateurs commutent, nous obtenons des résultats similaires en termes de prédictions numériques, dans la mesure où ils ne commutent pas, nous obtenons des explications différentes quant au champ responsable des effets.⁵⁴

Cette sous-détermination ne permet pas au *champ source* de ne jouer aucun rôle dans tous les effets, car dans l'émission spontanée, tous les ordonnancements impliquent au moins une certaine contribution de celui-ci. Cependant, en soi, la liberté d'ordonnancement laisse ouverte la possibilité que le *champ du vide* ne joue aucun rôle — et donc que ces effets ne fournissent aucune preuve de son existence.

Comme expliqué ci-dessus, il a été avancé que cette sous-détermination n'est que superficielle, car l'ordonnancement symétrique devrait être préféré à tous les autres choix et implique des contributions égales du champ de vide et du champ source. Les partisans de l'ordonnancement symétrique soutiennent qu'il devrait être préféré, car seul cet ordonnancement permet de constater que :

- La contribution du champ de vide à l'effet et la contribution du champ source à celui-ci sont représentées par des opérateurs hermitiques.
- Tout au long de la dérivation, les opérateurs pour le système et pour les champs sont des opérateurs hermitien (seuls les produits d'opérateurs hermitien sont impliqués, et non les produits d'opérateurs de création, par exemple).

Ces deux considérations sont ensuite élevées au rang d'exigences à imposer aux dérivations afin qu'elles aient un sens physique. L'argument est que, dans le NRQM, les observables étant représentés par des opérateurs hermitiques, exiger que les opérateurs soient hermitiques tout au long du calcul garantit qu'ils ont une signification physique tout au long du calcul.

Il convient de noter qu'il peut y avoir une certaine circularité dans le raisonnement, dans la mesure où ces deux « exigences » n'apparaissent pas comme des desiderata purement pour eux-mêmes, la levée de la sous-détermination n'étant qu'une conséquence heureuse de celles-ci. En fait, cette dernière, c'est-à-dire la suppression de l'ambiguïté, sert également d'argument en faveur d'un ordre plutôt qu'un autre, et motive la recherche de critères pour le faire. Le fait que les opérateurs soient hermitien ou non dépend en fin de compte de l'ordre, et a été considéré comme un bon critère à suivre.

Ce qui peut peut-être être considéré comme une conséquence heureuse de ces exigences est la mesure dans laquelle elles conduisent à une description intuitive et classique de

⁵⁴ Rappelons que ce sont les opérateurs du système et du champ *total* qui commutent, de sorte que quel que soit l'ordre dans lequel nous les classons, nous obtenons les mêmes résultats finaux globaux, mais que les opérateurs du vide et du champ source ne commutent pas séparément avec les opérateurs du système, de sorte que différents choix d'ordre impliquent différentes contributions des champs vide et source.

la relation entre le système et le champ de vide, selon laquelle on peut dire qu'ils se polarisent mutuellement sur la base des équations (27) et (28).

Cependant, tous les chercheurs ne s'accordent pas à dire que ces considérations justifient de considérer l'ordre symétrique comme *correct*, ou simplement *plus correct*, que les autres choix possibles. En fin de compte, la différence entre les différents points de vue réside dans le statut ontologique que l'on choisit d'attribuer à la partie positive (ou négative) de la fréquence d'un champ. Les physiciens qui défendent la supériorité des ordres symétriques affirment :

Il serait difficile d'élaborer une image physique à partir d'une expression impliquant uniquement la partie positive de la fréquence du champ qui n'est pas observable ([8], p. 9).

D'autres, notamment Peter Milonni, formulent la question en termes de distinction entre les caractéristiques quantiques et classiques, et s'abstiennent d'accorder une plus grande importance à ces dernières :

Divers [...] ordonnancements donnent des poids différents aux champs de vide et aux champs sources lorsque nous essayons d'interpréter les résultats d'un calcul. Pour mettre autant que possible en évidence les aspects classiques des champs de vide et des champs sources, nous pouvons choisir un ordonnancement symétrique à chaque étape d'un calcul ([18], p.7).

Dans cette optique, on en vient alors à décrire les phénomènes qui nous intéressent comme étant dus, par exemple, uniquement à la partie de fréquence positive du champ lorsque l'ordonnancement normal est choisi, et cette image physique doit être prise au sérieux au même titre que celle décrite par l'ordonnancement normal.⁵⁵

On peut également opposer la préoccupation de Dalibard et al. de limiter la dérivation aux observables au sens formel du terme, aux motivations opérationnalistes de Jaynes pour supprimer le champ du vide, c'est-à-dire le caractère au mieux indirect de son observabilité expérimentale. Jaynes déclare :

Il me semble que si vous dites que le rayonnement est « réel », vous devez entendre par là qu'il peut être détecté par un détecteur réel. Or, un pyromètre optique ne voit que le terme de Planck, et non le terme du point zéro, dans le rayonnement du corps noir.⁵⁶

⁵⁵ L'ordre normal correspond à :

$$E^{(-)} P + P E^{(+)} , \quad (34)$$

et $E^{(-)}$ agissant à gauche est équivalent à $E^{(+)}$ à droite.

Il convient de noter que malgré ces divergences d'opinion, Peter Milonni pense en réalité que le champ du vide existe tout autant que le champ source — simplement, il ne le pense pas sur la base des arguments qui favorisent l'ordre symétrique.

⁵⁶ [17], pp. 5-6. Pour l'intérêt du débat, Jaynes poursuit : « Bien sûr, un défenseur acharné de la théorie actuelle répondra immédiatement que ces objections ne reflètent que des préconceptions métaphysiques naïves de la réalité, semblables aux notions pré-relativistes de simultanéité absolue, que l'interprétation de Copenhague de la théorie quantique a reconnues et, à juste titre, éliminées de la science. ... C'est une ontologie souple qui suppose que les fluctuations du vide sont juste assez réelles pour déplacer

On a l'impression qu'en fin de compte, pour Dalibard et al. et Sciamma, c'est la description intuitive de la relation entre le système et le champ de vide qui les convainc de la justesse de l'ordre symétrique. C'est sans doute dans ce sens que l'on peut dire que la FDT fournit un argument solide en faveur de l'existence des fluctuations du vide, car elle offre un cadre classique dans lequel Dalibard et al. peuvent facilement interpréter leurs résultats.

5.2 S'agit-il vraiment d'un champ d' s du vide ?

Un point essentiel de l'argumentation de Jaynes était qu'il était facile de se tromper dans l'identification du champ de vide. Ce point est-il pertinent pour les arguments avancés par Dalibard et al. et Milonni ?

Comme indiqué ci-dessus, Jaynes lui-même a utilisé l'expression acceptée pour désigner l'énergie spectrale énergie du vide, « $\rho(\omega)$ », c'est-à-dire $\frac{\hbar\omega}{2\pi c^3}$. Ses travaux soulignaient que la simple

L'apparition de cette expression dans un résultat ne constitue pas une preuve de l'existence du champ de vide. Il semble juste de dire qu'il était d'accord pour dire que cette expression représentait la densité d'énergie que le champ de vide aurait effectivement par mode, si ces modes existaient. Il ne voyait simplement aucune raison dans l'apparition de cette expression pour affirmer qu'ils existent réellement. Dans son dérivation, l'élément de formalisme pertinent pour leur existence ou leur absence est la gamme de fréquences sur laquelle on intègre $\rho(\omega)$. Le champ du vide implique vraisemblablement une gamme infinie de fréquences, tandis que le rayonnement émis ne concerne que les fréquences émises par l'atome. Pour rendre compte des faits observables, Jaynes a fait valoir que la gamme de fréquences appropriée est cette dernière. Cela ne prouvait pas que le champ du vide ne pouvait pas être responsable, mais cela montrait que son existence n'était pas nécessaire.

L'argument de Milonni concernant les relations de commutation repose bien sûr sur

$\rho(\omega) = \frac{\hbar\omega^3}{2\pi c^3}$ comme densité spectrale du champ du vide. Il en va de même pour Jaynes invalide-t-il l'argument de Milonni ?

L'expression utilisée par Milonni implique une intégrale sur toutes les fréquences, et non simplement sur certaines d'entre elles, et son résultat repose sur la réalisation de cette intégrale.⁵⁷ Son résultat dépend donc bien de la gamme de fréquences qu'il considère. Cela suggère que l'argument de Jaynes (à savoir que le champ en question n'est peut-être pas le champ du vide) n'affecte pas le résultat de Milonni.

Cela ne semble pas non plus affecter les arguments de Dalibard et al., bien que les choses sont moins simples dans ce cas : la densité d'énergie spectrale $\frac{\hbar\omega^3}{2\pi c^3}$ ne jouent le rôle d'une définition opérationnelle du vide dans leurs travaux. Ils définissent plutôt le champ du vide par la solution homogène des équations du champ, c'est-à-dire le « champ libre ». Or, en électrodynamique classique, de telles solutions existent et ne correspondent en rien à ce que nous entendons par champ de vide : dans ce contexte, tous les champs ont été émis à un moment donné dans le passé par une source quelconque. Ainsi, si l'on considère

le niveau 2s de l'hydrogène de 4 microvolts ; mais pas suffisamment réel pour être visible à l'œil nu, bien que dans la bande optique, cela corresponde à un flux de plus de 100 kilowatts/cm². Néanmoins, l'œil adapté à l'obscurité, en regardant par exemple une étoile faible, peut voir un rayonnement réel de l'ordre de 10⁻¹⁵ watts/cm². »

⁵⁷ Et pas seulement en termes qui s'annulent dans l'intégrande, par exemple.

La solution homogène seule ne garantit évidemment pas que nous ayons affaire à un champ de vide. Ce qui la rend telle est l'ajout de la condition « aucun photon [réel] n'est présent initialement » ([8], p. 1618), qu'ils imposent par l'exigence que « le champ de rayonnement soit à l'état de vide au temps initial » ([8], p. 1622). Il semblerait donc que les préoccupations de Jaynes concernant l'identification du champ du vide ne s'appliquent pas non plus aux travaux de Dalibard et al.

5.3 Implications pour les relations d'

5.3.1 Préliminaire : interprétations des relations d'incertitude dans le contexte de la NRQM

Les relations d'incertitude ont été interprétées de différentes manières. Comme le notent Jan Hilgevoord et Jos Uffink ([21], p. 13) :

L'interprétation de ces relations a souvent fait l'objet de débats. Les relations de Heisenberg expriment-elles des restrictions sur les expériences que nous pouvons réaliser sur les systèmes quantiques, et donc des restrictions sur les informations que nous pouvons recueillir sur ces systèmes ? Ou bien expriment-elles des restrictions sur la signification des concepts que nous utilisons pour décrire les systèmes quantiques ? Ou bien s'agit-il de restrictions de nature ontologique, c'est-à-dire affirment-elles qu'un système quantique ne possède tout simplement pas de valeur définie pour sa position et sa quantité de mouvement en même temps ? La différence entre ces interprétations se reflète en partie dans les différents noms donnés à ces relations, par exemple « relations d'imprécision », « incertitude », « indétermination » ou « relations de flou ».

L'interprétation de l'UR que l'on a tendance à adopter est généralement liée à l'interprétation que l'on privilégie de la mécanique quantique dans son ensemble. C'est pourquoi Jan Hilgevoord et Jos Uffink considèrent comme « interprétation minimale » celle qui « semble être partagée à la fois par les partisans de l'interprétation de Copenhague et par les partisans d'autres points de vue ».⁵⁸

L'interprétation minimale est statistique, les incertitudes étant identifiées par des écarts de distributions de probabilité : on considère un très grand nombre de copies du système, toutes formellement décrites par le même vecteur d'état,⁵⁹ et on examine les résultats, par exemple, des mesures de position (ou de quantité de mouvement) sur toutes ces copies. Les UR sont alors interprétées comme signifiant que « les distributions de position et de quantité de mouvement ne peuvent pas être toutes deux arbitrairement étroites » ([21], pp. 34-35).

Ce qui permet à cette description de servir d'interprétation commune, c'est le fait qu'elle est purement épistémique, en ce sens qu'elle ne fait que restreindre ce que nous pouvons savoir des systèmes, tout en restant agnostique quant à la source de cette restriction. En plus de former un tel terrain d'entente, elle est

⁵⁸ [21], p. 34. Voir également notamment : [22], pp. 82-83, [23].

⁵⁹ Ce qui, d'un point de vue expérimental, implique qu'ils doivent tous être préparés dans le même état.

également minimal dans un second sens : il ne fait « guère plus que remplir la signification empirique » d'une inégalité largement considérée comme l'expression formelle de l'UR :

$$\Delta_{\Psi} p \Delta_{\Psi} q \geq \frac{\hbar}{2} \quad (35)$$

où $\Delta_{\Psi} p$ et $\Delta_{\Psi} q$ sont les écarts types de la quantité de mouvement et de la position, respectivement :

$$(\Delta_{\Psi} p)^2 = \langle \Psi | p^2 | \Psi \rangle - \langle \Psi | p | \Psi \rangle^2, \quad (36)$$

et de même pour la position.⁶⁰

Il s'agit d'un cas particulier d'un théorème plus général sur les opérateurs hermitien (A et B ici) :

$$\Delta_{\Psi} A \Delta_{\Psi} B \geq \frac{1}{2} |\langle \Psi | [A, B] | \Psi \rangle|. \quad (37)$$

Les manuels modernes identifient généralement l'UR à cette inégalité, qui est le résultat d'une dérivation formelle basée sur le formalisme quantique lui-même, contrairement à la formulation originale de Heisenberg à laquelle nous reviendrons bientôt. Comme le montre l'équation (37), l'incertitude impliquée est due aux relations de commutation : lorsque A et B commutent, le membre de droite disparaît et les écarts types peuvent être simultanément arbitrairement petits ; l'incertitude n'apparaît que lorsque ce n'est pas le cas. La manière dont Heisenberg a initialement abordé l'UR n'avait toutefois pas grand-chose à voir avec la propagation des distributions statistiques⁽⁶²⁾. Elle était plutôt liée à l'imprécision des instruments. Heisenberg et Bohr voyaient tous deux dans l'UR plus qu'une simple restriction de notre capacité à obtenir des informations sur les systèmes quantiques (bien sûr, selon eux, cela impliquait également cela, mais pas uniquement). Cela résultait d'un point de vue opérationnaliste : Heisenberg soutenait que les grandeurs physiques (position, impulsion) n'ont de sens que si nous pouvons spécifier une expérience permettant de les mesurer (ce que Hilgevoord et Uffink appellent le « principe mesure = signification »). Les restrictions de notre capacité à mesurer avec précision (ce qui est *en soi* une question épistémique) impliquent alors des imprécisions correspondantes dans la *signification* des grandeurs physiques :

Si, comme l'affirme Heisenberg, il n'existe aucune expérience permettant de mesurer simultanément et avec précision deux grandeurs conjuguées, alors ces grandeurs ne sont pas non plus bien définies simultanément.⁶³

⁶⁰ Ceci a été dérivé par Kennard en 1927, et Heisenberg l'a présenté dans ses propres travaux en 1930. [21], pp. 19-20, [24], [25].

⁶¹ Cette généralisation a été obtenue par Robertson en 1929. [21], p. 20.

⁶² Même s'il considérait le résultat de Kennard comme une preuve exacte de son UR, [21], p. 20.

⁶³ [21], pp. 6-8. Selon Heisenberg, les inexactitudes étaient liées à l'apparition de changements discontinus, c'est-à-dire dans l'expérience de pensée où il imagine essayer de mesurer la position d'un électron à l'aide de la lumière :

À l'instant où la position est déterminée, c'est-à-dire à l'instant où le photon est diffusé par l'électron, celui-ci subit un changement discontinu de sa quantité de mouvement. Ce changement est d'autant plus important que la longueur d'onde est petite.

Mais Heisenberg a parfois exprimé des opinions que l'on pourrait qualifier d'extrêmes, si nier *l'existence* d'une chose peut être considéré comme plus extrême que nier la possibilité de *définir* cette chose. Il a notamment écrit :

Je crois que l'on peut formuler de manière concise l'émergence de la trajectoire classique d'une particule comme suit : la trajectoire n'existe que parce que nous l'observons ([26], p. 185),

ce à quoi Hilgevoord et Uffink ont répondu :

Apparemment, selon lui, une mesure ne sert pas seulement à donner un sens à une quantité, elle crée une valeur particulière pour cette quantité. On peut appeler cela le principe de mesure = création. Il s'agit d'un principe ontologique, car il énonce ce qui est physiquement réel ([21], p. 12).

Cela suggère à son tour une interprétation de l'UR qui peut également être considérée comme ontologique :

Avant la mesure finale, le mieux que nous puissions attribuer à l'électron est une quantité d'impulsion imprécise ou floue. Ces termes sont ici utilisés dans un sens ontologique, caractérisant un attribut réel de l'électron.⁶⁴

La conception de Bohr concernant l'UR différerait encore davantage de celle de Heisenberg. Notamment, Bohr n'attribuait pas l'origine de l'UR à des changements discontinus de la quantité de mouvement, mais à son principe de complémentarité :

de la lumière utilisée, c'est-à-dire plus la détermination de la position est exacte. À l'instant où la position de l'électron est connue, son impulsion ne peut donc être connue que jusqu'à des grandeurs qui correspondent à ce changement discontinu ; ainsi, plus la position est déterminée avec précision, moins l'impulsion est connue avec précision, et inversement. [26], pp. 174-5, traduction anglaise [27], pp. 62-84. Voir également [21], p. 7.

⁶⁴ [21], p. 12. Hilgevoord et Uffink affirment cela après avoir résumé les opinions de Heisenberg

: Nous mesurons d'abord très précisément la quantité de mouvement de l'électron. En mesurant...

ment = signification, cela implique que le terme « momentum de la particule » est désormais bien défini. De plus, selon le principe de mesure = création, nous pouvons dire que ce momentum est physiquement réel. Ensuite, la position est mesurée avec une imprécision δq . À cet instant, la position de la particule devient bien définie et, là encore, on peut considérer qu'il s'agit d'un attribut physiquement réel de la particule. Cependant, la quantité de mouvement a maintenant changé d'un montant imprévisible d'un ordre de grandeur $|p_{(j)} - p'_{(j)}| \sim \frac{\hbar}{\delta q}$. La signification et la validité de cette affirmation peuvent être vérifiées par une mesure ultérieure de la quantité de mouvement.

La question est alors de savoir quel statut attribuer à la quantité de mouvement de l'électron juste avant sa mesure finale. Est-elle réelle ? Selon Heisenberg, non. Avant la mesure finale, le mieux que nous puissions attribuer à l'électron est une quantité de mouvement floue ou imprécise. Ces termes sont ici utilisés dans un sens ontologique, caractérisant un attribut réel de l'électron.

Ce n'est pas tant la perturbation inconnue qui rend la quantité de mouvement de l'électron incertaine, mais plutôt le fait que la position et la quantité de mouvement de l'électron ne peuvent être définies simultanément dans cette expérience.⁶⁵

La dérivation de l'UR proposée par Bohr en 1928, dans laquelle les incertitudes étaient définies en termes de caractéristiques d'un paquet d'ondes, est particulièrement intéressante⁽⁶⁶⁾. Δx et Δt représentaient les extensions spatiales et temporelles du paquet d'ondes, $\Delta \sigma$ et $\Delta \nu$ son intervalle de nombres d'ondes inverses ($\sigma = \frac{1}{\lambda}$) et ses fréquences.

Ce qui est remarquable dans cette dérivation, c'est qu'elle n'utilise pas les relations de commutation, contrairement à celle qui est désormais standard et qui a été discutée ci-dessus. Elle s'appuie plutôt sur l'analyse de Fourier pour obtenir le produit des incertitudes, responsable de la restriction dans la définition *simultanée* des quantités (car lorsqu'un facteur augmente, l'autre doit diminuer) ([21], p. 29) :

$$\Delta x \Delta \sigma \approx \Delta t \Delta \nu \approx 1. \quad (38)$$

Bohr a ensuite obtenu l'UR en substituant la fréquence et le nombre d'ondes dans ces relations, en utilisant respectivement la relation entre l'énergie et la fréquence $E = h\nu$, et entre la quantité de mouvement et la longueur d'onde $p = \frac{h}{\lambda}$ ([21], p.29) :

$$\Delta x \Delta p \approx \Delta t \Delta E \approx h. \quad (39)$$

Deux remarques s'imposent concernant l'utilisation de $E = h\nu$ et $p = \frac{h}{\lambda}$. Premièrement, c'est à travers elles que le caractère quantique des relations apparaît : $\Delta x \Delta \sigma \approx \Delta t \Delta \nu \approx 1$ sont obtenues à partir de considérations purement classiques. Cela ressort également du fait que c'est la substitution de $E = h\nu$ et $p = \frac{h}{\lambda}$ qui introduit h dans le résultat final, c'est-à-dire dans l'UR. Deuxièmement, l'équation (38) peut être considérée comme exprimant le principe de complémentarité si cher à Bohr, E et p étant considérés comme ayant trait à l'aspect particulaire des entités et ν et σ à leur aspect ondulatoire.

5.3.2 Relations d'incertitude, fluctuations du vide et particules virtuelles

Comme brièvement évoqué dans l'introduction, dans le NRQM, une particule dans un puits d'énergie « oscillateur harmonique » a un état fondamental, une énergie « zéro point » de $\frac{1}{2} \hbar \omega$, ce qui est généralement considéré comme une conséquence des relations d'urgence pour la position et la quantité de mouvement, qui impliquent des relations d'urgence correspondantes pour les énergies potentielle et cinétique respectivement.

Une dérivation reliant l'UR à l'énergie « zéro point » se présente comme suit. Si l'on considère l'énergie mécanique de la particule et :

- on remplace p par « l'incertitude de la quantité de mouvement » Δp dans le terme d'énergie cinétique et x par « l'incertitude de la position » Δx dans le terme d'énergie potentielle spécifique

⁶⁵ « Addition in Proof » à [26] ; [21], p. 24. L'expérience de pensée en question consiste à déterminer la position d'un électron en le faisant dévier par un photon, ce qui entraîne un changement discontinu de la quantité de mouvement de l'électron.

⁶⁶ Un paquet d'ondes est une superposition d'ondes de fréquences différentes, donc de nombres d'ondes différents.

à un oscillateur harmonique,⁶⁷

- on utilise la forme minimale de l'UR $\Delta x \Delta p \geq \frac{\hbar}{2}$, en remplaçant Δp par

- on minimise l'expression résultante de l'énergie mécanique par rapport à Δx ,

on constate que cette énergie a un minimum global (à $\Delta x = \frac{\hbar}{2m\omega}$), et que sa

valeur à cet endroit est $\frac{1}{2} \hbar \omega$.⁶⁸

Il convient de noter qu'aucune hypothèse n'a été formulée quant à la signification exacte de Δx et Δp , c'est-à-dire s'il faut les considérer comme représentant des imprécisions expérimentales dues aux limites des instruments, des écarts statistiques tels que les écarts types, ou s'ils sont associés à des paquets d'ondes.

Cette dérivation semble motiver une interprétation de l'UR qui va au-delà d'une interprétation épistémique, selon laquelle les relations ne limiteraient que notre connaissance de la position et de la quantité de mouvement (⁶⁹). Autrement dit, si nous interprétons le résultat de l'énergie du point zéro comme signifiant qu'une particule dans son état fondamental a *effectivement* une énergie de $\frac{1}{2} \hbar \omega$, cela impliquerait que l'UR nous dit quelque chose sur la valeur réelle d'une variable. La manière dont cela se produirait exactement semble moins claire. Une explication serait de dire que l'incertitude dans l'UR provient d'une indétermination intrinsèque des variables en jeu, c'est-à-dire qu'il n'y a en réalité aucun fait quant à la quantité exacte d'énergie potentielle (/cinétique) que possède la particule ; que par conséquent, les énergies potentielle et cinétique ne peuvent être à la fois réellement et exactement nulles, et que leur somme, l'énergie totale, ne peut donc pas l'être non plus. Dans cette histoire, ironiquement, le degré d'indétermination des énergies qui composent la somme détermine la valeur exacte de cette dernière. Cela correspondrait à une interprétation ontologique de l'UR telle que discutée par Hilgevoord et Uffink en relation avec les idées de Heisenberg, puisque les énergies potentielle et cinétique devraient être floues dans ce sens.

Si l'on considère plutôt (Δx) et (Δp) comme des écarts types, on peut se limiter à l'interprétation minimale de l'UR, qui est uniquement épistémique. Cela impliquerait que la valeur $\frac{1}{2} \hbar \omega$ pour l'énergie du point zéro devrait être interprétée de manière similaire.²

Nous venons d'examiner les énergies du point zéro dans le NRQM, nous revenons maintenant à la QFT et aux demi-quanta de son champ de vide. Comme indiqué précédemment, il s'agit de ce que l'on appelle les « fluctuations du vide », souvent décrites comme des « particules virtuelles » évanescences qui empruntent au vide suffisamment d'énergie pour devenir réelles pendant un laps de temps trop court pour être observées, grâce à l'UR entre le temps et l'énergie. (⁷⁰)

⁶⁷ En prenant l'énergie de l'état fondamental comme étant : $E_0 = \frac{1}{2} \hbar \omega = \frac{\hbar^2 \omega^2}{2m} (\Delta p)^2 + \frac{1}{2} m \omega^2 (\Delta x)^2$.

⁶⁸ Notez que cette dérivation repose sur l'utilisation de l'égalité dans l'UR, ce qui est possible ici car nous avons affaire à un potentiel d'oscillateur harmonique : la fonction d'onde de la particule dans un tel potentiel est une gaussienne, ce qui constitue un cas d'incertitude minimale.

⁶⁹ D'où les énergies potentielle et cinétique.

⁷⁰ Il convient peut-être de souligner que cette utilisation de l'expression « particules virtuelles » diffère de son usage standard et formel dans le contexte des diagrammes de Feynman, notamment : dans ce contexte, les « particules virtuelles » sont des particules « hors coquille », ce qui signifie qu'elles ne sont pas tenues de satisfaire la relation relativiste entre l'énergie et la quantité de mouvement. Les deux définitions ont toutefois un point commun : dans

Maintenant, la question de savoir si nous sommes fondés à parler d'énergies du point zéro en termes de « fluctuations », ce qui suggère un processus dynamique, et si nous avons des raisons d'accepter l'explication des « particules virtuelles » sont deux questions différentes. Elles sont toutefois liées : la seconde nécessite la première, car les « particules virtuelles » nécessitent que l'énergie varie légèrement avec le temps, augmentant et diminuant tour à tour, c'est-à-dire qu'elles nécessitent que l'énergie fluctue.

Ce qui rend la question de la justesse de cette explication fascinante, ce sont ses implications pour la signification des UR. Celles-ci jouent un double rôle dans ce schéma, car l'incertitude en matière d'énergie et l'incertitude en matière de temps jouent des rôles différents. Pour l'incertitude en matière d'énergie, les UR doivent concerner un fait concret du monde et nous renseigner sur les propriétés des entités, comme cela a déjà été dit dans le contexte des énergies du point zéro dans la NRQM. Cela exige donc que l'UR soit ontologique au sens où l'entendent Hilgevoord et Uffink. Mais on pourrait soutenir qu'elle est désormais ontologique dans un sens plus large : cette « incertitude » est responsable de l'existence même de certaines entités. En effet, le schéma décrit des particules virtuelles qui *existent* réellement parce qu'elles acquièrent suffisamment d'énergie pour devenir des quanta complets grâce à l'UR. La raison pour laquelle les implications des énergies du point zéro dans la NRQM ne sont pas aussi extrêmes est que cette dernière ne peut pas décrire la création de particules, de sorte que les quanta concernent l'énergie d'une particule que la NRQM doit considérer comme préexistante. Dans la QFT, en revanche, *l'existence* même d'une entité (une particule) dépend de la *valeur d'une propriété* d'une autre (le champ) : les quanta d'énergie appartiennent au champ, et non à une particule dès le départ ; leur présence ou leur absence dépend de la valeur d'une propriété de ce champ, c'est-à-dire de son énergie. Tant que l'énergie est un quantum complet, une particule existe réellement. Son caractère virtuel réside dans son incapacité à maintenir cette existence suffisamment longtemps pour être observée, ce qui nous amène au deuxième rôle joué par l'UR.

Alors que l'incertitude énergétique est responsable de « l'effet » lui-même (c'est-à-dire l'existence des particules), l'incertitude temporelle fournit une explication à ce qui pourrait autrement être considéré comme une contradiction avec les faits expérimentaux, à savoir notre incapacité à l'observer. Ainsi, dans cette explication, l'incertitude temporelle concerne à la fois un fait physique (les particules existent pendant un court instant) et une limite de notre connaissance (les particules virtuelles sont inobservables par principe).

Cependant, l'explication des « particules virtuelles » pose de sérieux problèmes.

En effet, le rôle que les demi-quanta d'énergie sont censés jouer dans cette description n'est pas clair. Leur relation avec les concepts de fluctuations du vide et de particules virtuelles semble certainement être médiée par l'UR : les particules virtuelles sont explicitement considérées comme le résultat direct de l'UR temps-énergie, les fluctuations sont clairement attribuées à l'incertitude (bien que la manière exacte reste encore floue), et les demi-quanta de champ sont considérés comme étant dus à l'UR par analogie avec l'énergie du point zéro de l'oscillateur harmonique NRQM.⁷¹ Cette analogie est suggérée par le fait que la particule NRQM et le champ électromagnétique QFT ont tous deux des énergies de base non nulles en raison de la non-commutation des opérateurs. Dans ce contexte, le champ et la quantité de mouvement canonique de la QFT sont les analogues de la position NRQM.

Dans les deux cas, les particules sont considérées comme inobservables par principe.

⁷¹ Particule dans un puits de potentiel d'oscillateur harmonique.

tion et la quantité de mouvement respectivement. Cependant, cette analogie pose problème. Un premier problème est que, pour que l'explication par les « particules virtuelles » ait un sens, il faut un analogue de l'UR énergie-temps du NRQM, et non de la position-quantité de mouvement. On peut se demander dans quelle mesure cela pose problème : étant relativiste, la QFT place d'un côté la position et le temps, et de l'autre l'énergie et la quantité de mouvement, sur un pied d'égalité. Par conséquent, le champ quantique et la quantité de mouvement canonique sont tous deux des fonctions du temps et de la position. Étant donné que leur relation de commutation implique autant le temps que la position, on pourrait soutenir que la relation de commutation entre le champ et la quantité de mouvement canonique est la plus pertinente. En outre, il a été avancé que « dans le contexte de la relativité restreinte, la forme énergie-temps pourrait être considérée comme une conséquence de la version position-impulsion, car x et t (ou plutôt ct) vont de pair dans le quadrivecteur position-temps » ([28], p. 112). On pourrait donc dire que la relation canonique entre le champ et la quantité de mouvement est dans la même relation avec les UR énergie-temps et position-quantité de mouvement.⁽⁷²⁾ Cependant, on ne sait pas clairement comment la théorie des « particules virtuelles » devrait être formulée pour refléter la relation canonique entre le champ et la quantité de mouvement. Un deuxième problème avec l'explication des « particules virtuelles », et même avec le simple fait de parler de fluctuations, est qu'elles font implicitement appel à un processus dynamique. Bien que la relation de commutation canonique entre le champ et la quantité de mouvement implique indirectement le temps, on peut difficilement dire qu'elle exprime une évolution dynamique. Les opérateurs peuvent dépendre du temps, mais ces relations de commutation sont généralement considérées à des instants égaux, il est donc difficile de voir comment une incertitude temporelle pourrait apparaître — sans parler de la manière dont cela pourrait ensuite donner lieu à une explication dynamique. De l'autre côté de la correspondance, comme l'a fait remarquer Robert Klauber, le fait que les états propres de l'énergie soient des demi-quanta ne cadre pas non plus avec une explication dynamique :

Ces demi-quanta semblent simplement « rester » dans le vide, et non « apparaître et disparaître » de celui-ci. Il n'existe aucun mécanisme apparent qui leur permette d'exister une partie du temps, mais pas tout le temps.⁷³

Un troisième problème, en partie lié au précédent et peut-être plus inquiétant, remettrait également en cause non seulement l'explication en termes de particules virtuelles, mais aussi la validité même du concept de fluctuation. Il existe une importante différence entre la relation de commutation NRQM et son équivalent QFT : contrairement à la position et à la quantité de mouvement dans la NRQM, les champs quantiques et leurs quantités de mouvement conjuguées ne sont pas observables, dans la mesure où leur valeur attendue est nulle. Pour cette raison, Klauber a fait valoir que l'UR que l'on peut associer à cette relation de commutation est en fait « dénuée de sens » ([29], p. 273).⁽⁷⁴⁾

En somme, il est donc difficile de voir comment les relations de commutation responsables des demi-quanta dans la QFT sont liées à l'UR temps-énergie invoquée dans le « virtuel ».

⁷² même si seule la position-impulsion est basée sur les relations de commutation.

⁷³ [29], p. 270. L'ouvrage de Klauber est avant tout un manuel sur la QFT, mais certaines parties, notamment le chapitre 10 intitulé « The vacuum revisited » (Le vide revisité), constituent une réflexion critique sur l'utilisation de l'expression « particules virtuelles ».

⁷⁴ Les mêmes remarques s'appliquent à la relation de commutation entre les opérateurs de création et d'annihilation.

particules » — et peut-être même nécessaire pour parler de fluctuations.

Peut-on néanmoins utiliser l'UR temps-énergie pour justifier cette explication, indépendamment de toute référence aux relations de commutation canonique champ-moment ? Après tout, les points de vue de Heisenberg et de Bohr sur l'UR *n'exigeaient* pas de relations de commutation ⁽⁷⁵⁾. Comme nous l'avons vu, Bohr a même fourni une dérivation de l'UR position-impulsion ainsi que de l'UR temps-énergie sans utiliser de relations de commutation. Et surtout, l'UR temps-énergie n'est de toute façon pas associée à une relation de commutation.

Cette approche peut sembler plus prometteuse, mais même dans ce cas, on ne voit pas clairement comment une explication dynamique pourrait émerger. On peut s'inspirer du fait que l'UR temps-énergie a été utilisée pour expliquer la période d'oscillation d'un système entre ses états stationnaires : la différence d'énergie entre deux états propres est liée à la période d'oscillation par cette UR. Cependant, cela nécessite que le système soit dans une superposition des deux états propres d'énergie, ce qui n'est pas le cas d'un champ à l'état de vide par définition. Une autre explication suggestive est la relation entre la durée de vie d'un type de particule et la dispersion observée dans sa masse (c'est-à-dire son énergie au repos) lors de mesures répétées, qui est également donnée par l'UR temps-énergie. Cependant, dans cette explication, ce que contrôle l'UR est le temps nécessaire aux particules pour cesser d'exister ; en revanche, ce dont nous avons besoin en priorité pour l'explication des « particules virtuelles », c'est qu'elle leur fournisse l'énergie nécessaire pour devenir des quanta complets. Il s'agit néanmoins d'une explication suggestive.

Une approche différente consisterait à rechercher une dynamique sans invoquer l'UR à cette fin. Cela ne fournirait pas beaucoup d'éléments pour étayer la théorie des particules virtuelles, mais justifierait probablement l'utilisation du terme « fluctuations ». À cet égard, la FDT pourrait être utile, car elle a été interprétée comme montrant que dans la limite $T = 0$, la force fluctuante est due au champ du vide ⁽⁷⁶⁾. Cela semblerait certainement exiger que ce dernier fluctue, au même sens dynamique que le champ de rayonnement de Planck dans la limite des hautes températures. D'une manière générale, dans les applications du FDT, les effets de point zéro et la limite de température élevée, les effets classiques ont la même relation avec la force fluctuante, chacun y contribuant par un terme. Il semble donc raisonnable que si les effets thermiques donnent lieu à une force fluctuante et dynamique, il en va de même pour les effets de point zéro.

Dans la mesure où l'on peut dire que le champ du point zéro associé provient de l'UR (ce qui est en soi discutable), cela semblerait en effet justifier d'interpréter ce dernier dans un sens plus qu'épistémique, mais aussi ontologique, c'est-à-dire comme affectant réellement la valeur des variables. Cette conclusion semble valable, que l'on considère ou non que ces fluctuations nécessitent la préexistence du champ de rayonnement, puisque l'UR implique un terme de point zéro dans les deux cas.

⁷⁵ La relation entre l'UR et les relations de commutation a été introduite par Kennard, comme discuté ci-dessus.

⁷⁶ C'est-à-dire que le terme de point zéro dans l'équation (15) pour $E(\omega, T)$ se manifeste sous forme de fluctuations.

6 Conclusions

Les effets traditionnellement attribués aux fluctuations du vide peuvent également être attribués à un champ source, et nous avons montré comment cette sous-détermination se produit : l'opérateur relatif au système commute avec l'opérateur pour le champ total (source d'+ s du vide), les dérivations donnent donc les mêmes résultats quel que soit l'ordre choisi, mais il ne commute pas avec l'opérateur pour le champ source ni pour le champ de vide, de sorte que des ordres différents donnent des contributions relatives différentes de ces deux champs. Nous avons vu que l'ambiguïté peut être levée en exigeant que seuls des opérateurs hermitiques soient utilisés tout au long de la dérivation, ce qui correspond au choix d'un ordre symétrique et implique que les deux champs contribuent — et donc que le champ du vide a un effet observable. Cette exigence se justifie par le fait que l'ordre symétrique est plus physique, dans la mesure où il implique des opérateurs hermitiques tout au long des calculs. Cependant, cela implique de considérer la partie positive (ou négative) de la fréquence du champ comme non physique.

Dans le même temps, les physiciens ont souvent fait appel au théorème de fluctuation-dissipation lorsqu'ils ont discuté des contributions relatives des champs sources et des champs de vide. Bien que le FDT donne une image physique intuitive claire dans la limite classique à haute température, son interprétation dans la limite $T = 0$ des contributions du point zéro n'est pas aussi simple. Si nous considérons que la force fluctuante qu'il décrit est exercée par la même entité que la force dissipative, nous devrions alors conclure qu'elle est due à un champ source. Cependant, le plus souvent, le FDT a été interprété comme suggérant que la force fluctuante est due au champ du vide. Dans tous les cas, le FDT nécessite la présence d'un champ source pour que des effets dissipatifs puissent se produire, et ce que nous entendons par champ du vide dans une telle situation nécessite une certaine prudence : en QFT, cette expression fait référence au champ électromagnétique dans son état fondamental. La présence même d'un champ source exclut cette possibilité. Ce qui est donc en jeu, ce sont les fluctuations du champ (excité) ayant la même origine que l'énergie du point zéro du champ du vide, c'est-à-dire les relations d'incertitude.

La plupart des physiciens ont interprété la FDT comme une preuve de la contribution des fluctuations du vide. Mais ce point de vue a été vivement critiqué à un moment donné, notamment par Edwin Jaynes, qui a interprété la FDT comme « montrant que les effets du champ source sont les mêmes que s'il y avait des fluctuations du vide », c'est-à-dire que la force fluctuante est exercée par le même champ que la force dissipative. Il a fait valoir que les gens avaient supposé à tort que les fluctuations du vide étaient responsables parce que l'expression de la densité spectrale d'énergie du champ du vide apparaissait dans leurs résultats, alors qu'en réalité, cette expression pouvait tout aussi bien être due au champ source.

Cependant, l'avertissement de Jaynes contre une interprétation abusive de certaines expressions comme représentant des effets de champ de vide ne s'applique pas à la recherche impliquant l'ordonnancement des opérateurs. Il ne semble pas non plus remettre en cause l'argument en faveur des effets de champ de vide concernant la cohérence de la théorie quantique du rayonnement : la prise en compte du champ de vide dans l'équation du mouvement de Heisenberg pour une charge en accélération est nécessaire pour que la relation de commutation position-impulsion

Par conséquent, à ce stade, il semble que le champ du vide contribue effectivement aux effets physiques en question.

La question de savoir s'il faut décrire le rôle du champ du vide comme impliquant des fluctuations au sens dynamique, ou parler de ces « fluctuations du vide » comme des particules évanescences, virtuelles, qui existent grâce à l'UR, est beaucoup plus discutable. La meilleure raison que nous ayons pour parler de fluctuations semble être la FDT, qui fait référence à une force fluctuante. À haute température, les contributions classiques à cette force sont certainement de nature dynamique, ce qui suggère que leurs contreparties au point zéro le sont également.

En outre, dans la mesure où l'on peut dire que ce champ de point zéro provient de l'UR, cela suggère qu'au moins dans ce contexte, ce dernier doit être interprété dans un sens plus que minimal et épistémique, mais aussi dans un sens ontologique.

7 Annexe : Sous-détermination et ordonnancement des opérateurs d'

Premier cas : Q donné par PE

Nous désignons par Q l'opérateur pour la quantité qui nous intéresse, par E l'opérateur de champ et par P l'opérateur pour la quantité qui décrit le système atomique. Si :

$$Q = PE, \quad (40)$$

et si nous nous abstenons de toute hypothèse concernant l'ordre de E et P dans Q , la forme la plus générale que ce dernier peut prendre peut s'écrire :

$$Q = \lambda P E + (1 - \lambda) EP, \quad (41)$$

où λ est arbitraire. Cela implique que les contributions à Q dues au vide et aux champs sources, c'est-à-dire respectivement Q_0 et Q_S peuvent à leur tour être écrites :

$$Q_0 = \lambda P E_0 + (1 - \lambda) E_0 P; \quad Q_S = \lambda P E_S + (1 - \lambda) E_S P. \quad (42)$$

P , E_S et E_0 sont hermitiennes, donc leurs conjugués sont :

$$P = P^*; \quad E_0 = E_0^*; \quad E_S = E_S^*, \quad (43)$$

et donc les conjugués hermitien de Q_0 et Q_S sont :

$$Q_0^* = \lambda E_0^* P^* + (1 - \lambda) P^* E_0^* = \lambda E_0 P + (1 - \lambda) P E_0 \quad (44)$$

$$Q_S^* = \lambda E_S^* P^* + (1 - \lambda) P^* E_S^* = \lambda E_S P + (1 - \lambda) P E_S. \quad (45)$$

Pour que Q_0 et Q_S soient hermitiennes, les conditions suivantes doivent être satisfaites :

$$Q_0 = Q_0^*; \quad Q_S = Q_S^* \quad (46)$$

$$\lambda P E_0 + (1 - \lambda) E_0 P = \lambda E_0 P + (1 - \lambda) P E_0; \quad \lambda P E_S + (1 - \lambda) E_S P = \lambda E_S P + (1 - \lambda) P E_S \quad (47)$$

$$\lambda = 1 - \lambda \Rightarrow \lambda = \frac{1}{2} \quad (48)$$

C'est-à-dire que Q_0 et Q_S hermitien implique que $\lambda = \frac{1}{2}$. Dans ce cas, les deux opérateurs prennent la forme :

$$Q_0 = \frac{1}{2} P E_0 + \frac{1}{2} E_0 P; \quad Q_S = \frac{1}{2} P E_S + \frac{1}{2} E_S P \quad (49)$$

Ceci correspond à l'expression suivante pour Q :

$$\begin{aligned} Q &= Q_0 + Q_S \\ &= \frac{1}{2} (P E_0 + E_0 P) + \frac{1}{2} (P E_S + E_S P) \\ &= \frac{1}{2} P E + \frac{1}{2} E P. \end{aligned} \quad (50)$$

Si nous développons maintenant les opérateurs de champ en fonction de leurs composantes de fréquence positive et négative :

$$Q = \frac{1}{2} P E + \frac{1}{2} E P = \frac{1}{2} P (E^+ + E^-) + \frac{1}{2} (E^+ + E^-) P, \quad (51)$$

ce qui correspond à un choix d'ordre symétrique des opérateurs de champ, comme annoncé.⁷⁷

Deuxième cas (atome à deux niveaux) : Q donné par des termes de la forme $E^- P^- + E^+ P^+$

Maintenant, Q n'est plus représenté par PE , mais par une somme de termes de la forme :

$$E^- P^- + E^+ P^+, \quad (52)$$

où E^- et P^- peuvent être ordonnés arbitrairement puisqu'ils commutent, et il en va de même pour E^+ et P^+ . Si nous choisissons ensuite de les ordonner différemment dans les deux termes, c'est-à-dire :

$$E^- P^- + P^+ E^+ \quad \text{ou} \quad P^- E^- + E^+ P^+ \quad (53)$$

nous remarquons que nous obtenons ces deux expressions dans le même ordre lorsque nous prenons leurs conjugués hermitien :

$$(E^- P^- + P^+ E^+)^\dagger = P^+ E^+ + E^- P^- = E^- P^- + P^+ E^+ \quad (54)$$

$$(P^- E^- + E^+ P^+)^\dagger = E^+ P^+ + P^- E^- = P^- E^- + E^+ P^+ \quad (55)$$

Si nous voulons maintenant exprimer cela en termes de contributions distinctes Q_0 et Q_s des champs de vide et de source, la même structure est conservée lorsque nous substituons $E^+ = E_0^+ + E_s^+$ et $E^- = E_0^- + E_s^-$. Par exemple :

$$Q_0^\dagger = (E_0^- P_0^- + P_0^+ E_0^+)^\dagger = P_0^+ E_0^+ + E_0^- P_0^- = E_0^- P_0^- + P_0^+ E_0^+ = Q_0. \quad (56)$$

Nous voyons donc qu'avec les deux choix d'ordre dans l'équation (53), Q_0 et Q_s sont hermitiennes. Cependant, ces deux choix ne sont pas équivalents : ils constituent des choix d'ordre différents qui conduisent à des contributions relatives différentes des champs de vide et de source.

Dans une telle situation, c'est-à-dire lorsque la variable d'intérêt Q est donnée par une somme de termes de la forme Eq.(52), supprimer la liberté d'ordre nécessite d'imposer une contrainte supplémentaire, en plus d'exiger que Q_0 et Q_s soient hermitiennes. Il faut également imposer que l'expression de Q soit réécrite sous une forme qui implique uniquement des produits de E avec P ,⁷⁸ et non des produits de E^+ ou E^- avec P^+ ou P^- comme ci-dessus. L'ordre qui correspond à cette exigence est à nouveau l'ordre symétrique.⁷⁹

⁷⁷ Rappelons que l'ordre normal consiste à placer les opérateurs d'annihilation à droite des opérateurs de création, l'ordre anti-normal est l'inverse, et l'ordre symétrique est la combinaison linéaire des deux en proportions égales.

⁷⁸ Ou des produits de combinaisons appropriées de E^+ , E^- et P^+ , P^- .

⁷⁹ Ces dérivations sont des versions clarifiées de ce qui peut être trouvé dans [8].

Références

- [1] R. Jaffe, « Casimir effect and the quantum vacuum », *Physical Review D* 72, 021301 (2005).
- [2] J. Bain, « Against particle/field duality: Asymptotic particle states and interpolating fields in interacting qft (or: Who's afraid of Haag's theorem?) », *Erkenntnis* 53, 375 (2000).
- [3] J. Earman, A. Arageorgis et L. Ruetsche, « Fulling non-uniqueness and the unruh effect: A primer on some aspects of quantum field theory », *Philosophy of Science* 70, 164 (2003).
- [4] D. Fraser, « The fate of particles in quantum field theories with interactions » (Le destin des particules dans les théories quantiques des champs avec interactions), *Studies in History and Philosophy of Science Part B: Studies in History and Philosophy of Modern Physics* 39, 841 (2008).
- [5] J. Bain, *Relativity and Quantum Field Theory* (Springer, 2010).
- [6] D. B. Malament, dans *Perspectives on quantum reality* (Springer, 1996), pp. 1–10.
- [7] H. Halvorson et R. Clifton, « No place for particles in relativistic quantum theories? », (2002).
- [8] J. Dalibard, J. Dupont-Roc et C. Cohen-Tannoudji, « Vacuum fluctuations and radiation reaction: identification of their respective contributions », *Journal de Physique* 43, 1617 (1982).
- [9] L. Allen, *Optical Resonance and Two-level Atoms*, vol. 28 (Courier Corporation, 1975).
- [10] J. Martin, « Tout ce que vous avez toujours voulu savoir sur le problème de la constante cosmologique (mais n'avez jamais osé demander) », *Comptes Rendus Physique* (2012).
- [11] J. Martin, « Tout ce que vous avez toujours voulu savoir sur le problème de la constante cosmologique (mais n'avez jamais osé demander) », prépublication arXiv arXiv:1205.3365 (2012).
- [12] P. Milonni, *The quantum vacuum: an introduction to quantum electrodynamics* (Academic Press, 1994), ISBN 9780124980808, URL <http://books.google.com/books?id=P83vAAAAAAAJ>.
- [13] R. P. Feynman, R. B. Leighton et M. Sands, *Les cours de physique de Feynman, vol. 2 : Principalement l'électromagnétisme et la matière* (Addison-Wesley, 1979).
- [14] D. W. Sciama, dans *The philosophy of vacuum*, édité par S. Saunders et H. R. Brown (Clarendon Press, Oxford, 1991).

- [15] H. B. Callen et T. A. Welton, « Irréversibilité et bruit généralisé », *Physical Review* 83, 34 (1951).
- [16] H. Nyquist, « Thermal agitation of electric charge in conductors », *Physical review* 32, 110 (1928).
- [17] E. Jaynes, dans *Coherence and Quantum Optics IV*, édité par L. Mandel et E. Wolf (Plenum Press, New York, 1978), URL <http://bayes.wustl.edu/etj/articles/electrodynamics.today.pdf>.
- [18] P. Milonni, « Différentes façons d'envisager le vide électromagnétique », *Physica Scripta* 1988, 102 (1988).
- [19] P. Milonni, « Les forces de Casimir sans champ de rayonnement du vide », *Physical Review A* 25, 1315 (1982).
- [20] P. Milonni, « Réaction au rayonnement et théorie non relativiste de l'électron », *Physics Letters A* 82, 225 (1981).
- [21] J. Hilgevoord et J. Uffink, « Le principe d'incertitude - Encyclopédie de philosophie de Stanford », (2006).
- [22] H. R. Pagels, *Le code cosmique : la physique quantique comme langage de la nature* (Courier Corporation, 1982).
- [23] C. Callender, « Rien à voir ici : Démystifier le principe d'incertitude », http://opinionator.blogs.nytimes.com/2013/07/21/nothing-to-see-here-demoting-the-uncertainty-principle/?_r=0 (21 juillet 2013).
- [24] W. Heisenberg, W. Heisenberg, W. Heisenberg et W. Heisenberg, *Die physikalischen prinzipien der quantentheorie*, vol. 1 (S. Hirzel Leipzig, 1930).
- [25] W. Heisenberg, « The physical principles of quantum mechanics », U. Chicago Press, Chicago p. 21 (1930).
- [26] W. Heisenberg, « Über den anschaulichen Inhalt der quantentheoretischen Kinematik und Mechanik », *Zeitschrift für Physik* 43, 172 (1927).
- [27] J. A. Wheeler et W. H. Zurek, *Quantum theory and measurement* (Princeton University Press, 1983).
- [28] D. J. Griffiths et E. G. Harris, *Introduction to quantum mechanics*, vol. 2 (Prentice Hall New Jersey, 1995).
- [29] R. D. Klauber, *Student friendly quantum field theory* (Sandtrove Press, 2013).