

Sur l'énergie propre de l'électron

Autor(en): **Mercier, André**

Objektyp: **Article**

Zeitschrift: **Helvetica Physica Acta**

Band (Jahr): **12 (1939)**

Heft I

PDF erstellt am: **26.04.2024**

Persistenter Link: <https://doi.org/10.5169/seals-110933>

Nutzungsbedingungen

Die ETH-Bibliothek ist Anbieterin der digitalisierten Zeitschriften. Sie besitzt keine Urheberrechte an den Inhalten der Zeitschriften. Die Rechte liegen in der Regel bei den Herausgebern.

Die auf der Plattform e-periodica veröffentlichten Dokumente stehen für nicht-kommerzielle Zwecke in Lehre und Forschung sowie für die private Nutzung frei zur Verfügung. Einzelne Dateien oder Ausdrucke aus diesem Angebot können zusammen mit diesen Nutzungsbedingungen und den korrekten Herkunftsbezeichnungen weitergegeben werden.

Das Veröffentlichen von Bildern in Print- und Online-Publikationen ist nur mit vorheriger Genehmigung der Rechteinhaber erlaubt. Die systematische Speicherung von Teilen des elektronischen Angebots auf anderen Servern bedarf ebenfalls des schriftlichen Einverständnisses der Rechteinhaber.

Haftungsausschluss

Alle Angaben erfolgen ohne Gewähr für Vollständigkeit oder Richtigkeit. Es wird keine Haftung übernommen für Schäden durch die Verwendung von Informationen aus diesem Online-Angebot oder durch das Fehlen von Informationen. Dies gilt auch für Inhalte Dritter, die über dieses Angebot zugänglich sind.

Sur l'énergie propre de l'électron

par André Mercier¹).

(9. XII. 38.)

Résumé. — L'énergie propre électrodynamique de l'électron provient de la possibilité d'échange de quantité de mouvement entre l'électron et le champ de fluctuation de la radiation du vide. On l'a calculée, au moyen de la méthode d'approximation de BORN, uniquement jusqu'à un terme qui fournit une expression divergente ayant $1/137$ en facteur (deuxième approximation). On a cherché des méthodes permettant de supprimer cette divergence. Or les méthodes proposées concernent la seconde approximation, mais rien ne permet d'affirmer que la série des diverses approximations auxquelles elles conduisent soit elle-même une expression finie, car on ne sait pas si cette série converge. C'est pour cette raison que nous avons calculé l'ordre de grandeur de chaque approximation, au moyen de la méthode d'approximation de BORN et en prenant $\sqrt{1/137}$ comme paramètre dans le développement en série. Les résultats sont les suivants: dans l'approximation d'ordre $2n$, le nombre de combinaisons possibles d'états intermédiaires est de l'ordre de $n!$; mais il ne nous a pas été possible de déterminer par une méthode générale le signe de chacun des termes. Dès lors il semble impossible de savoir si la somme de toutes les approximations converge. Nous faisons les calculs dans le cas d'un électron libre et dans celui de la Théorie des lacunes de DIRAC. Dans l'approximation d'ordre $2n$, l'énergie propre est proportionnelle à $Q^{2n}/137^n$ (où $Q \rightarrow \infty$) dans le cas d'un électron libre, et proportionnelle à

$$\frac{1}{137^n} \int_{Q \rightarrow \infty} d q_1 \frac{d q_2}{q_2} \frac{d q_3}{q_3} \dots \frac{d q_n}{q_n}$$

dans la théorie des lacunes.

Cela montre entre autres que la divergence logarithmique trouvée par WEISSKOPF dans la seconde approximation pour la théorie des lacunes est fortuite et qu'elle est due à une compensation de certains termes particulière à cette approximation et ne se produisant pas dans celles d'ordres supérieurs.

Il résulte de là qu'il ne suffit pas de considérer la seconde approximation: la théorie des lacunes conduit à des résultats beaucoup plus divergents que ne le suggère cette seconde approximation.

Enfin, en ce qui concerne le procédé de limitation des petites longueurs d'onde des photons, on ne peut pas non plus affirmer qu'il supprime toute divergence à cause de la divergence de la série-même des approximations, et il semble d'ailleurs actuellement difficilement justifiable. Peut-être la découverte de l'électron lourd permettra-t-elle de résoudre le problème²).

¹) Ce travail a été effectué lors d'un séjour de l'auteur à l'Institut de Physique Théorique de Copenhague. L'auteur désire remercier tout spécialement M. le prof. NIELS BOHR de son accueil et de ses conseils, ainsi que MM. CHR. MØLLER, L. ROSENFELD et V. WEISSKOPF.

²) L'auteur a publié avec M. GUSTAFSON une note sur l'énergie propre de l'électron parue dans les C. R. Acad. Sc. Paris, **206**, 1217, 1938.

1. Introduction.

La théorie des quanta attribue à l'électron une énergie propre d'origine électrodynamique qui n'a pas son correspondant dans la théorie classique. Ceci est dû au fait que, dans la théorie quantique, un électron est a priori en interaction avec un champ électromagnétique, et comme un électron est toujours susceptible d'émettre et d'absorber (virtuellement ou réellement) des quanta de lumière, il doit de ce fait posséder une énergie propre électro-dynamique. La radiation émissible par l'électron peut être considérée comme un champ ne consistant qu'en une fluctuation, mais pouvant être, comme tout champ électromagnétique, caractérisé par un potentiel vecteur $\vec{\Phi}$. Or dans la théorie de l'électron de Dirac, l'interaction entre l'électron et un champ est mesurée par le produit scalaire de l'opérateur $\vec{\alpha}$ associé à l'électron et du potentiel vecteur de ce champ. L'énergie propre d'un électron se calcule donc à partir de l'expression

$$\vec{\alpha} \cdot \vec{\Phi}.$$

On peut dire que cette énergie propre tire son origine de la possibilité d'échange de quantité de mouvement entre l'électron et ce champ de radiation virtuel. Etant donné un système soumis aux lois de la mécanique quantique, on a coutume de le traiter sous la forme d'un problème de perturbation et, celle-ci étant considérée comme petite, les diverses approximations sont données par des grandeurs proportionnelles aux puissances croissantes d'un paramètre. Pour ce dernier nous prendrons $\sqrt{1/137}$.

Les divers phénomènes (chocs entre particules, etc.) sont, suivant le cas, donnés par une approximation ou une autre. Deux approximations différentes correspondent à deux effets physiquement différenciés. Mais dans le problème particulier de l'énergie propre on a affaire à un système composé d'un électron unique, dont on calcule l'action sur lui-même, de sorte que chaque approximation contribue dans une certaine part à la valeur de l'énergie propre totale.

Jusqu'ici on ne s'est occupé à notre connaissance que de la deuxième approximation, et nous remarquons tout de suite qu'il ne faut pas confondre l'expression d'énergie propre employée dans la littérature (voir les diverses citations bibliographiques) avec celle de l'énergie propre totale qui nous intéresse dans ce travail. C'est pourquoi nous dirons toujours : énergie propre dans l'approximation d'ordre n .

Du fait que la théorie de l'électron suppose celui-ci ponctuel, son énergie propre est infinie, à moins que l'on ne prenne des pré-

cautions très particulières. Par précautions, il faut entendre l'introduction d'une hypothèse supplémentaire ayant pour but de corriger l'effet dû à la concentration de l'électron en un point (BORN, WATAGHIN, MARCH), soit celle d'un artifice mathématique au cours des calculs, du genre p. ex. d'un prolongement analytique (WENTZEL). Le procédé le plus commun qui consisterait à abolir sans autre toute émission de quanta lumineux ayant une impulsion supérieure à une certaine limite a le désavantage fatal de détruire l'invariance relativiste. A cause de ce fait, WATAGHIN¹⁾ a proposé d'introduire une modification dans la fonction hamiltonienne, ayant pour effet de rendre extrêmement rares les émissions de photons de grande fréquence, ce procédé étant choisi de manière à garder l'invariance relativiste. Mais il est naturellement arbitraire. WENTZEL²⁾ a montré qu'en procédant à un passage à la limite par un chemin détourné grâce à l'introduction du formalisme pluri-temporel de DIRAC-FOCK-PODOLSKI, on améliore les résultats.

Mais il n'est pas parvenu à supprimer toute divergence, et ses conclusions ne sont valables que dans le cadre d'une approximation déterminée (en particulier, la seconde). MARCH³⁾ a montré qu'on peut justifier une limitation supérieure des impulsions des photons émis en postulant que la géométrie ordinaire de MIN-KOWSKI n'est pas valable dans l'infiniment petit de la physique quantique, ou plus précisément que le ds^2 ne peut pas être mesuré avec une précision rigoureuse. Il pose

$$ds = \sqrt{\sum_i^4 dx_i^2} - \gamma$$

et son hypothèse revient à dire qu'il n'y a pas d'interaction entre l'électron et des photons dont la longueur d'onde est inférieure à γ . Cela permet une limitation qui ne détruit pas l'invariance relativiste. Mais, comme dans le procédé de WENTZEL, cette manière de faire est correcte dans une approximation, sans qu'il soit certain qu'elle l'est dans l'ensemble des approximations successives (énergie propre totale).

M. WEISSKOPF nous a fait remarquer l'intérêt qu'il y aurait à étudier le problème suivant: que peut-on dire de la convergence ou de la divergence des termes de l'énergie propre proportionnels

¹⁾ G. WATAGHIN, Zeitschr. f. Phys. (88, 92, 1934).

²⁾ G. WENTZEL, Zeitschr. f. Phys. (86, 479 et 635, 1933, et 87, 726, 1934).

³⁾ A. MARCH, Zeitschr. f. Phys. (104, 93 et 161, 1937, 105, 620, 1937, 106, 49, 291 et 532, 1937, 107, 144, 1937, 108, 128, 1938).

aux diverses puissances de $\sqrt{1/137}$, ainsi que de la série de ces termes ?

La réponse à cette question dépend sans doute de la définition que l'on donne du vide: on bien l'on considère l'électron comme absolument libre, c'est-à-dire capable (virtuellement) de passer par des états d'énergie positive ou négative quelconques, — ou bien, dans la théorie des lacunes de DIRAC, on le considère comme libre uniquement par rapport aux états d'énergie positive, puisque tous les états d'énergie négative sont supposés occupés par des électrons; de plus il se peut que l'électron envisagé entre en interaction avec les «électrons du vide». Ces deux faits: interaction possible et liberté restreinte, doivent avoir une influence telle que l'énergie propre (autant dans chaque approximation que dans sa totalité) diffère vraisemblablement considérablement suivant qu'on la calcule dans la théorie des lacunes ou dans l'hypothèse d'un électron absolument libre.

On sait que c'est le cas déjà en seconde approximation, dans laquelle WALLER¹⁾ a calculé l'énergie propre de l'électron libre, et WEISSKOPF²⁾ celle de l'électron dans la théorie des lacunes. Nous désignerons, dans n'importe quelle approximation ces deux possibilités par les expressions théorie ordinaire et théorie des lacunes.

La constante de structure fine $\sqrt{1/137}$ contient la charge e de l'électron en facteur. Or l'énergie propre ne peut dépendre des signes de la charge, de sorte que seules les approximations d'ordre pair (puissances de $1/137$) doivent fournir un résultat différent de zéro.

Dans ce travail, nous calculons tout d'abord avec quelque détail la quatrième approximation (terme proportionnel à $1/137$)²⁾ dans les deux théories. Puis nous cherchons ce qui se passe dans chacune de ces théories dans les approximations supérieures et discutons les avantages et les inconvénients de l'une et de l'autre.

2. Méthode de calcul.

On calcule les approximations successives au moyen de la formule de BORN, grâce à l'introduction d'un opérateur de perturbation exprimant la possibilité d'échange d'impulsion entre l'électron et le champ de radiation dont le vide pourrait être le siège. Cette radiation n'est due qu'à des fluctuations.

¹⁾ I. WALLER, Zeitschr. f. Phys. **62**, 673, 1930.

²⁾ V. WEISSKOPF, Zeitschr. f. Phys. **89**, 27, 1934 et **90**, 817, 1934.

Soit $\vec{\alpha}$ le vecteur (tridimensionnel) dont les composantes sont les opérateurs de DIRAC, et $\vec{\Phi}$ le potentiel vecteur du champ électromagnétique des fluctuations. Soit ε la charge de l'électron et ψ_a, ψ_b deux fonctions d'onde décrivant l'état du système électron-radiation. L'opérateur de perturbation qui fait passer le système de l'état b à l'état a s'écrit

$$V_{ab} = \varepsilon \int \psi_a^* \vec{\alpha} \cdot \vec{\Phi} \psi_b. \quad (1)$$

Le signe $\int$ désigne une intégration sur l'espace ordinaire, une sommation sur l'espace de configuration des photons, et une sur l'indice de spin. Les divergences qui apparaîtront dans la suite et qui sont le sujet de la discussion proviennent de la sommation sur les photons, parce qu'il n'y a en principe pas de limite à la grandeur de l'impulsion qu'un électron peut transmettre au champ de la radiation.

La fonction ψ d'un système composé d'un électron et de photons se développe comme suit:

$$\psi = \sum_{\mu} a(M_{\vec{p}}^{\mu}) u^{\mu}(\vec{p}) e^{\frac{i}{\hbar} \vec{p} \cdot \vec{x}} \prod_{\vec{q}} \delta(N_{\vec{q}})$$

p est l'impulsion de l'électron, $\vec{x}$ le vecteur de sa position. $\prod_{\vec{q}}$ est le produit des fonctions $\delta(N_{\vec{q}})$ associées aux photons présents. $N_{\vec{q}}$ est le paramètre de la distribution des photons et désigne le nombre de photons présents dont l'impulsion vaut $\vec{q}$.

Les opérateurs $a(M_{\vec{p}}^{\mu})$ satisfont aux conditions de la statistique de FERMÍ-DIRAC:

$$a^*(M_{\vec{p}}^{\mu}) a(M_{\vec{p}}^{\mu}) = M_{\vec{p}}^{\mu}, \quad a(M_{\vec{p}}^{\mu}) a^*(M_{\vec{p}}^{\mu}) = 1 - M_{\vec{p}}^{\mu},$$

où $M_{\vec{p}}^{\mu} = \left\{ \begin{smallmatrix} 1 \\ 0 \end{smallmatrix} \right\}$ est le nombre d'électrons d'impulsion $\vec{p}$ et dans l'état μ ($\mu = 1, 2, 3, 4$ désigne d'un même coup le signe de l'énergie et l'état du spin). Les $u^{\mu}(\vec{p})$ sont des spineurs. La fonction ψ est supposée normée de manière que $\int \psi^* \psi = 1$.

Pour introduire explicitement l'impulsion des photons, on développe le potentiel vecteur en série de FOURIER

$$\vec{\Phi}(\vec{x}) = \sum \frac{\vec{\sigma} T}{\sqrt{q}} \left[c(\vec{q}) e^{i \frac{\vec{q} \cdot \vec{x}}{\hbar}} + c^*(\vec{q}) e^{-i \frac{\vec{q} \cdot \vec{x}}{\hbar}} \right]$$

où $\vec{\sigma}$ est le vecteur de polarisation de l'onde partielle associée au photon d'impulsion $\vec{q}$, ($\vec{\sigma} \perp \vec{q}$), $c(\vec{q})$ et $c^*(\vec{q})$ sont les opérateurs de

la statistique de BOSE; il y correspond respectivement l'absorption et l'émission d'un photon. T a pour carré:

$$T^2 = \frac{h^2 c}{\Omega}$$

où Ω est le volume du vide.

La sommation sur l'espace d'impulsion se fait au moyen de l'introduction d'une fonction densité des états au point $\vec{q}$; soit $d\omega$ un angle solide infinitésimal d'axe $\vec{q}$, et $q = |\vec{q}|$. La sommation est à la limite une intégrale $\int d\vec{q} (...)$, où

$$d\vec{q} = \frac{\Omega q^2 dq d\omega}{h^3}.$$

Disons tout de suite que pour chaque émission suivie de l'absorption d'un quantum $\vec{q}$ arbitraire, il interviendra dans les calculs l'opération

$$\int \frac{\varepsilon^2 T^2 d\vec{q} (...)}{q},$$

et l'on a

$$\frac{\varepsilon^2 T^2 d\vec{q}}{q} = \frac{d\omega c^2 q dq}{137}.$$

Comme le calcul sera fait en détail en quatrième approximation, écrivons la formule de BORN pour ce cas et donnant l'énergie propre E_ε d'un électron ε dont l'état est fixé par l'indice l (la fonction d'onde ψ_l est celle d'un état où aucun photon n'est présent):

$$E_\varepsilon = \sum'_{mk i} \frac{V_{lm} V_{mk} V_{ki} V_{il}}{(E_i - E_l)(E_k - E_l)(E_m - E_l)} - \sum'_{mn} \frac{V_{lm} V_{ml} V_{ln} V_{nl}}{(E_n - E_l)^2 (E_m - E_l)} \quad (2)$$

i, k, m, n sont des indices d'états intermédiaires, $E_l, \dots, E_n$ l'énergie du système dans l'état $l, \dots, n$.

Par raison de commodité (ou si l'on préfère, par suite d'une transformation de LORENTZ convenable) nous choisissons pour état l celui d'un électron au repos:

$$E_l = mc^2 \text{ dans la théorie ordinaire,}$$

$E_l = mc^2 + \text{somme des énergies des électrons du vide, dans la théorie des lacunes.}$

Nous désignerons enfin par A une suite d'états intermédiaires rentrant dans la somme

$$\sum'_{mki} \frac{V_{lm} V_{mk} V_{ki} V_{il}}{(E_i - E_l)(E_k - E_l)(E_m - E_l)} \quad (3)$$

et par B une suite rentrant dans la somme

$$\sum'_{mn} \frac{V_{lm} V_{ml} V_{ln} V_{nl}}{(E_n - E_l)^2 (E_m - E_l)}. \quad (4)$$

3. Théorie ordinaire en quatrième approximation.

Désignons par $\varepsilon(\vec{p})$ un électron d'impulsion $\vec{p}$ et par $\varphi(\vec{q})$ un photon d'impulsion $\vec{q}$. Un électron libre peut passer par des suites d'états intermédiaires au nombre de deux du type A , que nous appelons A_1 et A_2 , et une du type B . A_1 est la série suivante:

A_1 : état fondamental	$\varepsilon(\vec{p} = 0)$
1er état interméd.	$\varepsilon(\vec{p}'), \varphi(\vec{q})$
2ième état interméd.	$\varepsilon(\vec{p}''), \varphi(\vec{q}), \varphi(\vec{q}')$
3ième état interméd.	$\varepsilon(\vec{p}'), \varphi(\vec{q})$
état final	$\varepsilon(\vec{p} = 0)$.

L'opérateur (1) qui fait passer de l'état l à l'état i est

$$\begin{aligned} V_{il} &= \frac{\varepsilon T}{\sqrt{q}} \sum_{\mu\nu} \int d\vec{x} \delta(1_{\vec{q}}) \delta(0_{\vec{q}}) \left\{ a^*(M_{\vec{p}'}^{\mu}) u^{\mu*}(\vec{p}') e^{-\frac{i}{\hbar} \vec{p}' \cdot \vec{x}}, \vec{\alpha} \cdot \vec{\sigma} c^*(q) \right. \\ &\quad \times e^{-\frac{i}{\hbar} \vec{q} \cdot \vec{x}} a(M_{\vec{p}}^{\nu}) u^{\nu}(\vec{p}) e^{\frac{i}{\hbar} \vec{p} \cdot \vec{x}} \left. \right\} \delta(0_{\vec{q}}) \delta(0_{\vec{q}}) \\ &= \text{const.} \delta(-\vec{p}' - \vec{q} + \vec{p}) \sum_{\mu\nu} \delta(0_{\vec{q}}) \delta(1_{\vec{q}}) a^*(M_{\vec{p}'}^{\mu}) a(M_{\vec{p}}^{\nu}) \\ &\quad \{u^{\mu*}, \vec{\alpha} \cdot \vec{\sigma} u^{\nu}\} \delta(1_{\vec{q}}) \delta(0_{\vec{q}}) \end{aligned}$$

où $\delta(-\vec{p}' - \vec{q} + \vec{p})$ est la fonction delta, et $\{u^*, u'\}$ le produit scalaire de deux spineurs u et u' . Si l'on calcule les autres $V...$, on obtient pour A_1 , après quelques transformations

$$\begin{aligned} A_1: \varepsilon^4 T^4 \sum_{\kappa\lambda\mu\nu} \sum_{\vec{q}\vec{q}'\vec{p}} \frac{M_{\vec{p}}^{\nu} [1 - M_{\vec{p}-\vec{q}}^{\kappa}] [1 - M_{\vec{p}-\vec{q}-\vec{q}'}^{\lambda}] [1 - M_{\vec{p}-\vec{q}'}^{\mu}]}{q q' (E_i - E_l) (E_k - E_l) (E_m - E_l)} \\ \{u^{\nu*}(\vec{p}), \vec{\alpha} \cdot \vec{\sigma} u^{\kappa}(\vec{p} - \vec{q})\} \{u^{\kappa*}(\vec{p} - \vec{q}), \vec{\alpha} \cdot \vec{\sigma}' u^{\lambda}(\vec{p} - \vec{q} - \vec{q}')\} \\ \times \{u^{\lambda*}(\vec{p} - \vec{q} - \vec{q}'), \vec{\alpha} \cdot \vec{\sigma}' u^{\mu}(\vec{p} - \vec{q})\} \{u^{\mu*}(\vec{p} - \vec{q}), \vec{\alpha} \cdot \vec{\sigma} u^{\nu}(\vec{p})\} \\ = \text{const.} \sum_{\kappa\lambda\mu\nu} \int q dq q' dq' d\omega d\omega' \\ \times \frac{\{u^{\nu*}, \vec{\alpha} \cdot \vec{\sigma} u^{\kappa}\} \{u^{\kappa*}, \vec{\alpha} \cdot \vec{\sigma}' u^{\lambda}\} \{u^{\lambda*}, \vec{\alpha} \cdot \vec{\sigma}' u^{\mu}\} \{u^{\mu*}, \vec{\alpha} \cdot \vec{\sigma} u^{\nu}\}}{(E_i - E_l) (E_k - E_l) (E_m - E_l)} \end{aligned}$$

$d\omega$ et $d\omega'$ sont les angles solides d'axes $\vec{q}$ et $\vec{q}'$.

Les valeurs de q, q' qui conduisent à des divergences éventuelles sont celles qui tendent vers l'infini. On a pour celles-ci approximativement

$$\left. \begin{aligned} E_i &= \pm c \sqrt{(-\vec{q})^2 + m^2 c^2} + cq \approx c(\pm q + q) \\ E_k &\approx c(\pm s + q + q') \quad (\vec{s} = \vec{q} + \vec{q}', s = |\vec{s}|) \\ E_m &= E_i. \end{aligned} \right\} \quad (5)$$

C'est donc le signe — qui conduit à la plus grande divergence; on a dans ce cas

$$E_i - E_l \approx -mc^2, \text{ etc.}$$

On effectue la sommation sur les indices de spin à l'aide de l'opérateur $\frac{1}{2} [1 \pm D(\vec{p}) / E(\vec{p})]$, où $D(\vec{p}) = c \vec{\alpha} \cdot \vec{p} + \beta mc^2$, et E est la valeur propre de D . Le produit des quatre accolades se ramène alors à une expression S_1 , dont le calcul est celui de la trace d'une matrice.

La série des états intermédiaires de A_2 et celle de B sont les suivantes:

$$\begin{array}{ll} A_2: & \varepsilon(\vec{p} = 0) \\ & \varepsilon(\vec{p}'), \varphi(\vec{q}) \\ & \varepsilon(\vec{p}''), \varphi(\vec{q}), \varphi(\vec{q}') \\ & \varepsilon(\vec{p}'''), \varphi(\vec{q}') \\ & \varepsilon(\vec{p} = 0) \\ B: & \varepsilon(\vec{p} = 0) \\ & \varepsilon(\vec{p}'), \varphi(\vec{q}') \\ & \varepsilon(\vec{p} = 0) \\ & \varepsilon(\vec{p}''), \varphi(\vec{q}'') \\ & \varepsilon(\vec{p} = 0). \end{array}$$

Le calcul des accolades fournit une expression S_2 pour A_2 et la valeur 1 pour B . Les dénominateurs de A_1 et A_2 sont les mêmes, et celui de B est $(mc^2)^3$. Il en résulte pour la valeur de l'énergie propre:

$$E_\varepsilon = \frac{8\pi^2}{137^2 \cdot m^2 c} \int_0^\pi \sin \vartheta d\vartheta \int_0^\infty dq \int_0^\infty dq' q q' \left\{ \frac{S_1 + S_2}{q + q' - mc - s} + \frac{1}{mc} \right\}$$

où $\vartheta = \angle(\vec{q}, \vec{q}')$.

Quelles que soient les valeurs de $\vec{q}$ et $\vec{q}'$, S_1 et S_2 restent bornés. On peut alors chercher l'ordre de grandeur de l'intégrale $\int \sin \vartheta d\vartheta \frac{S_1 + S_2}{q + q' - mc - s}$. Or $q + q' - mc - s$ peut s'annuler. Si l'on calcule la valeur principale de CAUCHY, on trouve que cette intégrale est de l'ordre de $1/q^1$). Cela montre que le terme $1/mc$

¹⁾ Le calcul rigoureux nous a été indiqué par M. T. GUSTAFSON.

dans l'accolade joue seul un rôle dans le calcul de la plus grande divergence de E_ε , et pour celle-ci on trouve donc

$$E_\varepsilon \sim \frac{(2\pi)^2 c}{137^2 m^3 c^3} / q^4 /_{q \rightarrow \infty}.$$

Donc dans la théorie ordinaire, E_ε diverge comme la quatrième puissance d'un paramètre tendant vers l'infini.

4. Théorie des lacunes en quatrième approximation.

Le vide, dans la théorie des lacunes, est défini par le système des électrons infiniment nombreux qui occupent les niveaux d'énergie négative. Soit E_V l'énergie propre de ce système. Si au vide ainsi défini, on ajoute un électron (désigné par ε_0) on obtient un système de « $\infty + 1$ » électrons dont l'énergie propre sera désignée par $E_{V+\varepsilon_0}$. On définit alors l'énergie propre de l'électron ε_0 par la relation

$$E_{\varepsilon_0} = E_{V+\varepsilon_0} - E_V. \quad (6)$$

Pour le calcul de (6), on emploie la formule de BORN (2). Mais il est cependant impossible d'évaluer $E_{V+\varepsilon_0}$ et E_V séparément, et il faut tout d'abord trouver une autre expression de E_{ε_0} que (6), où l'énergie du vide n'apparaisse plus explicitement. On y arrive par le raisonnement suivant: $E_{V+\varepsilon_0}$ se compose de trois parties, 1° l'énergie E'_{ε_0} de l'électron ε_0 considéré comme libre dans le demi-espace des énergies positives; 2° l'énergie C de couplage entre ε_0 et les électrons du vide; 3° l'énergie des électrons du vide corrigée pour tenir compte du fait que dans les états intermédiaires, les électrons n'ont pas le droit d'occuper la place de ε_0 , c'est-à-dire l'énergie du vide E_V diminuée de termes exprimant des sauts qui sont interdits aux électrons du vide; soit J la contribution à E_V de ces sauts interdits. (6) s'écrit donc

$$E_{\varepsilon_0} = E'_{\varepsilon_0} + C + (E_V - J) - E_V = E'_{\varepsilon_0} + C - J. \quad (7)$$

Si l'on cherchait à décomposer la formule (6) en portant son attention sur les valeurs que peuvent prendre les paramètres de distribution M_p^μ des électrons, on se persuaderait que la relation (7) est exactement la transposition de (6), étant convenu que dans chacun des trois termes de (7) on égale systématiquement à *un* les paramètres en question. En effet, soit $f(\dots M_p^\mu \dots)$ la fonction des paramètres de la distribution des électrons qui apparaît dans chaque terme par suite des conditions imposées aux opérateurs

$a(M_{\vec{p}}^{\mu}), a^*(M_{\vec{p}'}^{\nu})$. L'énergie du système des $(\infty + 1)$ électrons (électrons du vide et électron ε_0 considéré) est donnée en quatrième approximation par la somme de tous les termes comportant n'importe quels états intermédiaires au nombre de trois, mais où la valeur de $f(\dots M_{\vec{p}}^{\mu} \dots)$ doit être calculée rigoureusement. Or $f(\dots M_{\vec{p}}^{\mu} \dots)$ est précisément nulle pour les termes que nous avons interdits, de sorte qu'en soustrayant E_V de $E_{V+\varepsilon_0}$ dans cette définition, ces termes persistent avec le signe $-$; ce sont justement les termes $-J$.

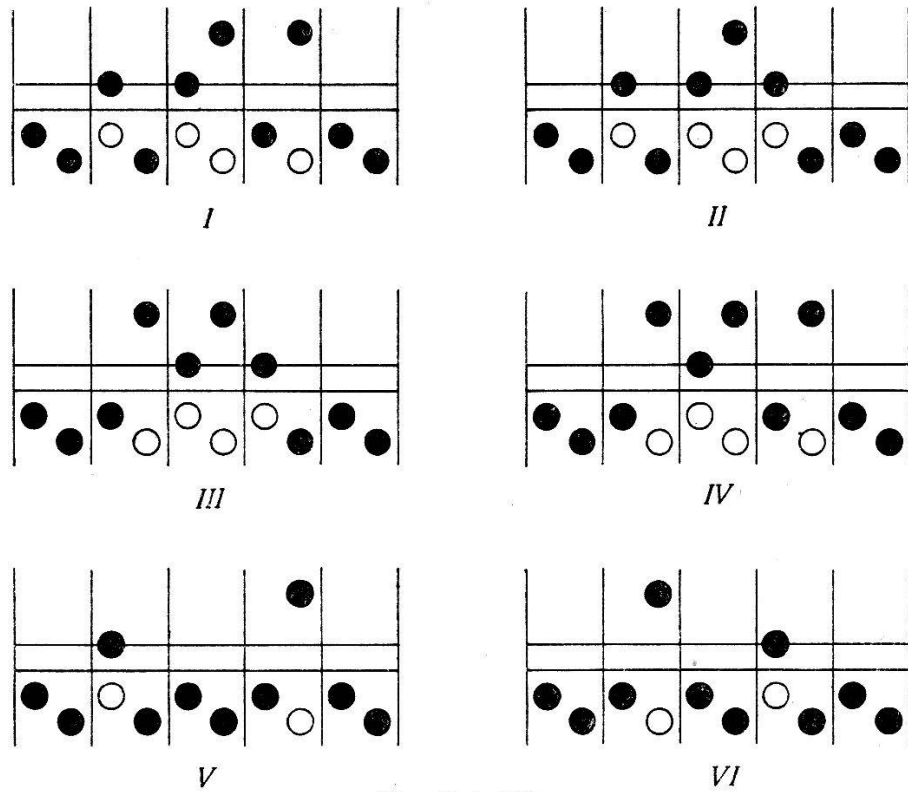


Fig. I à VI.

Schéma d'états intermédiaires (appartenant à J).

La relation (7) est beaucoup plus commode puisqu'elle ne contient plus l'énergie E_V . Elle est naturellement valable dans toutes les approximations. Considérons en particulier la quatrième. Nous séparons d'après (3) et (4) l'énergie propre en des termes A et des termes B.

Envisageons tout d'abord J , c'est-à-dire l'énergie à soustraire provenant des sauts interdits aux électrons du vide. Les figures I, II, III, IV représentent certains termes du type A et les figures V et VI d'autres du type B. Ces figures (cf. aussi la fig. 1) représentent par cinq colonnes l'état fondamental, les trois états intermédiaires et l'état final: ● désigne un électron; l'énergie de celui-ci y est mesurée par la hauteur au-dessus de mc^2 ou au-dessous de

— mc^2 . Un cercle $\circ$ signifie une lacune apparue dans le vide. Dans les séries I à VI, l'émission d'un photon par l'un des électrons du vide entrant en jeu est soumise à la restriction que cet électron se trouve dans un des états intermédiaires au niveau mc^2 (celui de ε_0). L'autre électron par contre émet un quantum d'impulsion arbitraire. Il en résulte que dans ces termes se présentent trois sommations: l'une sur l'impulsion contrainte du photon émis (et réabsorbé) par l'un des électrons, ou ce qui revient au même sur l'impulsion initiale de cet électron, la seconde sur l'impulsion arbitraire du second électron, et la troisième sur l'impulsion arbitraire du photon que ce dernier électron émet (et réabsorbe).

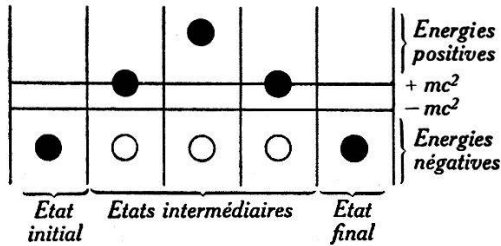
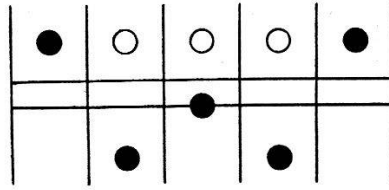

 Fig. 1. Electron: •, Lacune: $\circ$


Fig. 2.

Chacune de ces trois sommations, qui à la limite sont des intégrations, nécessite l'introduction d'un facteur $d\vec{p}$ proportionnel au volume Ω . Or l'émission suivie de l'absorption de deux photons au total n'introduit en dénominateur que Ω^2 , de sorte que certains termes J contiennent Ω à la première puissance.

Tous les termes de couplage sont représentés dans les figures i , ii , ... iv (type A), et v et vi (type B). Ils sont pour les mêmes raisons proportionnels à Ω .

Or il serait fort étonnant que l'énergie propre de l'électron ε_0 dépende du volume Ω de l'espace, car dans la définition (6) il n'est pas tenu compte des actions entre les électrons dépendant de leurs distances mutuelles, et d'autre part l'introduction de Ω ne sert qu'à permettre le dénombrement des états énergétiques possibles.

Nous allons montrer que les parties de J et de C qui sont proportionnelles à Ω sont identiquement nulles. Considérons tout d'abord J . Soit $\vec{p}$ l'impulsion d'un électron du vide, $\vec{p}'$ celle d'un autre, $\vec{q}$ et $\vec{q}'$ l'impulsion des photons émis. Dans le cas I par exemple, $\vec{p} = \vec{q}$. Posons

$$\left. \begin{aligned} c \sqrt{p^2 + m^2 c^2} + m c^2 + c q &= c P_1 \\ c \sqrt{p^2 + m^2 c^2} + c \sqrt{p'^2 + m^2 c^2} + m c^2 \\ &\quad + c \sqrt{(\vec{p}' - \vec{q}')^2 + m^2 c^2} + c q + c q' = c P_2 \\ c \sqrt{p'^2 + m^2 c^2} + c \sqrt{(\vec{p}' - \vec{q}')^2 + m^2 c^2} + c q' &= c P_3. \end{aligned} \right\} \quad (8)$$

Le terme I s'écrit au complet

$$I = \frac{c}{137^2} \sum_{\vec{p}} \sum_{\text{spin}} \int \frac{d\omega d\omega' q dq' q' dq'}{P_1 P_2 P_3} \{u^*(\vec{p}) \cdot v^t(\vec{p}'), \vec{\alpha} \cdot \vec{\sigma}' u(\vec{p}) v(\vec{r})\} \\ \times \{u^*(\vec{p}) v^*(\vec{r}), \vec{\alpha} \cdot \vec{\sigma} u(0) v(\vec{r})\} \{u^*(0) v^*(\vec{r}), \vec{\alpha} \cdot \vec{\sigma}' u(0) v(\vec{r}')\} \\ \times \{u^*(0) v^*(\vec{p}'), \vec{\alpha} \cdot \vec{\sigma} u(p) v(\vec{p}')\} \quad (9)$$

où u et v sont les spineurs des électrons considérés, $\vec{\sigma}$ et $\vec{\sigma}'$ les vecteurs de polarisation des photons et $d\omega, d\omega'$ les angles solides

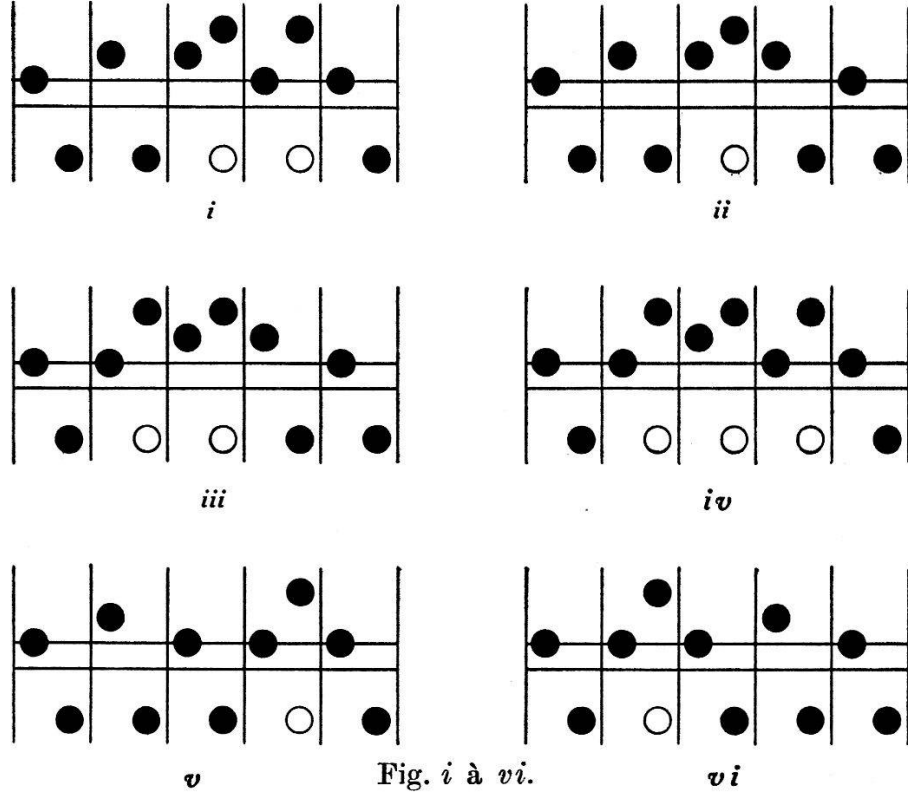


Fig. i à vi.

Schémas d'états intermédiaires (appartenant à C).

d'axes $\vec{q}$ et $\vec{q}'$. Pour II, III ... VI on trouve des expressions analogues à (9). Dans ces expressions, le produit des accolades se ramène à la trace $\mathfrak{S}$ d'une matrice. Du fait que deux électrons entrent en jeu dans les états intermédiaires, les traces $\mathfrak{S}$ sont toutes les mêmes pour les termes I à VI. Nous pouvons donc, en laissant de côté l'opération

$$\frac{c}{137^2} \sum_{\vec{p}'} \int d\omega d\omega' q dq' q' dq' \mathfrak{S}[\dots]$$

écrire

$$I \sim \frac{1}{P_1 P_2 P_3}, \quad II \sim \frac{1}{P_1^2 P_2}, \quad III \sim \frac{1}{P_1 P_2 P_3}, \quad IV \sim \frac{1}{P_2 P_3^2} \\ V \sim \frac{1}{P_1^2 P_3}, \quad VI \sim \frac{1}{P_1 P_3^2}.$$

Or en vertu des définitions (8), $P_2 = P_1 + P_3$. Donc

$$I + II - V \equiv 0, \quad III + IV - VI \equiv 0,$$

ce qui montre que les séries d'états intermédiaires interdits contenant Ω ne jouent aucun rôle.

Quant au couplage, il est représenté complètement par les figures *i* à *vi*. Les six termes contiennent la même trace $\mathfrak{S}$, et comme les dénominateurs s'obtiennent simplement à partir de ceux de I à VI en remplaçant m par $-m$, il en résulte que

$$C = 0.$$

Les figures 1 et 2 représentent des séries d'états intermédiaires pour les termes interdits de J où un seul électron du vide est mis en jeu. Elles sont du type A. Il existe deux autres possibilités du type A où l'ordre d'émission des deux photons n'est pas le même que celui de leur absorption. Il y a deux possibilités du type B. Ces 6 termes de J ne sont pas proportionnels à Ω .

Appelons J' les termes non nuls de J . L'énergie propre s'écrit maintenant en quatrième approximation

$$E_{\varepsilon_0} = E_{\varepsilon_0}' - J'. \quad (10)$$

Avant de calculer (10), il est indiqué de faire une remarque sur le calcul de l'énergie propre fait par WEISSKOPF en deuxième approximation¹⁾. Tout d'abord il est impossible qu'en seconde approximation il apparaisse des termes proportionnels à Ω , et il n'y a d'ailleurs pas de couplage entre ε_0 et le vide, vu qu'il n'y a qu'un seul état intermédiaire. L'énergie propre dans cette approximation, et dans la théorie des lacunes, se réduit donc d'emblée à une expression très semblable à (10), écrivons-la

$$E_{\varepsilon_0} = E_{\varepsilon_0}' - K \quad (11)$$

où il est entendu que E_{ε_0}' et E_{ε_0} n'ont pas la même signification que dans (10). Dans (11), E_{ε_0}' est représenté par le schéma de la figure 3, trait plein, indiquant le saut qu'effectue l'électron ε_0 à un niveau énergétique $\Delta + mc^2$ dans l'état intermédiaire. Le trait pointillé de la même figure est le schéma qui donne K , c'est-à-dire le saut particulier interdit aux électrons du vide. Or si $\Delta \gg mc^2$, on peut confondre approximativement mc^2 et $-mc^2$, en les égalant à zéro, de sorte que si l'électron du vide part d'un niveau d'énergie $-mc^2 - \Delta$, pour sauter au niveau $+mc^2$, on peut dire approximativement que les sauts de ε_0 et de l'électron du vide sont les

¹⁾ V. WEISSKOPF, loc. cit.

mêmes à un déplacement près. Or WEISSKOPF a montré, par une méthode de correspondance équivalente à celle des états intermédiaires, qu'il se produit une compensation considérable dans le calcul de E_{ε_0} ayant pour effet d'en réduire la divergence de linéaire (comme on doit s'attendre à la trouver tout d'abord) à logarithmique. C'est justement du fait de la similitude des deux sauts indiqués à la figure 3 et par conséquent du fait qu'on peut négliger mc devant des valeurs suffisamment grandes de l'impulsion du photon mis en jeu, que cette compensation se produit.

Ainsi s'explique le fait que dans la théorie des lacunes la divergence est réduite en seconde approximation. On peut se demander si un fait pareil se produit dans la quatrième. Pour examiner cette question, on se rapportera à la figure 4. Le trait

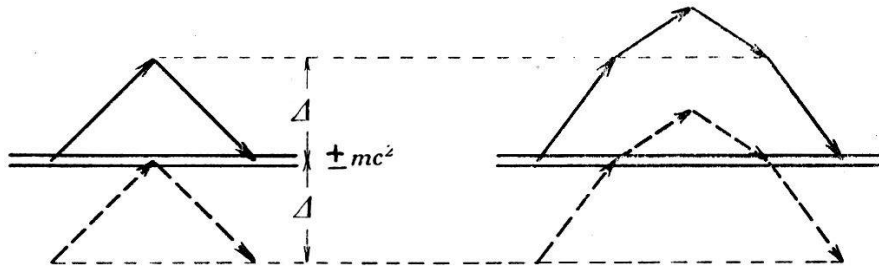


Fig. 3.

Compensation des termes en
seconde approximation.

Fig. 4.

Absence de compensation des termes
en quatrième approximation.

plein y donne une série du type A pour E'_{ε_0} , et le trait pointillé une série du même type pour J' , où l'un des quanta émis a une énergie (si elle est suffisamment grande) à peu près la même dans les deux cas. Jusqu'ici il y a analogie avec le cas de WEISSKOPF. Mais lors des passages du premier au deuxième état et du deuxième au troisième état intermédiaire, il n'y a pas de compensation possible entre le trait plein et le pointillé, parce que ce dernier seul passe par le niveau zéro. La réduction de la divergence est donc un phénomène propre à la deuxième approximation seulement. Toutefois, dans des séries du type B, la compensation se produit, parce que le type B en quatrième approximation est formellement la suite de deux de la seconde.

De cette discussion résulte qu'il suffit de considérer les termes du type A dans le calcul du second membre de la relation (10). La figure 4 montre aussi que les termes de E'_{ε_0} et ceux de $-J'$ dans (10) présentent une divergence du même ordre de grandeur. C'est pourquoi E_{ε_0} est à un facteur près, — d'ailleurs au plus égal à 6 (nombre total des séries), — de l'ordre d'un terme A_1' , par exemple, obtenu en calculant au lieu de A_1 (défini au § 3) de

la théorie ordinaire, ce que l'on obtient en choisissant dans les relations (5) le signe + pour l'énergie de l'électron dans les états intermédiaires, au lieu du signe —. On trouve donc en quatrième approximation dans la théorie des lacunes

$$E_{\varepsilon_0} \cong \frac{8 \pi^2 c}{137^2} \int_0^\pi \sin \vartheta d\vartheta \int_0^\infty dq \int_0^\infty dq' qq' \varphi(q, q')$$

où la fonction $\varphi(q, q')$ est de l'ordre de Q^{-3} lorsque q et q' tendent vers une valeur très grande Q . Il en résulte que la divergence contenue dans (12) n'atteint pas l'ordre de grandeur de Q^2 .

5. Approximations successives.

Nous avons déjà remarqué que ce sont seules les approximations d'ordre pair qui conduisent à un résultat différent de zéro. Le passage de l'approximation d'ordre $2n$ à celle d'ordre $2n + 2$ ¹⁾ signifie l'introduction de deux états intermédiaires supplémentaires, donc l'émission et l'absorption d'un photon. Dans l'approximation d'ordre $2n$, n photons sont mis en jeu. Pour chacun d'eux

on a dans l'expression de l'énergie propre un facteur $\frac{\varepsilon^2 T^2 d\vec{q}}{q}$.

L'approximation d'ordre $2n$ conduit donc à une énergie proportionnelle à $(137)^{-n}$; c'est un terme d'une progression géométrique. Dans chaque approximation on a une somme de termes, des types A et B, et dont le nombre est d'autant plus grand que n est grand; d'abord, autant dans la théorie ordinaire que dans la théorie des lacunes, parce qu'on peut combiner les émissions et les absorptions de manières d'autant plus nombreuses que n est grand; ensuite, mais dans la théorie des lacunes seulement, parce que les quanta mis en jeu sont émis éventuellement par divers électrons.

Le nombre de possibilités de ranger dans un ordre ou dans un autre les émissions et absorptions des photons est soumis à la condition que l'émission d'un certain photon ne soit jamais précédée de son absorption. On doit donc chercher à ranger n couples d'objets dans $2n$ cases, chaque couple étant composé d'une émission et d'une absorption. On calcule le nombre de possibilités de le faire comme suit: Plaçons le premier couple. Il y a $\frac{2n(2n-1)}{2}$ possibilités de le faire. Plaçons le deuxième couple dans l'une des $2n - 2$ cases restantes. Il y a $\frac{(2n-2)(2n-2-1)}{2}$ manières de le

¹⁾ Nous dirons pour simplifier le passage d'une approximation à la suivante.

faire, etc.; il y a $\frac{(2n-2p)(2n-2p-1)}{2}$ possibilités de placer le p ème couple dans $2n-2p$ cases. Le nombre de possibilités de placer les n couples est égal au produit de ces expressions, et il vaut $\frac{(2n)!}{2^n}$. En réalité, dans le problème que nous étudions, il n'est pas possible de différencier les n couples de photons qu'il s'agit de distribuer, autrement dit leurs permutations ne sont pas considérées comme fournissant des distributions différentes. Ces permutations sont au nombre de $n!$. D'où le nombre $\mathfrak{N}$ de possibilités de ranger les n quanta émis en émissions et absorptions lors des passages d'un état intermédiaire à l'autre:

$$\mathfrak{N} = \frac{(2n)!}{2^n n!}. \quad (13)$$

Pour $n = 1$, $\mathfrak{N} = 1$; il y a une seule possibilité d'émettre et absorber un quantum en deuxième approximation (cas de WALLER). Pour $n = 2$, $\mathfrak{N} = 3$; il y a trois manières de distribuer l'émission et l'absorption de deux photons en quatrième approximation, lorsqu'un électron unique est en jeu (ce sont les possibilités A_1 , A_2 et B dans la théorie ordinaire).

Théorie ordinaire. — Ces remarques préliminaires étant faites, considérons de plus près tout d'abord le cas de la théorie ordinaire. Nous avons vu que dans la quatrième approximation le terme B est d'un ordre de grandeur supérieur à A_1 et A_2 . Dans les approximations successives, la formule de BORN est plus compliquée que (2). (La formule (2) est d'ailleurs déjà simplifiée du fait qu'en troisième approximation l'énergie propre n'a pas signification); elle contient au second membre des termes de types C , D , ... F , G , où G , par exemple, est une somme de termes ayant pour dénominateur une expression de la forme

$$(E_m - E_l)^2 (E_n - E_l)^2 \dots (E_r - E_l)^2 (E_s - E_l).$$

Or il y a une seule possibilité pour G , celle où l'électron repasse par le niveau d'énergie mc^2 dans tous les états intermédiaires d'ordre pair. Pour les mêmes raisons que dans la quatrième approximation, ce sont les énergies négatives dans les états intermédiaires qui conduisent à la plus grande divergence. Le terme du type G contient alors approximativement $(mc^2)^{2n-1}$ en dénominateur, et diverge plus que les termes des autres types: l'une des $\mathfrak{N}$ possibilités contient la divergence la plus forte. L'énergie propre en $2n$ ème approximation est donc de l'ordre de grandeur

$$E_\varepsilon \propto \frac{c}{137^n (mc)^{2n-1}} \int d\omega, \dots d\omega_n q_1 dq_1 \dots q_n dq_n \quad (14)$$

dans la théorie ordinaire. Comme il ne se présente pas de fonction des angles entre les $\vec{q}_i$, l'intégration sur les angles donne $(4\pi)^n$. Donc

$$E_\varepsilon \sim mc^2 \left(\frac{2\pi Q^2}{137 m^2 c^2} \right)_{Q \rightarrow \infty}^n$$

Cette expression ne conduirait à aucune divergence si elle était finie et si elle était de plus le terme d'une progression géométrique, soumis à la condition $2\pi Q^2/137 \cdot m^2 c^2 < 1$, condition qu'on peut écrire

$$\frac{Q^2}{m} < \frac{137 mc^2}{2\pi}. \quad (15)$$

Il n'est toutefois pas certain que la condition (15) suffise pour assurer la convergence de la série des approximations successives qui donnent E_ε , vu que dans une pareille condition les termes des types $C, D, \dots F$ prennent tout autant d'importance que ceux du type G . On majore alors l'expression complète de E_ε en multipliant (14) par le nombre de termes possibles $\mathfrak{N}$ ce qui donne

$$E_\varepsilon < \frac{(2n)!}{n!} mc^2 \left(\frac{\pi Q^2}{137 m^2 c^2} \right)_{Q \rightarrow \infty}^n \quad (16)$$

Pour des valeurs grandes de n on simplifie cette expression au moyen de la formule de STIRLING, qui montre que $(2n)!/n!$ est de l'ordre de $n!$; cette divergence ne peut être réduite par aucun procédé, et cela laisse peu d'espoir que l'on trouve une majorante plus petite qui donne un sens à une condition du genre de (15).

En résumé, dans la théorie ordinaire, le passage d'une approximation à la suivante introduit en dénominateur sous un signe $\int$ une parenthèse $(E_f - E_i)$ qui n'atteint pas l'ordre de grandeur de q lorsque ce paramètre tend vers l'infini, et par conséquent il introduit au total une opération que l'on peut représenter symboliquement par

$$\frac{H}{m^2 c^2} \int q dq (\dots) \quad (17)$$

où H est peut-être proportionnel à n .

Théorie des lacunes. — Considérons maintenant la théorie des lacunes. A cause de l'hypothèse selon laquelle une infinité de niveaux énergétiques sont occupés, il se peut que le volume Ω dans lequel les électrons se trouvent ne soit pas compensé au cours des calculs. Si aucune compensation ne se produisait dans la $2n$ ième approximation, cela voudrait dire que l'énergie propre

dans la théorie des lacunes diverge non seulement proportionnellement à une certaine puissance d'un paramètre Q , mais aussi proportionnellement à Ω^{n-1} . Il est clair qu'aucun procédé se rattachant à l'attribution d'un rayon à l'électron ne supprime la divergence contenue dans Ω . Il nous semble d'autre part difficilement raisonnable que ce volume apparaisse dans l'expression de l'énergie propre d'un électron. Pour ces deux raisons, et dans l'idée que dans toutes les approximations il se produit une compensation semblable à celle que nous avons reconnue en quatrième approximation, nous ne considérerons dorénavant que des termes ne pouvant pas contenir Ω en facteur. Or les parenthèses $(E_i - E_l)$, ... qui interviennent dans la formule de BORN sont dans la théorie des lacunes toujours de l'ordre des impulsions q des photons apparus dans les états intermédiaires. Les termes ne contenant pas Ω présentent des intégrations soit sur ces impulsions, soit sur les impulsions d'électrons choisis dans le vide, mais l'hypothèse faite implique que chaque électron du vide mis en jeu passe par le niveau d'énergie mc^2 dans l'un ou l'autre des états intermédiaires. Or le passage d'une approximation à l'autre signifie l'introduction d'un nouveau photon, donc un facteur $1/q$, et de plus les deux nouveaux états intermédiaires introduisent deux parenthèses $(E_r - E_l)(E_s - E_l)$, donc un facteur de l'ordre de $1/q^2$. Comme une seule intégration supplémentaire intervient, il en résulte qu'une nouvelle approximation comparée à la précédente correspond symboliquement à l'opération

$$H' \int \frac{dq}{o(q)} (\dots) \quad (18)$$

où $o(q)$ est de l'ordre de q quand ce paramètre tend vers de grandes valeurs. Si l'on compare (18) à (17) on voit que l'augmentation de la divergence due uniquement au paramètre q est beaucoup plus considérable dans (17) que dans (18), c'est-à-dire plus considérable dans la théorie ordinaire que dans la théorie des lacunes. Mais nous ne savons pas comment divergent H et H' . Le nombre des séries d'états intermédiaires dans la théorie des lacunes est supérieur à celui de la théorie ordinaire, car les n photons mis en jeu dans la $2n$ ième approximation peuvent être distribués entre n , ou $(n-1)$, ou $(n-2)$... ou 1 électron. Toutefois le nombre de ces possibilités n'est pas si considérable qu'il n'y ait aucun espoir pour que la théorie des lacunes ne présente pas une divergence très grande. En effet, le nombre de manières de distribuer n quanta soit entre n , soit entre $n-1$, ... soit à un seul électron est égal à n^n , mais il faut naturellement tenir compte du fait qu'on ne peut

pas différencier les électrons du vide mis en jeu (qui sont au nombre de $n - 1$ au plus dans l'énergie de couplage et de n au plus dans les termes interdits), et qu'une certaine distribution des n quanta doit être regardée comme équivalente à celle obtenue en permutant les électrons provenant du vide. Les termes de couplage contiennent toujours Ω ; ils n'entrent donc plus en considération. Le nombre de possibilités de distribuer les n quanta à n électrons au plus est alors inférieur à l'expression $n^n/(n-1)!$, qui vaut pour des grandes valeurs de n approximativement $\sqrt{2\pi} e^n n^{\frac{1}{2}}$ et ne contient pas de factorielle. La plus petite valeur que les parenthèses $(E_f - E_i)$ puissent prendre est de l'ordre de $2cq$. Il en résulte que l'on peut choisir pour majorante en 2^{nième} approximation l'expression

$$E_{\varepsilon_0} < \frac{2 \sqrt{2\pi} e^n n^{\frac{1}{2}} (4\pi)^n}{2^{2n-1} 137^n} c \int \frac{dq_1}{q_1} \frac{dq_2}{q_2} \dots \frac{dq_{n-1}}{q_{n-1}} dq_n.$$

L'intégration sur les valeurs des impulsions ne peut diverger pour les petites valeurs des q_i . De sorte que cela n'aurait pas de sens de prendre pour limite inférieure de ces intégrations la valeur $q = 0$. Mais comme au dénominateur $(E_f - E_i)$ ne peut en réalité jamais s'annuler, il n'y a de toute manière aucune divergence provenant des valeurs très petites des q_i . L'inégalité ci-dessus s'écrit encore

$$E_{\varepsilon_0} < 2 \sqrt{2} n! e \left(\frac{2\pi e}{137} \right)^n c \int \frac{dq_1}{q_1} \dots \frac{dq_{n-1}}{q_{n-1}} dq_n \quad (19)$$

($e = 2.71828$).

Le second membre de (19) diverge moins que celui de (16), parce que l'intégration sur les q_i dans (16) fournit une puissance croissante de Q , tandis que (19) conduit à un produit de Q par des logarithmes. Mais comme ces relations expriment des inégalités qui majorent, l'une E_ε , l'autre E_{ε_0} , on ne peut tirer évidemment aucune conclusion décisive de la comparaison de ces deux grandeurs.

Ce qui est certain, c'est que la divergence due à

$$Q \rightarrow \infty$$

se fait sentir beaucoup plus dans la théorie ordinaire que dans la théorie des lacunes. Mais à cause du grand nombre de possibilités de ranger les quanta émis dans les états intermédiaires, et à cause aussi du fait qu'on ne peut pas estimer sans les calculer les divers termes correspondant à ces possibilités, en particulier pas leur signe (le calcul de traces telles que S_1 , S_2 , $\mathfrak{S}$, donne des

valeurs positives ou négatives suivant les cas), il semble très difficile de décider si, dans le problème de l'énergie propre, cela a un sens de considérer comme inexistants les photons dont les impulsions dépassent une certaine limite, autant dans la théorie ordinaire que dans la théorie des lacunes.

6. Conclusion.

Les conclusions que l'on peut tirer des calculs qui précèdent sont de nature négative, mais elles méritent toutefois quelque attention parce qu'elles doivent mettre en garde sur quelques difficultés du problème qui n'ont à notre connaissance pas été relevées jusqu'ici.

Comme nous l'avons déjà remarqué, les tentatives de rendre finies les expressions de l'énergie propre ne se rapportent qu'à une approximation déterminée, et, en particulier dans l'esprit des mémoires sur ce sujet, à la seconde. D'autre part, dans les approximations supérieures, on sait que les divers effets que l'on calcule dans la théorie des quanta ne sont en accord avec l'expérience que lorsque les longueurs d'onde mises en jeu sont supérieures à une limite du genre de celle donnée par la condition (15). Jusqu'à la découverte de l'électron lourd, on avait coutume de dire que la théorie des quanta ne convient pas pour des domaines pareils. Il semble par contre qu'on puisse se passer de cette hypothèse désagréable¹⁾. Toutefois, dans le problème de l'énergie, la difficulté persiste, et la tentative d'une limitation des longueurs d'onde garde son intérêt du point de vue théorique.

Le procédé proposé par WENTZEL (loc. cit.) a le désavantage d'être artificiel, mais il présente toutefois l'intérêt de montrer que les diverses manières d'affectuer les passages à la limite ne sont pas univoques. Mais il semble qu'à la lumière de nos remarques concernant la suite des approximations, et même déjà parce que, autant dans la théorie des lacunes que dans la théorie ordinaire, chaque approximation contient une divergence proportionnelle à Q au moins²⁾, il nous semble douteux que le procédé de WENTZEL conduise à la solution des difficultés de l'énergie propre.

Le chemin proposé par MARCH pour tourner la difficulté implique l'existence d'un postulat nouveau dans la physique quantique: le ds ne peut être déterminé qu'à une grandeur γ près. Il

¹⁾ Voir à ce sujet l'article de J. BHABHA dans les Proc. Roy. Soc. **164**, 257, 1938.

²⁾ Nous rappelons que la divergence logarithmique trouvée par WEISSKOPF est fortuite. Voir à ce sujet V. WEISSKOPF, Zeitschr. f. Phys. **89**, p. 27, Fussnote 5.

en résulte que la fonction *delta* doit être remplacée par une fonction dont la valeur est indéterminée dans le domaine correspondant, si bien que le développement en série de FOURIER perd son sens. MARCH modifie alors le développement en y introduisant une indétermination sur le paramètre q dans le domaine en question. Cette précaution étant prise, le procédé conduit effectivement à une valeur finie de l'énergie propre en seconde approximation. Mais la difficulté concernant la suite des approximations persiste.

Des procédés qui consistent à définir comme très rare l'émission de certains quanta de lumière ont le défaut d'être arbitraires.

La conclusion que nous devons tirer de nos calculs consiste en ceci: Si tel procédé ou tel autre a l'avantage de détruire la divergence de l'énergie propres, il ne le fait en principe que dans l'une ou l'autre des approximations successives, fournissant dans chacune d'elles une contribution finie

$$E_{\varepsilon}(2), E_{\varepsilon}(4), \dots E_{\varepsilon}(2n), \dots$$

Mais nous ne savons pas si la somme

$$\sum_{n=1}^N E_{\varepsilon}(2n)$$

converge ou non lorsque $N \rightarrow \infty$. Cela est peu satisfaisant, et la méthode d'approximation de BORN que nous avons suivie ne semble pas apte à conduire à une solution. On peut précisément craindre que les difficultés qui ont surgi au cours des calculs lui soient propre, ce qui condamnerait son emploi. On doit espérer qu'en remplaçant la méthode d'approximation de BORN par une autre, on soit conduit à des résultats plus précis. C'est avec cette idée que M. T. GUSTAFSON s'est proposé d'attaquer le problème différemment; dans un travail devant paraître ailleurs¹⁾, il recherche quelle est l'influence du procédé de limitation lorsqu'on emploie la méthode exacte de solution des équations en mécanique quantique qui consiste à résoudre un déterminant séculaire.

¹⁾ Arkiv för Mat., Astr. och Fysik (Bd. 26).