

INTERACTIONS de JAUGE

des PARTICULES ELEMENTAIRES

**Vincent BOURGES, agrégé de physique,
IPEST, BP 51 - 2070 LA MARSA - TUNISIE**

L'étude des particules élémentaires est intéressante à plus d'un titre: tout d'abord parce qu'elle nous livre les clés de la connaissance ultime de la matière et de la logique de cette construction, et, à ce titre, les quarks et l'interaction faible font partie de la culture de l'"homme homme" comme les quasars ou les trous noirs. Nos connaissances sur l'infiniment petit pourraient bien être à l'origine d'une révolution des esprits aussi grande que la révolution Copernicienne.

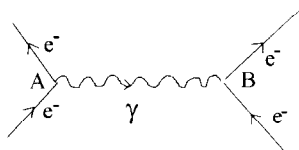
Elle nous intéresse aussi par les explications qu'elles pourraient donner à des questions jusqu'ici sans réponse: la physique des particules ne semble pas nous concerner directement, et pourtant elle nous parle de la symétrie gauche-droite, de l'origine de la masse

Les théories modernes qui essaient d'expliquer le classement et les interactions entre particules élémentaires sont des *théories de jauge*. Le cadre général de ces théories est la théorie quantique des champs, alors que, dans la mécanique quantique classique, les interactions sont décrites à l'aide de potentiels (coulombiens par exemple pour la gravitation et l'électromagnétisme) non quantifiés, la théorie des champs quantifie de façon symétrique particules et champs d'interaction dans un cadre relativiste. Le but de cet article est d'essayer de saisir la substance, la "philosophie" en quelque sorte des théories de jauge, et d'en comprendre le mécanisme.

LES INTERACTIONS

Dans le cadre de la théorie quantique des champs, les interactions entre particules ne sont pas instantanées, mais se font par l'intermédiaire d'autres particules, qui sont des particules d'échanges, appelées *bosons de jauge* ou *bosons intermédiaires*. Cet échange est décrit au moyen des diagrammes de Feynman

Par exemple, un électron émet un photon γ en A et recule (comme le recul d'un fusil).



Ce photon se déplace à la vitesse de la lumière de A à B. Il est capté en B par un autre électron qui recule également sous l'effet du choc. L'échange de photon se traduit donc par une répulsion électrostatique entre les deux charges de même signe. Le photon est le boson de jauge de l'interaction électromagnétique entre particules chargées.

Pourquoi les particules échangées sont-elles des *bosons* ? Cela vient du fait que les particules d'échange peuvent être créées ou détruites. Or, pour les particules, comme les leptons et les quarks, qui ont un spin demi-entier et sont donc des fermions, il existe une loi de conservation (du nombre leptonique ou baryonique), qui ne permet que la création et l'annihilation de paires particule-antiparticule. Le nombre de fermions (moins le nombre d'antifermions) reste donc constant, et l'on ne peut créer un fermion tout seul. Par contre le nombre de bosons (particules de spin entier) n'est pas constant, et il peut donc y avoir création ou destruction de bosons.

La portée de l'interaction (distance à laquelle se fait sentir celle-ci) dépend de la masse du boson de jauge. Si celle-ci est nulle (cas du photon), la portée est infinie et le comportement de l'interaction est le comportement géométrique en $1/r^2$.

Si la masse est non nulle (on verra que c'est le cas des bosons $W^\pm$ et Z^0 de l'interaction faible), la portée est liée à la masse m du boson de jauge. Au cours de l'interaction, il y a conservation de la quantité de mouvement, et il n'y a donc pas conservation de l'énergie. La particule échangée est une particule "virtuelle" que l'on ne peut pas observer et qui ne peut exister que pendant un temps permis par le principe d'incertitude de Heisenberg $\Delta t \Delta E \leq \hbar$ avec $\hbar = h/2\pi$, soit, avec $\Delta E = mc^2$, $\Delta t \leq \hbar/(mc^2)$. La particule, se déplaçant à une vitesse qui est toujours inférieure à celle de la lumière, ne pourra parcourir au plus que la distance $c\Delta t \leq \hbar/mc$. Cette quantité, appelée longueur d'onde de Compton, donne la portée de l'interaction.

Les bosons de jauge $W^\pm$ et Z^0 , ayant des masses de 80 et 90 GeV environ, la portée de l'interaction faible est donc très courte $\hbar/mc = \hbar c/mc^2 = 197 \text{ MeV fm}/80 \text{ GeV} = 2,46 \cdot 10^{-3} \text{ fm}$ (avec $1 \text{ fm} = 1 \text{ fermi} = 10^{-15} \text{ m}$, $\text{GeV} = 10^3 \text{ MeV} = 10^9 \text{ eV}$ et $1 \text{ eV} = 1 \text{ électron-volt}$).

La portée de l'interaction forte est également courte, de l'ordre de 1,5 fm. Cependant les bosons de jauge, qui sont au nombre de huit et constituent les *gluons*, ont une masse nulle. L'explication est ici différente et fait appel au confinement des quarks à l'intérieur des hadrons.

IDEE CENTRALE DES THEORIES DE JAUGE.

Pour essayer de comprendre ce qu'est une théorie de jauge, prenons l'exemple de l'électromagnétisme dont la théorie quantique s'appelle *électrodynamique quantique* (quantum electrodynamics Q.E.D.).

Une particule portant une charge électrique est décrite par une fonction d'onde complexe, caractérisée par son amplitude et sa phase. La fonction d'onde s'interprète par le fait que le carré du module de la fonction d'onde $|\psi|^2 = \psi^* \psi$ représente la densité de probabilité de présence de la particule en un point.

La phase de la fonction d'onde est inobservable en ce sens que celle-ci peut être tournée d'un angle constant partout dans l'espace sans modifier le comportement de la particule. C'est l'invariance globale de jauge.

Il semblerait naturel de pouvoir modifier la phase différemment aux différents points de l'espace, c'est alors l'invariance locale de jauge. Mais, pour cela, il faut *compenser* la variation locale de phase par un champ compensateur, appelé *champ de jauge*. A ce champ sont associées des particules, les *bosons de jauge*.

Imaginons que le chef du village veuille un village tout blanc, mais les habitants indisciplinés ont acheté des peintures de couleurs différentes. Le chef va envoyer des peintres, avec la couleur complémentaire de chacun des pots, afin qu'une fois le mélange effectué, tous les pots soient blancs. Les peintres sont les bosons de jauge, chargés de compenser les variations de couleur. L'envoi des bosons d'une particule à l'autre, constitue l'interaction.

L'interaction forte s'exerce entre quarks: les quarks possèdent une propriété appelée *couleur*; un quark donné peut se trouver dans une des trois couleurs R(red = rouge), G (green = vert), B (blue = bleu). Il ne s'agit en aucun cas des couleurs habituelles, mais seulement de différencier 3 façons d'être du quark. Un proton, composé de 2 quarks u et d'un quark d, doit être "sans couleur" (ou "blanc") en ce sens qu'il doit se composer d'un quark R, d'un quark G et d'un quark B. Si l'un des quarks change de couleur et passe par exemple de R à G, il faut que le quark G passe à R, de façon que le proton reste non coloré. Cette compensation de couleur s'effectue grâce à un champ de jauge, composé en fait de 8 sortes de gluons, messagers des changements possibles de couleur. Par exemple, le quark qui passe de R à G envoie un gluon portant à la fois les couleurs R et antiverte $\bar{G}$.

On a alors $R \rightarrow G + (\text{gluon noté } R\bar{G})$ et le gluon $R\bar{G}$ arrivant sur G le transforme en R suivant $R\bar{G} + G \rightarrow R$ en détruisant la couleur G et en mettant à la place R. Le proton, lui, reste non coloré grâce à l'interaction forte de couleur, véhiculée par les gluons, champs compensateurs des changements de couleur.

Une particule possède une double structure: spacio-temporelle et interne, comme la couleur pour les quarks. Si l'on veut que la structure interne de la particule soit identique en tout point et à tout instant, il faut introduire des champs compensateurs, qui sont les champs de jauge. Tel est le message central des théories de jauge.

Le prototype des théories de jauge est l'électrodynamique quantique (Q.E.D.). La théorie de l'interaction forte entre quarks est la *chromodynamique quantique* (Q.C.D.).

Carrés 1, 2 et 3, pages 20, 22, 24

L'INTERACTION FORTE ET LE CONFINEMENT

L'interaction forte est une interaction qui s'exerce entre quarks; elle est due à un échange de couleur, comme expliqué précédemment. Les quarks n'existent pas comme particules libres; on ne les trouve que par paquets de trois quarks qqq (ou trois antiquarks $\bar{q}\bar{q}\bar{q}$) formant les particules appelées baryons (ou antibaryons), ou par paquets d'un quark et d'un antiquark $q\bar{q}$ formant les mésons. Ceci est dû à une propriété, appelée *confinement* des quarks: les quarks sont confinés à l'intérieur des hadrons (baryons et mésons) et doivent former un ensemble non coloré. Or ceci est possible:

- soit par association des trois couleurs "BRG" (en fait la fonction d'onde antisymétrique bâtie sur RBG qui ne change pas dans une rotation des couleurs), on a alors un baryon. Par exemple, le proton uud, composé de deux quarks u et d'un quark d, a ses trois quarks de couleurs différentes.
- soit par association de trois anticouleurs " $\bar{R}\bar{B}\bar{G}$ ". On obtient un antibaryon (comme l'antiproton $\bar{p}$ ou l'antineutron $\bar{n}$).

- soit par association des couleurs avec les anticouleurs correspondantes sous la forme " $R\bar{R} + G\bar{G} + B\bar{B}$ " qui est une combinaison invariante par changement de couleur. C'est le cas des mésons.

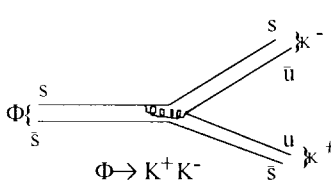
L'interaction forte étant due à un échange de couleur, on retrouve au cours de l'interaction les mêmes quarks avant et après, plus (ou moins) éventuellement des paires quark-antiquark de même saveur ($u\bar{u}, d\bar{d}$ etc...)

Exemples :

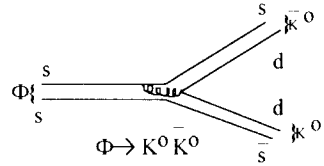
- $K^- + p \rightarrow \pi^+ + \pi^- + \Lambda$

$s\bar{u} \quad u\bar{u}d \quad u\bar{d} \quad \bar{u}d \quad uds$ Un méson K^- réagit avec un proton pour donner deux pions π^+ et π^- et une particule étrange Λ . On retrouve après l'interaction les mêmes quarks plus une paire $d\bar{d}$.

- $\Phi \rightarrow K^+ + K^-$ (noté $\Phi \rightarrow K^+K^-$) et $\Phi \rightarrow K^0\bar{K}^0$

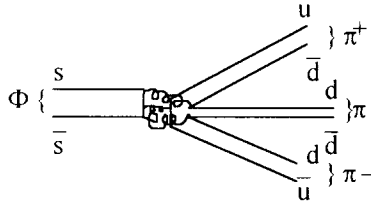


s émet un gluon qui se transforme en une paire $u\bar{u}$



création d'une paire $d\bar{d}$

- $\Phi \rightarrow \pi^+ \pi^- \pi^0$

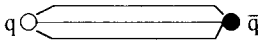


Les deux premiers diagrammes ($\Phi \rightarrow K^+K^-$ et $\Phi \rightarrow K^0\bar{K}^0$) interviennent dans 84% des cas alors que le dernier n'intervient que dans 15% des cas. Ceci est dû aux lignes de quark déconnectées et à l'émission de trois gluons qui est un phénomène moins fréquent que l'émission d'un seul gluon (règle de Zweig)

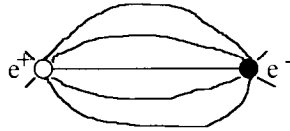
Quelle est l'origine physique du confinement ? Dans le cadre de la théorie quantique des interactions fortes (chromodynamique quantique ou quantum chromodynamics QCD), le confinement n'a pas encore été démontré, mais les physiciens pensent tout de même avoir une idée de l'explication du phénomène.

De même qu'il existe en électromagnétisme des champs électrique et magnétique, on peut définir dans QCD un champ électrique de couleur et un champ magnétique de couleur, et il existe entre un quark et un antiquark une force électrique de couleur. Mais à la différence des lignes de champ

électrique entre particules chargées, les lignes de champ électrique de couleur se concentrent dans un tube mince, car les gluons interagissent entre eux.

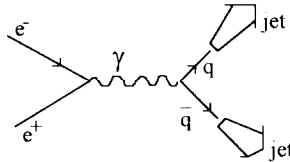


le champ de couleur qq avec
 $V(r)$ proportionnel à r

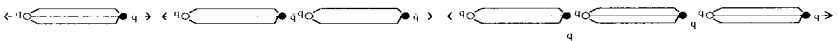


champ coulombien avec $V(r)$
proportionnel à $1/r$

la force entre q et $\bar{q}$ est pratiquement constante (les lignes de champ sont parallèles) et le potentiel (sauf à courte distance) est proportionnel à r , tandis que le potentiel coulombien est en $1/r$, la force étant en $1/r^2$.



Envoyons l'un contre l'autre un faisceau d'électrons et un faisceau de positrons. Les faisceaux s'annihilent, et, par l'intermédiaire de photons virtuels, peuvent se transformer en paires quark-antiquark qui partent avec des sens opposés. Le tube des lignes de couleur entre q et $\bar{q}$ s'allonge, l'énergie potentielle augmente, et, lorsqu'elle est suffisante, le tube se brise en donnant naissance à deux jets de particules se dirigeant en sens opposés.



La force qui lie les nucléons entre eux dans le noyau, qui fut la première motivation de l'interaction forte n'est en fait qu'une manifestation résiduelle des interactions entre les quarks des deux nucléons, de la même façon que la force de Van Der Waals entre molécules n'est qu'une manifestation des forces électriques entre particules chargées qui composent les molécules neutres.

Effet d'écran et liberté asymptotique

Parlons des constantes de couplage, et d'abord de la constante de couplage électromagnétique. La

loi de Coulomb s'écrit, $F = \frac{qq'}{4\pi\epsilon_0 r^2}$; il s'ensuit que $\frac{qq'}{4\pi\epsilon_0}$ a la dimension d'une force

multipliée par le carré d'une distance, qui est aussi celle du produit d'un moment cinétique par une vitesse. Mais deux constantes sont inscrites dans la structure du monde: la constante de Planck $\hbar$, qui caractérise le domaine de la mécanique quantique et la vitesse de la lumière c , qui caractérise les effets relativistes. On peut donc caractériser la force de l'interaction par un nombre sans

dimension $\frac{qq'}{4\pi\epsilon_0\hbar c}$, et, puisque la charge est quantifiée, $q = Qe$ et $q' = Q'e$, Q et Q' étant les "charges" sans dimension, quotient de q et q' par la charge du proton; on a alors $\frac{qq'}{4\pi\epsilon_0\hbar c} = \frac{QQ'e^2}{4\pi\epsilon_0\hbar c}$ et la force de l'interaction est donnée par la constante de couplage $\alpha = \frac{e^2}{4\pi\epsilon_0\hbar c}$ de valeur numérique $\alpha = 1/137$ (appelée encore constante de structure fine).

Dans les unités naturelles, où l'on pose $\hbar = c = 1$, on a $\alpha = \frac{e^2}{4\pi}$ (avec $\frac{1}{\epsilon_0} = 1$).

En fait, les choses sont un peu plus compliquées. Le vide de la mécanique quantique est rempli de paires particule-antiparticule virtuelles, qui ne peuvent exister réellement parce qu'il n'y a pas dans le vide une énergie suffisante pour les créer, mais qui peuvent exister comme particules virtuelles pendant un temps court donné par le principe d'incertitude.

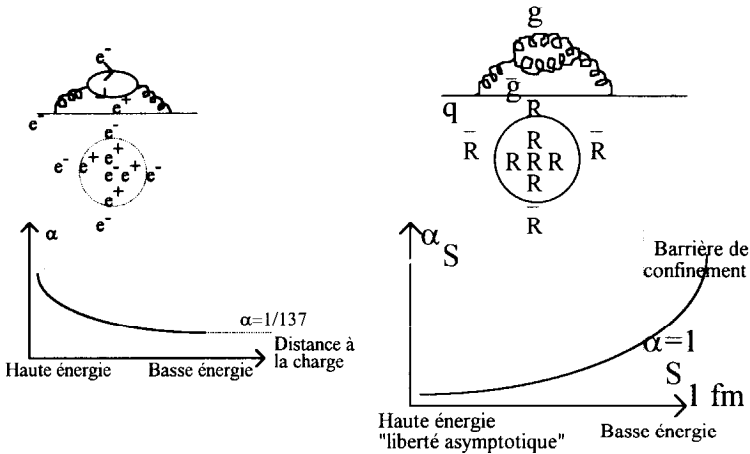
Si la particule réelle est introduite dans ce vide compliqué, elle va polariser les paires virtuelles électron-positron, et, si la charge introduite est positive, les électrons des paires virtuelles vont être attirés, tandis que les positrons seront repoussés. Le résultat est qu'une sphère centrée sur la particule positive contient une charge légèrement inférieure à celle de la charge introduite. C'est l'*effet d'écran*, et la charge mesurée expérimentalement sera un peu inférieure à la charge de la particule introduite. Suivant les calculs de QED, la grandeur de la charge e , et donc la constante de couplage α , vont décroître avec la distance, et la constante mesurée $\alpha = 1/137$ est en fait mesurée loin de la charge. α augmente lorsque la distance diminue (et l'énergie augmente car il faut une énergie de plus en plus forte pour s'approcher de plus en plus près), et α tendrait vers l'infini si l'on pouvait s'approcher infiniment près de la charge. En fait, pour les énergies que l'on sait faire actuellement, α reste proche de $1/137$.

La théorie de l'interaction forte QCD par interaction de couleur introduit de la même façon le couplage g_S (à la place de e) et la constante de couplage $\alpha_S = \frac{g_S^2}{4\pi}$ (dans les unités naturelles) (S est mis pour strong, fort).

La création de paires virtuelles quark-antiquark aurait tendance à faire diminuer α_S avec la distance, mais on peut avoir en plus création de paires de gluons, dont l'effet est inverse et plus important: une particule colorée s'entoure préférentiellement de particules de même couleur. Ce qui fait que l'on a globalement un *effet d'anti-écran* et α_S diminue lorsqu'on s'approche de la particule colorée.

L'interaction de couleur devient faible à courte distance (et forte énergie), de sorte que deux quarks très proches sont presque libres. C'est la propriété de "*liberté asymptotique*".

Résumons sur un diagramme:



Comme les calculs que l'on sait faire dans QED et QCD sont basés sur la théorie des perturbations et ne sont donc valables que si α est petit, on sait donc faire des calculs précis dans QED car α est petit et seulement pour les fortes énergies dans QCD car α_S devient petit. A basse énergie, on ne sait pas faire de calculs précis avec QCD, car α_S devient trop grand. Or, même avec les plus grands accélérateurs, on ne sait faire encore que des énergies qui restent faibles pour les calculs dans QCD. QCD est donc une théorie qualitativement bonne, mais qui manque encore de vérifications quantitatives précises.

La théorie de jauge de l'interaction forte est basée sur le groupe $SU(3)$ de couleur, groupe des transformations agissant sur les trois couleurs.

voir carré 4, pages 26

Elle se base sur le fait que le lagrangien des quarks doit être invariant par transformation locale de jauge de $SU(3)$.

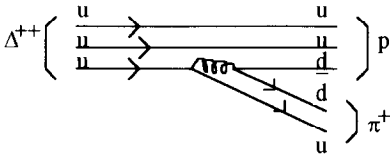
voir carré 5, page 29 et carré 6, pages 31, pour la signification géométrique des transformations de jauge.

INTERACTION FAIBLE

Présentation

La plupart des particules (sauf, entre autres, le proton, l'électron et le neutrino) sont instables et se désintègrent plus ou moins rapidement.

Prenons le cas de Δ^{++} , qui se désintègre en 10^{-23} s suivant la réaction $\Delta^{++} \rightarrow \pi^+ p$



Il s'agit d'une désintégration par interaction forte. Un quark u émet un gluon qui se transforme en une paire $d\bar{d}$.

La durée de vie moyenne est à peu près le temps que mettent π^+ et p pour se séparer d'une longueur égale à la portée de

l'interaction, de l'ordre de 1 fm: $\tau = R/c = 10^{-15} / 3 \cdot 10^8 \approx 10^{-23} \text{ s}$.

Si l'on considère la particule π^0 , elle ne peut se désintégrer par désintégration forte, car c'est le méson le plus léger. Elle se désintègre par interaction électromagnétique par le processus $\pi^0 \rightarrow \gamma\gamma$ en 10^{-16} s . Le fait que la désintégration soit plus lente s'explique par le fait que l'interaction électromagnétique est plus faible que l'interaction forte. On a $\tau_{\text{em}}/\tau_{\text{S}} \approx (\alpha_S/\alpha)^2 \approx 10^4$ à 10^6 , la constante α_S de l'interaction forte étant de l'ordre de 1 à une distance de 1 fm, pendant que $\alpha = 1/137$.

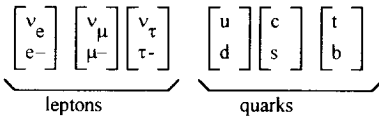
Que penser alors de la désintégration du neutron $n \rightarrow p + e^- + \bar{\nu}_e$, dont la durée de vie moyenne est de 15 minutes ? Ce n'est ni une désintégration forte, ni une désintégration électromagnétique. C'est donc qu'il existe une troisième sorte d'interaction, qui doit être très faible puisque la vie moyenne du neutron est élevée: c'est l'interaction faible, qui, le plus souvent, change la nature des particules qui interagissent, ce qui la différencie complètement des autres interactions. L'interaction faible apparaît davantage comme une réaction que comme une interaction.

On explique les interactions faibles par l'échange de bosons de jauge, qui, pour les interactions faibles, sont les bosons W^+ , W^- et Z^0 . Mais l'interaction faible ayant une très courte portée (de l'ordre de 10^{-18} m), il s'ensuit que les bosons faibles $W^\pm$ et Z^0 doivent être massifs, à la différence du photon et des gluons.

Les bosons $W^\pm$ sont les bosons de jauge des interactions à courant chargé (avec échange de charge) et Z^0 est le boson de jauge des interactions à courant neutre (sans échange de charge).

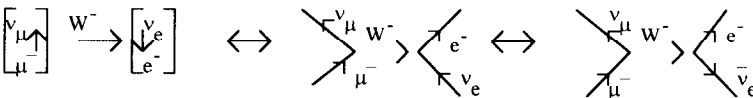
Mécanisme

Les leptons et les quarks se regroupent en doublets



L'interaction faible agit entre deux doublets en changeant la particule de chacun des doublets par l'autre particule du même doublet.

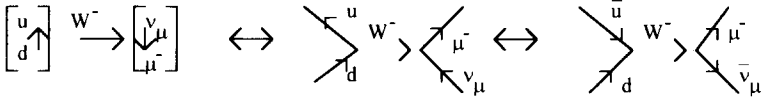
Prenons comme exemple la désintégration du muon μ^- : $\mu^- \rightarrow \nu_\mu e^- \bar{\nu}_e$ qui s'explique de la façon suivante:



Le dernier diagramme est obtenu en remplaçant ν_e entrant par $\bar{\nu}_e$ sortant.

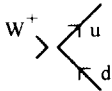
Il y a bien échange entre, d'une part μ^- et ν_μ , d'autre part entre ν_e et e^- , la charge de μ^- ayant été transférée, grâce à W^- , à e^- . La durée de vie moyenne du muon est de $2,2 \cdot 10^{-6}$ s.

Prenons comme autre exemple la désintégration de π^- : $\pi^- \rightarrow \mu^- \bar{\nu}_\mu$

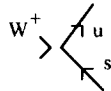


soit $\bar{u}d \rightarrow \mu^- \bar{\nu}_\mu$ ou $\pi^- \rightarrow \mu^- \bar{\nu}_\mu$. Sa durée de vie moyenne est de $2,6 \cdot 10^{-8}$ s.

Notons que, pour les quarks, on a parfois passage de u à s ou de c à d. On explique ces passages en remplaçant d et s dans les doublets par d' et s' qui sont des combinaisons linéaires de d et s, $d' = d \cos \theta_C + s \sin \theta_C$ et $s' = -d \sin \theta_C + s \cos \theta_C$ θ_C étant l'angle de Cabibbo, de valeur numérique proche de 13° . On peut donc observer les deux diagrammes :

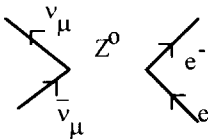


Transition "favorisée" de Cabibbo



Transition "supprimée" de Cabibbo

En 1973, on a découvert l'existence d'interactions avec $\bar{\nu}_\mu$ sans lepton chargé à l'état final, du type $\bar{\nu}_\mu e^- \rightarrow \bar{\nu}_\mu e^-$



Ce sont des interactions à courant neutre dont le boson de jauge est Z^0 . Z^0 ne change ni la charge, ni la particule.

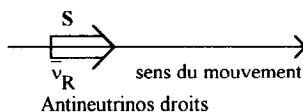
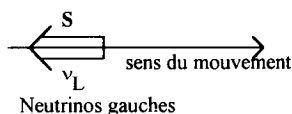
Théorie de jauge de l'interaction électrofaible

Expliquons d'abord ce que l'on entend par particule gauche et droite.

La projection du spin sur la direction du mouvement peut être dirigée dans le même sens, on dit alors que la particule est "droite" (R), ou en sens contraire, et la particule est "gauche" (L).

Une particule massive gauche peut être transformée en particule droite en l'arrêtant, et en l'accélérant en sens opposé sans changer son spin. Les particules massives peuvent donc être gauches ou droites. Cependant les neutrinos, ayant une masse nulle (ou très faible) voyagent à la vitesse de la lumière et restent gauches ou droites.

On constate expérimentalement que seuls les neutrinos gauches et les antineutrinos droits existent.



Cette propriété est importante car les interactions faibles sont associées à une symétrie $SU(2)_L$ qui agit différemment sur les composantes gauche et droite des diverses particules. Les composantes gauches font partie de doublets et les composantes droites de singlets (qui ne changent pas par transformation de $SU(2)_L$).

Si l'on se limite à la première génération de quarks et de leptons (ν_e , e^- , u et d), les singlets sont e^-_R , u_R , d_R et les antiparticules e^+_L , $\bar{d}_L$ et $\bar{u}_L$ et les doublets

$$\begin{bmatrix} \nu_{eL} \\ e^-_L \end{bmatrix}, \begin{bmatrix} u_L \\ d_L \end{bmatrix} \quad \text{et} \quad \begin{bmatrix} e^+_R \\ \bar{\nu}_R \end{bmatrix}, \begin{bmatrix} \bar{d}_R \\ \bar{u}_R \end{bmatrix}$$

La particule en haut d'un doublet a une charge électrique égale à celle de la particule du bas augmentée de une unité.

L'interaction faible est décrite par une théorie de jauge basée sur le groupe $SU(2)_L \times U(1)_Y$ des deux symétries $SU(2)_L$ et $U(1)_Y$ indépendantes l'une de l'autre, où Y est l'hypercharge faible.

voir pour plus de détails carré 7, pages 33

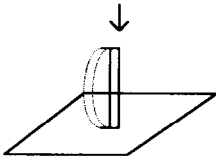
Mais cette description est insuffisante sur plusieurs points: d'abord l'interaction faible est à courte portée; il faut donc que les bosons $W^\pm$ et Z^0 soient des bosons massifs; or l'invariance de jauge est brisée si l'on introduit des termes de masse pour les bosons de jauge. De plus, si la symétrie $SU(2)$ se comportait comme la symétrie $SU(3)$ des interactions de couleur, il serait impossible de distinguer neutrino et électron gauche, qui formeraient une seule particule. Or elles ont, en fait, des propriétés très différentes. Pour expliquer ces faits, nous allons voir qu'on fait appel au phénomène de brisure spontanée de symétrie.

BRISURE SPONTANEE DE SYMETRIE

Présentation et exemples

Comment les théories de jauge peuvent-elles engendrer des forces à courte portée ? Ce problème a arrêté l'essor des théories de jauge pendant une dizaine d'années et a trouvé une

solution dans la brisure spontanée de symétrie: aux hautes énergies la symétrie $SU(2)$ est exacte, on ne peut différencier le neutrino et l'électron et entre les deux existe une force analogue à la force de couleur. Mais, aux environs de 100 GeV (correspondant à une température de $10^{15} K$ d'après $E = kT$, k étant la constante de Boltzman, égale à $1,38 \cdot 10^{-23} S.I.$), il y a un changement de phase, la symétrie se brise en ce sens que la loi de force reste symétrique, mais que le vide (qui est en mécanique quantique l'état de plus basse énergie) ne l'est plus.



Prenons quelques exemples: si l'on appuie fortement sur une aiguille initialement verticale, celle-ci finit par se tordre dans un plan déterminé. La symétrie initiale a été brisée par le choix d'un plan particulier pour la position de l'aiguille.

Autre exemple cité par Weinberg: des convives sont assis autour d'une table ronde; ils ont une serviette à leur droite et une à leur gauche; il suffit que l'un des convives touche une serviette, celle qui est à sa droite, par exemple, pour que tous les convives soient obligés de prendre la serviette à leur droite. La symétrie gauche-droite a été rompue par le choix particulier de l'un des convives. C'est une brisure de symétrie.

Prenons encore un exemple: dans un liquide, il y a une symétrie spatiale de rotation; l'aspect reste le même si l'on fait subir une rotation au liquide et les lois qui régissent son comportement sont invariantes par rotation. Lorsque le liquide cristallise, les lois décrivant le comportement du solide restent invariantes par rotation, mais l'état fondamental du solide n'est plus symétrique pour toutes les rotations car il existe des axes privilégiés. On dit que la symétrie est brisée spontanément car le choix de l'état fondamental a brisé la symétrie, bien que les lois restent symétriques.

Une particule qui vivrait à l'intérieur du cristal, sans pouvoir le fondre, verrait la structure cristalline comme une propriété fondamentale de l'espace et il établirait, par exemple une liste des vitesses de la lumière suivant les différents axes cristallographiques, qui serait pour lui des constantes fondamentales.

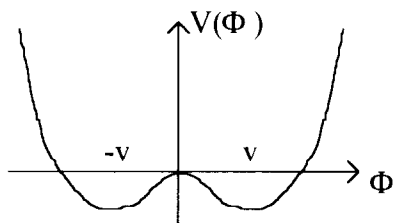
Notre espace physique est isotrope, donc symétrique pour les rotations. Mais l'espace des propriétés internes des particules est dans une phase à symétrie brisée, et nous attribuons, comme la particule dans le solide, des propriétés différentes aux interactions faible et électromagnétique, bien qu'il y ait une symétrie sous-jacente qui se retrouve dans les propriétés; cependant, au dessus de $10^{15} K$, température de la transition, il n'y a qu'une seule interaction électrofaible.

Un dernier exemple bien connu est le *ferromagnétisme*: au dessus de la température de Curie, tous les dipôles magnétiques sont orientés au hasard et il n'y a pas d'aimantation globale; le système possède la symétrie sphérique $SO(3)$ (il apparaît identique dans toutes les directions). Au dessous de la température de Curie, il y a aimantation spontanée et tous les dipôles s'alignent dans une direction donnée arbitraire (dépendant d'un défaut ou d'une fluctuation). Parmi tous les vides possibles (toutes les orientations possibles), la nature a choisi un vide particulier (une direction particulière): la symétrie sphérique a été spontanément brisée et la symétrie restante est axiale (autour de l'axe d'aimantation), bien que les lois restent à symétrie sphérique.

Mécanisme de la brisure de symétrie

Pour expliquer le mécanisme, on va faire une analogie avec un plasma. Soit donc un plasma constitué de charges négatives libres circulant dans un réseau de charges positives fixes de sorte qu'à l'équilibre, la densité moyenne de charge soit nulle. Supposons que, localement, il y ait un excès de charges négatives, celles-ci vont se repousser, la force étant proportionnelle au déplacement à partir de la position d'équilibre, comme pour un oscillateur harmonique. D'où des oscillations de plasma de pulsation ω dépendant des caractéristiques du système. Sous l'effet d'une onde électromagnétique, le plasma va subir des vibrations forcées dont l'amplitude sera maximale à la résonance. En fait, le plasma va jouer le rôle d'un filtre; il atténue fortement les ondes de pulsation inférieures ou supérieures à ω et ne laisse en fait passer que les ondes de pulsation ω . Lors d'une interaction électromagnétique à l'intérieur du plasma, seuls les photons virtuels de pulsation ω vont intervenir. Leur énergie est $\hbar\omega$ correspondant à une masse m telle que $\hbar\omega = mc^2$ et, dans ce modèle, les photons acquièrent une masse et les forces électriques une portée finie, puisque les photons de faible énergie, et donc de longue portée ont été éliminés.

Essayons maintenant d'utiliser ce modèle comme inspiration pour un mécanisme qui donnerait une faible portée aux interactions faibles et donc une masse aux particules $W^\pm$ et Z^0 . Comme analogue des charges libres dans le plasma, on postule l'existence d'un nouveau type de champ (différent des leptons, quarks et bosons de jauge) et on suppose que, dans le vide, ce champ a une valeur non nulle



Ce champ, appelé champ de Higgs, interagit avec lui-même et avec les particules; il possède une énergie potentielle dont la forme est donnée sur la figure ci-contre. A basse énergie, le champ de Higgs prend la valeur qui lui donne la plus faible énergie, soit $\Phi = \pm V$. La nature doit faire ici un choix entre $+V$ et $-V$, et la symétrie est brisée par ce choix.

A haute énergie, les fluctuations du champ dépassent le maximum obtenu à $\Phi = 0$, tout souvenir de la valeur initiale est perdu, et le champ oscille autour d'une valeur d'équilibre nulle; le système est dans la phase symétrique.

Dans la phase à symétrie brisée, lorsque Φ oscille autour de V (ou $-V$), le système se comporte au voisinage du minimum comme un oscillateur harmonique (on voit ceci d'après la forme de la courbe, ou, tout simplement parce que tout système, au voisinage de sa position d'équilibre, se comporte comme un oscillateur harmonique); le milieu (vide rempli de champ de Higgs) va, comme dans le plasma, filtrer les fréquences et donc donner une masse aux bosons de jauge $W^\pm$ et Z^0 . Le photon, lui, est lié à une symétrie exacte et reste sans masse.

Pour plus de détails, voir carré 8, pages 35 pour le phénomène de brisure de symétrie

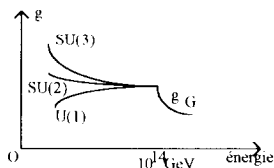
et carré 9, page 37, pour son application à l'interaction électrofaible.

UNIFICATION DES INTERACTIONS

Les interactions entre particules élémentaires sont correctement décrites jusqu'à des distances aussi petites que 10^{-18} m par une théorie de jauge $SU(3)_C \times SU(2)_L \times U(1)_Y$. C'est la théorie standard de la physique des particules : la chromodynamique quantique est la théorie de l'interaction forte et le modèle GSW (Glashow, Salam et Weinberg) donne l'explication des interactions faible et électromagnétique.

Cependant, le groupe $SU(3)_C \times SU(2)_L \times U(1)_Y$ est un produit de trois groupes indépendants de transformation de jauge : le groupe $SU(3)_C$ de couplage g_s , le groupe $SU(2)_L$ de couplage g et le groupe $U(1)_Y$ de couplage g' . Pour pouvoir relier ces trois couplages, il faut plonger $SU(3)_C \times SU(2)_L \times U(1)_Y$ dans un groupe plus étendu G , les trois couplages seront alors reliés au seul couplage g_G du groupe G .

Or l'expérience pousse dans cette voie.



Si l'on trace l'évolution en énergie des trois couplages, on constate que ceux-ci se rapprochent lorsque l'énergie augmente, et qu'ils semblent tendre vers la même valeur pour une énergie énorme de l'ordre de 10^{14} GeV. Au dessus de cette valeur, il y aurait donc unification des trois interactions et séparation de celles-ci pour les énergies inférieures.

La première et la plus simple des théories de grande unification (GUT) est basée sur le groupe $SU(5)$ (Géorgi et Glashow 1974) et, une fois ce groupe choisi, particules et interactions se mettent naturellement en place.

Il faut placer les particules dans des représentations de $SU(5)$ de façon à retrouver leurs nombres quantiques, et, pour obtenir le contenu $SU(3) \times SU(2)$ d'une représentation, on identifie les trois premiers indices avec les indices de couleur et les deux restants avec les indices $SU(2)_L$.

Le groupe $SU(5)$ s'applique aux spineurs à deux composantes décrivant des fermions sans masse soit gauches (L), soit droits (R). On montre que les représentations 5^* et 10 s'appliquent aux composantes L et 5 et 10^* aux composantes R .

La représentation $5^* = (3^*, 1) + (1, 2)$ décrit une antiparticule colorée, donc un antiquark gauche et un doublet de leptons gauches.

La représentation $10 = (3^*, 1) + (3, 2) + (1, 1)$ décrit un antiquark, un doublet de quarks et un lepton gauches.

Il s'agit d'intégrer dans ce schéma la première génération de quarks et de leptons, soit u , d , e^- et ν_e ; u et d ayant trois couleurs, cela fait pour les composantes gauches des particules et des antiparticules 16 états, d'où l'on retire ν_{eL} qui n'existe pas, d'où 15 états qui peuvent se placer dans la représentation $5^* + 10$.

Au fur et à mesure que l'énergie baisse, la nature brise les symétries et devient de moins en moins symétrique. Mais cette perte de symétrie est compensée par un gain en complexité. Les particules vont se diversifier, d'abord en quarks et en leptons, puis en trois sortes de quarks et leptons ; les interactions se différencient également et leur comportement est de plus en plus complexe. Enfin la nature choisit, à basse énergie, de relier entre elles les composantes gauche et droite des particules ayant même charge afin d'en faire une seule particule et de lui donner une masse.

La nature accomplit ce programme de la façon suivante: le groupe $SU(5)$ a 24 générateurs, dont 4 diagonaux: ce sont les valeurs propres de ces 4 générateurs qui donnent les nombres quantiques des particules et permettent donc de les distinguer.

En particulier l'hypercharge faible, Y , introduite arbitrairement dans la théorie GSW pour retrouver la charge, apparaît naturellement ici, comme égale, à un coefficient près, aux valeurs propres de l'un des opérateurs diagonaux. En choisissant le coefficient pour retrouver la charge nulle du neutrino, les particules se rangent alors naturellement avec les bons nombres quantiques (isospin faible, hypercharge faible et charge) dans les deux représentations

$$5^* = (\bar{d}^1, \bar{d}^2, \bar{d}^3, e^-, \nu_e)_L$$

$$10 = \begin{pmatrix} 0 & \bar{u}^3 & \bar{u}^2 & u_1 & d_1 \\ -\bar{u}^3 & 0 & \bar{u}^1 & u_2 & d_2 \\ \bar{u}^2 & -\bar{u}^1 & 0 & u_3 & d_3 \\ -u_1 & -u_2 & -u_3 & 0 & e^+ \\ -d_1 & -d_2 & -d_3 & -e^+ & 0 \end{pmatrix}$$

et

et l'on a toutes les composantes gauches de $u, d, \bar{u}, \bar{d}, e^-, e^+$ et ν_e ($\bar{\nu}_{eL}$ n'existe pas).

La valeur du coefficient calculé précédemment est relié à l'angle de Weinberg θ_W de l'interaction électrofaible: on trouve $\sin^2 \theta_W = 3/8$. C'est une prédiction de la théorie $SU(5)$, valable à l'énergie d'unification. Ramenée aux énergies habituelles par les équations du groupe de renormalisation, on trouve une valeur de $\sin^2 \theta_W$ voisine de 0,21, en bon accord avec la valeur expérimentale.

Pour plus de détails, voir carré 10, page 39.

Voyons maintenant le scénario prédit lorsque l'énergie baisse à partir de l'énergie d'unification où la symétrie $SU(5)$ est observée.

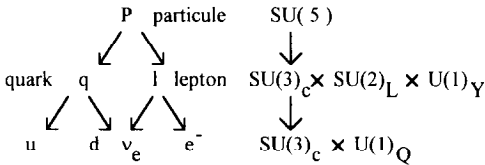
Il n'y a au départ qu'une seule sorte de particule (par génération), porteur d'une propriété interne (analogue à la couleur pour les quarks) pouvant prendre 5 valeurs. Par échange de cette propriété, on a une interaction entre les particules, de sorte que la propriété n'est pas perçue de l'extérieur, comme la couleur pour les quarks. A noter que la nature a déjà choisi entre gauche et droite en regroupant les composantes gauches des particules en doublets, les composantes droites restant sous forme de singulets, et il existe donc une théorie antérieure à découvrir.

Pour une énergie de l'ordre de 10^{14} GeV correspondant à une température fantastique de 10^{27} K, on assiste à la brisure de symétrie $SU(5) \rightarrow SU(3)_C \times SU(2)_L \times U(1)_Y$, la particule unique se sépare en un quark et un lepton. Lors de la brisure de symétrie, les bosons de jauge X et Y ont acquis une masse de l'ordre de l'énergie d'unification; on peut avoir passage de quark en lepton, mais avec une probabilité de plus en plus faible au fur et à mesure que l'énergie diminue. Le proton, dans cette théorie, peut se désintégrer, mais avec une durée de vie très grande ($\approx 10^{30}$ ans, à comparer à l'âge de l'univers $\approx 10^{10}$ ans). Vers 100 GeV (10^{15} K), nouvelle brisure de symétrie $SU(3)_C \times SU(2)_L \times U(1)_Y \rightarrow SU(3)_C \times U(1)_Q$.

C'est la brisure de symétrie électrofaible. Le quark devient un doublet $\begin{bmatrix} u \\ d \end{bmatrix}$ et le lepton également $\begin{bmatrix} \nu_e \\ e^- \end{bmatrix}$.

L'interaction électrofaible, qui permet le passage d'un membre du doublet à l'autre, devient faible à cause de la grande masse des bosons de jauge $W^\pm$ et Z^0 . Il ne reste finalement que la symétrie de couleur (interaction forte pour les quarks) et la symétrie $U(1)_Q$ (interaction électromagnétisme). En choisissant $U(1)_Q$ comme symétrie restante, la nature a choisi de regrouper composantes gauche et droite de même charge en une seule particule pouvant posséder une masse grâce au couplage avec les champs de Higgs. L'origine de la masse des fermions se trouve donc peut-être dans ce regroupement.

Résumons sur un schéma:



On s'est restreint ici à une seule génération, mais on sait qu'il existe, en fait, trois générations.

Les GUT ne nous apprennent rien de l'origine de ces trois générations de même qu'elle ne permet pas de prédire la masse des fermions.

La GUT basée sur SU(5) a tout de même de nombreux succès à son actif:

- elle prédit les nombres quantiques des particules et explique la quantification de la charge électrique (charge électrique fractionnaire des quarks due aux trois couleurs).
- elle introduit de façon naturelle l'hypercharge faible.
- elle donne l'évolution en énergie des trois couplages g_s , g et g' à partir d'une valeur unique à l'énergie de grande unification et la valeur de l'angle de Weinberg de la théorie électrofaible.
- elle prédit enfin la décomposition du proton, non encore observée à ce jour, qui constitue un test pour la théorie.

Enfin, la théorie inflatoire du début de l'univers, impliquant une expansion accélérée de l'univers entre 10^{-35} et 10^{-32} s et qui explique un certain nombre de faits, repose sur la théorie d'unification que nous avons tenté d'esquisser ici.

Carré 1 : classification des particules

Le proton et le neutron sont les deux membres les plus légers de la famille des hadrons: les hadrons sont les particules qui subissent les interactions fortes, et également les interactions électromagnétiques et faibles.

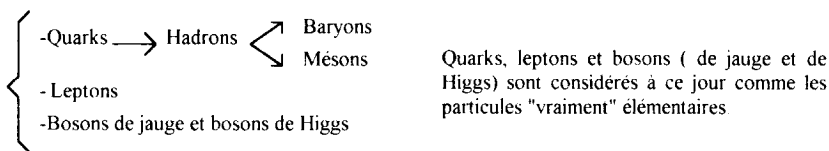
Ces hadrons sont de deux sortes: ceux qui ont un spin demi-entier et sont donc des fermions: ce sont les baryons ($p, n, \Lambda, \Xi, \Omega, \Delta \dots$) et ceux qui ont un spin entier et sont donc des bosons: ce sont les mésons ($\pi, K, \eta, D \dots$).

De leur côté, électrons et neutrinos ne subissent pas l'interaction forte, ce sont deux membres de la famille des leptons. Ils participent aux interactions faibles et électromagnétiques.

En 1964, Gellman et Zweig proposent le modèle des quarks: les hadrons (baryons et mésons) sont composés de quarks (et d'antiquarks).

Enfin les interactions sont véhiculées par les bosons de jauge, et il existe des bosons de Higgs, comme nous le verrons plus loin.

D'où une première classification:



Les quarks portent des charges électriques fractionnaires $+2e/3$ ou $-2e/3$.

Ils se distinguent par leur savueur et s'appellent u, d, s, c, b et t.

u et d ont presque même masse et sont regroupés en un doublet d'isospin ($I = 1/2$ avec la troisième composante $I_3 = +1/2$ pour u et $-1/2$ pour d.

quark

$Q/e = +2/3$: u, c, t

$Q/e = -1/3$: d, s, b

u "up" $I = 1/2$

d "down" $I = 1/2$ doublet avec u

s "strange" ($S = -1$)

c "charmed" ($C = +1$)

b "bottom" ($B^* = -1$)

antiquark

$Q/e = -2/3$: $\bar{u}, \bar{c}, \bar{t}$

$Q/e = +1/3$: $\bar{d}, \bar{s}, \bar{b}$

$m_u \approx m_d \approx 350 \text{ MeV}/c^2$

$m_s \approx 550 \text{ MeV}/c^2$

$m_c \approx 1800 \text{ MeV}/c^2$

$m_b \approx 4500 \text{ MeV}/c^2$

On suppose la présence d'un dernier quark, t "top", non encore découvert et encore plus lourd. Les quarks sont confinés à l'intérieur des hadrons et n'apparaissent jamais comme particules libres. Le quark s a pour nombre d'étrangeté $S = -1$, le quark c le nombre de charme $C = +1$ et le quark b le nombre de beauté $B^* = -1$ (à ne pas confondre avec le nombre baryonique B qui vaut $1/3$ pour chacun des quarks et $-1/3$ pour les antiquarks).

Les **leptons** portent des charges électriques entières 0 ou $\pm e$; les leptons apparaissent en doublets: à chaque lepton chargé correspond un neutrino sans masse (ou avec masse très faible) et sans charge.

leptons	antileptons	
$Q/e = -1 \quad e^-, \mu^-, \tau^-$	$Q/e = +1 \quad e^+, \mu^+, \tau^+$	$m_e = 0,511 \text{ MeV}/c^2$
$Q/e = 0 \quad \nu_e, \nu_\mu, \nu_\tau$	$\overline{\nu_e, \nu_\mu, \nu_\tau}$	$m_\mu = 105,6 \text{ MeV}/c^2$
		$m_\tau = 1800 \text{ MeV}/c^2$

Les neutrinos ont leur spin toujours dirigé en sens inverse de la vitesse.



Leur hélicité est négative et on dit qu'ils sont gauches. Les antineutrinos sont droits.

Les **bosons de jauge** sont les particules intermédiaires des interactions: ils comprennent le photon pour l'interaction électromagnétique, les bosons $W^\pm$ et Z^0 pour l'interaction faible, les huit gluons pour l'interaction forte et peut-être le graviton pour l'interaction gravitationnelle.

Les **bosons de Higgs** sont des particules hypothétiques associées aux changements de phase dus aux brisures de symétrie. On ne connaît ni leur nombre ni leur masse.

Carré 2 : équation de Dirac et formalisme lagrangien

Les particules sont caractérisées par leur masse, leur charge et leur moment cinétique propre (avec d'autres nombres quantiques internes) ; le moment cinétique propre, appelé spin de la particule, est quantifié et c'est un multiple entier ou demi-entier de la constante de Planck $\hbar$.

Lorsque le spin est un multiple entier de $\hbar$, on a affaire à un boson; s'il est nul, la fonction d'onde est un scalaire; s'il est égal à $\hbar$, la fonction d'onde est vectorielle.

Si le spin est un multiple demi-entier de $\hbar$, la particule est un fermion; les constituants élémentaires de la matière, quarks et leptons, ont tous un spin égal à $\hbar/2$; leur fonction d'onde a la structure géométrique d'un *spin*eur: en particulier, par rotation de 2π , le spineur change de signe.

Pour trouver l'équation à laquelle satisfait la fonction d'onde, on part de l'équation relativiste liant l'énergie et la quantité de mouvement: $E^2 = p^2c^2 + m^2c^4$ et on utilise le principe de correspondance qui fait passer de la physique classique à la physique quantique :

à E , on fait correspondre l'opérateur $i\hbar \frac{\partial}{\partial t}$ agissant sur la fonction d'onde

à $\mathbf{p}$ (vecteur à trois dimensions), on fait correspondre $\mathbf{P} = -i\hbar \nabla$

$$\text{où } \nabla = \left(\frac{\partial}{\partial x}, \frac{\partial}{\partial y}, \frac{\partial}{\partial z} \right) = \left(\frac{\partial}{\partial x^1}, \frac{\partial}{\partial x^2}, \frac{\partial}{\partial x^3} \right)$$

Si la fonction d'onde est scalaire, on aboutit à l'équation de Klein-Gordon, qui s'écrit, dans le système d'unités où $\hbar = c = 1$ (unités naturelles) :

$$(\square + m^2)\Phi = 0 \text{ où } \square \text{ est l'opérateur } \frac{\partial^2}{\partial t^2} - \left(\frac{\partial^2}{\partial x^2} + \frac{\partial^2}{\partial y^2} + \frac{\partial^2}{\partial z^2} \right)$$

Pour les fermions, Dirac cherche une équation linéaire en $\frac{\partial}{\partial t}$, de la forme $E\psi = (\alpha_i P_i + \beta m)\psi$ (somme sur l'indice i répété 2 fois avec $i = 1, 2, 3$)

On a alors

$$E^2\psi = (\alpha_i P_i + \beta m)(\alpha_j P_j + \beta m)\psi = (\alpha_i^2 P_i^2 + (\alpha_i \alpha_j + \alpha_j \alpha_i) P_i P_j + (\alpha_j \beta + \beta \alpha_j) m P_i + \beta^2 m^2)\psi$$

$$= 1 \qquad = 0 \qquad = 0 \qquad = 1$$

Ceci ne peut s'identifier à $E^2\psi = (P^2 + m^2)\psi$ que si α_i et β sont des matrices telles que $\alpha_1, \alpha_2, \alpha_3$ et β anticommulent et $\alpha_1^2 = \alpha_2^2 = \alpha_3^2 = \beta^2 = 1$.

Les matrices de dimension la plus faible satisfaisant à ces relations sont des matrices 4×4 .

L'équation s'écrit alors $i \frac{\partial \psi}{\partial t} = -i \alpha_i \frac{\partial \psi}{\partial x_i} + \beta m \psi$, ψ étant un spineur à 4 composantes.

En multipliant par β à gauche, $i \beta \frac{\partial \psi}{\partial t} = -i \beta \alpha_i \frac{\partial \psi}{\partial x_i} + m \psi$, et $(i \gamma^\mu \partial_\mu - m) \psi = 0$, équation de Dirac

avec les 4 matrices $\gamma^0, \gamma^1, \gamma^2, \gamma^3$ de Dirac, $\gamma^\mu = (\beta, \beta \alpha_1, \beta \alpha_2, \beta \alpha_3)$

$$\alpha_i = \begin{bmatrix} 0 & \sigma_i \\ \sigma_i & 0 \end{bmatrix} \quad \text{et} \quad \beta = \begin{bmatrix} 1 & 0 \\ 0 & -1 \end{bmatrix}, \quad I = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}$$

et les trois matrices $\sigma = (\sigma_1, \sigma_2, \sigma_3)$ sont les trois matrices de Pauli :

$$\sigma_1 = \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix}, \quad \sigma_2 = \begin{bmatrix} 0 & -i \\ i & 0 \end{bmatrix}, \quad \sigma_3 = \begin{bmatrix} 1 & 0 \\ 0 & -1 \end{bmatrix}$$

Ces équations peuvent s'exprimer dans le formalisme lagrangien :

En mécanique classique, un système à N degrés de liberté est décrit par N coordonnées

généralisées q_i et N vitesses généralisées $\dot{q}_i = \frac{dq_i}{dt}$.

On définit alors le lagrangien $L(q_i, \dot{q}_i, t) = T - V$, T et V étant l'énergie cinétique et l'énergie potentielle du système. En disant que la trajectoire suivie par le système est un extremum de

l'action $S = \int L dt$, on aboutit aux équations de Lagrange: $\frac{d}{dt} \left(\frac{\partial L}{\partial \dot{q}_i} \right) - \frac{\partial L}{\partial q_i} = 0$.

On passe d'un système discret, avec des coordonnées $q_i(t)$, à un système continu, c'est à dire un système dont les coordonnées varient de façon continue $\Phi(x_i, t)$ ($i = 1, 2, 3$) en remplaçant

$L(q_i, \dot{q}_i, t) \rightarrow \mathcal{L}(\Phi, \frac{\partial \Phi}{\partial x_\mu}, x_\mu)$, $\mathcal{L}$ est la densité lagrangienne et le lagrangien $L = \int \mathcal{L} d^3x$.

Les équations d'Euler-Lagrange deviennent $\frac{\partial}{\partial x^\mu} \left(\frac{\partial L}{\partial \partial_\mu \Phi} \right) - \frac{\partial L}{\partial \Phi} = 0$ avec $\partial_\mu \Phi = \frac{\partial \Phi}{\partial x^\mu}$.

Pour le champ scalaire, on a $\mathcal{L} = \frac{1}{2} (\partial_\mu \Phi) (\partial^\mu \Phi) - \frac{1}{2} m^2 \Phi^2$ qui mène à l'équation de Klein-Gordon $(\square + m^2) \Phi = 0$.

L'équation de Dirac provient du lagrangien $\mathcal{L} = i \bar{\psi} \gamma^\mu \partial_\mu \psi - m \bar{\psi} \psi$

Quant aux équations de Maxwell $\partial_\mu F^{\mu\nu} = j^\nu$, elles proviennent du lagrangien

$$\mathcal{L} = -\frac{1}{4} F_{\mu\nu} F^{\mu\nu} - j^\mu A_\mu \quad \text{où} \quad F^{\mu\nu} = \partial^\mu A^\nu - \partial^\nu A^\mu.$$

Carré 3 : électrodynamique quantique (Q.E.D.)

C'est la théorie quantique de l'électromagnétisme. Soit une particule de fonction d'onde $\psi(x, t)$ où x représente les variables d'espace et t le temps. Le carré de la norme $|\psi|^2 = \psi^* \psi$ représente la densité de probabilité de présence de la particule au point considéré et il s'ensuit que la phase d'une fonction d'onde est inobservable. Les équations du mouvement, obtenues à partir du lagrangien, contiennent des termes en $\psi^* \psi$ et des termes contenant des dérivées (du type

$\psi^* \partial_\mu \psi$) par exemple). Tous ces termes sont invariants par transformation de jauge globale,

c'est-à-dire par la transformation $\psi(x) \rightarrow e^{i\alpha} \psi(x)$ où α est une constante identique dans tout l'espace. Le groupe de ces transformations est le groupe $U(1)$ des transformations unitaires à 1 dimension.

La nature profite au maximum de sa liberté, et on aimerait pouvoir choisir α différemment en chaque point, c'est-à-dire $\alpha = \alpha(x)$. C'est ce que l'on appelle une transformation locale de jauge

$$\psi(x) \rightarrow \psi'(x) = e^{i\alpha(x)} \psi(x).$$

On a alors
$$\partial_\mu \psi'(x) = \partial_\mu (e^{i\alpha(x)} \psi(x)) = (i\partial_\mu \alpha(x) e^{i\alpha(x)} \psi(x) + e^{i\alpha(x)} \partial_\mu \psi(x)) = e^{i\alpha(x)} (\partial_\mu \psi(x) + i\psi(x) \partial_\mu \alpha) \neq e^{i\alpha(x)} \partial_\mu \psi(x).$$

Donc $\psi^*(x) \partial_\mu \psi$ n'est pas invariant par transformation locale de jauge.

Cependant, et c'est là que se trouve tout l'intérêt des théories de jauge, ce terme peut être rendu invariant à condition d'introduire un champ vectoriel sans masse (donc à longue portée) A_μ , appelé champ de jauge, la particule associée à ce champ sera un boson de jauge, le photon dans le cas considéré ici de l'électromagnétisme.

Pour un électron de charge $q = -e$, la dérivée est remplacée par la dérivée covariante

$$D_\mu = \partial_\mu + iqA_\mu = \partial_\mu - ieA_\mu \text{ qui va se transformer de façon covariante par transformation}$$

locale de jauge $D_\mu \psi \rightarrow e^{i\alpha(x)} D_\mu \psi$, le champ A_μ devant se transformer par transformation de

$$\text{jauge de deuxième espèce : } A_\mu \rightarrow A'_\mu = A_\mu + \frac{1}{e} \partial_\mu \alpha$$

$$\text{En effet } D_\mu \psi \rightarrow D'_\mu \psi' = (\partial_\mu - ieA'_\mu)(e^{i\alpha(x)} \psi(x)) = (\partial_\mu - ieA_\mu - i\partial_\mu \alpha)(e^{i\alpha(x)} \psi(x)) =$$

$$e^{i\alpha(x)} (i\partial_\mu \alpha \psi - ieA_\mu \psi + \partial_\mu \psi - i\partial_\mu \alpha \psi) = e^{i\alpha(x)} (-ieA_\mu \psi + \partial_\mu \psi) = e^{i\alpha(x)} D_\mu \psi(x)$$

et le terme $\psi^*(x) D_\mu \psi(x)$ est bien invariant par transformation locale de jauge.

Le fait d'utiliser D_μ à la place de ∂_μ spécifie la forme de l'interaction de la particule avec le champ. A_μ est ici le quadrivecteur potentiel électromagnétique, et la particule associée le photon, quantum du champ.

Pour un électron, le lagrangien (voir carré 2) s'écrit $\mathcal{L} = i\bar{\psi}\gamma^\mu\partial_\mu\psi - m\bar{\psi}\psi$; il mène à

l'équation de Dirac $(i\gamma^\mu\partial_\mu - m)\psi = 0$ où γ^μ représente les 4 matrices de Dirac, ψ étant un spineur à 4 composantes.

En remplaçant ∂_μ par D_μ , on obtient le lagrangien

$\mathcal{L} = i\bar{\psi}\gamma^\mu D_\mu\psi - m\bar{\psi}\psi = \bar{\psi}(i\gamma^\mu\partial_\mu - m)\psi + e\bar{\psi}\gamma^\mu\psi A_\mu$. Ce second terme est le lagrangien d'interaction d'un champ électromagnétique de photons avec l'électron. Le lagrangien final de QED s'obtient en ajoutant le lagrangien du champ électromagnétique libre qui s'écrit

$$-\frac{1}{4}F^{\mu\nu}F_{\mu\nu} \text{ avec } F_{\mu\nu} = \partial_\mu A_\nu - \partial_\nu A_\mu$$

$$\text{soit } \mathcal{L}_{\text{QED}} = \bar{\psi}(i\gamma^\mu\partial_\mu - m)\psi + e\bar{\psi}\gamma^\mu\psi A_\mu - \frac{1}{4}F^{\mu\nu}F_{\mu\nu}$$

Il n'y a pas de terme de masse pour le photon, car celui-ci ne serait pas invariant de jauge. On voit donc que l'invariance locale de jauge mène à l'électrodynamique quantique, qui contient tous les résultats du champ électromagnétique (équations de Maxwell) et de l'interaction fermion-photon (QED). Les résultats physiques s'obtiennent ensuite en utilisant la théorie des perturbations et se sont révélés en concordance extraordinaire avec les résultats expérimentaux.

Carré 4 : le groupe SU (3) de couleur et les gluons

Les quarks possèdent la propriété de couleur et peuvent se trouver dans un des trois états colorés R (red = rouge), G (green = vert), B (blue = bleu).

On peut représenter la couleur comme un vecteur dans un espace à trois dimensions complexe :

$$R = \begin{bmatrix} 1 \\ 0 \\ 0 \end{bmatrix}, G = \begin{bmatrix} 0 \\ 1 \\ 0 \end{bmatrix}, B = \begin{bmatrix} 0 \\ 0 \\ 1 \end{bmatrix}$$

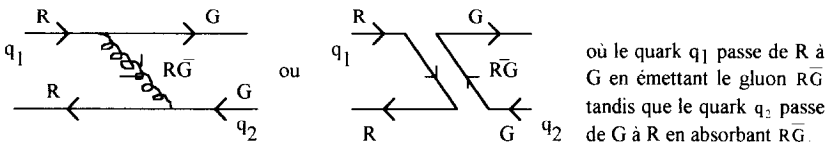
La couleur la plus générale est $\alpha R + \beta G + \gamma B$ avec $|\alpha|^2 + |\beta|^2 + |\gamma|^2 = 1$ pour normaliser à 1 la couleur.

On passe d'une couleur à une autre par une transformation du groupe SU (3) de couleur (SU(3)_C). Il s'agit du groupe spécial unitaire à 3 dimensions. Les transformations unitaires conservent la norme des vecteurs (la couleur étant normalisée à 1) ; spécial signifie que le déterminant de la transformation est égal à 1 (conservation de l'orientation de l'espace).

Aux trois couleurs correspondent trois anticouleurs, qui forment la représentation complexe conjuguée :

$$\bar{R} = (1, 0, 0), \bar{G} = (0, 1, 0), \bar{B} = (0, 0, 1)$$

Pour changer de couleur et passer par exemple de R à G, il faut un être bicolore qui enlève R et met à la place G (en enlevant l'anticouleur $\bar{G}$), suivant le schéma $R \rightarrow G + R\bar{G}$. Le diagramme de Feeyman correspondant est :



On peut former 8 combinaisons différentes entre une couleur et une anticouleur

$$R\bar{G}, R\bar{B}, G\bar{R}, G\bar{B}, B\bar{R}, B\bar{G}, \frac{1}{\sqrt{2}}(R\bar{R} - G\bar{G}), \frac{1}{\sqrt{6}}(R\bar{R} + G\bar{G} - 2B\bar{B}), \text{ la combinaison restante}$$

$$\frac{1}{\sqrt{3}}(R\bar{R} + G\bar{G} + B\bar{B}) \text{ ne porte pas de couleur et ne peut donc pas changer la couleur.}$$

Ces 8 combinaisons sont les 8 combinaisons bicolores des gluons (mais en fait, le gluon a également une structure spatio-temporelle et est un vecteur de l'espace-temps) ; la combinaison

$R\bar{G}$ correspond à la matrice produit de R par $\bar{G}$, soit

$$R \begin{bmatrix} 1 \\ 0 \\ 0 \end{bmatrix} \begin{bmatrix} 0 & 1 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{bmatrix} \bar{R}\bar{G} \\ \begin{bmatrix} 0 & 1 & 0 \end{bmatrix} \bar{G}$$

On obtient de cette façon 8 matrices qui sont 8 opérateurs du groupe $SU(3)$ (voir carré 5) (ce ne sont pas les générateurs habituellement utilisés); ces générateurs représentent la structure en couleur des 8 gluons G_i ($i = 1, 2, \dots, 8$).

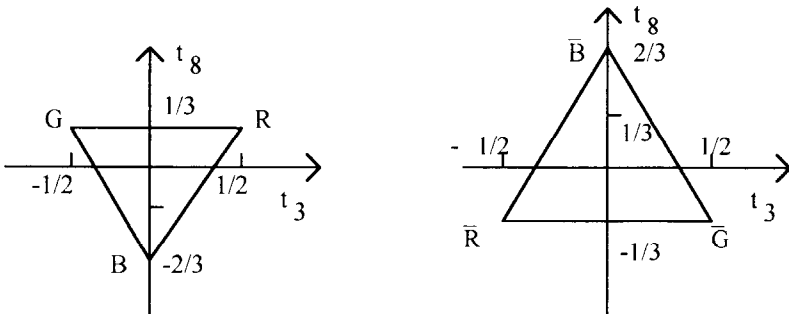
Il existe 2 gluons diagonaux correspondant aux deux dernières combinaisons; ce sont:

$$G_3 = \frac{1}{\sqrt{2}} \begin{bmatrix} 1 & 0 \\ 0 & -1 \\ 0 & 0 \end{bmatrix} \text{ et } G_8 = \frac{1}{\sqrt{3}} \begin{bmatrix} 1 & 0 \\ 0 & 1 \\ 0 & -2 \end{bmatrix}$$

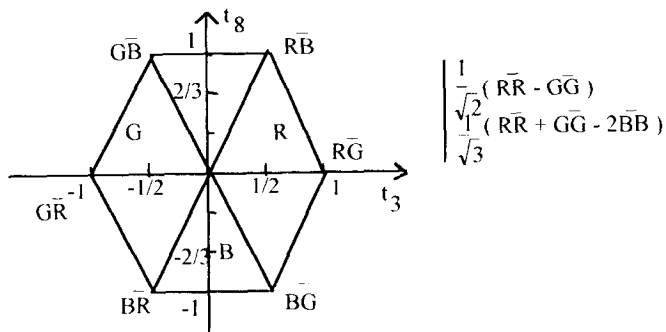
Définissons :

$$T_3 = \frac{1}{\sqrt{2}} G_3 = \frac{1}{2} \begin{bmatrix} 1 & 0 \\ 0 & -1 \\ 0 & 0 \end{bmatrix} \text{ et } T_8 = \frac{1}{\sqrt{3}} G_8 = \frac{1}{3} \begin{bmatrix} 1 & 0 \\ 0 & 1 \\ 0 & -2 \end{bmatrix}$$

Puisque ces deux opérateurs commutent, on peut définir les états colorés par les 2 valeurs propres t_3 et t_8 de ces deux opérateurs . Dans le plan (t_3 , t_8), cela donne pour les couleurs et les anticouleurs les diagrammes suivants sous forme de triangles:



Les valeurs propres des anticouleurs sont opposées à celles des couleurs et les anticouleurs dans le diagramme sont donc symétriques des couleurs par rapport à l'origine. Les valeurs propres des gluons bicolorés sont la somme des valeurs propres de la couleur et de l'anticouleur; pour les 3 gluons $R\bar{R}$, $R\bar{B}$ et $\bar{R}\bar{G}$, il faut donc décaler le triangle des anticouleurs pour amener le centre en R et donc $\bar{R}$ vient en O. D'où la construction:



Au cours des interactions de couleurs, on a donc conservation de t_3 et de t_8

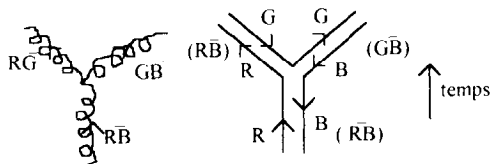
Exemple: $R \rightarrow B + \text{le gluon } R\bar{B}$

$$t_3 \quad 1/2 = 0 + 1/2$$

$$t_8 \quad 1/3 = -2/3 + 1$$

Les gluons possédant une couleur peuvent agir sur les autres gluons. On peut avoir par exemple

$$R\bar{B} \rightarrow R\bar{G} + G\bar{B}$$



Carré 5 : invariance de jauge non abélienne et QCD

La théorie de jauge des interactions fortes est la chromodynamique quantique (quantum chromodynamics QCD). Elle repose, comme nous l'avons vu, sur l'invariance par changement de couleur. QCD est bâtie sur la même idée que QED, mais le groupe de jauge $U(1)$ de QED va être remplacé par le groupe $SU(3)$ ($SU(3)$ de couleur). Les transformations de $SU(3)$ ne commutent pas entre elles, ce groupe est non abélien et QCD est une théorie de jauge non abélienne.

Une transformation U de $SU(3)$ s'écrit $e^{i\alpha^a(x)T_a}$ où a varie de 1 à 8; les $\alpha^a(x)$ sont les paramètres de la transformation et les 8 matrices T_a sont les générateurs du groupe. On prend généralement $T_a = \lambda_a/2$ avec :

$$\begin{aligned} \lambda_1 &= \begin{bmatrix} 0 & 1 & 0 \\ 1 & 0 & 0 \\ 0 & 0 & 0 \end{bmatrix} & \lambda_2 &= \begin{bmatrix} 0 & -i & 0 \\ i & 0 & 0 \\ 0 & 0 & 0 \end{bmatrix} & \lambda_3 &= \begin{bmatrix} 1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & 0 \end{bmatrix} & \lambda_4 &= \begin{bmatrix} 0 & 0 & 1 \\ 0 & 0 & 0 \\ 1 & 0 & 0 \end{bmatrix} \\ \lambda_5 &= \begin{bmatrix} 0 & 0 & -i \\ 0 & 0 & 0 \\ i & 0 & 0 \end{bmatrix} & \lambda_6 &= \begin{bmatrix} 0 & 0 & 0 \\ 0 & 0 & 1 \\ 0 & 1 & 0 \end{bmatrix} & \lambda_7 &= \begin{bmatrix} 0 & 0 & 0 \\ 0 & 0 & -i \\ 0 & i & 0 \end{bmatrix} & \lambda_8 &= \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & -2 \end{bmatrix} \end{aligned}$$

Une transformation U pouvant être considérée comme une suite de petites transformations, on va considérer le plus souvent une transformation infinitésimale: si $\alpha(x)$ petit, on a $U = e^{i\alpha^a(x)T_a} \approx 1 + i\alpha^a(x)T_a$ (on néglige les termes du deuxième ordre).

Les T_a obéissent à la relation de commutation $[T_a, T_b] = if_{abc}T_c$, où f_{abc} sont les constantes du groupe, réelles et antisymétriques par permutation de deux indices.

La fonction d'onde pour un quark dans l'espace des couleurs est

$$q = \begin{bmatrix} q_1 \\ q_2 \\ q_3 \end{bmatrix} \quad \text{où } q_1, q_2 \text{ et } q_3 \text{ représentent le quark rouge, vert ou bleu.}$$

Le lagrangien de Dirac valable pour les fermions de spin 1/2 est :

$$\mathcal{L}_0 = \bar{q}(i\gamma^\mu \partial_\mu - m)q = \bar{q}_j(i\gamma^\mu \partial_\mu - m)q_j \text{ avec } j = 1, 2, 3.$$

Ce lagrangien est invariant par transformation globale de jauge $U = e^{i\alpha^a T_a}$ (où α^a est constant dans tout l'espace), mais il n'est pas invariant par transformation locale car, pour une transformation infinitésimale

$$q \rightarrow (1 + i\alpha^a(x)T_a)q$$

$$\text{et } \partial_\mu q \rightarrow \partial_\mu [(1 + i\alpha^a(x)T_a)q] = (1 + i\alpha^a(x)T_a)\partial_\mu q + i\partial_\mu \alpha^a(x)T_a q.$$

Ce dernier terme brise l'invariance du lagrangien.

Procédons comme pour QED. On introduit 8 champs de jauge G_μ^a (ce sont les gluons) et on remplace la dérivée par la dérivée covariante $D_\mu = \partial_\mu + ig_S T_a G_\mu^a$, g_S étant le couplage de l'interaction forte; d'où le nouveau lagrangien:

$$\mathcal{L} = \bar{q}(i\gamma^\mu D_\mu - m)q = \bar{q}(i\gamma^\mu \partial_\mu - m)q - g_S \bar{q}\gamma^\mu T_a G_\mu^a q$$

Le second terme est le lagrangien d'interaction des quarks avec les gluons.

La transformation de jauge de deuxième espèce s'écrit $G_\mu^a \rightarrow G_\mu^a - \frac{1}{g} \partial_\mu \alpha^a(x) - f_{abc} \alpha^b G_\mu^c$, le dernier terme provenant du fait que les gluons, colorés, sont affectés par le changement de couleur même si α^a est indépendant de x (voir carré 6). On doit enfin ajouter un terme d'énergie cinétique pour les gluons, et le lagrangien final de QCD est

$$\mathcal{L} = \bar{q}(i\gamma^\mu \partial_\mu - m)q - g_S \bar{q}\gamma^\mu T_a G_\mu^a q - \frac{1}{4} G_{\mu\nu}^a G_a^{\mu\nu},$$

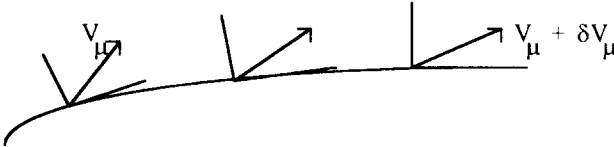
$$\text{où } G_{\mu\nu}^a = \partial_\mu G_\nu^a - \partial_\nu G_\mu^a - g_S f_{abc} G_\mu^b G_\nu^c$$

En plus du terme d'énergie cinétique, on remarque que le lagrangien du champ de jauge contient des termes d'autointeraction entre les bosons de jauge. Il faut noter enfin que les gluons sont sans masse.

Carré 6 : invariance de jauge et géométrie

Il existe un fondement géométrique à la notion de champ de jauge, qui fait appel à la théorie mathématique des espaces fibrés. Sans faire appel aux espaces fibrés, on va montrer, en prenant comme exemple les gluons du groupe $SU(3)$ de couleur de l'interaction forte, que ceux-ci ont une signification géométrique comme "connexion" dans l'espace des couleurs.

Revoyons tout d'abord rapidement quelques concepts géométriques de base pour un espace courbé. On a la notion de transport parallèle: pour comparer les valeurs d'un vecteur en deux points différents de l'espace, $V^\mu(x)$ et $V^\mu(x')$, il faut d'abord déplacer $V^\mu(x)$ de x à x' par transport parallèle de sorte que les deux vecteurs puissent être comparés dans le même système de coordonnées. Le transport parallèle est tel que le vecteur fait un angle fixe, au cours du transport, avec la tangente à la trajectoire.



Après transport, le vecteur devient $V^\mu + \delta V^\mu$, et si $x' = x + dx$, δV^μ est linéaire en dx^λ et V^μ de sorte que l'on pose $\delta V^\mu = -\Gamma_{\nu\lambda}^\mu V^\nu dx^\lambda$. Les coefficients $\Gamma_{\nu\lambda}^\mu$ sont les coefficients de connexion affine ou symboles de Cristoffel de l'espace courbé. La dérivation covariante est alors la différence du vecteur en x' , $V^\mu(x')$ et du vecteur $V^\mu(x)$ transporté parallèlement de x à x' , soit $DV^\mu = D_\lambda V^\mu dx^\lambda = V^\mu(x') - [V^\mu(x) + \delta V^\mu] = [V^\mu(x') - V^\mu(x)] - \delta V^\mu = [\partial_\lambda V^\mu + \Gamma_{\nu\lambda}^\mu V^\nu] dx^\lambda$.

La dérivée covariante est donc $D_\lambda V^\mu = \partial_\lambda V^\mu + \Gamma_{\nu\lambda}^\mu V^\nu$. La dérivée covariante d'un vecteur est un vecteur, alors que ce n'est pas le cas pour la dérivée ordinaire dans un espace courbé. L'opérateur de dérivation covariante D doit se transformer par changements de coordonnées comme un opérateur linéaire, et ceci donne les formules de transformations des symboles de Cristoffel.

Dans les théories de jauge, la fonction d'onde est donnée par un vecteur dans l'espace interne de

charge, $\psi(x)$ (Pour QCD, la fonction d'onde d'un quark est $q(x) = \begin{bmatrix} q_1(x) \\ q_2(x) \\ q_3(x) \end{bmatrix}$ dans l'espace interne des couleurs sur lequel agit $SU(3)_C$).

Ici, l'espace-temps est plat, mais par contre l'espace interne de charge peut changer d'un point de l'espace à l'autre; pour pouvoir comparer la fonction d'onde en deux points différents, il faut d'abord transporter parallèlement $\psi(x)$ de x à x' dans l'espace des couleurs avant de faire la comparaison. Ce transporté parallèle s'écrit $\psi(x) + \delta\psi(x)$; on définit alors la dérivée covariante de jauge $D\psi = \psi(x') - [\psi(x) + \delta\psi(x)] = d\psi - \delta\psi$.

Or la dérivée covariante de jauge se définit par : $D_\mu = \partial_\mu + igT_a G_\mu^a$ (voir carré 5). On identifie

donc $\delta\psi(x) = -igT_a G_\mu^a \psi(x) dx^\mu$ et $T_a G_\mu^a$ s'interprète géométriquement comme la "connexion" (c'est-à-dire le symbole Christoffel) de l'espace de charge interne.

Une telle interprétation est compatible avec les propriétés des transformations des champs de jauge G_μ^a , ou, plus généralement pour un groupe quelconque, des champs de jauge A_μ^a .

Si U est une transformation de l'espace de charge, un vecteur V de l'espace de charge se transforme suivant $V' = UV$, un covecteur comme $C' = CU^{-1}$ (on vérifie que $C'V' = CU^{-1}UV = CV$) et un opérateur linéaire $A = VC$ suivant $A' = V'C' = UVCU^{-1} = UAU^{-1}$.

La dérivée covariante D_μ doit se comporter comme un opérateur linéaire de l'espace de charge, donc, si l'on applique U , D_μ se transforme en $D'_\mu = UD_\mu U^{-1}$,

soit $\partial_\mu + igT_a A_\mu'^a = U(\partial_\mu + igT_a A_\mu^a)U^{-1}$ et $igT_a A_\mu'^a = U(igT_a A_\mu^a)U^{-1} + U\partial_\mu U^{-1}$

d'où $T_a A_\mu'^a = U(T_a A_\mu^a)U^{-1} - \frac{i}{g}U\partial_\mu U^{-1}$

Pour une transformation infinitésimale, $U(\theta) \approx 1 + iT_a \theta^a(x)$

$T_a A_\mu'^a = (1 + iT_a \theta^a)(T_b A_\mu^b)(1 - iT_a \theta^a) - \frac{i}{g}(1 + iT_a \theta^a)(T_a \partial_\mu \theta^a) =$

$T_a A_\mu^a + i\theta^a A_\mu^b [T_a, T_b] - \frac{1}{g}T_a \partial_\mu \theta^a$ en négligeant les termes du deuxième ordre en θ^a .

Or les relations de commutation des générateurs du groupe donnent $[T_a, T_b] = if_{abc}T_c$ où f_{abc} sont les constantes de structure du groupe, complètement antisymétriques.

D'où $T_a A_\mu'^a = T_a A_\mu^a - f_{abc}\theta^b A_\mu^c - \frac{1}{g}T_a \partial_\mu \theta^a$

En permutant circulairement les indices a, b, c du second terme, la constante de structure ne varie pas et

$A_\mu'^a = A_\mu^a - f_{abc}\theta^b A_\mu^c - \frac{1}{g}\partial_\mu \theta^a$ où l'on a identifié les termes en T_a .

On a retrouvé ainsi la loi de transformation de jauge de A_μ^a . Il résulte de même une corrélation

entre le tenseur de courbure d'un espace courbé $R_{\mu\alpha\beta}^\nu$ et le tenseur de second rang $F_{\mu\nu}^a$ qui

intervient dans le lagrangien du champ de jauge A_μ^a , donné par

$(D_\mu D_\nu - D_\nu D_\mu)\psi = ig(T_a F_{\mu\nu}^a)\psi$ où $F_{\mu\nu}^a = \partial_\mu A_\nu^a - \partial_\nu A_\mu^a - gf_{abc}A_\mu^b A_\nu^c$

Carré 7: théorie de jauge de l'interaction électrofaible

Cette théorie repose sur l'invariance de jauge par le groupe produit $SU(2)_L \times U(1)_Y$. Les particules de la première génération de leptons $[v_e, e^-]$ et de quarks $[u, d]$ forment des singulets

e_L^-, u_R et d_R et des doublets $l_L = \begin{bmatrix} v_L \\ e_L^- \end{bmatrix}$ et $q_L = \begin{bmatrix} u_L \\ d_L \end{bmatrix}$ pour les transformations de $SU(2)_L$.

Une transformation du groupe de jauge appliquée à un doublet (l_L ou q_L) donne pour l_L par exemple:

$l_L \rightarrow l'_L = e^{i\alpha^a(x)\tau_a} e^{i\beta(x)\frac{Y}{2}} l_L$, a variant de 1 à 3 puisque $SU(2)$ a trois générateurs

$\tau_a = \frac{\sigma_a}{2}$, les 3 matrices σ_a étant les matrices de Pauli (voir carré 2 ; pour un singulet, e_R par exemple, la transformation donne :

$e_R \rightarrow e'_R = e^{i\beta(x)\frac{Y}{2}} \times e_R$, le coefficient $\frac{1}{2}$ ayant été mis conventionnellement.

Pour pouvoir retrouver l'invariance de jauge de l'interaction électromagnétique, il faut que la

charge électrique Q soit reliée à l'hypercharge Y par la relation $Q = t_3 + \frac{Y}{2}$ de sorte que l'on a (t_3 étant la valeur propre associée à τ_3 et appelée isospin faible)

	t_3	Y	Q
v_L	1/2	-1	0
e_L^-	-1/2	-1	-1
e_R	0	-2	-1

	t_3	Y	Q
u_L	1/2	1/3	2/3
d_L	-1/2	1/3	-1/3
u_R		2/3	2/3
d_R		2/3	1/3

Considérons le lagrangien ne faisant intervenir que le doublet de v_e avec e^- (par exemple dans la réaction $v_e e^- \rightarrow e^- v_e$).

On part du lagrangien de Dirac des particules sans masse (car le terme de masse n'est pas invariant de jauge) :

$$\mathcal{L} = \overline{f_L}(i\gamma^\mu \partial_\mu)f_L + \overline{e_R}(i\gamma^\mu \partial_\mu)e_R.$$

En supposant l'invariance de jauge locale par le groupe $SU(2)_L \times U(1)_Y$, on est amené à

remplacer ∂_μ par $D_\mu = \partial_\mu + i g \tau_a W_\mu^a + \frac{g'}{2} Y B_\mu$ pour le premier terme et

$D_\mu = \partial_\mu + i \frac{g'}{2} Y B_\mu$ pour le second terme ; W_μ^a ($a = 1, 2, 3$) sont les trois champs de jauge

associés aux générateurs de $SU(2)_L$ et B_μ le champ de jauge associé à $U(1)_Y$, g et g' les deux couplages de $SU(2)_L$ et $U(1)_Y$.

En ajoutant les énergies cinétiques des champs W et B , le lagrangien devient :

$$\mathcal{L} = \overline{f}_L \gamma^\mu (i \partial_\mu - g \tau_a W_\mu^a - \frac{g'}{2} (-1) B_\mu) + \overline{e}_R \gamma^\mu (i \partial_\mu - \frac{g'}{2} (-2) B_\mu) e_R^- - \frac{1}{4} W_{\mu\nu}^a W_a^{\mu\nu} - \frac{1}{4} B_{\mu\nu} B^{\mu\nu},$$

$W_{\mu\nu}^a$ étant défini de la même façon que $G_{\mu\nu}^a$ pour l'interaction forte et $B_{\mu\nu}$ comme $F_{\mu\nu}$ de l'interaction électromagnétique.

Les interactions à courant neutre (c'est-à-dire sans transfert de charge) proviennent des termes

en W_μ^3 et B_μ . Les champs physiques sont le photon A_μ et le boson Z_μ^0 qui se déduisent de W_μ^3

et B_μ par une rotation d'angle ϑ_W , angle de Weinberg.

$$\begin{bmatrix} W_\mu^3 \\ B_\mu \end{bmatrix} = \begin{bmatrix} \cos \vartheta_W & \sin \vartheta_W \\ -\sin \vartheta_W & \cos \vartheta_W \end{bmatrix} \begin{bmatrix} Z_\mu^0 \\ A_\mu \end{bmatrix}$$

et

$$g \tau_3 W_\mu^3 + \frac{g'}{2} Y B_\mu = A_\mu (g \tau_3 \sin \vartheta_W + \frac{g'}{2} Y \cos \vartheta_W) + Z_\mu^0 (g \tau_3 \cos \vartheta_W - \frac{g'}{2} Y \sin \vartheta_W).$$

On identifie le premier terme avec l'interaction électromagnétique $e Q A_\mu$ et comme

$$Q = t_3 + \frac{Y}{2}, \text{ on doit avoir } g \sin \vartheta_W = g' \cos \vartheta_W = e.$$

Expérimentalement, l'angle de Weinberg est tel que $\sin^2 \vartheta_W = 0,234$.

Arrivés à ce point, on se heurte au problème de la masse des bosons de jauge et des fermions. Car le lagrangien de la théorie électrofaible suppose que la masse des particules est nulle, les termes de masse pour l'électron et les bosons de jauge n'étant pas invariants par $SU(2)_L$. Or les particules ont une masse, et l'interaction faible, de courte portée, nécessite des bosons massifs. La solution à ces problèmes a été trouvée dans la brisure spontanée de symétrie de $SU(2)_L \times U(1)_Y$ en

$U(1)_Q$, la symétrie électromagnétique, qui va engendrer les masses des bosons de jauge et des fermions.

Carré 8 : brisure spontanée de symétrie

Pour des particules scalaires libres satisfaisant à l'équation de Klein-Gordon, le lagrangien s'écrit

$$\mathcal{L} = \frac{1}{2}(\partial_\mu \Phi)^2 - \frac{1}{2}m^2\Phi^2, \text{ où } m \text{ désigne la masse de la particule (voir carré 2).}$$

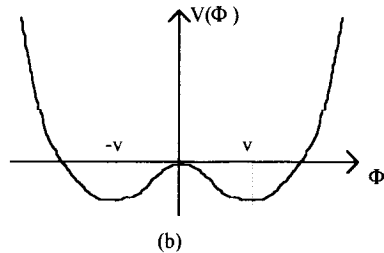
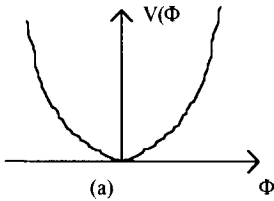
On va voir de manière simple sur un exemple comment, partant d'une particule sans masse dans une phase symétrique à haute température, le processus de brisure spontanée de symétrie engendre , ou mieux "révèle" la masse de la particule dans la phase à symétrie brisée à basse température.

Pour cela, considérons un champ de particules scalaires de Higgs décrit par le lagrangien (égal à

$$\text{l'énergie cinétique moins l'énergie potentielle) } \mathcal{L} = T - V = \frac{1}{2}(\partial_\mu \Phi)^2 - \left(\frac{1}{2}\mu^2\Phi^2 + \frac{1}{4}\lambda\Phi^4 \right)$$

avec $\lambda > 0$, où l'énergie potentielle est due à une interaction des particules entre elles.

$\mathcal{L}$ est invariant par l'opération de symétrie qui remplace Φ par $-\Phi$.



Le

cas (a) avec $\mu^2 > 0$ décrit un champ scalaire de masse μ . Le terme en Φ^4 indique que les particules du champ agissent les unes sur les autres.

Dans le cas (b) qui nous intéresse ici, $\mu^2 < 0$, le lagrangien a un terme de masse du mauvais signe, car le signe relatif du terme en Φ^2 et de l'énergie cinétique T est positif.

Dans ce cas, le potentiel V a deux minima, qui satisfont $\frac{\partial V}{\partial \Phi} = \Phi(\mu^2 + \lambda\Phi^2) = 0$,

$$\text{soit } \Phi = \pm v, \text{ avec } v = \sqrt{\frac{-\mu^2}{\lambda}}.$$

L'équilibre stable ayant lieu au minimum de l'énergie potentielle, posons $\Phi = v + \eta(x)$, où $\eta(x)$ représente les fluctuations quantiques autour du minimum. On obtient

$$\mathcal{L}' = \frac{1}{2}(\partial_\mu \eta)^2 - \lambda v^2 \eta^2 - \lambda v \eta^3 - \frac{1}{4}\lambda \eta^4 + V_{\min}.$$

Le champ η a un terme de masse du signe correct et, en identifiant λv^2 avec $\frac{1}{2} m_\eta^2$, on obtient

$$m_\eta = \sqrt{2\lambda v^2} = \sqrt{-2\mu^2}.$$

$\mathcal{L}$ et $\mathcal{L}'$ sont deux lagrangiens équivalents, qui renferment la même physique, mais, à haute température, il vaut mieux utiliser $\mathcal{L}$ et à basse température $\mathcal{L}'$.

En effet, à basse température, le champ de Higgs prend la valeur qui lui donne l'état de plus faible énergie, soit $\Phi = \pm v$. La nature doit faire ici un choix entre $+v$ et $-v$, et la symétrie de $\mathcal{L}$ est brisée par ce choix. (Si l'on prend $\Phi = -v$, on obtient un lagrangien différent.)

Quand l'énergie et la température montent, le champ de Higgs commence à fluctuer autour de la valeur minimale. Ces fluctuations correspondent, dans l'interprétation corpusculaire, à la présence de particules de Higgs massives. Quand la température dépasse une certaine valeur critique, les fluctuations deviennent si grandes que tout souvenir de la valeur initiale est perdue, le champ oscille autour d'une valeur d'équilibre nulle (c'est le maximum relatif $\Phi = 0$); la particule de Higgs dans cette phase n'a plus de masse et le système est dans sa phase symétrique décrite par le lagrangien $\mathcal{L}$.

La théorie est donc formulée de telle sorte que la symétrie existe lorsque les champs de Higgs ont une valeur nulle et se brise si le champ de Higgs prend une valeur non nulle.

Bien que ceci semble contraire à l'intuition, l'état de plus basse énergie est donc atteint dans un état où le champ de Higgs a une valeur non nulle parmi plusieurs valeurs possibles reliées entre elles par la symétrie du système.

Dans le cas d'une symétrie de jauge locale, le champ de Higgs satisfaisant à la symétrie considérée dans la phase symétrique à haute température, le passage à la phase de symétrie brisée donne une masse, non seulement à la particule de Higgs, mais aussi aux bosons intermédiaires et aux fermions (voir carré 9). L'interaction à symétrie brisée devient donc faible et à courte portée.

Carré 9 : masses des bosons $W^\pm$ et Z^0

Le phénomène de brisure spontanée de symétrie permet de donner une masse aux bosons de jauge $W^\pm$ et Z^0 de l'interaction électrofaible, ainsi qu'aux leptons et aux quarks. On suppose qu'il existe

$$\Phi = \frac{1}{\sqrt{2}} \begin{bmatrix} \Phi_1 + i\Phi_2 \\ \Phi_3 + i\Phi_4 \end{bmatrix}$$

des champs de Higgs, formant un doublet $SU(2)$, et on ajoute au lagrangien de l'interaction électrofaible (voir carré 7) le lagrangien invariant de jauge des champs de Higgs obtenu en remplaçant dans le lagrangien de Higgs (voir carré 8), ∂_μ par la dérivée covariante D_μ , soit $\mathcal{L} = (D_\mu \Phi)^\dagger (D_\mu \Phi) - \mu^2 \Phi^\dagger \Phi - \lambda (\Phi^\dagger \Phi)^2$ (où le + en exposant désigne la conjugaison hermitique).

Le minimum du potentiel est obtenu pour $\Phi^\dagger \Phi = \mu^2/2\lambda$ et on choisit le minimum particulier pour lequel $\Phi_1 = \Phi_2 = \Phi_4 = 0$ et $\Phi_3^2 = -\mu^2/\lambda = v^2$. Le vide particulier considéré est

$$\Phi_0 = \frac{1}{\sqrt{2}} \begin{bmatrix} 0 \\ v \end{bmatrix}$$

donc

Le terme de masse pour les bosons va être obtenu à partir de :

$$(D_\mu \Phi)^\dagger (D_\mu \Phi) = \left| \left(\partial_\mu + ig\tau_a W_\mu^a + i\frac{g'}{2} B_\mu \right) \Phi \right|^2 \quad (\text{la valeur de } Y \text{ pour le doublet de Higgs est}$$

$Y = +1$, car la formule $Q = t_3 + \frac{Y}{2}$ donne pour Φ_0 neutre, $-\frac{1}{2} + \frac{Y}{2} \Rightarrow Y = 1$), et, plus précisément, en remplaçant Φ par Φ_0 , le terme de masse provient de

$$\begin{aligned} \left| \left(ig\tau_a W_\mu^a + i\frac{g'}{2} B_\mu \right) \Phi_0 \right|^2 &= \frac{1}{8} \left| \begin{bmatrix} gW_\mu^3 + g'B_\mu & g(W_\mu^1 - iW_\mu^2) \\ g(W_\mu^1 + iW_\mu^2) & -gW_\mu^3 + ig'B_\mu \end{bmatrix} \begin{bmatrix} 0 \\ v \end{bmatrix} \right|^2 \\ &= \frac{1}{2} v^2 g^2 \left[(W_\mu^1)^2 + (W_\mu^2)^2 \right] + \frac{1}{8} v^2 (g'B_\mu - gW_\mu^3)^2 \end{aligned}$$

Le premier terme s'écrit $(\frac{1}{2}vg)^2 W_\mu^+ W^{-\mu}$ avec $W_\mu^+ = \frac{1}{\sqrt{2}}(W_\mu^1 - iW_\mu^2)$ et

$W_\mu^- = \frac{1}{\sqrt{2}}(W_\mu^1 + iW_\mu^2)$, et, comparant avec le terme de masse attendu pour un boson chargé,

$M_W^2 W^+ W^-$, on en tire $M_W = \frac{1}{2}vg$. La masse de Z^0 provient du second terme lorsque l'on

combine B_μ et W_μ^3 pour trouver le photon A_μ et Z^0 . On trouve $M_{Z^0} = \frac{1}{2}v\sqrt{g^2 + g'^2}$.

Les valeurs numériques donnent des résultats très proches des valeurs observées $M_W = 81 \text{ GeV}$ et $M_{Z^0} = 93 \text{ GeV}$, confirmant ainsi le modèle de l'interaction électrofaible, appelé modèle GSW (Glashow, Salam et Weinberg) ou modèle standard.

Le même doublet de Higgs qui engendre les masses de W^+ et Z^0 est suffisant pour donner des masses aux leptons et aux quarks. Prenons le cas de l'électron: si les champs de Higgs se couplent à l'électron, les nombres quantiques du doublet de Higgs sont juste ceux requis pour que le lagrangien

$$\mathcal{L} = -G_e \left[\bar{\nu}_L \bar{e}_L \right] \begin{bmatrix} \Phi_1 + i\Phi_2 \\ \Phi_3 + i\Phi_4 \end{bmatrix} e_R + \text{hermitique conjugué}$$

soit invariant de jauge (G_e est la constante de couplage de l'électron aux bosons de Higgs).

Dans le vide, ce lagrangien devient $-\frac{G_e v}{\sqrt{2}}(\bar{e}_L e_R + \bar{e}_R e_L) = -m_e \bar{e} e$ avec la masse de l'électron

$m_e = \frac{G_e v}{\sqrt{2}}$ Mais G_e n'est pas connu et cette relation ne permet pas de prédire la masse de l'électron. Il en est de même de la masse de la particule de Higgs prédite par la théorie. Cette particule n'a pas encore été découverte expérimentalement, et, si elle ne l'était pas après la mise en service des prochains accélérateurs de particules, il faudrait sans doute revoir la théorie standard GSW de l'interaction électrofaible.

Carré 10: théorie SU (5) de grande unification

Le groupe SU (n) a $n^2 - 1$ générateurs, le groupe SU (5) a donc 24 générateurs. Parmi ceux-ci, 4 sont diagonaux. L'hypercharge Y, introduite arbitrairement dans la théorie GSW pour retrouver la charge, apparaît ici comme égale, à un coefficient près, aux valeurs propres du générateur:

$$\lambda_0 = \frac{1}{\sqrt{15}} \begin{bmatrix} 1 & & & & \\ & 1 & & & \\ & & 1 & & \\ & & & -3/2 & \\ 0 & & & & -3/2 \end{bmatrix}$$

La dernière particule de la représentation $5^* = (3^*, 1) + (1, 2)$ est un lepton d'isospin faible $t_3 = -\frac{1}{2}$. On l'assimile donc au neutrino, dont la charge est nulle.

Or $Q = t_3 + \frac{Y}{2}$ (voir carré 7) et $Y = c\lambda_0$. On a donc pour le neutrino

$$0 = -\frac{1}{2} + \frac{1}{2} c \left(-\frac{3 \times 2}{2\sqrt{15}} \right) \Rightarrow c = -\frac{\sqrt{15}}{3} = -\sqrt{\frac{5}{3}}.$$

On constate, en calculant l'hypercharge et la charge, que les particules se classent naturellement, pour leur composante gauche, dans les représentations 5^* et 10 (voir texte).

On identifie la dérivée covariante relative à $U(1)_Y$ ($D_\mu = \partial_\mu + ig' B_\mu \frac{Y}{2}$) à la dérivée covariante relative à SU (5) lorsque seul le paramètre relié à λ_0 varie ($D_\mu = \partial_\mu + ig_G A_\mu^0 \frac{\lambda_0}{2}$).

On en tire $g'Y = g_G \lambda_0$ et comme $Y = -\sqrt{\frac{5}{3}} \lambda_0$, on a donc $g' = -\sqrt{\frac{3}{5}} g_G$. De la relation

$$g \sin \theta_W = g' \cos \theta_W \quad (\text{arré 7}), \text{ on tire } \tan \theta_W = \frac{g}{g'} \text{ et } \sin^2 \theta_W = \frac{g'^2}{g^2 + g'^2} = \frac{3}{8}.$$

Ceci est une des prédictions satisfaisantes de la théorie SU (5).

Aux 24 générateurs de SU (5) sont associés 24 bosons de jauge A_μ^a ($a = 0, 1, \dots, 23$). 8 sont associés aux gluons, 3 aux bosons W^1, W^2, W^3 et un au boson de jauge B_μ de $U(1)_Y$. Il en reste 12, qui forment un doublet de bosons colorés X et Y et le doublet de leurs antiparticules $\bar{X}$ et $\bar{Y}$, dont les charges sont $Q_X = -\frac{4}{3}$ et $Q_Y = -\frac{1}{3}$.

Lors de la brisure de symétrie $SU(5) \rightarrow SU(3)_C \times SU(2)_L \times U(1)_Y$, due à la valeur moyenne non nulle dans le vide d'un multiplet de champs de Higgs, les bosons X et Y acquièrent une grande masse, de l'ordre de l'énergie d'unification, qui sera, pendant longtemps encore, hors de portée des accélérateurs de particules.

Bibliographie

Perkins : introduction to high energy physics Cambridge University Press, 1982

Rabindra N. Mohapatra : unification and supersymmetry Springer-Verlag, 1986

Leader et Predazzi : An introduction to gauge theories and the new physics Cambridge University Press, 1982

Ta-Pei Cheng et Ling-Fong Li : gauge theory of elementary particles physics Oxford University Press, 1984

Elbaz : de l'électromagnétisme à l'électrofaible, Ellipses 1988

Grossetête et Vannucci : interactions et particules, Eyrolles 1991