

Theory of Fundamental Interactions

Leibniz University Hannover, Summer Semester 2020

TILL BARGHEER

Version: August 21, 2020

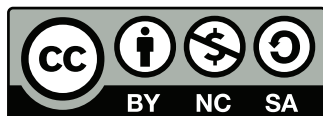
These lecture notes are based on a lecture course that I gave at Leibniz University Hannover during the summer semester 2020. Their structure largely follows the textbook “Modern Elementary Particle Physics” by Gordon Kane [1], with some additions, omissions, and modifications.

© 2020 Till Bargheer

This document and all of its parts are protected by copyright.

This work is licensed under the Creative Commons

“Attribution-NonCommercial-ShareAlike 4.0 International” License
(CC BY-NC-SA 4.0).



A copy of this license can be found at

<https://creativecommons.org/licenses/by-nc-sa/4.0/>

Contents

1	Introduction	5
1.1	The Framework of the Standard Model	6
1.2	The Forces	7
1.3	The Particles	8
1.4	Natural Units	10
2	Lagrangians, Conserved Currents, Interactions	11
2.1	Relativistic Notation	11
2.2	Lagrangians	12
2.3	Conserved Currents	14
2.4	Interactions	18
3	Gauge Invariance	21
3.1	Gauge Invariance in Electrodynamics	22
3.2	Gauge Invariance in Quantum Mechanics	22
3.3	Covariant Derivatives	23
4	Some Group Theory	25
4.1	Groups	25
4.2	Lie Groups and Algebras	26
4.3	Representations	27
4.4	Quantum Theory	29
4.5	$SO(3)$ and $SU(2)$ and Spin	29
5	Non-Abelian Gauge Theory	30
5.1	Strong Isospin	31
5.2	Non-Abelian Gauge Theories	34
5.3	Quarks and Leptons	36
6	Relativistic Fermions	37
6.1	The Dirac Equation	38
6.2	Conserved Current	43
6.3	Free Particle Solutions	44
6.4	Particles and Antiparticles	45
6.5	Chirality, Helicity, and Spin	47
6.6	The Dirac Lagrangian	51
6.7	Chirality in Field Theory	51
6.8	Coupling to the Photon	53
6.9	*Photon Polarizations and Helicity	56
7	The Standard Model Lagrangian	58
7.1	Lagrangians and Feynman Diagrams	58
7.2	The Gauge Group	62
7.3	Quark and Lepton States	63
7.4	The Quark and Lepton Lagrangian	68

8	The Electroweak Theory	69
8.1	Leptons: $U(1)$ and $SU(2)$ Terms	69
8.2	Neutral Currents	70
8.3	Charged Currents	73
8.4	Quark Terms and Further Families	74
8.5	The Fermion Gauge Boson Lagrangian	75
8.6	Historical Note	76
8.7	Masses?	77
9	Masses and the Higgs Mechanism	78
9.1	Spontaneous Symmetry Breaking	78
9.2	Breaking of a Continuous Symmetry	79
9.3	Abelian Higgs Mechanism	81
9.4	The Higgs Mechanism in the Standard Model	82
9.5	Fermion Masses	85
9.6	Summary of the Standard Model	87
10	Scattering and Decay	88
10.1	S-Matrix and Cross Sections	88
10.2	Lifetime and Decay Width	93
10.3	Scattering Through a Resonance	95
10.4	The W and Z Widths	100
11	More Aspects of the Standard Model	104
11.1	Measurements of Parameters	104
11.2	Higgs Production and Decay	105
11.3	Color Confinement, Jets, and Hadrons	107
11.4	Renormalization	110
11.5	Asymptotic Freedom.	113
11.6	Quark Mixing Angles	115
11.7	CP Violation	118
11.8	Parameters of the Standard Model	120
12	Beyond the Standard Model	121
12.1	Some Directions	121
12.2	Neutrino Masses	126
A	Numerical Data	130
	References	131

1 Introduction

The Forces. This course is about the theory of fundamental interactions. The terms “force” and “interaction” will be used interchangeably. At present, we know of four fundamental forces:

- Gravity: Gravity acts on all masses. It is the only relevant force on large scales (planetary systems, galaxies, the universe as a whole).
- Electromagnetism (electricity and magnetism).
- The weak nuclear force (subatomic).
- The strong nuclear force (subatomic).

All physical phenomena can ultimately be reduced to these four forces. The chemical elements, the light of the sun, the colors, all materials, life, everything can be brought down to the interplay between these four forces and the particles they act on. This is a truly amazing fact of physics.

The Standard Model. The last three of these four forces are described by the *Standard Model of Particle Physics*. Gravity is an outsider. It is extremely weak compared to the other forces, and plays no role at microscopic scales. The Standard Model of particle physics is the main subject of this lecture course. Some of its properties are:

- The Standard Model describes all forces and interactions between all elementary particles. Except gravity, which is however completely negligible at the microscopic scales of particle physics (except at enormous mass densities, for example near black holes, or shortly after the big bang, that will not be relevant for us).
- Although called a “model”, it is as fully a mathematical theory as there ever has been in the history of science.
- It is an incredibly successful theory that agrees with essentially all experimental data with extraordinary precision, across a huge range of scales.
- It was essentially completed by the 1980s, combining the results and achieving the goals of many centuries of physics. The theory completes the vision of the Greeks to reduce nature to its fundamental constituents. As far as all our experiments can see, all particles described by the Standard Model are indivisible and truly elementary.
- It is formulated in terms of *relativistic quantum field theory*, whose formulation began in the 1930s, and has been continuously developed since then.
- The Standard Model was fully completed by the experimental discoveries of the top quark at Fermilab in 1996 and of the Higgs boson at CERN in 2012.
- Just as electrodynamics depends on the electron mass and the strength of the electromagnetic coupling, the Standard Model depends on
 - the masses of the quarks, leptons, and gauge bosons, and
 - the coupling strengths of the electromagnetic, weak, and strong interactions.

Given these inputs, the Standard Model describes all known particle experiments! In its domain, it remains unchallenged.

- Despite all its success, the Standard Model does not solve *all* problems. There are still puzzles to solve!

Course Goals. The goals for this course are the following:

- We want to understand the way in which the Standard Model is formulated.
- We want to learn to understand (or estimate) some predictions of the theory, both to see how it works, and to understand the tests of the theory.
- We want to understand why it is now widely accepted that the Standard Model actually describes nature.
- We want to take a peek at what lies beyond the Standard Model.

Prerequisites. The level of this course is basic. The only required knowledge is introductory quantum mechanics (including spin), and some elementary classical mechanics and electrodynamics.

1.1 The Framework of the Standard Model

To understand the physical world, we need three kinds of knowledge:

- The particles everything is made of,
- The interactions (forces) among those particles,
- The rules for calculating the resulting world.

In *classical mechanics*, for any force F , the motion follows from Newton's law $F = ma$. The force could have various origins, for example gravity,

$$F = \frac{G_N m M}{r^2} , \quad (1.1)$$

or the Coulomb force

$$F = \frac{K q Q}{r^2} . \quad (1.2)$$

In *quantum theory* on the other hand, the dynamics are governed by the Schrödinger equation

$$H\psi = i \frac{\partial \psi}{\partial t} . \quad (1.3)$$

The Schrödinger equation replaces Newton's law, and holds for any *Hamiltonian* H . The Standard Model describes also massless particles that travel at the speed of light. Hence the description necessarily has to be *relativistic*! In relativistic theories, it is often useful to use *Lagrangians*, which are Lorentz scalars, rather than Hamiltonians, which are energies, that is components of a four-vector, and therefore behave non-trivially under Lorentz transformations.

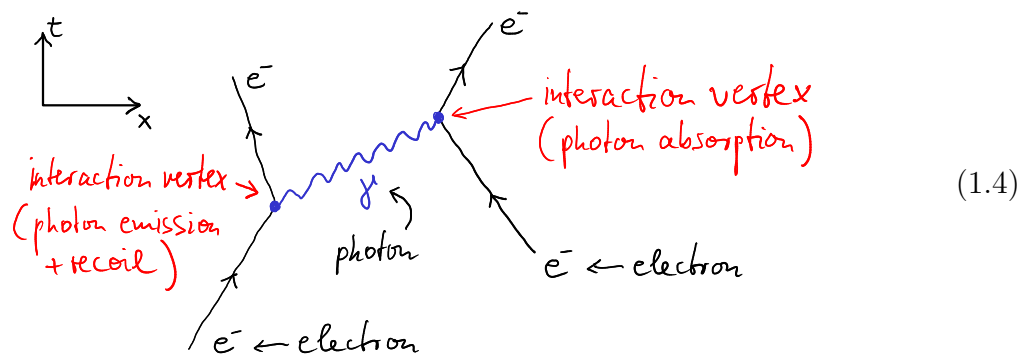
The formalism to compute “the motion” (that is, cross sections, decay rates, etc.) in a relativistic quantum field theory is to start with the appropriate Lagrangian, and extract

Feynman rules from this Lagrangian to write matrix elements / amplitudes. The squared matrix elements then yield transition probabilities, as is standard in quantum theory.

The combination of quantum theory with special relativity leads to quantum field theory. Intuitively, one can understand this as follows: Suppose there are various interacting particles. When one of these particles gets pushed, it will exert forces on the other particles. But the resulting interaction forces cannot produce instantaneous changes in their motions, since no signal can travel faster than the speed of light. Instead, the particle is the source of various fields that carry energy (and perhaps other quantities) through space. Eventually, those fields interact with the other particles. Because of quantum theory, the energy (and other quantities) is quantized, that is it is carried by discrete energy quanta. These quanta are identified as *particles transmitting the force*. Therefore, in a quantum field theory, *elementary interactions are interpreted in terms of exchanges of particles*.

The Standard Model is a *gauge theory*. Gauge theories are a special class of quantum field theory. They are based on an invariance principle that necessarily implies certain interactions among the particles. The interaction strength is proportional to a *charge*. This is familiar from electrodynamics: The charge e measures both the charge of the particles and the interaction strength.

Based on the above description, the basic view of particle interactions is as follows:



An electron emits a photon and recoils. The photon is absorbed by another electron (or other charged particle), which changes its motion in consequence. Such diagrams are useful pictures of what is happening. But they are more! Every such *Feynman diagram* can be converted to a mathematical expression for a matrix element / amplitude for a certain process. The rules for this conversion are called *Feynman rules*. The diagram / matrix element above gives Coulomb's law (in the non-relativistic limit).

1.2 The Forces

As stated at the very beginning, we know of four fundamental forces. Let us describe them in some more detail:

Gravity: Gravity is an attractive force between all masses. At the classical level, it is fully described by general relativity. No satisfactory *quantum* theory of gravity exists today (string theory is one approach). This is not a big problem in most cases, as the gravitational force is extremely weak compared to all other forces, and is therefore negligible for most of particle physics. Gravity is not part of the Standard Model.

Electromagnetism: The electric and magnetic interactions between all electrically charged particles and light are described by electromagnetism. This force is responsible for almost all physical processes of everyday matter (solids being solid,

heat, electrical conductivity, etc.). Electrodynamics is classically described by Maxwell's equations. The quantum theory is called *quantum electrodynamics* (QED). It was developed in the 1930s and 40s, and was awarded with the Nobel Prize in 1965 for Tomonaga, Schwinger, and Feynman. QED served as the model and template for all subsequent quantum field theories!

Weak force: The weak force is one of the two nuclear forces. It is an interaction between subatomic particles that is responsible for radioactive decay. Its range is limited to subatomic distances (less than the proton diameter of ≈ 1 fm). It is much weaker than the electromagnetic force and the strong force. Correspondingly, the timescale for weak interactions is 10^6 times larger than for electromagnetic interactions (10^{-13} s vs. 10^{-19} s). The weak force plays a critical role for energy creation in the sun, and for the generation of heavy elements in stars.

We do have a well-tested theory of the weak force (a quantum theory, there is no useful classical limit in this case!), and the weak and electromagnetic interactions are unified in the *electroweak theory*. The 1979 Nobel prize was awarded to Glashow, Salam, and Weinberg for the formulation of this theory.

Strong force: The strong force is another nuclear (subatomic) force that holds atomic nuclei (made of protons and neutrons) together, in spite of the electric repulsion of the protons. It is much stronger than the electromagnetic force. The elementary particles that experience the strong interaction are called *quarks*. The charge that is carried by the quarks and that responds to the strong force is called *color* (even though it has nothing to do with everyday colors). Quarks form color-neutral bound states that are held together by the strong force, just as electrons, protons, and neutrons form electrically neutral atoms. The quark bound states are called *hadrons* (protons, neutrons, pions, kaons, ...). Just as the residual electromagnetic field around neutral atoms causes them to form molecules, there is a residual strong force (color) field around hadrons, which is the nuclear force that lets protons and neutrons form nuclei.

The theory of the strong force is called *quantum chromodynamics* (QCD), a non-Abelian gauge theory.

The success of past unifications (electricity and magnetism into Maxwell's theory of electromagnetism, electromagnetism and the weak force into the electroweak theory) has prompted people to try to unify the electroweak theory with QCD to a "grand unified theory". Some approaches and candidates do exist, but none of them has been established as a real theory of nature to date.

The Standard Model is a combination of the electroweak theory and QCD.

1.3 The Particles

The particles that compose the world fall in two categories

- Matter particles (quarks and leptons), and
- Gauge bosons (these are the particles that transmit the forces in gauge theories).

All matter particles are fermions with spin $1/2$. By definition, all matter particles that carry color charge are quarks. All other matter particles are called *leptons*. So far, we

know of six different kinds of quarks (six quark “flavors”), and six different kinds of leptons (six lepton “flavors”).

Quarks. The six quarks are called (for no particular reason except history): *up*, *down*, *strange*, *charm*, *bottom*, *top*, or u, d, s, c, b, t for short. They naturally fall into pairs called *doublets* (we will later see why) that are called *families* or *generations*:

$$\begin{pmatrix} u \\ d \end{pmatrix}, \quad \begin{pmatrix} c \\ s \end{pmatrix}, \quad \begin{pmatrix} t \\ b \end{pmatrix}. \quad (1.5)$$

The quarks in the top row have electric charge $2/3 e$, the quarks in the bottom row have charge $-1/3 e$, where e is the absolute value of the electron’s charge (which equals $-e$).

Each quark flavor comes in three different colors. Quarks carry another quantum number called *baryon number* B . All quarks have $B = 1/3$. Baryon number is conserved (experimentally). Because of the strong force, all colored particles are normally bound inside colorless hadrons. Still, quarks exist as individual particles! Their masses cannot be calculated or derived theoretically, they have to be measured experimentally, and are put in as parameters of the theory. As quarks only occur in bound states, measuring their masses is a subtle analysis. Their masses are:

$$m_u = 2.16 \text{ MeV}, \quad m_c = 1.27 \text{ GeV}, \quad m_t = 173 \text{ GeV}, \quad (1.6)$$

$$m_d = 4.67 \text{ MeV}, \quad m_s = 93 \text{ MeV}, \quad m_b = 4.18 \text{ GeV}. \quad (1.7)$$

Leptons. Also the leptons are arranged in three families:

$$\begin{pmatrix} \nu_e \\ e \end{pmatrix}, \quad \begin{pmatrix} \nu_\mu \\ \mu \end{pmatrix}, \quad \begin{pmatrix} \nu_\tau \\ \tau \end{pmatrix}. \quad (1.8)$$

The leptons in the bottom row are the *charged leptons*: The *electron* e , the *muon* μ , and the *tau lepton* τ . All of them have electric charge $-e$. Each charged lepton comes with its own *neutrino*: The electron neutrino ν_e , the muon neutrino ν_μ and the tau neutrino ν_τ . All neutrinos are electrically neutral, that is have zero electric charge. As far as we know, the charged leptons e, μ, τ do not undergo transitions into each other.

As for quarks, the *lepton masses* have to be measured experimentally. The *neutrino masses* are not measured, but we know that at least two of them are non-zero, albeit extremely small. From cosmological arguments, we know that their sum must be smaller than $\approx 0.2 \text{ eV}$. The masses of the charged leptons are

$$m_e = 511 \text{ keV}, \quad m_\mu = 105.7 \text{ MeV}, \quad m_\tau = 1.777 \text{ GeV}. \quad (1.9)$$

Gauge Bosons. All force-carrying particles (quanta of the force fields) are bosons (with integer spin). Since the QFT of the Standard Model is a gauge theory, they are called *gauge bosons*.

- The electromagnetic force is mediated by *photons*.
- The weak force is mediated by $W^\pm$ and Z^0 bosons.
- The strong (color) force is mediated by *gluons*.
- Gravity is mediated by *gravitons*.

The following table summarizes the properties of the various gauge bosons:

Gauge boson	Interacts with	Mass	Spin
Graviton (gravity)	all particles	massless	2
Photon (EM force)	all electrically charged particles	massless	1
$W^\pm$, Z^0 (weak force)	quarks, leptons, $W^\pm$, Z^0 bosons	heavy	1
Gluons (strong force)	all colored particles (quarks and gluons)	massless	1

In a standard gauge theory, all gauge bosons are massless. Most of the Standard Model gauge bosons are indeed massless, but the weak gauge bosons $W^\pm$ and Z^0 are massive. This is explained by the Higgs mechanism in the context of the electroweak theory.

- Photons are familiar.
- Individual gravitons interact too weakly to be detected.
- The gluons were predicted to exist, and were observed first at the electron-positron collider PETRA in Hamburg in 1979.
- The $W^\pm$ and Z^0 bosons were predicted by the theory, and were observed at the proton-antiproton collider at CERN in 1983, with the expected properties ($m_W = 80.4 \text{ GeV}$, $m_Z = 91.2 \text{ GeV}$).

Higgs Boson. To make a consistent theory of particle masses and interactions, one more class of particles is needed: The spin zero (or scalar) *Higgs boson*. The electroweak theory requires one electrically neutral Higgs boson, but more could exist. The Higgs boson was discovered experimentally in 2012 at the CERN LHC.

The prediction and subsequent discoveries of gluons, the $W^\pm$ and Z^0 bosons, and of the Higgs boson, all with the expected properties, is part of the incredible success of the Standard Model.

Antiparticles. Lastly, each particle has an *antiparticle*, with the opposite values of electric charge, color charge, and flavor charge (weak charge) than the corresponding particle, but with the same mass and the same spin. Some particles are their own antiparticle (e.g. the photon). We denote antiparticles either by an inverted charge label (for example the positron e^+ is the antiparticle of the electron e^-), or by a bar (e.g. proton $p \leftrightarrow$ antiproton $\bar{p}$). Antiparticles are nothing special, they are just particles with specific properties.

1.4 Natural Units

We will mostly use a unit system called *natural units*, in which as many natural quantities as possible are set to unity. This will make formulas simpler and more readable. In natural units, $\hbar = 1$ and $c = 1$. Then energy (mc^2), momentum (mc), and mass (m) will all have dimension of mass, and will usually be stated in GeV ($1 \text{ GeV} \approx 1.6 \cdot 10^{-10} \text{ J}$). To convert any quantity back to SI units, one just has to multiply by the appropriate factors of $\hbar$ and c , where

$$\hbar = 6.6 \cdot 10^{-25} \text{ GeV s}, \quad c = 3 \cdot 10^{10} \text{ cm/s}. \quad (1.10)$$

Hence, for example

$$1 \text{ s} = 3 \cdot 10^{10} \text{ cm} , \quad (1.11)$$

$$\begin{aligned} 1 \text{ fermi} &= 10^{-13} \text{ cm} = 10^{-13} \text{ cm} \cdot \frac{1}{3} 10^{-10} \text{ s/cm} = \frac{1}{3} 10^{-23} \text{ s} \\ &= \frac{1}{3} 10^{-23} \text{ s} \cdot (6.6 \cdot 10^{-25} \text{ GeV s})^{-1} \\ &= \frac{1}{20} 100 \text{ GeV}^{-1} = 5 \text{ GeV}^{-1} , \end{aligned} \quad (1.12)$$

$$\begin{aligned} 1 \text{ GeV}^{-2} &= (6.6 \cdot 10^{-25} \text{ s})^2 = (6.6 \cdot 10^{-25} \cdot 3 \cdot 10^{10} \text{ cm})^2 \\ &= (20 \cdot 10^{-15} \text{ cm})^2 = 400 \cdot 10^{-30} \text{ cm}^2 = 0.4 \cdot 10^{-27} \text{ cm}^2 . \end{aligned} \quad (1.13)$$

2 Lagrangians, Conserved Currents, Interactions

We want to get to the mathematical formulation of the Standard Model. In this course, we want to describe the theory without going into the details of quantum field theory and difficult computations. Many results can be determined quite simply to a good approximation, which is fully sufficient to understand the physics. Nevertheless, before we can formulate the Standard Model, we need to collect some theoretical concepts in this and the next few sections.

2.1 Relativistic Notation

We will denote four-vectors by

$$a^\mu = (a^0; a^1, a^2, a^3) = (a^0; \mathbf{a}) , \quad (2.1)$$

for example, the spacetime coordinate vector is

$$x^\mu = (x^0; x^1, x^2, x^3) = (t; x, y, z) = (t; \mathbf{x}) , \quad (2.2)$$

and the momentum four-vector is

$$p^\mu = (p^0; p^1, p^2, p^3) = (E; p_x, p_y, p_z) = (E; \mathbf{p}) . \quad (2.3)$$

Four-vector indices are raised and lowered with the metric tensor

$$g_{\mu\nu} = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & -1 & 0 & 0 \\ 0 & 0 & -1 & 0 \\ 0 & 0 & 0 & -1 \end{pmatrix} , \quad (2.4)$$

that is

$$a_\mu = (a_0; a_1, a_2, a_3) = g_{\mu\nu} a^\nu = (a^0; -a^1, -a^2, -a^3) . \quad (2.5)$$

We use the Einstein summation convention: Repeated indices are implicitly summed over:

$$g_{\mu\nu} a^\nu \equiv \sum_{\nu=0}^3 g_{\mu\nu} a^\nu . \quad (2.6)$$

The Lorentz-invariant scalar product between two four-vectors a^μ and b^μ is given by

$$a_\mu b^\mu = a^\mu g_{\mu\nu} b^\nu = a^0 b^0 - a^1 b^1 - a^2 b^2 - a^3 b^3 = a^0 b^0 - \mathbf{a} \cdot \mathbf{b}. \quad (2.7)$$

Partial derivatives with respect to spacetime coordinates are collected in the covariant vector

$$\partial_\mu = \frac{\partial}{\partial x^\mu} = \left(\frac{\partial}{\partial t}, \frac{\partial}{\partial x}, \frac{\partial}{\partial y}, \frac{\partial}{\partial z} \right) = (\partial_0; \nabla), \quad (2.8)$$

such that

$$\partial^\mu = \frac{\partial}{\partial x_\mu} = (\partial^0; -\nabla), \quad (2.9)$$

and

$$\partial_\mu a^\mu = \frac{\partial a^0}{\partial t} + \nabla \mathbf{a}. \quad (2.10)$$

Finally, by definition,

$$\partial_\mu \partial^\mu = \frac{\partial^2}{\partial t^2} - \frac{\partial^2}{\partial x^2} - \frac{\partial^2}{\partial y^2} - \frac{\partial^2}{\partial z^2} = \partial_0^2 - \nabla^2 = \partial^\mu \partial_\mu. \quad (2.11)$$

2.2 Lagrangians

Lagrangians are central objects in field theory. We will first recall the role of Lagrangians in classical mechanics and electrodynamics, and only then consider Lagrangians for more general field theories.

Classical Mechanics. In classical mechanics, the Lagrangian L is a function of the phase space variables (positions and velocities). It is given by $L = T - V$, where T is the kinetic energy, and V is the potential energy. The action functional

$$S = \int_{t_0}^{t_1} L dt \quad (2.12)$$

takes a path in phase space as its argument. Extremizing (minimizing) the action S produces the Euler–Lagrange equation for the phase space variables, which are equivalent to Newton’s laws of motion.

Electrodynamics. To write the Lagrangian for classical electrodynamics, we first have to recall some definitions. The dynamical quantities are the electric and magnetic fields $\mathbf{E}$ and $\mathbf{B}$, which are sourced by the charge density ρ and the current density $\mathbf{J}$. The fields $\mathbf{E}$ and $\mathbf{B}$ are not independent. It is useful to write them in terms of the potentials V and $\mathbf{A}$ (which are also fields) as

$$\mathbf{E} = -\nabla V - \frac{\partial \mathbf{A}}{\partial t}, \quad \mathbf{B} = \nabla \times \mathbf{A}. \quad (2.13)$$

In components, these equations read

$$E^i = \partial^i A^0 - \partial^0 A^i, \quad B^i = \epsilon^{ijk} \partial_j A_k, \quad (2.14)$$

where ϵ^{ijk} is the totally anti-symmetric tensor. The potentials, as well as the charge and current densities, can be combined into Lorentz four-vectors:

$$A^\mu = (V, \mathbf{A}), \quad J^\mu = (\rho, \mathbf{J}). \quad (2.15)$$

One then defines the (anti-symmetric) field strength tensor

$$F^{\mu\nu} := \partial^\mu A^\nu - \partial^\nu A^\mu, \quad (2.16)$$

such that

$$F^{0i} = \partial^0 A^i - \partial^i A^0 = -E^i, \quad F^{ij} = \partial^i A^j - \partial^j A^i = \epsilon^{ijk} B^k, \quad (2.17)$$

that is

$$F^{\mu\nu} = \begin{pmatrix} 0 & -E^1 & -E^2 & -E^3 \\ +E^1 & 0 & -B^3 & +B^2 \\ +E^2 & +B^3 & 0 & -B^1 \\ +E^3 & -B^2 & +B^1 & 0 \end{pmatrix}. \quad (2.18)$$

Notice that under the transformation

$$A^\mu \rightarrow A^\mu + \partial^\mu \phi, \quad (2.19)$$

the field strength tensor transforms as

$$F^{\mu\nu} \rightarrow F^{\mu\nu} + \partial^\mu \partial^\nu \phi - \partial^\nu \partial^\mu \phi = F^{\mu\nu}, \quad (2.20)$$

that is the field strength tensor is *invariant*. This is the first example of a *gauge transformation*.

In a field theory such as electrodynamics, the dynamical variables are no longer discrete positions and velocities (or momenta), but rather fields that permeate all of space. The Lagrangian L is therefore a function of the field configuration. It is given by an integral over all of space,

$$L = \int \mathcal{L}(\mathbf{x}) d^3\mathbf{x}, \quad (2.21)$$

where $\mathcal{L}(\mathbf{x})$ is a function of the field variables at position $\mathbf{x}$ called the *Lagrangian density*. The action is again given by the integral of L over time,

$$S = \int_{t_0}^{t_1} dt L = \int_{t_0}^{t_1} dt \int d^3\mathbf{x} \mathcal{L}(t, \mathbf{x}). \quad (2.22)$$

In many cases, the boundary conditions are such that one can integrate from the infinite past to the infinite future, such that the integral runs over all of spacetime:

$$S = \int dt L = \int dt d^3\mathbf{x} \mathcal{L}(t, \mathbf{x}) = \int d^4x \mathcal{L}(x). \quad (2.23)$$

The Lagrangian density for electrodynamics reads

$$\mathcal{L} = \frac{1}{2} (\mathbf{E}^2 - \mathbf{B}^2) - \rho V + \mathbf{J} \cdot \mathbf{A}. \quad (2.24)$$

In terms of the field strength tensor, this reads (Problem 1.2)

$$\mathcal{L} = -\frac{1}{4} F_{\mu\nu} F^{\mu\nu} - J_\mu A^\mu. \quad (2.25)$$

Extremizing the action $S = \int \mathcal{L}(x) d^4x$ by variation produces the Euler–Lagrange field equations, which for this Lagrangian become Maxwell’s equations (Problem 1.3). We conclude that all physics of electromagnetism is contained in the Lagrangian density $\mathcal{L}$, which is written in terms of the fields and potentials. The first term $-1/4 F_{\mu\nu} F^{\mu\nu}$ is the kinetic term, it is quadratic in the fields and contains derivatives. The second term $-J_\mu A^\mu$ is an interaction term that couples the dynamical potential A to the external source J .

Particle Physics. Theories of particle physics are always defined in terms of a Lagrangian density. Since one mostly deals with the density $\mathcal{L}$, and only rarely with the integrated form L , one typically calls $\mathcal{L}$ simply “the Lagrangian”. Starting with the Lagrangian (density) $\mathcal{L}$, and using the rules of quantum field theory, all physical observables of particle physics can be calculated.

The Lagrangian is written in terms of the elementary fields, whose quanta are the fundamental particles. For electrodynamics, the *photon* is the quantum of the electromagnetic field, represented by A^μ , and the electron is the quantum of a fermion field ψ (we will learn about such fields later).

The kinetic parts of the Lagrangian $\mathcal{L}$ are completely determined by the field content, they only depend on the spins of the various fields/particles. The potential/interaction parts of $\mathcal{L}$ specify the forces. The Lagrangian is a single function that determines the dynamics of the theory. It must be a scalar in every relevant space. In particular, in a relativistically invariant theory, $\mathcal{L}$ must be Lorentz invariant, which ensures that all predictions computed from $\mathcal{L}$ will also be Lorentz invariant.

Real Scalar Field. Above, we wrote the Lagrangian for the electromagnetic field A . We will also need the Lagrangian for a scalar field ϕ . Such a field can be thought of as arising from some source, but as in electrodynamics, one can also consider ϕ just by itself, without sources. The Lagrangian for a real scalar field is simply

$$\mathcal{L} = \frac{1}{2} \left(\partial_\mu \phi \partial^\mu \phi - m^2 \phi^2 \right). \quad (2.26)$$

The Euler–Lagrange equation of motion that follows from this Lagrangian is (Problem 1.1)

$$\partial_\mu \partial^\mu \phi + m^2 \phi = 0, \quad (2.27)$$

where m is the mass of the field (and of the associated particle). This is as expected: The energy condition of a relativistic particle is $E^2 = m^2 + \mathbf{p}^2$. And in quantum theory, $E = i\partial_0$ and $\mathbf{p} = -i\nabla$, which implies

$$E^2 - \mathbf{p}^2 = -\partial_0^2 + \nabla^2 = -\partial_\mu \partial^\mu. \quad (2.28)$$

Hence the equation of motion is consistent with the relativistic energy condition:

$$(E^2 - \mathbf{p}^2 - m^2)\phi = 0. \quad (2.29)$$

Here, it is clear that $\partial_\mu \partial^\mu$ is the kinetic term ($\sim p^2$), and $m^2 \phi^2$ is a mass term. The Lagrangian has no interaction term: The field ϕ in this case is non-interacting.

2.3 Conserved Currents

Another important concept in field theory is that of a *conserved current*, which is always associated with a conservation law.

Quantum Mechanics. As a warm-up, consider ordinary quantum mechanics. Start with the Schrödinger equation

$$2mi \frac{\partial \psi}{\partial t} + \nabla^2 \psi = 0. \quad (2.30)$$

Multiply the Schrödinger equation by $i\psi^*$, and add the complex conjugate to get

$$-2m\psi^*\partial_t\psi - 2m\psi\partial_t\psi^* + i\psi^*\nabla^2\psi - i\psi\nabla^2\psi^* = 0. \quad (2.31)$$

Dividing by $-2m$, this simplifies to

$$\frac{\partial|\psi|^2}{\partial t} + \nabla\left(-\frac{i}{2m}(\psi^*\nabla\psi - \psi\nabla\psi^*)\right) = 0. \quad (2.32)$$

Now if we identify

$$\rho = |\psi|^2 \quad (2.33)$$

as a density, and

$$\mathbf{J} = -\frac{i}{2m}(\psi^*\nabla\psi - \psi\nabla\psi^*) \quad (2.34)$$

as a current, the above equation becomes the continuity equation

$$\boxed{\frac{\partial\rho}{\partial t} + \nabla\mathbf{J} = 0.} \quad (2.35)$$

To check whether this makes sense, consider a free particle with wave function

$$\psi = C \exp(i\mathbf{p} \cdot \mathbf{x} - i\omega t). \quad (2.36)$$

Then $\rho = |C|^2$ is the probability density, and $\mathbf{J} = \rho\mathbf{p}/m$ is the probability current density. Integrating the continuity equation, and assuming a closed system (such that boundary terms vanish), one finds the conservation law

$$\frac{\partial}{\partial t} \int \rho d^3\mathbf{x} = - \int \nabla\mathbf{J} d^3\mathbf{x} = 0, \quad (2.37)$$

which says that the total probability is preserved.

Complex Scalar. Now consider two free real scalar fields ϕ_1 and ϕ_2 , with identical masses m . The Lagrangian is

$$\mathcal{L} = \frac{1}{2}(\partial_\mu\phi_1\partial^\mu\phi_1 - m^2\phi_1^2) + \frac{1}{2}(\partial_\mu\phi_2\partial^\mu\phi_2 - m^2\phi_2^2). \quad (2.38)$$

Now define the complex field

$$\phi := \frac{\phi_1 + i\phi_2}{\sqrt{2}}, \quad \phi^* = \frac{\phi_1 - i\phi_2}{\sqrt{2}}. \quad (2.39)$$

In terms of this complex field, the Lagrangian reads

$$\mathcal{L} = \partial_\mu\phi^*\partial^\mu\phi - m^2\phi^*\phi, \quad (2.40)$$

which, accordingly, is the Lagrangian of a complex field with mass m .

Global Symmetry. Now comes an important observation: Nothing fixed the particular “direction” of ϕ within the two-dimensional space spanned by ϕ_1 and ϕ_2 . We could as well have used fields that are rotated by some constant angle α :

$$\phi'_1 = \phi_1 \cos \alpha + \phi_2 \sin \alpha, \quad \phi'_2 = -\phi_1 \sin \alpha + \phi_2 \cos \alpha, \quad (2.41)$$

and define the complex field

$$\phi' = \frac{\phi'_1 + i\phi'_2}{\sqrt{2}} = e^{-i\alpha} \phi, \quad \phi'^* = \frac{\phi'_1 - i\phi'_2}{\sqrt{2}} = e^{i\alpha} \phi^*. \quad (2.42)$$

Clearly, $\mathcal{L}$ is not affected by this transformation:

$$\mathcal{L}(\phi, \phi^*) = \mathcal{L}(\phi', \phi'^*). \quad (2.43)$$

This means that the physics is invariant under arbitrary rotations in the two-dimensional field space spanned by ϕ_1 and ϕ_2 .

We can extract very instructive implications from this example. Whenever a system is invariant under some transformation of the coordinates and/or fields, interesting results emerge. Consider a rotation by an infinitesimal angle α . One can then expand the exponentials, and finds

$$\begin{aligned} \phi' &= (1 - i\alpha)\phi = \phi + \delta\phi, & \delta\phi &= -i\alpha\phi, \\ \phi'^* &= (1 + i\alpha)\phi^* = \phi^* + \delta\phi^*, & \delta\phi^* &= +i\alpha\phi^*. \end{aligned} \quad (2.44)$$

The change in $\mathcal{L}$ under such an infinitesimal transformation is of course zero. But we will see that this zero can be written in a very useful way.

Conserved Current. Consider an arbitrary Lagrangian $\mathcal{L}(\phi, \partial_\mu \phi)$ that depends on some field ϕ and its derivative $\partial_\mu \phi$. For any infinitesimal transformation

$$\phi \rightarrow \phi + \delta\phi, \quad \phi^* \rightarrow \phi^* + \delta\phi^*, \quad (2.45)$$

the variation of $\mathcal{L}$ is

$$\begin{aligned} \delta\mathcal{L} &= \delta\phi \frac{\partial\mathcal{L}}{\partial\phi} + \delta(\partial_\mu \phi) \frac{\partial\mathcal{L}}{\partial(\partial_\mu \phi)} + (\phi \rightarrow \phi^*) \\ &= \delta\phi \frac{\partial\mathcal{L}}{\partial\phi} + \partial_\mu \left(\delta\phi \frac{\partial\mathcal{L}}{\partial(\partial_\mu \phi)} \right) - \delta\phi \left(\partial_\mu \frac{\partial\mathcal{L}}{\partial(\partial_\mu \phi)} \right) + (\phi \rightarrow \phi^*). \end{aligned} \quad (2.46)$$

Here we assumed that ϕ is complex. If it is real, the terms with ϕ^* are simply absent. Collecting terms, the above can be re-written as

$$\begin{aligned} \delta\mathcal{L} &= \delta\phi \left(\frac{\partial\mathcal{L}}{\partial\phi} - \partial_\mu \frac{\partial\mathcal{L}}{\partial(\partial_\mu \phi)} \right) + (\phi \rightarrow \phi^*) \\ &\quad + \partial_\mu \left(\delta\phi \frac{\partial\mathcal{L}}{\partial(\partial_\mu \phi)} + \delta\phi^* \frac{\partial\mathcal{L}}{\partial(\partial_\mu \phi^*)} \right). \end{aligned} \quad (2.47)$$

The term in parentheses in the first line is exactly the Euler–Lagrange equation, so the first line is zero by the equations of motion for ϕ (and ϕ^*). Now, if the transformation is a symmetry, that is if the Lagrangian $\mathcal{L}$ is invariant under the transformation, then

$$0 = \delta\mathcal{L} = \partial_\mu \left(\delta\phi \frac{\partial\mathcal{L}}{\partial(\partial_\mu \phi)} + (\phi \rightarrow \phi^*) \right). \quad (2.48)$$

This has the form of a conservation law! We can identify the conserved current

$$-\alpha S^\mu := \delta\phi \frac{\partial \mathcal{L}}{\partial(\partial_\mu \phi)} + (\phi \rightarrow \phi^*), \quad (2.49)$$

where we have factored out a parameter $-\alpha$. Every infinitesimal transformation will come with an infinitesimal (real) parameter, which we call α , as in the example of the complex scalar above. Then $\delta\phi \sim \alpha$, and $\delta\phi^* \sim \alpha$. By construction, the current S^μ satisfies

$$\boxed{\partial_\mu S^\mu = 0.} \quad (2.50)$$

Conserved Charge. What we have shown probably looks familiar: It is Noether's theorem. Invariance under a transformation implies a conserved charge. In components, the invariance reads

$$0 = \partial_\mu S^\mu = \frac{\partial S_0}{\partial t} + \nabla \mathbf{S}. \quad (2.51)$$

Defining the charge

$$Q := \int S_0(x) d^3x, \quad (2.52)$$

one finds

$$\frac{\partial Q}{\partial t} = \int \partial_t S_0(x) d^3x = - \int \nabla \mathbf{S} d^3x. \quad (2.53)$$

The last expression is the flow through the surface of spacetime. Assuming a closed system, this flow is zero, and hence

$$\boxed{\frac{\partial Q}{\partial t} = 0.} \quad (2.54)$$

The charge Q therefore is conserved, and S_0 is the associated charge density. Since Q is conserved, it is a good quantity to identify states ("quantum number"). This result is completely general. Familiar examples of symmetries and associated conserved charges in classical and quantum mechanics are:

Symmetry	$\leftrightarrow$	Conserved Charge	
rotational invariance	$\leftrightarrow$	angular momentum	(2.55)
translational invariance	$\leftrightarrow$	(linear) momentum	
time translation invariance	$\leftrightarrow$	energy	

Complex Scalar Example. Coming back to our example of the free complex scalar field with its two real components, the infinitesimal symmetry transformation was

$$\delta\phi = -i\alpha\phi, \quad (2.56)$$

hence the conserved current is

$$S_\mu = i(\phi \partial_\mu \phi^* - \phi^* \partial_\mu \phi). \quad (2.57)$$

We note that the current changes sign when ϕ is complex conjugated:

$$\phi \leftrightarrow \phi^* \quad \Rightarrow \quad S_\mu \rightarrow -S_\mu. \quad (2.58)$$

In particular, the charge density S_0 and therefore the conserved charge Q changes sign. Splitting the two degrees of freedom of the complex field ϕ into ϕ and ϕ^* , we see that ϕ

and ϕ^* have the same mass, same spin, but opposite charge. The component ϕ^* therefore describes the anti-particle of ϕ (and vice versa). We conclude that a complex scalar field *automatically* describes both particles and the corresponding anti-particles. We will see later that relativistic field theories inevitably include particles and the corresponding anti-particles.

Quantum Anomalies. The analysis of conserved currents so far was classical. When passing to quantum field theory, it can happen that quantum corrections lead to $\partial_\mu S^\mu \neq 0$ even though $\partial_\mu S^\mu = 0$ classically. Such violations of a classical invariance by quantum corrections are called *anomalies*.

Requiring the absence of anomalies can be a guiding principle to determining the “right” quantum theory. In particular, all *gauge anomalies* (quantum violations of gauge symmetry) must vanish. Otherwise the quantum gauge theory contains unphysical negative-norm states. The presence of symmetries (that remain symmetries at the quantum level) in a theory can imply that certain anomalies vanish.

String theory became exciting in 1984 because it was shown to be anomaly-free in ten spacetime dimensions (by Michael Green and John H. Schwarz), and therefore possibly a consistent quantum theory of gravity. No such theory had been found before. This discovery and the subsequent outburst of string theory research is called the “first superstring revolution”.

2.4 Interactions

So far, we have only considered free fields. Now we add interactions. Recall the Lagrangian for a free scalar field ϕ ,

$$\mathcal{L}_{\text{free}} = \frac{1}{2} \left(\partial_\mu \phi \partial^\mu \phi - m^2 \phi^2 \right), \quad (2.59)$$

Consider adding to this a term

$$\mathcal{L}_{\text{int}} = \phi \rho(\mathbf{x}, t), \quad \mathcal{L} = \mathcal{L}_{\text{free}} + \mathcal{L}_{\text{int}}, \quad (2.60)$$

where $\rho(\mathbf{x}, t)$ is some function of the spacetime coordinates. Then, by the Euler–Lagrange equations, the wave equation for ϕ gets a term,

$$\partial_\mu \partial^\mu \phi + m^2 \phi = \rho. \quad (2.61)$$

By analogy with electrodynamics, we can think of ρ as a source for the field ϕ . In electrodynamics, $\nabla \cdot \mathbf{E} = \rho$, and $\mathbf{E} = -\nabla V - \partial_t \mathbf{A}$, so in the static case, $-\nabla^2 V = \rho$.

Point Charge. To understand the effect of the new term, consider a time-independent point source,

$$\rho = g \delta^3(\mathbf{x}). \quad (2.62)$$

Here, $\delta^3(\mathbf{x})$ is the Dirac delta function, and g is the strength (magnitude) of the source. We want to solve the equation for ϕ . Since ρ is not time dependent, we only consider time-independent ϕ . Then (2.61) becomes

$$(-\nabla^2 + m^2)\phi = g \delta^3(\mathbf{x}). \quad (2.63)$$

We can solve this equation by Fourier transform. Recall the Fourier transformation from $\phi(\mathbf{x})$ to $\tilde{\phi}(\mathbf{k})$,

$$\phi(\mathbf{x}) = \frac{1}{(2\pi)^{3/2}} \int d^3k e^{i\mathbf{k}\cdot\mathbf{x}} \tilde{\phi}(\mathbf{k}), \quad \tilde{\phi}(\mathbf{k}) = \frac{1}{(2\pi)^{3/2}} \int d^3x e^{-i\mathbf{k}\cdot\mathbf{x}} \phi(\mathbf{x}).$$

To transform (2.63), we need the Fourier transform of the delta function:

$$\tilde{\delta}^3(\mathbf{k}) = \frac{1}{(2\pi)^{3/2}} \int d^3x e^{-i\mathbf{k}\cdot\mathbf{x}} \delta^3(x) = \frac{1}{(2\pi)^{3/2}} e^{-i\mathbf{k}\cdot 0} = \frac{1}{(2\pi)^{3/2}}. \quad (2.64)$$

The Fourier transform of (2.63) then becomes

$$(\mathbf{k}^2 + m^2) \tilde{\phi}(\mathbf{k}) = \frac{g}{(2\pi)^{3/2}} \Rightarrow \tilde{\phi}(\mathbf{k}) = \frac{1}{(2\pi)^{3/2}} \frac{g}{\mathbf{k}^2 + m^2}. \quad (2.65)$$

Transforming back, we obtain the solution for $\phi(\mathbf{x})$,

$$\phi(\mathbf{x}) = \frac{g}{(2\pi)^3} \int d^3k \frac{e^{i\mathbf{k}\cdot\mathbf{x}}}{\mathbf{k}^2 + m^2}. \quad (2.66)$$

The integral can be done by parametrizing $\mathbf{k}$ in spherical coordinates with the pole axis pointing in the direction of $\mathbf{x}$, such that $\mathbf{k} \cdot \mathbf{x} = kr \cos(\theta)$, with $k = |\mathbf{k}|$ and $r = |\mathbf{x}|$. Then

$$d^3k = k^2 \sin(\theta) dk d\phi d\theta = -k^2 dk d\phi d(\cos \theta) \quad (2.67)$$

and the integral becomes

$$\begin{aligned} \phi(\mathbf{x}) &= \frac{g}{(2\pi)^3} \int_0^\infty \frac{k^2 dk}{k^2 + m^2} \int_0^{2\pi} d\phi \int_{-1}^1 e^{ikr \cos \theta} d(\cos \theta) \\ &= \frac{g}{(2\pi)^3} \frac{2\pi}{i r} \int_0^\infty \frac{1}{k} \frac{k^2 dk}{k^2 + m^2} (e^{ikr} - e^{-ikr}) \\ &= \frac{g}{(2\pi)^2} \frac{1}{i r} \int_{-\infty}^\infty \frac{k dk}{k^2 + m^2} e^{ikr}. \end{aligned} \quad (2.68)$$

The integral can be done by residues: We can close the contour in the upper half plane, where the semicircle at infinity gives no contribution. The contour encloses a single pole at $k = im$. The integral then becomes

$$\phi(\mathbf{x}) = \frac{g}{(2\pi)^2} \frac{1}{i r} 2\pi i \frac{im}{2im} e^{-mr} = \frac{g}{4\pi} \frac{e^{-mr}}{r}. \quad (2.69)$$

[To read off the residue, write the denominator as $k^2 + m^2 = (k - im)(k + im)$.] Because of the exponential factor, for particles of mass m , the field is significant in a range of $r \sim 1/m$ around the source.

General static source. Since the equation (2.61) for ϕ is linear, the point-source solution (2.69) directly generalizes to a general time-independent source field $\rho_1(\mathbf{x})$. Namely,

$$\phi(\mathbf{x}) = \frac{1}{4\pi} \int d^3x' \rho_1(\mathbf{x}') \frac{e^{-m|\mathbf{x}-\mathbf{x}'|}}{|\mathbf{x}-\mathbf{x}'|}. \quad (2.70)$$

Interaction. Let us see how one source would interact with another through its force field. The interaction Hamiltonian between the force field ϕ and another source $\rho_2(\mathbf{x})$ is

$$H = - \int d^3x \phi(\mathbf{x}) \rho_2(\mathbf{x}) . \quad (2.71)$$

Plugging in the general solution (2.70), we find

$$H_{12} = - \frac{1}{4\pi} \int d^3x d^3x' \rho_1(\mathbf{x}) \rho_2(\mathbf{x}') \frac{e^{-m|\mathbf{x}-\mathbf{x}'|}}{|\mathbf{x}-\mathbf{x}'|} . \quad (2.72)$$

The interaction between two sources (charged particles) is therefore given by the potential

$$V(r) = - \frac{1}{4\pi} \frac{e^{-mr}}{r} . \quad (2.73)$$

Due to the exponential falloff, the interaction range is $\sim 1/m$ (in natural units), where m is the mass of ϕ .

Interpretation. The general interpretation is the following: The field ϕ acts as a force carrier that is sourced by the particles it interacts with (ρ in the interaction term), just as the electromagnetic field acts on and is sourced by electrically charged particles. In quantum theory, the energy transferred by the field ϕ is quantized, and the energy quanta form the (force) particles. This is the general quantum field theory picture: All interactions are due to the exchange of field quanta. The concepts of force and interaction are used interchangeably.

Remark: Propagator. In the equation for H_{12} (2.72), the interaction is written in position space (with coordinates x). In particle physics, matrix elements (transition amplitudes) are mostly written in momentum space. Looking at the expression (2.66) for the point-sourced field ϕ , we see that the denominator in momentum space is $\mathbf{k}^2 + m^2$. Had we included a possible time-dependence of ϕ , the denominator would be $-k_0^2 + \mathbf{k}^2 + m^2 = m^2 - k^2$, where $k^2 = k_\mu k^\mu$. This factor is very general: Whenever a particle of mass m and four-momentum k is exchanged in an interaction, this exchange is represented by a factor

$$1/(k^2 - m^2) . \quad (2.74)$$

This is called a *propagator*, and appears frequently in matrix elements. The full propagator also has a phase factor and a numerator that depends on the spin of the exchanged particle.

Meson theory. The interaction studied above is a model for the strong force among nucleons, proposed by Hideki Yukawa in 1934. Yukawa's theory predicts the existence of a new particle, the quantum of the force field ϕ , which is sourced by nucleons. He called these particles *mesons*, and the field ϕ the *meson field*, from the greek word “mesos”, which means “intermediate”. The reason was that the predicted mass of these mesons was between that of the electron and the proton. 13 years later, in 1947, mesons were indeed discovered experimentally, as short-lived subatomic particles that are produced in nucleon collisions. For his prediction, Yukawa was awarded the Nobel Prize in 1949. Today, we know that mesons are bound states of one quark and one antiquark, and there are various different kinds (pions, rho mesons, eta mesons, Kaons, etc). They are indeed the primary force carriers that hold atomic nuclei together.

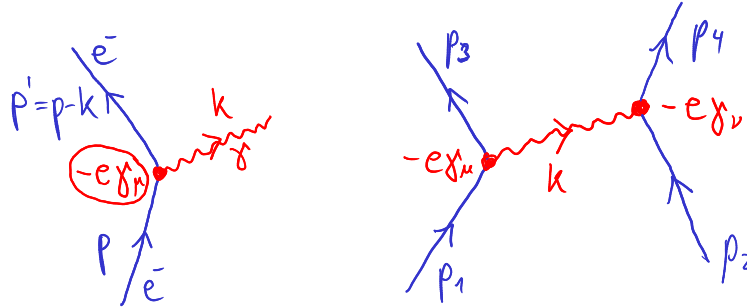
Feynman rules. For understanding the Standard Model (and quantum field / gauge theory in general), and especially for computing its predictions, we will need the Feynman rules for the Standard Model. These are used to compute matrix elements (transition amplitudes). To understand the general structure of the model, the interactions between fermions and bosons are the most important. Let us summarize the basic arguments, all of which were motivated in the discussion so far.

Consider the electromagnetic interaction. The interaction term in the Lagrangian is

$$\mathcal{L}_{\text{int}} = -J_\mu A^\mu = Q\bar{\psi}\gamma_\mu\psi A^\mu. \quad (2.75)$$

Here, Q is the electric charge (for electrons, $Q = -e$), $\bar{\psi}$ and ψ are final and initial electron states, and γ_μ is a spin factor that makes the combination $\bar{\psi}\gamma_\mu\psi$ a relativistically covariant four-vector. A^μ is the electromagnetic vector potential (photon field).

Suppose we want to describe an interaction where an electron with momentum p emits a photon of momentum k , thereby changing its momentum to $p' = p - k$.



The factor at the electron-electron-photon ($ee\gamma$) vertex is the interaction term in the Lagrangian with the initial and final states (wave functions) removed. In this case, the factor that remains is $-e\gamma_\mu$. To first approximation, the matrix element for any process can be constructed by

- writing the appropriate interaction factor for each vertex,
- putting a propagator factor $1/(k^2 - m^2)$ for any internal particle line of four-momentum k and mass m ,
- impose momentum conservation at every vertex.

The precise rules for arbitrary processes are a bit more complicated, but a good semi-quantitative understanding of the Standard Model and its tests and predictions can be obtained with these approximate rules. This is basically the Born approximation for transition amplitudes M in quantum theory: $M \simeq \langle f|V|i \rangle$ for initial state i , potential V , and final state f .

3 Gauge Invariance

Gauge invariance is one of the most important concepts of the Standard Model of particle physics. It is a symmetry principle that determines from the beginning how particles have to interact to make the theory consistent. For example, it explains why the electromagnetic interaction is due to a massless spin-one particle, the photon, that is being exchanged between electrically charged particles. More generally, if given types of matter particles exist and are to interact, gauge invariance predicts the existence of further particles that

mediate this interaction (by particle exchange), and it determines the properties of these particles as well as the precise form of the interaction terms in the Lagrangian.

This is very different from the historical situation where forces and interactions had to be postulated in a clever way to match experimental observations. Theories where the interactions are determined by gauge invariance are called *gauge theories* or *Yang–Mills theories* (after Chen Ning Yang and Robert Mills), and the force-carrying particles (quanta of the interaction field) are called *gauge bosons*. The theories for the electroweak and the strong interactions are of this type. The existence of the respective gauge bosons ($W^\pm$, Z^0 bosons and gluons) was predicted by gauge theory, and indeed they were all found experimentally. All their measured properties agree with the theoretical predictions.

We will first look at gauge invariance in classical electrodynamics, and then in quantum theory. Afterwards, we move on to Abelian gauge theory. Finally, we will look at non-Abelian gauge theory, which is the case of interest for the Standard Model.

3.1 Gauge Invariance in Electrodynamics

In classical electrodynamics, the electric and magnetic fields $\mathbf{E}$ and $\mathbf{B}$ are expressed in terms of the vector potential $\mathbf{A}$ and scalar potential V as

$$\mathbf{B} = \nabla \times \mathbf{A}, \quad \mathbf{E} = -\nabla V - \partial \mathbf{A} / \partial t. \quad (3.1)$$

If we transform the potentials $\mathbf{A}$ and V simultaneously as

$$\mathbf{A} \rightarrow \mathbf{A}' = \mathbf{A} + \nabla \chi, \quad V \rightarrow V' = V - \partial \chi / \partial t, \quad (3.2)$$

where χ is an arbitrary (differentiable) scalar, then the equations (3.1) for the fields are unchanged, or *invariant*. If we combine $\mathbf{A}$ and V into the four-vector $A^\mu = (V; \mathbf{A})$, then the transformation (3.2) can be written in the uniform way

$$A^\mu \rightarrow A'^\mu = A^\mu - \partial^\mu \chi. \quad (3.3)$$

These transformations are called *gauge transformations*. They have been known since the 1800s, but were largely viewed as a curiosity for a long time.

The uniform notation (3.3) emphasizes the correlation between the transformations for $\mathbf{A}$ and V . Turning the argument around, one could say that the electric and magnetic fields have to be related in the very specific way (3.1) in order to be invariant under gauge transformations. This is closer to the point of view we will adopt in the following.

3.2 Gauge Invariance in Quantum Mechanics

In quantum theory, gauge invariance takes a different form, which leads to the modern viewpoint. Observable quantities depend on wave functions ψ through probabilities $|\psi|^2$. Hence it is reasonable to demand that the theory is invariant under an overall phase change of the wavefunction,

$$\psi \rightarrow \psi' = e^{-i\alpha} \psi, \quad (3.4)$$

where α is a constant. This is called a *global gauge transformation* (since the change of phase is the same everywhere in space).

One would imagine that it should also be possible to change the phase of ψ differently at different points in space and time without affecting the theory (since the probabilities $|\psi|^2$ would still be unchanged by this). That is, the theory should be invariant under

$$\psi(\mathbf{x}, t) \rightarrow \psi'(\mathbf{x}, t) = e^{-i\chi(\mathbf{x}, t)} \psi(\mathbf{x}, t). \quad (3.5)$$

This is called a *local gauge transformation*.

Now consider the Schrödinger equation for a free particle described by the wavefunction ψ ,

$$i \frac{\partial}{\partial t} \psi(\mathbf{x}, t) = -\frac{1}{2m} \nabla^2 \psi(\mathbf{x}, t). \quad (3.6)$$

If ψ satisfies the Schrödinger equation, then the locally transformed ψ' for a general $\chi(\mathbf{x}, t)$ will *not* satisfy it, since the derivatives of χ do not cancel! The Schrödinger equation is therefore not invariant under local gauge transformations.

For electrically charged particles in the presence of an electromagnetic field, we know that the Schrödinger equation is modified to

$$(i\partial_t + eV)\psi = \frac{1}{2m} (-i\nabla + e\mathbf{A})^2 \psi, \quad (3.7)$$

where e is the absolute value of the charge of the electron. But if we now apply the local gauge transformations for the wave function and for the electromagnetic potentials *simultaneously*,

$$\begin{aligned} \psi(\mathbf{x}, t) &\rightarrow \psi'(\mathbf{x}, t) = e^{-i\chi(\mathbf{x}, t)} \psi(\mathbf{x}, t), \\ \mathbf{A}(\mathbf{x}, t) &\rightarrow \mathbf{A}'(\mathbf{x}, t) = \mathbf{A}(\mathbf{x}, t) + \frac{1}{e} \nabla \chi(\mathbf{x}, t), \\ V(\mathbf{x}, t) &\rightarrow V'(\mathbf{x}, t) = V(\mathbf{x}, t) - \frac{1}{e} \partial_t \chi(\mathbf{x}, t), \end{aligned} \quad (3.8)$$

then the modified Schrödinger equation (3.7) is unchanged, that is *invariant* under this transformation.

What happened? At first, the Schrödinger equation was (unexpectedly) not invariant under local gauge transformations, but once we add the electromagnetic field and properly transform it along with ψ , invariance under local gauge transformations is achieved. We can turn the logic around and say that *requiring* local gauge invariance *enforces* both the *presence* of the field $A^\mu = (V, \mathbf{A})$, and the *very specific way* it enters the Schrödinger equation. Since A^μ is a four-vector, it is a *vector field*, and its quanta (the photons) are *vector particles*, that is particles of spin one. Note that we made no use of any particular properties of ψ except that it satisfies the Schrödinger equation. Hence the effect applies to all particles in the same way: Whenever a particle is electrically charged, its interaction with photons is fixed by gauge invariance. The only quantity that remains undetermined is the charge (e in the above equations), it remains an input to the theory that has to be measured experimentally.

To summarize, the existence and specific form of electromagnetic interactions has in a certain sense been derived from requiring invariance under local gauge transformations. As we shall see, this principle generalizes to other interactions. Whenever a particle carries a certain charge, and the theory is invariant under certain “phase” transformations (generally called *gauge transformations*), then specific interaction fields of spin one (called *gauge fields*) must exist with their associated quanta (called *gauge bosons*), and the interaction Lagrangian is fixed, up to the numerical values of the charges.

3.3 Covariant Derivatives

We can rewrite the Schrödinger equation for a particle in an electromagnetic field in a form that makes its invariance under local gauge transformations very transparent. This form of the equation will generalize to other gauge theories. Define the differential operators

$$\mathbf{D} = -\nabla - ie\mathbf{A} \quad \text{and} \quad D^0 = \partial_t - ieV. \quad (3.9)$$

Then the Schrödinger equation (3.7) becomes

$$iD^0\psi = \frac{1}{2m} (i\mathbf{D})^2\psi. \quad (3.10)$$

Now let us see how $\mathbf{D}\psi$ changes when we apply a local gauge transformation:

$$\begin{aligned} \mathbf{D}'\psi' &= (-\nabla - ie\mathbf{A}')\psi' = (-\nabla - ie\mathbf{A} - i(\nabla\chi))e^{-i\chi}\psi \\ &= e^{-i\chi}(-\nabla - ie\mathbf{A})\psi = e^{-i\chi}(\mathbf{D}\psi). \end{aligned} \quad (3.11)$$

Similarly,

$$\begin{aligned} D'^0\psi' &= (\partial_t - ieV')\psi' = (\partial_t - ieV + i(\partial_t\chi))e^{-i\chi}\psi \\ &= e^{-i\chi}(\partial_t - ieV)\psi = e^{-i\chi}(D^0\psi). \end{aligned} \quad (3.12)$$

We see that both $\mathbf{D}\psi$ and $D^0\psi$ transform in the same way as the wavefunction ψ itself. The four-vector

$$D^\mu = (D^0; \mathbf{D}) \quad (3.13)$$

is called the *covariant derivative*. Remarkably, any equation for ψ written in terms of the covariant derivative will automatically be gauge invariant! When D^μ gets applied repeatedly, the result will still transform as a wave function: Since $D^\mu\psi$ transforms like a wave function, also $D_\mu D^\mu\psi$ does, etc.

This principle is very general: We can apply it for any interaction that is due to a “charge”, not only the electromagnetic one. Let us say we want our theory to be invariant under some transformation

$$\psi \rightarrow \psi' = U\psi, \quad (3.14)$$

where U is some operator acting on the wavefunction ψ . We want to define an operator

$$D^\mu = \partial^\mu - igA^\mu, \quad (3.15)$$

where g is a constant, and A^μ represents the interacting field that has to be added to make the theory invariant, but we do not know how A^μ transforms. What we want to impose is that $D^\mu\psi$ transforms in the same way as ψ , that is

$$D'^\mu\psi' = U(D^\mu\psi) \Leftrightarrow (\partial^\mu - igA'^\mu)U\psi = U(\partial^\mu - igA^\mu)\psi. \quad (3.16)$$

Solving this equation for A'^μ , we get

$$\begin{aligned} -igA'^\mu U\psi &= -\partial^\mu(U\psi) + U\partial^\mu\psi - igUA^\mu\psi \\ &= -(\partial^\mu U)\psi - igUA^\mu\psi. \end{aligned} \quad (3.17)$$

Since we want this equality to be true for any state ψ , it has to hold at the operator level, that is we can drop the ψ . Multiplying by U^{-1} from the right, one finds

$$A'^\mu = UA^\mu U^{-1} - \frac{i}{g}(\partial^\mu U)U^{-1}. \quad (3.18)$$

This is the transformation rule that A^μ has to follow for $D^\mu = \partial^\mu - igA^\mu$ to be a covariant derivative. Here, U was an arbitrary transformation operator acting on the state ψ . If U is a phase, $U = e^{-i\chi}$, then $UA^\mu U^{-1} = A^\mu$, and we recover the known result for the electrodynamic field, with $g = -e$. In general, g is the charge that couples to the gauge field A^μ . U could be (and will be) a matrix operator acting on an internal state space. Then also A^μ will be a matrix operator, and the ordering of factors in (3.18) matters.

4 Some Group Theory

Before we get to the non-Abelian gauge theories of the Standard Model, let us review some group theory. This will serve two purposes: Firstly, all matter particles are spin 1/2 fermions, so we will need the theory of spinors later on. Secondly, the electroweak gauge group is $SU(2) \times U(1)$, and the gauge group of QCD (the strong interactions) is $SU(3)$, which is a generalization of $SU(2)$. So we will need some basics about the representations of these groups to understand the Standard Model.

This review will be mostly mathematical and a bit technical. Don't be intimidated by this! When we continue with the physics part later on, we will for the most part only need the simplest examples of the concepts introduced in the following.

4.1 Groups

Groups. Let us recall the basic definition of a *group*. A group $(G, \cdot)$ is a set G of elements and a composition rule

$$\cdot : G \times G \rightarrow G, \quad (g, h) \mapsto g \cdot h = gh \quad (4.1)$$

which satisfies

- If $g, h \in G$, then also $g \cdot h \in G$.
- The composition is associative: $(g \cdot h) \cdot u = g \cdot (h \cdot u)$ for all $g, h, u \in G$.
- There is an identity element $1 \in G$, such that $g \cdot 1 = 1 \cdot g = g$ for all $g \in G$.
- For every $g \in G$, there is a unique inverse $g^{-1} \in G$, such that $g \cdot g^{-1} = g^{-1} \cdot g = 1$.

An example is the group of permutations of n elements, called S_n (a discrete group). Another is the set of complex phase factors $U(\theta) = e^{i\theta}$ for real θ , called $U(1)$ (a continuous group).

Groups as Transformations. In physics in general, and in particle physics in particular, groups appear mainly because

- Physical quantities behave in a specific way under some transformations,
- Physical systems may be invariant under certain transformations of (some or all) their constituents.

Some examples are

- Vectors in classical mechanics and electrodynamics transform in a specific (the well-known) way under rotations, represented by the rotation group $SO(3)$: $v \mapsto Rv$, $R \in SO(3)$.
- Quantum states are invariant under multiplication by a phase factor: $\psi \mapsto e^{i\theta}\psi$, all phase factors form the group $U(1) \ni e^{i\theta}$.
- Spinors transform in representations of the spin group $SU(2)$.

In all these examples, the groups are continuous, which means that the elements depend on continuous parameters that can take infinitely many different values. Discrete groups also play a role in physics (for example rotations of lattices by discrete angles), but will not be important for us.

4.2 Lie Groups and Algebras

Lie Groups. Continuous groups whose elements are differentiable functions of their continuous parameters (such that the group is at the same time a differentiable manifold) are called *Lie groups*. This is the case for all continuous transformation groups that we consider. The simplest example is $U(1)$, where $U(\theta) = e^{i\theta} \in U(1)$ is clearly a differentiable function of θ .

Exponential Map and Lie Algebras. It can be shown that for all Lie groups G , every element $g \in G$ can be written in the form

$$g = \exp(iX), \quad X = \sum_{k=1}^n \theta_k T_k. \quad (4.2)$$

The quantities T_i are called the *generators* of the Lie group, and for a Lie group that depends on n parameters (an n -dimensional Lie group), there are n such generators. The generators form an algebra called the *Lie algebra* of the Lie group. Expanding the exponential map

$$g = \exp(iX) = 1 + i \sum_{k=1}^n \theta_k T_k + \mathcal{O}(\theta_k^2), \quad (4.3)$$

one sees that the generators T_k define group elements $1 + i\epsilon T_k$ that are infinitesimally close to the identity $1 \in G$.

More on Lie Algebras. The Lie algebra $\mathfrak{g}$ is the tangent space of the Lie group G at the identity element $1 \in G$. The *exponential map* $\exp : \mathfrak{g} \rightarrow G$ maps $iX \in \mathfrak{g}$ to the integral curve in G whose derivative at $t = 0$ equals iX , that is $\partial/\partial t \exp(itX)|_{t=0} = iX$. When the group is represented by matrices, then the exponential map is equivalently defined by its Taylor series expansion:

$$\exp(iX) = \sum_{n=0}^{\infty} \frac{1}{n!} (iX)^n. \quad (4.4)$$

A Lie algebra $\mathfrak{g}$ can be defined abstractly, and independently of its corresponding Lie group, as a vector space with a bilinear product $[\cdot, \cdot] : \mathfrak{g} \times \mathfrak{g} \rightarrow \mathfrak{g}$ that obeys the two properties

$$\begin{aligned} \text{Antisymmetry: } & [X, Y] = -[Y, X], \\ \text{Jacobi identity: } & [X, [Y, Z]] + [Y, [Z, X]] + [Z, [X, Y]] = 0, \end{aligned} \quad (4.5)$$

for all $X, Y, Z \in \mathfrak{g}$. The basis elements of the vector space are called *generators*, and are typically denoted by T_a . The product is fully defined by the *structure constants* f_{ab}^c via

$$[T_a, T_b] = i f_{ab}^c T_c, \quad (4.6)$$

where $\{T_a\}$ is a basis of generators.

Example. Let us illustrate these concepts with an example. Arguably the most important transformation group for quantum theory is $SU(2)$, defined by

$$SU(2) = \{U \in \text{Mat}(2, \mathbb{C}) \mid UU^\dagger = U^\dagger U = 1, \det(U) = 1\}. \quad (4.7)$$

A general element of this group can be written as

$$U = \begin{pmatrix} \alpha & -\bar{\beta} \\ \beta & \bar{\alpha} \end{pmatrix}, \quad |\alpha|^2 + |\beta|^2 = 1. \quad (4.8)$$

This shows that $SU(2)$ is clearly a Lie group. One verifies easily that *infinitesimal group elements* $1 + i\varepsilon T_k$ are generated by the three matrices

$$T_1 = \frac{1}{2} \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}, \quad T_2 = \frac{1}{2} \begin{pmatrix} 0 & -i \\ i & 0 \end{pmatrix}, \quad T_3 = \frac{1}{2} \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}. \quad (4.9)$$

These are the *generators* of the Lie algebra $\mathfrak{su}(2)$, and are recognized as the familiar *Pauli matrices*, $T_k = \sigma_k/2$. The factor of $1/2$ is a convention that is chosen such that $\exp(2\pi T_k) = -1$.

The generators satisfy the commutation relations of angular momentum:

$$[T_a, T_b] = i\varepsilon_{abc} T_c, \quad (4.10)$$

so the *structure constants* are $f_{ab}^c = f_{abc} = \varepsilon_{abc}$ (the totally antisymmetric tensor). One can verify that the commutator (4.10) satisfies the Jacobi identity.

4.3 Representations

Representations. (Lie) groups and algebras are abstract objects that can be defined purely in terms of their algebraic properties. In practice, they take concrete realizations in terms of matrices. Such realizations are called *representations*. A *representation of a Lie group* G is a map R from G to the space of $n \times n$ complex matrices:

$$R : G \rightarrow \text{Mat}(n, \mathbb{C}), \quad R(g)R(h) = R(gh). \quad (4.11)$$

The condition on the right must hold for all elements $g, h \in G$. It says that R must be a *homomorphism* that preserves the product structure of the group.

Similarly, a *representation of a Lie algebra* $\mathfrak{g}$ is defined as a map ρ from $\mathfrak{g}$ to the space of $n \times n$ complex matrices that preserve the algebra product (commutator):

$$\rho : \mathfrak{g} \rightarrow \text{Mat}(n, \mathbb{C}), \quad [\rho(X), \rho(Y)] = \rho([X, Y]), \quad (4.12)$$

for all $X, Y \in \mathfrak{g}$.

The representations of a Lie group G and its Lie algebra $\mathfrak{g}$ are related in the obvious way: Given a representation R of G , the tangent space of $R(G)$ at the identity defines a representation ρ of the corresponding Lie algebra. Similarly, a representation ρ of $\mathfrak{g}$ induces a representation R of G by exponentiation. In other words:

$$\exp(iX) = g \quad \Leftrightarrow \quad \exp(i\rho(X)) = R(g). \quad (4.13)$$

Fundamental and Adjoint Representations. Because representations are bijective homomorphisms, every group can be *defined* by giving one of its representations. This is typically the case for Lie groups. For example, the group $SU(n)$ can be defined as the space of unitary $n \times n$ complex matrices with unit determinant, as we did for $SU(2)$ above. For $SU(n)$, this realization is called the *fundamental representation*. As we saw above, the generators of $\mathfrak{su}(2)$ in the fundamental representation are the Pauli matrices: $T_a = \sigma_a/2$.

Every Lie algebra $\mathfrak{g}$ with n generators is also an n -dimensional vector space. We can construct a representation of $\mathfrak{g}$ in terms of $n \times n$ matrices $\rho(X)$ that act on the vector space $\mathfrak{g}$ by the commutator $[X, \cdot]$ with X :

$$\rho(X) : \mathfrak{g} \rightarrow \mathfrak{g}, \quad Y \mapsto \rho(X) \cdot Y \equiv [X, Y]. \quad (4.14)$$

This is called the *adjoint representation*. In terms of the generators, it becomes

$$\rho(T_a) \cdot T_b = [T_a, T_b] = if_{ab}^c T_c. \quad (4.15)$$

In the generator basis, the generators vectors have entries $(T_a)_i = \delta_{ai}$, so the above equation becomes

$$[\rho(T_a)]^i_b = [\rho(T_a)]^i_j (T_b)^j = if_{ab}^c (T_c)^i = if_{ab}^i. \quad (4.16)$$

Writing everything in lower indices, and using the fact that f_{abc} is totally antisymmetric (this is true for all Lie algebras we consider), one finds

$$[T_a^{\text{adj}}]_{bc} \equiv [\rho(T_a)]_{bc} = -if_{abc}. \quad (4.17)$$

At the level of the *group*, the adjoint representation acts by conjugation:

$$R(g) : G \rightarrow G, \quad h \mapsto R(g)(h) \equiv ghg^{-1}. \quad (4.18)$$

Example: $\mathfrak{su}(2)$. To have a concrete example, let us look again at $\mathfrak{su}(2)$. As we saw above, the *fundamental* generators are $T_i^{\text{fund}} = \sigma_i/2$, where σ_i are the Pauli matrices. They satisfy the commutation relations

$$[T_a, T_b] = i\varepsilon_{abc} T_c, \quad (4.19)$$

and hence the *adjoint* generators T_a^{adj} have matrix elements

$$[T_a^{\text{adj}}] = -i\varepsilon_{abc}. \quad (4.20)$$

Let us check this explicitly:

$$[T_1, T_2] = iT_3, \quad [T_1, T_3] = -iT_2, \quad (4.21)$$

and similar relations hold for T_2 and T_3 . The adjoint generators are therefore:

$$T_1^{\text{adj}} = \begin{pmatrix} 0 & 0 & 0 \\ 0 & 0 & -i \\ 0 & i & 0 \end{pmatrix}, \quad T_2^{\text{adj}} = \begin{pmatrix} 0 & 0 & i \\ 0 & 0 & 0 \\ -i & 0 & 0 \end{pmatrix}, \quad T_3^{\text{adj}} = \begin{pmatrix} 0 & -i & 0 \\ i & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix}. \quad (4.22)$$

Example: $\mathfrak{su}(3)$. Another example that is important for the Standard Model is $\text{SU}(3)$, the group of 3×3 unitary matrices. For completeness and reference, we list some of its properties here. The Lie algebra $\mathfrak{su}(3)$ of this group has eight generators $T_a = \lambda_a/2$ with

$$\begin{aligned} \lambda_1 &= \begin{pmatrix} 0 & 1 & 0 \\ 1 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix}, & \lambda_2 &= \begin{pmatrix} 0 & -i & 0 \\ i & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix}, & \lambda_3 &= \begin{pmatrix} 1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & 0 \end{pmatrix}, \\ \lambda_4 &= \begin{pmatrix} 0 & 0 & 1 \\ 0 & 0 & 0 \\ 1 & 0 & 0 \end{pmatrix}, & \lambda_5 &= \begin{pmatrix} 0 & 0 & -i \\ 0 & 0 & 0 \\ i & 0 & 0 \end{pmatrix}, & \lambda_6 &= \begin{pmatrix} 0 & 0 & 0 \\ 0 & 0 & 1 \\ 0 & 1 & 0 \end{pmatrix}, \end{aligned}$$

$$\lambda_7 = \begin{pmatrix} 0 & 0 & 0 \\ 0 & 0 & -i \\ 0 & i & 0 \end{pmatrix}, \quad \lambda_8 = \frac{1}{\sqrt{3}} \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & -2 \end{pmatrix}. \quad (4.23)$$

Its structure constants, defined by

$$[\lambda_a, \lambda_b] = 2i f_{abc} \lambda_c \quad (4.24)$$

are

$$\begin{aligned} f_{123} &= 1, \\ f_{458} &= f_{678} = \sqrt{3}/2, \\ f_{147} &= f_{516} = f_{246} = f_{257} = f_{345} = f_{637} = 1/2, \end{aligned} \quad (4.25)$$

and all other values f_{abc} follow from the total antisymmetry.

4.4 Quantum Theory

Unitary Representations. In quantum theory, states are represented by elements of a Hilbert space. All (measurable) *observables* are expressed in terms of projections $\langle \phi | \psi \rangle$ of states onto other states, that is in terms of the *inner product* $\langle . | . \rangle$ on the Hilbert space. Any transformation $|\psi\rangle \mapsto |\psi^U\rangle = U|\psi\rangle$ that is supposed to be a *symmetry* of the quantum system therefore has to *preserve* the inner product. Because a general projection transforms as

$$\langle \phi | \psi \rangle \mapsto \langle \phi^U | \psi^U \rangle = \langle \phi | U^\dagger U | \psi \rangle, \quad (4.26)$$

this means that any symmetry U has to obey $U^\dagger U = 1$, or in other words $U^\dagger = U^{-1}$, which means that the transformation must be *unitary*. In particular, this implies that every *symmetry group* must be represented on quantum states by a *unitary representation* (a representation whose matrices are unitary).

Internal Rotations. The states may depend on continuous parameters (such as the position in space, or the momentum), but may also have discrete degrees of freedom. In such cases, a general state is represented as a vector, whose entries are the amplitudes for the various possible discrete values. A familiar example is spin: We can distinguish spin-up and spin-down states, and a general state is a superposition of the two:

$$|\uparrow\rangle = \begin{pmatrix} 1 \\ 0 \end{pmatrix}, \quad |\downarrow\rangle = \begin{pmatrix} 0 \\ 1 \end{pmatrix}, \quad |\psi\rangle = \begin{pmatrix} a \\ b \end{pmatrix}. \quad (4.27)$$

Consider more generally a quantum system that can be in one of two discrete states. The two states span a (complex) two-dimensional state space, and a general state is represented as a two-dimensional vector in this space, as in (4.27). Assume further that the system is *symmetric under rotations* in this space. By the above argument, any symmetry must be represented by unitary matrices. In the case of a two-state system, the relevant matrices are unitary 2×2 matrices, and such matrices form the group $SU(2)$. If a system has three discrete states, the relevant group would be $SU(3)$.

4.5 $SO(3)$ and $SU(2)$ and Spin

We will close this review of group theory by recalling the relevance of $SU(2)$ for spin and rotations in quantum systems.

Rotations. Rotations in space are represented by matrices $R \in \text{SO}(3)$, where $\text{SO}(3)$ is the group of orthogonal matrices. As we said above, any quantum system must transform in a unitary representation U of $\text{SO}(3)$ under spatial rotations R , such that $|\psi\rangle \mapsto U(R)|\psi\rangle$. Any rotation by an angle of 2π equals the identity operation, $R(2\pi) = 1$. Using the property

$$U(R)U(R') = U(RR') \quad \Rightarrow \quad U(R)U(1) = U(R \cdot 1) = U(R), \quad (4.28)$$

one finds that also $U(R(2\pi)) = U(1) = 1$, therefore every state will be mapped to back to itself under full 2π rotations.

And of course any reasonable *quantum system must* be invariant under rotations by an angle of 2π . But *quantum states* $|\psi\rangle$ are only well-defined up to arbitrary complex prefactors, so one could also allow transformations that map states to themselves *up to a constant prefactor*.

SU(2) and Spin. This is where $\text{SU}(2)$ comes into the game: $\text{SU}(2)$ is a *double cover* of $\text{SO}(3)$, which means that there is a mapping (a group homomorphism) $f : \text{SU}(2) \rightarrow \text{SO}(3)$ that is two-to-one: If $f(U) = R$, then also $f(-U) = R$. In particular, both $1 \in \text{SU}(2)$ and $-1 \in \text{SU}(2)$ are identified with $1 \in \text{SO}(3)$. And because the mapping is continuous, a full rotation in space is represented by $-1 \in \text{SU}(2)$.

So a state $|\psi\rangle$ transforming under $\text{SU}(2)$ under rotations gets mapped to $-|\psi\rangle$ under rotations by a full angle 2π . And this is admissible, since $|\psi\rangle$ and $-|\psi\rangle$ represent the same state. The upshot is that to represent spatial rotations on quantum states, we can also consider unitary representations of $\text{SU}(2)$!

And this is indeed what happens for spin.

Spin 1/2. A *spin-1/2 system* is represented by a two-component vector $|\psi\rangle$ that transforms in the fundamental representation of $\text{SU}(2)$ under spatial rotations. The fundamental generators for $\text{SU}(2)$ are the Pauli matrices, $T_i = \sigma_i/2$. So for a rotation by α :

$$|\psi\rangle \mapsto \exp(i\alpha \cdot \sigma/2) |\psi\rangle. \quad (4.29)$$

Spin One. For a *spin-one system*, states are represented by three-component vectors $|\psi\rangle$ that transform in a three-dimensional representation of $\text{SU}(2)$. That three-dimensional representation is the adjoint representation of $\text{SU}(2)$, whose generators we stated above in (4.22). So for a rotation by α :

$$|\psi\rangle \mapsto \exp(i\alpha \cdot T^{\text{adj}}) |\psi\rangle \quad (4.30)$$

Upon closer inspection, the generators T_i^{adj} are exactly the fundamental generators of ordinary $\text{SO}(3)$ rotations! So the adjoint representation of $\text{SU}(2)$ equals the fundamental representation of $\text{SO}(3)$, and spin-one states transform as regular vectors under spatial rotations.

5 Non-Abelian Gauge Theory

After having gone through a review of group theory, we are ready to move on to non-Abelian gauge theory.

5.1 Strong Isospin

Strong isospin

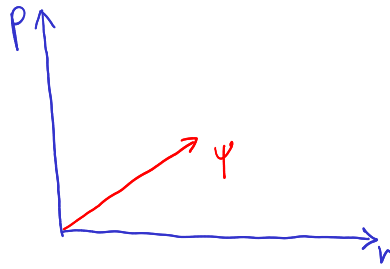
- is an approximate symmetry,
- had an important conceptual impact historically.
- There is also *weak isospin*, which is a more fundamental symmetry. Here, we consider *strong isospin* to get familiar with the concept.

Nucleons. Strong isospin is a symmetry between the two types of nucleons: Protons and neutrons. They have almost identical masses:

$$m_{\text{proton}} = 939.57 \text{ MeV} , \quad m_{\text{neutron}} = 938.27 \text{ MeV} . \quad (5.1)$$

The difference in masses is only $\sim 0.1\%$. The proton carries electric charge, the neutron is electrically neutral. Apart from that, they are very similar particles, especially from the point of view of the strong interaction. At nuclear scales, the strength of their EM interaction is only about $\sim 1\%$ of their strong interaction.

Consider the proton p and the neutron n as two states of the same object, the *nucleon* N . Imagine a 2d state space, the *strong isospin space*:



- Assume that forces that describe nucleon interactions (the strong force) are invariant under rotations in this space. This can only be approximately true due to the EM interactions, but we will neglect those for now.
- The EM charge is merely a label that distinguishes the p from the n state.
- Similar to spin, write general states as

$$N = \begin{pmatrix} p \\ n \end{pmatrix} \Rightarrow \text{Probabilities: } P_p = \frac{|p|^2}{|N|^2} = \frac{pp^*}{pp^* + nn^*} . \quad (5.2)$$

For properly normalized states: $pp^* + nn^* = |N|^2 = 1$.

- Invariance under rotations in 2d complex strong isospin space is equivalent to invariance under

$$N \mapsto UN , \quad U \in \text{SU}(2) . \quad (5.3)$$

Such unitary rotations preserve total probability (normalization). In other words, N forms an $\text{SU}(2)$ *doublet* (fundamental representation of $\text{SU}(2)$).

- The standard basis of SU(2) is given by the generators $\tau_i/2$, where τ_i are the Pauli matrices. We will call the Pauli matrices σ_i when they act on spinors, and τ_i when they act on some other space, like here.

The generator τ_3 is diagonal, hence the states p and n have definite eigenvalues $\pm 1/2$. Analogous to regular spin, we say that p has (strong) isospin $+1/2$, and n has (strong) isospin $-1/2$.

Pions. Besides p and n , do more particles form strong isospin multiplets, that is transform in representations of strong isospin SU(2)? *Yes!* For example *pions*. There are three species of pions: π^+ , π^- , and π^0 . The $\pi^\pm$ has electric charge ± 1 , the π^0 is neutral. Their masses are:

$$m^\pm = 139.57 \text{ MeV}, \quad m^0 = 134.96 \text{ MeV}. \quad (5.4)$$

Again, the three pion states have very similar strong interactions. The differences in their masses and interactions are $\sim 1\%$ EM-size effects. The three pions form an *isospin-one* state

$$\boldsymbol{\pi} = \begin{pmatrix} \pi_1 \\ \pi_2 \\ \pi_3 \end{pmatrix}, \quad (5.5)$$

which means that π transforms in the *spin-one*, or *adjoint*, representation of SU(2). This representation works as follows: First, one transforms $\boldsymbol{\pi}$ to a 2×2 matrix via

$$\pi_{ab} = \tau_{ab}^i \pi_i = \boldsymbol{\tau}_{ab} \cdot \boldsymbol{\pi}. \quad (5.6)$$

Then, the generators of strong isospin SU(2) act on this state by commutation:

$$\begin{aligned} T_{\text{adj}}^i &= [\tau^i, \cdot] \\ \Rightarrow T^i \pi &= T_{\text{adj}}^i \pi = [\tau^i, \pi] = [\tau^i, \tau^j] \pi_j = i \varepsilon^{ijk} \pi_j \tau^k. \end{aligned} \quad (5.7)$$

The eigenvalue of a state under T^3 is called I_3 (“strong isospin”). Experimentally, we know that the electric charge eigenstates $\pi^\pm$, π^0 have strong isospin $I_3^\pm = \pm 1$, $I_3^0 = 0$. Writing T^3 as a 3×3 matrix in the basis of (5.7), it reads

$$T^3 = \begin{pmatrix} 0 & i & 0 \\ -i & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix}. \quad (5.8)$$

This matrix indeed has eigenvalues $\{1, 0, -1\}$, with eigenstates

$$\pi^\pm = (\pi_1 \pm i\pi_2)/\sqrt{2}, \quad \pi^0 = \pi_3. \quad (5.9)$$

This shows the relationship between the *charge eigenstates* $\pi^\pm$, π^0 and the isospin-one multiplet (5.5).

The pions $\pi^\pm$ have opposite charges (electric and I_3), but are otherwise identical: π^- is the antiparticle of π^+ and vice versa. π^0 is its own antiparticle.

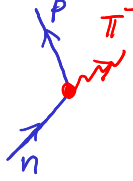
Symmetry. If strong isospin is a symmetry, also interactions must respect it. What is the most general symmetric Lagrangian for nucleon-pion interaction? Assume that the number of nucleons is preserved (but not the number of pions). We can build a Lagrangian out of the following operators:

$p^\dagger$	creates a proton / destroys an antiproton
p	destroys a proton / creates an antiproton
$n^\dagger$	creates a neutron / destroys an antineutron
n	destroys a neutron / creates an antineutron
π^+	creates a π^+ / destroys a π^-
$(\pi^+)^\dagger = \pi^-$	creates a π^- / destroys a π^+
π^0	creates/destroys a π^0

The most general nucleon-number preserving Lagrangian that also preserves I_3 then is

$$\mathcal{L}_{\text{int}} = g_{pn} p^\dagger n \pi^- + g_{np} n^\dagger p \pi^+ + g_{pp} p^\dagger p \pi^0 + g_{nn} n^\dagger n \pi^0, \quad (5.10)$$

where g_{xy} are coupling constants. Here, each term stands for a three-point interaction. For example, the first term stands for a process where a neutron n emits a pion π^- , turning into a proton p :



Even though it preserves I_3 , the Lagrangian $\mathcal{L}_{\text{int}}$ is not invariant under rotations in strong isospin space for general values of the couplings g_{xy} . For example, when we rotate p and n , then $p^\dagger p$ and $n^\dagger n$ mix with each other, and therefore g_{pp} and g_{nn} must be related.

How can we write $\mathcal{L}_{\text{int}}$ in a way that it is automatically – “manifestly” – invariant? The answer is, we have to make use of representation theory. We need to form a *singlet* (a trivial representation) of strong isospin $\text{SU}(2)$ out of a fundamental N , an anti-fundamental $N^\dagger = (p^\dagger, n^\dagger)$, and an adjoint π . This is exactly what the Pauli matrices τ_{ab}^i do! They carry one adjoint index i , and one fundamental/anti-fundamental index pair a, b . By contracting them with $N, N^\dagger$,

$$N^\dagger \tau^i N = (N^\dagger)^a \tau_{ab}^i N^b, \quad (5.11)$$

the fundamental index of N and the anti-fundamental index of $N^\dagger$ are transformed into an adjoint index i . Contracting this index with the adjoint index of π , we can form a *scalar* that is invariant under strong isospin $\text{SU}(2)$ rotations. The most general invariant Lagrangian therefore is

$$\mathcal{L}_{\text{int}} = g N^\dagger \boldsymbol{\tau} N \cdot \boldsymbol{\pi} = g (N^\dagger)^a \tau_{ab}^i N^b \pi_i. \quad (5.12)$$

More explicitly,

$$\begin{aligned} \boldsymbol{\tau} \cdot \boldsymbol{\pi} &= \pi_1 \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} + \pi_2 \begin{pmatrix} 0 & -i \\ i & 0 \end{pmatrix} + \pi_3 \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} \\ &= \begin{pmatrix} \pi_3 & \pi_1 - i\pi_2 \\ \pi_1 + i\pi_2 & -\pi_3 \end{pmatrix} = \begin{pmatrix} \pi^0 & \sqrt{2}\pi^- \\ \sqrt{2}\pi^+ & -\pi^0 \end{pmatrix}, \end{aligned} \quad (5.13)$$

hence the interaction Lagrangian becomes

$$\begin{aligned} \mathcal{L}_{\text{int}} &= g N^\dagger \boldsymbol{\tau} \cdot \boldsymbol{\pi} N = g \begin{pmatrix} p^\dagger & n^\dagger \end{pmatrix} \begin{pmatrix} \pi^0 & \sqrt{2}\pi^- \\ \sqrt{2}\pi^+ & -\pi^0 \end{pmatrix} \begin{pmatrix} p \\ n \end{pmatrix} \\ &= g (p^\dagger p \pi^0 + \sqrt{2} p^\dagger n \pi^- + \sqrt{2} n^\dagger p \pi^+ - n^\dagger n \pi^0). \end{aligned} \quad (5.14)$$

We see that all parameters g_{pp} , g_{nn} , g_{pn} , and g_{np} are related to the single coupling g .

Remark. As mentioned, strong isospin is *not* an exact symmetry, but this example illustrates the idea of invariance under internal “rotations”. Later, we will look at *weak isospin*, which is a true symmetry of the strong, weak, and electromagnetic interactions, and where the adjoint (isospin one) particle consists of the W bosons instead of pions.

5.2 Non-Abelian Gauge Theories

In the Abelian case, we considered states ψ that are invariant under phase transformations $\psi \mapsto e^{i\chi}\psi$, and local phase invariance implied the existence of and precise interaction with the gauge field A^μ , identified as the electromagnetic field. This is the essence of gauge theory.

Non-Abelian Transformations. For the example of strong isospin, general transformations are of the form

$$N \mapsto UN, \quad N = \begin{pmatrix} p \\ n \end{pmatrix}, \quad U = e^{i\boldsymbol{\varepsilon} \cdot \boldsymbol{\tau}/2} \in \text{SU}(2). \quad (5.15)$$

Here, invariance under phase transformations are generalized to invariance under $\text{SU}(2)$ transformations. This transformation is *non-Abelian*, since applying two successive transformations in different orders gives different results:

$$U, V \in \text{SU}(2) : \quad UVN \neq VUN \quad \Leftrightarrow \quad [U, V] \neq 0. \quad (5.16)$$

The non-commutativity is directly related to the non-trivial commutation relations of the Pauli matrices:

$$[\tau^i, \tau^j] = 2i \varepsilon^{ijk} \tau^k. \quad (5.17)$$

One could equally well consider particles transforming in representations of other groups and demand invariance. Besides $\text{SU}(2)$, another group that is important for the Standard Model is $\text{SU}(3)$. States $q = (a_1, a_2, a_3)$ that transform in the fundamental representation of $\text{SU}(3)$ transform as

$$q \mapsto Mq, \quad q = \begin{pmatrix} a_1 \\ a_2 \\ a_3 \end{pmatrix}, \quad M = e^{i\boldsymbol{\alpha} \cdot \boldsymbol{\lambda}/2} \in \text{SU}(3), \quad (5.18)$$

where $\boldsymbol{\lambda} = (\lambda_1, \dots, \lambda_8)$ are the eight $\text{SU}(3)$ generators, and $\boldsymbol{\alpha} = (\alpha_1, \dots, \alpha_8)$ are the corresponding eight transformation parameters. Quarks, the constituents of nucleons, indeed have such an $\text{SU}(3)$ degree of freedom called *color*.

Non-Abelian Gauge Theory. By promoting such non-Abelian symmetries to *local* symmetries, one obtains *non-Abelian gauge theory*. Consider a general state $\psi = (\psi_1, \dots, \psi_n)$ that transforms in a matrix representation of a general gauge group G :

$$\psi \mapsto U\psi, \quad \psi = \begin{pmatrix} \psi_1 \\ \psi_2 \\ \psi_3 \end{pmatrix}, \quad U = e^{i\boldsymbol{\omega} \cdot \boldsymbol{T}} = e^{i\omega^a T_a} \in G. \quad (5.19)$$

Here, to simplify the notation we do not distinguish between the abstract group G and its concrete representation: U is an element of G in a matrix representation. The gauge

group G will always be a Lie group, and $\mathbf{T} = (T_1, \dots, T_d)$ are the generators of the corresponding Lie algebra, where d is the dimension of the gauge group (and its associated algebra).

When we promote the transformations (5.19) to local symmetries, they are called *gauge transformations*, and the spacetime-dependent transformation parameters

$$\omega = \omega(t; \mathbf{x}) \quad (5.20)$$

are called *gauge parameters*. As in Abelian gauge theory, free particles by themselves cannot be gauge invariant, since the derivatives in the Schrödinger equation (or its relativistic equivalent) produce extra terms when acting on the gauge parameters $\omega(t; \mathbf{x})$. We again have to construct a *covariant derivative*

$$D^\mu = \partial^\mu - igA^\mu \quad (5.21)$$

to achieve invariance. Here, g is a coupling constant, and A^μ is a field that transforms in the adjoint representation of the gauge group. That is, A^μ takes values in the Lie algebra $\mathfrak{g}$ associated to the gauge Lie group G . It can therefore be expanded in the fundamental generators T^a :

$$A^\mu = A_a^\mu T^a, \quad A_a^\mu \in \mathbb{C}, \quad T^a \in \mathfrak{g}. \quad (5.22)$$

Earlier, we constructed the transformation rule of a gauge field A^μ for a general (in general non-commutative) transformation U :

$$A^\mu \mapsto A'^\mu = UA^\mu U^{-1} - \frac{i}{g} (\partial^\mu U) U^{-1}. \quad (5.23)$$

This transformation rule guarantees that the covariant derivative

$$D^\mu \psi \mapsto (D^\mu \psi)' = UD^\mu \psi \quad (5.24)$$

transforms in the same way as the state $\psi \mapsto U\psi$ under general gauge transformations U . We can find the infinitesimal form of the transformation (5.23) as follows: Write the transformation U in terms of generators, $U = e^{i\omega \cdot \mathbf{T}} = e^{i\omega_a T^a}$, and expand to linear order in ω_a :

$$\begin{aligned} A'^\mu &= (1 + i\omega \cdot \mathbf{T})A^\mu(1 - i\omega \cdot \mathbf{T}) - \frac{i}{g} i(\partial^\mu \omega) \cdot \mathbf{T} + \mathcal{O}(\omega^2) \\ &= A^\mu + i\omega \cdot [\mathbf{T}, A^\mu] + \frac{1}{g} (\partial^\mu \omega) \cdot \mathbf{T} + \mathcal{O}(\omega^2) \end{aligned} \quad (5.25)$$

Further expanding A^μ in terms of generators, this becomes

$$\begin{aligned} A_a'^\mu T^a &= A_a^\mu T^a + i\omega_a [T^a, T^b] A_b^\mu + \frac{1}{g} (\partial^\mu \omega_a) T^a + \mathcal{O}(\omega^2) \\ &= A_a^\mu T^a + i\omega_a i f^{ab}{}_c T^c A_b^\mu + \frac{1}{g} (\partial^\mu \omega_a) T^a + \mathcal{O}(\omega^2) \\ &= \left[A_a^\mu - \omega_c f^{cb}{}_a A_b^\mu + \frac{1}{g} (\partial^\mu \omega_a) \right] T^a + \mathcal{O}(\omega^2) \\ &= \left[A_a^\mu + f^{bc}{}_a A_b^\mu \omega_c + \frac{1}{g} (\partial^\mu \omega_a) \right] T^a + \mathcal{O}(\omega^2). \end{aligned} \quad (5.26)$$

The infinitesimal change of the components A_a^μ therefore is

$$A_a'^\mu = A_a^\mu + \delta A_a^\mu, \quad \delta A_a^\mu = f^{bc}{}_a A_b^\mu \omega_c + \frac{1}{g} (\partial^\mu \omega_a) \quad (5.27)$$

This shows that there is a consistent transformation for A_a^μ that makes the theory gauge invariant, by writing the Lagrangian in terms of the covariant derivative D^μ .

5.3 Quarks and Leptons

No theoretical principle predicts what internal spaces particles transform in. The invariance of the theory under such transformations is simply observed. The complete set of spaces needed for the Standard Model is $U(1)$, $SU(2)$ (the electroweak spaces), and $SU(3)$ (the color space).

Weak Isospin. In the Standard Model, both quarks and leptons can be put in $SU(2)$ doublets of a *weak isospin space*:

$$\begin{array}{cc} \text{quarks} & \text{leptons} \\ \begin{pmatrix} u \\ d \end{pmatrix}, \quad \begin{pmatrix} c \\ s \end{pmatrix}, \quad \begin{pmatrix} t \\ b \end{pmatrix}, & \begin{pmatrix} \nu_e \\ e \end{pmatrix}, \quad \begin{pmatrix} \nu_\mu \\ \mu \end{pmatrix}, \quad \begin{pmatrix} \nu_\tau \\ \tau \end{pmatrix}, \end{array} \quad (5.28)$$

and the theory is invariant under *local* weak isospin $SU(2)$ transformations. The weak isospin $SU(2)$ gauge transformations are

$$U = e^{i\boldsymbol{\varepsilon} \cdot \boldsymbol{\tau}/2} \in SU(2), \quad (5.29)$$

with the local (space-time dependent) gauge parameters

$$\boldsymbol{\varepsilon} = \boldsymbol{\varepsilon}(t; \mathbf{x}). \quad (5.30)$$

The covariant derivative D^μ and gauge field A^μ therefore are

$$D^\mu = \partial^\mu - ig_2 A^\mu, \quad A^\mu = \frac{\boldsymbol{\tau}}{2} \cdot \mathbf{W}^\mu = \frac{\tau_i}{2} W_i^\mu, \quad (5.31)$$

where g_2 is the coupling constant, and the components of A^μ are the *weak gauge bosons* W_i^μ . According to our general formula (5.27), the gauge bosons transform as

$$W_i^\mu \mapsto W_i^\mu + \delta W_i^\mu, \quad \delta W_i^\mu = \varepsilon_{jki} W_j^\mu \varepsilon_k + \frac{1}{g_2} (\partial^\mu \varepsilon_i), \quad (5.32)$$

where $\varepsilon_{jki} = f_{jki} = f^{jk}_i$ are the structure constants of $\mathfrak{su}(2)$.

Color. In addition, quark states transform in an $SU(3)$ triplet (fundamental representation) in *color space*:

$$q = \begin{pmatrix} r \\ g \\ b \end{pmatrix}, \quad q \in \{d, u, s, c, b, t\}. \quad (5.33)$$

Again, the theory is invariant under *local* color $SU(3)$ transformations. The story for $SU(2)$ repeats with different labels: The color $SU(3)$ gauge transformations are

$$U = e^{i\boldsymbol{\alpha} \cdot \boldsymbol{\lambda}/2} \in SU(3), \quad (5.34)$$

with the $SU(3)$ generators $\boldsymbol{\lambda} = (\lambda_1, \dots, \lambda_8)$ and the local (space-time dependent) gauge parameters

$$\boldsymbol{\alpha} = \boldsymbol{\alpha}(t; \mathbf{x}). \quad (5.35)$$

The covariant derivative D^μ and gauge field A^μ in this case are

$$D^\mu = \partial^\mu - ig_3 A^\mu, \quad A^\mu = \frac{\boldsymbol{\lambda}}{2} \cdot \mathbf{G}^\mu = \frac{\lambda_a}{2} G_a^\mu, \quad (5.36)$$

where g_3 is the coupling constant, and the components of A^μ are the *gluons* (strong gauge bosons) G_a^μ . According to the general formula (5.27), the gluons transform as

$$G_a^\mu \mapsto G_a^\mu + \delta G_a^\mu, \quad \delta G_a^\mu = f^{bc}_a G_b^\mu \alpha_c + \frac{1}{g_3} (\partial^\mu \alpha_a), \quad (5.37)$$

where f^{bc}_a are the structure constants of $\mathfrak{su}(3)$.

Covariant Derivative of the SM. By adding *several terms* to ∂_μ , we can make a covariant derivative D_μ that guarantees invariance under gauge transformations in several spaces (simultaneously or separately). The full covariant derivative of the Standard Model is

$$D^\mu = \partial^\mu - ig_1 \frac{Y}{2} B^\mu - ig_2 \frac{\tau_i}{2} W_i^\mu - ig_3 \frac{\lambda_a}{2} G_a^\mu \quad (5.38)$$

Here,

- Y is a generator of $U(1)$: *hypercharge*,
- $\frac{1}{2} \tau_i$ are generators of $SU(2)$: *weak isospin*,
- $\frac{1}{2} \lambda_a$ are generators of $SU(3)$: *color*.

The full space that D^μ acts on is a representation of the product group $U(1) \times SU(2) \times SU(3)$, which is the full gauge group of the Standard Model. The hypercharge generator Y acts trivially (as a unit matrix) in the $SU(2) \times SU(3)$ space. Similarly, τ_i acts trivially in $U(1) \times SU(3)$, and λ_a acts trivially in $U(1) \times SU(2)$.

Some comments:

- The three parameters g_1 , g_2 , and g_3 are arbitrary real numbers that have to be fixed by comparison to experiment.
- Various states (particles) have various charges (eigenvalues) under the generators Y , τ_i , and λ_a . These charges serve as state labels that identify the various particles (quarks and leptons).
- As for (Abelian) electrodynamics, the three terms in D^μ mean that several spin-one gauge bosons must exist: One B^μ , three W_i^μ , and eight G_a^μ . All these are confirmed experimentally.
- No one knows why the gauge group is $U(1) \times SU(2) \times SU(3)$ and not some other group.
- The equation for D^μ is the main equation of the Standard Model: It tells about the internal spaces, the gauge bosons that must exist, and their interactions, all based on the postulate of gauge invariance.

6 Relativistic Fermions

The elementary fields/particles of matter, the quarks and leptons, are all spin-1/2 fermions. To describe them, it is necessary to generalize the Schrödinger equation to include the spin degrees of freedom. Also, the theory should be relativistically invariant: Special

relativity dictates that the laws of physics must be the same in all inertial systems. In mathematical terms, this means that all physical equations must be invariant, that is should not change their form, when we apply Lorentz transformations that relate one inertial frame to another. This leads to the Dirac equation. The formulation will provide a concise notation that is useful for computations within the Standard Model.

6.1 The Dirac Equation

At the time of Dirac, physicists were looking for a relativistically invariant equation that describes the electron, including its spin. The Schrödinger equation

$$-i\partial_t\phi = -\frac{1}{2m}\nabla^2\phi \quad (6.1)$$

is linear in the time derivative, but quadratic in the spatial derivatives. Lorentz rotations mix spatial directions with the time direction, therefore time and spatial derivatives have to appear with identical degrees in a relativistically invariant equation.

Klein–Gordon Equation. A first guess would be to start with the relativistically invariant equation

$$E^2 = \mathbf{p}^2 + m^2 \quad (6.2)$$

and replace E and $\mathbf{p}$ by their quantum mechanical operators acting on a state ϕ :

$$-\partial_t^2\phi = (-\nabla^2 + m^2)\phi \quad \Leftrightarrow \quad (\partial_\mu\partial^\mu + m^2)\phi = 0. \quad (6.3)$$

But this results in a problem: A quantum mechanical state ϕ that satisfies the Schrödinger equation describes a *probability density*

$$\rho = \phi^*\phi \quad (6.4)$$

that obeys the continuity equation

$$\frac{\partial}{\partial t}\rho + \nabla \cdot \mathbf{J} = 0 \quad (6.5)$$

with the probability current

$$\mathbf{J} = -\frac{i}{2m}(\phi^*\nabla\phi - \phi\nabla\phi^*). \quad (6.6)$$

Writing the continuity equation in the form

$$\left(\frac{\partial}{\partial t}; \nabla\right) \cdot (\rho; \mathbf{J}) = 0, \quad (6.7)$$

it is clear that it *cannot* be relativistically invariant. To make it relativistically invariant, $(\rho; \mathbf{J})$ would have to be a relativistic four-vector. To complete $\mathbf{J}$ to a four-vector, the probability density ρ has to take the form

$$\rho = \frac{i}{2m}(\phi^*\partial_t\phi - \phi\partial_t\phi^*), \quad (6.8)$$

so that the full probability density four-current is relativistically covariant:

$$J^\mu = \frac{i}{2m}(\phi^*\partial^\mu\phi - \phi\partial^\mu\phi^*), \quad (6.9)$$

and the continuity equation becomes

$$\partial_\mu J^\mu = 0. \quad (6.10)$$

Everything is consistent with relativity now. But there is a problem: Because the equation (6.3) is of second order in time, the initial values of both ϕ and $\partial_t \phi$ may be freely chosen, and $\rho = J^0$ may become negative! That should be impossible for a probability density. So we cannot get a relativistically invariant generalization of the Schrödinger equation that is of second order in the time derivative.

The equation (6.3) is still relevant in quantum field theory, it is called the *Klein–Gordon equation*. When interpreted as a probability amplitude, its wave function ϕ (or the probability density $\phi^* \phi$) does not obey the laws of relativity. And the relativistic density $\rho \sim (\phi^* \partial_t \phi - \phi \partial_t \phi^*)$ can become negative and therefore cannot be a probability. But ρ can be interpreted as a *charge density*, for which negative values are admissible. The Klein–Gordon equation therefore can describe relativistic spinless particles with positive, negative, and zero charge, for example pions.

Dirac’s Ansatz. As we saw above, a relativistically invariant equation for a probability amplitude cannot be quadratic in time derivatives. Dirac was looking for an equation that was linear in the time derivative, and therefore, by relativistic invariance, had to be linear in spatial derivatives. He wrote down the most general such equation that is also linear in the state ψ :

$$i \frac{\partial}{\partial t} \psi = \left[-i \left(\alpha_1 \frac{\partial}{\partial x^1} + \alpha_2 \frac{\partial}{\partial x^2} + \alpha_3 \frac{\partial}{\partial x^3} \right) + \beta m \right] \psi. \quad (6.11)$$

The coefficients α_i and β are constrained by physics requirements. Using the relation $E^2 = \mathbf{p}^2 + m^2$ and replacing E and $\mathbf{p}$ by their quantum-mechanical operators, any solution ψ should satisfy the Klein–Gordon equation (6.3),

$$-\partial_t^2 \psi = (-\nabla^2 + m^2) \psi. \quad (6.12)$$

We will see later that this is not in contradiction to what we discussed above. Moving all terms in (6.11) to the same side, the equation takes the form $D\psi = 0$, which implies that also $D^2\psi = 0$. Multiplying this out, assuming that α_i and β are constant, and sorting terms, this becomes

$$-\frac{\partial^2}{\partial t^2} \psi = \left[-\alpha_i^2 \frac{\partial^2}{\partial x^{i2}} - \sum_{i < j} (\alpha_i \alpha_j + \alpha_j \alpha_i) \frac{\partial}{\partial x^i} \frac{\partial}{\partial x^j} - i m (\alpha_i \beta + \beta \alpha_i) \frac{\partial}{\partial x^i} + \beta^2 m^2 \right] \psi. \quad (6.13)$$

Requiring that this is compatible with the relativistic energy relation (6.12) gives the following constraints:

$$\begin{aligned} \alpha_i \alpha_j + \alpha_j \alpha_i &= 0 \quad (i \neq j), \\ \alpha_i \beta + \beta \alpha_i &= 0, \\ \alpha_i^2 &= \beta^2 = 1. \end{aligned} \quad (6.14)$$

These relations cannot be satisfied by numbers. But Dirac had just worked on Heisenberg’s new matrix mechanics, and immediately realized that α_i and β had to be matrices.

Massless Fermions. For massless fermions, the term with β is absent, and the conditions on α_i are

$$\alpha_i \alpha_j + \alpha_j \alpha_i = \{\alpha_i, \alpha_j\} = 2\delta_{ij}, \quad (6.15)$$

where

$$\{a, b\} \equiv ab + ba \quad (6.16)$$

is the anticommutator. These are exactly the anticommutation relations of the Pauli matrices, so one can set

$$\alpha_i = -\sigma_i, \quad (6.17)$$

where the minus sign is a convention. Then Dirac's ansatz (6.11) becomes the *massless Dirac equation*

$$i\partial_t \psi = \boldsymbol{\sigma} \cdot \mathbf{p} \psi, \quad p_i = i \frac{\partial}{\partial x^i}. \quad (6.18)$$

and ψ is a two-component spinor.

Massive Fermions. With a non-zero mass m , the first guess would be to find a 2×2 matrix β that satisfies the conditions (6.14) together with $\alpha_i = -\sigma_i$. A complete basis for all 2×2 matrices is formed by the Pauli matrices σ_i together with the unit matrix. The second condition $\alpha_i \beta + \beta \alpha_i = 0$ together with $\alpha_i^2 = 1$ shows that β cannot have α_i components. Setting $\beta = 1$, the unit matrix, also fails. Hence the conditions cannot be satisfied with 2×2 matrices.

It turns out that to satisfy the conditions (6.14), the matrices α_i, β have to be at least 4×4 in size. A particular solution to the constraints is

$$\alpha_i = \begin{pmatrix} 0 & \sigma_i \\ \sigma_i & 0 \end{pmatrix}, \quad \beta = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}, \quad (6.19)$$

where each matrix entry stands for a 2×2 submatrix. So for example,

$$\alpha_1 = \begin{pmatrix} 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 \\ 0 & 1 & 0 & 0 \\ 1 & 0 & 0 & 0 \end{pmatrix}. \quad (6.20)$$

Gamma Matrices. The choice of 4×4 matrices for α_i, β is not unique. We will multiply Dirac's ansatz (6.11) by β from the left:

$$i\beta \partial_t \psi = (-i\beta \boldsymbol{\alpha} \cdot \nabla + m)\psi, \quad (6.21)$$

and rename the coefficient matrices as

$$\gamma^0 = \beta, \quad \gamma^i = \beta \alpha_i. \quad (6.22)$$

These *gamma matrices* can be combined into a four-vector,

$$\gamma^\mu = (\gamma^0; \gamma^i). \quad (6.23)$$

The *massive Dirac equation* then can be written in the compact form

$$(i\gamma^\mu \partial_\mu - m)\psi = 0, \quad (6.24)$$

where ψ now is a four-component spinor. The conditions (6.14) on α_i and β are equivalent to the relations

$$\gamma^\mu \gamma^\nu + \gamma^\nu \gamma^\mu = \{\gamma^\mu, \gamma^\nu\} = 2g^{\mu\nu} \quad (6.25)$$

for the gamma matrices. Here, $g^{\mu\nu}$ are the components of the metric tensor

$$g = \begin{pmatrix} +1 & 0 & 0 & 0 \\ 0 & -1 & 0 & 0 \\ 0 & 0 & -1 & 0 \\ 0 & 0 & 0 & -1 \end{pmatrix}. \quad (6.26)$$

The relations (6.25) for the gamma matrices are the defining relations for a *Clifford algebra*. The particular solution for α_i and β (6.19) give the *standard representation*, also called *Dirac representation*, for the gamma matrices:

$$\gamma^0 = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}, \quad \gamma^i = \begin{pmatrix} 0 & \sigma_i \\ -\sigma_i & 0 \end{pmatrix}, \quad (6.27)$$

where again each entry stands for a 2×2 submatrix.

Relativistic Invariance. We have been calling γ^μ a four-vector, suggesting that it transforms as a vector under Lorentz transformations, such that $\gamma^\mu \partial_\mu$ is a Lorentz scalar. This cannot be correct if the gamma matrices are constant. What we really want is that the Dirac equation (6.24)

$$(i\gamma^\mu \partial_\mu - m)\psi = 0 \quad (6.28)$$

is relativistically invariant, meaning that it does not change its form when we apply a Lorentz transformation. We know how ∂_μ transforms under Lorentz transformations (as a covariant vector), but we do not yet know how γ^μ and ψ should transform.

If $\gamma^\mu \partial_\mu$ is to be invariant, γ^μ has to transform as a contravariant four-vector. Because Lorentz transformations are 4×4 matrices $\Lambda_L \in \text{SO}(1,3)$ that preserve the metric tensor $g^{\mu\nu}$, the transformed set γ'^μ will still satisfy the Clifford algebra relations. Now, it is a theorem that, if two sets of matrices γ'^μ and γ^μ satisfy the Clifford algebra, they must be related by a similarity transformation

$$\gamma'^\mu = \Lambda_S^{-1} \gamma^\mu \Lambda_S. \quad (6.29)$$

With $\gamma^\mu \partial_\mu = \gamma'^\mu \partial'_\mu$, the Dirac equation can be re-written as

$$\begin{aligned} (i\Lambda_S^{-1} \gamma^\mu \Lambda_S \partial'_\mu - m)\psi(x') &= 0 \\ \Leftrightarrow \Lambda_S^{-1} (i\gamma^\mu \partial'_\mu - m) \Lambda_S \psi(x') &= 0, \end{aligned} \quad (6.30)$$

where in the second step we have used that Λ_S is constant and $\Lambda_S^{-1} \Lambda_S = 1$. Now if we identify $\Lambda_S \psi$ as the transformed spinor ψ' ,

$$\psi' = \Lambda_S \psi, \quad (6.31)$$

then the Dirac equation becomes

$$\Lambda_S^{-1} (i\gamma^\mu \partial'_\mu - m) \psi'(x') = 0 \quad \Leftrightarrow \quad (i\gamma^\mu \partial'_\mu - m) \psi'(x') = 0, \quad (6.32)$$

which is of the same form as the original equation, with the same gamma matrices γ^μ , but with Lorentz-transformed coordinate x' , derivative ∂' , and spinor ψ' .

What have we shown? We have shown that the Dirac equation (with constant gamma matrices) is relativistically invariant provided that the spinor state ψ transforms according to (6.31), where Λ_S is the transformation (4×4 matrix) acting in *spinor space* that corresponds to the applied Lorentz transformation Λ_L via

$$\gamma'^\mu = (\Lambda_L)^\mu{}_\nu \gamma^\nu = \Lambda_S^{-1} \gamma^\mu \Lambda_S. \quad (6.33)$$

Lorentz Transformations. In order to understand the spinors ψ better, we have to note some facts about representations of the Lorentz group. The Lorentz group $\text{SO}(1, 3)$ comprises all transformations of space-time that relate different inertial frames. On contravariant four-vectors, like x^μ , Lorentz transformations act by multiplication with a matrix $\Lambda_L \in \text{SO}(1, 3)$ that preserves the metric tensor $g_{\mu\nu}$, such that $x^\mu \mapsto (\Lambda_L)^\mu{}_\nu x^\nu$. For example, a boost in the x^1 direction takes the form

$$\Lambda_L = \begin{pmatrix} \cosh \beta & \sinh \beta & 0 & 0 \\ \sinh \beta & \cosh \beta & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix} \Leftrightarrow K_1 = i \begin{pmatrix} 0 & -1 & 0 & 0 \\ -1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{pmatrix},$$

where the matrix on the right is the infinitesimal generator. A rotation in the (x^1, x^2) plane is represented by

$$\Lambda_L = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & \cos \theta & \sin \theta & 0 \\ 0 & -\sin \theta & \cos \theta & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix} \Leftrightarrow J_3 = i \begin{pmatrix} 0 & 0 & 0 & 0 \\ 0 & 0 & -1 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{pmatrix}.$$

We can arrange all generators of $\text{SO}(1, 3)$ (boosts and rotations) in a 4×4 matrix:

$$L^{\mu\nu} = \begin{pmatrix} 0 & K_1 & K_2 & K_3 \\ -K_1 & 0 & J_3 & -J_2 \\ -K_2 & -J_3 & 0 & J_1 \\ -K_3 & J_2 & -J_1 & 0 \end{pmatrix}, \quad (6.34)$$

where each entry is again a 4×4 matrix, for example $L^{23} = J_3$. The matrices $L^{\mu\nu}$ then satisfy the Lorentz algebra

$$[L^{\mu\nu}, L^{\rho\sigma}] = i(g^{\nu\rho} L^{\mu\sigma} - g^{\nu\sigma} L^{\mu\rho} - g^{\mu\rho} L^{\nu\sigma} + g^{\mu\sigma} L^{\nu\rho}). \quad (6.35)$$

The generators for the matrices Λ_S that represent Lorentz transformations on four-component (Dirac) spinors, and that we derived from the gamma matrices via (6.33), are given by

$$S^{\mu\nu} = \frac{i}{4} [\gamma^\mu, \gamma^\nu], \quad (6.36)$$

where we use the same notation as in (6.34): For fixed values of μ and ν , $S^{\mu\nu}$ is a 4×4 matrix. One can check that these generators indeed satisfy the Lorentz algebra (6.35) for any set of gamma matrices that satisfy the Clifford algebra relations (6.25).

Now, it is a fact that all finite-dimensional representations of the Lorentz group are non-unitary.¹ All unitary representations of the Lorentz group are infinite dimensional,

¹The reason is that the Lorentz group is non-compact. It can be understood by noting that the Lie algebra of the Lorentz group is $\mathfrak{su}(2) \oplus \mathfrak{su}(2)$, and therefore representations of the Lorentz group can be built from the representations of $\mathfrak{su}(2)$, which are familiar from non-relativistic spin (Pauli matrices). Those representations *are* unitary, but the mapping from pairs of $\text{SU}(2)$ representations to $\text{SO}(1, 3)$ representations is *complex*, which spoils the unitarity.

which means that states that transform in such a representation must have infinitely many (a continuum of) degrees of freedom. This implies that relativistic quantum states *must* be quantum *fields* with a continuum of degrees of freedom. In this sense, one can say that relativistic invariance *necessitates* quantum field theory.

Coming back to our representation $S^{\mu\nu}$ (6.36), one can indeed check that not all the matrices $S^{\mu\nu}$ are hermitian (neither are all the $L^{\mu\nu}$). For example, in the Dirac representation for the gamma matrices, γ^0 is hermitian, but γ^i are anti-hermitian:

$$(\gamma^0)^\dagger = \gamma^0, \quad (\gamma^i)^\dagger = -\gamma^i. \quad (6.37)$$

This means that $\psi^\dagger\psi$, contrary to naive intuition, is *not a Lorentz scalar*, since under Lorentz transformations:

$$\psi^\dagger\psi \mapsto \psi^\dagger \Lambda_S^\dagger \Lambda_S \psi, \quad (6.38)$$

which is not equal to $\psi^\dagger\psi$, since Λ_S is not unitary, and therefore $\Lambda_S^\dagger \neq \Lambda_S^{-1}$. In order to write Lorentz scalars, we introduce the combination

$$\bar{\psi} \equiv \psi^\dagger \gamma^0. \quad (6.39)$$

It can be shown from the properties of the gamma matrices and the definition $\Lambda_S = \exp(i\theta_{\mu\nu} S^{\mu\nu})$ that

$$\Lambda_S^\dagger \gamma^0 = \gamma^0 \Lambda_S^{-1}, \quad (6.40)$$

and therefore

$$\bar{\psi}\psi \quad (6.41)$$

is a Lorentz scalar. Similarly, one can show that $\bar{\psi}\gamma^\mu\psi$ transforms as a Lorentz vector, that is

$$\bar{\psi}\gamma^\mu\psi \mapsto (\Lambda_L)^\mu{}_\nu \bar{\psi}\gamma^\nu\psi \quad (6.42)$$

under Lorentz transformations.

For more details on representations of the Lorentz group and spinors, see for example Chapter 10 of Matthew Schwartz' excellent book "*Quantum Field Theory and the Standard Model*".

6.2 Conserved Current

We can construct a conserved current from ψ . Starting with the Dirac equation

$$(i\gamma^\mu \partial_\mu - m)\psi = 0, \quad (6.43)$$

we take its adjoint (hermitian conjugate):

$$-i\partial_\mu \psi^\dagger \gamma^{\mu\dagger} - m\psi^\dagger = 0, \quad (6.44)$$

and then multiply with γ^0 from the right, noting that

$$\gamma^{i\dagger}\gamma^0 = -\gamma^0\gamma^{i\dagger} = \gamma^0\gamma^i \quad \text{and} \quad \gamma^{0\dagger} = \gamma^0, \quad (6.45)$$

to obtain

$$\begin{aligned} 0 &= i\partial_\mu \psi^\dagger \gamma^0 \gamma^\mu + m\psi^\dagger \gamma^0 \\ &= i\partial_\mu \bar{\psi} \gamma^\mu + m\bar{\psi}. \end{aligned} \quad (6.46)$$

This equation is also sometimes written in the form

$$\bar{\psi}(i\gamma^\mu \overleftarrow{\partial}_\mu + m) = 0 \quad (6.47)$$

to resemble the Dirac equation for ψ more closely. Here, it is understood that the derivative operator $\overleftarrow{\partial}_\mu$ acts to the left. Now we multiply the original Dirac equation (6.43) by $\bar{\psi}$ from the left, and the conjugate equation (6.46) by ψ from the right, and sum the two equations. The mass terms cancel, and what remains is

$$0 = \bar{\psi}\gamma^\mu\partial_\mu\psi + (\partial_\mu\bar{\psi})\gamma^\mu\psi = \partial_\mu(\bar{\psi}\gamma^\mu\psi). \quad (6.48)$$

Hence, we found a current that is conserved:

$$j^\mu = \bar{\psi}\gamma^\mu\psi, \quad \partial_\mu j^\mu = 0. \quad (6.49)$$

Recall from the beginning of this section that the conserved current following from the Klein–Gordon equation could not be interpreted as a continuity equation for a probability density, because the time component of the conserved current was not positive definite. Hence the Klein–Gordon equation could not describe the probability amplitude of a relativistic particle. This problem is solved by the Dirac equation: The time component of the current is

$$j^0 = \bar{\psi}\gamma^0\psi = \psi^\dagger\psi \equiv \rho, \quad (6.50)$$

which is positive definite. It satisfies the continuity equation

$$\partial_t\rho = -\nabla\bar{\psi}\boldsymbol{\gamma}\psi, \quad \boldsymbol{\gamma} = (\gamma^1, \gamma^2, \gamma^3), \quad (6.51)$$

and its integral over all of space is preserved:

$$Q = \int \rho d^3x, \quad \partial_t Q = 0. \quad (6.52)$$

The current component ρ therefore has a consistent interpretation as a probability density for a fermion. The conserved charge Q is the total particle number, and the spatial components of j^μ describe the particle number flow.

6.3 Free Particle Solutions

After this mostly abstract treatment, we can get more familiar with the Dirac equation by looking at its solutions.

Wave Equation. The Dirac equation is a wave equation for massive relativistic particles with spin. The fact that it is a wave equation can be seen as follows: Multiply the Dirac equation

$$(i\gamma^\mu\partial_\mu - m)\psi = 0 \quad (6.53)$$

from the left with $(i\gamma^\nu\partial_\nu + m)$ to find

$$(\gamma^\nu\gamma^\mu\partial_\mu\partial_\nu + m^2)\psi = 0. \quad (6.54)$$

Since $\partial_\mu\partial_\nu$ is symmetric in μ and ν , we can replace also $\gamma^\mu\gamma^\nu$ by their symmetric combination,

$$\gamma^\mu\gamma^\nu \rightarrow \frac{1}{2}\{\gamma^\mu, \gamma^\nu\} = g^{\mu\nu}, \quad (6.55)$$

which, by the Clifford algebra, equals the metric tensor. So our equation becomes

$$(\partial_\mu \partial^\mu + m^2)\psi = 0. \quad (6.56)$$

This looks exactly like the Klein–Gordon equation, but now applied to the four-component spinor ψ . The differential operator is proportional to the identity matrix, therefore each component of ψ satisfies the Klein–Gordon equation. It is a wave equation, hence its solutions, and therefore also the solutions to the Dirac equation, can be expanded in plane waves.

Plane-Wave States. We have seen that solutions to the Dirac equation can be expanded in plane-wave states. These take the form

$$\psi = u(p) e^{-ip_\mu x^\mu}, \quad (6.57)$$

where $u(p)$ is a four-component spinor that may depend on p , but not on x . Plugging this plane-wave state into the Dirac equation directly gives an equation for the spinor $u(p)$:

$$(\gamma^\mu p_\mu - m) u(p) = 0. \quad (6.58)$$

Here, p^μ can take any value that is compatible with $E^2 = p^2 + m^2$. In particular, nothing prevents us from choosing $E = p_0 < 0$ *negative*. Because a particle with momentum p^μ is indistinguishable from an antiparticle with momentum $-p^\mu$, we re-interpret particles with momentum p^μ and $E < 0$ as antiparticles with momentum $-p^\mu$ and $E > 0$. In other words, when expanding a state in terms of plane waves, instead of summing over all momenta p^μ , we sum only over momenta with $E > 0$, but count particles *and* antiparticles for each value of p^μ . For antiparticles, we rename the spinor $u(p)$ to $v(p)$. Since the mass does not change sign, $v(p)$ satisfies the equation

$$(\gamma^\mu p_\mu + m) v(p) = 0. \quad (6.59)$$

6.4 Particles and Antiparticles

Four Particle States. We saw that spinors described by the Dirac equation have four components. Spin 1/2 particles have two components, so what is the meaning of these four components? The answer is that the Dirac equation automatically describes both particles and antiparticles. To understand this, let us again look at the free-particle solutions. As we saw earlier, the Dirac equation $(i\gamma^\mu \partial_\mu - m)\psi = 0$ for a plane-wave state $\psi = u(p)e^{-ip_\mu x^\mu}$ becomes an equation for the four-component spinor $u(p)$:

$$(\gamma^\mu p_\mu - m)u(p) = 0. \quad (6.60)$$

For a particle at rest, $\mathbf{p} = 0$, this becomes

$$(\gamma^0 E - m)u = 0. \quad (6.61)$$

Recall that in the Dirac basis,

$$\gamma^0 = \begin{pmatrix} \mathbf{1} & 0 \\ 0 & -\mathbf{1} \end{pmatrix}, \quad (6.62)$$

and therefore

$$\begin{pmatrix} E - m & 0 & 0 & 0 \\ 0 & E - m & 0 & 0 \\ 0 & 0 & E + m & 0 \\ 0 & 0 & 0 & E + m \end{pmatrix} u = 0. \quad (6.63)$$

This means that $E = m$ for two of the solutions, and $E = -m$ for the two other solutions. More explicitly, we set u_i to be the i 'th unit vector:

$$(u_i)_j = \delta_{ij} \quad \text{e. g.} \quad u_1 = \begin{pmatrix} 1 \\ 0 \\ 0 \\ 0 \end{pmatrix}. \quad (6.64)$$

we find the four solutions

$$\begin{aligned} E = +m : \quad \psi_{1,2} &= u_{1,2} e^{-imt}, \\ E = -m : \quad \psi_{3,4} &= u_{3,4} e^{+imt}. \end{aligned} \quad (6.65)$$

Two of these solutions have negative energies. This leads to a problem: For non-zero momentum, the energies of these states becomes $-\sqrt{m^2 + \mathbf{p}^2}$, which is not bounded from below. In particular, a particle could radiate away an infinity of energy by increasing its momentum. Physical energies are always bounded from below. So what are these states with $E = -m$? The resolution is that these states represent *antiparticles* with *positive* energies.

CPT Inversion. To understand this fully requires quantum field theory, but we can also get an intuitive understanding with more basic methods. The energy is the eigenvalue of the Hamiltonian, which is the generator of time translations. In other words, the energy is the eigenvalue of the time-shift operator $i\partial_t$. Now if we reverse the orientation of the time coordinate,

$$t \rightarrow t' = -t, \quad (6.66)$$

the time-shift operator, and therefore also the energy switches sign. Therefore, a negative-energy particle can be re-interpreted as a positive-energy particle for which time runs backwards: A *negative-energy* particle moving *forwards* in time is equivalent to a *positive-energy* particle moving *backwards* in time. A particle “moving backwards in time” makes no sense physically, it is merely a theoretical construct. There is no experimental evidence whatsoever of anything *actually* moving backwards in time. Now, a particle moving in some direction as we go back in time is equivalent to a particle moving in the opposite direction as we go forwards in time. Since moving backwards in time makes little sense, we compensate the time inversion by also inverting all spatial directions $\mathbf{x}$, which inverts the momentum $\mathbf{p}$ of our particle. In order to preserve the physics, we need to do one further transformation. A particle with some charge q moving with a velocity $\mathbf{v}$ produces a measurable current $\mathbf{j} = q\mathbf{v}$. When we invert the velocity (momentum), we need to invert also the charge to preserve the current. This is true for all types of charges.

All in all, we have applied a *CPT* transformation to our particle states, where

$$\begin{aligned} \text{C (charge conjugation)} : \quad & q \rightarrow -q \\ \text{P (parity)} : \quad & \mathbf{p} \rightarrow -\mathbf{p} \quad (\mathbf{x} \rightarrow -\mathbf{x}) \end{aligned}$$

$$T \text{ (time reversal)} : \quad E \rightarrow -E \quad (t \rightarrow -t). \quad (6.67)$$

As far as we know, the combined CPT transformation is an exact symmetry, not only of one-particle states, but of all of nature. Since all charges q are inverted, the transformation turns particles into anti-particles (and vice versa). This actually *defines* antiparticles: Under a CPT transformation, every particle turns into its anti-particle. Particle/antiparticle pairs have identical mass and spin, but opposite charges.

Interpretation. This is the *Feynman-Stückelberg* interpretation: The negative-energy states $\psi_{3,4}$ are re-interpreted as positive-energy *antiparticles* (that move *forward* in time). In view of this interpretation, it is sometimes convenient to invert the overall sign of p^μ for these states, such that $p_0 > 0$ for all states. One often rewrites the states $\psi_{3,4}$ as

$$\psi_{3,4} = u_{3,4}(p) e^{-ip_\mu x^\mu} \xrightarrow{p^\mu \rightarrow -p^\mu} \psi_{3,4} = v_{2,1}(p) e^{+ip_\mu x^\mu}, \quad (6.68)$$

where $p_0 > 0$ is *positive* in the right-hand expression, and

$$v_1(E, \mathbf{p}) = u_4(-E, -\mathbf{p}), \quad v_2(E, \mathbf{p}) = u_3(-E, -\mathbf{p}). \quad (6.69)$$

History. The inevitability of negative-energy states and their interpretation as (positive-energy) antiparticles led Dirac in 1931 to predict the existence of the electron's antiparticle, the *positron*. The positron was experimentally confirmed by Anderson (in cosmic ray measurements) in 1932, and this was one of the triumphs of quantum field theory at the time.

6.5 Chirality, Helicity, and Spin

Two further important concepts for elementary particles are *chirality* and *helicity*. Both are related to each other, and to spin.

Chirality. For the following discussion, it is useful to not consider the Dirac representation of the gamma matrices that we used above, but a slightly different representation called the *Weyl representation*, given by:

$$\gamma^0 = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}, \quad \gamma^i = \begin{pmatrix} 0 & \sigma_i \\ -\sigma_i & 0 \end{pmatrix}. \quad (6.70)$$

Again, every entry is a 2×2 submatrix. The matrices γ^i are the same as in the Dirac representation, only γ^0 is different. In this basis, the generators $S^{\mu\nu}$ of Lorentz transformations are block-diagonal:

$$S^{\mu\nu} = \frac{i}{4} [\gamma^\mu, \gamma^\nu] = \begin{pmatrix} * & 0 \\ 0 & * \end{pmatrix}, \quad (6.71)$$

$$\psi \mapsto e^{i\theta_{\mu\nu} S^{\mu\nu}} \psi = \begin{pmatrix} \exp((i\theta_i - \beta_i)\frac{1}{2}\sigma_i) & 0 \\ 0 & \exp((i\theta_i + \beta_i)\frac{1}{2}\sigma_i) \end{pmatrix} \psi.$$

Here, θ_i and β_i , $i = 1, 2, 3$ parametrize the three rotations and the three boosts that make up the Lorentz group. One can see that the first two components of ψ only transform among each other, and similarly the last two components. This means that the four-dimensional

representation $S^{\mu\nu}$ is *reducible*, it consists of two two-dimensional representations (that are irreducible). One can further see that the two blocks of $S^{\mu\nu}$ are *independent* fundamental representations of $\mathfrak{su}(2)$: $\theta_i^+ = i\theta_i - \beta_i$ and $\theta_i^- = i\theta_i + \beta_i$ are 3 + 3 independent parameters. This means that the Lorentz algebra $\mathfrak{so}(1, 3)$ of the Lorentz group $\text{SO}(1, 3)$ is a direct sum of two copies of the Lie algebra $\mathfrak{su}(2)$:

$$\mathfrak{so}(1, 3) = \mathfrak{su}(2) \oplus \mathfrak{su}(2). \quad (6.72)$$

The decomposition can be seen explicitly by combining the boost generators K_i and the rotation generators J_i into

$$J_i^+ := \frac{1}{2} (J_i + iK_i), \quad J_i^- := \frac{1}{2} (J_i - iK_i), \quad (6.73)$$

which separately satisfy the commutation relations of $\mathfrak{su}(2)$,

$$[J_i^+, J_j^+] = i\varepsilon_{ijk} J_k^+, \quad [J_i^-, J_j^-] = i\varepsilon_{ijk} J_k^-, \quad [J_i^+, J_j^-] = 0. \quad (6.74)$$

The finite-dimensional unitary representations of $\mathfrak{su}(2)$ are labeled by half integers

$$j = \{0, \frac{1}{2}, 1, \frac{3}{2}, 2, \dots\}, \quad (6.75)$$

and therefore the finite-dimensional unitary representations of the Lorentz algebra are labeled by pairs of half-integers (j_1, j_2) . The Weyl representation (6.71) is the $(\frac{1}{2}, 0) \oplus (0, \frac{1}{2})$ representation, since it consists of two independent spin-1/2 representations of the spin group $\text{SU}(2)$. The spinor ψ is correspondingly called a *Weyl spinor*. Because the representation is block-diagonal, the spinor ψ decomposes into two two-component spinors,

$$\psi = \begin{pmatrix} \psi_L \\ \psi_R \end{pmatrix}, \quad (6.76)$$

that transform separately under Lorentz transformations: ψ_L transforms in the $(1/2, 0)$ representation and is called a *left-handed* spinor, ψ_R in the $(0, 1/2)$ representation and is called a *right-handed* spinor. The handedness of a spinor is called its *chirality*. The fact that any four-component spinor ψ , and hence any fermion, can be separated into left-handed and right-handed parts ψ_L and ψ_R is one of the most important technical points in the structure of the Standard Model.

Equations of Motion. Let us look at the equations of motion for a free spinor. In the Weyl basis, the Dirac equation becomes

$$\begin{pmatrix} -m & i\sigma^\mu \partial_\mu \\ i\bar{\sigma}^\mu \partial_\mu & -m \end{pmatrix} \begin{pmatrix} \psi_L \\ \psi_R \end{pmatrix} = 0, \quad (6.77)$$

where

$$\sigma^\mu := (\sigma_0, \boldsymbol{\sigma}), \quad \bar{\sigma}^\mu := (\sigma_0, -\boldsymbol{\sigma}), \quad \sigma_0 := \mathbf{1}. \quad (6.78)$$

For plane-wave states $\psi = ue^{-ip_\mu x^\mu}$, this implies

$$\sigma^\mu p_\mu \psi_R = (E - \boldsymbol{\sigma} \cdot \mathbf{p}) \psi_R = m \psi_L, \quad (6.79)$$

$$\bar{\sigma}^\mu p_\mu \psi_L = (E + \boldsymbol{\sigma} \cdot \mathbf{p}) \psi_L = m \psi_R. \quad (6.80)$$

One can see that the mass term mixes left-handed and right-handed spinor states.

Massless Case: Helicity. In the *massless* case, left-handed and right-handed states do not mix with each other, and they are eigenstates of the *helicity* operator

$$\hat{h} = \frac{\boldsymbol{\sigma} \cdot \mathbf{p}}{|\mathbf{p}|} . \quad (6.81)$$

This operator measures the spin projection on the direction of motion (momentum), called *helicity*. In the massless case, the energy is $E = \pm|\mathbf{p}|$, and therefore the eigenvalues of $\hat{h}$ are ± 1 . Hence the left-handed and right-handed chirality eigenstates are also helicity eigenstates, with *opposite* eigenvalues. We therefore have the following:

	chirality	helicity
massless particles, $E = + \mathbf{p} $:	L	-1
	R	+1
massless antiparticles, $E = - \mathbf{p} $:	L	+1
	R	-1

(6.82)

Massive Case. In the *massive* case, the equations of motion mix left-handed and right-handed fields with each other. But the helicity operator still commutes with the Hamiltonian, hence helicity is conserved, and it can make sense to consider helicity eigenstates. But these are no longer the chirality eigenstates ψ_L and ψ_R . Also, the helicity of a massive particle can be changed by a Lorentz transformation by going to the rest frame and rotating. So helicity is not a good quantum label for massive particles.

Projection Operators. We have defined the left-handed and right-handed chirality states using the Weyl representation for the gamma matrices. Depending on the context / type of particles considered, other representations (like the Dirac representation) are more useful. We can define left-handed and right-handed chirality states independently of the basis by using the matrix

$$\gamma^5 := i\gamma^0\gamma^1\gamma^2\gamma^3 = \frac{i}{4!}\epsilon_{\mu\nu\sigma\tau}\gamma^\mu\gamma^\nu\gamma^\sigma\gamma^\tau . \quad (6.83)$$

In the Weyl basis,

$$\gamma^5 = \begin{pmatrix} -\mathbf{1} & 0 \\ 0 & \mathbf{1} \end{pmatrix} , \quad (6.84)$$

and therefore left-handed and right-handed spinors are eigenstates of γ^5 with eigenvalues ∓ 1 . We can define projection operators that project onto left-handed and right-handed states:

$$P_L = \frac{1 - \gamma^5}{2} , \quad P_R = \frac{1 + \gamma^5}{2} . \quad (6.85)$$

The fact that these are projection operators follows from the identities

$$P_L^2 = P_L , \quad P_R^2 = P_R , \quad P_L + P_R = 1 , \quad P_L P_R = 0 . \quad (6.86)$$

In the Weyl representation,

$$P_L = \begin{pmatrix} \mathbf{1} & 0 \\ 0 & 0 \end{pmatrix} , \quad P_R = \begin{pmatrix} 0 & 0 \\ 0 & \mathbf{1} \end{pmatrix} , \quad (6.87)$$

and therefore

$$P_L \begin{pmatrix} \psi_L \\ \psi_R \end{pmatrix} = \begin{pmatrix} \psi_L \\ 0 \end{pmatrix}, \quad P_R \begin{pmatrix} \psi_L \\ \psi_R \end{pmatrix} = \begin{pmatrix} 0 \\ \psi_R \end{pmatrix}. \quad (6.88)$$

Independently of the gamma matrix representation, the eigenstates of P_L and P_R are chirality eigenstates that transform in the $(\frac{1}{2}, 0)$ and the $(0, \frac{1}{2})$ representations of the Lorentz group. Hence the projectors P_L and P_R can always be used to separate spinors into their left-handed and right-handed parts.

Parity. The concept of chirality is related to *parity*. As we saw above, the parity operation inverts all spatial directions, $\mathbf{x} \rightarrow -\mathbf{x}$, but leaves the time direction unchanged. Under this operation, momenta switch sign: $\mathbf{p} \rightarrow -\mathbf{p}$. Recall from above that the equations of motion for a free massless particle were

$$(E + \boldsymbol{\sigma} \cdot \mathbf{p})\psi_L = 0, \quad (E - \boldsymbol{\sigma} \cdot \mathbf{p})\psi_R = 0. \quad (6.89)$$

Clearly the parity operation flips the signs in these two equations, and therefore interchanges the meaning of left-handed and right-handed states. Hence parity interchanges the $(\frac{1}{2}, 0)$ and $(0, \frac{1}{2})$ representations of the Lorentz group:

$$P: \quad (\tfrac{1}{2}, 0) \leftrightarrow (0, \tfrac{1}{2}), \quad \psi_L \leftrightarrow \psi_R. \quad (6.90)$$

Summary. To summarize, we have seen that the left-handed and right-handed *chirality* states ψ_L and ψ_R

- do not mix under Lorentz transformations: They transform in separate irreducible representations,
- each have two components, which are the two spin states of the spin 1/2 particle (e. g. the electron),
- are helicity eigenstates (only) in the massless limit.

Let us summarize the three spin-related quantities considered above:

Spin is a vector quantity, i. e. it has a direction. It is the eigenvalue of the spin operator $\mathbf{S} = \boldsymbol{\sigma}/2$ (for fermions). One also calls the spin s the number (scalar) that defines the eigenvalue $s(s+1)$ of $\mathbf{S}^2$. For example, fermions have $s = 1/2$ and are therefore called spin 1/2 particles.

Helicity is the projection of the spin on the direction of motion. Helicity eigenstates exist for any spin. The two helicity eigenstates of the photon (which is a massless spin-1 particle) are what we call circularly polarized light.

Chirality is a concept that only exists for spinors. More precisely, it only exists for representations (j_1, j_2) of the Lorentz group with $j_1 \neq j_2$. Something is *chiral* if it is not symmetric under spatial reflections (e. g. DNA molecules), that is under the *parity* transformation. The parity transformation interchanges the (j_1, j_2) with the (j_2, j_1) representation, and therefore exchanges left-handed spinors ψ_L and right-handed spinors ψ_R .

6.6 The Dirac Lagrangian

Field theories are typically described in terms of Lagrangians. Quarks and leptons are spin-1/2 fermions and therefore obey the Dirac equation of motion. The *Lagrangian* that gives rise to the Dirac equation is

$$\mathcal{L} = \bar{\psi}(i\gamma^\mu\partial_\mu - m)\psi \quad (6.91)$$

Even without knowing the Dirac equation, this Lagrangian is very constrained: Starting with a spinor ψ , the only CPT-invariant Lorentz scalars that one can write are $\bar{\psi}\psi$ and $\bar{\psi}\gamma^\mu\partial_\mu\psi$. So the only degree of freedom is the relative coefficient between the two terms. Now, we have seen earlier that the Dirac equation implies that each component of ψ satisfies the Klein–Gordon equation, which is equivalent to the relativistic energy condition $E^2 = \mathbf{p}^2 + m^2$. Imposing this condition thus fixes the relative coefficient in the Lagrangian. This in fact provides an alternative way of deriving the Dirac equation.

6.7 Chirality in Field Theory

Which of spin, helicity, or chirality is important depends on the physical context. For free massless spinors, the spin eigenstates are also helicity and chirality eigenstates. We call theories that are symmetric under parity transformations *non-chiral*. In such cases, it is mostly unnecessary to separately consider the chirality eigenstates. Theories that are not parity-symmetric are called *chiral theories*. In such theories, it is important to distinguish left-handed and right-handed states.

We can see how left-handed and right-handed states may interact by looking at typical terms that can appear in a Lagrangian.

Mass Terms. For example, mass terms have the form $m\bar{\psi}\psi$. we can expand such terms into left-handed and right-handed parts by using the chirality projection operators P_L and P_R :

$$\bar{\psi}\psi = \bar{\psi}(P_L^2 + P_R^2)\psi = \bar{\psi}P_LP_L\psi + \bar{\psi}P_RP_R\psi. \quad (6.92)$$

We note that

$$\bar{\psi}_L = (P_L\psi)^\dagger\gamma^0 = \psi^\dagger P_L\gamma^0 = \psi^\dagger\gamma^0 P_R = \bar{\psi}P_R, \quad \bar{\psi}_R = \bar{\psi}P_L. \quad (6.93)$$

The mass term hence becomes

$$\bar{\psi}\psi = \bar{\psi}_R\psi_L + \bar{\psi}_L\psi_R. \quad (6.94)$$

A mass term therefore flips the chirality of a particle. This is consistent with the fact that the mass term in the equations of motion mixes left-handed and right-handed spinors. Still, ψ_L and ψ_R appear symmetrically, hence *mass terms are non-chiral*.

Currents. How about interactions? Standard interactions between matter particles and gauge bosons occur through currents. We saw earlier that fermion currents are of the form $\bar{\psi}\gamma^\mu\psi$. Again, this can be expanded into left-handed and right-handed parts. Using the projection operators, current terms can be written as

$$\bar{\psi}\gamma^\mu\psi = \bar{\psi}(P_L + P_R)\gamma^\mu(P_L + P_R)\psi. \quad (6.95)$$

From the definition of γ^5 , it follows that $\{\gamma^5, \gamma^\mu\} = 0$. Therefore,

$$P_L \gamma^\mu = \gamma^\mu P_R, \quad P_R \gamma^\mu = \gamma^\mu P_L. \quad (6.96)$$

Using the projection operator property $P_L P_R = P_R P_L = 0$, we see that only the mixed terms survive, and the current hence becomes

$$\begin{aligned} \bar{\psi} \gamma^\mu \psi &= \bar{\psi} P_R \gamma^\mu P_L \psi + \bar{\psi} P_L \gamma^\mu P_R \psi \\ &= \bar{\psi}_L \gamma^\mu \psi_L + \bar{\psi}_R \gamma^\mu \psi_R. \end{aligned} \quad (6.97)$$

This means that *currents preserve the chirality*, and all interactions of the form $\bar{\psi} \gamma^\mu \psi$ are parity-symmetric. This is for example the case in electrodynamics, where the interaction term is

$$\bar{\psi} \gamma^\mu A_\mu \psi = A_\mu \bar{\psi} \gamma^\mu \psi = A_\mu (\bar{\psi}_L \gamma^\mu \psi_L + \bar{\psi}_R \gamma^\mu \psi_R). \quad (6.98)$$

Electrodynamics is therefore *non-chiral*, that is it is symmetric under the parity operation that swaps left-handed and right-handed spinors.

Chiral Currents. What happens when only the left-handed part of a current appears in an interaction? In this case, one has

$$\bar{\psi}_L \gamma^\mu \psi_L = \bar{\psi} P_R \gamma^\mu P_L \psi = \bar{\psi} \gamma^\mu P_L \psi = \frac{1}{2} \bar{\psi} \gamma^\mu (1 - \gamma^5) \psi \quad (6.99)$$

Let us see how the two terms transform under parity transformations. Parity inverts all spatial directions, and exchanges left-handed and right-handed spinors. In the Weyl representation, one can easily see that this means

$$\begin{aligned} \text{Parity:} \quad \psi(t; \mathbf{x}) &\rightarrow \gamma^0 \psi(t; -\mathbf{x}), \\ \bar{\psi}(t; \mathbf{x}) &= \psi^\dagger(t; \mathbf{x}) \gamma^0 \rightarrow \psi^\dagger(t; -\mathbf{x}) \gamma^0 \gamma^0 = \bar{\psi}(t; -\mathbf{x}) \gamma^0. \end{aligned} \quad (6.100)$$

This transformation rule is in fact independent of the representation. The first term in the current therefore transforms as

$$\bar{\psi} \gamma^\mu \psi(t; \mathbf{x}) \rightarrow \bar{\psi} \gamma^0 \gamma^\mu \gamma^0 \psi(t; -\mathbf{x}) = \bar{\psi} (\gamma^\mu)^\dagger \psi(t; -\mathbf{x}). \quad (6.101)$$

Recall that

$$(\gamma^0)^\dagger = \gamma^0, \quad (\gamma^i)^\dagger = -\gamma^i. \quad (6.102)$$

Therefore,

$$\bar{\psi} \gamma^0 \psi \rightarrow \bar{\psi} \gamma^0 \psi, \quad \bar{\psi} \gamma^i \psi \rightarrow -\bar{\psi} \gamma^i \psi, \quad (6.103)$$

which is the transformation rule of a vector, and hence $\bar{\psi} \gamma^\mu \psi$ is a true Lorentz vector. For the second term in the current, we need to use that $\{\gamma^5, \gamma^\mu\} = 0$, which follows from the definition $\gamma^5 = i\gamma^0 \gamma^1 \gamma^2 \gamma^3$. One finds

$$\bar{\psi} \gamma^\mu \gamma^5 \psi \rightarrow \bar{\psi} \gamma^0 \gamma^\mu \gamma^5 \gamma^0 \psi = -\bar{\psi} \gamma^0 \gamma^\mu \gamma^0 \gamma^5 \psi = -\bar{\psi} (\gamma^\mu)^\dagger \gamma^5 \psi. \quad (6.104)$$

Compared to (6.101), we find an additional sign. Hence, $\bar{\psi} \gamma^\mu \gamma^5 \psi$ transforms like a vector, but picks up an extra sign under parity transformations. It is therefore a *pseudovector*, also called *axial vector*. The left-handed current (6.99) is therefore also called V–A current.

Weak Interactions. Interactions with such V–A currents indeed appear in the weak interactions, hence the weak interaction is *chiral*. This shows that the Standard Model treats left-handed and right-handed fermions differently, and nature is not invariant under the parity transformation. In particular, there are left-handed neutrinos, but no right-handed neutrinos have ever been observed.

6.8 Coupling to the Photon

We have studied in detail the theory of a free spin-1/2 fermion. How do such particles (like the electron) interact with photons?

Covariant Dirac Equation. The photon is a gauge boson, and its interactions follow from gauge theory. As we learned earlier, to obtain a gauge-invariant theory, we have to use the covariant derivative. Under gauge transformations, spinors transform just like scalars, that is

$$\psi \rightarrow e^{-i\alpha} \psi. \quad (6.105)$$

Therefore, we have to use the same covariant derivative as for scalars,

$$D_\mu = \partial_\mu + ieA_\mu, \quad (6.106)$$

where e is the charge of the fermion (e. g. electron). Replacing the ordinary derivative in the Dirac Lagrangian with the covariant derivative, one finds

$$\mathcal{L} = \bar{\psi}(i\gamma^\mu D_\mu - m)\psi = \bar{\psi}(i\gamma^\mu \partial_\mu - e\gamma^\mu A_\mu - m)\psi. \quad (6.107)$$

Correspondingly, the Dirac equation becomes

$$0 = (i\gamma^\mu D_\mu - m)\psi = (i\gamma^\mu \partial_\mu - e\gamma^\mu A_\mu - m)\psi. \quad (6.108)$$

The term $-e\bar{\psi}\gamma^\mu A_\mu\psi$ provides the interaction between the spin 1/2 fermion and the electromagnetic field in the form of the vector potential A_μ .

Quadratic Equation. Something interesting happens when we compare this equation for a fermion to the Klein–Gordon equation for a *scalar* ϕ coupled to the photon field A_μ ,

$$\left((i\partial_\mu - eA_\mu)^2 - m^2\right)\phi = 0. \quad (6.109)$$

Multiplying (6.108) with $(i\gamma^\mu D_\mu + m)$ gives

$$0 = (i\gamma^\mu D_\mu + m)(i\gamma^\nu D_\nu - m)\psi = (\gamma^\mu \gamma^\nu iD_\mu iD_\nu + m^2)\psi. \quad (6.110)$$

In the first term, we can split both factors $\gamma^\mu \gamma^\nu$ and $iD_\mu iD_\nu$ into their symmetric and antisymmetric parts:

$$\begin{aligned} \gamma^\mu \gamma^\nu &= \frac{1}{2} \{\gamma^\mu, \gamma^\nu\} + \frac{1}{2} [\gamma^\mu, \gamma^\nu], \\ iD^\mu iD^\nu &= \frac{1}{2} \{iD^\mu, iD^\nu\} + \frac{1}{2} [iD^\mu, iD^\nu]. \end{aligned} \quad (6.111)$$

In the contraction $\gamma^\mu \gamma^\nu iD_\mu iD_\nu$, the cross terms vanish, leaving us with

$$0 = \left(\frac{1}{4} \{ \gamma^\mu, \gamma^\nu \} \{ iD_\mu, iD_\nu \} + \frac{1}{4} [\gamma^\mu, \gamma^\nu] [iD_\mu, iD_\nu] - m^2 \right) \psi \quad (6.112)$$

If the covariant derivatives D_μ were ordinary derivatives ∂_μ , the second term would vanish and we would recover the Klein–Gordon equation in the same form as for a scalar (6.109). But here, the second term becomes

$$\begin{aligned} [iD_\mu, iD_\nu] &= [i\partial_\mu - eA_\mu, i\partial_\nu - eA_\nu] \\ &= -e([i\partial_\mu, A_\nu] + [A_\mu, i\partial_\nu]) \\ &= -ei(\partial_\mu A_\nu - \partial_\nu A_\mu) = -eiF_{\mu\nu}, \end{aligned} \quad (6.113)$$

where we recognize the field strength tensor (cf. Problem 1.2)

$$F_{\mu\nu} = \partial_\mu A_\nu - \partial_\nu A_\mu. \quad (6.114)$$

Recalling that

$$\{ \gamma^\mu, \gamma^\nu \} = 2g^{\mu\nu}, \quad \frac{i}{4} [\gamma^\mu, \gamma^\nu] = S^{\mu\nu}, \quad (6.115)$$

we find

$$\begin{aligned} 0 &= (iD_\mu iD^\mu - eS^{\mu\nu} F_{\mu\nu} - m^2) \psi \\ &= ((i\partial_\mu - eA_\mu)^2 - eS^{\mu\nu} F_{\mu\nu} - m^2) \psi. \end{aligned} \quad (6.116)$$

Magnetic Moment. Compared to the equation (6.109) for a scalar, we find the extra term $-eS^{\mu\nu} F_{\mu\nu}$. What is this extra term? In the Weyl representation, the generators $S^{\mu\nu}$ take the form

$$S_{0i} = -\frac{i}{2} \begin{pmatrix} \sigma_i & 0 \\ 0 & -\sigma_i \end{pmatrix}, \quad S^{ij} = -\frac{1}{2} \varepsilon^{ijk} \begin{pmatrix} \sigma_k & 0 \\ 0 & \sigma_k \end{pmatrix}. \quad (6.117)$$

The components of the field strength tensor are

$$F_{0i} = -E_i = E^i, \quad F_{ij} = -\varepsilon_{ijk} B^k. \quad (6.118)$$

Hence the extra term expands to (using that $\varepsilon^{ijk} \varepsilon_{ijl} = 2\delta_l^k$)

$$\begin{aligned} -eS^{\mu\nu} F_{\mu\nu} &= -2eS^{0i} F_{0i} - eS^{ij} F_{ij} \\ &= e \begin{pmatrix} (\mathbf{B} + i\mathbf{E}) \cdot \boldsymbol{\sigma} & 0 \\ 0 & (\mathbf{B} - i\mathbf{E}) \cdot \boldsymbol{\sigma} \end{pmatrix}. \end{aligned} \quad (6.119)$$

The term $\mathbf{B} \cdot \boldsymbol{\sigma}$ represents a magnetic dipole term. (The term $i\mathbf{E}$ is also a magnetic dipole term, in a boosted frame). In the non-relativistic limit, the equation (6.116) indeed becomes the Schrödinger–Pauli equation that describes spin-1/2 particles in electromagnetic fields. In the conventional normalization, the magnetic moment μ_e (coefficient of $\mathbf{B} \cdot \boldsymbol{\sigma}/2$) is $\mu_e = 2\mu_B$, with the *Bohr magneton* $\mu_B = e/(2m)$. Historically, the Pauli equation was inspired by experimental data for the electron. Here, we see that this equation inevitably follows for all spin 1/2 particles from the Dirac equation plus gauge theory principles. We therefore have a testable prediction: All charged fermions should have a magnetic dipole moment of size $\mu = 2q/(2m)$. Experimentally, the magnetic moment of the electron is indeed $\sim 2.002\mu_B$. The factor 2.002 is a quantum effect that is beautifully matched by quantum field theory corrections (see below).

Quantum Electrodynamics. By promoting the derivative in the Dirac Lagrangian to the gauge covariant derivative, we have obtained the interaction between light and matter in the form of the interaction term $-e\bar{\psi}\gamma^\mu A_\mu\psi$. The only further term that we have to add is the kinetic term for the electromagnetic field. We found earlier (cf. Problem 1.2) that this has the form $-1/4 F_{\mu\nu}F^{\mu\nu}$. Adding this term to the covariant Dirac Lagrangian (6.107), we get

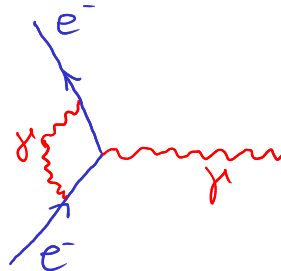
$$\begin{aligned}\mathcal{L} &= -\frac{1}{4}F_{\mu\nu}F^{\mu\nu} + i\bar{\psi}\gamma^\mu D_\mu\psi - m\bar{\psi}\psi, \\ &= -\frac{1}{4}F_{\mu\nu}F^{\mu\nu} + i\bar{\psi}\gamma^\mu\partial_\mu\psi - eA_\mu\bar{\psi}\gamma^\mu\psi - m\bar{\psi}\psi.\end{aligned}\tag{6.120}$$

This is the full Lagrangian of electrodynamics. The quantization of this Lagrangian provides a complete theoretical description of the interaction between light and electrons. It was awarded with the 1965 Nobel prize (Richard Feynman, Julian Schwinger, Shin'ichirō Tomonaga).

QED Precision. This theory is incredibly accurate: Experimental measurements and theoretical predictions agree to extreme precision. The most prominent example is the magnetic dipole moment μ_e of the electron. We saw above that the Dirac Lagrangian predicts

$$\mu_e = g_e\mu_B, \quad \mu_B \equiv \frac{e}{2m_e}, \quad g_e = 2,\tag{6.121}$$

with the *g-factor* g_e . Quantum effects lead to a correction of the *g-factor*, and therefore to the magnetic moment. The magnetic moment is one of the best measured physical quantities, it is known experimentally up to a relative uncertainty of $\sim 10^{-13}$. On the theoretical side, the first QED correction to g_e comes from the Feynman one-loop diagram


(6.122)

Further corrections come from further diagrams with more loops. Adding all diagrams with up to four loops (the current computational limit) gives a prediction for g_e in terms of the fine structure constant $\alpha = e^2/(4\pi) \approx 1/137.036$. The latter is not computable, so it must be measured and is an input to the theory.

The most precise experimental value of α is again determined from measuring the *g-factor* and fitting it to the theoretical value expressed in terms of α . Hence to non-trivially check the theoretical prediction, one needs two independent measurements of the *g-factor*. Comparing the two independent values of α obtained from the two measurements gives identical values up to a relative uncertainty of 10^{-8} . The uncertainty comes from our current limits on the computational as well as the experimental side. The precise agreement makes QED the best confirmed quantum theory ever devised.

6.9 *Photon Polarizations and Helicity

This is a side note explaining in some detail why the helicity eigenstates of the photon are identified with circularly polarized light. It is not essential for the remainder of these notes, and can therefore safely be skipped.

The electromagnetic field is described by the vector potential A_μ , which is a massless spin-1 field. Let us first consider how many degrees of freedom (physically independent components) the field A_μ has. We will see that the four components of A_μ are reduced to three independent components by the equations of motion. One of those is a gauge degree of freedom, which leaves us with two physically independent degrees of freedom. The Lagrangian for a free field A_μ is

$$\mathcal{L} = -\frac{1}{4} F_{\mu\nu}^2, \quad F_{\mu\nu} = \partial_\mu A_\nu - \partial_\nu A_\mu. \quad (6.122)$$

The equations of motion that follow from this Lagrangian are

$$\square A_\mu - \partial_\mu \partial_\nu A^\nu = 0. \quad (6.123)$$

where $\square = \partial_\mu \partial^\mu$. These are the Maxwell equations. Separating space and time components, the equations of motion expand to

$$-\nabla^2 A_0 + \partial_t \nabla \mathbf{A} = 0, \quad \square A_j - \partial_j (\partial_t A_0 - \nabla \mathbf{A}) = 0. \quad (6.124)$$

Importantly, the Lagrangian (6.122) is invariant under the *gauge transformations*

$$A_\mu(x) \mapsto A_\mu(x) + \partial_\mu \alpha(x) \quad (6.125)$$

for any function $\alpha(x)$: Two field configurations A_μ that differ by the derivative of a scalar are physically equivalent. We can use the freedom of re-defining A_μ via (6.125) to impose constraints on the components of A_μ . This procedure is called *gauge-fixing*. For the three-vector $\mathbf{A}$, the transformation (6.125) implies

$$\nabla \mathbf{A} \mapsto \nabla \mathbf{A} + \nabla^2 \alpha. \quad (6.126)$$

Hence we can always pick $\alpha(x)$ so that

$$\nabla \mathbf{A} = 0. \quad (6.127)$$

This is known as the *Coulomb gauge*. In this gauge, the equation of motion for A_0 becomes

$$\nabla^2 A_0 = 0. \quad (6.128)$$

Once we fix the Coulomb gauge, are there further gauge transformations that preserve this gauge? We could still apply gauge transformations with any α for which $\nabla^2 \alpha = 0$. This is the same constraint that A_0 satisfies. Within this class of functions, one can always find $\alpha(x)$ for which $\partial_t \alpha = -A_0$. Applying a further gauge transformation with this α gives

$$A_0 = 0. \quad (6.129)$$

Imposing this constraint in addition to the Coulomb gauge exhausts all gauge freedom, i. e. it completely fixes the gauge. The equation of motion for the spatial components A_j then becomes

$$\square A_j = 0. \quad (6.130)$$

Expanding in plane waves (by Fourier transformation),

$$A_\mu(x) = \int \frac{d^4 p}{(2\pi)^2} \varepsilon_\mu(p) e^{-ip_\nu x^\nu}, \quad (6.131)$$

with polarization vectors $\varepsilon_\mu(p)$, we find the constraints

$$\begin{aligned} p_\mu p^\mu &= 0 & (\text{equations of motion}), \\ \mathbf{p} \cdot \boldsymbol{\varepsilon} &= 0 & (\text{gauge choice}), \\ \varepsilon_0 &= 0 & (\text{gauge choice}). \end{aligned} \quad (6.132)$$

In particular, the second constraint implies that only polarizations transverse to the momentum $\mathbf{p}$ are physical. Without loss of generality, we consider the momentum

$$p^\mu = (E, 0, 0, E), \quad (6.133)$$

so that the particle moves along the z -direction. We see that the constraints (6.132) only have two independent solutions, and hence the field A_μ has only two physically independent components. A basis of solutions is

$$\varepsilon_1^\mu = (0, 1, 0, 0), \quad \varepsilon_2^\mu = (0, 0, 1, 0). \quad (6.134)$$

Using the expressions for the electric and magnetic fields

$$\mathbf{E} = \text{Re}(-\nabla V - \partial \mathbf{A} / \partial t), \quad \mathbf{B} = \text{Re}(\nabla \times \mathbf{A}), \quad (6.135)$$

one finds that plane-wave solutions

$$A_\mu(x) = \varepsilon_\mu(p) e^{-ip_\nu x^\nu} \quad (6.136)$$

with the above polarizations yield linearly polarized light. Another common basis are the circular polarizations

$$\varepsilon_L^\mu = \frac{1}{\sqrt{2}} (0, 1, -i, 0), \quad \varepsilon_R^\mu = \frac{1}{\sqrt{2}} (0, 1, +i, 0). \quad (6.137)$$

Plane waves with these polarizations yield circularly polarized light. Let us see how they are related to helicity. The helicity operator is

$$\hat{h} = \frac{\mathbf{S} \cdot \mathbf{p}}{|\mathbf{p}|}, \quad (6.138)$$

where $\mathbf{S}$ is the spin operator. Photons have spin one, which means that they transform in the 3×3 matrix representation of the spin group $\text{SU}(2)$. This is called the *vector* or *adjoint* representation of $\text{SU}(2)$. The generators in this representation are simply the $\text{SO}(3)$ rotation generators J_i , that is $\mathbf{S} = \mathbf{J}$. Since $\mathbf{p} = (0, 0, E)$, the helicity operator is

$$\hat{h} = J_3 = \begin{pmatrix} 0 & -i & 0 \\ i & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix}. \quad (6.139)$$

One now easily sees that the circular polarization vectors are at the same time the two helicity eigenstates:

$$\hat{h} \boldsymbol{\varepsilon}_L = -\boldsymbol{\varepsilon}_L, \quad \hat{h} \boldsymbol{\varepsilon}_R = +\boldsymbol{\varepsilon}_R. \quad (6.140)$$

7 The Standard Model Lagrangian

We are now ready to go into the formulation of the Standard Model. In the preceding sections, we have learned

- the rules of spinors and the Dirac equation, which are needed for spin-1/2 fermions (i. e. all matter particles),
- the idea of internal symmetry as a symmetry group that acts on multiplets of particles,
- the basic principles of gauge theory, which tell us to begin with a free-particle Lagrangian and rewrite it with a covariant derivative to find the interactions between matter particles and gauge bosons (force-carrying particles).

These are all the ingredients we need to write down and understand all terms in the Standard Model Lagrangian. Before we do that, let us briefly recall the basic meaning and features of Lagrangians, in particular how they give rise to the Feynman diagrams that describe particle interactions.

7.1 Lagrangians and Feynman Diagrams

Lagrangians. The dynamics of a local field theory are encoded in its action

$$S = \int L(t) dt , \quad L(t) = \int \mathcal{L}(t; x) d^3x , \quad (7.1)$$

where $\mathcal{L}$ is the Lagrangian. More accurately, L is the Lagrangian, and $\mathcal{L}$ is the Lagrangian density. For simplicity, $\mathcal{L}$ is often simply just called the Lagrangian. The Lagrangian (density) $\mathcal{L}$ is a local function of all fields and their derivatives, as well as all parameters of the theory. It is typically a polynomial consisting of various terms that describe the propagation and interaction of field quanta (particles). To understand this completely requires a course in quantum field theory. But we can still outline the basic qualitative picture without going into all the QFT details.

The most direct way to understand how the Lagrangian is connected to probability amplitudes is via the *path integral*. In quantum mechanics, the probability to find a system in a final state $|\psi_F\rangle$ at time t , assuming that it was in some initial state $|\psi_I\rangle$ at an earlier time t_0 is given by the transition amplitude

$$\langle \psi_F | e^{-i\hat{H}T} | \psi_I \rangle , \quad (7.2)$$

where $T = t - t_0$, and we assumed that the Hamiltonian $\hat{H}$ is time-independent. Inserting two complete bases of position states $1 = \int dq |q\rangle \langle q|$, one finds as a central object the *propagator*

$$\langle q_F | e^{-i\hat{H}T} | q_I \rangle . \quad (7.3)$$

In the quantum mechanical description of a particle, this propagator is the transition probability amplitude for the particle to travel from position q_I at time t_0 to position q_F at time t . Now, Feynman famously showed in 1948 that the propagator can be expressed as the *path integral*

$$\langle q_F | e^{-i\hat{H}T} | q_I \rangle = \int \mathcal{D}q e^{i/\hbar S[q]} = \int \mathcal{D}q e^{i/\hbar \int dt L(q(t), \dot{q}(t))} . \quad (7.4)$$

Here, $\int \mathcal{D}q$ is the integral over *all paths* $q(t)$ between $q(t_0) = q_I$ and $q(t) = q_F$, and $S[q]$ is the action of the theory evaluated on the path q . This path integral expression for the propagator can be obtained by discretizing $T = N \cdot \delta t$ into small time intervals δt , and taking the $\delta t \rightarrow 0$ limit. The path integral expression shows that all paths leading from point q_I to point q_F contribute to the amplitude with equal magnitude, only with different phase factors. In the classical limit, where the action S is much larger than $\hbar$, the classical path that extremizes S strongly dominates, as all other paths are suppressed by interference. But in the quantum regime, where $S \sim \hbar$, virtually all paths become important. The intuitive picture is that the quantum system explores *all possible histories* between the two states at times t_0 and t .

The path integral generalizes to quantum field theory: The transition amplitude from an initial field configuration $\phi(t_0, \mathbf{x}) = \phi_I(\mathbf{x})$ at time t_0 to a final field configuration $\phi(t, \mathbf{x}) = \phi_F(\mathbf{x})$ at time t is given by

$$\langle \phi(t, \mathbf{x}) | e^{-i\hat{H}T} | \phi(t_0, \mathbf{x}) \rangle = \int \mathcal{D}\phi e^{i/\hbar S[\phi]} = \int \mathcal{D}\phi e^{i/\hbar \int d^4x \mathcal{L}(\phi, \partial\phi)}, \quad (7.5)$$

where now $\int \mathcal{D}\phi$ is the integral over all field configurations $\phi(t, x)$ with boundary conditions given by the initial and final field configurations: $\phi(t_0, \mathbf{x}) = \phi_I(\mathbf{x})$ and $\phi(t, \mathbf{x}) = \phi_F(\mathbf{x})$, and the action S is the integral over the Lagrangian density $\mathcal{L}$. Again, the system explores all possible histories between the initial and final states, and all these possible histories contribute to the total transition amplitude.

Feynman Diagrams. How the path integral translates to probabilities for all kinds of processes via Feynman rules is the subject of quantum field theory. It goes beyond the scope of this course, but we can outline the general schematics: For a free-field Lagrangian $\mathcal{L}_{\text{free}}$, it is a theorem (Wick's theorem) that the amplitude between a number of incoming field quanta and a number of outgoing field quanta is given by the sum of all ways of connecting the incoming and outgoing quanta pairwise with propagator factors. For example, with two incoming and two outgoing particles, the amplitude is the sum

$$\langle p_3, p_4 | p_1, p_2 \rangle = \begin{array}{c} 3 \\ | \\ 1 \end{array} \begin{array}{c} 4 \\ | \\ 2 \end{array} + \begin{array}{c} 3 \quad 4 \\ \diagdown \quad \diagup \\ 1 \quad 2 \end{array} \quad (7.6)$$

In the free theory, summing over these propagator connections exhausts all possible histories of the system. The free Lagrangian $\mathcal{L}_{\text{free}}$ contains the kinetic terms ($\partial_\mu \phi \partial^\mu \phi$ for scalars, $i\bar{\psi} \gamma^\mu \partial_\mu \psi$ for fermions) and the mass terms ($m^2 \phi^2$ for scalars, $m\bar{\psi}\psi$ for fermions), and these are accounted for by the propagators in (7.6). The propagators for scalars and fermions are

$$\langle \phi_2 | \phi_1 \rangle = \frac{\delta^4(p_1 - p_2)}{p_1^2 - m^2}, \quad \langle \psi_2 | \psi_1 \rangle = \frac{\delta^4(p_1 - p_2)(\gamma^\mu p_\mu + m)}{p^2 - m^2}. \quad (7.7)$$

The delta functions imply that the amplitude is only non-zero if the set of incoming momenta equals the set of outgoing momenta.

Besides the free part $\mathcal{L}_{\text{free}}$, the Lagrangian of any *interacting* theory $\mathcal{L} = \mathcal{L}_{\text{free}} + \mathcal{L}_{\text{int}}$ has an interaction part $\mathcal{L}_{\text{int}}$ that is composed of terms that are at least cubic in the fields. For example, we saw that the Lagrangian of quantum electrodynamics contains a term

$$e A_\mu \bar{\psi} \gamma^\mu \psi, \quad (7.8)$$

which is an interaction between the fermion ψ (the electron) and a gauge boson A_μ (the electromagnetic field / photon). In quantum field theory, all fields in the Lagrangian become creation or annihilation operators. For example, ψ becomes an operator that creates a quantum of the electron field, and $\bar{\psi}$ becomes the corresponding annihilation operator. The interaction term (7.8) therefore absorbs a fermion (via the annihilation operator $\bar{\psi}$), emits a fermion (via the creation operator ψ), and absorbs or emits a photon – since the photon is its own antiparticle, the operator A_μ is both creation and annihilation operator for the photon field. Graphically, this can be depicted as:

$$eA_\mu\bar{\psi}\gamma^\mu\psi \sim \begin{array}{c} \bar{\psi} \\ \uparrow \\ \text{e}\gamma^\mu \\ \uparrow \\ \psi \end{array} \begin{array}{c} \text{red wavy line} \\ \downarrow \\ A_\mu \end{array} \quad \text{or} \quad \begin{array}{c} \bar{\psi} \\ \uparrow \\ \text{e}\gamma^\mu \\ \uparrow \\ \psi \end{array} \begin{array}{c} \text{red wavy line} \\ \uparrow \\ A_\mu \end{array} \quad (7.9)$$

In the amplitude (7.5)

$$\int \mathcal{D}\phi e^{i/\hbar \int d^4x (\mathcal{L}_{\text{free}} + \mathcal{L}_{\text{int}})}, \quad (7.10)$$

the interaction part can be Taylor-expanded:

$$e^{i \int d^4x \mathcal{L}_{\text{int}}} = 1 + i \int d^4x \mathcal{L}_{\text{int}} + \frac{1}{2} \left(i \int d^4x \mathcal{L}_{\text{int}} \right)^2 + \dots \quad (7.11)$$

Each term in the interaction Lagrangian comes with a coefficient called the *coupling*, and the Taylor expansion (7.11) makes sense as long as this coupling constant is small. For electrodynamics, the coupling constant is e , the charge of the electron. The various powers of $\mathcal{L}_{\text{int}}$ in the Taylor expansion (7.11) imply that any number of interaction terms (at least cubic in the fields) can be inserted into the possible histories of the system. As in the free case of propagator connections (7.6), the possible histories can be represented graphically in terms of *Feynman diagrams*. Each interaction term becomes an *interaction vertex* as in (7.9), and all interaction vertices are connected by propagators (from the free part $\mathcal{L}_{\text{free}}$). Each Feynman diagram then consists of any number of interaction vertices connected by propagators in any way possible. Each such diagram represents one possible history of the system. The total transition amplitude between two states is therefore given by the sum of all possible Feynman diagrams (histories) between these two states. For example, to describe electron-electron scattering, we have to compute the transition amplitude between two two-electron states, which expands as follows:

$$\langle \psi_3, \psi_4 | \psi_1, \psi_2 \rangle = \begin{array}{c} \begin{array}{ccc} \begin{array}{c} 3 \\ | \\ 1 \end{array} & \begin{array}{c} 4 \\ | \\ 2 \end{array} & \begin{array}{c} 3 \\ | \\ 1 \end{array} \end{array} + \begin{array}{c} \begin{array}{ccc} \begin{array}{c} 3 \\ \diagdown \\ 1 \end{array} & \begin{array}{c} 4 \\ \diagup \\ 2 \end{array} & \begin{array}{c} 3 \\ \diagup \\ 1 \end{array} \end{array} + \begin{array}{c} \begin{array}{ccc} \begin{array}{c} 3 \\ \diagdown \\ 1 \end{array} & \begin{array}{c} 4 \\ \diagup \\ 2 \end{array} & \begin{array}{c} 3 \\ \diagup \\ 1 \end{array} \end{array} \\ + \begin{array}{c} \begin{array}{ccc} \begin{array}{c} 3 \\ | \\ 1 \end{array} & \begin{array}{c} 4 \\ | \\ 2 \end{array} & \begin{array}{c} 3 \\ | \\ 1 \end{array} \end{array} + \begin{array}{c} \begin{array}{ccc} \begin{array}{c} 3 \\ \diagdown \\ 1 \end{array} & \begin{array}{c} 4 \\ \diagup \\ 2 \end{array} & \begin{array}{c} 3 \\ \diagup \\ 1 \end{array} \end{array} + \dots \end{array} \quad (7.12)$$

Since each interaction vertex is weighted by the (small) coupling constant, terms with more vertices are suppressed by small numerical prefactors. Hence the terms with the fewest possible number of vertices are the dominant terms.

Summary of Free Lagrangians. To close this general discussion of Lagrangians, let us recall the Lagrangians for the various types of free particles:

- For a *real spin-zero field of mass m* , the Lagrangian is

$$\mathcal{L} = \frac{1}{2} (\partial_\mu \phi \partial^\mu \phi - m^2 \phi^2). \quad (7.13)$$

- For a *complex spin-zero field of mass m* , the Lagrangian is

$$\mathcal{L} = (\partial_\mu \phi)^* (\partial^\mu \phi) - m^2 \phi^* \phi \quad (7.14)$$

- For a *spin 1/2 fermion field of mass m* , the Lagrangian is

$$\mathcal{L} = \bar{\psi} (i\gamma^\mu \partial_\mu - m) \psi. \quad (7.15)$$

- For an *Abelian vector field B_μ of mass m* , we use the field strength tensor

$$F_{\mu\nu} = \partial_\mu B_\nu - \partial_\nu B_\mu. \quad (7.16)$$

The Lagrangian then is

$$\mathcal{L} = -\frac{1}{4} F_{\mu\nu} F^{\mu\nu} + \frac{1}{2} m^2 B_\mu B^\mu, \quad (7.17)$$

- A *non-Abelian vector field W_a^μ* carries an additional internal index a . In the case of non-Abelian gauge fields (such as the gluons of QCD), this index is an adjoint index of the gauge group: The gauge field takes values in the Lie algebra of the gauge group, and hence can be expanded in the generators T^a of the gauge group as $W^\mu = W_a^\mu T^a$. In this case, the field strength tensor gets an additional term:

$$W_a^{\mu\nu} = \partial^\mu W_a^\nu - \partial^\nu W_a^\mu + g f_a^{bc} W_b^\mu W_c^\nu, \quad (7.18)$$

where g is the coupling constant, and f_a^{bc} are the structure constants of the Lie algebra. The Lagrangian for a field of mass m is a generalization of the Abelian case:

$$\mathcal{L} = -\frac{1}{4} W_a^{\mu\nu} W_{\mu\nu}^a + \frac{1}{2} m^2 W_a^\mu W_\mu^a. \quad (7.19)$$

Because of the extra term in the field strength tensor that is quadratic in the field W , the Lagrangian contains cubic and quartic terms, which means that the gauge field is self-interacting. This term is inevitable to make the theory gauge and Lorentz invariant, which shows that a non-Abelian gauge field is necessarily self-interacting.

This closes our general discussion of Lagrangians. Let us now take a look at the gauge group of the Standard Model.

7.2 The Gauge Group

The basic principle of gauge theory tells us that we have to start with the Lagrangian of free fermions (and/or possibly scalars), and promote the ordinary derivatives in the kinetic terms to covariant derivatives of the respective gauge groups. There is no principle that prescribes what the correct gauge groups are, these have to be inferred from experimental data. As discussed earlier, the gauge group of the Standard Model is the product group

$$U(1) \times SU(2) \times SU(3). \quad (7.20)$$

Let us look at the individual factors in turn:

- The $U(1)$ group is the group of rotations in the complex plane (by phase factors) that the Standard Model is invariant under. This $U(1)$ invariance is not directly the $U(1)$ invariance of electromagnetism, but is related to it. The gauge boson required by invariance under $U(1)$ transformations will be called B_μ . The symbol A_μ will be reserved for the photon field. The precise relation between this $U(1)$ and electromagnetism as well as the relation between B_μ and A_μ will be described later.
- The second factor in the gauge group is an $SU(2)$ group, and invariance under its transformations is responsible for the *weak* interactions. The invariance under the weak $SU(2)$ works analogously to the strong isospin invariance explained earlier. For this reason, the weak $SU(2)$ gauge transformations of the Standard Model are also called *weak isospin* transformations. The associated gauge bosons required for the invariance of the theory are the weak gauge bosons W_i^μ . As for all gauge fields, the index μ is required for the field to transform in the same way as the partial derivative under Lorentz transformations, which means that it has to be a spin-one field (this is true for all gauge fields).

There is one gauge boson for each of the three generators of $SU(2)$, hence $i = 1, 2, 3$. As in the case of strong isospin, the vector $(W_1^\mu, W_2^\mu, W_3^\mu)$ forms an adjoint multiplet (triplet) of the weak isospin $SU(2)$. In other words, the gauge boson W^μ has *weak isospin one* (a doublet would have weak isospin $1/2$). Keep in mind that weak isospin is unrelated to Lorentz spin! Weak isospin $SU(2)$ transformations are *internal* transformations, they act on an *internal structure* of the particles. Lorentz spin on the other hand is also an internal property of the particle, but it is an *angular momentum* that transforms if and only if we apply an *external* spacetime transformation.

As in the case of strong isospin, we pick a basis of particle states that are eigenstates of the third weak isospin generator T^3 . The eigenvalue of any particle under this generator is called the *weak isospin charge* of that particle. In the adjoint representation, the generators are 3×3 matrices, and the eigenvalues of T^3 are ± 1 and 0 . The eigenstates are

$$W^\pm = (-W_1 \pm iW_2)/\sqrt{2}, \quad W^0 = W_3. \quad (7.21)$$

As their names suggest, the states $W^\pm$ have weak isospin charges ± 1 , while W^0 has weak isospin charge 0 . As in the case of strong isospin, the W boson weak isospin eigenstates are at the same time the states with definite *electromagnetic charge*. The states $W^\pm$ have electromagnetic charges ± 1 (in units of the electron's charge), while W^0 has electromagnetic charge 0 .

Note that one must distinguish various types of charges: Electromagnetic charge, U(1) charge, weak isospin SU(2) charge, and “color” SU(3) charge. Every type of particle has a unique set of charges that, taken together, fully describe its internal (non-spacetime) properties.

- The third gauge group factor is SU(3), which gives another, independent set of internal transformations that are responsible for the strong nuclear interactions. Invariance under these transformations again require the existence of the associated gauge bosons. They are called *gluons* and are labeled G_a^μ , where $a = 1, 2, \dots, 8$, since there is one Gauge boson for each of the eight generators of SU(3). The theory of gluon interactions is called *quantum chromodynamics* (QCD).

The force that arises due to interactions with gluons is called the *strong force* or sometimes also *color force*, and it gives rise to the strong nuclear interactions. The associated internal charge that each particle carries and that determines its interaction with gluons is called the *color charge*. Since SU(3) is non-Abelian, its gauge bosons (the gluons) are necessarily self-interacting, which means that the gluons themselves carry color charge. (Just as the weak gauge bosons $W^\pm$ carry weak isospin charge.)

The strong force and the color charge have nothing to do with everyday color. The name “color” is used because some of the properties of color interactions are similar to properties of everyday color. For example, protons and neutrons each consist of three particles (quarks) that carry three different color charges in such a way that both protons and neutrons have no color charge (are color neutral). This is reminiscent of white light being composed of three primary colors.

7.3 Quark and Lepton States

Next, we have to specify how the various matter particles (quarks and leptons) transform under the different gauge group factors. In other words, we need to know in *which representations* of the three different gauge group factors the various matter particles transform, that is how they behave when acted upon by the gauge group generators. This determines the form of the new interaction terms that arise in the Lagrangian when the ordinary derivative is promoted to the covariant derivative.

Weak Representations. How particles behave under the weak SU(2) transformations is familiar from the theory of Lorentz (spacetime) spin, since particles with different spin also transform in different representations of SU(2). Remember that the SU(2) of spacetime spin is different from the weak SU(2): They are mathematically the same groups, but their physical transformations are completely independent from each other! To distinguish the two transformations groups, we could denote them as $SU(2)_{\text{spin}}$ and $SU(2)_W$. Coming back to the different representations of $SU(2)_{\text{spin}}$: Particles with *spin zero* are singlets, they transform trivially (i.e. do not transform) under $SU(2)_{\text{spin}}$. Particles with *spin 1/2* transform as doublets ($\uparrow, \downarrow$) in the fundamental representation of $SU(2)_{\text{spin}}$. Particles with *spin one* form triplets, they transform in the adjoint representation of $SU(2)_{\text{spin}}$. In the basis where the third generator $T^3 = J_z$ is diagonal, the three entries of the triplet have charges (eigenvalues) $J_z = (+1, -1, 0)$.

For the weak isospin $SU(2)_W$, we have already seen an example of a triplet: The weak gauge bosons W_i . How all of the observed particles making up our world transform in

the weak $SU(2)_W$ is an experimental question, it has to be determined by measurements. All known quarks and leptons are observed to be either weak $SU(2)_W$ singlets or parts of weak isospin doublets. The way in which the various particles are assigned to weak isospin singlets and doublets is a subtle and important aspect of the Standard Model.

Weak Multiplets: Leptons. We will start by looking at the leptons. Consider the electron state, given by a four-component Dirac spinor ψ_e . We have learned previously that every Dirac spinor can be separated into left-handed and right-handed chirality components, using the projectors P_L and P_R . Hence we can define

$$e_R^- = P_R \psi_e, \quad e_L^- = P_L \psi_e, \quad (7.22)$$

such that

$$\psi_e = e_L^- + e_R^-. \quad (7.23)$$

Similar separations into left-handed and right-handed parts are made for all other fermions. The superscript “ $-$ ” indicates the (negative) electric charge. Remarkably, the right-handed and left-handed states transform differently under the weak isospin $SU(2)$! Right-handed electrons are weak isospin singlets:

$$e_R^- = (SU(2) \text{ singlet}). \quad (7.24)$$

But left-handed electrons are components of weak isospin doublets

$$L := \begin{pmatrix} \nu_e \\ e^- \end{pmatrix}_L \equiv \begin{pmatrix} \nu_{e,L} \\ e_L^- \end{pmatrix}. \quad (7.25)$$

The doublet partner of the left-handed electron e_L^- is the left-handed electron neutrino $\nu_{e,L}$. Note that the L on the left-hand side of (7.25) stands for *lepton* (doublet), whereas the subscript L stands for left-handed. When L points “up” in weak isospin $SU(2)_W$ space, it represents the electron neutrino $\nu_{e,L}$, when it points down, it represents the left-handed electron e_L^- . Since the third generator of $SU(2)_W$ in the doublet representation (fundamental representation) is

$$T^3 = \frac{\tau^3}{2} = \frac{1}{2} \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}, \quad (7.26)$$

the electron neutrino $L_1 = \nu_{e,L}$ carries weak isospin charge $+1/2$, while the electron $L_2 = e_L^-$ carries weak isospin charge $-1/2$.

Weak isospin $SU(2)_W$ transformations rotate left-handed electrons into electron neutrinos and vice versa. The W gauge bosons are $\mathfrak{su}(2)_W$ algebra-valued, in particular the combinations $W^\pm$ act as raising and lowering operators that connect the two components of the lepton state L , just as angular momentum raising and lowering operators connect the spin-up and spin-down states. Physically, this means that W bosons can turn left-handed electrons into electron neutrinos and vice versa! This matches with their charges. For example, the W^+ has weak isospin charge $+1$ and electric charge $+1$, and can therefore lift the electron state with charges $(-1/2, -1)$ to the electron neutrino state with charges $(+1/2, 0)$.

The right-handed electron state e_R^- is a singlet and therefore not affected by weak $SU(2)$ rotations. Remarkably, *right-handed neutrino states have never been observed*. To our knowledge, they do not exist. Recall that left-handed and right-handed components of

any Dirac fermion ψ are only coupled to each other through the mass terms. Neutrinos have no (or at most an extremely tiny) mass, so it is consistent to only have left-handed but no right-handed neutrinos. We therefore often drop the subscript “L” for “left-handed” and just write ν_e instead of $\nu_{e,L}$.

Weak Multiplets: Quarks. Next, we consider the up and down quarks, represented by four-component spinors u and d . Again, we split them into left-handed and right-handed components

$$u_L = P_L u, \quad u_R = P_R u, \quad d_L = P_L d, \quad d_R = P_R d. \quad (7.27)$$

As for the leptons, the left-handed quark parts are combined in a weak isospin quark doublet

$$Q_L := \begin{pmatrix} u \\ d \end{pmatrix}_L \equiv \begin{pmatrix} u_L \\ d_L \end{pmatrix} \quad (7.28)$$

that transforms in the fundamental representation of weak isospin SU(2). The right-handed parts

$$u_R, \quad d_R \quad (7.29)$$

are weak isospin SU(2) singlets. So, as for the leptons, the left-handed up quark u_L has weak isospin charge $+1/2$, while the left-handed down quark d_L has weak isospin charge $-1/2$. The right-handed quarks have no weak isospin charge. Hence again the $W^\pm$ bosons can convert left-handed up quarks into left-handed down quarks and vice versa.

Beta Decay. We will discuss some physical implications of the weak interactions later, but can already make one remark: We noted that W bosons can convert up quarks to down quarks and electrons to neutrinos (and vice versa). This is how β -decay works! Neutrons have quark content (u, d, d) . By emitting a W^- , one of the down quarks can turn into an up quark, turning the neutron (u, d, d) into a proton (u, d, u) . The W^- then decays into an electron and an anti-neutrino:

β decay :

$$\quad (7.30)$$

Here, the right-handed anti-neutrino $\bar{\nu}_{e,R}$ is the anti-particle of the left-handed neutrino $\nu_e = \nu_{e,L}$. It is part of the anti-lepton doublet

$$\bar{L} := \begin{pmatrix} \bar{\nu}_e \\ e^+ \end{pmatrix}_R \equiv \begin{pmatrix} \bar{\nu}_{e,R} \\ e_R^+ \end{pmatrix}. \quad (7.31)$$

The second component $\bar{L}_2$ is the right-handed positron e_R^+ , which is the anti-particle of the left-handed electron. It has weak isospin $+1/2$, whereas $\bar{\nu}_{e,R}$ has weak isospin $-1/2$. The anti-particles not only have opposite charges as the corresponding particles, but also opposite chirality. We can understand this by recalling that any particle turns into its anti-particle under a CPT transformation (by definition). In particular, this involves a parity transformation (P), which, as we saw earlier, exchanges left-handed and right-handed chirality states.

Parity. Because the left-handed and right-handed components of the quark and lepton states transform in different multiplets of the weak isospin gauge group, parity symmetry is clearly broken: The Standard Model is a *chiral theory*. This is in accordance with experimental data (we will see that in more detail later), so nature is not invariant under parity transformations. In the Standard Model, this fact is beautifully accounted for by putting the particles in the respective multiplets. But of course this does not fundamentally *explain why* parity symmetry is broken. Understanding why parity symmetry is broken at a deeper level is still desirable and remains a goal of particle physics.

Color: Quarks. Next, let us turn to the color gauge group $SU(3)$. Besides the weak gauge group $SU(2)$, quarks also transform non-trivially under color $SU(3)$ transformations. Just as the fundamental $SU(2)$ representation is a doublet, the fundamental $SU(3)$ representation is a triplet. Both the up and the down quark form fundamental triplet representations, so they carry an additional index α that can take values $\alpha \in \{1, 2, 3\}$ and that makes both u and d three-dimensional vectors in “color space”. The three components are also called r, g, and b, for “red”, “green”, and “blue”. So the quark states really are

$$Q_{L,\alpha} = \begin{pmatrix} u_{L,\alpha} \\ d_{L,\alpha} \end{pmatrix}, \quad u_{R,\alpha}, \quad d_{R,\alpha}, \quad (7.32)$$

with $\alpha \in \{1, 2, 3\}$, or equivalently $\alpha \in \{r, g, b\}$. The color group $SU(3)$ has 8 generators. Two of these can be simultaneously diagonalized. In other words, the maximal Abelian subalgebra of the Lie algebra $\mathfrak{su}(3)$ is two-dimensional. In the standard basis, the two diagonal generators are

$$\frac{\lambda_3}{2} = \frac{1}{2} \begin{pmatrix} 1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & 0 \end{pmatrix}, \quad \frac{\lambda_8}{2} = \frac{1}{2\sqrt{3}} \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & -2 \end{pmatrix}. \quad (7.33)$$

The color charge of any vector in the three-dimensional color space therefore has two components, which are the eigenvalues under $\lambda_3/2$ and $\lambda_8/2$. The basis vectors

$$\mathbf{r} := \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}, \quad \mathbf{g} := \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix}, \quad \mathbf{b} := \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix} \quad (7.34)$$

have charges

$$q_r = \left(+\frac{1}{2}, +\frac{1}{2\sqrt{3}} \right), \quad q_g = \left(-\frac{1}{2}, +\frac{1}{2\sqrt{3}} \right), \quad q_b = \left(0, -\frac{1}{\sqrt{3}} \right). \quad (7.35)$$

All leptons (electrons and neutrinos) are color singlets, which means they transform trivially (not at all) under the color $SU(3)$ group, and hence they do not carry a color index α .

Gauge-invariant combinations of states have to be colorless, that is they have to have “white” color $(0, 0)$. In other words, they must form *color singlets*, i. e. trivial representations of the color group. Also, there are no solitary quarks, they always form bound states that are colorless. This phenomenon is called *confinement*. For example, protons and neutrons each consist of three quarks, one red, one green, and one blue, such that the total nucleon charge is $q_r + q_g + q_b = (0, 0)$.

For every quark, there is a corresponding anti-quark with opposite charges and opposite chirality. The anti-up and anti-down quark states therefore are

$$\bar{Q}_R := \begin{pmatrix} \bar{u}_{R,\bar{\alpha}} \\ \bar{d}_{R,\bar{\alpha}} \end{pmatrix}, \quad \bar{u}_{L,\bar{\alpha}}, \quad \bar{d}_{L,\bar{\alpha}}. \quad (7.36)$$

The anti-quarks transform in the *anti-fundamental* representation of SU(3), with generators

$$T_{\text{anti-fund}}^a = -(T_{\text{fund}}^a)^* = -\lambda_a^*/2, \quad (7.37)$$

where $T_{\text{fund}}^a = \lambda_a/2$ are the fundamental generators under which the up and down quarks transform. The color index of the anti-quarks takes values “anti-red”, “anti-green”, or “anti-blue”, that is $\bar{\alpha} \in \{\bar{r}, \bar{g}, \bar{b}\}$. The charges of the three basis anti-quarks are opposites of the corresponding quark charges (because the relevant generators are $T_{\text{anti-fund}}^3 = -\lambda_3/2$ and $T_{\text{anti-fund}}^8 = -\lambda_8/2$):

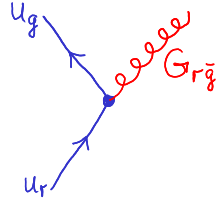
$$q_{\bar{r}} = \left(-\frac{1}{2}, -\frac{1}{2\sqrt{3}}\right), \quad q_{\bar{g}} = \left(+\frac{1}{2}, -\frac{1}{2\sqrt{3}}\right), \quad q_{\bar{b}} = \left(0, +\frac{1}{\sqrt{3}}\right). \quad (7.38)$$

In particular, this means that colorless objects can also be composed from a quark and the corresponding anti-quark. Such particles are called *mesons*. For example, a meson may contain a red quark and an anti-red anti-quark, whose color charges add up to zero.

Color: Gluons. Gluons are responsible for the strong interactions. They are the gauge bosons G^μ of the color gauge group factor SU(3), and therefore they transform in the adjoint representation of the color group, which means that they can be expanded in the generators $\lambda_a/2$ of the color group SU(3):

$$G^\mu = G_a^\mu \frac{\lambda_a}{2}. \quad (7.39)$$

The color group has 8 generators, hence there are 8 different gluons. We have seen that quarks and leptons can emit and absorb weak gauge bosons W^μ and thereby change their weak isospin charge. Similarly, quarks can emit and absorb gluons and thereby change their color charge. For example, a red quark may turn into a green quark by emitting a red/anti-green gluon:



$$(7.40)$$

Further Families. In all of the above discussion, we have only considered one family of leptons: (ν_e, e) , and one family of quarks: (u, d) . It is by now well established that there are two more families of both leptons and quarks. The two further lepton families are the muon with the muon neutrino (ν_μ, μ) and the tau with the tau neutrino (ν_τ, τ) . On the quark side, the second family (c, s) consists of the charm and the strange quark, and the top and bottom quarks form the third family (t, b) .

Remarkably, all of the discussion above exactly replicates for the two further families of particles. The second and third families fall into exactly the same weak and strong

multiplets as the first family, with the same weak and strong charges. In particular, *all three families interact via the same set of weak and strong gauge bosons*. The only distinction between the families are the different masses of the various particles.

A further remarkable fact is that the known universe is made entirely of particles of the first family: Electrons, up quarks, and down quarks. (Neutrinos are a bit special since they interact only extremely weakly with all other matter). All the heavier particles of the second and third families are only created at accelerators, or in cosmic ray collisions. They have very short lifetimes of $<10^{-6}$ s, and quickly decay into particles of the first family. No fundamental reason for the existence of the second and third families has yet been found.

In the following lectures, we will mostly restrict the discussion to the first family of particles. Since the other families behave identically (except for their masses), the theory is easily generalized to include the second and third families.

7.4 The Quark and Lepton Lagrangian

Now that we have reviewed the Standard Model gauge group, and have learned in which representations of the various gauge group factors each particle transforms, we are ready to state the Lagrangian for quarks and leptons. Quarks and leptons are fermions, so we have to take the kinetic energy term of the Dirac Lagrangian, and replace the ordinary derivative ∂_μ with the covariant derivative D_μ :

$$\bar{\psi}\gamma^\mu\partial_\mu\psi \rightarrow \bar{\psi}\gamma^\mu D_\mu\psi. \quad (7.41)$$

Here, ψ can be any fermion state. We will collectively denote the fermion states of the Standard Model by f . For the first family of leptons and quarks, f can be any of

$$f \in \{L, e_R, Q_L, u_R, d_R\}, \quad (7.42)$$

and we have to sum over terms of the form (7.41) for each possible f . Recall that the covariant derivative of the Standard Model has one term for each local gauge symmetry factor:

$$D_\mu = \partial_\mu - ig_1 \frac{Y}{2} B_\mu - ig_2 \frac{\tau^i}{2} W_{i\mu} - ig_3 \frac{\lambda^a}{2} G_{a\mu}. \quad (7.43)$$

Here, B_μ , $W_{i\mu}$, and $G_{a\mu}$ are the spin-one gauge fields required to maintain invariance under local U(1), SU(2), and SU(3) gauge transformations, respectively. $Y/2$ is the generator of U(1) gauge transformations, $\tau^i/2$, and $\lambda^a/2$ are the generators of SU(2) and SU(3) gauge transformations. As usual, the repeated indices i and a are summed over. Gauge invariance determines the form of each term in D_μ , but not its overall strength, which is represented by the *coupling constants* g_1 , g_2 , and g_3 . These have to be measured to match with experiment.

In order to write the Lagrangian in a compact form, we can make use of the following fact: Whenever the gauge boson terms in D_μ act on fermion states that form *singlets* of the respective gauge group factor, they give zero, by definition. Hence the gauge boson terms only give non-zero contributions when they act on fundamental multiplets of the respective gauge group factor. For example, $\tau^i W_i$ is a 2×2 matrix in SU(2) space, so it gives a non-zero result when acting on the fundamental doublets L and Q_L , but it gives zero when acting on the singlets e_R , u_R , and d_R . Similarly, $\lambda^a G_a$ is a 3×3 matrix in color space, so it acts non-trivially on the quark states Q_L , u_R , and d_R , but it gives zero when

acting on the color-neutral leptons L and e_R . With this convention, we can write the Standard Model Lagrangian for quarks and leptons of the first family in the compact form

$$\mathcal{L}_{\text{ferm}} = \sum_{\substack{f=L, e_R, \\ Q_L, u_R, d_R}} \bar{f} i \gamma^\mu D_\mu f. \quad (7.44)$$

To include the second and third families, we have to simply add two more terms of the same form, where the lepton and quark contents (e, ν_e, d, u) are replaced by (μ, ν_μ, s, c) and (τ, ν_τ, b, t) , respectively. Starting from this Lagrangian, all quark and lepton interactions can be calculated. All presently known experimental information on these interactions is consistent with the predictions from $\mathcal{L}_{\text{ferm}}$.

8 The Electroweak Theory

The Lagrangian $\mathcal{L}_{\text{ferm}}$ contains a lot of physical information. To extract this information and its connection to experimental observations, we study it piece by piece. The U(1) and SU(2) terms have the same form for leptons and quarks. Here, we first focus on the leptons.

8.1 Leptons: U(1) and SU(2) Terms

U(1) Terms. The U(1) interaction terms for the first family of leptons are

$$\mathcal{L}_{\text{leptons}}^{\text{U(1)}} = \bar{L} i \gamma^\mu \left(-i g_1 \frac{Y_L}{2} B_\mu \right) L + \bar{e}_R i \gamma^\mu \left(-i g_1 \frac{Y_R}{2} B_\mu \right) e_R. \quad (8.1)$$

Here, U(1) has only a single generator Y , which is just a number, called the (*weak*) *hypercharge*. However, this number may have different values for different fermions, which is why we introduced different labels Y_L and Y_R for the different terms. Since $g_1 Y_L B_\mu$ is just a number in SU(2) space, we can commute it past $\bar{L} \gamma^\mu$. The current $\bar{L} \gamma^\mu L$ then expands to

$$\bar{L} \gamma^\mu L = \bar{\nu}_L \gamma^\mu \nu_L + \bar{e}_L \gamma^\mu e_L. \quad (8.2)$$

The Lagrangian can hence be written as

$$\mathcal{L}_{\text{leptons}}^{\text{U(1)}} = \frac{g_1}{2} \left(Y_L (\bar{\nu}_L \gamma^\mu \nu_L + \bar{e}_L \gamma^\mu e_L) + Y_R \bar{e}_R \gamma^\mu e_R \right) B_\mu. \quad (8.3)$$

Before we can interpret this, we have to take a look at the SU(2) part as well.

SU(2) Terms. The term $\tau^i W_i$ is a 2×2 matrix that only has a non-trivial action on the left-handed doublet L :

$$\mathcal{L}_{\text{leptons}}^{\text{SU(2)}} = \bar{L} i \gamma^\mu \left(-i g_2 \frac{\tau^i}{2} W_{i\mu} \right) L. \quad (8.4)$$

Expanding the 2×2 matrix gives

$$\mathcal{L}_{\text{leptons}}^{\text{SU(2)}} = \frac{g_2}{2} \begin{pmatrix} \bar{\nu}_L & \bar{e}_L \end{pmatrix} \gamma^\mu \begin{pmatrix} W_{3\mu} & W_{1\mu} - i W_{2\mu} \\ W_{1\mu} + i W_{2\mu} & -W_{3\mu} \end{pmatrix} \begin{pmatrix} \nu_L \\ e_L \end{pmatrix}. \quad (8.5)$$

Going to the charge eigenstates $W^\pm = (-W_1 \pm iW_2)/\sqrt{2}$ and $W^0 = W_3$, this becomes

$$\begin{aligned}\mathcal{L}_{\text{leptons}}^{\text{SU}(2)} &= \frac{g_2}{2} \begin{pmatrix} \bar{\nu}_L & \bar{e}_L \end{pmatrix} \gamma^\mu \begin{pmatrix} W_\mu^0 & -\sqrt{2}W^+ \\ -\sqrt{2}W^- & -W_\mu^0 \end{pmatrix} \begin{pmatrix} \nu_L \\ e_L \end{pmatrix} \\ &= \frac{g_2}{2} \left(\bar{\nu}_L \gamma^\mu \nu_L W_\mu^0 - \sqrt{2} \bar{\nu}_L \gamma^\mu e_L W_\mu^+ \right. \\ &\quad \left. - \sqrt{2} \bar{e}_L \gamma^\mu \nu_L W_\mu^- - \bar{e}_L \gamma^\mu e_L W_\mu^0 \right).\end{aligned}\tag{8.6}$$

The full lepton Lagrangian is the sum of U(1) and SU(2) terms:

$$\mathcal{L}_{\text{leptons}} = \mathcal{L}_{\text{leptons}}^{\text{U}(1)} + \mathcal{L}_{\text{leptons}}^{\text{SU}(2)},\tag{8.7}$$

and it fully describes all interactions involving leptons (except for Higgs interactions, which will be explained later).

8.2 Neutral Currents

Let us see how the lepton interactions match with what we know from observations. We first focus on the *neutral current* terms, which are all terms involving the electrically neutral B_μ or W_μ^0 . The reason is that the electromagnetic interaction is also mediated by an electrically neutral field, the photon field A_μ . We know that the electromagnetic interactions of the electron are captured by an interaction term

$$-eA_\mu \bar{\psi}_e \gamma^\mu \psi_e = -eA_\mu (\bar{e}_L \gamma^\mu e_L + \bar{e}_R \gamma^\mu e_R),\tag{8.8}$$

where $-e$ is the charge of the electron. Several terms in our Lagrangian $\mathcal{L}_{\text{leptons}}$ are of this form. They somehow have to combine to recover the electromagnetic interaction. But first we note that there are also terms in $\mathcal{L}_{\text{leptons}}$ that couple the *neutrino current* $\bar{\nu}_L \gamma^\mu \nu_L$ to neutral gauge bosons, namely

$$\mathcal{L}_{\text{leptons}}^{\bar{\nu}\gamma\nu} = \left(\frac{g_1}{2} Y_L B_\mu + \frac{g_2}{2} W_\mu^0 \right) \bar{\nu}_L \gamma^\mu \nu_L.\tag{8.9}$$

But we know that neutrinos do not interact with the electromagnetic field A_μ . So unless we set $g_1 = g_2 = 0$ (which makes the theory useless), we must assume that the electromagnetic field A_μ is a combination of B_μ and W_μ^0 that is *orthogonal* to the combination in front of the neutrino current. In other words, we want to define a new basis in the (B_μ, W_μ^0) field space, such that one of the new basis fields A_μ has no coupling to the neutrino current, and can therefore be identified as the electromagnetic field. Orthonormality of the old fields (B, W^0) and the new fields (A, Z) requires that the two sets are related by an SO(2) rotation,

$$\begin{pmatrix} A_\mu \\ Z_\mu \end{pmatrix} = \begin{pmatrix} \cos \theta_W & \sin \theta_W \\ -\sin \theta_W & \cos \theta_W \end{pmatrix} \begin{pmatrix} B_\mu \\ W_\mu^0 \end{pmatrix}.\tag{8.10}$$

The rotation angle is called the *electroweak mixing angle* θ_W . Identifying the combination that couples to the neutrino current $\bar{\nu}_L \gamma^\mu \nu_L$ as Z_μ , comparing coefficients and normalizing gives

$$\sin \theta_W = -\frac{g_1 Y_L}{\sqrt{g_1^2 Y_L^2 + g_2^2}}, \quad \cos \theta_W = \frac{g_2}{\sqrt{g_1^2 Y_L^2 + g_2^2}}.\tag{8.11}$$

Let us see what this implies for the electron interactions. Collecting all terms in $\mathcal{L}_{\text{leptons}}$ that contain electron currents, we find

$$\mathcal{L}_{\text{leptons}}^{\bar{e}\gamma e} = \left(\frac{g_1}{2} Y_L B_\mu - \frac{g_2}{2} W_\mu^0 \right) \bar{e}_L \gamma^\mu e_L + \left(\frac{g_1}{2} Y_R B_\mu \right) \bar{e}_R \gamma^\mu e_R. \quad (8.12)$$

The old fields B_μ and W_μ^0 are expressed in terms of the new fields A_μ and Z_μ through the inverse transformation

$$\begin{pmatrix} B_\mu \\ W_\mu^0 \end{pmatrix} = \begin{pmatrix} \cos \theta_W & -\sin \theta_W \\ \sin \theta_W & \cos \theta_W \end{pmatrix} \begin{pmatrix} A_\mu \\ Z_\mu \end{pmatrix}. \quad (8.13)$$

Inserting this into (8.12) and collecting terms gives

$$\begin{aligned} \mathcal{L}_{\text{leptons}}^{\bar{e}\gamma e} = & A_\mu \left(\frac{g_1 g_2 Y_L}{\sqrt{g_2^2 + g_1^2 Y_L^2}} \bar{e}_L \gamma^\mu e_L + \frac{g_1 g_2 Y_R}{2\sqrt{g_2^2 + g_1^2 Y_L^2}} \bar{e}_R \gamma^\mu e_R \right) \\ & + Z_\mu \left(\frac{g_1^2 Y_L^2 - g_2^2}{2\sqrt{g_2^2 + g_1^2 Y_L^2}} \bar{e}_L \gamma^\mu e_L + \frac{g_1^2 Y_R Y_L}{2\sqrt{g_2^2 + g_1^2 Y_L^2}} \bar{e}_R \gamma^\mu e_R \right). \end{aligned} \quad (8.14)$$

To match this with the known electromagnetic interaction (8.8), we must identify both coefficients in the first line with the electric charge $-e$ of the electron, that is

$$-e = \frac{g_1 g_2 Y_L}{\sqrt{g_2^2 + g_1^2 Y_L^2}} = \frac{g_1 g_2 Y_R}{2\sqrt{g_2^2 + g_1^2 Y_L^2}}. \quad (8.15)$$

In particular, this fixes

$$Y_R = 2Y_L. \quad (8.16)$$

Since Y_L always appears with coefficient g_1 , we can without loss of generality set

$$Y_L = -1, \quad (8.17)$$

since any rescaling of Y_L can be absorbed by a re-definition of g_1 . Then

$$e = \frac{g_1 g_2}{\sqrt{g_2^2 + g_1^2}}, \quad (8.18)$$

and we have indeed recovered the known electromagnetic interaction of the electron, with a neutral (not electromagnetically interacting) neutrino! What we have also found is *another neutral gauge boson* Z_μ that interacts both with electrons and with neutrinos. Before we study those interactions, let us take another look at the coupling constants. With the fixed values of Y_R and Y_L , we have

$$\sin \theta_W = \frac{g_1}{\sqrt{g_1^2 + g_2^2}}, \quad \cos \theta_W = \frac{g_2}{\sqrt{g_1^2 + g_2^2}}, \quad (8.19)$$

such that

$$g_1 = \frac{e}{\cos \theta_W}, \quad g_2 = \frac{e}{\sin \theta_W}. \quad (8.20)$$

So the previously unknown couplings g_1 and g_2 have been replaced by the known electron charge e and the unknown *electroweak mixing angle* θ_W . The mixing angle θ_W has to be determined by comparison to measurements. Its experimental value is $\sin^2 \theta_W \approx 0.23$. The

relations between the couplings and the mixing angle can be visualized in the following triangle:

(8.21)

Let us re-examine the couplings of the neutrino and the electron to the new field Z . The neutrino term (8.9) becomes

$$\mathcal{L}_{\text{leptons}}^{\bar{\nu}\gamma\nu} = \frac{\sqrt{g_1^2 + g_2^2}}{2} Z_\mu \bar{\nu}_L \gamma^\mu \nu_L = \frac{g_2}{2 \cos \theta_W} Z_\mu \bar{\nu}_L \gamma^\mu \nu_L. \quad (8.22)$$

Hence the interaction vertex between Z_μ and the neutrino current $\bar{\nu}_L \gamma^\mu \nu_L$ comes with a factor $g_2/(2 \cos \theta_W)$. We can think of this factor as the “electroweak charge” of the neutrino. For the electron current, the coupling to Z_μ with our values for Y_R and Y_L becomes

$$\mathcal{L}_{\text{leptons}}^{\bar{e}\gamma e} = Z_\mu \left(\frac{g_1^2 - g_2^2}{2\sqrt{g_2^2 + g_1^2}} \bar{e}_L \gamma^\mu e_L + \frac{g_1^2}{\sqrt{g_2^2 + g_1^2}} \bar{e}_R \gamma^\mu e_R \right). \quad (8.23)$$

We can re-write the prefactors in terms of e and θ_W , using the identity

$$\sqrt{g_1^2 + g_2^2} = \frac{e}{\cos \theta_W \sin \theta_W}. \quad (8.24)$$

Then:

$$\begin{aligned} \frac{g_1^2 - g_2^2}{2\sqrt{g_2^2 + g_1^2}} &= \frac{e}{\cos \theta_W \sin \theta_W} \left(-\frac{1}{2} + \sin^2 \theta_W \right), \\ \frac{g_1^2}{\sqrt{g_2^2 + g_1^2}} &= \frac{e}{\cos \theta_W \sin \theta_W} \left(+\sin^2 \theta_W \right). \end{aligned} \quad (8.25)$$

We have written these couplings in a suggestive form: They can be unified to

$$\frac{e}{\sin \theta_W \cos \theta_W} \left(T_{3f} - Q_f \sin^2 \theta_W \right), \quad (8.26)$$

where T_{3f} is the eigenvalue of the diagonal generator $T_3 = \tau_3/2$, and Q_f is the electric charge of the respective fermion f (in units of e). Recall that in our convention, $T_3 = 0$ for the right-handed SU(2) singlets. In fact, one can check that the expression (8.26) also reproduces the correct coupling for the neutrinos. Moreover, since the couplings of the left-handed and right-handed quarks to the U(1) and SU(2) gauge bosons look exactly the same as for the leptons, their couplings to the rotated gauge fields A_μ and Z_μ work out in exactly the same way. Here, one has to take into account that the values for Y_L and Y_R can be different for the quarks than for the leptons. For generic eigenvalues T_{3f} and hypercharges Y_f , the coupling to A_μ (identified as the electromagnetic charge) works out to

$$eQ_f, \quad \text{with} \quad Q_f = T_{3f} + \frac{Y_f}{2}, \quad (8.27)$$

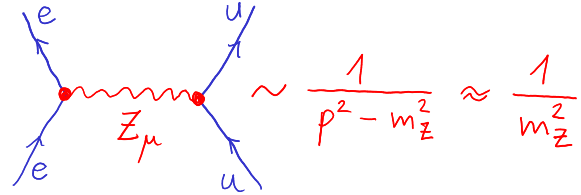
and the coupling to Z_μ works out to (8.26). Hence, the expression (8.26) gives the unified *electroweak charge* for all fermions.

Summarizing all terms, the full electroweak Lagrangian for the coupling of all quarks and leptons of the first family to the neutral bosons A_μ and Z_μ has the compact form

$$\mathcal{L}_{\text{EW}}^{\text{neutral}} = \sum_{\substack{f=L,e_R, \\ Q_L,u_R,d_R}} \left(e A_\mu \bar{f} Q_f \gamma^\mu f + \frac{e}{\sin \theta_W \cos \theta_W} Z_\mu \bar{f} (T_{3f} - Q_f \sin^2 \theta_W) f \right), \quad (8.28)$$

where $Q_f = T_{3f} + Y_f/2$, and keep in mind that $T_{3f} = 0$ for all right-handed fermions.

Let us summarize what we have accomplished. The electroweak theory contains the known electromagnetic interaction, *and* predicts an additional photon-like particle Z_μ called the *Z-boson* that interacts with any fermion f having non-zero electric charge Q_f or weak isospin T_{3f} . The strength of this interaction is not small. In fact, $1/(\sin \theta_W \cos \theta_W) > 1$, so the Z_μ interaction is stronger than the photon interaction! If the theory is correct, then why were the Z-boson and its interactions not discovered long ago? The only possible explanation is that the Z-boson is massive. The larger its mass m_Z , the more energy is needed to produce it, and the smaller are its effects at low energies: Whenever a Z_μ particle is exchanged between two fermions, the process is suppressed by a factor $\sim 1/m_Z^2$ from the Z propagator:



$$\sim \frac{1}{p^2 - m_Z^2} \approx \frac{1}{m_Z^2} \quad (8.29)$$

We will see later that the Z-boson can acquire a mass through the Higgs mechanism. Indeed, its mass can be predicted, and the Z-boson was detected in 1983 at exactly the expected mass of $m_Z \approx 91.2 \text{ GeV}$. In comparison, the proton mass is 938 MeV , so the Z-boson is ~ 100 times more heavy than the proton.

8.3 Charged Currents

We have analyzed all U(1) and SU(2) interaction terms with the neutral bosons B_μ and W_μ^0 . All these terms are diagonal in the fermions, that is the interaction does not change the fermion involved in the interaction ($\nu_L \rightarrow \nu_L$, $e_L \rightarrow e_L$, $e_R \rightarrow e_R$). What is left are the non-diagonal SU(2) terms. These are interactions with the charged bosons $W^\pm$. These act as raising and lowering operators on the fermion doublets. For the leptons, these terms are

$$\mathcal{L}_{\text{leptons}}^{\text{charged}} = -\frac{g_2}{\sqrt{2}} \left(\bar{\nu}_L \gamma^\mu e_L W_\mu^+ + \bar{e}_L \gamma^\mu \nu_L W_\mu^- \right). \quad (8.30)$$

Note that only left-handed electrons e_L are involved in the interaction with $W^\pm$, that is only e_L can make a transition to a neutrino by emitting a W^- or absorbing a W^+ . The right-handed part e_R does not interact with $W^\pm$ at all. This is the parity violation of the weak interactions. Recall that

$$\bar{\nu}_L \gamma^\mu e_L = \bar{\nu} \gamma^\mu P_L e = \frac{\bar{\nu} \gamma^\mu e - \bar{\nu} \gamma^\mu \gamma_5 e}{2} \quad \left(P_L = \frac{1 - \gamma_5}{2} \right). \quad (8.31)$$

So the interaction is with a sum of vector (γ^μ) and axial vector ($\gamma^\mu\gamma_5$) currents. It is therefore called a *V-A charged current interaction*. The best known interaction of this kind is the neutron beta decay mentioned earlier. There are terms of exactly the same form for the quarks:

$$\mathcal{L}_{\text{quarks}}^{\text{charged}} = -\frac{g_2}{\sqrt{2}} \left(\bar{u}_L \gamma^\mu d_L W_\mu^+ + \bar{d}_L \gamma^\mu u_L W_\mu^- \right). \quad (8.32)$$

Hence a down quark d can turn into an up quark u by emitting a W^- that in turn decays into an electron and an anti-neutrino. This way a neutron (with quark content udd) can turn into a proton (with quark content uud).

Just as for the neutral current, the interactions with charged $W^\pm$ are of comparable strength as the electromagnetic interactions:

$$\left(\frac{g_2}{\sqrt{2}} \right)^2 = \frac{e^2}{2 \sin^2 \theta_W} \approx 2e^2 \quad \left(\sin^2 \theta_W \approx 0.23 \right). \quad (8.33)$$

In contrast, we observe that the weak interactions are *several magnitudes smaller* than the electromagnetic interactions. As for the Z-boson, the resolution is that the charged bosons $W^\pm$ must be massive. This reduces the strength of the transition rates (e.g. beta decay) due to the mass in the denominator of the propagator. The $W^\pm$ bosons were discovered experimentally in 1983 at CERN, confirming this prediction. Their mass is $m_W \approx 80.4 \text{ GeV}$, so they are a bit lighter than the Z-boson.

8.4 Quark Terms and Further Families

Quarks. As mentioned above, the U(1) and SU(2) interaction terms for quarks take exactly the same form as for leptons. Hence all the conclusions for leptons similarly hold for quarks. They couple to the same gauge bosons A_μ , Z_μ , and $W_\mu^\pm$ as the leptons. The coupling to the neutral Z_μ occurs with the same universal strength (electroweak charge) (8.26)

$$\frac{e}{\sin \theta_W \cos \theta_W} \left(T_{3f} - Q_f \sin^2 \theta_W \right) \quad (8.34)$$

for all (left-handed and right-handed) quarks.

In addition, quarks are the only particles that couple non-trivially to the SU(3) gauge bosons (gluons). That is, the $\lambda^a G_{a\mu}$ term in the covariant derivative D_μ only gives a non-zero contribution when acting on a quark state q . The matrices λ^a have size 3×3 , hence quark states carry a “color” index α that runs from 1 to 3. The SU(3) terms then read

$$\begin{aligned} \mathcal{L}_{\text{QCD}} &= \sum_{q=u,d} \bar{q}_\alpha i \gamma^\mu \left(-ig_3 \frac{\lambda_{\alpha\beta}^a}{2} G_{a\mu} \right) q_\beta \\ &= \frac{g_3}{2} \sum_{q=u,d} \bar{q}_\alpha \gamma^\mu \lambda_{\alpha\beta}^a G_{a\mu} q_\beta. \end{aligned} \quad (8.35)$$

For the electroweak theory, we wrote out the matrices $\tau^i W_i$ to recover the known electromagnetic interactions, and to identify the charged weak bosons $W^\pm$ and the neutral boson Z . Here, the G_a are eight gluons that are all electrically neutral. They interact with the quarks in a way that is somewhat similar to a photon interaction. The major difference is that, since the generators λ^a are not all diagonal, quarks can arbitrarily change their “color” by emitting or absorbing gluons. Since the color structure is hard to observe directly, because quarks and gluons are confined in hadrons, we do not study the structure of these terms as explicitly as the electroweak terms.

Fermion Wave Functions. One point that might be confusing is that the fermion states q , e , etc. have non-trivial structure in several spaces. For example, consider a quark state q . As is familiar from quantum mechanics, the total quark wave function is a product of factors

$$q = \eta_{\text{space}} \chi_{\text{spin}} \phi_{\text{U}(1)} \psi_{\text{SU}(2)} \xi_{\text{color}}. \quad (8.36)$$

Each factor is parametrized by some labels, indices, and/or coordinates, and describes the wave function in the respective space. When we act for example with an $\text{SU}(2)$ generator on q , then only the $\psi_{\text{SU}(2)}$ factor is affected, all other factors are left invariant. Similarly, when we perform a spacetime transformation (e.g. a rotation), then only η and χ are affected, and the wave functions ϕ , ψ , and ξ in the internal spaces are left invariant. Orthonormality of the wave function holds for each factor separately. Hence for example in the product $\bar{q}\tau^i W_i q$, all wave function factors except ψ give a trivial factor 1.

The Second and Third Families. So far, we have restricted ourselves to the first families of leptons (ν_e, e) and quarks (u, d) . As mentioned earlier, the physics exactly replicates for the other two families, so in all of the above, we can replace (ν_e, e) , (u, d) by (ν_μ, μ) , (c, s) or by (ν_τ, τ) , (t, b) to recover the interactions of the second and third families. The only differences between the three families are the different masses of the respective particles.

8.5 The Fermion Gauge Boson Lagrangian

We can now bring together all interaction terms of quarks and leptons with photons, electroweak bosons $W^\pm$ and Z , and gluons. All interactions of quarks and leptons arise from these terms.

$$\begin{aligned} \mathcal{L}_{\text{ferm}}^{\text{gauge}} = & e \sum_{f=\nu_e, e, u, d} Q_f (\bar{f} \gamma^\mu f) A_\mu \\ & + \frac{g_2}{\cos \theta_W} \sum_{f=\nu_e, e, u, d} \bar{f} \gamma^\mu (T_{3f} - Q_f \sin^2 \theta_W) f Z_\mu \\ & + \frac{g_2}{\sqrt{2}} \left[(\bar{\nu}_{eL} \gamma^\mu e_L + \bar{u}_L \gamma^\mu d_L) W_\mu^+ + (\bar{e}_L \gamma^\mu \nu_{eL} + \bar{d}_L \gamma^\mu u_L) W_\mu^- \right] \\ & + \frac{g_3}{2} \sum_{q=u, d} \bar{q}_\alpha \gamma^\mu \lambda_{\alpha\beta}^a q_\beta G_{a\mu}. \end{aligned} \quad (8.37)$$

Here, the neutral current Z_μ term expands to

$$\begin{aligned} \bar{f} \gamma^\mu (T_{3f} - Q_f \sin^2 \theta_W) f Z_\mu = \\ \left[\bar{f}_L \gamma^\mu (T_{3f} - Q_f \sin^2 \theta_W) f_L + \bar{f}_R \gamma^\mu (-Q_f \sin^2 \theta_W) f_R \right] Z_\mu. \end{aligned} \quad (8.38)$$

The respective interaction terms for the second and third families are obtained by replacing $(\nu_e, e, u, d) \rightarrow (\nu_\mu, \mu, c, s)$ or (ν_τ, τ, t, b) . What remains is to specify the charges of the

various fermions:

Particle	T_3	Y	$Q = T_3 + \frac{Y}{2}$	SU(3)
ν_e, ν_μ, ν_τ	$+\frac{1}{2}$	-1	0	0
e_L, μ_L, τ_L	$-\frac{1}{2}$		-1	0
e_R, μ_R, τ_R	0	-2	-1	0
u_L, c_L, t_L	$+\frac{1}{2}$	$\frac{1}{3}$	$+\frac{2}{3}$	$\square$
d_L, s_L, b_L	$-\frac{1}{2}$	$\frac{1}{3}$	$-\frac{1}{3}$	$\square$
u_R, c_R, t_R	0	$+\frac{4}{3}$	$+\frac{2}{3}$	$\square$
d_R, s_R, b_R	0	$-\frac{2}{3}$	$-\frac{1}{3}$	$\square$

(8.39)

Here, the box $\square$ means that the particle transforms in the fundamental representation of SU(3), while a zero means that it is a singlet of the respective group. The various coupling constants have approximate values

$$\begin{aligned}
g_1 &= \frac{e}{\cos \theta_W}, & g_2 &= \frac{e}{\sin \theta_W}, & \sin^2 \theta_W &\approx 0.23, \\
\alpha &:= \frac{e^2}{4\pi} \approx \frac{1}{137}, & \alpha_1 &:= \frac{g_1^2}{4\pi} \approx \frac{1}{100}, \\
\alpha_2 &:= \frac{g_2^2}{4\pi} \approx \frac{1}{30}, & \alpha_3(m_Z) &:= \frac{g_3(m_Z)^2}{4\pi} \approx 0.12.
\end{aligned}
\tag{8.40}$$

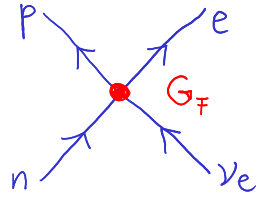
The values of the couplings α , α_1 , α_2 , and α_3 depend on the energy scale of the interaction. This phenomenon is due to the influence of spontaneously created particle-antiparticle pairs on the respective interaction, which is an effect that depends on the energy scale. This is called *running of the coupling* and is described by the *renormalization group* in quantum field theory. The values of α , α_1 , and α_2 given above are for interactions in the range of a few GeV or below, and they vary very slowly. For the strong coupling α_3 , the value is given at the Z-boson mass $m_Z \approx 91.2 \text{ GeV}$. It varies slowly above this scale, but it increases at lower energies, reaching ≈ 0.3 at about 1 GeV, and grows quickly to $\alpha_3 > 1$ below 1 GeV, causing quarks to bind into hadrons.

8.6 Historical Note

Now that we have learned how the electroweak theory unifies electromagnetic and weak interactions, one may ask: How did this theory come to be? *Beta decay* was discovered around 1900, and was the first hint for the weak interactions. At first, beta decay was observed as the mutation of an atomic nucleus due to emission of an electron. In the next ~ 30 years, it became clear that the energy spectrum of the emitted electron was not consistent with energy conservation. In 1933, in a famous letter to the ETH Zürich, Wolfgang Pauli proposed the existence of an extremely light neutral particle (now called neutrino) that is emitted along with the electron but could not be observed at the time. The neutrino was indeed detected experimentally in 1956 (Nobel Prize Frederick Reines 1995).

The first theoretical description of beta decay was by Enrico Fermi through *Fermi's interaction*, also in 1933. He proposed a four-fermion contact interaction between the

neutron (which was only discovered one year earlier), the proton, the electron, and the neutrino:


(8.41)

Fermi's interaction already put the neutron and proton into a doublet, and the interaction Hamiltonian was given in terms of raising and lowering operators acting on this doublet. The strength of the interaction is parametrized by the *Fermi coupling constant* G_F . In modern terms, it is given by

$$\frac{G_F}{\sqrt{2}} = \frac{g_2^2}{8m_W^2} \approx 1.166 \cdot 10^{-5} \text{ GeV}^{-2}, \quad (8.42)$$

where m_W is the mass of the $W^\pm$ boson. Fermi's theory described the weak interaction remarkably well. But the interaction probability grows as the square of the energy, $\sigma \sim G_F^2 E^2$. Since it grows without bound, the theory is invalid at energies $\gtrsim 100 \text{ GeV}$. The four-fermion interaction therefore had to be replaced by a more complete theory. The runaway of the interaction probability is avoided if the interaction is due to the exchange of a particle of mass $\sim 100 \text{ GeV}$ – the $W^\pm$ boson.

In 1956, it was discovered experimentally (by Chien-Shiung Wu) that the weak interaction *violates parity* symmetry (Nobel Prize 1957 to Tsung-Dao Lee and Chen-Ning Yang for the theory of parity violation and proposal of the experiment). This result prompted a search to relate the weak and electromagnetic interactions. The electroweak theory was proposed by Sheldon Glashow, Abdus Salam, and Steven Weinberg in 1968 (Nobel Prize 1979). In their theory, the photon and the $W^\pm$ as well as the new Z boson, including their masses, arise from the *spontaneous symmetry breaking* of $U(1) \times SU(2)$, which will be discussed in the next section. The $W^\pm$ and Z bosons were detected experimentally in 1983 at the CERN Super Proton Synchrotron (Nobel Prize Carlo Rubbia and Simon van der Meer 1984).

8.7 Masses?

So far, we have treated all fermions and gauge bosons as massless. Although everything we said about the interactions of quarks and leptons with electroweak gauge bosons and gluons is correct, *all fermions are in fact massive*, and also the $W^\pm$ and Z bosons are massive.

Could we just add mass terms to the Lagrangian “by hand”? That will not work! Any mass term would break gauge invariance. For fermions, mass terms would be of the form $m\bar{\psi}\psi$, but we know that

$$\begin{aligned} m\bar{\psi}\psi &= m\bar{\psi}(P_L + P_R)\psi \\ &= m\bar{\psi}P_L P_L \psi + m\bar{\psi}P_R P_R \psi \\ &= m(\bar{\psi}_R \psi_L + \bar{\psi}_L \psi_R). \end{aligned} \quad (8.43)$$

However, the left-handed fermions transform in $SU(2)$ doublets, while the right-handed fermions are $SU(2)$ singlets. Hence $\bar{\psi}_R \psi_L$ and $\bar{\psi}_L \psi_R$ are not $SU(2)$ singlets and would not

give an SU(2)-invariant Lagrangian. Similarly, mass terms for the gauge bosons, e. g.

$$\frac{1}{2} m_B^2 B_\mu B^\mu, \quad (8.44)$$

are clearly not gauge invariant. The only way to preserve gauge invariance of the Lagrangian is to set the masses of all fermions and gauge bosons to zero. If mass terms are put in by hand, gauge invariance is lost, and the theory produces unphysical infinities.

There is a way to solve this problem, which is the *Higgs mechanism*. In the resulting Lagrangian, the U(1) and SU(2) gauge invariance is broken, but in a subtle way that preserves the good effects of the gauge symmetry.

9 Masses and the Higgs Mechanism

In the following, we will develop the mechanism that consistently gives masses to the weak gauge bosons as well as all fermions. The fundamental assumption is that the universe is filled with a spin-zero field called the *Higgs field*. Its essential properties are that it is a doublet in the weak SU(2) space, and also carries a non-zero U(1) hypercharge, but is a singlet under the color SU(3). The gauge bosons and fermions interact with the Higgs field, and in its presence, they acquire a mass. Because the Higgs field carries non-trivial U(1) and SU(2) quantum numbers, a universe-filling “background” Higgs field effectively breaks these gauge symmetries. The symmetry is present in the Lagrangian, but is broken by the vacuum state. This situation is called *spontaneous symmetry breaking*. Several examples of spontaneous symmetry breaking were the subject of the exercise problems. Here, we briefly recapitulate these examples, which will prepare us to understand the Higgs mechanism in the Standard Model at the end of this section.

9.1 Spontaneous Symmetry Breaking

We start with the simplest example. Consider a Lagrangian for a scalar field ϕ ,

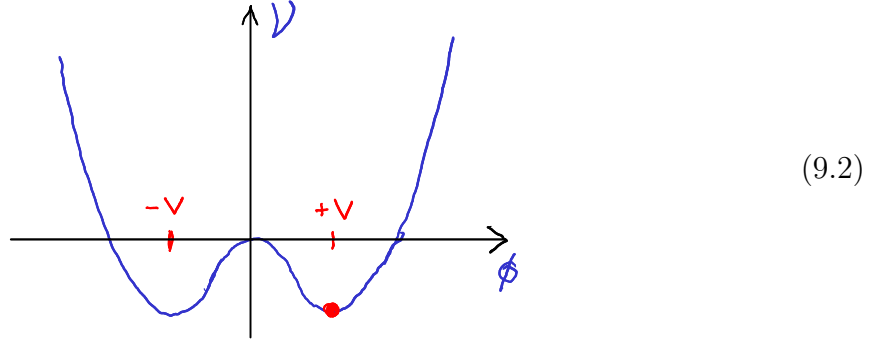
$$\mathcal{L} = \mathcal{T} - \mathcal{V}, \quad \mathcal{T} = \frac{1}{2} \partial_\mu \phi \partial^\mu \phi, \quad \mathcal{V} = \frac{1}{2} \mu^2 \phi^2 + \frac{1}{4} \lambda \phi^4. \quad (9.1)$$

Here, μ and λ are external parameters of the potential $\mathcal{V}$. We require $\lambda > 0$, since otherwise the potential would not be bounded from below. The theory has a *symmetry*: It is invariant under $\phi \mapsto -\phi$.

To find the spectrum of the theory, one starts by finding the minimum of the potential. The field configuration that minimizes $\mathcal{V}$ is the classical ground state of the system. Then one expands the fields around this ground state to determine the possible excitations. This is familiar from quantum mechanical perturbation theory. In field theory, one calls the ground state the *vacuum*, and the excitations are the *particles*. Their masses are determined by the form of the potential near the classical minimum, by comparison with the Lagrangian of a free massive scalar field.

For $\mu^2 > 0$, the vacuum is $\phi = 0$. In that case, we identify $-\mu^2 \phi^2/2$ as a mass term, hence ϕ is a massive field with mass $m_\phi = \mu$. The $\lambda \phi^4$ terms represents an interaction of

the field ϕ with itself. If on the other hand $\mu^2 < 0$, the potential takes the form



Its minima are at

$$\phi = \pm v, \quad v := \sqrt{\frac{-\mu^2}{\lambda}}. \quad (9.3)$$

Since the configuration $\phi = \pm v$ is constant, it also minimizes the kinetic energy $\mathcal{T}$ of the system, and is therefore a valid ground state, or vacuum. The value v is also called the *vacuum expectation value* of the system. A variation of this is also what happens to the Higgs field, as we will see.

To find the particle spectrum, we study the theory near the vacuum, that is we set

$$\phi(x) = v + \eta(x), \quad (9.4)$$

with $\eta(x)$ small. We could have equally set $\phi = -v + \eta$, but the physics would be the same, since the theory is invariant under $\phi \mapsto -\phi$. Writing $\mathcal{L}$ in terms of $\eta(x)$ gives

$$\mathcal{L} = \frac{1}{2} \partial_\mu \eta \partial^\mu \eta - \left(\lambda v^2 \eta^2 + \lambda v \eta^3 + \frac{1}{4} \lambda \eta^4 \right) + \text{constant}. \quad (9.5)$$

The η^2 term has the correct sign, so it can be interpreted as a mass term. The Lagrangian describes a real scalar field $\eta(x)$ with mass

$$m_\eta^2 = 2\lambda v^2 = -2\mu^2. \quad (9.6)$$

We identify the mass as the curvature of the potential at its minimum. There are two interaction terms: A cubic one of strength λv , and a quartic one of strength $\lambda/4$.

If the theory is solved exactly, the two descriptions (one in terms of ϕ , the other in terms of η) must be equivalent. But in a perturbative description, one has to perturb around the ground state (vacuum), otherwise the perturbative expansion does not converge. And choosing one of the two possible vacua ($\phi = +v$ here) breaks the symmetry of the theory, in the sense that for small excitations $\eta(x)$ around $\phi = v$, the original symmetry is no longer present. A memory of the original symmetry is preserved in the η^3 term, but not in an obvious way. This situation is called *spontaneous symmetry breaking*, and it frequently occurs in physical systems.

9.2 Breaking of a Continuous Symmetry

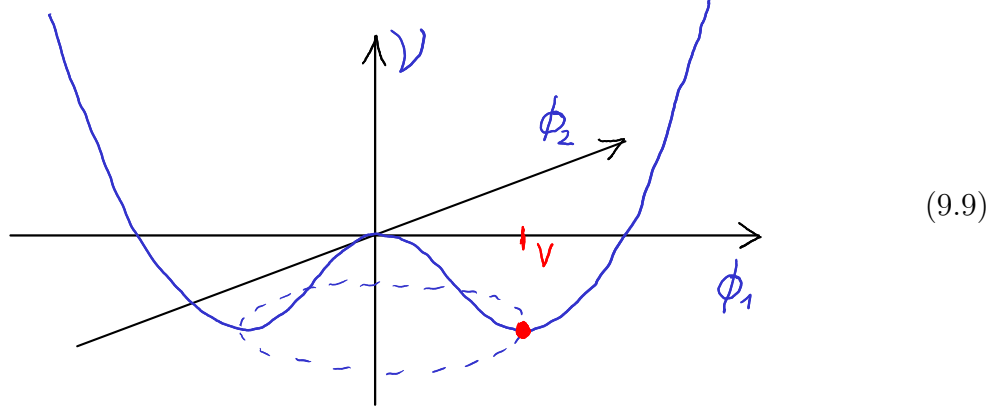
Now consider a complex scalar field $\phi = (\phi_1 + i\phi_2)/\sqrt{2}$, with Lagrangian

$$\mathcal{L} = (\partial_\mu \phi)^* (\partial^\mu \phi) - (\mu^2 \phi^* \phi + \lambda (\phi^* \phi)^2). \quad (9.7)$$

This theory is invariant under a continuous family of global U(1) transformations

$$\phi \mapsto \phi' = e^{ix}\phi. \quad (9.8)$$

For $\mu^2 > 0$, the potential (and total energy) are minimized by $\phi \equiv 0$. In this case, the theory describes a massive complex field with mass $m_\phi^2 = \mu^2$, and with a quartic interaction of strength λ . For $\mu^2 < 0$, the potential has the form



This potential is often called the “*Mexican hat*” potential. The minimum of the potential (and total energy) is along a circle of radius

$$\phi_1^2 + \phi_2^2 = v^2, \quad v := \sqrt{\frac{-\mu^2}{\lambda}}. \quad (9.10)$$

Because of the U(1) symmetry, all points on the minimizing circle are equivalent. To do perturbation theory, we arbitrarily, but without loss of generality, choose $\phi_1 = v$, $\phi_2 = 0$, and write ϕ as

$$\phi = \frac{v + \eta(x) + i\rho(x)}{\sqrt{2}}, \quad (9.11)$$

with η and ρ real. With this substitution, we find again that the Lagrangian can be interpreted in terms of particles and their interactions:

$$\begin{aligned} \mathcal{L} = & \frac{1}{2} (\partial_\mu \eta)^2 + \frac{1}{2} (\partial_\mu \rho)^2 + \mu^2 \eta^2 - \lambda v (\eta \rho^2 + \eta^3) \\ & - \frac{\lambda}{2} \eta^2 \rho^2 - \frac{\lambda}{4} \eta^4 - \frac{\lambda}{4} \rho^4 + \text{constant}. \end{aligned} \quad (9.12)$$

There are two real scalar fields, η and ρ . The η^2 term is a mass term for η . We can read off the mass as

$$m_\eta^2 = -2\mu^2. \quad (9.13)$$

There is no ρ^2 term, so we must conclude that ρ is massless! This massless field is called a *Goldstone boson*. There is a general theorem that says whenever a continuous global symmetry (as the U(1) in this case) is spontaneously broken (that is the Lagrangian is still invariant, but the choice of vacuum breaks the symmetry), then the spectrum around the vacuum will contain a massless field, called the Goldstone boson.

It is intuitively clear why the massless boson arises: The minimum of the energy lies along a circle in field space (or along some other curve for a more generic symmetry). Perturbations around the vacuum in the radial direction (or more generally in a perpendicular direction to the minimum) must be pushed up the potential, and the curvature

of the potential at the minimum is identified with the mass of the perturbation. But when we perturb along the direction of the minimum, the potential is flat, so its curvature vanishes and the associated perturbation is massless: It only costs kinetic energy to push the system along the minimum, but no potential energy.

9.3 Abelian Higgs Mechanism

Now we get closer to the kind of spontaneous symmetry breaking that occurs in the Standard Model. Above, we considered a global continuous symmetry. Now we promote this to a local gauge symmetry. The spontaneous breaking of this local gauge symmetry will lead to a mass for the (previously massless) gauge boson. As is by now familiar, invariance under a field transformation

$$\phi(x) \mapsto \phi'(x) = e^{i\chi(x)}\phi(x), \quad (9.14)$$

with a local parameter $\chi(x)$ requires the introduction of a covariant derivative with a vector (spin one) field A_μ ,

$$\partial_\mu \rightarrow D_\mu = \partial_\mu - igA_\mu. \quad (9.15)$$

The gauge field A_μ transforms as

$$A_\mu \mapsto A'_\mu = A_\mu - \frac{1}{g} \partial_\mu \chi(x). \quad (9.16)$$

The gauge-invariant Lagrangian is

$$\mathcal{L} = (D_\mu \phi)^*(D^\mu \phi) - \frac{1}{4} F_{\mu\nu} F^{\mu\nu} - \mathcal{V}, \quad \mathcal{V} = \mu^2 \phi^* \phi + \lambda (\phi^* \phi)^2. \quad (9.17)$$

Here, $\frac{1}{4} F_{\mu\nu} F^{\mu\nu}$ is the kinetic term for the gauge field. It will not play a role in the analysis. For $\mu^2 > 0$, this Lagrangian describes a scalar particle with mass μ that interacts with the electromagnetic field A_μ with charge g . As always in standard gauge theory, A_μ is massless.

The potential $\mathcal{V}$ is the same as in the previous case. Hence when $\mu^2 < 0$, the total energy is again minimized for non-zero constant ϕ :

$$\phi = \frac{v}{\sqrt{2}}, \quad v = \sqrt{\frac{-\mu^2}{\lambda}}. \quad (9.18)$$

The field ϕ is complex, but we can use the gauge transformations $\phi(x) \mapsto e^{i\chi(x)}\phi(x)$ to rotate ϕ to real values everywhere. Hence we can expand ϕ around the vacuum as

$$\phi(x) = \frac{v + h(x)}{\sqrt{2}}, \quad (9.19)$$

with a real field $h(x)$. We could not make ϕ real everywhere in the previous case, because there the theory was only invariant under global phase rotations, not local ones. Writing the Lagrangian in terms of $h(x)$ instead of $\phi(x)$, one finds

$$\begin{aligned} \mathcal{L} = & \frac{1}{2} (\partial_\mu h)(\partial^\mu h) - \lambda v^2 h^2 - \frac{1}{4} F_{\mu\nu} F^{\mu\nu} + \frac{1}{2} g^2 v^2 A_\mu A^\mu \\ & + g^2 v h A_\mu A^\mu - \lambda v h^3 - \frac{\lambda}{4} h^4 + \frac{1}{2} g^2 h^2 A_\mu A^\mu \end{aligned} \quad (9.20)$$

As before, the real field $h(x)$ (it was called η in the previous example) is massive, with a mass

$$m_h^2 = 2\lambda v^2 = -2\mu^2 \quad \Rightarrow \quad m_h = \sqrt{2\lambda}v = \sqrt{-2\mu^2}. \quad (9.21)$$

Surprisingly, also the gauge field acquires a mass! Its squared mass is the coefficient of $A_\mu A^\mu/2$, hence

$$m_A = gv. \quad (9.22)$$

The mass term arises from the term in $(D_\mu\phi)^*(D^\mu\phi)$ that is quadratic in A_μ . The terms in the second line of the Lagrangian are various interaction terms between h and A_μ .

In the unbroken phase of the theory ($\mu^2 > 0$), the gauge boson is massless, and therefore has two spin degrees of freedom. In the broken phase ($\mu^2 < 0$), it becomes massive, and therefore has three spin degrees of freedom. What happened here is that the massless Goldstone boson ρ of the previous case has become the third degree of freedom of the gauge boson A_μ . This phenomenon is also referred to as the Goldstone boson being “eaten” by the gauge boson. The effect can be seen more explicitly by doing the computation without the simplifying gauge transformation step, i. e. by setting

$$\phi = \frac{v + h(x) + i\rho(x)}{\sqrt{2}} \quad (9.23)$$

with $h(x)$ and $\rho(x)$ real. Then one would find a term $\sim A_\mu \partial^\mu \rho$ in $\mathcal{L}$, which would mean that the gauge field can turn into a ρ field as it propagates. This means that the fields A_μ and ρ are not diagonalized. Properly diagonalizing them amounts to a gauge transformation that then eliminates ρ altogether. Note that before and after the symmetry breaking, the total number of field components is 4: Before the breaking, the complex scalar field and the gauge field each have 2 components. After the breaking, the real scalar h has one component, and the massive gauge field has 3 components.

The phenomenon we have just studied is called the *Higgs mechanism*: Spontaneously breaking a gauge symmetry leads to a non-zero vacuum value for ϕ that in turn makes the gauge boson massive. The extra degree of freedom of a massive vector boson is what used to be the Goldstone boson. What is left of the scalar field $\phi(x)$ are real scalar excitations $h(x)$ around the vacuum value v . The particle h is called the *Higgs boson*. Note that the mass of the gauge boson is fixed if g and v are known, but the mass of the Higgs boson in addition depends on the possibly unknown parameter λ .

9.4 The Higgs Mechanism in the Standard Model

To get to the Higgs mechanism in the Standard Model, we only need one more bit of complexity. In the Standard Model, the Higgs field forms a doublet of the weak isospin $SU(2)_W$:

$$\phi = \begin{pmatrix} \phi^+ \\ \phi^0 \end{pmatrix}, \quad (9.24)$$

where ϕ^+ and ϕ^0 are each complex fields,

$$\phi^+ = \frac{\phi_1 + i\phi_2}{\sqrt{2}}, \quad \phi^0 = \frac{\phi_3 + i\phi_4}{\sqrt{2}}, \quad (9.25)$$

whose complex phases are rotated by the $U(1)$ gauge symmetry of the Standard Model. The Lagrangian has to be a $U(1) \times SU(2)_W$ singlet. Such a singlet is formed by

$$\phi^\dagger \phi = \left((\phi^+)^* \quad (\phi^0)^* \right) \begin{pmatrix} \phi^+ \\ \phi^0 \end{pmatrix} = (\phi^+)^* \phi^+ + (\phi^0)^* \phi^0. \quad (9.26)$$

In terms of the real components,

$$\phi^\dagger \phi = \frac{\phi_1^2 + \phi_2^2 + \phi_3^2 + \phi_4^2}{2} . \quad (9.27)$$

The form of the Lagrangian for ϕ is analogous to the previous cases,

$$\mathcal{L}_\phi = (\partial_\mu \phi)^\dagger (\partial^\mu \phi) - \mathcal{V} , \quad (9.28)$$

with the same potential

$$\mathcal{V} = \mu^2 \phi^\dagger \phi + \lambda (\phi^\dagger \phi)^2 . \quad (9.29)$$

Again, for $\mu^2 < 0$ the potential is minimized when

$$\phi^\dagger \phi = \frac{v^2}{2} , \quad v := \sqrt{\frac{-\mu^2}{\lambda}} . \quad (9.30)$$

There are many ways to satisfy this condition: The field ϕ has four real components, and this is only one (real) condition, hence it is satisfied in a three-dimensional subspace. But since the theory is invariant under arbitrary $\text{SU}(2)_W$ transformations

$$\phi \mapsto \phi' = U \phi , \quad U = e^{i\alpha_i \tau^i / 2} \in \text{SU}(2) , \quad (9.31)$$

all points in this space are equivalent. We pick

$$\phi_0 := \frac{1}{\sqrt{2}} \begin{pmatrix} 0 \\ v \end{pmatrix} , \quad (9.32)$$

that is $\phi_1 = \phi_2 = \phi_4 = 0$, $\phi_3 = v$, and expand $\phi(x)$ around this vacuum as

$$\phi(x) = \frac{1}{\sqrt{2}} \begin{pmatrix} 0 \\ v + h(x) \end{pmatrix} . \quad (9.33)$$

As in the previous example of the Abelian Higgs mechanism, the $\mu^2 \phi^\dagger \phi$ term in the potential leads to a mass term for the *Higgs boson* h , with mass

$$m_h^2 = 2\lambda v^2 = -2\mu^2 \quad (9.34)$$

Similarly as in the Abelian Higgs mechanism, we can always bring ϕ to this form by applying an $\text{SU}(2)_W$ transformation (to set the first component $\phi^+ = 0$), followed by an $\text{U}(1)$ phase rotation to make the second component ϕ^0 real. In other words, we “gauged away” three field components. And for doing so, we have to use three gauge degrees of freedom, so we have *broken three gauge symmetries*. By the Goldstone theorem, there should be three massless Goldstone bosons. But because the symmetries are *gauge* symmetries, instead *three gauge bosons should become massive*, and thereby absorb the three Goldstone degrees of freedom. This is exactly what we need to give masses to the electroweak gauge bosons Z_μ and $W_\mu^\pm$.

Before we study the properties of the Lagrangian with the vacuum ϕ_0 , we note that the lower component ϕ^0 of the Higgs field has to be electrically neutral. Otherwise, the vacuum would be charged, which is not what we observe. For example, a charged vacuum could absorb and emit photons. Because of the relation

$$Q = T_3 + \frac{Y}{2} \quad (9.35)$$

between electric charge, weak isospin, and hypercharge, and because ϕ^0 , as the lower component of a weak isospin $SU(2)_W$ doublet, has $T_3 = -1/2$, we must assign

$$Y_H = 1 \quad (9.36)$$

for the Higgs doublet ϕ .

We said that the choice of vacuum breaks three gauge symmetries. The original gauge group $U(1) \times SU(2)$ has four dimensions, so the vacuum must preserve one symmetry. Which one is it? Clearly, the vacuum ϕ_0 does not preserve any $SU(2)_W$ symmetry. Also, because $Y_H \neq 0$, it also does not preserve the $U(1)$ symmetry. However, if we act with the electric charge operator,

$$Q\phi_0 = (T_3 + Y/2)\phi_0 = (-1/2 + 1/2)\phi_0 = 0, \quad (9.37)$$

we find that it annihilates ϕ_0 and therefore the vacuum is invariant under a transformation

$$\phi_0 \mapsto e^{i\alpha(x)Q}\phi_0. \quad (9.38)$$

The generator Q of this particular $U(1)'$ transformation is a particular linear combination of the original $U(1)$ and $SU(2)_W$ gauge transformations. Of course, this $U(1)'$ has to be the $U(1)$ of electromagnetism, as we identified its generator as the electric charge. Hence we found that the vacuum is invariant under electromagnetic $U(1)$ transformations, which means that of the original four gauge bosons, the one combination that remains massless must be the photon field A_μ .

Let us see what the choice of vacuum $\phi = \phi_0$ implies for the Lagrangian. We know that we must replace the ordinary derivative ∂_μ by the covariant derivative

$$D_\mu = \partial_\mu - ig_1 \frac{Y}{2} B_\mu - ig_2 \frac{\boldsymbol{\tau}}{2} \mathbf{W}_\mu. \quad (9.39)$$

Then the Lagrangian acquires extra terms

$$\phi^\dagger \left(-ig_1 \frac{Y}{2} B_\mu - ig_2 \frac{\tau^i}{2} W_{i\mu} \right)^\dagger \left(-ig_1 \frac{Y}{2} B_\mu - ig_2 \frac{\tau^i}{2} W_{i\mu} \right) \phi. \quad (9.40)$$

Setting $Y = Y_H = 1$, inserting the explicit 2×2 Pauli matrices for τ^i , and setting $\phi = \phi_0$, this becomes

$$\begin{aligned} & \frac{1}{8} \left| \begin{pmatrix} g_1 B_\mu + g_2 W_{3\mu} & g_2(W_{1\mu} - iW_{2\mu}) \\ g_2(W_{1\mu} + iW_{2\mu}) & g_1 B_\mu - g_2 W_{3\mu} \end{pmatrix} \begin{pmatrix} 0 \\ v \end{pmatrix} \right|^2 \\ &= \frac{1}{8} v^2 g_2^2 ((W_{1\mu})^2 + (W_{2\mu})^2) + \frac{1}{8} v^2 (g_1 B_\mu - g_2 W_{3\mu})^2. \end{aligned} \quad (9.41)$$

We can immediately recognize the first term as mass terms for the gauge bosons W_1 and W_2 , with a common mass

$$m_W = \frac{v}{2} g_2. \quad (9.42)$$

Passing to the charge eigenstates $W^\pm$, the first term reads $m_W^2 W_\mu^+ W^{-\mu}$, which is the common way to write a mass term for a charged field. So we conclude that the charged bosons $W^\pm$ indeed acquire a mass m_W through the Higgs mechanism!

The second term is not diagonal in the fields, so we have to perform a rotation in the space of $(B_\mu, W_{3\mu})$ fields to find the fields with definite masses. In fact, we already

know what the rotation must be, since what appears in the second term is exactly the combination that we called Z_μ ,

$$Z_\mu = \frac{g_1 B_\mu - g_2 W_{3\mu}}{\sqrt{g_1^2 + g_2^2}} . \quad (9.43)$$

Hence the second term reads

$$\frac{1}{2} m_Z^2 Z_\mu Z^\mu , \quad m_Z = \frac{v}{2} \sqrt{g_1^2 + g_2^2} , \quad (9.44)$$

where we have recognized the mass of the Z_μ boson. The electromagnetic gauge boson A_μ remains massless, as expected. Using the relation

$$\cos \theta_W = \frac{g_2}{\sqrt{g_1^2 + g_2^2}} , \quad (9.45)$$

we find the identity

$$\frac{m_W}{m_Z} = \cos \theta_W . \quad (9.46)$$

This is an important identity: Once the electroweak mixing angle θ_W is measured in some way, the ratio of the Z_μ and $W_\mu^\pm$ masses is a clear prediction of the Standard Model. And it perfectly matches with experimental data: The experimental ratio $\cos \theta_W m_Z / m_W$ equals unity with an uncertainty of $\approx 0.1\%$.

9.5 Fermion Masses

Now that we have introduced the Higgs field ϕ , it is possible to write gauge invariant interaction terms between the Higgs doublet and the fermions.

Leptons. For example, for the leptons one can introduce a term

$$\mathcal{L}_{\text{int}}^{\phi e} = g_e (\bar{L} \phi e_R + \bar{e}_R \phi^\dagger L) . \quad (9.47)$$

Here, the $\text{SU}(2)_W$ indices of $\bar{L}$ and ϕ are contracted: $\bar{L} = (\nu_{eL} \ e_L)$ is a row vector, and ϕ is a column vector – their product is an $\text{SU}(2)_W$ singlet. The factor e_R is an $\text{SU}(2)_W$ singlet by itself, so the whole first term is an $\text{SU}(2)_W$ singlet. The factor e_R has to be present because of Lorentz invariance: The spin structures of $\bar{L}$ and e_R combine to form a Lorentz scalar. The second term in the Lagrangian is the hermitian conjugate of the first, so the same symmetry considerations apply. The overall coefficient g_e is the coupling constant of the interaction, it is arbitrary at this point. To find the consequences of this term, we replace ϕ by its vacuum value and the Higgs particle h :

$$\phi \rightarrow \frac{1}{\sqrt{2}} \begin{pmatrix} 0 \\ v + h \end{pmatrix} . \quad (9.48)$$

This gives

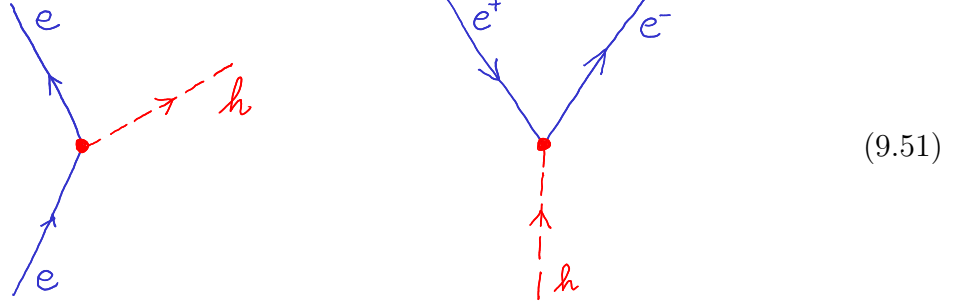
$$\mathcal{L}_{\text{int}}^{\phi e} = \frac{g_e v}{\sqrt{2}} (\bar{e}_L e_R + \bar{e}_R e_L) + \frac{g_e}{\sqrt{2}} (\bar{e}_L e_R + \bar{e}_R e_L) h . \quad (9.49)$$

Recalling that $(\bar{e}_L e_R + \bar{e}_R e_L) = \bar{e} e$, the first term has exactly the form of a mass term $m_e \bar{e} e$ for the electron, and we can identify the electron mass as

$$m_e = \frac{g_e v}{\sqrt{2}} . \quad (9.50)$$

Hence the coupling of the electron to the Higgs doublet ϕ introduces a mass term for the electron, without breaking the gauge invariance of the theory!

The second term in $\mathcal{L}_{\text{int}}^{\phi e}$ says that there is an electron-Higgs interaction vertex in the theory, of coupling strength $g_e/\sqrt{2} = m_e/v$:



Hence we find that an electron can emit or absorb a Higgs particle h . Similarly, a Higgs particle can decay into an electron-anti-electron pair (e^-, e^+). This interaction vertex is used to compute the probability of producing a Higgs particle in particle collisions. Writing g_e in terms of m_e , the interaction Lagrangian becomes

$$\mathcal{L}_{\text{int}}^{\phi e} = m_e \bar{e}e + \frac{m_e}{v} \bar{e}e h. \quad (9.52)$$

Note that the coupling strength between the fermion (electron) and the Higgs particle is proportional to the mass of the fermion.

Importantly, the interaction Lagrangian $\mathcal{L}_{\text{int}}^{\phi e}$ produces no mass term for the neutrino ν_{eL} , and moreover it is not possible to write an interaction term that does produce such a mass term. The reason is that, by assumption, the theory does not contain a right-handed neutrino ν_{eR} , so one cannot produce a mass term $\bar{\nu}_{eR}\nu_{eL}$. This in particular implies that the neutrino does not interact with the Higgs particle h . If there was a right-handed neutrino ν_{eR} , it would be hard to detect, since it would have $T_3 = 0$ and $Q = 0$, so it would couple to none of the electroweak gauge bosons $W_\mu^\pm$, Z_μ , or A_μ . We will come back to the neutrino mass question and the possible existence of ν_{eR} later on.

Quarks. The same mechanism that works for the leptons equally works for the quarks. In the case of quarks, there *is* a right-handed partner to both components of the left-handed doublet Q_L , so, unlike in the electron/neutrino case, we *can* produce mass terms for both the up and the down quark. The mass of the down quark (and the corresponding interaction with the Higgs particle h) is produced in exactly the same way as for the electron. For the up quark, we use the fact that from any SU(2) doublet $(a \ b)^T$, one can construct another “conjugate” SU(2) doublet $(-b^* \ a^*)^T$ (this other doublet is equivalent to the anti-fundamental representation). We can therefore use the conjugate Higgs doublet

$$\phi_c := -i\tau_2 \phi = \begin{pmatrix} -\phi^{0*} \\ \phi^- \end{pmatrix}, \quad \phi^- := \phi^{+*}. \quad (9.53)$$

In terms of the vacuum v and the Higgs particle h , this becomes

$$\phi_c \rightarrow \frac{1}{\sqrt{2}} \begin{pmatrix} -(v+h) \\ 0 \end{pmatrix}. \quad (9.54)$$

Combining this with the quark doublet Q_L gives the SU(2)_W singlet

$$\bar{Q}_L \phi_c \sim \bar{u}_L (v+h). \quad (9.55)$$

Therefore, to give masses to both the up and the down quark, we add interaction terms

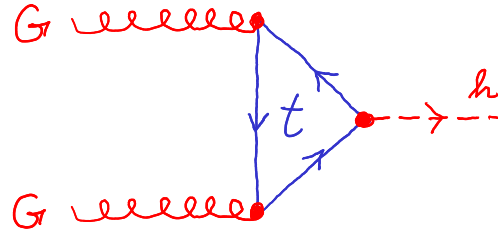
$$\mathcal{L}_{\text{int}}^{q\phi} = g_d \bar{Q}_L \phi d_R + g_u \bar{Q}_L \phi_c u_R + \text{h.c.}, \quad (9.56)$$

where “h.c.” stands for the hermitian conjugate of the first two terms. Writing ϕ and ϕ_c in terms of v and h , one finds, exactly as for the leptons,

$$\mathcal{L}_{\text{int}}^{q\phi} = m_d \bar{d}d + \frac{m_d}{v} \bar{d}dh + m_u \bar{u}u + \frac{m_u}{v} \bar{u}uh, \quad (9.57)$$

where again the couplings g_d and g_u have been eliminated in favor of the quark masses m_d and m_u , following the same steps as for the electron. We conclude that the Higgs mechanism also succeeds in giving both the up and down quarks masses. As for the electron, the values of the masses are not predicted by the theory, but are free parameters (in the form of the couplings g_d and g_u). They have to be determined by measurements.

Further Families. All of the above analysis replicates for the other two families of quarks and leptons, just as it did for the interactions between fermions and gauge bosons. The masses of all fermions are free parameters of the Standard Model. They remain to be explained by a more fundamental theory. Since the coupling of each fermion to the Higgs particle h is proportional to its mass, the Higgs particle interacts most strongly with the heaviest particles. By far the heaviest fermion is the top quark (173 GeV). The second-heaviest fermion, the bottom quark at ~ 4 GeV, is already much lighter. Therefore, the physics of the Higgs particle is dominated by its interaction with the top quark. For example, the dominant process for Higgs production at the LHC is gluon fusion via a “top triangle”:



The diagram shows two incoming gluons (G) represented by red wavy lines on the left. They meet at two vertices (red dots) which are connected by a top quark (t) loop, shown as a blue triangle with arrows indicating the flow of the quark. The top quark loop connects to a single vertex (red dot) from which a Higgs particle (h) is emitted, represented by a red dashed line. The entire diagram is labeled (9.58) on the right.

9.6 Summary of the Standard Model

Our description of the basic structure of the Standard Model is now complete. The basic ingredients are

- Gauge invariance under the $U(1) \times SU(2) \times SU(3)$ gauge group.
- Three families of quarks and leptons that interact with the three types of gauge bosons (electroweak gauge bosons and gluons).
- The scalar Higgs field with a “Mexican hat” potential that lets the Higgs field acquire a non-zero vacuum expectation value v . The Higgs vacuum v breaks the electroweak $U(1) \times SU(2)$ gauge symmetry to a residual $U(1)$, which is the electromagnetic gauge symmetry. The three gauge bosons $W^\pm$ and Z acquire masses due to the three broken gauge symmetries.
- Couplings between the Higgs field and all fermions (except neutrinos) lead to (i) mass terms for these fermions, and (ii) couplings between these fermions and the Higgs particle h . Through the covariant derivative D_μ in the Higgs kinetic energy, the Higgs particle also couples to the massive gauge bosons $W^\pm$ and Z .

Summing up all terms, the full Standard Model Lagrangian at this point reads (before spontaneous symmetry breaking)

$$\begin{aligned}\mathcal{L}_{\text{SM}} = & -\frac{1}{4} B^{\mu\nu} B_{\mu\nu} - \frac{1}{2} W_i^{\mu\nu} W_{\mu\nu}^i - \frac{1}{2} G_a^{\mu\nu} G_{\mu\nu}^a \\ & + (\partial_\mu \phi)^\dagger (\partial^\mu \phi) + \lambda (v^2 \phi^\dagger \phi - (\phi^\dagger \phi)^2) \\ & + \sum_{\substack{f=L,e_R, \\ Q_L,u_R,d_R}} \bar{f} i \gamma^\mu D_\mu f \\ & + \frac{\sqrt{2}}{v} (m_e \bar{L} \phi e_R + m_d \bar{Q}_L \phi d_R + m_u \bar{Q}_L \phi u_R + \text{h.c.}),\end{aligned}\tag{9.59}$$

$$D_\mu = \partial_\mu - i g_1 \frac{Y}{2} B_\mu - i g_2 \frac{\tau^i}{2} W_{i\mu} - i g_3 \frac{\lambda^a}{2} G_{a\mu}.\tag{9.60}$$

Here, we chose to express the couplings μ^2 , g_e , g_d , and g_u in terms of the Higgs vacuum value v and the masses m_e , m_d , and m_u . As usual, the terms in the last two lines of the Lagrangian replicate for the second and third families of fermions.

10 Scattering and Decay

Now that we have understood the basic structure of the Standard Model, we want to learn how to compute its predictions, which can then be compared to measurements.

10.1 S-Matrix and Cross Sections

S-Matrix. Every Lagrangian quantum field theory provides the rules to compute matrix elements (transition probability amplitudes) between different states. For particle physics experiments (e. g. in particle colliders), one is interested in *scattering* between states that consist of various incoming and outgoing particles. An important underlying assumption is that all incoming and outgoing particles can be treated as *free particles* as long as they are sufficiently far from the interaction region. The appropriate Hilbert space for such states far from the interaction region therefore is the Fock space that is spanned by arbitrary *multi-particle states*, which are defined as *products of free one-particle states*:

$$\mathcal{H} = \bigoplus_{n=0}^{\infty} \mathcal{H}_n, \quad \mathcal{H}_n = \mathcal{H}_1 \otimes \cdots \otimes \mathcal{H}_1 = \mathcal{H}_1^{\otimes n}.\tag{10.1}$$

In a scattering event, the incoming state is typically a pure state consisting of two particles. As the particles approach each other, they cease being free, but start interacting. Their quantum mechanical interaction is completely described by the Lagrangian of the theory, which dictates the time evolution of the combined state in the scattering region. After the scattering, what emerges is in general a *linear combination (superposition)* of many different free multi-particle states. One of these states is what will be measured in the detector. The quantum operator that maps incoming free multi-particle states to superpositions of outgoing free multi-particle states is the *S-matrix or scattering matrix* S ,

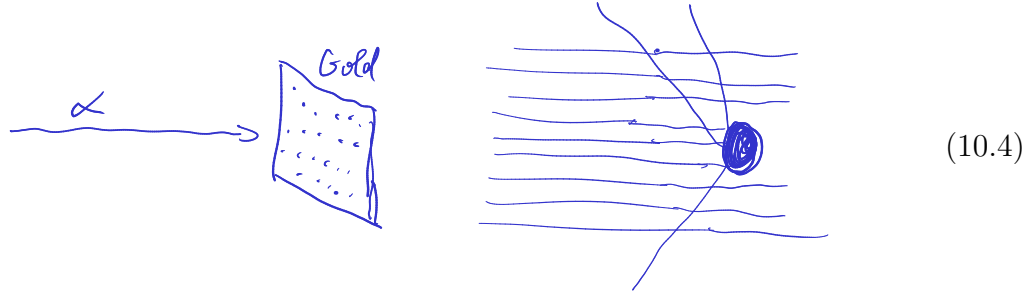
$$S : \mathcal{H} \rightarrow \mathcal{H}.\tag{10.2}$$

The S-matrix is a central object in every quantum field theory. What one typically computes are *S-matrix elements*

$$\langle f | S | i \rangle\tag{10.3}$$

between fixed incoming states $|i\rangle$ and outgoing states $|f\rangle$. Such matrix elements can be computed via Feynman rules and Feynman diagrams that follow from the Lagrangian.

Cross Sections. On the other hand, what one typically can measure in an experiment are *cross sections*. Classically, the cross section of an object is the area of its profile in the plane perpendicular to an incident beam. For example, think of Rutherford's experiment, where a beam of α -particles is directed at a gold foil, and picture the gold nuclei as balls of radius r :



The classical cross section of a nucleus is its area: $\sigma = \pi r^2$. In terms of the beam properties,

$$\sigma = \frac{\text{number of scattered particles}}{\text{time} \times \text{particle flux}} = \frac{N}{T \Phi} ,$$

$$\Phi = \text{flux} = \text{number density in beam} \times \text{beam velocity} . \quad (10.5)$$

Here, T is the total time duration of the experiment, N is the number of particles that got scattered, and Φ is the incoming flux of particles. T and Φ depend on the setup of the experiment, but N depends on the microscopic interactions of the beam with the target. It is also natural to measure the differential cross section $d\sigma/d\Omega$, which gives the number of scattered particles in a certain solid angle $d\Omega$. Classically, this gives us information on the shape of the object, or its potential.

In quantum theory, we can only know the *probability* for a particle to scatter or not. Classically, the probability is $P = N/N_{\text{inc}}$, where N is the number of particles that scatter, and N_{inc} is the total number of incident particles. It is then natural to define the *quantum mechanical cross section* as

$$\sigma = \frac{1}{T \Phi} P , \quad d\sigma = \frac{1}{T \Phi} dP , \quad (10.6)$$

where P is now the quantum mechanical probability for a particle to scatter, and the flux Φ is now normalized as if the beam has just one particle. The differential quantities $d\sigma$ and dP depend on the kinematics of the final particles, such as their angles and energies. The differential number of scattering events per unit time measured in a collider experiment is

$$\frac{dN}{dt} = L \sigma , \quad (10.7)$$

where L is the *luminosity*, which is defined by this equation. The luminosity is the rate of scattering events that are measured per unit of cross section, and depends on the properties of the beam and target. The integral of L over time is called the *integrated luminosity*,

$$L_{\text{int}} = \int L dt = \int \frac{dN}{\sigma} . \quad (10.8)$$

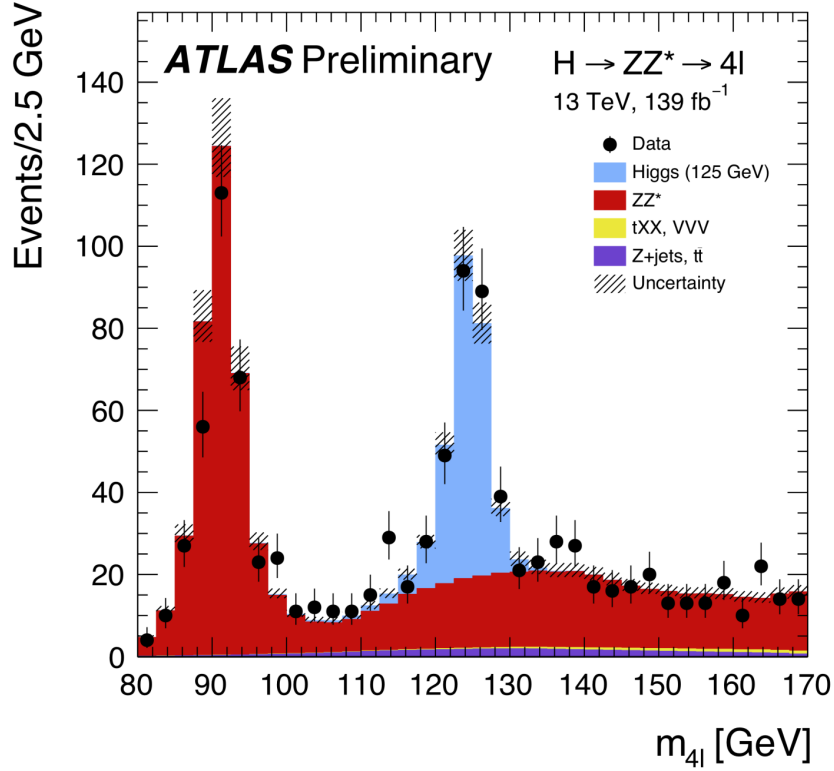


Figure 1: ATLAS 2019 four-lepton final state data. (Image: ATLAS Collaboration, CERN)

The integrated luminosity is the total number of measured events per unit cross section, and is an important measure for any scattering experiment: The higher the integrated luminosity, the more data the experiment produces.

Cross sections are typically expressed in “barn” (b), where $1 \text{ b} = 10^{-28} \text{ m}^2$. The origin of this term is that inducing nuclear fission by hitting ^{235}U with neutrons is as easy as hitting the broad side of a barn. The neutron- ^{235}U scattering cross section is around 1 barn. The luminosity L has dimension $\text{m}^{-2} \text{ s}^{-1}$, or $\text{b}^{-2} \text{ s}^{-1}$. The integrated luminosity L_{int} has dimension m^{-2} or b^{-1} . Typical accelerators achieve luminosities in the range of a few fb^{-1} .

In practice, experimental data is often presented as the total number of events seen for a given integrated luminosity. For example, Figure 1 shows the differential cross section for final states with four leptons from colliding proton initial states, as measured by the ATLAS experiment at the LHC. The data shown in the figure is differential in the Lorentz invariant mass of the four leptons $m_{4l}^2 = (p_1 + p_2 + p_3 + p_4)^2$. Each data point shows the number of events where the measured mass lies in the respective 2.5 GeV interval. As indicated in the figure, the experiment has an integrated luminosity of $L_{\text{int}} = \int L dt = 139 \text{ fb}^{-1}$ at a center-of-mass energy of 13 TeV. To compare this data to theory, one would calculate $d\sigma/dm_{4l}$ at the given energy (13 TeV) from quantum field theory, and multiply by the luminosity. The four-lepton final state can come from various sources, whose theoretical predictions are shown in the solid histogram. For example, the red histogram shows the contribution to the cross section from a pair of intermediate Z bosons. The sum of histograms agrees very well with the data if a 125 GeV Higgs boson is included in the theory as a possible intermediate state (light blue).

Computing Cross Sections. We want to relate the differential cross section measured by experiments to the S-matrix elements computable from field theory. Let us focus on the case where two particles collide, that is the initial state $|i\rangle$ is a two-particle state. We thus consider differential cross sections for $2 \rightarrow n$ processes:

$$p_1 + p_2 \rightarrow \sum_{j=3}^{n+2} p_j . \quad (10.9)$$

In the rest frame of one of the colliding particles, the flux (for a single particle) is the velocity of the incoming particle divided by the volume:

$$\Phi = \frac{|\mathbf{v}|}{V} . \quad (10.10)$$

In a different frame, e. g. the center-of-mass frame, particles come in from both sides, and the flux is determined from the difference between the particle's velocities:

$$\Phi = \frac{|\mathbf{v}_1 - \mathbf{v}_2|}{V} . \quad (10.11)$$

The cross section is therefore

$$d\sigma = \frac{1}{T\Phi} dP = \frac{V}{T|\mathbf{v}_1 - \mathbf{v}_2|} dP \quad (10.12)$$

In quantum theory, probabilities are given by squares of amplitudes. The normalized differential probability dP therefore is

$$dP = \frac{|\langle f|S|i\rangle|^2}{\langle f|f\rangle\langle i|i\rangle} d\Pi . \quad (10.13)$$

Here, $d\Pi$ is the region of phase space (final state momenta) that we are looking at. The differential region $d\Pi$ is proportional to the differential momentum $d^3\mathbf{p}_j$ of each final state j , and must integrate to 1. It therefore must be

$$d\Pi = \prod_{j=1}^n \frac{V}{(2\pi)^3} d^3\mathbf{p}_j \quad (10.14)$$

This integrates to $\int d\Pi = 1$, because (for a one-dimensional integral)

$$\int \frac{dp}{2\pi} = \frac{1}{L} , \quad (10.15)$$

where V is the total volume. This relation follows from

$$\begin{aligned} \int_L dx = L , \quad \int_L \delta(x) dx = 1 & \Rightarrow \delta(x=0) = \frac{1}{L} , \\ \delta(x) = \int \frac{dp}{2\pi} e^{ipx} & \Rightarrow \delta(x=0) = \int \frac{dp}{2\pi} . \end{aligned} \quad (10.16)$$

The normalization factors $\langle f|f\rangle$ and $\langle i|i\rangle$ in the denominator of (10.13) are necessary because the one-particle states may not be normalized to $\langle f|f\rangle = \langle i|i\rangle = 1$. In fact, such a normalization would not be Lorentz-invariant and hence impractical. Instead, in

quantum field theory, one-particle states $|p\rangle$ of four-momentum $p = (E; \mathbf{p})$ in a volume V are normalized such that

$$\langle p|p\rangle = 2EV. \quad (10.17)$$

For the initial state $|i\rangle = |p_1\rangle|p_2\rangle$ and final state $|f\rangle = \prod_{j=3}^{n+2}|p_j\rangle$, we therefore find

$$\langle i|i\rangle = (2E_1V)(2E_2V), \quad \langle f|f\rangle = \prod_{j=3}^{n+2}(2E_jV). \quad (10.18)$$

We will see that all volume factors V will drop out at the end, such that the $V \rightarrow \infty$ limit becomes trivial.

The only thing left is the S-matrix element $\langle f|S|i\rangle$. These S-matrix elements are usually computed perturbatively via Feynman diagrams. In a free theory without interactions, the S-matrix is simply the identity operator $\mathbf{1}$. One typically splits off this part,

$$S = \mathbf{1} + i\mathcal{T}, \quad (10.19)$$

where $\mathcal{T}$ is called the *transfer matrix* and describes deviations from the free theory. The factor of i is a convention, motivated by writing S as $S = e^{i\hat{T}}$ (even though $\mathcal{T}$ is not exactly $\hat{T}$). Since the S-matrix vanishes unless the initial and final states have the same total four-momentum, it is helpful to factor out an overall momentum-conserving delta function,

$$\mathcal{T} = (2\pi)^4 \delta^4(\Sigma p) \mathcal{M}, \quad \delta^4(\Sigma p) := \delta^4(P_i^\mu - P_f^\mu), \quad (10.20)$$

where P_i^μ is the sum of all the initial particles' momenta, and P_f^μ is the sum of all the final particles' momenta. We therefore have

$$\langle f|S|i\rangle = \delta_{fi} + i(2\pi)^4 \delta^4(\Sigma p) \langle f|\mathcal{M}|i\rangle. \quad (10.21)$$

The non-trivial part of the S-matrix is $\mathcal{M}$. In quantum field theory, “matrix elements” usually means $\langle f|\mathcal{M}|i\rangle$. The case where initial and final state are identical is special and needs to be treated differently. Here we focus on the case where the final and initial states are different, $|f\rangle \neq |i\rangle$, which is the interesting part that goes into the differential cross section. For this part, the identity operator contributes nothing, $\delta_{fi} = 0$. For the cross section, we need the absolute square of the second term. It looks worrisome to take the square of a delta function,

$$\delta^4(\Sigma p) \delta^4(\Sigma p) = \delta^4(0) \delta^4(\Sigma p). \quad (10.22)$$

However, this is actually simple to resolve, as long as we work with a finite volume V . We only take the $V \rightarrow \infty$ limit at the end (it will be trivial by then). In one dimension, the delta function satisfies

$$2\pi \delta(p) = \int e^{ipx} dx \quad \Rightarrow \quad \delta(p=0) = \frac{1}{2\pi} \int dx = \frac{L}{2\pi}. \quad (10.23)$$

In three spatial dimensions, this generalizes to

$$\delta^3(p=0) = \frac{1}{(2\pi)^3} \int d^3x = \frac{L^3}{(2\pi)^3} = \frac{V}{(2\pi)^3}. \quad (10.24)$$

Adding time as the fourth dimension, we get

$$\delta^4(p=0) = \frac{1}{(2\pi)^4} \int d^4x = \frac{TV}{(2\pi)^4}. \quad (10.25)$$

The absolute square of the matrix element for $|f\rangle \neq |i\rangle$ therefore is

$$|\langle f|S|i\rangle|^2 = (2\pi)^8 \delta^4(0) \delta^4(\Sigma p) |\mathcal{M}|^2 = (2\pi)^4 TV \delta^4(\Sigma p) |\mathcal{M}|^2, \quad (10.26)$$

where $|\mathcal{M}|^2 := |\langle f|\mathcal{M}|i\rangle|^2$. Putting everything together, we find for the differential probability (10.13)

$$\begin{aligned} dP &= \frac{|\langle f|S|i\rangle|^2}{\langle f|f\rangle\langle i|i\rangle} d\Pi \\ &= \frac{(2\pi)^4 TV \delta^4(\Sigma p) |\mathcal{M}|^2}{(2E_1 V)(2E_2 V)} \prod_{j=3}^{n+2} \left(\frac{1}{2E_j V} \frac{V}{(2\pi)^3} d^3\mathbf{p}_j \right) \\ &= \frac{T}{V} \frac{|\mathcal{M}|^2}{(2E_1)(2E_2)} d\Pi_{\text{LIPS}}, \end{aligned} \quad (10.27)$$

where $d\Pi_{\text{LIPS}}$ is the *Lorentz-invariant phase space* differential

$$d\Pi_{\text{LIPS}} := (2\pi)^4 \delta^4(\Sigma p) \prod_{\substack{\text{final} \\ \text{states } j}} \left(\frac{1}{2E_j} \frac{d^3\mathbf{p}_j}{(2\pi)^3} \right) \quad (10.28)$$

For the differential cross section, we therefore find

$$\boxed{d\sigma = \frac{V}{T |\mathbf{v}_1 - \mathbf{v}_2|} dP = \frac{|\mathcal{M}|^2}{(2E_1)(2E_2) |\mathbf{v}_1 - \mathbf{v}_2|} d\Pi_{\text{LIPS}}.} \quad (10.29)$$

All factors of V and T have dropped out, so it is now trivial to take the limit $V \rightarrow \infty$ and $T \rightarrow \infty$. Recall that $\mathbf{v} = \mathbf{p}/E = \mathbf{p}/p_0$.

Decay Rates. Consider an unstable particle, that is a particle that can decay into two (or more) other particles. Let dP be the differential probability that the one-particle state decays into a given multi-particle state during a time period T . The *differential decay rate* $d\Gamma$ then is

$$d\Gamma = \frac{1}{T} dP. \quad (10.30)$$

From a quantum field theory viewpoint, a decay is a $1 \rightarrow n$ scattering process. Following the same steps as for the differential cross section, the decay rate for a state $|p_1\rangle$ can be written as

$$\boxed{d\Gamma = \frac{|\mathcal{M}|^2}{2E_1} d\Pi_{\text{LIPS}}.} \quad (10.31)$$

Note that a decay rate can never be Lorentz-invariant. The above is the decay rate in the rest frame of the particle. If the particle is moving at relativistic speeds, it will decay much slower due to time dilation. The decay rate in a boosted frame can be calculated with special relativity.

10.2 Lifetime and Decay Width

Lifetime. We want to consider the lifetime of an unstable particle that is at rest. The wave function of a single-particle state should be of the form

$$\psi(t) = \psi(0)e^{-iEt}. \quad (10.32)$$

If the energy E is purely real, then $|\psi(t)|^2 = |\psi(0)|^2$, so there is no transition to another state, and hence no decay, which is not satisfactory. Instead, we expect

$$|\psi(t)|^2 = |\psi(0)|^2 e^{-t/\tau}, \quad (10.33)$$

where τ is the mean *lifetime* of the particle described by the state ψ . After one lifetime, the probability that the particle has not decayed is $1/e$. The above formula suggests to consider a complex energy

$$\hat{E} = E_0 - i\Gamma/2, \quad (10.34)$$

with E_0 and Γ real. By comparing with the definition of the lifetime τ , we find

$$\Gamma = \frac{1}{\tau}. \quad (10.35)$$

This means that Γ is the *decay rate* of the particle (which justifies the symbol).

Decay Width. What does a complex energy mean? To see one important implication, let us Fourier transform the time variable to energy modes,

$$\begin{aligned} \tilde{\psi}(E) &= \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\infty} e^{iEt} \psi(t) dt \\ &= \frac{\psi(0)}{\sqrt{2\pi}} \int_{-\infty}^{\infty} e^{i(E-E_0)t - \Gamma t/2} dt. \end{aligned} \quad (10.36)$$

Now we consider the case that the particle came into existence at time $t = 0$, for example due to the collision of two other particles, or by emission. We hence assume that $\psi(t) = 0$ for $t < 0$. The integral then becomes

$$\tilde{\psi}(E) = \frac{\psi(0)}{\sqrt{2\pi}} \int_0^{\infty} e^{i(E-E_0)t - \Gamma t/2} dt. \quad (10.37)$$

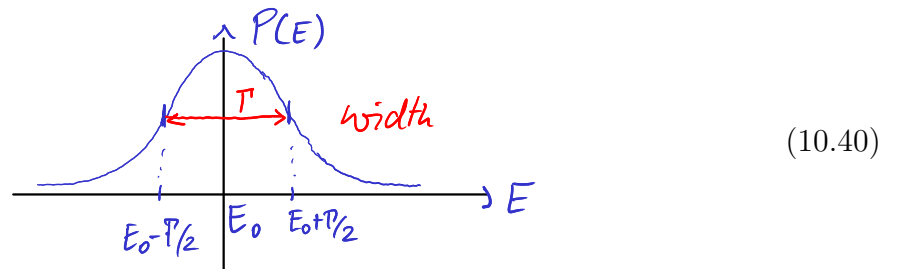
The integral is easy to evaluate. The contribution from $t = \infty$ vanishes. What remains is

$$\tilde{\psi}(E) = \frac{-i\psi(0)}{\sqrt{2\pi}} \frac{1}{E - E_0 + i\Gamma/2}. \quad (10.38)$$

The probability to find the state at energy E therefore is

$$P(E) = |\tilde{\psi}(E)|^2 = \frac{|\psi(0)|^2}{2\pi} \frac{1}{(E - E_0)^2 + \Gamma^2/4}. \quad (10.39)$$

This shows that an unstable state does not have a definite energy, but rather exists in a range of energies E , spread about a central value E_0 :



It is clear that $P(E)$ falls to half its maximal value when $E - E_0 = \pm\Gamma/2$. The *width* of the energy distribution around E_0 is therefore Γ . For this reason, Γ is also called the *decay width*. The terms “decay width” and “decay rate” are used interchangeably.

The energy distribution makes perfect sense from a quantum theory viewpoint: A central element of quantum theory is Heisenberg’s uncertainty principle, which for the dual variables energy E and time t reads

$$\sigma_E \sigma_t \geq \frac{\hbar}{2} \quad \Rightarrow \quad \Delta E \Delta t \gtrsim \hbar, \quad (10.41)$$

where σ_E and σ_t are the standard deviations of energy and time. As is well-known, the uncertainty principle is saturated (the uncertainty is minimized) for a Gaussian normal distribution: A Gaussian distribution $f(E)$ in energy of width ΔE corresponds (via Fourier transformation) to a Gaussian distribution $\tilde{f}(t)$ in time of width $\Delta t = \hbar/\Delta E$. The distribution (10.40) is similar to a Gauss distribution near its peak (though the tails are flatter). We can identify its width as the uncertainty in energy, $\Delta E = \Gamma$. Similarly, we can identify the width of the corresponding probability distribution in time as the average lifetime $\tau = \Delta t$. And indeed the two are related by (reinstating the factor $\hbar$ which we usually set to 1):

$$\Delta E = \Gamma = \frac{\hbar}{\tau} = \frac{\hbar}{\Delta t}, \quad (10.42)$$

in agreement with the uncertainty principle.

The take-home message is that the width Γ of the energy distribution of a particle is inversely proportional to its lifetime τ . Here, energy equals mass, since we consider a particle at rest. Short-lived particles have broad energy distributions. Conversely, the more stable a particle, the more narrow its energy distribution. The energy distribution of a perfectly stable particle is a delta function, i. e. the particle has a definite energy (and therefore a definite mass).

10.3 Scattering Through a Resonance

Resonances. New physics often appears through the production of a new particle, that then decays into other particles. For example, consider the process

$$A + B \rightarrow R \rightarrow C + D. \quad (10.43)$$

For example, A and B as well as C and D could be e^+e^- , or a quark pair, the intermediate particle R could be a $W^\pm$ or Z boson. To understand this process, we look at the amplitude

$$A + B \rightarrow C + D. \quad (10.44)$$

In quantum field theory, this amplitude receives many contributions. If the theory contains another particle R that couples to the initial and final particles through interaction terms

$$g_{AB}ABR \quad \text{and} \quad g_{CD}CDR \quad (10.45)$$

in the Lagrangian, then the leading-order contribution to this amplitude is given by Feynman diagrams of the form

$$\mathcal{M} \sim \frac{g_{AB} g_{CD}}{p^2 - m^2}, \quad (10.46)$$

where m is the mass of the intermediate particle R , and $p^\mu = p_A^\mu + p_B^\mu$ is the internal four-momentum. Its Lorentz-invariant square is

$$s := p^2 = (E_A + E_B)^2 - (\mathbf{p}_A + \mathbf{p}_B)^2. \quad (10.47)$$

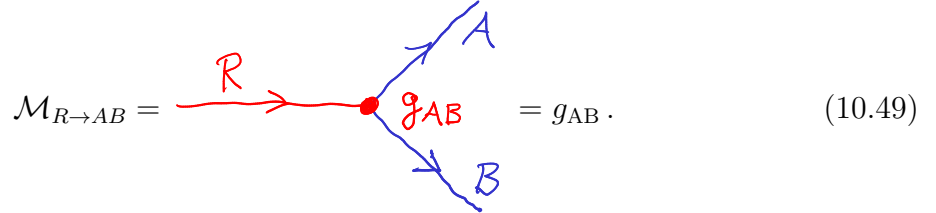
Evaluating it in the center-of-mass frame where $\mathbf{p}_A + \mathbf{p}_B = 0$ gives

$$p^2 = E^2 = (E_A + E_B)^2 \geq (m_A + m_B)^2. \quad (10.48)$$

Consider first the case that the mass m of the intermediate particle is smaller than the total mass $m_A + m_B$ of the initial state. (For example, this would be the case if all particles A, B, C, D, R are identical.) In this case, the propagator $1/(p^2 - m^2)$ in (10.46) is always finite, and the internal state is called “*off-shell*”, because it does not satisfy its mass-shell condition $p^2 = m^2$. Such a state is referred to as a “*virtual particle*”, because it cannot exist as a freely propagating state, as real particles do.

However, if the internal mass m is bigger than the initial-state mass $m_A + m_B$, then there exists physical initial momenta for which $p^2 = (p_A + p_B)^2 = m^2$. In this case, the internal state becomes “*on-shell*”, which means that it exists as a real particle (that can propagate). For such initial states, the propagator and therefore the amplitude becomes infinite. An infinite amplitude is not physical and thus cannot be correct, so we must be missing something!

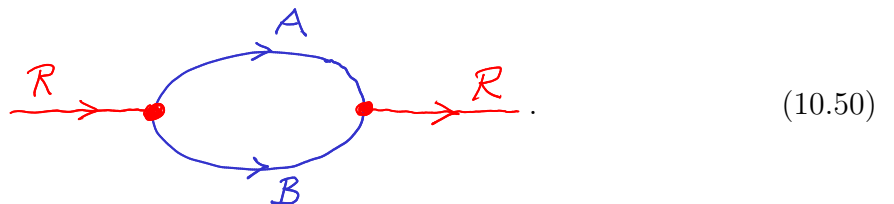
What happens once $m > m_A + m_B$ is that the intermediate state R becomes *unstable*: Due to its large mass, it can decay into particles of lower masses, namely into a pair A, B . The amplitude for this decay is given by the diagram



$$\mathcal{M}_{R \rightarrow AB} = g_{AB} = g_{AB}. \quad (10.49)$$

Note that this amplitude is independent of the particles' masses. But as long as $m < m_A + m_B$, the momentum delta function $\delta^4(p_R - p_A - p_B)$ has no support and hence the amplitude vanishes. Only when $m > m_A + m_B$ does the phase space open up and the decay rate becomes non-zero.

We saw previously that unstable particles have complex energies. Let us see how this comes about. The eigenvalues of any quantum operator crucially depend on the boundary conditions. In quantum mechanics, we usually require the wave function to asymptote to zero at spatial infinity. The energy operator (Hamiltonian) then has real eigenvalues. For a decay process on the other hand, there must be a non-trivial outgoing wave function at infinity. That wave function will be complex, and then also the energy eigenvalues become complex. In quantum field theory, not only transition amplitudes receive quantum corrections, but also eigenvalues, and in particular particle masses. Corrections to the mass of a particle arise due to Feynman diagrams that correct the propagator of the respective field. The simplest such diagram is



$$(10.50)$$

Such corrections lead to a *renormalized mass* m_R , which is the mass that one actually measures experimentally. The mass parameter m in the Lagrangian is called the *bare mass*. And what one finds is that when $m > m_A + m_B$, the diagram above has a non-zero imaginary part. In fact, the *optical theorem* states that the imaginary part of the mass resulting from such “bubble” diagrams is always equal to $-\Gamma_R/2$, where Γ_R is the *total decay rate* of R . One therefore has in general

$$m \rightarrow m_R - i\Gamma_R/2, \quad (10.51)$$

where m_R is the renormalized mass. This is exactly the complex energy that we encountered before, and we have sketched how it arises from quantum field theory.

The formula $g_{AB} g_{CD}/(p^2 - m^2)$ for the amplitude remains correct, but m is replaced by $m_R - i\Gamma_R/2$. The process gets particularly interesting when $\Gamma_R \ll m_R$: In this case, the intermediate state has a lifetime $\tau_R = 1/\Gamma_R$ that is much bigger than the Compton wavelength $1/m_R$. Hence the state R can be considered as a real intermediate particle that is produced in the collision, physically lives for some time, and then decays. In this case, the intermediate state R is called a *resonance*. In the resonance approximation $\Gamma_R \ll m_R$, we can approximate

$$(m_R - i\Gamma_R/2)^2 \approx m_R^2 - im_R\Gamma_R, \quad (10.52)$$

and the amplitude becomes

$$\mathcal{M}_{2 \rightarrow 2} \approx \frac{g_{AB} g_{CD}}{E^2 - m_R^2 + im_R\Gamma_R}, \quad (10.53)$$

where E is the total center-of-mass energy. We see that now at $E = m_R$ the amplitude is no longer divergent: The imaginary part in the denominator regulates the divergence. The amplitude still has a maximum at $E = m_R$, but is no longer infinite. Near $E = m_R$, we can further approximate

$$E^2 - m_R^2 = (E - m_R)(E + m_R) \approx 2m_R(E - m_R), \quad (10.54)$$

and the amplitude becomes

$$\mathcal{M}_{2 \rightarrow 2} \approx \frac{1}{2m_R} \frac{g_{AB} g_{CD}}{E - m_R + i\Gamma_R/2}. \quad (10.55)$$

Resonance Cross Section. Let us use this amplitude to compute the cross section of the process. We earlier derived the general formula for the cross section:

$$d\sigma = \frac{|\mathcal{M}|^2}{2E_1 2E_2 |\mathbf{v}_1 - \mathbf{v}_2|} d\Pi_{\text{LIPS}}. \quad (10.56)$$

For the special case of $2 \rightarrow 2$ scattering, we can go to the center-of-mass frame. Using momentum conservation, the four-momenta can be written as

$$\begin{aligned} p_1 &= (E_1, \mathbf{p}), & p_3 &= (E_3, \mathbf{p}'), \\ p_2 &= (E_2, -\mathbf{p}), & p_4 &= (E_4, -\mathbf{p}'), \end{aligned} \quad (10.57)$$

where $\mathbf{p}'$ and $\mathbf{p}$ are related by energy conservation. The phase space differential then can be simplified to (see also Problem 6.1)

$$d\Pi_{\text{LIPS}} = \frac{1}{16\pi^2} \frac{|\mathbf{p}'|}{\sqrt{s}} d\Omega. \quad (10.58)$$

Here, $s = (p_1 + p_2)^2 = (E_1 + E_2)^2$ is the square of the center-of-mass energy, $d\Omega = \sin\theta d\theta d\phi$ is the solid angle differential, and $|\mathbf{p}'|$ is determined by the conservation laws $\delta^4(p_1 + p_2 - p_3 - p_4)$:

$$|\mathbf{p}'| = \frac{1}{2\sqrt{s}} \sqrt{s^2 - 2s(m_3^2 + m_4^2) + (m_3^2 - m_4^2)^2}. \quad (10.59)$$

We can also simplify the denominator in the expression for $d\sigma$,

$$E_1 E_2 |\mathbf{v}_1 - \mathbf{v}_2| = E_1 E_2 \left| \frac{\mathbf{p}_1}{E_1} - \frac{\mathbf{p}_2}{E_2} \right| = |\mathbf{p}_1 E_2 - \mathbf{p}_2 E_1| = \sqrt{s} |\mathbf{p}|. \quad (10.60)$$

The $2 \rightarrow 2$ cross section therefore simplifies to (in the center-of-mass frame)

$$\boxed{d\sigma = \frac{|\mathcal{M}|^2}{64\pi^2 s} \frac{|\mathbf{p}'|}{|\mathbf{p}|} d\Omega.} \quad (10.61)$$

Here $|\mathbf{p}|$ is given by the same formula as $|\mathbf{p}'|$ with $(m_3, m_4) \rightarrow (m_1, m_2)$. We can now plug in our formula for the amplitude $\mathcal{M}_{2 \rightarrow 2}$ near a resonance. Noting that $\sqrt{s} = E \approx m_R$ near the resonance, this gives

$$\begin{aligned} d\sigma &\approx \frac{1}{64\pi^2 m_R^2} \left| \frac{1}{2m_R} \frac{g_{AB} g_{CD}}{E - m_R + i\Gamma_R/2} \right|^2 \frac{|\mathbf{p}'|}{|\mathbf{p}|} d\Omega \\ &= \frac{g_{AB}^2 g_{CD}^2}{(16\pi)^2 m_R^4} \frac{|\mathbf{p}'|}{|\mathbf{p}|} \frac{1}{(E - m_R)^2 + \Gamma_R^2/4} d\Omega. \end{aligned} \quad (10.62)$$

We can also compute the width Γ_R in the denominator, which is the decay rate of the resonance R . The general decay rate formula derived earlier,

$$d\Gamma = \frac{|\mathcal{M}|^2}{2E_1} d\Pi_{\text{LIPS}}, \quad (10.63)$$

for a $1 \rightarrow 2$ process also simplifies: In the center-of-mass frame (where the initial particle is at rest), the decay rate for the process $R \rightarrow C + D$ becomes (see Problem 6.1):

$$\boxed{\Gamma_{R \rightarrow CD} = \frac{S_{CD} |\mathbf{p}'|}{8\pi m_R^2} |\mathcal{M}_{R \rightarrow CD}|^2 \approx \frac{S_{CD} |\mathbf{p}'| g_{CD}^2}{8\pi m_R^2}.} \quad (10.64)$$

Here, we used that the amplitude $\mathcal{M}_{R \rightarrow CD} \approx g_{CD}$, and S_{CD} is a symmetry factor that equals 1/2 if the particles C and D are identical, and is 1 otherwise. Similarly, we find for the decay process $R \rightarrow A + B$:

$$\Gamma_{R \rightarrow AB} \approx \frac{S_{AB} |\mathbf{p}| g_{AB}^2}{8\pi m_R^2}. \quad (10.65)$$

The *total decay rate* Γ_R of the particle R is the sum of decay rates of all decay channels of the particle:

$$\Gamma_R = \sum_{\substack{\text{all possible} \\ \text{final states } f}} \Gamma_{R \rightarrow f}, \quad (10.66)$$

and it is this total decay rate that stands in the denominator of the cross section near the resonance. The resonance R may have further decay channels besides AB and CD , hence

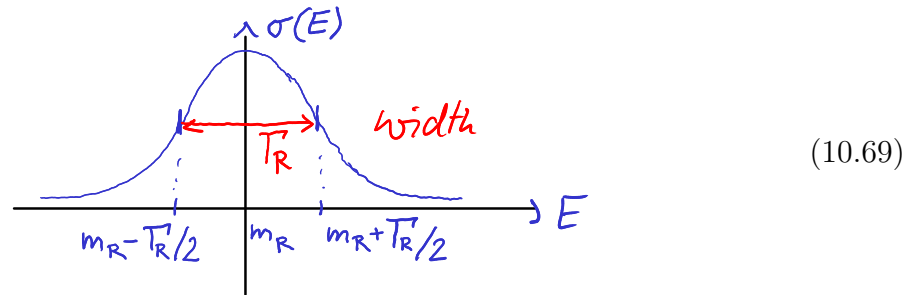
we leave Γ_R in the denominator as it is. But we can eliminate g_{AB} and g_{CD} in favor of $\Gamma_{R \rightarrow AB}$ and $\Gamma_{R \rightarrow CD}$, since the decay rates are the quantities that can be observed, whereas the couplings could be effective couplings with no fundamental significance, for example if the resonance R is a bound state of more fundamental components. Hence we re-write the cross section as

$$d\sigma \approx \frac{1}{4S_{AB}S_{CD}} \frac{1}{|\mathbf{p}|^2} \frac{\Gamma_{R \rightarrow AB} \Gamma_{R \rightarrow CD}}{(E - m_R)^2 + \Gamma_R^2/4} d\Omega. \quad (10.67)$$

Performing the angular integrations, we obtain for the total cross section

$$\sigma(E) \approx \frac{\pi}{S_{AB}S_{CD}|\mathbf{p}|^2} \frac{\Gamma_{R \rightarrow AB} \Gamma_{R \rightarrow CD}}{(E - m_R)^2 + \Gamma_R^2/4}. \quad (10.68)$$

This is called the *Breit–Wigner distribution* (after Gregory Breit and Eugene Wigner). We have seen its functional form before, when we looked at the decay probability for an unstable particle. Its graph is bell-shaped, similar to a Gaussian distribution,



The full width of the graph at half its maximal value, or “full-width-half-maximum” (FWHM) is Γ_R . For this reason, Γ_R is also called the *resonance width*. At the same time, Γ_R is the decay rate (decay width) of the unstable particle R responsible for the resonance.

Spin and Color Sums. In all of the above, we have neglected spin degrees of freedom. If the resonance has spin j , we must *sum* over the $2j + 1$ possible spin states of the resonance, which adds an overall factor of $2j + 1$. Moreover, if the initial particles A and B have spins s_A and s_B , and we know their spin states, we have to use the partial width $\Gamma_{R \rightarrow i}$ for these specific spin states. More commonly, one does not know the initial spin states, that is the incoming particles are described by a *completely unpolarized* density matrix

$$\rho_{\text{spin}} = \frac{1}{2s + 1} \sum_{m=-s}^s |m\rangle \langle m|. \quad (10.70)$$

In this case, one has to *average* over the initial spin states by inserting

$$\frac{1}{(2s_A + 1)(2s_B + 1)} \sum_{m_A=-s_A}^{s_A} \sum_{m_B=-s_B}^{s_B}. \quad (10.71)$$

By convention, the partial decay rate $\Gamma_{R \rightarrow f}$ into some final state f is defined as the sum over the decay rates into all possible spin configurations of f (because one does not care about partial rates into specific spin polarizations). The sums over m_A and m_B are therefore *absorbed* into the definition of $\Gamma_{R \rightarrow AB}$. Similarly, the cross section σ is summed over final spin states, and this sum is absorbed into $\Gamma_{R \rightarrow CD}$. Similar arguments apply to

color degrees of freedom: One has to sum over all possible internal color states, and the cross section is averaged over initial and summed over final color states, with the state sums over initial and final states being absorbed into the decay widths. The general form of the resonance cross section therefore is

$$\sigma(E) \approx \frac{(2j+1)c_R}{(2s_A+1)c_A(2s_B+1)c_B} \frac{\pi}{S_{AB}S_{CD}|\mathbf{p}|^2} \frac{\Gamma_{R \rightarrow AB} \Gamma_{R \rightarrow CD}}{(E - m_R)^2 + \Gamma_R^2/4}. \quad (10.72)$$

Here, c_R , c_A , and c_B are the color multiplicities of the internal state and the two initial states. The color multiplicity of a quark is 3, for a gluon it is 8, and for color singlets it is 1. The width Γ_R in the denominator is the *total width* (decay rate) of R , that is the sum over decay rates into all possible final states. In the computation of the *partial widths* $\Gamma_{R \rightarrow AB}$ and $\Gamma_{R \rightarrow CD}$, all final spin and color states should be summed over, as described above.

Example. As an example, consider the scattering of an electron e^- with a positron e^+ , where the outgoing state is another fermion anti-fermion pair $f\bar{f}$:

$$e^- e^+ \rightarrow f \bar{f}. \quad (10.73)$$

When the center-of-mass energy $E = \sqrt{s}$ is close to the mass of the Z -boson, the cross section $\sigma(E)$ will show a peak that has the form of the Breit–Wigner distribution, where the intermediate state (resonance) is the Z boson. In this case, $s_A = s_B = 1/2$, while $J = 1$. All states involved are color singlets, that is $c_A = c_B = c_R = 1$. The two particles in the initial and final states are distinguishable, so $S_{AB} = S_{CD} = 1$. Finally, $m_Z \gg m_e$, and therefore $|\mathbf{p}|^2 \approx m_Z^2/4$. The cross section near the resonance will therefore have the form

$$\sigma(E) \approx \frac{3}{2 \cdot 2} \frac{4\pi}{m_Z^2} \frac{\Gamma_{Z \rightarrow e^+ e^-} \Gamma_{Z \rightarrow f \bar{f}}}{(E - m_Z)^2 + \Gamma_Z^2/4}. \quad (10.74)$$

We see that the cross section as a measurable function of the center-of-mass energy $E = \sqrt{s}$ contains a lot of information: The location of the resonance peak yields the mass m_Z of the Z -boson, the width of the distribution yields its total decay width/rate Γ_Z , and the height of the peak gives us information on the partial decay rates/widths in the numerator.

10.4 The W and Z Widths

In a hadron collider (such as the LHC), $W^\pm$ and Z bosons are produced from collisions of quarks, which are the constituents of all hadrons. The $W^\pm$ and Z bosons then decay into all kinds of particles allowed by the Standard Model.

W Width. Let us compute the decay width for the $W^\pm$ bosons. By looking at the Standard Model Lagrangian, we find that the possible decays of the W^+ boson into two-particle states are

$$\begin{aligned} W^+ &\rightarrow e^+ \nu_e, & W^+ &\rightarrow \mu^+ \nu_\mu, & W^+ &\rightarrow \tau^+ \nu_\tau, \\ W^+ &\rightarrow u \bar{d}, & W^+ &\rightarrow c \bar{s}. \end{aligned} \quad (10.75)$$

The decay into $t\bar{b}$ is excluded by energy conservation, because the top quark ($m_t \approx 173$ GeV) is heavier than the $W^\pm$ ($m_W \approx 80.4$ GeV). First consider the $e^+ \nu_e$ decay channel. The

coupling between W^+ , e^+ , and ν_e is given by the following interaction term (vertex) in the Lagrangian:

$$\frac{g_2}{\sqrt{2}} W_\mu^+ \bar{\nu}_e \gamma^\mu P_L e = \text{diagram} = \text{diagram} \quad (10.76)$$

At leading order, the matrix element $\mathcal{M}$ for the transition $W^+ \rightarrow e^+ \nu_e$ is given by this interaction vertex, where the field operators are replaced by the wave functions of the respective particles. Therefore,

$$\mathcal{M} = \frac{g_2}{\sqrt{2}} W_\mu^+ \bar{\nu}_e \gamma^\mu P_L e, \quad (10.77)$$

where W_μ^+ is the wave function of the W^+ boson, $\bar{e}$ is the wave function of the positron e^+ , and ν_e is the wave function of the electron neutrino.

For the decay rate, we have to sum $|\mathcal{M}|^2$ over final spin states, and average over initial spin states. Instead of doing the detailed computation, we will make a simple estimate. Because of the projector P_L onto left-handed spinors, only one spin state of each e^+ and ν_e interacts with W^+ . The W_μ^+ has three polarization states, but we average over them. So the sum and average over spin states just gives a trivial factor of one. The decay rate must have dimension (1/time). Looking at the general formula (10.64) for Γ , this implies that $|\mathcal{M}|^2$ must have dimension (mass)². The neutrino and electron masses are extremely small compared to the W mass. Neglecting them, the only quantity that can supply the mass dimension is m_W . We do not know the numerical prefactor, but the naive choice is to just replace the product of wave functions for given spin states by m_W , that is $W_\mu^+ \bar{\nu}_e \gamma^\mu P_L e \rightarrow m_W$. Since summing and averaging over spin states gives a trivial factor one, this would result in $|\mathcal{M}|^2 = g_2^2 m_W^2 / 2$. The correct answer is

$$|\mathcal{M}_{W^+ \rightarrow e^+ \nu_e}|^2 = \frac{g_2^2 m_W^2}{3}, \quad (10.78)$$

so our estimate is off by a factor 2/3. We got pretty close with our naive approximation! Using the general formula (10.64), we find for the decay rate

$$\Gamma_{W^+ \rightarrow e^+ \nu_e} \approx \frac{S_{e^+ \nu_e} |\mathbf{p}'|}{8\pi m_W^2} |\mathcal{M}_{W^+ \rightarrow e^+ \nu_e}|^2 = \frac{|\mathbf{p}'|}{8\pi} \frac{g_2^2}{3}. \quad (10.79)$$

Neglecting the electron and neutrino masses, we can identify $|\mathbf{p}'| \approx m_W/2$. Therefore

$$\Gamma_{W^+ \rightarrow e^+ \nu_e} \approx \frac{m_W}{16\pi} \frac{g_2^2}{3} = \frac{1}{12} \alpha_2 m_W, \quad \alpha_2 = \frac{g_2^2}{4\pi}. \quad (10.80)$$

The computation did not depend on the final particles' masses. For all decay channels, the final state is a pair of spin 1/2 fermions. Hence the partial decay width for all decay channels is the same. For the quark anti-quark final states, we also have to sum over all *color states*. Since the W^+ is a color singlet, also the final state must be a color singlet. With three possible colors, a quark anti-quark pair can form three possible singlets: $r\bar{r}$,

$g\bar{g}$, and $b\bar{b}$. Hence we have to count the quark final states with a factor of 3. Counting all possible final states, we find for the total decay width

$$\Gamma_{W^+} = (1 + 1 + 1 + 3 + 3)\Gamma_{W^+ \rightarrow e^+ \nu_e} = \frac{3}{4}\alpha_2 m_W \approx 2 \text{ GeV}. \quad (10.81)$$

An identical computation gives the same result for the charge conjugate W^- : $\Gamma_{W^-} \approx 2 \text{ GeV}$. This is indeed very close to the experimental value $\Gamma_W = 2.085(42) \text{ GeV}$. The $W^\pm$ is a rather narrow resonance, with $\Gamma_W/m_W \approx 1/40$.

Note that the result depends on the *number* N_c of possible colors for each quark, and that the result matches experiment with $N_c = 3$. Final quark and lepton states can clearly be distinguished from each other, and one finds that a $W^\pm$ decays twice as often into quarks as it does into leptons. So even though the measured process does not involve the strong force directly, it confirms that leptons are colorless, and that the number of quark colors must be three.

The Z Width. The width of the Z boson can be computed in a very similar way. In this case, the possible decay channels within the first family are

$$Z \rightarrow e^+ e^-, \quad Z \rightarrow \nu_e \bar{\nu}_e, \quad Z \rightarrow u\bar{u}, \quad Z \rightarrow d\bar{d}, \quad (10.82)$$

and there are similar decay channels for the other two families. Essentially the only difference compared to the $W^\pm$ decay is that one has to sum over left-handed and right-handed fermions, and the squared coupling $(g_2/\sqrt{2})^2$ has to be replaced by the square of the electroweak charge

$$\left(\frac{g_2}{\cos \theta_W} (T_{3f} - Q_f \sin^2 \theta_W) \right)^2. \quad (10.83)$$

The total decay width works out to (see Problem 6.2)

$$\Gamma_Z \approx 2.4 \text{ GeV}. \quad (10.84)$$

This is again very close to the experimental value $\Gamma_Z = 2.4952(23) \text{ GeV}$. The Z boson is the most precisely measured resonance, see Figure 2. The Large Electron-Positron Collider (LEP) at CERN (the predecessor of the LHC, in the same 27 km tunnel), was able to measure both the Z width and the Z mass $m_Z = 91.1875(21) \text{ GeV}$ to an accuracy of $\sim 10 \text{ MeV}$. For the Z mass, this is a relative precision of $\sim 0.002\%$.

To get to this level of precision, the LEP group had to take into account all kinds of environmental effects. For example, tidal forces from the moon distort the rock around the accelerator, changing the 4.3 km radius of the accelerator by $\pm 0.15 \text{ mm}$; this changes the beam energy by $\sim 10 \text{ MeV}$, so the data has to be corrected accordingly! Another effect is the TGV railway line, which leaks currents that return to earth via the nearby Versoix river and the LEP ring: Each time a train passes by, a small current circulates the ring, slightly changing the magnetic field, which again changes the beam energy by $\sim 10 \text{ MeV}$. These examples give an idea of the accuracy to which such experiments are performed.

Lifetimes. From the equation for the decay lifetime, we note that the Z and $W^\pm$ lifetimes are about $\tau_W = 1/\Gamma_W \approx 1/2 \text{ GeV}^{-1}$. In our natural units where $\hbar = c = 1$, we have $1 \text{ GeV}^{-1} = 6.6 \cdot 10^{-25} \text{ sec}$, hence the lifetime is $\tau_W \approx 3 \cdot 10^{-25} \text{ sec}$. In its lifetime, a $W^\pm$ can travel a distance of $\gamma c \tau_W$, where $\gamma = E/m_W$ is the Lorentz factor. Since $E/m_W \leq 100$ for any foreseeable machine, and $c = 3 \cdot 10^{10} \text{ cm/sec}$, the distance a $W^\pm$ can travel is

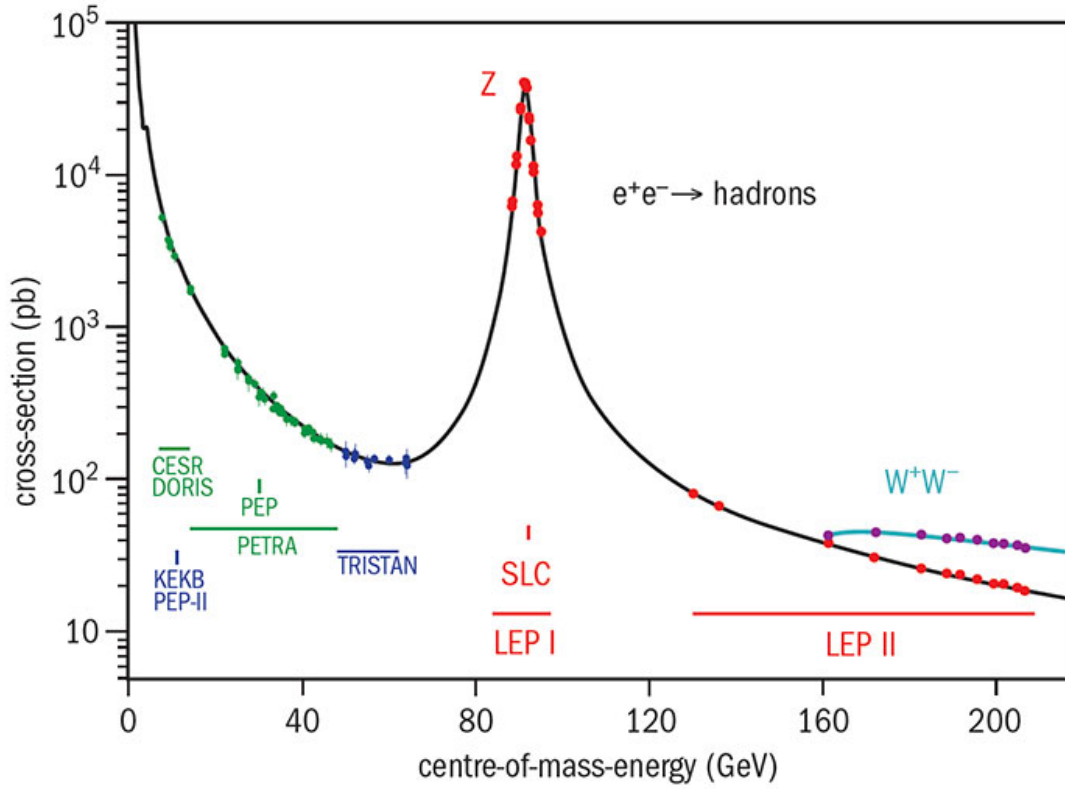


Figure 2: The measured total cross section of processes $e^+e^- \rightarrow \text{hadrons}$ shows a clear resonance peak. The peak is centered at the Z boson mass $m_Z = 91.2 \text{ GeV}$, the width of the peak is the Z boson width Γ_Z . The figure shows data combined from various experiments. (Image from <https://cerncourier.com/a/revisiting-the-b-revolution/>)

$\leq 10^{-12} \text{ cm}$. Detectors can resolve distances of $\sim 10^{-2} \text{ cm}$, so a $W^\pm$ (and similarly a Z) will always decay before it can be detected, and its presence can only be deduced from its decay products.

Invisible Width and Neutrinos. One of the most important aspects of the Z width is that any new particle that has a non-zero weak isospin or non-zero electric charge will couple to the Z boson, and will therefore appear in Z decays. For example, if additional families of particles would exist (besides the three known families), Z could decay into their neutrinos. The partial decay rate into a $\nu\bar{\nu}$ pair is $\approx 160 \text{ MeV}$. Therefore every new neutrino would increase the total decay rate Γ_Z by $\approx 160 \text{ MeV}$, which is clearly in conflict with the measured width Γ_Z .

Charged particles leave clear signatures in detectors, but neutrinos are basically impossible to detect. Nonetheless they contribute to the total width Γ_Z of the Z resonance cross section. The partial decay rates $\Gamma_{Z \rightarrow q\bar{q}}$ and $\Gamma_{Z \rightarrow \ell^+\ell^-}$ into quarks (hadrons) and leptons can be measured separately (from the height of the cross section at the resonance peak). From this data, one can compute the *invisible decay width*

$$\Gamma_Z^{\text{inv}} := \Gamma_Z - \Gamma_{Z \rightarrow \ell^+\ell^-} - \Gamma_{Z \rightarrow q\bar{q}} \quad (10.85)$$

Calling $\Gamma_Z^{\nu\bar{\nu}}$ the expected invisible width for one family of neutrinos, the current experi-

mental result from LEP is $\Gamma_Z^{\text{inv}}/\Gamma_Z^{\nu\bar{\nu}} = 3.04(4)$, which clearly accounts for the three decays $Z \rightarrow \bar{\nu}_e\nu_e$, $Z \rightarrow \bar{\nu}_\mu\nu_\mu$, and $Z \rightarrow \bar{\nu}_\tau\nu_\tau$, and *rules out further families* with light neutrinos with a large confidence level. This result also tells us that *heavy neutrinos cannot exist*, unless they are heavier than $m_Z/2 \approx 45.6$ GeV.

This results seems unsurprising, given our current knowledge. But before LEP started in 1987, the number of quark and lepton families was unknown. In fact, the upper bound on the number of light neutrinos was very weak: There could have been as many as ~ 6000 different neutrinos. These would have completely washed out the Z resonance, which was a realistic fear at the time. But it took only a few weeks of measurements and analysis in 1989 to determine that there are no more than three neutrinos. For an interesting tale of the LEP Z boson measurements and the number of families, see [2].

11 More Aspects of the Standard Model

11.1 Measurements of Parameters

To test more predictions of the Standard Model, one must determine the numerical values of its parameters: Masses of the gauge bosons and fermions, α , $\sin^2 \theta_W$, the QCD coupling.

- The electromagnetic coupling $\alpha = e^2/4\pi$ can be measured in many ways that do not require particle physics.
- The electroweak mixing angle θ_W appears in the electroweak neutral current coupling ($T_3 - Q_3 \sin^2 \theta_W$), and so many different $2 \rightarrow 2$ fermion cross sections depend on it. The fact that all these cross sections are consistent with a unique value of θ_W is a strong consistency check of the Standard Model.
- The coupling $g_2 = e/\sin \theta_W$ can also be measured from muon decay $\mu^- \rightarrow \nu_\mu e^- \bar{\nu}_e$ via a W^- boson:


(11.1)

The decay rate of this process can be computed,

$$\Gamma_\mu = \frac{G_F^2 m_\mu^5}{192\pi^3}, \quad \frac{G_F}{\sqrt{2}} = \frac{g_2^2}{8m_W^2}, \quad (11.2)$$

where G_F is the Fermi coupling. The lifetime of the muon can be measured very accurately,

$$\tau_\mu = \frac{1}{\Gamma_\mu} = 2.196\,981\,1(22) \cdot 10^{-6} \text{ sec}. \quad (11.3)$$

This gives (including quantum corrections to Γ_μ):

$$G_F = 1.166\,378\,7(6) \cdot 10^{-5} \text{ GeV}^{-2}. \quad (11.4)$$

Using $g_2 = e/\sin \theta_W$, $\alpha = e^2/4\pi = 1/137$, and the definition of G_F , this can be written as

$$m_W \approx \frac{37}{\sin \theta_W} \approx 77 \text{ GeV}, \quad (11.5)$$

where we put in $\sin^2 \theta_W = 0.23$ to get the numerical value. This was the historical prediction of the W mass. To detect the W and Z bosons, the collider experiments were designed to be most sensitive in this energy region.

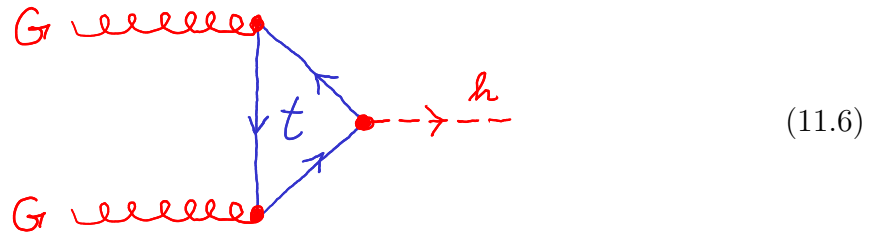
11.2 Higgs Production and Decay

A central ingredient of the Standard Model is the scalar “Higgs field”, whose non-zero vacuum expectation value spontaneously breaks the electroweak gauge symmetry, thereby explaining the masses of the $W^\pm$ and Z bosons. At the same time, the Higgs field explains how all fermions can acquire masses. In turn, the Higgs mechanism firmly predicts the existence of a scalar “Higgs boson” h . Since the Standard Model is consistent with all measured particle phenomena to extraordinary precision, it was firmly believed that the Higgs particle must indeed exist, long before it was confirmed experimentally.

The Higgs boson is difficult to observe because its couplings to the other particles is small (proportional to their masses), and also because its mass m_h was largely unconstrained by the theory: We saw that m_h depends on the parameter λ in the Higgs potential, whose value was unknown. Hence the search for the Higgs boson had to be carried out across a large range of energies. The search was one of the main motivations to build the world’s biggest and most complex experimental facility, the Large Hadron Collider (LHC) at CERN.

The discovery of the Higgs boson was finally announced in 2012/13. All measurements since then show that its interactions and decays are exactly as predicted by the Standard Model. Current studies investigate whether it has all the predicted properties to even higher precision, or whether the theory has to be modified. For example, some theories predict the existence of multiple Higgs bosons.

Higgs Production. The dominant process for Higgs production at the LHC is gluon fusion via a top loop:



The LHC collides protons. Each proton consists of three quarks, that are held together by gluon exchange (the strong force). Two such gluons can produce a Higgs boson through the above process. The particle in the loop could be any quark (any particle coupling to both gluons and the Higgs), but because the coupling to the Higgs is proportional to the particle’s mass, the top quark channel strongly dominates. Other Higgs production

processes are

W/Z fusion:

$t/\bar{t}$ fusion:

W/Z Bremsstrahlung:

(11.7)

But these processes are all much more rare than the top-loop process.

By the time the Higgs particle was found, the total number of Higgs particles produced was about one million. That sounds a lot, but the number of events that were actually detected is much lower: The Higgs boson is observed through its decay products, and isolating those events from all the other processes happening in the collision is difficult (the raw event rate at the LHC is ~ 600 million per second).

Higgs Decay. The most obvious decays are into fermion anti-fermion pairs, and since the Higgs coupling is proportional to the fermion mass, the most frequent decays are into the heaviest fermions:

(11.8)

The Higgs boson can also decay into two W or Z bosons, which will then further decay into quarks or leptons (since the W and Z are too massive, at least one of them will be off-shell):

(11.9)

A Higgs boson could also decay into two gluons, by reversing the gluon-fusion process that dominates its production. By the same process, it can also decay to two photons:

(11.10)

There are two more processes by which a Higgs boson can decay into two photons:

$$(11.11)$$

The total decay rate of the 125 GeV Higgs boson splits into the following *branching ratios* $\Gamma_{h \rightarrow f}/\Gamma_h^{\text{tot}}$ (a branching ratio is the ratio between a partial decay rate $\Gamma_{h \rightarrow f}$ into a specific final state f and the total decay rate Γ_h of the particle):

Final state f	$\Gamma_{h \rightarrow f}/\Gamma_h^{\text{tot}}$	Observed
$b\bar{b}$	0.57	yes
W^+W^-	0.21	yes
2 gluons	0.09	no?
$\tau^+\tau^-$	0.06	yes
$c\bar{c}$	0.03	no?
ZZ	0.03	yes
2 photons	0.002	yes

$$(11.12)$$

The $b\bar{b}$ decay is so dominant because for fermionic decays $\Gamma \sim |\mathcal{M}|^2 \sim m^2$, where m is the mass of the final fermion, and for quarks there is an extra factor of three for the color degree of freedom.

The most common decays are difficult to detect due to a lot of background “noise” from other processes. Even though the branching ratio for the two-photon final state is very small, this channel has a clear signature in the detector, and here the Higgs boson was first seen. At the time of the discovery, the number of relevant events was about $0.002 \cdot 1\,000\,000 = 2000$.

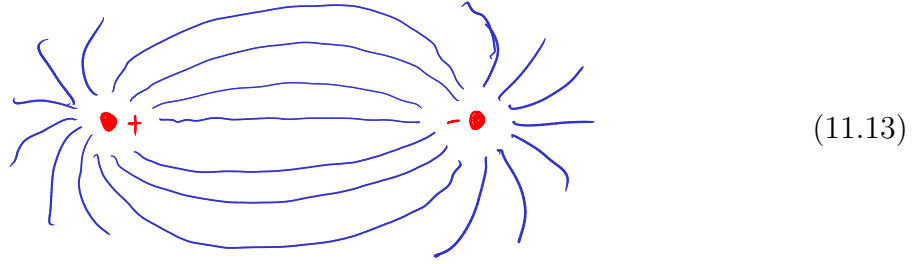
Any scalar particle can decay into two photons, so even though the peak in the two-photon cross section was in the expected Higgs mass range at 125 GeV, it does not clearly identify itself as a Higgs boson resonance. But the decays to W^+W^- and ZZ , as well as to $b\bar{b}$ and $\tau^+\tau^-$ strongly confirm that the newly found particle is indeed the Higgs boson. All of these decays are separately measurable. The W^+W^- and ZZ decays confirm that the Higgs field has a non-zero vacuum expectation value $v \neq 0$ that breaks the electroweak symmetry (recall that the hW^+W^- vertex comes from the $\phi^\dagger W^\dagger W \phi$ term in the Higgs field Lagrangian after symmetry breaking $\phi \rightarrow v + h$). The fermion channels $b\bar{b}$ and $\tau^+\tau^-$ confirm that the Higgs coupling is proportional to the particles’ masses, and that the same mechanism works both for quarks and for leptons.

After the discovery of the Higgs boson, the 2013 Nobel Prize was awarded to François Englert and Peter Higgs, for the theoretical discovery of the Higgs mechanism.

11.3 Color Confinement, Jets, and Hadrons

Color Confinement. The strong force behaves very differently from the more familiar electromagnetic force. The reason is that the force-carrying gluons themselves are color-charged, unlike photons (the carriers of the electromagnetic force) which are electrically

neutral. The EM field between two opposite charges spreads in all spatial directions as the two charges separate:

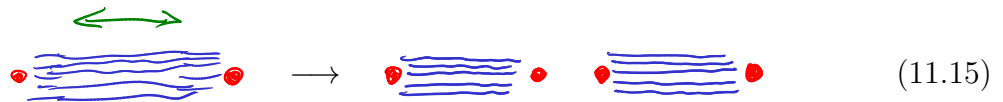


As a result, the EM force decreases as $\sim 1/r^2$ with the distance r between the two charges. On the contrary, because gluons are charged, the color force field lines between two color-charged objects (such as quarks) stick together and form a narrow “color flux tube” that extends between the two charges:



For this reason, the force between two color-charged objects is *constant*, that is *independent of their separation*! For this reason, color-charged objects can never be separated, since this would require an infinite amount of energy. This feature of the strong force is called *color confinement*. It is not fully understood theoretically (in fact there is a \$1 000 000 prize for its theoretical proof), but the above considerations are strongly supported by all observations, as well as by numerical simulations (lattice gauge theory).

Jets. When two color-charged objects (quarks or gluons) are produced in a scattering experiment, and move in different directions at a high energy (with large momentum), the energy in the color flux tube increases until it reaches a level where quark anti-quark pair production becomes energetically favorable. Such a pair can split the color flux tube in two:



This process repeats until most of the energy is absorbed in pair creations. Each high-energy quark (or gluon) hence fragments into a bunch of color-neutral *hadrons* (particles made of quarks and gluons) that all move in roughly the same direction. This process is called *hadronization*. The resulting bunch of particles is called a *jet* and is what one observes in particle detectors. Experimentally, a 10 GeV quark fragments into ~ 7 hadrons, while a 100 GeV quark fragments into ~ 15 or so hadrons. Since the lightest hadrons are pions, they form the majority in jets.

Hadrons. Since quarks and gluons can only occur in color-neutral combinations, they always form bound states, called *hadrons*. There are several ways to form color-neutral states. The most obvious is to combine one color with its anti-color, for example $r\bar{r}$ (red anti-red). Such states are formed from two quarks, and are called *mesons*. The lightest of these are pions, with quark content $u\bar{u}$, $d\bar{d}$, $u\bar{d}$ or $\bar{u}d$.

Another possibility is to combine three quarks in the totally antisymmetric combination $\varepsilon_{ijk}q_1^iq_2^jq_3^k$, where $i, j, k \in \{r, g, b\}$. Such states consisting of an odd number of quarks are

called *baryons*. The lightest and most stable (and most familiar) of these are the *proton* and the *neutron*. Of course, more complex states could be formed, such as $qq\bar{q}\bar{q}$, or $qqqq\bar{q}$. Such *tetraquarks* and *pentaquarks* have higher energies and are unstable, but some of them have in fact been observed recently at the LHCb experiment at CERN.

Color singlets can also be formed from gluons. Such states are called *glueballs*. For example, two gluons can be combined by symmetrically summing over all colors (this is the singlet in the product $\mathbf{8} \times \mathbf{8}$ of two gluon color octet/adjoint representations). Such states are also mesons. Glueballs have not been directly confirmed experimentally. They mix with the $q\bar{q}$ mesons, and sometimes have the same quantum numbers. They hence contribute to the total number of mesons states, which could be measurable.

Status of QCD. Experimentally, all predicted low-lying mesons $q\bar{q}$ and baryons qqq are observed (some dozens of states in total). Some extra meson states have been observed, which have the quantum numbers expected for glueballs. No states have been observed that were not predicted by the theory. All properties of mesons and baryons are consistent with the quark picture of QCD.

On the theoretical side, the bound-state spectrum and all low-energy dynamics of QCD (like hadronization of quarks) is basically impossible to compute analytically. The reason is that the QCD coupling constant $\alpha_3 := g_3^2/4\pi$ is large (>1) at low energies, and hence no perturbative expansion in terms of Feynman diagrams is possible. Only at high energies >1 GeV (as in high-energy collisions) does QCD become weakly-coupled ($\alpha_3 \ll 1$), and perturbation theory gives reasonable answers. For the bound-state spectrum and low-energy dynamics, one has to resort to numerical lattice simulations (lattice QCD). Figure 3 shows the agreement between the experimentally measured hadron spectrum and various lattice QCD computations.

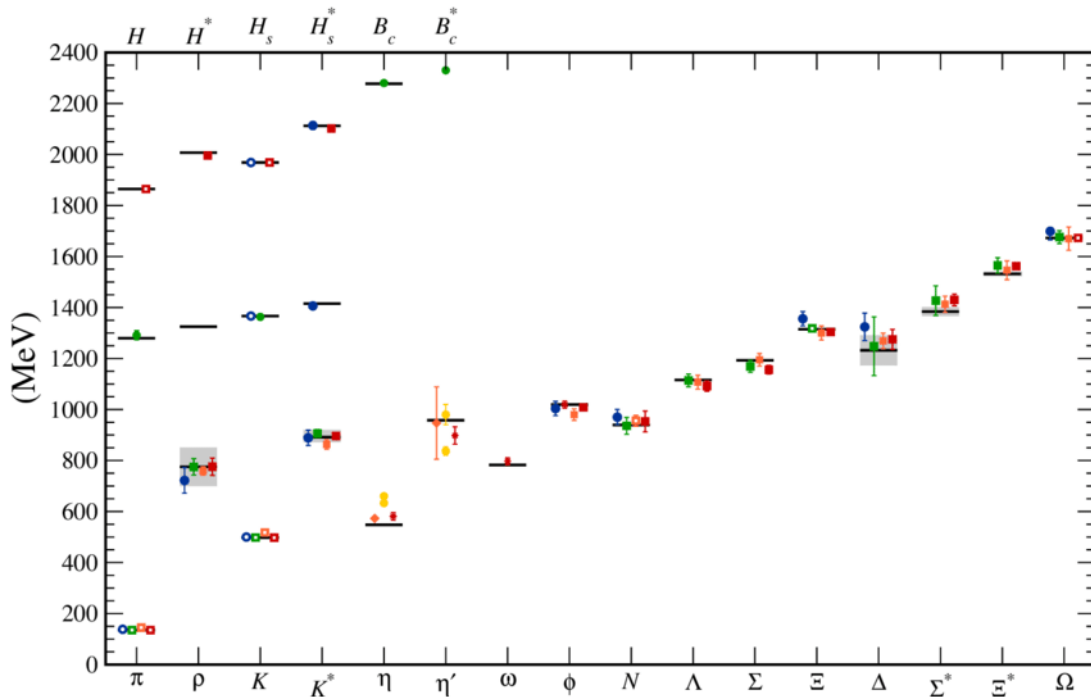


Figure 3: Hadron Spectrum from various lattice QCD computations (colored symbols) versus experimental values (black lines). Figure from [3].

11.4 Renormalization

General Idea. It was mentioned at several points that the coupling “constants” α_i , $i \in \{1, 2, 3\}$ of the electromagnetic, weak, and strong interactions are not really constant, but rather *depend on the energy scale of the interaction process*. We want to understand this point a bit better. The basic mechanism is particle pair creation. Picture two electrons interacting via photon exchange:

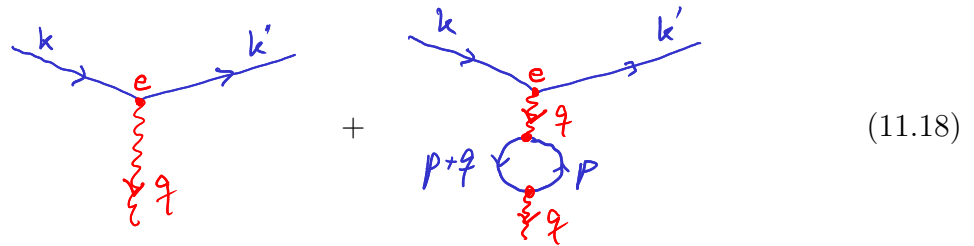


The higher the energy concentration, the higher is the probability that e^+e^- pairs will form. Any number of such pairs could form, and they will have an influence on the photon exchange interaction between the two “probe” electrons. In terms of Feynman diagrams, we have to include the following diagram:



This diagram is subleading in the coupling constant, but due to the propagators of the virtual electrons in the loop, its contribution depends on the energy scale: When the energy of the virtual photon is large, it will give a larger contribution than when the energy is small.

The Amplitude. Let us first consider the simpler process of photon emission. The leading and first subleading diagrams for this process are



We will not go into the details of the computation, but summing over all possible intermediate states (momenta and spins), the two terms work out to

$$e\bar{u}(k')\gamma^\mu u(k)\epsilon_\mu(1 + I(q^2)). \quad (11.19)$$

Here, $u(k)$ and $\bar{u}(k')$ are the spinors of the incoming and outgoing electrons, and ϵ_μ is the polarization vector of the photon. $I(q^2)$ is the loop correction, and is an integral that depends on the four-momentum q of the emitted photon,

$$I(q^2) = \frac{\alpha}{3\pi} \int_{m^2}^{\infty} \frac{d\rho^2}{\rho^2} - \frac{2\alpha}{\pi} \int_0^1 dx x(1-x) \log\left(1 - \frac{q^2 x(1-x)}{m^2}\right), \quad (11.20)$$

where m is the electron mass. The integrals stem from the integration over all (unconstrained) possible momenta running inside the virtual electron loop.

Notably, the first integral gives a logarithmic divergence $\sim \log(\infty)$ from the upper integration limit! This is one of the typical infinities that arise in quantum field theory. For the moment, we regulate this infinity by replacing the ∞ in the upper integration limit by some large but finite value Λ . The first term then gives

$$\frac{\alpha}{3\pi} \log\left(\frac{\Lambda^2}{m^2}\right). \quad (11.21)$$

The second integral is finite and can be computed analytically. We are most interested in the short-distance / large momentum-transfer behavior, where q^μ is large and space-like. Hence we make the approximation $-q^2/m^2 \gg 1$. In this case, we can approximate

$$\log\left(1 - \frac{q^2 x(1-x)}{m^2}\right) \approx \log\left(\frac{-q^2}{m^2}\right). \quad (11.22)$$

Using $\int_0^1 dx x(1-x) = 1/6$, the contribution $I(q^2)$ becomes

$$I(q^2) = \frac{\alpha}{3\pi} \log\left(\frac{\Lambda^2}{m^2}\right) - \frac{\alpha}{3\pi} \log\left(\frac{-q^2}{m^2}\right) = \frac{\alpha}{3\pi} \log\left(\frac{\Lambda^2}{-q^2}\right). \quad (11.23)$$

The dependence on m^2 has dropped out, as expected for a high-energy limit. Attaching the other electron current $e\bar{u}(p')\gamma_\mu u(p)$ that absorbs the photon, we find for the $2 \rightarrow 2$ amplitude

$$\begin{aligned} \mathcal{M} &= \text{[tree-level diagram]} + \text{[one-loop diagram]} \\ &\approx e^2 \left(1 - \frac{\alpha}{3\pi} \log\left(\frac{\Lambda^2}{-q^2}\right)\right) \bar{u}(k')\gamma_\mu u(k) \bar{u}(p')\gamma_\mu u(p). \end{aligned} \quad (11.24)$$

So far we have summed the terms with zero and one fermion loop. We can in fact add all terms where the virtual photon is dressed with any number of loops in a chain. Since all intermediate photons will have the same momentum q , the loops factor, giving a geometric series $1 - \epsilon + \epsilon^2 - \epsilon^3 + \dots$ that sums to $1/(1 + \epsilon)$. Therefore, the full coefficient should be

$$\frac{e^2}{1 + \frac{\alpha}{3\pi} \log\left(\frac{\Lambda^2}{-q^2}\right)}. \quad (11.25)$$

Running Coupling. After this computation, now comes the main physics point. So far, we have been assuming that $\alpha = e^2/4\pi \approx 1/137$. But the value of α (or e) that we actually measure in experiments necessarily *includes all terms with any number of fermion loops*, so we should get the measured value $1/137$ only after adding all these loop contributions. And the answer of this computation depends on the transferred momentum q^2 ! The measurement of α has to be performed at a certain value of q^2 (that depends on the experiment). Let us call the value of q^2 in the measurement $-\mu^2$, so that $\alpha = 1/137$ at $q^2 = -\mu^2$. Now call the value of e that appears in the Lagrangian the *bare coupling* e_0 , and

$\alpha_0 = e_0^2/4\pi$. Then the physical amplitude is given by the sum of terms with any number of loops,

$$= \text{tree-level } e_0 + \text{one-loop } e_0 + \dots \quad (11.26)$$

and the value of α measured at $q^2 = -\mu^2$ is

$$\alpha(\mu^2) = \frac{\alpha_0}{1 + \frac{\alpha_0}{3\pi} \log\left(\frac{\Lambda^2}{\mu^2}\right)}. \quad (11.27)$$

At any other value of q^2 , we find

$$\alpha(q^2) = \frac{\alpha_0}{1 + \frac{\alpha_0}{3\pi} \log\left(\frac{\Lambda^2}{-q^2}\right)} = \frac{\alpha_0}{1 + \frac{\alpha_0}{3\pi} \left(\log\left(\frac{\Lambda^2}{\mu^2}\right) + \log\left(\frac{\mu^2}{-q^2}\right)\right)}. \quad (11.28)$$

Now we can use (11.27) to simplify the denominator:

$$\alpha(q^2) = \frac{\alpha_0}{\frac{\alpha_0}{\alpha(\mu^2)} + \frac{\alpha_0}{3\pi} \log\left(\frac{\mu^2}{-q^2}\right)} = \frac{\alpha(\mu^2)}{1 + \frac{\alpha(\mu^2)}{3\pi} \log\left(\frac{\mu^2}{-q^2}\right)}. \quad (11.29)$$

Remarkably, the dependence on α_0 and on the momentum cutoff Λ has disappeared! Only finite, physical quantities enter this equation: $\alpha(\mu^2)$ is the measured value of the coupling at some particular momentum $-\mu^2$, and q^2 is another physical momentum. The coupling $\alpha(q^2)$ is called a *running coupling constant*.

Further Fermions. In our computation, we have corrected the photon propagator with electron loops. Similar loops are contributed by muons, tau leptons, or quarks. The correction terms have to be summed over all particles that can run in the loops. If all fermions satisfy $|q^2| \gg m^2$, the log term in the denominator should be multiplied by a factor

$$n_\ell + 3\left(\frac{4}{9}\right)n_u + 3\left(\frac{1}{9}\right)n_d, \quad (11.30)$$

where n_ℓ is the number of charged leptons, n_u is the number of up-type quarks with charge $2/3e$, n_d is the number of down-type quarks with charge $1/3e$, and all quarks come with a factor of three for color. Each term comes with a factor (charge)² since it couples to a photon at each side of the loop.

If $-q^2$ is not large enough, some heavy fermions might give a reduced contribution because the fermion mass in the propagators will suppress the momentum integral. Hence a full computation will give threshold effects as $-q^2$ increases. If quarks and leptons only occur in complete families, then N families contribute

$$N\left(1 + \frac{4}{3} + \frac{1}{3}\right) = \frac{8N}{3}. \quad (11.31)$$

Loops with $W^\pm$ have to be included as well if $|q^2| \geq m_W^2$. If a particle is much heavier than the exchanged momentum, $m^2 \gg |q^2|$, then its effect on the running coupling drops off as $\sim 1/m^2$; this is called “decoupling” of the particle.

Qualitative Picture. The sign between the two terms in the denominator of

$$\alpha(q^2) = \frac{\alpha(\mu^2)}{1 + \frac{\alpha(\mu^2)}{3\pi} \log\left(\frac{\mu^2}{-q^2}\right)} \quad (11.32)$$

is very important: As $|q^2|$ increases, the logarithm decreases, and hence $\alpha(q^2)$ increases. Conversely, $\alpha(q^2)$ decreases as $|q^2|$ decreases. Physically, this is a *screening effect*: Imagine a negative charge at the origin. Close to this charge, fermion anti-fermion pairs will spontaneously form. These pairs can be thought of as the loops that correct the photon propagator. For each pair, the positively charged particle will be attracted by the negative charge at the origin, whereas the negatively charged particle will be repelled:



This situation is similar to the classical situation of an electron inside a dielectric medium: Through the polarized “molecules” (particle anti-particle pairs), the electric charge of the electron gets partially screened (dielectric screening). For a probe particle at some distance, the negative charge at the origin will be shielded by some net positive charge. As the probe gets closer to the origin (smaller distance $\leftrightarrow$ larger momentum), it sees less screening charge, and therefore a larger net negative charge, which lets α increase.

The asymptotic value of α at low energies is the familiar $\alpha = 1/137$. At the scale of the weak bosons $W^\pm$ and Z , about 90 GeV, the effective value is $\alpha(m_Z^2) \approx 1/127$. So the effect is not negligible.

11.5 Asymptotic Freedom.

Quark and Gluon Loops. A very similar effect occurs for QCD, but there is a new feature with remarkable consequences. The gluon emission process is corrected by two types of diagrams:

The quark loop in the second diagram gives the same contribution for each quark flavor, since the quark-gluon coupling is flavor-independent. This diagram gives the same contribution as the loop correction in the QED case, only, due to color factors, the coefficient $\alpha(\mu^2)/3\pi$ changes to $\alpha_3(\mu^2)/6\pi$ for each flavor.

The third diagram provides the new feature. It has the same space-time structure as the second diagram, but gives an important numerical factor. Since eight gluons contribute (and only six quarks), and the color charge of a gluon is larger than that of a quark, the

third diagram contributes more than the second. More importantly, its contribution has *the opposite sign*. The reason is that gluons carry color charge. Qualitatively, this is easy to understand: Consider for example a “red” quark q_b at the origin. The quark can emit a gluon, for example $q_r \rightarrow q_b + G_{\bar{b}r}$. Then a probe would not see the blue color charge concentrated at the origin, but somewhat moved out into the surrounding gluon cloud. Hence there is an anti-screening effect, because radiating gluons dilute the charge. This did not happen with photons in QED, since photons have zero (electric) charge. The higher the energy $|q^2|$, the more gluon radiation, and the more charge dilution occurs. In the high-energy limit, the charge is completely diluted, and the point-like quark is effectively colorless and therefore effectively *free*!

Asymptotic Freedom. This property is called *asymptotic freedom*, and it was first observed experimentally in the scattering of electrons from quarks in hadrons, before the strong force was explained by the QCD theory. But when it was later discovered that asymptotic freedom comes out of the QCD theory, this was a major factor that led to the acceptance of QCD as the correct theory of the strong force, and in particular of the idea that quarks are real particles and not mere theoretical constructs.

The result of combining the quark and gluon loop diagrams is the replacement (compared to the QED case)

$$\frac{\alpha(\mu^2)}{3\pi} \rightarrow -\frac{\alpha_3(\mu^2)}{4\pi} \left(11 - \frac{2}{3}n_f\right), \quad (11.35)$$

where n_f is the total number of quark flavors. The renormalization of the QCD coupling constant is therefore given by

$$\boxed{\alpha_3(q^2) = \frac{\alpha_3(\mu^2)}{1 - \frac{\alpha_3(\mu^2)}{12\pi} (33 - 2n_f) \log\left(\frac{\mu^2}{-q^2}\right)}}. \quad (11.36)$$

There is a running of the coupling constant, as for QED. The formula is consistent with decay and collider experiments at various energies. Unlike in QED, as long as $33 - 2n_f$ is positive (recall that presumably $n_f = 6$), when $|q^2|$ *increases*, the denominator increases, and $\alpha_3(q^2)$ *decreases*. This is the behavior of asymptotic freedom: At large energies, the coupling becomes small.

Confinement. At the other end of the spectrum, for small $|q^2|$, the two terms in the denominator have opposite signs, so α_3 becomes large. At some value $-q^2 = \Lambda_{\text{QCD}}^2$, the denominator can vanish, so the QCD force apparently becomes infinitely strong. This is of course unphysical, and shows that our approximations are not valid in this regime. Still, the QCD coupling will become very large at this energy scale. Solving for Λ_{QCD} ,

$$\begin{aligned} \log \frac{\mu^2}{\Lambda_{\text{QCD}}^2} &= \frac{12\pi}{\alpha_3(\mu^2)(33 - 2n_f)} \\ \Rightarrow \Lambda_{\text{QCD}} &= \mu \exp\left(-\frac{6\pi}{\alpha_3(\mu^2)(33 - 2n_f)}\right). \end{aligned} \quad (11.37)$$

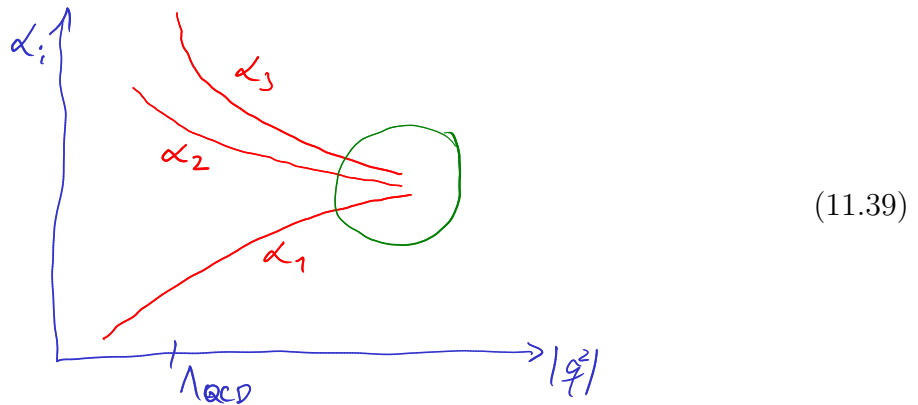
Suppose at some large μ^2 , for example $\mu = 10 \text{ GeV}$, that $\alpha_3 \approx 0.2$ (this is roughly correct) and $n_f = 5$ (the top quark is too heavy to contribute at this energy). Then we can infer

from the above formula:

$$\Lambda_{\text{QCD}} \approx \exp\left(-\frac{19}{23/5}\right) \cdot 10 \text{ GeV} \approx 180 \text{ MeV}. \quad (11.38)$$

One can therefore expect that QCD becomes very strong at a scale of a few times Λ_{QCD} . In this strongly-coupled regime, all colored particles will form color-neutral bound states. Hence all quarks and gluons have to become bound into colorless hadrons at this scale. This is called *color confinement*, and is exactly what we observe (protons and neutrons have energies of $\approx 1 \text{ GeV}$).

Comparison of Couplings. We can sketch the values of the three couplings as functions of the momentum transfer $|q^2|$:



We saw that α increases with $|q^2|$, and α_3 decreases. The result for α_2 is similar to α_3 : The gauge boson loops dominate because the electroweak charge of the W bosons is larger than that of the fermions. The result is the same as for α_3 (11.36), with the number 33 replaced by a slightly smaller number.

As becomes clear from the preceding discussion, all forces associated with non-Abelian gauge symmetries will be asymptotically free at high energies, but are strongly coupled at low energies. Forces with larger gauge symmetry groups will be more strongly coupled than forces with smaller gauge groups. Forces arising from an Abelian $U(1)$ gauge symmetry are weak at low energies, and therefore not confining. This is the reason that only the electromagnetic force is weak and long-ranged at low energies, whereas the weak and strong forces are very short-ranged due to their confining nature. The behavior sketched above suggests that the three forces have similar strengths at large $|q^2|$ than at the more familiar, lower values of $|q^2|$.

11.6 Quark Mixing Angles

Eigenstates. There is one ingredient to the Standard Model Lagrangian that we did not incorporate yet. Notice that, with the Lagrangian that we have written so far, the heavy down-type quarks (strange and bottom) are stable! They do not couple to any lighter quarks, and hence have no decay channels. This is not in accordance with observations. The reason for this situation is an implicit assumption that we made: We assumed that the quark flavors that form the components of the left-handed $SU(2)_W$ doublets are identical to the quark states of definite mass. In other words, we assumed that the eigenstates of the electroweak Hamiltonian are at the same time eigenstates of the (massive) kinetic Hamiltonian. We have no reason to make that assumption, and in fact it is wrong.

Cabibbo Angle. First suppose that there are only two families of quarks. Denote by u , d , c , and s the mass eigenstates, that is the free states with definite energies. The charged current we wrote earlier that couples to the $W^\pm$ bosons can then be written as

$$J_{\text{ch}}^\mu = \begin{pmatrix} \bar{u} & \bar{c} \end{pmatrix} \gamma^\mu P_L \begin{pmatrix} d \\ s \end{pmatrix} = \bar{u} \gamma^\mu P_L d + \bar{c} \gamma^\mu P_L s, \quad (11.40)$$

where we have used row and column vectors in flavor space, and we have used the mass eigenstates. But in fact the weak eigenstates (eigenstates of hypercharge Y and isospin T_3) could be different from the mass eigenstates (in fact they are). So in the current, we should replace d and s by d' and s' , where q' are the weak eigenstates. The new eigenstates can be written as linear combinations of the old,

$$\begin{pmatrix} d' \\ s' \end{pmatrix}_L = V \begin{pmatrix} d \\ s \end{pmatrix}_L, \quad (11.41)$$

where V is a 2×2 unitary matrix. Every such matrix can be written in terms of three real angles θ , α , and γ as

$$V = \begin{pmatrix} \cos \theta e^{i\alpha} & \sin \theta e^{i\gamma} \\ -\sin \theta e^{-i\gamma} & \cos \theta e^{-i\alpha} \end{pmatrix}. \quad (11.42)$$

The angles γ and α can be absorbed into the definitions of the quark states:

$$d' \rightarrow e^{-i\alpha} d', \quad s' \rightarrow e^{i\gamma} s', \quad s \rightarrow e^{-i(\gamma-\alpha)} s. \quad (11.43)$$

Then the transformation is

$$V = \begin{pmatrix} \cos \theta & \sin \theta \\ -\sin \theta & \cos \theta \end{pmatrix}, \quad (11.44)$$

and the charged current that we should have used from the beginning becomes

$$\begin{aligned} J_{\text{ch}}^\mu &= \begin{pmatrix} \bar{u} & \bar{c} \end{pmatrix} \gamma^\mu P_L \begin{pmatrix} d' \\ s' \end{pmatrix} = \begin{pmatrix} \bar{u} & \bar{c} \end{pmatrix} \gamma^\mu P_L V \begin{pmatrix} d \\ s \end{pmatrix} \\ &= \bar{u} \gamma^\mu P_L d \cos \theta + \bar{u} \gamma^\mu P_L s \sin \theta - \bar{c} \gamma^\mu P_L d \sin \theta + \bar{c} \gamma^\mu P_L s \cos \theta. \end{aligned} \quad (11.45)$$

There are two new terms, both multiplied by $\sin \theta$, and the old terms are reduced by a factor $\cos \theta$. The angle θ is called the *Cabibbo angle*, its experimental value is $\theta \approx 13^\circ$. If θ was zero, the s quark would be stable. Now, with non-zero θ , it can decay to a u quark by emitting a W^- , via the second term in the current. The electroweak vertices are now

$$(11.46)$$

In this computation, we only rotated the down-type quarks. This is in fact the most general rotation: Had we also rotated the up-type quarks, the current would have been

$$J_{\text{ch}}^\mu = \begin{pmatrix} \bar{u} & \bar{c} \end{pmatrix} V_{\text{up}}^\dagger \gamma^\mu P_L V_{\text{down}} \begin{pmatrix} d \\ s \end{pmatrix}. \quad (11.47)$$

But the product of the two rotations $V_{\text{up}}^\dagger V_{\text{down}}$ is again a rotation, so we can replace it by a single rotation V that only acts on the down-type quarks.

Next, let us check what happens to the neutral current that couples to the Z boson. It has the form

$$J_{\text{neu}}^\mu = \sum_{f=u,d,c,s} \left(\bar{f}_L \gamma^\mu [T_3 - Q \sin^2 \theta_W] f_L + \bar{f}_R \gamma^\mu [0 - Q \sin^2 \theta_W] f_R \right). \quad (11.48)$$

The down-type quarks that we rotate have identical charges T_3 and Q . In other words, the square brackets are proportional to the identity matrix in $(d \ s)$ space. Therefore the rotation V does not change the neutral current at all: The neutral current is diagonal both in mass eigenstates and in weak eigenstates.

This observation has an important consequence: There is no interaction vertex of the form $\bar{s}dX$ (where X could be any particle). Therefore, decays involving $s \rightarrow d$, called *flavor-changing neutral currents* are much less common than $s \rightarrow u$ (charged) decays: An s can only turn to a d through some non-trivial intermediate state. Flavor-changing neutral currents are interesting since they are possible probes of new interactions.

CKM Matrix. We can generalize these results to all three quark families. Including all quarks, the charged current becomes

$$J_{\text{ch}}^\mu = \begin{pmatrix} \bar{u} & \bar{c} & \bar{t} \end{pmatrix} \gamma^\mu P_L V \begin{pmatrix} d \\ s \\ b \end{pmatrix}, \quad (11.49)$$

where V is now a 3×3 unitary matrix. Such a matrix in general has 9 real parameters. We can re-define the phases of five quarks (an overall phase of all quarks does not change anything), so 4 parameters remain. An orthogonal matrix describing real 3×3 rotations has only 3 real parameters, so one parameter of V could still be a complex phase.

The matrix V is called the *Cabibbo–Kobayashi–Maskawa (CKM) matrix*. Its entries, or rather their magnitudes, all have been measured, they are

$$V_{\text{mag}} = \begin{pmatrix} V_{ud} & V_{us} & V_{ub} \\ V_{cd} & V_{cs} & V_{cb} \\ V_{td} & V_{ts} & V_{tb} \end{pmatrix} = \begin{pmatrix} 0.974 & 0.225 & 0.0035 \\ 0.225 & 0.973 & 0.04 \\ 0.009 & 0.04 & 0.999 \end{pmatrix}. \quad (11.50)$$

One can see that transition from one family to the next (along the mass scale) are small, and transitions between the lightest and the heaviest family are very small.

The real matrix V_{mag} accounts for three of the four parameters of the full CKM matrix V . The fourth parameter is a complex phase. It can enter in various ways, as it can be shifted around by re-defining the quark phases. The common way to include it is to split the matrix V into three subsequent Euler rotations,

$$V = V_{23}(\theta_{23}) \cdot V_{13}(\theta_{13}) \cdot V_{12}(\theta_{12}) \quad (11.51)$$

and to modify the rotation V_{13} by a complex phase $e^{i\delta}$:

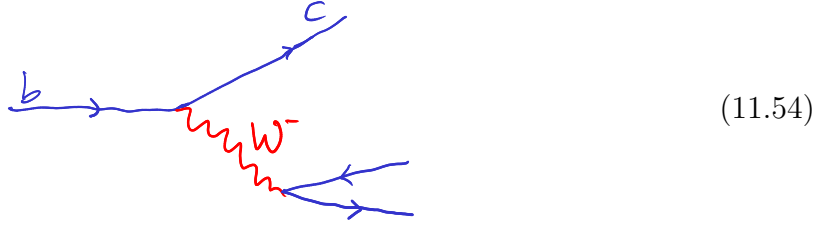
$$V_{13} \rightarrow \begin{pmatrix} \cos \theta_{13} & 0 & \sin \theta_{13} e^{-i\delta} \\ 0 & 1 & 0 \\ -\sin \theta_{13} e^{i\delta} & 0 & \cos \theta_{13} \end{pmatrix}, \quad (11.52)$$

where θ_{13} is a real angle, and δ is the phase. δ is measured to be $\delta = 1.20(8)$ radians, which is large. This non-zero phase means that terms $\sim W_\mu J_{\text{ch}}^\mu$ can be complex, which means that the theory will not be invariant under time reversal! By CPT symmetry, time reversal is equivalent to a CP transformation (charge conjugation and parity), hence the non-zero phase δ is important for understanding CP violation, as we shall see below.

Due to the non-trivial CKM matrix, also b quarks can decay: One of the terms in the current is

$$\bar{c}\gamma^\mu P_L V_{cb} b, \quad (11.53)$$

hence the bottom quark can decay into a charm quark and a W boson (which in turn decays into quark or lepton pairs):



Since $V_{cb} \approx 0.04$ is small, the bottom quark is relatively long-lived, considering its mass. For its width, one finds

$$\frac{\Gamma_b}{\Gamma_\tau} \approx 0.4, \quad (11.55)$$

so the b quark lives about 2.5 times longer than the tau lepton. This lifetime is long enough to observe the bottom quark directly, by the separation of production and decay vertices in the detector.

Lepton Mixing. We could similarly rotate the weak lepton eigenstates relatively to the lepton mass eigenstates. However, assuming that all neutrinos are massless, such a rotation would make no difference. In the case of quarks, we chose to rotate the three down-type quarks, but we could as well have rotated the up-type quarks instead. For leptons, this would mean to rotate the neutrinos among each other. But we can anyhow not tell the neutrinos apart by their mass. In other words, in the three-dimensional neutrino state space, we can pick the neutrino mass eigenstates arbitrarily. By definition, we just use the weak eigenstates as the mass eigenstates.

11.7 CP Violation

CP Symmetry. Many physical phenomena are invariant under parity transformations P that invert all spatial coordinates, $\mathbf{x} \rightarrow -\mathbf{x}$. By now, we understand well that the Standard Model is not parity invariant: The parity transformation exchanges left-handed and right-handed spinors. The weak interactions couple left-handed and right-handed fermions differently, so it manifestly breaks parity symmetry.

However, we can combine parity with charge conjugation C , which turns all particles into their anti-particles. C by itself is not a symmetry of the Standard Model, as the weak interactions couple left-handed fermions differently than left-handed anti-fermions. However, the combined CP transformation turns left-handed fermions into right-handed anti-fermions, and these do appear symmetrically in the weak interaction Lagrangian, so the weak interactions are symmetric under CP transformations. In the 1950's, it was proposed that CP -symmetry could be a true symmetry of fundamental physics.

CP Violation. In 1964, Cronin and Fitch observed that CP symmetry is broken (violated) in the process of neutral kaon decay (neutral kaons K^0 are mesons consisting of a down and a strange quark). This was a completely unexpected surprise, and opened the door to questions that are still central to particle physics and cosmology today. The discovery was awarded with the 1980 Nobel Prize.

To understand how CP violation arises from the Standard Model, we recall that the combination of charge conjugation, parity, and time reversal (CPT) is a true symmetry for all quantum field theories. Hence a violation of CP is equivalent to a violation of time reversal T . Quantum mechanically, time reversal transforms the Hamiltonian as

$$H \rightarrow THT^{-1}. \quad (11.56)$$

Hence time reversal symmetry requires

$$H \stackrel{!}{=} THT^{-1}. \quad (11.57)$$

In quantum theory, the time reversal operator is an *anti-unitary* operator: $T = UK$, where U is unitary, and K is the complex conjugation operator. This can be understood by considering the canonical commutator $[x, p] = i\hbar$. Time reversal does not change coordinates x , that is $TxT^{-1} = x$, but it reverses momenta p , that is $TpT^{-1} = -p$. Therefore, T must be anti-unitary, that is $TiT^{-1} = -i$.

As a result, the Hamiltonian H can only be time-reversal invariant when it is real. Conversely, a complex Hamiltonian cannot be time-reversal invariant, $THT^{-1} \neq H$, and therefore breaks CP symmetry. We saw above that the CKM matrix V that enters the charged current via

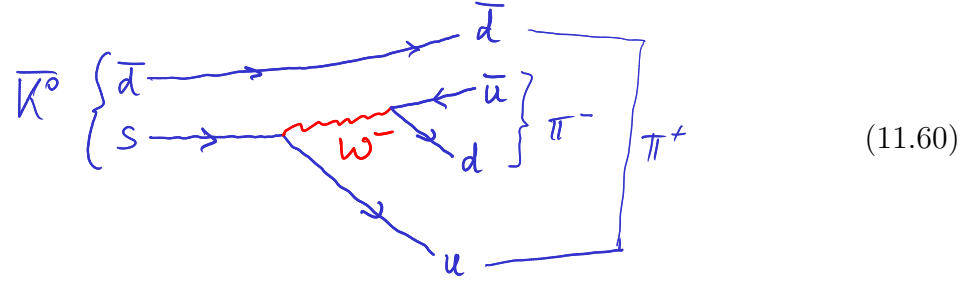
$$\begin{pmatrix} \bar{u} & \bar{c} & \bar{t} \end{pmatrix} \gamma^\mu P_L V \begin{pmatrix} d \\ s \\ b \end{pmatrix} \quad (11.58)$$

is indeed complex, due to the non-zero phase δ . Hence the Standard Model indeed incorporates CP violation. Specifically, CP symmetry is violated by interactions between charged quark currents and $W^\pm$ bosons.

Neutral Kaon Decay. CP violation has been observed in decays of kaons and of b quarks. There are two neutral kaon states, $K^0 = (d\bar{s})$ and its anti-particle $\bar{K}^0 = (\bar{d}s)$. It is more useful to consider the CP eigenstates

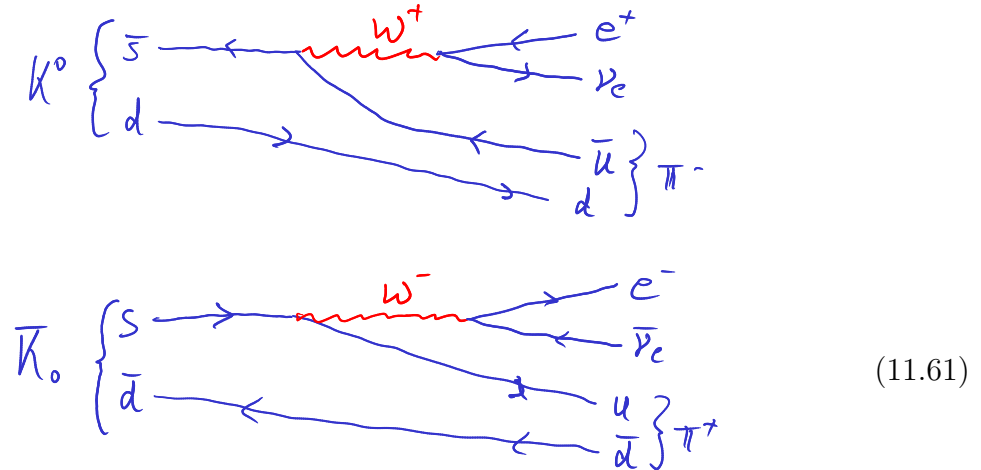
$$K_S = K^0 + \bar{K}^0, \quad K_L = K^0 - \bar{K}^0. \quad (11.59)$$

The state K_S is even under CP transformations, while K_L is odd. The most common decay mode for neutral kaons is $K \rightarrow \pi\pi$ to two pions. For example:



The $\pi\pi$ final state is an even CP eigenstate. It could be $\pi^0\pi^0$, or $\pi^+\pi^-$, both are even under CP transformations. Assuming that CP-symmetry is respected, the K_S state can decay into $\pi\pi$, but K_L cannot. Since all other decay modes have much smaller widths, the K_L state has a much longer lifetime than the K_S state (the subscript L stands for “long-lived”). This is indeed observed, showing that CP is at least approximately a symmetry.

To study a possible CP violation, one can compare two decay channels that are CP conjugates of each other. The most common are the semi-leptonic decays



If CP-symmetry is preserved, the two decays should occur with identical probabilities. Measuring the decay rates for these two channels gives

$$\frac{\Gamma(K_L \rightarrow \pi^- e^+ \nu_e) - \Gamma(K_L \rightarrow \pi^+ e^- \bar{\nu}_e)}{\Gamma(K_L \rightarrow \pi^- e^+ \nu_e) + \Gamma(K_L \rightarrow \pi^+ e^- \bar{\nu}_e)} = 0.00333(14). \quad (11.62)$$

Thus one observes indeed a small CP violation. The violation is small, but clearly non-zero.

11.8 Parameters of the Standard Model

With the quark mixing matrix, our formulation of the Standard Model is truly complete. To conclude its description, we list its free parameters, whose values have to be measured and used as an input to the model, in Table 1. There are 19 such parameters in total. Instead of the electroweak couplings g_1 and g_2 , one could also use the electron charge e and the electroweak mixing angle θ_W , via the relations

$$g_1 = \frac{e}{\cos \theta_W}, \quad g_2 = \frac{e}{\sin \theta_W}. \quad (11.63)$$

Symbol	Description	Value
m_e, m_μ, m_τ	Electron, muon, tau masses	511 keV, 105.7 MeV, 1.777 GeV
m_u, m_c, m_t	Up, charm, top quark masses	2.16 MeV, 1.27 GeV, 173 GeV
m_d, m_s, m_b	Down, strange, bottom quark masses	4.67 MeV, 93 MeV, 4.18 GeV
θ_{12}	CKM 12-mixing (Cabibbo) angle	13.1°
θ_{23}	CKM 23-mixing angle	2.4°
θ_{13}	CKM 13-mixing angle	0.2°
δ	CKM phase	0.995
g_1 or g'	U(1) gauge coupling	0.357 at $\mu = m_Z$
g_2 or g	SU(2) gauge coupling	0.652 at $\mu = m_Z$
g_3 or g_s	SU(3) gauge coupling	1.221 at $\mu = m_Z$
v	Higgs vacuum expectation value	246 GeV
m_h	Higgs mass	125.09(24) GeV
θ_{QCD}	QCD vacuum angle	0

Table 1: The 19 unfixed parameters of the Standard Model, whose values have to be determined experimentally, and used as an input to the theory.

Instead of the Higgs field vacuum expectation value v and Higgs boson mass m_h , one could also use the two parameters λ and μ in the Higgs potential:

$$m_h^2 = -2\mu^2, \quad v^2 = \frac{-\mu^2}{\lambda}. \quad (11.64)$$

The masses of the W and Z bosons are given by

$$m_W = \frac{v g_2}{2}, \quad m_Z = \frac{m_W}{\cos \theta_W}. \quad (11.65)$$

The only parameter we have not discussed yet is the QCD vacuum angle. It is the parameter of a hypothetical term in the QCD Lagrangian that would be allowed by gauge invariance, but is apparently absent.

12 Beyond the Standard Model

To round off this course, let us briefly touch on some subjects that go beyond the Standard Model. The Standard Model is a very successful theory, whose predictions agree with essentially all experimental particle physics data to high precision. Nevertheless, there are some open questions and puzzles which indicate that the current formulation of the Standard Model, as a theory of particle physics (excluding gravity), might not be complete.

12.1 Some Directions

In the following, we will consider some questions or puzzles the Standard Model leaves open, and also some possible modifications or extensions of the Standard Model that could address these questions.

A More Fundamental Theory? Looking at the Standard Model, perhaps the most obvious question is: What fixes the various parameters of the model? Is there some underlying principle that explains the specific values we observe? Also, what fixes the particle content? Why are there three families of quarks and leptons, not more (or fewer)? Why is the gauge group $U(1) \times SU(2) \times SU(3)$? Stepping back even further, one could ask: Why does the universe have $3 + 1$ dimensions?

Some of these questions could be addressed by a more fundamental theory of particle physics, from which the Standard Model would arise as an effective low-energy description. The prevalent candidate for such a theory is *string theory*, which has almost no input parameters, and which, at least at a qualitative level, provides mechanisms to determine some properties and parameters of the Standard Model.

Existence of Matter. Some of the input to the Standard Model is constrained by consistency conditions, for example from cosmology. For instance, one fundamental question of cosmology is: Why is there more matter than anti-matter in the observed universe? At first sight, the Standard Model treats matter and anti-matter symmetrically. But in fact the theory does admit non-perturbative processes that break the matter anti-matter symmetry, and favor the creation of matter over anti-matter under certain conditions. Such processes can create an asymmetry in the matter anti-matter balance in the early universe, such that the matter density is slightly larger than the anti-matter density, by a relative factor of $1 + 10^{-10}$. As the universe expanded and cooled, most of the matter and anti-matter annihilated, leaving only the small imbalance as the remaining matter in today's universe.

The process of matter generation is called *baryogenesis/leptogenesis*. The point to make for the Standard Model is that the processes that create the matter anti-matter asymmetry crucially require that CP symmetry is broken. We have seen that the electroweak theory indeed breaks CP symmetry, through the complex phase in the CKM matrix. But to incorporate this complex phase, the matrix has to be at least three-dimensional. So within the current models of particle physics and cosmology, the existence of matter requires that there are at least three families of quarks and leptons.

Two very concrete puzzles raised by the Standard Model are the *strong CP problem* and the *hierarchy problem*.

Strong CP Problem. We have seen that the electroweak theory does not preserve CP symmetry (due to the complex CKM matrix), and that the broken CP symmetry is necessary to explain the matter content of the universe. But the theory of the strong force (QCD), as it stands, *preserves CP symmetry*. This is puzzling for the following reason: One could include a term

$$\theta_{\text{QCD}} G_{\mu\nu} \tilde{G}^{\mu\nu}, \quad \tilde{G}_{\mu\nu} := \epsilon_{\mu\nu\rho\sigma} G^{\rho\sigma} \quad (12.1)$$

in the QCD Lagrangian, where $\tilde{G}$ is the dual of the QCD field strength tensor. Such a term would not preserve CP symmetry. Since it is gauge-invariant and Lorentz-invariant, nothing prevents its presence in the Lagrangian. In fact it is the only term compatible with gauge and Lorentz symmetry that is *not* present in the Lagrangian. But for some unknown reason, nature chose to set $\theta_{\text{QCD}} = 0$. The current experimental bound is $\theta < 10^{-10}$ (from the vanishing of the neutron electric dipole moment). The parameter θ_{QCD} could have taken any generic value. Among all possibilities, only the single value $\theta_{\text{QCD}} = 0$

preserves CP symmetry, and this is the actual value θ_{QCD} takes. Such a situation is called a *fine-tuning problem*: θ_{QCD} is precisely tuned (to many digits) to a very specific (non-generic) value.

One long-standing proposal for the resolution of the strong CP problem is *Peccei-Quinn* theory (formulated in 1977), which explains the value $\theta_{\text{QCD}} = 0$ by introducing a new complex scalar field a , with a coupling

$$\sim a G_{\mu\nu} \tilde{G}^{\mu\nu} . \quad (12.2)$$

The resulting theory is invariant under a global U(1) symmetry, with the effective potential for a given by the above interaction term. However, Peccei and Quinn showed that the vacuum expectation value (VEV) of a is

$$a \sim \theta , \quad (12.3)$$

such that the CP-violating terms $G_{\mu\nu} \tilde{G}^{\mu\nu}$ cancel each other, naturally leading to the observed effective value $\theta_{\text{QCD}} = 0$. The non-trivial VEV of a spontaneously breaks the U(1) symmetry. The excitations of the field a around this vacuum lead to a new particle called the *axion*. This new particle is not exactly massless, but astrophysical constraints imply that its mass is very small,

$$m_a \leq 10^{-5} \text{ eV} . \quad (12.4)$$

In fact, the axion is one candidate for the *dark matter* in the universe, since it is very light and very weakly coupled.

Hierarchy Problem. The hierarchy problem is a problem of energy or mass scales: When one includes quantum corrections to the mass of the Higgs boson, their contributions will be large, and will raise the Higgs mass by many orders of magnitude, presumably all the way to the next physically meaningful energy scale. Something similar would happen to the W and Z masses, and also to the fermion masses. The only way to prevent this is an incredible fine-tuning that leads to cancellations between almost all quantum corrections and the bare mass (that stands in the Lagrangian).

The question can be rephrased to: Why is there such a huge gap between the energy scale of the electroweak theory ($\sim 100 \text{ GeV}$) and the next-higher physically relevant energy scale. Staying within the Standard Model, this next-higher scale would be gravity, whose energy scale is the Planck mass

$$M_{\text{Pl}} = \sqrt{\hbar c / G} \approx 1.2 \cdot 10^{19} \text{ GeV} , \quad (12.5)$$

which is gigantic. Including possible beyond-the-standard-model physics, the next-higher energy scale could also be the grand unification scale (see below), or the mass scale of (hypothetical) heavy neutrinos, which are still much higher than the electroweak scale. The huge gap between the energy scales is equivalent to the huge gap between the strengths of the electroweak forces and the gravitational force: Defining the gravitational coupling constant α_G in a similar way as the other coupling constants, one finds

$$\alpha_G = \frac{G m_e^2}{\hbar c} = \frac{m_e^2}{M_{\text{Pl}}^2} \approx 1.75 \cdot 10^{-45} , \quad (12.6)$$

which is 43 orders of magnitude smaller than $\alpha_2 \approx 1/30$.

Supersymmetry. There are several candidate solutions of the hierarchy problem. One is *supersymmetry*, which is a symmetry that relates bosons and fermions, and thereby associates a fermionic *superpartner* to each boson, and conversely a bosonic superpartner to each fermion. Due to the symmetry, the quantum corrections from fermion loops and boson loops have the same magnitude, but opposite sign, and therefore cancel each other. This would provide a mechanism to stabilize the tiny Higgs mass (compared to the gravity scale) without any fine-tuning.

Extra Dimensions. Another possible resolution would be extra spacetime dimensions that are compact (have small extent), which would allow gravity to be much stronger at microscopic scales, but still be weak at everyday scales (as observed). The mechanism is the following: The total gravitational flux through a closed surface C surrounding a mass m in d dimensions is

$$\int_C \mathbf{g} \cdot d\mathbf{A} = -S_{d-2} G_d m, \quad (12.7)$$

where S_{d-2} is the surface area of a $(d-2)$ -dimensional unit sphere, and G_d is Newton's constant in d dimensions, which is defined by this equation. Now imagine a $d = 4 + n$ dimensional spacetime, where n dimensions are compactified in an n -dimensional volume $V_n = L^n$. L is the size of the extra dimensions. At small distances $r < L$, the magnitude of the gravitational field is the total flux divided by the surface area $A_r = S_{d-2} r^{d-2}$ of a sphere of radius r ,

$$|\mathbf{g}| = \frac{S_{d-2} G_d m}{A_r} = \frac{G_d m}{r^{d-2}}, \quad (12.8)$$

which is just the usual gravitational field of a point mass. But at distances $r \gg L$, the flux has fewer dimensions to spread: The flux evenly distributes across a surface at uniform distance r in *four dimensions* that covers the entire volume V_n of the compactified dimensions. Such a surface has area $4\pi r^2 V_n = 4\pi r^2 L^n$, and therefore

$$|\mathbf{g}| = \frac{S_{n+2} G_{4+n} m}{4\pi r^2 L^n} = \frac{G_4 m}{r^2}. \quad (12.9)$$

We have equated the result with the familiar gravity law in four dimensions. The four-dimensional Newton constant G_4 is now effective, and related to the more fundamental Newton constant in d dimensions via

$$\frac{G_4}{G_d} = \frac{S_{n+2}}{4\pi L^n}. \quad (12.10)$$

To convert this to energies, note that G_4 has dimension $1/\text{mass}^2$, whereas G_{4+n} has dimension $1/\text{mass}^{2+n}$. In natural units, we therefore have

$$G_4 = \frac{1}{M_{\text{Pl},4}^2}, \quad G_{4+n} = \frac{1}{M_{\text{Pl},4+n}^{2+n}}, \quad (12.11)$$

where $M_{\text{Pl},4+n}$ is the fundamental Planck mass in $d = 4 + n$ dimensions, and $M_{\text{Pl},4}$ is the resulting effective Planck mass in four dimensions. We can therefore relate back the known four-dimensional $M_{\text{Pl},4}$ to a more fundamental value by extra dimensions of appropriate size L . If we want the fundamental Planck mass $M_{\text{Pl},4+n}$ to be close to the weak scale, e. g. $M_{\text{Pl},4+n} = 1000 \text{ GeV}$, we need (neglecting the numerical constants)

$$L^n = \frac{M_{\text{Pl},4}^2}{M_{\text{Pl},4+n}^{2+n}} \approx \frac{(10^{19} \text{ GeV})^2}{(10^3 \text{ GeV})^{2+n}} \approx 2^n \cdot 10^{32-19n} \text{ m}^n, \quad (12.12)$$

where we have used that $2 \text{ GeV} \approx 10^{16} \text{ m}^{-1}$ in natural units. For some interesting values of n , this gives

$$L \approx \begin{cases} 2 \cdot 10^{13} \text{ m} & n = 1, \\ 2 \text{ mm} & n = 2, \\ 10 \text{ nm} & n = 3, \\ 40 \text{ fm} & n = 6. \end{cases} \quad (12.13)$$

An extra dimension of 10^{13} m would be noticeable astronomically, but extra dimensions of sub-millimeter size are within experimental bounds. For comparison, the proton size is $\approx 1 \text{ fm}$.

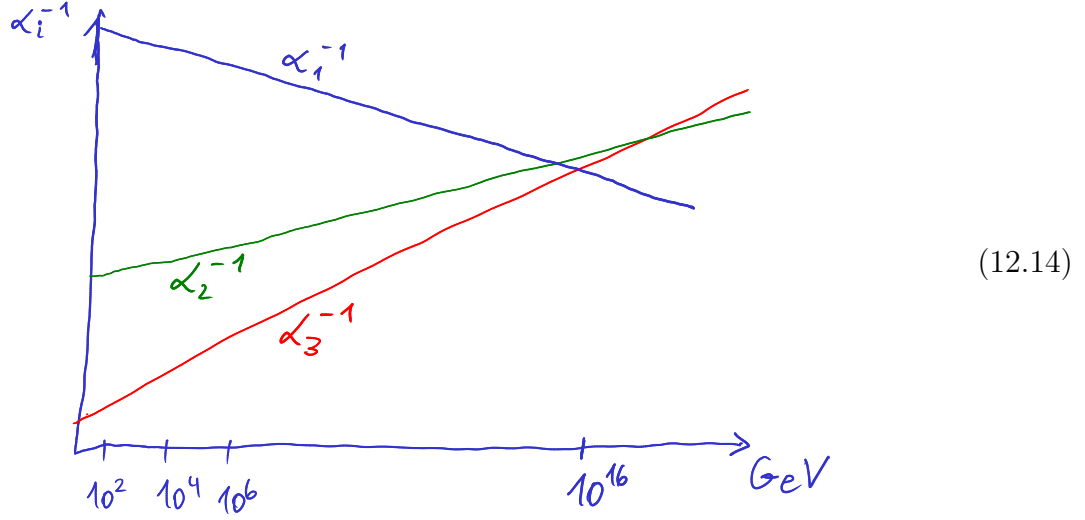
Both supersymmetry and extra dimensions are natural in string theory. The upper bound on the number of extra dimensions in superstring theory is $n = 6$.

Dark Matter. Based on astronomical observations of gravitational effects, for example in galaxies, about 25 % of the mass-energy density of the universe consists of *dark matter*, which does not emit any form of visible radiation. There seems to be no dark matter candidate in the Standard Model. Even though this is not completely certain, it is motivation enough to look for Standard Model extensions that include dark matter candidates. Dark matter might consist at least partly of small black holes. Also some form of quark matter is not entirely excluded. Other candidates are the axions of the Peccei–Quinn theory, or heavy neutrinos.

Grand Unification. One way to address some of the questions raised above are *grand unified theories* (GUTs). These assume that the weak and strong gauge groups $U(1)$, $SU(2)_W$, and $SU(3)$ are unified into a bigger gauge group. The smallest simple Lie group that contains the Standard Model gauge groups is $SU(5)$. Other candidates are $SO(10)$, or E_6 . Such a larger gauge group would mean that there exist further gauge bosons, which can convert quarks into leptons and vice versa. To make this consistent with our observations, this larger symmetry group must be spontaneously broken to the $U(1) \times SU(2)_W \times SU(3)$ that we observe.

We have seen that the three coupling constants “run” with the energy scale: The non-Abelian gauge couplings decrease at higher energies, while the $U(1)$ electromagnetic coupling constant increases. If the gauge groups are unified in a single $SU(5)$ group, then there is only a single coupling constant, and not several different coupling constants, so the couplings should approach the same value at some energy scale, called the *GUT scale*, or *grand unification energy* Λ_{GUT} . Indeed, perturbative computations show that the three coupling constants of the Standard Model nearly, but not quite, meet at the same point at

an energy scale of $\Lambda_{\text{GUT}} \approx 10^{16}$ GeV:



If nature is indeed described by a GUT, then this would be the energy scale at which the unified gauge symmetry is broken to the residual $U(1) \times SU(2)_W \times SU(3)$ that we observe. All GUTs modify the running of the couplings by additional particles that can run in the quantum loop corrections. In some GUTs, in particular the supersymmetric $SU(5)$ GUT, the couplings in fact meet at a single point much more accurately. This non-trivial result could be a coincidence, but can also be taken as an indication that our world is indeed supersymmetric at some high energy scale. This is especially encouraging since supersymmetry can also solve the hierarchy problem by stabilizing the Higgs mass against large quantum corrections.

12.2 Neutrino Masses

So far, we have treated neutrinos as massless particles. In the Standard Model, neutrinos are the only massless fermions. The reason is that there are no right-handed neutrinos.

Dirac Mass. Recall that for all other fermions, we introduced interaction terms

$$g \bar{f}_L \phi f_R \quad (12.15)$$

with the Higgs field ϕ . Such terms become mass terms $m \bar{f}_L f_R$ when the Higgs field acquires a vacuum expectation value. A mass term of this form is called a *Dirac mass*, since the fermion spinor f is a Dirac spinor. Such mass terms evidently require the existence of a right-handed component f_R .

Right-Handed Neutrinos. Right-handed neutrinos have never been observed, so we might assume that they do not exist. One should be careful though: If right-handed $SU(2)_W$ singlet neutrinos do exist, they are difficult to produce and to detect, since they would not interact with $W^\pm$ bosons (only $SU(2)_W$ doublets do), not with Z bosons or photons (both T_3 and Q would be zero), and not with gluons. If right-handed Dirac neutrinos do exist (and we just have not observed them yet), a Dirac mass term would be possible.

Majorana Mass. But even if no right-handed neutrinos exist, there is another possibility for a mass term, one that we did not discuss so far. Namely, neutrinos could be by *Majorana fermions*, which are fermions that are their own anti-particles. This is only possible for neutrinos, since all other fermions have non-zero charges and therefore cannot be their own anti-particles. Majorana fermions are described by *Majorana spinors*, which are four-component Dirac spinors ψ that satisfy an extra condition

$$\psi^c = \psi, \quad (12.16)$$

where ψ^c is the *charge-conjugate* spinor, defined by

$$\psi^c = C\psi^*, \quad (12.17)$$

where C is a 4×4 matrix that satisfies

$$C^\dagger C = 1 \quad \text{and} \quad C^\dagger \gamma^\mu C = -(\gamma^\mu)^*. \quad (12.18)$$

This definition ensures that ψ^c satisfies the Dirac equation provided ψ does, and transforms in the same way under Lorentz transformations as ψ does. Due to the Majorana condition $\psi^c = \psi$, Majorana spinors have only two degrees of freedom. They describe fermions that are their own anti-particles. Recall that all Dirac spinors can be split into two two-component (left-handed and right-handed) Weyl spinors,

$$\psi = \begin{pmatrix} \psi_L \\ \psi_R \end{pmatrix}. \quad (12.19)$$

For Majorana spinors, the Majorana condition implies that the right-handed component equals the charge-conjugate of the left-handed component, and vice versa,

$$\psi = \psi^c = \begin{pmatrix} \psi_L \\ \psi_R \end{pmatrix} = \begin{pmatrix} \psi_L \\ \psi_L^c \end{pmatrix} = \begin{pmatrix} \psi_R^c \\ \psi_R \end{pmatrix} \quad (12.20)$$

A Majorana spinor can therefore describe a left-handed neutrino ν_L and its right-handed anti-neutrino partner $\nu_R = \nu_L^c$. In particular, one can write a Lorentz-invariant mass term

$$m\bar{\nu}_L\nu_L^c \quad (12.21)$$

that does not involve right-handed neutrinos or left-handed anti-neutrinos. At present, it is experimentally not ruled out that neutrinos are indeed Majorana fermions.

We see that one can always write neutrino mass terms, whether right-handed neutrinos exist or not. And in quantum field theory, any interaction term that can be written (i. e. that is not forbidden by symmetries), *will have to be added to the Lagrangian* from quantum corrections at higher orders in the perturbative expansion. Hence from a quantum field theoretical viewpoint, there is no reason or principle that sets neutrino masses to zero.

Neutrino Oscillations. Suppose that some or all neutrinos get non-zero masses by some mechanism. Then, as for quarks, there is no reason to assume that the weak lepton eigenstates are the same as the mass eigenstates. For simplicity, consider only two families of leptons. We label the weak eigenstates as ν_e and ν_μ , and the mass eigenstates as ν_1 and

ν_2 . The weak eigenstates are some linear combinations of the mass eigenstates at time $t = 0$,

$$\begin{pmatrix} \nu_\mu(0) \\ \nu_e(0) \end{pmatrix} = \begin{pmatrix} \cos(\alpha) & \sin(\alpha) \\ -\sin(\alpha) & \cos(\alpha) \end{pmatrix} \begin{pmatrix} \nu_1(0) \\ \nu_2(0) \end{pmatrix}. \quad (12.22)$$

In a neutrino beam, the mass eigenstates ν_i are free particles with definite energies E_i , hence their time evolution is

$$\nu_i(t) = e^{-iE_i t} \nu_i(0). \quad (12.23)$$

Therefore,

$$\nu_\mu(t) = \cos(\alpha) \nu_1(0) e^{-iE_1 t} + \sin(\alpha) \nu_2(0) e^{-iE_2 t}. \quad (12.24)$$

Expanding $\nu_i(0)$ in terms of $\nu_e(0)$ and $\nu_\mu(0)$, this becomes

$$\begin{aligned} \nu_\mu(t) &= e^{-iE_1 t} \cos(\alpha) (\cos(\alpha) \nu_\mu(0) - \sin(\alpha) \nu_e(0)) \\ &\quad + e^{-iE_2 t} \sin(\alpha) (\sin(\alpha) \nu_\mu(0) + \cos(\alpha) \nu_e(0)) \\ &= (e^{-iE_1 t} \cos^2(\alpha) + e^{-iE_2 t} \sin^2(\alpha)) \nu_\mu(0) \\ &\quad + \sin(\alpha) \cos(\alpha) (e^{-iE_2 t} - e^{-iE_1 t}) \nu_e(0). \end{aligned} \quad (12.25)$$

Since

$$E_1 = \sqrt{m_1^2 + p^2}, \quad E_2 = \sqrt{m_2^2 + p^2}, \quad (12.26)$$

where p is the momentum of the states, we see that $E_1 \neq E_2$. In this case, a state that begins as a pure weak eigenstate ν_μ at time $t = 0$ has some ν_e mixed in. The probability to find the state ν_e at time t (for example via weak interactions) is

$$\begin{aligned} P(\nu_\mu \rightarrow \nu_e) &= |\langle \nu_e(0) | \nu_\mu(t) \rangle|^2 \\ &= \sin^2(\alpha) \cos^2(\alpha) |e^{-iE_2 t} - e^{-iE_1 t}|^2 \\ &= \frac{\sin^2(2\alpha)}{2} [1 - \cos((E_2 - E_1)t)]. \end{aligned} \quad (12.27)$$

We see that the probability to find ν_e in a beam that initially was pure ν_μ oscillates with time. This effect is called *neutrino oscillation*.

Data and PMNS Matrix. The effect of neutrino oscillations has been measured by several experiments, using neutrinos created in particle colliders, neutrinos from decaying pions produced by cosmic rays in the atmosphere, and neutrinos emitted by the sun. The data shows that

$$\begin{aligned} m_2^2 - m_1^2 &= 7.6(2) \cdot 10^{-5} \text{ eV}^2, \\ |m_3^2 - m_2^2| &= 2.3(1) \cdot 10^{-3} \text{ eV}^2. \end{aligned} \quad (12.28)$$

So the three neutrinos have different masses, which means that at least two of them must have non-zero masses!

The mixing of the neutrino flavors can be described by a matrix, similar to the CKM matrix of the quark flavor mixing. For the leptons, the matrix is called the *PMNS matrix*, after Pontecorvo–Maki–Nakagawa–Sakata. Parametrizing it in terms of three Euler angles

θ_{ij} and a complex phase δ (like the CKM matrix), the experimental values for the angles and the phase are

$$\theta_{12} = 33.6(8)^\circ, \quad \theta_{23} = 47(4)^\circ, \quad \theta_{13} = 8.54(15)^\circ, \quad \delta = 234(43)^\circ. \quad (12.29)$$

The non-zero value of δ in particular shows that there is CP violation also in the lepton sector.

The measurements of neutrino oscillations only give constraints on the differences of squared masses, hence they do not rule out large but very similar masses. However, cosmological data sets strong limits on the sum of all three masses. The lighter the neutrinos, the less likely they get bound in galaxies in the early universe. Cosmic microwave background and structure formation data implies that

$$\sum_i m_{\nu_i} < 0.2 \text{ eV}. \quad (12.30)$$

The data on mass differences show that at least one mass is $\geq 0.05 \text{ eV}$.

Seesaw Mechanism. Since the neutrino masses are non-zero experimentally, their masses have to be included in the Standard Model in some way. If right-handed neutrinos exist, it is natural to include a Dirac mass term

$$\mathcal{L}_D = -m_D \bar{\nu} \nu = -m_D (\bar{\nu}_L \nu_R + \bar{\nu}_R \nu_L) \quad (12.31)$$

to the Lagrangian, where the Dirac mass m_D is the product of a Higgs coupling and the Higgs vacuum expectation value. The observed neutrino masses imply that m_D is many orders of magnitude smaller than all other lepton masses. That is possible, but seems unnatural.

Because neutrinos are electrically neutral and form their own anti-particles, we saw above that they could be Majorana fermions. In this case, the right-handed neutrino and its anti-partner would form a Majorana spinor (ν_R^c, ν_R) , and the left-handed neutrino and its anti-partner would form another, independent Majorana spinor (ν_L, ν_L^c) . Hence one can add a Majorana mass term $\mathcal{L}_M$ for only the right-handed neutrino to the Lagrangian,

$$\mathcal{L}_M = -\frac{1}{2} M (\bar{\nu}_R^c \nu_R + \bar{\nu}_R \nu_R^c). \quad (12.32)$$

Such a term is independent of the Dirac mass term. A similar Majorana mass term for the left-handed neutrino is excluded because it would have non-zero hypercharge and hence would not be gauge invariant. The right-handed neutrino is assumed to be an $SU(2)_W$ isospin singlet, as all other right-handed fermions, and hence its Majorana mass term is gauge invariant.

Using that $\bar{\nu}_L = \bar{\nu}_R^c$ and $\nu_R = \nu_L^c$, one can write the Dirac mass term as

$$\mathcal{L}_D = -\frac{1}{2} m_D (\bar{\nu}_L \nu_R + \bar{\nu}_R^c \nu_L^c + \text{h. c.}) \quad (12.33)$$

The Dirac and Majorana mass terms then combine to

$$\mathcal{L}_{\nu, \text{mass}} = -\frac{1}{2} \begin{pmatrix} \bar{\nu}_L & \bar{\nu}_R^c \end{pmatrix} \begin{pmatrix} 0 & m_D \\ m_D & M \end{pmatrix} \begin{pmatrix} \nu_L^c \\ \nu_R \end{pmatrix} + \text{h. c.} \quad (12.34)$$

To find the mass eigenstates, we have to diagonalize the mass matrix, its eigenvalues will be the measured neutrino masses m_ν . The eigenvalues satisfy

$$0 = \det \begin{pmatrix} -m_\nu & m_D \\ m_D & M - m_\nu \end{pmatrix} = m_\nu^2 - Mm_\nu - m_D^2. \quad (12.35)$$

We know that the measured neutrino masses are small. Assuming that the Dirac mass m_D stems from a Higgs coupling, it is natural to expect that its value is comparable to the Higgs expectation value $v \approx 250 \text{ GeV}$. In this case, the small neutrino mass is obtained if $M \gg m_D$. Then the two eigenvalues are

$$m_\nu^{\text{light}} \approx \frac{m_D^2}{M}, \quad m_\nu^{\text{heavy}} \approx M. \quad (12.36)$$

For example, if $m_D \approx 100 \text{ GeV}$, then a neutrino mass of $m_\nu^{\text{light}} = 0.01 \text{ eV}$ requires

$$M \approx \frac{(100 \text{ GeV})^2}{0.01 \text{ eV}} = 10^{15} \text{ GeV}. \quad (12.37)$$

This is called the *seesaw mechanism*: As M goes up, m_ν goes down. It naturally leads to the small neutrino masses that we observe, assuming that there is a physically natural energy scale at $\approx 10^{15} \text{ GeV}$. This is near the grand unification scale Λ_{GUT} , which is encouraging. The neutrino states with $m_\nu \approx M$ are so heavy that we naturally do not observe them.

If the heavy neutrinos have zero hypercharge, zero weak isospin, and zero color charge, they do not interact with any of the gauge bosons, and are called *sterile neutrinos*. Such sterile neutrinos are another candidate for dark matter.

A Numerical Data

Unit conversion factors in natural units where $\hbar = c = 1$:

$$1 \text{ GeV}^{-1} = 6.6 \cdot 10^{-25} \text{ sec} = 2 \cdot 10^{-16} \text{ m}. \quad (\text{A.1})$$

Coupling constants:

$$\alpha = \frac{e^2}{4\pi} \approx \frac{1}{137}, \quad \alpha_2 := \frac{g_2^2}{4\pi} \approx \frac{1}{30}, \quad G_F = \frac{g_2^2}{4\sqrt{2}m_W^2} = 1.166\,378\,7(6) \cdot 10^{-5} \text{ GeV}^{-2}, \quad (\text{A.2})$$

$$\alpha_3 := \frac{g_3^2}{4\pi} \approx \begin{cases} >1 & \text{at } \ll 1 \text{ GeV} \\ 0.3 & \text{at } \approx 1 \text{ GeV} \\ 0.12 & \text{at } m_Z \approx 91.2 \text{ GeV} \\ \ll 1 & \text{at } > 91.2 \text{ GeV} \end{cases}. \quad (\text{A.3})$$

Higgs vacuum expectation value:

$$v = \frac{2m_W}{g_2} \approx 250 \text{ GeV}. \quad (\text{A.4})$$

Particles and their masses:

Particles	Masses	
ν_e, ν_μ, ν_τ	$m_\nu \leq 1.2 \text{ eV}$	
e^-, μ^-, τ^-	$m_e = 511 \text{ keV}, m_\mu = 105.7 \text{ MeV}, m_\tau = 1.777 \text{ GeV}$	
u, c, t	$m_u = 2.16 \text{ MeV}, m_c = 1.27 \text{ GeV}, m_t = 173 \text{ GeV}$	(A.5)
d, s, b	$m_d = 4.67 \text{ MeV}, m_s = 93 \text{ MeV}, m_b = 4.18 \text{ GeV}$	
$W^\pm, Z, \gamma$	$m_W = 80.4 \text{ GeV}, m_Z = 91.2 \text{ GeV}, m_\gamma < 10^{-18} \text{ eV}$	
h	$m_h = 125.09(24) \text{ GeV}$	

Lifetimes / decay widths:

$$\begin{aligned}\tau_\mu &= \frac{1}{\Gamma_\mu} = 2.196\,981\,1(22) \cdot 10^{-6} \text{ sec}, \\ \Gamma_h &\approx 4 \text{ MeV}.\end{aligned}\tag{A.6}$$

References

- [1] G. Kane, “*Modern Elementary Particle Physics (Second Edition)*”, Cambridge University Press (2017), Cambridge, UK.
- [2] A. Blondel, “*The Number of Neutrinos and the Z Line Shape*”, in: “*The Standard Theory of Particle Physics*”, Advanced Series on Directions in High Energy Physics, pp. 145-160, ed.: L. Maiani and L. Rolandi, World Scientific (2016), Singapore.
- [3] A. S. Kronfeld, “*Twenty-first Century Lattice Gauge Theory: Results from the QCD Lagrangian*”, Ann. Rev. Nucl. Part. Sci. 62, 265 (2012), [arxiv:1203.1204](#).