

A Child's Guide to Spinors

William O. Straub

Pasadena, California 91104

December 31, 2016

(Expanded and slightly revised January 11, 2017)

No one fully understands spinors. Their algebra is formally understood, but their geometrical significance is mysterious. In some sense they describe the “square root” of geometry and, just as understanding the concept of $\sqrt{-1}$ took centuries, the same might be true of spinors. — Sir Michael Atiyah, British mathematician

The fermionic particles that make up all ordinary matter in the universe are described not by scalars, vectors or tensors, but by mathematical quantities known as *spinors*. Scalars and vectors are taught in high school, and tensors normally get introduced in undergraduate school, but spinor formalism is generally postponed until graduate school. The reason seems to be that spinors, which fully incorporate special relativity—including the Lorentz group of rotations and boosts—are deemed more complicated than other, more familiar mathematical objects despite their apparent greater simplicity. When students ask what exactly is a spinor, they invariably hear that it's something like the “square root” of a vector, a two-component vector-like object that has special transformation properties. When the student then looks to the available literature, she finds detailed information presented only in quantum field theory or abstract mathematical descriptions that defy comprehension to the uninitiated.

Here we derive the spinor formalism at a very elementary level, intended as an introduction for the typical undergraduate student who wants to pursue directly related subjects, especially the Dirac relativistic electron equation.

Why I Hate Spinors

I recall the time, many years ago, when I was first introduced to the Dirac relativistic electron equation,

$$i\hbar \gamma^\mu \partial_\mu \Psi(x, t) = mc \Psi(x, t) \quad (1)$$

Profound, intuitive, easy to derive, and expressed in beautiful covariant notation, it all made perfect sense to me: the partial derivative ∂_μ is a covariant vector, the 4×4 Dirac gamma matrices γ^μ obviously represent a contravariant vector, and the wave function $\Psi(x, t)$ is a scalar (even though it has four components, like a 4-vector). No, said the professor, although the derivative is indeed a vector, γ^μ is not, so Lorentz invariance of the equation depends on how the wave function (which is also not a scalar) transforms. He then added, “*The wave function is a complex, multi-component vector-like quantity that has special transformation properties.*” The professor did not elaborate.

As the course went on, I noted that the gamma matrices were invariably treated as though they constituted a true vector—they could be expressed as either γ^μ or γ_μ , as the index could be raised or lowered by the metric tensor $g_{\mu\nu}$ or $g^{\mu\nu}$, and the matrices could also be transformed into flat-space Lorentz form using tetrads ($\gamma^a = e^a_\mu \gamma^\mu$). Furthermore, I learned that the metric tensor itself is composed of the gamma matrices via

$$g^{\mu\nu}(=\eta^{\mu\nu}) = \frac{1}{2}(\gamma^\mu \gamma^\nu + \gamma^\nu \gamma^\mu)$$

All this meant that the γ^μ must also be functions of the coordinates, which led me to think that Ψ really is a scalar after all, and that the Dirac equation must also be valid in curved space.

The professor went on to sketch the solution of the Dirac equation for the hydrogen atom (which I found was considerably more difficult than solving the non-relativistic Schrödinger equation). In all that time, the professor never mentioned the word “spinor.”

When I discovered that I was wrong about everything, I quickly learned to hate spinors. They're non-intuitive, far more complicated than they have any right to be, and come either with two components or four (like the Dirac

spinor, which is actually a stack of two spinors called a *bispinor*), both versions being equally hard to comprehend. They come to you innocently disguised as either scalars or vectors, but they do not act like scalars or vectors at all, especially when you try to calculate their inner and outer products. Since they involve complex variables, spinors are most comfortable in the world of quantum mechanics, although they can be dragged kicking and screaming into general relativity, where they cause all kinds of trouble. They figure prominently in the preposterously complicated theory known as *supersymmetry*, where they wear all manner of bewildering notational disguises involving subscripts, superscripts, daggers, hats, dots, tildes, squiggles and overbars, and even then the notation usually comes up short.

Spinors were originally introduced by the French mathematician Élie Cartan in 1913 (and subsequently greatly expanded upon by Hermann Weyl, Richard Brauer and Oswald Veblen). Unfortunately, Cartan's book *The Theory of Spinors* is quite formal and of little use to the undergraduate (I still find it incomprehensible). But spinors were later found to be absolutely indispensable in describing the behavior of *fermions*, particles that make up all the common stuff you're familiar with—electrons, protons, neutrons and the like.

Never believe that the Creator didn't have a sense of humor when she invented something as damnable as spinors to describe all the ordinary matter in the universe, including you.

1. Introduction

Again, fermions like protons, electrons and quarks comprise all the ordinary matter in the universe. But despite this nearly universal dominance in Nature, they do not obey the behavior typical of scalar, vector and tensor quantities, but that of spinors. So what is a spinor? It's essentially a two-component vector-like quantity in which rotations and Lorentz boosts are built into the overall formalism. In that simplistic sense, a spinor is not the "square root" of a vector, but much more like half of a 4-vector. The fact that they exist at all is remarkable, and if they were of only mathematical interest they might easily be ignored.

The notion of electron spin was first surmised in 1922 when the German physicists Otto Stern and Walther Gerlach noticed that a stream of silver atoms (each having a single electron in the outer 5s orbital) could be separated by a non-uniform magnetic field into two streams, "up" and "down." At the time, no one knew what to make of this odd behavior, but three years later George Uhlenbeck and Samuel Goudsmit proposed that electrons could exist in two spin states, $\pm 1/2$, each with the units of angular momentum $\hbar = h/2\pi$. Still, the concept of electron spin varied between a simplistic physical "spinning" of the electron about some axis (like a child's top) and a kind of internal angular momentum. Today we look upon electrons and other spin-1/2 particles as having an intrinsic angular momentum with no classical counterpart. This gave rise to the notion of such particles living in "spin space," an abstract two-dimensional internal space that requires a description beyond that of scalars, vectors and tensors. Cartan's spinor formalism was found to be appropriate for this description.

Many attempts have been made over the years to explain spinors at an intuitive, elementary level, but the simplest approach remains an appeal to basic Lorentz group theory. This is rather a pity, because undergraduate students often express an aversion to group theory because of its mathematical nature. But the theory of Lorentz rotations and boosts is relatively simple, and it has the nice property of being relativistic from the start. In addition, it neatly admits a formalism that underlies that of the 4-vectors it ordinarily applies to, which is where spinors come in. Even better, Lorentz theory confirms the intuitive notion that if a spinor represents half of a 4-vector (rather than the square root), then there should be two kinds of spinor: one comprising the upper half and another representing the lower half. This observation is critical, since a single two-component spinor can be shown to violate odd-even parity in quantum physics, and it takes two spinors acting together to preserve it. Thus, the Dirac bispinor—a four-component object consisting of two stacked spinors—fully preserves parity.

Our approach will therefore be based on the Lorentz transformations of rotations and boosts. There are, however, a few subtleties that other, more advanced treatments either gloss over or assert by inference, and I will try to explain these in a more straightforward and understandable manner as we go along.

2. Getting Started

Many physics professors, when questioned about the nature of spinors, will simply respond with "They're two-component, vector-like quantities with special transformation properties." The key word here is

transformation. A transformation is something like a change in appearance that somehow doesn't affect any underlying physics. But how on Earth can a simple quantity like

$$\xi = \begin{bmatrix} 1 \\ 0 \end{bmatrix}$$

(which is a spinor) possibly transform? To answer this, we resort to the standard formalism of vector transformations: we find ourselves a square matrix that takes a vector and turns it into another vector that's equivalent in some sense. Some care needs to be exercised here—a transformation might take a vector and physically turn it in a new direction (in which case it's a different vector), or it might just express the same vector in some other, possibly more convenient, coordinate system. It's mainly this latter case we'll be concerned with here.

Ordinary 3- and 4-vectors typically get transformed by matrices with real components. Such matrices reflect the real aspects of the vectors they transform, so we typically see 3×3 and 4×4 matrices that are themselves symmetric, antisymmetric, or orthogonal, with real components. But for mathematical reasons 2×2 matrices in physics need to be *unitary*, meaning that their conjugate transposes are equal to their inverses, or $U^\dagger = U^{-1}$. Therefore, unitary matrices are in general not real but *complex*, and they act on vectors that are also complex. It also means that the spinors they transform will have complex character. So, while 2×2 matrices and spinors might appear to be simpler than 4×4 matrices and 4-vectors, this simplicity is largely illusory, and that's probably what makes spinors so hard to understand.

Fortunately, it's not difficult to build the most general 2×2 matrix possible, which I exhibit here without proof. If a and b are any two complex variables, then the matrix

$$U = \begin{bmatrix} a & b \\ -b^* & a^* \end{bmatrix} \quad (2.1)$$

is unitary, subject to the condition that $aa^* + bb^* = 1$. Thus there are $4 - 1$ independent real variables in U , which will turn out to be very convenient. A related quantity we'll need later is the conjugate transpose of U , designated as $U^\dagger$, which also happens to be the inverse matrix U^{-1} . It is

$$U^\dagger = \begin{bmatrix} a^* & -b \\ b^* & a \end{bmatrix} \quad (2.2)$$

An arbitrary spinor transforms according to $\xi' = U\xi$, or

$$\begin{bmatrix} \xi'_1 \\ \xi'_2 \end{bmatrix} = \begin{bmatrix} a & b \\ -b^* & a^* \end{bmatrix} \begin{bmatrix} \xi_1 \\ \xi_2 \end{bmatrix}$$

or

$$\xi'_1 = a\xi_1 + b\xi_2, \quad \xi'_2 = -b^*\xi_1 + a^*\xi_2$$

This property of spinor transformation is actually quite complicated. Nevertheless, we'll need to determine a and b if we're to understand how spinors transform. But first an aside about the Pauli matrices, a set of 2×2 matrices that you should already be familiar with.

2.1. The Pauli Matrices

The set of 2×2 Pauli matrices plays a key role in spinor formalism, but only in establishing a connection with Lorentz rotations and boosts. The four Pauli matrices σ_μ (sometimes called a *quaternion*) are

$$\sigma_0 = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}, \quad \sigma_x = \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix}, \quad \sigma_y = \begin{bmatrix} 0 & -i \\ i & 0 \end{bmatrix}, \quad \sigma_z = \begin{bmatrix} 1 & 0 \\ 0 & -1 \end{bmatrix} \quad (2.1.1)$$

These matrices form what is called a *Lorentz algebra* defined by the commutation relation

$$\frac{1}{2}\sigma_x \cdot \frac{1}{2}\sigma_y - \frac{1}{2}\sigma_y \cdot \frac{1}{2}\sigma_x = \left[\frac{1}{2}\sigma_x, \frac{1}{2}\sigma_y \right] = \frac{1}{2}i\sigma_z \quad (2.1.2)$$

along with cyclic permutations of x, y, z . A more formal way of expressing this is

$$\left[\frac{1}{2}\sigma_i, \frac{1}{2}\sigma_j \right] = \frac{1}{2}i\epsilon_{ijk}\sigma_k,$$

where ϵ_{ijk} is the fully-antisymmetric *structure constant* of the Pauli algebra, with $\epsilon_{123} = 1$.

2.2. Basic Approach

In the previous section I exhibited the most general 2×2 unitary matrix possible. Here I'll use the Pauli matrices to develop a link between spinors and rotations of 4-vectors, and this link will ultimately provide the needed identities for the a and b quantities in the matrix U , at least as far as rotations are concerned.

Let us take the Cartesian spacetime vector

$$\begin{bmatrix} t \\ x \\ y \\ z \end{bmatrix}$$

(where $c = 1$), and dot it into the Pauli matrices:

$$x^\mu \sigma_\mu = t\sigma_0 + x\sigma_x + y\sigma_y + z\sigma_z$$

Calling this sum H , we then have the hermitian matrix

$$H = \begin{bmatrix} t+z & x-iy \\ x+iy & t-z \end{bmatrix} \quad (2.2.1)$$

whose determinant is the invariant quantity $t^2 - x^2 - y^2 - z^2$. One nice property of unitary matrices is that they preserve the structural form of the quantities they operate on. The matrix H thus transforms according to $H' = UHU^\dagger$, or

$$H' = \begin{bmatrix} t' + z' & x' - iy' \\ x' + iy' & t' - z' \end{bmatrix} \quad (2.2.2)$$

This means that $UHU^\dagger$ will take the t, x, y, z in H and convert them to t', x', y', z' , and from that we can figure out the identities needed for the components a, b in U . Sound complicated? It's really not too bad. But before we do that, I need to backtrack a bit again to understand how 4-vectors behave under Lorentz transformations.

3. Lorentz Transformations

The Lorentz group of transformations describes ordinary rotations in space as well as relativistic transformations of space and time when going from one inertial frame to another (called "boosts"). Since boosts involve the four dimensions x, y, z and $x^0 = ct$, it is desirable to express ordinary rotations in four dimensions as well (in which case the time dimension just goes along for the ride). Therefore, we begin by demanding that all vectors be expressed as 4-vectors, while the matrices that transform those vectors be 4×4 matrices. The rotation matrices are all orthogonal (that is, the transpose and the inverse are the same) with determinants equal to unity, and in four dimensions they comprise a group called $SO(4)$, while the boost matrices are all symmetric, again with determinants equal to unity. Again, the only 4-vector we'll need is the Cartesian spacetime vector

$$\begin{bmatrix} t \\ x \\ y \\ z \end{bmatrix}$$

The student should already be fairly aware of the Lorentz transformation matrix Λ , which transforms one 4-vector into another,

$$(x^\mu)' = \Lambda^\mu_\nu x^\nu$$

The components of Λ_j^i representing rotations are antisymmetric, while the Λ_j^0 boost terms are symmetric. However, when the components are lowered using the metric tensor $g_{\mu\nu}$ the entire matrix becomes antisymmetric. Later, we will refer to the parameters (θ or ϕ) of these lowered components as $\omega_{\mu\nu} (= -\omega_{\nu\mu})$.

Once we have grounded ourselves on the formalism of Lorentz transformations, we'll show how they're related in a very intimate way with the a, b components of the unitary matrix U used to define spinor behavior.

3.1. Generators of Lorentz Rotations

In three dimensions, the counterclockwise rotation of some 3-vector V_k about the z axis by the angle θ is given by $V' = R_z V$, or

$$\begin{bmatrix} V'_x \\ V'_y \\ V'_z \end{bmatrix} = \begin{bmatrix} \cos \theta & \sin \theta & 0 \\ -\sin \theta & \cos \theta & 0 \\ 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} V_x \\ V_y \\ V_z \end{bmatrix}$$

where V'_k represents the vector in the primed or rotated system and R_z is the associated (orthogonal) rotation matrix. This is easily converted into four-dimensional form by simply writing

$$\begin{bmatrix} V'_0 \\ V'_x \\ V'_y \\ V'_z \end{bmatrix} = \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & \cos \theta & \sin \theta & 0 \\ 0 & -\sin \theta & \cos \theta & 0 \\ 0 & 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} V_0 \\ V_x \\ V_y \\ V_z \end{bmatrix}$$

(As mentioned, the time component of the vector does not change under rotation, and the component just goes along for the ride). Similar expressions hold for rotations about the x and y axes, which are respectively $V' = R_x V$ and $V' = R_y V$. Summarizing, we can write

$$R_x = \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & \cos \theta & \sin \theta \\ 0 & 0 & -\sin \theta & \cos \theta \end{bmatrix}, \quad R_y = \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & \cos \theta & 0 & -\sin \theta \\ 0 & 0 & 1 & 0 \\ 0 & \sin \theta & 0 & \cos \theta \end{bmatrix}, \quad R_z = \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & \cos \theta & \sin \theta & 0 \\ 0 & -\sin \theta & \cos \theta & 0 \\ 0 & 0 & 0 & 1 \end{bmatrix} \quad (3.1.1)$$

We'll now convert the above 4×4 rotation matrices into rotation *generators*, the meaning of which will soon become easily apparent. Let's assume an infinitesimal rotation for the matrix R_z , which is

$$R_z(d\theta) = \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & d\theta & 0 \\ 0 & -d\theta & 1 & 0 \\ 0 & 0 & 0 & 1 \end{bmatrix} = I + i \begin{bmatrix} 0 & 0 & 0 & 0 \\ 0 & 0 & -i & 0 \\ 0 & i & 0 & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix} d\theta$$

where I is the 4×4 unit matrix. We *define* the matrix

$$J_z = \begin{bmatrix} 0 & 0 & 0 & 0 \\ 0 & 0 & -i & 0 \\ 0 & i & 0 & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix}$$

to be the generator of the z -rotation. To generate a finite angle θ , we need the product $[R_z(d\theta)]^N$, where $N \rightarrow \infty$:

$$R_z(\theta) = \left(I + \frac{iJ_z\theta}{N} \right)^N \rightarrow e^{iJ_z\theta}$$

The exponential term is actually a 4×4 matrix, which is easily expanded into

$$e^{iJ_z\theta} = \cos \theta I + i \sin \theta J_z \quad (3.1.2)$$

The generators J_x and J_y follow. Summarizing, we have

$$J_x = \begin{bmatrix} 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & -i \\ 0 & 0 & i & 0 \end{bmatrix}, \quad J_y = \begin{bmatrix} 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & i \\ 0 & 0 & 0 & 0 \\ 0 & -i & 0 & 0 \end{bmatrix}, \quad J_z = \begin{bmatrix} 0 & 0 & 0 & 0 \\ 0 & 0 & -i & 0 \\ 0 & i & 0 & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix} \quad (3.1.3)$$

3.2. Generators of Lorentz Boosts

In undergraduate school (and some high schools), students typically associate the Lorentz transformation with the variation of 4-vectors undergoing shifts from one inertial frame to another (often depicted as rocket ships flying past one another). They learn about primed and unprimed reference frames, time dilation, length contraction, the twin paradox and all that, but the word “boost” is rarely used to describe this kind of transformation. Well, it is called a boost, but more importantly it forms just half of the Lorentz formalism, the other half being the set of rotations, which you just learned about. Fortunately, the boost formalism is nearly identical to that of rotations.

Unlike rotations, which can be conveniently described with 3×3 matrices, boosts require 4×4 matrices right from the start (which is why I decided to express rotations with four-dimensional matrices). Imagine a frame of reference passing parallel to a fixed frame in the x -direction with velocity v . Then the two frames are related according to the transformation

$$\begin{bmatrix} t' \\ x' \\ y' \\ z' \end{bmatrix} = \begin{bmatrix} \gamma & \beta\gamma & 0 & 0 \\ \beta\gamma & \gamma & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} t \\ x \\ y \\ z \end{bmatrix}$$

where

$$\beta = v/c, \quad \gamma = \frac{1}{\sqrt{1-\beta^2}}$$

and where v is the velocity of the primed frame with respect to the unprimed (fixed) frame. The identity $\gamma^2 - \beta^2\gamma^2 = 1$ prompts the convenient identification

$$\cosh \phi = \gamma, \quad \sinh \phi = \beta\gamma$$

where ϕ is the “angle” associated with the boost. We can now proceed exactly as we did before with infinitesimal rotations by considering infinitesimal boosts. We summarize the associated boost generators K_i with

$$K_x = \begin{bmatrix} 0 & -i & 0 & 0 \\ -i & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix}, \quad K_y = \begin{bmatrix} 0 & 0 & -i & 0 \\ 0 & 0 & 0 & 0 \\ -i & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix}, \quad K_z = \begin{bmatrix} 0 & 0 & 0 & -i \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ -i & 0 & 0 & 0 \end{bmatrix} \quad (3.2.1)$$

Like the rotations, the 4×4 boost transformation matrices can be written simply as $e^{iK \cdot \phi}$. Note that the K_i are now all symmetric, as opposed to the antisymmetry of the (hermitian) rotation generators.

4. Putting It All Together

Let us now return to the transformation I set up back in Section 2.2, which was $H' = UH U^\dagger$, or

$$\begin{bmatrix} t' + z' & x' - iy' \\ x' + iy' & t' - z' \end{bmatrix} = \begin{bmatrix} a & b \\ -b^* & a^* \end{bmatrix} \begin{bmatrix} t + z & x - iy \\ x + iy & t - z \end{bmatrix} \begin{bmatrix} a^* & -b \\ b^* & a \end{bmatrix}$$

Brute force expansion gives

$$t' = (aa^* + bb^*) t = t \quad (4.1)$$

(by virtue of $aa^* + bb^* = 1$). We also have

$$x' = \frac{1}{2}(aa^* + a^*a^* - bb - b^*b^*)x - \frac{1}{2}i(aa - a^*a^* + bb - b^*b^*)y - (ab - a^*b^*)z \quad (4.2)$$

$$y' = \frac{1}{2}i(aa - a^*a^* - bb - b^*b^*)x + \frac{1}{2}(aa + a^*a^* + bb + b^*b^*)y - i(ab - a^*b^*)z \quad (4.3)$$

$$z' = (a^*b + ab^*)x + i(a^*b - ab^*)y + (aa^* - bb^*)z \quad (4.4)$$

The transformation thus preserves the time component, which means that the Lorentz boosts cannot participate in the approach we’re using. However, if we set

$$a = \cos \frac{1}{2}\theta, \quad b = i \sin \frac{1}{2}\theta$$

we see that (4.2) gives

$$x' = x, \quad y' = y \cos \theta + z \sin \theta, \quad z' = -y \sin \theta + z \cos \theta$$

which is the set of Lorentz rotations about the x -axis. Similarly, setting

$$a = \cos \frac{1}{2} \theta, \quad b = \sin \frac{1}{2} \theta$$

in (4.3) we get

$$x' = x \cos \theta - z \sin \theta, \quad y' = y, \quad z' = x \sin \theta + z \cos \theta$$

which is the set of rotations about the y -axis. Finally, setting

$$a = e^{\frac{1}{2}i\theta}, \quad b = 0$$

in (4.4) we have

$$x' = x \cos \theta + y \sin \theta, \quad y' = -x \sin \theta + y \cos \theta, \quad z' = z$$

which describes a rotation about the z -axis. (The student might ask if there are any other useful identities for a, b . There are none; that's it.)

The appearance of half-angles in all this is highly significant. In ordinary vector space, a 360° rotation brings the vector back to itself, but in spinor space a full 720° rotation is needed. In that sense, a spinor is rather like an arbitrary vector lying on a Möbius strip; it has to go around the strip twice to get back where it started.

The unitary transformation matrix U thus depends on the choice of a, b . However, the student should have no difficulty showing that

$$U_k = e^{\frac{1}{2}i\sigma_k \theta},$$

an identity that is sometimes expressed in texts as

$$U = e^{\frac{1}{2}i\sigma \cdot \theta} \quad (4.5)$$

when an arbitrary rotation axis is chosen.

Again, let me emphasize that U does not transform the time component of a spinor. This is a damned shame, because the set of Lorentz boosts cannot be expressed for spinors using this or any other equivalent formalism. The boosts will have to be introduced in another, somewhat artificial manner using what is referred to as a *representation*. We now turn to that approach.

5. Representations

Now comes perhaps the hardest part of all this (at least, it was and still is for me), and the main subtlety I mentioned at the beginning. The Lorentz algebra of the rotation generators J_i is

$$[J_x, J_y] = iJ_z$$

and its cyclic counterparts. Similarly, the Lorentz algebra of the Pauli matrices is

$$\left[\frac{1}{2}\sigma_x, \frac{1}{2}\sigma_y \right] = \frac{1}{2}i\sigma_z,$$

etc. Thus, the two matrices J_i and $\frac{1}{2}\sigma_i$ have exactly the same algebra. This cannot be a coincidence; it means there is some kind of fundamental correspondence between the matrices, in spite of the fact that one is orthogonal and 4×4 with unit determinant, a group that we call $SO(4)$, and the other is unitary and 2×2 , also with unit determinant, which is called $SU(2)$. This correspondence is given the representation $SO(4) = SU(2)$, where the equal sign is not to be taken literally. This representation is also called the $SO(4) \cdot SU(2)$ “double cover,” perhaps only in the sense that the “double” refers to the fact that the rotation dimension is double that of the spinor dimension. (Note that if I had left the rotation matrices in 3-dimensional form, as many texts do, none of this would make any sense.)

What all this implies is that a 4×4 orthogonal rotation matrix operating on a real 4-component vector is conceptually “equivalent” in some sense to a 2×2 unitary matrix operating on a complex 2-component spinor (most texts include statements like $J_i = \frac{1}{2}\sigma_i$, which is confusing as hell). Whether this makes any sense to you or not belies the fact that it not only works, but works brilliantly. It alone is responsible for opening up an entire new way of looking at how Nature operates regarding particles and fields in which the spinor formalism applies.

Still, setting $SO(4) = SU(2)$ looks somehow mathematically inadequate. Shouldn’t it be $SO(4) = SU(2) + SU(2)$, or something like that? If you considered this possibility, then you are way ahead of most people, because you are exactly right. One spinor accounts for only half the overall formalism, so there must be two distinct kinds of spinor. To see this, we’re going to do an easy calculation that Anthony Zee has rightly called “one of the most significant calculations in the history of twentieth century physics.”

Using our identities for the K_i , it’s a simple matter to show that these matrices neither commute nor form a Lorentz algebra; instead we have the commutator

$$[K_x, K_y] = -iJ_z$$

along with its cyclic counterparts. This identity has two unusual properties. One, it demonstrates that two boosts in different directions result in a rotation (a phenomenon responsible for the *Thomas precession* of an electron in a magnetic field). And two, there is a minus sign that turns out to be all-important in the overall scheme of things. What it means is that if we append $\pm i$ to any K_i , we get the commutator

$$[\pm iK_x, \pm iK_y] = iJ_z,$$

which gives precisely the same algebra as that for the J_i . Thus, if the J_i are assigned the representation $J_i \rightarrow \frac{1}{2}\sigma_i$, then we can also assign the similar representation $\pm iK_i \rightarrow \pm \frac{1}{2}\sigma_i$. That is,

$$e^{K_i} \rightarrow e^{\pm \frac{1}{2}i\sigma_i}$$

We can therefore write the complete unitary 2×2 transformation matrix for spinorial rotations and boosts as either of two combined quantities,

$$U = e^{\frac{1}{2}i\sigma \cdot \theta + \frac{1}{2}\sigma \cdot \phi} \quad (5.1)$$

or

$$U = e^{\frac{1}{2}i\sigma \cdot \theta - \frac{1}{2}\sigma \cdot \phi} \quad (5.2)$$

(Note the lack of the imaginary i on the boost term; it was killed off when we used $K \rightarrow \pm i\frac{1}{2}\sigma$.) So there are indeed two kinds of spinor: one gets transformed under the unitary matrix in (5.1), and the other transforms according to (5.2), with the overall formalism now denoted as $SO(4) = SU(2) \oplus SU(2)$. The spinor associated with (5.1) is traditionally called a “right handed” spinor (and given the label φ_R), while the other is a “left handed” spinor, called φ_L (these spinors are called “Weyl spinors” in honor of the man who elucidated many of their properties). One of the amazing facts of physics is that all neutrinos in the universe are left handed (because their spinor descriptions are of the left-handed type). Nature is left handed!

6. The Dirac Equation

When it was discovered that electrons can exist in both $+\frac{1}{2}$ and $-\frac{1}{2}$ spin states, the Austrian physicist Wolfgang Pauli suggested in 1927 that the scalar wave function Ψ in Schrödinger’s equation be replaced by a two-component spinor, each component representing one of the electron’s allowed spin states. But while Pauli’s approach worked, it presented a problem having to do with *parity*; that is, the sign reversal operation $\Psi(x, t) \rightarrow \Psi(-x, t)$ of a spinor gave inconsistent results, violating the notion that Nature should be mirror-image invariant.

However, in 1928, at the age of just 25, the great British mathematical physicist Paul Adrienne Maurice Dirac made a monumental discovery, perhaps the greatest discovery in all of modern physics. The student can look up the details, but what Dirac did was essentially take the square root of the relativistic energy-mass equation

$$E^2 = m^2c^4 + c^2p^2$$

He ended up with the set of four partial differential equations in (1) involving four 4×4 matrices γ^μ (called the gamma matrices, of course), along with a new four-component wave function $\Psi_a(x, t)$, where $a = 1, 2, 3, 4$. It was soon realized that Dirac's Ψ was a *bispinor*, a four-component mathematical object consisting of two Weyl-like spinors:

$$\Psi = \begin{bmatrix} \psi_1 \\ \psi_2 \\ \psi_3 \\ \psi_4 \end{bmatrix} = \begin{bmatrix} \phi_R \\ \phi_L \end{bmatrix}$$

Dirac's bispinor was found to preserve parity under the sign reversal operation $\Psi(x, t) \rightarrow \Psi(-x, t)$. Far more importantly, the spinor ϕ_R effectively represents the spin-up and spin-down components of an ordinary electron, while ϕ_L represents the spin-up and spin-down components of an anti-electron (known as a *positron*)—Dirac's work thus predicted the existence of antimatter (the positron was subsequently discovered in 1932, for which the Caltech experimental physicist Carl Anderson won the Nobel Prize). Dirac's relativistic electron equation also explained electron spin as a form of intrinsic angular momentum called S . Thus, the angular momentum L of an electron alone is not conserved—instead, it is $L + S$ that is conserved. Dirac's equation and its underlying mathematics today represent the foundation of much of modern quantum field theory.

Although I displayed Dirac's electron equation at the beginning, I won't bother to derive it here. I did mention that Dirac took the square root of the relativistic energy-mass equation $E^2 = m^2 c^4 + c^2 p^2$, arriving at a set of 4×4 matrices and a four-component spinor that had all the markings of a fully Lorentz-covariant equation. It is remarkable that the Dirac equation can also be derived from scratch using just the 2×2 transformation matrix for the two Weyl spinors φ_R and φ_L , but it's a rather laborious computation (you're referred to my 2005 paper for the gory details). I'll leave it as an exercise for the student to accept (1) as written (it's not hard to derive), because what I really want to do now is focus on the Lorentz covariance of the Dirac equation, which follows from what we've seen earlier.

Before we begin, however, there's a little technical detail you need to be aware of first. It has to do with the fact that there are two versions available for the gamma matrices, only one of which is considered appropriate for the Dirac equation. It will turn out that the Lorentz formalism we've developed so far is actually applicable only to Weyl spinors, and we'll have to make some adjustments when describing the Lorentz covariance of the Dirac equation.

6.1. Two Kinds of Gamma Matrices

I noted that Dirac arrived at a set of matrices γ^μ that he used in deriving his equation. They all exhibit the following properties:

$$(\gamma^0)^2 = I, \quad (\gamma^i)^2 = -I, \quad \gamma^\mu \gamma^\nu + \gamma^\nu \gamma^\mu = 2\eta^{\mu\nu} \quad (6.1.1)$$

where I is the 4×4 identity matrix and $\eta^{\mu\nu}$ is the flat-space form of the metric tensor $g^{\mu\nu}$, with signature $(+, -, -, -)$. Consequently, we also have the condition

$$\gamma^\mu \gamma^\nu + \gamma^\nu \gamma^\mu = 0 \quad \text{for } \mu \neq \nu \quad (6.1.2)$$

Anyway, the first set of gamma matrices we want is called the *Weyl representation*,

$$\gamma^0 = \begin{bmatrix} 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \\ 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \end{bmatrix}, \quad \gamma^1 = \begin{bmatrix} 0 & 0 & 0 & -1 \\ 0 & 0 & -1 & 0 \\ 0 & 1 & 0 & 0 \\ 1 & 0 & 0 & 0 \end{bmatrix}, \quad \gamma^2 = \begin{bmatrix} 0 & 0 & 0 & i \\ 0 & 0 & -i & 0 \\ 0 & -i & 0 & 0 \\ i & 0 & 0 & 0 \end{bmatrix}, \quad \gamma^3 = \begin{bmatrix} 0 & 0 & -1 & 0 \\ 0 & 0 & 0 & 1 \\ 1 & 0 & 0 & 0 \\ 0 & -1 & 0 & 0 \end{bmatrix}$$

while the other set is called the *Dirac representation*,

$$\gamma^0 = \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & -1 & 0 \\ 0 & 0 & 0 & -1 \end{bmatrix}, \quad \gamma^1 = \begin{bmatrix} 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 \\ 0 & -1 & 0 & 0 \\ -1 & 0 & 0 & 0 \end{bmatrix}, \quad \gamma^2 = \begin{bmatrix} 0 & 0 & 0 & -i \\ 0 & 0 & i & 0 \\ 0 & i & 0 & 0 \\ -i & 0 & 0 & 0 \end{bmatrix}, \quad \gamma^3 = \begin{bmatrix} 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & -1 \\ -1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \end{bmatrix}$$

You can check for yourself that both sets of matrices satisfy the conditions in (6.1.1) and (6.1.2).

Both sets of gamma matrices are used in quantum field theory, and their equivalence can be shown by way of the similarity transformation

$$\gamma_d^\mu = A \gamma_w^\mu A^{-1}$$

where A is the unitary matrix

$$A = A^{-1} = \frac{1}{\sqrt{2}} \begin{bmatrix} 1 & 1 & 1 & 1 \\ 1 & 1 & 1 & 1 \\ 1 & 1 & -1 & -1 \\ 1 & 1 & -1 & -1 \end{bmatrix} \quad (6.1.3)$$

This matrix can also be represented in compressed form using

$$A = \frac{1}{\sqrt{2}} \begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix}$$

6.2. Lorentz Transformation of the Weyl Spinors in 4D

We showed earlier how the Weyl spinors φ_R and φ_L transform under Lorentz rotations and boosts:

$$\begin{bmatrix} \varphi'_R \\ \varphi'_L \end{bmatrix} = e^{\frac{1}{2}i\sigma \cdot \phi} \begin{bmatrix} e^{\frac{1}{2}\sigma \cdot \phi} & 0 \\ 0 & e^{-\frac{1}{2}\sigma \cdot \phi} \end{bmatrix} \begin{bmatrix} \varphi_R \\ \varphi_L \end{bmatrix} \quad (6.2.1)$$

where I've set out the rotation matrix for convenience. In order to display this in full 4×4 notation, we first note that

$$\sigma \cdot \phi = \sigma \cdot \mathbf{n} \phi = \begin{bmatrix} n_3 & n_- \\ n_+ & -n_3 \end{bmatrix} \phi$$

where n_i is the unit vector representing the i th direction of the angle ϕ and $n_{\pm i} = n_1 \pm i n_2$. In full 4D notation, (6.2.1) is then

$$\begin{bmatrix} \varphi'_{R1} \\ \varphi'_{R2} \\ \varphi'_{L1} \\ \varphi'_{L2} \end{bmatrix} = e^{\frac{1}{2}i\sigma \cdot \phi} \begin{bmatrix} \cosh \frac{1}{2}\phi + n_3 \sinh \frac{1}{2}\phi & n_- \sinh \frac{1}{2}\phi & 0 & 0 \\ n_+ \sinh \frac{1}{2}\phi & \cosh \frac{1}{2}\phi - n_3 \sinh \frac{1}{2}\phi & 0 & 0 \\ 0 & 0 & \cosh \frac{1}{2}\phi - n_3 \sinh \frac{1}{2}\phi & -n_- \sinh \frac{1}{2}\phi \\ 0 & 0 & -n_+ \sinh \frac{1}{2}\phi & \cosh \frac{1}{2}\phi + n_3 \sinh \frac{1}{2}\phi \end{bmatrix} \begin{bmatrix} \varphi_{R1} \\ \varphi_{R2} \\ \varphi_{L1} \\ \varphi_{L2} \end{bmatrix}$$

It is interesting to note that this matrix can be expressed much more conveniently using

$$\begin{bmatrix} \cosh \frac{1}{2}\phi + n_3 \sinh \frac{1}{2}\phi & n_- \sinh \frac{1}{2}\phi & 0 & 0 \\ n_+ \sinh \frac{1}{2}\phi & \cosh \frac{1}{2}\phi - n_3 \sinh \frac{1}{2}\phi & 0 & 0 \\ 0 & 0 & \cosh \frac{1}{2}\phi - n_3 \sinh \frac{1}{2}\phi & -n_- \sinh \frac{1}{2}\phi \\ 0 & 0 & -n_+ \sinh \frac{1}{2}\phi & \cosh \frac{1}{2}\phi + n_3 \sinh \frac{1}{2}\phi \end{bmatrix} = \cosh \frac{1}{2}\phi I + \gamma^0 \gamma \cdot \mathbf{n} \sinh \frac{1}{2}\phi$$

where the gamma matrices are in the Weyl representation. (For convenience, I've completely dropped the rotation term; more about that later.) This begs the compact identification

$$S_w = \cosh \frac{1}{2}\phi I + \gamma^0 \gamma \cdot \mathbf{n} \sinh \frac{1}{2}\phi = e^{\frac{1}{2}\gamma^0 \gamma \cdot \mathbf{n}} \quad (6.2.2)$$

where we are now calling the transformation matrix S_w , since the Weyl representation of the gamma matrices was used in its construction. Remember that the exponential term is itself a 4×4 matrix.

6.3. Lorentz Transformation of the Dirac Bispinor

Although S_w as expressed by (6.2.2) is completely valid, for technical reasons the Weyl representation is not appropriate for the Dirac equation, where a diagonal form for γ^0 is desirable. We must therefore transform S_w into the Dirac form S_d . This is easily accomplished using the unitary matrix A in (6.1.3). We then have

$$S_d = A S_w A^{-1}$$

which is easily (well, it takes some work) shown to be

$$S_d = \begin{bmatrix} \cosh \frac{1}{2}\phi & 0 & n_3 \sinh \frac{1}{2}\phi & n_- \sinh \frac{1}{2}\phi \\ 0 & \cosh \frac{1}{2}\phi & n_+ \sinh \frac{1}{2}\phi & -n_3 \sinh \frac{1}{2}\phi \\ n_3 \sinh \frac{1}{2}\phi & n_- \sinh \frac{1}{2}\phi & \cosh \frac{1}{2}\phi & 0 \\ n_+ \sinh \frac{1}{2}\phi & -n_3 \sinh \frac{1}{2}\phi & 0 & \cosh \frac{1}{2}\phi \end{bmatrix} \quad (6.3.1)$$

Remarkably, this matrix can be cast into the exact same form as S_w in (6.2.2):

$$S_d = \cosh \frac{1}{2}\phi + \gamma^0 \gamma^i \cdot n \sinh \frac{1}{2}\phi = e^{\frac{1}{2}\gamma^0 \gamma \cdot \mathbf{n}} \quad (6.3.2)$$

where the gamma matrices are now in the Dirac representation.

To get the Dirac bispinor, we also have to transform the stack of Weyl spinors using A . In compressed form, this is

$$\begin{bmatrix} \varphi_R \\ \varphi_L \end{bmatrix} \rightarrow \frac{1}{\sqrt{2}} \begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix} \begin{bmatrix} \varphi_R \\ \varphi_L \end{bmatrix} = \frac{1}{\sqrt{2}} \begin{bmatrix} \varphi_R + \varphi_L \\ \varphi_R - \varphi_L \end{bmatrix} \quad (6.3.4)$$

This is the Dirac bispinor, which is normally just relabeled as

$$\Psi_d = \begin{bmatrix} \phi_R \\ \phi_L \end{bmatrix}$$

Note that each component of the Dirac bispinor is a mixture of left and right Weyl spinors, which guarantees its parity invariance.

6.4. What About Rotations?

You will notice that I eventually dropped all reference to Lorentz rotations in the above argument. You may feel that since spinors are somehow all about spin, leaving rotations out of the formalism would seem to be very inappropriate. While you can carry around the $\exp(\frac{1}{2}i\sigma \cdot \theta)$ factor if you like, it simply isn't interesting—it guarantees Lorentz rotation invariance, alright, but it's the boosts that generate all the physics in the Dirac formalism (most texts drop the rotation term off pretty early in the argument). That's because the boost parameters $\cosh \frac{1}{2}\phi$ and $\sinh \frac{1}{2}\phi$ are actually tied to the energy and momentum of spinor particles:

$$\cosh \phi = \gamma, \quad \cosh \frac{1}{2}\phi = \sqrt{\frac{\gamma+1}{2}} = \sqrt{\frac{E+mc^2}{2mc}},$$

$$\sinh \phi = \beta\gamma, \quad \sinh \frac{1}{2}\phi = \sqrt{\frac{\gamma-1}{2}} = \sqrt{\frac{E-mc^2}{2mc}},$$

which are useful identities that you should verify for yourself. We use a Lorentz boost to take a spin-1/2 particle at rest, $\Psi(0)$, and push it to one having a non-zero momentum $\Psi(\mathbf{p})$ in some particular direction (again, I provide the details in my 2005 write-up). The four components of the Dirac bispinor at rest are a snap to write down:

$$\Psi_1(0) = \begin{bmatrix} 1 \\ 0 \\ 0 \\ 0 \end{bmatrix}, \quad \Psi_2(0) = \begin{bmatrix} 0 \\ 1 \\ 0 \\ 0 \end{bmatrix}, \quad \Psi_3(0) = \begin{bmatrix} 0 \\ 0 \\ 1 \\ 0 \end{bmatrix}, \quad \Psi_4(0) = \begin{bmatrix} 0 \\ 0 \\ 0 \\ 1 \end{bmatrix}$$

For free spin-1/2 particles, the boosted Dirac bispinor provides an accurate description of both electrons and positrons.

7. Sundry Final Comments

Since we used the algebra of Lorentz rotations to derive the components of the unitary matrix U , you might be asking yourself why we couldn't also use H to derive the boost factors out of the same algebra. We could have

then dispensed with all this conceptual talk about “representations,” because the quantity $e^{\pm \frac{1}{2} \sigma \cdot \phi}$ might then have automatically popped out of the formalism. However, H always leaves the zeroeth component of a 4-vector invariant, whereas a boost requires that time also be transformed, so H just won’t work. Sadly, no independent or alternative scheme seems to be available for the boosts (at least I haven’t found any), so we’re stuck with the representation argument which, as I noted earlier, is probably the hardest part of understanding all this.

Like the Pauli matrices, the gamma matrices (in any representation) have truly fascinating properties that have direct application to quantum field theory. The student is encouraged to manipulate products of these matrices (along with powers of the products) to see how expansion of the exponentials in S_d and S_w generates the cosine, sine, cosh and sinh functions. They constitute the basis of what is known as *Clifford algebra*, and many books have been written about them.

It is also interesting to note that the gamma matrices play a key role in what is known as *spacetime algebra*, which is the 4D version of *geometrical algebra*. In these algebras, quantities like $\gamma^1 \gamma^3$ play the role of the imaginary i ; that is, $(\gamma^1 \gamma^3)^2 = -1$, so the mathematics of these algebras is entirely real, not complex.

In the transform of the Dirac bispinor $\Psi' = S_d \Psi$, the textbooks invariably display the matrix S_d as

$$S_d = e^{-\frac{i}{4} \omega_{\mu\nu} \sigma^{\mu\nu}}$$

where

$$\sigma^{\mu\nu} = \frac{i}{2} (\gamma^\mu \gamma^\nu - \gamma^\nu \gamma^\mu)$$

and where $\omega_{\mu\nu}$ is an antisymmetric matrix representing the boost and rotation parameters ϕ and θ in the full Lorentz transformation matrix Λ (for example, a rotation about the x -axis is represented by a θ in the ω_{23} position, while a boost in the z -direction is represented by a ϕ in the ω_{03} position). Most students find the formal derivation of S_d to be frustrating, not because of its complexity but because the approach that texts tend to use is sparse on details. Furthermore, the notation is redundant: why have two imaginary i ’s in the formula when they simply cancel? We thus have the equivalent expression

$$S_d = e^{\frac{1}{8} \omega_{\mu\nu} (\gamma^\mu \gamma^\nu - \gamma^\nu \gamma^\mu)}$$

But this is still redundant: since $\omega_{\mu\nu}$ is antisymmetric, we can write this as

$$S_d = e^{\frac{1}{4} \omega_{\mu\nu} \gamma^\mu \gamma^\nu}$$

Whoops, it’s *still* redundant! If we restrict the summation using $\mu < \nu$, we can finally write it as

$$S_d = e^{\frac{1}{2} \omega_{\mu\nu} \gamma^\mu \gamma^\nu}, \quad (\mu < \nu)$$

To test this, let’s assume a boost in the x -direction represented by $\omega_{01} = \phi$, with all other terms equal to zero. This gives

$$S_d = e^{\frac{1}{2} \gamma^0 \gamma^1 \phi} = \cosh \frac{1}{2} \phi + \gamma^0 \gamma^1 \sinh \frac{1}{2} \phi$$

in confirmation of (6.3.2). The S_d for rotations works exactly the same way. For example, a rotation about the y -axis (no boosts) has the parameter $-\theta$ sitting in the ω_{13} position, so we have

$$S_d = e^{-\frac{1}{2} \gamma^1 \gamma^3 \theta} = \cos \frac{1}{2} \theta - \gamma^1 \gamma^3 \sin \frac{1}{2} \theta$$

(Again, the student should verify these formulas by expanding the exponentials in series.)

So now you know what a spinor is, and when asked you can confidently reply, like my professor did, that *a spinor is a complex, multi-component vector-like quantity that has special transformation properties*. No wonder he didn’t elaborate!

References

1. J. Bjorken and S. Drell, *Relativistic quantum mechanics*, McGraw-Hill, 1964.
2. E. Cartan, *The theory of spinors*, Dover Publications, 1996.
3. D. Griffiths, *Introduction to elementary particles*, Wiley and Sons, 1987.
4. L. Ryder, *Quantum field theory*, 2nd ed., Cambridge University Press, 1996.
5. W. Straub, *Weyl spinors and the Dirac equation*, <http://www.weylmann.com/weyl-dirac.pdf>, 2005.
6. A. Zee, *Quantum field theory in a nutshell*, 2nd ed., Princeton University Press, 2010.