

2 The Supersymmetry Algebra

The purpose of this section is to describe, in mathematical terms, what supersymmetry actually is. Usually in physics, we think of symmetries as associated to groups. But, at least for continuous symmetries, these groups have an underlying algebra and often that contains all the information that we need. So it is with supersymmetry. We will describe the algebra that underlies supersymmetry and start to explore some of its representations.

I should warn you that this section will be a little dry in flavour. There will be few fields and certainly no dynamics. These will come in later sections. But this section lays the necessary groundwork for the stories that are to come.

2.1 The Lorentz Group

Minkowski space $\mathbb{R}^{1,3}$ is the stage for relativistic quantum field theory. This space comes equipped with the Minkowski metric

$$\eta_{\mu\nu} = \text{diag}(+1, -1, -1, -1)$$

The set of symmetries of Minkowski space include Lorentz transformations of the form $x^\mu \rightarrow \Lambda^\mu{}_\nu x^\nu$ where

$$\Lambda^T \eta \Lambda = \eta$$

Embedded among these are a couple of discrete transformations: parity with $\Lambda = \text{diag}(1, -1, -1, -1)$ and time reversal with $\Lambda = \text{diag}(-1, 1, 1, 1)$. The transformations that are continuously connected to the identity have $\det \Lambda = 1$ and $\Lambda^0{}_0 > 0$ and form the *Lorentz group* $SO(1, 3)$. (The restriction to $\Lambda^0{}_0 > 0$ is sometimes written as $SO^+(1, 3)$.)

Our main goal in this section is to spell out some properties of the spinor representations of the Lorentz group. In fact, strictly speaking the group $SO(1, 3)$ doesn't have any spinor representations. However, there is a closely related group called $\text{Spin}(1, 3)$ that does admit spinors. This is the double cover, in the sense that

$$SO(1, 3) \cong \text{Spin}(1, 3)/\mathbb{Z}_2$$

where that $\mathbb{Z}_2$ is the famous minus sign that spinors pick up under a 2π rotation, a minus sign that vectors like x^μ are oblivious to. The fact that there are spinors in our world is the statement that the true symmetry group is $\text{Spin}(1, 3)$ rather than $SO(1, 3)$.

When we introduced spinors in the [Quantum Field Theory](#) course, we did so by first looking at the algebra $so(1,3)$ that is shared by both groups $\text{Spin}(1,3)$ and $SO(1,3)$. A Lorentz transformation acting on a 4-vector can be written as

$$\Lambda = \exp \left(-\frac{i}{2} \omega_{\mu\nu} M^{\mu\nu} \right) \quad (2.1)$$

where $\omega_{\mu\nu}$ are six numbers that specify what Lorentz transformation we're doing, while $M^{\mu\nu} = -M^{\nu\mu}$ are a choice of six 4×4 anti-symmetric matrices that generate the different Lorentz transformations. The matrix indices are suppressed in the above expressions; in their full glory we would write $(M^{\mu\nu})^\rho{}_\sigma$. So, for example

$$(M^{01})^\rho{}_\sigma = i \begin{pmatrix} 0 & 1 & 0 & 0 \\ 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{pmatrix} \quad \text{and} \quad (M^{12})^\rho{}_\sigma = i \begin{pmatrix} 0 & 0 & 0 & 0 \\ 0 & 0 & -1 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{pmatrix} \quad (2.2)$$

(Note that the generators differ by a factor of i from those defined in the [Quantum Field Theory](#) lectures. This is compensated by an extra factor of i in the exponent (2.1).) The matrices generate the algebra $so(1,3)$,

$$[M^{\mu\nu}, M^{\rho\sigma}] = i (\eta^{\nu\rho} M^{\mu\sigma} - \eta^{\nu\sigma} M^{\mu\rho} + \eta^{\mu\sigma} M^{\nu\rho} - \eta^{\mu\rho} M^{\nu\sigma}) \quad (2.3)$$

In the lectures on [Quantum Field Theory](#), we then constructed the spinor representations by first looking at the Clifford algebra of gamma matrices, $\{\gamma^\mu, \gamma^\nu\} = 2\eta^{\mu\nu}$ and, from these, constructing a new representation of the Lorentz algebra (2.3). Here, we'll take a slightly different path. It will be useful to first extract a little more information from the algebra (2.3).

The six different Lorentz transformations naturally decompose into three rotations J_i and three boosts K_i , defined by

$$J_i = \frac{1}{2} \epsilon_{ijk} M_{jk} \quad \text{and} \quad K_i = M_{0i}$$

where these $j, k = 1, 2, 3$ indices are summed over, and $\epsilon_{123} = +1$. The rotation matrices are Hermitian, with $J_i^\dagger = J_i$ while the boost matrices are anti-Hermitian with $K_i^\dagger = -K_i$. This ensures that the rotations in (2.1) give rise to a compact group while the boosts are non-compact. From the Lorentz algebra, we find that these generators obey

$$[J_i, J_j] = i\epsilon_{ijk} J_k \quad , \quad [J_i, K_j] = i\epsilon_{ijk} K_k \quad , \quad [K_i, K_j] = -i\epsilon_{ijk} J_k$$

The rotations form an $su(2)$ sub-algebra. That, of course, is to be expected and is related to the fact that $SO(3) \cong SU(2)/\mathbb{Z}_2$.

We can, however, find two mutually commuting $su(2)$ algebras sitting inside $so(1,3)$. For this we take the linear combinations

$$A_i = \frac{1}{2}(J_i + iK_i) \quad \text{and} \quad B_i = \frac{1}{2}(J_i - iK_i)$$

Both of these are Hermitian, with $A_i^\dagger = A_i$ and $B_i^\dagger = B_i$. They obey

$$[A_i, A_j] = i\epsilon_{ijk}A_k \quad , \quad [B_i, B_j] = i\epsilon_{ijk}B_k \quad , \quad [A_i, B_j] = 0 \quad (2.4)$$

But we know about representations of $SU(2)$: they are labelled by an integer or half-integer $j \in \frac{1}{2}\mathbb{Z}$ which, in the context of rotations, we call “spin”. The dimension of the representation is then $2j+1$. The fact that we can find two $su(2)$ sub-algebras of the Lorentz algebra tells us that all representations must carry two such labels

$$(j_1, j_2) \quad \text{with} \quad j_1, j_2 \in \frac{1}{2}\mathbb{Z} \quad (2.5)$$

and has dimension $(2j_1+1)(2j_2+1)$. We’ll flesh out the meaning of these representations more below. But for now, we can identify the simplest such representations just by counting: we have

$$\begin{aligned} (0, 0) : & \quad \text{scalar} \\ (\tfrac{1}{2}, 0) : & \quad \text{left-handed Weyl spinor} \\ (0, \tfrac{1}{2}) : & \quad \text{right-handed Weyl spinor} \\ (\tfrac{1}{2}, \tfrac{1}{2}) : & \quad \text{vector} \\ (1, 0) : & \quad \text{self-dual 2-form} \\ (0, 1) : & \quad \text{anti-self-dual 2-form} \end{aligned}$$

We see that the smallest representations of the Lorentz group are the left- and right-handed Weyl spinors. What we call the physical spin of a particle is the quantum number under rotations $\vec{J}$: this is $j = j_1 + j_2$.

There’s something a little odd about the our discovery of two $su(2)$ sub-algebras. After all, it certainly isn’t true that the Lorentz group is isomorphic to two copies of $SU(2)$. This is because $SU(2)$ is a compact group: keep doing a rotation and you will eventually get back to where you started. Indeed, two copies of the group $SU(2)$ give rotation group of Euclidean space $\mathbb{R}^4$.

$$\text{Spin}(4) \cong SU(2) \times SU(2) \quad \text{with} \quad SO(4) \cong \text{Spin}(4)/Z_2$$

In contrast, the Lorentz group is non-compact: keep boosting and you get further and further from where you started. How does this manifest itself in the two $su(2)$ algebras that we’ve found in (2.4)?

The answer is a little subtle and is to be found in the reality properties of the generators A_i and B_i . Recall that all integer, $j \in \mathbb{Z}$, representations of $SU(2)$ are real, while all half-integer spin, $j \in \mathbb{Z} + \frac{1}{2}$, are pseudoreal (which means that, while not actually real, the representation is isomorphic to its complex conjugate). However, the A_i and B_i in (2.4) do *not* have these properties. You can see in (2.2) that both J_i and K_i are pure imaginary. This, in turn, means that the generators A_i and B_i are complex conjugates of each other

$$(A_i)^\star = -B_i$$

This is where the difference lies that distinguishes $SO(4)$ from $SO(1, 3)$. The Lie algebra $so(1, 3)$ does not contain two, mutually commuting copies of the real Lie algebra $su(2)$, but only after a suitable complexification. This means that certain complex linear combinations of the Lie algebra $su(2) \times su(2)$ are isomorphic to $so(1, 3)$. To highlight this, the relationship between the two is sometimes written as

$$so(1, 3) \cong su(2) \times su(2)^\star$$

For our purposes, it means that the complex conjugate of a representation (j_1, j_2) exchanges the two quantum numbers

$$(j_1, j_2)^\star = (j_2, j_1)$$

Both the scalar representation $(0, 0)$ and the vector representation $(\frac{1}{2}, \frac{1}{2})$ are real, while the left- and right-handed Weyl spinors $(\frac{1}{2}, 0)$ and $(0, \frac{1}{2})$ are exchanged under complex conjugation. This last statement will be important as we proceed. In the context of quantum field theory, if a field appears in a theory then so too does its complex conjugate. This means that if you have a left-handed spinor, you also have a right-handed complex conjugated spinor.

2.1.1 Spinors and $SL(2, \mathbb{C})$

There is another way to discover spinors, this time one that doesn't involve going through the algebra. We will use the fact that there is an isomorphism between two groups

$$\text{Spin}(1, 3) \cong SL(2, \mathbb{C}) \tag{2.6}$$

To see this, we first note that we can write a point x^μ in Minkowski space as a 2×2 Hermitian matrix,

$$X = x_\mu \sigma^\mu = \begin{pmatrix} x_0 + x_3 & x_1 - ix_2 \\ x_1 + ix_2 & x_0 - x_3 \end{pmatrix}$$

where we've introduced the 4-vector of 2×2 matrices,

$$\sigma^\mu = (1, \sigma^i) \quad \text{with} \quad \sigma^1 = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}, \quad \sigma^2 = \begin{pmatrix} 0 & -i \\ i & 0 \end{pmatrix}, \quad \sigma^3 = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} \quad (2.7)$$

The σ^i are, of course, the Pauli matrices. The matrix X is Hermitian: $X = X^\dagger$. Moreover, there is clearly a one-to-one map between 4-vectors x^μ and 2×2 Hermitian matrices. The Minkowski inner product is particularly natural in this language: it is

$$\det X = (x_0)^2 - (x_1)^2 - (x_2)^2 - (x_3)^2 = x_\mu x^\mu$$

Now consider an $SL(2, \mathbb{C})$ transformation that acts as

$$X \rightarrow X' = SX S^\dagger \quad (2.8)$$

with $S \in SL(2, \mathbb{C})$. We have $(X')^\dagger = X'$ and $\det X' = \det X$ since $\det S = 1$. This means that the map (2.8) must be a Lorentz transformation.

In fact, it is not hard to see that we can implement all Lorentz transformations this way and we'll give an explicit construction of the generators shortly. For now, we can just do some simple counting. A general complex 2×2 matrix has 4 complex entries. The requirement that its determinant is 1 reduces this to 3 complex parameters, or 6 real parameters. This agrees with the dimension of the Lorentz group: 6 = 3 rotations + 3 boosts. Moreover, the $SL(2, \mathbb{C})$ transformation $S = -1$ does not act on X , which is the reason why $SL(2, \mathbb{C})$ coincides with the double cover (2.6).

It is clear that the fundamental representation of $SL(2, \mathbb{C})$ is not a 2×2 matrix: it is a 2-component, complex object $\psi_\alpha = (\psi_1, \psi_2)$ that transforms as

$$\psi_\alpha \rightarrow S_\alpha^\beta \psi_\beta \quad \alpha, \beta = 1, 2$$

Clearly it is a complex two-dimensional representation. In terms of our previous classification (2.5), we take it to correspond to $(\frac{1}{2}, 0)$: it is what we call a *left-handed Weyl spinor*.

Given any complex representation of a Lie group, we can always form another representation by taking the conjugate. This is equivalent to the original if we can find a matrix C for which $S^* = CSC^{-1}$. In the present case, no such C exists and the matrix S and its conjugate S^* are inequivalent representations. We denote the complex conjugate as

$$(\psi_\alpha)^\dagger = \bar{\psi}_{\dot{\alpha}}$$

We've adopted two notational flourishes to distinguish the two representations. First, we use different indices $\alpha, \beta = 1, 2$ and $\dot{\alpha}, \dot{\beta} = 1, 2$ for the two different representations. This is useful because the two indices are telling us that the objects transform in different ways. In addition, we also add a bar over any object, like $\bar{\psi}$, that transforms in the conjugate representation. This allows us to identify these objects even when we suppress the indices. (Note that a bar on a Weyl spinor simply means complex conjugation while, as we learned in the [Quantum Field Theory](#) lectures, a bar on a Dirac spinor means complex transpose together with multiplication by γ^0 .) The complex conjugate spinor then transforms as

$$\bar{\psi}_{\dot{\alpha}} \rightarrow (S^*)_{\dot{\alpha}}^{\dot{\beta}} \bar{\psi}_{\dot{\beta}} \quad \dot{\alpha}, \dot{\beta} = 1, 2$$

In our previous classification (2.5) it is the representation $(0, \frac{1}{2})$. It is a *right-handed Weyl spinor*.

Some of the index conventions above (and below) differ from what you may have seen in other contexts and it's worth quickly explaining why. Suppose that we've got a vector u that transforms in the fundamental of $SU(N)$. We write the components as u_a with $a = 1, \dots, N$. The vector $u^\dagger$ transforms in the conjugate representation and we would write these components as $(u^\dagger)^a$, with the index raised and no dots in sight. This reflects the fact that we can contract $u^\dagger$ and u to form a singlet: $(u^\dagger)^a u_a$. However, the representations of $SL(2, \mathbb{C})$ have a different structure and, as we'll see shortly, you can't contract a spinor and its conjugate to get a singlet. That's why we introduce the strange looking dotted indices, rather than raising the index, to distinguish the conjugate representation

Building Scalars from Spinors

The group $SL(2, \mathbb{C})$ has the following invariant tensors

$$\epsilon^{\alpha\beta} = \epsilon^{\dot{\alpha}\dot{\beta}} = \begin{pmatrix} 0 & 1 \\ -1 & 0 \end{pmatrix} \quad \text{and} \quad \epsilon_{\alpha\beta} = \epsilon_{\dot{\alpha}\dot{\beta}} = \begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix}$$

Note that the $\epsilon_{\alpha\beta}$ with indices lowered differs by a minus sign from $\epsilon^{\alpha\beta}$. This ensures that one is the inverse of the other: $\epsilon^{\alpha\beta} \epsilon_{\beta\gamma} = \delta_\gamma^\alpha$. This, in turn, means that when we use epsilon symbols to raise and lower indices (as we will below) then if we choose to raise an index and subsequently lower it again then we don't get a minus sign for our troubles.

Given, say, two left-handed Weyl fermions ψ and χ , we can use the epsilon tensors to form invariants. We define

$$\psi\chi := \epsilon^{\alpha\beta}\psi_\beta\chi_\alpha = \psi_2\chi_1 - \psi_1\chi_2$$

To see that these are, indeed, invariants under $SL(2, \mathbb{C})$, we just need to perform a transformation

$$\psi\chi \rightarrow S_\alpha^\gamma S_\beta^\delta \epsilon^{\alpha\beta}\psi_\delta\chi_\gamma = (\det S)\epsilon^{\gamma\delta}\psi_\delta\chi_\gamma = \psi\chi \quad (2.9)$$

where, in the first equality we've used the fact that $S_\alpha^\gamma S_\beta^\delta \epsilon^{\alpha\beta} = \det S \epsilon^{\gamma\delta}$, which you can confirm simply by checking all the cases $\gamma, \delta = 1, 2$. In the second equality we've used the fact that $\det S = 1$.

In some ways, the ϵ symbols play a role for spinors that is akin to role played by the metric $\eta^{\mu\nu}$ for vectors. Of course, one key difference is that $\epsilon^{\alpha\beta}$ is anti-symmetric, but this tallies nicely with the fact that, in quantum field theory, spinors are anti-commuting Grassmann variables. We then have

$$\psi\chi = \psi_2\chi_1 - \psi_1\chi_2 = -\chi_1\psi_2 + \chi_2\psi_1 = \chi\psi$$

In particular, $\psi\psi = 2\psi_2\psi_1$ is non-vanishing.

We can do something similar for right-handed fermions. However, a fiddly minus sign rears its head. We define

$$\bar{\psi}\bar{\chi} := \epsilon^{\dot{\alpha}\dot{\beta}}\bar{\psi}_{\dot{\alpha}}\bar{\chi}_{\dot{\beta}} = \bar{\psi}_1\bar{\chi}_2 - \bar{\psi}_2\bar{\chi}_1 \quad (2.10)$$

With anti-commuting spinors, we again have $\bar{\psi}\bar{\chi} = \bar{\chi}\bar{\psi}$. Note that the ordering of the indices in (2.10) differs from (2.9). The reason for choosing this different ordering, resulting in a minus sign difference in the definitions, is that it ensures that $(\psi\chi)^\dagger = \bar{\psi}\bar{\chi}$, since

$$(\psi\chi)^\dagger = (\psi_2\chi_1 - \psi_1\chi_2)^\dagger = \bar{\chi}_1\bar{\psi}_2 - \bar{\chi}_2\bar{\psi}_1 = \bar{\psi}\bar{\chi}$$

We can use the ϵ symbols to raise and lower spinor indices, just as we use the Minkowski metric to raise and lower vector indices. We have

$$\psi^\alpha = \epsilon^{\alpha\beta}\psi_\beta \quad , \quad \psi_\alpha = \epsilon_{\alpha\beta}\psi^\beta \quad \text{and} \quad \bar{\psi}^{\dot{\alpha}} = \epsilon^{\dot{\alpha}\dot{\beta}}\bar{\psi}_{\dot{\beta}} \quad , \quad \bar{\psi}_{\dot{\alpha}} = \epsilon_{\dot{\alpha}\dot{\beta}}\bar{\psi}^{\dot{\beta}}$$

In this notation, the Lorentz scalars (2.10) become

$$\psi\chi = \psi^\alpha\chi_\alpha \quad \text{and} \quad \bar{\psi}\bar{\chi} = \bar{\psi}_{\dot{\alpha}}\bar{\chi}^{\dot{\alpha}}$$

Our fiddly minus sign difference between (2.9) and (2.10) has now transmuted into the following rule: for left-handed spinors we should contract (undotted) indices in the direction $\searrow$, while for right-handed spinors we should contract (dotted) indices in the direction $\nearrow$.

We can ask how these new objects ψ^α and $\bar{\psi}^{\dot{\alpha}}$ fare under Lorentz transformations. We have

$$\begin{aligned}\psi^\alpha &\rightarrow \epsilon^{\alpha\beta} S_\beta^\gamma \psi_\gamma = (S^{-1T})^\alpha_\beta \psi^\beta \\ \bar{\psi}^{\dot{\alpha}} &\rightarrow \epsilon^{\dot{\alpha}\dot{\beta}} (S^\star)_{\dot{\beta}}^{\dot{\gamma}} \bar{\psi}_{\dot{\gamma}} = (S^{-1\dagger})^{\dot{\alpha}}_{\dot{\beta}} \bar{\psi}^{\dot{\beta}}\end{aligned}\tag{2.11}$$

where the equality follows from the following algebra

$$S_\alpha^\gamma \epsilon^{\alpha\beta} S_\beta^\delta = \epsilon^{\gamma\delta} \quad \Rightarrow \quad (S^T)^\gamma_\alpha \epsilon^{\alpha\beta} S_\beta^\delta = \epsilon^{\gamma\delta} \quad \Rightarrow \quad \epsilon^{\alpha\beta} S_\beta^\delta = (S^{-1T})^\alpha_\gamma \epsilon^{\gamma\delta}$$

with similar manipulations for the right-handed spinor. The matrices S^{-1T} don't form a new representation of $SL(2, \mathbb{C})$; they are equivalent to the fundamental representation since, from above, we have $\epsilon S \epsilon^{-1} = S^{-1T}$. This means that the covariant and contravariant left-handed spinors ψ_α and ψ^α transform in equivalent representations. Similarly, the right-handed spinors $\bar{\psi}_{\dot{\alpha}}$ and $\bar{\psi}^{\dot{\alpha}}$ transform in equivalent representations.

Building Vectors from Spinors

A key take-away from our discussion above is that if you want to form a Lorentz scalar then you need a pair of left-handed fermions or a pair of right handed fermions. Suppose that we instead have one object of each type, say a left-handed spinor ψ_α and a right-handed spinor $\bar{\chi}_{\dot{\alpha}}$. What kind of object can we then build? The answer is clear from the quantum numbers of these representations:

$$(\tfrac{1}{2}, 0) \otimes (0, \tfrac{1}{2}) = (\tfrac{1}{2}, \tfrac{1}{2})$$

This is the vector representation of the Poincaré group.

To explicitly construct the vector, we sandwich the Pauli matrices

$$(\sigma^\mu)_{\alpha\dot{\alpha}} = (1, \sigma^i)_{\alpha\dot{\alpha}}$$

between two spinors. We write

$$\psi \sigma^\mu \bar{\chi} = \psi^\alpha (\sigma^\mu)_{\alpha\dot{\alpha}} \bar{\chi}^{\dot{\alpha}}$$

Note that, as shown above, the Pauli matrices σ^μ should come with an index of each type – one undotted, and one dotted – and both subscripts. Taking the conjugate, we have $(\psi \sigma^\mu \bar{\chi})^\dagger = \chi \sigma^\mu \bar{\psi}$.

To see that the object does indeed transform as a 4-vector, we can contract this with any other 4- vector x^μ to give $\psi X \bar{\chi}$ with $X = x_\mu \sigma^\mu$. But we know from (2.8) and (2.11) how each of these transforms: we then have

$$\psi X \bar{\chi} = \psi^\alpha X_{\alpha\dot{\alpha}} \bar{\chi}^{\dot{\alpha}} \rightarrow (\psi^\beta (S^{-1})_\beta^\alpha) (S_\alpha^\delta X_{\delta\dot{\delta}} S_{\dot{\alpha}}^{\star\dot{\delta}}) (\bar{\chi}^{\dot{\beta}} (S^{\star-1})_{\dot{\beta}}^{\dot{\alpha}}) = \psi X \bar{\chi}$$

The fact that $\psi X \bar{\chi}$ forms a singlet shows that $\psi \sigma^\mu \bar{\chi}$ must transform as a vector. In fancy maths words, we say that the Pauli matrices act as the intertwiner between the different representations.

We can use the epsilon symbols to raise the spinor indices on the Pauli matrices $\sigma_{\alpha\dot{\alpha}}^\mu$. This gives us a closely related set of matrices that we denote

$$(\bar{\sigma}^\mu)^{\dot{\alpha}\alpha} = \epsilon^{\alpha\beta} \epsilon^{\dot{\alpha}\dot{\beta}} \sigma_{\beta\dot{\beta}}^\mu$$

The bar on $\bar{\sigma}$ doesn't denote anything to do with complex conjugation. The $\bar{\sigma}^\mu$ are simply a different set of 2×2 matrices from σ^μ . Note that the indices have not only been raised, but also switched: σ^μ has the undotted index first, while $\bar{\sigma}^\mu$ has the dotted index first. If we define $\epsilon = i\sigma^2$ then, viewed as matrix multiplication, we have $\bar{\sigma} = \epsilon \sigma^T \epsilon^T$. A quick calculation shows that

$$(\bar{\sigma}^\mu)^{\dot{\alpha}\alpha} = (1, -\sigma^i)^{\dot{\alpha}\alpha}$$

We can then similarly construct the vector

$$\bar{\chi} \bar{\sigma}^\mu \psi = \bar{\chi}_{\dot{\alpha}} (\bar{\sigma}^\mu)^{\dot{\alpha}\alpha} \psi_\alpha$$

This isn't a new object: you can check that $\psi \sigma^\mu \bar{\chi} = -\bar{\chi} \bar{\sigma}^\mu \psi$.

Generators of $SL(2, \mathbb{C})$

Finally we can give a description of the generators of $SL(2, \mathbb{C})$. We define the anti-symmetrised product of sigma matrices,

$$(\sigma^{\mu\nu})_\alpha^\beta = \frac{i}{4} (\sigma^\mu \bar{\sigma}^\nu - \sigma^\nu \bar{\sigma}^\mu)_\alpha^\beta$$

These are linearly independent and so can be taken as a generators of $SL(2, \mathbb{C})$. Because of the anti-symmetry in μ and ν , there are six such generators which is the dimension of the Lorentz group. Indeed, we can see explicitly that these generate the Lorentz group by computing the commutator

$$[\sigma^{\mu\nu}, \sigma^{\rho\sigma}] = i (\eta^{\nu\rho} \sigma^{\mu\sigma} - \eta^{\nu\sigma} \sigma^{\mu\rho} + \eta^{\mu\sigma} \sigma^{\nu\rho} - \eta^{\mu\rho} \sigma^{\nu\sigma})$$

This reproduces the algebra of the Lorentz group (2.3) as promised. A left-handed spinor then transforms as

$$\psi_\alpha \rightarrow \exp\left(-\frac{i}{2}\omega_{\mu\nu}\sigma^{\mu\nu}\right)_\alpha^\beta \psi_\beta \quad (2.12)$$

where $\omega_{\mu\nu}$ are the same set of six numbers that specify the Lorentz transformation (2.1).

The conjugate representation is generated by

$$(\bar{\sigma}^{\mu\nu})^{\dot{\alpha}}_{\dot{\beta}} = \frac{i}{4}(\bar{\sigma}^\mu\sigma^\nu - \bar{\sigma}^\nu\sigma^\mu)^{\dot{\alpha}}_{\dot{\beta}}$$

These too satisfy the algebra of the Lorentz group. Correspondingly, a right-handed spinor transforms as

$$\bar{\psi}^{\dot{\alpha}} \rightarrow \exp\left(-\frac{i}{2}\omega_{\mu\nu}\bar{\sigma}^{\mu\nu}\right)^{\dot{\alpha}}_{\dot{\beta}} \bar{\psi}^{\dot{\beta}} \quad (2.13)$$

Note that, from the positioning of the indices of $\bar{\sigma}^{\mu\nu}$, these act naturally as generators on $\bar{\psi}^{\dot{\alpha}}$, with the index raised.

2.1.2 Lagrangians for Spinors

We can now describe how to construct Lagrangians from a Weyl spinor. Suppose that we have just a single left-handed Weyl spinor ψ to play with. This necessarily comes with its conjugate, a right-handed spinor $\bar{\psi} = \psi^\dagger$. We can then form a kinetic term

$$\mathcal{S}_{\text{Weyl}} = - \int d^4x \, i\bar{\psi}\bar{\sigma}^\mu\partial_\mu\psi \quad (2.14)$$

Upon quantisation, this theory gives a single massless, left-handed fermion of helicity $-\frac{1}{2}$ and massless right-handed anti-particle of helicity of $+\frac{1}{2}$. The theory has a global $U(1)$ symmetry under which $\psi \rightarrow e^{i\alpha}\psi$; if the left-handed fermion has charge $+1$ then the right-handed fermion has charge -1 , as befits an anti-particle.

We can add a mass term for a single Weyl fermion. This is known as a *Majorana mass*,

$$\mathcal{S}_{\text{Maj}} = \int d^4x \, \frac{m}{2}\psi\psi + \frac{m^*}{2}\bar{\psi}\bar{\psi} \quad (2.15)$$

In general, we can take $m \in \mathbb{C}$ although any complex phase of m can be absorbed into ψ and, upon quantisation, the resulting particle has mass $|m|$. Importantly, the Majorana mass explicitly breaks the global $U(1)$ symmetry, so there is no quantum number to distinguish particle from anti-particle. Upon quantisation, the theory consists of a single massive spin $\frac{1}{2}$ particle that is now its own anti-particle.

Because the Majorana mass term explicitly breaks the $U(1)$ symmetry, it is not allowed if the $U(1)$ is gauged. Relatedly, it's not possible to write down such a term for any fermion ψ that transforms in a complex representation of a gauge group. It is, however, possible to write down such terms for fermions in real representations.

Recovering Dirac Spinors

All this discussion of spinors and, so far, not a gamma matrix or Clifford algebra in sight! Yet these played a central role in the discussion of spinors that we met in the [Quantum Field Theory](#) lectures. What's going on?

The Dirac spinor is *not* an irreducible representation of the Lorentz group in $d = 3+1$ dimensions. Instead, it consists of independent left- and right-handed spinors. In our earlier notation:

$$\left(\frac{1}{2}, 0\right) \oplus \left(0, \frac{1}{2}\right) : \text{ Dirac spinor}$$

We write a Dirac spinor as a 4-component object, consisting of a left-handed Weyl fermion ψ_α and a right-handed Weyl fermion $\bar{\chi}^{\dot{\alpha}}$ (note the index up),

$$\Psi = \begin{pmatrix} \psi_\alpha \\ \bar{\chi}^{\dot{\alpha}} \end{pmatrix}$$

We also introduce the chiral basis of gamma matrices

$$\gamma^\mu = \begin{pmatrix} 0 & \sigma^\mu \\ \bar{\sigma}^\mu & 0 \end{pmatrix} \tag{2.16}$$

These obey the Clifford algebra $\{\gamma^\mu, \gamma^\nu\} = 2\eta^{\mu\nu}$. In the [Quantum Field Theory](#) lectures, we showed that the generators of Lorentz transformations for a Dirac spinor are

$$S^{\mu\nu} = \frac{i}{4}[\gamma^\mu, \gamma^\nu] = \begin{pmatrix} \sigma^{\mu\nu} & 0 \\ 0 & \bar{\sigma}^{\mu\nu} \end{pmatrix}$$

(As with our earlier definition of $M^{\mu\nu}$, this differs by a factor of i from the conventions in the [Quantum Field Theory](#) lectures.) Under a Lorentz transformation, a Dirac spinor transforms as $\Psi \rightarrow \exp(-\frac{i}{2}\omega_{\mu\nu}S^{\mu\nu})\Psi$. This reproduces the transformations of Weyl spinors that we saw in [\(2.12\)](#) and [\(2.13\)](#).

The Dirac action that we met in our [Quantum Field Theory](#) lectures is

$$S_{\text{Dirac}} = - \int d^4x \ i \bar{\Psi} \gamma^\mu \partial_\mu \Psi - M \bar{\Psi} \Psi$$

where, for a Dirac spinor (but not a Weyl spinor!) the bar notation means $\bar{\Psi} = \Psi^\dagger \gamma^0$. Decomposed in terms of Weyl fermions, it becomes

$$S_{\text{Dirac}} = - \int d^4x \ i \bar{\psi} \bar{\sigma}^\mu \partial_\mu \psi + i \chi \sigma^\mu \partial_\mu \bar{\chi} - M(\chi \psi + \bar{\psi} \bar{\chi}) \quad (2.17)$$

The first term coincides with the kinetic term (2.14) for a left-handed fermion. The second term is simply a different way of writing this, with the derivative now acting on a right-handed fermion; if you play around lowering and raising indices then the second term can be massaged to look like the first.

The mass term in (2.17) is *not* of the Majorana type (2.15). First, the mass is necessarily real, $M \in \mathbb{R}$, although it can be positive or negative. Second, because the mass term involves two distinct Weyl fermions it preserves a $U(1)$ symmetry, under which the phase of ψ and χ rotate oppositely. The result is that, upon quantisation, the action (2.17) gives a particle of spin $+\frac{1}{2}$ and charge $+1$, together with a distinct anti-particle of spin $+\frac{1}{2}$ and charge -1 , both with mass $|M|$.

It is possible to restrict the Dirac fermion Ψ to have the same content as a single Weyl fermion. In a general basis of gamma matrices, we do this by introducing a charge conjugation matrix. But in the chiral basis (2.16), it's particularly simple: we just restrict $\bar{\chi} = \bar{\psi} \equiv \psi^\dagger$. A Dirac spinor with such a restriction is called a *Majorana spinor*.

Throughout these lectures, we will have no need to resort to 4-component spinors. We will write everything in terms of 2-component Weyl fermions.

2.1.3 The Poincaré Group and its Extensions

The continuous symmetries of Minkowski space comprise of Lorentz transformations together with spacetime translations. Combined, these form the *Poincaré group*. Spacetime translations are generated, as usual, by the momentum 4-vector P^μ . Their commutation relations with themselves and with the Lorentz generators $M^{\mu\nu}$ are given by

$$[P^\mu, P^\nu] = 0 \quad \text{and} \quad [M^{\mu\nu}, P^\sigma] = i (P^\mu \eta^{\nu\sigma} - P^\nu \eta^{\mu\sigma}) \quad (2.18)$$

The latter of these is equivalent to the statement that P^μ transforms as a 4-vector under Lorentz transformations. These commutation relations should be considered in conjunction with the Lorentz algebra (2.3),

$$[M^{\mu\nu}, M^{\rho\sigma}] = i(\eta^{\nu\rho} M^{\mu\sigma} - \eta^{\nu\sigma} M^{\mu\rho} + \eta^{\mu\sigma} M^{\nu\rho} - \eta^{\mu\rho} M^{\nu\sigma}) \quad (2.19)$$

Together, (2.18) and (2.19) form the algebra of the Poincaré group.

It's not unusual for quantum field theories to exhibit further continuous symmetries. Say, a global $U(1)$ symmetry that rotates the phase of a complex field, or perhaps a non-Abelian $SU(N)$ symmetry under which a multiplet of fields transforms. The generators of these symmetries – which we'll denote collectively as T – correspond to some conserved charge or isospin and are always Lorentz scalars. This means that they necessarily commute with the Poincaré generators,

$$[P^\mu, T] = [M^{\mu\nu}, T] = 0$$

One could ask: is it possible for something less trivial to happen, with the new generators transforming in some interesting fashion under the Poincaré group? For example, this would happen if the additional generators T themselves carried some spacetime index. If this were possible, the Poincaré group would be subsumed into a larger group. And that sounds interesting.

A theorem due to Coleman and Mandula greatly restricts this possibility. Roughly speaking, the theorem states that, in any spacetime dimension greater than $d = 1 + 1$, the symmetry group of any interacting quantum field theory must factorise as

$$\text{Poincaré} \times \text{Internal} \quad (2.20)$$

We won't prove the Coleman-Mandula theorem here¹. The gist of the proof is that Poincaré invariance already greatly restricts what can happen in, say, 2 to 2 scattering, with only the scattering angle left undetermined. Any internal symmetries that factorise, as in (2.20), put restrictions on the kinds of interactions that are allowed, for example enforcing conservation of electric charge. But if the generators T were to carry a spacetime index then they would put further constraints on the scattering angle itself and that would be overly restrictive, at best allowing scattering to occur only at discrete angles. But if one assumes that the scattering amplitudes are analytic functions of the angle then the amplitude must vanish for all angles and the theory is free.

¹The original Coleman-Mandula paper is from 1967 and entitled “[All Possible Symmetries of the S-matrix](#)”. Witten's “[Introduction to Supersymmetry](#)” lectures give a clear intuitive explanation of the theorem. A full proof can be found Weinberg vol III.

Like all no-go theorems in physics, the Coleman-Mandula theorem comes with a number of underlying assumptions. Some of these are eminently reasonable, such as locality and causality. But it may be possible to relax other assumptions to find interesting loopholes to the Coleman-Mandula theorem. Two such loopholes have proven to be extremely important.

- **Conformal Invariance:** The Coleman-Mandula theorem assumes that the theory has a mass gap, meaning that all particles are massive. Indeed, it studies symmetries of the S-matrix which is really only well defined for massive particles where we don't have to worry about IR divergences. For theories of massless particles something interesting can, and often does, happen.

The first interesting thing is that interacting massless theories typically exhibit scale invariance. This means that physics is unchanged under the symmetry $x^\mu \rightarrow \lambda x^\mu$. The associated symmetry generator is called D for “dilatation”. This can only be a symmetry of a theory that has no dimensionful parameters. In particular, no masses.

The second interesting thing is more surprising. For reasons that are not entirely understood, theories that exhibit scale invariance also exhibit a further symmetry known as *special conformal transformations* of the form

$$x^\mu \rightarrow \frac{x^\mu - a^\mu x^2}{1 - 2a \cdot x + a^2 x^2}$$

This transformation depends on a vector parameter a^μ and the associated generator is a 4-vector K^μ . The resulting conformal algebra extends the Poincaré algebra (2.18) and (2.19) with the non-trivial commutators

$$\begin{aligned} [D, K^\mu] &= -iK^\mu \quad , \quad [D, P^\mu] = iP^\mu \\ [K^\mu, P^\nu] &= 2i(D\eta^{\mu\nu} - M^{\mu\nu}) \\ [M^{\mu\nu}, K^\sigma] &= i(K^\nu\eta^{\mu\sigma} - K^\mu\eta^{\nu\sigma}) \end{aligned}$$

Interacting conformal field theories crop up in many places in physics. In their Euclidean incarnation, they describe critical points, or second order phase transitions, that were the focus of our lectures on [Statistical Field Theory](#). In $d = 1 + 1$ dimensions the conformal group has rather more structure and a detailed introduction can be found in the lectures on [String Theory](#). We'll meet examples of supersymmetric conformal field theories later in Section 6.4 when we discuss the low-energy physics of certain gauge theories.

- **Supersymmetry:** The second loophole to the Coleman-Mandula theorem is supersymmetry. As you may by now have guessed, exploiting this loophole will be the topic of the rest of these lectures.

2.2 The Supersymmetry Algebra

Supersymmetry evades the Coleman-Mandula no-go theorem because it is a different kind of symmetry. In contrast to the symmetries discussed above, it is not characterised by a Lie algebra. Instead it is characterised by a mathematical structure known as a $\mathbf{Z}_2$ -graded Lie algebra. For our purposes, this simply means that the algebra contains both commutation and anti-commutation relations.

A generalisation of the Coleman-Mandula theorem to graded Lie algebras was given by Haag, Lopuszanski and Sohnius. Roughly speaking, it says that the only possibility is supersymmetry. We will now, finally, explain what this means.

Supersymmetric theories have a new conserved charge that is a left-handed Weyl spinor Q_α , together with its right-handed counterpart $\bar{Q}_{\dot{\alpha}}$. This is known as the *supercharge*. It is possible to have multiple supercharges, a situation known as extended supersymmetry. We will discuss this in Section 2.4 and, for now, stick to just a single complex supercharge. This is known as $\mathcal{N} = 1$ supersymmetry.

At the heart of the supersymmetry algebra is the anti-commutation relation

$$\{Q_\alpha, \bar{Q}_{\dot{\alpha}}\} = 2\sigma_{\alpha\dot{\alpha}}^\mu P_\mu \quad (2.21)$$

It is no surprise that a spinor should have an anti-commutator. But the structure of this relation is interesting: it tells us that the supercharges should be viewed as the square-root of spacetime translations! Our goal in these lectures is to understand what, exactly, this means.

The full supersymmetry algebra comprises of commutation relations (2.18) and (2.19) of the Poincaré group, which remain unchanged, together with the (anti)-commutation relations of the supercharges. The first of these is

$$[M^{\mu\nu}, Q_\alpha] = (\sigma^{\mu\nu})_\alpha{}^\beta Q_\beta \quad \text{and} \quad [M^{\mu\nu}, \bar{Q}^{\dot{\alpha}}] = (\bar{\sigma}^{\mu\nu})^{\dot{\alpha}}{}_{\dot{\beta}} \bar{Q}^{\dot{\beta}} \quad (2.22)$$

This is simply the statement that the supercharges transform under a Lorentz transformation in the manner expected of operators that are Weyl fermions. To see this, first recall from (2.12) that any spinor like Q_α transforms as $Q_\alpha \rightarrow U_\alpha{}^\beta Q_\beta$ where $U = \exp(-\frac{i}{2}\omega_{\mu\nu}\sigma^{\mu\nu})$. But Q_α is also an operator acting on a Hilbert space and, viewed through this lens, we get a different expression for how it transforms. Any

state in the Hilbert space transforms as $|\phi\rangle \rightarrow V|\phi\rangle$ with $V = \exp(-\frac{i}{2}\omega_{\mu\nu}M^{\mu\nu})$. Here, $M^{\mu\nu}$ is the abstract generator of Lorentz transformations and its action on any state depends on the quantum number of that state. Correspondingly, operators $\mathcal{O}$ transform as $\mathcal{O} \rightarrow V\mathcal{O}V^\dagger$ since this ensures that the matrix elements $\langle\phi'|\mathcal{O}|\phi\rangle$ remains unchanged. Equating these two ways in which the supercharge transforms, we have $VQ_\alpha V^\dagger = (UQ)_\alpha$. The algebra (2.22) is the infinitesimal version of this transformation law.

The remaining commutation relations are somewhat less interesting, although no less important

$$[Q_\alpha, P^\mu] = \{Q_\alpha, Q_\beta\} = 0 \quad (2.23)$$

There are, however, reasons why these commutators take this boring form.

First, why do we necessarily have $[Q_\alpha, P^\mu] = 0$? Clearly the right-hand side should be something with α and μ indices so that the commutator is covariant under Lorentz transformations. But that leaves the option for $[Q_\alpha, P^\mu] = c(\sigma^\mu)_{\alpha\dot{\alpha}}\bar{Q}^{\dot{\alpha}}$ for some $c \in \mathbf{C}$. What forces us to have $c = 0$?

The answer to this lies in the Jacobi identity

$$[P^\mu, [P^\nu, Q_\alpha]] + [P^\nu, [Q_\alpha, P^\mu]] + [Q_\alpha, [P^\mu, P^\nu]] = 0$$

Clearly the last term vanishes, as $[P^\mu, P^\nu] = 0$. If we choose $[Q_\alpha, P^\mu] = c(\sigma^\mu)_{\alpha\dot{\alpha}}\bar{Q}^{\dot{\alpha}}$ and, correspondingly, $[\bar{Q}^{\dot{\alpha}}, P^\mu] = c^*(\bar{\sigma}^\mu)^{\dot{\alpha}\beta}Q_\beta$ then the Jacobi identity becomes

$$-c\sigma_{\alpha\dot{\alpha}}^\nu[P^\mu, \bar{Q}^{\dot{\alpha}}] + c\sigma_{\alpha\dot{\alpha}}^\mu[P^\nu, \bar{Q}^{\dot{\alpha}}] = |c|^2(\sigma^\nu\bar{\sigma}^\mu - \sigma^\mu\bar{\sigma}^\nu)_\alpha{}^\beta Q_\beta = 0$$

This requires $c = 0$.

There is a similar reason for why we must have $\{Q_\alpha, Q_\beta\} = 0$. Once again, there is an alternative since if we just try to pair up indices then we might think that $\{Q_\alpha, Q_\beta\} = c'(\sigma^{\mu\nu})_\alpha{}^\beta M_{\mu\nu}$ would be acceptable for any $c' \in \mathbf{R}$. But if we take the commutator with P^ρ then, from the argument above, the left-hand-side must vanish which, because $[P^\rho, M^{\mu\nu}] \neq 0$, tells us that $c' = 0$.

(An aside: there's actually a subtlety in this last discussion. While it is true that $\{Q_\alpha, Q_\beta\} = 0$ when sandwiched between any finite energy states, some supersymmetric theories have multiple ground states and it turns out that $\{Q_\alpha, Q_\beta\}$ can be non-vanishing when evaluated on the infinite energy domain walls that interpolate between these ground states. This subtlety is interesting, at least if you care about domain walls, but somewhat beyond the scope of these lectures.)

2.2.1 R-Symmetry

We started this section by noting that all internal symmetries must commute with the spacetime symmetries of the Poincaré group. But must they also commute with the supercharge Q_α ? The answer is: almost.

All internal symmetries must commute with Q_α with one exception: it may be that theories admit an internal $U(1)$ symmetry that acts as

$$Q_\alpha \rightarrow e^{-i\lambda} Q_\alpha \quad \text{and} \quad \bar{Q}_{\dot{\alpha}} \rightarrow e^{i\lambda} \bar{Q}_{\dot{\alpha}} \quad (2.24)$$

This $U(1)$ symmetry is known as an *R-symmetry* and is sometimes denoted $U(1)_R$. If we denote the generator as R then it has commutation relations

$$[R, Q_\alpha] = -Q_\alpha \quad \text{and} \quad [R, \bar{Q}_{\dot{\alpha}}] = +\bar{Q}_{\dot{\alpha}} \quad (2.25)$$

When we turn to theories of extended supersymmetry in Section 2.4, we'll see different R-symmetry groups arising. But for theories with $\mathcal{N} = 1$ symmetry we have only $U(1)_R$. Nonetheless, this will play an important role when we come to analyse the dynamics of supersymmetric theories in later sections. We'll see this, for example, in Section 3.3.

This, then, is the supersymmetry algebra: it comprises of the algebra of the Poincaré group (2.18) and (2.19), together with the algebra of the supercharges (2.21), (2.22) and (2.23) and, finally, the R-symmetry (2.25). The next question is: what can we do with it?

2.2.2 A Consequence: Energy is Positive

Even before we write down any field theories, we can derive one feature of supersymmetric theories from the algebra alone. This follows from the key algebraic relation (2.21),

$$\{Q_\alpha, \bar{Q}_{\dot{\alpha}}\} = 2\sigma_{\alpha\dot{\alpha}}^\mu P_\mu \quad (2.26)$$

If we compute the expectation of the left-hand side in any state $|\phi\rangle$ then we find that it is necessarily positive

$$\langle\phi|Q_\alpha\bar{Q}_{\dot{\alpha}} + \bar{Q}_{\dot{\alpha}}Q_\alpha|\phi\rangle = |(Q_\alpha)^\dagger|\phi\rangle|^2 + |Q_\alpha|\phi\rangle|^2 \geq 0 \quad (2.27)$$

The same must be true of the right-hand side

$$\sigma_{\alpha\dot{\alpha}}^\mu \langle\phi|P_\mu|\phi\rangle \geq 0$$

If we set $\alpha = \dot{\alpha}$ and sum over $\alpha = 1, 2$ then we make use of the fact that $\text{tr } \sigma^0 = 2$ and $\text{tr } \sigma^i = 0$. This then reduces to the statement that the energy of any state in a supersymmetric theory is necessarily positive

$$\langle \phi | P_0 | \phi \rangle \geq 0$$

This is curious. Usually in physics, we don't care about the overall value of the energy: if you add an overall constant to all energies, then physics remains unchanged. There are two places where this state of affairs no longer holds. The first is in gravity where the energy of the vacuum contributes as a cosmological constant. The second is, as we've seen above, in supersymmetric theories where energies are necessarily positive definite.

Physically, it's far from clear if there is any deep relation between these two ideas. In fact, as we will see later in these lectures, the energy of the ground state acts as an order parameter for the breaking of supersymmetry. This means that the ground state energy is zero if supersymmetry is exact, otherwise it is non-zero. In our world, it's clear that there is no supersymmetry visible at the TeV scale, while the cosmological constant is many orders of magnitude smaller, at 10^{-3} eV. This makes it difficult to see how supersymmetry can help alleviate the [cosmological constant problem](#).

However, at the formal mathematical level, the relationship between supersymmetry and gravity has proven rather useful. For example, there exists a greatly simplified proof of the positive energy theorem in general relativity, due to Witten, that uses ideas of supersymmetry.

There is one further piece of physics hiding in (2.26). For any other symmetry in field theory, we can think about gauging it. This means that we try to construct theories in which the symmetry is realised locally. Supersymmetry is no different. One can construct theories in which the associated infinitesimal parameter for supersymmetry transformations depends on x^μ . From (2.26), we see that such theories necessarily enjoy a symmetry in which you do different translations at different points in space. But such transformations are diffeomorphisms and are the characteristic feature of general relativity. In other words, theories of local supersymmetry are necessarily theories of gravity! Such theories are known as *supergravity*, usually shortened to the ugly acronym “sugra”. We will mention supergravity only very briefly in this section. In subsequent sections our interest will be entirely on theories with global supersymmetry.

2.3 Representations on Particle States

Given an algebra, our next task is to explore its representations. There are different ways that we could approach this. Ultimately, we will be interested in quantum field theories that enjoy supersymmetry and this means understanding the way supersymmetry acts on fields. This we will do in later sections. Here, to build some intuition, we will understand how supersymmetry acts on single particle states in the Hilbert space.

Without doing any work, we can guess that something interesting is going on. The supercharge Q_α is a fermionic operator, both in the sense that it carries spin $\frac{1}{2}$ and in the sense that it is naturally anti-commuting as in (2.21). This means that, schematically, we must have

$$Q|\text{fermion}\rangle = |\text{boson}\rangle \quad \text{and} \quad Q|\text{boson}\rangle = |\text{fermion}\rangle \quad (2.28)$$

This is the defining feature of supersymmetry.

In fact, it is straightforward to show that any representation of the supersymmetry algebra must have an equal number of bosonic and fermionic states. To this end, we introduce the *fermionic number operator* $(-1)^F$. This acts on bosonic states as

$$(-1)^F|B\rangle = |B\rangle \quad \text{and} \quad (-1)^F|F\rangle = -|F\rangle$$

Because Q_α swaps a bosonic state for a fermionic state, we necessarily have

$$(-1)^F Q_\alpha = -Q_\alpha (-1)^F \quad \Rightarrow \quad \{(-1)^F, Q_\alpha\} = 0$$

The result that we now want follows straightforwardly from the algebra $\{Q_\alpha, \bar{Q}_{\dot{\alpha}}\} = 2\sigma_{\alpha\dot{\alpha}}^\mu P_\mu$. Suppose that we have a finite collection of one-particle states that form a representation of the supersymmetry algebra. We can take the following trace over elements of this multiplet

$$\begin{aligned} \text{tr} [(-1)^F \{Q_\alpha, \bar{Q}_{\dot{\alpha}}\}] &= \text{tr} [(-1)^F Q_\alpha \bar{Q}_{\dot{\alpha}} + (-1)^F \bar{Q}_{\dot{\alpha}} Q_\alpha] \\ &= \text{tr} [-Q_\alpha (-1)^F \bar{Q}_{\dot{\alpha}} + (-1)^F \bar{Q}_{\dot{\alpha}} Q_\alpha] = 0 \end{aligned}$$

Here the second equality we've used the fact that $\{(-1)^F, Q_\alpha\} = 0$ while the final equality uses the cyclicity of the trace. The supersymmetry algebra then tells us that

$$\sigma_{\alpha\dot{\alpha}}^\mu \text{tr} [(-1)^F P_\mu] = 0$$

Note that $\sigma_{\alpha\dot{\alpha}}^\mu$ sits outside the trace over states: it's just a bunch of numbers as far as the trace is concerned. Meanwhile P_μ sits inside the trace because it is an operator

acting on states. We can choose these states to be momentum eigenstates, so that $P_\mu|\text{any state}\rangle = p_\mu|\text{any state}\rangle$. We then simply have

$$\sigma_{\alpha\dot{\alpha}}^\mu p_\mu \operatorname{tr}(-1)^F = 0$$

But $\operatorname{tr}(-1)^F$ simply counts the number of bosonic states n_B minus the number of fermionic states n_F ,

$$\operatorname{tr}(-1)^F = n_B - n_F = 0$$

The number of such states must be equal. The quantity $\operatorname{tr}(-1)^F$ is called the *Witten index*.

There's actually a loophole in the discussion above. It may be that Q_α and $\bar{Q}_{\dot{\alpha}}$ annihilate states in the supersymmetry multiplet. From the supersymmetry algebra (and the positivity conditions (2.27) that follows from it) this can only happen for states of zero energy which are necessarily the ground states of the system. This means that there may be a mismatch between the number of bosonic and fermionic ground states of a system. It is in studying such ground states that the *Witten index* really finds its teeth and we'll revisit this in Section 3.4.2. More sophisticated examples can be found in the lectures on [Supersymmetric Quantum Mechanics](#).

We now know that supersymmetry requires an equal number of bosonic and fermionic states. The next step is to understand exactly what kind of fermion is paired with what kind of boson.

2.3.1 Representations of the Poincaré Group

To set the scene, let's first recall how we construct the irreducible representations of the Poincaré group. In fact, let's start even more simply: how do we construct irreducible representations of the rotation group?

We work with the algebra $so(3) \cong su(2)$ rather than the group. This is, of course, defined by the familiar commutation relations

$$[J_i, J_j] = i\epsilon_{ijk}J_k$$

To construct representations, the first thing we do is look to the Casimirs. These are operators that commute with all generators of the group. For $su(2)$, there is just a single Casimir,

$$C = \sum_{i=1}^3 J_i^2$$

Irreducible representations are labelled by their eigenvalue of the Casimir. For $su(2)$, the eigenvalue of J^2 is $j(j+1)$ with the spin j taking values in $j = 0, \frac{1}{2}, 1, \dots$. Each representation has dimension $2j+1$, with the states within a multiplet identified by their eigenvalue under, say, J_3 whose eigenvalue lies in $|j_3| \leq j$. The result is the familiar one from quantum mechanics: states are labelled by $|j, j_3\rangle$.

Now let's turn to the Poincaré group. The irreducible representations are what we call “particles”. Again, they are characterised by the Casimirs. I won't tell you how to construct Casimirs, but will instead just present you the result: the Poincaré group has two Casimirs, given by

$$C_1 = P_\mu P^\mu \quad \text{and} \quad C_2 = W_\mu W^\mu$$

Here $W^\mu = \frac{1}{2}\epsilon^{\mu\nu\rho\sigma}P_\nu M_{\rho\sigma}$ is the *Pauli-Lubański vector*. It can be thought of as a relativistic version of angular momentum.

Representations of the Poincaré group are then labelled by the eigenvalues of C_1 and C_2 . The first of these is simply the mass m of a particle: $C_1 = m^2$. What happens next is a little different depending on whether the particles are massive or massless.

- Massive Particles: In this case, we can always boost to the rest frame of the particle so that $P^\mu = (m, 0, 0, 0)$. In this frame, the Pauli-Lubański vector is

$$W^0 = 0 \quad \text{and} \quad W^i = -mJ^i$$

with J^i the generators of rotations. This means that $C_2 = -m^2J^2$ and so is specified by the eigenvalue of J^2 . We find the familiar fact that massive particles are characterised by their mass m and spin j .

- Massless Particles: Now $C_1 = m^2 = 0$. There are some subtleties that we sweep under the rug here, but it turns out that the most interesting representations also have $C_2 = W^2 = 0$, so both Casimirs vanish. To characterise the representation, we choose a frame such that, say, $P^\mu = (E, 0, 0, E)$. There, we have $W^\mu = M_{12}P^\mu$, so the constant of proportionality between W and P is determined by the eigenvalue of the $U(1)$ rotation in the (x^1, x^2) -plane. The eigenvalue of this rotation is the *helicity*, $h = 0, \frac{1}{2}, 1, \dots$. We learn that massless particles are characterised by (obviously) $m = 0$ and their helicity h .

Although the results are different for $m = 0$ and $m \neq 0$, the strategy is the same. In each case, we boost to a preferred frame of the particle which is then characterised by how it transforms under the surviving symmetry group. This surviving symmetry — $SU(2)$ for a massive particle, $U(1)$ for a massless one — is called the *little group*.

There is a slight twist to the story when it comes to realising these representations on the Hilbert space of single particle states. For massive particles, the states take the form

$$|p_\mu; j, j_3\rangle \tag{2.29}$$

where the momentum is restricted to obey $p_\mu p^\mu = m^2$ while the azimuthal angular momentum takes values in $j_3 \leq |j|$. This fills out the $2j + 1$ dimensional set of spin sets. However, for massless particles, there is just a single state $|p_\mu; h\rangle$. This is because the helicity describes the representation of the Abelian group $U(1)$ generated by M^{12} rather than the non-Abelian group $SU(2)$ and irreducible representations of Abelian groups are one-dimensional.

The problem is that we know that massless particles also have internal degrees of freedom. For example, the photon necessarily has two polarisation states. Clearly we're missing something. What we're missing is the additional requirement that the spectrum of states is invariant under CPT. For massive particles, this doesn't buy us anything new: the set of states (2.29) is already invariant under CPT. However, for massless particles CPT flips $h \mapsto -h$ and tells us that massless states must come in pairs

$$|p_\mu; h\rangle \quad \text{and} \quad |p_\mu, -h\rangle$$

This is the origin of the two polarisation states of the photon or graviton, or the two helicities of a massless Weyl spinor. Note that a massless scalar has helicity $h = 0$ and so is CPT self-conjugate. This means that there's no requirement from CPT to add an additional degree of freedom in this case.

2.3.2 Massless Representations

We now turn to the representations of the $\mathcal{N} = 1$ supersymmetry algebra. The simple observation (2.28) tells us that we should expect representations to contain particles of different spin and this will turn out to be true. Once again we need to treat massless and massive particles separately.

The supersymmetry algebra also has two Casimirs. The first is familiar:

$$C_1 = P_\mu P^\mu$$

The fact that this is a Casimir tells us that all particles in a supersymmetric multiplet must have the same mass, $C_1 = m^2$.

In contrast, the other Casimir of the Poincaré group, $W_\mu W^\mu$, is *not* a Casimir of the supersymmetry algebra. This is because $[W_\mu, Q_\alpha] \neq 0$ which, in turn, can be traced to the commutation relation $[M_{\mu\nu}, Q_\alpha] \neq 0$. But it was $W_\mu W^\mu$ that told us that representations of the Poincaré group are characterised by the spin of a particle. The fact that $W_\mu W^\mu$ is no longer a Casimir means that representations of the supersymmetry algebra can contain particles of different spin.

It is possible to construct a new Casimir. First define

$$Y_\mu = W_\mu - \frac{1}{4} \bar{Q}_{\dot{\alpha}} \bar{\sigma}_\mu^{\dot{\alpha}\beta} Q_\beta$$

Then the second Casimir of the supersymmetry algebra turns out to be

$$\tilde{C}_2 = (Y_\mu P_\nu - Y_\nu P_\mu)(Y^\mu P^\nu - Y^\nu P^\mu)$$

However, in what follows we won't need this result. Instead we will build up a representation of the supersymmetry algebra more directly. Our strategy is to start from a particle (i.e. a representation of the Poincaré group) and then act on it with successive supersymmetry generators until we build up a representation of the full algebra.

It turns out that things are slightly simpler for massless representations. Consider a state $|p_\mu, h\rangle$ of a massless particle of helicity h . We can again boost to a frame in which $p_\mu = (E, 0, 0, E)$. Restricted to act on such states, the supersymmetry algebra becomes

$$\{Q_\alpha, \bar{Q}_{\dot{\alpha}}\} = 2\sigma_{\alpha\dot{\alpha}}^\mu P_\mu = 2E(1 + \sigma^3)_{\alpha\dot{\alpha}} = 4E \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix}$$

From the positivity condition (2.27), we see that Q_2 and $\bar{Q}_2$ necessarily annihilate this state,

$$\langle p_\mu, h | \{Q_2, \bar{Q}_2\} | p_\mu, h \rangle = 0 \quad \Rightarrow \quad Q_2 | p_\mu, h \rangle = \bar{Q}_2 | p_\mu, h \rangle = 0$$

To build a representation of the full supersymmetry algebra, we only need consider the action of Q_1 and $\bar{Q}_1$. But these act just like fermionic creation and annihilation operators. Specifically, if we rescale the operators to become

$$a = \frac{Q_1}{\sqrt{4E}} \quad \text{and} \quad a^\dagger = \frac{\bar{Q}_1}{\sqrt{4E}} \quad \Rightarrow \quad \{a, a^\dagger\} = 1 \quad \text{and} \quad \{a, a\} = \{a^\dagger, a^\dagger\} = 0$$

The representations of this algebra are straightforward: they consist of two states $|0\rangle$ and $|1\rangle$ such that $a|0\rangle = 0$ and $|1\rangle = a^\dagger|0\rangle$. This ensures that $a^\dagger|1\rangle = 0$. For us, this means that we can start by taking a state which, by assumption, is annihilated by a ,

$$a|p_\mu, h\rangle = 0$$

The full supersymmetry multiplet then consists of $|p_\mu, h\rangle$ and $a^\dagger|p_\mu, h\rangle$. The question is: what is the helicity of this second state? This follows from the commutation relation (2.22)

$$[M^{\mu\nu}, Q_\alpha] = (\sigma^{\mu\nu})_\alpha{}^\beta Q_\beta \quad \text{and} \quad [M^{\mu\nu}, \bar{Q}^{\dot{\alpha}}] = (\bar{\sigma}^{\mu\nu})^{\dot{\alpha}}{}_{\dot{\beta}} \bar{Q}^{\dot{\beta}} \quad (2.30)$$

Restricting to rotations in the (x^1, x^2) plane, which is what we mean by helicity, we have

$$\begin{aligned} [M^{12}, Q_1] &= \frac{1}{2} Q_1 & \text{and} & & [M^{12}, Q_2] &= -\frac{1}{2} Q_2 \\ [M^{12}, \bar{Q}^1] &= \frac{1}{2} \bar{Q}^1 & \text{and} & & [M^{12}, \bar{Q}^2] &= -\frac{1}{2} \bar{Q}^2 \end{aligned}$$

The first equation tells us that Q_1 raises the helicity by $\frac{1}{2}$. This suggests that the adjoint $\bar{Q}_1$ lowers the helicity by $\frac{1}{2}$. To see that this is the case, we need to remember that, after lowering an index, $\bar{Q}_1 = -\bar{Q}^2$ so we have

$$[M^{12}, \bar{Q}_1] = -\frac{1}{2} \bar{Q}_1$$

So $\bar{Q}_1$ does indeed lower the helicity by $\frac{1}{2}$ as anticipated. We learn that the massless representations of the supersymmetry algebra consist of just two states:

$$|p_\mu, h\rangle \quad \text{and} \quad |p_\mu, h - \frac{1}{2}\rangle = \frac{\bar{Q}_1}{\sqrt{4E}} |p_\mu, h\rangle$$

As we saw above, for massless states we must also add their CPT conjugates. The different representations of the supersymmetry algebra then arise by picking different starting helicities h . There are three representations that are most important:

- If we start with $h = \frac{1}{2}$ then we have

h	$-\frac{1}{2}$	0	$+\frac{1}{2}$
multiplicity	1	2	1

This is the matter content that we get from quantising a single Weyl spinor together with a *complex* scalar. This is known as a *chiral multiplet*.

The chiral multiplets should be thought of as matter particles. We will devote Section 3 to studying field theories associated to chiral multiplets. Here we make a quick comment. The fact that any other internal symmetry generator must commute with Q_α means that the fermion and scalar in a given chiral multiplet must experience the same force. In particular, if one is charged under a gauge group then so is the other. We'll see this explicitly when we construct supersymmetry gauge theories in Section 4.

- If we start with $h = 1$ then we have

h	-1	$-\frac{1}{2}$	$+\frac{1}{2}$	$+1$
multiplicity	1	1	1	1

This is the matter content of a photon together with a single Weyl spinor. It is known as the *gauge multiplet* or *vector multiplet*.

We will devote Section 4 to the study of vector multiplets. There we will see that we can construct supersymmetric versions of Yang-Mills theory with gauge group G by taking $\dim G$ vector multiplets. As usual, the $h = 1$ gauge bosons transform in the adjoint of the gauge group. But now, so too, must its fermionic supersymmetric partner. In this context, the fermion is called a *gaugino*.

- If we start with $h = 2$ then we have

h	-2	$-\frac{3}{2}$	$+\frac{3}{2}$	$+2$
multiplicity	1	1	1	1

This is the matter content of a graviton together with a helicity $\frac{3}{2}$ spinor, sometimes known as a *Rarita-Schwinger* field or, in this context, the *gravitino*. They combine to form the *supergravity multiplet*.

If we keep going, we get massless fields with helicity $h > 2$. But there are strong restrictions that prohibit the existence of interacting theories with massless fields of such high helicity. (This statement is true in Minkowski spacetimes; there are remarkable "higher spin" theories that include an infinite tower of massless states in de Sitter or anti de Sitter spacetimes.) We also skipped the $h = \frac{3}{2}$ multiplet for similar reasons; it turns out that the existence of a massless helicity $\frac{3}{2}$ particle implies the existence of a local supersymmetry which, in turn, requires that the theory is coupled to gravity.

2.3.3 Massive Representations

We next turn to massive representations of the supersymmetry algebra. In the rest frame of a particle we have $p_\mu = (m, 0, 0, 0)$. Acting on such states, the supersymmetry algebra becomes

$$\{Q_\alpha, \bar{Q}_{\dot{\alpha}}\} = 2\sigma^\mu_{\alpha\dot{\alpha}}P_\mu = 2m\sigma^0_{\alpha\dot{\alpha}} = 2m \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} \quad (2.31)$$

This time, after rescaling, both Q_1 and Q_2 act as fermionic creation/annihilation operators

$$a_\alpha = \frac{Q_\alpha}{\sqrt{2m}} \quad \text{and} \quad a_\alpha^\dagger = \frac{\bar{Q}_{\dot{\alpha}}}{\sqrt{2m}} \quad \Rightarrow \quad \{a_\alpha, a_{\dot{\alpha}}^\dagger\} = \delta_{\alpha\dot{\alpha}}$$

with $\{a_\alpha, a_\beta\} = \{a_\alpha^\dagger, a_\beta^\dagger\} = 0$. We start with a state $|\Omega\rangle = |p_\mu; j, j_3\rangle$ that we assume to be annihilated by $a_\alpha|\Omega\rangle = 0$. Then the full supermultiplet consists of four states

$$\begin{aligned} &|\Omega\rangle \\ &a_1^\dagger|\Omega\rangle \quad \text{and} \quad a_2^\dagger|\Omega\rangle \\ &a_1^\dagger a_2^\dagger|\Omega\rangle \end{aligned}$$

Again, the question is: what is the spin of these other states. We could use the commutation relations (2.30) to understand how the new states transform under the $SU(2)$ little group but it's a little fiddly while the end result is intuitive and straightforward. The initial state $|\Omega\rangle$ has spin j . The states $a_\alpha^\dagger|\Omega\rangle$ then sit in the tensor product of representations $j \otimes \frac{1}{2} = (j + \frac{1}{2}) \oplus (j - \frac{1}{2})$. The final state can be written as $a_1^\dagger a_2^\dagger|\Omega\rangle = \frac{1}{2}\epsilon^{\alpha\beta} a_\alpha^\dagger a_\beta^\dagger|\Omega\rangle$, where the $\epsilon^{\alpha\beta}$ now contracts the creation operators to be a spin singlet. This means that the state $a_1^\dagger a_2^\dagger|\Omega\rangle$ once again has spin j .

The upshot is that a massive supermultiplet contains two particles of spin j , a particle of spin $j - \frac{1}{2}$ and a particle of spin $j + \frac{1}{2}$. Note that the degeneracy of the two particles of spin j is precisely equal to the degeneracies of the other two particles:

$$2 \times (2j + 1) = \left[2 \left(j + \frac{1}{2} \right) + 1 \right] + \left[2 \left(j - \frac{1}{2} \right) + 1 \right]$$

This is simply that statement that we saw previously: a supermultiplet must have an equal number of bosonic and fermionic degrees of freedom.

There are just two massive supermultiplets that will be of interest

- If we start with $j = 0$, we have

j	0 $\frac{1}{2}$
multiplicity	2 1

This is the matter content of a massive complex scalar with a single massive Weyl fermion. We recognise it as the same matter content as the chiral multiplet that we met previously, now of course with all particles having a mass.

- If we start with $j = \frac{1}{2}$, we have

j	0	$\frac{1}{2}$	1
multiplicity	1	2	1

In other words, we have a massive spin 1 particle, two massive Weyl fermions, and a massive spin 0 particle. This is now more states than we found in the massless gauge multiplet. In fact, this collection of states is equivalent to a massless gauge multiplet *and* a massless chiral multiplet. But that makes sense. In quantum field theory, a massless gauge boson can become massive only through the Higgs mechanism, in which the gauge boson “eats” a scalar. The supersymmetric extension of this is that a massless vector multiplet “eats” a chiral multiplet to become the massive vector multiplet described above.

There’s one further subtlety that is worth flagging up. This is how parity acts on the two scalars in the massive chiral multiplet. It turns out that one of them is a scalar and the other a pseudoscalar. Here, the meaning of a “pseudoscalar” is that it picks up a minus sign under parity. This statement follows, like everything else in this section, from the supersymmetry algebra. We denote the parity operator as $\hat{\mathcal{P}}$ to distinguish it from the momentum operator P^μ . By definition, we must have

$$\hat{\mathcal{P}}P^\mu\hat{\mathcal{P}}^{-1} = (P^0, -P^i)$$

Meanwhile, parity also exchanges left-handed and right-handed spinors. This means that parity must exchange some combination of Q_α and $\bar{Q}_{\dot{\alpha}}$. One can check that the supersymmetry algebra remains unchanged if we take

$$\hat{\mathcal{P}}Q_\alpha\hat{\mathcal{P}}^{-1} = (\sigma^0)_{\alpha\dot{\alpha}}\bar{Q}^{\dot{\alpha}} \quad \text{and} \quad \hat{\mathcal{P}}\bar{Q}^{\dot{\alpha}}\hat{\mathcal{P}}^{-1} = -(\sigma^0)^{\dot{\alpha}\alpha}Q_\alpha$$

(More generally one can include a complex phase in these relations but it will not affect our discussion here.)

Now our two scalar states in the massive chiral multiplet are $|\Omega\rangle$ and $|\Omega'\rangle = a_1^\dagger a_2^\dagger |\Omega\rangle \sim \bar{Q}_1 \bar{Q}_2 |\Omega\rangle$. They obey $Q_\alpha |\Omega\rangle = \bar{Q}_{\dot{\alpha}} |\Omega'\rangle = 0$. Since parity exchanges Q_α and $\bar{Q}_{\dot{\alpha}}$, it must also exchange $|\Omega\rangle$ and $|\Omega'\rangle$. This means that the parity eigenstates are

$$\hat{\mathcal{P}}(|\Omega\rangle \pm |\Omega'\rangle) = \pm(|\Omega\rangle \pm |\Omega'\rangle)$$

and we have one scalar (with the + sign) and one pseudoscalar (with the - sign) as advertised.

2.4 Extended Supersymmetry

It is possible for theories to exhibit more than one supersymmetry. This means that there is a collection of $\mathcal{N}$ supercharges

$$Q_\alpha^I \quad \text{and} \quad \bar{Q}_{\dot{\alpha}}^I \quad I = 1, \dots, \mathcal{N}$$

Each of these supercharges retains the same commutation relations with the generators of the Poincaré group,

$$[M^{\mu\nu} Q_\alpha^I] = (\sigma^{\mu\nu})_\alpha^\beta Q_\beta^I \quad \text{and} \quad [P^\mu, Q_\alpha^I] = 0$$

and the key part of the supersymmetry algebra holds for each generator separately

$$\{Q_\alpha^I, \bar{Q}_{\dot{\alpha}}^J\} = 2\sigma_{\alpha\dot{\alpha}}^\mu P_\mu \delta^{IJ}$$

However, there are two novelties. The first is that the anti-commutator of the supercharges with themselves can be more interesting

$$\{Q_\alpha^I, Q_\beta^J\} = \epsilon_{\alpha\beta} Z^{IJ} \quad \text{and} \quad \{\bar{Q}_{\dot{\alpha}}^I, \bar{Q}_{\dot{\beta}}^J\} = \epsilon_{\dot{\alpha}\dot{\beta}} (Z^\dagger)^{IJ} \quad (2.32)$$

Here $Z^{IJ} = -Z^{JI}$ is a *central charge*, meaning that it commutes with all other elements of the algebra. The exact nature of these central charges depends on the precise theory that we consider, but they must be constructed from other conserved quantities that are at hand. We'll see the role that these central charges play shortly.

The second novelty is the R-symmetry group. Recall that for $\mathcal{N} = 1$ we had a $U(1)_R$ symmetry (2.24) that rotates the phase of the supercharge. For $\mathcal{N} > 1$, the R-symmetry rotates the supercharges among themselves. For reasons that will become clear shortly, our primary interest will be in $\mathcal{N} = 2$ and $\mathcal{N} = 4$ supersymmetry. Here the R-symmetries are:

- $\mathcal{N} = 2$: The R-symmetry group is $U(2)_R \cong U(1)_R \times SU(2)_R$.
- $\mathcal{N} = 4$: A priori, the R-symmetry group is $U(4)$. However, it turns out that only $SU(4)$ is realised on fields. This is equivalent to $SU(4) \cong \text{Spin}(6)$. (This is sometimes written, a little inaccurately, as $SO(6)$ but the supercharges transform in the spinor representation of $\text{Spin}(6)$ which is not a representation of $SO(6) = \text{Spin}(6)/\mathbf{Z}_2$.)

Theories with extended supersymmetry are a subset of those theories with $\mathcal{N} = 1$ supersymmetry. This means that the representations of theories with $\mathcal{N} > 1$ must be constructed by joining together the $\mathcal{N} = 1$ supermultiplets that we described above. In the rest of this section, we explain how this works.

2.4.1 Massless Representations

For representations on states $|p^\mu, h\rangle$ of massless particles, we proceed as before. We boost to a frame with $p_\mu = (E, 0, 0, E)$ and restrict attention to the algebra on such states. We then have

$$\{Q_\alpha^I, \bar{Q}_{\dot{\alpha}}^J\} = 4E \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix} \delta^{IJ}$$

As previously, we have $Q_2^I |p^\mu, h\rangle = \bar{Q}_2^I |p^\mu, h\rangle = 0$. From (2.32), we then have $Z^{IJ} |p^\mu, h\rangle = 0$ which tells us that the central charges play no role for the massless states. We're left, as before, just with the Q_1^I and $\bar{Q}_1^I$ operators to deal with. These now form a collection of $\mathcal{N}$ fermionic creation and annihilation operators

$$a^I = \frac{Q_1^I}{\sqrt{4E}} \quad \text{and} \quad a^{I\dagger} = \frac{\bar{Q}_1^I}{\sqrt{4E}} \quad \Rightarrow \quad \{a^I, a^{J\dagger}\} = \delta^{IJ} \quad \text{and} \quad \{a^I, a^J\} = \{a^{I\dagger}, a^{J\dagger}\} = 0$$

We now start with some fiducial state $|\Omega\rangle = |p^\mu, h\rangle$ satisfying $a^I |\Omega\rangle = 0$ and build up the full representation by acting with successive creation operators. The end result is a collection of states

$$\begin{aligned} & |\Omega\rangle \\ & a^{I\dagger} |\Omega\rangle \\ & a^{I\dagger} a^{J\dagger} |\Omega\rangle \\ & \dots \\ & a^{1\dagger} \dots a^{\mathcal{N}\dagger} |\Omega\rangle \end{aligned}$$

Our initial state $|\Omega\rangle$ has helicity h . If we act with p of the $a^\dagger$ excitation operators then there are $\binom{\mathcal{N}}{p}$ different states, each of which has helicity $h - p/2$. The full multiplet consists of $2^\mathcal{N}$ different states. If we add the CPT conjugate states then we have $2^{\mathcal{N}+1}$ states overall. Let's now look at some specific examples.

$\mathcal{N} = 2$ Supersymmetry

Again, the different multiplets arise by considering initial states $|\Omega\rangle$ with different helicities. We'll deal with each in turn.

- If we start with $h = \frac{1}{2}$ then there are two states in the first level, $a^{I\dagger} |\Omega\rangle$, each with $h = 0$, and a single state in the final level, $a^{1\dagger} a^{2\dagger} |\Omega\rangle$, with $h = -\frac{1}{2}$. After adding the CPT conjugate we end up with

h	$-\frac{1}{2}$	0	$+\frac{1}{2}$
multiplicity	2	4	2

This is called a *hypermultiplet*. It consists of two chiral multiplets or, equivalently, two complex scalars and a Dirac fermion (i.e. two Weyl fermions).

You might wonder why we needed to add the CPT conjugate in this case. After all, starting with $h = +\frac{1}{2}$ gave a single chiral multiplet which is already CPT self-conjugate. The answer to this is buried in the details of the $SU(2)_R$ symmetry which acts on the scalars $a^I| \Omega \rangle$ as a doublet. But this means that each of these scalars must be complex and that, in turn, requires that we add the CPT conjugate.

- If we start with $h = 0$ then we get two additional states with $h = -\frac{1}{2}$ and one with $h = -1$. Adding the CPT conjugate gives

h	-1	$-\frac{1}{2}$	0	$+\frac{1}{2}$	$+1$
multiplicity	1	2	2	2	1

This is the $\mathcal{N} = 2$ vector multiplet, comprising of an $\mathcal{N} = 1$ vector multiplet and $\mathcal{N} = 1$ chiral multiplet.

- If we start with $h = 2$ then, after adding the CPT conjugate, we end up with

h	-2	$-\frac{3}{2}$	-1	$+1$	$+\frac{3}{2}$	$+2$
multiplicity	1	2	1	1	2	1

This is the $\mathcal{N} = 2$ supergravity multiplet. It comprises of an $\mathcal{N} = 1$ supergravity multiplet together with an $\mathcal{N} = 1$ vector multiplet.

There's one important feature of the spectrum above that is worth highlighting. The fermions now come in pairs, meaning that they can be viewed as Dirac fermions rather than Weyl fermions. This puts restrictions on the kind of supersymmetric theories that we can build. In particular, it's not possible to construct a *chiral gauge theory* with $\mathcal{N} > 1$ supersymmetry. Here a chiral theory is one in which left- and right-handed fermions experience different forces, like in the Standard Model. Such theories are possible with $\mathcal{N} = 1$ supersymmetry (or, indeed, $\mathcal{N} = 0$ supersymmetry as in our world!). But any extended supersymmetry forces the theories to be vector-like.

$\mathcal{N} = 4$ Supersymmetry

We can play the same game with $\mathcal{N} = 4$ supersymmetry.

- If we start with $h = 1$ then we get the following multiplet

h	-1	$-\frac{1}{2}$	0	$+\frac{1}{2}$	$+1$
multiplicity	1	4	6	4	1

This consists of an $\mathcal{N} = 2$ vector multiplet with an $\mathcal{N} = 2$ hypermultiplet and is the unique $\mathcal{N} = 4$ multiplet that does not include gravity. Note that there is now no longer a distinction between forces and matter: once you specify the gauge group, all matter content is also fixed. Furthermore, all matter fields necessarily transform in the adjoint representation of the gauge group.

For once, we did not need to add the CPT conjugate to the above multiplet: it's already CPT self-conjugate. As we saw above, it was almost possible to achieve this for the $\mathcal{N} = 2$ matter representation but we fell at the last hurdle when we considered how the $SU(2)_R$ symmetry acts on the scalars. But now we have no such concern. The scalars are the set of 6 states $a^{I\dagger}a^{J\dagger}|\Omega\rangle$ and transform in the **6** of the $SU(4)$ R-symmetry. But this is a real representation and there is no need to add the CPT conjugate.

- If we start with $h = 2$ then, after adding the CPT conjugate multiplet, we have

h	-2	$-\frac{3}{2}$	-1	$-\frac{1}{2}$	0	$+\frac{1}{2}$	$+1$	$+\frac{3}{2}$	$+2$
multiplicity	1	2	2	2	2	2	2	2	1

This is the $\mathcal{N} = 4$ supergravity multiplet, comprising of an $\mathcal{N} = 2$ supergravity multiplet and $\mathcal{N} = 2$ vector multiplet.

You may have noticed that we jumped straight from $\mathcal{N} = 2$ to $\mathcal{N} = 4$, missing out $\mathcal{N} = 3$ in the middle. If you try to build a multiplet of single particle states with $\mathcal{N} = 3$ supersymmetry starting from, say, $h = \frac{1}{2}$ or $h = 1$ then you'll find that you're obliged to add the CPT conjugate representation and you just end up with $\mathcal{N} = 4$ supersymmetry after all. This observation is the key element of a proof that says any *perturbative* theory with $\mathcal{N} = 3$ global supersymmetry necessarily has $\mathcal{N} = 4$ supersymmetry.

The word “perturbative” is important in the above statement. This means that the theory is weakly coupled and the single particle states that we’re considering here are a good approximation to the spectrum of the theory. It turns out $\mathcal{N} = 3$ supersymmetry can be realised in strongly coupled, interacting quantum field theories, with no perturbative regime.

$\mathcal{N} = 8$ Supersymmetry

If we go beyond $\mathcal{N} = 4$ supersymmetry then we no longer have multiplets with helicities $h \leq 1$. This means that we are now necessarily in the realm of local supersymmetry and supergravity. Furthermore, by the time we get beyond $\mathcal{N} = 8$ supersymmetry the multiplets have particles with helicity $h > 2$. As we mentioned before, such theories are always free in Minkowski space and therefore of limited interest. In this sense, $\mathcal{N} = 8$ is the maximum number of supersymmetries possible. The theory has a unique supergravity multiplet with the following degeneracies

h	-2	$-\frac{3}{2}$	-1	$-\frac{1}{2}$	0	$+\frac{1}{2}$	$+1$	$+\frac{3}{2}$	$+2$
multiplicity	1	8	28	56	70	56	28	8	1

$\mathcal{N} = 8$ supergravity has some interesting properties and plays a role in string theory. However, we won’t discuss it further in this course.

2.4.2 Massive Representations and BPS Bounds

Rather than repeating the whole story for massive representations, we will instead just focus on the novelty. This arises from the central charges Z^{IJ} that appear in the supersymmetry algebra

$$\{Q_\alpha^I, Q_\beta^J\} = \epsilon_{\alpha\beta} Z^{IJ}$$

For reasons that we now explain, this is where much of the power of extended supersymmetry comes from.

Our goal is to understand representations of this algebra, in conjunction with the original supersymmetry algebra which, in the rest frame of the particle, reads (2.31)

$$\{Q_\alpha^I, \bar{Q}_{\dot{\alpha}}^J\} = 2m \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} \delta^{IJ}$$

We’ll illustrate the story with $\mathcal{N} = 2$ supersymmetry, although the general idea holds for any theory with extended supersymmetry. With $\mathcal{N} = 2$, the anti-symmetric central charge is necessarily just a complex number Z

$$Z^{IJ} = 2\epsilon^{IJ} Z$$

For simplicity, we take Z to be real. (Typically it's not but we'll dodge this issue for now and state the full result below.) We then define the following combination of creation and annihilation operators

$$a_\alpha = \frac{1}{\sqrt{2}} \begin{pmatrix} Q_1^1 + \bar{Q}_2^2 \\ Q_2^1 - \bar{Q}_1^2 \end{pmatrix} \quad \text{and} \quad b_\alpha = \frac{1}{\sqrt{2}} \begin{pmatrix} Q_1^1 - \bar{Q}_2^2 \\ Q_2^1 + \bar{Q}_1^2 \end{pmatrix}$$

Note that we've mixed up α and $\dot{\alpha}$ indices. This is acceptable because we're working in the rest frame of the particle and so have already broken Lorentz invariance. The choice of a and b operators is designed to disentangle the mass and central charge Z , so their commutation relations read

$$\{a_\alpha, a_\beta^\dagger\} = 2(m + Z)\delta_{\alpha\beta} \quad \text{and} \quad \{b_\alpha, b_\beta^\dagger\} = 2(m - Z)\delta_{\alpha\beta}$$

with all other anti-commutators vanishing. The $\{a_\alpha, a_\beta^\dagger\}$ and $\{b_\alpha, b_\beta^\dagger\}$ are both positive definite, so the corresponding right-hand sides must be too. But this is only true if the masses are bounded by the central charges,

$$m \geq |Z|$$

This formula also holds if Z is complex; we just need to redefine the operators a and b using a phase to derive the same result. This formula is interesting. Although we haven't seen yet any specific examples, recall that the central charge Z is some combination of conserved charges in the quantum field theory. We learn that the masses of particles is bounded by the charges. This is known as the *BPS bound* although in the present context the name *Witten-Olive bound* would be more appropriate.

What about the representation theory of the algebra? Crucially, this depends on whether $m > |Z|$ or $m = |Z|$.

If $m > |Z|$, then we are in a situation very similar to the massive representation theory that we saw before. Both $a_\alpha^\dagger$ and $b_\alpha^\dagger$ act as creation operators and the result is that we have a multiplet comprising of 16 states. This is known as a *long multiplet*. We can also repeat this story with $\mathcal{N}$ supersymmetries to find that long multiplets have $2^{2\mathcal{N}}$ states.

More interesting is what happens when $m = |Z|$. In this case, half of the creation operators do nothing. For example, when $m = Z$, the b_α operators must just vanish on all states in the multiplet. Now we're back to the situation we met when discussing massless representations, with only $a_\alpha^\dagger$ acting as creation operators. The result is the hypermultiplet or vector multiplet that we saw above, each with 8 states, but now with a mass $m = Z$. This is known as a *short multiplet*.

The existence of short multiplets, whose mass is fixed to be $m = |Z|$, turns out to be a wonderfully powerful tool in the study of quantum field theories with extended supersymmetry. The basic idea is that one can usually solve quantum field theories at weak coupling. There we can identify the various states and understand the spectrum of long and short multiplets. As one moves into the strong coupling realm, we typically lose control over the dynamics. However, the short multiplets are special because their mass is pinned to be $m = |Z|$. The mass can't deviate from $|Z|$ because this would need there to be extra states in the Hilbert space and these can't magically appear from nowhere as some parameter, like a coupling constant, is varied. The only way that the short multiplets can free themselves from this constraint is if two or more short multiplets become degenerate and then combine to become a long multiplet whose mass is no longer protected. By understanding when this can (or, better yet, can't) happen we get a precious handle on the strong coupling dynamics of certain quantum field theories.

In this way, the study of short BPS multiplets shines a rare light into what happens at strong coupling. It allows us to effectively solve the dynamics of $\mathcal{N} = 2$ and $\mathcal{N} = 4$ gauge theories. It also allows us to understand the strong coupling limits of string theory, including the existence of M-theory, and to compute the microscopic entropy of certain BPS black hole solutions. It is, in short, a very useful tool.

The BPS trick is not available for $\mathcal{N} = 1$ theories and so we won't be wielding it for much of these lectures. (Actually, it can be used to compute the tension of domain walls and vortex strings in certain $\mathcal{N} = 1$ theories, but not the masses of particle states.)

2.4.3 Supersymmetry in Other Dimensions

Throughout these lectures, we will restrict ourselves to supersymmetric theories in $d = 3 + 1$ spacetime dimensions. There are, however, many interesting things to say about supersymmetric theories in other dimensions. Here we merely make a few very simple comments.

Supersymmetric Gauge Theories in Different Dimensions

We've seen that the vector multiplet of $\mathcal{N} = 1$ supersymmetry has a photon paired with a single massless Weyl spinor. This works because both have two internal degrees of freedom in $d = 3 + 1$ dimensions. We can ask: in what other spacetime dimensions might we be able to pair a photon with a fermion?

The number of polarisation states of a photon is $d - 2$. So the question really is: in what dimensions does a spinor have $d - 2$ degrees of freedom? We will see that we

can have a supersymmetric theory in which a photon pairs with a single fermion in $d = 3, 4, 6$ and 10 Lorentzian spacetime dimensions.

The story is simplest in $d = 3 + 1$ and $d = 5 + 1$. In even spacetime dimension d , a Dirac spinor has $2^{d/2}$ complex components. But the irreducible representations of the Lorentz group are Weyl spinors with $2^{(d-2)/2}$ complex components. While a complex scalar has two degrees of freedom, a complex spinor has the same number of degrees of freedom as the number of components. This is because the Dirac equation (or Weyl equation) is first order so these components include both “position” and “momentum”. This means that if we want the number of degrees of freedom of a Weyl spinor to match those of a photon then we need to solve the equation

$$2^{(d-2)/2} = d - 2$$

The solutions are $d = 4$ and $d = 6$ as advertised.

In $d = 3 + 1$ dimensions we can choose to impose *either* a Majorana condition *or* a chiral projection to a Weyl fermion. However in $d = 2 \bmod 8$ spacetime dimensions, it is possible to impose both a Majorana and Weyl condition. This halves the number of degrees of freedom of a Weyl fermion. Attempting to match the degrees of freedom of a Majorana-Weyl fermion to a photon we have

$$2^{(d-4)/2} = d - 2 \quad \text{with } d = 2 \bmod 8$$

The unique solution is $d = 10$.

Finally we’re left searching solutions in odd spacetime dimensions. It is not hard to see that there is just one possibility. In $d = 2 + 1$ dimensions, a photon has just a single polarisation state. Meanwhile, a Dirac spinor in $d = 2 + 1$ has two complex components. However we can impose a Majorana condition to make the spinor real. (For example, we can take the real Clifford algebra $\gamma^0 = i\sigma^2$, $\gamma^1 = \sigma^2$ and $\gamma^2 = \sigma^3$.) So a Majorana spinor in $d = 2 + 1$ has two real components and, correspondingly, one degree of freedom, matching that of the photon.

If we’re not in the magic spacetime dimension $d = 3, 4, 6$ or 10 then we can still have supersymmetric theories that relate a photon to a fermion. But now we need to include extra scalar degrees of freedom as well to make up the numbers.

The fact that the number of fermion degrees of freedom increases exponentially with d , while the number of bosonic degrees of freedom increases only linearly, suggests that there may be a maximum spacetime dimension in which supersymmetry is possible.

Indeed this is the case. If we don't wish to get our hands dirty with supergravity then $d = 9 + 1$ dimensions is the highest we can go. If we're happy to include gravity in the mix then there is a unique supersymmetry theory in $d = 10 + 1$ dimensions known, reasonably enough, as *eleven dimensional supergravity*. It is extremely interesting and describes the low-energy behaviour of M-theory.

Extended Supersymmetry and Higher Dimensions

There is a close relationship between supersymmetric theories in higher dimensions and extended supersymmetry. In particular, theories with $\mathcal{N} = 2$ supersymmetry naturally descend from $d = 5 + 1$ dimensions while those with $\mathcal{N} = 4$ supersymmetry come from $d = 9 + 1$ dimensions. (This statement, taken at face value, is true only at the classical level. But there are also a myriad of subtle and wonderful connections at the quantum level, none of which will be touched upon in these lectures.)

To see this, we will briefly jump ahead of ourselves slightly and use the language of fields, rather than the language of single particle quantum states that we've invoked until now. The relationship between theories in different dimensions involves a process known as *dimensional reduction*. This means that we take the fields in a higher dimension and state, by fiat, that they are independent of certain spatial coordinates. For example, consider a gauge field A_M in, say, $d = 5 + 1$ dimensions. This means that $M = 0, 1, \dots, 5$. Upon dimensional reduction, we insist that this gauge field only depends on x^μ with $\mu = 0, 1, 2, 3$. The gauge field itself then decomposes as

$$A_M \rightarrow (A_\mu, \phi_4, \phi_5)$$

That is, we get a $d = 3 + 1$ dimensional gauge field A_μ together with two real scalars ϕ_4 and ϕ_5 . But this is precisely the bosonic content of the $\mathcal{N} = 2$ vector multiplet that we found above. A $d = 5 + 1$ Weyl fermion decomposes into two $d = 3 + 1$ Weyl fermions in a similar fashion (although you have to work a little harder playing around with the gamma matrices to see this).

Playing the same game with a $d = 9 + 1$ gauge field, we find a $d = 3 + 1$ gauge field together with $10 - 4 = 6$ scalars. This is the bosonic content of the $\mathcal{N} = 4$ vector multiplet that we found above. Decomposing a $d = 9 + 1$ Majorana-Weyl fermion completes the story, giving four $d = 3 + 1$ Weyl fermions.

Finally, if you dimensionally reduce eleven dimensional supergravity you find $\mathcal{N} = 8$ supergravity in $d = 3 + 1$ dimensions.

Counting Supersymmetries

The way in which we count supersymmetries in different dimensions can be rather bewildering when you first meet it. In $d = 3 + 1$ we count supersymmetries by the number of Weyl spinor supercharges Q_α^I with $I = 1, \dots, \mathcal{N}$. But this is clearly specific to 4d. In other dimensions the counting depends on what kinds of minimal spinors we can construct. Moreover, if we dimensionally reduce then what is a minimal supersymmetry in a higher dimension typically becomes an extended supersymmetry in a lower dimension.

To avoid this confusion, it can be useful to count the number of *components* of the supercharges. We count these as N (rather than the calligraphic $\mathcal{N}$.) These components are, sadly, also referred to as supercharges! Because spinors can be real in some dimensions, we count the number of real components or, equivalently, twice the number of complex components. This means that, in $d = 3 + 1$ dimensions, $\mathcal{N} = 1$ supersymmetry has four supercharges, $\mathcal{N} = 2$ has eight supercharges, and so on.

To orient you, here is a list of some of the most interesting classes of supersymmetric theories and how they are labelled in various dimensions. The list is by no means complete but gives some sense of the more compelling supersymmetric stories out there.

The maximum number of supercharges is $N = 32$. These are all supergravity theories and can exist in any dimension $d = 10 + 1$ and below. Upon dimensional reduction, the number of minimal spinor supercharges $\mathcal{N}$ in various dimensions is

<u>N=32 supercharges:</u>	Dimension d	11	10	6	4
	Supersymmetry $\mathcal{N}$	1	(1,1)	(2,2)	8

This is not an exhaustive list: supersymmetric theories with $N = 32$ supercharges exist in all dimension $d \leq 11$. But the dimensions listed above are, for various reasons, the most interesting and well studied.

Note the strange (n, n) notation in $d = 5 + 1$ and $d = 9 + 1$. This is because of one more subtlety of representations of the Clifford algebra. When $d = 2 \bmod 4$, the two types of Weyl spinor are *not* related by complex conjugation in Lorentzian signature. This means that you can have a spinor of one chirality without necessarily having the other. In contrast, when $d = 0 \bmod 4$ (including, as we saw in great detail, in $d = 3 + 1$) the complex conjugate of a left-handed spinor is a right-handed spinor, so if you have one then you always have the other. The notation (n, n) tells us how many left- and right-handed spinor supercharges we have.

There is another supergravity theory in $d = 9 + 1$ dimension which has also 32 supercharges but with $\mathcal{N} = (2, 0)$ supersymmetry. This is more commonly known as Type IIB supergravity, with the $\mathcal{N} = (1, 1)$ theory known as Type IIA. They are the low-energy descriptions of Type IIA and IIB string theories.

Theories with $N = 16$ supercharges can exist in dimensions $d = 9 + 1$ and below. Upon dimensional reduction, the associated supersymmetry is:

<u>N=16 supercharges:</u>	Dimension d	10	6	4	3	2
	Supersymmetry $\mathcal{N}$	(1,0)	(1,1)	4	8	(8,8)

The most famous and well studied of these is the Yang-Mills theory associated to the $\mathcal{N} = 4$ vector multiplet in $d = 3 + 1$. It has many remarkable properties, including electromagnetic duality and the fact that, at strong coupling, it can be viewed as a theory of quantum gravity through the AdS/CFT correspondence. There are also interesting stories to tell about the quantum dynamics of the theories in $d = 2 + 1$ and $d = 1 + 1$ dimensions.

There is one further interesting theory with 16 supercharges. This is a strongly interacting superconformal quantum field theory in $d = 5 + 1$ dimensions with $\mathcal{N} = (2, 0)$ supersymmetry. In some ways, it can be viewed as the grandfather of all quantum field theories. Given its importance, it has a remarkably rubbish name: it is simply called the $(2, 0)$ theory.

Theories with $N = 8$ supercharges exist in $d = 5 + 1$ dimensions and below. Upon dimensional reduction, the names of the supersymmetries that one finds are

<u>N=8 supercharges:</u>	Dimension d	6	4	3	2
	Supersymmetry $\mathcal{N}$	(1,0)	2	4	(4,4)

Again, the theories with $\mathcal{N} = 2$ supersymmetry in $d = 3 + 1$ dimensions are the best studied and were first solved by Seiberg and Witten.

Theories with $N = 4$ supercharges exist in $d = 3 + 1$ dimensions and below. Upon dimensional reduction, this becomes

<u>N=4 supercharges:</u>	Dimension d	4	3	2
	Supersymmetry $\mathcal{N}$	1	2	(2,2)

Much of the focus of these lectures notes will be on understanding the dynamics of $\mathcal{N} = 1$ theories in $d = 3 + 1$ dimensions. But there are many beautiful stories in lower

dimensions as well. In particular, the study of superconformal $\mathcal{N} = (2, 2)$ theories in $d = 1 + 1$ dimensions is where one can first find the mathematical study of mirror symmetry. There are also interesting 2d theories with $\mathcal{N} = (0, 4)$ supersymmetry.

Finally, theories with $N = 2$ supercharges exist in $d = 2 + 1$ dimensions and below. The dimensional reduction to $d = 1 + 1$ gives

<u>N=2 supercharges:</u>	Dimension d	3	2
	Supersymmetry $\mathcal{N}$	1	(1,1)

There are also $\mathcal{N} = (0, 2)$ theories that do not descend from $d = 2 + 1$ dimensions. Note that these are usually written as $(0, 2)$ rather than $(2, 0)$ to give an extra hint that we're talking about 2d theories rather than the 6d theory mentioned above.

I've not included $d = 0 + 1$ theories in the above list, also known as quantum mechanics, but it's not for want of things to say. You can read about supersymmetric quantum mechanics in the [companion lecture notes](#).