

LECTURES ON SUPERSYMMETRY

by

Ingenfar Bengtsson

Fall, 1988

1.	Introduction.....	1
2	Spinors.....	4
2.1	Spinors in four-dimensional space-time.....	4
2.1.1	Geometric theory of spinors.....	4
2.1.2	Algebraic theory of spinors.....	6
2.1.3	The four-component formalism.....	10
2.2	Spinors in Euclidean spaces.....	13
2.2.1	SO(4) spinors.....	13
2.2.2	SU(2) spinors.....	15
2.3	Spinors in higher dimensions.....	16
3	Space-time supersymmetry.....	17
3.1	No-go theorems.....	17
3.1.1	All possible even symmetries of the S-Matrix.....	17
3.1.2	All possible odd symmetries of the S-Matrix.....	20
3.2	Representations of supersymmetry.....	22
4	Supersymmetric Actions.....	24
4.1	The nature of supersymmetry.....	24
4.2	The Wess-Zumino model.....	26
4.3	Supersymmetric Yang-Mills theory.....	28
4.3.1	All possible N=1 models.....	28
4.3.2	Dimensional reduction and the N=4 model.....	30
4.4	Other models.....	32
5	Superspace.....	33
5.1	Grassman numbers and all that.....	33
5.2	Some simple superfields.....	37
5.2.1	Some generalities about fields.....	37
5.2.2	Chiral superfields.....	39
5.2.3	The Wess-Zumino model in superspace.....	42
5.3	Interlude: Covariant derivatives.....	46
5.4	Gauge theories in superspace.....	48
5.4.1	Supersymmetric QED.....	48
5.4.2	Supersymmetric Yang-Mills.....	50
5.5	Survey of superspaces.....	53
6	Supergraphs.....	55
6.1	Feynman rules in superspace.....	55
6.1.1	Sample calculations.....	58
6.2	Non-renormalization theorems.....	61
6.2.1	Finite models.....	64
6.3	On the absence of quadratic divergencies.....	65
7	Supersymmetric quantum mechanics.....	68
8	Supergravity.....	68

1. INTRODUCTION

Supersymmetry has been a subject of intense interest among particle physicists since 1974 or thereabouts. So I cannot cover the whole subject in these lectures. It is difficult to make selections however, because there is no experimental evidence whatsoever that supersymmetry is relevant for particle physics. Hence I do not know which parts of the subject that may be important in the future. I will simply choose some topics according to fancy.

You are supposed to know what a symmetry is, and also that the properties of a symmetry - apart from some global details - are succinctly summarized by a Lie algebra, i.e. a vector space equipped with a bilinear, antisymmetric bracket operation obeying the Jacobi identity. Schematically,

$$[E, E] = E, \quad (1)$$

where E is a generic symbol for the elements of the vector space. We will call them even elements, and then introduce odd elements O which obey what is called a super-Lie algebra, of the generic form

$$\begin{aligned} [E, E] &= E \\ [E, O] &= O \end{aligned} \quad (2)$$

$$[O, O] = E.$$

The curly bracket operation is postulated to be bilinear and symmetric - if $[,]$ is a commutator, $\{ , \}$ is an anti-commutator - and one requires a generalization of the Jacobi identity, namely

$$\begin{aligned} [[E_1, E_2], E_3] + [[E_3, E_1], E_2] + [[E_2, E_3], E_1] &= 0 \\ [[O_1, O_2], E] + [[E, O_1], O_2] - [[O_2, E], O_1] &= 0 \end{aligned} \quad (3)$$

(and a few more, which should be obvious - eq. (3) is fairly obvious too, once you learn how to keep track of the sign; there is one each time you interchange the order of two odd elements).

As you know, all compact Lie algebras were classified by Cartan, and you now learn that all compact super-Lie algebras were classified by Kac*. We will not go into that, since we are interested in a very particular kind of

supersymmetry. Remember that the Hamiltonian is a symmetry generator for all isolated physical systems, and moreover that it is bounded from below in healthy models. This suggests that one should try to express it as a sum of squares, which is achieved if there is a supersymmetry generator Q in the model, such that

$$\{Q, Q\} = H. \quad (4)$$

It is such supersymmetries that I will talk about. They were first considered by Golfand and Likhman in 1971*.

It is not too hard to give an example of a quantum mechanical model which has this property. Consider a non-relativistic particle of spin $1/2$, moving on a line. The wave function will be a two component spinor Ψ . For Q , we simply make a guess, and then we will define the Hamiltonian as $1/2 Q^2$. So:

$$Q = 1/2(\tau_1 p + \tau_2 W(x)) \quad (5)$$

where $p = -i\hbar \frac{d}{dx}$, $W(x)$ is some arbitrary function and the τ_i are the Pauli matrices. We find that

$$H = 1/2(p^2 + W^2 + \hbar \tau_3 \frac{dW}{dx}) \quad (6)$$

If we look closer, we find that there actually are two supersymmetry generators;

$$\begin{aligned} Q_1 &= 1/2(\tau_1 p + \tau_2 W) \\ Q_2 &= 1/2(\tau_2 p - \tau_1 W) \end{aligned} \quad (7)$$

obey the algebra

$$\begin{aligned} \{Q_i, Q_j\} &= \delta_{ij} H \\ [H, Q_i] &= 0. \end{aligned} \quad (8)$$

So now we have a simple quantum mechanical model which has $N=2$ supersymmetry, of the kind we wanted. It has some quite interesting properties, but for now it will be enough to make a single remark: One could change the coefficients in the Hamiltonian without destroying its

* V.G. Kac, A Sketch of Lie Superalgebra Theory, Comm. Math. Phys. 53 (1977) 31.

* Y.A. Golfand and E.P. Likhman. Extension of the Algebra of Poincaré Group Generators and Violation of P Invariance, JETP Lett. 12 (1971) 323.

positivity, but then supersymmetry would be lost - hence supersymmetric Hamiltonians are very special. Since this is so, one can study supersymmetry for various reasons:

1. Supersymmetric quantum field theories may be so special that one can construct them - this would be interesting in the same sense that two-dimensional quantum field theories are interesting, as toy models.

2. The mere fact that a theory admits a supersymmetric extension may allow conclusions to be drawn, just as one can prove theorems about real functions by going out in complex plane.

3. There may be very special physical systems - in statistical physics, say - which exhibit supersymmetry.

4. If there is a unified field theory, it is presumably very special. Perhaps it is supersymmetric.

2 and 3 have been justified already (the most dramatic instance of the former is an alternative proof of the positive energy conjecture in general relativity), while 1 and 4 remain hopes. Anyway, we have seen that a supersymmetric model exist. However, in particle physics we want our models to be Poincaré invariant, and this means that the Hamiltonian $H = P_0$ is part of a four vector. As a result, some extra cleverness has to go into the construction of the supersymmetry algebra. It is necessary to decide how the supersymmetry generator is to transform under the Poincaré group. Eq. (4) is too simple; a suitable substitute is in fact

$$\{Q_a, \bar{Q}^b\} = 2(\gamma_a P^a)_a{}^b \quad (9)$$

Here Q_a is a spinor with suitable properties, and the γ is a gamma-matrix.

You do not have to learn much mathematics to understand an average paper on supersymmetry, but there are some things you have to learn. Spinors, and gamma-matrices, come first.

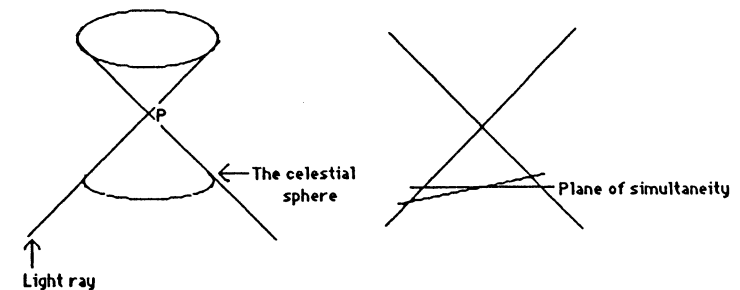
2. SPINORS.

2.1 Spinors in four-dimensional space-time.

2.1.1 Geometric theory of spinors.

You are supposed to know what vectors, and tensors, are. Spinors are similar in some respects, but they are very different in others (and they do not exist on an arbitrary manifold - to admit a "spinor structure" is a non-trivial property for a manifold). Their detailed properties depend very much on the dimension and signature of the space. For this reason I will first describe how they work in a four-dimensional space-time, and then mention only briefly what happens elsewhere. The standard treatise on spinors in space-time is the book by Penrose and Rindler*.

We begin by looking at the sky. We observe the celestial sphere. A little thought will convince you that we are looking out along the light cone from the point P, and that all the stars can be regarded as laying at some fixed distance d from P, as far as the visual impression is concerned. We can draw a picture of the situation, with one space dimension suppressed, as follows (actually, spinors in a three-dimensional space-time are similar in some respects to those in four, so you can take the picture to be accurate and change some of the words, if you like - I leave that as an exercise).



Suppressing one more dimension for clarity, and performing a Lorentz boost, the plane of simultaneity - i.e. in this context the location of the celestial sphere on the light cone - changes as shown in the second figure. A little thought will convince you that this amounts to a projective transformation of the sphere. So do rotations, obviously. In fact, there is a one-to-one correspondence between Lorentz transformations in spacetime

and projective transformations of the sphere. Now if we think of the celestial sphere as the Riemann sphere, the projective transformations are given by Möbius transformations:

$$z \sim z' = \frac{\alpha z + \beta}{\gamma z + \delta} \quad ; \quad \alpha\delta - \beta\gamma = 1. \quad (1)$$

In projective geometry, it is often useful to introduce homogeneous coordinates. For us, who are studying CP^1 (i.e. the Riemann sphere), which is the space of lines in the two complex dimensional space C^2 , this means that we should study the latter space directly. So we now have two complex coordinates, and we can recover the single complex coordinate on the Riemann sphere by forming the quotient

$$z = \frac{z_1}{z_2} \quad (2)$$

To every Möbius transformation of z corresponds exactly two transformations in the group $SL(2,C)$:

$$\begin{pmatrix} z_1 \\ z_2 \end{pmatrix} \sim \begin{pmatrix} z'_1 \\ z'_2 \end{pmatrix} = \begin{pmatrix} \alpha & \beta \\ \gamma & \delta \end{pmatrix} \begin{pmatrix} z_1 \\ z_2 \end{pmatrix} \quad (3)$$

These two component objects are called spinors, and the two complex dimensional space in which they live is called spin-space. We have seen that every spinor determines a lightlike direction in space-time; actually four real numbers are needed to specify a spinor completely, and only two to determine the direction of a light ray, so there are two extra pieces of information in the spinor to account for. It is clear that one should be able to define some kind of modulus for the spinor, which will correspond to the length of some light like vector. There will then be a phase factor left, so that to every light like vector corresponds a one-parameter family of spinors.

In some sense, then, spinors are more "fundamental" than vectors, since they can be used as building blocks for the latter. Moreover, they represent the "end of the road": We have established that $SL(2,C)$ gives a two-fold covering of the Lorentz group, and the process ends there, since $SL(2,C)$ is simply connected. It is called the universal covering group of the Lorentz group, and spinors give linear representations of the universal covering group.

2.1.2 Algebraic theory of spinors.

Let me now develop the theory of spinors purely algebraically, as the theory of linear representations of $SL(2,C)$. First of all, starting from two-component objects Ψ^A , where the index A runs from 1 to 2, we can form multi-component spinors $\Psi^{AB} \dots T$, just as we would form tensors from vectors. There will also be a dual spin space, which gives rise to spinors with indices downstairs, and contraction of spinor indices can be performed, just as with tensors. Moreover, there is an invariant two-index spinor which, because it is invariant, can be used as a "metric" in spin space:

$$L^A_C L^B_D \epsilon^{CD} = \epsilon^{AB} \quad (1)$$

where the L 's are $SL(2,C)$ matrices and

$$\epsilon^{AB} = \begin{pmatrix} 0 & 1 \\ -1 & 0 \end{pmatrix} \quad (2)$$

It is a "symplectic" metric, since it is anti-symmetric. We can use this metric, and its inverse, to raise and lower indices, i.e. to establish an isomorphism between spin space and its dual. Since ϵ^{AB} is anti-symmetric, it is important to get the conventions clear at this point. We define

$$\Psi^A = \epsilon^{AB} \Psi_B \quad ; \quad \Psi_A = \Psi^B \epsilon_{BA} \quad (3)$$

(where ϵ_{AB} is the transpose of the inverse of ϵ^{AB}). Note that

$$\begin{aligned} \Psi_A^A &= \epsilon^{AB} \Psi_{AB} = -\epsilon^{BA} \Psi_{AB} = -\Psi^A_A \\ \epsilon_A^B &= -\epsilon^B_A = \delta_A^B. \end{aligned} \quad (4)$$

Since spin space is two dimensional only, the theory of multi-component spinors will be very simple. It will be enough to consider totally symmetric objects, since all anti-symmetric pieces can be separated out using :

$$\Psi^{AB} = \Psi^{(AB)} + \Psi^{[AB]} = \Psi^{(AB)} + 1/2 \epsilon^{AB} \Psi^C_C \quad (5)$$

This is a very useful formula.

There is one further operation that we can perform on our spinors, which has no analogue in a real vector space, and that is to take the complex conjugate of a spinor. The result of that operation can not be another element of spin space, because if it was, we could define real and imaginary spinors, and so the spinors would cease to be on an equal footing - the $SL(2, \mathbb{C})$ covariance would be lost. Hence the complex conjugate of a spinor Ψ^A is an element $\Psi^{A'}$ - with a primed index - of a new space, anti-isomorphic to spin space in the sense that

$$\overline{w\Psi^A + z\omega^A} = \overline{w}\Psi^{A'} + \overline{z}\omega^{A'} \quad (6)$$

The theory of primed spinors is of course exactly analogous to the theory of the un-primed ones; there is an ϵ -spinor, we can form multi-component spinors, and so on. We can also form multi-component spinors with mixed indices, as follows:

$$\Psi_{AB} \dots T_{A'B'} \dots S' \quad (7)$$

(the order between a primed and an un-primed index does not matter, unless we are thinking in terms of explicit matrices.) Contractions involving one primed and one un-primed index is not allowed.

Of particular interest are multi-component spinors with an equal number of primed and un-primed spinors, since such spinors can be defined to be real. The subspace of real spinors with one index of each sort is a four dimensional vector space which turns out to be isomorphic to the space of real Minkowski space vectors (with timelike metric - the anti-hermitian subspace has a spacelike metric, and some authors prefer that). In fact, we will identify these spaces. Schematically,

$$V^a \equiv V^{AA'} \quad , \quad \eta_{ab} \equiv \epsilon_{AB} \epsilon_{A'B'} \quad (8)$$

We could have developed the theory of spinors "abstractly", in which case the equality signs in equation (8) would be true equalities. However, we have been thinking in terms of specific coordinate systems in spin space, and explicit components of the spinors, and so we have to take the equality signs here as something of a metaphor only. Or, to phrase it in another way, there is only one space of world vectors, but it admits different bases. One in which the vector appears as a four component column vector, and one in which it appears as a two by two hermitean (due to the reality condition) matrix. Both ways of thinking about a vector are equally correct.

To exhibit the announced isomorphism, we introduce the Infeld-van der Waerden symbols

$$\begin{aligned} \sigma_0^{AB'} &= 2^{-1/2} \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} = \sigma^0_{AB'} = \bar{\sigma}^0_{A'B'} & \sigma_1^{AB'} &= 2^{-1/2} \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} = \sigma^1_{AB'} = \bar{\sigma}^1_{A'B'} \\ \sigma_2^{AB'} &= 2^{-1/2} \begin{pmatrix} 0 & i \\ -i & 0 \end{pmatrix} = -\sigma^2_{AB'} = \bar{\sigma}^2_{A'B'} & \sigma_3^{AB'} &= 2^{-1/2} \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} = \sigma^3_{AB'} = \bar{\sigma}^3_{A'B'} \end{aligned} \quad (9)$$

These matrices obey

$$\sigma_a^{AB'} \sigma_b^{AB'} = \delta_a^{b'} \quad ; \quad \sigma_{aAB'} \sigma^{aCD'} = \delta_A^C \delta_{B'}^{D'} \quad (10)$$

Now we can rewrite eq. (8) as

$$\begin{aligned} V^a &= V^{AA'} \sigma_a^{AA'} \quad , & \eta_{ab} &= \epsilon_{AC} \epsilon_{B'D'} \sigma_a^{AB'} \sigma_b^{CD'} \\ V^{AA'} &= V^a \sigma_a^{AA'} \quad , & \epsilon_{AB} \epsilon_{C'D'} &= \eta_{ab} \sigma_a^{AC'} \sigma^{b}_{CD'} \end{aligned} \quad (11)$$

All tensors can be converted to multi-spinors in this way. If one makes use of eq. (5), it is a convenient way to decompose a tensor into its irreducible parts. As an (important!) example, an anti-symmetric tensor with two indices can be reexpressed by means of a symmetric two index spinor, as follows:

$$F^{[ab]} \equiv F^{(AB)} \epsilon^{A'B'} + \epsilon^{AB} F^{(A'B')} \quad (12)$$

The "self-duality" condition

$$*F^{ab} \equiv 1/2 \epsilon^{ab}_{cd} F^{cd} = i F^{ab} \quad (13)$$

becomes simply

$$F^{ab} = \epsilon^{AB} \Psi^{A'B'} \quad (14)$$

(which is necessarily complex - in Euclidean four dimensional space one can have real self-dual tensors, however).

For later use, we note that

$$\begin{aligned} \sigma_a^{AC'} \bar{\sigma}_{bCB'} + \sigma_b^{AC'} \bar{\sigma}_{aCB'} &= \eta_{ab} \delta_B^A \\ \bar{\sigma}_{aAC'} \sigma_b^{CB'} + \bar{\sigma}_{bAC'} \sigma_a^{CB'} &= \eta_{ab} \delta_A^{B'} \end{aligned} \quad (15)$$

Finally, we return to the connection between lightlike vectors and spinors. Obviously, every vector of the form

$$\pm \Psi^A \bar{\Psi}_{A'} \quad (16)$$

is real and lightlike, because

$$\Psi^A \Psi_A = 0 \quad (17)$$

(provided the spinor is made out of commuting numbers - later we will deal a lot with spinors made out of anti-commuting numbers, for which this does not hold).

Conversely, every real lightlike vector can be written in this form (with an obvious phase ambiguity), since

$$\det V^{AA'} = 1/2 V^2 \quad (18)$$

and when the determinant vanishes, the two-by-two matrix splits in an outer product, as claimed. The sign ambiguity in (16) is also important;

$$\text{Tr } V^{AA'} = \sqrt{2} P^0. \quad (19)$$

Hence the choice of sign determines whether the lightlike vector points into the future or into the past. This clearly requires that space-time is time-orientable, otherwise spinors could not exist. Here we have an example of a condition that a manifold has to fulfill in order to admit a spinor structure (there are further conditions).

2.1.3. The four-component formalism,

The spinor formalism that I developed above is called, for obvious reasons, the two component formalism. That such a formalism exists is closely related to the fact that $SO(4)$ is locally isomorphic to a direct product of two $SU(2)$'s - as will become obvious when we come to spinors in Euclidean spaces later - and this in its turn is an important and absolutely unique property of four dimensional manifolds. (Partial analogues of the two-component formalism exist in some special higher dimensional cases.) There is another formalism available, however, which generalizes more smoothly to higher dimensions, and which is called the four component formalism. Moreover, it was in this formalism that spinors were originally discovered, by Dirac. You are actually supposed to know this already, but I will run through it quickly.

We start by writing down the Clifford algebra

$$\{\gamma_a, \gamma_b\} = 2 \eta_{ab}. \quad (1)$$

This can be represented in terms of four by four matrices; for instance

$$(\gamma_a)_a{}^b = \sqrt{2} \begin{pmatrix} 0 & \sigma_a \\ \bar{\sigma}_a & 0 \end{pmatrix}. \quad (2)$$

These matrices are to act on the four-component spinors. Two further matrices of interest are γ_5

$$\gamma_5 = i \gamma_0 \gamma_1 \gamma_2 \gamma_3 = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} \quad (3)$$

and the charge conjugation matrix C^{ab} (which is the analogue of the ϵ -spinor; we will use it to raise and lower the γ -matrix indices $a, b, \dots$):

$$C^{ab} = \begin{pmatrix} -\epsilon_{AB} & 0 \\ 0 & \epsilon^{A'B'} \end{pmatrix}. \quad (4)$$

They obey respectively

$$\{\gamma_5, \gamma_a\} = 0, \quad \gamma_5^2 = 1; \quad C \gamma_a C^{-1} = -(\gamma_a)^T. \quad (5)$$

The charge conjugation matrix is always antisymmetric, whatever the representation.

The spinor is a column "vector" with complex entries. The Lorentz group acts on the spinors through the Lorentz generators

$$S_{ab} = -2i \gamma_{ab} = -i[\gamma_a, \gamma_b] \quad (6)$$

which obey

$$[S_{ab}, S_{cd}] = i(\eta_{bd}S_{ac} - \eta_{ac}S_{bd} + \eta_{ad}S_{bc} - \eta_{bc}S_{ad}) \quad (7)$$

(this can be shown directly from the Clifford algebra (1)). In the γ -matrix representation that we use, we can write

$$\Psi_a = \begin{pmatrix} \Psi^A \\ \bar{\Psi}_{A'} \end{pmatrix} ; \quad S_{ab} = \begin{pmatrix} I_B & 0 \\ 0 & I_{A'B'} \end{pmatrix} \quad (8)$$

where the I 's are $SL(2, C)$ -generators. There is also a conjugate spinor

$$\bar{\Psi}^a = (\Psi^\dagger \gamma_0) = (\omega_A \bar{\Psi}^{A'}) \quad (9)$$

We have seen that the spinor formalism can be developed with objects that have only two components. In the four-component formalism this means that it is possible to impose algebraic conditions on the four-component spinors, without destroying the linearity of the representation. One such condition is the Weyl condition

$$(1 - \gamma_5) \Psi = 0 \quad (10)$$

(or the anti-Weyl condition, with the sign reversed). This takes us back to the two-component formalism. Another possibility is the Majorana condition

$$(\bar{\Psi})^T = C \Psi. \quad (11)$$

In our representation, this means that

$$\Psi_a = \begin{pmatrix} \Psi^A \\ \bar{\Psi}_{A'} \end{pmatrix} \quad (12)$$

With slightly different conventions (e.g. spacelike metric) it is possible to find a real representation of the γ -matrix algebra; the Majorana condition then means that the spinor is real.

Finally we come to the *sine qua non* of supersymmetry in the four-

component formalism: Fierz identities. The starting point is the observation that any four-by-four matrix can be expressed as a linear combination of the sixteen matrices

$$1, \quad \gamma_a, \quad \gamma_{ab} = \gamma_{[ab]}, \quad \gamma_{a5} = \gamma_a \gamma_5, \quad \gamma_5, \quad (13)$$

all of which, except for the identity matrix, is traceless. In particular,

$$\begin{aligned} \Psi_a \bar{\lambda}^b = & -1/4 \bar{\lambda} \Psi \delta_a^b - 1/4 \bar{\lambda} \gamma_a \Psi (\gamma^a)^b + 1/4 \bar{\lambda} \gamma_{ab} \Psi (\gamma^{ab}) + \\ & + 1/4 \bar{\lambda} \gamma_{a5} \Psi (\gamma^a_5) - 1/4 \bar{\lambda} \gamma_5 \Psi (\gamma_5). \end{aligned} \quad (14)$$

This is a Fierz identity. To prove it, take the trace of both sides, then multiply with γ_b and again take the trace, and so on. I have assumed that the spinors are *anti*-commuting. For commuting spinors, change signs.

For Weyl spinors you have to be careful with factors of 1/2; also some of the terms on the right hand side go away. Simplifications also occur for Majorana spinors. To see this you need to know that

$$\begin{aligned} C^{ab}, (C\gamma_{a5})^{ab}, (C\gamma_5)^{ab} & \text{ are antisymmetric in } a, b. \\ (C\gamma_a)^{ab}, (C\gamma_{ab})^{ab} & \text{ are symmetric in } a, b. \end{aligned} \quad (15)$$

For anti-commuting Majorana spinors, this means that

$$\begin{aligned} \bar{\lambda} \Psi = -\lambda_a C^{ab} \Psi_b = -\Psi_b C^{ba} \lambda_a = \bar{\Psi} \lambda \\ \bar{\lambda} \gamma_a \Psi = -\bar{\Psi} \gamma_a \lambda \end{aligned} \quad (16)$$

and so on. The Fierz identity, and some elementary γ -matrix algebra, can then be used to prove the useful identities

$$\gamma_a \Psi_1 \bar{\Psi}_2 \gamma^a \Psi_3 + \gamma_a \Psi_3 \bar{\Psi}_1 \gamma^a \Psi_2 + \gamma_a \Psi_2 \bar{\Psi}_3 \gamma^a \Psi_1 = 0 \quad (17)$$

for anti-commuting Majorana spinors, and

$$\gamma_a \Psi \bar{\Psi} \gamma^a \Psi = 0 \quad (18)$$

for commuting ones (which is spinors and lightlike vectors all over again).

2.2 Spinors in Euclidean spaces.

2.2.1 SO(4) spinors.

The geometric picture of spinors that I presented above required a Lorentzian space-time, but spinors exist whatever the signature of the space. Moreover, the subject still has a strong flavour of projective geometry. Easiest to generalize, however, is the algebraic theory, in which we simply look for linear representations of the universal covering group of SO(n,m). We will look at spinors in Euclidean space only briefly, since I will not use them much in these lectures.

To see how spinors in space-time are related to spinors in Euclidean space, it is instructive to begin with complex vectors in a four complex dimensional space C^4 , which we will write in matrix form

$$x^{AA'} = \begin{pmatrix} x^{00'} & x^{01'} \\ x^{10'} & x^{11'} \end{pmatrix} \quad (1)$$

We introduce a norm, namely the one we used in Minkowski space earlier:

$$\|x^{AA'}\|^2 = 2 \det x^{AA'} \quad (2)$$

The group which preserves this norm is obviously $SL(2,C) \times SL(2,C)$, acting through matrix multiplication

$$x \rightsquigarrow Ax B \quad (3)$$

where A and B are independent $SL(2,C)$ -matrices. Clearly, space-time is a real slice of C^4 , consisting of the subspace of hermitean matrices. Its symmetry group is the "diagonal subgroup" $SL(2,C)$, where the matrices in eq. (3) are subject to the condition $B = A^\dagger$. It is related by analytic continuation to Euclidean space, which is a different real slice of C^4 (i.e. a set left invariant by a mapping t which obeys $t^2 = -1$, but I will skip the details). In fact, x belongs to Euclidean space if it is a matrix of the form

$$x^{AA'} = \begin{pmatrix} z & y \\ -\bar{y} & \bar{z} \end{pmatrix} \quad (4)$$

As you can check, the subgroup of $SL(2,C) \times SL(2,C)$ which preserves this form is $SU(2) \times SU(2)$, which is the universal covering group of $SO(4)$.

Much of the algebraic theory of space-time spinors can then be carried over directly to Euclidean space; the difference is that the primed and the unprimed indices are now completely unrelated, whereas they were

related by complex conjugation in the space-time case.

Perhaps it is as well to remark that the simple relation between Euclidean and Lorentzian spaces that I have described holds for topologically trivial spaces only - in general, the complexification of a real manifold with Euclidean signature does not possess a different real slice with a Lorentzian metric.

2.2.2 SU(2) spinors.

Let me now describe "spatial" spinors, i.e. SU(2) spinors (SU(2) is the covering group for the rotation group). Such spinors will be useful for Hamiltonian formulations, so it is natural to approach them through a "3+1 decomposition" of SL(2,C) spinors.

Begin by fixing a timelike vector of length $\sqrt{2}$:

$$n_{AA'} = \sqrt{2} \sigma_{0A'A} . \quad (1)$$

The normalization is such that

$$n^{AA'} n_{AB'} = \delta_B^{A'} . \quad (2)$$

The subgroup of SL(2,C) which leaves $n_{AA'}$ invariant is precisely SU(2). We can use it to define a positive definite hermitean inner product for the spinors:

$$\langle \phi, \psi \rangle = \bar{\phi}^{A'} n_{A'A} \psi^A . \quad (3)$$

Because of this extra structure, we can convert all primed indices to unprimed indices via

$$\bar{\psi}^{A'} n_{A'A} = \bar{\psi}_A . \quad (4)$$

Hence only one kind of indices is needed for SU(2) spinors.

The space of hermitean two index spinors which obey

$$x^A_A = 0 \quad (5)$$

is a three dimensional vector space isomorphic to, and to be identified with, E^3 . The metric is

$$q_{ij} = \text{Tr } \sigma_i \sigma_j = \sigma_i^A_B \sigma_j^B_A \quad (6)$$

i.e.

$$x \cdot y = x^i q_{ij} y^j = x^A_B y^B_A . \quad (7)$$

The epsilon-tensor is

$$\epsilon_{ijk} = i\sqrt{2} \text{Tr } \sigma_i \sigma_j \sigma_k \quad (8)$$

2.3 Spinors in higher dimensions.

Unlike tensors, spinors care about the dimension of space-time. For this reason, one might expect that spinors for spaces of dimension higher than four have nothing to do with physics. Nevertheless, much research in supersymmetry has been concerned with higher dimensional "space-times", and a smattering of knowledge about spinors in higher dimensions is desirable if you want to learn about supersymmetry. (There are many readable summaries*.) In even dimensions, the four component formalism becomes a $2^{D/2}$ -formalism, where D is the dimension of space-time. (The two-component formalism does not exist outside four dimensions.) The Weyl condition can always be imposed, but the Majorana condition is available only in some cases. If the dimension of space-time is 1+1 modulo eight (9+1 for instance), one can impose both conditions at once.

Odd dimensions are quite different, since the SO(2N) groups are quite different from the SO(2N+1) groups. The γ_5 -matrix becomes an ordinary γ -matrix, and hence the Weyl condition is no longer available.

The connection between spinors and lightlike vectors holds true in 3, 4, 6 and 10 dimensions. The reason is that eq. (18) on page 13 holds in these dimensions, if the spinor is taken to be Majorana (3 and 4), Weyl (6), or Majorana-Weyl (10), respectively. Eq. (17) holds under the same conditions, and is of importance in supersymmetry, as I will explain later. The underlying reason why these equations hold is the existence of the four "Hurwitz algebras" - real numbers, complex numbers, quaternions and octonions. These are "number fields" which can be used to define a plane on which a projective geometry can be set up. The result is that partial analogues of the two-component formalism can be set up in 6 and 10 dimensions, using quaternions and octonions, respectively. However, quaternions do not commute and octonions do not associate, and this tends to diminish the flexibility of the formalism in 6 and 10 dimensions.

* T. Kugo and P. Townsend, Supersymmetry and the Division Algebras, Nucl. Phys. B221 (1983) 357.

3. SPACE-TIME SUPERSYMMETRY.

3.1 No-go theorems.

3.1.1 All possible even symmetries of the S-matrix.

Since the supersymmetry that we are looking for is a non-trivial extension of the Poincaré group, it is only fitting to begin by describing known results that severely constrain the possibilities in this direction. A physicist conditioned by unification might try to find a comprehensive group containing both the Poincaré group and the internal symmetry groups of elementary particle physics as subgroups. (A symmetry is called internal if all matrix elements connecting states with different momenta, or spins, are zero.) There are various no-go theorems in the way, however. If, for the moment, we restrict ourselves to symmetries generated by even generators (obeying commutators), then these symmetries have to commute with all Poincaré transformations - in other words, the symmetry group of a relativistic system always splits (locally) into a direct product of space-time symmetries and internal symmetries. This does not mean that the internal symmetries in Nature have nothing to do with space-time symmetries; the existence of the internal symmetry group in Yang-Mills theory, for instance, is basically forced upon you by Lorentz invariance if you try to introduce self-interactions for massless spin one-fields. Moreover, there may still be some mixing between internal symmetries and rotations, for instance, since the group does not have to be a direct product globally. Nevertheless, the restrictions are real, and now I will describe how they come about.

The two most famous no-go theorems concerning the unification of internal and space-time symmetries are O'Raifeartaigh's theorem and the Coleman-Mandula theorem*. The former states that any two particles that belong to a symmetry multiplet of some kind must have the same mass. (This is still true for supersymmetry multiplets, as we will see.) The latter, which is an even more powerful result, is stated as a theorem about S-matrices. This means that it is hard to speak with absolute confidence about massless particles, since the asymptotic states may be difficult to define in that case. Probably, the conclusions go through unaffected in the massless case as well (at least if "Poincaré group" is replaced with "conformal group"), but for definiteness I will discuss theories with a massgap only in this section.

The basic idea behind the Coleman-Mandula theorem is this: Consider two particles that scatter against each other. When conservation of linear and angular momentum is taken into account, only the scattering angle is

left as an undetermined parameter. It is then easy to imagine that the existence of a further symmetry with non-trivial space-time properties might constrain the scattering process so severely that no solution exists, except for the trivial solution of no scattering at all. To make this idea precise, one defines a symmetry of the S-matrix as a symmetry which transforms one-particle states into one-particle states, and transforms many-particle states as tensor products. It is not obvious, but true, that these assumptions imply locality; attempts have been made to derive a version of the theorem also for non-local symmetries, but as yet the conclusions are not quite definitive. Also the generator of the symmetry is assumed to commute with the S-matrix, and it has reasonable continuity properties. The S-matrix is taken to be an analytic function of the Mandelstam variables. Analyticity is in fact necessary in the proof; there are counterexamples to the theorem, in which non-trivial scattering occurs, but only in the backward and forward directions. In particular, the theorem fails completely for 1+1 dimensional models, for which the S-matrices are not analytic. A non-trivial S-matrix in 1+1 dimensions simply involves a few phase shifts.

Next comes an assumption which is perhaps questionable, namely that the number of one-particle states with mass less than any given number M is finite. Perhaps it could be derived from the requirement that empty space should have finite heat capacity. If it is dropped, one can find counterexamples to the theorem. (A possibly interesting theory which would violate particle finiteness, if it exists, is massless higher spin theory.) It is further assumed that the symmetry group of the S-matrix contains the Poincaré group as a subgroup (if the Galilei group is chosen instead, the theorem fails), and that the S-matrix is non-trivial in the strong sense that any two particles scatter against each other, except possibly for isolated values of the momenta. From these assumptions it can be proved that the most general symmetry group possible is, locally, a direct product of the Poincaré group with some internal symmetry group.

A similar theorem could probably be proved in classical field theory as well; the "S-matrix" of the classical theory would be a transformation of free field data in the remote past to similar data in the remote future.

The Coleman-Mandula theorem is obviously very important. It is common belief that it kills all hopes to find integrable relativistic field theories in 3+1 dimensions. In an integrable model on a $2n$ -dimensional phase space, one can find n conserved charges that commute with the Hamiltonian. For a field theory to be integrable, one would need an infinite set of conserved currents, and the S-matrix of such a theory could probably be proved to be trivial by means of Coleman-Mandula type arguments, except in 1+1 dimensions, where integrable field theories are known. Still, there is no theorem that precludes field theories from being integrable in some sense; there is in fact a supersymmetric field theory, called $N=4$ supersymmetric Yang-Mills theory, which seems to share

* S. Coleman and J. Mandula, All Possible Symmetries of the S-Matrix, Phys. Rev. **159** (1967) 1251.

certain properties with integrable models. I will return to this topic.

3.1.2 All possible odd symmetries of the S-matrix.

Now we are ready to write down the super-Poincaré algebra. The even part will consist of the Poincaré algebra, together with some internal symmetry group. The most general form of the odd part of a supersymmetry algebra which is to be a symmetry algebra of a relativistic S-matrix was written down by Haag, Lopuszanski and Sohnius*. For theories with a massgap, it is

$$\begin{aligned}
 \{Q^I_A, Q^J_B\} &= \epsilon^{AB} Z^{IJ} & \{\bar{Q}_I^A, \bar{Q}_J^B\} &= \epsilon^{AB} Z_{IJ} \\
 \{Q^I_A, \bar{Q}_J^A\} &= 2 \delta^I_J P^{AA'} = \delta^I_J 2^{1/2} \begin{pmatrix} P^0 + P^3 & P^1 + iP^2 \\ P^1 - iP^2 & P^0 - P^3 \end{pmatrix} \\
 [P^{AA'}, Q^B] &= [P^{AA'}, \bar{Q}^B] = 0 \\
 [J^{ab}, Q^A] &= \sigma^{ab} A_B Q^B \\
 [B_i, Q^{JA}] &= s_i^J K^{KA}
 \end{aligned} \tag{1}$$

Here the Z's are called central charges, since they commute with everything, I, J are internal symmetry indices running from 1 to N, and the B's are generators of the internal symmetry, which can be as large as U(N) unless the Z's are non-zero, in which case it is some subgroup of U(N) (since the requirement that the B's commute with the Z's gives a condition on the former). The s's and σ 's are representation matrices for the B's and the Lorentz generators, respectively. The algebra is called, for obvious reasons, the N-extended supersymmetry algebra.

Note that

$$\text{Tr} \{Q^A, \bar{Q}^A\} = 2\sqrt{2} P^0, \tag{2}$$

and since the left hand side can be written as a sum of squares, this is really a supersymmetry of the kind we were looking for. We also see that the odd part of the algebra, given above, in fact strings together the Poincaré group with an internal symmetry group in a non-trivial manner. However, it is hard to see what such an internal symmetry group might have to do with the internal symmetry groups that we actually know about from experiment.

Let me also stress what Haag, Lopuszanski and Sohnius found to be impossible. First, one cannot get the Lorentz generators on the right hand

* R. Haag, J.T. Lopuszanski and M. Sohnius, All Possible Generators of Supersymmetries of the S-Matrix, Nucl. Phys. B88 (1975) 257.

side of some anti-commutator. Second, the only odd elements that are allowed in a symmetry algebra of the S-matrix are spinors of rank 1; spinors of higher rank cannot appear. I should say that in the massless case, a somewhat tighter structure, namely the superconformal algebra, could in principle appear as a symmetry group of the S-matrix. I will not go into this; normally one would say that any non-trivial quantum theory has non-vanishing β -function, and therefore conformal symmetries are always broken by anomalies. Actually, there are counterexamples to this among supersymmetric theories, so maybe I ought to discuss superconformal symmetries, but anyway I will not.

3.2 Representations of supersymmetry.

So what is the particle content of a supersymmetry multiplet? You are supposed to know how one gets the unitary irreducible representations of the Poincaré group. All these representations are infinite-dimensional, a typical representation space being the space of positive frequency solutions to Maxwell's equations, and at first sight it seems like a difficult task to classify all such representations. But Wigner figured out how to do it. The idea is to fix a particular momentum vector - timelike, lightlike or spacelike - and then work out the unitary representations of the "little group" of Lorentz transformations that leave this vector invariant. One then relies on the theory of induced representations - developed by Wigner especially for the purpose - and concludes that the representations of the Poincaré group stand in one-to-one correspondence with the representations of the little groups. The representations of the supersymmetry algebra can be obtained in the same way.

Let us begin with the algebra (1)* with the central charges $Z=0$. In the massless case, we choose a lightlike stability vector $P^a = P^+$ (where I use "light front notation" $P^\pm = 2^{-1/2}(P^0 \pm P^3)$). The supersymmetry algebra (for $N=1$) that pertains to the little group is simply

$$\begin{aligned} \{Q^1, \bar{Q}^1\} &= 2 P^+ & \{Q^2, \bar{Q}^2\} &= 0 \\ \{Q^A, Q^B\} &= \{\bar{Q}^A, \bar{Q}^B\} = 0. & & (1) \\ [J^{12}, Q^1] &= -1/2 Q^1 & [J^{12}, \bar{Q}^1] &= 1/2 \bar{Q}^1 \end{aligned}$$

(and a few more, which need not concern us). So the Q 's behave like a pair of creation- and annihilation operators. This means that if we start from a state with helicity λ ;

$$J^{12}|\lambda\rangle = \lambda|\lambda\rangle \quad (2)$$

we get one more state, with helicity $\lambda + 1/2$, when we act on that state with $\bar{Q}^1$. Since the Q 's are odd, it stops there, and we conclude that a supersymmetry multiplet consists of two massless states, with helicity differing by a half-integer. For $N>1$, we evidently get N independent creation- and annihilation operators, and hence we get 2^N states, with helicities ranging from λ to $\lambda + N/2$. A table of all supersymmetry multiplets with maximum helicity equal to one is given below. For such multiplets, $N \leq 4$. In the table I have added the PCT-conjugate multiplet as well, since field theories are always PCT-invariant. Hence they contain irreducible supermultiplets only if the latter are PCT-self conjugate. For representations with maximum helicity equal to two - i.e. multiplets

	N=1	N=2	N=1	N=2	N=3	N=4
Helicity=1			1	1	1	1
Helicity=1/2	1	1	1	2	3+1	4
Helicity=0	1+1	2		1+1	3+3	6
Helicity=-1/2	1	1	1	2	1+3	4
Helicity=-1			1	1	1	1

Helicity content of a few massless supermultiplets (with CPT-conjugate multiplets added where appropriate).

which include gravitons, but no higher spins - the maximum value of N is 8.

So much for massless representations. If the stability vector is taken to be timelike, we obtain massive representations (tachyonic representations of supersymmetry do not exist). For the little group, we get

$$\{Q^A, \bar{Q}^A\} = 2^{1/2} P^0 \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} \quad (3)$$

So we get two independent pairs of creation- and annihilation operators for each supercharge. Hence the number of states in a massive supermultiplet is 2^{2N} .

Representations of the algebra when the central charges do not vanish are always massive, since the central charges have dimension of mass. To work out the representations in this case, one has to perform linear combinations of the supercharges in a suitable way. It turns out that for particular values of the central charges, the phenomenon of "multiplet shortening" occurs; this means that the number of states, and the range of spins, is less than in a generic massive multiplet. This fact is of some importance in connection with symmetry breaking in gauge theories. Suppose that we have a supersymmetric gauge theory in which the gauge symmetry is broken via the Higgs mechanism, while supersymmetry remains intact. Then the number of states in the theory should remain the same as in the massless phase, and yet they are required to fill out some massive supermultiplet, which generically contains more states than the massless one. The paradox can be evaded if the algebra develops a central charge when the symmetry breaks*.

* E. Witten and D. Olive, Supersymmetry Algebras That Include Topological charges, Phys. Lett. 78B (1978) 97.

4. SUPERSYMMETRIC ACTIONS.

4.1 The nature of supersymmetry.

In this chapter I will give some simple examples of field theoretic models which exhibit supersymmetry. It is useful to sit back and think first, because certain features of these models can be guessed without calculation. First of all, it is clear from the supersymmetry algebra, for instance from

$$\{Q^1, \bar{Q}^1\} = 2 P^+ \quad (1)$$

(where I am using "light front notation";

$$\sqrt{2} P^\pm = P^0 \pm P^3, \sqrt{2} P = P^1 + iP^2, \sqrt{2} \bar{P} = P^1 - iP^2) \quad (2)$$

that two supertranslations in succession result in a translation along the light cone, forward in time. This means that supersymmetry is a dynamical symmetry, which depends on the equations of motion in critical way. (Other dynamical symmetries are generated by the Lorentz boost generators and the Hamiltonian itself.) Symmetries which only result in a reshuffling of initial data, on the other hand, are called kinematical (rotation and translation in space are examples, together with various internal symmetries). The canonical generators of dynamical symmetries - in interacting, non-linear theories - differ from generators of kinematical symmetries in an important respect, namely that they contain non-linear terms. From Noether's theorem we know that this can happen only if the Lagrangian of the model transforms into a total derivative under a dynamical symmetry transformation, since that is the only way in which a Noether charge can get non-linear terms. The symmetry is truly a symmetry only if the boundary conditions are such that this total derivative integrates to zero, of course.

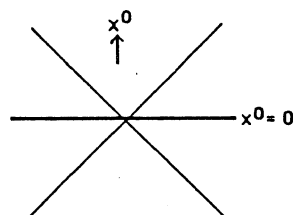
Another property of dynamical symmetries is that, in general, one expects that the transformations that they generate close to the appropriate algebra only if information about the time development of the dynamical variables is supplied, i.e. only if the equations of motion hold.

We will see in this chapter that the supersymmetric actions that we study do have these properties. In chapter 5 we will see that it is possible to add auxiliary fields to the models in such a way that linear representations of supersymmetry are obtained, and the algebra closes also "off-shell", i.e. when the equations of motion do not hold.

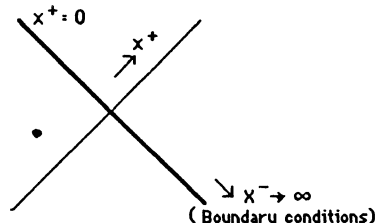
It is of course much more awkward to deal with non-linear symmetry generators, than with kinematical ones. Going back to eq. (1), we see that

there is a way of making some of the supersymmetry charges kinematical. The way to do it is to use what Dirac called "the front form of dynamics", which is what I call "light front dynamics" and everybody else, incorrectly, call "the light cone gauge". The point is that, in relativistic theories, there is some freedom in choosing the surface on which initial data for the relativistic equations are set. The usual choice is a space like surface, defined by $x^0 = 0$, say. Dirac called this "the instant form of dynamics", and if this is what you do, the supercharges are indeed dynamical, since they transform the system out of the initial data surface. However, another possible choice of surface is a surface tangent to the light cone.

Two possible initial data surfaces:



The instant form of dynamics



The front form of dynamics

The coordinate x^+ then plays the role of "time". If we use such an initial data surface, P^+ , together with Q^1 and $\bar{Q}^1$, become kinematical symmetry generators. This is one reason why light front dynamics is useful for certain questions in supersymmetry*. This kind of initial value problem also has advantages for gauge theories in general (it is possible to solve all constraints explicitly, and express the whole theory in terms of free initial data); unfortunately, there are drawbacks, since it is difficult to deal with the boundary conditions at $x^- \rightarrow \infty$ in a rigorous way. Light front dynamics is especially useful for theories that are formulated in first quantized language (i.e. strings). I will not use the light front formulation in these lectures, but you should know that it is there for you to use.

It is perhaps worth noting that in the non-relativistic limit, when the light cone "collapses outwards" and merges into the "instant" initial data surface, all the supersymmetry generators become kinematical. Supersymmetry becomes like an ordinary internal symmetry.

*L. Brink, O. Lindgren and B.E.W. Nilsson, N=4 Yang-Mills Theory on the Light Cone, Nucl. Phys. B212 (1983) 401.

4.2 The Wess-Zumino model.

Now it is time to present an actual example of a supersymmetric action. Looking through the list of representations on page 23, it is clear that the easiest example must be a model which contain two scalar states, and two spin 1/2 states (spin up and spin down). For this we need an action which contains one complex scalar field, and one two-component spinor (or one Majorana spinor, if we work in the four component formalism). In the free, massless case, the action will be

$$S_0 = \int -\bar{\psi} \not{\partial} \psi - i \lambda^A \partial_{AA'} \bar{\lambda}^{A'} \quad (1)$$

(I will be careless about surface terms throughout this chapter). It is rather simple to see that there actually is a symmetry here, which connects bosons and fermions, namely

$$\delta A = -\epsilon_A \lambda^A \quad \delta \lambda^A = -2i \partial^{AA'} \bar{\epsilon}_{A'} A. \quad (2)$$

The action is invariant, up to a surface term, under this transformation:

$$\delta S_0 = \int \partial_a (\bar{\epsilon}^{A'} \bar{\lambda}_{A'} \partial^a A - \epsilon_A \lambda^A \partial^a \bar{A} + \bar{A} \partial^a \epsilon_A \lambda^A). \quad (3)$$

Moreover, the transformations do close to the supersymmetry algebra, provided that we use the equations of motion for the spinor field:

$$\begin{aligned} [\delta_1, \delta_2] A &= 2i \epsilon_2^A \partial_{AA'} \bar{\epsilon}_1^{A'} A - (1 \leftrightarrow 2) \\ [\delta_1, \delta_2] \lambda^B &= 2i \epsilon_2^A \partial_{AA'} \bar{\epsilon}_1^{A'} \lambda^B + \epsilon_1^A \bar{\epsilon}_2^{A'} \partial_{AC} \lambda^C - (1 \leftrightarrow 2). \end{aligned} \quad (4)$$

This is so far so good, but not yet exciting, since free theories always have a large number of symmetries that are completely uninteresting, because they are not shared by any interaction terms. A symmetry is of interest only if it can be realized in some model with a non-trivial S-matrix. It is not particularly easy to find interaction terms which exhibit supersymmetry, since the very form of the supersymmetry transformations has to change in the interacting theory (that is to say, unless we take the approach of the next chapter, and add suitable auxiliary fields). It is not particularly hard, either, provided you are convinced that it is possible. Suppose we look for a renormalizable interaction. Then, there can be no derivatives in the interaction terms, which means that there will be no scalar field momenta in the canonical supercharges. This means that the transformation of the scalar field will be the same in both the free and the interacting model. Experimenting a bit, you find that the action

$$S = S_0 - \int g^2 \bar{A}^2 A^2 + 1/2g(A \lambda_A \lambda^A + \bar{A} \bar{\lambda}^A \bar{\lambda}_A) \quad (5)$$

is invariant, up to surface terms, under the transformation

$$\delta A = -\epsilon_A \lambda^A \quad \delta \lambda^A = -2i \partial^{AA'} \bar{\epsilon}_{A'} A + \sqrt{2} g \epsilon^A \bar{A}^2. \quad (6)$$

These transformations close to the supersymmetry algebra provided that you use the equations of motion that follow from the action (5). This model certainly appears to have a non-trivial S-matrix (assuming it has an S-matrix at all, which, of course, might be disputed by a constructive field theorist). We will study it in detail later. It exhibits already what, perhaps, is the most interesting property of supersymmetric field theories, namely some rather remarkable cancellations among ultra violet divergent Feynman diagrams. It is called the Wess-Zumino model, because it was first studied by Wess and Zumino, in a paper which touched off an explosion of interest in the subject*.

Actually, the way in which I constructed the action for the Wess-Zumino model here is much inferior to the superspace method that I will describe in the next chapter. The latter method has its limitations, though, in that it is not applicable to all possible models. The method used here, clumsy as it is, has the advantage that it always works.

4.3 Supersymmetric Yang-Mills theories.

4.3.1 All possible N=1 models.

Now we turn to supersymmetric gauge theories with maximum helicity 1*. According to the group theoretical analysis, there ought to be four models of this kind, with N = 1,2,3 and 4 supersymmetries respectively. However, since a field theory always has a CPT self conjugate spectrum, we see from the table on page 25 that the N=3 model has the same particle content as the N=4 model, so these two ought to be identical when considered as field theories.

Begin by looking for the N=1 model. First we write down an action containing massless spin 1 and 1/2 fields, working in four component formalism for a change, because we want to keep the dimension of space time arbitrary for the time being. Since the supersymmetry commutes with the internal colour symmetry, it is clear that the spinor field has to transform under the same (i.e. the adjoint) representation of the colour group as the vector field does. So, the action is

$$S = \int -1/4 F_{ab} F^{ab} + i \bar{\lambda}^r \gamma^r D \lambda^r \quad (1)$$

where r,s,t... are colour indices and

$$F_{ab}^r = \partial_a A_b^r - \partial_b A_a^r + f^{rst} A_a^s A_b^t$$

$$D_a \lambda^r = \partial_a \lambda^r - f^{rst} \lambda^s A_a^t \quad (2)$$

If this action is to be supersymmetric, there is almost no choice involved in the form of the transformations. On dimensional grounds, they have to be

$$\delta A_a^r = i \bar{\epsilon} \gamma_a \lambda^r - i \bar{\lambda}^r \gamma_a \epsilon \quad \delta \lambda^r = \gamma_{ab} F^{ab} \epsilon. \quad (3)$$

Up to a surface term, it turns out that the action transforms under (3) into

$$\int f^{rst} (\bar{\lambda}^r \gamma_a \lambda^s \bar{\epsilon} \gamma^a \lambda^t - \bar{\lambda}^r \gamma_a \lambda^s \bar{\lambda}^t \gamma^a \epsilon) \quad (4)$$

So the expression inside the brackets has to vanish (both terms must vanish separately) if the action is to be supersymmetric. But this happens in precisely four cases. Suppose first that the dimension of space-time is four. Since the f^{rst} 's are totally antisymmetric, we see from the Fierz

* J. Wess and B. Zumino, Supergauge Transformations in Four Dimensions, Nucl. Phys. **70B** (1974) 477.

*L. Brink, J.H. Schwarz and J. Scherk, Supersymmetric Yang-Mills Theories, Nucl. Phys. **B121** (1977) 77.

identity on page 12 that (4) is indeed identically zero, provided the spinors are chosen to obey the Majorana condition. But this Fierz identity is also true in $D=3, 6$ and 10 , provided the spinors are chosen to be Majorana, Weyl, and Majorana-Weyl, respectively. For all other choices of D , it is false. Hence, we conclude that an $N=1$ supersymmetric Yang-Mills model, with the above action, exists in precisely these dimensions.

One can of course check that the transformations (3) close to the supersymmetry algebra, provided that the equations of motion hold. Actually, this is true only up to gauge transformations; but that is enough.

Let us count the degrees of freedom in the action for the cases that it is supersymmetric. The vector field has $2x(D-2)$ degrees of freedom (the factor of 2 follows from the fact that the vector field obeys a second order differential equation). The spinor field has $2^{D/2}$ complex components when D is even, and 2 components in $D=3$. The Majorana condition restricts the number of independent degrees of freedom with a factor of 2, and so does the Weyl condition. Hence there are 2, 4, 8 and 16 real degrees of freedom in the spinor field in $D=3, 4, 6$ and 10 , respectively. So we see that the numbers match, and also that this happens only for special choices of D .

This counting is interesting from another point of view, since 4, 8 and 16 are precisely the number of degrees of freedom that we expect from an N -extended supermultiplet with maximum helicity 1 in four dimensions, when $N=1, 2$ and 4 , respectively. This suggests that there is a direct connection between (say) the existence of a super Yang-Mills model in 10 dimensions and the $N=4$ Yang-Mills model in 4 dimensions. This is in fact so, as we will see in more detail in the next section.

4.3.2 Dimensional reduction and the $N=4$ model.

To see how the $N=4$ model is hidden inside the $N=1$ ten dimensional model, we start out by declaring that nothing depends on the superfluous dimensions:

$$A_A(x) = A_a(x^0, x^1, x^2, x^3, 0, \dots, 0) \quad (1)$$

and similarly for the spinor field. ($A, B, \dots$ denotes ten dimensional tensor indices, in this section only.) This condition explicitly breaks the $SO(1, 9)$ symmetry down to $SO(1, 3) \times SO(6)$. $SO(6)$, which is locally isomorphic to $SU(4)$, is precisely the internal symmetry which is allowed into the $N=4$ supersymmetry algebra by the theorem of Haag, Lopuszanski and Sohnius.

Since $SO(1, 9)$ is no longer with us, we may as well give new names to six of the components of the vector potential. This will be done with some cleverness, as follows:

$$A_A = A_a \quad A=0, 1, 2, 3 \quad (2)$$

$$\varphi_{i4} = 1/\sqrt{2}(A_{i+3} + iA_{i+6}) ; i=1, 2, 3 \quad \varphi^{jk} = 1/2 \epsilon^{jklm} \varphi_{lm} = (\varphi_{jk})^*$$

The φ 's are scalars under $SO(1, 3)$ and transform as a sixplet under $SU(4)$. There are six of them, which is just right for the $N=4$ model in four dimensions. To take the spinor apart, we choose a suitable representation of the ten dimensional gamma-matrices (denoted by Γ ; γ will be reserved for four dimensional gamma matrices):

$$\Gamma^A = \gamma^a \times 1 \quad A=0, 1, 2, 3$$

$$\Gamma^i = \gamma_5 \times \begin{pmatrix} 0 & \rho^{ij} \\ \rho_{ij} & 0 \end{pmatrix} \quad i, j=1, 2, 3, 4 \quad C_{10} = C \times \begin{pmatrix} 0 & 1_4 \\ 1_4 & 0 \end{pmatrix} \quad (3)$$

where the ρ 's are antisymmetric $SU(4)$ matrices and C_{10} is the ten dimensional charge conjugation matrix. In this representation, a ten dimensional Majorana-Weyl spinor assumes the form

$$\lambda = \begin{pmatrix} \chi^i \\ \tilde{\chi}_i \end{pmatrix} \quad \tilde{\chi}_i = (C\tilde{\chi})^T \quad (4)$$

where the $SU(4)$ index i runs from one to four, and $\chi^i(\tilde{\chi}_i)$ is a four dimensional Weyl (anti-Weyl) spinor. There are four of each; again the right number for the $N=4$ model.

All that remains to do is to insert the vector and spinor fields, so labelled, into the ten dimensional Lagrangian density. The result is:

$$\begin{aligned} \mathcal{L} = & -1/4 F_{ab}^r F^{ab} + 1/2 D_a \phi_{ij}^r D^a \phi^{ij} + i \bar{\chi}^r \gamma \cdot D \chi^r - \\ & - i/2 f^{rst} (\bar{\chi}^r \chi^s \phi_{ij}^t - \bar{\chi}^s \chi^r \phi_{ij}^t) \\ & - 1/4 f^{rst} \phi_{ij}^s \phi_{kl}^t f^{ruv} \phi^{u\bar{ij}} \phi^{v\bar{kl}} . \end{aligned} \quad (5)$$

This is indeed the Lagrangian density of the N=4 supersymmetric Yang-Mills model. You can check the N=4 supersymmetry by relabelling the ten dimensional formulæ for the supersymmetry transformations, in the same manner as I did for the action.

This model turns out to have some quite amazing properties; for instance, its β -function vanishes, at least perturbatively, which means that the scale invariance of the classical model survives quantum corrections. Nevertheless, there is a scale available in the model. As you can see, the scalar potential does not have a unique minimum; rather, it vanishes along "vacuum valleys" in the space of values for the scalar fields. Hence the vacuum expectation values of the scalars might serve as a scale in the quantum theory.

4.4 Other models.

We have seen how to construct interacting field theories corresponding to all representations of supersymmetry with maximum helicity one (the N=2 model can be obtained by dimensional reduction of the six dimensional model). In the list on page 23, only the N=2 model with maximum spin one half - called the hypermultiplet - is missing. Actually, it so happens that no renormalizable interacting model of this type exists. Interacting models with maximum helicity three halves do not exist at all*, and bringing in helicity two (the graviton) is something I will consistently avoid, until I come to chapter 8.

Of course, many more models can be obtained by coupling different models together. For instance, the hypermultiplet can be coupled to the N=2 Yang-Mills model. The result is precisely the N=4 model, possibly modified by - say - the inclusion of a mass term for the hypermultiplet. A supersymmetric version of QED can be obtained by coupling the Wess-Zumino multiplet to a vector multiplet, and so on. On the other hand, it is impossible to assign the electron and the photon to the same supermultiplet, since they differ in their internal quantum numbers.

The latter remark is rather serious from the point of view of phenomenology. The world is clearly not supersymmetric, since particles in the real world have different masses. However, one might hope that the observed particles could be unified in the framework of spontaneously broken supersymmetry, where mass differences will develop. Unfortunately, spontaneous breakdown of supersymmetry will not change internal quantum numbers, and when one scans the list of observed particles one realizes that their quantum numbers are such that no two observed particles can be "superpartners" of each other. I will show later that it is nevertheless possible to make a case for supersymmetry in phenomenology. But this will then force you to postulate the existence of a new, hitherto unobserved, superpartner for every existing particle, and then to invent explanations for why they have not been observed.

These unobserved superpartners go under names such as photinos, gluinos, squarks, sneutrinos, shiggses and what nots. On the basis of this, you may be tempted to formulate a comment on the amount of æsthetic sense that has gone into this business so far.

For extended supersymmetry, an additional problem appears with regard to phenomenology: If the helicity 1/2 fermions sit in the same supermultiplet as those with helicity -1/2, then both kinds of fermions must have the same internal quantum numbers. But this is not true for the fermions in the standard model.

* Anders Bengtsson, Spin, Supersymmetry and Interacting Field Theories, PhD thesis, Göteborg 1984.

5. SUPERSPACE.

Superspace - an invention due to Salam and Strathdee* - is a space, some of whose dimensions are spanned by anti-commuting numbers. It can be regarded as just a clever book-keeping device, which enables one to write down supersymmetric actions without effort, and then to compute Feynman diagrams in a very efficient manner. It has its limitations, though; for reasons which (I think) are poorly understood, it works smoothly mostly for $N=1$ supersymmetry in four dimensions, and in dimensions lower than that. It could be that there is more to the story. Time will tell. The subject is highly developed technically, and there are a number of good reviews**. I will consciously try to avoid to sound like these reviews, because what they do I can not do better.

5.1 Grassman numbers and all that.

I have used anti-commuting numbers here and there already, and now it is time to formalize them a little bit. A Grassman algebra with N generators (N can be infinite, if you like) is an associative algebra with a unit, over real or complex numbers, generated by N quantities which obey

$$\xi_i \xi_j + \xi_j \xi_i = 0 \quad (1)$$

Any further relation among the generators is a consequence of eq. (1). Now the rule is this: Take anything that you can do with real numbers, like integration, say, and try to generalize it to Grassman numbers. Then you call the result "super-whatever"; superintegration, in this case.

Let us begin by writing down the most general member of a Grassman algebra:

$$z = \sum a_{ijk\dots} \xi_i \xi_j \xi_k \dots \quad (2)$$

where the coefficients $a_{ijk\dots}$ are real (or complex) numbers. This is called a supernumber; the first term, which is an ordinary number, is called its body, and the rest is called its soul. Those terms which contain an even number of generators are called the even part; it has "Grassman parity 0", and commutes with everything. The rest of the supernumber is its odd part; it has "Grassman parity 1" and it anticommutes with every other Grassman odd number. Superaddition and supersubtraction are defined in

* A. Salam and J. Strathdee, Supergauge Transformations in Four Dimensions, Nucl. Phys. B76 (1974) 477.

** J. Wess and J. Bagger, Supersymmetry and Supergravity, Princeton U.P., 1983.
M. Sohnius, Introducing Supersymmetry, Phys. Rep. 128 (1985) 39.

the obvious way. As for multiplication, you should note that - as long as N is finite, which is the case which we will deal with below - any supernumber with vanishing body is nilpotent, i.e. you can find an integer n such that $z^n = 0$. Division does not generalize so well; the inverse of a supernumber exists - and is then unique - if and only if its body is non-zero (because one cannot define the inverse of a nilpotent quantity).

The complex conjugate of a product is defined as

$$(\xi_1 \xi_2)^* = \xi_2^* \xi_1^* \quad (3)$$

Let us now turn to superanalysis. Differentiation is straightforward, of course:

$$\frac{d}{d\xi} \xi = 1 \quad (4)$$

Taylor expansions are particularly easy, and involve just two terms (or a finite number, if you expand in several variables):

$$f(\xi) = a + b\xi \quad (5)$$

Integration is more tricky. It is defined formally, as a linear functional which has nothing to do with measure theory. You should not sniff at it, though; it turns out to be very useful. Here is the definition:

$$\int d\xi (a + b\xi) = a \int d\xi 1 + b \int d\xi \xi = b \quad (6)$$

There are no endpoints. We are not integrating from one number to another, we just have this definition. (Actually, differentiating and integrating yields the same result.) Experimenting a bit, you see that we can define a delta-function which obeys

$$\delta(\xi' - \xi) = \xi' - \xi \quad \delta(-\xi) = -\delta(\xi) \quad (7)$$

Multiple integrals are defined in the obvious way. Note that the $d\xi_i$'s anticommute with each other.

You may wish to change variables in a superintegral. In an ordinary integral, you can do this if you know how to compute the Jacobian determinant. The superdeterminant is called the Berezinian. So we are looking at matrices of the general form

$$A = \begin{pmatrix} m_1 & n_1 \\ n_2 & m_2 \end{pmatrix} ; m_i \text{ even, } n_i \text{ odd} \quad (8)$$

A sensible definition of a superdeterminant should guarantee that

$$\text{Ber } AB = \text{Ber } A \text{ Ber } B \quad (9)$$

It turns out that this demand is obeyed by

$$\text{Ber } A = \text{Det } A (\text{Det } (D - CA^{-1}B))^{-1}. \quad (10)$$

(For a matrix made out of ordinary numbers, the same formula applies provided that the last factor is raised to the power +1).

The supertrace is defined as

$$\text{STr } A = \text{Tr } m_1 - \text{Tr } m_2. \quad (11)$$

The minus sign is there to guarantee that the supertrace of a commutator always vanishes. The Berezinian can then be expressed in terms of the supertrace as

$$\text{Ber } A = e^{\text{STr}(\ln A)} \quad (12)$$

This works out as it should, i.e.

$$\begin{aligned} \text{Ber } AB &= \text{Ber } e^{\ln A + \ln B + 1/2(\ln A, \ln B) + \dots} = e^{\text{STr}(\ln A + \ln B)} = \\ &= e^{\text{STr} \ln A} e^{\text{STr} \ln B} = \text{Ber } A \cdot \text{Ber } B. \end{aligned} \quad (13)$$

This is already everything I need for these lectures. As you begin to use Grassman numbers in calculations, all sort of funny sign errors will happen. All you need to sort them out is a little bit of care and common sense. There is much more to "supermathematics", though. I will mention a few things very briefly - you can find out more about it in a book written by Berezin*, who was the pioneer in this field.

First of all, Grassman numbers, and a supersymmetry, occur naturally, already at the classical level, in gauge theories. It is now understood that the "correct" phase space of a gauge theory includes ghosts, which are variables having a Grassman parity opposite to that of the constraints which define the physical subspace of the system. The point about extending the phase space with these ghosts is that different ways of writing the constraints that define the physical subspace turn out to be equivalent up to canonical transformations in this extended phase space. Moreover, there is a symmetry, called BRST symmetry, in this phase space, which is actually a supersymmetry, since it mixes the original phase space variables with the ghosts, which have the opposite Grassman parity.

* F.A. Berezin, *Superanalysis* (edited by A.A. Kirillov), Reidel 1987.

It is not a supersymmetry of the kind to which these lectures are devoted, though, so I have nothing more to say about the BRST supersymmetry.

The "classical limit" of a quantum theory which contains fermions also uses Grassman numbers. But this means that we must define infinite dimensional symplectic supermanifolds carefully (more or less; much less in practice) before we can study gauge theories, or theories containing fermions, in depth. The "symplectic" part is easy. It just means that we have to define Poisson brackets, in the obvious way:

$$\{A, B\} = A \overleftarrow{\partial}_i \omega^i_j \overrightarrow{\partial}_j B \quad (14)$$

The only difference, compared to the usual case, is that the Poisson bracket is now symmetric, as opposed to anti-symmetric. Again, a little bit of common sense will enable you to work out all the rest.

Supermanifolds constitute a difficult subject, however. The last word on things like the topology of supermanifolds is yet to be said, as far as I understand. They are defined using "sheafs", in an algebraic manner. An ordinary manifold - its topology, differentiable structure, the whole works - can be defined in terms of the algebras of real valued functions from subsets of the abstract set that is going to be the manifold, and homomorphisms of these algebras that arise when one set is a subset of another. It is said to have dimension n if the algebra of functions at any given point has n generators - which is just another way of saying that any function can be expressed as a function of n coordinates. (You can find out more about this in the book by Penrose and Rindler.) A supermanifold can be defined rigorously in the same manner. It is said to have dimension (p, q) if the algebras of functions have $p + q$ generators, of which p take values among ordinary numbers and q are Grassman valued.

5.2 Some simple superfields.

5.2.1 Some generalities about fields.

Ordinary Minkowski space is the homogeneous space that you get from the Poincaré group by dividing out the Lorentz group. What this means is that for any point in Minkowski space, you can find a group, isomorphic to the Lorentz group, which leaves that point invariant. Moreover, you can go from any given point to any other point by means of a translation. Superspace is the homogeneous space that you get from the super Poincaré group by dividing out the Lorentz group in a similar manner. A general element of the supertranslation group will be parametrized as

$$e^{ix \cdot P + i\theta_A \hat{Q}^A + i\bar{\theta}^{\dot{A}} \hat{\bar{Q}}_{\dot{A}}} \quad (1)$$

The "extra" dimensions, spanned by the θ 's, that you get in this way are "Grassman dimensions", since they correspond to supertranslations. Since the latter do not anticommute with each other, flat superspace has torsion. You can keep these things in the back of your mind, if you want, but I will introduce superfields in a more pedestrian manner. The important thing to notice is that the idea is very simple, even if the details tend to be a bit messy.

So, a superfield should be a function of x^a , θ_A and $\bar{\theta}_{\dot{A}}$, and the hope is that we can find superfields that can be used to describe the models that we constructed in chapter 4. Suppose we try a scalar superfield. To see what sits inside it, we perform a Taylor expansion in the θ 's:

$$\begin{aligned} \Phi(x, \theta, \bar{\theta}) = & a(x) + \theta \lambda_1(x) + \bar{\theta} \bar{\lambda}_2(x) + \theta \theta b_1 + \bar{\theta} \bar{\theta} b_2 + \\ & + \theta V(x) \bar{\theta} + \bar{\theta} \bar{\theta} \theta \chi_1(x) + \theta \theta \bar{\theta} \bar{\theta} \chi_2(x) + \theta \theta \bar{\theta} \bar{\theta} c(x) \end{aligned} \quad (2)$$

where

$$\theta \lambda = \theta_A \lambda^A \quad \bar{\theta} \bar{\lambda} = \bar{\theta}^{\dot{A}} \bar{\lambda}_{\dot{A}} = (\theta \lambda)^* \quad \theta V \bar{\theta} = \theta^A V_{AA'} \bar{\theta}^{\dot{A}'} \quad (3)$$

I will use this notation for anti-commuting spinors only.

We see that the number of bosonic and fermionic degrees of freedom are equal, but apart from this it does not look like any of the unitary representations of the supersymmetry algebra that we studied in chapters 3 and 4.. Before we think about this further, it is useful to think a little bit about how ordinary fields work, just to remind you that they are hard to understand, too. So, suppose we want to describe a massive particle with

spin 1. It can exist in three different states, so we expect to use a field with three degrees of freedom. On the other hand, we insist that we should use a linear, finite dimensional representation of the Lorentz group to describe it. This is a strange idea to start with, because a finite dimensional representation of a non-compact group can not be unitary. Moreover, there are no three dimensional representations of that kind around. We insist on it anyway, and will say that the massive spin 1 - particle is described in a "manifestly covariant" manner if we succeed. The reason why we insist on it is partly a matter of convenience, and partly philosophical. The point is that, to a mathematician, a tensor is not a bunch of components that transform in a specific way under changes of coordinate systems, it is something which exists independently of any coordinate system. Starting from the "algebraic" way of looking at a manifold that I sketched in section 5.1, one defines a tangent vector at a point as a derivation of the algebra of functions at the point, and tensors of higher valence are then defined in terms of the vectors. One can then compute what the tensor looks like in any coordinate system, and derive their transformation properties. The point is that you can run this definition backwards - once your model is expressed in terms of linear tensor representations of the Lorentz group, you know that it is independent of any coordinate system, which, from a philosophical point of view, is a very important property. Of course, the philosophical motivation for using superfields is less clear, since we have no very clear preconceptions about the meaning, if any, of superspace.

Now, the only reasonable candidate for a Lorentz tensor representation which can describe a massive spin 1-particle is a field A_a , transforming as a vector under the Lorentz group. In order to make contact with the representation theory of the Poincaré group, it is necessary to build an action for this field. The variational principle will then give us equations of motion and constraint equations that supply additional conditions on the degrees of freedom that make up the field. In this particular case, it turns out that the fourth component - A_0 - is connected to the spatial components of the vector field through second class constraints, so there are only three independent degrees of freedom in the vector field after all. The lesson that we draw from this is that the representation content of a field, as far as the Poincaré group is concerned, is determined both by the tensor character of the field (its "off-shell" properties), and by the action principle. In general, the number of "off-shell" degrees of freedom is greater than the number of "on-shell" degrees of freedom. Hence the fact that our scalar superfield seems to contain far too many degrees of freedom to describe a supermultiplet is not necessarily a problem. It still needs some polishing, though, as I will show in the next section.

5.2.2 Chiral superfields.

Let us now try to represent the supersymmetry algebra on the scalar superfield. We are not yet concerned with unitary representations - that will come with the action principle - but only with "off-shell" representations. Notice that the supersymmetry transformations now have to form a closed algebra without help from the equations of motion, since we do not have any equations of motion yet. We can write down the following set of operators, which furnish a linear representation of the super Poincaré algebra when they act on the scalar superfield:

$$\begin{aligned} P_A &= i \partial_A \\ J_{ab} &= i(x_{AA'} \partial_{BB'} - x_{BB'} \partial_{AA'} - \theta_{(A} \partial_{B)} \epsilon_{A'B'} - \bar{\theta}_{(A'} \bar{\partial}_{B')} \epsilon_{AB}) \\ Q_A &= \partial_A - i \partial_{AA'} \bar{\theta}^{A'} \\ \bar{Q}_{A'} &= \bar{\partial}_{A'} - i \partial_{AA'} \theta^A \end{aligned} \quad (1)$$

There are two signs that you should notice here. In the first place, Q and $\bar{Q}$ obey the supersymmetry algebra with the sign reversed. This is not a mistake; the canonical generators that you build from the fields will get the correct sign in this way. The second is a peculiar "Grassman-sign":

$$\partial_A = \frac{\partial}{\partial \theta^A}; \quad (\partial_A)^* = -\bar{\partial}_{A'} \quad (2)$$

The representation in terms of component fields is determined as follows:

$$\delta \Phi = (\epsilon Q - \bar{\epsilon} \bar{Q}) \Phi = \delta a + \theta \delta \lambda_1 + \bar{\theta} \delta \bar{\lambda}_2 + \dots \quad (3)$$

The Lorentz properties of the component fields also work out correctly.

However, we do not yet have an irreducible representation. One way of getting such a representation is to restrict Φ to be real, a property which is preserved by super Poincaré transformations, and which obviously will cut down the number of independent fields inside it quite a bit. As it happens, a real scalar superfield turns out to be suitable for describing super Yang-Mills theories. There is another way to get an irreducible representation, however. You notice that the operators $D_A, \bar{D}_{A'}$, defined as

$$D_A = \partial_A + i \partial_{AA'} \bar{\theta}^{A'} \quad \bar{D}_{A'} = \bar{\partial}_{A'} + i \partial_{AA'} \theta^A \quad (4)$$

obey

$$\{D_A, \bar{D}_{A'}\} = 2i \partial_{AA'}; \quad \{D, D\} = \{\bar{D}, \bar{D}\} = \{D, Q\} = \{D, \bar{Q}\} = \{\bar{D}, Q\} = \{\bar{D}, \bar{Q}\} = 0 \quad (5)$$

These operators are called covariant derivatives, because they obey Leibniz' rule, and because the covariant derivative, when applied to a superfield, results in another superfield, which also supplies a linear representation of the super Poincaré group (this is a consequence of eq. (5)). That such covariant derivatives exist in superspace, but not in Minkowski space, is a consequence of the non-commutativity of the Q 's. The point is that the action of a supertranslation on a general element of the supertranslation group can take place either from the left or from the right; the D 's then anti-commute with the Q 's as a consequence of group associativity:

$$\begin{aligned} g(x', \theta', \bar{\theta}') &= e^{i(\epsilon Q + \bar{\epsilon} \bar{Q})} g(x, \theta, \bar{\theta}) = e^{i(\epsilon D + \bar{\epsilon} \bar{D})} g(x, \theta, \bar{\theta}) = g(x, \theta, \bar{\theta}) e^{i(\epsilon Q + \bar{\epsilon} \bar{Q})}; \\ (e^{i(\epsilon Q + \bar{\epsilon} \bar{Q})} g) e^{i(\zeta Q + \bar{\zeta} \bar{Q})} &= e^{i(\epsilon Q + \bar{\epsilon} \bar{Q})} (g e^{i(\zeta Q + \bar{\zeta} \bar{Q})}) \\ e^{i(\zeta D + \bar{\zeta} \bar{D})} e^{i(\epsilon Q + \bar{\epsilon} \bar{Q})} g &= e^{i(\epsilon Q + \bar{\epsilon} \bar{Q})} e^{i(\zeta D + \bar{\zeta} \bar{D})} g \end{aligned} \quad (6)$$

Anyway, this means that we can impose a condition on the superfield which preserves its character as a superfield, viz.

$$\bar{D}_{A'} \Phi = 0 \quad (7)$$

A superfield which obeys this condition is called a chiral superfield (or an antichiral superfield, if you impose $D_A \Phi = 0$), and this is irreducible, too.

One can solve the chirality condition, as follows

$$\Phi = e^{i\theta\partial\bar{\theta}} (\sqrt{2} A(x) + \sqrt{2}\theta\lambda(x) + \theta\theta F(x)) \quad (8)$$

where the exponential function is defined by means of its Taylor expansion (and the square roots of two have been inserted for later convenience). The representation content is now reasonably close to the Wess-Zumino multiplet, and we will see in the next section, when we construct an action involving such a field, that it gives precisely the Wess-Zumino multiplet. Since any (interesting) superfield has to contain a spinor, which has two complex components, and since there must be an equal number of bosonic components in the superfield, it is anyway clear that this is as close to the Wess-Zumino multiplet as we can get when we work with "off-shell" fields.

One more comment is in order. The condition (7), which guarantees the irreducibility of the superfield, has no analogue for ordinary Minkowski space fields. A massive spin 2 field, for instance, is described by a symmetric tensor, which carries a reducible representation of the Lorentz group. However, any attempt to separate the trace from the traceless part

before writing down the action will violate locality, and so any constraint on the field must come from the action principle itself. For the superfield, there is no problem; the condition (7) does constrain the superfield, but its remaining component fields (A, F and λ) are unconstrained.

5.2.3 The Wess-Zumino model in superspace.

So how do we define an action directly in superspace ? Since the supersymmetry algebra is represented linearly on the superfields, it is clear that polynomials in superfields, such as

$$\Phi^2, \Phi^3, \dots, \quad \bar{\Phi}\Phi, \dots \quad (1)$$

are superfields, too. Moreover, expressions involving only Φ are still chiral superfields. This is already something. Next we notice an important property of any superfield, namely that its last component - the term in the Taylor expansion which multiplies $\theta\theta\theta\bar{\theta}$ - always transform into a total derivative under supersymmetry transformations, just like the Lagrangian density of a supersymmetric model does. Hence the last components of the superfields in eq. (1) are suitable candidates for Lagrangian densities. To pick out these components from the superfields, we can use Berezin integration, as follows

$$\int d^4x d^4\theta (\bar{\Phi}\Phi + \Phi^2 + \bar{\Phi}^2 + \dots) \quad (2)$$

where

$$d^4\theta = d^2\theta d^2\bar{\theta}; \quad \int d^2\theta \theta\theta = 1, \quad \int d^2\bar{\theta} \bar{\theta}\bar{\theta} = 1. \quad (3)$$

Only terms which contain four θ 's will survive the integration.

The action looks a little bit funny, because no derivatives are apparent, but you should remember that there are derivatives hidden inside the superfield. It is not yet satisfactory, though, because it is obvious from eq. (8) in the last section that the chiral terms in eq. (2) will contribute total derivatives to the Lagrangian density only, since the $\theta\theta\theta\bar{\theta}$ -term in a chiral superfield is a total derivative. Pondering eq. (8) a bit further, you realize that, in a chiral superfield, already the $\theta\theta$ -term transforms into a total derivative under supertranslations. Hence we can try the following action for the Wess-Zumino multiplet, obviously supersymmetric (as long as surface terms can be ignored, as they can in massive models), real, and obviously renormalizable, since there are no coupling constants with positive dimension in units of length ($\dim[x]=1$, $\dim[\theta]=1/2$):

$$\int d^4x d^4\theta \bar{\Phi}\Phi + \left(\int d^4x d^2\theta (\mu \Phi^2 + g/6 \Phi^3) + \text{c.c.} \right) \quad (4)$$

As we will see later, when we discuss quantum perturbation theory in superspace, the fact that some of the terms in the action involve an integration over $d\theta^2$ only has very important consequences. For now, let

me simply evaluate the Berezin integral:

$$S = \int d^4x (- \bar{A} \not{D} A - i \lambda \bar{\lambda} + F \bar{F} + \mu (2\sqrt{2} A F + 2\sqrt{2} \bar{A} \bar{F} - \lambda \lambda - \bar{\lambda} \bar{\lambda}) + g (A^2 F - 1/2 A \lambda \lambda + \bar{A}^2 \bar{F} - 1/2 \bar{A} \bar{\lambda} \bar{\lambda})) \quad (5)$$

We see that the "extra" scalar field F is an auxiliary variable, which will give rise to a very simple second class constraint when we analyze the action:

$$F + 2\sqrt{2} \mu \bar{A} + g \bar{A}^2 = 0. \quad (6)$$

Inserting this directly into the action (5), as we are allowed to do, it becomes:

$$S = \int d^4x (- \bar{A} \not{D} A - i \lambda \bar{\lambda} - 8\mu^2 \bar{A} A - \mu (\lambda \lambda + \bar{\lambda} \bar{\lambda}) - g^2 \bar{A}^2 A^2 - 2\sqrt{2} \mu g (\bar{A}^2 A + \bar{A} A^2) - 1/2 g (A \lambda \lambda + \bar{A} \bar{\lambda} \bar{\lambda})) \quad (7)$$

(The fermion masses equal the boson masses, as they should.)

But this is precisely the action we studied in the last chapter. Extracting the supersymmetry transformations that we get for the component fields, and using eq. (6), we find that

$$\begin{aligned} \delta \Phi &= (\epsilon Q - \bar{\epsilon} \bar{Q}) \Phi = \\ &= \sqrt{2} (- \epsilon \lambda) + \sqrt{2} \theta_A (- \sqrt{2} \epsilon^A F - 2i \not{D}^{AA'} \bar{\epsilon}_{A'} A) + \theta \theta (- i \sqrt{2} \bar{\epsilon}^{A'} \partial_{AA'} \lambda^A) \\ \delta A &= - \epsilon \lambda \quad F = - i \sqrt{2} \bar{\epsilon}^{A'} \partial_{AA'} \lambda^A \\ \delta \lambda^A &= - 2i \not{D}^{AA'} \bar{\epsilon}_{A'} A - \sqrt{2} \epsilon^A F = - 2i \not{D}^{AA'} \bar{\epsilon}_{A'} A + 4\mu \epsilon^A \bar{A} + \sqrt{2} g \epsilon^A \bar{A}^2 \end{aligned} \quad (8)$$

Again precisely what we had in the last chapter. We see that it is the auxiliary field F which ensures that the supersymmetry algebra closes when the equations of motion do not hold, and also that it enables the representation to become linear. Dynamically, it is trivial, since it is connected to the other variables through a second class constraint.

So, what happened ? It is clear that the superspace technique enables one to write down supersymmetric actions with very little labour (although it may become necessary to devote some care to surface terms in models involving massless fields). This is already something. The real worth of the technique will not become apparent until we discuss superspace perturbation theory, however.

It is instructive to derive the equations of motion directly in superspace. So we will vary the action with respect to Φ . This is slightly tricky, however. Consider the kinetic term. Clearly

$$S_{\text{KIN}} = \int d^4x d^4\theta \delta \Phi \Phi \quad \not{D} \frac{\delta S}{\delta \Phi} = \Phi \quad (9)$$

This is because Φ is anti-chiral, and therefore the functional derivative has to be anti-chiral, too. In order to define the functional derivative correctly, it is useful to observe first that

$$\int d^2\theta = - 1/4 DD + \text{surface term}. \quad (10)$$

Ignoring the surface term, we can therefore write the variation of the action in the form

$$\begin{aligned} \delta S &= \int d^4x d^2\theta \delta \bar{\Phi} (- 1/4 DD \Phi + 2\mu \bar{\Phi} + g/2 \bar{\Phi}^2) \\ \frac{\delta S}{\delta \Phi(x, \theta, \bar{\theta})} &= - 1/4 DD \Phi + 2\mu \bar{\Phi} + g/2 \bar{\Phi}^2. \end{aligned} \quad \Leftrightarrow \quad (11)$$

The term in brackets is anti-chiral, and is indeed the equation of motion for the superfield. There is another way to say this. Introduce the following projection operators:

$$\begin{aligned} \Pi_{0+} &= - 1/8 \frac{DD\bar{D}\bar{D}}{\square} & \Pi_{0-} &= - 1/8 \frac{\bar{D}\bar{D}DD}{\square} & \Pi_{1/2} &= 1/4 \frac{D\bar{D}\bar{D}D}{\square} \\ \Pi_{0+} + \Pi_{0-} + \Pi_{1/2} &= 1 \\ \Pi_{+} \bar{\Phi} &= \bar{\Phi} & \Pi_{-} \Phi &= \Phi \end{aligned} \quad (12)$$

We can in fact rewrite the action as an integral over the entire superspace, provided that we allow non-local operators, as follows:

$$\begin{aligned} S &= \int d^4x d^4\theta \bar{\Phi} \Phi + \int d^4x d^2\theta (\mu \bar{\Phi}^2 + g/6 \bar{\Phi}^3) + \dots = \\ &= \int d^4x d^4\theta \bar{\Phi} \Phi - 1/8 \int d^4x d^2\theta \frac{DD\bar{D}\bar{D}}{\square} \bar{\Phi} (\mu \bar{\Phi} + g/6 \bar{\Phi}^2) + \dots = \\ &= \int d^4x d^4\theta (\bar{\Phi} \Phi + \mu/2 \frac{\bar{D}\bar{D}}{\square} \bar{\Phi} \bar{\Phi} + \frac{DD}{\square} \Phi \Phi) + g/12 (\frac{\bar{D}\bar{D}}{\square} \bar{\Phi} \bar{\Phi}^2 + \frac{DD}{\square} \Phi \Phi^2). \end{aligned} \quad (13)$$

If we define

$$\frac{\delta \Phi(x, \theta, \bar{\theta})}{\delta \bar{\Phi}(x', \bar{\theta}', \theta')} = - 1/4 \bar{D}\bar{D} \delta^4(x - x') \delta^4(\theta - \theta') \quad (14)$$

(and similarly for the complex conjugate), we recover precisely the functional derivative of the action as given in eq. (11). The definition (14) will be useful later.

Finally, a rule of thumb: Whenever you have a nice book-keeping device such as superspace available, you should make as much use of it as you can. Think in terms of superfields, and avoid component fields as much as possible. (There is another rule of thumb, which says that you should not be afraid of expanding in components whenever that becomes convenient - which rule of thumb you should apply depends on the context.)

5.3 Interlude: Covariant derivatives.

Before we tackle gauge theories in superspace, I would like to insert a few words on covariant derivatives, torsion, curvature, and all that. I will be very sketchy; I just want to outline the idea. I will call an operator acting on an algebra (of functions) a derivative if it obeys Leibnitz' rule, i.e. if

$$\partial_a (fg) = \partial_a f g + f \partial_a g \quad (1)$$

where a is some index which I will take to be a vector index (or more generally a superspace index a, A, A'). Differentiating with respect to some coordinate gives an example of a derivative, of course, but a general derivative need not be expressible in that form. In fact, the derivative operator need not be commutative. The commutator is a derivative also, however, so I can write

$$[\partial_a, \partial_b]f = T_{ab}^c \partial_c f \quad (2)$$

By definition, the object on the right hand side is called the torsion (the torsion tensor, since a, b, c are vector indices). I have already shown you an example of a derivative with non-zero 'torsion', namely the superspace covariant derivative D_A .

Now for covariant derivatives. Suppose there are a second kind of objects $k, l, \dots$ on which the derivative can act, and that these objects transform under a structure group of some sort. In general, the object $\partial_a k$ will not transform in the same way as k does (for instance, ∂_a might be a coordinate derivative and the elements of the structure group may be x -dependent). A derivative ∇_a is called a covariant derivative if $\nabla_a k$ in fact does transform in the same way as k .

In Riemannian geometry, things are slightly more complicated, in that the covariant derivative is supposed to take tensors into tensors of a different type, but the definition I just gave is appropriate for the much easier Yang-Mills case. Then you can regard the covariant derivative as being supplied with additional matrix indices (for the structure group to act on):

$$k^i \rightarrow (\nabla_a)^i_j k^j. \quad (3)$$

These indices are rarely written out explicitly.

In this more general situation, the commutator of two covariant derivatives takes the form

$$[\nabla_a, \nabla_b]k = (F_{ab} + T_{ab}^c \nabla_c)k. \quad (4)$$

F_{ab} is called the curvature tensor; like the torsion tensor, it is matrix valued, although I have suppressed the indices. Of course, you are familiar with such curvature tensors from Yang-Mills theory. It turns out that, for the covariant derivative that we need for Yang-Mills theory in superspace, both torsion and curvature are present.

If you write out the statement that

$$\nabla_{[a}\nabla_b]\nabla_c = \nabla_{[a}\nabla_b\nabla_c] = \nabla_{[a}\nabla_{[b}\nabla_{c]}] \quad (5)$$

explicitly, you will prove the very important Bianchi identity

$$\nabla_{[a}F_{bc]} + T_{[ab}{}^d F_{c]d} = 0 \quad (6)$$

which must be obeyed by any tensors that aspire to play the roles of torsion and curvature.

Selecting a suitable basis, one may always write the covariant derivative in the form

$$\nabla_a f = \partial_a f, \quad \nabla_a k = (\partial_a + A_a)k \quad (7)$$

where A_a , which is a matrix valued function, is called the connection. Using eq. (7), you can obtain formulae for the torsion and curvature tensors which automatically obey the Bianchi identity.

Let me be a little bit more explicit about how things transform under some element t of the structure group:

$$\begin{aligned} k &\rightarrow tk & \partial_a k &\rightarrow t \partial_a k + \partial_a t k & \nabla_a k &\rightarrow t \nabla_a k \\ F_{ab} &\rightarrow t F_{ab} t^{-1} & A_a &\rightarrow t A_a t^{-1} + t \partial_a t^{-1} \end{aligned} \quad (8)$$

Note also that, at least in the absence of torsion,

$$F_{ab} = 0 \quad \Leftrightarrow \quad A_a = t \partial_a t^{-1} \quad (9)$$

There may be several structure groups in the problem. In superspace gauge theories, for instance, there are both supersymmetry and colour to take account of. A derivative which is covariant with respect to one of the structure groups need not be covariant with respect to the others.

5.4 Gauge theories in superspace.

5.4.1 Supersymmetric QED.

The easiest kind of gauge theories to deal with is, of course, the Abelian ones, so I will begin with a few words on supersymmetric QED. First we have to select a suitable kind of superfield for the photon multiplet. A real scalar one seems like a possible choice, since it contains a vector field:

$$V(x, \theta, \bar{\theta}) = a + \theta\chi + \bar{\theta}\bar{\chi} + \theta\theta b + \bar{\theta}\bar{\theta}\bar{b} + \theta A\bar{\theta} + \bar{\theta}\bar{\theta}\bar{\lambda} + \bar{\theta}\bar{\theta}\theta\lambda + \theta\theta\bar{\theta}\bar{\theta} D \quad (1)$$

What kind of gauge invariance are we looking for? Well, the kinetic term of the action for a "matter" - i.e. Wess-Zumino - multiplet

$$\int d^4x d^4\theta \bar{\Phi}\Phi \quad (2)$$

is invariant under the rigid transformation $\Phi \rightarrow e^{i\Lambda} \Phi$, and the only reasonable local form of such a gauge transformation is

$$\Phi \rightarrow e^{i\Lambda(x)} \Phi \quad (3)$$

where Λ is a chiral super field (it has to be a superfield). Then the action

$$\int d^4x d^4\theta \bar{\Phi} e^V \Phi \quad (4)$$

is gauge invariant, provided that

$$V \rightarrow V + i(\bar{\Lambda} - \Lambda). \quad (5)$$

At first sight, the action (4) looks unpalatable in the extreme, since it is non-polynomial, but inspection of the gauge transformation (5) shows that it is possible to set all component fields in V to zero by means of gauge transformations, excepting A_a , D - which will become an auxiliary field, just like F in the Wess-Zumino model - and λ . In this gauge, which is called the Wess-Zumino gauge, e^V becomes a polynomial function, so the situation is not that bad.

It is easy to construct gauge invariant "field strengths" from V , viz.

$$W_A = \bar{D}\bar{D}D_A V, \quad \bar{W}_A = D\bar{D}\bar{D}_A V. \quad (6)$$

W_A is obviously chiral; moreover these objects obey

$$DW - \bar{D}\bar{W} = 0. \quad (7)$$

When written out in terms of component fields, this identity becomes the Bianchi identity for F_{ab} .

With the field strengths in hand, it is easy to write down a gauge invariant action for the vector multiplet. Since W_A is chiral, the action involves an integration over half of superspace only, but it is possible, and useful, to rewrite it (remembering eq. (10) on page 44, and ignoring surface terms) as an integral over the entire superspace:

$$\begin{aligned} S &= 1/4 \int d^4x \{ \int d^2\theta WW + \int d^2\bar{\theta} \bar{W}\bar{W} \} = \\ &= 1/4 \int d^4x \{ \int d^2\theta \bar{D}\bar{D}(D_A V W^A) + \int d^2\bar{\theta} D D(\bar{D}^A \bar{V} \bar{W}_{\bar{A}}) \} = \\ &= - \int d^4x d^4\theta (\bar{D}_A D_A V \bar{D}^A D^A V + D_A \bar{D}_A V D^A \bar{D}^A V) \end{aligned} \quad (8)$$

If you write this out in terms of component fields, you will find that the action contains higher derivatives (than two). You can escape this conclusion by means of redefinitions of the component fields.

Varying with respect to V , we find the equation of motion

$$DW + \bar{D}\bar{W} = 0 \quad (9)$$

Non-Abelian gauge theory can be discussed in the same fashion, but in the next section I will treat it using much heavier machinery, which will be good for your education.

5.4.2 Supersymmetric Yang-Mills.

The approach in the last section was the "minimal" one - starting from a guess about what kinds of superfields one should use, one fiddles out what the gauge theory has to look like. In this section, we will try a "maximal" approach - labourious, but instructive. First we collect all the superspace indices into a single Swedish index $a = (a, A, A')$. Then we write down superfields which are to play the role of curvature and torsion tensors, i.e. they are subject to the Bianchi identity

$$\nabla_{[a} F_{\bar{a}\bar{b}]} + T_{[a\bar{a}}{}^{\bar{c}} F_{\bar{c}\bar{b}]} = 0 \quad (1)$$

where the bracket denotes antisymmetrization of even indices and symmetrization of odd ones. A priori there are no more restrictions on these tensors. Of course, supersymmetric Yang-Mills theory does not contain that many gauge covariant component fields, so we expect to be able to impose a few additional constraints. A first restriction is suggested by insisting that the covariant derivative ∇_a should be written in the form

$$\nabla_a = (\partial_a + A_a, D_A + A_A, \bar{D}_{A'} + \bar{A}_{A'}) \quad (2)$$

(Note that as yet there are no reality conditions implied, so A_A and $\bar{A}_{A'}$ are independent fields. Also remember that all the fields are matrix valued, with indices appropriate to the internal symmetry group that we are gauging.) Eq. (2) leads naturally to constraints on the torsion tensor. The only non-zero piece of the torsion is assumed to be the unavoidable

$$T_{AA'}{}^a = 2i \sigma_{AA'}{}^a \quad (3)$$

More constraints are needed. Suppose that, eventually, we want to couple our supersymmetric Yang-Mills model to a "gauge covariantly chiral" matter superfield. This suggests a consistency condition

$$\nabla_A \Phi = 0 \quad \Rightarrow \quad [\nabla_A, \nabla_B] \Phi = F_{AB} \Phi = 0 \quad (4)$$

In fact, we will require that

$$F_{AB} = F_{A'B'} = F_{AA'} = 0 \quad (5)$$

At the moment, this is simply a suggestion, to be checked for usefulness (and to be supplemented with reality conditions). In a later section, we will see that there is a kind of interpretation of these constraints, which

suggests that more can be said about them, but for the moment we will proceed to investigate what additional requirements these constraints imply via the Bianchi identities.

Taking the Swedish index apart into Latin indices, we find that the following non-zero Bianchi identities remain:

$$\begin{aligned}
 \nabla_a F_{bc} + \nabla_c F_{ab} + \nabla_b F_{ca} &= 0 \\
 \nabla_A F_{bc} + \nabla_c F_{Ab} + \nabla_b F_{cA} &= 0 & \bar{\nabla}_A F_{bc} + \nabla_c F_{Ab} + \nabla_b F_{cA} &= 0 \\
 -\nabla_c F_{ab} + \nabla_b F_{ca} &= 0 & -\bar{\nabla}_c F_{aB} + \bar{\nabla}_B F_{ca} &= 0 \\
 \sigma_{CA}{}^d F_{Bd} + \sigma_{BC}{}^d F_{Ad} &= 0 & \sigma_{CA}{}^d F_{B'd} + \sigma_{B'C}{}^d F_{A'd} &= 0 \\
 -\bar{\nabla}_c F_{aB} + \nabla_B F_{Ca} + 2i \sigma_{BC}{}^d F_{ad} &= 0.
 \end{aligned} \tag{6}$$

The second to last line implies that

$$F_{Aa} = i \sigma_{aAA'} \bar{W}^{A'} \quad F_{A'a} = i \sigma_{aAA'} W^A. \tag{7}$$

The last line then allows us to solve for F_{ab} in terms of F_{Aa} and $F_{A'a}$; looking carefully at this and the other identities, one sees that the general solution of eqs. (6) is obtained by constraining the spinorial superfield W_A in the following fashion:

$$\bar{\nabla}_A W_A = 0 = \nabla_A \bar{W}_A, \quad \nabla W - \bar{\nabla} \bar{W} = 0. \tag{8}$$

So the curvature and torsion tensors that I started out with have now boiled down to W_A and $\bar{W}_A$. It is clear that we have found a generalization of the Abelian field strengths from the last section.

However, in order to build an action for supersymmetric Yang-Mills, we have to figure out what the superspace constraints imply for the gauge potentials (and also we have to make decisions concerning reality conditions). From

$$[\nabla_A, \bar{\nabla}_{A'}] = F_{AA'} + T_{AA'}{}^{\dot{A}} \nabla_{\dot{A}} \tag{9}$$

and our choice of constraints, you can see that it is possible to solve for A_a in terms of A_A and $\bar{A}_A$. This is already something. The condition $F_{AB} = F_{A'B} = 0$ implies that the spinorial gauge potentials have the form

$$A_A = e^V D_A e^{-V} \quad \bar{A}_A = e^U \bar{D}_A e^{-U}. \tag{10}$$

where e^V and e^U are group elements and scalar superfields. They are called prepotentials; there are two of them, unlike the single prepotential we had in the last section. This suggests, correctly, that it should be possible to set one of them to zero as a partial choice of gauge. To this end, we examine the gauge symmetries present in the problem. Firstly, we see that

$$F_{\dot{A}\dot{A}} \rightarrow e^{iX} F_{\dot{A}\dot{A}} e^{-iX} \quad \Rightarrow \quad e^V \rightarrow e^{iX} e^V \quad e^U \rightarrow e^{iX} e^U \tag{11}$$

However, there is an extra gauge symmetry present, which $F_{\dot{A}\dot{A}}$ does not feel at all, namely

$$e^V \rightarrow e^V e^{i\bar{\Lambda}} \quad e^U \rightarrow e^U e^{i\Lambda} \tag{12}$$

where Λ is chiral and $\bar{\Lambda}$ antichiral (and independent of Λ , since there are as yet no reality conditions). We can use the gauge symmetry in eq. (11) to set $U = 0 = \bar{A}_A$; there is still some of that gauge symmetry left, since the gauge fixing is unaffected by gauge transformations of the form (11) when X is restricted to be chiral.

Now all that remains to do is to impose a reality condition of some sort. A suitable one is clearly to demand that V should be real. Then we have indeed recovered a set up which reduces to that in the last section, once the colour group is chosen to be Abelian. The remaining gauge symmetry is given by

$$e^V \rightarrow e^{-i\bar{\Lambda}} e^V e^{i\bar{\Lambda}} \tag{13}$$

(where $\bar{\Lambda}$ is now the complex conjugate of Λ , because of the reality condition). The action is given by

$$S = 1/4 \text{Tr} \int d^4x d^2\theta WW + \text{c.c.}, \tag{14}$$

where the trace is over the group matrices, and W_A and its complex conjugate is to be expressed in terms of the prepotential V .

5.5 Survey of superspaces.

One of the things that should be fairly obvious from the last section is that it is by no means straightforward, technically, to find superspace actions. However, as far as $N=1$ superspace in four dimensional space-time is concerned - the Wess-Zumino model, supersymmetric Yang-Mills and supergravity - all problems have been solved already, and the resulting formalism has proved to be very useful, both when it comes to proving general theorems and when explicit calculations are to be performed. (I will show some examples of this in the next chapter.) Of course, models with extended supersymmetry can be formulated in $N=1$ superspace, too. For instance, $N=4$ super Yang-Mills is described by a $N=1$ superspace Yang-Mills model coupled to three chiral superfields. However, it seems natural to introduce extended superspaces, spanned by N different θ 's.

For extended supersymmetry, described in superspaces with N different θ 's, the number of component fields in a superfield grows exponentially with N , and it is not surprising that matters such as choosing the correct constraints on superspace field strengths do become considerably much more involved. For $N=2$ superspace, these matters have nevertheless been sorted out, but the resulting formalism is very cumbersome. For $N>2$, it turns out that something rather surprising happens; it has been proven* that once $N>2$ - with one exception, which has to do with a model that can be obtained by dimensional reduction from ten dimensions - it is impossible to write down an action in terms of N -extended superfields. Any attempt to impose constraints on the field strengths, in order to get rid of superfluous degrees of freedom, turns out to "put the theory on shell"; in other words, the constraints

$$F_{AB}^I + F_{BA}^I = F_{IA'B'}^I + F_{IB'A'}^I = F_{JAA'}^I = 0 \quad (1)$$

say, are equivalent to the equations of motion for $N=3$ super Yang-Mills. This means that such superfields can not be used to formulate action principles. Of course, similar statements can be made for superfields in space-times with dimension higher than four.

The proofs are based on counting arguments, and deeper reasons for this phenomenon are not known. There are some suggestive observations, on the other hand. In some ways, constraints such as those in eq. (1) are reminiscent of the self-duality condition

$${}^*F_{ab} = F_{ab} \quad (2)$$

* B.E.W. Nilsson, Off-Shell Fields for the 10-Dimensional Supersymmetric Yang-Mills Theory, Göteborg preprint 81-6, February 1981.

V.O. Rivelles and J.G. Taylor, Off-Shell No-Go Theorems for Higher Dimensional Supersymmetries and Supergravities, Phys. Lett. **121B** (1983) 37.

in pure Yang-Mills theory, in as much as both types of constraints are algebraic constraints on the field strengths. The self-duality condition also implies the equations of motion, although it is much stronger. In fact, the self-dual Yang-Mills equations, also in Euclidean space where it is non-trivial, can be solved exactly, and they are, in one sense of the word, integrable, although they can not be put in Hamiltonian form. The latter property more or less amounts to the statement that they can not be derived from a reasonable action. (The superspace analogue of the self duality condition is to set $W_A = 0$; provided that we are in Euclidean space, or that the reality condition is relaxed, W_A may still be non-zero.)

The solubility of the self-dual Yang-Mills equations is intimately related to the fact that they can be obtained as compatibility conditions for a system of linear equations. It seems far to optimistic to nurture any hopes about a method which would enable one to construct the general solution of the full Yang-Mills equations. Nevertheless, there is a linear system yielding compatibility conditions which are the superspace constraint equations, and hence - in the $N=3$ case - the full set of equations of motion*. It is

$$\pi^A \nabla_A \Psi = 0 \quad \bar{\mu}^A \bar{\nabla}_{JA'} \Psi = 0 \quad \pi^A \bar{\mu}^A \nabla_{AA'} \Psi = 0 \quad (3)$$

where $\pi^A, \bar{\mu}^A$ are commuting spinors. The fact that two commuting spinors appear, as "spectral parameters", rather than just one, makes the present set up hard to deal with. Nothing much has been achieved in this direction yet, beyond writing down this linear system.

* E. Witten, An Interpretation of Classical Yang-Mills Theory, Phys. Lett **77B** (1978) 394.
J. Avan, Superconformally Covariant Linear System for $N=3,4$ Supersymmetric Yang-Mills Theory in Four Dimensions, Phys. Lett. **190B** (1987) 110.

6. SUPERGRAPHS.

In the last chapter a lot of effort was expended on formalism, but not very much came out of it. Actually, very much has come out of superspace. The main application has been to quantum perturbation theory. Moreover, the very special ultraviolet behaviour which supersymmetric models exhibit in perturbation theory has always been in the center of interest for the subject. So this ought to be the most important chapter in these notes. However, due to my limitations, I will be very brief, and I will avoid entirely some of the tricky points, such as questions concerning regularization (these questions are in fact so difficult that a major mistake was once made by 'tHooft). My aim is simply to show that superspace methods are useful, and then to mention some of the results obtained. You can find out more about "supergraphs", for instance in various publications by Grisaru, Siegel and Rocek *, who developed the superspace Feynman rules into their present form.

6.1 Feynman rules in superspace.

I will confine myself to the Wess-Zumino model, and only the massless case, for additional simplicity. Remembering the manipulations that I did at the end section 5.2.3, we can write the action, coupled to a chiral source J, in either of the forms:

$$S = \int d^4x d^4\theta \bar{\Phi}\Phi + \int d^4x d^2\theta (g/6\Phi^3 + 2^{-1/2}\Phi J) + \dots = \quad (1)$$

$$= \int d^4x d^4\theta (\bar{\Phi}\Phi + (g/12\Phi^2 + 2^{-3/2}J)\frac{DD}{\square}\Phi + (g/12\bar{\Phi}^2 + 2^{-3/2}\bar{J})\frac{\bar{D}\bar{D}}{\square}\bar{\Phi}) .$$

The idea now is to derive the Feynman rules directly in superspace. Evidently, they will look more or less like the Feynman rules for the scalar ϕ^3 -theory, except that there will be a few superspace covariant derivatives hanging around. Concentrating on the free action, we can rewrite the path integral in the usual manner as

$$Z_0[J] \propto \int d\Phi d\bar{\Phi} \exp \int (\bar{\Phi}\Phi + 2^{-3/2}\Phi\frac{DD}{\square}J + 2^{-3/2}\bar{\Phi}\frac{\bar{D}\bar{D}}{\square}\bar{J}) = \quad (2)$$

$$= \int d\Phi d\bar{\Phi} \exp \int ((\bar{\Phi} + 2^{-3/2}\frac{DD}{\square}J)(\Phi + 2^{-3/2}\frac{\bar{D}\bar{D}}{\square}\bar{J}) - 1/8\frac{DD}{\square}J\frac{\bar{D}\bar{D}}{\square}\bar{J}) \propto$$

$$\propto \exp \int J \frac{1}{\square} \bar{J} \quad (2)$$

Since

$$\frac{\delta J(x, \theta, \bar{\theta})}{\delta J(x', \theta', \bar{\theta}')} = -1/4 \bar{D}\bar{D} \delta^4(x-x') \delta^4(\theta - \theta') \quad (3)$$

the propagator becomes

$$\langle T(\Phi(-k, \theta_1, \bar{\theta}_1) \bar{\Phi}(k, \theta_2, \bar{\theta}_2)) \rangle = 1/16 \bar{D}\bar{D} \frac{\delta_{12}}{k^2} DD \quad (4)$$

where

$$\delta_{12} = \delta^2(\theta_1 - \theta_2) \delta^2(\bar{\theta}_1 - \bar{\theta}_2) = (\theta_1 - \theta_2)(\theta_1 - \theta_2)(\bar{\theta}_1 - \bar{\theta}_2)(\bar{\theta}_1 - \bar{\theta}_2) \quad (5)$$

(note that Fourier transformations are done for x only, not for the θ 's). In the massive case, a rather more complicated expression results, since there are $\Phi\Phi$ -terms in the action.

Explicit calculation yields the following important information, to be used later:

$$\begin{aligned} \delta_{12}\delta_{12} &= \delta_{12}D\delta_{12} = \delta_{12}\bar{D}\delta_{12} = \delta_{12}DD\delta_{12} = \delta_{12}D\bar{D}\delta_{12} = \delta_{12}\bar{D}\bar{D}\delta_{12} = \\ &= \delta_{12}DDD\delta_{12} = \delta_{12}D\bar{D}\bar{D}\delta_{12} = 0 \\ \delta_{12}DD\bar{D}\bar{D}\delta_{12} &= 16\delta_{12} \\ D_A(k, \theta_1)\delta_{12} &= -D_A(-k, \theta_2)\delta_{12} \end{aligned} \quad (6)$$

Now we turn to the interacting case:

$$Z[J] = \exp g/3! \int d^4x d^2\theta \left(\frac{\delta}{\delta J} \frac{\delta}{\delta J} \frac{\delta}{\delta J} \right) + \text{c.c.} \exp \int J \frac{1}{\square} \bar{J} . \quad (7)$$

At this point, I will make a drastic simplification, which will save myself a lot of time, although it is not quite fair to you. I will ignore all numerical factors.

Now it is clear that, with every chiral (anti-chiral) vertex, there will be associated an integral over $d^2\theta$ ($d^2\bar{\theta}$). It is convenient to rearrange things a bit, however, before writing down the Feynman rules. There is a propagator between every pair of vertices, and it turns out to be convenient to move the covariant derivatives in eq. (4) from the propagator to the vertices. Hence with every chiral (anti-chiral) vertex there will come a factor $\bar{D}\bar{D}$ (DD) multiplying every line that leaves the vertex; one of which will be used to convert the integration to an

* M.T. Grisaru, W. Siegel and M. Rocek, Improved Methods for Supergraphs, Nucl. Phys. B159 (1979) 429.

M.T. Grisaru, Four Lectures on Supergraphs, Spring School on Supergravity, Trieste 1981.

integration over $d^4\theta$. Similarly, we see from eq. (2) that every external chiral (anti-chiral) line comes with a factor $DD/\square$ ($\bar{D}\bar{D}/\square$), which will cancel an additional $\bar{D}\bar{D}$ (DD). So, we end up with the following Feynman rules - with all numerical factors conveniently suppressed - for how to compute the effective action:

Propagator $\Phi \text{---} \bar{\Phi} \quad \frac{\delta_{12}}{k^2}$

Chiral vertex $\Phi \text{---} \bar{D}\bar{D} \text{---} \Phi \quad \int d^4\theta$

Anti-chiral vertex $\bar{\Phi} \text{---} DD \text{---} \bar{\Phi} \quad \int d^4\bar{\theta}$

For each external chiral (anti-chiral) line, drop a factor $\bar{D}\bar{D}$ (DD).

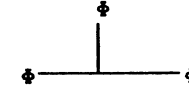
Integrations: $\int \prod_{\text{loops}} d^4p \prod_{\text{ext}} d^4p \delta^4(\Sigma p)$

The next section will be devoted to some explicit calculations. The reason why we compute the effective action, rather than some S-matrix element, is that the former calculation yields just a number. An S-matrix element with N legs, on the other hand, is given by an expression which is a superfield in N different θ 's; it can be obtained by functionally differentiating the effective action with respect to N superfields.

The reason why I do not discuss supersymmetric Yang-Mills models here is that such a discussion is rather involved; it requires a good grasp of the background field method (unless one takes recourse to light front superspace methods).

6.1.1 Sample calculations.

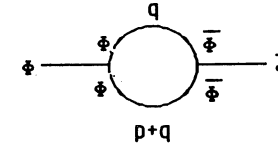
Now I will show how the Feynman rules work in practice. First a tree level calculation, which would have served as a check on the normalization, if I had bothered about numerical factors:



This gives

$$\begin{aligned} & \int d^4p_1 d^4p_2 d^4p_3 d^4\theta \delta^4(\Sigma p_i) \frac{DD}{p_i^2} \Phi(p_1, \theta) \Phi(p_2, \theta) \Phi(p_3, \theta) = \\ & = \int d^4p_1 d^4p_2 d^4p_3 d^2\theta \delta^4(\Sigma p_i) \Phi(p_1, \theta) \Phi(p_2, \theta) \Phi(p_3, \theta). \end{aligned} \quad (1)$$

Next, a loop diagram:



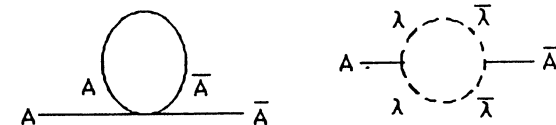
We get:

$$\int d^4p d^4q d^4\theta_1 d^4\theta_2 \Phi(-p, \theta_1) DD \frac{\delta_{12}}{q^2} DD \frac{\delta_{12}}{(p+q)^2} \Phi(p, \theta_2). \quad (2)$$

Using the identities (5) from the previous section, this turns out to be

$$\int d^4p d^4q d^4\theta \bar{\Phi}(-p, \theta) \Phi(p, \theta) \frac{1}{q^2(p+q)^2}. \quad (3)$$

The integral over q is logarithmically divergent. But this is amazing; if we had done the calculation without the superspace technique - starting from the action in section 4.2 - we would certainly have encountered quadratic divergencies, from the diagrams

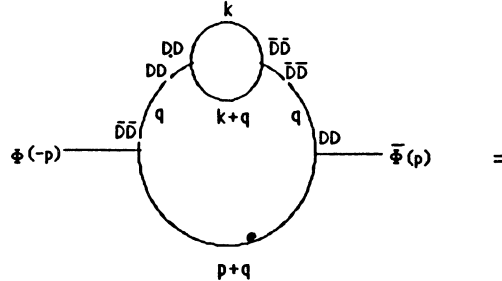


However, since a fermion loop contributes a minus sign, these quadratic divergencies come with the opposite sign. At this point, one might recall an idea from the 30's, which has spent half a century in the dust bin. It

was once suggested that the divergencies in spinor QED might go away, once the contributions from the newly discovered scalar mesons were included, since the meson loops have opposite sign from the electron loops. That did not work out, of course, but it is precisely this mechanism which operates in supersymmetric models. The contributions from the diagrams above are actually included in eq. (3); however, the requirement of supersymmetry has adjusted their relative strengths in precisely such a manner that the quadratic divergencies cancel.

There is still a logarithmic divergence, but one might hope that other supersymmetric models are completely finite, order by order in perturbation theory. There does in fact exist such models, but for the moment we ought to regularize the divergence at hand. It is clear that dimensional regularization is somewhat against the spirit of the problem, since all the superspace spinor algebra was carried out in four dimensions. All I have to say about this is that the issue is tangled, and that I am certainly not the person to say something about it.

While we are at it, we might as well go on to two loops. There is actually only one diagram to compute (this would not be true in the massive model, where there are Φ - Φ -propagators to take account of as well). This, if anything, ought to convince you about the power of supergraphs, since the number of component diagrams that "sit inside" the single supergraph is rather large. Anyway, the diagram is



$$= \int d^4 p d^4 q d^4 k d^4 \theta_1 d^4 \theta_2 d^4 \theta_3 d^4 \theta_4 \bar{\Phi}(-p, \theta_1) D D \frac{\delta_{12}}{q^2} \bar{D} \bar{D} \bar{D} \bar{D} \frac{\delta_{23}}{k^2} D D \times$$

$$\times D D \frac{\delta_{34}}{q^2} \bar{D} \bar{D} \frac{\delta_{23}}{(k+q)^2} \frac{\delta_{14}}{(p+q)^2} \Phi(p, \theta_4) . \quad (4)$$

By means of partial integrations and eq. (6) from the previous section, we can manipulate this expression as follows:

$$\dots = \int d^4 p d^4 q d^4 k d^4 \theta_1 d^4 \theta_2 d^4 \theta_3 d^4 \theta_4 \bar{\Phi}(-p, \theta_1) D D \frac{\delta_{12}}{q^2} \bar{D} \bar{D} \frac{1}{k^2} \times$$

$$\times D D \frac{\delta_{34}}{q^2} \bar{D} \bar{D} \frac{\delta_{23}}{(k+q)^2} \frac{\delta_{14}}{(p+q)^2} \Phi(p, \theta_4) =$$

$$= \int d^4 p d^4 q d^4 k d^4 \theta_1 d^4 \theta_2 \bar{\Phi}(-p, \theta_1) D D \frac{\delta_{12}}{q^2} \bar{D} \bar{D} \frac{1}{k^2} D D \frac{\delta_{21}}{\bar{D} \bar{D}} \times$$

$$\times \frac{1}{(k+q)^2} \frac{1}{(p+q)^2} \Phi(p, \theta_1) = \quad (5)$$

$$= \int d^4 p d^4 q d^4 k d^4 \theta_1 d^4 \theta_2 \bar{\Phi}(-p, \theta_1) \frac{\delta_{12}}{q^2} D D \bar{D} \bar{D} D D \bar{D} \bar{D} \frac{\delta_{21}}{q^2} \times$$

$$\times \frac{1}{k^2} \frac{1}{(k+q)^2} \frac{1}{(p+q)^2} \Phi(p, \theta_1) =$$

$$= \int d^4 p d^4 q d^4 k d^4 \theta \bar{\Phi}(-p, \theta) \frac{1}{q^2} \frac{1}{k^2} \frac{1}{(k+q)^2} \frac{1}{(p+q)^2} \Phi(p, \theta) .$$

This easy calculation automatically sums all the component Feynman diagrams that contribute to the two point function at two loops. There is more to superspace perturbation theory than mere calculational convenience, however, as I will discuss in the next section.

6.2 Non-renormalization theorems.

If you have followed the calculations in the previous section, you will believe the central theorem of superspace perturbation theory, which is that the effective action Γ is an expression of the form

$$\Gamma = \int d^4x_1 \dots d^4x_N d^4\theta G(x_1, \dots, x_N) \times (\text{Polynomial in the fields and their derivatives}) \quad (1)$$

What matters here is the way in which the θ 's appear. The theorem says that eqs. (6) of sect. 6.1, together with partial integrations, are enough to ensure that Γ always takes the form of an integral over a single θ , and always an integral over the entire superspace - as opposed to the chiral subspace - and moreover that the only θ -dependence in the integrand comes from the external fields. A number of striking conclusions about how supersymmetric models (in four dimensions - some of the conclusions are invalid in, say, two dimensions, where superspace works somewhat differently) behave in perturbation theory follow immediately:

1. All vacuum bubbles vanish identically. This follows since they contain no external fields, and hence no θ 's in the integrand, so that they are killed by the integration over θ . It means that the energy density of the physical vacuum is zero, as compared with the bare vacuum, while it is - naively, at least - minus infinity in generic models. This is not quite the same thing as saying that there is no energy difference between the physical and the bare vacua; however, it turns out that the level shift is indeed zero, perturbatively, provided that there are at least two supersymmetries in the model*.

2. Chiral terms in the superspace action - i.e. terms which involve an integration over $d^2\theta$ only - will receive no quantum corrections. This has important consequences. Consider the action for the Wess-Zumino model as an example:

$$\int d^4x d^4\theta \bar{\Phi} \Phi + \left(\int d^4x d^2\theta (\mu \Phi^2 + g/6 \Phi^3) + \text{c.c.} \right). \quad (2)$$

One would normally expect three independent renormalizations to be necessary for this model:

$$\Phi_0 = Z^{1/2} \Phi \quad g_0 = Z_g g \quad m_0 = Z_m m \quad (3)$$

However, the non-renormalization theorem just mentioned implies that

$$Z^{3/2} Z_g = 1 \quad Z Z_m = 1. \quad (4)$$

*I. Bengtsson and O. Lindgren, Extended Supersymmetry and the Vacuum, Phys. Lett. **127B** (1983) 65.

Hence only one independent renormalization constant is needed for the model. (This is of course reminiscent of Yang-Mills theory, as renormalized using the background field method.) Note that this means that the β -function of the model can be computed by computing the two-point function, so that the more elaborate calculation of the three point function can - for this purpose - be avoided.

3. Since $\dim [d^4\theta] = 2$, simple power counting leads to the conclusion that there will be at most logarithmic divergencies in perturbation theory, provided that there are at most quadratic divergencies among the component diagrams. One would expect even more dramatic cancellations of divergencies among the component diagrams in models with extended supersymmetry, since then there are "more θ 's"; however, since extended superspace does not work smoothly, it was for quite some time an open question whether there indeed exists a model without any divergencies in perturbation theory. I will say a few words about this in the next section.

4. Provided that supersymmetry is unbroken at the tree level, it will not be spontaneously broken in perturbation theory either. Of course, this very important conclusion can not be obvious unless you know what the conditions for spontaneous breakdown of supersymmetry are, and I have not told you that, so I just ask you to believe this statement. (A good place to learn about this topic is two papers by Witten*.)

These conclusions are clearly remarkable; they are also enough to provide some grounds for belief in the relevance of supersymmetry for particle physics. I will spend a few words on this presently.

* E. Witten, Dynamical Breaking of Supersymmetry, Nucl. Phys. **B185** (1981) 513.
E. Witten, Constraints on Supersymmetry Breaking, Nucl. Phys. **B202** (1982) 253.

6.2.1 Finite models.

Returning to point 3 of the last section, it is clear that the mechanism for removing divergencies in perturbation theory that we uncovered - any Feynman diagram contains a superspace measure having a certain dimension of length - will work only if the coupling constant of the model is dimensionless; it will be quite helpless in the face of gravity, although it took some time before this point was understood. The best candidate for a finite model should be the $N=4$ supersymmetric Yang-Mills model. It has been checked by explicit calculation that this model is indeed finite (in the ultraviolet - infrared divergencies are another matter) up to the three loop level. A number of different proofs - using light front superspace*, $N=2$ superspace, and consideration of anomalies respectively - that this property holds to all orders have also been published. Unfortunately, there are weak points in all of these proofs - points which are rarely discussed in the papers which give the proofs - which means that it is difficult for a non-expert to tell whether the result actually has been proven. However, at least the light front proof has stood up to later scrutiny**, and the following statements*** are hardly in doubt.

The $N=4$ model is finite order by order in perturbation theory. You can regard the $N=4$ model as an $N=2$ supersymmetric Yang-Mills model coupled to an $N=2$ hypermultiplet in a specific way, and it turns out to be possible to change the way in which the hypermultiplet couples to the Yang-Mills multiplet in certain ways, without disturbing finiteness, so that finite models having $N=2$ supersymmetry only result. Moreover it turns out to be possible to add certain "soft" terms (not necessarily supersymmetric), i.e. mass terms and terms cubic in the scalar fields, in such a way that finiteness is preserved.

Note that perturbative finiteness of a quantum field theory implies that its β -function is zero, perturbatively. For a Yang-Mills model, this means that the conformal symmetry of the classical action is preserved by quantum corrections.

One can turn the argument around: Starting out with a fairly general gauge theory containing spinor and scalar fields, and demanding that perturbative divergencies should cancel, one finds that the model has to be supersymmetric up to possible "soft" terms. The same conclusion follows, again in a large class of models, from the weaker requirement that all

* S. Mandelstam, Light-Cone Superspace and the Ultraviolet Finiteness of the $N=4$ Model, Nucl. Phys. B213 (1983) 149.

L. Brink, O. Lindgren and B.E.W. Nilsson, The Ultraviolet Finiteness of the $N=4$ Yang-Mills Theory, Phys. Lett. 123B (1983) 323.

**J.C. Taylor and H.C. Lee, The Light-Cone Gauge and Finiteness of the $N=4$ Supersymmetric Theory, Phys. Lett. 185B (1987) 363.

***P.C. West, Supersymmetry and Finiteness, Shelter Island II, June 1983.

again in a large class of models, from the weaker requirement that all quadratic divergencies, caused by the scalars, should cancel*.

* N.G. Deshpande, R.J. Johnson and E. Ma, Does the Cancellation of Quadratic Divergencies Imply Supersymmetry?, Phys. Rev. D29 (1984) 2851.
W. Lucha and H. Neufeld, Finite Quantum Field Theories, Phys. Rev. D34 (1986) 1089.

6.3 On the absence of quadratic divergencies.

(I do not fully understand the following argument, nor am I completely sure that it can be understood.)

The standard model predicts that there exists something which has never been observed: Fundamental scalar particles. These are usually thought to have masses below 1 TeV, say; otherwise their couplings have to be so strong that perturbation theory breaks down, which is at least undesirable from the point of view of the physicist. Now 1 TeV is a very low energy indeed compared to the Planck energy, which is often - very glibly - thought to be the "next" energy scale in Nature. Therefore an explanation for the lightness of these scalars - the Higgses - is being sought for.

Renormalized quantum field theory can not predict these masses, since the value of the physical masses are free input parameters in such a theory. However, from a physical point of view, it seems to make a certain amount of sense to regard the bare masses as input parameters. The theory then contains a parameter Λ which "regulates" the theory. The bare parameters depend on Λ in a very precise way, so that all physical parameters come out to be polynomials in Λ^{-1} , which means that physical quantities do not depend on Λ when the limit $\Lambda \rightarrow \infty$ is taken. This is what it means for a theory to be renormalizable. An intuitively appealing choice of Λ is to say that all momentum integrals should be "cut off" at some energy scale where the model is assumed to have lost its validity, and a different theory takes over. So we assume that the cut off (Λ) is in fact the Planck energy.

Now it is considered "natural" to demand that the physical masses are not all that sensitive to what values one chooses for the bare masses. However, this requirement is inconsistent with the presence of quadratic divergencies in the perturbation expansion. Suppose that the input parameters are the dimensionless quantities

$$\mu_0 = m_0/\Lambda \quad (1)$$

and that the quantum corrections to the masses take the form

$$m^2 = m_0^2 + \Lambda^2 g_0^2. \quad (2)$$

Then we find that

$$\mu_0^2 = m^2/\Lambda^2 - g_0^2. \quad (3)$$

If the physical mass $m \approx 1$ TeV and $\Lambda \approx 10^{16}$ TeV, this means that μ_0 has to

be adjusted to within one part to the 10^{32} in order to give the correct value of m , which is clearly unnatural from the point of view adopted. A logarithmic dependence on Λ is not nearly as bad, and therefore we conclude that a "natural" renormalizable theory has only logarithmic divergencies in its perturbation expansion.

Two different ways to avoid quadratic divergencies in the Higgs sector of the standard model have been suggested. The first starts with the observation that scalar particles - mesons - occur in QCD as well, but there they are not accompanied by quadratic divergencies since they are bound states of quarks. So the suggestion is that the Higgs particle is actually a bound state of fermions, which is kept together by strong "technicolour" forces. The other way starts with the observation that fundamental scalars are present in supersymmetric models, but again without quadratic divergencies, since supersymmetry ensures that the latter cancel out.

There is a slightly different line of argument which leads to the conclusion that supersymmetry might explain why very light - as compared to the Planck mass - fundamental scalars exist. Usually, an explanation of the smallness of some mass amounts to finding a symmetry which requires that the mass is exactly zero, and then an argument for why this symmetry is very softly broken. For fermions, chiral symmetry might do, but no ordinary symmetry is known to require a scalar mass to vanish. Supersymmetry gets around this by requiring that the scalar mass should equal the mass of the fermion. The mass of the scalar can then become non-zero only when supersymmetry breaks down. However, since spontaneous breakdown of supersymmetry does not occur in perturbation theory, it has to have a non-perturbative origin (unless it happens already at the tree level, of course). So a suggested explanation for the small mass of the Higgs particle is that the Higgs particle is related to a massless fermion by a supersymmetry, which is broken by a very weak non-perturbative effect.

You don't have to believe this if you don't want to.

7. SUPERSYMMETRIC QUANTUM MECHANICS.*

8. SUPERGRAVITY.

Ask Bengt.

* P. Salomonson and J.W. van Holten, Fermionic Coordinates and Supersymmetry in Quantum Mechanics, Nucl. Phys. **B196** (1982) 509.