

Model independent sparticle mass measurements at ATLAS

Christopher Gorham Lester

St John's College



A dissertation submitted to the University of Cambridge
for the degree of Doctor of Philosophy
December 2001

Model independent sparticle mass measurements at ATLAS

Christopher Gorham Lester

Abstract

This thesis presents a method by which masses of supersymmetric particles might be deduced from end points of kinematic distributions obtained by the ATLAS experiment. An assumption at the core of the analysis is that $m_{\tilde{\chi}_1^0} < m_{\tilde{l}_R} < m_{\tilde{\chi}_2^0} < m_{\tilde{q}}$. Care is taken to keep the number of *additional* model-dependent assumptions as small as possible, so that the otherwise model-*independent* nature of the method may be emphasized. Strengths and weaknesses of the method, as well as its degree of model independence, are investigated by applying it to events simulated in accordance with two different supersymmetric models, and corresponding to an integrated luminosity of 100 fb^{-1} . The method is shown to be largely successful in both models; principally constraining mass differences, but also allowing absolute masses to be determined with a statistical and energy scale error of 20 to 30 GeV in the models considered. It is observed that special care should be taken when interpreting results concerning the slepton masses. Under some circumstances the model independence of the method can result in the light slepton mass failing to be *uniquely* determined. However, it is also demonstrated that, when this problem does not arise, it can even be possible to obtain a “bonus” measurement of the mass of a heavier slepton such as $\tilde{l}_L$.

Declaration

This dissertation is the result of my own work, except where explicit reference is made to the work of others, and has not been submitted for another qualification to this or any other university.

Christopher Gorham Lester

Acknowledgements

I would like to thank the particle physics and astronomy research council (PPARC) for providing financial support during the first nine terms of this degree (studentship number PPA/S/S/1998/02665). I would also like to thank the United Kingdom registered companies (numbers 03447836, 00177136, 01519749, 03796653, 00411181, 04006695, 03712991, 03822566, 03655231 and 03822970 (now insolvent)) whose on- and off-shore operations funded the tenth term.

Contents

1	Introduction	1
2	Theory	5
2.1	Introduction	5
2.2	The Unbroken Standard Model	6
2.3	The Standard Model	7
2.4	Shortcomings of the Standard Model	10
2.5	Supersymmetry	11
2.6	MSSM	13
2.7	Supersymmetry breaking in the MSSM	15
2.8	Minimal supergravity models	16
2.9	Other models	17
2.9.1	String theory	18
2.9.2	Superstrings	20
2.9.3	The transition to the effective quantum field theory	21
2.9.4	Optimized string models	23
2.10	Summary	25
3	The ATLAS Experiment	27
3.1	The LHC	27
3.2	The ATLAS experiment	28
3.2.1	Coordinate systems	31

3.2.2	Inner Detector	32
3.2.2.1	Magnetic field of the inner detector	34
3.2.2.2	Pixel detector	35
3.2.2.3	Semiconductor tracker	35
3.2.2.4	Transition radiation tracker	36
3.2.3	Calorimeters	38
3.2.3.1	Electromagnetic Calorimeter	39
3.2.3.2	Hadronic calorimeter	41
3.2.4	Muon spectrometer	41
3.3	Summary	42
4	The ATLFast Detector Simulation	45
4.1	Introduction	45
4.2	Simulation and reconstruction method	45
4.2.1	“Calorimetric cells” and “clusters”	45
4.2.2	Muons, electrons and photons	47
4.2.3	Jets	52
4.2.4	Missing transverse momentum	53
4.3	Resolutions	54
4.3.1	Electron and photon resolution	54
4.3.2	Muon resolution	55
4.3.3	Resolution for jets and unclustered calorimetric cells	59
5	Analysis	61
5.1	Introduction	61
5.1.1	(Bayesian) measurements of scale	62
5.1.2	“Show you might be right”	64
5.1.3	“Model independent” techniques	65
5.1.4	“Assume a model with few parameters”	68

5.1.5	Summary	70
5.2	Aim of the analysis	70
5.3	Summary of method	72
5.4	Test models: S5 and O1.	74
5.5	The decay chains	77
5.6	Edges, endpoints and thresholds	80
5.6.1	The l^+l^- edge	80
5.6.2	The l^+l^-q edge	80
5.6.3	The l^+l^-q threshold	82
5.6.4	The hq and Zq edges (Xq edges)	82
5.6.5	Near, far, high and low $l^\pm q$ edges	84
5.6.6	Dislepton production, $M_{T2}^{\max}(\chi)$, ΔM and the $\tilde{l}\tilde{l}$ edge	87
5.6.7	Notes concerning interpretation of endpoint positions	93
5.7	Event generation and pile-up	95
5.8	Event selection cuts	96
5.8.1	Cut summaries, grouped by observable	99
5.9	Edge resolutions	103
5.9.1	Resolutions for “boost invariant” edges	103
5.9.2	ΔM edge resolution	108
5.10	Modelling the endpoint measurement process	116
5.10.1	Reconstruction of sparticle masses	117
5.11	Basic reconstruction	119
5.12	Extended reconstruction	124
5.13	Reconstruction using $\Delta M^{\max}$	132
6	Conclusions	137
A	The Supersymmetry Scale	141
B	Higgs doublet multiplicities	145

C	Comparing the supersymmetric and SM Higgs potentials	149
D	Rapidity and Pseudorapidity	153
E	Edges	155
E.1	Successive two body decays	155
E.2	l^+l^- edge	157
E.3	Xq edge	157
E.4	l^+l^-q edge	157
E.4.1	$0 < \lambda < 1$	159
E.4.2	$\lambda = 0$	160
E.4.3	$\lambda = 1$	160
E.4.4	Summary	161
E.5	l^+l^-q threshold	162
E.6	High, low, near and far $l^\pm q$ edges	164
E.7	$\tilde{l}\tilde{l}$ edge	169
E.8	Phase space for $c \rightarrow Zb$ followed by $b \rightarrow al^+l^-$	171
F	Abbreviations	173

List of Figures

2.1	One-loop contributions to the higgs mass	11
3.1	The LHC above and below ground	29
3.2	Higgs discovery significance	31
3.3	Cross section through inner detector	33
3.4	Magnetic field inside the inner detector	34
3.5	Electron-pion separation in the TRT	37
3.6	ATLAS calorimeters	38
3.7	ATLAS electromagnetic calorimeter	39
3.8	Radiation lengths in the ECAL	40
3.9	Side view of muon spectrometer	43
3.10	End view of muon spectrometer	44
4.1	Acceptance of simulated lepton pairs	50
4.2	Simplified charged particle track in constant B -field	56
4.3	Standalone muon momentum resolution	57
5.1	Supersymmetric mass scale	63
5.2	Two successive two-body decays	66
5.3	Sequential squark decay.	77
5.4	Branched squark decays through Z and higgs bosons.	78
5.5	Dislepton production process	78
5.6	Simulated invariant mass distributions at S5 after cuts	104

5.7	Simulated invariant mass distributions at O1 after cuts	105
5.8	The line shapes fitted to boost invariant edges	107
5.9	Signal dilepton events	110
5.10	SM background for ΔM	111
5.11	Supersymmetric background for ΔM	112
5.12	ΔM distributions, all events	113
5.13	Slepton/neutralino resolution, “basic” analysis	119
5.14	Slepton/neutralino resolution, “cut-down” analysis	121
5.15	Slepton/neutralino resolution, “extended” analysis	124
5.16	Sparticle mass resolutions at S5	126
5.17	Sparticle mass resolutions at O1	127
5.18	Breakdown of sparticle mass resolutions at O1	129
5.19	Slepton/neutralino resolution; the ΔM constraint	132
A.1	Size of loop corrections to higgs mass-squared in a broken supersymmetry	143
C.1	Parts of the supersymmetric Higgs potential	149
E.1	Two successive two-body decays (1)	155
E.2	Two successive two-body decays (2)	156
E.3	Three successive two-body decays	158

List of Tables

2.1	The Standard Model fermions	8
3.1	LHC beam and accelerator parameters	30
4.1	Muon, electron and photon reconstruction parameters	48
5.1	Comparison of sparticle and higgs spectrum of O1 and S5	76
5.2	Kinematic edge positions	86
5.3	Standard Model and supersymmetric production cross sections.	95
5.4	Endpoints and associated fit uncertainties for S5 and O1	106
5.5	Widths of sparticle mass distributions	131

Chapter 1

Introduction

Supersymmetry, introduced in Section 2.5, is the most commonly discussed extension of the Standard Model (SM). Since supersymmetry began to emerge in the early 1970s, a large base of knowledge has grown up around it, and the theory can now be said to have a strong theoretical footing. Sadly, for the theory, the theoretical grounding is not matched by experimental confirmation. So far, no experiment has yet produced any compelling evidence for the existence of “sparticles”. Proponents of supersymmetry have called for colliders with higher beam energies, arguing that searches failed in the past simply because the sparticles were too heavy to be produced by the experiments. The large hadron collider (LHC) is at present under construction at the European laboratory for particle physics (CERN), and by colliding two 7 TeV proton beams, it is expected to address the need for higher energies. Will the energy available at the LHC be sufficient to make a statement one way or the other about supersymmetry? Most commentators are prepared to answer this question with a “yes”, for reasons which will be discussed later,^a and so the LHC is regarded as defining “make or break” time for supersymmetry. It is thus of great interest to investigate methods for constraining supersymmetry and,

^aSee Section 2.5.

in particular, methods for measuring the masses of sparticles.

This thesis presents a method by which masses of supersymmetric particles might be deduced by the ATLAS experiment, which is one of the two general purpose particle detectors which are expected to observe collisions at the LHC. The ATLAS detector and the LHC are described in Chapter 3.

The method used in this thesis to measure sparticle masses takes “end points of kinematic distributions” as its primary observable quantities. The selling point of this method is its “model independence”, as only the existence of a small set of decay chains needs to be assumed. The key assumption at the core of the analysis is that

$$m_{\tilde{\chi}_1^0} < m_{\tilde{l}_R} < m_{\tilde{\chi}_2^0} < m_{\tilde{q}}. \quad (1.1)$$

It should be remembered, though, that model independence is a term which must be used carefully. What is seen as model specific in one context may be regarded as a standard assumption in many others; discussion of this and related issues takes place in Chapter 5. Having motivated the core method of this thesis on model independent grounds, Chapter 5 goes on to emphasise the importance of not polluting this model independence by the inclusion of model dependent assumptions via a back door. In particular, it is stressed that (beyond the hierarchical assumption in (1.1)) no further assumptions should be made concerning the relationships of the sparticle masses. The identification and removal of such assumptions where they have existed in the past is a central theme of the thesis.

Another innovative component of this thesis concerns a “test” of the claimed model independence. The strengths and weaknesses of the method, as well as its degree of model independence, are investigated by applying it to events simulated in accordance with *two* different supersymmetric models; one a standard minimal supergravity

(mSUGRA) point, and the other an “optimised string model” featuring heavier sleptons. The method is shown to be largely successful in both models; principally constraining mass differences, but also allowing absolute masses to be determined with a statistical and energy scale error of 20 to 30 GeV in the models considered. Although the performance of the method is indeed found to be similar in the two models, it is the subtle differences which are more interesting. It is observed that special care should be taken when interpreting results concerning the slepton masses. Under some circumstances the model independence of the method can result in the light-slepton’s mass failing to be *uniquely* determined. However, it is also demonstrated that, when this problem does not arise, it can even be possible to obtain a “bonus” measurement of the mass of a heavier slepton such as $\tilde{l}_L$.

Finally, it is noted that while the analysis presented in this thesis was under study, the author was also involved in other work with direct consequences for the present day ATLAS collaboration. This work was primarily in the field of “ATLAS Generic Detector Description” (AGDD). Various experimental activities require accurate descriptions of the detector: simulation, reconstruction, visualisation and calibration to name just a few. Each has different needs, but all must use consistent data. The role of the AGDD group is to develop the framework and all the services required to provide detector descriptions, in appropriate forms, to applications which need them. The primary exchange format for the AGDD data is defined by the “AGDD-XML” syntax. In collaboration with^b Stan Bentvelsen, Christian Arnault, Jean-Francois Laporte, Steven Goldfarb, Marc Virchaux, Tom LeCompte and others, the author was involved in defining the AGDD-XML syntax [1]. In addition, the author developed the first representation of the Semiconductor Tracker (SCT) to use this format. Also, working with Szymon Gadomski, a study was performed in which an AGDD-XML based description of SCT modules in a testbeam was used to compare testbeam data with the results of montecarlo. This served the dual

^b(in no particular order)

purpose of validating montecarlo simulation of testbeam data, while simultaneously acting as a testbed application for AGDD-XML.

Chapter 2

Theory

2.1 Introduction

At present, the SM is the physical theory which accounts for the widest range of experimental data in the field of High Energy Physics. It is only recently, when studies of atmospheric and solar neutrinos have begun to provide evidence of neutrino oscillations, that the predictions of the SM have been shown to be in disagreement with experiment.

This chapter will briefly summarise some of the theory behind the SM, and follow this with an outline of some of the reasons one might be dissatisfied with the SM despite its many successes. Subsequently, supersymmetry will be championed as a candidate for a better theory.

The SM is a gauge field theory which describes nature in terms of matter particles, force carriers and a higgs boson.^a To understand how it is extended to include super-

^aWith regard to the matter of the capitalization of personal names and words which are derived from them, I shall use upper case when the name refers to the person himself (“Joseph Louis Lagrange was born in Turin in 1736”), upper case in the genitive (“The development of the calculus of variations was primarily the result of Lagrange’s efforts”), upper case in adjectival constructions (“Put simply, Keynesian macroeconomics is the theory that shows how a market-based capitalist economy may reach

symmetry it is instructive to first return to the simpler unbroken theory from which the SM may be derived.

2.2 The Unbroken Standard Model

The unbroken SM is a gauge field theory in which two types of matter (massless spin- $\frac{1}{2}$ quarks and leptons) interact via three forces arising from massless vector boson exchange. The three charges (gauge quantum numbers) associated with each of the forces in this theory are known as hypercharge, weak isospin and colour. Associated with each of these charges, in turn, are the gauge groups: $U(1)$, $SU(2)_L$ and $SU(3)$, under whose local transformations the action of the theory is required to remain invariant. As with any gauge theory, the number of gauge bosons required to mediate any particular force is precisely the number of generators of the corresponding symmetry group. Thus in the unbroken SM we have: one $U(1)$ “hypercharge boson” (B); three $SU(2)$ “isospin bosons” (W_1 , W_2 and W_3); and eight $SU(3)$ “colour” gluons (g_i).

Of course, the unbroken SM is no good at describing nature because experiments observe *massive* particles. Notably, experiments demand particularly massive vector bosons (~ 100 GeV) for the weak interaction, massive leptons, and quarks with masses up to 174 GeV. One might hope to fix the fermion and gauge boson masses by adding mass terms for them to the lagrangian of the unbroken SM, but unfortunately it can be shown that this is not possible without breaking gauge invariance and consequently the renormalizability of the theory. Finding an acceptable way of fitting masses into the framework of the unbroken SM is thus of great importance. The generally accepted

equilibrium with large scale unemployment and how government spending may be used to raise it out of this to a new equilibrium at the full-employment level of output” or “Many texts on Lagrangian mechanics were published in Victorian England”) but in lower case when it refers to a noun which takes its name from a proper name (“Neither iron nor einsteinium are gases at room temperature”, “The oxygen atom is nothing like the higgs boson”, “Enrico Fermi’s name has been incorporated into the term fermion” and “Partial derivatives of the lagrangian occur in Euler-Lagrange equations”).

method is that of Higgs [2] which introduces to the theory an additional scalar field. Provided this scalar field somehow obtains a non-zero expectation in the vacuum state, it imparts effective masses to other particles via their interactions with it.

But the appropriate time to discuss Higgs's method is later as, just before the introduction of this mechanism, the routes to supersymmetry and the broken (i.e. usual) SM begin to diverge. As shall be seen later, the *immediate* introduction of a higgs field leads to the “usual” SM, while in supersymmetric theories it is better if this so-called electro-weak symmetry breaking (EWSB) is delayed until *after* extra supersymmetric particles have been introduced. It may be helpful to note that it is this “early” separation (i.e. pre-EWSB) between the SM and supersymmetric models, which will make phrases like “the superpartner of the B -boson” much easier to give meaning to than “the superpartner of the photon”.

Table 2.1 lists the non-mass quantum numbers of all the matter fields which are required to be present in a good theory of nature and which are also present in the unbroken SM. The contents of this table will be joined by others as the model is extended further. Note that the right handed quarks and leptons have no isospin, and thus do not couple to isospin bosons.

2.3 The Standard Model

As mentioned earlier, the particle content of the full SM is completed by the addition of a scalar higgs field ϕ to those already present in the unbroken SM. The field which is added to the model is, in actual fact, a single $SU(2)$ doublet containing two complex scalars having four degrees of freedom in all. During the process of EWSB, however, three of those degrees of freedom are “converted into masses for the W^+ , W^- and Z ”. It is the single remaining scalar field which is known as the higgs boson. In order to fix

Fermion	Hypercharge Y	Isospin	I_3	$C \times I \times Y^W$	Electric charge $Q = I_3 + Y^W$
u_R	$+\frac{4}{3}$	0	0	$u = (\bar{\mathbf{3}}, \mathbf{1}, +\frac{2}{3})$	$+\frac{2}{3}$
d_R	$-\frac{2}{3}$	0	0	$d = (\bar{\mathbf{3}}, \mathbf{1}, -\frac{1}{3})$	$-\frac{1}{3}$
$\begin{pmatrix} u \\ d \end{pmatrix}_L$	$+\frac{1}{3}$	$+\frac{1}{2}$	$+\frac{1}{2}$ $-\frac{1}{2}$	$Q = (\mathbf{3}, \mathbf{2}, +\frac{1}{6})$	$+\frac{2}{3}$ $-\frac{1}{3}$
e^-_R	-2	0	0	$e = (\mathbf{1}, \mathbf{1}, -1)$	-1
$\begin{pmatrix} \nu_e \\ e^- \end{pmatrix}_L$	-1	$+\frac{1}{2}$	$+\frac{1}{2}$ $-\frac{1}{2}$	$L = (\mathbf{1}, \mathbf{2}, -\frac{1}{2})$	0 -1

Table 2.1: *The quantum numbers of the SM fermions.*

the transformation properties of the higgs field under the $U(1)$ symmetry it is necessary to assign it a hypercharge. Within the SM, ϕ is conventionally assigned a hypercharge of $Y = +1$, and thus ϕ^* takes hypercharge -1 . This choice is not much more than convention, though, as the opposite assignments could also be made to work, provided that the roles of the ϕ and ϕ^* were also interchanged.^b Naturally, accommodation of ϕ within the lagrangian of the SM requires the addition of new terms which permit, at the very least, ϕ to propagate (for example Equation (2.2) later). Part of the beauty of this extension to the theory is that by merely requiring that the kinetic term is gauge invariant (by using the covariant derivative D_μ), the lagrangian density naturally acquires extra mass terms for the gauge bosons proportional to the vacuum expectation value (VEV) of the higgs field, should it have one. Finally, when the lagrangian is completed by the addition of the remaining renormalizable terms involving the higgs field, it is reassuring to find that there are indeed terms which permit the higgs field to acquire a non-zero VEV, as desired. The higgs terms in the SM can be broken down

^bThe relevance of this point will become clearer later when supersymmetry will be seen to require two higgs doublets. One of these will be analogous to ϕ and the other to ϕ^* , and so one will have hypercharge $+1$ and the other -1 .

into three parts, $\mathcal{L}_{\text{higgs}} = \mathcal{L}_{\text{Higgs potential}} + \mathcal{L}_{\text{kinetic/gauge}} + \mathcal{L}_{\text{fermion masses}}$, given by:

$$\mathcal{L}_{\text{Higgs potential}} = \mu^2 \phi^\dagger \phi - \lambda (\phi^\dagger \phi)^2, \quad (2.1)$$

$$\mathcal{L}_{\text{kinetic/gauge}} = (D_\mu \phi)^\dagger (D^\mu \phi) \quad \text{and} \quad (2.2)$$

$$\begin{aligned} -\mathcal{L}_{\text{fermion masses}} &= \sum_{Y_f + Y_{f'} + Y_\phi = 0} \bar{\psi}_{f,L} g_{ff'} \phi \psi_{f',R} + \text{h.c.} \\ &+ \sum_{Y_f + Y_{f'} + Y_{\phi^*} = 0} \bar{\psi}_{f,L} g'_{ff'} i \sigma_2 \phi^* \psi_{f',R} + \text{h.c.}, \end{aligned} \quad (2.3)$$

where $\psi_{f,L/R} = \frac{1}{2}(1 \mp \gamma_5)\psi$ are the left/right handed components of the fermion fields and ϕ is the higgs field. The higgs-fermion-fermion couplings, $g_{ff'}$ and $g'_{ff'}$, are only non-zero for those combinations making hypercharge-neutral vertices. (See also discussion in Appendix B.) Together, these make up all the possible renormalizable gauge- and lorentz-invariant terms involving the higgs field which the lagrangian will admit. As far as EWSB is concerned, $\mathcal{L}_{\text{Higgs potential}}$ and $\mathcal{L}_{\text{kinetic/gauge}}$ are arguably the most important terms: $\mathcal{L}_{\text{Higgs potential}}$ allows the higgs field to obtain a non-zero vacuum expectation value $\phi = \begin{pmatrix} 0 \\ v \end{pmatrix}$, while $\mathcal{L}_{\text{kinetic/gauge}}$ contains (non-diagonal) mass terms for the hypercharge and isospin bosons as well as the terms which allow the higgs boson to propagate. After a change of basis, constructed so as to bring the mass matrix into diagonal form, $\{W^1, W^2\} \rightarrow \{W^+, W^-\}$ and $\{B, W^3\} \rightarrow \{A, Z\}$, the gauge fields can then be brought to the more familiar mass eigenstates of the massive $W^\pm$ and Z and the massless photon.

Although the mechanism by which the terms in $\mathcal{L}_{\text{fermion masses}}$ generate masses for the “heavy” fermions is left to a discussion in Appendix B, it should be noted that the form of the SM which has been outlined *here* cannot describe massive neutrinos (notably the particle content of the model lacks a right handed neutrino). There is no shortage of ways of incorporating neutrino masses and neutrino mixing into the SM lagrangian directly, but none will be followed here, partly for simplicity and partly because neutrino mixing might emerge just as well as a side effect of a more fundamental theory, such as

supersymmetry.^c Nevertheless, the approximation should have no effect on the results of this thesis as most allowable mixing models predict no observable effects at the LHC. There are, of course, some interesting exceptions such as [4].

2.4 Shortcomings of the Standard Model

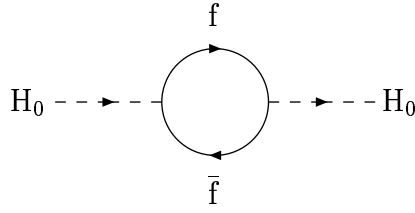
As it stands, the SM is a good but not a fundamental theory.

Perhaps the most obvious missing element is gravity, which is not addressed by the SM at all, despite being the most important long range force whose effects will determine the fate of our universe. Although supersymmetry sometimes plays a supporting role in theories which attempt to describe gravity (for example “supergravities”, which are field theories with *local* supersymmetry) these issues will not be addressed here. There are plenty of other interesting questions, such as “Is there a reason for having *three* families of quarks and leptons?”, which are also not answered by the SM. Supersymmetry, however, is equally tight lipped on most of these issues, and so they will not be developed further here either.

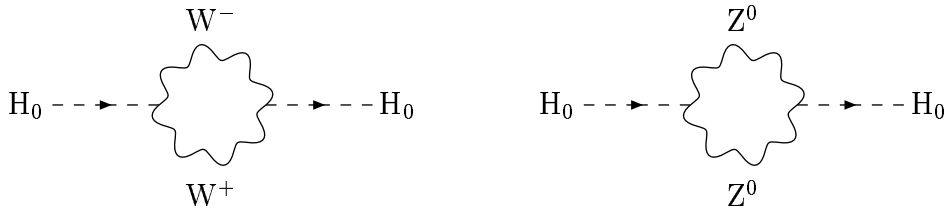
Perhaps the most unpopular feature of the SM is the so-called “naturalness” or “hierarchy” problem which concerns the renormalization of the higgs mass. Within the SM the one loop corrections to the higgs mass come from the diagrams shown in Figure 2.1.^d These diagrams introduce corrections to the higgs-mass-squared which are quadratic in the cutoff for the loop momenta. If this cutoff were to be placed at 10^{14} GeV (roughly the GUT scale) then the bare higgs mass would be forced to the same scale. Therein lies the problem; many physicists are not happy with the idea that the bare higgs

^cFor example, neutrino mixing might occur within supersymmetry (described later) as a result of loop diagrams permitted by non-zero R -parity and lepton number violating couplings. [3]

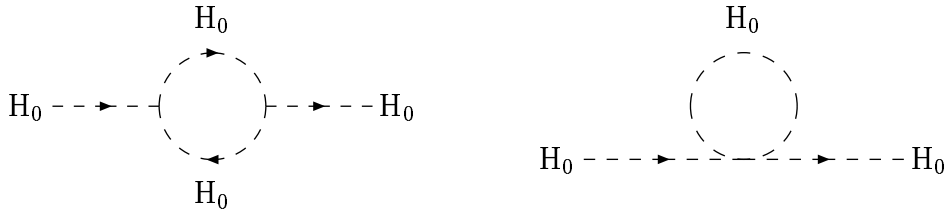
^dFor Figure 2.1 I am grateful to Peter Richardson from whose PhD thesis [5] it was reproduced with kind permission.



(a) Fermion loops.



(b) W and Z boson loops.



(c) Higgs loops.

Figure 2.1: *One-loop contributions to the higgs mass.*

mass would then have to be tuned to about 12 decimal places in order to reproduce the observed masses of the W and Z bosons. Supersymmetry provides a solution to the hierarchy problem which will be discussed later.

2.5 Supersymmetry

Described in very general terms, supersymmetry proposes that there is an additional as-yet-undiscovered symmetry of nature which links fermions and bosons. General re-

views of supersymmetry may be found in [6–8]. A brief discussion will follow in which supersymmetry will be seen to address the hierarchy problem of Section 2.4.

From a phenomenological point of view, the effect of requiring supersymmetry in a quantum field theory (QFT) is that for every boson (or fermion) degree of freedom (DOF) with non-mass quantum numbers $\{q_i\}$ there exists a fermion (or boson) DOF with the same non-mass quantum numbers. The particle corresponding to either member of such a boson-fermion pairing is known as the “superpartner” of the other. An “exact” supersymmetry is one in which every particle is mass degenerate with its superpartner. It is often said^e that one of the great successes of supersymmetry is that half of the particles that it predicts have already been discovered! However, it is the *extra* particle content which helps to solve the hierarchy problem.

For each of the diagrams shown in Figure 2.1 supersymmetry requires another where the new loop particle is the superpartner of the old one. The leading term in the cutoff in two such diagrams will have a relative minus sign coming from the interchange fermion $\leftrightarrow$ boson (recall the Feynman rule that requires a -1 be inserted for every fermion loop) and so these leading corrections to the higgs mass will all cancel out exactly. When supersymmetry is broken later, in order to make supersymmetric particles heavier than their SM counterparts, one of the constraints on the allowable supersymmetry breaking interactions will be that they do not spoil this leading term cancellation. The name given to this is “soft” supersymmetry breaking. The remaining correction to the higgs-mass-squared in the case of the SM fermions and their superpartners in an unbroken theory is of order $|g_f|^2 m_f^2 \ln(\Lambda/m_f)$ (where Λ is the loop momentum cutoff, g_f is the coupling of the fermion to the higgs, and m_f is the mass of the fermion in question). This has achieved the desired aim of solving the hierarchy problem by reducing the dependence of the correction to the higgs-mass-squared from the Λ^2 of the SM to a more

^emost frequently by Gordon Kane

manageable $m_f^2 \ln(\Lambda/m_f)$. Requiring that this higgs mass correction be of the order of the electroweak scale 100 GeV (or smaller) thus implies that the masses of the particles in the supermultiplets be of order 1 TeV or less (see Appendix A).

There are two types of supersymmetric theories: R -parity conserving ($R_P C$) and R -parity violating ($R_P V$). $R_P V$ theories allow sparticles to be produced singly and allow them to decay entirely into SM particles. In contrast, $R_P C$ theories only permit vertices involving *pairs* of sparticles. As a consequence, $R_P C$ sparticles can only be produced in pairs. Moreover, changes to $R_P C$ sparticles can only take place in one of two ways: 1) by decay to a lighter sparticle via the emission of one or more SM particles, or 2) by annihilation with an antiparticle (which is extremely unlikely at a hadron collider). In these models, the lightest supersymmetric particle (LSP) is effectively stable. The classic characteristic of an event involving sparticles under an $R_P C$ model, then, is that the final state will contain exactly two LSPs and a number of SM decay products.

2.6 MSSM

The MSSM is the name given to the “simplest” model which is supersymmetric, $R_P C$ and which has as a subset of its particle content all the quarks, leptons and gauge bosons of the SM. The meaning of “simplest” in this context will become clearer with later discussion.

The first stage in the construction of the MSSM requires that the particle content of the unbroken SM be represented within the framework of “supermultiplets”.

A supermultiplet may be thought of as a two-component object which contains pairs of bosonic and fermionic fields, whose spins differ by one-half of a unit, and which may be related to each other via a supersymmetric transformation. In all other respects, the two components are identical, in that they must have the same representation under

the gauge groups. In [9] it is remarked that after suitable transformations, “ $N = 1$ ” supersymmetric theories may always be written down in terms of just two different types of supermultiplet, and these are known as chiral and gauge supermultiplets.^f Chiral supermultiplets contain two spin- $\frac{1}{2}$ fermions and two scalars, while gauge supermultiplets contain two spin- $\frac{1}{2}$ fermions and one massless spin-1 boson.^g In which types of multiplet should the SM particles be placed? Clearly the scalar higgs-like particles must go into chiral supermultiplets and the spin-1 vector bosons must go into gauge supermultiplets. What about the quarks and leptons? Critically, the two types of supermultiplet differ in their treatment of their fermionic components. In the chiral supermultiplets the left and right components of the spin- $\frac{1}{2}$ fermions transform differently under Lorentz transformations, whereas in gauge supermultiplets they do not. Since the left and right components of quarks and leptons also transform differently under Lorentz transformations, they must reside in chiral supermultiplets. Consequently the superpartners of the SM scalars have spins half a unit *greater* than their counterparts, whereas the superpartners of the SM fermions and gauge bosons have spins half a unit *less*.

The name of the MSSM implies that it introduces the minimum number of new particles into the theory. In fact this is not quite the case, as it introduces an entirely new superpartner for each and every SM particle. This is perhaps surprising as, in principle, a single chiral supermultiplet could contain both the (ν_L, e_L) doublet of $SU(2)$ as well as a higgs doublet with hypercharge -1 (since both of these fields transform identically under the gauge groups) which would in effect assert that the sneutrino is none other than the higgs boson. However, despite its evident aesthetic appeal, it seems that adoption of this idea by a supersymmetric theory leads to “insoluble phenomenological problems ... including lepton number violation and a mass for one of the neutrinos in gross violation

^fIn the literature, chiral supermultiplets are often also known as “scalar supermultiplets” or “matter supermultiplets”, while gauge supermultiplets may also be known as “vector supermultiplets”.

^gStrictly speaking, the “two fermions” in each of the two types of supermultiplet should be described as a single Weyl fermion with two helicity states and thus two degrees of freedom.

of experimental bounds” [6] and thus nature seems to force the MSSM to be less simple than it ‘could’ be. In the end, every SM particle must acquire its own supermultiplet, and so supersymmetrizing the SM doubles the particle content of the theory.

The second stage in the construction of the MSSM requires that masses for the standard gauge bosons, quarks and leptons are generated using a Higgs mechanism analogous to that used in the SM. The main difference between the EWSB mechanisms of the SM and supersymmetry is that in supersymmetry, two higgs doublets rather than one are required. Since the reasons for this are not critical at this point, discussion is left to Appendix B.^h.

2.7 Supersymmetry breaking in the MSSM

Since the scalar partners of the quarks and leptons of the SM have not been discovered yet, if they do exist, they must be more massive than their SM counterparts. The Higgs mechanism used to effect EWSB is quite capable of giving masses to MSSM particles via their Yukawa interactions, but it cannot give different masses to superpartners. Although the “final theory” is expected to be fully supersymmetric, some as-yet-unknown form of spontaneous symmetry breaking is expected to be the cause of the broken supersymmetry we (hope to) observe. Rather than make assumptions about the possible forms which this mechanism could take, the attitude of the MSSM is to promote the lagrangian to an effective one, while at the same time adding to it *all* the supersymmetry breaking mass terms which do not break gauge invariance or spoil the solution to the hierarchy problem. This way, the precise details of symmetry breaking are able to remain hidden, but we are still left with a general broken MSSM, valid below the breaking scale. Whatever the mechanism of supersymmetry breaking, we ought to be able to model it by careful

^hThe interested reader is also referred to [10] for a good explanation of the higgs sector within supersymmetry.

choice of breaking parameters.

Supersymmetric theories are often unfairly criticised for having large numbers of free parameters. Ironically it is the supersymmetry *breaking* terms which do the real damage. Whereas the unbroken MSSM has approximately the same number of parameters as the SM, it is reported in [6] that: “A careful count reveals that there are 105 masses, phases and mixing angles in the MSSM lagrangian which cannot be rotated away by redefining the phases and flavor basis for the quark and lepton supermultiplets. Thus, in principle, supersymmetry appears to introduce a tremendous arbitrariness to the lagrangian.”

2.8 Minimal supergravity models

In order to have a concrete model to work with, and one with a more manageable set of parameters than those of the fully broken MSSM, it is necessary to make some simplifying assumptions. The class of simplified models which we shall be concerned with here are known as mSUGRA models.

Crudely put, the simplicity of mSUGRA models is little more than a trick. All supersymmetry breaking scalar mass matrices are made proportional to identity matrices, with a common constant of proportionality, m_0 , at the GUT scale. The supersymmetry breaking gaugino masses are also unified at the GUT scale, taking the common value $m_{1/2}$. The supersymmetry breaking Yukawa couplings are all made proportional to their supersymmetry conserving counterparts via another constant A_0 , while the remaining supersymmetry breaking parameter (which couples H_u to H_d) requires one final DOF. Once the supersymmetry conserving parameters of the MSSM are input, it is possible to use the renormalization group equations (RGEs) to generate particle masses at the electroweak scale. If the additional requirement is then made: that EWSB should successfully reproduce the masses for the W - and Z -bosons; it is possible to eliminate

some of the redundancy from the input parameters. As a result, mSUGRA models are eventually characterised by four real parameters and one sign, these being:

$$m_0, m_{1/2}, A_0, \tan \beta \text{ and } \text{sgn} \mu, \quad (2.4)$$

the only new quantities here being μ and $\tan \beta$. $\tan \beta$ is the ratio of the two higgs VEVs, while μ is the supersymmetric analogue of the SM's higgs mass parameter μ which was first introduced in (2.1). The role of μ is described in more detail in Appendix C in the context of a discussion highlighting the differences between the Higgs potentials of mSUGRA and the SM. In connection with the definition of mSUGRA, however, we only need to observe that μ is a parameter which is determined, up to a sign, by the requirements of correct EWSB.

2.9 Other models

As has already been pointed out, neither the SM nor supersymmetry as described aboveⁱ is expected to be a “final answer”. Most notably, neither makes any attempt to describe gravity. Nevertheless, each is a valuable theory, if only insofar as it might be a perfectly good “low energy effective field theory” approximation to a more fundamental theory.

To date, much effort has been expended in attempts to discover whether string theory can be just such a theory.

A detailed discussion of string theory would be beyond the scope of this document, but since an optimised string model (OSM) will be used in the analysis, a few comments prefacing its definition (see Section 2.9.4) would seem appropriate.

ⁱThat is to say, supersymmetry defined as a global spacetime symmetry of a QFT.

2.9.1 String theory

At its core, string theory sees the world composed not of point particles, but of little “strings”. If the truth be told, the strings are not really very much like the kind of strings we usually imagine, most notably because they exist in ten or twenty-six dimensional space. In fact, the only thing they *do* have in common with ordinary strings, is that the region of spacetime swept out by each is a two dimensional surface. This is the same as saying that if you are an ant crawling around on such a string, you have the choice of whether to walk forwards or backwards, but not sideways.

Formally speaking, then, a complete description of a string is a matter of specifying the coordinates $X^\mu(\sigma, \tau)$ of each bit of it. The parameters σ and τ should be thought of as parameters on the string’s two dimensional world sheet. For example, σ might be “distance along the string from one end”, while τ might represent “the time at which some observer sees that part of the string”. Anyway, we don’t really care much about the precise details of what σ and τ are, the main point is that the worldsheet be parametrized by *two* parameters, thus *making it a string*. The real interest lies in the X^μ values, for $\mu = 0, \dots, D - 1$, for it is these coordinates in the D -dimensional space which actually tell us *where* the string is and *what* shape it has.

Notice this means that strings can do some important things that point particles and plane waves cannot: they can wobble, they can stretch, and they can tie themselves in knots. This leads to one of the biggest differences between string theories and QFTs; every time you need to describe a new particle in a QFT you have to introduce a new field, whereas in string theory you can just try to associate a new particle with a more complicated string “mode” (i.e. way of oscillating or being knotted). The other major difference between string theory and QFT is the nature of spacetime itself.^j In QFT, on

^j(at least before either theory gets quantized, that is)

the one hand, we imagine *at each point in spacetime* a little collection of fields. Taken together across some finite region of spacetime, the field values might be seen to represent the passing of a plane wave perhaps. On the other hand the spacetime of string theory is largely empty. At almost every point there is nothing interesting happening at all. Where string theory's spacetime *is* interrupted, it is by walls “of string”, one might say. In principle, “string” might look very much the same everywhere, and only the different modes or configurations would serve to distinguish states. However some branches of string theory (so called heterotic strings) define strings not only by their shape $X^\mu(\sigma, \tau)$, but also by an equal number of fermionic degrees of freedom $\Psi^\mu(\sigma, \tau)$. Classically, one can think of these degrees of freedom as fermionic fields “decorating” the world-sheet of the string. Either way, fields where they exist are restricted to the string, while the bulk remains empty.^k

At some point, of course, an action must be written down. The main quantities which feature in the lagrangian may be thought of as equivalent to $X^\mu(\sigma, \tau)$ and $\Psi^\mu(\sigma, \tau)$, although in practice other fields are used. For example, other choices are better able at making explicit the invariance of the action under changes in the (fairly arbitrary) definitions of σ and τ .

At this point the number of branches of string theory begin to multiply, as various constraints may now be imposed on these actions by different groups of model builders. One such branch is “superstring theory” which, being of relevance to models discussed later, will now be discussed.

^kNote that this description of an empty bulk is only valid classically. At some point during quantization, all possible string configurations must be considered for the path integral, and so in some sense, Ψ^μ behaves very much like a field in the target space. It should also be noted that the comments above have glossed over a myriad of complications concerning, for example, the formulation of string theories in curved spacetimes. However it is hoped that the pictures being formed are nonetheless useful.

2.9.2 Superstrings

Superstrings come from the branch of heterotic string theory in which the string action is required to be invariant under local supersymmetry transformations involving X^μ and Φ^μ .

At first sight one is tempted to say “Oh good! We will now have a string theory which will describe supersymmetric particles!”, but in reality the situation is more complex. When supersymmetry was introduced in Section 2.5, the emphasis was entirely on its phenomenological effects. As a consequence, the picture of supersymmetry which has been presented up to now has been a narrow one, differing substantially from viewpoints biased toward group theory, for example. From the group theoretic point of view, supersymmetry concerns itself purely with the transformation properties of fields under the action of the generators Q_α , $\bar{Q}_{\dot{\alpha}}$ and P_μ of the supersymmetry algebra. (See for example [11].) However, for the purposes of this thesis, it will not be necessary to use this approach. It will only be important to appreciate that the *underlying* definition of supersymmetry is one that is connected with the relationships between bosonic and fermionic fields and nothing, *a priori*, to do with quarks and squarks. It is only in the context of “ordinary” QFTs, in which the fermionic and bosonic fields really *do* represent quarks and squarks, that the interpretation of Section 2.5 happens to apply. This type of supersymmetry is known as a “spacetime supersymmetry”. In contrast, the supersymmetry of superstrings is known as a “worldsheet supersymmetry”, as it relates the fields X^μ and Φ^μ . (Remember that although $X^\mu(\sigma, \tau)$ really represents the coordinates of parts of a string, in the context of the worldsheet supersymmetry, X^μ and Φ^μ may be thought of as fields in the two dimensional (σ, τ) -space of the worldsheet.) Thus two-dimensional worldsheet supersymmetry is different from spacetime supersymmetry not only because it has, *prima facie*, nothing to do with quark and squark fields, but because it does not even act in a space with the same number of dimensions as Minkowski space

or the bulk! Nonetheless, it will be noted later that, despite their differences, worldsheet supersymmetry does indeed help to secure the much desired spacetime supersymmetry for the final theory.

So by this stage we have a picture of superstrings, conceived as a basic strings, but subject to the additional constraint of worldsheet supersymmetry invariance.

2.9.3 The transition to the effective quantum field theory

What do superstring theories predict?

Extracting reliable predictions from string theory has proved to be particularly tricky. The problems tend to stem from the complexities of theory; not only is there the question of how to compactify to four dimensions, but there are even two regimes of perturbative expansion to worry about: one describing the interactions of the fields on the two dimensional worldsheet and the other describing the number of ways the worldsheets themselves can knot and form loops of their own. Be that as it may, by aiming to find *effective* quantum field theories, there is a way of sidestepping some of these complexities. The SM is seen as an effective theory: an approximation to some other unknown theory valid at higher energies. Superstring phenomenologists are faced with the opposite problem. They know the high energy theory, not the low energy one. Given a particular string theory, they have to try to write down an effective field theory which, they hope, approximates it at low energies.

Although the discussion of Section 2.9.2 highlighted the differences between worldsheet and spacetime supersymmetry, the authors of [12] state that the existence of a particular charge quantization condition in the worldsheet is sufficient to make the one imply the other. Consequently, the first requirement of effective field theories seeking to describe string theories satisfying this extra condition is that they exhibit $N = 1$

spacetime supersymmetry. In addition, the subset of these string models concerning us has been shown (e.g. in [9]) to possess no light states which cannot be described using gauge, chiral or graviton-gravitino supermultiplets.

This is excellent news, as the form of the most general action, able to describe an $N = 1$ supergravity theory coupled only to gauge and chiral supermultiplets, has been known for some time [13]. For a good summary of *all* the components of the lagrangian, the reader is again referred to [9]. However, the only part we shall need to mention here is the “Kähler potential”. The Kähler potential is simply an arbitrary real function $K(z, \bar{z})$ of the chiral fields z of the theory. Its principle role is to generate that part of the lagrangian corresponding to the kinetic terms for the chiral supermultiplets. The rule is

$$\mathcal{L}_{\text{chiral kinetic part}} = K_{z\bar{z}} \partial_\mu z \partial^\mu \bar{z} \quad (2.5)$$

where $K_{z\bar{z}} = \partial^2 K / \partial z \partial \bar{z}$. We have already seen that the squarks, sleptons and higgses of the MSSM live in chiral supermultiplets. We should not be surprised to discover later, therefore, that the masses of these particles can be sensitive to the form of the Kähler potential.

One of the tasks of the superstring phenomenologist is to calculate, given a particular string model, the form of the Kähler potential.¹ In the case of superstring models with *orbifold* compactifications, the tree-level Kähler potential has been found to have the form

$$K = -\log(S + S^*) - 3\log(T + T^*) + \sum_i (T + T^*)^{n_i} \phi_i \phi_i^*. \quad (2.6)$$

We shall ignore the terms containing only dilaton and moduli fields (S and T), preferring instead to concentrate on the sum of terms involving the chiral superfields ϕ_i . It will be noted that these terms are multiplied by factors which include the only free parameters,

¹The other tasks include calculating the remaining parts of the general $N = 1$ supergravity lagrangian; namely the superpotential $W(z)$ and the gauge-kinetic function $f_{ab}(z)$.

n_i , present in the potential. These n_i are known as “modular weights”, and the values they may take are usually determined by or dependent upon the particular string model in question. Typical values are negative integers in the range $-5 \leq n_i \leq -1$ (see [14]). With additional assumptions, the authors of [14] conclude that, at the string scale, the soft scalar masses of the effective theory have the following dependence on the modular weights:

$$m_i^2 = m_{3/2}^2(1 + n_i \cos^2 \theta), \quad (2.7)$$

where $m_{3/2}$ is the gravitino mass and θ is the “goldstino angle” (a supersymmetry breaking parameter). Note that, to ensure positive scalar masses, the definitions in Equations 2.8 and 2.9 will later require that the restriction $\cos^2 \theta \leq 1/2$ also be imposed.

2.9.4 Optimized string models

Having described some general results for superstring theories with orbifold compactification, we are now in a position to be able to describe the particular class of string model which will be used later on in this thesis for a comparison with a benchmark mSUGRA model.

“Optimized” string models, proposed in [14], are defined by their modular weights:

$$n_{Q_L} = n_{d_L^c} = n_{u_L^c} = -2, \quad (2.8)$$

$$n_{L_L} = n_{e_L^c} = -1. \quad (2.9)$$

As is readily seen from Equation 2.7, the effect of giving the sleptons larger (more positive) modular weights than the squarks, is to make the sleptons heavier than in models with universal weights. At first sight, the larger slepton weight assignments (at the string scale) would also appear to make the sleptons heavier than the squarks, but this relationship is not preserved by the RGEs during the run to lower scales. Nevertheless,

the heavier starting point of the “optimized model” sleptons makes it easier to run them to heavier electroweak scale masses than in similar mSUGRA models.

Aside from providing an alternative to mSUGRA, optimized models might be interesting in their own right. One of their original motivations was a desire for models in which the physical minimum of the Higgs potential is also the *global* minimum. Theorists refer to this as “avoiding unbounded from below (UFB) directions in the effective potential”. In such models, one removes the frightening possibility that the universe might catastrophically tunnel out of its present state into a lower minimum with a different vacuum! Still, the jury remains out on whether UFB problems should worry us, especially if the timescale associated with the tunneling is much larger than the age of the universe. Perhaps more compellingly, optimised models were also designed to avoid dangerous charge and colour breaking minima, which is to say that the global minimum of the whole potential $V(H_1, H_2, \tilde{q}_i, \tilde{l}_i)$ is required to occur at a point at which $\tilde{q}_i = \tilde{l}_i = 0$, thus preserving colour and electric charge as unbroken symmetries [14, 15].

2.10 Summary

In this chapter, support has been given to the suggestion that supersymmetry might play a vital role in extending the SM. The two classes of supersymmetric model which will be used later in the the analysis (mSUGRA and OSMs) have been introduced and motivated, and an explanation has been offered as to why the main differences between the two classes of models should be expected to be seen in the slepton masses. These slepton mass differences will later play a part in demonstrating the strengths and weaknesses of the model independent approach used in the analysis.

Chapter 3

The ATLAS Experiment

3.1 The LHC

The best particle-detector in the world will fail to discover new particles, supersymmetric or otherwise, if too few of them are presented to it for measurement. A plentiful supply of new physics events will not assure success either, if the new particles are always obscured by detritus from unrelated interactions. The accelerator which will supply physics events to ATLAS and which must steer its course between the competing demands of high luminosity, high energy and event simplicity, will be the LHC. The LHC will collide protons from two beams at a centre of mass energy of 14 TeV and at luminosities starting from $0.1 \times 10^{34} \text{cm}^{-2} \text{s}^{-1}$ and rising to $1.0 \times 10^{34} \text{cm}^{-2} \text{s}^{-1}$. Table 3.1 summarises some of the other main properties expected of the accelerator when using proton beams. In principle, higher collisions energies could be obtained by accelerating heavier ions, but the increasingly complex nature of the resultant collisions makes them unsuited to a discovery machine.

At the time of writing, the LHC is still under construction. Most of the LHC will occupy the tunnel of the old large electron-positron (LEP) collider (see Figure 3.1).

The official ATLAS/LHC schedule of 9th February 2001 expects that the accelerator will produce a stable beam in February or March of 2006, in time to produce its first collisions at the start of April. Lagging somewhat behind, however, the complete ATLAS detector is not expected to be ready by that stage. Consequently, colliding beams are not scheduled to go beyond ‘pilot run’ of four weeks. During this time ATLAS, lacking its third pixel layer, the forward end caps of the transition radiation tracker (TRT) and some of the muon stations at high rapidity, would nonetheless be able to take some useful data for physics studies. A three month shutdown is then foreseen, in which the remainder of ATLAS would be installed ready for the first ‘real’ physics run starting in August 2006 and lasting for seven months. During this run, it is expected that the LHC will be able to deliver an integrated luminosity of at least 10 fb^{-1} . Due to the short space of time in which this must be delivered, a luminosity of at least $2 \text{ nb}^{-1}\text{s}^{-1}$ is expected. This luminosity is around twice that which has typically been assumed in historical investigations into physics performance in the “year of low-luminosity running” at ATLAS.

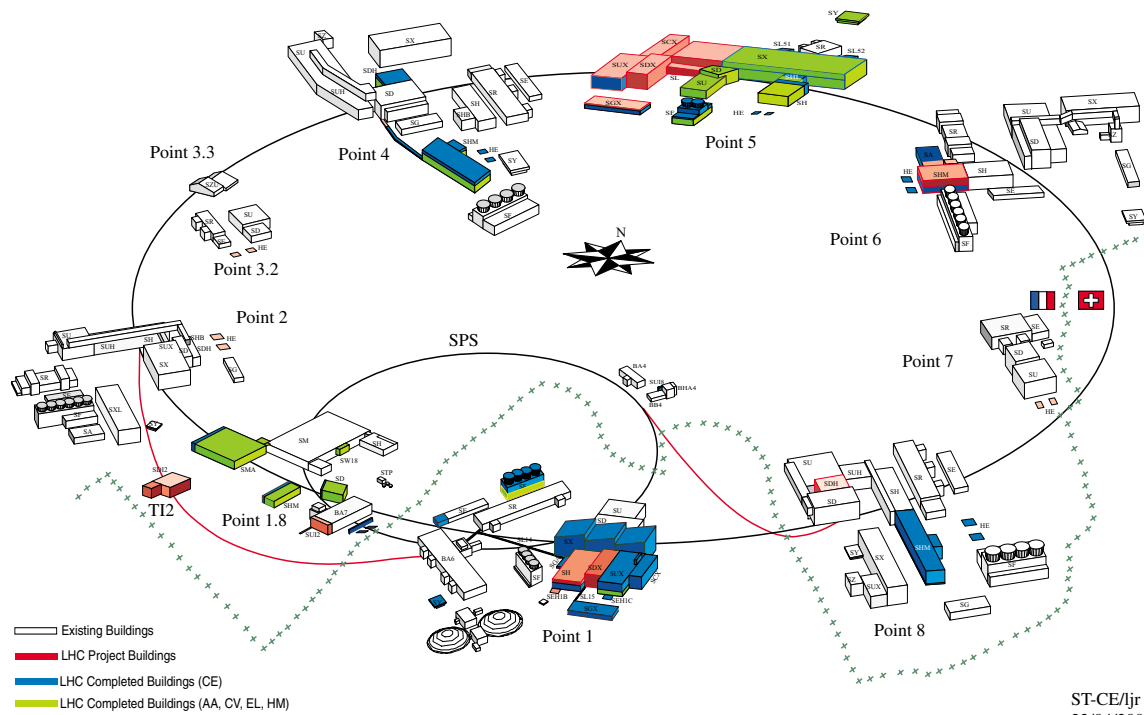
3.2 The ATLAS experiment

The ATLAS collaboration can be said to have come into being in October 1992, when the publication of its letter of intent [16] sealed the union of the EAGLE and ASCOT groups [17, 18]. It is not surprising that the first thing which ATLAS inherited from EAGLE was a highly contrived acronym (see list of abbreviations, Appendix F).

The principle aims of ATLAS are to be able to make a better-than- 5σ discovery of the higgs boson over the mass range 100 GeV to 1000 GeV, to probe for signs of physics beyond the SM such as supersymmetry, and to make competitive measurements in B -physics. There is good evidence that, if all goes to plan, the higgs target should be met.

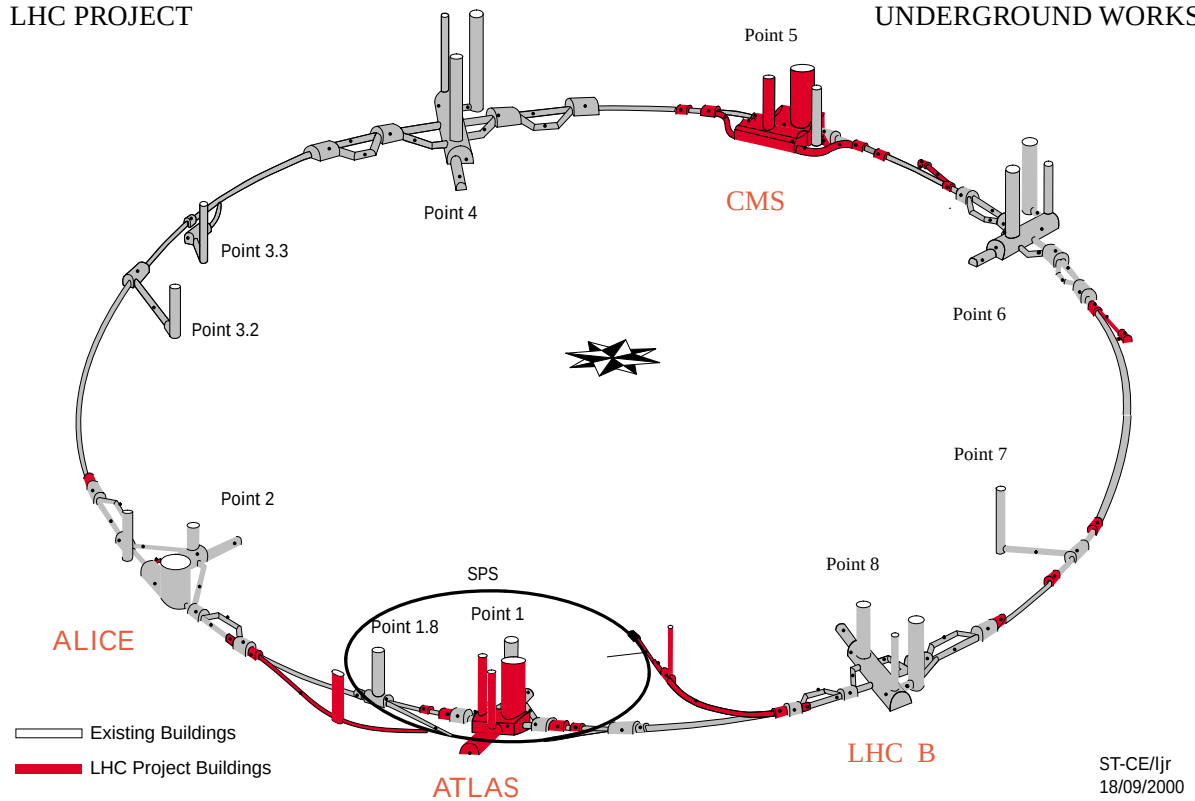
LHC PROJECT

SURFACE BUILDINGS



LHC PROJECT

UNDERGROUND WORKS

Figure 3.1: *The LHC above and below ground.*

Particle energy at collision	7	TeV
Particle energy at injection	450	GeV
Luminosity (commissioning)	1	nbarn ⁻¹ s ⁻¹
Luminosity (nominal)	10	nbarn ⁻¹ s ⁻¹
Luminosity (maximum)	23	nbarn ⁻¹ s ⁻¹
DC beam current	0.56	A
Bunch spacing	7.48	m
Number of bunches	2808	
Bunch separation	24.95	ns
Number of particles per bunch	1.1×10^{11}	
Total crossing angle	300	μ rad
Beam size at IP1 and IP5 (r.m.s.)	15.9	μ m
Energy loss per turn	7	keV
Total radiated power per beam	3.8	kW
Stored energy per beam	350	MJ
Inelastic cross section	60	mbarn
Inelastic event rate per IP	600	MHz

Table 3.1: *LHC beam and accelerator parameters.*

Figure 3.2 shows the expected discovery significance over a wide range of potential higgs masses. It is much harder to produce a plot which makes a similarly definitive statement about the performance of ATLAS in relation to supersymmetry - the parameter space is just too big. In order to stretch its supersymmetric reach to the limit, the ATLAS detector must rely on maximizing its all round performance as a general-purpose particle detector. The investigations in this thesis rely heavily on assumptions about the accuracy with which ATLAS will identify isolated jets and leptons within events and measure their momenta. Reliance will also be placed on its ability to distinguish the lepton charges, identify lepton flavour, measure missing transverse momentum, and to a lesser extent identify *b*-tag jets coming from *b*-quarks.

Although a short description of the each of the ATLAS detector's subsystems will now follow, the analyses performed in this thesis rely solely on simulated rather than real data, and so a greater emphasis will be given to describing the model used for fast simulation of the detector, than to describing the detector itself. While it is hoped

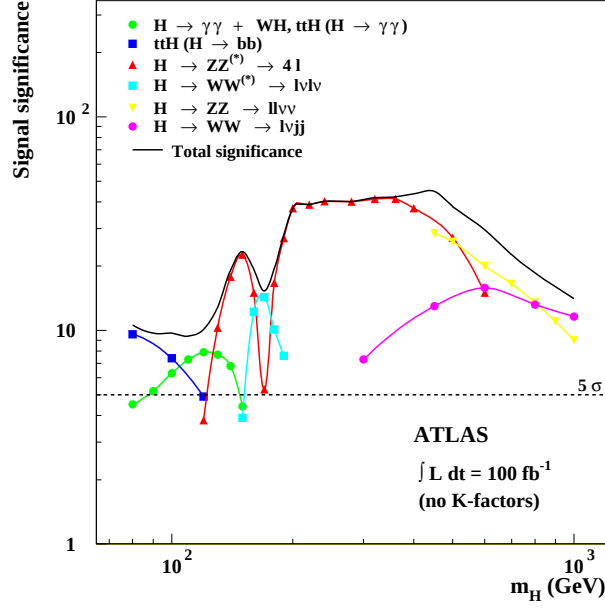


Figure 3.2: This figure, taken from [19], illustrates the likely significance of an ATLAS Higgs discovery as a function of Higgs mass. The significance shown corresponds to 100 fb^{-1} integrated luminosity.

that the method of fast simulation will faithfully model the performance of the real experiment when data-taking begins, there is nothing quite as convincing as real data. Until real data begins to arrive, the conclusions of this thesis will have to stand, in the first instance, in the context of predictions for a detector modelled by the fast simulation, rather than as predictions for the ATLAS experiment.

3.2.1 Coordinate systems

The ATLAS detector is, to first approximation, cylindrically symmetric about the mean beam line at the interaction point. Accordingly, the coordinate system commonly used is cylindrical about that axis, and its principle coordinates are r , measuring distance from that axis, and ϕ , measuring the azimuthal angle about that axis. Where used, x and y measure positions in the plane transverse to the beam axis, such that $x = r \cos \phi$ and $y = r \sin \phi$. The z -axis completes the right handed triad. At hadron colliders, where

particles radiate from the origin and have a centre of mass with an unknown longitudinal boost, the most useful set of coordinates for describing physics, is one that is symmetric under rotations about and boosts along the beam axis. One such coordinate, ϕ , selecting angle in the transverse plane, has been seen already. The other, making a longitudinal measurement, is the pseudorapidity η (see Appendix D) which is related to the polar angle θ from the positive z -axis. These two, η and ϕ , combined with momentum in the transverse plane, are the coordinates most commonly used to describe particle momenta.

3.2.2 Inner Detector

At the heart of ATLAS lies the inner detector. In the words of [20];

The task of the inner detector is to reconstruct the tracks and vertices in the event with high efficiency, contributing, together with the calorimeter and muon systems to the electron, photon and muon recognition, and supplying the important extra signature for short-lived particle decay vertices.

The inner detector is actually three independent detectors nested within each other. In order of increasing distance from the beam pipe they are the pixel detector, the semiconductor tracker (SCT) and the TRT.

All parts of a given sub-detector share a common detection technology, but the form of the layout is always different in the central region to that at higher values of $|\eta|$ as may be seen in Figure 3.3. The term “barrel” refers to those parts of the detector in the central region, and “forward” to those at high and low η .

Both the pixel detector and the SCT use reverse biased silicon diodes to detect charged particles. During their passage through the silicon, passing charged particles create electron hole pairs. Some of these electrons and holes are created in the depletion

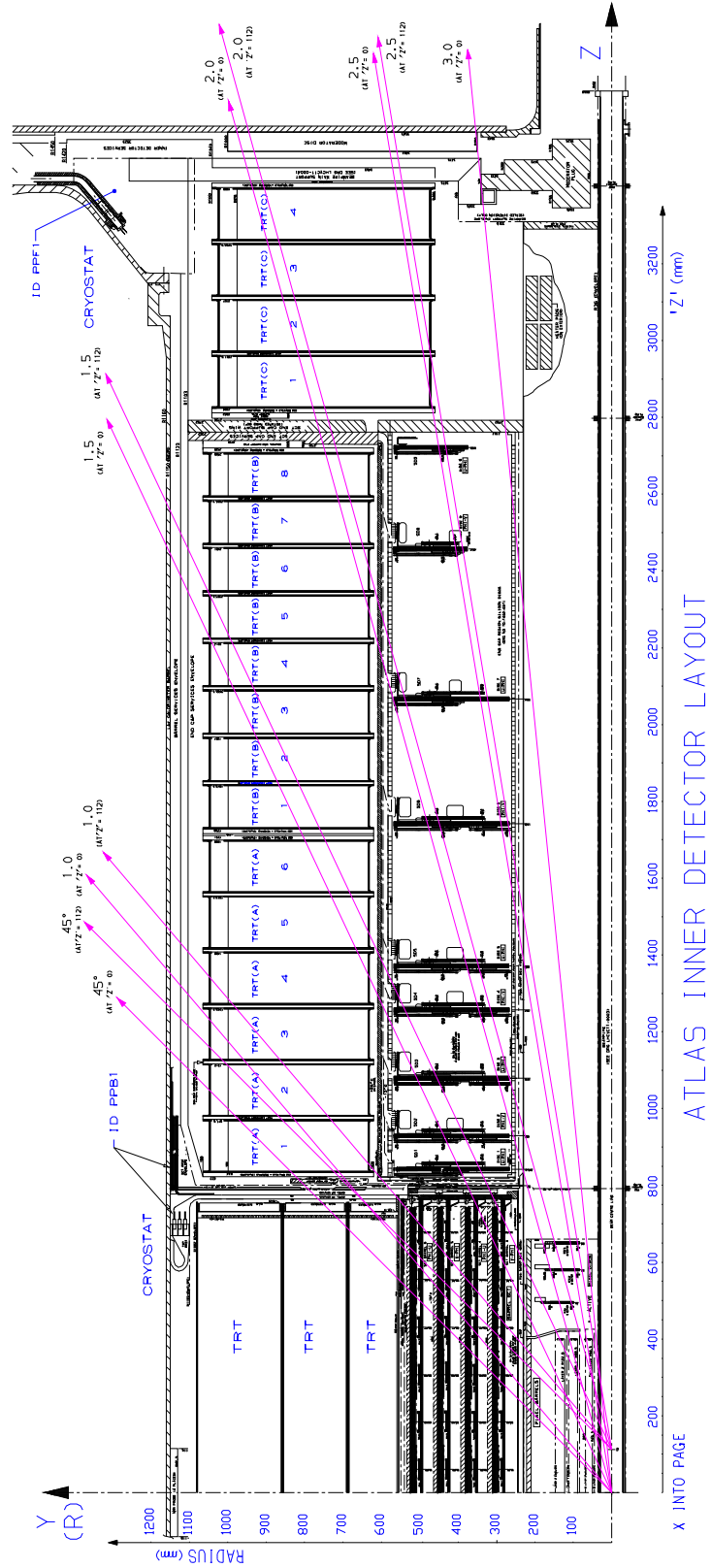


Figure 3.3: Cross section through one quarter of the inner detector.

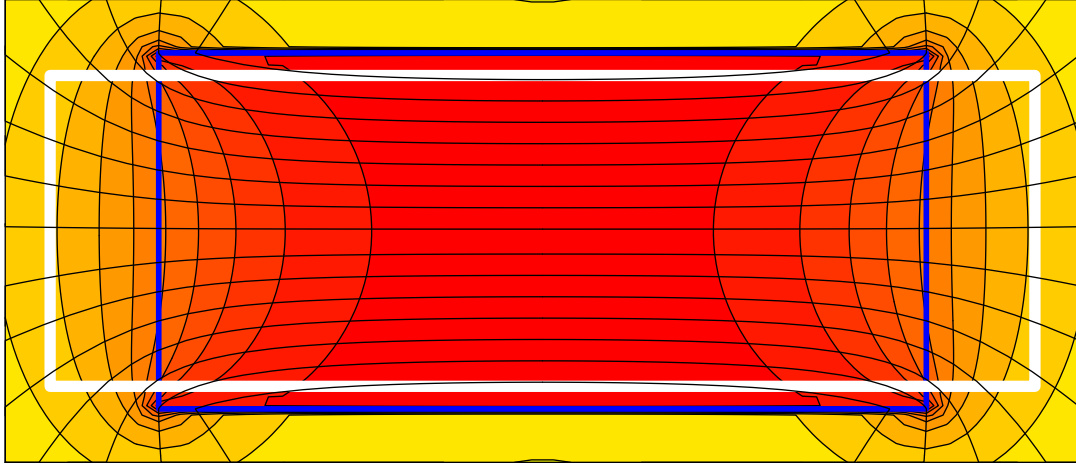


Figure 3.4: *The magnetic field of the inner detector. The relative dimensions of the solenoid and the inner detector are represented by the blue and white rectangles respectively. Contours of magnetic field strength are shown at 10% intervals, and the regions between the contours are coloured accordingly, ranging from dark red for the strongest fields in the central region of the solenoid to pale yellow on the outside. Magnetic field lines are also shown running through the solenoid from left to right.*

regions of diodes within the detector elements, and these are swept by the bias voltage to conductors on the surface of the silicon where the charge collected may be read out. The sizes and spacings of the diodes determine the space-point resolution of the detector. The pixel detector and the SCT are purely digital devices, and only generate yes/no answers to the question of whether a charged particle passed near enough to a detector element to trigger it. The TRT is a more complicated analogue device and will be described in more detail shortly.

3.2.2.1 Magnetic field of the inner detector

The magnetic field within the inner detector is generated by a solenoidal coil just outside the tracker volume and contained within the cryostat of the electromagnetic calorimeter visible in Figure 3.3. An approximate map of the field is shown in Figure 3.4. The

approximation neglects the finite thickness of the solenoid (22 mm) and does not simulate the flux return structures in the girders of the calorimeters. The solenoid is 5.30 m long and has inner and outer radii of 1.218 and 1.240 m respectively, giving it an aspect ratio of 2.16 to 1. Notably, the outermost parts of the TRT are to be found 0.72 m (27%) beyond the end of the solenoid in a region where the field strength is approximately 20% of its value at the centre, and where the field lines are almost parallel to the trajectory of a particle radiated at the same η . This degrades the momentum resolution of the inner detector over $|\eta| > 2.5$ (see also Equation (4.7)).

3.2.2.2 Pixel detector

The pixel detector, being closest to the interaction point and behind the least dead material, plays the most important role in vertex resolution. To obtain the highest possible ϕ resolution, the individual diodes are typically $50 \times 400 \mu\text{m}$ in size and run parallel to z in the barrel, and parallel to r in the disc regions. Nearly 50,000 of these pixels are grouped together into $16 \times 60 \text{ mm}$ arrays on single silicon wafers, and over 2000 of these make up the barrel and wheel structures which comprise the pixel detector.

3.2.2.3 Semiconductor tracker

At the radius of the SCT, p_T resolution demands that measurement of the azimuthal position of charged tracks is more important than measurement of η . This, combined with the greater area which the SCT has to cover, strongly influences its design. The individual diodes of the SCT are 60 mm long strips, aligned along $\pm\eta$, and arrayed in ϕ at a pitch which in the barrel is $80 \mu\text{m}$. Each array is deposited on a single $60 \times 60 \text{ mm}$ silicon wafer, which is sometimes also referred to as a “detector”. Except in some of the forward modules, end-to-end connection of pairs of these wafers actually increases the effective length of most diodes by a factor of two. Whether 60 or 120 mm long, hits

in these strip-arrays form the basis of the data read out from the SCT. Physically, the strip-arrays are mounted in back-to-back pairs called “modules”. In order to achieve *some* resolution in η , the strips of the front and back of each module are subject to a deliberate misalignment of 20 mrad, leading to a total crossing angle of 40 mrad. A single rhombus of intersection of two such $80\ \mu\text{m}$ strips thus has an intrinsic $(r\phi) \times z$ resolution of $\frac{1}{\sqrt{24}}(80\ \mu\text{m} \times 4\ \text{mm}) = 16\ \mu\text{m} \times 820\ \mu\text{m}$.

3.2.2.4 Transition radiation tracker

The TRT is a collection of over 370,000 conducting straw tubes, each containing a central sense wire surrounded by xenon gas. A potential difference of 1.6 kV is maintained between the wire and the straw, and used to collect charge liberated by the passage of charged particles through the gas. In addition, the straws are embedded in a medium whose abruptly varying refractive index causes traversing particles to emit radiation along the track direction. It is these soft x-rays of transition radiation (TR) which give the TRT its name. Discharges caused by transition radiation photons are also picked up by the straws.

Each straw has a diameter of about four millimetres, allowing the packing density to be such that a stiff track will pass through an average of 32 straws. A hit in a given straw produces both timing information, from which track displacement may be obtained, as well as pulse shape information from which the energy deposited may be measured. It is possible to discriminate at some level between hits due to ionization and hits due to TR using the difference in the energy spectra for the two types of deposit. Analysis of pulse shape can in principle provide further discrimination.

The amount of TR radiated at a discontinuity by a particle travelling with Lorentz factor γ and charge z is proportional to γz^2 but is independent of its mass. Thus a measurement of the magnitude of the TR deposited in a straw is equivalent to a

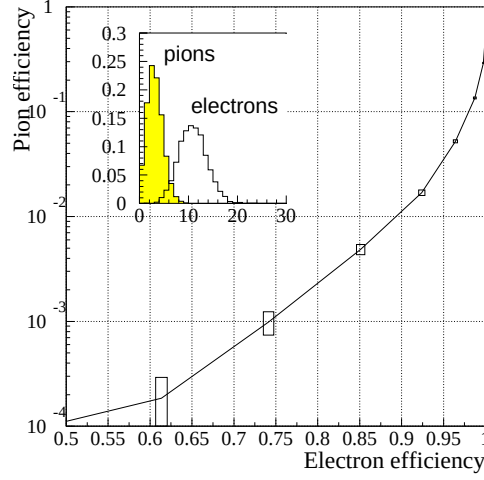


Figure 3.5: *Pion efficiency versus electron efficiency at 20 GeV as measured using the TRT end-cap sector prototype. The probabilities to observe a given number of TR hits (using a 6 keV TR threshold) for pions and electrons are shown in the top left-hand corner. (Figure taken from [20].)*

measurement of $\beta\gamma$ for that track. On the other hand, the track momentum as measured by the track shape is a measurement of $\beta\gamma m$. For two tracks of identical momentum, then, the one corresponding to the lighter particle can be expected to have produced more TR. Thus the percentage of TR hits on a given track, combined with a measurement of the track momentum, permit the inner detector, on its own, to perform elementary particle identification tasks such as discrimination of electrons from pions on the basis of their mass.

Unfortunately charge collected due to TR can not be measured in isolation from the charge coming from ionization as TR photons are radiated very close to the direction of motion. Consequently TRTs measure a combination of the charge due to both effects, and this limits their discriminatory power. In practice, TR is quantified by setting a threshold, and defining any straw with an energy deposition above this threshold to be a “TR hit”. With a threshold of about 6 keV, a TRT end-cap sector prototype has demonstrated a rejection of 70 against charged pions for an electron efficiency of 90%. Pion efficiencies at other electron efficiencies may be seen in Figure 3.5.

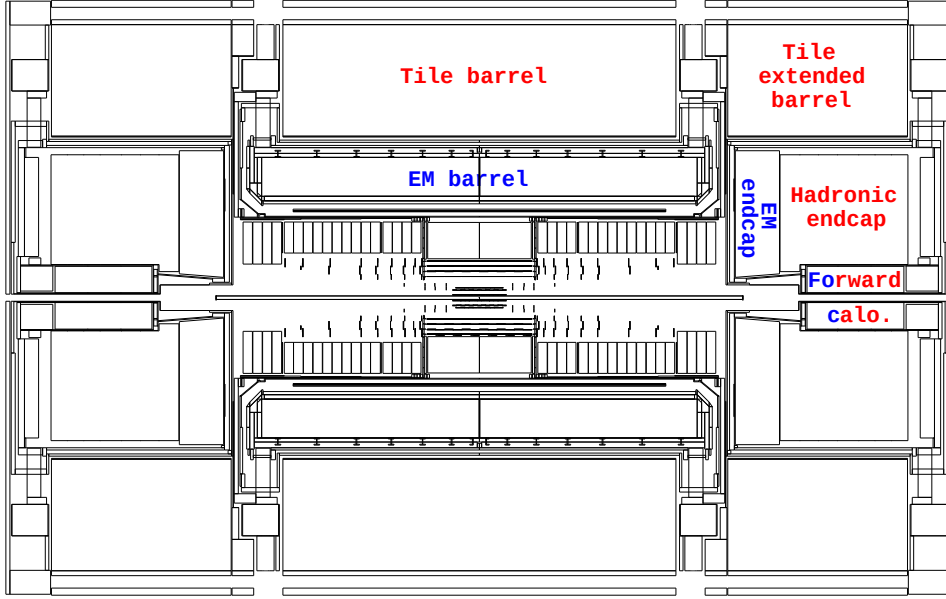


Figure 3.6: *ATLAS calorimeters. Cross section through the centre of ATLAS showing the main calorimeter regions. Hadronic calorimeters are labelled in red, and the electromagnetic calorimeters in blue. (Adapted from [21])*

3.2.3 Calorimeters

In common with most general-purpose particle detectors, the ATLAS detector will be equipped with electronic and hadronic calorimeters. The main purpose of the calorimeters is to measure electron, photon and jet energies, and provide a major contribution to the measurement of the missing transverse momentum of events.

ATLAS will only possess *sampling* calorimeters, and so their resolution will be dominated by sampling fluctuations, leading to fractional resolutions scaling as $1/\sqrt{E}$.

The ATLAS calorimetry is of necessity segmented into the six different regions: the EM barrel, the EM endcap, the tile barrel, the tile “extended barrel”, the hadronic endcap and the forward calorimeter. These regions are shown in Figure 3.6. Although the geometry varies, the electromagnetic calorimeter (ECAL) uses a common technology (lead-liquid argon (LAr)) throughout. In contrast, the hadronic calorimeter (HCAL)

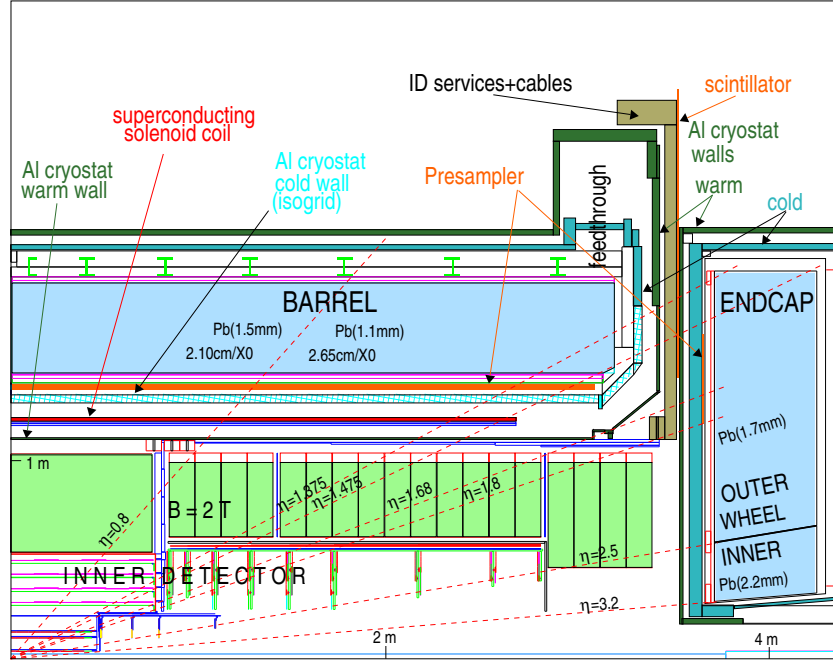


Figure 3.7: A cross section showing the two main components of the ECAL. In addition, a small part of the ECAL is contained within the forward calorimeter (see Figure 3.6). (From [21])

uses different technologies in each of its three main rapidity regimes. Brief descriptions of the calorimeters follow.

3.2.3.1 Electromagnetic Calorimeter

As described in the ATLAS calorimeter performance document [21], the ECAL ...

... is a lead-LAr detector with accordion-shaped Kapton electrodes and lead absorber plates over its full coverage. The accordion geometry provides complete ϕ -symmetry without azimuthal cracks.

The total coverage of the ECAL, including the forward calorimeter, goes up to $|\eta| = 4.9$. The two main parts of the ECAL cover $|\eta| < 3.2$. Over the central region ($|\eta| < 2.5$) the ECAL is segmented into three different layers known as “samplings”. Particles traversing

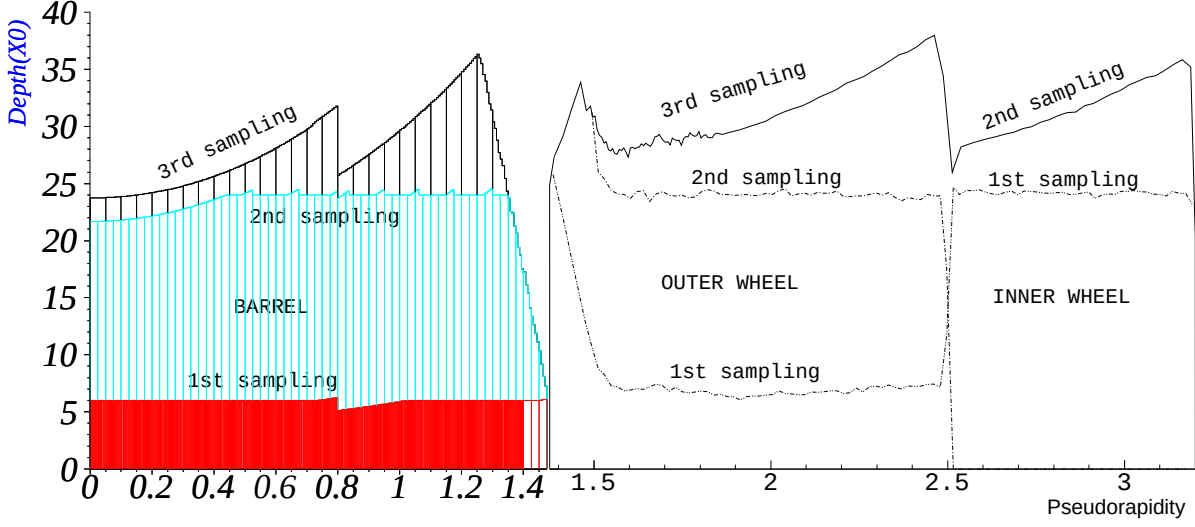


Figure 3.8: *Distribution of radiation lengths within the ECAL. The sizes of the calorimeter towers are indicated within the barrel section only. (Adapted from [21])*

the ECAL see each of these samplings in turn:

The first sampling ... is equipped with narrow strips with a pitch of ~ 4 mm in the η direction. This component acts as a “preshower” detector, enhancing particle identification (γ/π^0 , e/π separation, etc.) and providing a precise measurement in η . The second sampling is transversely segmented into square towers of size $\Delta\eta \times \Delta\phi = 0.025 \times 0.025$ ($\sim 4 \times 4$ cm² at $\eta = 0$). ... The third compartment has a twice coarser granularity in η .

Most of the ECAL’s 24 to 36 radiation lengths (see Figure 3.8) are to be found in the region of the second and third samplings. This thickness is expected to be sufficient to keep to an acceptable level the contributions to the error on the missing momentum measurement, when parts of high energy showers (> 500 GeV) manage to pass all the way through the calorimeter.

3.2.3.2 Hadronic calorimeter

ATLAS contains hadronic calorimeters of three different designs. All three are sampling calorimeters. In the barrel and extended barrel regions, the calorimeters are composed of iron absorber plates interspersed with plastic scintillator plates which are read out via wavelength shifting fibres. In the hadronic end caps the design consists of thick copper absorber plates separated by instrumented gaps filled with LAr. The forward HCAL is denser still. There, an absorber matrix of tungsten slugs fills the space between an array of grounded copper tubes, down whose centres run solid tungsten electrodes. The electrodes are separated from the walls of the copper tubes by a ~ 0.25 mm thick “active” layer of LAr. Remaining gaps in the forward HCAL are filled with “inactive” LAr.

3.2.4 Muon spectrometer

Designed with the intention of being able to measure muon momenta above 300 GeV with an accuracy of at least $\Delta(\frac{1}{p}) < 1 \times 10^{-4} \text{ GeV}^{-1}$, the muon spectrometer is the largest of the detector subsystems. With rapidity coverage up to $|\eta| = 2.7$ (and triggering capability over $|\eta| < 2.4$) the muon spectrometer should reach most of its physics goals, for example, a kinematic acceptance of 62% for $H(150 \text{ GeV}) \rightarrow \mu^+ \mu^- \mu^+ \mu^-$ decays.

The magnetic field of the muon spectrometer is toroidal, and maintained by a number of superconducting magnets arrayed radially about the beam with an eightfold symmetry. Although every effort is made to keep fields and material distributions uniform in ϕ , inevitably the distortion of the field in the vicinity of the eight barrel and sixteen endcap torroids introduces a ϕ dependence of the momentum resolution which has to be modelled properly by the fast simulation.

For the determination of momentum, locating the track position in the *bending plane*

of the magnetic field is more important than in the *out of plane* direction. Accordingly the majority of the muon spectrometer's precision measurements come from arrays of drift tubes^a aligned with $\eta \approx \text{const.}$ It is reasonable to expect ≥ 6 precision η measurements (with a single wire precision of $\sim 80 \mu\text{m}$) on each stiff muon track. To assist with the process of track fitting and pattern recognition, an orthogonal co-ordinate is also needed, typically ϕ . Since it is also the case that the drift tubes' charge collection times exceed the bunch crossing period of 25 ns, a trigger system of a different technology is also required. Resistive plate chambers in the barrel and thin gap or cathode strip chambers in the end caps step into this vacancy, both acting as trigger chambers and supplying an orthogonal measurement. Most stiff tracks pass through three of these trigger chambers, whose locations and positions are optimised so as to match their acceptance to that of the precision chambers.

3.3 Summary

The construction of the ATLAS detector has now been described in some detail, but with the turn on date for the LHC still some years in the future, we must now turn to discussion of the means by which the detector is *simulated* for the study which follows.

^aThe arrays of drift tubes are often referred to as “monitored drift tube stations” (MDT stations). Here, the word “monitored” does not refer to the fact that they are read out, but instead indicates that large scale deformations in the shape of the entire array are monitored by an internal alignment system.

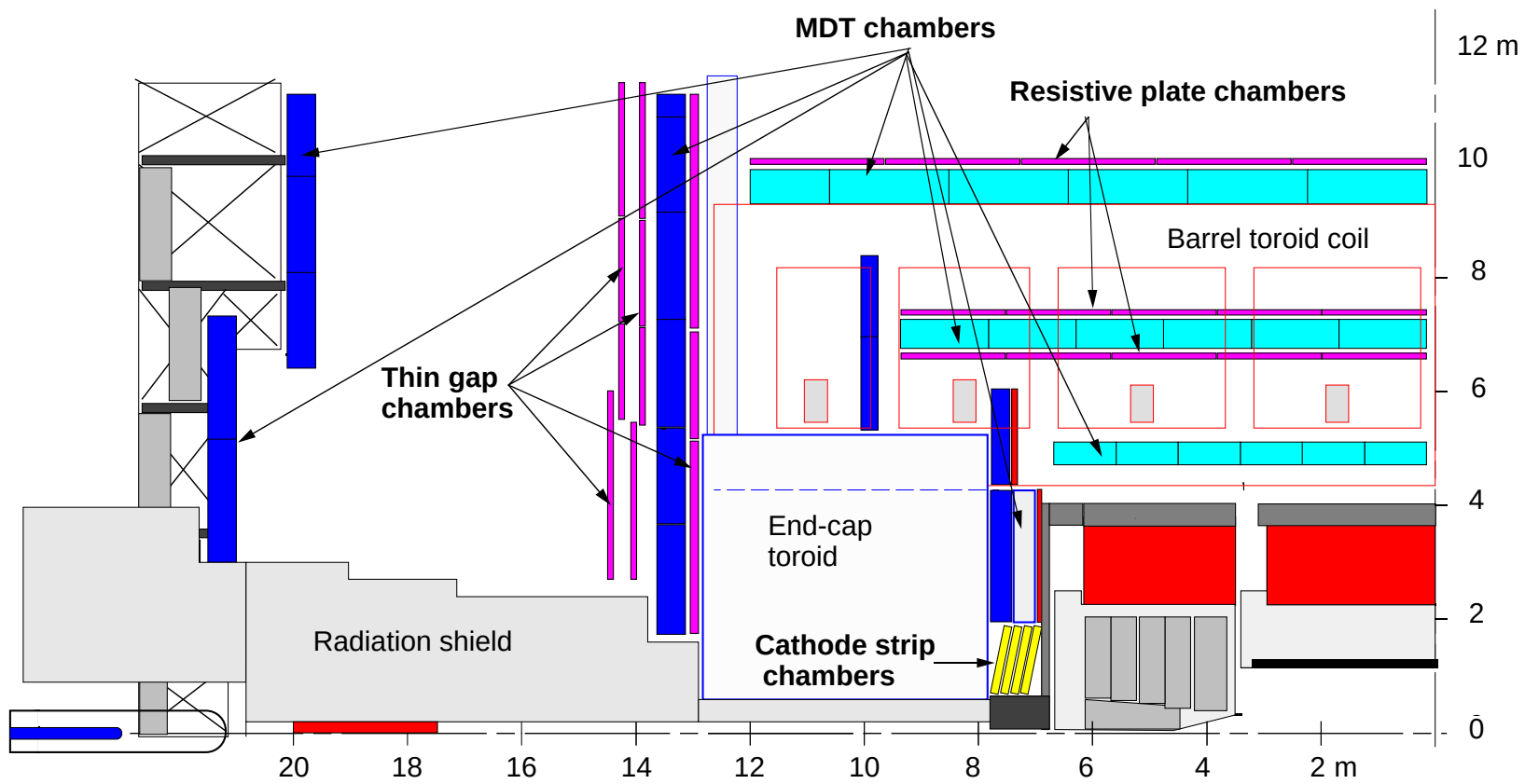


Figure 3.9: Side view of one quadrant of the muon spectrometer. (From [22])

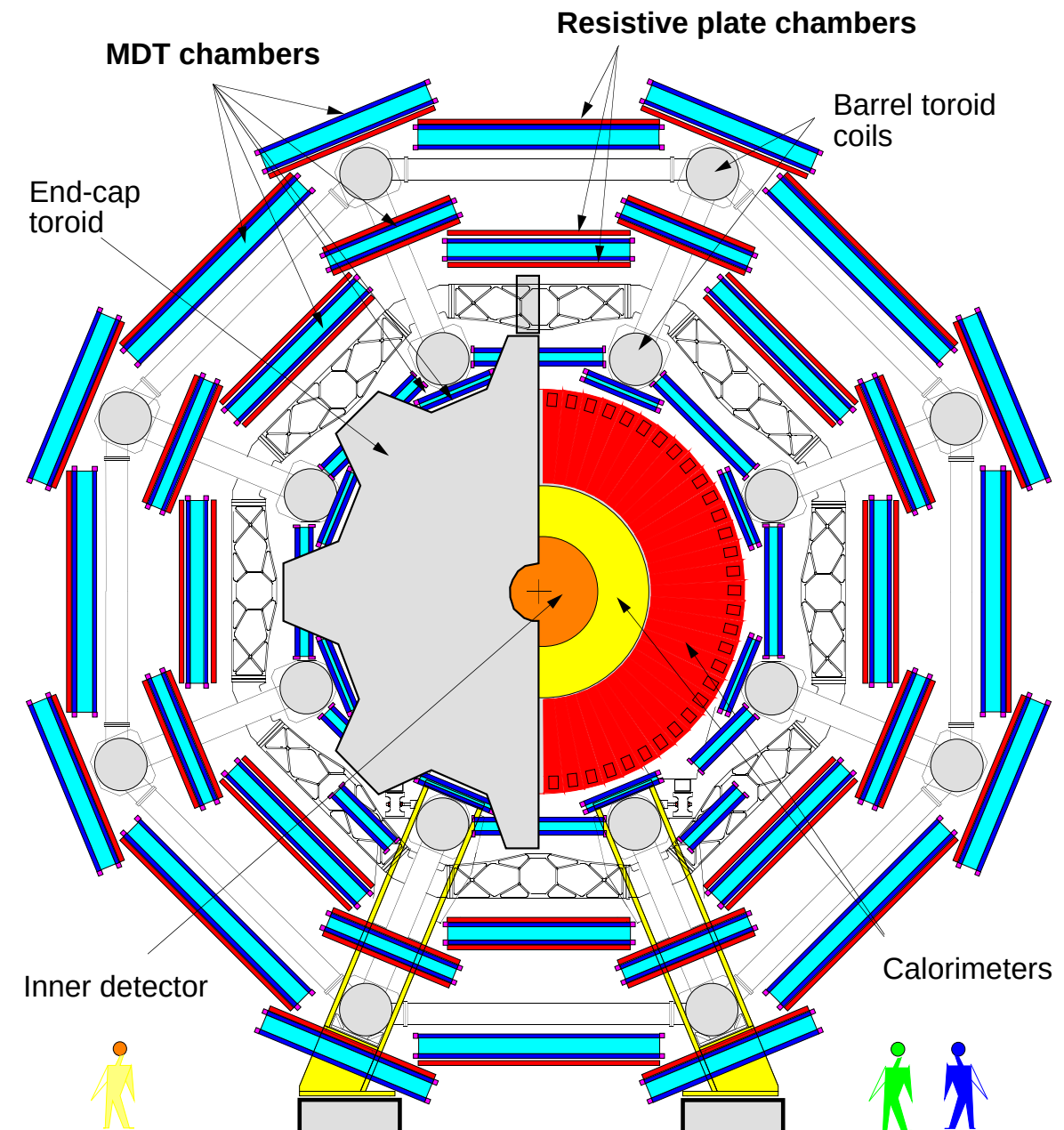


Figure 3.10: *End view of the barrel of the muon spectrometer. (From [22])*

Chapter 4

The ATLFAST Detector Simulation

4.1 Introduction

The ATLAS detector is simulated using the ATLFAST-2.16 computer program [23]. ATLFAST is a “fast” simulation which trades the accuracy of full simulation for significant gains in processing speed. ATLFAST achieves times of much less than one second per event, making event generation the speed limiting process.

4.2 Simulation and reconstruction method

4.2.1 “Calorimetric cells” and “clusters”

At the core of the ATLFAST simulation and reconstruction model are “calorimetric cells” and “clusters”.

As far as the simulation is concerned, a calorimetric cell is a somewhat idealized object which makes no distinction between electromagnetic and hadronic calorimetry.

It is a computational tool not a piece of the detector. Each cell spans a particular η and ϕ range, and may have a quantity of deposited energy associated with it. Within $|\eta| < 3$ the cells have size $\Delta\eta \times \Delta\phi = 0.1 \times 0.1$, while outside this range $\Delta\eta \times \Delta\phi = 0.2 \times 0.2$. When modelling ATLAS’s response to an event, ATLFASST begins by visiting the final state particles in turn. Unless a particle is a muon, a neutrino or the LSP, and provided that it is also within the detector acceptance, the particle’s energy is added to that of the calorimeter cell at the same η and ϕ . No randomized smearing of particle momenta takes place at this stage – but some smearing is induced by the process of dumping energy into discrete η and ϕ bins. In particular, the particle’s energy, direction and momentum all retain their “truth” values. The effects of magnetic fields are also neglected.

Once all appropriate cells have been filled, clustering begins. This process groups a sub-set of cells into jet-like^a clusters according to the following cone algorithm. Cluster formation is initiated at cells with more than 1.5 GeV of p_T (the highest first), and continues via the incorporation of all additional (unclustered) cells within a radius $r_{\text{clus}} = 0.4$ of the initiator. Although the cluster *cone* always remains centred on the initiator cell, the cone axis should not be thought of as the “centre” of the cluster. That role is taken by the cluster’s *barycentre*, whose position is independently calculated from the individual cell deposits. For example, “distance to cluster” will in future be taken to mean “distance to cluster barycentre”. Provided that the total p_T accumulated by the cluster then exceeds $p_T^{\text{min}}_{\text{clus}}$ (10 GeV), the cluster is accepted and the cells will not be considered for any further clusters. At this stage there has still been no direct or “deliberate” smearing.

Next, ATLFASST begins the task of identifying individual particles, and applying realistic resolutions to their “reconstructed” momenta. This begins with muons, continues with electrons and photons, and finishes with jets and taus.

^aClusters are “jet-like” in that they are produced by a clustering algorithm, but they are certainly not all “jets” as they can be produced by isolated electrons, for example.

4.2.2 Muons, electrons and photons

Muons, electrons and photons are treated as follows: (numerical values for cuts the mentioned may be found in Table 4.1)

1. Identification and smearing:

Final state muons, electrons and photons are sought, in that order, in the Monte Carlo event record. To mimic η -, ϕ - and momentum-dependent detector resolutions, an ‘elementary reconstruction’ is performed by smearing the four-momenta of those final state particles according to a set of parametrizations described in Section 4.3. The ‘true’ momenta, having been usurped by the new ‘smeared’ momenta, play no further part in either the simulation or this thesis and so may be forgotten. In the sections which follow, the qualifier “smeared” will be omitted in connection with muon, electron and photon momenta, but should be assumed to be present.

2. Acceptance:

Having been selected, each particle is then required to satisfy basic acceptance criteria: $|\eta| < 2.5$ and $p_T > 5$ GeV (6 GeV for muons). Those that survive pass on to the next step.

3. Isolation:

The muon detector allows muons to be observed in circumstances where the tracker and calorimeters alone would have trouble; such as when muons are buried in jets. Electrons and photons, on the other hand, may only be seen easily when they are isolated from other activity. Accordingly, within **ATLFAST**, muons are reconstructed in two varieties (termed isolated and non-isolated) while electrons and photons are only reconstructed if isolated. There are two separate requirements for isolation: isolation from clusters and isolation from activity.

Requirement	Isolated			Non-isolated		
	μ	e	γ	μ	e	γ
Minimum p_T required in order to be considered for “reconstruction”	6.0 GeV	5.0 GeV	5.0 GeV	6.0 GeV	Not reconstructed	Not reconstructed
Maximum $ \eta $ required in order to be considered for “reconstruction”	2.5	2.5	2.5	2.5		
Radius of cone, r_{dep} , over which energy is summed when isolation is assessed	0.2	0.2	0.2	–		
Maximum energy ($p_{T\text{ dep}}^{\text{max}}$) which may be deposited, by other sources, within r_{dep} of the particle	10.0 GeV	10.0 GeV	10.0 GeV	–		
$r_{\text{xj}}^{\text{min}}$. Minimum distance permitted between particle and <i>any</i> cluster	0.4	–	–	–		
Maximum distance permitted between particle and <i>closest</i> cluster	–	–	–	0.4		
$r_{\text{xj}}^{\text{min}}$. Minimum distance between the particle and <i>any cluster with the exception of the closest</i> , but with this exception being conditional upon the closest cluster being within $r_{\text{xc}}^{\text{max}}$ of the particle	–	0.4	0.4	–		
$r_{\text{xc}}^{\text{max}}$. Maximum permitted separation between the particle and its associated calorimetric cluster	–	0.15	0.15	–		

Table 4.1: *This table summarises the parameters which ATLFASST uses to determine which final state muons, electrons and photons from the montecarlo event will be “reconstructed” by the simulation, as if “measured” by ATLAS. The table’s requirements compare pseudo-measured particles (which is to say particles stolen from the montecarlo truth but which have been smeared according to resolutions described elsewhere in this document) with quantities derived from the simulation’s “calorimetric cells” and “clusters” (see Section 4.2.1). The p_T and $|\eta|$ requirements define the basic acceptance criteria which particles must satisfy in order to be “reconstructed”. The remaining requirements determine whether or not activity elsewhere in the event will cause the reconstruction to fail, and also determine whether the particles will be considered isolated or not. Lengths are measured in the dimensionless units of (η, ϕ) space.*

- **from clusters:**

A jet can throw secondary particles some distance from its core, particularly where heavy flavour leptonic decays are involved. Accordingly, isolated muons are required to be separated from all cluster barycentres by at least $r_{\text{xj}}^{\text{min}}$. This requirement is weakened a little for electrons and photons, since they are able to deposit energy into the simulation’s calorimetric cells, and have a habit of forming their *own* clusters: if the nearest cluster to an electron or photon is found to be within $r_{\text{xc}}^{\text{max}}$, the cluster is tagged as such, and overlap is permitted with this and only this cluster.

- **from activity:**

As well as being distant from other clusters, an isolated particle must find itself surrounded by a total deposit of less than $p_{T\text{ dep}}^{\text{max}}$ of transverse momentum within a range of r_{dep} . This removes, for example, particles obscured by general activity from the underlying event and leptons in broad jets. For muons, this energy deposit is calculated by simply summing the transverse momenta in all cells within r_{dep} of the particle. For electrons and photons, the same sum is made, but the p_T of the particle itself (which would have contributed to that sum) is then subtracted to leave the “additional” momentum.

4. Reconstruction:

A particle passing all these tests may then officially be considered “reconstructed”, allowing its momentum to be recorded. In addition, where a cluster was found to belong to an electron or a photon in Step 3, it is removed from the list of clusters. Thus, once assigned to a reconstructed particle, a cluster can neither be appropriated by another particle, nor play a further role in preventing reconstruction of others on grounds of “isolation from clusters”. This will prove to be important later, as the clusters remaining at the end of the whole process will form the starting point for reconstructed jets.

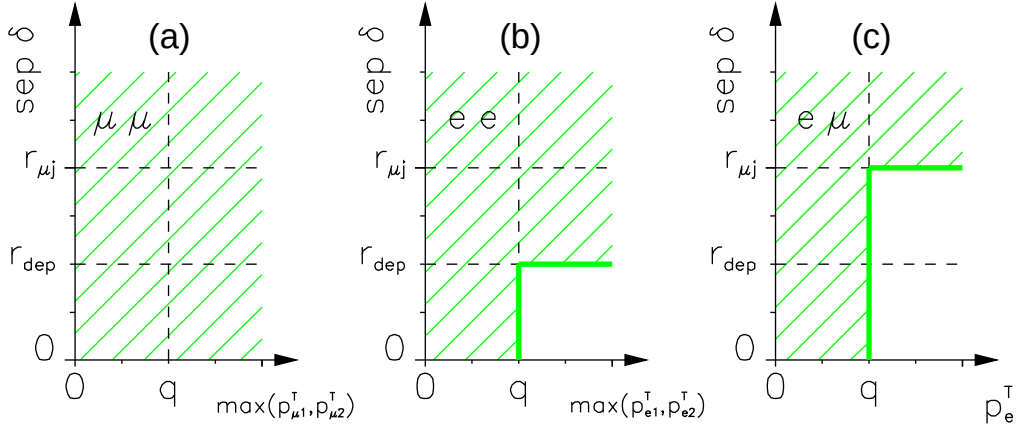


Figure 4.1: *Acceptance of simulated lepton pairs. The shaded regions indicate the separations and momenta for which the detector simulation is able to reconstruct both leptons of a pair. In all cases, the separation in (η, ϕ) -space of the two leptons lies on the vertical axis, while the relevant momentum variable lies on the horizontal axis. Plots are shown for the cases: (a) $\mu\mu$, (b) ee and (c) $e\mu$. In the $e\mu$ case there is no dependence on p_μ^T . The momentum cut value, q , denotes $p_{T\text{dep}}^{\text{max}}$ and $p_{T\text{clus}}^{\text{min}}$ which share the same value (10 GeV) throughout this analysis.*

5. And repeat:

The algorithm then returns to Step 1 in order to consider the next muon from the event record. When all muons have been found, the whole process repeats for electrons, and finally for photons.

No other effects likely to modify the efficiency for lepton reconstruction are simulated, beyond those described above (detector acceptance, detector thresholds, isolation effects *etc*). It should be noted that cluster removal in Step 4 has a side effect: reconstruction of pairs of adjacent leptons occurs in a different way for each of the three different flavour combinations. For example:

- Consider a pair of arbitrarily close muons. Both will be classified as isolated, regardless of their separation, since they neither produce clusters nor deposit energy in calorimetric cells and so cannot interfere with each other. (Figure 4.1(a))

- Consider, on the other hand, a pair of arbitrarily close electrons. Both of these will be reconstructed only if each satisfies $p_T < p_{T\text{dep}}^{\text{max}}$. If one or either has more p_T than this, they will have to have a finite separation of at least r_{dep} , in (η, ϕ) -space, in order for both to reconstruct. (Figure 4.1(b))
- Consider now an arbitrarily close electron and muon. The electron will always be reconstructed (since the muon neither generates a cluster nor deposits energy in calorimetric cells and so cannot interfere with the electron). The muon, however, will only be considered isolated if both of the following are true: (Figure 4.1(c))
 1. *Either* the leptons are separated by more than r_{xj} , *or* the electron’s energy is less than $p_{T\text{clus}}^{\text{min}}$ (preventing it from being energetic enough to form a cluster by itself) *or both*.
 2. *Either* the leptons are separated by more than r_{dep} , *or* the electron’s energy is less than $p_{T\text{dep}}^{\text{max}}$ *or both*.^b

It is important to consider these “small separation” effects because much of the later analyses will be concerned with end points of invariant mass distributions of groups of leptons. While it is true that higher invariant masses of *pairs* of leptons occur at *large* opening angles, the same can no longer be said once the case of three or more leptons is considered. There, the configuration of highest invariant mass is often one in which (in a particular particle’s rest frame) two leptons are parallel while the other is anti-parallel. This leads to the possibility that end points are sensitive to the treatment of leptons with small separation. Note that all the cases in Figure 4.1 treat low momenta even handedly. It is only at momenta above $p_{T\text{dep}}^{\text{max}}$ or $p_{T\text{clus}}^{\text{min}}$ that differences appear. Whether edges in a given process are likely to be affected significantly will clearly depend on the available

^bWithin this thesis, the second condition happens to be unnecessary as it is automatically implied by the first. This is because $p_{T\text{dep}}^{\text{max}}$ and $p_{T\text{clus}}^{\text{min}}$ are, as it happens, both equal to 10 GeV, and because $r_{\text{dep}} < r_{\text{xj}}$.

energy (Are any of the decays only-just above threshold?) and on the boost of the initial particle. However, it is clear that more work needs to be done to understand whether this effect ever becomes a significant systematic limiting simulation of $e^+e^- + \mu^+\mu^- - e^\pm\mu^\pm$ background subtraction, or measurements of slepton nonuniversality.

It is estimated in [24] that it should be possible to reduce the error on the absolute muon momentum scale to 0.1% via analysis of $Z \rightarrow \mu^+\mu^-$ events. For electrons the target uncertainty is even lower: 0.02%. At this level it should be possible to make maximum use of available statistics from $W \rightarrow e\nu$ events, to provide a competitive measurement of the W mass. Throughout this thesis, however, muons and electrons are treated on an equal footing, and so a conservative error for the absolute lepton momentum scale of 0.1% will be assumed. This error is an order of magnitude smaller than the error on the absolute jet energy scale (see next section) and so is not significant in the subsequent analysis.

4.2.3 Jets

The starting point for jet reconstruction is the set of clusters which remain after those belonging to electrons and photons have been removed. There are three stages: cluster smearing, muon recombination and rejection. During the cluster smearing stage, each cluster is visited in turn. All the components of a given cluster's momentum are then rescaled by a common factor distributed such that $\Delta E_e/E_e$ has a width given later by Equation (4.8). Once all clusters have been scaled in this way, muon recombination begins. All the muons that fell into the “non-isolated” category during the reconstruction phase are examined in order of decreasing p_T . Those muons within r_{clus} of the nearest cluster barycentre are then defined to be part of that jet, and the cluster's momentum has the muon's added to it. As a consequence the cluster's barycentre may move. After all appropriate non-isolated muons have been added to clusters, the modified clusters are

inspected, and those still within acceptance ($|\eta| \leq 5.0$) and whose transverse momentum exceeds $p_T^{\min}$ (15 GeV) are officially promoted to the status of “reconstructed jets”.

It is estimated in [24] that it should be possible to reduce the error on the absolute jet energy scale (for jets with transverse momenta between 70 GeV and ~ 500 GeV) to “below 1%”. The primary means of determining this scale would be by analysis of $W \rightarrow jj$ decays in $t\bar{t}$ events. At higher rapidities, calibration may be better achieved by balancing the transverse momenta of single jets recoiling against leptonically decaying Z -bosons. Study of this method (also in [24]) suggests that for transverse momenta above 120 GeV the energy scale might be determined to 0.4%. Throughout this thesis, an error on the jet energy scale of 1.0% will be assumed.

4.2.4 Missing transverse momentum

The final stage implemented in ATLFASST’s detector simulation is the missing transverse momentum, or $\cancel{p}_T$, “measurement”. By definition, $\cancel{p}_T$ is the momentum vector pointing in the opposite direction to the visible transverse momentum. Accordingly, all visible momentum must be summed. Visible momentum can be divided into two categories. The first is momentum carried by particles which have already been reconstructed. The second, “the rest”, is largely a mixture of soft pions, leakage from jets, gluon radiation, and particles below reconstruction thresholds. It is easy to sum momenta of the first kind. The full list of objects going into this sum is as follows: jets, non-used clusters, isolated muons, non-isolated muons which were not added to clusters, isolated electrons, and isolated photons. This leaves momenta of the second kind, and these are to be found in the calorimetric cells which were never incorporated into clusters. Each calorimetric cell is visited in turn, and the energy contained therein is smeared according to the resolution defined in Equation (4.8) as described in Section 4.3.3. After smearing, calorimetric cells found to contain energies above a certain value are retained,

and their transverse energies are summed vectorially. The cut is placed at 0 GeV in this analysis in keeping with the usual procedure. The sum of this transverse vector and the one obtained from the individually reconstructed particles is (up to a sign) the missing transverse momentum vector returned by ATLFASST.

4.3 Resolutions

4.3.1 Electron and photon resolution

Within the range $-4 < \eta < 4$, photon and electron energies are smeared according to

$$\frac{\Delta E_\gamma}{E_\gamma} = \begin{cases} \frac{0.10 \text{ GeV}^{\frac{1}{2}}}{\sqrt{E}} \oplus \frac{0.245 \text{ GeV}}{p_T} \oplus 0.007 & \text{at low luminosity, or} \\ \frac{0.10 \text{ GeV}^{\frac{1}{2}}}{\sqrt{E}} \oplus \frac{0.245 \text{ GeV}}{p_T} \oplus 0.007 \oplus \frac{r^{e,\gamma}(\eta)}{p_T} & \text{at high luminosity} \end{cases} \quad (4.1)$$

and

$$\frac{\Delta E_e}{E_e} = \begin{cases} \frac{0.12 \text{ GeV}^{\frac{1}{2}}}{\sqrt{E}} \oplus \frac{0.245 \text{ GeV}}{p_T} \oplus 0.007 & \text{at low luminosity, or} \\ \frac{0.12 \text{ GeV}^{\frac{1}{2}}}{\sqrt{E}} \oplus \frac{0.245 \text{ GeV}}{p_T} \oplus 0.007 \oplus \frac{r^{e,\gamma}(\eta)}{p_T} & \text{at high luminosity} \end{cases} \quad (4.2)$$

where

$$r^{e,\gamma}(\eta) = \begin{cases} 0.320 \text{ GeV} & \text{for } |\eta| \leq 0.6 \\ 0.295 \text{ GeV} & \text{for } 0.6 < |\eta| \leq 1.4 \\ 0.270 \text{ GeV} & \text{for } 1.4 < |\eta| \end{cases} \quad (4.3)$$

is a quantity designed to parametrize the effects of pile-up and electronic noise at high luminosity as determined in [25]. In addition, the polar angle “ θ ” of each photon is

smeared according to an η dependent resolution $\Delta\theta = \frac{\sigma_\theta(\eta)}{\sqrt{E}}$ where^c

$$\sigma_\theta(\eta) = \begin{cases} 0.065 \text{ GeV}^{\frac{1}{2}} & \text{for } |\eta| \leq 0.8 \\ 0.050 \text{ GeV}^{\frac{1}{2}} & \text{for } 0.8 < |\eta| \leq 1.4 \\ 0.040 \text{ GeV}^{\frac{1}{2}} & \text{for } 1.4 < |\eta| \leq 2.5 \\ 0.000 \text{ GeV}^{\frac{1}{2}} & \text{for } 2.5 < |\eta| \end{cases} \quad (4.4)$$

4.3.2 Muon resolution

Contributions to the measurement of muon momenta within ATLAS come from two sources. The first is the muon spectrometer itself, and the second the inner detector. In both cases, it is the shape of the track in regions of strong magnetic field which provides a handle on the particle momenta. The model which **ATLFAST** uses to parametrize muon momentum resolution is best illustrated with a detector which is simpler than ATLAS. In reality the ATLAS B -field is a complicated object, and the track shape is the result of a fit to many precision space points. However, the form of the momentum dependence of the resolution may be obtained more simply by considering a spectrometer, of known length and containing a uniform B -field, which makes a single measurement of the track sagitta with a Gaussian distributed error.

As illustrated in Figure 4.2, a measurement of the sagitta with Gaussian errors is equivalent to a Gaussian measurement of $1/r$. Since a particle’s momentum is proportional to its radius of curvature, $1/p$ rather than p is also measured with Gaussian errors.

Consequently, **ATLFAST** smears a muon’s momentum by “performing two independent measurements of $1/p$ ” (one in the muon system and one in the inner detector) each of

^cNote that the discrepancy between the parametrization listed in Equation (4.4) for $1.4 < |\eta| \leq 2.5$ and that contained in Section C.3 of the **ATLFAST** documentation [23] is due to a simple typographical error on their part.

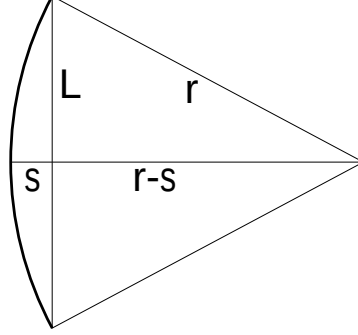


Figure 4.2: The bold line represents the path of a charged particle in a constant B -field. Also labelled are r (the local radius of curvature of the track), L (the length scale associated with the particle track) and the sagitta, s . We note $r^2 = (r - s)^2 + L^2$ from which $r = (L^2 + s^2)/2s$. Thus $r \propto 1/s$ when $s^2 \ll L^2$.

which are subject to smearing according to their own eta and phi dependent resolutions $\Delta_{\text{muon}}(1/p)$ and $\Delta_{\text{inner}}(1/p)$. These two measurements are then combined into a single average, each weighted according to its own resolution, i.e.:

$$\frac{1}{p_{\text{avg}}} = \frac{w_m(1/p_{\text{muon}}) + w_i(1/p_{\text{inner}})}{w_m + w_i} \quad (4.5)$$

where $w_m = 1/\Delta_{\text{muon}}^2(1/p)$ and $w_i = 1/\Delta_{\text{inner}}^2(1/p)$.

In the case of the inner detector measurement, ATLFASST implements a simple parametrization of $\Delta_{\text{inner}}(1/p)$ of the form

$$\Delta_{\text{inner}}(1/p) = a_{\text{track}} \oplus \frac{0.012}{p_T} \quad (4.6)$$

where the first term

$$a_{\text{track}} = 0.0005 \left(1 + \frac{|\eta|^{10}}{7000} \right) \text{ GeV}^{-1} \quad (4.7)$$

may be thought of as the sagitta resolution. The $|\eta|$ part crudely models the sharp drop off in the inner detector's momentum resolution at the end of the solenoid, where the divergent B -field is both reduced in strength and becomes parallel to the particle

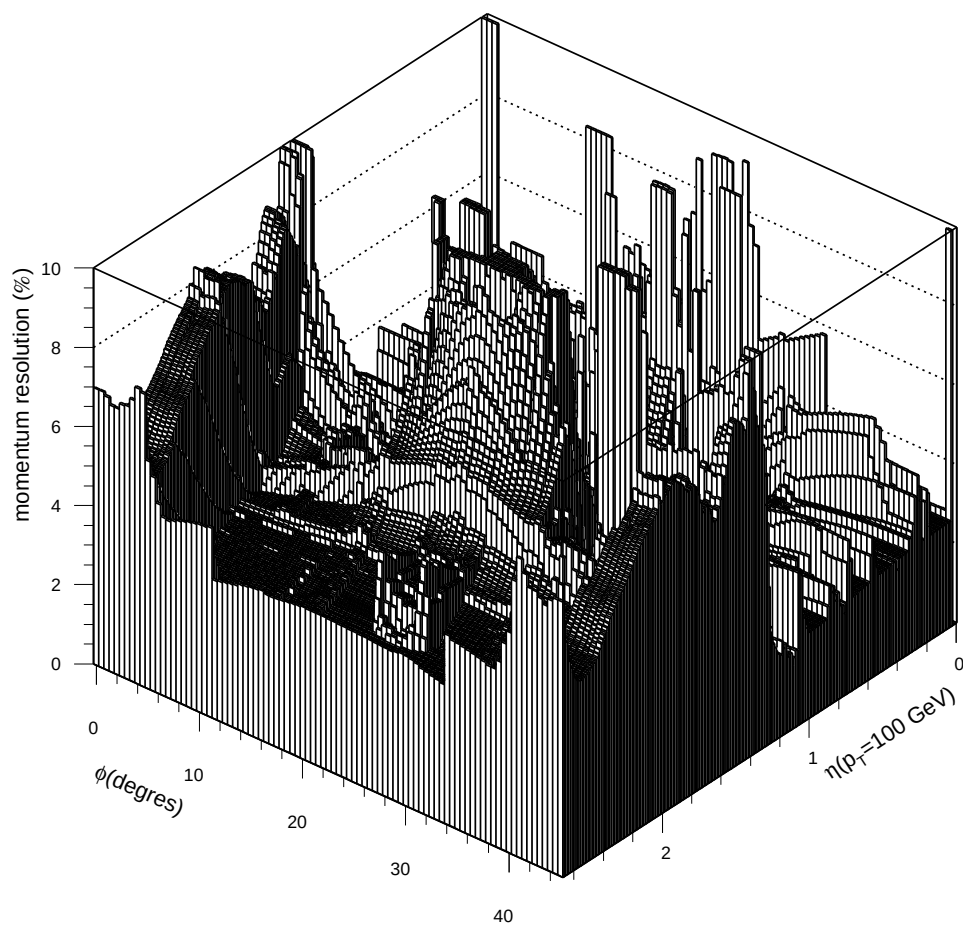


Figure 4.3: *Standalone momentum resolution (for a $p_T = 100$ GeV muon) for the muon spectrometer as a function of η and ϕ . (Figure is taken from Chapter 6 of [24])*

tracks (see Figure 3.4). This $|\eta|$ term approaches 1 at around $|\eta| \approx 2.4$. The second term in (4.6) models effects which become important at low momenta, such as multiple scattering which limit the resolution of the spectrometer. Both the parametrization of the inner detector momentum resolution and the momentum weighting are described in [26].

In calculating $\Delta_{\text{muon}}(1/p)$ for the muon spectrometer, **ATLFAST** employs a more sophisticated parametrization. Detailed studies fully simulating the muon spectrometer were used to map out the muon stand-alone momentum resolution as a function of η and ϕ . The studies were performed first for 100 GeV and then for 1000 GeV muons. Resolutions for muons of other momenta are obtained by interpolation. Mapping the resolution as a function of ϕ is particularly important, as the eight barrel toroid magnets significantly affect the amount of material through which muons travelling in their directions have to pass. The results of this full simulation are stored in a database within **ATLFAST**, and are sampled to generate $\Delta_{\text{muon}}(1/p)$ appropriately. Figure 4.3 shows how the momentum resolution of the muon spectrometer is expected to vary as a function of η and ϕ .

4.3.3 Resolution for jets and unclustered calorimetric cells

ATLFAST parametrizes hadron resolution according to

$$\frac{\Delta E_j}{E_j} = \begin{cases} \frac{0.50 \text{ GeV}^{\frac{1}{2}}}{\sqrt{E}} \oplus s^{\text{had}}(\eta) & \text{at low luminosity, or} \\ \frac{0.50 \text{ GeV}^{\frac{1}{2}}}{\sqrt{E}} \oplus s^{\text{had}}(\eta) \oplus \frac{r_{\text{had}}(r_{\text{param}})}{p_T} & \text{at high luminosity} \end{cases} \quad (4.8)$$

where

$$r_{\text{had}}(r_{\text{param}}) = \begin{cases} 0.4 \text{ GeV} & \text{for } r_{\text{param}} < 0.4 \\ 7.5 \text{ GeV} & \text{for } r_{\text{param}} = 0.4 \\ 12.0 \text{ GeV} & \text{for } 0.4 < r_{\text{param}} \leq 0.5 \\ 18.0 \text{ GeV} & \text{for } 0.5 < r_{\text{param}} \leq 0.7 \\ 20.0 \text{ GeV} & \text{for } 0.7 < r_{\text{param}} \end{cases} \quad (4.9)$$

is a quantity which seeks to account for the effects on calorimeter performance of noise and pile up at high luminosity, and where the small constant term

$$s^{\text{had}}(\eta) = \begin{cases} 0.030 & \text{for } |\eta| < 3.0 \\ 0.070 & \text{for } |\eta| \geq 3.0 \end{cases} \quad (4.10)$$

accounts for the inhomogeneity of the calorimeter. When used to smear jet energies, r_{param} takes the value r_{clus} (defined in Section 4.2.1) whereas when used in the missing momentum measurement to smear unclustered calorimetric cells (see Section 4.2.4) it takes the value $r_{\text{param}} = 0$.

Chapter 5

Analysis

5.1 Introduction

The biggest problem facing experimenters wishing to study supersymmetry is the size of the parameter space.^a While this problem has not been *ignored* in the past, there has perhaps never been a better time to expect a discovery of supersymmetry than in the coming years at the LHC, so it is time to ask “How can supersymmetry searches actually be performed and particle masses measured?” Indeed, is it realistic to believe that they can be performed at all?

To put the methods used in this thesis into context, it will be helpful to precede their description with an outline of the main ways in which the task of constraining supersymmetry has been approached in the field as a whole. The emphasis in this discussion will be on the general styles of approach which have been used, rather than on particular analyses, although specific mention will be made of one (namely [27]). For the purposes of the discussion, the styles of analysis used have been grouped under four

^aOne could also say that the sheer size of the parameter space is a bit of an embarrassment to many theorists!

headings, but in practice it is rare for a particular analysis to fall exclusively into only one of the categories.

5.1.1 (Bayesian) measurements of scale

One way of attempting to acquire information about supersymmetry dispenses with the idea of performing any type of analysis whose success is dependent on something specific happening in a subset of the events. This approach acknowledges that the space of possible models is at present so large as to make the conclusions of analyses based on only small fractions of it a little unsatisfying. The methods available to this approach are then usually restricted to the identification of observables which correlate well with one or more properties having a common meaning in all supersymmetric models; for example the total supersymmetric cross section, or a sufficiently well defined supersymmetric mass scale. Typically this correlation will be demonstrated by choosing points “at random” from a space of supersymmetric models, such as the MSSM, and by then obtaining at each point an estimate of both the “observable” and the characteristic of the model which the observable is designed to measure.

One such analysis is [27], which tests five observables for correlation with two definitions of the supersymmetric mass scale, and another for correlation with the total supersymmetric cross section. Figure 5.1, taken from [27], shows scatter plots of $M_{eff}^{susy} = M_{susy} - m_{\tilde{\chi}_1^0}^2/M_{susy}$ versus the best observable M_{est} for 100 points chosen at random from each of the three spaces of models indicated. Note that M_{est} is defined by

$$M_{est} = \not{p}_T + \sum_i |\mathbf{p}_T^{j_i}|, \quad (5.1)$$

where the sum extends over all jets with $p_T^j > 50$ GeV. The parameter M_{susy} is defined to be a weighted mean of the sparticle masses in each model; the weights being chosen to be

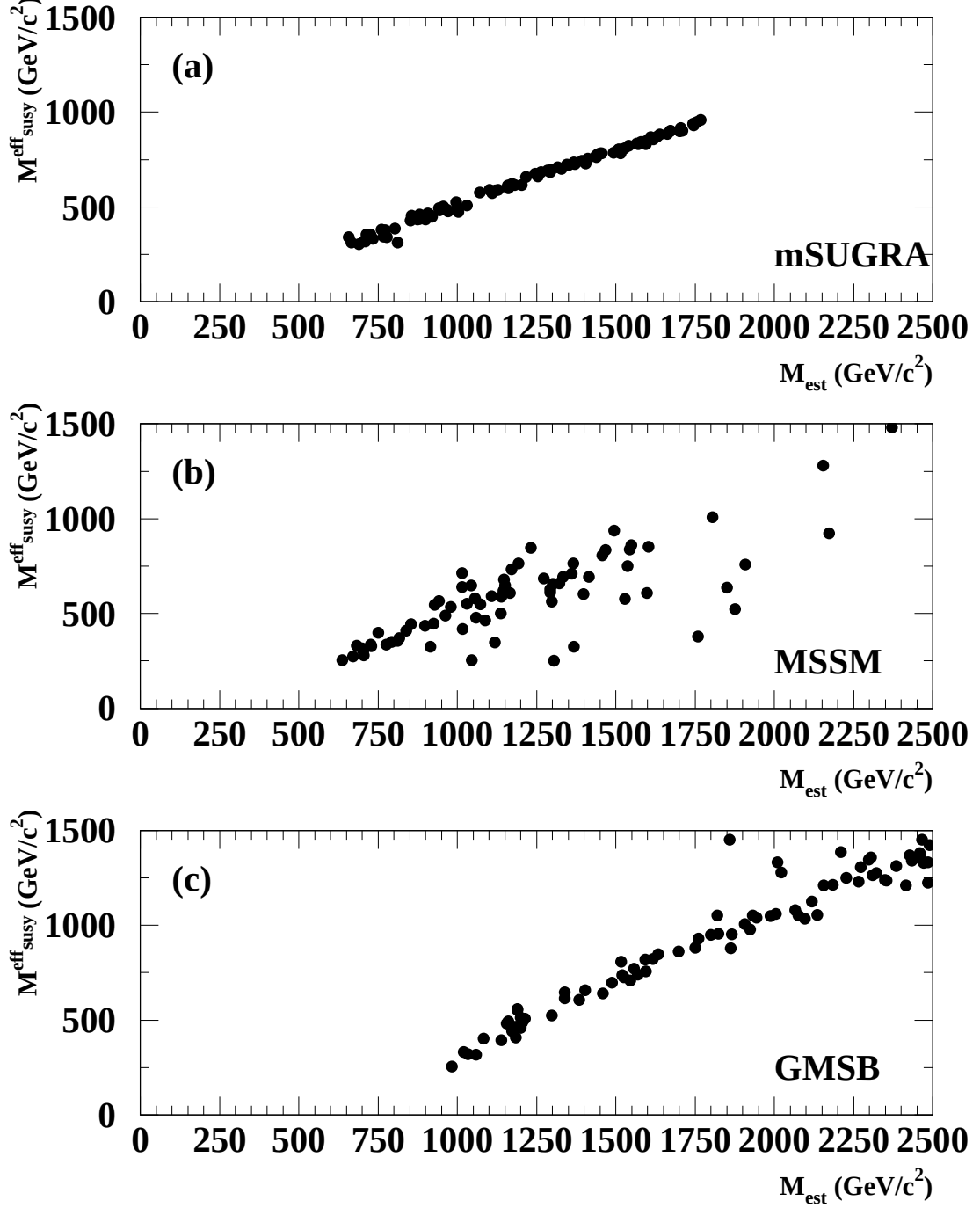


Figure 5.1: *This figure is taken from [27]. It shows scatter plots of $M_{\text{eff}}^{\text{susy}}$ against M_{est} for 100 points drawn “at random” from (a) mSUGRA models, (b) the MSSM, and (c) GMSB models (see [28]). The variables $M_{\text{eff}}^{\text{susy}}$ and M_{est} are defined in Section 5.1.1.*

proportional to relative production cross sections for each type of sparticle. The excellent performance of M_{est} in the mSUGRA and GMSB models is evident from Figures 5.1(a) and (c), but the non-universality of the correlation may readily be seen by comparing the differences between the scatter points for these two sets of models at low values of M_{est} . This observation is confirmed in Figure 5.1(b) in which points were drawn from the MSSM, and in which the spread of $M_{eff}^{susy} - M_{est}$ values is seen to be much broader. Nevertheless, there is still a clear difference between regions which points *are* able to “hit” and others where this is impossible, and so even at the MSSM these techniques could be used to establish useful allowed/excluded regions for (in this case) M_{eff}^{susy} .

Of course, the drawing of points “at random” from the set of all MSSM models, even if points are first checked to ensure that they have not already been excluded by an earlier study, is not something which can be achieved without the introduction of some prior distribution, according to which the points are chosen. This is not necessarily a bad thing, as physical insight and intuition can, at times, provide a good guide to sensible regions of parameter space. Nevertheless, an error attached to a measurement obtained by this approach, in a situation like that shown in Figure 5.1(b), is likely to depend very strongly on the choice of prior.

5.1.2 “Show you might be right”

This is probably the most common and at the same time the least satisfactory approach of them all. In a typical “show you might be right” study, a physicist has probably simulated events using a known (supersymmetric) model, has then predicted the existence of a salient feature which might be sought in the data, has duly found the said salient feature at the desired location, and has patted himself on the back. Although a somewhat dismissive tone has been used here to describe this approach, it should be remembered that the method itself is none other than that which has successfully been

used in the majority of the tests of the SM since the measurements of the masses of the W and Z bosons. When testing a theory, like the SM, whose parameters are very well pinned down and which is able to make strong experimentally-falsifiable predictions, there is very little about the approach which is contentious. But therein lies the problem as far as a highly unconstrained theory like supersymmetry is concerned; for while the approach is excellent for knocking established theories on the head, it is of very little use in identifying where in a vast uncharted parameter space one actually lies. In short, the discovery in nature of a phenomenon X , tells you less about the status of a model M which implies X , than the failure to discover X would!

5.1.3 “Model independent” techniques

These would be valuable, if only they existed. In contrast to the “assume a model with few parameters” approach, (described in the next section), the proponent of the “model independent” technique hopes to convince the reader that his or her conclusions are applicable in the broadest possible set of circumstances. He or she hopes to measure a property, such as a particle’s mass, which has a consistent interpretation in a large number of models. But, inevitably, papers using “model independent” techniques contain, either explicitly or implicitly, a large number of conditions and assumptions upon whose validity the success of the technique depends. Regrettably, it is very often unclear whether it would be possible for the experiment to confirm, independently, that these conditions are actually satisfied. Indeed, whether the conditions are likely or reasonable, at all, may even be open to question.^b

The classic tool of the “model independent” analysis is relativistic kinematics. By trying to rely only on four-momentum conservation, the physicist hopes that his or her

^bThis point is particularly pertinent when account is taken of the human tendency to neglect to publish or report on unexciting or null results.

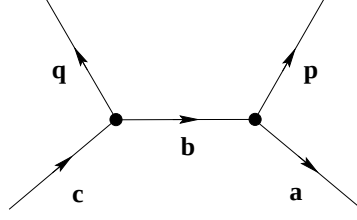


Figure 5.2: *Two successive two-body decays. Arrows indicate particles' directions of motion.*

results will retain the widest applicability. By way of an example, consider the two stage process $C \rightarrow QB$ followed by $B \rightarrow PA$, with A unobservable as illustrated in Figure 5.2. The kinematic feature which might be sought in this process is the upper end point of the invariant mass spectrum of $p_P^\mu + p_Q^\mu$. In particular, it is easy to show that, when P and Q are massless (or effectively massless),

$$(m_{PQ}^{\max})^2 = (m_C^2 - m_B^2)(m_B^2 - m_A^2)/m_B^2. \quad (5.2)$$

The above result is proved in the Appendix in Section E.1. Evidently, measurement of this endpoint provides a valuable constraint on the masses of the particles involved. However, the problem with using the measurement to pin down supersymmetric masses lies in the interpretation. Firstly, there is the problem of determining which masses were constrained. Secondly, one must ascertain whether the masses were constrained in the way which was first thought — in other words determine the kinematic structure *within* the process.

To give a concrete example, which will be used in earnest later on, the process above might represent the two stage decay of a heavy neutralino, C , into a lighter neutralino, A , and an opposite-sign same-lepton-family (OSSF) dilepton pair (P, Q) via an internal on-shell right slepton, B . In full, this example would be written as $\tilde{\chi}_2^0 \rightarrow l^+ \tilde{l}_R \rightarrow l^+ l^- \tilde{\chi}_1^0$. For reasons of space, however, it will be more convenient to use a form in which the SM

particles are omitted, and thus the above example will be termed $\tilde{\chi}_2^0 \rightarrow \tilde{l}_R \rightarrow \tilde{\chi}_1^0$ for short. Note, however, that in most mSUGRA models $m_{\tilde{l}_L} > m_{\tilde{l}_R}$, and both of these are trivially greater than $m_{\tilde{\chi}_1^0}$ if the LSP is uncharged. This means that if the second half of the decay chain is open for $\tilde{l}_R$ it is probably open for $\tilde{l}_L$ too. Consequently, in addition to the above, there might be significant branching ratios for any, or even all, of the following chains, each of which share the same final state: $\tilde{\chi}_4^0 \rightarrow \tilde{l}_L \rightarrow \tilde{\chi}_1^0$, $\tilde{\chi}_4^0 \rightarrow \tilde{l}_R \rightarrow \tilde{\chi}_1^0$, $\tilde{\chi}_3^0 \rightarrow \tilde{l}_L \rightarrow \tilde{\chi}_1^0$, $\tilde{\chi}_3^0 \rightarrow \tilde{l}_R \rightarrow \tilde{\chi}_1^0$, and $\tilde{\chi}_2^0 \rightarrow \tilde{l}_L \rightarrow \tilde{\chi}_1^0$. In fact there are still more chains capable of producing dilepton pairs distributed with the same kinematic structure and which are potentially confusable with the desired chain, namely: $\tilde{\chi}_4^0 \rightarrow \tilde{l}_R \rightarrow \tilde{\chi}_3^0$, $\tilde{\chi}_4^0 \rightarrow \tilde{l}_L \rightarrow \tilde{\chi}_3^0$, $\tilde{\chi}_4^0 \rightarrow \tilde{l}_R \rightarrow \tilde{\chi}_2^0$, $\tilde{\chi}_4^0 \rightarrow \tilde{l}_L \rightarrow \tilde{\chi}_2^0$, $\tilde{\chi}_3^0 \rightarrow \tilde{l}_R \rightarrow \tilde{\chi}_2^0$ and $\tilde{\chi}_3^0 \rightarrow \tilde{l}_L \rightarrow \tilde{\chi}_2^0$. A given experiment would presumably have to look for leptons and missing energy, and so would have trouble distinguishing events coming from any of these channels. The length of the list illustrates the problems involved in deciding, given the result of an experimental search, which the correct masses to participate in Equation 5.2 should be.

Assuming for the moment that the initial and final state particles themselves may be reliably identified, there still remains the second problem of understanding the kinematic structure *within* the process. In the example under consideration, were one or both of the sleptons to become *heavier* than the initial neutralino, the occurrence of an internal on-shell slepton would no longer be permitted. It is, nonetheless, often the case that decays to the same final state would still be possible, but now through diagrams involving *virtual* sleptons. The resultant three body decay, $C \rightarrow APQ$, has an end point for m_{PQ} located not at the position given in Equation 5.2, but instead at $m_C - m_A$, even if P and Q are massive. Thus, even if a search for leptons and missing energy reveals a distribution with an end point, it is by no means clear what constraint this implies for the particle masses, even if the identities of the initial and final state particles are themselves well known.

Note also that, while it is correct to point out that altering the internal kinematic structure of the decay chain will also alter the *shape* of the resultant kinematic distributions, a variety of other (often *model-dependent*) spin-related or matrix element effects will also modify the shape of the distributions with respect to phase space alone. Making use of this information would thus, in addition to requiring a greater understanding of the detector itself, force the abandonment of the original goal of “model independence”.

By way of a final caution, it should be pointed out that the chain discussed so far is probably the simplest non-trivial case one can imagine. The problems multiply as the decay chains lengthen, even if only by one stage. Consider, for example, a decay chain consisting of three successive two body decays, such as $\tilde{q} \rightarrow \tilde{\chi}_2^0 \rightarrow \tilde{l}_R \rightarrow \tilde{\chi}_1^0$. This chain can be mimicked by a two body decay followed by a three body decay, or a three body decay followed by a two body decay.

5.1.4 “Assume a model with few parameters”

The simplest and arguably most consistent way of making the supersymmetric parameter space more manageable, is to give up the pretence of generality, and consider instead any one of a highly constrained subset of models. Various models exist (mSUGRA, GMSB, anomaly mediated supersymmetry breaking (AMSB)) which manage to reduce the number of dimensions of parameter space to no more than five, and all see frequent use.

Given a manageable parameter space, $\mathcal{P}$, it is possible to imagine the “perfect” analysis; a “total fit to all data”. The “perfect” analysis would take a large number of independent distributions from the experimental data, and then try to fit them with high-statistics montecarlo data drawn from points spread throughout $\mathcal{P}$. Once obtained from the best fit, the model parameters could be used, in principle, to answer any

question about the model.

As far as I am aware, no study has yet been performed which can claim to have evaluated the likely performance of ATLAS using exactly this technique.^c The largest stumbling block is the expense of generating montecarlo over the whole of $\mathcal{P}$, even if $\mathcal{P}$ is only, say, four-dimensional. When faced with real data, it is probably true that a bit of careful thought and intuition will be able to narrow the region of interest in $\mathcal{P}$ by a significant amount, and so the hurdle of montecarlo generation might not be as high as it appears at first. Nevertheless, the studies [29, 30] which have, thus far, investigated ATLAS’s ability to measure mSUGRA parameters have found it necessary to avoid montecarloing at points in $\mathcal{P}$ *altogether*. Instead, the parameter space of a particular model is scanned *only* in order to generate sparticle spectra (lists of sparticle masses). That’s it. Once the mass spectra have been obtained, analysis follows the same path as the “show you might be right” method, which is to say that:

1. A particular point in $\mathcal{P}$ is chosen for detailed investigation. This point is thus “special” and a set of events are generated using it as a model for physics.
2. The set of “model independent” observables (usually kinematic end points) which best shed light on sparticle masses *at that particular point* are cherry-picked from the larger list of available observables.
3. The expressions for the positions of these observables, which are typically only valid *in the vicinity of the special point*, are re-written, approximately, in terms of particle masses alone.
4. Using the approximations to the observables, defined in step 3, the mass spectra from the scanned points in $\mathcal{P}$ are each checked for consistency with edges generated in step 1. A favoured region in $\mathcal{P}$ is thus determined.

^cBut see discussion of [29, 30] later.

Unfortunately, the approximations made in step 3 mean that little can be deduced about the status of points at large distances from the special one. This is disappointing as we would like to be able to rule out “wrong” regions of parameter space in addition to being able to rule in “right” ones.

5.1.5 Summary

In summary, there is much encouraging work on supersymmetry in connection with ATLAS, but many questions remain unanswered. For sufficiently general questions (for example [27] discussed in Section 5.1.1) the scope for misinterpreting the results of experiment seems small. The more specific the questions become, however, the more one finds that the analyses needed to answer them become open to a variety of interpretations.

In part, this thesis attempts to address a few of these concerns. Although it is hoped that the analysis presented demonstrates *some* improvement over existing work, there is no doubt that the majority of the earlier criticisms will continue to apply. Which these criticisms are, and where improvements have been made, will become more clear as the methods are described.

5.2 Aim of the analysis

The principle aim of the following analysis is to evaluate prospects for “model independent” measurements of sparticle masses at ATLAS.

In order to make flesh out the above statement, it is necessary to first outline the general approach which is taken:

A “model independent” method for measuring a small set of sparticle masses is proposed. The method is based upon the work of others, [31–33] who have investigated

the measurement of squark, slepton and neutralino masses at ATLAS. These analyses were “model independent” in that they relied principally on kinematic end points from multi-stage sparticle decay chains. However, it is arguable that these analyses are guilty of assuming too much knowledge about the mechanics of the underlying process — information that the experimenter either would not know, or might have to work hard to find. The number of these assumptions is reduced, but not entirely eliminated. The performance of the new method is investigated at *two* test points from the space of $R_P C$ supersymmetric models. At each point, there is discussion of the accuracy with which the sparticle masses would be reconstructed.

The initial aims of the analysis may therefore be subdivided as follows:

- To highlight shortcomings of the investigations upon which this analysis is based.
- To suggest modifications which may be made to the original approach to counter the above shortcomings, make the method more general and so increase the “model independence”.
- To draw attention to new problems which the above changes introduce.
- To propose how the method might be extended so as to address the new problems.
- To examine the implications which the above changes have for the accuracy with which sparticle mass measurements may be performed for the two test models.

There is an additional aim which, although related to the above sequence, is not directly part of it. This is the aim of testing the “model independence” of the method by applying it at *two* test points. Although two models can not be said to be a representative sample of all possibilities, this sort of check can provide valuable insights into the implicit assumptions of the technique, and is rarely employed.

5.3 Summary of method

Section 5.2 has listed the aims of the analysis, and summarised the overall method in a single paragraph. In this section, the methods used to “measure sparticle masses”, and “evaluate the accuracy with which the measurements may be made” will be described in a little more detail. The methods may be thought of as having three stages of development: “original”, “basic” (my interpretation of the original analysis) and “extended”. In this section, the parts of the method which are common to all three stages will be outlined. Distinctions between them will only be developed in later sections.

The method used to measure sparticle masses is primarily the study of kinematic endpoints of invariant mass distributions resulting from decay chains of various forms. A general description of analysis of kinematic endpoints has already been given in Section 5.1.3.

The steps which are used in this thesis to simulate the measurement of sparticle masses at ATLAS, and which thereby evaluate the prospects for such measurements, may be broken down into the following stages:

Generation and simulation of events

Montecarlo event generators are used to simulate proton-proton collisions at LHC energies. These events are passed to the **ATLFAST** detector simulation, which parametrises the response of ATLAS, and produces a list of “observed” quantities, such as particle momenta and missing energy *etc.* Events are generated from all processes contributing to the production of supersymmetric particles. As well as containing signal events, these also contain the majority of the background. SM processes are also generated where they form a potential background to the searches.

Selection of events

Cuts are applied to the results of the simulation in order to generate invariant mass distributions, whose endpoints may be used to draw conclusions about the masses of the sparticles present in the decay chains.

Estimation of endpoint errors

Fits to the edges of the invariant mass distributions are performed, but more attention is paid to the statistical contribution to the resolution on the endpoint, than is paid to the value obtained for the endpoint itself. The actual positions of the endpoints are affected by many factors, including systematic effects originating within the detector, and systematic effects caused by supersymmetric backgrounds in the vicinity of each edge. The cost which would be required to undertake an investigation of the supersymmetric effects at every edge, would clearly outweigh benefits which would be gained from doing so at this stage.^d The dominant detector induced systematic (the error on the absolute jet energy scale) is, however, taken into account. As this investigation chooses to concentrate on the limits placed upon sparticle mass measurements by the scarcity of events, it will make the assumption that statistical effects (from edge fitting) and detector induced systematics will dominate the eventual errors on the edge positions. Whether this turns out to be the case will determine whether the final errors quoted for reconstructed sparticle masses will represent true errors, or just lower bounds for their magnitudes.

Modelling the *endpoint* measurement process itself

Given the statistical error associated with the measurement of each endpoint, and having made the assumption that non-detector systematic effects do not dominate,

^dOf course this statement would not be true if the events being examined were from *real* data; there it would be imperative to investigate *all* edge effects. But such investigations can only meaningfully take place once the idiosyncrasies of the detector have been determined, after it has been running for long enough for calibrations to take place.

it is possible to mimic the whole process described above stochastically. In this way, the “measurement” of n endpoints, may be “simulated” by randomly smearing n exact endpoint values according to their respective statistical and systematic errors. Any such “measurement”, generates a set of endpoints of the kind that a fully calibrated ATLAS might be expected to obtain during its lifetime. These “endpoint measurements” may represent best case scenarios, as the inclusion of supersymmetric systematic effects can only degrade performance.

Simulation of *mass* measurements

Once a set of “measured” endpoints has been obtained, the sparticle masses may be estimated. This is done by performing a chi-squared fit to the edge locations as a function of the sparticle masses. The advantage of the synthetic measurement process is that it may be repeated many times, each time with different random smearings of the endpoints. By repeatedly performing the endpoint measurements and mass reconstructions (effectively simulating an ensemble of ATLAS experiments) it is possible to observe the spread of reconstructed sparticle masses.

Interpretation of results

The spread of reconstructed masses is a measure of the resolution of the method when used under favourable conditions at ATLAS. The results show the limit which would be obtained if physics related systematic errors were to become negligible in comparison with detector effects and statistical errors.

5.4 Test models: S5 and O1.

One of the main criticisms of the majority of “model independent” analyses of the past has also been that, despite their model independent claims, they have rarely been tried at more than one point in model space! Indeed, in line with the main criticism of the

“show you might be right” approach, the analyses have usually only been tested at the very point in parameter space for which they were developed. It was thus thought interesting to test the claim of model independence *directly* by actually trying to use a single sufficiently general analysis, on physics governed by *two different models*.^e The two models chosen for study were LHC mSUGRA Point 5 (mirroring [31–33]) and a point in the space of OSMs. OSMs are described in Section 2.9.4, while mSUGRA models may be found in Section 2.8.

mSUGRA Point 5 (S5) is characterised by the following values of the model parameters:

$$m_0 = 100 \text{ GeV}, m_{\frac{1}{2}} = 300 \text{ GeV}, A_0 = 300 \text{ GeV}, \tan \beta = 2.1 \text{ and } \mu > 0. \quad (5.3)$$

The significance of the above values is that they make S5 interesting cosmologically. The suggestion is that the neutralinos might be able to form part of the dark matter of the universe, provided that their relic density is “consistent with theoretical prejudices and astrophysical observations” [33]. One of the constraints that this puts on S5 is that the sleptons should be relatively light, in order that neutralino annihilation (through t -channel sleptons) can take place at the desired rate.

The particular optimised model used here (O1) is defined by

$$M_{\frac{3}{2}} = 250 \text{ GeV}, \tan \beta = 10 \text{ and } \theta = \frac{\pi}{4}, \quad (5.4)$$

and was first defined in [34]. The above values were chosen largely to make the production cross sections for supersymmetric particles approximately similar across the two

^eFrom a historical perspective, the desire to test two models with a single analysis came first. It was the unfortunate failure of a naive application of the analysis, when used at the second point, which lead to the decision to investigate further.

M_g	m_{u_L}	$m_{u_R^c}$	m_{d_L}	$m_{d_R^c}$	m_{b_1}	m_{b_2}	m_{t_1}	m_{t_2}	
747	660	632	664	630	608	636	494	670	
733	654	631	657	628	600	629	460	671	
m_{ν_e}	m_{e_L}	$m_{e_R^c}$	m_{ν_τ}	m_{τ_1}	m_{τ_2}				
273	284	217	271	209	285				
230	239	157	230	157	239				
$m_{\chi_1^0}$	$m_{\chi_2^0}$	$m_{\chi_3^0}$	$m_{\chi_4^0}$	$m_{\chi_1^+}$	$m_{\chi_2^+}$	m_{h^0}	m_{H^0}	m_{A^0}	m_{H^+}
125	234	371	398	233	398	114	450	449	456
122	233	499	523	232	520	94	612	607	612

Table 5.1: *Comparison of sparticle and higgs spectrum of O1 and S5 (in bold type). These spectra were calculated using ISASUGRA and all masses are quoted in GeV. The masses of the second family of sparticles are approximately equal to those of the first.*

models. As sparticle production is dominated by squarks and gluinos in both models, the effect of the tuning is to make the squark mass scale similar in each. The effect of the difference in the squark and slepton modular weights at O1, however, is to drive its sleptons heavier than those of S5. Table 5.1 provides a comparison of the sparticle and higgs masses for the two models.

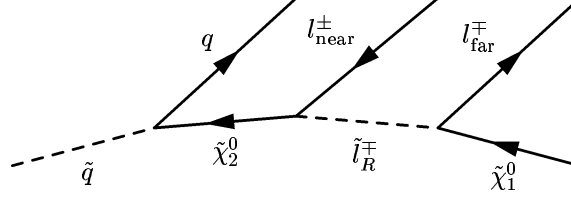


Figure 5.3: “Sequential” squark decay.

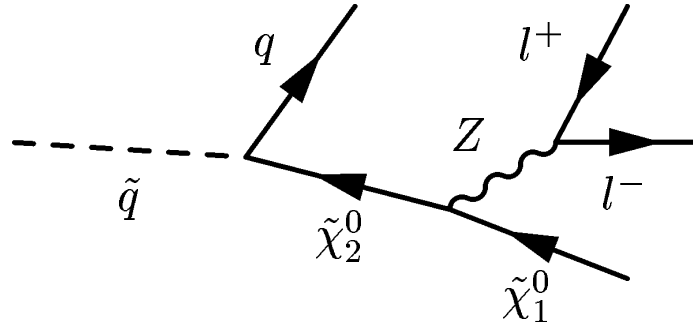
5.5 The decay chains

There are three principal types of decay chain or production process used in this analysis:

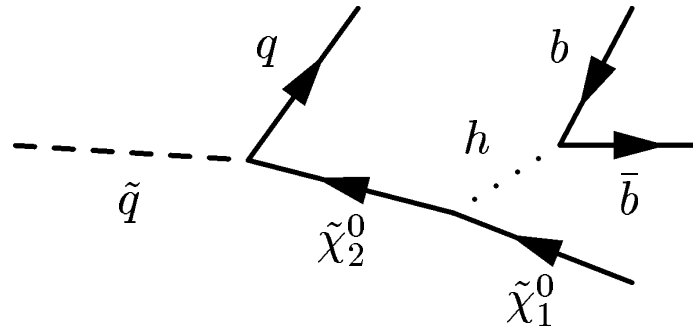
- “Sequential” squark decay: $\tilde{q}_L \rightarrow \tilde{\chi}_2^0 q \rightarrow \tilde{l}_R^\pm l_{\text{near}}^\pm q \rightarrow \tilde{\chi}_1^0 l_{\text{far}}^\pm l_{\text{near}}^\pm q$ as shown in Figure 5.3. Note that the names l_{near} and l_{far} are given to the leptons to distinguish them. The lepton radiated by the $\tilde{\chi}_2^0$ is called l_{near} (in the standard diagram it is the one “nearer to” the quark) while the lepton radiated by the slepton is called l_{far} (in the standard diagram it is the one “further from” the quark).
- “Branched” squark decays: $\tilde{q}_L \rightarrow \tilde{\chi}_2^0 q \rightarrow \tilde{\chi}_1^0 X q \rightarrow \tilde{\chi}_1^0 l^\pm l^\mp q$ as shown in Figure 5.4. Here X stands for the SM particle radiated by the decaying $\tilde{\chi}_2^0$, which in this thesis is taken to be a higgs or a Z -boson.
- Dislepton production processes: $pp \rightarrow Y \tilde{l}_R^\pm \tilde{l}_R^\mp \rightarrow Y l^\pm l^\mp \tilde{\chi}_1^0 \tilde{\chi}_1^0$ as shown in Figure 5.5. Here Y denotes whatever initial- or final-state radiation the $\tilde{l}_R^\pm \tilde{l}_R^\mp$ recoils against.

The sequential squark decay and the branched decay were both used in the original analyses [31–33], while the use of the dislepton production process is novel.

Notice that, by presupposing the existence of the above decay chains, implicit assumptions are already beginning to creep into the analysis. For example, for the “sequential” decay mode to be allowed, the $\tilde{\chi}_2^0$ must be heavier than the $\tilde{l}_R$. While this relationship is satisfied in the two models being compared, it is easy to find others where



(a) Z mode



(b) higgs mode

Figure 5.4: “Branched” squark decays through on-shell Z and higgs bosons.

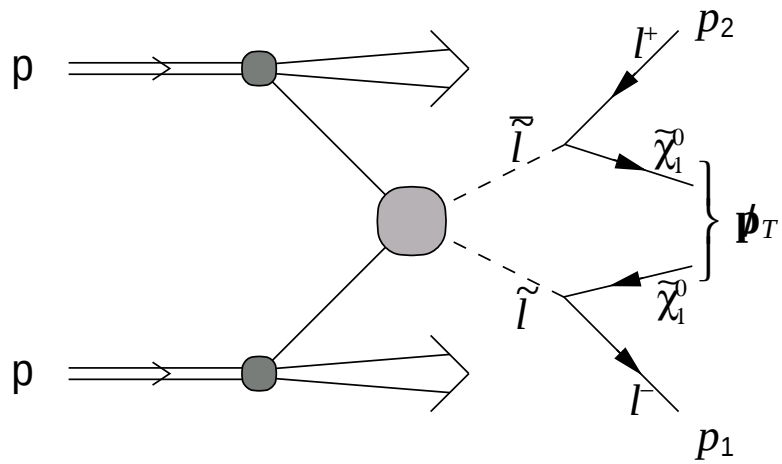


Figure 5.5: Dislepton production process.

this is not the case, such as mSUGRA Point 3. This point has $m_0 = 200$ GeV, $m_{1/2} = 100$ GeV, $A_0 = 0$, $\tan\beta = 2$, $\mu < 0$, and so has particularly light neutralinos. Even though S3 lacks the “sequential” decay chain, it is quite capable of producing the same final state via a virtual slepton, $\tilde{q}_L \rightarrow \tilde{\chi}_2^0 q \rightarrow \tilde{\chi}_1^0 l^\mp l^\pm q$. It is not at all obvious that one may reliably discriminate between these two situations without recourse to model *dependent* methods, and so the consequences for applying the analysis here to S3 are less than clear. Nevertheless, some assumptions must be made or there can be no progress. There remains the hope that when the number of observables becomes large enough, application of methods at inappropriate points should eventually result in inconsistencies, or poor chi-squareds, allowing discrimination between good and bad models.

Another assumption which is implicit in this analysis is that supersymmetry is $R_P C$. The huge differences between the phenomena associated with $R_P C$ and $R_P V$ models makes this assumption very defensible, at least when the LSP is only able to interact via the weak force. Although uncontentious, this assumption has large consequences for the analysis technique; for as discussed in Section 2.5, two LSPs (of unknown mass) can be expected to go missing from each supersymmetric event. The momenta of the LSPs cannot be measured on an event by event basis, and their masses are unknown *a priori*. This means that it is not possible to reconstruct entire decay chains “from the bottom up”. Thus, endpoints in kinematic variables constructed from the visible products of the decay chains (or *invisible* in the case of missing transverse momentum) alone may be used. In the calculations which follow, the masses of the SM leptons and quarks are neglected, except where stated otherwise.

5.6 Edges, endpoints and thresholds

The edges which were used to measure sparticle masses in [31–33] are introduced in Sections 5.6.1 to 5.6.4. In Sections 5.6.5 and 5.6.6, the edges which are new to this analysis are introduced.

5.6.1 The l^+l^- edge

This picks out, from the “sequential” decays, the position of the very sharp edge in the dilepton invariant mass spectrum caused by $\tilde{\chi}_2^0 \rightarrow l\tilde{l}$ followed by $\tilde{l} \rightarrow l\tilde{\chi}_1^0$. It is easy to see (Appendix E) that the maximum is attained when the lepton momenta are back to back in the rest frame of the slepton, and is given by the value in Table 5.2. When considering events below the maximum, the only free kinematic variable of interest in this frame is the *uniformly distributed* spherical polar parameter $\cos\theta$, defining the angle θ between the first and second lepton. Since the square of the invariant mass of the pair is given by $m_{ll}^2 = 2p_1p_2(1 - \cos\theta)$, the distribution of m_{ll}^2 is uniform. As a consequence, m_{ll} has a linear distribution over $[0, m_{ll}^{\max}]$ which goes to zero at the origin. This result is not subject to spin corrections since the slepton is a scalar particle. The sudden cutoff in the m_{ll} distribution at its upper end, makes $m_{ll}^{\max}$ the endpoint whose position is obtainable with the best precision.

5.6.2 The l^+l^-q edge

This is the upper edge in the llq invariant mass spectrum of “sequential” decays. In general, the proportion of phase space close to an endpoint decreases as the number of particles under consideration increases. For this reason the l^+l^-q edge is not a step (as the l^+l^- edge is) but is typically an edge with a linear dropoff to zero, at least in the

vicinity of the endpoint. In [31–33], the theoretical position of $(m_{llq}^{\max})^2$ was evaluated as a function of the sparticle masses, and determined to be

$$(m_{\tilde{q}}^2 - m_{\tilde{\chi}_2^0}^2)(m_{\tilde{\chi}_2^0}^2 - m_{\tilde{\chi}_1^0}^2)/m_{\tilde{\chi}_2^0}^2. \quad (5.5)$$

However, while this expression gives the correct location of the endpoint in the O1 and S5 models, the formula is not general. Bear in mind that the eventual aim is to measure sparticle masses by finding those sets which are consistent with the observed endpoint positions. It is vital, then, to ensure that no combination of masses which is consistent with a given endpoint is missed, otherwise allowed regions of parameter space might be falsely excluded. Having presupposed the existence of the “sequential” chain, we are interested in knowing the position of the endpoint subject only to the hierarchy of masses

$$0 < m_{\tilde{\chi}_1^0} < m_{\tilde{l}_R} < m_{\tilde{\chi}_2^0} < m_{\tilde{q}}. \quad (5.6)$$

Careful calculation (see Appendix E) shows that, where this hierarchy exists, there are actually four possible locations for the endpoint:

$$(m_{llq}^{\max})^2 = \begin{cases} (m_{\tilde{q}}^2 - m_{\tilde{\chi}_2^0}^2)(m_{\tilde{\chi}_2^0}^2 - m_{\tilde{\chi}_1^0}^2)/m_{\tilde{\chi}_2^0}^2 & \text{iff } m_{\tilde{\chi}_2^0}^2 < m_{\tilde{\chi}_1^0} m_{\tilde{q}}, \\ (m_{\tilde{q}}^2 - m_{\tilde{l}}^2)(m_{\tilde{l}}^2 - m_{\tilde{\chi}_1^0}^2)/m_{\tilde{l}}^2 & \text{iff } m_{\tilde{\chi}_1^0} m_{\tilde{q}} < m_{\tilde{l}}^2, \\ (m_{\tilde{q}}^2 m_{\tilde{l}}^2 - m_{\tilde{\chi}_2^0}^2 m_{\tilde{\chi}_1^0}^2)(m_{\tilde{\chi}_2^0}^2 - m_{\tilde{l}}^2)/(m_{\tilde{\chi}_2^0}^2 m_{\tilde{l}}^2) & \text{iff } m_{\tilde{l}}^2 m_{\tilde{q}} < m_{\tilde{\chi}_1^0} m_{\tilde{\chi}_2^0}^2, \\ (m_{\tilde{q}} - m_{\tilde{\chi}_1^0})^2 & \text{otherwise.} \end{cases} \quad (5.7)$$

Note that the first of these cases corresponds to expression (5.5). Note also that all four conditions in (5.7) are mutually exclusive when taken together with the hierarchy of (5.6). The conditions may also be summarised in the equivalent way shown in Table 5.2.

5.6.3 The l^+l^-q threshold

This end point could be said to be “artificial” in origin, as it arises as the result of an arbitrary constraint upon the lepton momenta. The l^+l^-q threshold is the non-zero minimum in the “sequential” llq invariant mass spectrum formed from the subset of events in which the invariant mass of the two leptons is greater than a particular fixed value. The value used here (following [31–33]) is $m_{ll}^{\text{max}}/\sqrt{2}$, which requires the measurement of the l^+l^- edge to have been performed first. There is nothing particularly special about this choice of this cut, other than that it slightly simplifies the maths, by having the effect of requiring that the angle between the two leptons (in the centre of mass frame of the slepton) is greater than $\pi/2$.

The location of the l^+l^-q threshold may be found in Table 5.2 and is derived in Appendix E. There are no special cases to consider; the formula given is valid for all combinations of masses satisfying (5.6). Although seemingly arbitrary, the cut is useful as it provides further information about mass relationships, having a little more sensitivity to *sums* of sparticle masses than the first two edges.

5.6.4 The hq and Zq edges (Xq edges)

In general, the Xq edge is the upper edge of the distribution of the invariant mass of three visible particles in the “branched” decays of Figure 5.4. The location of the endpoint of this edge (derived in Appendix E) is again determined by two-body kinematics, and is shown in Table 5.2. Depending on the higgs mass, the mass difference between the $\tilde{\chi}_2^0$ and the $\tilde{\chi}_1^0$, and the relative gaugino/higgsino composition of the neutralinos, one of the two “branched” decay chains will in general be strongly suppressed with respect to the other, so typically only one edge will be visible. At S5, where the neutralinos are predominantly gaugino, the higgs mode dominates. In contrast, the Z mode is dominant

at O1, as the $m_{\tilde{\chi}_2^0} - m_{\tilde{\chi}_1^0}$ mass difference (109 GeV) is less than the mass of the light higgs (114 GeV). In the context of a particular model, then, the Xq edge is either referred to as the hq edge or the Zq edge.

There are two places in this analysis where the techniques applied at O1 differ from those applied at S5. These are points to watch for, as they bring into question the degree to which the analysis is guilty of falling foul of its own criticisms of “model independent” methods.^f The first of these places is the treatment of the Xq edges:

At S5, where the higgs mode dominates, no attempt will be made to search for (and quantify the lack of significance of) a resonant Z -peak in the dilepton spectrum of l^+l^-q events. There will consequently be no attempt to plot the m_{llq} invariant mass distribution for events whose leptons lie in this peak. When the analysis is performed at this point, it will be implicitly assumed that the hq edge rather than the Zq edge should be sought.

Conversely, at O1, the $m_{b\bar{b}q}$ spectrum will be ignored in favour of l^+l^-q events.

Here it is argued that the difference between the way the hq and Zq edges are treated at the two test points is not a cause for concern. Firstly, the two analyses have entirely different cuts and examine entirely different states. There is thus no reason why both analyses could not be performed at both points, or indeed at general points. The only complication would involve deciding how to treat the null results at points where one or other of the modes was suppressed. Almost by definition, these null results would emerge from distributions containing only background events, and so the details of the dismissal of each point would depend very much on the particular model. One way of avoiding the problem altogether would be to “symmetrize” the analysis, requiring that: “Given real data, preliminary investigations of both Xq edges should be performed, but only

^fSee comments near the beginning of Section 5.4.

the edge derived from the larger resonant peak should be taken to completion”. This way, the difference between the two methods would be the result of a “spontaneously” broken symmetry. This approach would encapsulate the idea that at least one of the two modes is to be expected at most points. Finally it should be noted that the “kind of unequal treatment” which arises in the case of the Xq edges (where one has to choose which of two largely exclusive physics processes to study) is of a very different *type* to the “kind of unequal treatment” which arises as the result of interpreting a particular observable (such as an edge position) with an implicit bias (e.g. by assuming a particular relationship among masses). While discussion of both types of issue have their place, it is the feeling of the author that the former is less important than the latter, at least if one is worried about falsely excluding parts of parameter space.

5.6.5 Near, far, high and low $l^\pm q$ edges

In “sequential” decays there are three observable outgoing particles; one quark and two leptons, l_{near} and l_{far} (see Figure 5.3 and the definitions of l_{near} and l_{far} in Section 5.5). There are only four different ways of grouping these particles together, and so only four different invariant masses may be formed from them: m_{ll} , m_{llq} , $m_{l_{\text{near}}q}$ and $m_{l_{\text{far}}q}$. The first two of these, m_{ll} and m_{llq} , may be formed from the observed momenta without knowledge of which lepton was l_{near} and which was l_{far} ; only the *total* four-momentum of the leptons is required. Edges in these invariant mass distributions have already been described. If it were possible to tag the near- and far-leptons separately, the third and fourth invariant mass combinations could also be formed unambiguously, allowing the positions of two more endpoints, $m_{l_{\text{near}}q}^{\text{max}}$ and $m_{l_{\text{far}}q}^{\text{max}}$, to play a part in the final fit.

Unfortunately, such tagging is impossible. If further information is to be gathered from these decays *in a model independent way*, it is then necessary to look for functions of $m_{l_{\text{near}}q}$ and $m_{l_{\text{far}}q}$ which *are* observable. On an event-by-event basis, then, it is necessary to define

$$m_{lq(\text{high})} = \max(m_{l^+q}, m_{l^-q}) \equiv \max(m_{l_{\text{near}}q}, m_{l_{\text{far}}q}) \quad (5.8)$$

and

$$m_{lq(\text{low})} = \min(m_{l^+q}, m_{l^-q}) \equiv \min(m_{l_{\text{near}}q}, m_{l_{\text{far}}q}). \quad (5.9)$$

Although the new masses are clearly related to $m_{l_{\text{near}}q}$ and $m_{l_{\text{far}}q}$, it is not necessary to know which lepton is which in order to construct them. Where $m_{l_{\text{near}}q}$ and $m_{l_{\text{far}}q}$ failed, $m_{lq(\text{high})}$ and $m_{lq(\text{low})}$ constitute observable invariant masses whose distributions and endpoints *can* be studied. The author's predictions for the positions of the corresponding observable endpoints, $m_{lq(\text{high})}^{\text{max}}$ and $m_{lq(\text{low})}^{\text{max}}$, are listed in Table 5.2, having been derived in Appendix E. For completeness, the same table includes the positions of the unobservable near- and far-edges. Note that although the position of the high-edge is just the higher of $m_{l_{\text{near}}q}^{\text{max}}$ and $m_{l_{\text{far}}q}^{\text{max}}$, the position of the low-edge has a more interesting form. This asymmetry arises because it is not always kinematically possible for the invariant mass of the lepton-quark pair coming from the lower of the near- and far-distributions to approach $\min[m_{l_{\text{near}}q}^{\text{max}}, m_{l_{\text{far}}q}^{\text{max}}]$ arbitrarily closely, while simultaneously requiring that this invariant mass is less than that of the other lepton-quark pair.

Related edge	Kinematic endpoint
l^+l^- edge	$(m_{ll}^{\max})^2 = (\tilde{\xi} - \tilde{l})(\tilde{l} - \tilde{\chi})/\tilde{l}$
l^+l^-q edge	$(m_{llq}^{\max})^2 = \begin{cases} \max \left[\frac{(\tilde{q}-\tilde{\xi})(\tilde{\xi}-\tilde{\chi})}{\tilde{\xi}}, \frac{(\tilde{q}-\tilde{l})(\tilde{l}-\tilde{\chi})}{\tilde{l}}, \frac{(\tilde{q}-\tilde{\xi})(\tilde{\xi}-\tilde{l})}{\tilde{\xi}\tilde{l}} \right] \\ \text{except for the special case in which } \tilde{l}^2 < \tilde{q}\tilde{\chi} < \tilde{\xi}^2 \\ \text{and } \tilde{\xi}^2\tilde{\chi} < \tilde{q}\tilde{l}^2 \text{ where one must use } (m_{\tilde{q}} - m_{\tilde{\chi}_1^0})^2. \end{cases}$
Xq edge	$(m_{Xq}^{\max})^2 = X + \frac{(\tilde{q}-\tilde{\xi})}{2\tilde{\xi}} \left[\tilde{\xi} + X - \tilde{\chi} + \sqrt{(\tilde{\xi} - X - \tilde{\chi})^2 - 4X\tilde{\chi}} \right]$
l^+l^-q threshold	$(m_{llq}^{\min})^2 = \begin{cases} [& 2\tilde{l}(\tilde{q}-\tilde{\xi})(\tilde{\xi}-\tilde{\chi}) + (\tilde{q}+\tilde{\xi})(\tilde{\xi}-\tilde{l})(\tilde{l}-\tilde{\chi}) \\ & -(\tilde{q}-\tilde{\xi})\sqrt{(\tilde{\xi}+\tilde{l})^2(\tilde{l}+\tilde{\chi})^2 - 16\tilde{\xi}\tilde{l}^2\tilde{\chi}}]/(4\tilde{l}\tilde{\xi}) \end{cases}$
$l_{\text{near}}^{\pm}q$ edge	$(m_{l_{\text{near}}q}^{\max})^2 = (\tilde{q} - \tilde{\xi})(\tilde{\xi} - \tilde{l})/\tilde{\xi} \quad (\text{unobservable})$
$l_{\text{far}}^{\pm}q$ edge	$(m_{l_{\text{far}}q}^{\max})^2 = (\tilde{q} - \tilde{\xi})(\tilde{l} - \tilde{\chi})/\tilde{l} \quad (\text{unobservable})$
$l^{\pm}q$ high-edge	$(m_{lq(\text{high})}^{\max})^2 = \max \left[(m_{l_{\text{near}}q}^{\max})^2, (m_{l_{\text{far}}q}^{\max})^2 \right]$
$l^{\pm}q$ low-edge	$(m_{lq(\text{low})}^{\max})^2 = \min \left[(m_{l_{\text{near}}q}^{\max})^2, (\tilde{q} - \tilde{\xi})(\tilde{l} - \tilde{\chi})/(2\tilde{l} - \tilde{\chi}) \right]$
$\tilde{l}\tilde{l}$ edge	$\Delta M^{\max} = (m_{\tilde{l}}^2 - m_{\tilde{\chi}_1^0}^2)/(2m_{\tilde{l}})$

Table 5.2: *The absolute kinematic endpoints of invariant mass quantities formed from decay chains of the types mentioned in the text for known particle masses. The following shorthand notation has been used: $\tilde{\chi} = m_{\tilde{\chi}_1^0}^2$, $\tilde{l} = m_{\tilde{l}_R}^2$, $\tilde{\xi} = m_{\tilde{\chi}_2^0}^2$, $\tilde{q} = m_{\tilde{q}}^2$ and X is m_h^2 or m_Z^2 depending on which particle participates in the “branched” decay. The positions of endpoints are derived in Appendix E.*

5.6.6 Dislepton production, $M_{T2}^{\max}(\chi)$, ΔM and the $\tilde{l}\tilde{l}$ edge

One of the novel parts of this analysis concerns the introduction of ΔM , a variable designed to better constrain the slepton-LSP mass difference.

It is easiest to describe ΔM in terms of a related variable, $M_{T2}(\chi)$ proposed by the author (in collaboration) in [35].

$M_{T2}(\chi)$ may be used to study events where:

- the event contains two identical two-body decays (i.e. $X \rightarrow Y a_1 a_2$ followed by $a_1 \rightarrow b_1 c_1$ and $a_2 \rightarrow b_2 c_2$),
- $m_A = m_{a_1} \equiv m_{a_2}$ is unknown,
- $m_C = m_{c_1} \equiv m_{c_2}$ is unknown,
- $m_B = m_{b_1} \equiv m_{b_2}$ is known,
- particles of type c are undetectable,
- the longitudinal momentum of the incoming particles is also unknown, and
- Y does not contain any unobservable particles such as neutrinos.

In such cases, the variable provides a kinematic constraint on the masses of particles of type a and type c . In this analysis, $M_{T2}(\chi)$ will be applied to dislepton events of the form $pp \rightarrow Y \tilde{l}_R^\pm \tilde{l}_R^\mp \rightarrow Y l^\pm l^\mp \tilde{\chi}_1^0 \tilde{\chi}_1^0$ (Figure 5.5) using the assignment $a = \tilde{l}$, $b = e/\mu$ and $c = \tilde{\chi}_1^0$. These events are used to constrain $m_{\tilde{l}}$ and $m_{\tilde{\chi}_1^0}$ in the following way:

Suppose, for the moment, that the mass of the neutralino were known, and denoted by $m_{\tilde{\chi}}$ (the numbers in $\tilde{\chi}_1^0$ have been dropped to avoid confusion with subscripts indicating which decay the neutralino came from). Then, if it were possible to observe the momenta

of each neutralino, and correctly identify its associated lepton, it would be possible to compute the invariant mass of either slepton using the exact relation:

$$m_i^2 = m_l^2 + m_{\tilde{\chi}}^2 + 2(E_T^l E_T^{\tilde{\chi}} \cosh(\Delta y) - \mathbf{p}_T^l \cdot \mathbf{p}_T^{\tilde{\chi}}) \quad (5.10)$$

where $E_T^l = \sqrt{\mathbf{p}_T^{l^2} + m_l^2}$, $E_T^{\tilde{\chi}} = \sqrt{\mathbf{p}_T^{\tilde{\chi}^2} + m_{\tilde{\chi}}^2}$ and Δy is the difference between the rapidities (see Appendix D) of the lepton and the neutralino. This definition includes the lepton masses for completeness, although these are neglected in all computations. Of course, observing the neutralinos is impossible at ATLAS, so the difference in rapidities, Δy , will certainly not be known. Nevertheless, using the fact that $\cosh \Delta y \geq 1$, it is possible to make the reduced claim:

$$m_i^2 \geq m_T^2(\mathbf{p}_T^l, \mathbf{p}_T^{\tilde{\chi}}, m_{\tilde{\chi}}) \equiv m_l^2 + m_{\tilde{\chi}}^2 + 2(E_T^l E_T^{\tilde{\chi}} - \mathbf{p}_T^l \cdot \mathbf{p}_T^{\tilde{\chi}}), \quad (5.11)$$

or even the stronger claim:

$$m_i^2 \geq \max\{m_T^2(\mathbf{p}_T^{l_1}, \mathbf{p}_T^{\tilde{\chi}_1}, m_{\tilde{\chi}}), m_T^2(\mathbf{p}_T^{l_2}, \mathbf{p}_T^{\tilde{\chi}_2}, m_{\tilde{\chi}})\} \quad (5.12)$$

using both of the decays. But neither $m_T^2(\mathbf{p}_T^l, \mathbf{p}_T^{\tilde{\chi}}, m_{\tilde{\chi}})$ nor the right hand side of (5.12) is an observable quantity, as each relies upon knowledge of the transverse components of the momenta of each neutralino. Under the prevailing assumptions, the only thing known about the neutralinos is the *sum* of their transverse momenta, given by the missing momentum of the event:

$$\not{p}_T = \mathbf{p}_T^{\tilde{\chi}_1} + \mathbf{p}_T^{\tilde{\chi}_2}. \quad (5.13)$$

Although the form of the splitting (5.13) is not known, it is possible to turn the right hand side of (5.12) into a quantity $M_{T_2}^2(m_{\tilde{\chi}})$ which *is* observable, by trying all possible

splittings. This motivates the definition:

$$M_{T2}^2(m_{\tilde{\chi}}) \equiv \min_{\not{\mathbf{p}}_1 + \not{\mathbf{p}}_2 = \not{\mathbf{p}}_T} \left[\max \{m_T^2(\mathbf{p}_T^{l_1}, \not{\mathbf{p}}_1, m_{\tilde{\chi}}), m_T^2(\mathbf{p}_T^{l_2}, \not{\mathbf{p}}_2, m_{\tilde{\chi}})\} \right], \quad (5.14)$$

which must therefore satisfy

$$m_{\tilde{l}}^2 \geq M_{T2}^2(m_{\tilde{\chi}}). \quad (5.15)$$

In [35], the authors note that it *is* kinematically possible for equality to be reached in (5.15), and that this occurs when $\Delta y_1 = \Delta y_2 = 0$ and $\mathbf{v}_T^{l_1} - \mathbf{v}_T^{\tilde{\chi}_1}$ is parallel (note: not anti-parallel) to $\mathbf{v}_T^{l_2} - \mathbf{v}_T^{\tilde{\chi}_2}$, where $\mathbf{v}_T = \mathbf{p}_T/E_T$. Thus, by construction, $M_{T2}(m_{\tilde{\chi}})$ has the property that, for a large enough number of signal events in a perfect detector,

$$M_{T2}^{\max}(m_{\tilde{\chi}_1^0}) \equiv \max_{\text{events}} [M_{T2}(m_{\tilde{\chi}_1^0})] = m_{\tilde{l}}, \quad (5.16)$$

which is equivalent to saying that the endpoint of the $M_{T2}(m_{\tilde{\chi}})$ distribution occurs at the slepton mass.

Note that up to this point it has been necessary to assume knowledge of the neutralino mass, $m_{\tilde{\chi}} \equiv m_{\tilde{\chi}_1^0}$, which is a little unfair, given that it is one of the quantities which the analysis seeks to determine. Formally, then, the M_{T2} distribution is not a single distribution, but a family of distributions: one for each possible value of the neutralino mass. For example; suppose it is proposed that the neutralino mass is equal to “ χ ” (not necessarily close to the correct value). This would motivate plotting the distribution of $M_{T2}(\chi)$, generated from dilepton events under the hypothesis that the neutralino mass may be taken to be χ rather than $m_{\tilde{\chi}}$ in definition (5.14). Having plotted this distribution, the location $M_{T2}^{\max}(\chi)$ of its endpoint could be ascertained, and the interpretation of $M_{T2}^{\max}(\chi)$ would be “the slepton mass favoured by the data given the assumed neutralino mass”.

Thus the overall result obtained from a full set of $M_{T2}(\chi)$ distributions, is a function

$M_{T2}^{\max}(\chi)$, describing the most favoured slepton mass as a function of supposed neutralino mass χ .

In Section E.7 of the Appendix it is shown that the function $M_{T2}^{\max}(\chi)$ is in fact dependent on only one unknown parameter $p > 0$; namely the modulus of the three-momentum of the lepton (or the neutralino) in the rest frame of the slepton. In general p is given by

$$p^2 = \frac{m_l^4 + m_{\tilde{\chi}_1^0}^4 + m_l^4 - 2(m_l^2 m_{\tilde{\chi}_1^0}^2 + m_l^2 m_l^2 + m_{\tilde{\chi}_1^0}^2 m_l^2)}{4m_l^2}. \quad (5.17)$$

In Section E.7 the general form of the dependence of $M_{T2}^{\max}(\chi)$ upon p is shown to be:

$$M_{T2}^{\max}(\chi) = \sqrt{m_l^2 + p^2} + \sqrt{\chi^2 + p^2}. \quad (5.18)$$

Both (5.17) and (5.18) assume that the effects of slepton boosts may be neglected. Reasons supporting this approximation are given in the Appendix. Significantly, (5.18) may be inverted for arbitrary χ yielding:

$$p^2 = \frac{\tilde{L}^4 + \chi^4 + m_l^4 - 2(\tilde{L}^2 \chi^2 + \tilde{L}^2 m_l^2 + \chi^2 m_l^2)}{4\tilde{L}^2} = \text{“independent of } \chi\text{”}, \quad (5.19)$$

where $M_{T2}^{\max}(\chi)$ has been abbreviated to $\tilde{L}$. The reader is encouraged to compare (5.19) with (5.17). The result (5.19) is important for two reasons. Firstly it tells us that we no longer need to worry about choosing a particular value of χ in the region of $m_{\tilde{\chi}_1^0}$, and secondly it shows us that p from (5.17) is the quantity which $M_{T2}^{\max}(\chi)$ is *really* constraining. In fact, the most elegant way of extracting p from $M_{T2}(\chi)$ measurements actually *exploits* the freedom to choose an arbitrary value of χ . Surprisingly, the best⁸ choice for χ turns out to be neither 0 nor $m_{\tilde{\chi}_1^0}$ (which one might have expected from the original definition of χ) but is in fact m_l . Using this value one may define the quantity

⁸In this context, “best” means “able to give the final answer a physical interpretation”.

$\Delta M \geq 0$ by

$$(\Delta M)^2 \equiv \frac{1}{4} (M_{T2}(m_l))^2 - m_l^2 \quad (5.20)$$

and thereby obtain from (5.18) the significant result that:

$$\Delta M^{\max} = p = p(m_{\tilde{l}}, m_{\tilde{\chi}_1^0}, m_l). \quad (5.21)$$

As a consequence of this important result, it is distributions of ΔM which are chosen to be plotted to extract model independent information from dilepton events. The result (5.21) shows us that the location of the ΔM edge (i.e. $\Delta M^{\max}$) is a direct measurement of p given by (5.17).

Thus far, the only approximation which has been made was the presumption that transverse slepton boosts have a negligible effect on the prediction for $\Delta M^{\max}$ given in (5.21). Within this thesis, however, there is additionally a motivation for neglecting the lepton masses which will now be described.

In order to maximise the size of the signal, the analysis presented in this thesis will seek both e^+e^- and $\mu^+\mu^-$ events. The presence of the lepton mass in (5.21) (or equivalently in the definition of p) suggests that one ought to analyse events with different leptons differently. Note, however, that in cases where the $\Delta M^{\max}$ is observed to lie at a mass much in excess of m_l , one may immediately deduce that $m_l \ll p$. The assumption $m_l \ll p$, together with definition (5.17), may be used to show that $m_l \ll m_{\tilde{l}} - m_{\tilde{\chi}_1^0}$, which itself permits p to be expanded in terms of $\epsilon = m_l/(m_{\tilde{l}} - m_{\tilde{\chi}_1^0})$ as follows:

$$\Delta M^{\max} = p = p_{\text{approx}} \left\{ 1 - \frac{m_{\tilde{l}}^2 + m_{\tilde{\chi}_1^0}^2}{(m_{\tilde{l}} + m_{\tilde{\chi}_1^0})^2} \epsilon^2 - \frac{2m_{\tilde{l}}^2 m_{\tilde{\chi}_1^0}^2}{(m_{\tilde{l}} + m_{\tilde{\chi}_1^0})^4} \epsilon^4 + O(\epsilon^6) \right\} \quad (5.22)$$

where

$$p_{\text{approx}} \equiv \frac{m_{\tilde{l}}^2 - m_{\tilde{\chi}_1^0}^2}{2m_{\tilde{l}}}. \quad (5.23)$$

This shows that when $\Delta M^{\max} \gg m_l$, then one may take $\Delta M^{\max}$ to be really just a measurement of p_{approx} given by (5.23).

The ΔM distributions arising from two test models examined in this thesis do in fact happen to satisfy $\Delta M^{\max} \gg \max[m_e, m_\mu]$ at the $100 \text{ GeV} \gg 100 \text{ MeV}$ level, and so within this thesis $\Delta M^{\max}$ is taken to be equivalent to p_{approx} as defined in (5.23). This means that within this part of the analysis it is safe to take the leptons to be massless and to analyse both the e^+e^- and the $\mu^+\mu^-$ events together.

To summarise the analysis of the dilepton edge:

- The distributions which will be plotted and whose edges will be identified, will be distributions of ΔM .
- The observation that $\Delta M^{\max} \gg m_\mu$ for the models under study permits us to: neglect lepton masses, analyse both e^+e^- and $\mu^+\mu^-$ events together, and interpret $\Delta M^{\max}$ as a measurement of $p_{\text{approx}} = (m_{\tilde{l}}^2 - m_{\tilde{\chi}_1^0}^2)/(2m_{\tilde{l}})$.
- Were the analysis to be performed at another point where $\Delta M^{\max}$ was much less than (or order of) m_l , it would no longer be possible to treat $\Delta M^{\max}$ as a measurement of p_{approx} . $\Delta M^{\max}$ would instead constrain p . Because p depends on m_l , one would have to consider e^+e^- events separately from $\mu^+\mu^-$ events.

The ΔM edge can thus be included in the analysis to provide an additional constraint on the slepton and neutralino masses. The name ‘ ΔM ’ was originally chosen because, to first order in $(m_{\tilde{l}} - m_{\tilde{\chi}_1^0})/m_{\tilde{l}}$, the value of $\Delta M^{\max} \approx p_{\text{approx}}$ is given by the slepton/neutralino mass difference. Nowhere does the analysis expect or depend on this approximation, though.

In practice the ΔM edge can be distorted by the finite resolution of the detector or missing energy from soft underlying events. Additionally, standard model (SM) back-

grounds provide constraints on minimum detectable slepton/neutralino mass differences (see Section 5.9.2). However the most important factor to control is the ability to correctly identify which production processes are contributing to an observed ΔM edge – particularly since it is possible to have multiple edges in the *non* standard model contributions to ΔM distributions. At S5, for example, where both the right- and left-sleptons are lighter than at O1 (see Table 5.1), both $\tilde{l}_R^+ \tilde{l}_R^-$ and $\tilde{l}_L^+ \tilde{l}_L^-$ events pass the cuts, and two edges are generated. The edge coming from the (lighter) right-slepton falls well under the SM background and so is not measurable, while the (heavier) left-slepton still has a cross section for pair-production high enough to let it form a good edge of its own at $\Delta M = (m_{\tilde{l}_L}^2 - m_{\tilde{\chi}_1^0}^2)/(2m_{\tilde{l}})$. It is important that this edge is not mistaken for one produced by $\tilde{l}_R$, so methods are needed to dismiss it. Discussion of these issues are delayed until Section 5.13.

5.6.7 Notes concerning interpretation of endpoint positions

Were all squarks to have the same mass, then (with perfect detector resolution and with infinite statistics) the endpoints of the edges listed above would be found at the positions listed above and summarised in Table 5.2. Note that these positions depend on four unknown parameters, namely the masses of the squark, the slepton and the two neutralinos participating in each decay. The mass of the lightest higgs boson also appears, but it will be assumed that the higgs mass will already be known or will be obtained by other methods (e.g. [19]) to within 2%.

In a more realistic “non-degenerate squark masses” scenario, every distribution with a squark related edge will actually be a superposition of many underlying distributions (one for each squark mass), each having its edge at a slightly different location. This results in some smearing of the edges, even before detector effects (resolutions, acceptances, jet energy calibrations etc.) are taken into account.

It is not possible to measure each of the squark masses separately. Consequently the quantity $m_{\tilde{q}}$ in Table 5.2 will, after being obtained from the “smeared” edges, represent a squark mass *scale* rather than a specific squark mass. In all squark related edges, except the l^+l^-q threshold, the contribution to the outermost part of each edge is provided by the heaviest squarks. In the case of the l^+l^-q threshold, the “true” endpoint is set by the lightest squark. However, if (as at O1 and S5) the other eleven squarks are much heavier (see Table 5.1), it is easier to measure the contribution coming from them. A strong correlation between $m_{\tilde{q}}$ and the mass of the heaviest squark would therefore be expected.

More work would be required to fully understand the “theoretical” systematic errors including the use of a full simulation of the expected edge positions. Better edge models than those used later in this analysis will definitely be needed in order to permit measurements of edge positions to be better correlated with functions of the particle masses. Other sources of systematic errors which require further analysis are the detector effects mentioned above, combinatorial backgrounds near the edges and possible cut biases. Such systematic errors can only meaningfully be studied when real data are available.

It has already been said that, within this thesis, it will be assumed that the above investigations could be performed so as to leave the eventual edge resolutions determined only by statistics and detector resolution. In this analysis, then, all edges are fitted with simple shapes (see Section 5.9.1) with the intention of extracting only an estimate of the *uncertainty* on the edge position, and not to obtain the edge position itself.

5.7 Event generation and pile-up

All events, except those for $q\bar{q} \rightarrow W^+W^-$ background processes, were simulated by HERWIG-6.0 [36]. The W -pair events were generated by ISAJET-7.42 [37].

At high luminosity, approximately 20 minimum bias events (“pile-up” events) are expected to occur in each bunch crossing. No pile-up events were generated for this analysis, but their effects will have been accounted for within the ATLFast simulation, which alters its reconstruction resolutions in the different luminosity environments appropriately. But resolutions are not the only things which change at high luminosity. For example, the number of jets per bunch crossing with more than a given p_T rises with luminosity, but this effect would not be reproduced by the simulation method described so far. Consequently all cuts which are expected to perform under high luminosity must be checked for susceptibility to the effects which ATLFast does not consider.

Except where stated otherwise, events were generated and studied in batches corresponding to 100 fb^{-1} in the high luminosity environment. This is expected to correspond to approximately one year of high luminosity running.

Standard Model		S5			O1		
$t\bar{t}$	W^+W^-	sparton	gaugino	slepton	sparton	gaugino	slepton
5.7E5	2.5E3	2.0E4	1.7E3	2.1E2	1.9E4	1.6E3	1.0E2

Table 5.3: *Production cross sections for signal and background processes at S5 and O1, and for SM backgrounds. All cross sections are in fb.*

5.8 Event selection cuts

Section 5.8.1 summarises the cuts employed by this analysis to obtain the invariant mass distributions whose edges underpin it. But before presenting the summaries, it will be useful to make some general comments about the cuts.

The cuts used to select events for the l^+l^- edge, l^+l^-q edge, l^+l^-q threshold and hq edge do not differ significantly from those used in the comparison study [31–33]. These cuts were all designed with S5 in mind, and so are not necessarily optimal for O1. The cuts used for the Zq edge are essentially the same as those used in [38] which was performed at LHC mSUGRA Point 2 (S2) where $m_0 = 400$ GeV, $m_{1/2} = 400$ GeV, $A_0 = 0$, $\tan\beta = 10$, $\mu > 0$.

None of the cuts above were developed with O1 in mind and one of them does not even consider S5. It is possible, therefore, that the analysis would benefit from an optimisation of the cuts, designed to equalise performance at S5 and O1. However, adequate performance in both test models is observed without recourse to optimisation, and the temptation to optimise has been resisted. It is hoped that this stance strengthens the claim of the analysis to the label “model independent”.

Although the way $l^\pm q$ events are treated in this analysis differs substantially from [31–33] (see Section 5.6.5) the basic purpose of the cuts designed to *obtain* those events, is common to both analyses. For reasons which will become clear later, however, it was found necessary to make some alterations to them. The most significant change is the relaxation of the “splitting requirement”. The original splitting requirement tried to guarantee that the jet which came from the quark produced in association with the observed dilepton pair was correctly identified. It achieved this by insisting that:

- both the dilepton pair and one of the two hardest jets j_1 and j_2 (ranked by p_T) was **consistent** with being the decay products of a squark (i.e. $m_{llj_i} < m_{\text{cutoff}}$), and
- the invariant mass m_{llj_j} of the two leptons and the other of the two highest p_T jets was **inconsistent** with these being the decay products of a squark (i.e. $m_{llj_j} > m_{\text{cutoff}}$).

Although the above demand for inconsistency increased the purity of the signal, it had a significant detrimental effect on the efficiency. In order to observe sufficient events at O1 (where the number of signal events passing the cuts was an order of magnitude less than the number of background events doing the same) it was found necessary to improve the cuts. For this thesis, the consistency requirement was retained but the inconsistency requirement dropped. Instead, m_{ll} was required to lie within the region given by $0 \leq m_{ll} \leq m_{ll}^{\text{max}}$. (The use of m_{ll}^{max} requires the l^+l^- edge measurement to have already been performed.) Application of the new cuts at O1 did not affect the rejection against background events, but improved the signal efficiency by a factor of six. At S5 a factor ten increase in both the number of signal and the number of background events passing cuts was observed. Background subtraction is also performed, modeling most of the OSSF backgrounds by the distributions produced by opposite-sign different-lepton-family (OSDF) events which pass the same cuts.

The production cross section for dilepton events is very small in both models (see Table 5.3). In fact, the total number of events available for the dilepton analysis is less than in any of the other channels. In addition, the $\tilde{l}\tilde{l}$ edge cuts are relatively soft, as there can be comparatively little visible energy in clean dilepton events. To show that the small number of dilepton events passing cuts is not in danger of being swamped by small variations in backgrounds or the cuts themselves, two set of cuts referred to as “hard” and “soft” are developed. The “hard” cuts attempt to maximise the signal

to background ratio in the vicinity of the dilepton edge in the $\Delta M(\chi)$ spectrum, while the much looser “soft” cuts try to maximise statistics at the edge at the expense of allowing in additional supersymmetric backgrounds. While it is presently expected that both sets of cuts could be used at high luminosity, the greater tolerance of jet activity from pile-up events exhibited by the “soft” cut is the basis of the conservative decision to base the conclusions of the analysis only on results obtainable with these cuts.

Most of the cuts listed in summary in the following section are self explanatory, but some general remarks will serve to motivate some of them:

- In general, leptons and jets are subject to minimum p_T requirements to ensure that they are well reconstructed.
- All cuts feature some sort of missing transverse momentum requirement, primarily to select for signal (missing LSPs) but also to reject against SM backgrounds (which can only form missing momentum via neutrinos). The $\cancel{p}_T$ cuts may be set higher in cuts looking for observables among the “boost invariant edges” (i.e. those looking at invariant masses), as these rely on information coming from decay chains on only “one side” of the primary interaction. In these events, correlations between the momenta of the event’s two LSPs are unimportant, events are in plentiful supply, and the long chains ensure that the neutralinos may readily attain large boosts. In contrast, the $\cancel{p}_T$ cut has to be set at a fairly low level when looking for dilepton events; here boosts are smaller, events are fewer in number, and the $\cancel{p}_T$ vector is actually used as an input to the ΔM observable.
- When events containing multi-stage decay chains are sought, high jet-multiplicity and/or large effective mass requirements select for events containing squarks or gluinos, while ensuring that SM backgrounds are reduced to negligible levels.

5.8.1 Cut summaries, grouped by observable

This section summarises the cuts used in the analysis. To avoid cumbersome repetition in the cuts which follow, suffixes attached to reconstructed leptons and jets will be used to indicate ranking of these objects by p_T . For example, j_2 will denote the reconstructed jet *with the second highest* p_T . Similarly the transverse momentum of this jet would be denoted by $\mathbf{p}_T^{j_2}$. Except where stated otherwise, cuts are given in the positive sense. Thus “ $p_T > 100$ GeV” is a *requirement* which an event passing the cut must satisfy, not a reason to cut it out.

l^+l^- edge

$n_{\text{leptons}} = 2$, both leptons are OSSF and $p_T^{l_1} \geq p_T^{l_2} \geq 10$ GeV.

$n_{\text{jets}} \geq 2$ and $p_T^{j_1} \geq p_T^{j_2} \geq 150$ GeV. $p_T \geq 300$ GeV.

The kinematics of OSSF leptons coming from background processes which produce uncorrelated leptons (for example tau decays) are modelled well by OSDF lepton combinations. Consequently, edge resolution is improved by “flavour subtracting” OSDF event distributions from OSSF event distributions.

l^+l^-q edge

$n_{\text{leptons}} = 2$, both leptons are OSSF and $p_T^{l_1} \geq p_T^{l_2} \geq 10$ GeV.

$n_{\text{jets}} \geq 4$, $p_T^{j_1} \geq 100$ GeV, and $p_T^{j_1} \geq p_T^{j_2} \geq p_T^{j_3} \geq p_T^{j_4} \geq 50$ GeV.

$p_T \geq \max(100 \text{ GeV}, 0.2M_{\text{effective}})$ and $M_{\text{effective}} \geq 400$ GeV, where $M_{\text{effective}}$ is defined by the scalar sum:

$$M_{\text{effective}} = \cancel{E}_T + p_T^{j_1} + p_T^{j_2} + p_T^{j_3} + p_T^{j_4}. \quad (5.24)$$

Since the desired edge is a maximum, only the smaller of the two m_{llq} combinations which can be formed using j_1 or j_2 is used.

Flavour subtraction is employed here as at the l^+l^- edge.

l^+l^-q threshold

All the cuts for the l^+l^- edge as above.

In addition $m_{ll}^{\max}/\sqrt{2} \leq m_{ll} \leq m_{ll}^{\max}$, where the value of $m_{ll}^{\max}$ would be obtained from the l^+l^- edge in practice, but was determined theoretically in this analysis.

Since the desired edge is a minimum, only the larger of the two m_{llq} combinations which can be formed using j_1 or j_2 is used.

hq edge

$n_{\text{leptons}} = 0$ and $p_T > 300$ GeV.

Exactly two b jets with $p_T^{j_{b1}} \geq p_T^{j_{b2}} \geq 50$ GeV. No other b jets, regardless of p_T .

At least two non- b jets with $p_T^{j_{q1}} \geq p_T^{j_{q2}} \geq 100$ GeV, with at least one inside $-2.0 \leq \eta^j \leq 2.0$.

m_{bb} within 17 GeV of higgs peak in m_{bb} spectrum.

Since the desired edge is a maximum, the non- b jet (i.e. j_{q1} or j_{q2}) chosen to form the m_{hq} invariant mass is that which minimises m_{hq} .

Zq edge

$n_{\text{leptons}} = 2$, both leptons are OSSF and $p_T^{l1} \geq p_T^{l2} \geq 10$ GeV.

m_{ll} within 2.5 GeV of the centre of the Z-mass peak in the m_{ll} spectrum.

At least two non- b jets exist with $p_T^j \geq 100$ GeV and $p_T \geq 300$ GeV.

Since the desired edge is a “maximum”, the jet chosen to form the m_{Zq} invariant mass is the one (from those with $p_T^j \geq 100$ GeV) which minimises m_{Zq} .

Again, flavour subtraction is used.

$l^{\pm}q$ high and $l^{\pm}q$ low edges

All the cuts for the l^+l^-q edge above are required.

Additionally, events consistent with the l^+l^- edge measurement are selected by asking for $m_{ll} \leq m_{ll}^{\max} + 1$ GeV.

To choose the jet from which to form $m_{l^{\pm}q}$ we select $j_i = j_1$ or j_2 , whichever gives the smaller value of m_{lj_i} , and require $m_{lj_i} < m_{\text{cutoff}}$, where m_{cutoff} is chosen to be above the l^+l^-q edge, but is otherwise arbitrary. For easy comparison with [31], $m_{\text{cutoff}} = 600$ GeV was used in this analysis.

Finally, the two invariant mass combinations, $m_{l^{\pm}q}$, are assigned to $m_{lq(\text{high})}$ and $m_{lq(\text{low})}$ as defined earlier by Equations (5.8) and (5.9).

 $\tilde{l}\tilde{l}$ edge (hard cuts)

Events are required to have exactly one OSSF pair of isolated leptons satisfying $p_T^{l_1} > 50$ GeV and $p_T^{l_2} > 30$ GeV. Events are not permitted to contain any other isolated leptons.

$\delta_T < 20$ GeV is required, where $\delta_T = |\mathbf{p}_T^{l_1} + \mathbf{p}_T^{l_2} + \mathbf{p}_T|$. Large δ_T in signal events is indicative of an unidentifiable transverse boost to the centre of mass frame, perhaps due to initial state radiation. Large δ_T in other (not necessarily signal) events simply suggests an inconsistency with the desired event topology.

Events containing one or more jets with $p_T^j > 40$ GeV are vetoed. This cut also helps to reduce standard model backgrounds, notably $t\bar{t}$. The lower this cut is placed, the better for the background rejection. However the cut cannot be placed too low (especially at high luminosity) due to the significant number of jets coming from the underlying event and other minimum bias events in the same bunch crossing. For order 25 minimum bias events per bunch crossing, [39] estimates that about 10% (1%) of bunch crossings will include a jet from the underlying event with a p_T greater than 40 GeV (50 GeV).

Events with $|m_{l_1 l_2} - m_Z| < 5 \text{ GeV}$ are vetoed to exclude lepton pairs from Z bosons.

$m_{l_1 l_2} > 80 \text{ GeV}$ and $\not{p}_T > 50 \text{ GeV}$ are also required.

$\tilde{l}\tilde{l}$ **edge (soft cuts)** These are as above, but

- the upper limit for δ_T is extended from 20 to 90 GeV, and
- the dilepton invariant mass cut is removed altogether.

5.9 Edge resolutions

5.9.1 Resolutions for “boost invariant” edges

By applying at S5 the cuts listed in Section 5.8.1 it is possible to reproduce the m_{ll} and m_{llq} distributions whose edges are analyzed in [31–33] (see Figure 5.6: a, b and d). A second application confirms that the same cuts may also be used to generate edges of similar quality at O1 (see Figure 5.7: a, b and d). In addition, plots c1 and c2 in Figures 5.6 and 5.7 display the $m_{lq(\text{high})}$ and $m_{lq(\text{low})}$ distributions generated from the modified $l^\pm q$ edge cuts also listed above.

The distributions seen at S5 and O1 are broadly similar to each other. The most obvious difference between the two sets of data is the peak at the Z mass present in Figure 5.7(a) but absent in Figure 5.6(a). This peak comes from direct $\tilde{\chi}_2^0 \rightarrow \tilde{\chi}_1^0 Z \rightarrow \tilde{\chi}_1^0 l^+ l^-$ decays, which have a branching fraction of 39% at O1 compared with 0.65% at S5. The most common direct neutralino decay mode at S5 is through the h_0 (65%) of which only a tiny proportion (0.02%) decay to the two light leptons as the partial decay width $\Gamma(h_0 \rightarrow ff)$ goes approximately as m_f^2 . Scenarios in which the $l^+ l^-$ edge happens to coincide with the Z peak would require special treatment and are not considered here.

It has already been mentioned that the primary reason for plotting these distributions is not to recover the locations of the edges themselves, but to get an idea of the limit which finite statistics imposes on the resolution of each endpoint. These errors will serve as the foundation for assessing the method as a whole, under the assumption that supersymmetric systematic effects are small in comparison with detector systematics and statistical errors.

In order to extract estimates for the statistical errors on the endpoints, simple fits have been made to the distributions of Figures 5.6 and 5.7. The shapes fitted to the

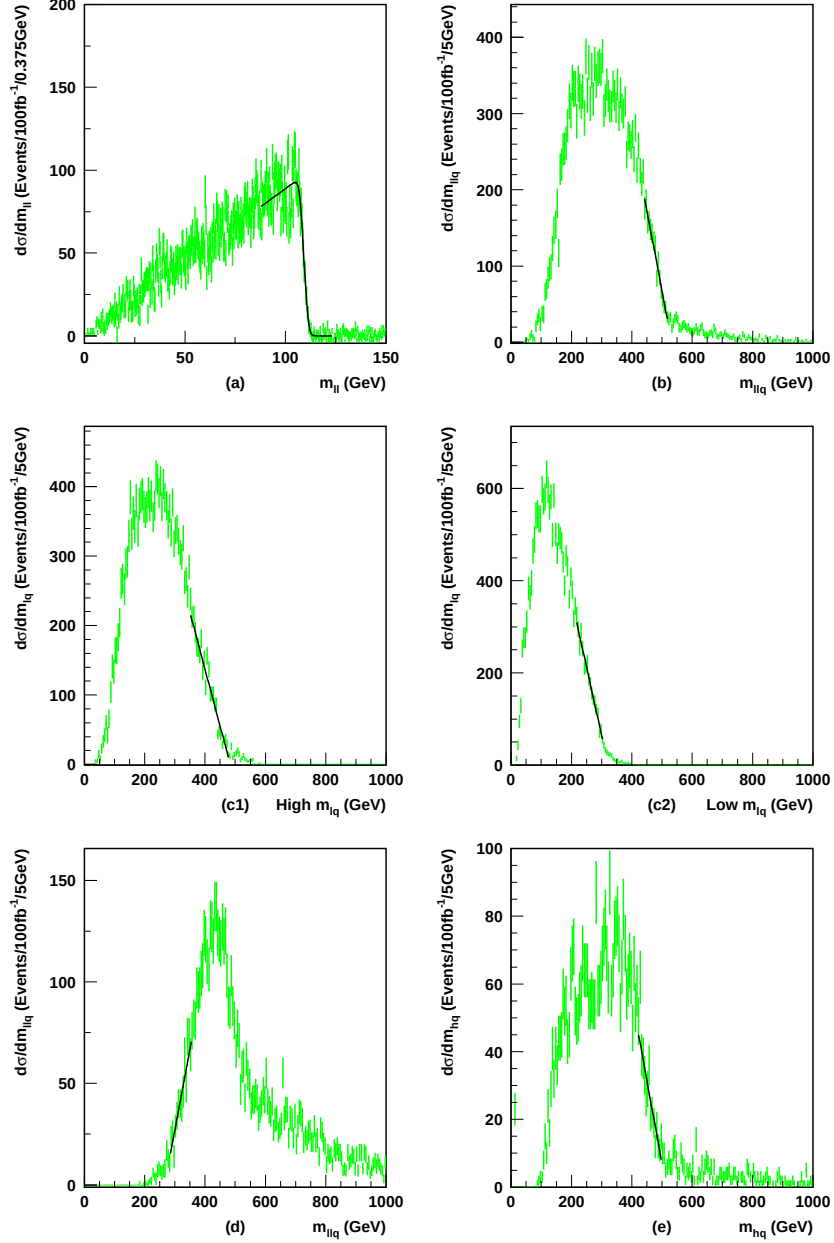


Figure 5.6: The $S5$ distributions whose end points are described in Table 5.2: a) the l^+l^- edge, b) the l^+l^-q edge, c1) the $l^\pm q$ high-edge c2) the $l^\pm q$ low-edge, d) the l^+l^-q threshold and e) the hq edge. Plots were produced with the cuts described in Section 5.8.1. The number of events corresponds to 100 fb^{-1} of high luminosity running.

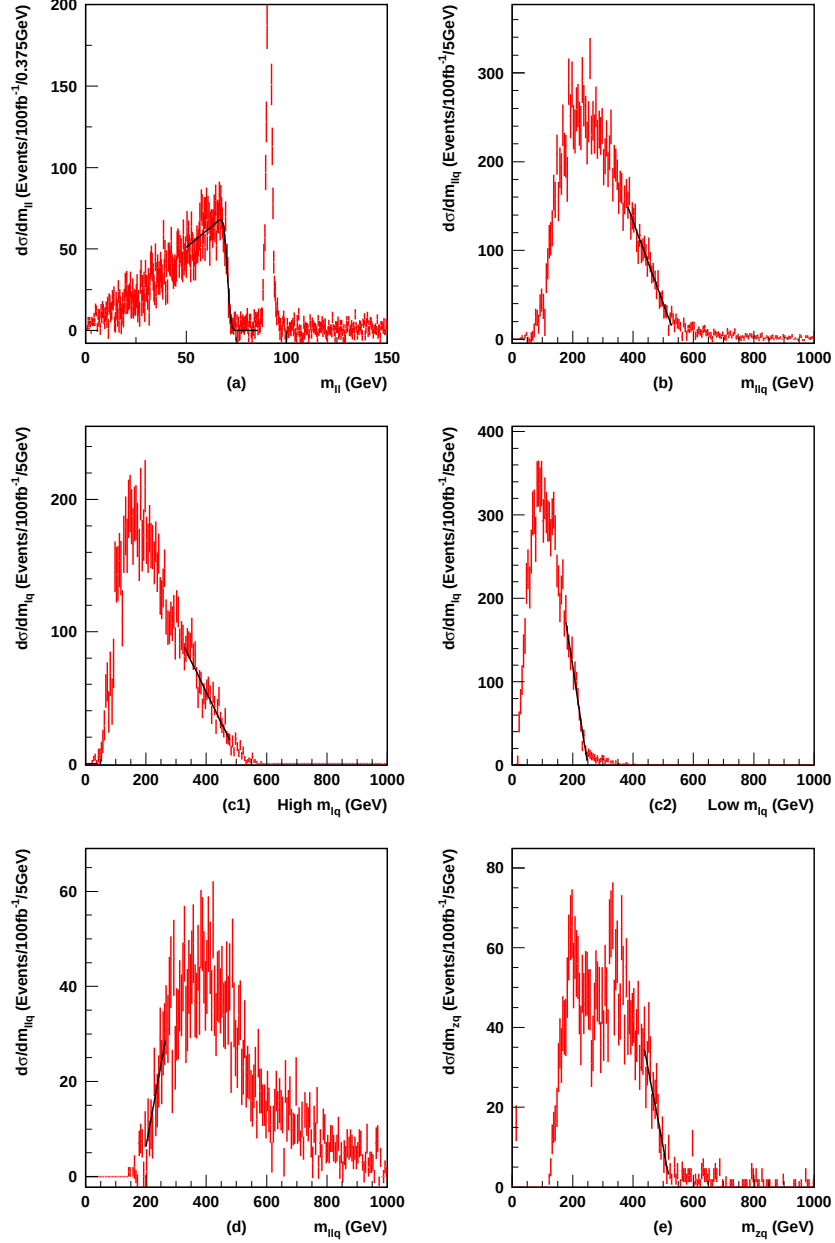


Figure 5.7: The $O1$ distributions whose end points are described in Table 5.2: a) the l^+l^- edge, b) the l^+l^-q edge, c1) the $l^\pm q$ high-edge c2) the $l^\pm q$ low-edge, d) the l^+l^-q threshold and e) the Zq edge. Plots were produced with the cuts described in Section 5.8.1. The number of events corresponds to 100 fb^{-1} of high luminosity running.

Endpoint	S5		O1		Table 5.2 values	
	Fit	Fit error	Fit	Fit error	S5	O1
l^+l^- edge	109.10	0.13	70.47	0.15	109.13 (109.13)	70.50 (70.5)
l^+l^-q edge	532.1	3.2	544.1	4.0	536.7 (494.9)	530.5 (496.1)
$l^\pm q$ high-edge	483.5	3.8	515.8	7.0	464.2 (428.1)	513.6 (480.3)
$l^\pm q$ low-edge	321.5	2.3	249.8	1.5	337.0 (310.7)	231.3 (216.2)
l^+l^-q threshold	266.0	6.4	182.2	13.5	264.9 (264.1)	168.1 (158.2)
Xq edge	514.1	6.6	525.5	4.8	509.2 (471.0)	503.4 (471.8)
ΔM edge: $\tilde{l}_R$	—	—	—	10%	31.6 (31.6)	72.8 (72.8)
ΔM edge: $\tilde{l}_L$	—	10%	—	—	88.5 (88.5)	114.8 (114.8)

Table 5.4: *Endpoints and associated fit uncertainties. The results of the naive fits shown in Figures 5.6 and 5.7 may be seen in the ‘fit’ and (one-sigma) ‘fit error’ columns. For the reasons outlined in Section 5.6.7 there is presently a lack of good theoretical predictions for these edge positions. (In this context a “good prediction” is one capable of taking into account, up to a level compatible with the statistical/fit errors, either the naivety of the fits, or the effects introduced by edge distortions and the presence of backgrounds.) The best guides presently available are the quantities listed in Table 5.2, evaluated at $m_{\tilde{q}}$ equal to the largest squark mass. These values are shown for comparison purposes. Also indicated (in parentheses) are the same quantities evaluated at $m_{\tilde{q}}$ equal to the average squark mass. All values are in GeV except where stated otherwise.*

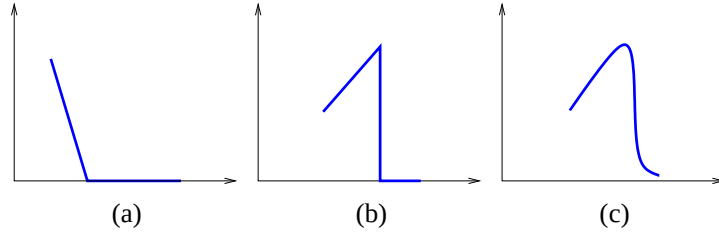


Figure 5.8: *Line-shapes used for fits to the mass distributions. Shape (a) (or its reflection in the vertical axis) is the standard “straight line” used to extract the l^+l^-q edge, $l^\pm q$ high-edge, $l^\pm q$ low-edge and l^+l^-q threshold resolutions. Shape (b) models the expected shape of the sharp l^+l^- edge in the absence of detector effects. The l^+l^- edge is actually fitted with shape (c), which is identical to (b) except that it is smeared with a Gaussian resolution whose width is a free parameter of the fit.*

data (see Figure 5.8) and the algorithms for determining the boundaries of the fitted regions have been kept as simple and generic as possible, with the intention of making the fit results both conservative and simple to interpret. Only two shapes were used to fit the edges; a straight line for most edges, and a triangular distribution smeared by a Gaussian resolution for the l^+l^- edge. The use of the special method at the l^+l^- edge is justified by the unique step-like nature of its edge (which may not be approximated by a straight line of finite gradient) and also by the fact that the expected shape of the edge is simple to interpret theoretically (see Section 5.6.1) and is unlikely to introduce unwanted model dependence.

To try to avoid bias toward one or either model or toward one or either edge, all the distributions fitted with a *straight line* (ten in total) were fitted by *the same algorithm*. This algorithm was given the tasks of both identifying the range of bins over which to fit, as well as the endpoint of the edge and its gradient. There were a number of inputs to the algorithm common to all distributions, describing “what to look for in an edge” (such as a maximum height which the top part of the edge may reach as a fraction of the height of the bin containing the most events, or a minimum number of events required to be present in each bin over the fit range) but only one input which was permitted to

vary from edge to edge; namely the instruction indicating whether to look for a threshold (i.e. left-hand edge) or a forwards edge (i.e. right-hand edge).

The resolutions obtained from the fits to the distributions in Figure 5.6 and Figure 5.7 are listed in Table 5.4.

Given real data, it would be worth putting substantial effort into understanding how detector and physics effects affect the shape of each distribution in turn, particularly in the neighbourhood of each edge. This might significantly improve particle mass resolutions, or at least improve confidence in the errors associated with each edge. But as the events being processed here are entirely synthetic, the fit errors listed in Table 5.4 will be used as they are, providing a first order estimate of the magnitude of the eventual fit error.

5.9.2 ΔM edge resolution

Determining the likely resolution for the ΔM edge requires a slightly different approach to that used for the “boost invariant” edges above. Whereas the cuts used to obtain the above invariant mass distributions have high SM rejections (primarily due to the presence of at least one high mass p_T or $M_{\text{effective}}$ requirement), the ΔM cuts cannot be set as hard, primarily because the desired dilepton events have very little jet activity. (With the dilepton production cross sections typically two orders of magnitude smaller than the squark/gluino production cross sections (Table 5.3) a low efficiency is not affordable.) There are also irreducible SM backgrounds (primarily $W^+W^- \rightarrow l^+l^-\nu\bar{\nu}$ and $t\bar{t} \rightarrow b\bar{b}W^+W^- \rightarrow jjl^+l^-\nu\bar{\nu}$ in cases where jets are below the reconstruction threshold or are outside detector acceptance) which have signatures identical to dilepton events. The smallness of the signal and the presence of these backgrounds would cause problems for naive straight-line fitting techniques of the type used above, and so estimation of the

edge precision requires the following discussion.

Figure 5.9 shows the ΔM distributions obtained at S5 and O1 after applying the cuts listed in Section 5.8.1 to a 100 fb^{-1} sample of signal dilepton events. Events from lighter sleptons ($\tilde{e}_R$ and $\tilde{\mu}_R$) occupy the unhatched region in each plot, while the events from heavier left-sleptons are cross hatched. It will be noted that events from both light and heavy sleptons succeed in passing the cuts in both models. In principle, then, there are two edges to be observed in each of the models: one for the lighter slepton and one for the heavier. We note, however, that as the slepton masses increase, their production is strongly suppressed, and so there are very few heavy slepton events at O1 where there are in fact none within 10 GeV of the kinematic limit. It is readily observed that at the three remaining edges, where statistics are higher, there is good agreement between the theoretical prediction and the observed endpoint of each distribution.

Looking next at the SM backgrounds in Figure 5.10(a) and Figure 5.10(b) we see that the number of $t\bar{t}$ and W^+W^- events passing the “soft” cuts is about two orders of magnitude larger than the number of signal events passing cuts. Fortunately, although large in number, these events are easy to understand. As the source of leptons and missing energy in these events are not sleptons but W -bosons, we expect the edge of the SM background to fall at half the W mass (as this is the modulus of the three momentum of the decay products of the W in its rest frame). From the “zoomed out” histograms in Figures 5.10(a) and 5.10(b) it may be seen that the W -pair background does indeed have an endpoint at $m_W/2$. The same is true of the $t\bar{t}$ events passing the hard cuts, although under the soft cuts the $t\bar{t}$ edge does begin to wander up to 50 GeV, probably due to the effects of poorer $\cancel{p}_T$ resolution when more hadronic energy is permitted in events as well as the effect of allowing a larger range of transverse boosts for the W -bosons. Nevertheless, the sheer number of SM events passing both sets of cuts, combined with the finite detector resolution, allows significant numbers of events to reach 50 GeV (hard

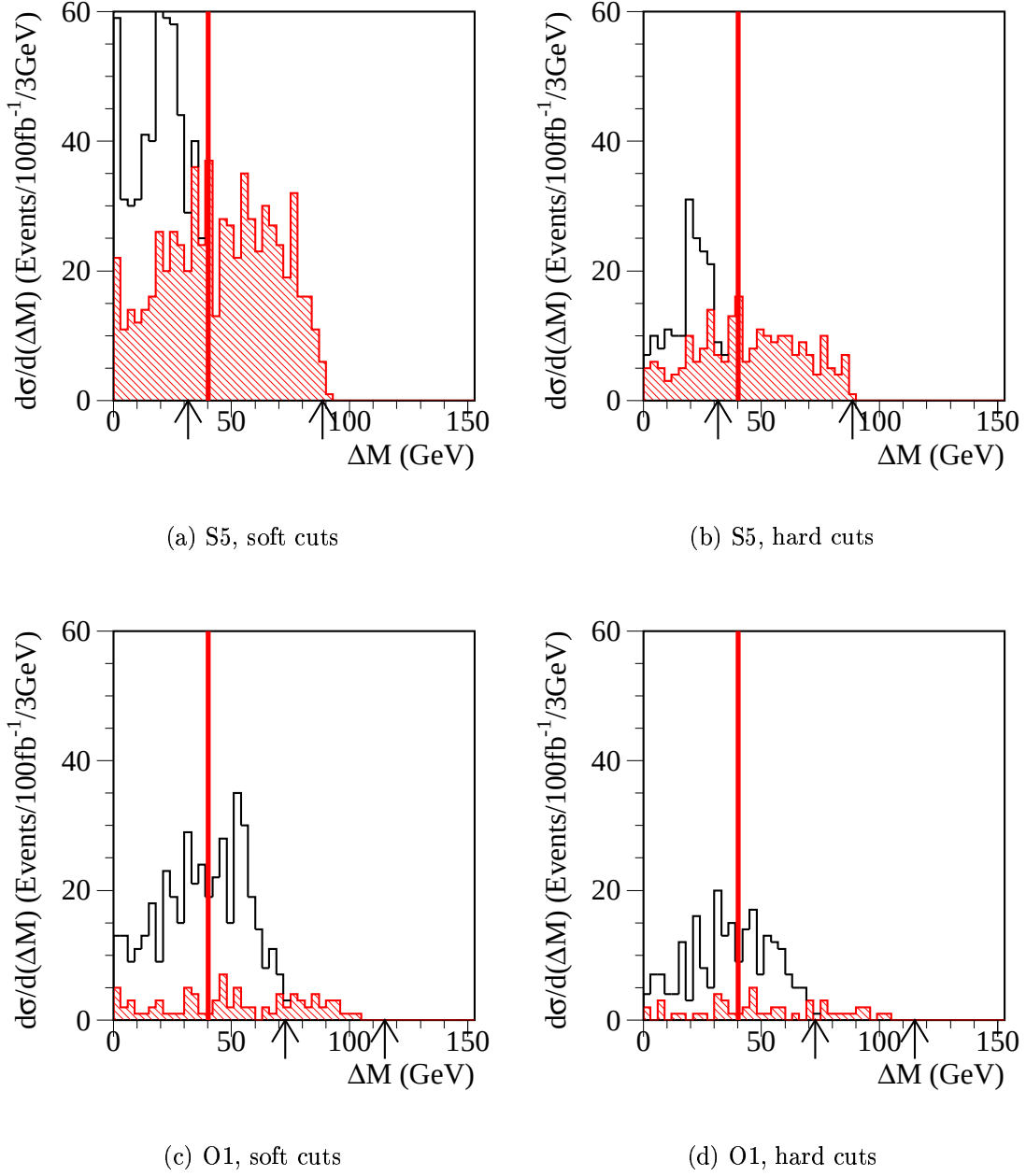


Figure 5.9: ΔM distributions obtained at S5 and O1 after applying the cuts listed in Section 5.8.1 to a 100 fb^{-1} sample of $\tilde{e}_R \tilde{e}_R$, $\tilde{\mu}_R \tilde{\mu}_R$, $\tilde{e}_L \tilde{e}_L$ and $\tilde{\mu}_L \tilde{\mu}_L$ dilepton events. Events from light right-sleptons (unhatched) are stacked on top of those from heavier left-sleptons (hatched). Only events from OSSF leptons combinations are shown. The plots are generated without OSDF background subtraction, but were it to be performed, no significant differences would be apparent as only 4 (12) signal events are able to pass OSDF soft cuts at S5 (O1). Arrows indicate the values of ΔM^{max} predicted by theory for the two types of slepton in each model. A red vertical line is drawn through each plot at half the W mass.

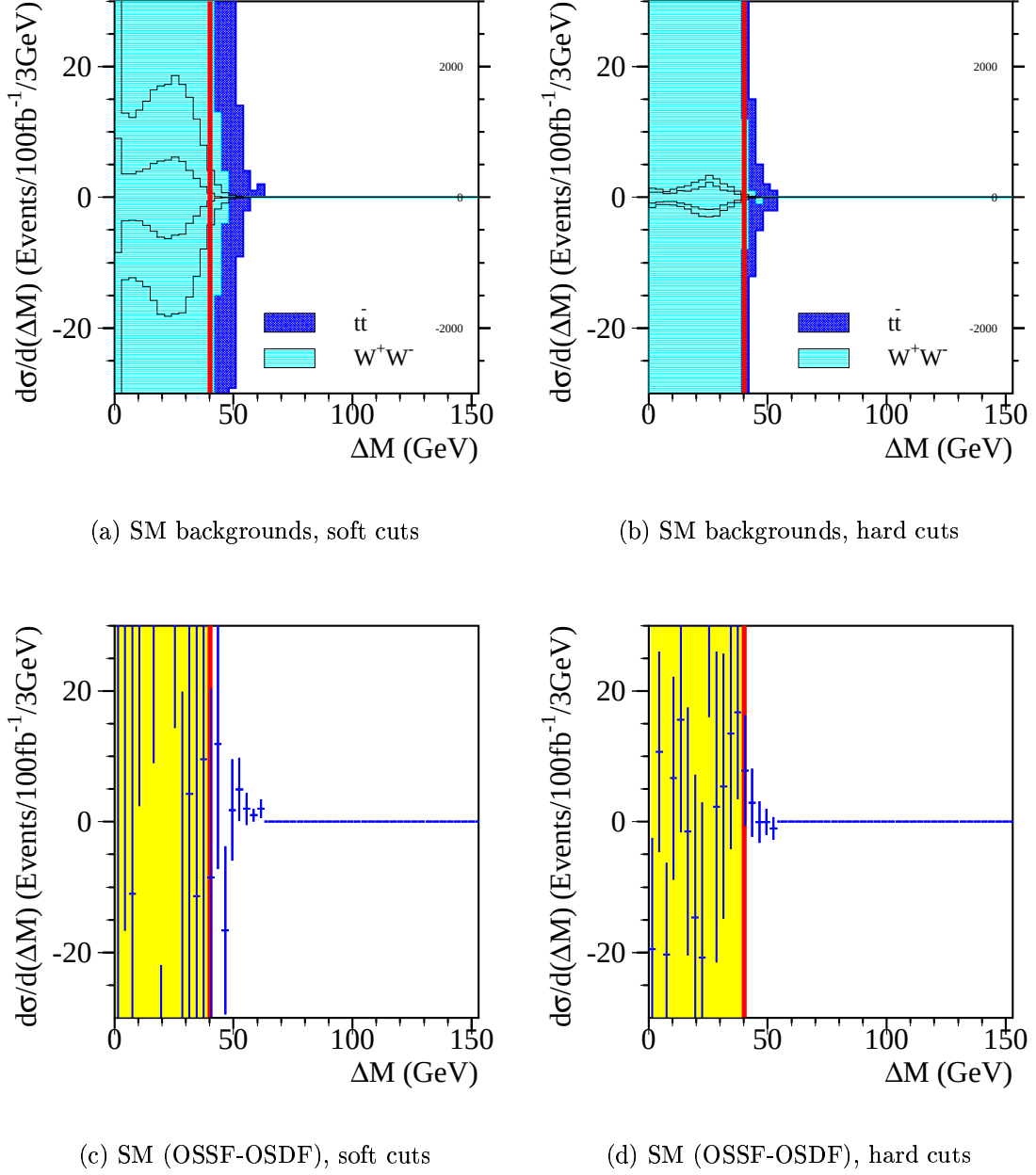
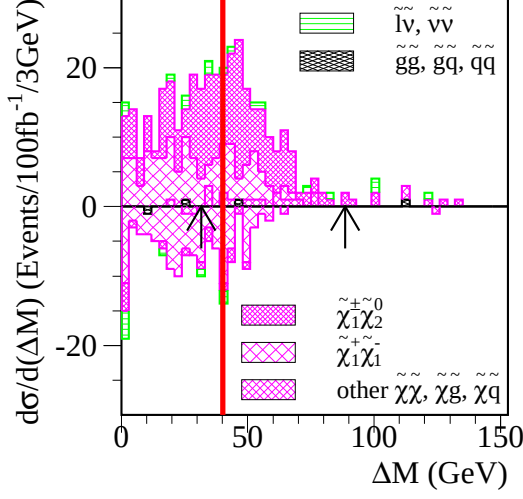
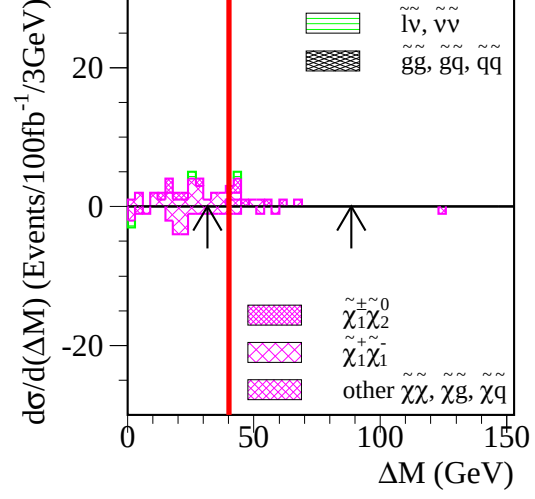


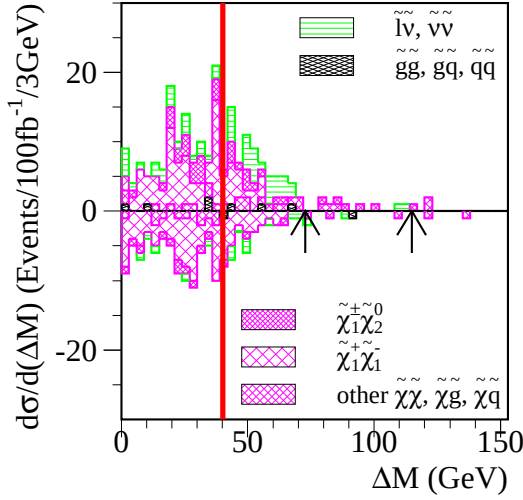
Figure 5.10: *SM background to ΔM measurement. The plots above show the ΔM distributions obtained after cuts from 100 fb^{-1} of W -pair and $t\bar{t}$ background samples. In (a) and (b), events passing the OSSF cuts are stacked with positive weight, while events passing the OSDF cuts are stacked with negative weight. Because the number of events passing cuts in (a) and (b) is very large, the outlines of the stacked histograms are overlapped after being downscaled by a factor of 100 (see scales on right hand side of each plot). Plots (c) and (d) show the result of subtracting the OSDF events passing the cuts from the OSSF events passing the cuts. A red vertical line is drawn through each plot at half the W mass. Plots (c) and (d) are shaded to the left of this line in order to draw attention to the performance of the background subtraction above this point.*



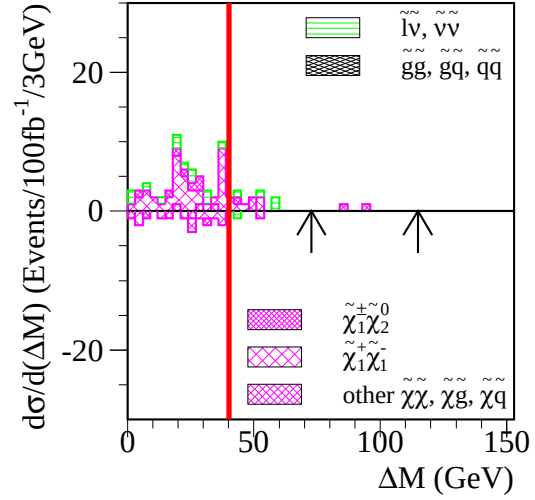
(a) S5, soft cuts



(b) S5, hard cuts



(c) O1, soft cuts



(d) O1, hard cuts

Figure 5.11: ΔM distributions obtained at S5 and O1 after applying the cuts listed in Section 5.8.1 to the remainder of the supersymmetric cross section not generated for Figure 5.9. Events passing the OSSF cuts are stacked by primary parton pair with positive weight, while events passing the OSDF cuts are stacked with negative weight. A red vertical line is drawn through each plot at half the W mass. The arrows of Figure 5.9 are provided for comparison.

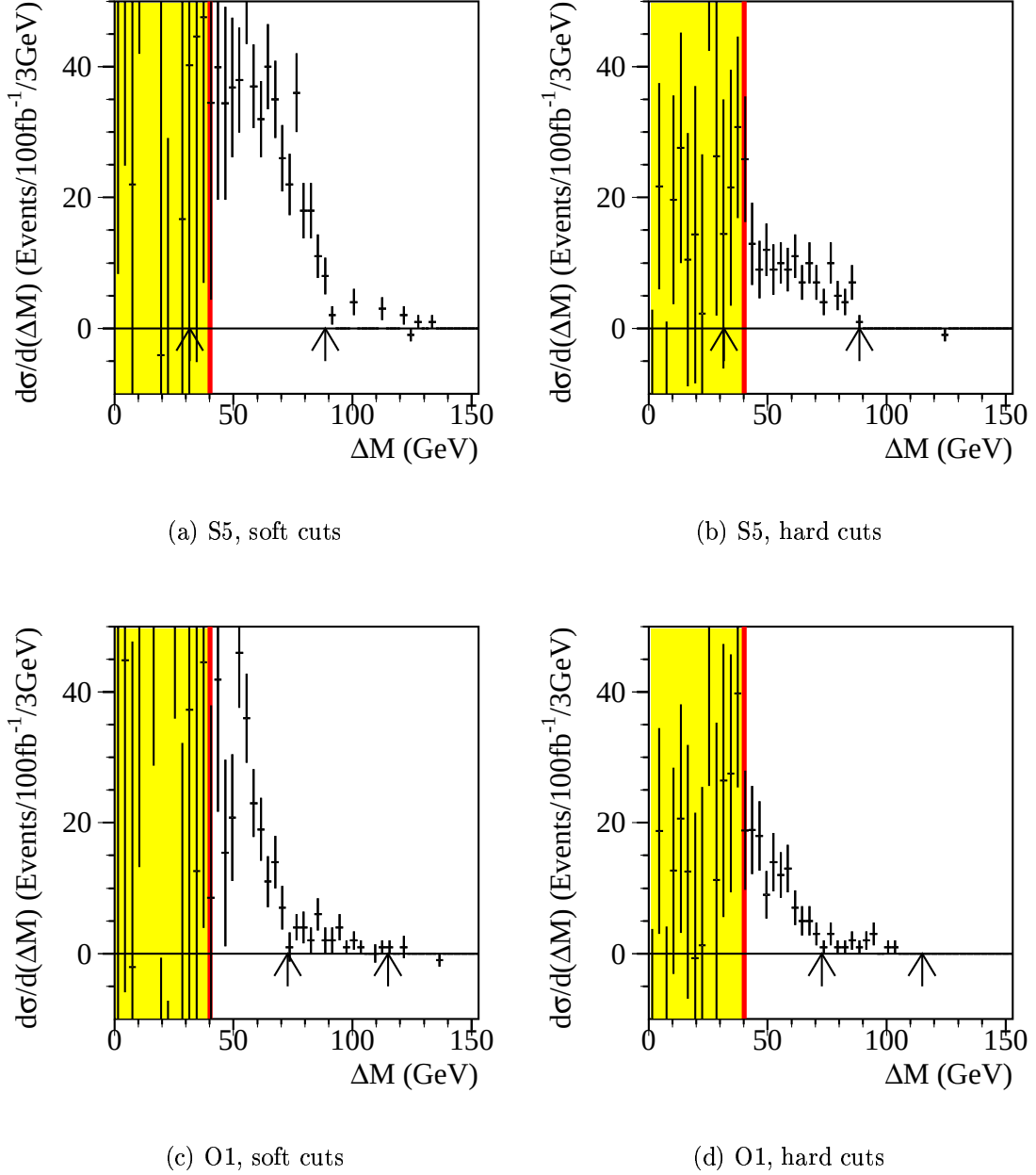


Figure 5.12: *Flavour-subtracted ΔM distributions for combined signal and background at S5 and O1 after applying the cuts listed in Section 5.8.1. A red vertical line is drawn through each plot at half the W mass. Plots are shaded to the left of this line in order to draw attention to the events which reconstruct above this point. Compare these plots with those of Figure 5.9 which contain only signal events. The arrows of Figure 5.9 are provided for comparison.*

cuts) or 60 GeV (soft cuts). This leakage threatens to further restrict the minimum position at which a supersymmetric edge in ΔM might be observed. Fortunately, the OSSF SM background may be seen to be well modelled by its OSDF counterpart. Since the signals from dilepton pair production are expected to be purely OSSF (because of lepton-number conservation) there is a strong motivation for using different-family background subtraction. The quality of this subtraction may be seen in Figures 5.10(c) and (d). Across the whole range of $\Delta M > 0$ the SM background is consistent with zero after subtraction. If these Figures are compared with the signal distributions shown in Figure 5.9, it may be seen that after subtraction, the SM backgrounds have been reduced to a level where they are very much less than the signal in all bins above $m_W/2$, except perhaps in the two lowest bins in the case of the soft cuts. It is reasonable to assume, therefore, that in the real experiment; the OSDF ΔM distribution could be used to place limits on the number of OSSF events which, if observed above a given value of ΔM , would provide significant evidence of physics beyond the SM. Because of the interplay between signal edge position and signal cross section, one cannot make a definitive statement of the form “if a supersymmetric model has a dilepton edge above X GeV then it will be observable”. However, it is clear that observation of dilepton edges below $m_W/2$ is highly unlikely, observation of edges in the range $m_W/2 < \Delta M < 50$ GeV will be possible if the number of signal events is significantly larger than the number of SM background events, and observation of an edge above 50 GeV is likely to be limited only by signal cross section itself. From these considerations we see that the light slepton edge at S5 is rendered unobservable by the presence of the SM background, but that the statuses of the other edges (including those of O1) remain unaffected.

Having looked at the signal and the SM processes, there remains only the supersymmetric backgrounds to consider. Figure 5.11 shows the remaining non-signal events from supersymmetric production processes, after both the OSSF and the OSDF versions of the cuts. We notice here that, with the single exception of the soft cuts at S5, the

OSSF distributions are well modelled by the OSDF events passing the same cuts. Even where the background subtraction is not perfect, we again see that the supersymmetric backgrounds remain small in comparison with the signal distributions (see Figure 5.9) and are either absent or constant in the vicinity of the slepton edges.

All events (signals and backgrounds) are combined in Figure 5.12 after different-family background subtraction. The reader is encouraged to compare these plots with those from Figure 5.9 showing the desired signal shapes. As expected, all signal shape information is lost to the left of $m_W/2$ due to obliteration by the SM backgrounds. To the right of this point, at least one clear edge is observable in both models (the left-slepton edge at S5, and the right-slepton edge at O1) and in both sets of cuts. The hard cuts are able to suppress the supersymmetric backgrounds to such a degree that there is even compelling evidence at O1 for the existence of *two* edges, although the lack of statistics in the higher edge limits the precision with which the endpoint may be located. By the same token, the lack of a high tail at S5 (under the hard cuts) might be interpreted as evidence that the salient edge was itself due to the heaviest left-sleptons. These conjectures will not be pursued in this thesis, however. It will be sufficient for the analysis in hand to draw the following conservative conclusion from the plots: It will be assumed that where a signal dilepton edge occurs in the range $50 \text{ GeV} < \Delta M^{\text{max}} < 100 \text{ GeV}$, then the location of the endpoint can be determined using the soft cuts alone to a precision of order 10%.^h The (lowest) edge measured in this way will be referred to as the “principal ΔM edge”. At O1 it comes from lighter right-sleptons, while at S5 it comes from heavier left-sleptons.

Finally, we note that if the $t\bar{t}$ background were to prove larger than anticipated, the employment of a b -jet veto could be expected to reduce the number of $t\bar{t}$ events passing

^hThe value of 10% is chosen somewhat arbitrarily, but plays no further role in the thesis, save to indicate in Figure 5.19 the regions of parameter space favoured by ΔM measurements in the two test models.

the soft cuts by a factor of order the inverse of the b -tagging efficiency; most of these events contain one b -jet from the top decay (the other being below the reconstruction threshold or else outside the detector acceptance).

5.10 Modelling the endpoint measurement process

Given the non-trivial relationships between the expected endpoint locations and the sparticle masses, the approach used to evaluate the eventual errors for the reconstructed sparticle masses is non trivial. The approach used in this thesis, consistent with the assumption of negligible supersymmetric systematic error, will require “modelling the endpoint measurement process itself”. This will allow mass reconstructions by an “ensemble” of ATLAS experiments to be performed in parallel, the spread of whose measurements will form an estimate of the accuracy obtainable by any one measurement on its own.

Using the fit errors, σ_i , of Table 5.4 as estimates of the statistical error associated with a typical endpoint measurement, the whole process of edge fitting, calibration and endpoint measurement is mimicked using the following prescription:

1. Select n endpoints to “measure”.
2. Identify the values $O_i(M_{\text{model}})$ ($i = 1, \dots, n$) which the observables (which is to say endpoint positions) would be expected to take if a perfect measurement could be performed. Clearly the observables would depend on the masses of the sparticles in the model in question, and this is indicated by the functional dependence of the $O_i(M_{\text{model}})$ on the masses, M_{model} .
3. Construct “measured” observables, $\Omega_i^{\text{sm}}(M_{\text{model}})$, by smearing the $O_i(M_{\text{model}})$ according to $\Omega_i^{\text{sm}} = O_i(M_{\text{model}})(1 + \kappa_i Y) + \sigma_i X_i$, where $X_i \sim N(0, 1)$ are n random

variables from the Normal Distribution (with zero mean and unit standard deviation), σ_i is the anticipated statistics component of the fit error for the i^{th} observable, $Y \sim N(0, \sigma_{JetScale})$ is another random variable from a Normal Distribution, and κ_i is equal to 1 in those observables scaling with jet momenta and equal to 0 in those observables which do not.ⁱ The role of κ_i , Y and $\sigma_{JetScale}$ is to model the (correlated) dependence of the observables on the absolute jet energy scale. $\sigma_{JetScale}$ is taken to be 0.01 for the reasons given in Section 4.2.3.

The n observables Ω_i^{sm} may thus be thought of as endpoint locations which *might have been obtained at ATLAS after collection of 100 fb^{-1} of data*, were it possible to remove systematic effects arising from poor understanding of supersymmetric backgrounds in the vicinity of the edges. As the generation of the Ω_i^{sm} only requires the selection of n random numbers, many such “measurements” can be made very quickly.

5.10.1 Reconstruction of sparticle masses

Given a set of “measured” observables Ω_i^{sm} , sparticle masses are reconstructed by finding the particular set of masses, M_{fit} , which minimises the chi-squared of the fit of the predictions $O_i(M_{\text{fit}})$ to the measurements Ω_i^{sm} .

The main parameters M of the fit are $m_{\tilde{\chi}_1^0}$, $m_{\tilde{l}}$, $m_{\tilde{\chi}_2^0}$ and $m_{\tilde{q}}$. The masses m_h and m_Z also appear as fit parameters, but they are strongly constrained (particularly in the case of the Z) by existing LEP measurements and the expected LHC higgs mass error (0.0077% and 2.0% respectively).

ⁱThe observables taken to scale with jet momenta were the l^+l^-q edge, the l^+l^-q threshold, the $l^\pm q$ low-edge, the $l^\pm q$ high-edge, the Zq edge and the hq edge.

Formally, then, the chi-squared is formulated as:

$$\chi^2(M) = \sum_{i,j=1}^n (\Omega_i^{\text{sm}} - O_i(M)) \mathbf{H}_{ij} (\Omega_j^{\text{sm}} - O_j(M)) \quad (5.25)$$

where the matrix $\mathbf{H}$ is the inverse of the covariance matrix whose elements are given by

$$(\mathbf{H}^{-1})_{ij} \equiv \langle (\Omega_i^{\text{sm}} - \overline{\Omega_i^{\text{sm}}}) (\Omega_j^{\text{sm}} - \overline{\Omega_j^{\text{sm}}}) \rangle \quad (5.26)$$

$$= \delta_{ij} \sigma_i^2 + \kappa_i \kappa_j \sigma_{\text{JetScale}}^2 O_i(M_{\text{model}}) O_j(M_{\text{model}}) \quad (5.27)$$

$$\approx \delta_{ij} \sigma_i^2 + \kappa_i \kappa_j \sigma_{\text{JetScale}}^2 \Omega_i^{\text{sm}} \Omega_j^{\text{sm}} \quad (5.28)$$

where δ_{ij} is the Kronecker delta and where σ_i , κ_i and σ_{JetScale} have the same meanings as before (see Section 5.10). The results of the fit, M_{fit} , are thus the particular masses which minimise the chi-squared. M_{fit} may be interpreted as “the set of masses most consistent with Ω_i^{sm} ”, where Ω_i^{sm} are a set of measurements which could have been made by ATLAS after 100fb^{-1} of high luminosity running. In order to determine the accuracy with which this reconstruction can be performed, the above fitting process is repeated many times for different values of the Ω_i^{sm} , so as to observe the spread in the reconstructed masses.

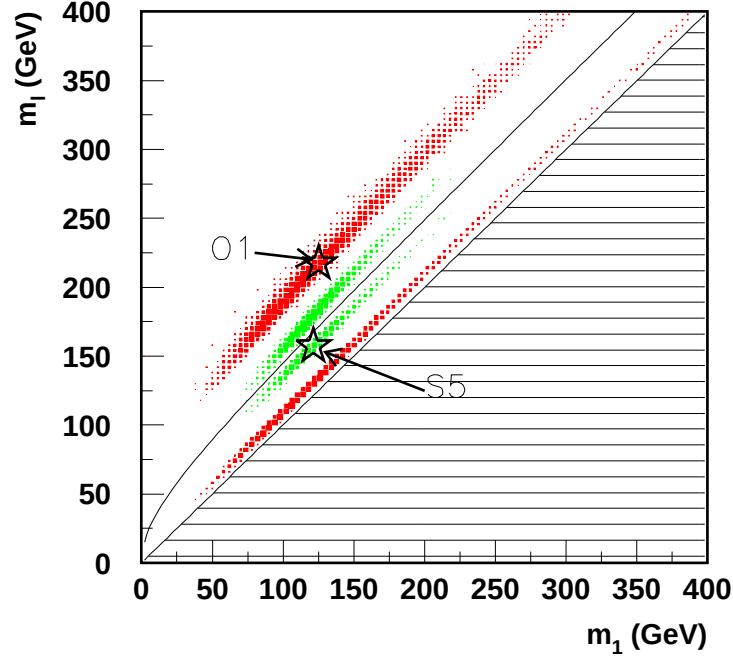


Figure 5.13: The spread of $\tilde{l}_R$ and $\tilde{\chi}_1^0$ masses reconstructed by an ensemble of simulated ATLAS experiments, each using the “basic” analysis (see Section 5.11). Reconstructions in which the underlying model was S5 are shown in light green, while the results of reconstruction based on O1 are shown in dark red. The stars indicate the input values of the masses in each of the models. The thin black curve denotes the “fault line” defined by Equation (5.31) and described in the text. The region forbidden by the mass hierarchy assumed in Equation (5.6) is shown hatched.

5.11 Basic reconstruction

The role of each edge in the analysis is best demonstrated by an incremental approach, in which edges are incorporated one at a time. The first form of the analysis, including only the l^+l^- edge, the l^+l^-q edge, the $l^\pm q$ high-edge, the Xq edge and the l^+l^-q threshold, might be termed the “basic analysis”.

Of all the forms of the analysis which will be described here, the basic analysis is the one most similar to that used in [31–33]. Common to both, are the number of edges and the types of particles involved in them. The very important difference between

them, however, is that the analysis described here does not presuppose any *special* relationships between the masses of the sparticles in the decay chains, as is evidenced by the generalisations of the l^+l^-q edge and $l^\pm q$ high-edges. As a consequence, the basic analysis does not constrain sparticle masses as much as the original did.

Figure 5.13 shows the spread of $\tilde{l}_R$ and $\tilde{\chi}_1^0$ masses that are reconstructed by use of the basic analysis at the two test models. The two most important features to observe in Figure 5.13 are the doubling up of the reconstructed regions for each model, and the strong positive correlation between reconstructed neutralino and slepton masses. While it is reassuring to see that sparticle masses stand a “fair” chance of being reconstructed in a region close to their input values, it is evidently the case that there are comparably sized regions which are equally consistent with the edge measurements, but which are not centred on the desired values.

To state the problem more precisely, for each model there appear to be two slepton/neutralino mass differences which are consistent with the measured edge positions. Both mass differences are measured with an error which is tiny in comparison with the error on the overall mass scale, but there is no way of choosing between the two possibilities. It is not surprising that mass differences are better constrained than absolute masses, since the energies available to the visible particles radiated in decay chains are dependent primarily on the mass differences between the invisible sparticles. The unfortunate twofold ambiguity concerning the choice of the *correct* slepton/neutralino mass difference arises as a direct consequence of the generalised treatment of the l^+l^-q edge and $l^\pm q$ high-edge. In particular, the formula for the $l^\pm q$ high-edge (see Table 5.2) places the endpoint at the larger of $m_{l_{\text{near}q}}^{\text{max}}$ and $m_{l_{\text{far}q}}^{\text{max}}$. It is easy to show from their definitions that these two quantities are equal when

$$m_l^2 = m_{\tilde{\chi}_1^0} m_{\tilde{\chi}_2^0}, \quad (5.29)$$

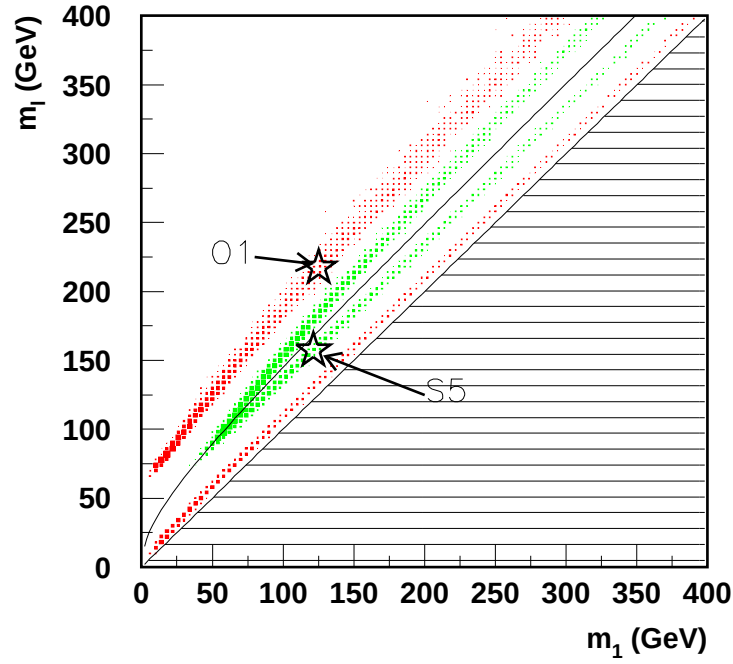


Figure 5.14: The spread of $\tilde{l}_R$ and $\tilde{\chi}_1^0$ masses reconstructed by an ensemble of simulated ATLAS experiments, each using a cut-down version of the “basic” analysis (see Section 5.11) in which no use was made of the l^+l^-q threshold. Note how the overall mass scale here is constrained much less than in Figure 5.13. Colours as in Figure 5.13. The thin black curve denotes the “fault line” defined by Equation (5.31) and described in the text. The region forbidden by the mass hierarchy assumed in Equation (5.6) is shown hatched.

and so (5.29) describes a surface in sparticle-mass space separating regions in which the edge positions have different functional forms. Since (to first order) we can treat most of the edges as constraints of mass differences, we can translate (5.29) into a line in “reconstructed $(m_{\tilde{\chi}_1^0}, m_{\tilde{l}})$ -space” by making the following approximation among *reconstructed* masses

$$m_{\tilde{\chi}_2^0} = \delta m_{21} + m_{\tilde{\chi}_1^0} \quad (5.30)$$

in which $\delta m_{21} = m_{\tilde{\chi}_2^0} - m_{\tilde{\chi}_1^0}$ stands for the *true* neutralino mass difference. Having substituted this into (5.29) we obtain

$$m_{\tilde{l}} = \sqrt{m_{\tilde{\chi}_1^0}(\delta m_{21} + m_{\tilde{\chi}_1^0})} \quad (5.31)$$

which we can expect to describe a “fault line” in the reconstructed sparticle mass space across which ambiguities can be expected to arise. This line is plotted in Figure 5.13 (and subsequent scatter plots) and is indeed seen to run between the pairs of ambiguous regions. It is clear that more work must now be done to constrain the masses better.

The role of the l^+l^-q threshold

Before demonstrating an improved analysis, it is interesting to note that without the l^+l^-q threshold the situation would be even worse, as shown in Figure 5.14. In this figure one can see that, in addition to false solutions being present, the absolute sparticle masses are not even bounded above. Indeed, only at S5 could there even be said to be *lower* bounds for the sparticle masses. It should be noted that in this “cut-down” version of the basic analysis, the number of edge constraints is equal to the number of free sparticle masses (four), so independent constraints should leave the masses entirely determined (no degrees of freedom). That this does not happen, is proof that the constraints coming from the edges are not completely independent. To first order they

are principally constraints on the mass differences, and they only constrain the overall mass scale at second order. It is thus possible to confirm one of the claims of [31–33], that the introduction of the l^+l^-q threshold is crucial in fixing the overall mass scale.

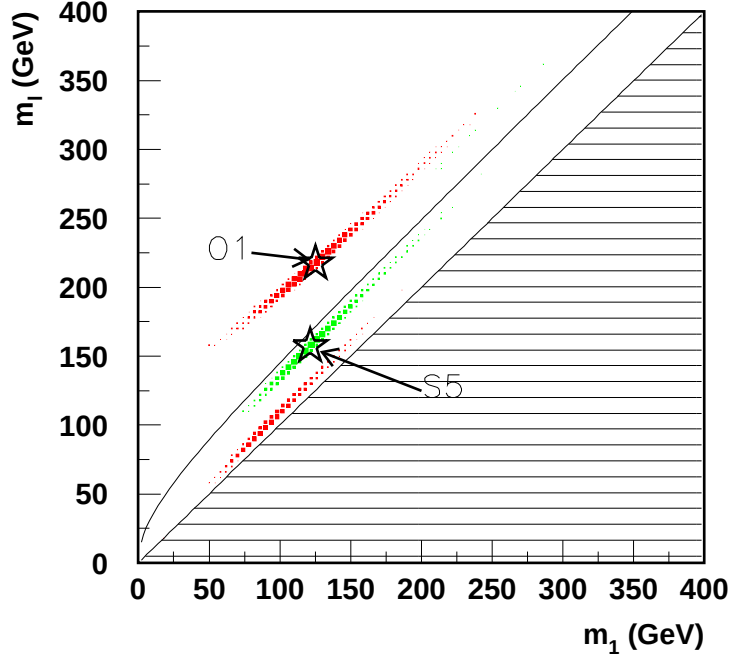


Figure 5.15: The spread of $\tilde{l}_R$ and $\tilde{\chi}_1^0$ masses reconstructed by an ensemble of simulated ATLAS experiments, each using the “extended” analysis (see Section 5.12). Colours as in Figure 5.13. The thin black curve denotes the “fault line” defined by Equation (5.31) and described in the text. The region forbidden by the mass hierarchy assumed in Equation (5.6) is shown hatched.

5.12 Extended reconstruction

An “extended analysis” may be defined by incorporating the novel $l^\pm q$ low-edge within the basic analysis of the previous section. The extended analysis thus makes maximal use of the boost invariant edges described in this thesis.

Figure 5.15 shows the reconstruction regions obtained after an application of the extended analysis at the two test points. As before, there remains a strong correlation between reconstructed slepton and neutralino masses. However, the reconstruction is clearly much improved by the inclusion of the new edge. Not only is the overall mass scale now much better constrained (the regions are smaller) but crucially, the undesired high

slepton-mass region has disappeared entirely at S5. At O1 the improvement consists only of an improvement in resolution,^j as the undesired low slepton-mass region still remains.

Although the results are presented in terms of an *ensemble* of experiments, a *single* experiment observing edges like those of S5 would also find only one region of the $(m_{\tilde{\chi}_1^0}, m_{\tilde{l}_R})$ -plane consistent with the edge positions. Effectively this means that, for S5, the analysis can stop at this point. This is not to say that the smallest possible sparticle mass errors have been achieved yet; the point is only that the probability distributions for reconstructed masses are now approximately gaussians (and centred on the correct values) rather than spread over two disconnected regions. Indeed, it will shortly be seen that further analysis at S5 (using $\Delta M^{\max}$) can in fact produce a valuable new result. Nevertheless, the difference in character between the “single consistent region” of S5 and the “pair of consistent regions” remaining at O1, mean that the way subsequent results are *interpreted* will differ in the two cases. It will be possible to make stronger claims when faced with edges like those of S5, where the problem is already better constrained.

In order to remove the ambiguities associated with the correct evaluation of $m_{\tilde{l}_R} - m_{\tilde{\chi}_1^0}$, it is necessary to introduce the $\Delta M^{\max}$ observable. This will take place in Section 5.13.

It is interesting to look at the effect which the steady development of the method has had on the sparticle mass reconstructions themselves. Figures 5.16 and 5.17 show how the distributions of reconstructed sparticle masses (in S5 and O1 respectively) have improved as each new edge has been included in the reconstruction process. All distributions have been normalised to unit area. The scatter plots of the “extended” analysis at S5 have shown that there the reconstructed masses are distributed about their true values, and so Figure 5.16(c) may be interpreted as displaying the probability distributions for reconstructed masses which would be expected in an S5-like situation

^jand a more subtle reduction in skewness described later in this section

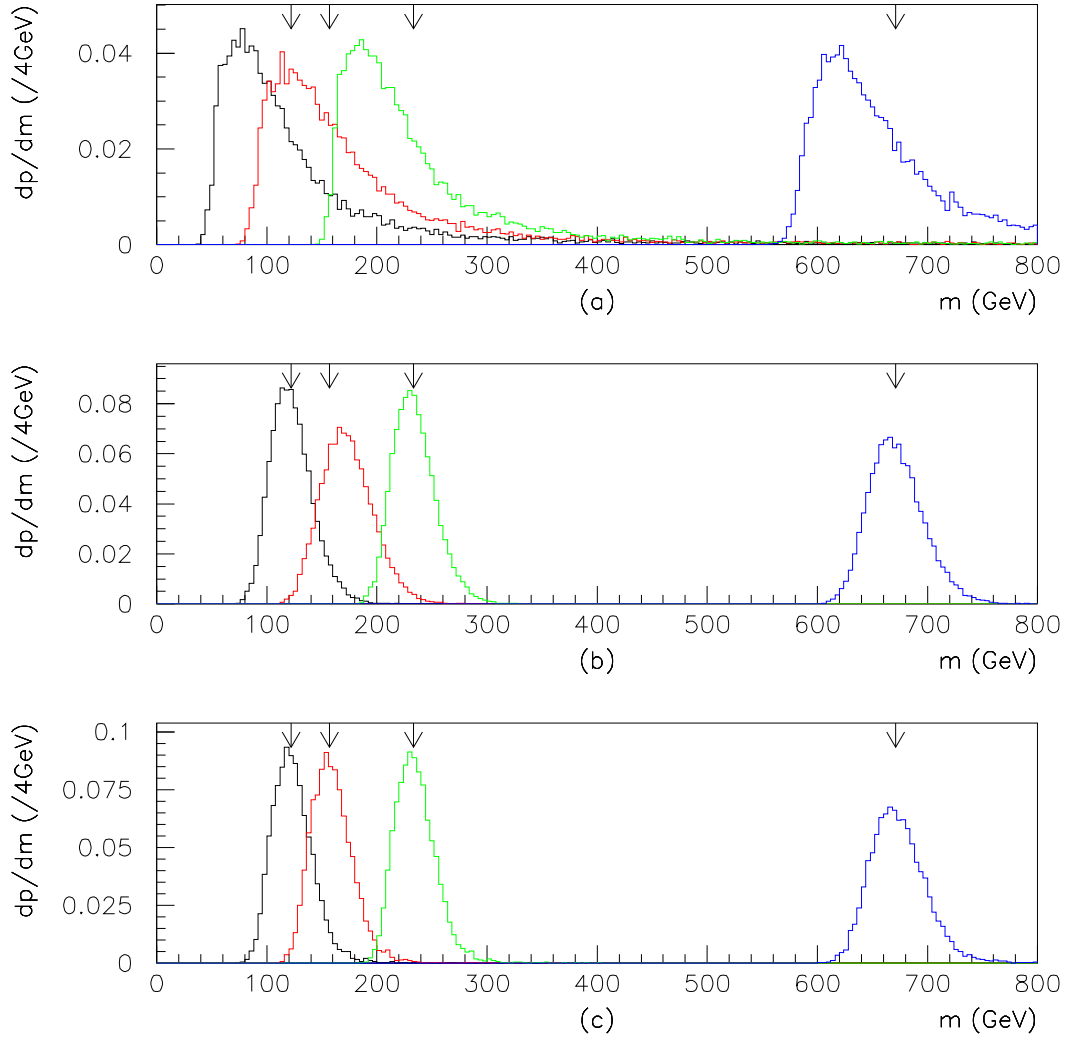


Figure 5.16: The plots above illustrate how the sparticle mass reconstructions improve with the addition of each new edge at $S5$. The distributions in (a) are generated with the “cut-down” analysis, in (b) with the “basic” analysis, and in (c) with the “extended” analysis (see definitions in the text). The histograms display (from left to right) the distributions of reconstructed $\tilde{\chi}_1^0$ (black), $\tilde{l}$ (red), $\tilde{\chi}_2^0$ (green) and $\tilde{q}$ (blue) masses obtained from an ensemble of atlas experiments. In the slepton and lightest neutralino sector, the above plots are just projections of the scatter plots in Figures 5.14, 5.13 and 5.15 respectively, and so contain no new information, although the widths of the distributions are now more readily apparent. Small arrows indicate the sparticle masses of $S5$ which were used as the starting point for the simulated reconstructions. Note that the projections lose information connected with correlations between the reconstructed masses which were evident in the scatter plots. Distributions have been normalised to unit area, but see note on interpretation as probability distributions in the text.

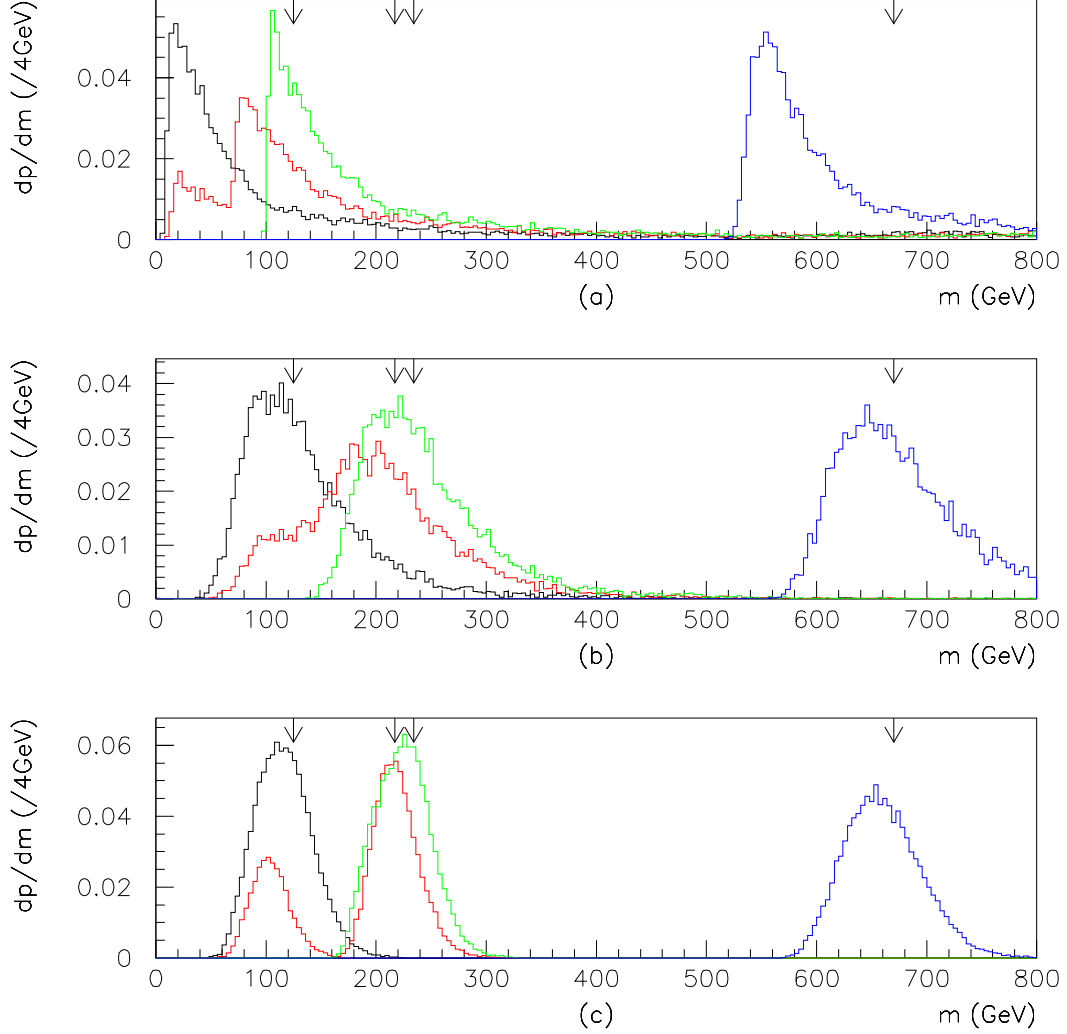


Figure 5.17: The plots above illustrate how the sparticle mass reconstructions improve with the addition of each new edge at $O1$. The distributions in (a) are generated with the “cut-down” analysis, in (b) with the “basic” analysis, and in (c) with the “extended” analysis (see definitions in the text). The histograms display (from left to right) the distributions of reconstructed $\tilde{\chi}_1^0$ (black), $\tilde{l}$ (red), $\tilde{\chi}_2^0$ (green) and $\tilde{q}$ (blue) masses obtained from an ensemble of atlas experiments. In the slepton and lightest neutralino sector, the above plots are just projections of the scatter plots in Figures 5.14, 5.13 and 5.15 respectively, and so contain no new information, although the widths of the distributions are now more readily apparent. Small arrows indicate the sparticle masses of $O1$ which were used as the starting point for the simulated reconstructions. Note that the projections lose information connected with correlations between the reconstructed masses which were evident in the scatter plots. Distributions have been normalised to unit area, but see note on interpretation as probability distributions in the text.

in which statistical and energy scale effects dominate the reconstruction error. At O1 (Figure 5.17) and in the simpler analyses at S5 (Figures 5.16(a) and (b)), however, the presence of two equally good reconstruction regions, evidenced by the scatter plots, makes the “probability distribution interpretation” less applicable in these cases. For example, no significant information is carried by the relative heights of the two slepton peaks in Figure 5.17(c). That the slepton mass peak at higher masses contains more experiments than that at lower masses is just an artefact of the initial conditions of the fit and the method of fitting chosen.^k Accordingly, the skewness of the other mass peaks in Figure 5.17(c) is also dependent on the same artefact of the fitting process. Figure 5.18 will shortly be introduced with the intention of showing results that *are* independent of these artefacts, but before getting to that stage it is still worth noting that: the fact that the non slepton-mass distributions in Figure 5.17 do not exhibit the “dual peak” structure displayed by the corresponding slepton mass distribution, indicates that their central values are largely insensitive (at the level of the width of the squark distribution, for example) to the ambiguity in the choice of slepton mass. It is for this reason that these combined plots have been shown.

Figure 5.18 takes a closer look at the distributions shown in Figures 5.17(b) and (c) to provide support for some of the assertions of the preceding paragraph. The principal difference between Figures 5.17 and 5.18 is that the latter splits the reconstructed distributions up into two parts; those *below* the “fault line” described by Equation (5.31) (corresponding to the “light” solutions), and those *above* the “fault line” (corresponding to “heavy” solutions). Figure 5.18(a) shows this decomposition for the reconstructed masses of the “basic analysis”. The two sets of distributions have been overlaid on top of each other. The main thing to observe here is that the distributions for the neutralinos and the squark-scale *really are* independent of whether the solution is light or heavy.

^kEffectively it just means that the volume of parameter space draining into the lower mass solutions was smaller than that draining into the higher mass solutions.

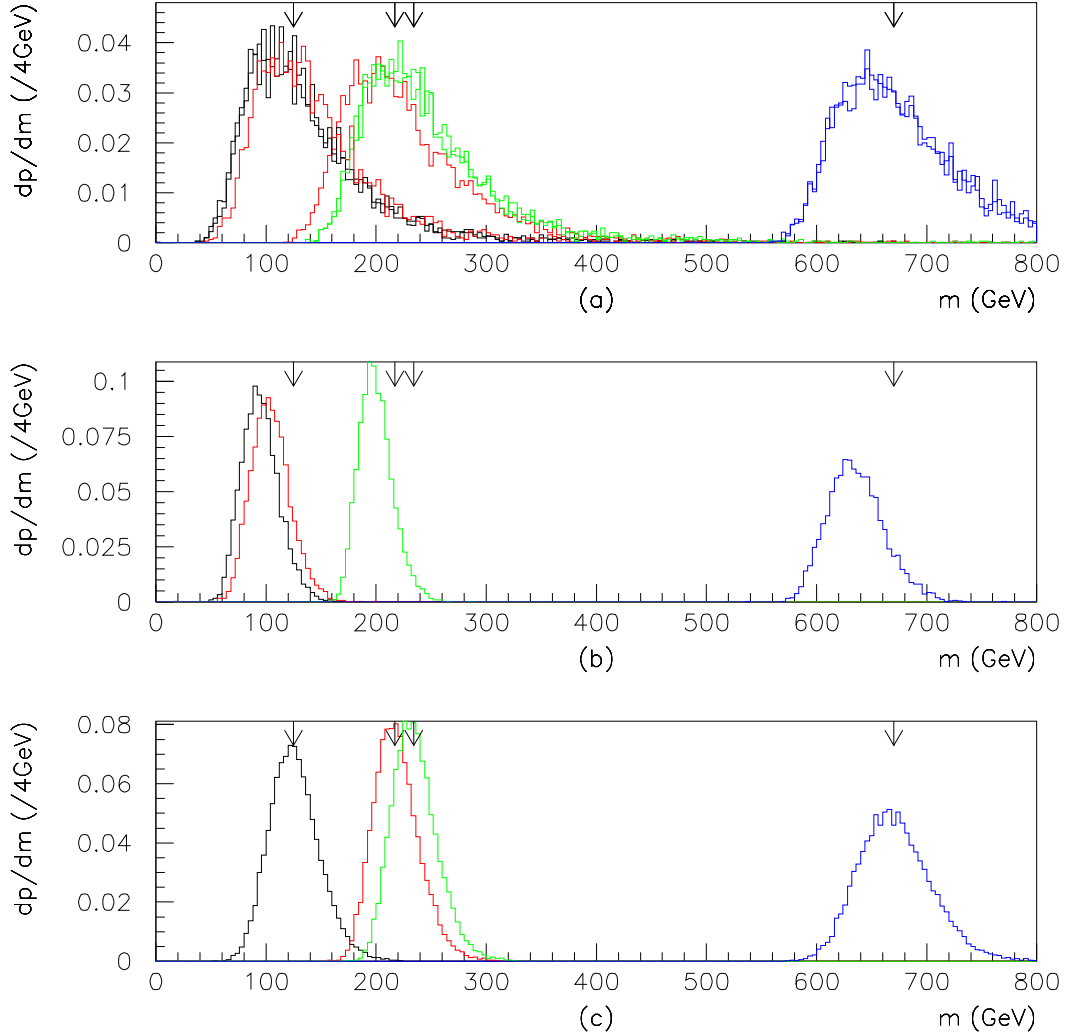


Figure 5.18: This figure shows a breakdown of the sparticle mass distributions shown in Figures 5.17(b) and (c). In this figure, the reconstructed distributions are split up into two parts; those below the “fault line” (described by Equation (5.31)), and those above it. Figure (a) shows the results of the “basic analysis” at O1. Both components of the mass distributions may be seen overlaid on each other. Figures (b) and (c) show the same data, but this time for the “extended analysis”. In Figures (b) and (c) the two components of the mass distributions have been plotted separately; the component from “below the fault line” is shown in (b), and the component from “above the fault line” is shown in (c). Each component of each distribution has been normalised to unit area. The small arrows and colours have the same meanings as in Figure 5.17.

Only the slepton-mass distributions are seen to split. Note that while such independence would in general be a good thing, the significant skewness of the mass distributions in (a) and their failure to centre on the desired values (indicated by the tick marks) leave much to be desired. Contrast this situation with that of the “extended analysis” shown in Figures 5.18(b) and (c). (This time the light and heavy distributions have been plotted separately for clarity.) Here, the reconstructed mass distribution *for any particle* has a different form in the (false) “light” solution (b) from that which is has in the (true) “heavy” solution (c). The main difference between the distributions of the two solutions is that those of the light solution *all* have central values below the true answers (indicated by the tick marks) whereas those corresponding to the heavy solution do not. In particular the distributions of the neutralino masses and the squark mass-scale are some 20 GeV below their preferred positions. This situation is very different from that in the “basic analysis” seen in (a) where the only distributions seen to differ in the two solutions were those of the slepton masses.

This shows us that at the “extended analysis” is capable of making a very particular improvement over the “basic analysis” (i.e. an improvement beyond the more readily apparent fact that the widths of any of its reconstructed mass distributions are smaller than those in the preceding analyses). The “extended analysis” is actually able to take the skewed and off-centre sparticle mass distributions from the “basic analysis”, and turn them (in the case of the correct solution) into relatively unskewed distributions centred on the correct values.

The price which has been paid for this improvement is that, instead of having just one solution for every non-slepton sparticle mass, we now have two; the lower mass corresponding to the “light slepton” solution, and the larger mass to the “heavy slepton” solution. But this price need not necessarily be regarded as unwelcome, for at least the analysis is now producing progressively more specific statements about the sparticle

model	sparticle	m_{true}	$\overline{m}_{recon}$	$m_{true} - \overline{m}_{recon}$	RMS: $\sqrt{(m_{recon} - \overline{m}_{recon})^2}$
S5	$\tilde{\chi}_1^0$	121.52	123.73	2.2	19
S5	$\tilde{l}_R$	157.21	159.52	2.3	20
S5	$\tilde{\chi}_2^0$	233.02	235.27	2.2	19
S5	$\tilde{q}$	670.68	673.34	2.7	25
O1	$\tilde{\chi}_1^0$	124.90	116.66 126.19	- 8.2 1.3	26 23
O1	$\tilde{l}_R$	217.35	183.67 218.87	-33.7 1.5	56 21
O1	$\tilde{\chi}_2^0$	233.77	224.46 235.31	- 9.3 1.5	25 21
O1	$\tilde{q}$	669.74	660.56 671.43	- 9.2 1.7	35 33

Table 5.5: *This table contains statistics which summarise the performance of the mass reconstructions achieved by the “extended analysis” at S5 and O1. The statistics are derived from the S5 reconstructed mass distributions shown in Figures 5.16(c), the O1 reconstructed mass distributions shown in 5.17(c) and the O1 “above the fault line” reconstructed mass distributions shown in 5.18(c). The statistics for the “above the fault line” distributions at O1 are shown in bold. The numbers listed comprise the following: The first column of numbers lists the “true” sparticle masses which were supplied as inputs to the algorithm which simulated the ensemble of experiments. These numbers correspond to the locations of the small arrows in the figures mentioned above. The second column of numbers lists, for each sparticle, the mean value of the reconstructed mass distribution. The third column of numbers lists the difference in the two previous quantities. The final column of numbers contains the root mean square deviation of the reconstructed masses from their mean values. All values are in GeV.*

masses, increasing the prospects for any subsequent attempt to rule out one of the solutions using other methods. This point will be touched upon in the next section.

Table 5.5 lists statistics (derived from Figures 5.16(c), 5.17(c) and 5.18(c)) which summarise the performance of the “extended analysis” at S5 and O1. These show that the statistical errors on the edge fits and the uncertainty on the overall jet energy scale are together likely to limit the precision of the sparticle mass measurement made by the “extended analysis” to around 19 to 25 GeV at S5 and 21 to 33 GeV at O1.

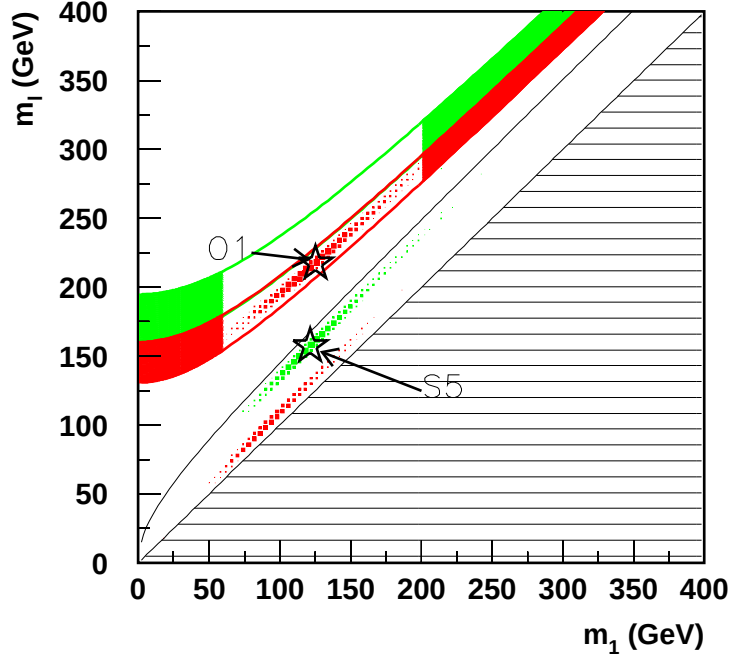


Figure 5.19: *The ΔM constraint. This figure is identical to Figure 5.15, except that it also shows the one-sigma regions favoured by 10% measurement of the “principal ΔM edge” in each of the two models. A window has been left in the centre section of each of the one-sigma regions so that they do not obscure the underlying scatter points. Colours as in Figure 5.13. The thin black curve denotes the “fault line” defined by Equation (5.31) and described in the text. The region forbidden by the mass hierarchy assumed in Equation (5.6) is shown hatched.*

5.13 Reconstruction using $\Delta M^{\max}$

Up to this point, edges have made definite (albeit ambiguous) statements about $m_{\tilde{\chi}_1^0}$, $m_{\tilde{\chi}_2^0}$, $m_{\tilde{l}_R}$ and the squark mass scale $m_{\tilde{q}}$. As we have already seen, the ΔM edge is not so kind to us. While we have assumed we might measure the position of the principal ΔM edge to $\sim 10\%$, we cannot be sure (from the edge alone) whether it is due to light ($\tilde{l}_R$) or heavy ($\tilde{l}_L$) sleptons. This can be both a blessing and a curse.

The best situation to find oneself in is one in which “other” measurements have already determined $m_{\tilde{l}_R} - m_{\tilde{\chi}_1^0}$ to some precision, while at the same time the ΔM edge is

seen to favour an even larger lepton/neutralino mass difference. In these circumstances one has a strong claim to have obtained exciting and compelling evidence for a new heavier slepton (which we will identify with $\tilde{l}_L$) whose relation in mass to the mass of the LSP may be inferred from the location of the ΔM edge. We are fortunate enough to find ourselves in this position at S5. This may be seen in Figure 5.19, which shows the scatter points from the “extended analysis” (of the last section) overlaid with the region favoured by a 10% measurement of the “principal ΔM edge”. The region favoured by the ΔM measurement is seen to lie many standard deviations above the region determined by the preceding analysis. Far from signalling inconsistency, this information is welcomed as a sign that the ΔM measurement has managed to measure the mass of the as-yet-undiscovered left-slepton. The mass it would attribute to the left-slepton is given by $m_{\tilde{l}_L} = \Delta M^{\text{max}} + \sqrt{m_{\tilde{\chi}_1^0}^2 + (\Delta M^{\text{max}})^2}$. The true value of ΔM^{max} (abbreviated by M in what follows) for the principal edge at S5 is 88.5 GeV. To first order in $M/m_{\tilde{\chi}_1^0}$, then, the error on the heavy-slepton/neutralino mass difference would be given by

$$\frac{\Delta(m_{\tilde{l}_L} - m_{\tilde{\chi}_1^0})}{m_{\tilde{l}_L} - m_{\tilde{\chi}_1^0}} = \left(1 + \frac{M}{2m_{\tilde{\chi}_1^0}}\right) \left(\frac{\Delta M}{M}\right) \oplus \left(\frac{M}{2m_{\tilde{\chi}_1^0}}\right) \left(\frac{\Delta m_{\tilde{\chi}_1^0}}{m_{\tilde{\chi}_1^0}}\right) \quad (5.32)$$

$$\approx 1.4 \left(\frac{\Delta M}{M}\right) \oplus 0.4 \left(\frac{\Delta m_{\tilde{\chi}_1^0}}{m_{\tilde{\chi}_1^0}}\right). \quad (5.33)$$

This suggests that, were it possible to measure $m_{\tilde{\chi}_1^0}$ to 25%, an error of greater than 7% on ΔM^{max} would dominate the error on the mass difference when only these two sources are considered.

At O1 the situation is somewhat different. Firstly, as has already been mentioned, the extended analysis generates two regions in the sparticle mass plane which are consistent with the observed edges at O1. Secondly, the masses favoured by the ΔM edge at O1 are not “inconsistent” with the results of the earlier analysis (as was the case at S5). On the contrary, they actually overlap the upper of the two remaining regions, as Figure 5.19 shows. Being in the privileged position of knowing what the montecarlo was instructed

to generate, the coincidence does not surprise us; after all, *we* know that the lower region is the “fake” while the upper region and the ΔM constraint are both measurements of $m_{\tilde{l}_R}$. Putting this knowledge to one side, however, what is our position?

If the ΔM constraint had fallen *between* the two regions we could have immediately accredited it with having both measured the mass of a heavy slepton, while at the same time settling the ambiguity associated with the interpretation of the two regions. Even if the ΔM edge constraint had fallen above the upper region, we could still have claimed a measurement of a heavy slepton’s mass, even if the ambiguity concerning the lighter slepton’s mass remained unsolved. However, the place at which the ΔM constraint *does* fall at O1 is the most tantalizing. It tells us only one thing outright, namely that there is at least one slepton with a mass lying in the region covered by the ΔM constraint. What we do not learn is:

- whether the consistency with the upper region entitles to rule out the lower region (thus letting us claim a “double” measurement of the light slepton mass, with consequently improved errors);
- whether the consistency is just an accident; a “chance coincidence” between “a ΔM constraint favouring a heavy slepton mass” and “an unphysical region of parameter space which was regrettably not ruled out by the earlier analysis”. In this situation we would regard the lower O1 region as a measurement of the light-slepton mass; or
- whether the light and heavy sleptons might both be degenerate in mass (or nearly so) at the value indicated by the upper region and the ΔM constraint, taken together.

The situation at O1 is clearly not “pretty” but it may at least be summarised in a positive way: the good news is that at O1 we either obtain an *improved* measurement of

the light slepton mass (while learning nothing about heavier sleptons) or we are able to claim measurement of both light *and* heavy slepton masses. The catch is that the edges alone (at the given resolutions) cannot distinguish between these two possibilities.

What methods might we use to resolve the interpretation issues arising in connection with the slepton masses at O1? The first and most obvious suggestion is to try to use information from cross sections. Since we need to resolve problems with slepton masses, dislepton production would appear to be the production process we need to study. But unfortunately, the cuts which the counting of dislepton events require are basically those used to obtain events for the ΔM edge. We have already seen that these cuts are fine for counting events containing sleptons which are heavy enough to produce ΔM values that may be seen above the SM backgrounds, but these are precisely the events we are *not* interested in. The principal ΔM edge already gives us an estimate of such slepton masses which is likely to be more reliable (less model dependent) than one derived from a counting argument using the same events. Our problem is not that we do not know the value of *this* slepton mass, but that at O1 we are unsure about whether there is a *lighter* slepton at the mass suggested by the lower of O1's regions in Figure 5.19. The ability to use cross sectional information to detect the existence of such lighter sleptons requires cuts which are able to pull out a signal of the size and shape of S5's light slepton events (Figure 5.9(a) and (b)) from underneath the SM backgrounds which are two orders of magnitude larger (Figure 5.10), and there is no evidence to suggest that such cuts are likely to be forthcoming.

If cross sections are to be abandoned in the search for means to remove the remaining issues at O1 we are left with branching ratios, the use of new decay chains, and the creation of new observables for existing chains. All these options will be left for future investigations, although it is remarked that the possibilities for forming new "artificial" edges (like the l^+l^-q threshold, which we recall is but one of an otherwise unexplored

family of edges governed by the choice of lower limit for m_{ll}) have by no means been exhausted yet.

Chapter 6

Conclusions

If supersymmetry is to play a vital role in extending the SM, then at the next generation of hadron colliders, the ability to measure sparticle masses impartially will be paramount. This thesis has extended our understanding of how to perform, and interpret the results of, model independent sparticle mass measurements based on the endpoints of kinematic variables. This study has been performed in the context of the ATLAS experiment. New kinematic variables have been proposed, and the importance of avoiding dangerous and unnecessary assumptions about sparticle mass hierarchies has been stressed. To test the degree of model independence and assess the strengths and weaknesses of the approach used, a single analysis was applied to ensembles of ATLAS experiments simulated in two contexts; one in which the physics was that of a standard mSUGRA point (S5), and another in which nature was described by an Optimised String Model (O1). These models differed principally in their slepton sectors, with the sleptons of O1 heavier than those of S5. The performance of the analysis was found to be broadly the same in both analyses, with by far the best constrained quantities being sparticle mass differences. The different sparticle masses in the two models also made clear the ways in which the results of the analysis could differ between the two models. It was discovered that,

while all the reconstructed masses at S5 were clustered about the correct values, at O1 the generalised nature of the approach resulted in an unavoidable ambiguity in the reconstructed slepton mass; with two disconnected regions of sparticle-mass space being consistent with the observables. Nevertheless, it was shown that one of these regions was indeed centred on the correct masses, and that the sparticle mass resolutions obtained there were consistent with those found at S5. The root mean square widths of the sparticle mass distributions at S5 (O1) were found to be ~ 20 GeV (23 GeV) for the non strongly-interacting sparticles, and ~ 25 GeV (~ 33 GeV) for the squarks. These widths take into account the uncertainties in the measured endpoints arising from the fall-off of statistics in the vicinity of the edges (which limits the precision of any fit to the endpoint position), and also approximately take into account the effect of a systematic 1% uncertainty in the jet energy scale.

No account is yet taken of uncertainties in the endpoints arising from the presence of unsubtracted supersymmetric backgrounds in the vicinity of each edge, and this represents an important avenue for further study. The importance of assessing the relationship between the extracted squark mass-scale and the squark masses in the models themselves has also been mentioned as an important “theoretical” systematic to control. The large number of variables affecting these sorts of considerations makes it unlikely that, at this early stage, all these effects could be fruitfully taken into account. Nevertheless, they are potentially important, and their future study in the light of real data will not only be imperative, but hopefully also tractable. The absence of the modelling of these effects within this thesis means that the reported errors represent limits, below which the resolution of the (unmodified) method is unlikely to improve.^a

It was shown that a new kinematic variable ($\Delta M^{\max}$), if added to the “extended analysis”, was able to offer valuable new information concerning the slepton/neutralino

^aIt is noted that an addition, in quadrature, of 1% to the fractional fit errors on all boost-invariant edge positions results in an overall increase of the reconstructed widths by a factor of 2.5.

mass differences. While not *proving* which of the two solutions at O1 was the true one, it was certainly able to lend strong “circumstantial” support to the correct one. Indeed, at S5, where ΔM^{max} was not strictly needed, it turned out to provide a “bonus” measurement of the mass of the left-slepton l_L . Suggestions of ways that the analysis could be extended in the future have been made, and it looks likely that this technique will have a role to play at ATLAS in the years after turn on.

Appendix A

The Supersymmetry Scale

When any quantum field theory is renormalized, corrections to the bare masses of particles must be calculated by the inclusion of loop diagrams into their propagators. In the case of the higgs particle the corrections

$$\delta M_h^2 = M_h^2 - M_{h^0}^2 \quad (\text{A.1})$$

can often be very large indeed. However in supersymmetric theories it is possible to consider the loop diagrams in pairs, such that the looping particle in one is the superpartner of the looping particle in the other. If this is done, then the order of the correction from any such pair of diagrams is (see [6]) only of order

$$\delta M_h^2 \approx \frac{m_{\text{soft}}^2}{16\pi^2} \ln(\Lambda/m_{\text{soft}}) \quad (\text{A.2})$$

where m_{soft} is of order of the mass difference of the two superpartners in the loop and Λ is the loop momentum cutoff. Figure A.1 shows the variation of δM_h^2 with m_{soft} for GUT-scale and planck-scale values of Λ . The plot shows that, with only very slight dependence on the scale of the cutoff, the loop corrections become commensurate with

reasonable electroweak scales around $m_{\text{soft}} = 200$ GeV and then proceed to grow beyond the most generous electroweak scales from around $m_{\text{soft}} = 700$ GeV. It is for this reason the scale for supersymmetry is often said to lie around 1 TeV. At values of m_{soft} much in excess of this, obtaining the correct value for the renormalized higgs mass would require ever finer tunings of the bare higgs mass, and in doing so remove one of the reasons for wanting to introduce supersymmetry in the first place.

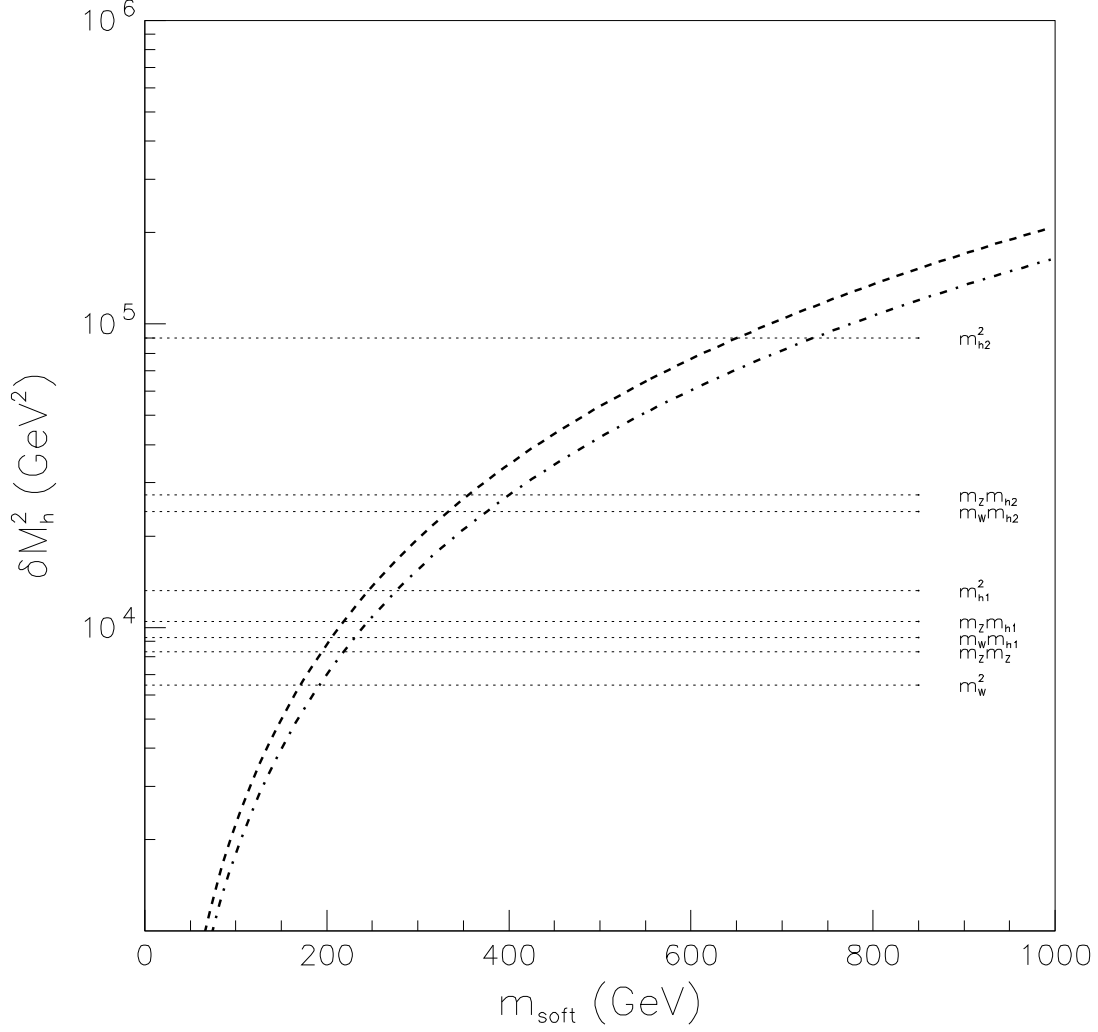


Figure A.1: Approximate size δM_h^2 of loop corrections to the higgs mass-squared M_h^2 for a supersymmetric theory in which the largest mass splitting between superpartners is order m_{soft} . Two curves of δM_h^2 are shown for different values of the loop momentum cutoff Λ . In the upper (dashed) curve $\Lambda = 10^{17}$ GeV, while in the lower (dash-dot) curve $\Lambda = 10^{14}$ GeV. Horizontal (dotted) lines indicate a selection of electroweak mass-squared scales for comparison purposes. $m_{h1} = 115$ GeV and $m_{h2} = 300$ GeV were chosen to be suggestive of the present lower and upper bounds of the SM higgs mass.

Appendix B

Higgs doublet multiplicities

Where the SM has only one higgs doublet field, the MSSM has two. There are effectively two reasons for this. The first is a technical result which comes from theory (see [6] for a full account) while the second is more down to earth.

The theory result establishes that the most general set of renormalizable non-gauge interactions for a supersymmetric theory can always be written down in a lagrangian of the form

$$\mathcal{L}_{\text{int}} = F(\Phi, W(\Phi)) + \text{h.c.}, \quad (\text{B.1})$$

where the so-called “superpotential”, W , is an analytic function of the supermultiplets $\Phi = \{\Phi_i\}$ of the theory, and where F is a prescription for turning the superpotential and the supermultiplets into interaction terms. The key point is that W *must* be analytic in the fields it is a function of, and so if W is a function of a particular supermultiplet field Φ_i , then it cannot also be a function of the complex (or Hermitian) conjugate field Φ_i^* .

Next it will be shown that it is not possible to write down a complete set of SM mass terms of the form $\mathcal{L}_{\text{mass}} = \mathcal{L}_{\text{half}} + \text{h.c.}$ in which $\mathcal{L}_{\text{half}}$ does not contain, at least once, both

a field and its complex (or Hermitian) conjugate. The implication of this for the MSSM, will be that the successful supersymmetrization of the SM depends on the removal of this ‘duplication of conjugated fields’ by the replacement of one of each duplicated pair with a new independent superfield carrying the same quantum numbers.

Were it not for the fact that they would break gauge invariance, the mass terms of the SM would *ideally* be incorporated directly into the lagrangian of the unbroken SM in the form $m_q(\bar{q}_L q_R + \bar{q}_R q_L)$, or equivalently $m_q \bar{q}_L q_R + \text{h.c.}$. Such a scheme, were it possible, would have a certain appeal through avoiding the need to introduce a new particle (the higgs) into the theory to account for fermion masses. But for better or for worse, gauge invariance requires that the SM fermion mass terms are instead those of (2.3).

Considering for the moment just the first generation of fermions, (2.3) requires that all the possible hypercharge-neutral combinations of fermion fields $\psi \in \{L, e, Q, u, d\}$ be substituted into

$$\bar{\psi}_{f,L} g_{ff'} \phi \psi_{f',R} + \bar{\psi}_{f,L} g'_{ff'} i \sigma_2 \phi^* \psi_{f',R} + \text{h.c.} \quad (\text{B.2})$$

Recalling the hypercharge assignments $\bar{Q} = -\frac{1}{3}$, $u = +\frac{4}{3}$, $d = -\frac{2}{3}$, $\bar{L} = +1$, $e = -2$, $\phi = +1$ and $\phi^* = -1$ we find that only the combinations

$$\{\bar{Q} i \sigma_2 \phi^* u, \bar{Q} \phi d, \bar{L} \phi e\} + \text{h.c.} \quad (\text{B.3})$$

are allowed, while all of the remainder:

$$\{\bar{L} \phi d, \bar{L} \phi u, \bar{L} i \sigma_2 \phi^* d, \bar{L} i \sigma_2 \phi^* u, \bar{L} i \sigma_2 \phi^* e, \bar{Q} \phi u, \bar{Q} \phi e, \bar{Q} i \sigma_2 \phi^* d, \bar{Q} i \sigma_2 \phi^* e\} + \text{h.c.} \quad (\text{B.4})$$

are banned. (Incidentally this has the fortuitous side effect of retaining separate lepton

and baryon number conservation within the SM!) The allowed mass terms (B.3) may be written out as

$$\left\{ \begin{aligned} & g_u^{ij} \left(\begin{array}{cc} \bar{u}_L & \bar{d}_L \end{array} \right)_i \cdot i \left(\begin{array}{cc} 0 & -i \\ i & 0 \end{array} \right) \cdot \left(\begin{array}{c} \phi^- \\ \phi^{0*} \end{array} \right) (u_R)_j \\ & + g_d^{ij} \left(\begin{array}{cc} \bar{u}_L & \bar{d}_L \end{array} \right)_i \left(\begin{array}{c} \phi^+ \\ \phi^0 \end{array} \right) \cdot (d_R)_j \\ & + g_e^{ij} \left(\begin{array}{cc} \bar{\nu}_L & \bar{e}_L \end{array} \right)_i \left(\begin{array}{c} \phi^+ \\ \phi^0 \end{array} \right) \cdot (e_R)_j \end{aligned} \right\} + \text{h.c.} \quad (\text{B.5})$$

which if we approximate ϕ^+ and ϕ^0 by their respective vacuum expectation values 0 and v become

$$g_u^{ij} v^* \bar{u}_L^i u_R^j + g_d^{ij} v \bar{d}_L^i d_R^j + g_e^{ij} v \bar{e}_L^i e_R^j + \text{h.c.}, \quad (\text{B.6})$$

mimicking the banned terms which were required in the first place.

Now that we have established the form of the SM mass terms, we can take stock of the situation. The most important thing to **ignore** is that the “+h.c.” term, which fixes the lagrangian to be real, causes *every* field to be present in both its normal and its Hermitian conjugated forms *somewhere* in the lagrangian. However, we are not interested in the entirety of the mass terms, but only in one half of them. Putting the “h.c.” part to one side, it can be seen that the remaining terms in (B.3), say, are constructed from $\bar{Q}$, u , d , $\bar{L}$, e , ϕ and ϕ^* alone, and we note that the higgs field (only) is repeated. Of course, in place of (B.3) we could have ignored convention and used the equivalent set $\{\bar{u}\phi i\sigma_2 Q, \bar{Q}\phi d, \bar{L}\phi e\} + \text{h.c.}$. This scheme has the advantage that the non-h.c. terms no longer contain occurrences of ϕ^* , but the price for this was that the fields Q and $\bar{u}$ had to be introduced. Whatever alterations we make, we are always faced with some repetition of conjugate fields. We have the choice of duplicating the higgs field ϕ , of duplicating Q and u , or of duplicating Q , d , L and e .

Faced with these choices, the MSSM opts to avoid introducing three as-yet-unobserved

generations of quarks and/or leptons into the theory, and instead creates a new independent higgs doublet to replace $i\sigma_2\phi^*$.

Appendix C

Comparing the supersymmetric and SM Higgs potentials

One of the common objections to the SM is that its Higgs potential (2.1) is somewhat arbitrary. Although the ratio μ^2/λ is determined by the requirement of correct EWSB, the potential provides us with no clue to the magnitudes of μ and λ , and so the higgs mass is largely independent of the rest of the SM.

The supersymmetric Higgs potential is more complicated. The MSSM lagrangian

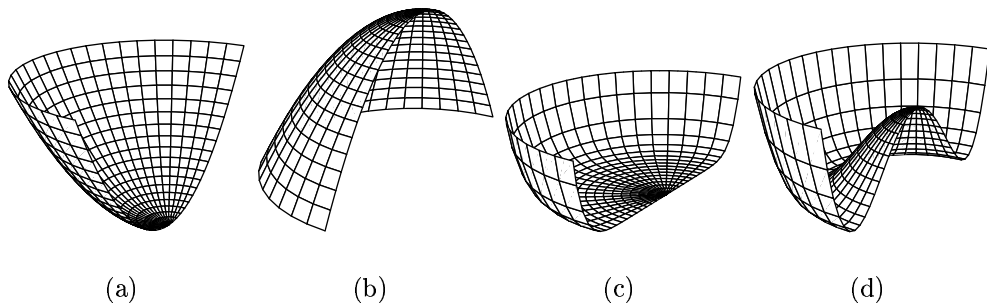


Figure C.1: *Diagrammatic representations of parts of the supersymmetric Higgs potential.*

includes the terms

$$\mathcal{L}_{Higgs\ potential} = \mu^2 H_1^2 + \mu^2 H_2^2 \quad (C.1)$$

$$+ G(H_1^2 - H_2^2)^2 \quad (C.2)$$

$$+ \mu B H_1 H_2 \quad (C.3)$$

$$+ M_{h_1}^2 H_1^2 + M_{h_2}^2 H_2^2 \quad (C.4)$$

$$+ \text{“supertrace terms”}. \quad (C.5)$$

The first of these, (C.1), comes directly from the supersymmetry-respecting term $W_H = \mu H_1 H_2$ in the superpotential. On its own, (C.1) provides a contribution to the MSSM Higgs potential of the quadratic form shown in Figure C.1(a). In principle, such a term could also yield a contribution like that shown in Figure C.1(b), but in practice the RGEs never run μ^2 negative. Term (C.2) is also strongly motivated, arriving as a consequence of the terms which require gauge invariance. It promotes a quartic restraint (Figure C.1(c)) except in the so-called “quartic flat” directions $H_1 = \pm H_2$, however those directions are restrained by the term (C.3) which is one of the soft supersymmetry breaking terms whose inclusion in the lagrangian was motivated by the maxim “*can add; must add*”. So, by this stage, we have a bounded from below (BFB) potential, but probably not a broken one, since μ^2 tends never to run negative. This is where term (C.4) comes in. This term contains the “soft” higgs masses, M_{hi} . In the general MSSM one has complete freedom to assign arbitrary values (positive or negative) to the squares of these masses, and this freedom permits breaking of the supersymmetric Higgs potential. In this way, the soft higgs masses play very much the role of μ in the SM. By choosing their values appropriately, the correct EWSB may be accomplished.

Thus far there is little to choose between the supersymmetric and the SM approaches. However, the same reasons which lead us to believe that the MSSM is likely to be only a low energy approximation to a theory incorporating gravity, lead us to expect an or-

ganising principle behind the parameters of the MSSM. Typically, then, in models like mSUGRA, the soft higgs masses are not given arbitrary values, but are taken to unify with the other scalar masses at the GUT scale. Although this might be aesthetically pleasing, at first sight it appears to jeopardize the success of EWSB which demands imaginary, not real, soft higgs masses. Fortunately, from the GUT scale to the electroweak scale, the soft higgs mass-squareds are *run negative* by the RGEs and are usually able to break electroweak symmetry anyway.^a

The overall effect is that a universal scalar mass M_0 , at the GUT scale, ends up setting not only the mass of the higgses, but also the masses of the sleptons and the squarks. One might argue that this is still not much of an improvement over the SM; after all; there was only one scalar particle in the SM (the higgs) and μ^2 was quite able to specify *its* mass. Nevertheless, the “soft supersymmetry breaking” requirement that supersymmetry not be broken *too* badly (see Appendix A) puts limits on the permitted difference between M_0 and the universal fermion mass, $M_{1/2}$. These limits, in turn, ensure that supersymmetric higgses in universal models are relatively light (in contrast to their highly unconstrained SM counterparts).

^aOn rare occasions, for example in S2, the squares of the soft higgs masses do not run sufficiently negative to counteract the larger positive running of μ^2 , and so in these cases EWSB can only be performed by the supertrace terms (C.5).

Appendix D

Rapidity and Pseudorapidity

Under a lorentz boost $\Lambda^\mu{}_\nu = \begin{pmatrix} \gamma & -\gamma\beta \\ -\gamma\beta & \gamma \end{pmatrix}$ along the z -axis, the four momentum $p^\mu = \begin{pmatrix} E \\ p_z \end{pmatrix}$ (with x and y components suppressed) transforms to $p'^\mu = \begin{pmatrix} \gamma(E - \beta p_z) \\ \gamma(p_z - \beta E) \end{pmatrix}$. As a result, the rapidity y of p^μ defined by $y = \frac{1}{2} \log \left[\frac{E+p_z}{E-p_z} \right]$ transforms to

$$y' = \frac{1}{2} \log \left[\frac{\gamma(E - \beta p_z + p_z - \beta E)}{\gamma(E - \beta p_z - p_z + \beta E)} \right] \quad (\text{D.1})$$

$$= \frac{1}{2} \log \left[\frac{(1 - \beta)(E + p_z)}{(1 + \beta)(E - p_z)} \right] \quad (\text{D.2})$$

$$= y + \frac{1}{2} \log \left[\frac{1 - \beta}{1 + \beta} \right]. \quad (\text{D.3})$$

Thus, under a boost along the z -axis, a set of rapidities $\{y_1, y_2, \dots\}$ will be subject to a boost-dependent translation, but their separations will remain the same. It is the invariance of rapidity differences under lorentz boosts which makes rapidity a useful quantity at colliders, such as the LHC, where the longitudinal momentum of the collision centre of mass is never known. The factor of $\frac{1}{2}$ in front of the definition of y is merely there to ensure that y and ϕ have the same normalization in the central region.

However, it is much harder to measure E and p_z than the more directly observed

polar angle θ . Fortunately, it is usually the case that $m^2 \ll p_T^2 + p_z^2$, and so the approximation $E \approx \sqrt{p_T^2 + p_z^2}$ allows us to produce an approximation for rapidity:

$$y \approx \frac{1}{2} \log \left[\frac{\sqrt{p_T^2 + p_z^2} + p_z}{\sqrt{p_T^2 + p_z^2} - p_z} \right] \quad (\text{D.4})$$

$$= \frac{1}{2} \log \left[\frac{1 + p_z/\sqrt{p_T^2 + p_z^2}}{1 - p_z/\sqrt{p_T^2 + p_z^2}} \right] \quad (\text{D.5})$$

$$= \frac{1}{2} \log \left[\frac{1 + \cos \theta}{1 - \cos \theta} \right] \quad (\text{D.6})$$

$$= \frac{1}{2} \log \left[\frac{2 \cos^2 \theta/2}{2 \sin^2 \theta/2} \right] \quad (\text{D.7})$$

$$= -\log(\tan \theta/2) \quad (\text{D.8})$$

which is known as the pseudorapidity, η . For its ease of measurability and its approximate transformation properties under boosts, the pseudorapidity is used in place of rapidity as the main experimental variable for describing the longitudinal separation of particles and momenta.

Appendix E

Edges

E.1 Successive two body decays

Most of the edges considered in this thesis come from processes containing successive two-body decays. The simplest such case, two successive two-body decays, is depicted in Figure E.1. To evaluate the maximum invariant mass which may be attained by p and q , it is best to move to the rest frame of b . Having changed frame, the situation is as depicted in Figure E.2. In this frame, it is clear that the only kinematic freedom is the angle θ between p and q . The momenta (and thus the energies) of the particles may

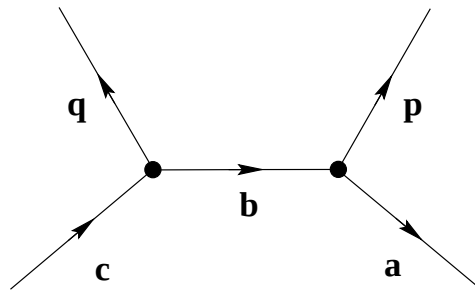


Figure E.1: *Two successive two-body decays are depicted. Arrows indicate particles' directions of motion.*

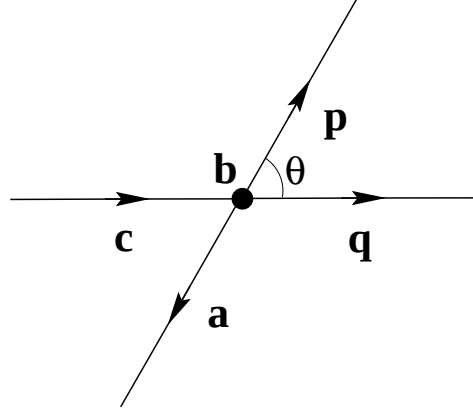


Figure E.2: *The two successive two-body decays of Figure E.2 are depicted in rest frame of b. Arrows indicate particles' directions of motion.*

be obtained by four-momentum conservation, and are given by

$$\mathbf{p}^2 = \mathbf{a}^2 = [m_p^2, m_b^2, m_a^2], \quad (\text{E.1})$$

$$\mathbf{q}^2 = \mathbf{c}^2 = [m_q^2, m_b^2, m_c^2], \quad (\text{E.2})$$

$$E_p^2 = m_p^2 + \mathbf{p}^2 \quad \text{and} \quad (\text{E.3})$$

$$E_q^2 = m_q^2 + \mathbf{q}^2 \quad (\text{E.4})$$

where

$$[x, y, z] \equiv \frac{x^2 + y^2 + z^2 - 2(xy + xz + yz)}{4y}. \quad (\text{E.5})$$

In general, the invariant mass of p and q is given by

$$m_{pq}^2 = (p + q)^2 \quad (\text{E.6})$$

$$= m_p^2 + m_q^2 + 2p^\mu q_\mu, \quad (\text{E.7})$$

$$= m_p^2 + m_q^2 + 2(E_p E_q - |\mathbf{p}| |\mathbf{q}| \cos \theta), \quad (\text{E.8})$$

and thus m_{pq} reaches its maximum or minimum values when p and q are back-to-back or collinear in the rest frame of b . Thus,

$$(m_{pq}^{\max})^2 = m_p^2 + m_q^2 + 2(E_p E_q + |\mathbf{p}||\mathbf{q}|) \quad \text{and} \quad (\text{E.9})$$

$$(m_{pq}^{\min})^2 = m_p^2 + m_q^2 + 2(E_p E_q - |\mathbf{p}||\mathbf{q}|). \quad (\text{E.10})$$

E.2 l^+l^- edge

The location of the l^+l^- edge may thus be obtained from (E.9) by making the substitutions: $p, q \rightarrow l^+l^-$, $c \rightarrow \tilde{\chi}_2^0$, $b \rightarrow \tilde{l}_R$ and $a \rightarrow \tilde{\chi}_1^0$, and by treating the leptons as massless. The result of this substitution may be found in Table 5.2.

E.3 Xq edge

The location of the Xq edge may also be obtained from (E.9) by using the following substitutions: $p \rightarrow \text{quark}$, $q \rightarrow X$, $c \rightarrow \tilde{\chi}_2^0$, $b \rightarrow \tilde{l}_R$ and $a \rightarrow \tilde{\chi}_1^0$, and by treating the quark as massless. The result of this substitution may be found in Table 5.2.

E.4 l^+l^-q edge

The l^+l^-q edge is a consequence of three successive two-body decays as shown in Figure 5.3. In order to find $m_{l\bar{l}q}^{\max}$, it is helpful to treat two of the ejected particles as an “effective particle” with a mass equal to the invariant mass of its constituents. This picture is represented in Figure E.3. From (E.9) and (E.10) we can deduce that m_q may

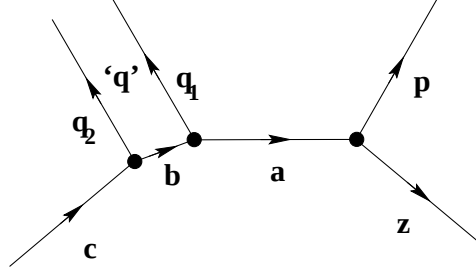


Figure E.3: *Three successive two-body decays are depicted. Arrows indicate particles' directions of motion.*

take values given by

$$m_q^2 \equiv q = \lambda(c-b)(b-a)/b, \quad \lambda \in [0, 1], \quad (\text{E.11})$$

where the notation $q = m_q^2$, $c = m_c^2$, $b = m_b^2$ and $a = m_a^2$ has been introduced. Later, $z = m_z^2$ will also be used.

For a particular value of q , and taking m_p to be zero, the maximum possible invariant mass of p , q_1 and q_2 combined is given, as before, by:

$$m_{q_1 q_2 p}^{\max(q)} = q + 2(p^\mu q_\mu)^{\max} \quad (\text{E.12})$$

which will again happen when p and q are back-to-back in the rest frame of a . A second application of (E.9) shows that

$$m_{q_1 q_2 p}^{\max(q)} = q + \frac{(a-z)}{2a} \left\{ \sqrt{[q - (a+c)]^2 - 4ac} + \sqrt{[q - (c-a)]^2} \right\}. \quad (\text{E.13})$$

To find the *true* maximum of $m_{q_1 q_2 p}$ (not just the maximum for a given value of q) it is necessary to maximize (E.13) over the range of q allowed by (E.11). There are three possibilities; the maximum can occur for either one of the extreme values of λ , or it can occur somewhere in between. Each case must be considered separately.

E.4.1 $0 < \lambda < 1$

Differentiating (E.13) with respect to λ (or equivalently q), setting the result equal to zero, and solving for q yields:

$$q = (a + c) - \sqrt{\frac{c(4aj_q + (a - z)^2)}{j_q}} \quad (\text{E.14})$$

where

$$j_q = a + (a - z) \operatorname{sgn}[q - (c - a)] \quad (\text{E.15})$$

$$= \begin{cases} 2a - z & \text{when } q > c - a \\ z & \text{when } q < c - a. \end{cases} \quad (\text{E.16})$$

However, $q > c - a$ may be shown to imply $q > c - a > (\sqrt{c} - \sqrt{a})^2$, which itself contradicts the requirement $q \leq q^{\max} = (c - b)(b - c)/b$ from (E.11). Thus (E.16) implies that, if a stationary point of (E.13) exists, it can only occur at

$$q = \frac{(m_c m_z - m_a^2)(m_c - m_z)}{m_z}, \quad (\text{E.17})$$

provided that it also satisfies

$$q < c - a. \quad (\text{E.18})$$

The conditions (if any) under which q , given by (E.17), lies in the range permitted by (E.11) must now be determined. The $\lambda > 0$ bound ($q > 0$) clearly implies that $m_a^2 < m_z m_c$. The $\lambda < 1$ bound requires more work, but yields the additional constraints: $m_z m_c < m_b^2$ and $m_z m_b^2 < m_a^2 m_c$. These conditions can also be shown to enforce (E.18) which is required for consistency.

If the value of q given by (E.17) is substituted into (E.13), one finds that: $m_{q_1 q_2 p}^{\max(q)} = m_c - m_z$. Strictly speaking, thus far it has only been shown that $m_{q_1 q_2 p}^{\max(q)}$ is station-

ary. Given the value which has now been obtained, however, it is clear from energy conservation that the value is a maximum. From the fact that there were no other stationary values, we can conclude that the conditions under which there is a maximum within $0 < \lambda < 1$ are also sufficient conditions for the maximum to be *global*. Thus, in summary,

$$m_{q_1 q_2 p}^{\max} = m_c - m_z \quad \text{iff} \quad (E.19)$$

$$m_a^2 < m_z m_c < m_b^2 \quad \text{and} \quad m_z m_b^2 < m_a^2 m_c. \quad (E.20)$$

E.4.2 $\lambda = 0$

Substituting $q = 0$ into (E.13) yields

$$(m_{q_1 q_2 p}^{\max(q)})^2 = \frac{(c-a)(a-z)}{a}. \quad (E.21)$$

E.4.3 $\lambda = 1$

Substituting $q = q^{\max}$ into (E.13) yields

$$(m_{q_1 q_2 p}^{\max(q)})^2 = \frac{2a(c-b)(b-a) + (a-z)\{|b^2 - ac| + b^2 + ac - 2ab\}}{2ab} \quad (E.22)$$

which thus leads to two cases, depending on the sign of $b^2 - ac$ as follows:

$$(m_{q_1 q_2 p}^{\max(q)})^2 = \begin{cases} \frac{(ca-bz)(b-a)}{ab} & \text{where } b^2 > ac \\ \frac{(c-b)(b-z)}{b} & \text{where } b^2 < ac. \end{cases} \quad (E.23)$$

Conveniently, the difference between the two solutions can be shown to be a positive constant multiplied by $b^2 - ac$, thus (E.23) is equivalent to

$$(m_{q_1 q_2 p}^{\max(q)})^2 = \max \left[\frac{(ca - bz)(b - a)}{ab}, \frac{(c - b)(b - z)}{b} \right]. \quad (\text{E.24})$$

E.4.4 Summary

In Section E.4.1 a necessary and sufficient condition was found for the global maximum of $m_{q_1 q_2 p}$ to occur at a non-extremal value of q . If these conditions are not satisfied, the global maximum will be the larger of the two $m_{q_1 q_2 p}$ values obtained for the extremes of q . These were found in Sections E.4.2 and E.4.3. Putting these results together, the global maximum of $m_{q_1 q_2 p}$ is found to lie at:

$$(m_{q_1 q_2 p}^{\max})^2 = \begin{cases} (m_c - m_z)^2 & \text{where } m_a^2 < m_z m_c < m_b^2 \text{ and } m_z m_b^2 < m_a^2 m_c \\ \max \left[\frac{(c-a)(a-z)}{a}, \frac{(ca-bz)(b-a)}{ab}, \frac{(c-b)(b-z)}{b} \right] & \text{otherwise.} \end{cases} \quad (\text{E.25})$$

Subject only to the substitutions $m_z \rightarrow m_{\tilde{\chi}_1^0}$, $m_a \rightarrow m_{\tilde{l}_R}$, $m_b \rightarrow m_{\tilde{\chi}_2^0}$ and $m_c \rightarrow m_{\tilde{q}}$, (E.25) now describes the position of the l^+l^-q edge. At this point an answer exists and no more work is required. However, it is interesting to note that, mathematically, one may rewrite (E.25) in the equivalent form:

$$(m_{llq}^{\max})^2 = \begin{cases} (m_{\tilde{q}}^2 - m_{\tilde{\chi}_2^0}^2)(m_{\tilde{\chi}_2^0}^2 - m_{\tilde{\chi}_1^0}^2)/m_{\tilde{\chi}_1^0}^2 & \text{iff } m_{\tilde{\chi}_2^0}^2 < m_{\tilde{\chi}_1^0} m_{\tilde{q}}, \\ (m_{\tilde{q}}^2 - m_{\tilde{l}}^2)(m_{\tilde{l}}^2 - m_{\tilde{\chi}_1^0}^2)/m_{\tilde{l}}^2 & \text{iff } m_{\tilde{\chi}_1^0} m_{\tilde{q}} < m_{\tilde{l}}^2, \\ (m_{\tilde{q}}^2 m_{\tilde{l}}^2 - m_{\tilde{\chi}_2^0}^2 m_{\tilde{\chi}_1^0}^2)(m_{\tilde{\chi}_2^0}^2 - m_{\tilde{l}}^2)/(m_{\tilde{\chi}_2^0}^2 m_{\tilde{l}}^2) & \text{iff } m_{\tilde{l}}^2 m_{\tilde{q}} < m_{\tilde{\chi}_1^0} m_{\tilde{\chi}_2^0}^2, \\ (m_{\tilde{q}} - m_{\tilde{\chi}_1^0})^2 & \text{otherwise} \end{cases} \quad (\text{E.26})$$

which is nicer as the conditions are now all explicit.

E.5 l^+l^-q threshold

In the absence of detector effects, the invariant mass m_{llq} from decays of the form shown in Figure 5.3, can take any value between 0 and $m_{llq}^{\max}$ (given by (E.26)). These maximum and minimum values can be brought closer together by considering a constrained subset of the available l^+l^-q events, rather than an unconstrained sample as above. Requiring that $m_{ll} \geq m_{ll}^{\max}/\sqrt{2}$ (implying that the angular separation of the two leptons in the rest frame of the slepton be not less than a right angle) causes the minimum value of m_{llq} to become somewhat greater than zero. In particular, $m_{llq}^{\min}$ must then be *at least* as big as $m_{ll}^{\max}/\sqrt{2}$, but it is usually *bigger still* as the quark can only be at rest in the frame of the slepton when $m_{\tilde{q}}^2 + m_{\tilde{l}}^2 = 2m_{\tilde{\chi}_2^0}^2$. The location of $m_{llq}^{\min}$ under these circumstances is known as the l^+l^-q threshold.

Having imposed the constraint $m_{ll} \geq m_{ll}^{\max}/\sqrt{2}$, one intuitively expects the minimum l^+l^-q invariant mass to occur in events where m_{ll} actually equals $m_{ll}^{\max}/\sqrt{2}$. Although this intuition turns out to be correct, formal justification of it requires a little argument:

Without loss of generality we may work in the rest frame of the slepton, taking the lepton (l_{near}) from $\tilde{\chi}_2^0 \rightarrow \tilde{l}_R^\pm l^\pm$ to be emitted along the positive x -axis and the other lepton (l_{far}) to be emitted in the second quadrant of the (x, y) -plane.

It is easy to show that the minimum invariant mass is obtained when the quark lies in the positive- y half of the plane defined by the two leptons (consider rotation about the x -axis of arbitrary quark momenta, and the effect this has on the angle between the quark and each lepton). If the minimum invariant mass occurs when the quark is in the *first* quadrant (between the two leptons) then it follows that the leptons will adopt the configuration of least invariant mass, and the intuition is proved correct.

What if the invariant mass were to be minimised for a quark in the *second* quadrant of the (x, y) -plane? If the quark were to lie *between* the two leptons, rotation of l_{far} toward

the quark would reduce the invariant mass. Similarly, by boosting back into the frame of the $\tilde{\chi}_2^0$ (where arbitrary rotations of the quark momentum may be performed at constant modulus) it may be seen that if the quark lies *outside* the two leptons, rotating the quark towards l_{far} would again reduce the invariant mass of all three. Thus, if the invariant mass is minimised with the quark in the second quadrant, the minimum configuration has the quark parallel to l_{far} . In this special configuration (and treating the SM particles as massless) it is the case that $m_{llq}^2 = m_{l_{\text{near}}q}^2 + m_l^2$. To prove the correctness of the intuition, then, it is sufficient to show that $m_{l_{\text{near}}q}$ and m_{ll} are monotonic increasing functions of the angle between l_{near} and the pair of collinear particles. But in (respectively) the rest frames of the $\tilde{\chi}_2^0$ and the slepton, these requirements are self evident.

The result follows that, when calculating m_{llq}^{min} , it is sufficient to consider the case where the leptons are emitted at right angles in the rest frame of the slepton.^a

Treating the two leptons as a single “pseudo particle” with momentum q^μ and mass-squared q given by

$$q = \frac{1}{2}(m_l^{\text{max}})^2 = \frac{1}{2}(b-a)(a-z)/a, \quad (\text{E.27})$$

and denoting the quark momentum by p^μ , an application of (E.10) permits us to obtain the analogue of (E.13):

$$(m_{llq}^{\text{min}})^2 = q + \frac{(c-b)}{2b} \left\{ (q+b-z) - \sqrt{[q-(b+z)]^2 - 4bz} \right\}. \quad (\text{E.28})$$

^aNote that, were an alternative cut on m_{ll} to be proposed, the same arguments could be used to show that (for the purpose of calculating m_{llq}^{min}) consideration of the configuration(s) with m_{ll} *at the cut value* would be sufficient.

Once the value of q specified in (E.27) has been substituted into (E.28), the following result is obtained:

$$(m_{llq}^{\min})^2 = \frac{(c+b)(b-a)(a-z) + 2a(c-b)(b-z) - (c-b)\sqrt{(b+a)^2(a+z)^2 - 16ba^2z}}{4ba}. \quad (\text{E.29})$$

Note that if the position of the threshold is required for a different value of the $m_{ll}^{\min}$ cut, one only has to alter the value of q in (E.27).

E.6 High, low, near and far $l^\pm q$ edges

$l_{\text{near}}^\pm q$ edge: $m_{l_{\text{near}}q}^{\max}$

The maximum value of $m_{l_{\text{near}}q}$ is obtained from (E.9) by making the substitutions: $p \rightarrow$ quark, $q \rightarrow l^\pm$, $c \rightarrow \tilde{q}$, $b \rightarrow \tilde{\chi}_2^0$ and $a \rightarrow \tilde{l}$, and by treating the leptons as massless. The result of this substitution may be found in Table 5.2.

$l_{\text{far}}^\pm q$ edge: $m_{l_{\text{far}}q}^{\max}$

The maximum value of $m_{l_{\text{far}}q}$ has to occur when l_{far} and the quark are back-to-back in the rest frame of the slepton. The only kinematic degree of freedom over which it is necessary to maximise $m_{l_{\text{far}}q}$ is thus the angle ψ between the quark and l_{near} in the rest frame of the $\tilde{\chi}_2^0$. It can be shown that

$$(m_{l_{\text{far}}q})^2 = \frac{4sq}{\sqrt{b}}(E + p \cos \psi) \quad (\text{E.30})$$

where $p = (b-a)/(2\sqrt{a})$ and $q = (a-z)/(2\sqrt{a})$ are (respectively) the three momentum of l_{near} and l_{far} in the rest frame of the slepton; $E = \sqrt{b+p^2}$ is the energy of the $\tilde{\chi}_2^0$ in

the same frame; and $s = (c - b)/(2\sqrt{b})$ is the three momentum of the quark in the rest frame of the $\tilde{\chi}_2^0$. Once again, z , a , b and c are the squares of the sparticle masses as defined in Section E.4. Equation (E.30) shows that $m_{l_{\text{far}}q}$ is maximised by taking $\psi = 0$, and so after simplification one obtains the result:

$$(m_{l_{\text{far}}q}^{\text{max}})^2 = \frac{(c - b)(a - z)}{a}. \quad (\text{E.31})$$

$l^\pm q$ high-edge: $m_{lq(\text{high})}^{\text{max}}$

From the definition of $m_{lq(\text{high})}$ (see Section 5.6.5) it immediately follows that:

$$m_{lq(\text{high})}^{\text{max}} = \max_{\text{possible decays}} \{ \max(m_{l_{\text{near}}q}, m_{l_{\text{far}}q}) \} \quad (\text{E.32})$$

$$= \max\{m_{l_{\text{near}}q}^{\text{max}}, m_{l_{\text{far}}q}^{\text{max}}\}. \quad (\text{E.33})$$

$l^\pm q$ low-edge: $m_{lq(\text{low})}^{\text{max}}$

The $l^\pm q$ low-edge is a tricky beast. By definition:

$$m_{lq(\text{low})}^{\text{max}} \equiv \max_{\text{possible decays}} \{ \min(m_{l_{\text{near}}q}, m_{l_{\text{far}}q}) \}, \quad (\text{E.34})$$

and so it is not possible to consider the $l_{\text{near}}^\pm q$ and $l_{\text{far}}^\pm q$ events separately, as was possible with the $l^\pm q$ high-edge above. The best one can say is that:

It is true that the maximum value of $m_{lq(\text{low})}$ must either:

1. occur in a configuration in which $m_{l_{\text{near}}q} = m_{l_{\text{far}}q}$, or
2. be local maximum of $m_{l_{\text{near}}q}$ which is achievable in a configuration in which $m_{l_{\text{near}}q} < m_{l_{\text{far}}q}$, or

3. be a local maximum of $m_{l_{\text{far}}q}$ which is achievable in a configuration in which $m_{l_{\text{far}}q} < m_{l_{\text{near}}q}$.

The maximum value of $m_{lq(\text{low})}$ will be the larger (or largest) of the values permitted by the above Cases 1, 2 and 3.

Each of the Cases 1, 2 and 3 will be dealt with in turn. For Case 1 it is helpful make the following definitions:

Let the abstract quantity $\mathbf{k} \in \mathcal{K}$ denote a particular decay configuration^b from the space of all possible decays $\mathcal{K}$, then we may then view $m_{l_{\text{near}}q}$ and $m_{l_{\text{far}}q}$ as functions of $\mathbf{k}$. Define the subset of events $\mathcal{K}' \subset \mathcal{K}$ to be those where all SM particles are co-planar in the rest frame of the slepton. Then define the region $\mathcal{S} \subset \mathcal{K}'$ by the requirement that it contain only events where $m_{l_{\text{near}}q}(\mathbf{k}) = m_{l_{\text{far}}q}(\mathbf{k})$. Finally define the set of points $\mathcal{P} \subset \mathcal{S}$, to be those that are local maxima in $m_{l_{\text{near}}q}(\mathbf{k})$ (or indeed in $m_{l_{\text{far}}q}(\mathbf{k})$, since these are equivalent in $\mathcal{S}$). The determination of $m_{lq(\text{low})}^{\text{max}}$ then runs as follows.

Firstly, working in the rest frame of the slepton, and choosing (without loss of generality) to have l_{near} run along the positive x -axis and the quark to lie in the upper half of the (x, y) -plane, one notes that l_{far} can be emitted in any direction, but always with constant momentum. Thus, given an arbitrary event $\mathbf{k} \in \mathcal{K}$, it is possible to spatially rotate l_{far} about the quark's direction until it lies in the (x, y) -plane. This rotation changes neither $m_{l_{\text{near}}q}$ nor $m_{l_{\text{far}}q}$, and so one can see that in order to consider all pairs $(m_{l_{\text{near}}q}, m_{l_{\text{far}}q})$ it is sufficient to consider only events in $\mathcal{K}'$. This simplifies the analysis of all three Cases.

^bFor example, $\mathbf{k}$ might be an abstract vector containing the spherical polar angles for each of the sparticles in the decay chain, up to an overall rotation. It will be assumed that a 'sensible' metric exists on $\mathcal{K}$, and that $\mathcal{K}$ is compact.

Case 1

By dividing $\mathcal{P}$ up into four regions, according to whether the quark is (or is not) back-to-back with each of the two leptons, and by considering each region separately, one may show that configurations in $\mathcal{P}$, if there are any, must have the quark back-to-back with at least one of the leptons. The other possibilities are ruled out by showing that their existence would lead to a contradiction. Moreover, in the cases which *are* allowed, it may be shown that the ratio p/q determines which of the leptons should oppose the quark; in the case that $p \leq q$, the back-to-back $l^\pm q$ pair must contain l_{near} , while if $p \geq q$ it must contain l_{far} . These conditions are complementary to each other, and are also sufficient for $\mathcal{P}$ to be non-empty. Thus (for a given choice of sparticle masses) one can show that $\mathcal{P}$ contains *exactly one* configuration.

Although $\mathcal{P}$ has been shown to contain exactly one configuration, its form is not yet specified completely. We have not yet determined the angle ϕ between the quark and the lepton which it is **not** back-to-back with. This ϕ is a function of the sparticle masses, and is fixed by the requirement that $m_{l_{\text{far}}q} = m_{l_{\text{near}}q}$. By explicitly determining ϕ , and then evaluating the invariant mass in that configuration, it may be shown that if $\mathbf{k}_P$ is the element in $\mathcal{P}$, then:

$$M_{\mathcal{P}}^2 \equiv m_{l_{\text{near}}q}^2(\mathbf{k}_P) = m_{l_{\text{far}}q}^2(\mathbf{k}_P) = \begin{cases} (m_{l_{\text{near}}q}^{\text{max}})^2 & \text{when } p \leq q \\ \frac{(c-b)(a-z)}{2a-z} & \text{when } p \geq q. \end{cases} \quad (\text{E.35})$$

Thus far we have found that if the maximum value of $m_{lq(\text{low})}$ occurs in “Case 1” then it will take the value given by $M_{\mathcal{P}}$ defined in (E.35).

Case 2

There is only one unconstrained local maximum for $m_{l_{\text{near}}q}$, and it occurs when the quark is back-to-back with l_{near} in the rest frame of the slepton. We have already determined the value of $m_{l_{\text{near}}q}^{\text{max}}$, and so now we are interested in finding out whether this configuration (in which the direction of l_{far} is still free) allows $m_{l_{\text{far}}q}$ to exceed $m_{l_{\text{near}}q}^{\text{max}}$. Since $m_{l_{\text{far}}q}$ is maximised when it too is opposite to the (fixed) squark it is evident that $p < q$ is the condition for $m_{l_{\text{near}}q}^{\text{max}}$ to be added to the list of possible candidates for $m_{lq(\text{low})}^{\text{max}}$. Note however, that the condition controlling its inclusion is precisely the condition which causes it to be added under “Case 1” anyway. Therefore “Case 2” adds no new information and may be ignored.

Case 3

There is only one unconstrained local maximum for $m_{l_{\text{far}}q}$, and it occurs when the quark is pointing along l_{near} and back-to-back with l_{far} in the rest frame of the slepton. Note that, in contrast with “Case 2” all the momenta are now fixed. Moreover, this configuration has fixed $m_{l_{\text{near}}q} = 0$, which means that $m_{l_{\text{far}}q}$ in this state is most certainly not smaller than $m_{l_{\text{near}}q}$. Thus “Case 3” generates no masses at all.

Summary

Putting all the results together, the following result is obtained:

$$(m_{lq(\text{low})}^{\text{max}})^2 = \begin{cases} (m_{l_{\text{near}}q}^{\text{max}})^2 & \text{when } p \leq q \\ \frac{(c-b)(a-z)}{2a-z} & \text{when } p \geq q. \end{cases} \quad (\text{E.36})$$

E.7 $\tilde{l}\tilde{l}$ edge

It was mentioned in Section 5.6.6 that a sufficient condition for an event to yield a maximal value of $M_{T2}(m_{\tilde{\chi}_1^0})$ is that $\Delta y_1 = \Delta y_2 = 0$ and $\mathbf{v}_T^{l_1} - \mathbf{v}_T^{\tilde{\chi}_1^0}$ is parallel (note: not anti-parallel) to $\mathbf{v}_T^{l_2} - \mathbf{v}_T^{\tilde{\chi}_2^0}$. The definitions of Δy_i and $\mathbf{v}_T^X$ may be found by referring back to Section 5.6.6. This condition was derived in the context of a special value of χ , namely the true neutralino mass, $m_{\tilde{\chi}_1^0}$. What about values of χ ? A study using a toy monte-carlo has indicated that the evaluation of $M_{T2}(\chi)$ at values of χ different than $m_{\tilde{\chi}_1^0}$ can introduce a slight dependence of the edge position upon the maximum transverse boost attained by the sleptons. For example, it has been observed that in a situation where *both* sleptons are able to be emitted collinearly and with transverse momenta of 20 GeV (or 100 GeV or 1000 GeV), then the value of $\Delta M^{\max}$ predicted by Equation (5.21) for heavy (left) sleptons at S5 would underestimate the true value by about 2% (or 8% or 14%) as a result of ignoring the boosts. As expected, transverse boosts are observed to have no effect on $M_{T2}^{\max}(m_{\tilde{\chi}_1^0})$. This effect is still under investigation,^c but within this analysis it may safely be ignored for two reasons. Firstly, dilepton production is threshold dominated, and so sleptons typically have small transverse momenta. Secondly, and more importantly, the cut limiting jet activity to “at most one with $p_T < 40$ GeV” puts a strong constraint on the momentum available for the sleptons to recoil against. This restricts the washing out of the edge to a 2% effect, which is far less than the order 10% error on the edge position due to statistics. In situations where the effect cannot be neglected, systematic error could be minimised by evaluating $M_{T2}(\chi)$ at the best available guess for $m_{\tilde{\chi}_1^0}$.

^cPreliminary results from a toy monte-carlo study of the situation in which both sleptons acquire large transverse boosts in the same direction suggest that the maximal value of $M_{T2}(0)$ is obtained in events where (in addition to the original conditions on Δy_1 , Δy_2 , $\mathbf{v}_T^{l_1} - \mathbf{v}_T^{\tilde{\chi}_1^0}$ and $\mathbf{v}_T^{l_2} - \mathbf{v}_T^{\tilde{\chi}_2^0}$) the two leptons are emitted with minimal transverse momenta in the laboratory frame. In general, it seems likely that the original conditions are both necessary and sufficient for maximal $M_{T2}(m_{\tilde{\chi}_1^0})$, but that when $M_{T2}(\chi)$ is evaluated at values of χ other than $m_{\tilde{\chi}_1^0}$, the sufficiency of the conditions is lost when the sleptons are permitted to acquire transverse boosts.

To establish the value of $M_{T2}^{\max}(\chi)$ (in a situation where the effects of any slepton boosts may be neglected) it is thus sufficient to calculate the value of $M_{T2}(\chi)$ in an event satisfying the condition described above. The simplest such event contains two sleptons, at rest in the frame of the detector, and with each decaying into a lepton travelling along the negative x -axis and a neutralino along the positive x -axis. In the frame of the sleptons, both the neutralino and the lepton will have a three momentum $\mathbf{p}$ with modulus p satisfying

$$p^2 = \frac{m_l^4 + m_{\tilde{\chi}_1^0}^4 + m_l^4 - 2(m_l^2 m_{\tilde{\chi}_1^0}^2 + m_l^2 m_{\tilde{\chi}_1^0}^2 + m_l^2 m_l^2)}{4m_l^2}. \quad (\text{E.37})$$

In order to calculate $M_{T2}(\chi)$ in accordance with Definition (5.14) we need to evaluate

$$M_{T2}^2(\chi) \equiv \min_{\mathbf{p}_1 + \mathbf{p}_2 = \mathbf{p}_T} \left[\max \{ m_T^2(\mathbf{p}_T^{l_1}, \mathbf{p}_1, \chi), m_T^2(\mathbf{p}_T^{l_2}, \mathbf{p}_2, \chi) \} \right]. \quad (\text{E.38})$$

For the particular configuration described above we can say that

$$m_T^2(\mathbf{p}_T^{l_1}, \mathbf{p}_1, \chi) = \chi^2 + m_l^2 + 2 \left(\sqrt{m_l^2 + p^2} \sqrt{\chi^2 + (p + \delta)^2 + q^2} + p^2 + p\delta \right) \quad (\text{E.39})$$

$$m_T^2(\mathbf{p}_T^{l_2}, \mathbf{p}_2, \chi) = \chi^2 + m_l^2 + 2 \left(\sqrt{m_l^2 + p^2} \sqrt{\chi^2 + (p - \delta)^2 + q^2} + p^2 - p\delta \right) \quad (\text{E.40})$$

where $\mathbf{p}_1$ and $\mathbf{p}_2$ have been parametrised by the two real numbers δ and q according to

$$\mathbf{p}_1 = \begin{pmatrix} p + \delta \\ q \end{pmatrix} \quad \text{and} \quad \mathbf{p}_2 = \begin{pmatrix} p - \delta \\ -q \end{pmatrix} \quad (\text{E.41})$$

in order to ensure that $\mathbf{p}_1 + \mathbf{p}_2 = \mathbf{p}_T$ is automatically satisfied. Note that we may choose both $q \geq 0$ and $\delta \geq 0$ without loss of generality, and having done so we may then deduce that the right-hand-side of (E.39) is always greater than or equal to the right-hand-side of (E.40). With this fact in mind, the two values of m_T may be substituted into (E.38)

yielding

$$M_{T2}^2(\chi) = \min_{q \geq 0, \delta \geq 0} \left[\chi^2 + m_l^2 + 2 \left(\sqrt{m_l^2 + p^2} \sqrt{\chi^2 + (p + \delta)^2 + q^2} + p^2 + p\delta \right) \right]. \quad (\text{E.42})$$

Minimisation over q is trivial and is achieved by setting $q = 0$. The derivative with respect to δ of the argument of the “max” in (E.42) may also be evaluated, and can be shown to be positive. The minimisation with respect to δ is thus achieved by setting δ also equal to zero. We now have all that is required to determine $M_{T2}(\chi)$ for this particular configuration, which by design provides us with the general formula for $(M_{T2}^{\max}(\chi))^2$:

$$(M_{T2}^{\max}(\chi))^2 = \chi^2 + m_l^2 + 2p^2 + 2\sqrt{m_l^2 + p^2}\sqrt{\chi^2 + p^2}, \quad (\text{E.43})$$

from which we observe that

$$M_{T2}^{\max}(\chi) = \sqrt{m_l^2 + p^2} + \sqrt{\chi^2 + p^2}. \quad (\text{E.44})$$

Notice that if χ is set equal to the real neutralino mass, $m_{\tilde{\chi}_1^0}$, then we recover the expected result $M_{T2}^{\max}(m_{\tilde{\chi}_1^0}) = E_l + E_{\tilde{\chi}_1^0} = E_{\tilde{l}} = m_{\tilde{l}}$.

E.8 Phase space for $c \rightarrow Zb$ followed by $b \rightarrow al^+l^-$

Neglecting lepton masses, the (phase space only) distribution of M_{Zll}^2 from the decay chain $c \rightarrow Zb$ followed by $b \rightarrow al^+l^-$ has the form:

$$\frac{d\sigma}{d(M_{Zll}^2)} \propto W(M_{Zll}^2), \quad (\text{E.45})$$

where the “constant” of proportionality may depend on m_a , m_b and m_c but not on M_{Zl} , and where

$$W(M_{Zl}^2) = \begin{cases} w^+ & \text{if } M_{Zl}^2 \leq (m_c - m_a)^2 \text{ and } w^+ \geq 0 \geq w^- \\ w^+ - w^- & \text{if } M_{Zl}^2 \leq (m_c - m_a)^2 \text{ and } w^- \geq 0 \\ 0 & \text{otherwise} \end{cases} \quad (\text{E.46})$$

where w^+ and w^- are defined by

$$\begin{aligned} w^\pm &= 2(m_a^2 + m_b^2)m_c^2 \pm [m_c^2, m_b^2, m_z^2] [m_c^2, m_a^2, M_{Zl}^2] \\ &\quad - (m_c^2 + m_b^2 - m_z^2)(m_c^2 + m_a^2 - M_{Zl}^2) \end{aligned} \quad (\text{E.47})$$

in which the following notation has been used:

$$[x, y, z] \equiv \sqrt{x^2 + y^2 + z^2 - 2(xy + xz + yz)}. \quad (\text{E.48})$$

Appendix F

Abbreviations

ATLAS	A Toroidal LHC Apparatu S
AMSB	anomaly mediated supersymmetry breaking
BFB	bounded from below
CERN	European laboratory for particle physics
DOF	degree of freedom
EAGLE	Experiment for Accurate Gamma, Lepton and Energy measurements
ECAL	electromagnetic calorimeter
EWSB	electro-weak symmetry breaking
GMSB	gauge mediated supersymmetry breaking
HCAL	hadronic calorimeter
LAr	liquid argon
LEP	large electron-positron
LHC	large hadron collider
LSP	lightest supersymmetric particle
MSSM	minimal supersymmetric standard model
mSUGRA	minimal supergravity
O1	optimised string model Point 1
OSDF	opposite-sign different-lepton-family
OSM	optimised string model
OSSF	opposite-sign same-lepton-family
PPARC	the particle physics and astronomy research council
QFΘ	quantum field theory
RGEs	renormalization group equations
R_P C	R -parity conserving
R_P V	R -parity violating
S2	mSUGRA Point 2
S3	mSUGRA Point 3
S5	mSUGRA Point 5
SCT	semiconductor tracker
SM	Standard Model
TR	transition radiation
TRT	transition radiation tracker
UFB	unbounded from below
VEV	vacuum expectation value

Bibliography

- [1] C. Arnault et al. *A generic approach to the detector description in ATLAS* (Feb 2000) Published in “Padua 2000, Computing in high energy and nuclear physics” 269-272.
- [2] P. Higgs, *Broken symmetries, massless particles and gauge fields*, *Phys. Lett.* **12** (1964) 132–133.
- [3] R. Hempfling, *Neutrino masses and mixing angles in SUSY-GUT theories with explicit R-parity breaking*, *Nucl. Phys.* **B478** (1996) 3–30, [[hep-ph/9511288](#)].
- [4] M. Hirsch, M. A. Diaz, W. Porod, J. C. Romao, and J. W. F. Valle, *Neutrino masses and mixings from supersymmetry with bilinear R-parity violation: A theory for solar and atmospheric neutrino oscillations*, *Phys. Rev.* **D62** (2000) 113008, [[hep-ph/0004115](#)].
- [5] P. Richardson, *Simulations of R-parity violating SUSY models*, [hep-ph/0101105](#).
- [6] S. P. Martin, *A supersymmetry primer*, [hep-ph/9709356](#).
- [7] J. Bagger and J. Wess, *Supersymmetry and Supergravity*. 1992. ISBN 0691085560.
- [8] H. E. Haber and G. L. Kane, *The search for supersymmetry: Probing physics beyond the standard model*, *Phys. Rept.* **117** (1985) 75.
- [9] F. Quevedo, *Lectures on superstring phenomenology*, [hep-th/9603074](#).
- [10] J. F. Gunion, H. E. Haber, G. L. Kane, and S. Dawson, *The higgs hunter’s guide*. SCIPP-89-13.

- [11] D. Balin and A. Love, *Supersymmetric Gauge Field Theory and String Theory*. 1994. Chapter 2: Lagrangians for Chiral Superfields.
- [12] T. Banks, L. J. Dixon, D. Friedan, and E. J. Martinec, '*Phenomenology and conformal field theory*' or '*Can string theory predict the weak mixing angle?*', *Nucl. Phys.* **B299** (1988) 613–626.
- [13] E. Cremmer, S. Ferrara, L. Girardello, and A. Van Proeyen, *Coupling supersymmetric Yang-Mills theories to supergravity*, *Phys. Lett.* **B116** (1982) 231.
- [14] J. A. Casas, A. Ibarra, and C. Munoz, *Phenomenological viability of string and M-theory scenarios*, *Nucl. Phys.* **B554** (1999) 67, [[hep-ph/9810266](#)].
- [15] J. A. Casas, *Charge and color breaking*, [hep-ph/9707475](#). Published in "Perspectives on Supersymmetry", edited by G.L. Kane, World Scientific.
- [16] **ATLAS** Collaboration, *ATLAS: Letter of intent for a general purpose p p experiment at the large hadron collider at CERN*. 1992. CERN-LHCC-92-4.
- [17] P. Jenni, *EAGLE: experiment for accurate gamma, lepton and energy measurements*. In [40], 5-8 Mar, 1992. Presented at General Meeting on LHC Physics and Detectors Towards the LHC Experimental Program, Evian-les-Bains, France.
- [18] P. R. Norton, *ASCOT: A Detector for the LHC: Expression of interest*. In [40], 5-8 Mar, 1992. Presented at General Meeting on LHC Physics and Detectors Towards the LHC Experimental Program, Evian-les-Bains, France.
- [19] **ATLAS** Collaboration, *Higgs Bosons*, ch. 19. In [41], May, 1999.
- [20] **ATLAS** Collaboration, *ATLAS Inner Detector Technical Design Report*. No. CERN/LHCC/97-16. April, 1997.
- [21] **ATLAS** Collaboration, *ATLAS Calorimeter Performance Technical Design Report*. No. CERN/LHCC/96-40. December, 1996.

- [22] **ATLAS** Collaboration, *ATLAS muon spectrometer: Technical design report*. No. CERN-LHCC-97-22. 1997.
- [23] E. Richter-Was, D. Froidevaux, and L. Poggioli, *ATLFAST 2.0: a fast simulation package for ATLAS*, Tech. Rep. ATL-PHYS-98-131, 1998.
- [24] **ATLAS** Collaboration, *ATLAS Detector and Physics Performance Technical Design Report 1*. No. CERN-LHCC-99-014 ATLAS-TDR-14. May, 1999.
- [25] **ATLAS** Collaboration, W. W. Armstrong *et. al.*, *ATLAS: Technical proposal for a general-purpose $p p$ experiment at the large hadron collider at CERN*, 1994. CERN-LHCC-94-43.
- [26] L. Poggioli, D. Froidevaux, and C. Guyot, *Physics performance for various muon system configurations*, Tech. Rep. ATL-PHYS-95-076, 1995.
- [27] D. R. Tovey, *Measuring the SUSY mass scale at the LHC*, *Phys. Lett.* **B498** (2001) 1–10, [[hep-ph/0006276](#)].
- [28] M. Dine, W. Fischler, and M. Srednicki, *Supersymmetric technicolor*, *Nucl. Phys.* **B189** (1981) 575–593.
- [29] I. Hinchliffe, F. E. Paige, M. D. Shapiro, J. Soderqvist, and W. Yao, *Precision SUSY measurements at LHC*, *Phys. Rev.* **D55** (1997) 5520–5540, [[hep-ph/9610544](#)].
- [30] **ATLAS** Collaboration, *Fitting minimal SUGRA parameters*. In [41], May, 1999. Section 20.2.9.
- [31] **ATLAS** Collaboration, *Supergravity Models*. In [41], May, 1999. Sections 20.2.3-20.2.6.
- [32] H. Bachacou, I. Hinchliffe, and F. E. Paige, *Measurements of masses in SUGRA models at LHC*, *Phys. Rev.* **D62** (2000) 015009, [[hep-ph/9907518](#)].

- [33] G. Polesello, L. Poggioli, E. Richter-Was, and J. Söderqvist, “*Precision SUSY measurements with ATLAS for SUGRA point 5.*” ATLAS Internal Note, 1997. ATL-PHYS-97-111.
- [34] B. C. Allanach, C. G. Lester, M. A. Parker, and B. R. Webber, *Measuring sparticle masses in non-universal string inspired models at the LHC*, *JHEP* **09** (2000) 004, [[hep-ph/0007009](#)].
- [35] C. G. Lester and D. J. Summers, *Measuring masses of semi-invisibly decaying particle pairs produced at hadron colliders.*, *Phys. Lett.* **B463** (1999) [[hep-ph/9906349](#)].
- [36] G. Corcella *et. al.*, *HERWIG 6.1 release note*, [hep-ph/9912396](#).
- [37] F. E. Paige, S. D. Protopopescu, H. Baer, and X. Tata, *ISAJET 7.40: A monte carlo event generator for pp , $p\bar{p}$, and e^+e^- reactions*, [hep-ph/9810440](#).
- [38] E. Richter-Was, D. Froidevaux, and J. Söderqvist, “*Precision SUSY measurements with ATLAS for SUGRA points 1 and 2.*” ATLAS Internal Note, 1997. ATL-PHYS-97-108.
- [39] G. Ciapetti and A. Di Ciaccio, *Monte carlo simulation of minimum bias events at the LHC energy*, in *Aachen LHC Workshop* (G. Jarlskog and D. Rein, eds.), vol. 2, pp. 155–162, Dec, 1990.
- [40] *Proceedings of the General Meeting on LHC Physics and Detectors Towards the LHC Experimental Program, Evian-les-Bains, France.* 5-8 Mar, 1992.
- [41] **ATLAS** Collaboration, *ATLAS Detector and Physics Performance Technical Design Report 2*. No. CERN-LHCC-99-015 ATLAS-TDR-15. May, 1999.