

Luminosity determination and searches for supersymmetric sleptons and gauginos at the ATLAS experiment

Thesis submitted for the degree of
Doctor of Philosophy

by

Anders Floderus



LUND
UNIVERSITY

DEPARTMENT OF PHYSICS
LUND, 2014



Abstract

This thesis documents my work in the luminosity and supersymmetry groups of the ATLAS experiment at the Large Hadron Collider. The theory of supersymmetry and the concept of luminosity are introduced. The ATLAS experiment is described with special focus on a luminosity monitor called LUCID. A data-driven luminosity calibration method is presented and evaluated using the LUCID detector. This method allows the luminosity measurement to be calibrated for arbitrary inputs. A study of particle counting using LUCID is then presented. The charge deposited by particles passing through the detector is shown to be directly proportional to the luminosity. Finally, a search for sleptons and gauginos in final states with exactly two oppositely charged leptons is presented. The search is based on 20.3 fb^{-1} of pp collision data recorded with the ATLAS detector in 2012 at a center-of-mass energy of $\sqrt{s} = 8 \text{ TeV}$. No significant excess over the Standard Model expectation is observed. Instead, limits are set on the slepton and gaugino masses.

Populärvetenskaplig sammanfattning

Partikelfysiken är studien av naturens minsta beståndsdelar — De så kallade *elementarpartiklarna*. All materia i universum består av elementarpartiklar. Den teori som beskriver vilka partiklar som finns och hur de uppför sig heter *Standardmodellen*. Teorin har historiskt sett varit mycket framgångsrik. Den har gång på gång förutspått existensen av nya partiklar innan de kunnat påvisas experimentellt, och klarar av att beskriva många experimentella resultat med imponerande precision. Men Standardmodellen är inte hela sanningen. Det finns fenomen som den inte kan förklara, och partikelfysiker bygger därför experiment för att testa olika utökningar av modellen. En av de mest lovande kallas för *supersymmetri*. Denna teori postulerar att den finns en ny, ännu oupptäckt, partikel för varje partikel inom Standardmodellen. De nya partiklarna tros vara mycket tyngre än de partiklar som hittills har observerats, och skulle till exempel kunna förklara förekomsten av fenomenet mörk materia.

För att skapa tunga partiklar krävs det stora mängder energi. I CERN-laboratoriet utanför Geneve har man byggt världens största partikelaccelerator, vars syfte är just att förse partiklar med hög kinetisk energi. Acceleratorn kallas för LHC (Large Hadron Collider). Den accelererar två protonstrålar till en energi på 7 TeV, eller *teraelektronvolt* per proton, och låter sedan protonerna kollidera i mitten av en partikeldetektor. I dessa kollisioner skapas nya partiklar ur protonernas kinetiska energi. Detektorns uppgift är att mäta egenskaperna hos de partiklar som bildas. Exempel på sådana egenskaper är massa och elektrisk laddning. Genom att studera data från detektorn kan partikelfysiker bestämma huruvida några fenomen bortom Standardmodellen har ägt rum.

En viktig pusselbit i jakten på ny fysik är mätningen av protonstrålens *luminositet*. Detta är ett mått på det totala antalet kollisioner som äger rum. Det är nämligen inte tillräckligt att bara *observera* en ny typ av partikel — Man vill även mäta hur ofta den produceras! Det är inte ovanligt att olika utökningar av Standardmodellen gör snarlika förutsägelser i detta avseende, eller har fria parametrar som gör att antalet producerade partiklar inte kan bestämmas i förväg. För att testa en modell jämförs den teoretiska sannolikheten (även känt som *tvärsnittet*) för att producera en viss partikel med en experimentell mätning. Alla modeller som inte stämmer överens med den experimentella mätningen kan sedan uteslutas. För att mäta ett tvärsnitt divideras helt enkelt antalet kollisioner som gav upphov till partikeln i fråga med det totala antalet kollisioner — Det vill säga luminositeten! Förmågan att mäta luminositet med hög precision är också viktig för att kunna göra meningsfulla simuleringar av experimentet. Innan en simulering kan jämföras med riktig data så måste antalet simulerade kollisioner normaliseras till antalet kollisioner som faktiskt ägt rum.

Denna avhandling beskriver den allra största partikeldetektorn vid LHC — Det så kallade ATLAS-experimentet. Därefter diskuteras olika metoder för att mäta luminositeten i ATLAS med hjälp av en av experimentets subdetektorer vid namn LUCID. En metod använder data från protonstrålar som tagits under speciella förutsättningar, och korregerar med hjälp av denna data för det icke-linjära beteende som observeras hos LUCID-detektorn vid höga luminositetsnivåer. En annan metod är baserad på att mäta den totala energin hos partiklarna som passerar genom LUCID. Denna energi visas vara proportionell mot luminositeten, och metoden visas ha vissa egenskaper som är att föredra framför mer konventionella metoder vad gäller mätningar när luminositeten är hög. Slutligen beskrivs en analys vars mål är att hitta två olika typer av supersymmetriska partiklar i de kollisioner som ägt rum inom ATLAS-detektorn. Inga nya partiklar observeras, och resultatet av analysen är därför gränsvärden på vilka modeller av supersymmetri som kan uteslutas.

To my family

Contents

Preface	1
Acknowledgements	3
1 The Standard Model	4
1.1 Overview	4
1.2 Gauge theories	5
1.2.1 Lagrangians	5
1.2.2 Gauge invariance	6
1.3 Quantum electrodynamics	7
1.3.1 The QED Lagrangian	7
1.3.2 Vacuum polarization	9
1.4 Quantum chromodynamics	9
1.4.1 The quark model	10
1.4.2 The QCD Lagrangian	10
1.4.3 Long range behavior	12
1.5 Weak interactions	13
1.5.1 Chiral fermions	14
1.5.2 Weak iso-doublets	14
1.5.3 A weak Lagrangian	15
1.5.4 Electroweak unification	15
1.5.5 Spontaneous symmetry breaking	16
1.5.6 The Higgs mechanism	19
1.5.7 Yukawa couplings	22
1.5.8 The electroweak Lagrangian	23
1.6 An incomplete theory	24
1.6.1 Dark matter	24
1.6.2 The hierarchy problem	24
1.6.3 Grand unification	24
1.6.4 Neutrino masses	24
2 Supersymmetry	27
2.1 Overview	27
2.2 Solving the hierarchy problem	27
2.3 Particles in the MSSM	29
2.3.1 Sparticle mixing	30
2.4 R-parity	31
2.5 Supersymmetry breaking	31
2.5.1 Gravity mediated	32
2.5.2 Gauge mediated	33
2.6 Sparticle decays	34
2.7 The pMSSM	36

3	The Large Hadron Collider	37
3.1	Overview	37
3.2	Accelerator system	38
3.3	Experiments	39
4	The ATLAS experiment	41
4.1	Overview	41
4.2	Coordinate system	43
4.3	Magnet system	45
4.4	Inner detector	46
4.4.1	Pixel detector	47
4.4.2	Semiconductor tracker	48
4.4.3	Transition radiation tracker	48
4.5	Calorimeters	48
4.5.1	Electromagnetic calorimeter	50
4.5.2	Hadronic calorimeter	50
4.6	Muon spectrometer	53
4.7	Forward detectors	54
4.7.1	BCM	55
4.7.2	MBTS	55
4.7.3	LUCID	55
4.7.4	ZDC	56
4.7.5	ALFA	56
4.8	Trigger and Data Acquisition	56
5	The LUCID detector	58
5.1	Overview	58
5.2	Detector design	59
5.3	Readout electronics	62
5.3.1	The local stream	63
5.3.2	LUMAT	64
5.4	Calibration	65
5.5	Upgrade	66
6	Luminosity	70
6.1	Overview	70
6.2	Luminosity monitoring	71
6.3	Calibration	72
6.4	Algorithms	73
6.4.1	EventOR	74
6.4.2	EventAND	76
6.4.3	Hit counting	77
6.5	The Online Luminosity Calculator	77
7	The LUCID pileup method	78
7.1	Introduction	78
7.2	Method	78
7.2.1	Obtaining a single-interaction sample	78
7.2.2	Pileup	79

7.3	Data	81
7.4	Migration effects	83
7.5	Statistical uncertainty	83
7.6	Systematic uncertainty	84
7.6.1	Thresholds	84
7.6.2	Consistency	87
7.7	Conclusions	89
8	The LUCID particle counting method	90
8.1	Introduction	90
8.2	Charge measurement	90
8.2.1	Static baselines	90
8.2.2	PreTrain baselines	91
8.2.3	Dynamic baselines	91
8.3	Calibration	91
8.4	Systematic uncertainties	93
8.5	Runs with low HV	94
8.6	Results	96
8.7	Conclusions	97
9	A search for sleptons and gauginos in final states with two leptons	99
9.1	Introduction	99
9.2	Signal models	100
9.2.1	Simplified models	100
9.2.2	pMSSM	102
9.3	Event reconstruction and selection	103
9.3.1	Object definitions	103
9.3.2	Overlap removal	104
9.3.3	Triggers	105
9.3.4	Relative missing transverse energy	110
9.3.5	Stranverse mass	111
9.4	Signal regions	113
9.4.1	Optimization technique	113
9.4.2	SR- m_{T2}	114
9.4.3	SR-WW	115
9.5	Background estimation	115
9.5.1	WW	116
9.5.2	ZV	119
9.5.3	Top	120
9.5.4	Non-prompt	123
9.5.5	Others	126
9.5.6	Simultaneous fit	126
9.6	Background uncertainty	127
9.6.1	Experimental uncertainties	127
9.6.2	WW modelling uncertainty	129
9.6.3	ZV modelling uncertainty	134
9.6.4	Top modelling uncertainty	136
9.7	Results	136
9.7.1	SR- m_{T2}	136

9.7.2	SR- WW	138
9.8	Interpretation	138
9.8.1	Simplified models	140
9.8.2	pMSSM	140
9.9	Conclusions	146
9.10	Outlook	146
A	Additional particle counting tables	148
B	Monte Carlo generators	156
C	Limit setting techniques	159
C.1	Hypothesis testing	159
C.2	Constructing the test-statistic	160
C.3	Setting the limits	163
C.4	The CL_s technique	165

Preface

Physics is about understanding nature. The subfield of *particle physics* reasons that if nature can be understood on its most fundamental level, then the rest will more or less follow. Simply put, particle physicists use experiments in order to deduce the mathematical equations that govern the universe. These equations describe nature's elementary building blocks, the *particles*, and how they interact with each other. But the particle physicist is not content with studying just *any* old interaction. Things get interesting only after giving the particles gratuitous amounts of *energy*! To that effect, the physics community has built the Large Hadron Collider (LHC). This is a machine designed to provide groups of particles with as much energy as is technically feasible and then *smash them together* millions of times every second. With all that energy at their proverbial fingertips, the particles have the freedom to interact in ways that would normally not be possible — The heaviest particles are created and the rarest processes occur! The high-energy frontier is the uncharted territory of nature, and only in such an environment can her fundamental laws truly be studied in detail.

The instruments used to observe and record the interactions that take place in collisions at the LHC are called *particle detectors*. I began working at one such detector, the ATLAS experiment, during my MSc project in 2009. My task was to calibrate and monitor the data quality of one of the experiment's subdetectors. The subdetector, which goes by the name LUCID, is a *luminosity monitor*. This essentially means that it counts the total number of particle collisions that take place in the LHC beam. Having the ability to measure luminosity accurately is crucial when calculating the probabilities, or *cross-sections*, of various physics processes. It is easy to understand why — A process cross-section is obtained by dividing the number of collisions in which the process occurred by the total number of collisions. Precise luminosity measurements are also needed when modelling an experiment using Monte Carlo simulation, since meaningful predictions can be made only after normalizing the number of collisions generated in the simulation to the number of collisions recorded in real life.

After completing my MSc, I eventually returned to ATLAS and LUCID as a PhD student. This was in December 2010. Like all PhD students, I was given two tasks: To qualify as an ATLAS author by completing a one-year technical project and to carry out a physics analysis for my dissertation. On the technical side, it was natural for me to continue working with luminosity. I spent the first year of my PhD at CERN, where I took over the development of an application called the Online Luminosity Calculator (OLC). This is a piece of C++ code that runs in the ATLAS online environment (the TDAQ for those in the know). It collects information from various subdetectors, from which it calculates luminosity and generates beam diagnostics. The OLC was impressive, to be sure, but not without its flaws. Let me illustrate this point with an example — My very first job was to implement the formula used by LUCID to convert a raw event count into a luminosity. This functionality, although basic, was not in the program. Or as my supervisor famously put it, *the bloody thing can't even calculate a logarithm!* I began by restructuring the way OLC handles luminosity calibration and went on to implement many new functionalities, such as a customizable alarm framework and various plug-ins to correct the lumi measurement online. The way that ATLAS calculates its official luminosity [1] relies heavily on this infrastructure.

In parallel, I continued to work with the LUCID detector. My position as the OLC developer put me in a good spot to try out new, experimental calibration techniques. Using LUCID's private data stream, I developed a data-driven calibration method to tackle the non-linearities experienced by the detector at high luminosities. Mere logarithms are yesterday's news — Now we solve arbitrary equations using numerical methods instead! If I managed to pique your interest, the analysis is documented in Section 7.

Even after completing the technical project and leaving CERN, I kept being *the guy* who knew how to work with LUCID's private data. This data comes straight from the detector electronics and is mostly used for diagnostic purposes. Whenever something was odd about the lumi measurement, I had a look at

the private data stream to understand what was going on. I think many physicists would agree with me that it is these kinds of cross-checks, trying to *understand what's going on*, that really take up most of our time. Sadly (or perhaps thankfully), they usually do not lead to very interesting documentation, which is why no one will ever read about them in any dissertations. However, I would come to use the private data stream for one last lumi analysis interesting enough to present in its own right. Section 8 describes a new calibration method designed to be used during *run II* of the LHC, i.e. when the accelerator ramps up close to its nominal collision energy starting from spring 2015. The method was developed and evaluated using lower-energy *run I* data as recorded by LUCID in 2012.

Upon returning home to Lund late 2011 as a newly qualified ATLAS author, I began to work on my second task: The physics analysis. My choice fell upon the search for supersymmetry. In particle physics, the currently accepted theoretical framework is called the *Standard Model*. Supersymmetry is an extension of this framework. What supersymmetry postulates is that every particle in the Standard Model has a yet-to-be-discovered partner that differs only in terms of mass and spin. The challenge is to figure out what kind of signature those new particles would give rise to in the ATLAS detector and then *find* them. I joined forces with a handful of other physicists in designing a search for the electroweak production of sleptons (a type of new particle) in final states with exactly two leptons. We soon merged our analysis with another team of physicists who were looking for the electroweak production of gauginos (another type of new particle), also in final states with exactly two leptons. Thus was born the electroweak 2-lepton group.

I quickly discovered that working in the 2-lepton group was very different from working in the luminosity group. While I enjoyed many fruitful (and essential) discussions with my luminosity colleagues, I always carried out *the analyses themselves* entirely on my own. Not so in the 2-lepton group. We had dozens of people working on the same search, often with multiple people reproducing each other's results. This redundancy was not due to bad communication between the physicists. Rather, it served as a rigorous cross-check of the analysis. There is no better way to expose bugs and inconsistencies than two people independently trying to come up with the same result!

The 2-lepton analysis spawned several papers throughout its lifetime, including two journal publications [2, 3] and an article that I wrote for the proceedings of the Hadron Collider Physics Symposium (HCP) in Kyoto 2012 [4]. Following my credo that *the best way to learn something is to do it*, it was my desire to be involved in many different facets of the analysis. My first task was the evaluation of the systematic uncertainty of the diboson background, a responsibility that I retained throughout the lifetime of the project. As I found programming to be one of the aspects of the analysis that I enjoyed most, I went on to write software packages for use by the other members of the group. I created a package implementing our trigger matching scheme, which I later generalized and shared with the 3-lepton group. For a time I was also responsible for the trigger reweighting code, to which I made major contributions, but as my focus shifted I eventually handed the package over to a colleague. Wishing to broaden my horizons, I went on to provide the background estimate for the WW process and contributed to the development of its control region. I also cross-checked the final results of the analysis by reproducing the limit calculations. The 2-lepton search is described in Section 9. Rather than starting from an outdated analysis and listing the changes made over time, I have chosen to present the analysis directly in its final form. At the time of writing, the results presented in this text (and in the journal publication, of course) represent the strongest ATLAS 2-lepton limits on electroweak supersymmetry production to date. They supersede any results from previous iterations of the analysis.

The remainder of this thesis provides the background and context necessary to understand the analyses presented in Sections 7 to 9. I begin by discussing some particle physics theory. The Standard Model is introduced in Section 1 while supersymmetry is covered in Section 2. The subsequent sections are more technical. Section 3 briefly introduces the LHC. Section 4 presents the ATLAS experiment, with Section 5 discussing the LUCID detector in particular. Luminosity is then covered in Section 6.

Acknowledgements

The world of particle physics can be unforgiving at times. Be it due to perplexing theories, complicated experimental techniques or the sheer workload of an analysis that simply *has* to be finished before the next conference, no one manages to get by entirely on their own. That's why we collaborate. The best thing about working in a big collaboration is that when in peril, you can always count on your colleagues and friends to help you. I have relied on many people over the course of my PhD. First among them is my thesis project supervisor, Vincent Hedberg. Throughout my stay in ATLAS, Vincent's skill and no-nonsense approach to science have always been a source of inspiration. I could not have asked for a more dependable advisor.

In the LUCID group, I have had the pleasure of working with many of Bologna's finest. There's Carla Sbarra, who taught me everything there is to know (and then some) about the LUCID data stream. There's Benedetto Giacobbe, whose insights and input proved invaluable during the development of new calibration algorithms. There's Davide Caforio, who tirelessly watched over the LUCID running conditions and reconfigured them for the needs of my analyses, and there's Marco Bruschi, whose knowledge of the detector hardware is unmatched! You, and everyone else in the group, made luminosity fun.

When I joined the 2-lepton analysis, I knew almost nothing about supersymmetry or how to search for it. I'm grateful to the many people who, over the course of countless weekly meetings, forum threads and email conversations, patiently answered all of my excruciatingly detailed questions. There are much too many of you to even begin listing. However, I would especially like to thank Serhan Mete and Geraldine Conti, who in their roles as the editors of our last paper had to burden the brunt of my ignorance.

Most of my PhD was spent at the Lund University division of particle physics. This section would not be complete without a tribute to my fellow PhD students and friends. To Anthony, my trusty partner in crime. To Lene, who knows everything. And to Sasha, who is the only one worthy of inheriting LUCID now that I'm gone. To Tuva, Martin and Vytautas from the ALICE office, although ALICE is *clearly* the inferior experiment. The dinners, board games and random discussions that we shared were some of the best times of my life. To everyone at the division — You were the reason I could go to work every day with a smile on my face!

Last but not least, I want to thank my family, who is always there when I need them. And now the reader must beware, for the rest of this text will be physics.

1 The Standard Model

1.1 Overview

The Standard Model (SM) of particle physics is a theory that describes the currently known fundamental particles and how they interact. A particle is characterized by a set of quantum numbers, such as spin and electric charge, that define how the particle behaves. A distinction can be made between *fermions*, which carry half-integer spin, and *bosons*, which carry integer spin. Fermions are the fundamental building blocks of matter. They bind together to form atoms, which in turn make up the macroscopic objects encountered in everyday life. Bosons are responsible for mediating the forces that bind the fermions to each other. The Standard Model describes three such forces. They are called the electromagnetic force, the strong force and the weak force. The force of gravity, although it is known to exist in nature, is not accounted for. This section briefly discusses each particle and force that is present in the SM. While the overview should be accessible to anyone regardless of prior knowledge, subsequent sections require the reader to have some familiarity of four-vectors and Lagrangians in particle physics. More complete discussions can be found in many documents and textbooks, e.g. Refs. [5–7].

There are two basic types of fermions. They are called leptons and quarks. Table 1 summarizes their masses and quantum numbers [8]. The layout of the table emphasizes the generational structure of the particle families. In total, there are six types (or *flavors*) of leptons and six types of quarks. They are divided into three generations with each generation containing a pair of leptons and a pair of quarks. The second generation is simply a copy of the first generation in that their particles have identical quantum numbers and therefore interact via the same forces. The only difference is that the particles in the second generation are heavier. In the same fashion, the third generation is a heavier copy of the other two.

Table 1: The table shows the fermions of the Standard Model. The mass and electric charge of each particle is also shown. Neutrinos are massless in the SM, and the numbers reflect experimental upper limits. For each particle listed here, the model also includes an antiparticle with opposite quantum numbers.

	Leptons				Quarks		
Generation	Flavor	Charge [e]	Mass		Flavor	Charge [e]	Mass
I	e	-1	0.511 MeV		d	$-1/3$	$4.8^{+0.7}_{-0.3}$ MeV
	ν_e	0	< 2 eV		u	$2/3$	$2.3^{+0.7}_{-0.5}$ MeV
II	μ	-1	105.7 MeV		s	$-1/3$	95 ± 5 MeV
	ν_μ	0	< 0.19 MeV		c	$2/3$	1275 ± 25 MeV
III	τ	-1	1777 MeV		b	$-1/3$	4180 ± 30 MeV
	ν_τ	0	< 18.2 MeV		t	$2/3$	173.5 ± 0.6 GeV

The leptons of the first generation are the electron e and its associated neutrino ν_e . The electron is electrically charged. It was long thought that the charge carried by an electron was elementary, and charges in particle physics are therefore measured in units of e . The electron itself carries a charge of $-e$. In contrast, the neutrino is electrically neutral. It follows that the electron interacts via the electromagnetic force while the neutrino does not. They both interact via the weak force but not via the strong force. The remaining leptons are the muon μ , the muon neutrino ν_μ , the tau τ and finally the tau neutrino ν_τ .

The quarks of the first generation are the up quark u and the down quark d . As it turns out, quarks

carry electric charge in fractions of e . The charge of the up quark is $\frac{2}{3}e$ and the charge of the down quark is $-\frac{1}{3}e$. Uniquely among Standard Model particles, quarks interact via every force. The remaining quarks are the charm quark c , the strange quark s , the top quark t and finally the bottom quark b . In addition to every fermion mentioned here, the Standard Model includes an antiparticle with opposite quantum numbers. For example, the antiparticle of the electron is the positron with charge $+e$. Particles and antiparticles have the same mass.

The fundamental forces of nature are mediated by force particles. Table 2 summarizes their properties. Each force particle couples to a certain kind of charge. The photon γ , which is the carrier of the electromagnetic force, couples to electric charge. Any electrically charged particle may emit or absorb photons. This results in a change of momentum, which is equivalent to a force. The strong force is mediated by the gluon g . It couples to color charge. For technical reasons to be discussed in Section 1.4, there are eight distinct kinds of gluons. The weak force is mediated by three different particles denoted W^+ , W^- and Z . They couple to a type of charge called weak isospin. As discussed in Section 1.5, the weak force does not conserve flavor. For example, an electron that emits a W^- will turn into an electron neutrino in the process. In addition to the force carriers, which are all spin 1 bosons, the Standard Model contains a spin 0 boson that does not mediate a force. This particle is called the Higgs. It is discussed separately in Section 1.5.4.

Table 2: The table shows the force carriers of the Standard Model. They are all spin 1 bosons.

Boson	Force	Couples to	Mass [GeV]
γ	Electromagnetic	Electric charge	0
g	Strong	Color charge	0
W^+	Weak	Weak isospin	80
W^-	Weak	Weak isospin	80
Z	Weak	Weak isospin	91

1.2 Gauge theories

1.2.1 Lagrangians

The Standard Model is written in the language of Lagrangians. A Lagrangian L is a function from which the equations of motion for a given system can be derived. It is instructive to start by considering the Lagrangian of classical mechanics.

$$L = T - V \quad (1)$$

This Lagrangian describes a system, such as a particle, with kinetic energy T and potential energy V . The potential energy of the particle depends on its position, and the kinetic energy of the particle depends on its velocity. It follows that L must be a function of the particle's coordinates q_i as well as their time derivatives $\dot{q}_i$. The equations of motion are obtained by plugging L into the Euler-Lagrange equation [9]:

$$\frac{d}{dt} \left(\frac{\partial L}{\partial \dot{q}_i} \right) = \frac{\partial L}{\partial q_i} \quad (i = 1, 2, 3) \quad (2)$$

The calculation is straightforward. Starting with the x coordinate ($i = 1$), the kinetic energy of the particle is

$$T = \frac{mv_x^2}{2} \quad (3)$$

so that

$$\frac{\partial L}{\partial \dot{q}_1} = \frac{\partial T}{\partial v_x} = mv_x \quad (4)$$

and

$$\frac{\partial L}{\partial q_1} = -\frac{\partial V}{\partial x}. \quad (5)$$

Inserting Eqs. (5) and (4) into Eq. (2) yields

$$m \frac{dv_x}{dt} = -\frac{\partial V}{\partial x} \quad (6)$$

which is simply the x -component of Newton's second law, $\mathbf{F} = m\mathbf{a}$. The calculation can be repeated for the remaining coordinates to obtain the y - and z -components. Thus the Lagrangian formulation of classical mechanics leads to the familiar Newtonian equations of motion.

The Standard Model is a relativistic quantum field theory. Whereas the classical Lagrangian deals with localized particles, the SM Lagrangian deals with *fields* that occupy some region of space. Particles are treated as excitations of the underlying fields. A field ψ is a function of both position and time. There are different kinds of fields. For example, a scalar field associates each point in space with a scalar value (such as the temperature of a room or the pressure in a tank of water) while a vector field associates each point in space with a vector (such as the magnitude and direction of a gravitational force). The Standard Model Lagrangian $\mathcal{L}$, which is technically a Lagrangian *density* but will be referred to as a Lagrangian for brevity, describes a system of fields. It is a function of the fields ψ_i and their x, y, z and t derivatives.

Unlike classical mechanics, a relativistic theory makes no sharp distinction between time and space. The time coordinate is treated on equal footing with the space coordinates. In the four-vector formalism, the derivatives are given by

$$\partial_\mu \psi_i \equiv \frac{\partial \psi_i}{\partial x^\mu} \quad (7)$$

and the Euler-Lagrange equation generalizes to:

$$\partial_\mu \left(\frac{\partial \mathcal{L}}{\partial (\partial_\mu \psi_i)} \right) = \frac{\partial \mathcal{L}}{\partial \psi_i} \quad (i = 1, 2, 3, 4) \quad (8)$$

The different parts of the Standard Model Lagrangian will be discussed in Sections 1.3, 1.4 and 1.5.

1.2.2 Gauge invariance

The Lagrangian for a given system is not uniquely defined. Multiplying $\mathcal{L}$ by an arbitrary constant leads to the same equations of motion. In the same manner, adding an arbitrary constant to $\mathcal{L}$ does not change the final result. The Lagrangian is said to be *invariant* under such transformations. A transformation under which the Lagrangian is invariant is called a *gauge* transformation.

For example, consider a Lagrangian $\mathcal{L}$ that is a function of a field ψ and its derivatives. Then $\mathcal{L}$ is invariant under the transformation

$$\psi \rightarrow \psi' = e^{-i\alpha} \psi \quad (9)$$

where α is any real number. This can be shown by plugging ψ' into Eq. (8). The phase factor $e^{-i\alpha}$ (which is just a constant) factors out of the derivative according to

$$\partial_\mu \psi' = \partial_\mu (e^{-i\alpha} \psi) = e^{-i\alpha} \partial_\mu \psi \quad (10)$$

which leaves Eq. (8) unchanged.

$$\partial_\mu \left(\frac{\partial \mathcal{L}}{\partial (\partial_\mu \psi'_i)} \right) = \frac{\partial \mathcal{L}}{\partial \psi'_i} \Leftrightarrow \frac{1}{e^{-i\alpha}} \cdot \partial_\mu \left(\frac{\partial \mathcal{L}}{\partial (\partial_\mu \psi_i)} \right) = \frac{1}{e^{-i\alpha}} \cdot \frac{\partial \mathcal{L}}{\partial \psi_i} \Leftrightarrow \partial_\mu \left(\frac{\partial \mathcal{L}}{\partial (\partial_\mu \psi_i)} \right) = \frac{\partial \mathcal{L}}{\partial \psi_i} \quad (11)$$

Equation (9) is called a *global* phase transformation because it affects ψ in the same way at every point in spacetime. In a *local* phase transformation, the phase α is allowed to vary with the spacetime coordinates x_μ .

$$\psi \rightarrow \psi' = e^{-i\alpha(x_\mu)}\psi \quad (12)$$

The Lagrangian is not necessarily invariant under local phase transformations because of the (now non-vanishing) extra term in the derivative of α .

$$\partial_\mu \psi' = \partial_\mu (e^{-i\alpha}\psi) = e^{-i\alpha}\partial_\mu \psi - i(\partial_\mu \alpha)e^{-i\alpha}\psi \quad (13)$$

However, by adding terms to $\mathcal{L}$ that exactly cancel out the extra terms picked up from the derivative, the invariance can be restored. The added terms would then represent other fields interacting with ψ . Incidentally, it turns out that local gauge invariance is realized in nature, and the “other fields” describe the forces through which ψ interacts. This is fittingly called the *principle of local gauge invariance*. Sections 1.3, 1.4 and 1.5 show how Lagrangians describing the forces of nature can be obtained by starting from the free Lagrangian for a field and demanding that it be invariant under local phase transformations.

1.3 Quantum electrodynamics

The theory of quantum electrodynamics (QED) describes how electrically charged particles interact via the electromagnetic force. Because QED is the mathematically simplest part of the Standard Model, it lends itself to a slightly more in-depth discussion. This section will introduce many mathematical tools and concepts that are used when describing fundamental forces in general.

1.3.1 The QED Lagrangian

The complete QED Lagrangian can be derived by starting from the Lagrangian for a free spinor (spin $\frac{1}{2}$) field ψ and imposing local gauge invariance. The starting Lagrangian is called the Dirac Lagrangian, and spin $\frac{1}{2}$ particles are sometimes referred to as Dirac particles. The Dirac Lagrangian is given by

$$\mathcal{L}_{\text{Dirac}} = i\bar{\psi}\gamma^\mu\partial_\mu\psi - m\bar{\psi}\psi \quad (14)$$

where $\bar{\psi}$ is the adjoint of ψ and γ^μ are the Dirac matrices.[†] Under the local phase transformation

$$\psi \rightarrow \psi' = e^{-i\alpha(x_\mu)}\psi \quad (\bar{\psi} \rightarrow \bar{\psi}' = e^{i\alpha(x_\mu)}\bar{\psi}) \quad (15)$$

the Dirac Lagrangian transforms according to

$$\mathcal{L}_{\text{Dirac}} \rightarrow \mathcal{L}'_{\text{Dirac}} = \mathcal{L}_{\text{Dirac}} + (\bar{\psi}\gamma^\mu\psi)\partial_\mu\alpha \quad (16)$$

where the new term is produced by the derivative analogously to Eq. (13). Note that the phase transformation does nothing to the mass term since $e^{i\alpha}e^{-i\alpha} = 1$. The appearance of an additional term in Eq. (16) means that the free Lagrangian is *not* invariant under local phase transformations. However, the principle of local gauge invariance *demand*s that any real Lagrangian be invariant under Eq. (15). In order to obtain an invariant Lagrangian, something must be added to cancel out the additional term. Consider the new Lagrangian

$$\mathcal{L}_{\text{QED}} = [i\bar{\psi}\gamma^\mu\partial_\mu\psi - m\bar{\psi}\psi] - e(\bar{\psi}\gamma^\mu\psi)A_\mu \quad (17)$$

[†]The Dirac matrices incorporate information about the spin of the particles. Their definition is outside the scope of this text but can be found in any introductory textbook to particle physics.

where e is the electric charge of the interacting particle and A_μ is a vector (spin 1) field that transforms with Eq. (15) according to:

$$A_\mu \rightarrow A'_\mu = A_\mu + \frac{1}{e} \partial_\mu \alpha \quad (18)$$

Strictly speaking, every instance of e should be replaced by Qe to accommodate for the fractional charges of quarks, but the Q will be omitted for simplicity. The new Lagrangian in Eq. (17) is locally invariant because the additional term in Eq. (16) is cancelled by the additional term in Eq. (18). Written out explicitly, $\mathcal{L}_{\text{QED}}$ transforms according to:

$$\mathcal{L}_{\text{QED}} \rightarrow \mathcal{L}'_{\text{QED}} = \mathcal{L}_{\text{QED}} + (\bar{\psi} \gamma^\mu \psi) \partial_\mu \alpha - e (\bar{\psi} \gamma^\mu \psi) \frac{1}{e} \partial_\mu \alpha = \mathcal{L}_{\text{QED}} \quad (19)$$

Putting the equations into words, it is possible to obtain a locally invariant Lagrangian by adding a new vector field A_μ to the theory. This field is nothing but the electromagnetic potential, and the corresponding particle is the photon. The added term, $e (\bar{\psi} \gamma^\mu \psi) A_\mu$, says that Dirac particles (such as electrons and quarks) can couple to photons, and the strength of that coupling is e . The appearance of photons in the theory was a direct consequence of applying the principle of local gauge invariance. However, because there is now a new particle (the photon) present in the Lagrangian, another term must be added that describes how it behaves.

A free vector field A^μ is described by the Proca Lagrangian

$$\mathcal{L}_{\text{Proca}} = -\frac{1}{4} F^{\mu\nu} F_{\mu\nu} + \frac{1}{2} m_A^2 A^\nu A_\nu \quad (20)$$

where the electromagnetic field strength tensor

$$F^{\mu\nu} \equiv (\partial^\mu A^\nu - \partial^\nu A^\mu) \quad (21)$$

is a convenient shorthand. Naively, completing the Lagrangian for quantum electrodynamics should simply be a matter of adding $\mathcal{L}_{\text{Proca}}$ to the present version of $\mathcal{L}_{\text{QED}}$. But doing so would yield a Lagrangian that is not locally gauge invariant because $A^\nu A_\nu$ is not invariant under Eq. (18). The only way to retain all the pieces of the theory while obeying the principle of local gauge invariance is to set $m_A = 0$ so that the offending term vanishes. In other words, the theory requires photons to be massless. This agrees with experimental observations. The complete Lagrangian for quantum electrodynamics is therefore:

$$\mathcal{L}_{\text{QED}} = i \bar{\psi} \gamma^\mu \partial_\mu \psi - m \bar{\psi} \psi - \frac{1}{4} F^{\mu\nu} F_{\mu\nu} - e (\bar{\psi} \gamma^\mu \psi) A_\mu \quad (22)$$

As shown by Eq. (13), the difference between local and global phase transformations is the appearance of an extra term in the derivatives of the fields. In this chapter, it was shown that the extra term can be cancelled out by adding a new vector field A_μ to the Lagrangian. A convenient way of including A_μ in the theory is by defining the covariant derivative

$$\mathcal{D}_\mu = \partial_\mu + ie A_\mu \quad (23)$$

and then replacing every instance of ∂_μ with $\mathcal{D}_\mu$ in the free Lagrangian. In fact, it can be shown that a Lagrangian written in terms of covariant derivatives will always be locally gauge invariant. In terms of a covariant derivative, then, the complete Lagrangian for quantum electrodynamics reads:

$$\mathcal{L}_{\text{QED}} = i \bar{\psi} \gamma^\mu \mathcal{D}_\mu \psi - m \bar{\psi} \psi - \frac{1}{4} F^{\mu\nu} F_{\mu\nu} \quad (24)$$

The discussion of QED offers a good place to introduce some of the group theory that is used to describe the Standard Model. The phase transformations under which the QED Lagrangian is invariant can be written

$$\psi \rightarrow \psi' = U \psi \quad (25)$$

where U is any 1×1 matrix such that $U^\dagger U = 1$ (thereby ensuring that the overall normalization of the wavefunction is preserved). Matrices that fulfill this requirement are referred to as *unitary*, and the group of all such matrices is called $U(1)$. QED is therefore said to be symmetric under $U(1)$ gauge transformations, and $U(1)$ is referred to as the *symmetry group* of QED.

1.3.2 Vacuum polarization

Consider two electrically charged particles that are interacting with each other through the exchange of photons. According to QED, the space separating those particles is filled with virtual electron-positron pairs which are produced by the electromagnetic field as shown in Figure 1. Each such electron-positron pair acts as an electric dipole. Since the dipoles exist within an external electromagnetic field, they will tend to align themselves with that field. This phenomenon is called *vacuum polarization*.

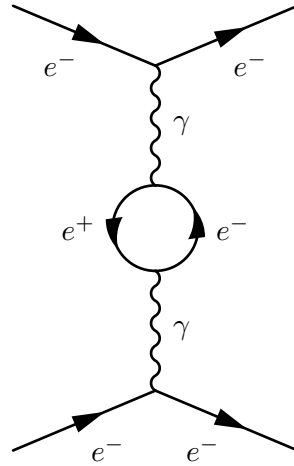


Figure 1: An electron emits a photon that fluctuates into a virtual electron-positron pair before being absorbed by another electron. The virtual loop could in principle contain any electrically charged particle-antiparticle pair.

The polarization of the vacuum leads to a screening effect — When a positive charge q is observed from a distance, there is on average more negative than positive charge in the space separating q and the observer. A reduced charge $q_{\text{eff}} < q$ is therefore observed. It follows that the strength of the electromagnetic force depends on the separation of the charges. The strength is typically expressed in terms of a coupling constant

$$\alpha_e(q^2) = \frac{\alpha_e(0)}{1 - \frac{\alpha_e(0)}{3\pi} \cdot \ln\left(\frac{|q^2|}{\Lambda^2}\right)} \quad (|q^2| \gg \Lambda^2) \quad (26)$$

where $\alpha_e(0)$ and Λ are constants, and q^2 is the momentum transfer between the two participating particles. It is clear that the strength of the electromagnetic force increases as two charges are brought closer together (larger $|q^2|$). Calling α_e a “constant” is therefore somewhat of a misnomer, and sometimes the term *running* coupling constant is used instead. The relation between the coupling constant and the fundamental electric charge is $e = \sqrt{4\pi\alpha_e}$.

1.4 Quantum chromodynamics

The theory of quantum chromodynamics (QCD) describes how particles that carry color charge interact via the strong force. This section will introduce color in the context of the quark model, discuss the QCD Lagrangian and finally describe the behavior of the strong force at long ranges.

1.4.1 The quark model

Atomic nuclei are made up of protons and neutrons. These protons and neutrons are not elementary particles. They are bound states of quarks [10] (see Table 1 for a list). The quantum numbers of a bound state are obtained by summing or multiplying the quantum numbers of the constituent quarks as appropriate. For example, the neutron consists of an up quark and two down quarks (udd). Its electric charge becomes $\frac{2}{3} - \frac{1}{3} - \frac{1}{3} = 0$ as expected.

In the early days of the quark model, it was difficult to account for the experimental observation of bound states consisting of three same-flavor quarks. The Pauli exclusion principle states that no two fermions may occupy the same quantum state simultaneously. Quarks, being fermions, must follow this principle. Particles like the neutron (udd) were fine. The two down quarks are not necessarily identical because they can have opposite spin. But the observation of particles such as the Δ^{++} (uuu) presented a problem. The solution was to introduce an entirely new kind of quantum number called *color*. If the quarks that make up the Δ^{++} have different colors, they are no longer in the same quantum state and the Pauli exclusion principle has no objection.

The color hypothesis [11] states that every quark carries a color charge. They can be either red, green or blue. Antiquarks carry the respective anticolors. Furthermore, all naturally occurring particles must be *colorless*. A particle is considered colorless if it contains either no net color or all colors in equal amounts. Particles made up of two quarks ($q\bar{q}$) are called *mesons*. In mesons, the colors carried by the quark and antiquark (e.g. red and antired) cancel out. Particles made up of three quarks (qqq) or three antiquarks ($\bar{q}\bar{q}\bar{q}$), one of each color, are called *baryons*. Mesons and baryons are collectively referred to as *hadrons*.

Although describing hadrons as bound states of two or three quarks is a useful simplification, it is not correct on a more fundamental level. The quarks stay bound by constantly emitting and absorbing gluons, and each such gluon can fluctuate into a virtual quark-antiquark pair before being reabsorbed. At any given moment in time, a hadron contains a whole sea of gluons and quarks that are continuously popping in and out of existence. Those quarks and gluons are collectively referred to as *partons*. Realistic models describe hadrons as collections of partons. The partons share the total energy of the hadron, and the fractional momentum carried by each type of parton is given by a *parton distribution function* (PDF).

1.4.2 The QCD Lagrangian

The complete QCD Lagrangian can be derived by starting from the free-quark Lagrangian and imposing local gauge invariance. This procedure is analogous to the one that was used to derive the QED Lagrangian. As explained in Section 1.4.1, there are three distinct versions of each quark flavor corresponding to three different colors. They can be arranged into a color triplet. It will be shown that the strong force allows quarks to change their color, or put differently, convert from one member of the color triplet to another. When written out explicitly, the free Lagrangian for a particular flavor is the sum of three one-particle Lagrangians.

$$\mathcal{L}_{\text{QCD}} = \left[i\bar{\psi}_r \gamma^\mu \partial_\mu \psi_r - m\bar{\psi}_r \psi_r \right] + \left[i\bar{\psi}_g \gamma^\mu \partial_\mu \psi_g - m\bar{\psi}_g \psi_g \right] + \left[i\bar{\psi}_b \gamma^\mu \partial_\mu \psi_b - m\bar{\psi}_b \psi_b \right] \quad (27)$$

where r, g and b denote red, green and blue. It is possible to write $\mathcal{L}_{\text{QCD}}$ more compactly by introducing the triplets

$$\Psi = \begin{pmatrix} \psi_r \\ \psi_g \\ \psi_b \end{pmatrix} \quad \text{and} \quad \bar{\Psi} = (\bar{\psi}_r \quad \bar{\psi}_g \quad \bar{\psi}_b) \quad (28)$$

so that:

$$\mathcal{L}_{\text{QCD}} = i\bar{\Psi} \gamma^\mu \partial_\mu \Psi - m\bar{\Psi} \Psi \quad (29)$$

This three-particle quark Lagrangian looks exactly like the one-particle Dirac Lagrangian from Section 1.3.1. Just as the one-particle Lagrangian is symmetric under global U(1) gauge transformations, the three-particle Lagrangian is symmetric under global U(3) gauge transformations. In other words, it is invariant under transformations of the form

$$\Psi \rightarrow \Psi' = U\Psi \quad (\bar{\Psi} \rightarrow \bar{\Psi}' = \bar{\Psi}U^\dagger) \quad (30)$$

where U is any 3×3 matrix such that $U^\dagger U = 1$. As usual, the physics emerges upon demanding that the global invariance should hold locally. The first step is to rewrite Eq. (30) in a more familiar way. Any unitary matrix U can be written on the form [12]

$$U = e^{iH} \quad (31)$$

where H is *Hermitian*, meaning it fulfills $H^\dagger = H$. Equation (31) is essentially a generalization of the statement that any complex number can be written on the form $e^{i\theta}$ where θ is real. In the case of QCD, U and H are both 3×3 matrices. The most general Hermitian 3×3 matrix can be expressed in terms of eight real numbers $a_1, a_2, \dots, a_8$ and a phase factor θ :

$$H = \theta 1 + \lambda \cdot \mathbf{a} \quad (32)$$

where 1 is the 3×3 unit matrix and $\lambda_1, \lambda_2, \dots, \lambda_8$ are the Gell-Mann matrices.[†] A word on the notation is in order. In this text, the dot product $\lambda \cdot \mathbf{a}$ is used interchangeably with the repeated index notation $\lambda_j a_j$ depending on which is more convenient. They both represent the same sum.

$$\lambda \cdot \mathbf{a} \equiv \lambda_j a_j \equiv \lambda^j a^j \equiv \lambda_1 a_1 + \lambda_2 a_2 + \dots + \lambda_8 a_8 \quad (33)$$

Plugging (32) into (31) yields:

$$U = e^{i\theta} e^{i\lambda \cdot \mathbf{a}} \quad (34)$$

The factor $e^{i\theta}$ represents the familiar U(1) gauge transformations already explored in Section 1.3.1. The calculations need not be repeated here. What's *new* in QCD is the factor $e^{i\lambda \cdot \mathbf{a}}$. It represents a subgroup of U(3) that contains all unitary 3×3 matrices with determinant 1. This group is the symmetry group of QCD. It is referred to as the “special” unitary group, SU(3). In summary, $\mathcal{L}_{\text{QCD}}$ is invariant under global SU(3) gauge transformations of the form

$$\Psi \rightarrow \Psi' = e^{i\lambda \cdot \mathbf{a}} \Psi \quad (35)$$

and should now be modified so that it becomes invariant also under local SU(3) gauge transformations.

$$\Psi \rightarrow \Psi' = e^{i\lambda \cdot \mathbf{a}(x_\mu)} \Psi \equiv e^{-i\lambda \cdot \mathbf{a}(x_\mu)/2} \Psi \quad (36)$$

Similarly to QED, the solution is to introduce massless vector fields G_μ to cancel the extra terms picked up by the derivatives. Eight new fields are required — One for each of the SU(3) group generators $\lambda_1, \lambda_2, \dots, \lambda_8$. The algebra is more complicated, though, because 3×3 matrices do not commute. In group theoretical language, SU(3) is *non-Abelian*. A detailed discussion of non-Abelian theories is beyond the scope of this text but can be found in e.g. Ref. [13]. The difference with respect to Abelian theories such as QED shows up in the definition of the eight gluon field strength tensors $G^{\mu\nu}$. As discussed below, this turns out to have striking consequences for how the force behaves. The complete QCD Lagrangian becomes:

$$\mathcal{L}_{\text{QCD}} = i\bar{\Psi}\gamma^\mu \mathcal{D}_\mu \Psi - m\bar{\Psi}\Psi - \frac{1}{4}G_j^{\mu\nu}G_{j\mu\nu} \quad (37)$$

[†]The Gell-Mann matrices are the generators of SU(3). Their definition is outside the scope of this text but can be found in any introductory textbook to particle physics.

The field strength tensors are given by

$$G_j^{\mu\nu} \equiv (\partial^\mu G_j^\nu - \partial^\nu G_j^\mu) + g_s f_{jkl} G_k^\mu G_l^\nu \quad (38)$$

where g_s is the fundamental color charge and f_{jkl} are the structure constants of SU(3). Finally, the covariant derivative is defined as:

$$\mathcal{D}_\mu = \partial_\mu + i \frac{g_s}{2} \lambda^j G_\mu^j \quad (39)$$

The vector fields G_μ correspond to the eight gluons that mediate the strong force. They play a role analogous to that of the photon in QED with one important difference. The photon is electrically neutral, and therefore it does not interact with other photons. The gluon, in contrast, carries its own color charge. Gluons can therefore emit and absorb *other* gluons. This is reflected in the QCD Lagrangian, which contains terms of the form:

$$g_s f_{jkl} G_k^\mu G_l^\nu \quad (40)$$

Those terms correspond to gluon-gluon couplings with strength proportional to g_s . Expanding the factor $G_j^{\mu\nu} G_{j\mu\nu}$ in $\mathcal{L}_{\text{QCD}}$ shows that it contains both terms with three factors of G^μ and terms with four factors of G^μ . As illustrated in Figure 2, the three-gluon vertex corresponds to a gluon emitting or absorbing another gluon, and the four-gluon vertex corresponds to elastic gluon scattering.

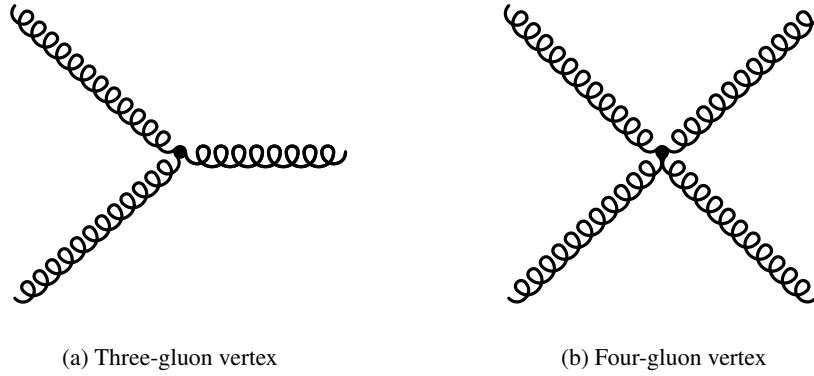


Figure 2: Feynman diagrams illustrating the gluon-gluon couplings present in the QCD Lagrangian.

The fact that gluons carry color charge has one more consequence. When a colored particle emits or absorbs a gluon, the color of that particle can change. The gluon carries away the difference. Gluons therefore carry both color and anticolor simultaneously. For example, if a red quark emits a gluon that is red and antiblue, the quark will turn blue in the process. The color doesn't *have* to change, however. If a red quark emits a gluon that is red and antired, the quark will of course remain red.

1.4.3 Long range behavior

Just as the vacuum is polarized by the electromagnetic field in QED, it is polarized by the color field in QCD. However, due to the presence of gluon-gluon couplings, the effect of the polarization is different. Figure 3 shows two possible processes. Virtual quark loops lead to a screening effect as usual, but virtual gluon loops turn out to have the opposite effect. They result in a competing *antiscreening* effect that causes the color field to appear stronger at greater distances from the source. The strong coupling constant α_s is given by

$$\alpha_s(q^2) = \frac{12\pi}{(11n - 2f) \cdot \ln\left(\frac{|q^2|}{\Lambda_s^2}\right)} \quad (|q^2| \gg \Lambda_s^2) \quad (41)$$

where $n = 3$ is the number of colors present in the Standard Model, $f = 6$ is the number of flavors, Λ_s is a constant and q^2 is the momentum transfer of the interaction. The coupling constant is related to the fundamental color charge by $g_s = \sqrt{4\pi\alpha_s}$. Since $11n > 2f$, the antiscreening effect dominates and the strength of the strong force decreases as the charges are brought closer together. Quarks that are in close proximity of each other, such as the quarks inside of a hadron, behave approximately as free particles. This phenomenon is called *asymptotic freedom*.

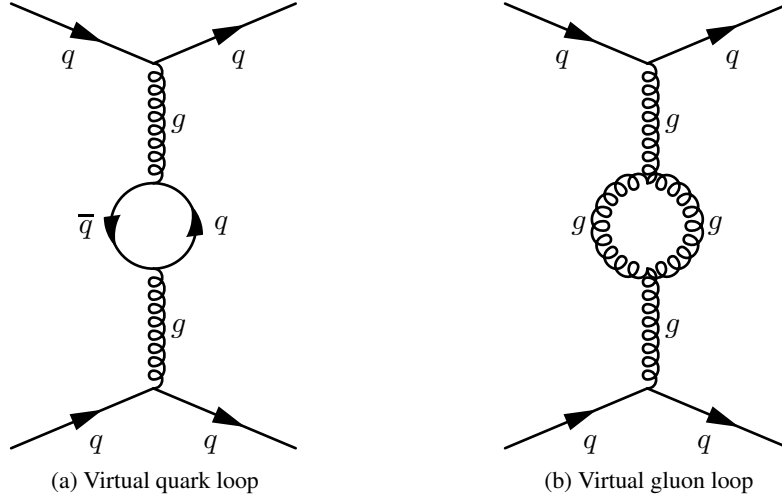


Figure 3: In QCD, virtual quark loops lead to a screening effect while virtual gluon loops lead to an antiscreening effect.

Attempting to create free quarks by separating the quarks that make up a hadron results in new hadrons being produced instead. When the distance between the quarks increases, so does the strength of the force that binds them together. The energy required to overcome the attractive force is stored in the gluon field. Eventually, this energy becomes large enough to produce new hadrons. The process is illustrated in Figure 4. The hadrons typically emerge in narrow showers, called *jets*, along the trajectories of the original quarks. Due to this *hadronization* process, free quarks are never observed experimentally.

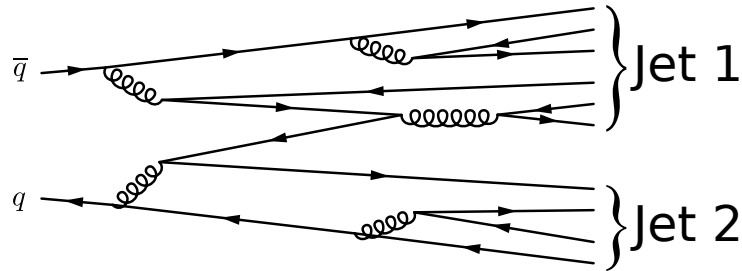


Figure 4: A quark and an antiquark hadronizing into two jets.

1.5 Weak interactions

The weak force makes up the most complicated part of the Standard Model Lagrangian. Weak gauge bosons couple to fermions that carry a quantum number called weak isospin. This section introduces weak isospin in the context of chiral fermions and then goes on to discuss the problems that arise when

naively trying to construct a Lagrangian using the same method as in Section 1.4. Finally, solutions to those problems are presented.

1.5.1 Chiral fermions

The Standard Model distinguishes between particles of different *chirality*. The chirality of a particle determines whether or not it interacts via the weak force. Particles that interact via the weak force have “left-handed” chirality. Conversely, particles that do not interact via the weak force have “right-handed” chirality. Every massive fermion in the Standard Model comes in two versions with opposite chiralities. For example, the Standard Model contains two electrons. The left-handed electron can interact weakly while the right-handed electron cannot. Furthermore, the two particles can transform into one another. This can be understood by expanding the fermion mass term in its left-handed and right-handed components.

$$m\bar{\psi}\psi = m(\bar{\psi}_R\psi_L + \bar{\psi}_L\psi_R) \quad (42)$$

When read as an interaction term (see Section 1.5.7), it represents couplings between the left- and right-handed fermions. The coupling strength is proportional to m . This means that massless fermions have no such coupling. A massless fermion that is created left-handed remains left-handed forever. Since the Standard Model treats neutrinos as massless, and neutrinos are only created in weak interactions, it follows that all neutrinos are left-handed by definition.

1.5.2 Weak iso-doublets

Left-handed fermions are arranged into pairs called weak iso-doublets. Members of a doublet can convert into each other by emitting or absorbing a weak gauge boson. This is, of course, completely analogous to how members of color triplets convert into each other by emitting or absorbing gluons. The lepton doublets L_L and quark doublets Q_L are

$$L_L = \begin{pmatrix} \psi_{\nu_e} \\ \psi_e \end{pmatrix}_L, \begin{pmatrix} \psi_{\nu_\mu} \\ \psi_\mu \end{pmatrix}_L, \begin{pmatrix} \psi_{\nu_\tau} \\ \psi_\tau \end{pmatrix}_L \quad \text{and} \quad Q_L = \begin{pmatrix} \psi_u \\ \psi_{d'} \end{pmatrix}_L, \begin{pmatrix} \psi_c \\ \psi_{s'} \end{pmatrix}_L, \begin{pmatrix} \psi_t \\ \psi_{b'} \end{pmatrix}_L \quad (43)$$

where a subscripted L means left-handed. The symmetry group of the weak force is denoted $SU(2)_L$, where the L serves as a reminder that only left-handed fields transform under this gauge group. Right-handed fermions do not interact via the weak force and are therefore put in iso-singlets. The lepton singlets L_R , up-type quark singlets U_R and down-type quark singlets D_R are given by

$$L_R = (\psi_e)_R, (\psi_\mu)_R, (\psi_\tau)_R, \quad U_R = (\psi_u)_R, (\psi_c)_R, (\psi_t)_R \quad \text{and} \quad D_R = (\psi_d)_R, (\psi_s)_R, (\psi_b)_R. \quad (44)$$

The primed quarks in Eq. (43) denote mixtures of all down type quarks. For example, d' is a mixture of d , s and b . The weak force can therefore convert a u -quark to a d -, s - or b -quark with certain probabilities. Put differently, the quark mass eigenstates are different from the weak eigenstates. They are related by the Cabibbo-Kobayashi-Maskawa (CKM) matrix [14]:

$$\begin{pmatrix} d' \\ s' \\ b' \end{pmatrix} = \begin{pmatrix} V_{ud} & V_{us} & V_{ub} \\ V_{cd} & V_{cs} & V_{cb} \\ V_{td} & V_{ts} & V_{tb} \end{pmatrix} \begin{pmatrix} d \\ s \\ b \end{pmatrix} \quad (45)$$

The probability for a transition to occur is given by the square of the corresponding matrix element. From experiment, the absolute values are [8]:

$$\begin{pmatrix} |V_{ud}| & |V_{us}| & |V_{ub}| \\ |V_{cd}| & |V_{cs}| & |V_{cb}| \\ |V_{td}| & |V_{ts}| & |V_{tb}| \end{pmatrix} = \begin{pmatrix} 0.9743 & 0.2253 & 0.0035 \\ 0.2252 & 0.9734 & 0.0412 \\ 0.0087 & 0.0404 & 0.9991 \end{pmatrix} \quad (46)$$

1.5.3 A weak Lagrangian

Naively trying to construct a Lagrangian for the weak force following the same procedure as for QED and QCD leads to a theory that does not agree with experimental observations. This section describes where in such a naive Lagrangian the disagreement appears, and later sections will describe how those problems are resolved. The starting point is a Lagrangian for SU(2) derived analogously to Eq. (37). Since the weak force treats left-handed and right-handed fermions differently, the first step is to expand the fermionic part of the Lagrangian in terms of ψ_L and ψ_R . Noting that

$$i\bar{\psi}\gamma^\mu \mathcal{D}_\mu \psi + m\bar{\psi}\psi = i\bar{\psi}_L\gamma^\mu \mathcal{D}_\mu \psi_L + i\bar{\psi}_R\gamma^\mu \mathcal{D}_\mu \psi_R + m(\bar{\psi}_R\psi_L + \bar{\psi}_L\psi_R) \quad (47)$$

it is clear that such an expansion would preserve gauge invariance for the propagator term, since it keeps ψ_L and ψ_R separate. The mass term, on the other hand, mixes ψ_L and ψ_R together. This would violate gauge invariance because only left-handed fields transform under SU(2)_L. Setting $m = 0$ restores the invariance by getting rid of the offending term, but experiments clearly show that fermions can have mass. The solution to this problem is to introduce *Yukawa mass terms* as explained in Section 1.5.7.

By analogy with the strong force, the covariant derivative in the naive weak Lagrangian is

$$\mathcal{D}_\mu \equiv \partial_\mu + i\frac{g_2}{2}\tau^j W_\mu^j \quad (48)$$

where g_2 is the fundamental weak charge, and τ_1 , τ_2 and τ_3 are the Pauli matrices.[†] The three vector fields W_μ^1 , W_μ^2 and W_μ^3 represent massless gauge bosons that couple to left-handed fermions. They are the mediators of the weak force. This should be compared to experimental observations, which suggest that the weak force is mediated by three *massive* gauge bosons W^+ , W^- and Z . The problem is solved by the *Higgs mechanism*, which is discussed in Section 1.5.6. Through the Higgs mechanism, the vector bosons can acquire masses in a gauge invariant way. Finally, the mass of the Z boson is different from the mass of W^+ and W^- . The Z boson also appears to couple to right-handed fermions whereas W^+ and W^- do not. These observations are explained by electroweak unification.

1.5.4 Electroweak unification

Although the electromagnetic and weak interactions appear as two distinct phenomena at low energies, they are in fact intimately related. The Standard Model describes them as two manifestations of a single, more fundamental interaction called the *electroweak* force. It is invariant under the gauge group SU(2)_L $\otimes$ U(1)_Y. The SU(2)_L part gives rise to the three W^j bosons. They couple to weak isospin I_3^W . The U(1)_Y part gives rise to another boson, denoted B . It couples to weak hypercharge Y . These bosons then mix together to produce the mass eigenstates that are observed experimentally. The W^+ and W^- are linear combinations of W^1 and W^2 .

$$\begin{aligned} W_\mu^+ &= \frac{1}{\sqrt{2}}(W_\mu^1 - iW_\mu^2) \\ W_\mu^- &= \frac{1}{\sqrt{2}}(W_\mu^1 + iW_\mu^2) \end{aligned} \quad (49)$$

In the same manner, γ and Z are linear combinations of B and W^3 . The associated fields mix according to

$$\begin{aligned} A_\mu &= B_\mu \cos(\theta_w) + W_\mu^3 \sin(\theta_w) \\ Z_\mu &= -B_\mu \sin(\theta_w) + W_\mu^3 \cos(\theta_w) \end{aligned} \quad (50)$$

[†]The Pauli matrices are the generators of SU(2). Their definition is outside the scope of this text but can be found in any introductory textbook to particle physics.

where θ_w is a parameter called the weak mixing angle. The field A_μ is, of course, the electromagnetic potential from Section 1.3. At this point it makes sense that the Z boson has a different mass than W^+ and W^- . The latter contain no B in the mix. It is also clear why the Z boson couples to right-handed particles. It is really the B -component that mediates the interaction, and all right-handed fermions turn out to carry weak hypercharge. Table 3 shows the fermions of the first generation and their associated electroweak quantum numbers. The quantum numbers of the second and third generations are analogous to those of the first.

Table 3: The table shows the weak isospin I_3^W , weak hypercharge Y and electric charge Q of the first-generation fermions. Upper and lower members of weak iso-doublets are assigned a weak isospin of $\frac{1}{2}$ and $-\frac{1}{2}$, respectively. The weak hypercharge is then assigned such that $Q = I_3^W + Y/2$ (see Section 1.5.6).

Fermion	I_3^W	Y	Q
$\begin{pmatrix} \nu_e \\ e \end{pmatrix}_L$	$\frac{1}{2}$	-1	0
$\begin{pmatrix} u \\ d \end{pmatrix}_L$	$-\frac{1}{2}$	-1	-1
e_R	$\frac{1}{2}$	$\frac{1}{3}$	$\frac{2}{3}$
u_R	$-\frac{1}{2}$	$\frac{1}{3}$	$-\frac{1}{3}$
e_R	0	-2	-1
u_R	0	$\frac{4}{3}$	$\frac{2}{3}$
d_R	0	$-\frac{2}{3}$	$-\frac{1}{3}$

The Standard Model does not predict the value of the weak mixing angle θ_w , but it does predict several relations where the angle is involved. The masses of the $W^\pm$ and Z bosons are related according to

$$\cos \theta_w = \frac{m_W}{m_Z} \quad (51)$$

where m_W and m_Z are the $W^\pm$ and Z masses, respectively. The theory also requires the $U(1)_Y$ coupling constant g_1 and the $SU(2)_L$ coupling constant g_2 (analogous to e and g_s for the electromagnetic and strong forces) to be related according to:

$$g_1 \cos(\theta_w) = g_2 \sin(\theta_w) = e \quad (52)$$

The experimental value of θ_w is 28.74° [8].

At low energies, it is electric charge Q , and not I_3^W or Y , that is a good quantum number. Somehow the $SU(2)_L \otimes U(1)_Y$ symmetry is *broken* and manifests itself as the familiar $U(1)$ symmetry of electromagnetism (which is sometimes denoted $U(1)_{EM}$ to emphasize the distinction from $U(1)_Y$). The process by which this happens is called electroweak symmetry breaking. Section 1.5.5 introduces symmetry breaking in general. Section 1.5.6 explains how the electroweak symmetry is broken in the Standard Model.

1.5.5 Spontaneous symmetry breaking

In a quantum field theory, such as the Standard Model, particles are treated as excitations (the quanta) of underlying physical fields. This approach is based on perturbation theory. Lagrangians are written in terms of the underlying fields, and excited states are obtained by expanding the system in terms of fluctuations around its ground state. The ground state, sometimes called the *vacuum*, is the field configuration that yields the minimum energy. Finding this configuration is often trivial — Simply set all fields to zero. This has been the case for all Lagrangians encountered so far. But it doesn't *have* to

be. As a simple but instructive example, consider a scalar (spin 0) field ϕ described by the Lagrangian

$$\mathcal{L}_{\text{Disc}} = T - V = \frac{1}{2} \partial_\mu \phi \partial^\mu \phi - \left(\frac{1}{2} \mu^2 \phi^2 + \frac{1}{4} \lambda \phi^4 \right) \quad (53)$$

where μ and λ are constants. To help guide the eye, $\mathcal{L}_{\text{Disc}}$ has been separated into a kinetic term and a potential term just like in Eq. (1). The kinetic term vanishes provided that ϕ is constant. The minimum of the potential term is found by solving:

$$\frac{\partial V}{\partial \phi} = \phi (\mu^2 + \lambda \phi^2) = 0 \quad (54)$$

The solution to Eq. (54) depends on the sign of μ^2 as well as the sign of λ . However, if $\mathcal{L}_{\text{Disc}}$ is meant to describe a physical system, it must hold that $\lambda > 0$ or else the potential would not be bounded from below as $\phi \rightarrow \infty$. There are no such constraints on μ^2 , which can be either positive or negative. Figure 5 shows the shape of V in each case. If $\mu^2 > 0$, then the system is in its ground state when $\phi = 0$. A more interesting situation arises when $\mu^2 < 0$. This system has two ground states. The minima are located at:

$$\phi_{\text{min}} = \pm \sqrt{\frac{-\mu^2}{\lambda}} \quad (55)$$

This result has profound consequences. When the system is in its ground state, ϕ is not zero. Evidently, the field ϕ exists everywhere in space and has the value ϕ_{min} . This is called the vacuum expectation value of the field.

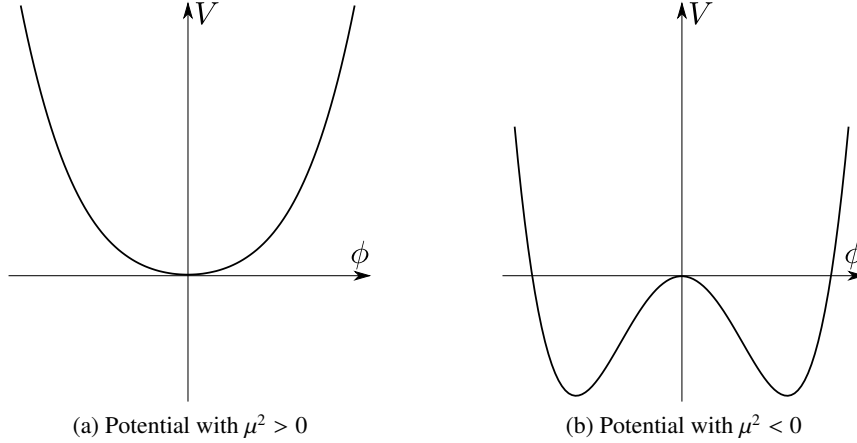


Figure 5: The shape of the potential term in Eq. (53). If μ^2 is positive (a), there is only one minimum. If μ^2 is negative (b), there are two minima. Since they are located off-center, ϕ acquires a vacuum expectation value.

The mass and interactions of the particle described by $\mathcal{L}_{\text{Disc}}$ are found by expanding the Lagrangian in terms of fluctuations around one of its ground states. There are two such states to choose from, but all physics conclusions remain the same no matter what state is used. Picking $v = \sqrt{-\mu^2/\lambda}$, the Lagrangian can be rewritten in terms of a new field η such that:

$$\phi = v + \eta \quad (56)$$

Substituting Eq. (56) into Eq. (53) and using the fact that all terms linear in η must vanish (because the vacuum expectation value for η is zero) yields:

$$\mathcal{L}_{\text{Disc}} = \frac{1}{2} \partial_\mu \eta \partial^\mu \eta - \left(\frac{1}{2} \mu^2 \eta^2 + \lambda v \eta^3 + \frac{1}{4} \lambda \eta^4 \right) + \text{constant} \quad (57)$$

At this point another observation can be made. The original Lagrangian, Eq. (53), is symmetric under transformations of the form

$$\phi \rightarrow -\phi \quad (58)$$

but the modified Lagrangian, Eq. (57), is not. It should be noted that the two Lagrangians describe exactly the same physical system. However, ϕ is not the fundamental particle of the theory. To identify the real particle, $\mathcal{L}_{\text{Disc}}$ must be expanded in terms of fluctuations around its ground state. But there are two ground states. Picking one state over the other *breaks* the symmetry that was present in the original Lagrangian. This is a basic example of the phenomenon called *spontaneous symmetry breaking*. It happens whenever a field in the theory acquires a vacuum expectation value.

A distinction can be made between *discrete* symmetries and *continuous* symmetries. The symmetry exhibited by $\mathcal{L}_{\text{Disc}}$ as written in Eq. (53) is discrete, which is reflected by the fact that the potential has exactly two minima. The remainder of this section will discuss what happens when a continuous symmetry is broken. Consider the Lagrangian

$$\mathcal{L}_{\text{Cont}} = T - V = \frac{1}{2} \partial_\mu \phi^* \partial^\mu \phi - \left(\mu^2 \phi^* \phi + \lambda (\phi^* \phi)^2 \right) \quad (59)$$

where

$$\phi = \frac{\phi_1 + i\phi_2}{\sqrt{2}} \quad (60)$$

is now a complex scalar field. When written on the form in Eq. (59), it is easy to see that the Lagrangian is invariant under global U(1) gauge transformations.

$$\phi \rightarrow \phi' = e^{-i\alpha} \phi \quad (61)$$

The shape of the potential can be understood by expressing $\mathcal{L}_{\text{Cont}}$ in terms of ϕ_1 and ϕ_2 .

$$\mathcal{L}_{\text{Cont}} = \frac{1}{2} \partial_\mu \phi_1 \partial^\mu \phi_1 + \frac{1}{2} \partial_\mu \phi_2 \partial^\mu \phi_2 - \left(\frac{1}{2} \mu^2 (\phi_1^2 + \phi_2^2) + \frac{1}{4} \lambda (\phi_1^2 + \phi_2^2)^2 \right) \quad (62)$$

Only the combination $(\phi_1^2 + \phi_2^2)$ appears in the potential term, so the shape should be circular. Figure 6 shows V as a function of ϕ_1 and ϕ_2 . As before, a negative μ^2 is assumed. The minima lie along a circle that fulfills:

$$\frac{\partial V}{\partial (\phi_1^2 + \phi_2^2)} = \frac{1}{2} \mu^2 + \frac{1}{2} \lambda (\phi_1^2 + \phi_2^2) = 0 \Leftrightarrow (\phi_1^2 + \phi_2^2) = \frac{-\mu^2}{\lambda} \equiv v \quad (63)$$

There is an infinite number of such minima, reflecting the fact that the symmetry exhibited by $\mathcal{L}_{\text{Cont}}$ is continuous.

To find the masses and interactions of the particles described by $\mathcal{L}_{\text{Cont}}$, it should be expanded around one of its ground states. A simple choice is

$$\begin{aligned} \phi_1 &= v \\ \phi_2 &= 0 \end{aligned} \quad (64)$$

so that the Lagrangian should be rewritten in terms of the new fields η and ρ which are defined according to:

$$\begin{aligned} \phi_1 &= v + \eta \\ \phi_2 &= \rho \end{aligned} \quad (65)$$

Inserting Eq. (65) into Eq. (62) yields:

$$\mathcal{L}_{\text{Cont}} = \frac{1}{2} \partial_\mu \eta \partial^\mu \eta + \frac{1}{2} \partial_\mu \rho \partial^\mu \rho + \mu^2 \eta^2 - \lambda v (\eta \rho^2 + \eta^3) - \frac{1}{2} \lambda \eta^2 \rho^2 - \frac{1}{4} \lambda \eta^4 - \frac{1}{4} \lambda \rho^4 + \text{constant} \quad (66)$$

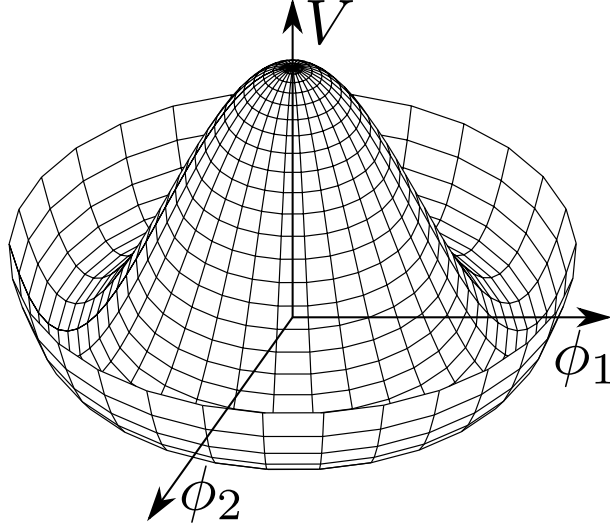


Figure 6: The shape of the potential term in Eq. (62) assuming $\mu^2 < 0$. There is an infinite number of minima located in a ring around the center.

This Lagrangian deserves a closer look. The first two terms represent the kinetic energy of η and ρ . The third term reveals that the particle corresponding to the η -field has a mass given by[†]:

$$m_\eta = \sqrt{-2\mu^2} \quad (67)$$

All remaining terms contain factors of both η and ρ . They are therefore interaction terms of various strengths. The key point is that the Lagrangian does not contain a mass term for ρ . Evidently, the corresponding particle must be massless. This is no coincidence. Goldstone's theorem [15] states that whenever a continuous global symmetry is spontaneously broken, the theory will contain a massless scalar boson (called a *Goldstone boson*) for each broken generator of the original symmetry group. In the present case, the original Lagrangian was symmetric under global U(1) gauge transformations. This symmetry was spontaneously broken when ϕ_1 acquired a vacuum expectation value.

Figure 6 invites an intuitive physical interpretation of the Goldstone boson. In the above discussion, the vacuum was chosen according to Eq. (64). At the point $(v, 0)$, there is no resistance to excitations in the ρ direction because all the points along the ring of minima are at the same potential. Goldstone bosons therefore correspond to transitions between different ground states. They are massless because all ground states are at the same potential. This has a special implication if the theory happens to be locally gauge invariant. Choosing the vacuum is equivalent to choosing a gauge. Goldstone bosons can transform the system into states that are degenerate with the vacuum. However, since the theory is locally gauge invariant, it is always possible to apply a gauge transformation that takes the system back to its original state. It is often convenient to choose the vacuum such that all fields that correspond to Goldstone bosons vanish and then fix the gauge. This procedure is implicitly followed in Section 1.5.6.

1.5.6 The Higgs mechanism

In the Standard Model, the electroweak $SU(2)_L \otimes U(1)_Y$ gauge symmetry is spontaneously broken when a scalar field Φ , called the *Higgs field*, acquires a vacuum expectation value. The Higgs field carries weak isospin and weak hypercharge. It therefore transforms under both $SU(2)_L$ and $U(1)_Y$. As always, the principle of local gauge invariance applies. The transformations can be made locally gauge invariant

[†]The mass term for scalar fields is $-\frac{1}{2}m^2\phi^2$.

by the usual machinery of writing the Higgs Lagrangian in terms of a covariant derivative involving B_μ and W_μ^j fields. Those fields then mix together. After picking a particular ground state for the Higgs field, mass terms emerge for the linear combinations of B_μ and W_μ^j that correspond to the $W^\pm$ and Z bosons. This is the Higgs mechanism. Fermions can also acquire masses by interacting with the Higgs field, albeit in a different way. The topic of this section is the Higgs field and its interaction with vector bosons.

The Higgs Lagrangian is given by

$$\mathcal{L}_{\text{Higgs}} = T - V = (\partial_\mu \Phi)^\dagger (\partial^\mu \Phi) - \left(\mu^2 \phi^\dagger \phi + \lambda (\phi^\dagger \phi)^2 \right) \quad (68)$$

where the Higgs field Φ is an $\text{SU}(2)$ doublet of complex scalar fields.

$$\Phi = \begin{pmatrix} \phi^+ \\ \phi^0 \end{pmatrix} = \begin{pmatrix} \frac{\phi_1 + i\phi_2}{\sqrt{2}} \\ \frac{\phi_3 + i\phi_4}{\sqrt{2}} \end{pmatrix} \quad (69)$$

It carries a weak hypercharge $Y = 1$. The upper and lower member of Φ is assigned a weak isospin of $\frac{1}{2}$ and $-\frac{1}{2}$, respectively. When written on the form in Eq. (68), it is easy to see that $\mathcal{L}_{\text{Higgs}}$ is invariant under both $U(1)_Y$ and $\text{SU}(2)_L$ global gauge transformations.

$$\Phi \rightarrow \Phi' = e^{-i\alpha Y/2} \Phi \quad (70)$$

$$\Phi \rightarrow \Phi' = e^{-i\tau \cdot \alpha/2} \Phi \quad (71)$$

To understand the shape of the potential, it is perhaps simpler to write it in terms of ϕ_1, ϕ_2, ϕ_3 and ϕ_4 .

$$V = \frac{\mu^2 (\phi_1^2 + \phi_2^2 + \phi_3^2 + \phi_4^2)}{2} + \frac{\lambda (\phi_1^2 + \phi_2^2 + \phi_3^2 + \phi_4^2)^2}{2} \quad (72)$$

Assuming a negative μ^2 , the minima are the field configurations that fulfill:

$$\frac{\partial V}{\partial (\phi_1^2 + \phi_2^2 + \phi_3^2 + \phi_4^2)} = \frac{\mu^2}{2} + \lambda (\phi_1^2 + \phi_2^2 + \phi_3^2 + \phi_4^2) = 0 \Leftrightarrow (\phi_1^2 + \phi_2^2 + \phi_3^2 + \phi_4^2) = \frac{-\mu^2}{2\lambda} \equiv v^2 \quad (73)$$

The appropriate choice of ground state is

$$\begin{aligned} \phi_1 &= 0 & \phi_2 &= 0 \\ \phi_3 &= v & \phi_4 &= 0 \end{aligned} \quad (74)$$

so that the vacuum becomes:

$$\Phi_0 = \frac{1}{\sqrt{2}} \begin{pmatrix} 0 \\ v \end{pmatrix} \quad (75)$$

Expanding the Lagrangian around the vacuum in terms of a new field H yields a mass term

$$m_H = \sqrt{-2\mu^2} \quad (76)$$

which is obtained analogously to the mass of η in Section 1.5.5. The field H describes a physical particle called the Higgs, whose observation in 2012 marked the completion of the Standard Model. The Higgs has a mass of roughly 126 GeV [16, 17].

When the Higgs field obtains a vacuum expectation value, the $SU(2)_L$ invariance of $\mathcal{L}_{\text{Higgs}}$ is broken. The $U(1)_Y$ invariance is also broken. However, there is a particular combination of quantum numbers that remains invariant given Φ_0 . The electric charge operator Q yields

$$Q\Phi_0 = (I_3^W + Y/2)\Phi_0 = 0 \quad (77)$$

so the vacuum is still invariant under the following $U(1)$ gauge transformation:

$$\Phi_0 \rightarrow \Phi'_0 = e^{i\alpha(x_\mu)Q}\Phi_0 \quad (78)$$

The corresponding symmetry group is, of course, the familiar $U(1)_{\text{EM}}$ of electromagnetism. Charge remains a good quantum number, and the electroweak $SU(2)_L \otimes U(1)_Y$ symmetry is broken to $U(1)_{\text{EM}}$.

In order to make $\mathcal{L}_{\text{Higgs}}$ locally gauge invariant, the vector fields B_μ and W_μ^j are introduced in terms of a covariant derivative. The new Lagrangian reads

$$\mathcal{L}_{\text{Higgs}} = T - V = (\mathcal{D}_\mu \Phi)^\dagger (\mathcal{D}^\mu \Phi) - (\mu^2 \phi^\dagger \phi + \lambda (\phi^\dagger \phi)^2) \quad (79)$$

where the covariant derivative is defined according to:

$$\mathcal{D}_\mu \equiv \partial_\mu + i\frac{g_1}{2}YB_\mu + i\frac{g_2}{2}\tau^j W_\mu^j \quad (80)$$

To see how the B and W^j bosons mix and acquire masses, the corresponding pieces of $\mathcal{D}_\mu$ can be written out explicitly under the chosen vacuum. The Pauli matrices multiply with the W_μ^j fields according to

$$\tau^j W_\mu^j = \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix} W_\mu^1 + \begin{bmatrix} 0 & -i \\ i & 0 \end{bmatrix} W_\mu^2 + \begin{bmatrix} 1 & 0 \\ 0 & -1 \end{bmatrix} W_\mu^3 = \begin{bmatrix} W_\mu^3 & W_\mu^1 - iW_\mu^2 \\ W_\mu^1 + iW_\mu^2 & -W_\mu^3 \end{bmatrix} \quad (81)$$

so that:

$$\begin{aligned} & \Phi_0^\dagger \left(i\frac{g_1}{2}YB_\mu + i\frac{g_2}{2}\tau^j W_\mu^j \right)^\dagger \left(i\frac{g_1}{2}YB^\mu + i\frac{g_2}{2}\tau_j W_j^\mu \right) \Phi_0 \\ &= \\ & \frac{1}{8} \begin{bmatrix} 0 & \nu \end{bmatrix} \begin{bmatrix} g_1 B_\mu + g_2 W_\mu^3 & g_2 (W_\mu^1 - iW_\mu^2) \\ g_2 (W_\mu^1 + iW_\mu^2) & g_1 B_\mu - g_2 W_\mu^3 \end{bmatrix} \begin{bmatrix} g_1 B^\mu + g_2 W_3^\mu & g_2 (W_1^\mu - iW_2^\mu) \\ g_2 (W_1^\mu + iW_2^\mu) & g_1 B^\mu - g_2 W_3^\mu \end{bmatrix} \begin{bmatrix} 0 \\ \nu \end{bmatrix} \\ &= \\ & \frac{1}{8} \nu^2 g_2^2 \left((W_\mu^1)^2 + (W_\mu^2)^2 \right) + \frac{1}{8} \nu^2 (g_1 B_\mu - g_2 W_\mu^3)^2 \end{aligned} \quad (82)$$

Using Eqs. (49), (50) and (52) to rewrite Eq. (82) in terms of the physical fields $W_\mu^\pm$ and Z_μ yields:

$$\left(\frac{1}{2} \nu g_2 \right)^2 W_\mu^+ W^{-\mu} + \frac{1}{2} \left(\frac{1}{2} \nu \sqrt{g_1^2 + g_2^2} \right)^2 Z_\mu Z^\mu \quad (83)$$

Without breaking gauge invariance, mass terms for the $W^\pm$ and Z bosons have appeared in the Lagrangian. The masses[†] are proportional to the vacuum expectation value ν of the Higgs field:

$$\begin{aligned} m_W &= \frac{1}{2} \nu g_2 \\ m_Z &= \frac{1}{2} \nu \sqrt{g_1^2 + g_2^2} \end{aligned} \quad (84)$$

[†]The mass term is $m^2 A_\mu A^\mu$ for charged vector fields and $\frac{1}{2} m^2 A_\mu A^\mu$ for neutral vector fields.

No mass term appears for the combination of fields corresponding to the photon, which is therefore left massless. This is no coincidence, but a consequence of the choice of vacuum in Eq. (75). Another choice would have resulted in a massive photon. However, since the photon is experimentally observed to be massless, the choice made in Eq. (75) must be correct. Replacing ν by $\nu + H$ in Eq. (83) results in interaction terms between the Higgs boson and the $W^\pm$ and Z bosons. Again, no interaction terms between the Higgs boson and the photon are present in the theory. The Higgs boson only couples to particles that have mass.

No Goldstone bosons appear in the theory, because the gauge is fixed according to Eq. (74) and the corresponding fields vanish. This might seem inconsistent. Every field is associated to a certain number of degrees of freedom. The total number of degrees of freedom present in $\mathcal{L}_{\text{Higgs}}$ should not change simply because a particular gauge is chosen. As a consistency check, the total number of degrees of freedom before and after electroweak symmetry breaking occurs can be counted. Recall that scalar fields (both massless and massive) carry one degree of freedom, massless vector fields carry two degrees of freedom and massive vector fields carry three degrees of freedom. Table 4 confirms that the theory is consistent. After symmetry breaking, the three degrees of freedom associated to ϕ_1 , ϕ_2 and ϕ_4 have transferred to the W^+ , W^- and Z bosons, which become massive as a result. The process has been humorously described as the $W^\pm$ and Z bosons “eating” the Goldstone bosons, thereby gaining weight. Finally, the degree of freedom associated to ϕ_3 manifests itself as the Higgs boson.

Table 4: The table shows the particles present in $\mathcal{L}_{\text{Higgs}}$ before and after electroweak symmetry breaking. The total number of degrees of freedom does not change.

Before electroweak symmetry breaking			After electroweak symmetry breaking		
Particle	Type	Degrees of freedom	Particle	Type	Degrees of freedom
ϕ_1	Massless scalar	1	H	Massless scalar	1
ϕ_2	Massless scalar	1	γ	Massless vector	2
ϕ_3	Massless scalar	1	W^+	Massive vector	3
ϕ_4	Massless scalar	1	W^-	Massive vector	3
B	Massless vector	2	Z	Massive vector	3
W^1	Massless vector	2			
W^2	Massless vector	2			
W^3	Massless vector	2			

1.5.7 Yukawa couplings

Once the Higgs field is present in the theory, it becomes possible to write down Lagrangians of the form

$$\mathcal{L}_{\text{Yukawa}} = -\lambda_\ell \bar{L}_L \Phi L_R - \lambda_d \bar{Q}_L \Phi D_R - \lambda_u \bar{Q}_L \Phi^\dagger U_R + \text{h.c.} \quad (85)$$

where λ_ℓ , λ_d and λ_u are arbitrary constants and h.c. denotes the Hermitian conjugate. An interaction between scalar and Dirac fields is called a *Yukawa* interaction. The iso-doublets L_L , U_L and D_L as well as the Higgs field Φ transform under $\text{SU}(2)_L$, so Eq. (85) is invariant under such transformations. Inserting Eq. (75) and expanding around the vacuum yields terms of the form

$$\mathcal{L}_{\text{Yukawa}} \supset -\frac{\lambda_\nu}{\sqrt{2}} (\bar{\psi}_L \psi_R + \bar{\psi}_R \psi_L) - \frac{\lambda}{\sqrt{2}} (\bar{\psi}_L \psi_R + \bar{\psi}_R \psi_L) H = -\frac{\lambda_\nu}{\sqrt{2}} \bar{\psi} \psi - \frac{\lambda}{\sqrt{2}} \bar{\psi} \psi H \quad (86)$$

where ψ can be any fermion except a neutrino. This is because of the Standard Model's assumption that there are no right-handed neutrinos. The first term in Eq. (86) corresponds to a mass:

$$m = \frac{\lambda v}{\sqrt{2}} \quad (87)$$

The second term corresponds to an interaction between the fermion and the Higgs boson. It is precisely this interaction that couples left-handed and right-handed fermions together, as alluded to in Section 1.5.1. Both the mass and the coupling strength are proportional to λ , whose value is not predicted by the theory and must be measured experimentally. Every fermion flavor has a different value of λ .

1.5.8 The electroweak Lagrangian

The electroweak Lagrangian before symmetry breaking can be divided into four parts.

$$\mathcal{L}_{\text{EW}} = \mathcal{L}_{\text{F}} + \mathcal{L}_{\text{G}} + \mathcal{L}_{\text{Higgs}} + \mathcal{L}_{\text{Y}} \quad (88)$$

The different parts describe, in order, the fermions, the gauge bosons, the Higgs field and the Yukawa couplings. This section will summarize each part while making frequent references to previous sections. The fermion part is given by

$$\mathcal{L}_{\text{F}} = \sum_{\text{Fermions}} \left(i\bar{\Psi}_L \gamma^\mu \mathcal{D}_\mu \Psi_L + i\bar{\Psi}_R \gamma^\mu \mathcal{D}_\mu \Psi_R \right) \quad (89)$$

where Ψ_L and Ψ_R are any of the iso-doublets and singlets from Eqs. (43) and (44), respectively. The gauge part is given by:

$$\mathcal{L}_{\text{G}} = -\frac{1}{4} B^{\mu\nu} B_{\mu\nu} - \frac{1}{4} W_j^{\mu\nu} W_{j\mu\nu} \quad (90)$$

The field strength tensors are shorthand for

$$B^{\mu\nu} \equiv \partial^\mu B^\nu - \partial^\nu B^\mu \quad (91)$$

and

$$W_j^{\mu\nu} \equiv \left(\partial^\mu W_j^\nu - \partial^\nu W_j^\mu \right) + g_2 \epsilon_{jkl} W_k^\mu W_l^\nu \quad (92)$$

where ϵ_{jkl} are the structure constants of SU(2). The Higgs part is simply Eq. (79), and the Yukawa part is a sum over the Yukawa mass term for each fermion such that

$$\mathcal{L}_{\text{Y}} = \sum_{\text{Fermions}} \mathcal{L}_{\text{Yukawa}} \quad (93)$$

where $\mathcal{L}_{\text{Yukawa}}$ is given by Eq. (85). Finally, the definition of the covariant derivative used in Eq. (89) differs for left-handed and right-handed particles (because only the former transform under SU(2)). When acting on left-handed particles, the covariant derivative is defined according to

$$\mathcal{D}_\mu \equiv \partial_\mu + i\frac{g_1}{2} Y B_\mu + i\frac{g_2}{2} \tau^j W_\mu^j \quad (94)$$

but when acting on right-handed particles the definition is:

$$\mathcal{D}_\mu \equiv \partial_\mu + i\frac{g_1}{2} Y B_\mu \quad (95)$$

1.6 An incomplete theory

The Standard Model does not provide a complete description of nature. Although the agreement between theory and experiment is impressive, there are clear signs that parts of the model are missing. *Supersymmetry* (see Section 2) is an extension of the Standard Model that more than doubles the number of particles present in the theory. This section will briefly discuss three problems for which supersymmetry might be the solution and finish with the unrelated issue of how the Standard Model could be extended to account for massive neutrinos.

1.6.1 Dark matter

Evidence from galaxy rotation curves [18] and gravitational lensing [19] suggests that the amount of luminous matter present in a typical galaxy only makes up a small fraction of the total matter content. The rest of the matter, which apparently does not interact via the electromagnetic force, is called *dark matter*. There is no particle in the Standard Model that could account for this observation. In fact, Standard Model particles only make up 4.9% of the total mass/energy content of the universe with dark matter contributing another 26.8% [20]. The remainder is so-called dark energy, which is a topic beyond the scope of this text. Supersymmetry predicts the existence of a massive, stable particle that interacts only via the weak force and gravity. Such a particle could be a viable dark matter candidate.

1.6.2 The hierarchy problem

To first order, the mass of the Higgs boson is predicted to be $m_H = \sqrt{-2\mu^2}$ as discussed in Section 1.5.6. However, a more detailed calculation (see Section 2.2) shows that the Higgs mass also receives contributions from higher-order effects. Every particle that couples to the Higgs contributes to its mass through loop diagrams. Bosons make positive contributions while fermions make negative contributions. It turns out that the higher-order contributions are divergent, which in principle should result in an infinite Higgs mass. Since the Standard Model is most likely incomplete, it can be argued that the mass should only be pushed up to some *ultraviolet cut-off scale* Λ_{UV} where new physics starts to alter the behavior of the theory. In any case, the experimentally observed mass of 126 GeV seems unlikely. It suggests that, for no apparent reason, the positive and negative contributions from bosons and fermions cancel out almost exactly. This is termed a *fine tuning* problem. Supersymmetry solves the problem by introducing a new fermion for every boson in the Standard Model, as well as a new boson for every fermion in the Standard Model. The positive and negative contributions to the Higgs mass then cancel by design rather than coincidence.

1.6.3 Grand unification

Eqs. (26) and (41) define running coupling constants for $U(1)_{EM}$ and $SU(3)$. In the same manner, running coupling constants α_1 and α_2 can be defined for $U(1)_Y$ and $SU(2)_L$. Figure 7 (dashed lines) shows how the coupling constants for the currently known gauge interactions vary with q in the Standard Model. *Grand unification* is the idea that they should all converge at some high energy scale Λ_{GUT} [21]. Models that allow for this are called *grand unified theories* (GUTs). They describe all forces as different manifestations of a single interaction characterized by a larger gauge group. As shown by the solid lines in Figure 7, the unification of forces is possible within the framework of supersymmetry.

1.6.4 Neutrino masses

Experiments on neutrino oscillations indicate that neutrinos have small but nonzero masses [23]. This is not explained by the Standard Model, which describes neutrinos as massless particles. Masses for a

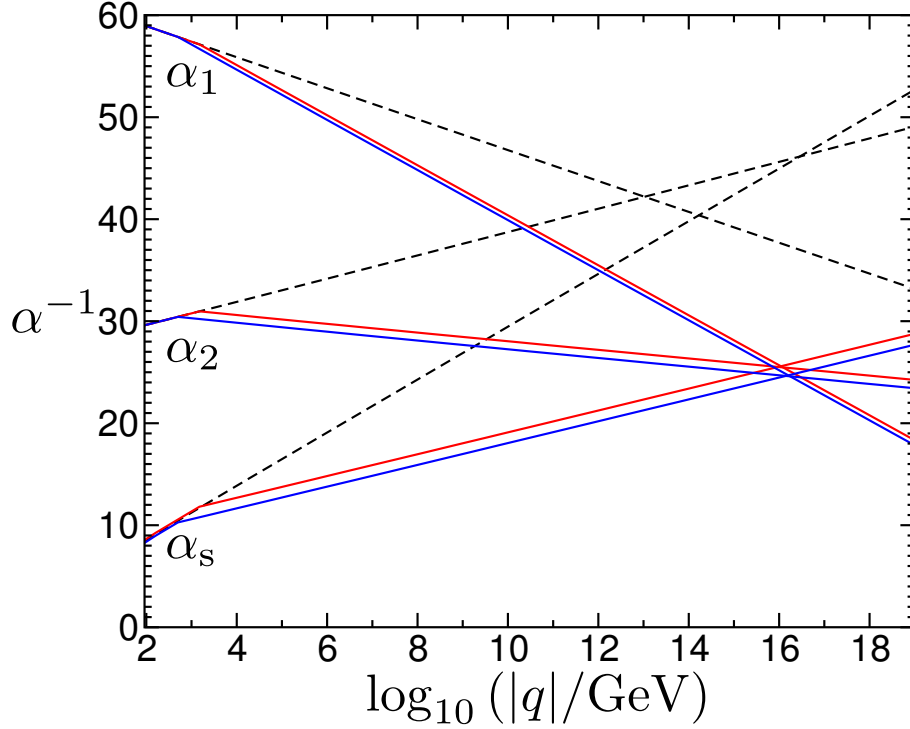


Figure 7: The dashed lines show the evolution of α_1 , α_2 and α_s in the Standard Model. They do not converge. The solid lines show the same parameters in a particular supersymmetric scenario, where they do converge. Figure adapted from Ref. [22].

given generation can be accounted for by adding a right-handed neutrino field ψ_{ν_R} to the theory. Such a field would be an SU(2) singlet and therefore not interact via the weak force. Neither would it interact via the electromagnetic and strong forces. In fact, right-handed neutrinos would only interact gravitationally and are sometimes referred to as a *sterile neutrinos* for this reason. There are two possible ways to write down a mass term for ψ_{ν_R} . The first possibility is the Dirac mass term. It comes from the familiar Yukawa interaction (see Section 1.5.7) and is given by

$$\mathcal{L}_D = -m_D \bar{\psi}_{\nu_R} \psi_{\nu_L} + \text{h.c.} \quad (96)$$

where m_D denotes the Dirac mass. If the entire mass of the neutrino comes from Eq. (96), then its Yukawa coupling must be extremely small — About seven orders of magnitude smaller than the Yukawa coupling of the electron and thirteen orders of magnitude smaller than that of the tau. It could be argued that such a discrepancy is not natural, and that the leptons of a given generation ought to have similar masses on some fundamental level (just like the quarks of a given generation do). The second possibility is to write down a Majorana mass term. It is given by

$$\mathcal{L}_M = -\frac{m_R}{2} \bar{\psi}_{\nu_R} \psi_{\nu_R}^c + \text{h.c.} \quad (97)$$

where m_R is the Majorana mass of the right-handed neutrino and $\psi_{\nu_R}^c$ denotes the charge conjugate of ψ_{ν_R} . A Majorana mass term can be read as an interaction that changes a particle into its own antiparticle. This means that charged leptons are not allowed to have Majorana mass terms, because such terms would violate charge conservation. In particular, left-handed neutrinos cannot have Majorana masses because they are charged under both SU(2) and U(1). Right-handed neutrinos are not (they are sterile), and

Majorana masses are consequently allowed. The full neutrino mass Lagrangian becomes:

$$\mathcal{L}_{\text{Full}} = \mathcal{L}_D + \mathcal{L}_M = -\frac{1}{2} \begin{bmatrix} \bar{\psi}_{\nu_L}^c & \bar{\psi}_{\nu_R} \end{bmatrix} \begin{bmatrix} 0 & m_D \\ m_D & m_R \end{bmatrix} \begin{bmatrix} \psi_{\nu_L} \\ \psi_{\nu_R}^c \end{bmatrix} + \text{h.c.} \quad (98)$$

The eigenvalues of the neutrino mass matrix are given by

$$m_{\pm} = \frac{m_R \pm \sqrt{m_R^2 + 4m_D^2}}{2} \quad (99)$$

so that the physical neutrinos gain masses m_+ and m_- . It is possible that m_D is fundamentally on the order of the other lepton masses while m_R , not being tied to any other Standard Model particle by gauge interactions, is very large. In the limit where $m_R \gg m_D$, the physical neutrinos would gain masses

$$\begin{aligned} m_+ &\approx m_R \\ m_- &\approx \frac{m_D^2}{m_R} \end{aligned} \quad (100)$$

so that the *heavier* the heavy neutrinos are, the *lighter* the light neutrinos become as a result. This is called the *see-saw mechanism*, and it is a popular theory for explaining why the observed neutrinos have such tiny masses compared to the other leptons. The heavy neutrinos, with masses on the GUT scale, would be too massive to produce experimentally.

2 Supersymmetry

2.1 Overview

Supersymmetry (SUSY) [24–32] is an extension of the Standard Model of particle physics. It is based on the idea that there exists a symmetry between fermions (particles with half-integer spin) and bosons (particles with integer spin). The main prediction of SUSY is that every fermion in the Standard Model should have a *superpartner* that is a boson. In the same way, every boson in the Standard Model should have a superpartner that is a fermion. Superpartners have the same quantum numbers as their Standard Model counterparts except for their spin, which differs by $\pm\frac{1}{2}$. For reasons of internal consistency, SUSY also requires the existence of an additional Higgs field. Whereas the Standard Model contains *one* complex scalar doublet Φ (see Section 1.5.6), SUSY predicts *two* complex scalar doublets plus the corresponding superpartners.

The particles mentioned so far constitute the Minimal Supersymmetric Standard Model (MSSM). It is called *minimal* because it is impossible to build a model with fewer particles while still realizing supersymmetry. This section will introduce SUSY in the context of the MSSM. The discussion will be mostly non-mathematical with a focus on describing the particle content of the theory rather than the underlying Lagrangian. A more robust mathematical treatment can be found in Ref. [22].

No particle predicted by supersymmetry has ever been experimentally observed. This suggests that all superpartners, if they exist, must be rather heavy. If they were light, past collider experiments would have discovered them long ago. It is not obvious why the mass of a superpartner should be different from its Standard Model counterpart. Indeed, if the symmetry were exact, they ought to have *exactly the same* masses. The usual way around this problem is to suppose that supersymmetry is somehow broken in a way that increases the masses of the superpartners. A general Lagrangian that contains all possible SUSY breaking terms is written down, and the task of experiments becomes to constrain the parameters of that Lagrangian. The MSSM contains a total of 105 free parameters that are not present in the Standard Model. All but one of them arise due to supersymmetry breaking.

2.2 Solving the hierarchy problem

Section 1.6.2 introduced the hierarchy problem. This section will describe the problem in more detail, and explain how it is solved by introducing supersymmetry. Recall from Section 1.5.7 that the Higgs boson couples to fermions via Yukawa interaction terms of strength λ . A Higgs boson can therefore fluctuate into a virtual fermion-antifermion pair as shown in Figure 8 (a). This process contributes to the self-energy, and therefore the mass, of the Higgs boson. The total contribution is obtained by summing over all the possible ways in which the process can occur. Put more technically, the total contribution is the integral over all the allowed momentum states of the fermion-antifermion pair. Since the fermions are virtual particles, they are allowed to have *any* momentum. This means that there is an infinite number of ways in which the process can happen, and the integral becomes divergent.

As usual in quantum field theory, infinities can be dealt with by regularization. The divergent integral is only evaluated up to some ultraviolet cut-off scale Λ_{UV} , which can be interpreted as the energy scale where new physics enters the picture and alters the low-energy behavior of the theory. In terms of Λ_{UV} , the leading order correction to the Higgs mass from virtual fermion loops becomes

$$\Delta m_H^2 = -\frac{|\lambda_f|^2}{8\pi^2} \Lambda_{UV}^2 + \dots \quad (101)$$

where λ_f is the coupling strength to the fermion in question. In the same manner, the Higgs mass receives quantum corrections from virtual scalar loops such as the one shown in Figure 8 (b). The leading order

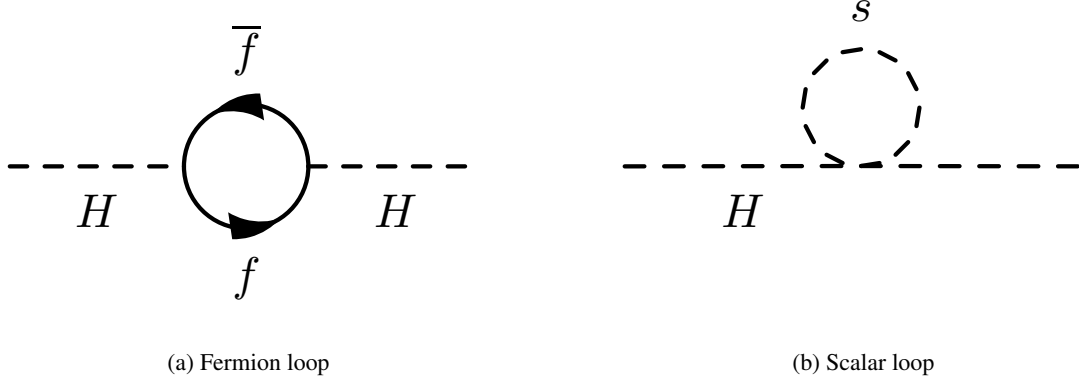


Figure 8: The Higgs mass receives quantum corrections from fermion (a) and scalar (b) loop diagrams.

contribution is

$$\Delta m_H^2 = \frac{\lambda_s}{16\pi^2} \Lambda_{UV}^2 + \dots \quad (102)$$

where λ_s is the coupling strength to the scalar in question. The relative minus sign between Eqs. (101) and (102) allows the radiative corrections to partially cancel, but naively there is no reason to expect the cancellation to be complete. Instead, the corrections should push the Higgs mass towards the cut-off scale Λ_{UV} . This scale could in principle be very large, possibly approaching the Planck scale $\Lambda_{\text{Planck}} \approx 10^{19}$ GeV. The experimentally observed mass $m_H = 126$ GeV therefore seems very unlikely. This is known as the hierarchy problem.

By design, all supersymmetric theories contain two scalars for every fermion.[†] Furthermore, in unbroken supersymmetry, their couplings are such that:

$$|\lambda_f|^2 = \lambda_s \quad (103)$$

The leading order terms of Eqs. (101) and (102) therefore cancel exactly, which solves the hierarchy problem. In the MSSM, it is assumed that the SUSY breaking mechanism, whatever it is, happens in such a way that Eq. (103) remains satisfied. Otherwise the solution to the hierarchy problem, which is the main motivation for SUSY in the first place, would be lost. The idea that SUSY is broken without altering any of the Higgs couplings present in the theory is referred to as *soft* supersymmetry breaking. It is expressed by separating the MSSM Lagrangian into two parts.

$$\mathcal{L}_{\text{MSSM}} = \mathcal{L}_{\text{SUSY}} + \mathcal{L}_{\text{Soft}} \quad (104)$$

The first part, $\mathcal{L}_{\text{SUSY}}$, contains all the gauge and Yukawa interactions present in the theory (and therefore covers all terms that contain λ_f or λ_s). This part is assumed to preserve supersymmetry invariance. The second part, $\mathcal{L}_{\text{Soft}}$, contains all the terms that break supersymmetry and grant higher masses to the superpartners.

Although the cancellation of the leading order terms in Eqs. (101) and (102) is exact, higher order terms can still contribute to the Higgs mass. The size of these contributions can be written as

$$\Delta m_H^2 = m_{\text{Soft}}^2 \left(\frac{\lambda}{16\pi^2} \ln(\Lambda_{UV}/m_{\text{Soft}}) + \dots \right) \quad (105)$$

[†]In SUSY, the two scalar degrees of freedom are assembled into a single complex scalar field. Rather than saying that SUSY introduces one boson for each fermion, it is perhaps more precise to say that it introduces one bosonic degree of freedom for each fermionic degree of freedom.

where λ is a coupling strength and m_{Soft} is the highest mass scale associated with the SUSY breaking Lagrangian $\mathcal{L}_{\text{Soft}}$. If supersymmetric particles are very heavy, the contribution made to Δm_H^2 again becomes much greater than the experimentally observed Higgs mass. In practice, this means that the solution to the hierarchy problem is effectively lost unless SUSY is found at or below the TeV scale.

2.3 Particles in the MSSM

This section will list the particles present in the MSSM. It will also introduce some nomenclature and notational conventions. In general, supersymmetric particles are referred to as *sparticles*. They are represented by putting a tilde above the symbol that denotes the corresponding Standard Model particle. Table 5 shows the first-generation leptons and quarks together with their superpartners. The sparticles in this table are all complex scalars, but other than their spin they have quantum numbers identical to their Standard Model counterparts. The name of an individual scalar sparticle is obtained by putting an “s” in front of the Standard Model name. For example, the superpartner of the left-handed electron e_L is the left-handed *selectron* $\tilde{e}_L$ and the superpartner of the right-handed top t_R is the right-handed *stop* $\tilde{t}_R$. As a group, supersymmetric leptons and quarks are called *sleptons* and *squarks*. Taken all together they make up the *sfermions*.

Table 5: The table shows the Standard Model leptons and quarks together with their scalar superpartners. Only the first generation is shown, but the corresponding second and third generations of particles and superpartners also exist.

Spin 0	Spin $\frac{1}{2}$	Sparticle name
$\tilde{\nu}_L$	ν_L	Left-handed and right-handed sleptons
$\tilde{e}_L$	e_L	
$\tilde{e}_R$	e_R	
$\tilde{u}_L$	u_L	Left-handed and right-handed squarks
$\tilde{d}_L$	d_L	
$\tilde{u}_R$	u_R	
$\tilde{d}_R$	d_R	

Whereas the Standard Model contains only one complex scalar Higgs doublet (for a total of four scalar degrees of freedom), the MSSM predicts the existence of *two* complex scalar Higgs doublets (for a total of eight scalar degrees of freedom). The doublets are denoted

$$\begin{pmatrix} H_u^+ \\ H_u^0 \end{pmatrix} \quad \text{and} \quad \begin{pmatrix} H_d^0 \\ H_d^- \end{pmatrix}. \quad (106)$$

After electroweak symmetry breaking, three degrees of freedom are eaten by the $W^\pm$ and Z , which become massive as a result. The remaining five scalar fields gain their own mass terms, and the corresponding particles should be possible to observe experimentally. The MSSM contains three neutral scalar Higgses h^0 , H^0 and A^0 , as well as two electrically charged scalar Higgses H^+ and H^- . If supersymmetry exists, then the Higgs particle observed by ATLAS and CMS was actually the lightest MSSM Higgs h^0 . In large parts of the MSSM phase space, h^0 turns out to be nearly indistinguishable from a Standard Model Higgs.

Table 6 shows the Higgses of the MSSM together with their fermionic superpartners. The name of a fermionic sparticle is obtained by appending “ino” to the name of the corresponding particle. Supersymmetric Higgses are therefore called *Higgsinos*. Finally, Table 7 shows the Standard Model gauge bosons together with their superpartners. Since the sparticles in this table are all fermions, their names are again

obtained by appending “ino” to the corresponding Standard Model particle. The gluinos, winos and the bino are collectively referred to as *gauginos*.

Table 6: The table shows the Higgses of the MSSM together with their fermionic superpartners.

Spin 0	Spin $\frac{1}{2}$	Sparticle name
H_u^+	$\widetilde{H}_u^+$	Higgsinos
H_u^0	$\widetilde{H}_u^0$	
H_d^0	$\widetilde{H}_d^0$	
H_d^-	$\widetilde{H}_d^-$	

Table 7: The table shows the Standard Model gauge bosons together with their fermionic superpartners.

Spin $\frac{1}{2}$	Spin 1	Sparticle name
$\widetilde{g}$	g	Gluino
$\widetilde{W}^j$	W^j	Wino
$\widetilde{B}$	B	Bino

2.3.1 Sparticle mixing

Tables 5-7 showed the *gauge eigenstates* of the particles in the MSSM. These are the states that participate in gauge interactions. If some particles have compatible quantum numbers, they can mix with each other. The resulting *mass eigenstates* are the particles that actually propagate through time and space, and show up in experiments. They are linear combinations of the gauge eigenstates. Table 8 summarizes the particle mixing that takes place within the MSSM.

Table 8: Sparticles in the MSSM can mix with each other if they have compatible quantum numbers. The table shows what gauge eigenstates mix into what mass eigenstates.

Gauge eigenstates	Mass eigenstates
$\widetilde{W}^1, \widetilde{W}^2, \widetilde{H}_u^+, \widetilde{H}_d^-$	$\rightarrow \widetilde{\chi}_1^+, \widetilde{\chi}_1^-, \widetilde{\chi}_2^+, \widetilde{\chi}_2^-$
$\widetilde{B}, \widetilde{W}^3, \widetilde{H}_u^0, \widetilde{H}_d^0$	$\rightarrow \widetilde{\chi}_1^0, \widetilde{\chi}_2^0, \widetilde{\chi}_3^0, \widetilde{\chi}_4^0$
$\widetilde{\tau}_L, \widetilde{\tau}_R$	$\rightarrow \widetilde{\tau}_1, \widetilde{\tau}_2$
$\widetilde{t}_L, \widetilde{t}_R$	$\rightarrow \widetilde{t}_1, \widetilde{t}_2$
$\widetilde{b}_L, \widetilde{b}_R$	$\rightarrow \widetilde{b}_1, \widetilde{b}_2$

The effects of electroweak symmetry breaking causes the Higgsinos and electroweak gauginos to mix into the *neutralinos* $\widetilde{\chi}_i^0$ and the *charginos* $\widetilde{\chi}_i^\pm$. To understand how $\widetilde{W}^1$ and $\widetilde{W}^2$ can mix with the electrically charged $\widetilde{H}_u^+$ and $\widetilde{H}_d^-$, recall that the former are really the components of $\widetilde{W}^+$ and $\widetilde{W}^-$. Meanwhile, the electrically neutral $\widetilde{W}^3$ and $\widetilde{B}$ mix with the remaining components of the Higgs fields. By convention, the gaugino mass eigenstates are numbered such that the lightest particles have the lowest index. For example, $\widetilde{\chi}_1^0$ is the lightest neutralino while $\widetilde{\chi}_2^+$ and $\widetilde{\chi}_2^-$ are the heaviest charginos. The positive and negative charginos are always mass-degenerate (i.e. they have the same mass). It is worth noting that, while the exact amount of mixing depends on various free parameters in the MSSM, some rules of thumb apply. It is often the case that $\widetilde{\chi}_1^0$ is taken to be “mostly bino” while $\widetilde{\chi}_1^\pm$ and $\widetilde{\chi}_2^0$ are taken to be “mostly wino”. The heavier gauginos, then, are mostly Higgsino.

Left- and right-handed squarks and sleptons also mix with each other. In phenomenologically viable variants of the MSSM, it turns out that the amount of mixing for a given sfermion is proportional to its Yukawa coupling. The third-generation sfermions therefore mix the most. Mixing within the first two generations tends to be small in comparison and is often neglected completely. Gluinos, being the only gauginos to carry color charge, do not mix at all with any other sparticles.

2.4 R-parity

In a supersymmetric Lagrangian, it is possible to write down terms that violate leptons number conservation as well as terms that violate baryon number conservation. This is problematic as the corresponding processes have not been observed experimentally. For example, the experimental lower bound on the proton decay process $p \rightarrow \pi^0 + e^+$ is above 10^{33} years [8]. If baryon and lepton number violating processes were allowed in SUSY, they would mediate rapid proton decay via interactions such as the one shown in Figure 9. Therefore, there must exist some conservation law that forbids these interactions.

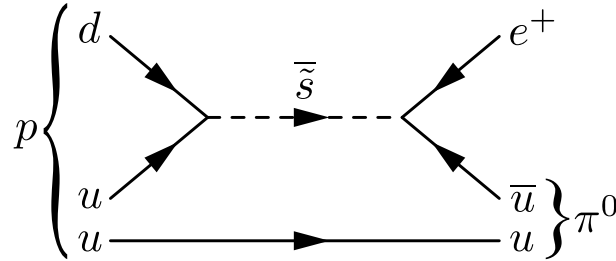


Figure 9: Proton decay mediated by the superpartner of the strange antiquark. This process would lead to a very short proton lifetime. It is forbidden by the introduction of R-parity.

Rather than taking baryon and lepton number conservation to be fundamental symmetries of nature (which is known not to be the case [33]), a new multiplicatively conserved quantum number called *R-parity* is introduced. A particle has an R-parity given by

$$R_p = (-1)^{2s+3B+L} \quad (107)$$

where s is the spin of the particle, B is its baryon number and L is its lepton number. By plugging in the appropriate numbers, it is easy to see that all Standard Model particles have $R_p = 1$. Meanwhile, superpartners differ from their Standard Model counterparts by $\Delta s = \pm \frac{1}{2}$ and therefore have $R_p = -1$. This has three major phenomenological consequences.

- 1) In interactions between Standard Model particles, SUSY particles are always produced in pairs.
- 2) SUSY particles can only decay into states that contain an odd number of other SUSY particles.
- 3) The lightest supersymmetric particle (LSP) is stable.

The process shown in Figure 9 is forbidden by both items 1) and 2). Finally, the fact that the LSP cannot decay makes it a good dark matter candidate in many models [34, 35].

2.5 Supersymmetry breaking

If the symmetry between fermions and bosons were exact, then supersymmetric particles would have the same mass as their Standard Model counterparts. This is clearly ruled out by experiment, so the symmetry must be broken in the vacuum state chosen by nature. However, finding a supersymmetry

breaking mechanism that reproduces experimental results is not so simple. Within the MSSM, there is no field that could develop a vacuum expectation value in a way that leads to a phenomenologically viable model. Whatever field is responsible for breaking SUSY is therefore presumed to exist *beyond* the MSSM. The idea is illustrated schematically in Figure 10. Supersymmetry is broken in some *hidden sector* that contains particles that are decoupled from the *visible sector* where the MSSM resides. Either the hidden particles are much too heavy to be produced in a collider experiment, or their interactions with particles in the visible sector have negligible coupling strength — But as long as the hidden particles are allowed to interact with the visible particles on *some* level, symmetry breaking effects can be mediated to the MSSM. Whatever interaction is responsible, it is often assumed to be *flavor blind* such that it couples in the same way to each generation of squarks and sleptons. This simplifies the final model considerably.

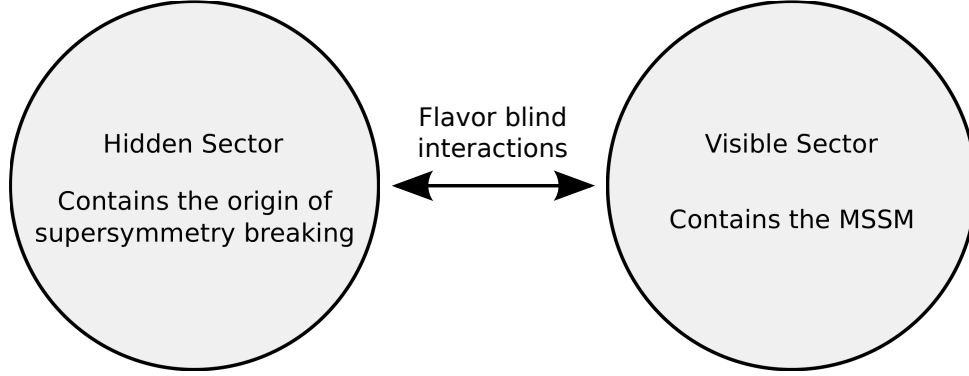


Figure 10: Schematic diagram showing the presumed mechanism behind supersymmetry breaking. SUSY is broken in a hidden sector and the effects are mediated to the MSSM by flavor blind interactions.

Since the exact symmetry breaking mechanism is not known, the usual approach is to simply write down the most general supersymmetry breaking Lagrangian possible. Given that the idea of soft SUSY breaking applies (see Section 2.2), such a Lagrangian contains 105 free parameters without counterpart in the Standard Model. Luckily, making certain assumptions about the nature of the symmetry breaking mechanism can greatly reduce this parameter space. The basis of such assumptions, other than experimental constraints, is the idea that relationships between different physical constants are simpler at some high energy scale. Because the value of a physical observable depends on the energy scale at which it is measured, quantities that are the same at a high energy scale may be very different at a lower energy scale. The energy evolution of a parameter is dictated by its renormalization group equation. While a detailed discussion is outside of the scope of this text, the concept as such should already be familiar from e.g. Section 1.6.3.

2.5.1 Gravity mediated

The interactions that mediate supersymmetry breaking could be of gravitational origin. Supersymmetric theories that account for gravity (not all do) are called *supergravity* theories. They include a spin 2 graviton and introduce its superpartner, the spin $\frac{3}{2}$ gravitino. The simplest such models, which are also the historically most popular and well-studied, are known as minimal supergravity (mSUGRA) models. They tend to result in a phenomenology where the LSP is the lightest neutralino $\tilde{\chi}_1^0$. The 105 free parameters in the MSSM are reduced to just five by making a series of simplifying assumptions. Some of them are illustrated in Figure 11, which shows the renormalization group running of various parameters in a particular version of mSUGRA. All the squark and slepton masses are assumed to converge at a common value m_0 at the GUT scale. This is the first free parameter. The bino mass M_1 , the wino mass M_2 and the gluino mass M_3 are assumed to converge at a common value $m_{1/2}$, which is the second

free parameter. Finally, H_u and H_d denote the running quantities $(\mu^2 + m_{H_u}^2)^{1/2}$ and $(\mu^2 + m_{H_d}^2)^{1/2}$ which appear in the MSSM Higgs potential. They are assumed to converge at a common value $(\mu^2 + m_0^2)^{1/2}$. The parameter μ is called the Higgsino mass parameter. It can be thought of as the supersymmetric version of the Higgs boson mass and is ultimately related to the masses of the MSSM Higgses. The third free parameter in mSUGRA is the sign of μ . The fourth free parameter $\tan\beta$ is given by:

$$\tan\beta = \frac{\langle H_u^0 \rangle}{\langle H_d^0 \rangle} \quad (108)$$

The fifth and final free parameter is the trilinear coupling A_0 : The supersymmetric version of the Yukawa interaction is the coupling of an MSSM Higgs to a left- and right-handed sfermion. Since this is an interaction between three scalars, the coupling strength is referred to as a trilinear scalar coupling. In mSUGRA, the strength of this coupling is assumed to be proportional to the corresponding Yukawa coupling, and A_0 is simply the constant of proportionality.

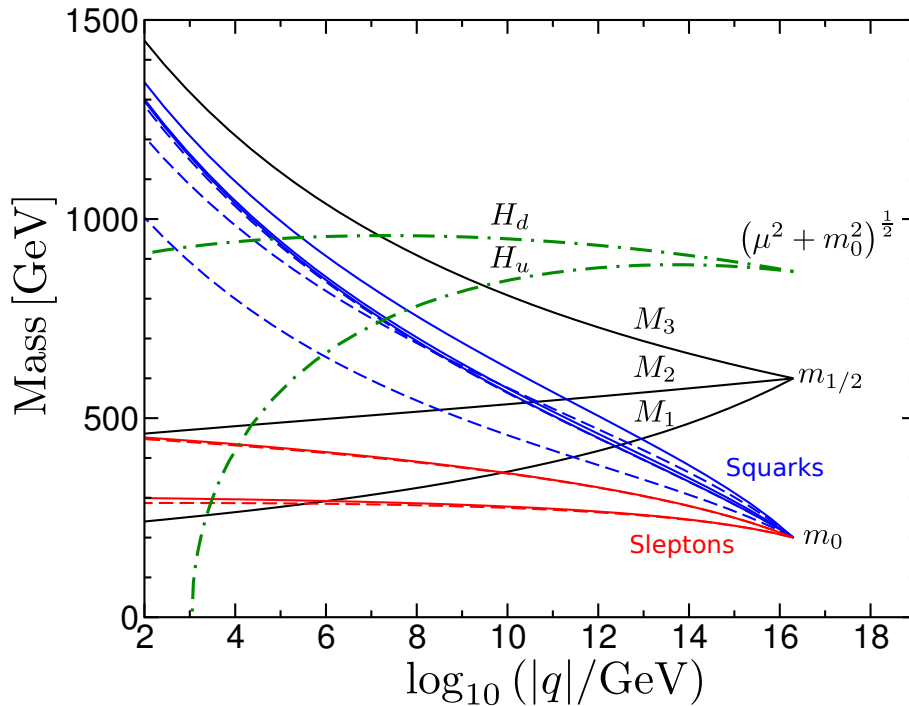


Figure 11: The figure shows the renormalization group running of the sfermion masses, the gaugino masses as well as the quantities $H_u = (\mu^2 + m_{H_u}^2)^{1/2}$ and $H_d = (\mu^2 + m_{H_d}^2)^{1/2}$ in a particular mSUGRA model. Dashed sfermion lines denote third-generation masses while solid lines denote first- and second-generation masses. The renormalization group equations are such that third-generation sfermions tend to become the lightest due to their large Yukawa couplings. The fact that H_u runs negative ultimately provokes electroweak symmetry breaking. Figure adapted from Ref. [22].

2.5.2 Gauge mediated

The interactions that mediate supersymmetry breaking could be the usual $SU(3) \times SU(2)_L \times U(1)_Y$ gauge interactions. This scenario is referred to as gauge mediated supersymmetry breaking (GMSB). It introduces a set of very heavy messenger particles that couple directly to the source of symmetry breaking, and indirectly to particles in the MSSM via virtual loop corrections. An example of such loops is shown in Figure 12. The messenger particles are taken to be part of some larger gauge group such as $SU(5)$.

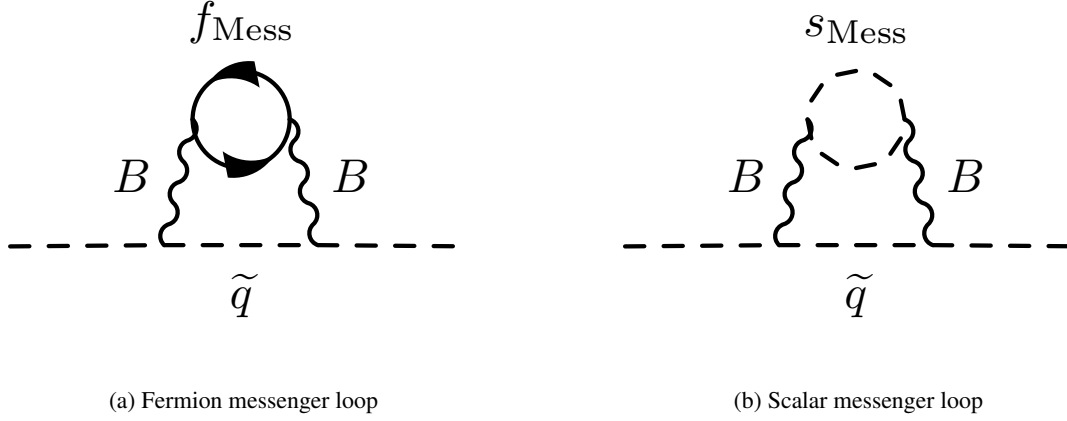


Figure 12: GMSB models introduce a set of very heavy messenger particles. These messengers interact with gauge bosons and therefore couple to particles in the MSSM via quantum loops. The figure shows a squark coupling to a fermionic (a) and scalar (b) messenger via a Standard Model B boson.

GMSB models contain gravity and generally predict the LSP to be a gravitino with a mass on the order of one eV. Such a particle, although light, could have avoided detection due to its feeble interaction strength. It should be noted that GMSB models necessarily contain some amount of gravity mediated supersymmetry breaking. However, its effect is usually neglected since gauge interactions are so much stronger than interactions mediated by gravity.

Free parameters in GMSB models vary but may include the messenger particle mass scale M_{Mess} , the supersymmetry breaking scale F_m and a gravitino mass scaling parameter C_{Grav} that essentially sets the lifetime of the next-to-lightest supersymmetric particle (NLSP). The sign of μ and $\tan\beta$ also appear, and have the same meaning as in mSUGRA models. Finally, it is possible to introduce more than one set of SU(5) messengers. The number of *sets* of messenger particles N_5 is a free parameter as well.

2.6 Sparticle decays

This section will summarize the most important sparticle decay modes. Whenever a decay is listed, it is implicitly assumed that it is kinematically allowed. That is, the mass of the decaying particle must be greater than the total mass of the decay products. Whether or not this is true depends on the mass hierarchy of the model in question. Examples of what mass hierarchies might look like in different mSUGRA models are shown in Figure 13. In particular, Figure 13 (a) shows the same model as Figure 11.

Table 9 shows some of the decays that could occur within the MSSM. Right-handed sleptons decay predominantly into $\tilde{\chi}_1^0$ if it is mostly bino. This is because right-handed particles do not couple to winos. In a similar manner, left-handed sleptons prefer to decay into $\tilde{\chi}_1^\pm$ and $\tilde{\chi}_2^0$ if they are mostly wino. Although left-handed particles can interact with both binos and winos, the former interaction wins because the SU(2) coupling strength g_2 is much greater than the U(1) coupling strength g_1 (see Figure 7). Squarks decay into gluinos if kinematically allowed. If not, they follow the same decay pattern as sleptons. Decays to the heavy Higgses are not listed since they are dominated by decays into the lightest Higgs h^0 . Because of R-parity conservation, all decay chains eventually end up in a final state involving only Standard Model particles and the LSP.

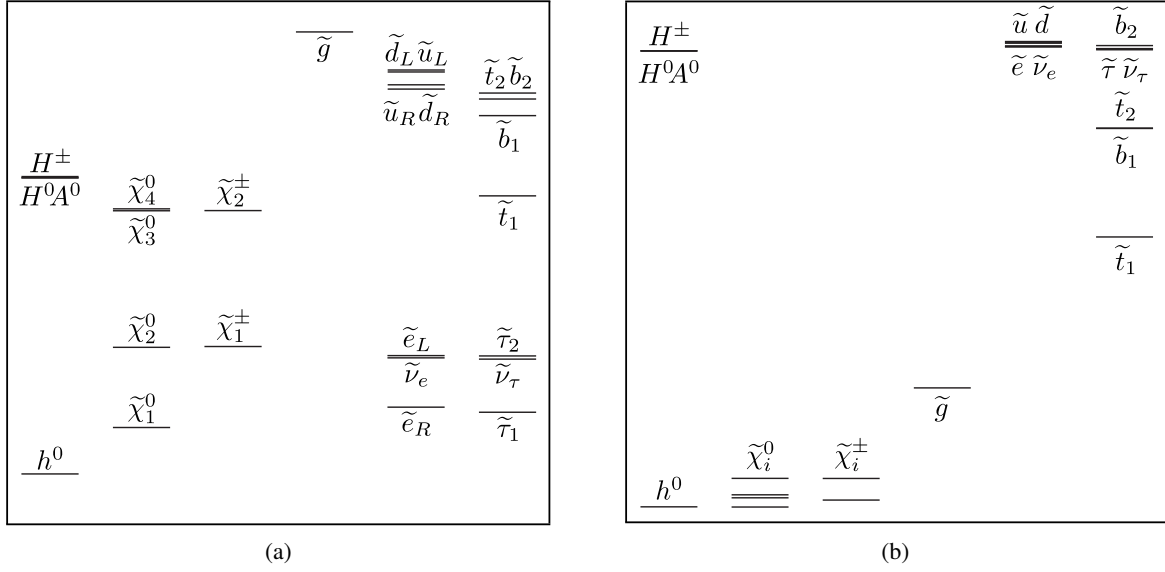


Figure 13: The figure shows mass hierarchies from two sample mSUGRA models. Heavier particles are higher up in the spectrum. One model (a) has $m_0 = 200$ GeV, $m_{1/2} = -A_0 = 600$ GeV, $\tan\beta = 10$ and $\mu > 0$, which yields heavy squarks and light sleptons. The other model (b) shows a more compressed scenario with $m_0 = 3200$ GeV, $m_{1/2} = -A_0 = 320$ GeV, $\tan\beta = 10$ and $\mu > 0$. These parameters yield a spectrum where left- and right-handed sleptons are nearly mass-degenerate, and all the gauginos are considerably lighter than the sfermions. Note that neither model is likely to describe reality, and the scale of the y-axis is deliberately not shown.

Table 9: The table contains a summary of the decays that could occur within the MSSM. Quarks and charged leptons are generically denoted by q and ℓ , with their actual charges suppressed to avoid clutter. Gaugino subscripts are such that $i > j$.

Particle	Decay products
$\tilde{\chi}_i^0$	$\rightarrow Z\tilde{\chi}_j^0, W^\pm\tilde{\chi}_j^\pm, h_0\tilde{\chi}_j^0, \ell\tilde{\ell}, \nu\tilde{\nu}, q\tilde{q}$
$\tilde{\chi}_i^\pm$	$\rightarrow W^\mp\tilde{\chi}_1^\pm, Z\tilde{\chi}_1^\pm, h_0\tilde{\chi}_1^\pm, \ell\tilde{\nu}, \nu\tilde{\ell}, q\tilde{q}'$
$\tilde{\ell}$	$\rightarrow \ell\tilde{\chi}_i^0, \nu\tilde{\chi}_i^\pm$
$\tilde{\nu}$	$\rightarrow \nu\tilde{\chi}_i^0, \ell\tilde{\chi}_i^\pm$
$\tilde{q}$	$\rightarrow q\tilde{g}, q\tilde{\chi}_i^0, q'\tilde{\chi}_i^\pm$
$\tilde{g}$	$\rightarrow q\tilde{q}$

2.7 The pMSSM

In recent years, the LHC (see Section 3) community has put significant work towards searching for supersymmetry in constrained models like mSUGRA and GMSB. With no hints of SUSY so far, it is reasonable to question whether the assumptions made when constructing such models are well-founded in reality. An alternative approach is to make no assumptions about physics beyond the MSSM, and try to reduce the phase space based only on experimental constraints. This yields the *phenomenological* Minimal Supersymmetric Standard Model (pMSSM) [36]. It should be thought of as the minimal supersymmetric model that is still phenomenologically viable at the LHC.

The pMSSM is based on three assumptions. The first assumption is that supersymmetry does not introduce any new sources of CP violation. Experimental constraints on sources of CP violation from beyond the Standard Model come from measurements of the electron and neutron electric moments in the Kaon system. Eliminating any potential CP violation from SUSY is done by setting all phases in $\mathcal{L}_{\text{Soft}}$ to zero. The second assumption is that supersymmetry does not introduce any flavor-changing neutral currents, which is a process that has never been experimentally observed. This is achieved by setting all non-diagonal elements of the sfermion mass and trilinear coupling matrices to zero. The third assumption is that the first- and second-generation sfermions are mass-degenerate. Their mass splitting is subject to severe experimental constraints from K^0 - $\bar{K}^0$ mixing. This assumption also implies that the trilinear couplings of the first two generations can be set to zero without phenomenological consequences, since only the third (and lightest) sfermion generation will be relevant in the context of an experimental search. Table 10 shows the remaining free parameters. There are 19 in total, but only a subset will be relevant for any given search.

Table 10: The table shows the free parameters of the pMSSM as well as their meanings.

Parameter	Meaning
$\tan\beta$	See Eq. (108)
μ	Higgsino mass parameter
M_A	A^0 mass
M_1	Bino mass
M_2	Wino mass
M_3	Gluino mass
$m_{\tilde{q}}, m_{\tilde{u}_R}, m_{\tilde{d}_R}$	First- and second-generation squark masses
$m_{\tilde{\ell}}, m_{\tilde{e}_R}$	First- and second-generation slepton masses
$m_{\tilde{Q}}, m_{\tilde{t}_R}, m_{\tilde{b}_R}$	Third-generation squark masses
$m_{\tilde{L}}, m_{\tilde{\tau}_R}$	Third-generation slepton masses
A_t, A_b, A_τ	Third-generation trilinear couplings

3 The Large Hadron Collider

3.1 Overview

The Large Hadron Collider (LHC) is a particle accelerator operated by the European Organization for Nuclear Research (CERN). It is the largest and most powerful of its kind, designed to collide protons at a center-of-mass energy of $\sqrt{s} = 14$ TeV [37]. The collisions take place in designated interaction points (IPs) which are surrounded by large particle detectors. New particles are created from the kinetic energy carried by the protons as they collide. The function of the detectors is to measure the properties of those particles. This section will briefly introduce the LHC and the particle detectors that it hosts.

The LHC is located outside the city of Geneva, crossing the border between Switzerland and France. It is hosted in an underground tunnel that is circular in shape with a circumference of 27 km. Figure 14 shows a schematic diagram of the LHC and its experiments. Also shown is the chain of pre-accelerators [38] used to inject protons into the LHC. The chain starts with a linear accelerator called LINAC 2. It is followed by the Proton Synchrotron Booster, the Proton Synchrotron (PS) and the Super Proton Synchrotron (SPS), which supplies the LHC with 450 GeV protons. The LHC then accelerates them to their final collision energy. This energy varied considerably from $\sqrt{s} = 900$ GeV (no acceleration) during the startup in 2008 to $\sqrt{s} = 8$ TeV prior to the first long shutdown at the end of 2012. This period of time is referred to as *run I*. Run II will commence at the start of 2015, at which point the LHC plans to operate at $\sqrt{s} = 13$ TeV with the goal of eventually ramping up to the design energy of $\sqrt{s} = 14$ TeV.

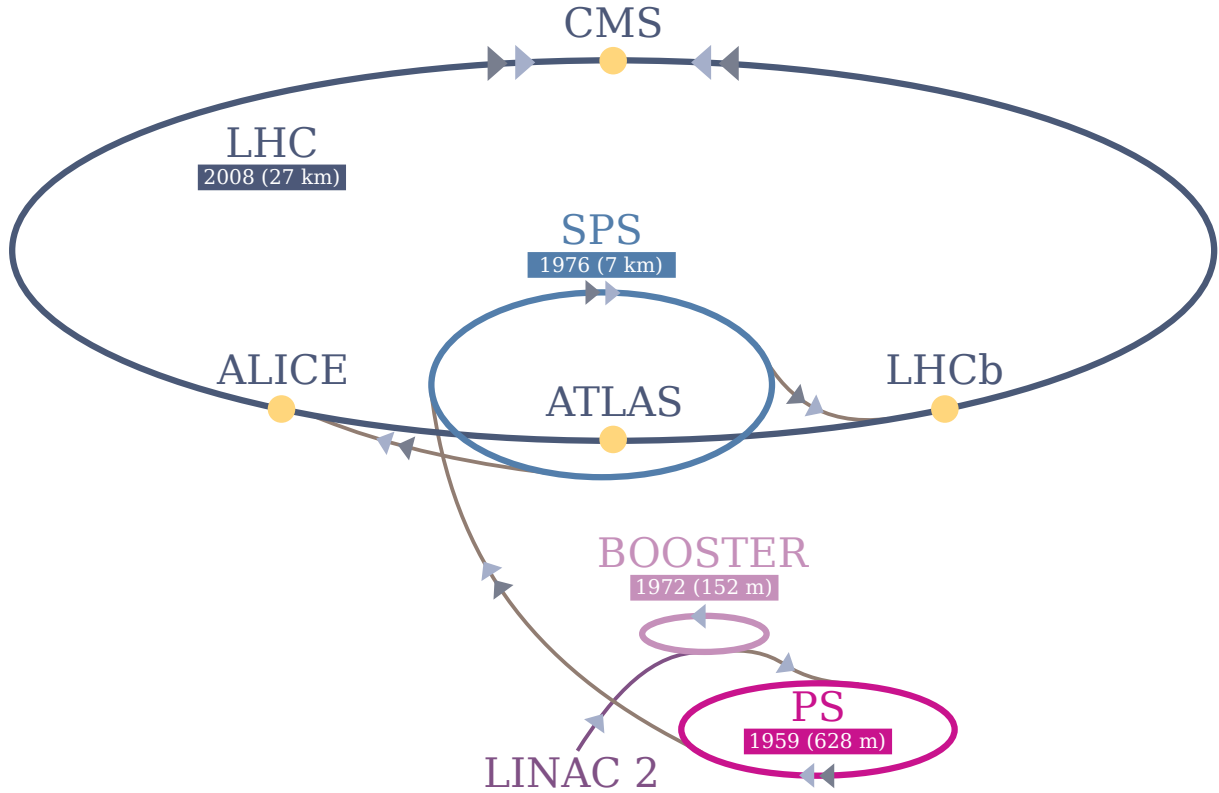


Figure 14: Schematic diagram of the LHC, its injector chain and main experiments. The boxes show the commissioning date and circumference of each accelerator.

The tunnel that now hosts the LHC was previously used to host the Large Electron-Positron (LEP) Collider. It is instructive to highlight some of the differences between the two accelerators. As the name

suggests, LEP collided electrons and positrons. The maximum center-of-mass energy of the collisions was steadily increased throughout the lifetime of the accelerator, reaching $\sqrt{s} \approx 200$ GeV in the final stages of its operation. The advantage of colliding electrons and positrons is that they are elementary. The entire center-of-mass energy is then available for the creation of new particles. This is not the case for proton–proton collisions. Protons are not elementary particles. They are made out of partons (see Section 1.4.1). A collision between two protons is really a collision between two partons, and each parton carries only a fraction of the total proton momentum. For a given collision, the size of that fraction is unknown. The collision energy is therefore *random* and *strictly less* than the center-of-mass energy of the proton–proton system. This means that electron-positron collisions, where the energy is known exactly, are often better suited for precision measurements. The advantage of proton–proton colliders lies elsewhere. Due to synchrotron radiation, charged particles lose energy at a rate proportional to $1/m^4$ as they travel around the accelerator ring. This severely limits the maximum achievable center-of-mass energy when accelerating light particles such as electrons and positrons. For heavier particles, like protons, synchrotron radiation is not a limiting factor. This is the main reason why the LHC is able to reach greater collision energies than LEP.

Finally, it should be noted that the LHC is not exclusively a proton collider. During part of the year, it operates as a lead ion collider instead. Lead ion experiments are, however, outside the scope of this text and will not be discussed further.

3.2 Accelerator system

The structure of the LHC ring alternates between curved sections, called *arcs*, and straight sections called *insertions*. There are eight arcs and eight insertions in total. Each arc hosts dipole magnets that are used to bend the beam, keeping it locked in a circular orbit. The magnetic field strength of the dipoles is a key design parameter as it effectively sets the maximum proton energy. The bending radius r of a proton with momentum p is related to the magnetic field strength B according to

$$Br = \frac{p}{e} \quad (109)$$

where e is the elementary electric charge. As the protons in the beam gain momentum, the magnetic field strength of the bending magnets is ramped up in direct proportion, eventually reaching 8.33 T at $\sqrt{s} = 14$ TeV. Higher-order multipole magnets, such as quadrupoles and sextupoles, are also positioned around the ring. Their purpose is to keep the beam focused in the transverse direction. All magnets are superconducting and cooled down to 1.9 K by liquid helium.

Some insertions host cavities over which a time-varying electric field is applied. The angular frequency of this field is matched to the revolution frequency of the protons. To achieve acceleration, the field frequency is tuned so that protons passing through a cavity always encounter a positive field. Figure 15 shows three protons arriving in time to be accelerated. Cavities of this type are used in all manners of particle accelerators. They are called *RF cavities* because the electric field is typically in the radio frequency (RF) range.

The frequency of the electric field is increased as the protons gain energy (and therefore velocity). This is done in such a way as to keep the arrival time of the protons constant with respect to the zero crossing of the RF field. Otherwise the protons would arrive too early and encounter a negative electric field, leading to deceleration. Protons whose revolution frequency is exactly matched to the RF field are said to be in *synchronous orbit*. Since every particle in the beam has a slightly different velocity, this condition can not be fulfilled for all protons simultaneously.

RF cavities have a natural compensating mechanism that corrects for differences in proton velocity and leads to *bunching* of the beam. This is also illustrated in Figure 15. The figure shows three protons A, B and C with different arrival times. Assume that the frequency of the field has been adjusted such

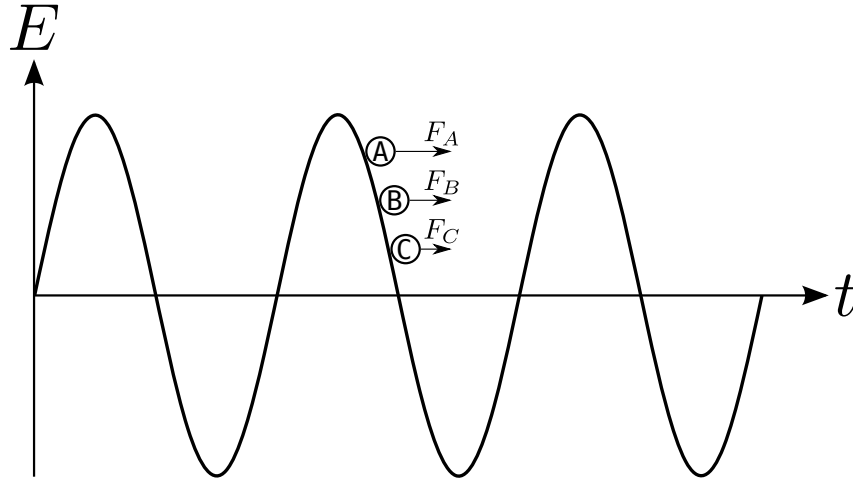


Figure 15: A radio frequency electric field is used to accelerate the protons. The figure shows how the field within an RF cavity varies as a function of time, as well as the arrival time of three test protons A, B and C. Proton A arrives first, followed by B and C. The accelerating forces that act on the protons are indicated.

that proton B is in synchronous orbit. The arrival time of the other protons is related to their energy. Protons with *less* energy will bend more in the arcs, thereby taking a shorter path around the ring. They eventually arrive *earlier* than the synchronous proton. This case corresponds to proton A. As a result, proton A encounters a stronger electric field, which compensates for the fact that it had less energy. Conversely, protons with *more* energy will bend less in the arcs, thereby taking a longer path around the ring. They eventually arrive *later* than the synchronous proton. This case corresponds to proton C. As a result, proton C encounters a weaker electric field, which compensates for the fact that it had more energy. Over time, proton A and proton C will walk back and forth along the time axis of Figure 15, oscillating around the synchronous proton B. This mechanism causes an initially continuous beam of protons to form bunches and prevents existing bunches from dispersing. It is a form of longitudinal focusing, working in tandem with the transverse focusing provided by the magnet system.

The revolution frequency of the protons is 11.245 kHz and the frequency of the RF field is 400.8 MHz. The RF field therefore completes 35640 periods during the time it takes for a bunch to make a full turn around the accelerator ring. This means that there is room for exactly 35640 bunches in the LHC. Each slot is referred to as an *RF bucket*. However, due to the beam structure of the SPS (which is responsible for injecting protons into the LHC) only one out of every ten buckets is actually filled. Each such bucket is assigned its own ID number. There are 3564 *bunch crossing IDs* (BCIDs) in total. Two neighboring BCIDs are separated in time by 25 ns. A typical beam structure alternates between regions with many filled BCIDs (called *bunch trains*) and regions with many empty BCIDs. Most physics runs in 2012 used a 50 ns bunch spacing within a train and had 1368 filled bunches in total. It is foreseen that runs in 2015 will use a 25 ns bunch spacing (thereby filling every consecutive bunch within a train) and have up to 2808 filled bunches in total.

3.3 Experiments

Figure 14 shows the four main LHC experiments. ATLAS (A Toroidal LHC Apparatus) [39] and CMS (Compact Muon Solenoid) [40] are general-purpose detectors. Their main goal is to look for any sort of new physics. Supersymmetry falls into this category, and so did the Higgs boson before it was independently discovered by each experiment in 2012. Another goal is to provide precision measurements

of various Standard Model parameters such as the weak mixing angle and the mass of the top quark. ATLAS is described in detail in Section 4. LHCb (Large Hadron Collider beauty) [41] looks for CP violation in the decays of B-mesons. CP violation is necessary to explain the excess of matter over anti-matter in the universe, but models suggest that the level of CP violation present in the SM is not sufficient to account for the full difference. ALICE (A Large Ion Collider Experiment) [42] is designed to study a phase of matter called quark-gluon plasma, in which quarks and gluons behave as free particles. The LHC creates quark-gluon plasma when colliding lead ions. It is a state of extremely high temperature and pressure thought to be similar to the conditions of the early universe. The goal of ALICE is to gain a better understanding of QCD. In addition to the four main experiments, several smaller ones exist and yet more have been proposed. The minor experiments are not discussed in this text.

4 The ATLAS experiment

4.1 Overview

ATLAS [43–45] is the largest particle detector in the world. It measures roughly 25 m in height and 44 m in length, with a total weight of 7000 tonnes. The ATLAS collaboration involves about 3000 scientists from 175 labs and universities in 38 countries. This section will briefly introduce the main detector subsystems and their functions. Each subsystem is treated in more detail later on.

As explained in Section 3.1, the LHC collides protons in designated interactions points around the accelerator ring. ATLAS is built around one such interaction point. The purpose of the detector is to measure the charge, momentum, kinetic energy and various other properties of the particles that are created in the collisions. ATLAS can be thought of as a large cylinder that is divided into layers. Each layer is a separate subsystem designed to measure a specific property of the particles that pass through it. As the particles travel outwards from the IP, they traverse the layers in turn starting with the innermost one. Some particles, such as electrons and protons, stop before reaching the outermost layer while other particles, such as muons and neutrinos, escape the detector entirely. Figure 16 shows an overview of ATLAS as seen from the outside. A slice of the detector has been removed to reveal the subsystems within.

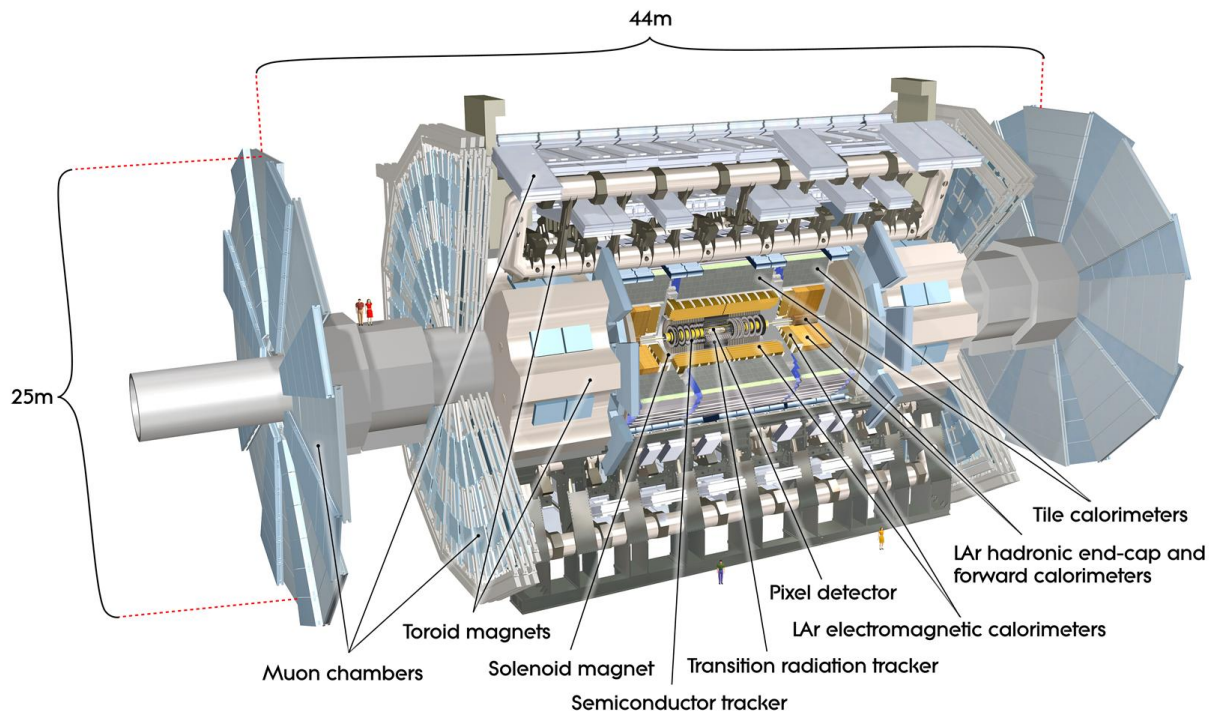


Figure 16: An overview of the ATLAS detector [45].

Figure 17 shows a cross-sectional view of ATLAS with the various layers indicated by different colors. The innermost layer is the pixel detector. This is a tracking detector whose purpose is to record the passage of charged particles. It is followed by two additional tracking detectors called the semiconductor tracker (SCT) and transition radiation tracker (TRT). These three trackers are collectively referred to as the inner detector (ID). The inner detector is surrounded by a solenoid that generates a powerful magnetic field. Charged particles traversing the ID therefore travel in curved paths. The trackers can determine

the momentum of a charged particle by recording its track and measuring the radius of curvature (see Eq. (109)). The very same measurement is used to determine the *sign* of the charge, as positively and negatively charged particles bend in opposite directions. Electrically neutral particles leave no track in the ID.

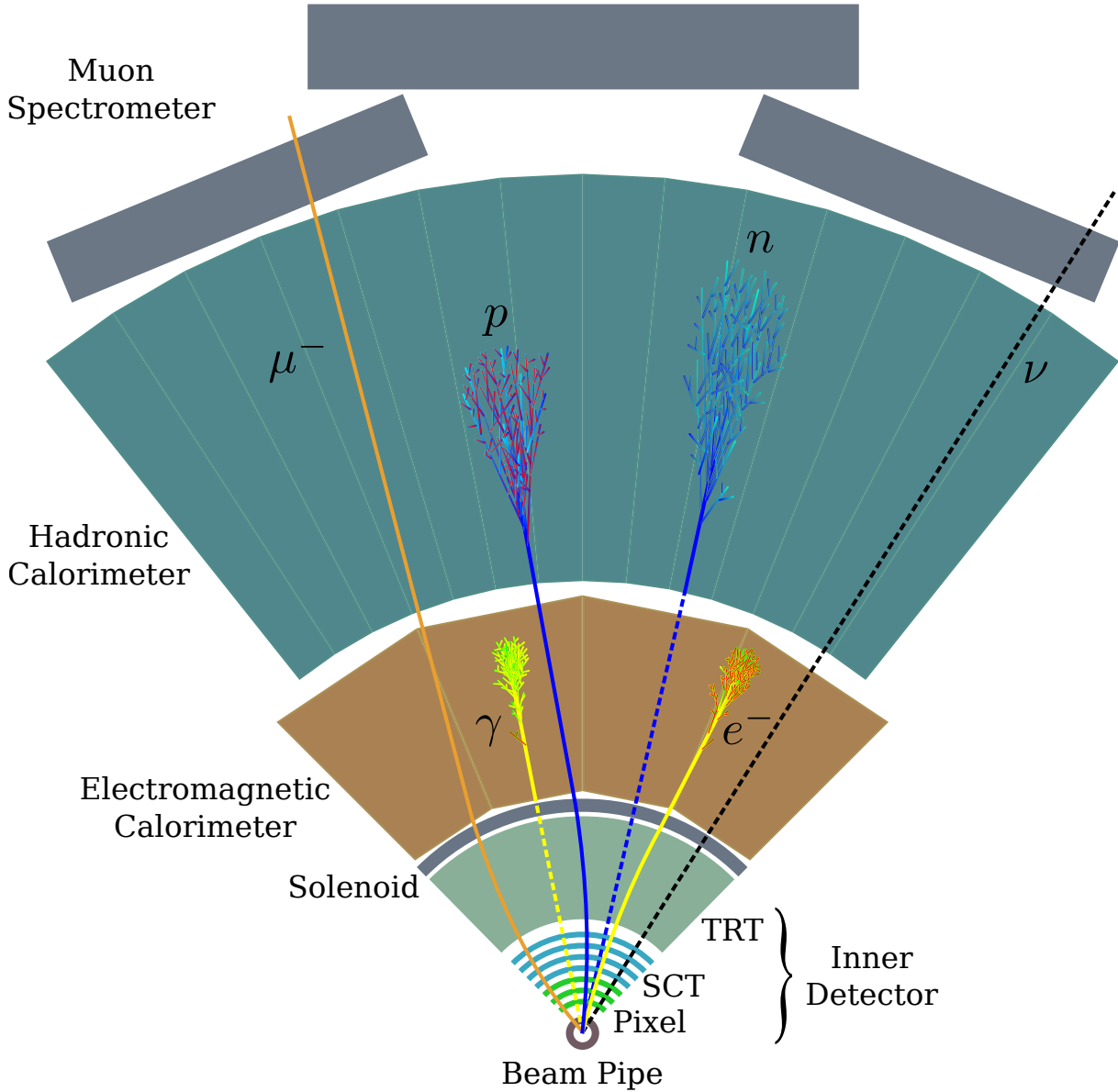


Figure 17: This schematic shows a cross-sectional view of ATLAS with the main subsystems indicated by different colors. Also shown are the signatures left by various particles as they traverse the detector. A dashed line indicates that the particle leaves no track. The figure is not to scale.

The solenoid is followed by the electromagnetic (EM) and hadronic calorimeters. A calorimeter is a device that measures the kinetic energy of particles. As the name suggests, the electromagnetic calorimeter is designed to measure the energy of particles that interact with matter primarily via the electromagnetic force. This category includes electrons, positrons and photons. The material of the EM calorimeter is chosen such that electrons and positrons have a high probability of emitting photons in the form of bremsstrahlung. Those photons subsequently undergo pair production yielding additional

electrons and positrons, and the process repeats. The resulting cascade of particles is called an *electromagnetic shower*. Since the shower continuously loses energy by interacting with the calorimeter material, it eventually stops. The calorimeter can determine the kinetic energy of the initial particle by measuring the total energy deposited by the shower.

All charged particles deposit some amount of energy in the EM calorimeter, but only electrons, positrons and photons are actually stopped. Hadrons, which interact with matter primarily via the strong force (i.e. with atomic nuclei rather than electron shells), are able to reach the hadronic calorimeter. Examples of particles that fall in this category are protons, neutrons, pions and kaons. The material of the hadronic calorimeter is chosen to promote the development of *hadronic showers*. Incoming hadrons collide with the nuclei of the calorimeter material, and new particles (mostly protons, neutrons and pions) are created. Those particles then collide with other nuclei, and the process repeats. Since neutral pions decay rapidly via the process $\pi^0 \rightarrow \gamma\gamma$, hadronic showers also develop an electromagnetic component. The kinetic energy of the initial particle is again obtained by measuring the total energy of the shower.

Muons are able to penetrate both the electromagnetic and hadronic calorimeters, making it all the way to the muon spectrometer. This is the outermost layer of ATLAS. Similarly to the ID, the muon spectrometer is a tracking detector. It is immersed in a magnetic field provided by the toroids (see Figure 16). The muons, being charged particles, bend in the magnetic field. Equation (109) is used to calculate their momenta based on the radius of curvature of their tracks. The reason why muons do not shower and stop in the EM calorimeter is that the probability for a charged particle to emit bremsstrahlung goes as $1/m^2$, and muons are about 200 times heavier than electrons.

A combination of subdetector responses is used to identify and distinguish between different types of particles. Electrons can be distinguished from photons since only charged particles leave tracks in the ID. Protons are distinguished from neutrons in the same way. Muons are identified using the fact that no other type of particle gives rise to a track in the muon spectrometer. Finally, neutrinos do not leave a track in any of the layers and escape undetected. Their only signature is in the form of *missing energy*. Since momentum is a conserved quantity, the vector sum of all the particle momenta in the plane transverse to the beam should always be zero. If the transverse momenta of the detectable particles in an event do not add up to zero, the difference is attributed to undetectable particles such as neutrinos.

4.2 Coordinate system

ATLAS uses a right-handed coordinate system with the z-axis pointing along the beam direction, the y-axis pointing vertically upwards and the x-axis pointing towards the middle of the accelerator ring. The origin is placed at the interaction point. The side of the detector that lies along the positive z-axis is called “Side A” while the side that lies along the negative z-axis is called “Side C”. Figure 18 illustrates how various angles and distances are defined in the ATLAS coordinate system. The polar angle θ is measured from the beam axis while the azimuthal angle ϕ is measured from the x-axis in the plane perpendicular to the beam. Another way to quote the polar angle is in terms of the pseudorapidity η , which is defined according to:

$$\eta = -\ln\left(\tan\frac{\theta}{2}\right) \quad (110)$$

Figure 19 shows a graphic representation of the relationship between η and θ . The region defined by large η is called the *forward region*. The start of the forward region is generally taken to be somewhere around $\eta = 2.5$.

The direction of particle tracks is typically given in η - ϕ space. Distances in this space are denoted ΔR and calculated according to:

$$\Delta R = \sqrt{(\Delta\eta)^2 + (\Delta\phi)^2} \quad (111)$$

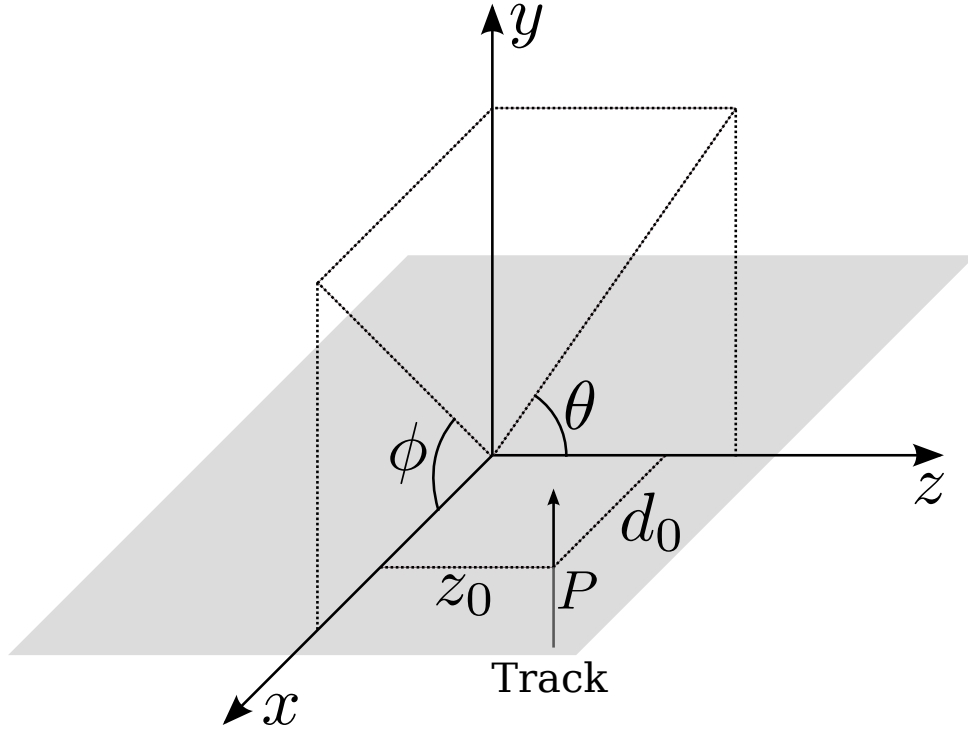


Figure 18: The figure shows how various angles and distances are measured in the ATLAS coordinate system. If two protons just collided at the origin, then the particles emerging from the vertex would have polar angles θ and azimuthal angles ϕ . A close-by track would have impact parameters z_0 and d_0 with respect to the vertex.

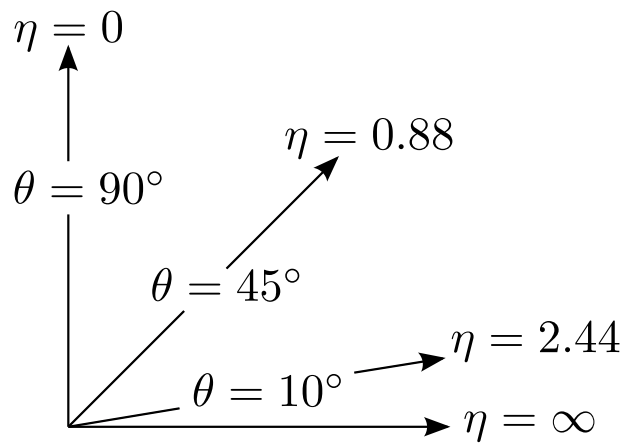


Figure 19: The figure illustrates the relationship between the polar angle θ and the pseudorapidity η .

This should not be confused with a plain R , which is sometimes used to denote the radial distance from the beam pipe. The cartesian distance between a track and a vertex (collision point) is expressed in terms of *impact parameters*. Let P denote the point of closest approach. The longitudinal and transverse impact parameters z_0 and d_0 measure the distance between P and the vertex along the beam axis and in the plane perpendicular to the beam axis, respectively.

4.3 Magnet system

The ATLAS magnet system [46] consists of a central solenoid, a barrel toroid and two end-cap toroids. Their structure is shown in Figure 20. A second look at Figure 16 may also be useful to put them into context. The central solenoid surrounds the inner detector and provides it with a 2 T magnetic field. Having a strong field in the ID results in a good momentum resolution for the trackers, but it also means that low-energy particles are unable to escape the solenoid. Instead, they curl up and cause many unnecessary hits in the inner detector. This leads to ambiguities when matching tracks to particles. The chosen magnetic field strength is a tradeoff between the two effects. It implies a minimum escape energy of 335 MeV for particles emitted at $\eta = 0$.

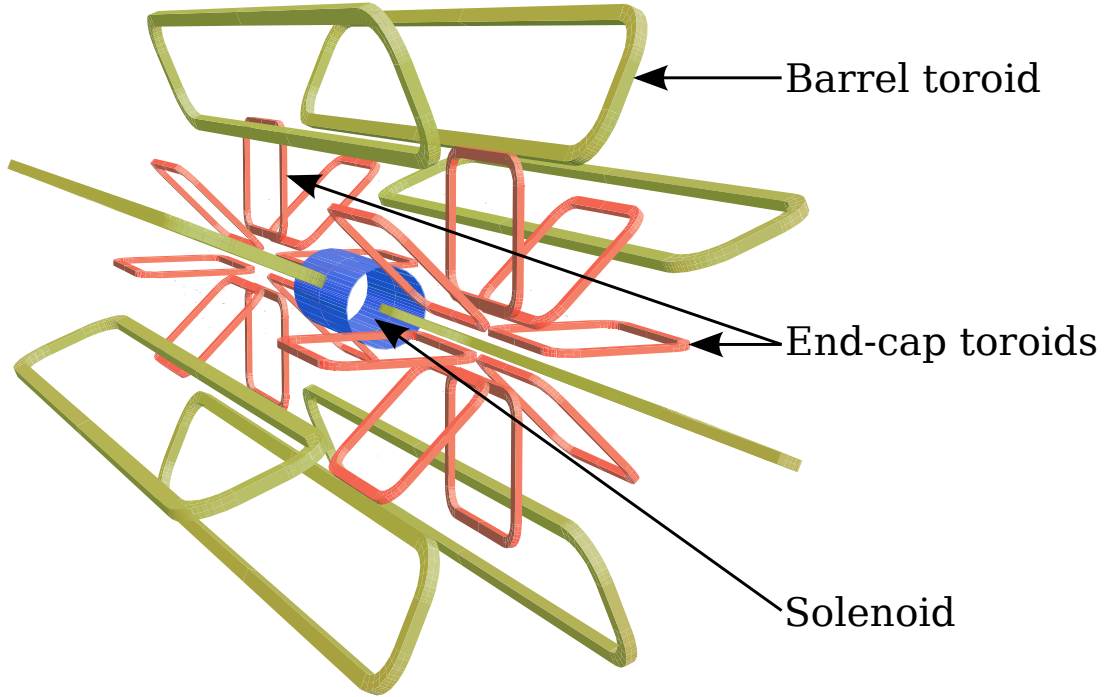


Figure 20: The ATLAS magnet system consists of a central solenoid, a barrel toroid and two end-cap toroids. The solenoid provides a magnetic field for the inner detector. The toroids provide a magnetic field for the muon spectrometer. Figure adapted from Ref. [45].

The solenoid has a length of 5.8 m, a diameter of 2.56 m and a thickness of a mere 5 cm. Keeping the thickness at a minimum was one of its main design goals. This is because some particles have to pass through the solenoid before reaching the calorimeters. In doing so, they lose energy by interacting with the solenoid material. Since the calorimeter measurement relies on the particles depositing their full energy within the calorimeter volume, any energy deposited elsewhere leads to a less accurate measurement.

The barrel toroid consists of 8 coils assembled symmetrically around the beam axis. The coils have a length of 25.3 m, an inner diameter of 9.4 m and an outer diameter of 20.1 m. This means that they

are located within the space occupied by the muon barrel spectrometer. The end-cap toroids also consist of 8 coils each. They are situated about 10 m away from the interaction point and have an outer diameter of 10.7 m, which places them within the space occupied by the muon end-cap spectrometers. In terms of η , the barrel toroid covers the region $|\eta| < 1.6$ while the end-caps cover the region $1.4 < |\eta| < 2.7$. The average magnetic field provided by the barrel and end-cap toroids is 0.5 T and 1 T, respectively, although the exact value is very location-dependent. In particular, the overlap region in the range $1.4 < |\eta| < 1.6$ has a much lower field as the barrel and end-cap contributions tend to cancel each other out. The toroid field lines run azimuthally around the ATLAS cylinder such that the bending force exerted on charged particles is directed along the z -axis.

4.4 Inner detector

The inner detector [47, 48] is the part of ATLAS located closest to the beam pipe. It is a tracker, which means that it records the passage of charged particles. One of the main design concerns when constructing a tracker is to minimize the amount of detector material. Precision trackers are made up of several layers of high-granularity detectors. Each layer contributes one space point to the final track. Adding more layers increases the precision of the measurement. However, since a particle loses energy whenever it interacts with the detector material, any additional layers also decrease the accuracy of the subsequent calorimeter measurement. The aim, therefore, is to use as little material as possible to achieve a given tracking precision.

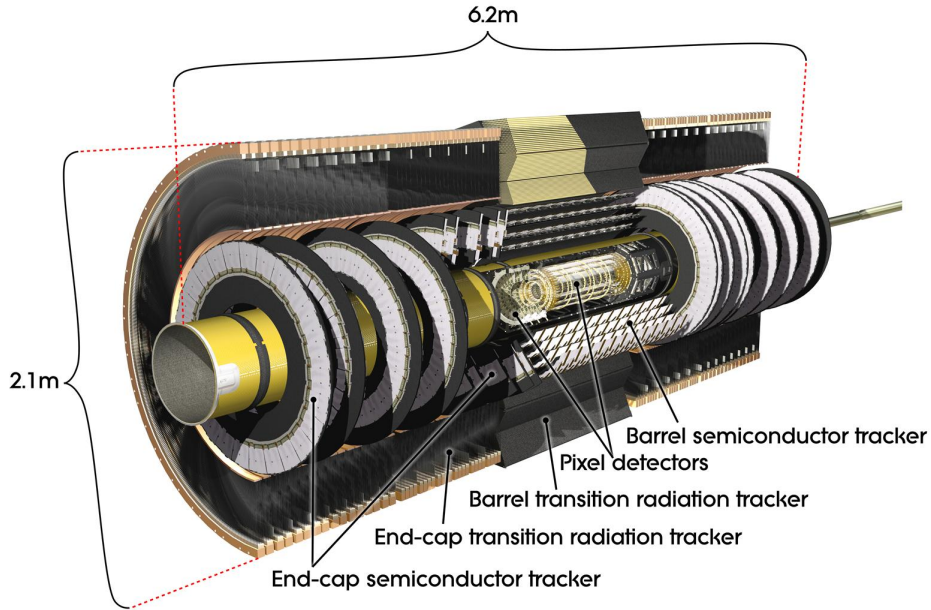


Figure 21: An overview of the inner detector [45].

The particle density in ATLAS is at a maximum close to the IP. It then drops off as $1/R^2$, where R is the radial distance from the collision point. This means that the granularity of a tracker must be high close to the IP, whereas further away a lower granularity might suffice. Very precise trackers are also prohibitively expensive in large volumes. These concerns have guided the design of the ATLAS inner detector. Figure 21 shows an overview of the ID, which is actually made up of three separate subdetectors. The innermost one is called the pixel detector. It consists of three layers both in the barrel and the end-caps. The pixel detector is followed by four layers of silicon microstrip detectors

in the barrel and nine in the end-caps. This is the semiconductor tracker (SCT). The outermost part of the ID is called the transition radiation tracker (TRT). It is made up of gaseous straw tubes interleaved with transition radiation material and provides lower granularity continuous tracking. Figure 22 shows a detailed schematic of the inner detector and its most important parts. The ID as a whole covers the range $|\eta| < 2.5$.

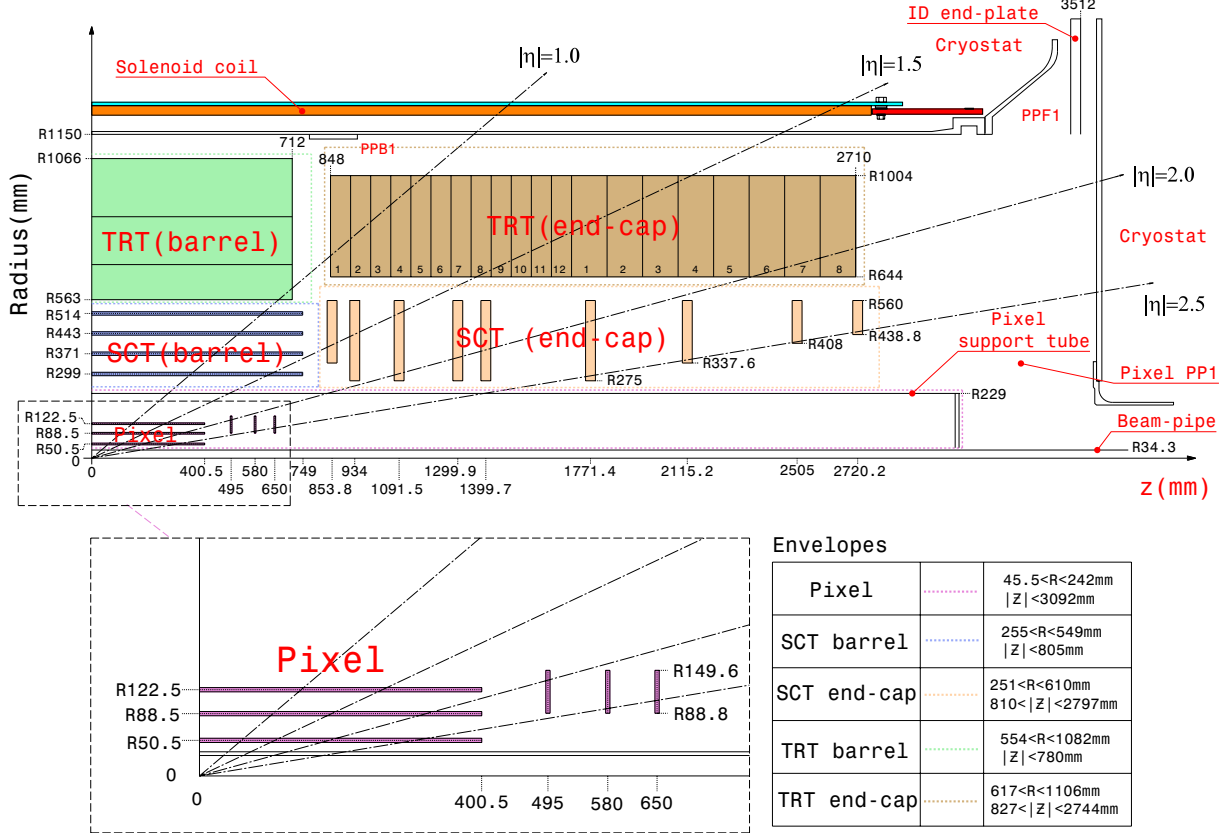


Figure 22: Detailed schematic of the inner detector [45]. The positions of the main detector components as well as the overall subdetector envelopes are shown.

4.4.1 Pixel detector

The innermost part of the ID is a collection of very small semiconductor detectors called *pixels* [49]. The pixels have a nominal surface area of $50 \mu\text{m} \times 400 \mu\text{m}$ and a thickness of $250 \mu\text{m}$. The precision of the tracking, however, is much better than that. This is because particles produce a signal in many neighboring pixels as they travel through the detector. Calculating the centroid of those signals yields the most likely position of the track. The barrel layers of the pixel detector have an accuracy of $10 \mu\text{m}$ in $R-\phi$ and $115 \mu\text{m}$ in z , while the disks have an accuracy of $10 \mu\text{m}$ in $R-\phi$ and $115 \mu\text{m}$ in R . The pixels represent more than half of the readout channels in all of ATLAS, with a total of 80.4 million individual detector elements.

The pixel detector plays a crucial role in identifying jets originating from b -quarks. Such jets are called b -jets and the technique used to identify them is referred to as b -tagging. Hadrons containing b -quarks have typical lifetimes on the order of a picosecond and therefore travel somewhere around $300 \mu\text{m}$ before decaying. If a primary vertex (from a pp collision) is reconstructed with a secondary vertex $300 \mu\text{m}$ away, then the secondary vertex is probably due to the decay of a b -hadron, and a jet

originating from this vertex can be tagged as a b -jet. Many physics analyses rely on the pixel detector to distinguish between b -jets and jets from light-flavor quarks.

4.4.2 Semiconductor tracker

The SCT is an arrangement of silicon microstrip detectors. Each strip has a length of 6 cm, a width of 80 μm and a thickness of 285 μm . Naturally, because the length of the strips is several orders of magnitude greater than the width, their tracking precision is much better in the width dimension. For this reason, each layer of the SCT consists of two back-to-back planes of strips. The planes are rotated with respect to each other so that the lack of precision in the length dimension of one plane is compensated for by the other plane. Using a rotation angle of 90° would yield the same accuracy in both the length and width dimensions, while any other angle would increase the accuracy in one dimension at the expense of the other. Since accuracy in the R - ϕ plane is more important for physics analyses, the chosen angle is only 40 mrad. This choice yields a precision of 17 μm in R - ϕ and 580 μm in z for the barrel layers. The disks have a precision of 17 μm in R - ϕ and 580 μm in R . There is a total of 6.3 million readout channels in the SCT.

4.4.3 Transition radiation tracker

The TRT is made up of approximately 0.3 million drift tubes filled with a gaseous mixture of 70% Xe, 27% CO_2 and 3% O_2 . They have a diameter of 4 mm and a length that depends on their position. The tubes in the barrel are 144 cm long and oriented parallel to the beam, while the tubes in the end-caps are 37 cm long and lie in the plane perpendicular to the beam. Due to their dimensions, the tubes are often referred to as *straws*. The TRT occupies the majority of the space in the inner detector.

A gold-plated tungsten wire is suspended at the center of each tube, and the inside of the tube wall is coated with an electrically conductive material. The wire is kept at a 1500 V potential difference with respect to the wall. Charged particles traveling through the tubes ionize the gas, liberating electrons which are then collected by the anode wires. Since a charged particle typically causes several ionizations per tube, the relative drift times of the collected electrons can be used to extract information about the position of the track. The TRT only provides tracking information in the R - ϕ plane. It has a precision of 130 μm per straw.

The space between the tubes is packed with either fibers (in the barrel) or sheets (in the end-caps) made of polypropylene–polyethylene. This is a material with a high ($n \approx 1.5$) refractive index. When charged particles pass between two media with different refractive indices, they emit *transition radiation* in the form of photons. These photons also ionize the gas in the drift tubes. Because the intensity of the transition radiation is proportional to the Lorentz factor γ of the charged particle, the signal can be used for purposes of particle identification. In particular, it is used to distinguish between electrons and pions. Electrons are much lighter than pions and are typically produced with significantly greater values of γ at LHC energies. In cases where the information from other subdetectors is ambiguous, the amount of transition radiation produced by a particle may reveal its identity. The TRT provides electron identification in the range $|\eta| < 2.0$.

4.5 Calorimeters

The ATLAS calorimeter system [50–52] serves several functions. It provides positional information complementary to that of the inner detector, is essential for purposes of particle identification and works as a shield that prevents non-muons from producing signals in the muon spectrometer. The main purpose of the calorimeter system, however, is to measure the kinetic energy of particles emerging from the IP. The calorimeter volume is therefore filled with materials in which the particles are likely to interact.

Every time the particles interact with the detector material, they deposit some of their energy. The aim is to completely stop the particles. A fraction of the deposited energy is then measured and related to the total energy by a calibration procedure. Different types of radiation interact with matter in different ways, and there is no single type of calorimeter suitable for all of them. ATLAS therefore employs a variety of calorimeter designs. They can be divided into two principal categories. Electromagnetic calorimeters are used to measure the energy of electrons, positrons and photons while hadronic calorimeters are used to measure the energy of hadrons.

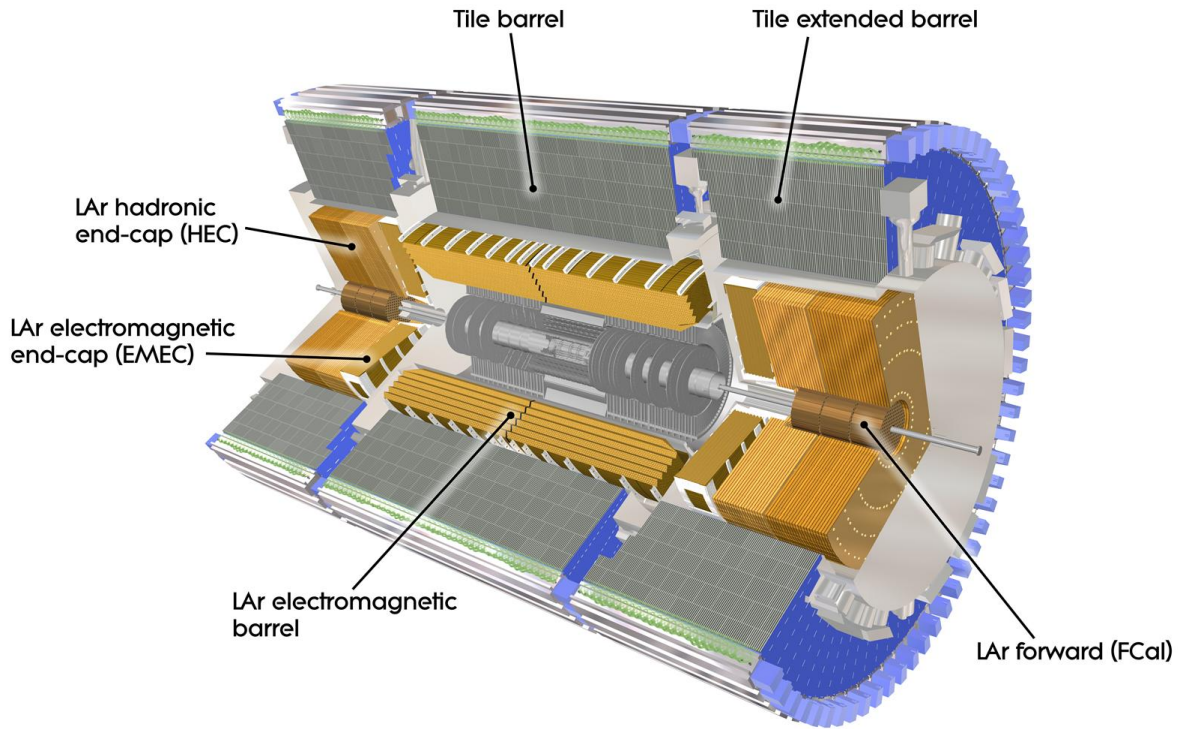


Figure 23: An overview of the calorimeter system [45]. The electromagnetic barrel and end-cap calorimeters measure the kinetic energy of electrons, positrons and photons. The tile and hadronic end-cap calorimeters measure the energy of hadrons. The forward calorimeter measures the energy of all particles, but is typically grouped together with the hadronic calorimeters.

Figure 23 shows an overview of the various calorimeters present in ATLAS. They are all so-called *sampling* calorimeters. A sampling calorimeter consists of alternating layers of *absorbers* and *active material*. The absorber material is chosen to promote showering of the incoming radiation. In electromagnetic calorimeters, charged particles interact by emitting bremsstrahlung which subsequently undergoes pair production. In hadronic calorimeters, the incoming radiation interacts mainly with the atomic nuclei of the absorber. Some of the particles created in the shower enter the active material. This is the part of the calorimeter that is responsible for the actual energy measurement. In the tile calorimeter, the active material consists of plastic scintillator tiles that are read out by photomultipliers. In the other calorimeters, the active material is liquid argon (LAr). Each liquid argon layer works similarly to a proportional chamber. Particles traversing the layer leave a trail of ionization, thereby producing free electrons. These electrons are accelerated by a potential difference that is applied across the argon by two electrodes. This results in a signal that is subsequently read out.

Although the calorimeters are made thick enough to stop even the most energetic of particles, only a fraction of that energy is deposited within the active material. Some is deposited in the absorbers, and some is lost to processes that do not result in a measurable signal (like neutrino production). A calorimeter is calibrated by finding the constant of proportionality between the energy deposited in the active material and the total energy carried by the original particle. This number is referred to as the *energy scale* of the measurement. It is known with an uncertainty of 1–2% for electrons and photons and 5–10% for jets. Because the interaction between matter and radiation is an inherently random process, the calorimeters also have a limited resolution. Each individual particle showers in a different way, and there is no guarantee that two identical particles will deposit the same amount of energy in the active material. The associated uncertainty is referred to as the *energy resolution* of the measurement. The relative uncertainty scales as $1/\sqrt{E}$ and therefore gets better for high-energy particles. A resolution of $10\%/\sqrt{E}$ has been achieved for the electromagnetic calorimeters, where E is given in units of GeV. For the hadronic calorimeters, the resolution is $100\%/\sqrt{E}$ in the forward region and $50\%/\sqrt{E}$ everywhere else.

4.5.1 Electromagnetic calorimeter

The electromagnetic calorimeter is divided into a barrel part that covers the range $|\eta| < 1.475$ and two identical end-caps that cover $1.375 < |\eta| < 3.2$. Each end-cap is further divided into two coaxial wheels. The outer wheel, which covers $1.375 < |\eta| < 2.5$, has a higher granularity than the inner wheel. This division is done because the inner detector only covers $|\eta| < 2.5$. Particles outside of this range are often not used by physics analyses, so a lower accuracy can be afforded.

The structure of the detector is illustrated by Figure 24, which shows a detailed schematic of one of the barrel modules. The module has three layers. They are referred to as *samplings*. The first sampling is located closest to the beam pipe. It has a high granularity, which not only allows it to produce precise η measurements, but also makes it useful for purposes of particle identification. For example, the first sampling can distinguish between a single photon and a π^0 decaying into two photons that emerge almost collinearly. The thickness of the first sampling is $4.6X_0$, where X_0 denotes radiation lengths. The second sampling makes up the bulk of the calorimeter with a thickness of $16X_0$. This is where the main shower takes place. Finally, the third sampling has a thickness that varies between $2X_0$ and $12X_0$. Since electromagnetic showers are likely to be contained within the second sampling while hadronic showers will extend well beyond, the third sampling can be used to discriminate between them. In the end-cap outer wheel, the modules look similar to the barrel. In the inner wheel, the granularity is coarser ($\Delta\eta \times \Delta\phi = 0.1 \times 0.1$) and there are only two samplings.

All parts of the electromagnetic calorimeter use lead absorbers and employ liquid argon as the active material. Liquid argon is chosen primarily for its radiation hardness. The absorbers, which are covered by electrodes made of Kapton, are folded into an accordion shape as shown in Figure 24. This choice of geometry ensures complete ϕ symmetry without any azimuthal cracks. The argon resides in the space between the absorbers.

The amount of material in the inner detector and the solenoid corresponds to approximately $2.3X_0$ at $\eta = 0$. This means that many particles start to shower before even reaching the calorimeters, which degrades the quality of the subsequent energy measurement. In order to partially recover the energy resolution, a *presampler* is used in the range $|\eta| < 1.8$. The presampler detector contains a thin (11 mm in the barrel and 5 mm in the end-caps) layer of liquid argon situated in front of the first sampling.

4.5.2 Hadronic calorimeter

The hadronic calorimeter can be divided into three subsystems. They are called the tile calorimeter (TileCal), hadronic end-cap calorimeter (HEC) and forward calorimeter (FCal). TileCal is further di-

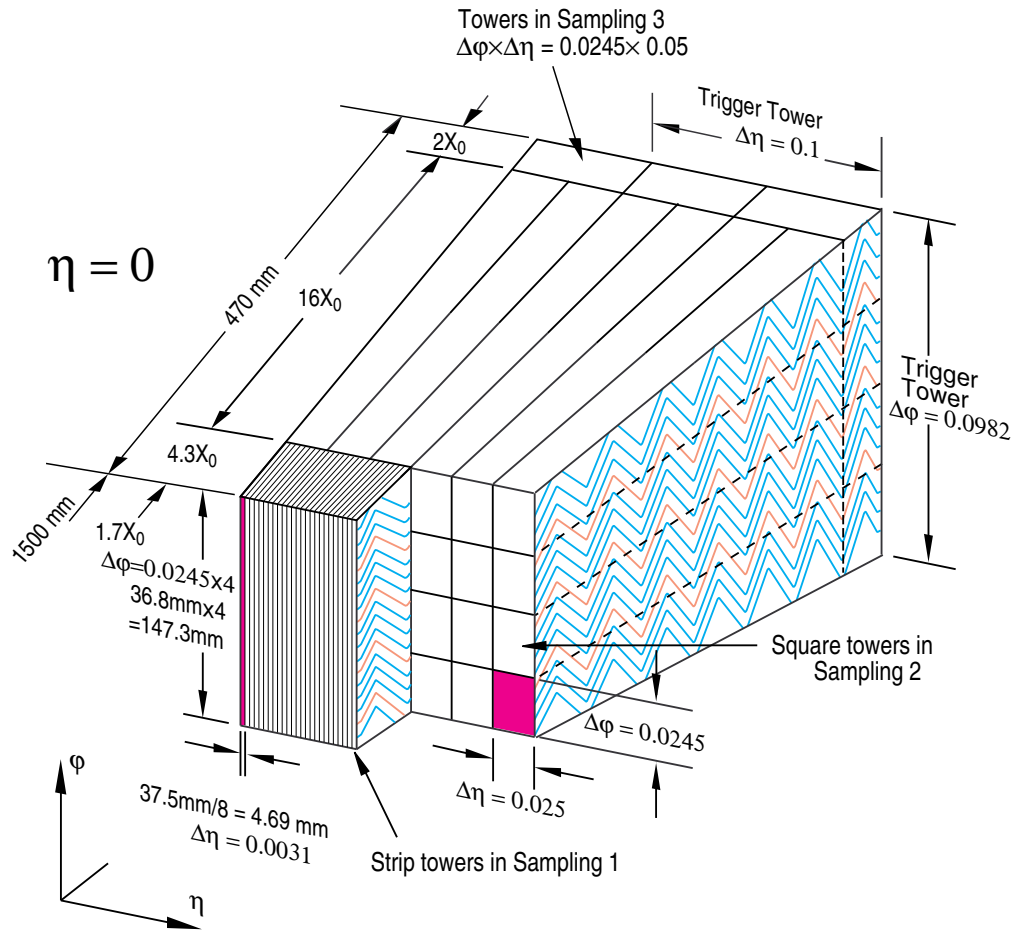


Figure 24: Detailed schematic of an electromagnetic calorimeter barrel module [50]. It is segmented into three sampling layers with granularity and thickness as indicated. The figure also shows how the calorimeter towers are grouped into so-called *trigger towers*. The trigger towers are used for the Level 1 trigger as described in Section 4.8.

vided into a barrel module that covers the range $|\eta| < 1.0$ and two extended barrel modules that cover $0.8 < |\eta| < 1.7$. Figure 25 shows a schematic diagram of one of the TileCal modules. It uses absorbers made of steel, while plastic scintillator tiles serve as the active material. This sets TileCal apart from the rest of the ATLAS calorimeter system, where the active material is liquid argon. However, because TileCal is shielded by the electromagnetic calorimeter and because it is situated in the barrel region, it is also exposed to the least amount of radiation. This motivates the choice of the less radiation-hard plastic. The tiles are oriented along the R - ϕ plane in order to maximize the granularity of the detector in η , yielding a precision of up to $\Delta\eta \times \Delta\phi = 0.1 \times 0.1$. Readout of the scintillators is done by photomultiplier tubes (PMTs) via wavelength shifting fibers. The total thickness of TileCal is 7.4λ , where λ denotes interaction lengths. For comparison, the thickness of the entire ATLAS detector at the outer edge of TileCal is 9.7λ .

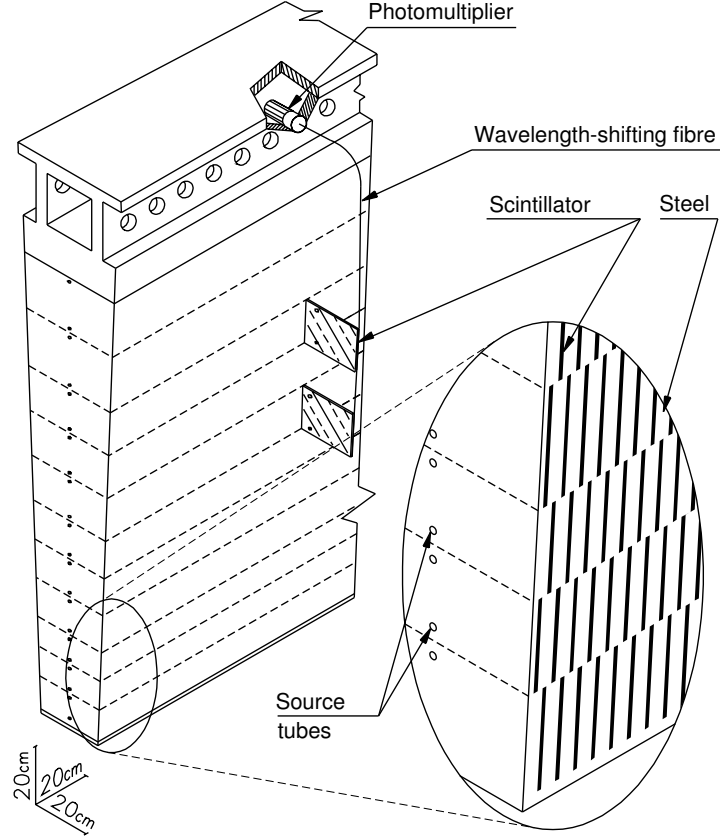


Figure 25: Schematic diagram of a TileCal module [45]. The plastic scintillators are read out by photomultipliers via wavelength shifting fibers.

The hadronic end-cap calorimeter resides in the end-caps and covers $1.5 < |\eta| < 3.2$. Each end-cap wheel is then further divided into a front and back wheel as shown schematically by Figure 26. The front wheel uses thinner absorber plates and is therefore able to provide a more precise positional measurement than the back wheel. The granularity is $\Delta\eta \times \Delta\phi = 0.1 \times 0.1$ in the region $|\eta| < 2.5$ and $\Delta\eta \times \Delta\phi = 0.2 \times 0.2$ elsewhere. All HEC absorbers are made of copper, and the active material is liquid argon.

FCal sits closest to the beam pipe, covering the region $3.1 < |\eta| < 4.9$. It experiences the highest particle flux out of all the calorimeters, which calls for special considerations when it comes to the choice of material. As shown by Figure 26, FCal is made up of three modules. The innermost module is built for electromagnetic calorimetry and uses absorbers made of copper to optimize heat removal. The two outer modules are optimized for hadronic calorimetry and use dense tungsten absorbers in order to

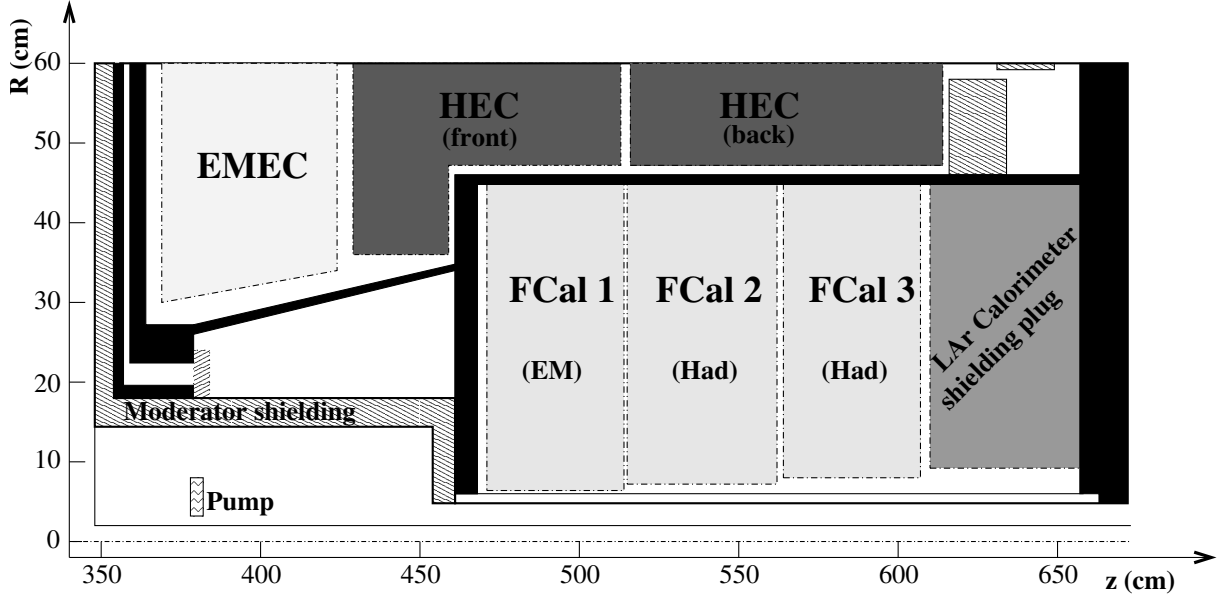


Figure 26: Schematic diagram showing a side-view of the calorimeters in the forward region [45].

minimize the lateral spread of hadronic showers. The active material in all the modules is liquid argon. Although small, FCal alone has a thickness of 10λ .

4.6 Muon spectrometer

The muon spectrometer [53] contains two categories of detectors. The first category is the trackers. They measure the position of the muons that traverse the spectrometer. Since the muons are charged particles, they bend in the magnetic field provided by the toroid (see Section 4.3). The muon momentum p is related to the bending radius r according to $p = eBr$ where e is the elementary charge and B is the magnetic field strength. A measurement of the muon sagitta can therefore be translated into a measurement of the muon momentum.

The second category is the trigger detectors. They are responsible for determining which bunch crossing a particular muon came from and making crude (but fast) measurements of the muon momentum. This information is then used by the Level 1 trigger as described in Section 4.8. In addition, the trigger detectors complement the trackers by measuring the muon position in the plane where the trackers are not sensitive. Figure 27 shows an overview of the muon spectrometer. It is by far the largest part of ATLAS.

In the barrel, tracking is handled by Monitored Drift Tube (MDT) chambers. Each chamber contains a number of gas-filled tubes in which a central wire is kept at a potential difference with respect to the tube wall. When muons traverse the tubes, they liberate electrons which drift towards the anode wires. The drift velocity is roughly constant over the size of the tube. Close to the anode, where the electric field is high, the electrons cause an avalanche and induce a measurable current in the wire. A trigger detector is used to determine the arrival time of a given muon in the tubes. The electrons liberated by that muon then require a time t_{drift} to drift to the wire of a tube. Since the electron drift velocity is known, a measurement of t_{drift} can be translated into a coordinate along the tube radius. This is the working principle of a drift chamber. There are three layers of chambers, which surround the beam pipe in concentric shells at radii of approximately 5 m, 7.5 m and 10 m. This places the first shell ahead of the toroid field, the second shell within the toroid and the third shell outside of the toroid. The trigger measurement is handled by a set of Resistive Plate Chambers (RPCs). These are gas-filled chambers in

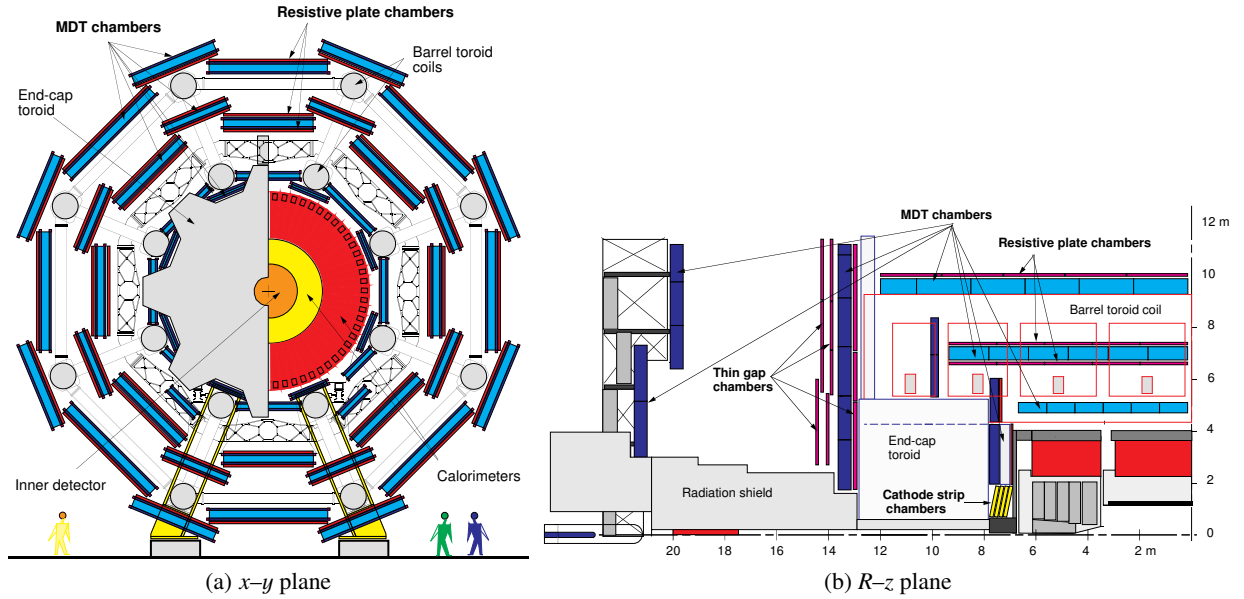


Figure 27: An overview of the muon spectrometer in the x - y plane (a) and R - z plane (b) [53].

which two parallel plates are kept at a potential difference with respect to each other. There are three layers of RPCs. They are mounted on the MTDs as indicated by Figure 27. The barrel covers $|\eta| < 1.05$.

Tracking in the end-caps is handled by a combination of Monitored Drift Tubes and Cathode Strip Chambers (CSCs). A CSC is a multiwire proportional chamber in which the cathode has been segmented into strips. They have a higher granularity than MDTs and are therefore used in the forward region close to the beam pipe, where the particle flux is high. The end-cap MDTs are arranged into three wheels located 7.4 m, 14 m and 21.5 m away from the IP. This places the innermost wheel ahead of the toroid and the two outermost wheels outside of the toroid. The wheels cover the range $1.05 < |\eta| < 2.7$, except for the innermost one which only goes up to $|\eta| < 2.0$. The remainder is covered by the Cathode Strip Chambers. The trigger measurement in the end-caps is done by a set of Thin Gap Chambers (TGCs) that surround the second MDT wheel. A TGC is essentially a multiwire proportional chamber with the characteristic that the wire-to-wire distance is smaller than the wire-to-cathode distance.

Due to the size of the muon spectrometer, alignment of the different subsystems becomes a major issue. The position of each detector element must be known with a precision better than $30 \mu\text{m}$ to fulfill the design goal of the tracking measurement. This is achieved by the use of approximately 12000 precision-mounted optical alignment sensors. Another important parameter is the magnetic field strength B , since any uncertainty in B immediately translates into an uncertainty in p . The magnetic field strength is measured by the use of 1800 Hall sensors distributed throughout the spectrometer volume. Table 11 summarizes the properties of each muon spectrometer subsystem.

4.7 Forward detectors

ATLAS contains a number of detectors dedicated to monitoring activity in the forward (high η) region. They are called, in order of increasing distance from the interaction point, BCM (Beam Conditions Monitor), MBTS (Minimum Bias Trigger Scintillators), LUCID (Luminosity Measurement Using a Cherenkov Integrating Detector), ZDC (Zero Degree Calorimeter) and ALFA (Absolute Luminosity For ATLAS). Every forward detector can be used to measure a property of the LHC beam called *luminosity*. Some of them are dedicated luminosity detectors, while others have more specialized tasks and perform luminosity measurements only as a secondary function. Luminosity is discussed further in Section 6.

Table 11: The table shows the intrinsic position and time resolution of each muon spectrometer subsystem. The values do not include uncertainties in the chamber alignment. Also shown is the number of measurement points per track.

Type	Function	Chamber resolution				Points per track	
		z	R	ϕ	t	Barrel	End-cap
MDT	Tracking	35 μm	–	–	–	20	20
CSC	Tracking	–	40 μm	5 mm	7 ns	–	4
RPC	Trigger	10 mm	–	10 mm	1.5 ns	6	–
TGC	Trigger	–	2-6 mm	3-7 mm	4 ns	–	9

4.7.1 BCM

The quality of the LHC beam depends on several parameters such as the settings of the bending magnets and RF system, the vacuum pressure in the beam pipe and the position of the ATLAS collimators. The effects interplay in complex ways that are difficult to simulate. Dedicated detectors are therefore used to monitor the beam conditions. BCM is such a detector. It is made up of two stations, one on either side of the interaction point, located within the inner detector at $z = \pm 184$ cm and $R = 5.5$ cm (corresponding to $|\eta| = 4.2$). Each station contains a set of fast diamond sensors that produce signals with a rise time of approximately 1 ns. Since the bunches are separated by at least 25 ns, this is sufficient to monitor beam background rates on a BCID-by-BCID basis.

A situation might arise in which the LHC beam is bent out of place and begins hitting the collimators in front of the ATLAS detector. This would produce intense particle showers traveling from one side of ATLAS to the other. Such showers could, due to their high instantaneous intensity, cause damage to the ATLAS subdetectors. In the event of such an accident, BCM is responsible for quickly triggering a beam abort. This is done by comparing the arrival time of signals in the two detector stations. Particles originating from the IP arrive at the stations simultaneously while particles traversing ATLAS from one side to the other require approximately 12.5 ns to travel the 368 cm between them. Therefore, if the arrival time of the signals from the two BCM stations differs by 12.5 ns, the beam may be aborted.

4.7.2 MBTS

Most triggers (see Section 4.8) are designed to select events that contain interesting physics. They apply a set of criteria in order to decide which events to keep for further analysis and which events to ignore. Sometimes, in order to minimize the model dependence of a measurement, it is desirable to select events using criteria that are as inclusive as possible. This is called a *minimum-bias* selection. The Minimum Bias Trigger Scintillators [54] are 32 scintillator tiles (16 on each side of the IP) mounted in front of the end-cap calorimeter cryostats at $z = \pm 3.6$ m. They are very sensitive to activity within a bunch crossing and can therefore be used to trigger in an inclusive way. However, due to their high sensitivity they saturate even at moderate luminosities and are therefore only used during special runs. MBTS covers the range $2.09 < |\eta| < 3.84$.

4.7.3 LUCID

LUCID is a dedicated luminosity monitor. It consists of two modules which are mounted at $z = \pm 17$ m and cover $5.6 < |\eta| < 5.9$. Since LUCID is one of the main topics of this text, it is discussed separately in Section 5.

4.7.4 ZDC

The Zero Degree Calorimeter is made up of one detector station on each side of the interaction point. The stations are located at $z = \pm 140$ m. This is precisely the point where the two counter-rotating proton beams, which are joined into a single beam pipe for collisions in ATLAS, split back into two independent pipes. ZDC is placed after the splitting point, directly in front of the original beam. It covers the range $|\eta| > 8.3$ and is designed to detect neutrons emitted in the very forward direction. Such neutrons are of special interest in lead ion collisions, where they are related to a quantity called centrality [55]. The calorimeter itself consists of quartz rods interleaved with a tungsten absorber. The particles created as neutrons shower in the tungsten emit Cherenkov light in the quartz. This light is subsequently read out by photomultipliers.

4.7.5 ALFA

The optical theorem [56] is a general law in scattering theory that relates the forward scattering amplitude of a process to the total cross-section of that process. At the LHC, the theorem relates the number of protons scattered in the forward direction to the total proton–proton interaction cross-section. ALFA is designed to detect and count the forward protons, which combined with a measurement of the total interaction rate yields the absolute luminosity (see Section 6). In the context of the optical theorem, forward protons means protons scattered at zero angles. It is of course impossible to directly detect such protons as it would require putting a detector in the path of the beam. It is, however, possible to detect the number of protons scattered at extremely small angles and then extrapolate that measurement down to zero.

In order to get as close as possible to the beam, the ALFA detector employs the so-called *roman pot* technique in which the detectors are connected to the inside of the beam pipe by vacuum bellows. When the detector is in use, it is separated from the beam by a distance of only 1 mm. The principle is illustrated schematically in Figure 28. ALFA modules are located at $z = \pm 240$ m on each side of the interaction point and can detect protons scattered at angles of about $3 \mu\text{rad}$. Since this is less than the nominal beam divergence, ALFA is only used during dedicated runs with special beam optics. The sensitive area of the detector consists of scintillating fibers that are read out by photomultipliers.

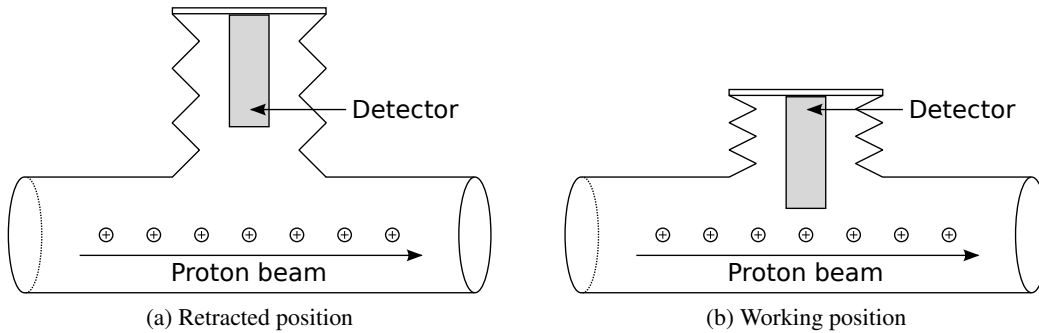


Figure 28: The ALFA detectors are connected to the inside of the beam pipe via movable devices called roman pots. They are retracted during nominal running (a) but can be brought very close to the beam when the detectors are in use (b).

4.8 Trigger and Data Acquisition

The LHC bunch crossing rate is 40 MHz. For each bunch crossing, the ATLAS subdetectors produce about 1.3 MB worth of data that could potentially be used for offline analysis. This equates to 52 TB of

ATLAS data per second. Since it is not technically feasible to save all the data to disk, a trigger is used to select only events that look like they contain interesting physics. The ATLAS trigger is divided into three stages. They are called Level 1, Level 2 and the Event Filter. Each stage applies a set of increasingly stringent selection criteria to the incoming events. Events that fail the selection are rejected, while events that pass the selection are kept for the subsequent stage.

The Level 1 trigger [57] is the first stage. It uses reduced-granularity information from the calorimeters and the muon trigger systems (RPCs and TGCs) in order to select events that either contain particles with high transverse momentum or have other special features such as a large amount of missing energy. For each selected event, the trigger also defines one or more *regions of interest*. These are the η - ϕ coordinates of regions in ATLAS where interesting features were observed. The actual trigger decision is taken by the Central Trigger Processor (CTP). This is a piece of hardware that combines and evaluates the information gathered from the various subdetectors. The time between a bunch crossing and the corresponding trigger decision is designed to be no greater than 2.5 μ s. Note that this is not the processing time required per bunch crossing but the time required for the information to travel via cables from the subdetectors to the CTP. Unprocessed events are kept in buffers until a trigger decision is made. After the Level 1 trigger, the initial 40 MHz event rate is reduced to 75 kHz.

The second stage is the Level 2 trigger. It applies refined selection criteria using information from the regions of interest defined at Level 1. Whereas the Level 1 trigger decision is based on reduced-granularity data, Level 2 has access to the full precision of ATLAS. The selection reduces the event rate to approximately 3.5 kHz, which consequently becomes the input rate for the Event Filter. This is the last stage of the trigger. It uses full-granularity data from the entire detector (rather than just the regions of interest) and applies a final selection that reduces the event rate to 200 Hz. The time between a bunch crossing and the corresponding Event Filter trigger decision is typically around four seconds.

The Data Acquisition system (DAQ) is the framework in which the trigger operates. The DAQ receives and buffers data from the various ATLAS subsystems, and upon a Level 1 trigger accept it transmits the data from the regions of interest to the Level 2 trigger. If an event passes also the Level 2 selection, the DAQ builds the full event in a standard ATLAS format and sends it to the Event Filter. Accepted events are then saved to disk for offline analysis. The trigger and the DAQ are collectively referred to as the Trigger and Data Acquisition system (TDAQ).

5 The LUCID detector

5.1 Overview

The purpose of LUCID (Luminosity Measurement Using a Cherenkov Integrating Detector) [58, 59] is to monitor the luminosity delivered to ATLAS by the LHC. The ATLAS luminosity, as discussed in Section 6, is a quantity that is proportional to the number of particles created due to pp collisions in the ATLAS interaction point. This number, in turn, is directly proportional to the number of particles detected by LUCID (or any other detector for that matter). By monitoring the particle flux, LUCID is therefore able to measure the luminosity.

Figure 29 shows the position of LUCID within ATLAS. The detector consists of two modules, one on each side of the interaction point, located at $z = \pm 17$ m. The modules are mounted approximately 10 cm from the beam and cover the range $5.6 < |\eta| < 5.9$. Each module is made up of 20 mechanically polished aluminium tubes. The tubes can be divided into two categories. Sixteen of the tubes have photomultipliers mounted directly at the end. The PMT windows are made of radiation-hard quartz. The remaining four tubes are connected to PMTs via bundles of optical quartz fibers, such that the signal is guided away from the detector before being read out. The remote PMTs are located in a low-radiation area on top of a thick piece of iron shielding called the *monobloc*.

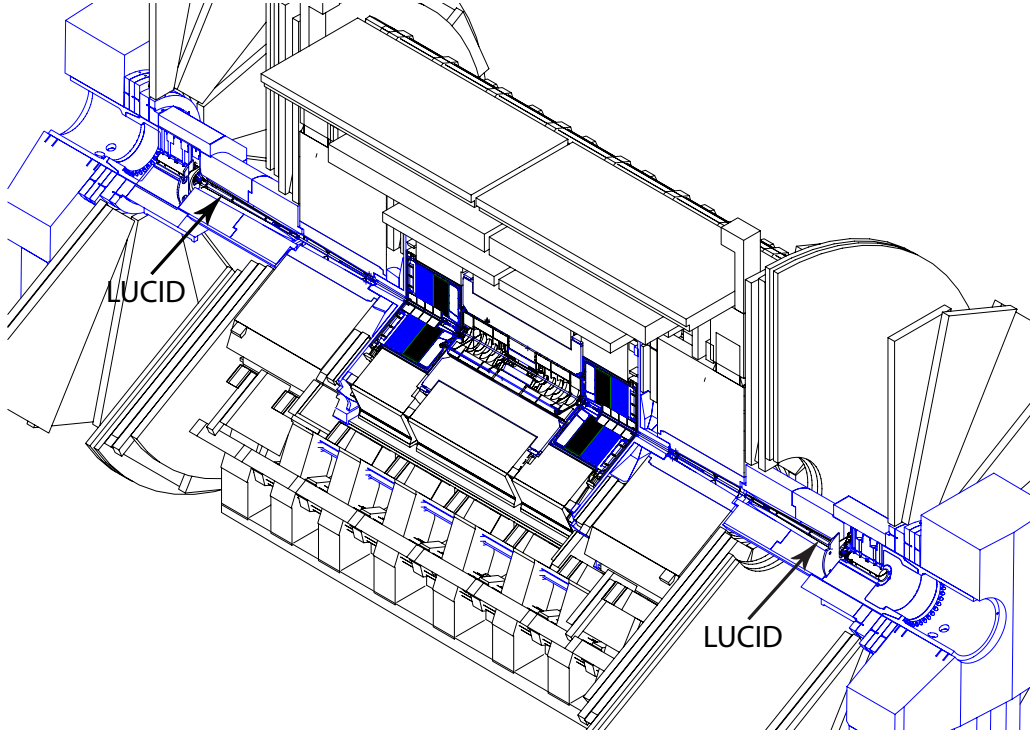


Figure 29: The position of LUCID within ATLAS. Each detector module is located 17 m away from the interaction point.

The tubes all reside within an air-tight gas vessel that can be either empty or filled. When the vessel is filled, the gas contained within acts as a Cherenkov medium. A relativistic charged particle traveling through a tube continuously emits Cherenkov light. This light is reflected on the inside of the tube wall until it reaches the PMT. In addition, the particle produces a burst of Cherenkov light as it passes through the quartz window (or quartz fiber bundle) at the end of the tube. This means that LUCID can be operated even without gas in the vessel, in which case the quartz is the sole source of Cherenkov light.

5.2 Detector design

Inelastic pp collisions at the LHC result in a large number of particles being produced at small angles to the beam. The LUCID detector is placed in the forward region and designed to have a good acceptance for such events. This allows it to quickly perform luminosity measurements with a negligible statistical uncertainty. The main design requirements for a luminosity monitor in the forward region is radiation hardness and speed. Both requirements are fulfilled by the choice of a Cherenkov counter. Since the detector body is made of aluminium, it is radiation hard (unlike a scintillation counter). Since Cherenkov light is emitted promptly, a fast response time is ensured (as opposed to gaseous proportional counters, which are associated with long ion drift times). LUCID has a time resolution that is significantly better than the bunch spacing of the LHC. This allows it to perform luminosity measurements bunch by bunch.

Figure 30 shows a schematic diagram of the two types of tubes used by LUCID. The standard tube is 149.5 cm long, 1 mm thick and has an inner diameter of 14 mm. A photomultiplier is mounted at the far end of the tube via a cylinder. The PMT (model Hamamatsu R762) has a window made of quartz, which is a material chosen for its radiation hardness. The quartz window is 1.2 mm thick and has a diameter of 15 mm. This is greater than the inner tube diameter, ensuring that there is no gap where Cherenkov light could be lost. Each of the two LUCID modules contains 16 standard tubes for a total of 32 tubes. All of the luminosity measurements that LUCID delivered to ATLAS during run I were based on this readout method.

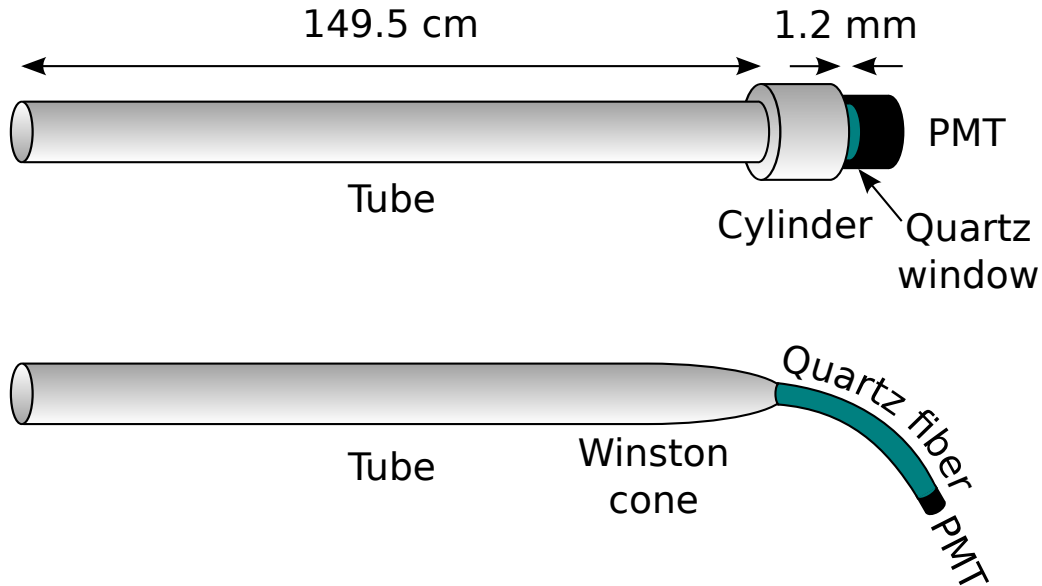


Figure 30: Schematic diagram showing the two types of tubes used by LUCID. Each detector module contains 16 tubes that have PMTs mounted directly at the far end and 4 tubes that are connected to PMTs via bundles of optical quartz fibers.

The second type of tube is equipped with a Winston cone at the far end. This cone focuses the incident Cherenkov light into a bundle of optical quartz fibers. Each fiber in the bundle has a quartz core, a silicone buffer and a Tefzel jacket with diameters of 0.8 mm, 0.95 mm and 1.3 mm respectively. This results in a numerical aperture of 0.4. The overall bundle consists of 37 individual fibers and has a diameter of 6 mm. Light captured by the fiber bundle is guided away from the detector to a remote PMT (again of model Hamamatsu R760) located on top of the monobloc. The fiber-based readout channels were not used for any official luminosity measurement during run I. Instead, they served as a test bench for the planned run II upgrade (see Section 5.5). LUCID is located close to the beam pipe where the

particle density is high. Radiation damage and the high current drawn by the PMT anodes cause the standard PMTs to age and eventually become unusable. The PMTs served by fibers, on the other hand, are protected from radiation by the monobloc. They also see less signal compared to the standard PMTs due to the small diameter of the fiber bundles. They therefore have a longer lifetime, which is expected to become an important factor owing to the increased particle densities that will be encountered during run II. Each of the two LUCID modules contains 4 fiber bundles for a total of 8 bundles in the detector overall. All tubes, both standard and fiber, are slightly tilted so that they point towards the interaction point. This helps suppress background when the detector is filled with gas. Only particles originating from the IP travel the full length of the tubes. Such particles therefore produce more Cherenkov light than background would.

The LUCID gas vessels, which define the overall size of the detector, are mounted on the inside of a cone-shaped support structure called the *VJ cone*. The gas vessels begin at $z = \pm 1672$ cm and have a length of 153.2 cm. Throughout the first half of run I, they were filled with C_4F_{10} gas ($n = 1.00137$) at a pressure of 1.1 bar. C_4F_{10} is a non-flammable, commercially available gas with good transparency to photons in the ultraviolet energy range. Running with gas increases the detector efficiency, which is advantageous when the luminosity is low. In high luminosity runs, however, having a high efficiency can lead to saturation of the luminosity algorithms (see Section 5.3.2). Gas also contributes to migration (see Section 7.4), which is a nonlinear effect that results in an overestimation of the luminosity. The gas vessels were therefore emptied on July 30:th 2011. With the exception of a few special runs, LUCID was operated without gas throughout the remainder of run I.

Figure 31 illustrates the overall layout of the LUCID detector. The figure shows a cross-sectional view of the area around the beam taken at the back end of one of the LUCID gas vessels. In the center is the beam pipe. It is covered by an isolating layer of aerogel. Surrounding the beam pipe are the inner and outer walls of the gas vessel. The tubes are contained within the gas vessel. They are arranged into two concentric rings centered 105 mm and 125 mm from the beam, with the fibers being part of the inner ring. The entire detector is mounted on the inside of the VJ cone. Holes in the cone allow access to LUCID. Table 12 shows the dimensions of all components present in the figure. Finally, a photograph of the LUCID modules prior to their installation in ATLAS is shown in Figure 32.

Table 12: The table shows the length, diameter and thickness of the components present in Figure 31. For hollow objects, thickness is measured in the radial direction and the diameter refers to the *inner* diameter.

Component	Length	Diameter	Thickness
Beam pipe	–	120 mm	1.5 mm
Aerogel isolation	–	121.5 mm	4 mm
Gas vessel inner	153.2 cm	170 mm	3.2 mm
Gas vessel outer	153.2 cm	294 mm	3.2 mm
VJ cone	–	374 mm	5 mm
Tube	149.5 cm	14 mm	1 mm
Cylinder	–	25 mm	–
Quartz window	1.2 mm	15 mm	–
PMT	–	19 mm	–
Fiber bundle	5 m	6 mm	–

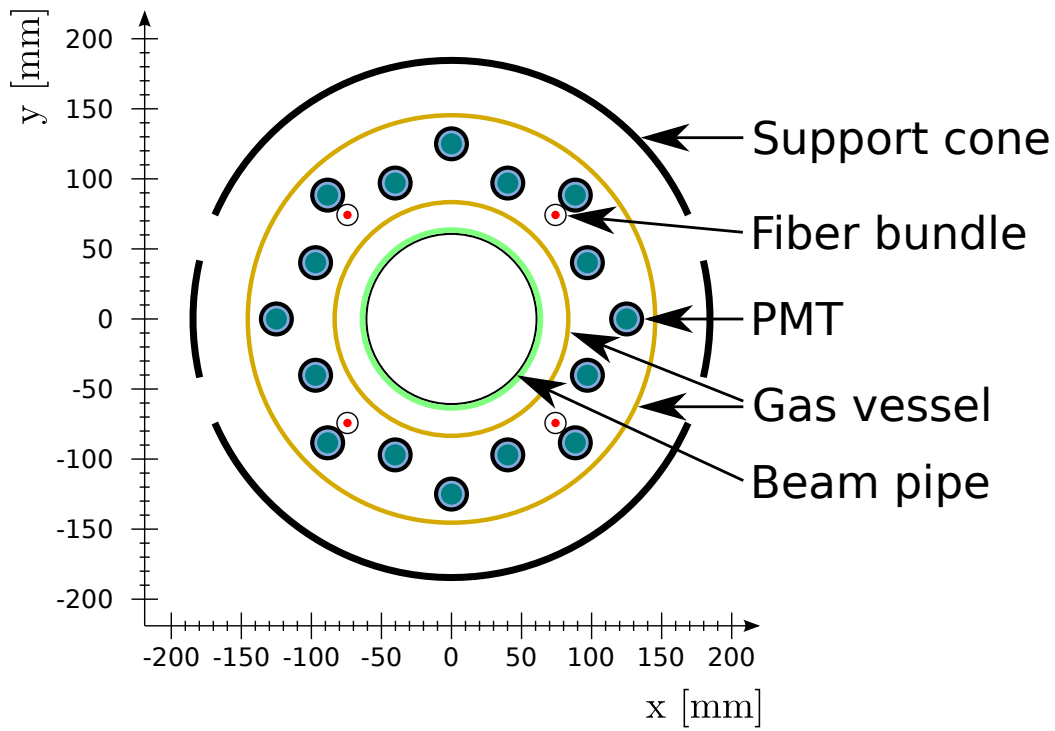


Figure 31: A cross-sectional view of the area around the beam pipe taken at the back of one of the LUCID gas vessels.

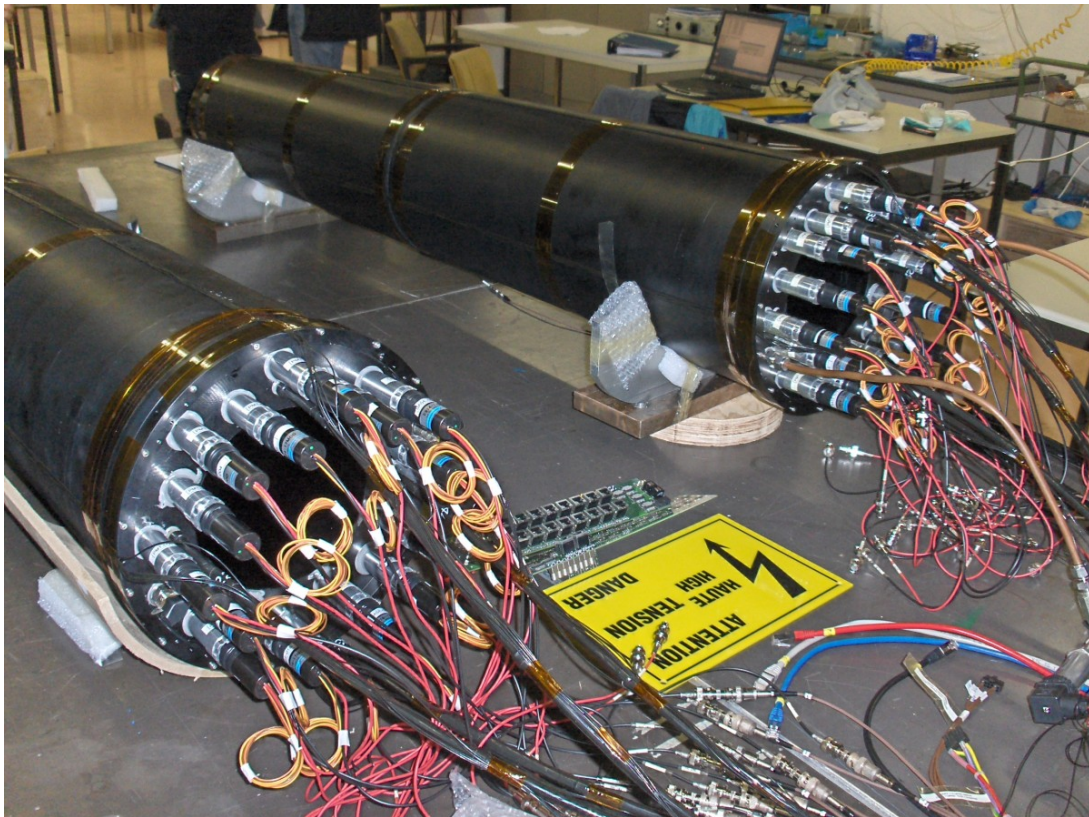


Figure 32: A photograph of the two LUCID modules prior to their installation in ATLAS.

5.3 Readout electronics

Figure 33 shows a schematic representation of the LUCID readout system. This section will describe how the system works. It should be noted that the aim is not to provide a complete description of the electronics. Several components have been omitted from Figure 33 for the sake of simplicity. These components are, however, not necessary in order to understand how LUCID operates as a luminosity monitor. It should also be noted that some PMTs have broken over time. This section will describe the system as it was originally designed and introduce any modifications as they become relevant to the discussion.

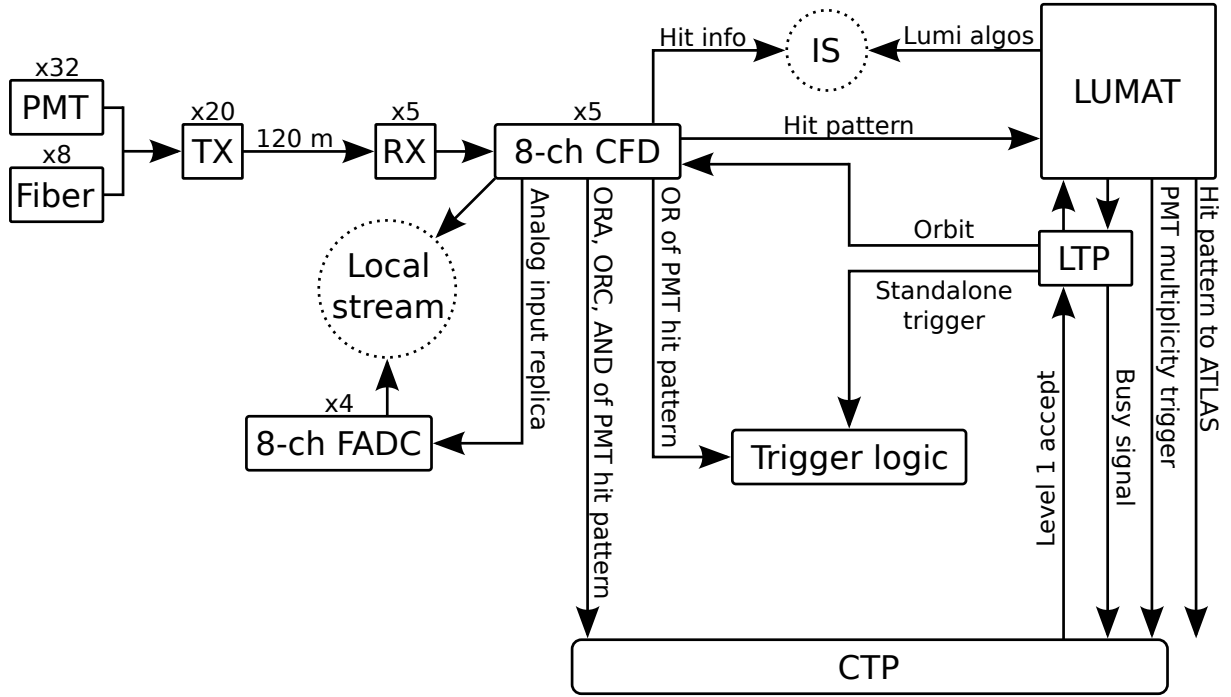


Figure 33: The LUCID readout system.

Each tube in LUCID represents one readout channel. Therefore, there is a total of 40 inputs provided by the 32 standard tubes and the 8 fiber bundles. While the PMTs themselves are located inside the ATLAS detector, most of the readout components reside in a nearby service area called USA15. In order to reduce the impact of electromagnetic interference, the PMT signals are transmitted to the readout electronics differentially. A set of transmitter (TX) cards convert the PMT inputs to a differential form and send them over 120 m of twisted-pair cable to USA15. Once there, a set of receiver (RX) cards convert the signals back to a single-ended form. The RX cards also provide a zero-pole compensation in order to correct for cable losses. A constant fraction discriminator (CFD) then receives the signals and compares the pulse heights to an internal set of threshold values. The thresholds are adjusted such that a signal is considered to be above threshold if at least 15 photoelectrons were emitted from the photocathode of the corresponding PMT.

Upon receiving the signals, the CFD does many things. First, it provides an amplified replica of the analog signals to a flash analog to digital converter (FADC). The FADC records the pulse shape. If the signal fulfills some trigger condition, data from the CFD and the FADC is written to disk. The trigger condition can be defined in several ways, all of which are independent of the global Level 1 trigger decision made by ATLAS. The data stream is therefore referred to as the *local stream*. It is discussed further in Section 5.3.1. The second task that the CFD performs is to provide the hit pattern

to the LUMAT board (see Section 5.3.2). For simplicity, the hit pattern can be thought of as a number represented by as many bits as there are readout channels. A bit is set if the corresponding readout channel was above threshold. LUMAT uses the hit information as input for the luminosity calculation. Third, the CFD integrates the number of hits (readout channels above threshold) over time and publishes the information to a TDAQ infrastructure called the *Information Service* (IS) where it becomes visible to other systems in ATLAS. The number of full revolutions of the LHC beam (also called the number of *orbits*) over which the hits were integrated is published as well. Fourth, the CFD sends information to other components about the part of the hit pattern that represents the standard tubes. If there was a hit anywhere, a signal is sent from a front panel to the LUCID trigger logic, which decides whether or not to write information about the event to the local stream. Three signals are sent from a backplane to the CTP for use in the Level 1 trigger decision. If there was at least one hit on side A, the ORA signal is sent. Similarly, the ORC signal is sent if there was at least one hit on side C. The AND signal is sent if there was at least one hit on side A and at least one hit on side C.

LUMAT receives the hit pattern from the CFD and publishes the derived luminosity information to IS. It also looks at the part of the hit pattern that represents the standard tubes and calculates how many of them were above threshold. If this number exceeds a predefined value, a trigger signal (called the multiplicity trigger) is sent to the CTP for use in the Level 1 trigger decision. Any other communication between the CTP and the LUMAT board is done via a Local Trigger Processor (LTP). The LTP serves several functions. It makes sure that the LUMAT clock is synchronized with the global ATLAS time and generates an orbit signal that can be used to calculate how many bunch crossings occurred in a given time frame. The very same orbit signal is also sent to the CFD, which requires the information for its IS publications. Whenever a Level 1 trigger accept decision is made, the LTP is responsible for communicating that decision to LUMAT. This prompts LUMAT to send the hit pattern to the global ATLAS data flow. The LTP also receives the LUCID busy signal and forwards it to the CTP. The CTP looks at similar signals from many subdetectors in ATLAS and may decide to halt data taking if certain detectors are busy. Since LUCID is fast, it rarely causes ATLAS to be put in a busy state. Finally, the LTP can be used to define custom trigger conditions for the local stream. For example, it is possible to trigger on a specific BCID regardless of whether or not it produced a hit in LUCID.

5.3.1 The local stream

For a given event, the local data stream contains the pulse shapes recorded by the FADC, the hit pattern from the CFD, the luminosity block number (see Section 6) and some information relevant to detector calibration. Data from the local stream is mainly analyzed offline in order to study the behavior of the LUCID detector and to calibrate the PMTs, but it is also used for online monitoring of the data quality. As indicated by Figure 33, there is a total of 32 FADC channels available. This means that only a subset of the LUCID readout channels can be monitored at any given time. For the majority of run I, the local stream monitored all the standard PMTs except for two that were broken. The broken PMTs were replaced by orbit signals. Towards the end of run I, another eight standard PMTs were removed from the local stream and replaced by fibers. This allowed the fiber signals to be studied in detail prior to the LUCID upgrade. Finally, three additional standard PMTs and one fiber (one readout channel per FADC module) were removed from the local stream and replaced by synchronized clock signals. Since all bunch positions are known with respect to the clock, this allowed the BCID boundaries in a given event to be identified unambiguously. Knowing the exact position of the boundaries is crucial when studying individual bunches in runs where the bunch spacing is small.

Figure 34 shows a typical signal as recorded by the FADC. The FADC samples the signal amplitude every 4 ns and stores the information in an internal buffer until read out. The buffer length is a configurable parameter. In Figure 34 it was set to 84 samples, which corresponds to 336 ns. Only one peak is visible in the spectrum, indicating that it was obtained from a run where the bunch spacing was large.

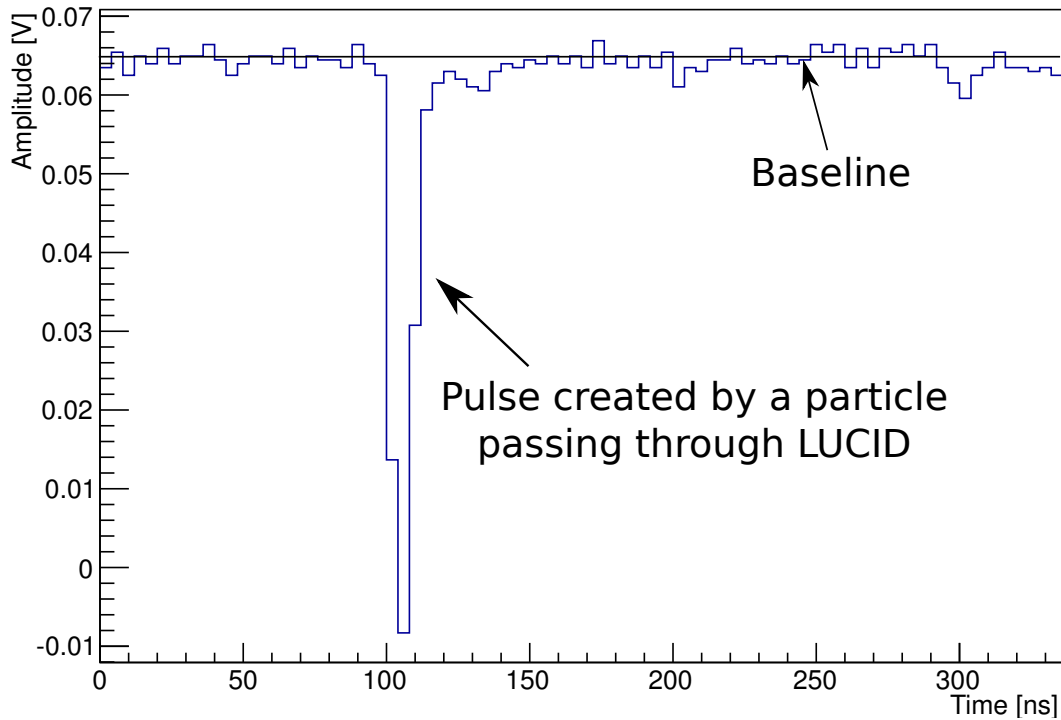


Figure 34: A flash ADC readout from one of the standard PMTs. The approximate position of the signal baseline is indicated.

Two different methods were used to trigger the local stream during physics data taking in run I. The first method is based on the OR of the standard PMT hit pattern. This signal is provided by the CFD. A trigger is issued whenever at least one of the standard PMTs is above threshold. Since reading out the local stream information involves the slow operation of writing data to disk, it is not possible to trigger at the bunch crossing rate of the LHC. Instead, the trigger is prescaled so that only a subset of the events are actually saved. The second trigger method is based on the LTP. A specific BCID is chosen to be the triggering BCID. Whenever the triggering BCID collides in ATLAS, a random number is generated to decide whether or not the local stream should be read out. This results in a dataset that is not biased by any trigger conditions imposed by LUCID.

5.3.2 LUMAT

The LUMAT (Luminosity And Trigger) board is the heart of the LUCID luminosity measurement. It provides the CTP with a trigger based on the hit multiplicity in LUCID, it sends the hit pattern to the ATLAS data flow whenever there is a Level 1 trigger accept, and it uses the hit pattern to derive information about the luminosity. This information is periodically published to IS where it is picked up by the Online Luminosity Calculator (see Section 6.5). It is also saved to disk for offline processing. LUMAT receives one hit pattern per bunch crossing. It scans the part of the hit pattern that corresponds to the standard tubes and checks if any of the following four conditions are fulfilled.

- There was at least one hit anywhere. This is the OR condition.
- There was at least one hit on side A. This is the ORA condition.
- There was at least one hit on side C. This is the ORC condition.
- There was at least one hit on side A and at least one hit on side C. This is the AND condition.

Bunch crossings that fulfill at least one condition are called *events*. Depending on which conditions were fulfilled, LUMAT increments a collection of internal scalers. There are six sets of scalers representing six different counting algorithms. Each set contains one scaler per BCID for a total of 3564 scalers per set. Within a given set, each scaler is associated with exactly one BCID. The number stored by the scaler contains information about the luminosity of that BCID. The six algorithms work as follows.

- If the OR condition is met, the associated scaler is incremented by one.
- If the ORA condition is met, the associated scaler is incremented by one.
- If the ORC condition is met, the associated scaler is incremented by one.
- If the AND condition is met, the associated scaler is incremented by one.
- If the OR condition is met, the associated scaler is incremented by the total number of hits.
- If the AND condition is met, the associated scaler is incremented by the total number of hits.

The algorithms are referred to as, respectively, EventOR, EventORA, EventORC, EventAND, HitOR and HitAND. The scalers are periodically reset, and the numbers are published to IS together with the total number of orbits over which they were accumulated. Section 6 explains how the numbers stored by the scalers are related to the luminosity of the LHC. In addition to the six algorithms associated with the standard tubes, there are another six associated with the fiber tubes. They work in the same way, except the input is obtained from the part of the hit pattern that represents the fiber readout channels.

5.4 Calibration

The calibration procedure has two goals. The first goal is to find the constant of proportionality between the signal charge recorded by the FADC and the number of photoelectrons (PEs) emitted from the photocathode of a PMT. The second goal is to equalize the gain between the photomultipliers so that they all produce the same output for a given input.

The charge carried by a particular signal is measured by integrating over the pulse shape as recorded by the FADC. Figure 35 illustrates the principle. The integral is obtained by summing up the amplitude of each FADC sample in a given integration range. This range can be chosen arbitrarily as long as it is wide enough to cover the signal when present and as long as it is kept constant throughout the calibration process.

The first step of the calibration procedure is to measure the charge distribution in the absence of light. Figure 36 (a) shows an example of such a distribution. Since no light is incident on the detector, the charge measurements are simply sampling the pedestal and the random dark current processes that occur in the PMT. This measurement is used to obtain the position and width of the pedestal distribution. The second step is to measure the charge distribution in the presence of a well-defined signal. A light-emitting diode (LED) is used to pulse the PMTs with monochromatic light of tunable intensity. A total of four LEDs with different wavelengths are available, but only one is used at any given time. The LEDs are located on top of the monobloc, and the light is guided to the tubes via optical quartz fibers. During calibration runs, the trigger frequency of the local stream is set to match the pulse frequency of the LEDs. When the LED light hits the photocathode of a PMT, a number of photoelectrons N are emitted. The number of photons that arrive at the photocathode in a given pulse is a Poisson-distributed variable. Each photon then excites a photoelectron from the photocathode with a probability given by the quantum efficiency of the PMT. Since this is a simple binary process (either it happens or it doesn't happen), N is a Poisson-distributed variable as well. The LED intensity is tuned such that approximately one photoelectron is emitted on average, and a charge distribution is recorded. An example of such a charge distribution is shown in Figure 36 (b).

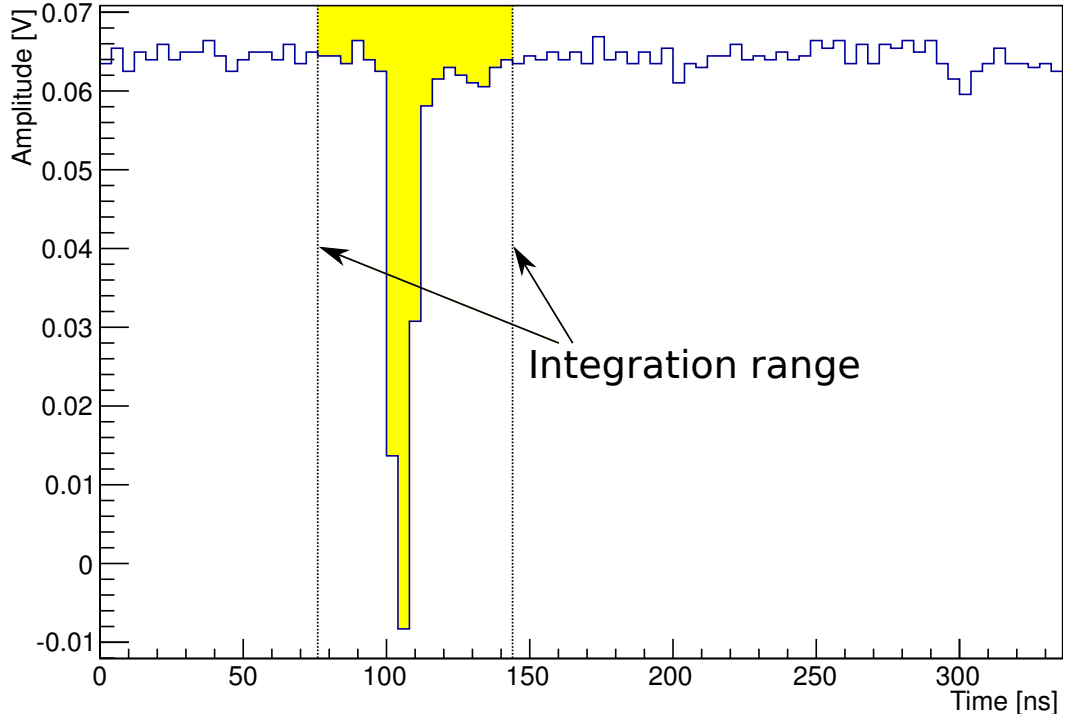


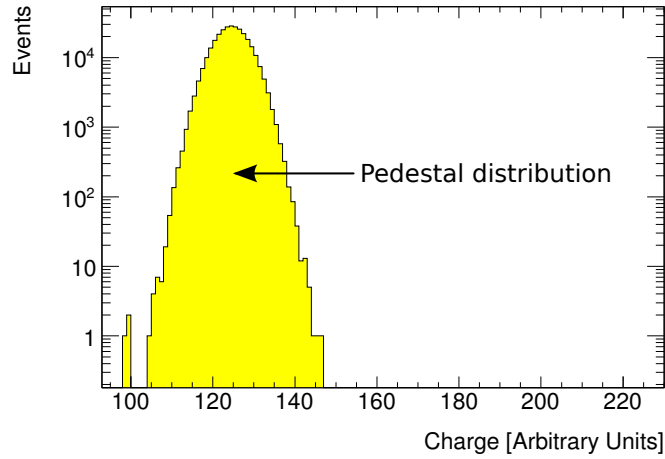
Figure 35: The figure shows how the charge carried by a signal can be obtained by integrating over the pulse shape.

Each photoelectron is amplified by the PMT with an average gain that depends on the PMT high voltage (HV). The amount of charge that results from the amplification of a photoelectron is described by a Gaussian distribution. The charge distribution in Figure 36 (b) can therefore be modelled as a sum of Gaussians. This is illustrated in Figure 36 (c). A given Gaussian corresponds to N photoelectrons being emitted from the photocathode of the PMT. The peak value of that Gaussian is proportional to the Poisson probability to get N PEs given a mean of 1 PE. After subtracting the charge of the pedestal, the mean of the Gaussian is simply N times the average gain of the PMT. The last step of the calibration procedure is to fit the charge distribution using the model just described. The constant of proportionality between charge and photoelectrons is taken to be the difference between the mean of the pedestal distribution and the Gaussian corresponding to 1 PE. The high voltage supplied to each PMT is tuned until this difference is the same for every tube. Finally, Figure 36 (c) shows a small distribution labeled “Non-Poisson pedestal”. It represents the fraction of pedestal events that could not be accounted for using Poisson statistics.

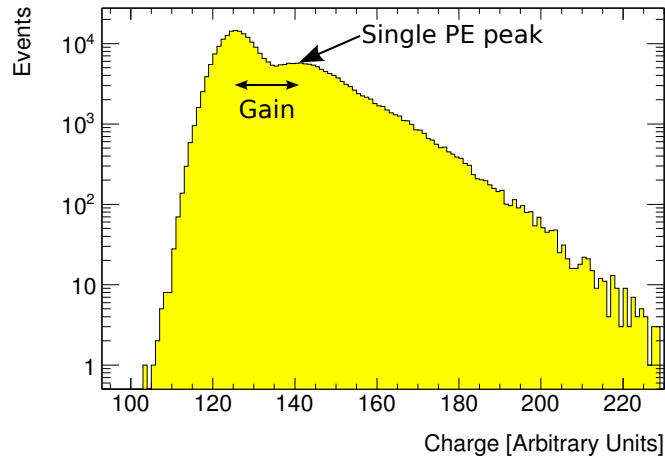
5.5 Upgrade

In run II, the LHC will operate at $\sqrt{s} = 13$ TeV with a nominal bunch spacing of 25 ns, reaching higher levels of luminosity than were ever achieved during run I. The design of LUCID must change in order to cope with these new running conditions, and upgrade work is ongoing. Figure 37 shows a cross-sectional view of the new detector layout. Figure 38 shows a side-view of the detector and the surrounding support structure.

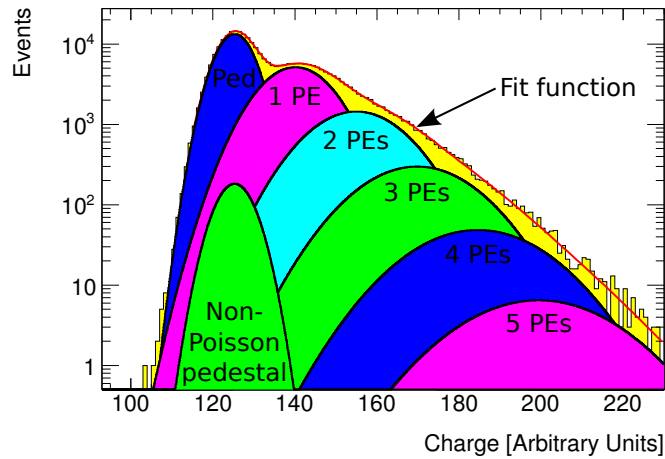
After the upgrade, LUCID will no longer use gas as a Cherenkov medium. The tubes and the gas vessel have therefore been removed and replaced with a carbon fiber support cylinder on which the photomultipliers and fibers are mounted. The PMT windows and the fiber bundles, which are still made of quartz, are the sole sources of Cherenkov light. The window diameter of the new PMTs (model



(a) Pedestal distribution



(b) Single photoelectron distribution



(c) Fitted distribution

Figure 36: The first step of the calibration procedure is to record a charge distribution in the absence of light (a). The second step is to record a charge distribution with the LED tuned to produce one photoelectron on average (b). The last step is to fit a function to the data (c).

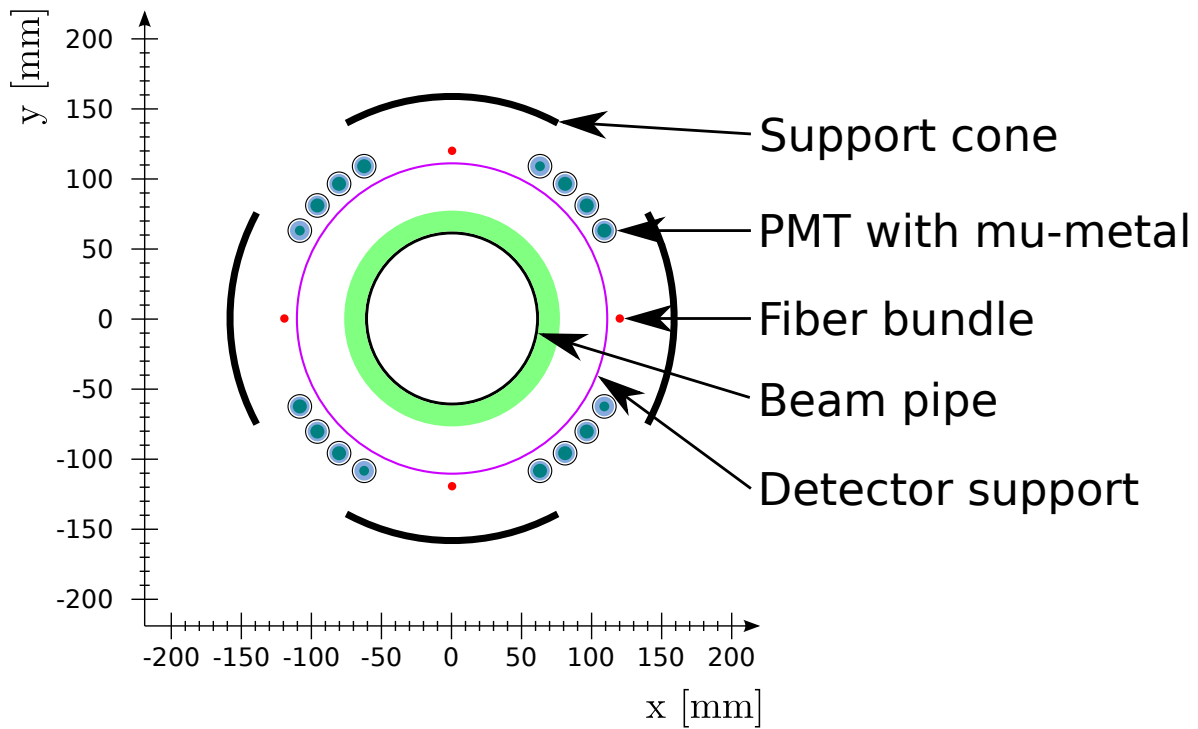


Figure 37: A cross-sectional view of the upgraded LUCID detector.

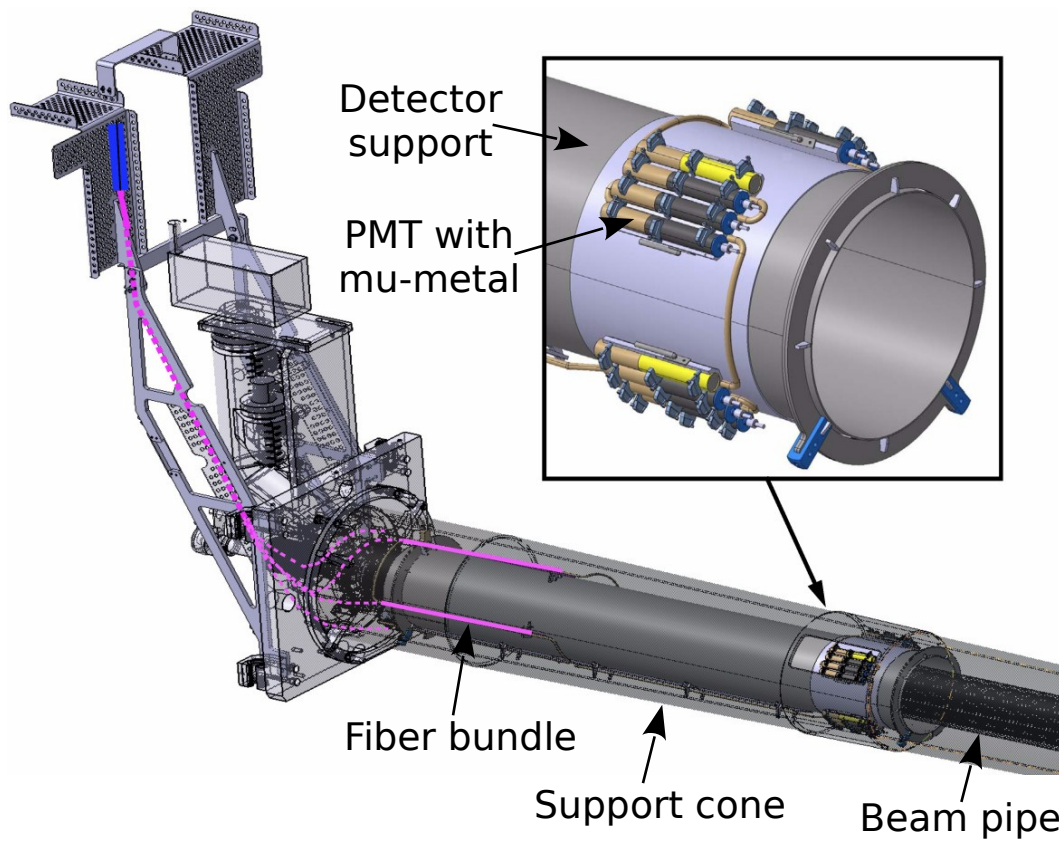


Figure 38: A side-view of the upgraded LUCID detector and the surrounding support structure.

Hamamatsu R760) has been decreased to 10 mm in order to reduce the detector acceptance. This is done to avoid saturation problems, as the particle density in run II will be much higher than in run I. One fourth of the PMTs are further modified such that their active window diameter is reduced to 7 mm. Each PMT is enclosed in a mu-metal cylinder that provides shielding against external magnetic fields. The new VJ cone has larger windows than the old one, and the PMTs are grouped beneath the windows for easy access. The PMTs are located at $z = \pm 17$ m, while the fiber bundles start a bit further down the beam pipe at $z = \pm 18$ m. Table 13 lists some of the dimensions used in the new design. Finally, the upgraded detector will use a new method for monitoring luminosity as discussed in Section 8. New readout electronics have been installed in order to implement the method.

Table 13: The table shows the length, diameter and thickness of the components present in Figure 37. For hollow objects, thickness is measured in the radial direction and the diameter refers to the *inner* diameter.

Component	Length	Diameter	Thickness
Beam pipe	–	120 mm	2.2 mm
Aerogel isolation	–	124.4 mm	14.8 mm
Detector support	155 cm	220 mm	1.5 mm
VJ cone	–	312 mm	5 mm
Quartz window	1.2 mm	10 mm	–
Modified window	1.2 mm	7 mm	–
PMT	–	14 mm	–
Mu-metal	–	16 mm	0.8 mm
Fiber bundle	2 m	6 mm	–

6 Luminosity

6.1 Overview

Whenever two protons collide at the LHC, they may interact with each other in a myriad of different ways. The probability for a particular process to occur is given by its cross-section σ_{proc} . Intuitively, the rate R_{proc} at which a certain process occurs in ATLAS should be directly proportional to the rate at which protons collide in the interaction point. This observation invites the definition of a new quantity $\mathcal{L}$ called *luminosity*.

$$\mathcal{L}\sigma_{\text{proc}} = R_{\text{proc}} \quad (112)$$

When two bunches of protons collide, each and every proton in the first bunch has the opportunity to interact with each and every proton in the second bunch. The luminosity therefore scales as the *product* of the number of protons in the two bunches. In this sense, the luminosity delivered by the LHC is simply the number of potential pp interactions per second. If the beams are offset so that the bunches do not collide head-on, $\mathcal{L}$ is reduced by a geometrical factor that describes the bunch overlap. Luminosity carries units of $\text{cm}^{-2}\text{s}^{-1}$ and cross-sections carry units of cm^2 . Since cross-sections in particle physics are typically very low, areas are conveniently measured in *barns* where 1 b equals 10^{-24} cm^2 .

Integrating Eq. (112) over the course of a run yields

$$\sigma_{\text{proc}} \int_{t_{\text{start}}}^{t_{\text{stop}}} \mathcal{L} dt = N_{\text{proc}} \quad (113)$$

where N_{proc} is the number of events of a given process that took place during the run. The amount of data recorded in a run, or over the course of multiple runs, is thus measured in terms of *integrated* luminosity. During physics runs, ATLAS identifies which processes take place in each bunch crossing and stores the information to disk. This yields an experimental value of N_{proc} . At the same time, dedicated luminosity monitors measure $\mathcal{L}$. If a previously unknown physics process is observed by ATLAS, the integrated luminosity is used together with the observed number of events in order to calculate the cross-section of the process. The uncertainty of the luminosity measurement is then reflected in the uncertainty of the experimentally measured σ_{proc} . Minimizing this uncertainty is crucial when comparing the experimental value of the cross-section to theoretical predictions. Another situation in which a precise luminosity measurement is important is when modelling the ATLAS experiment based on Monte Carlo (MC) simulations. Physics analyses often do this in order to estimate N_{proc} for various background and signal processes (whose cross-sections are calculated from theory). In order to compare the simulated N_{proc} to the number of events observed in real ATLAS data, the Monte Carlo prediction must be normalized to the integrated luminosity of the ATLAS dataset. Any uncertainty in the luminosity measurement directly translates into an uncertainty of the Monte Carlo prediction.

All ATLAS runs are divided into short segments of time called Luminosity Blocks (LBs). The length of a Luminosity Block is defined by the CTP. This length is normally set to one minute, but it may vary for special runs. In case an ATLAS subsystem happens to experience technical problems such that the data quality is not optimal, the Luminosity Block during which the problem occurred is discarded and not used for any physics analyses. The full ATLAS dataset is then defined as the sum of all good Luminosity Blocks. Figure 39 shows how the size of the $\sqrt{s} = 8 \text{ TeV}$ dataset that was recorded by ATLAS during 2012 increased over time. The preliminary uncertainty of the luminosity measurement is estimated to be 2.8% following the methodology detailed in Ref. [1].

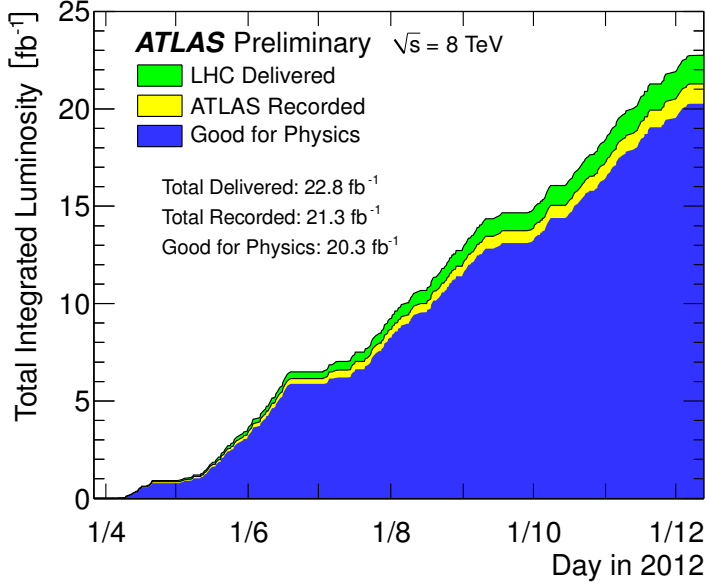


Figure 39: The figure shows the total integrated luminosity delivered by the LHC during $\sqrt{s} = 8$ TeV runs in 2012 [60]. The amount of luminosity recorded by ATLAS reflects the DAQ efficiency. A fraction of the recorded luminosity is then discarded due to problems with data quality.

6.2 Luminosity monitoring

In practice, the luminosity is measured individually for every bunch pair in the LHC. The total luminosity is then obtained by summing over the bunches according to

$$\mathcal{L} = \sum_{n_b} \mathcal{L}_b \quad (114)$$

where n_b is the number of colliding bunch pairs and $\mathcal{L}_b$ is the luminosity of each bunch. The number of inelastic proton–proton interactions n that take place when two bunches collide is a random variable. It is a basic assumption of any luminosity measurement that n is Poisson-distributed with a mean of μ , so that μ denotes the average number of inelastic pp interactions per bunch crossing. In terms of μ , the luminosity of a given bunch pair becomes

$$\mathcal{L}_b = \frac{R_{\text{inel}}^b}{\sigma_{\text{inel}}} = \frac{\mu f_r}{\sigma_{\text{inel}}} \quad (115)$$

where R_{inel}^b is the rate of inelastic scattering events from the bunch pair in question, σ_{inel} is the inelastic scattering cross-section and f_r is the revolution frequency of the LHC. Since luminosity monitors have a limited acceptance and efficiency, they are only able to observe a subset of the inelastic scattering events. The product of the detector acceptance and the efficiency of the measurement method is denoted ϵ . It is typically referred to simply as *the efficiency* for brevity. The average number of inelastic pp interactions per bunch crossing that are visible to a particular luminosity monitor is given by $\mu_{\text{vis}} = \epsilon\mu$. Rewriting Eq. (115) in terms of μ_{vis} yields

$$\mathcal{L}_b = \frac{\epsilon\mu f_r}{\epsilon\sigma_{\text{inel}}} = \frac{\mu_{\text{vis}} f_r}{\sigma_{\text{vis}}} \quad (116)$$

where σ_{vis} denotes the visible cross-section.

A distinction can be made between luminosity monitoring and direct luminosity measurements. A luminosity monitor is a detector that measures μ_{vis} . This is an experimentally observable quantity that is

proportional to the luminosity. Direct measurements of $\mathcal{L}$ are called *absolute* luminosity measurements. They are more complicated and can only be carried out under special beam conditions. A luminosity monitor is calibrated by performing an absolute luminosity measurement and a measurement of μ_{vis} at the same time. This is equivalent to measuring σ_{vis} . Calibration is the topic of Section 6.3, while Section 6.4 discusses measurements of μ_{vis} .

6.3 Calibration

Before a luminosity monitor can be used to measure $\mathcal{L}$, it must be calibrated. This is done by measuring μ_{vis} with the luminosity monitor while simultaneously measuring the absolute luminosity with an independent method. ATLAS uses the *van der Meer* (vdM) method [61, 62] to determine the absolute luminosity. This section contains a brief introduction to vdM scans.

The luminosity $\mathcal{L}_b$ of a single bunch pair colliding at a zero crossing-angle can be expressed in terms of beam parameters as

$$\mathcal{L}_b = f_r n_1 n_2 \int \hat{\rho}_1(x, y) \hat{\rho}_2(x, y) dx dy \quad (117)$$

where f_r is the revolution frequency of the LHC, n_1 is the number of protons in the first bunch and n_2 is the number of protons in the second bunch. The integral describes the bunch overlap at the interaction point, with $\hat{\rho}_1$ and $\hat{\rho}_2$ denoting the normalized particle density of each bunch in the x - y plane. A standard assumption is that the two-dimensional particle density of a bunch can be factorized into two one-dimensional components that independently describe the density in the horizontal and vertical planes.

$$\hat{\rho}(x, y) = \rho_x(x) \rho_y(y) \quad (118)$$

The integral in Eq. (117) can then be separated into two independent parts, where one part depends only on x and the other part depends only on y . Defining the beam-overlap integrals

$$\begin{aligned} \Omega_x(\rho_{x1}, \rho_{x2}) &= \int \rho_{x1}(x) \rho_{x2}(x) dx \\ \Omega_y(\rho_{y1}, \rho_{y2}) &= \int \rho_{y1}(y) \rho_{y2}(y) dy \end{aligned} \quad (119)$$

yields

$$\mathcal{L}_b = f_r n_1 n_2 \Omega_x(\rho_{x1}, \rho_{x2}) \Omega_y(\rho_{y1}, \rho_{y2}). \quad (120)$$

In the vdM method, the beam-overlap integrals are measured by scanning one beam across the other. For example, to measure Ω_x , the two beams are positioned such that they are centered in y but offset in x . Since the beams are offset, the bunches do not collide (the beam-overlap integral is zero) and no luminosity is recorded. One beam is then scanned along x while the other is kept stationary, and the luminosity is monitored as a function of the beam separation δ . This yields a scan curve such as the one shown in Figure 40. The beam-overlap integral can be expressed as

$$\Omega_x(\rho_{x1}, \rho_{x2}) = \frac{R_x(0)}{\int R_x(\delta) d\delta} \quad (121)$$

where R represents any observable that is proportional to the luminosity (such as μ_{vis}). The quantity of interest is the *width* of the scan curve. The horizontal beam width Σ_x is defined according to:

$$\Sigma_x = \frac{1}{\sqrt{2\pi}} \frac{\int R_x(\delta) d\delta}{R_x(0)} \quad (122)$$

Equation (122) has straightforward interpretation if $R_x(\delta)$ is Gaussian, in which case Σ_x is simply the standard deviation of that Gaussian. The vertical beam width Σ_y is defined analogously. It is obtained by performing a vdM scan in y . In terms of the beam widths, the luminosity is given by:

$$\mathcal{L}_b = \frac{f_r n_1 n_2}{2\pi \Sigma_x \Sigma_y} \quad (123)$$

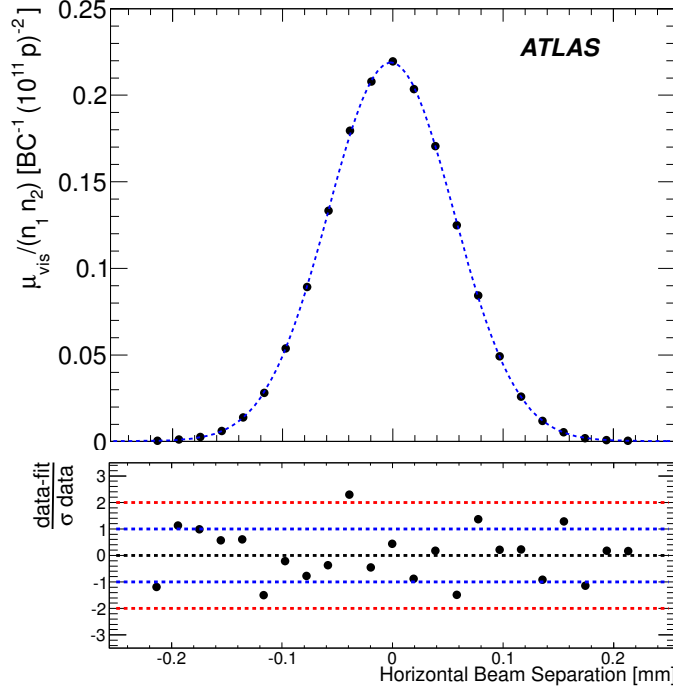


Figure 40: The top plot shows μ_{vis} as a function of beam separation as recorded during a vdM scan [1]. All measurements are normalized to the bunch population product. A Gaussian on top of a constant background is used to fit the scan curve. The bottom plot shows the residuals normalized to the statistical uncertainty of each point.

A luminosity monitor is calibrated by equating Eq. (123) to Eq. (116). This yields

$$\sigma_{\text{vis}} = \mu_{\text{vis}}^{\text{MAX}} \frac{2\pi \Sigma_x \Sigma_y}{n_1 n_2} \quad (124)$$

where $\mu_{\text{vis}}^{\text{MAX}}$ is the visible interaction rate observed at the peak of the scan curve by the luminosity monitor in question. If the x and y scans yield different $\mu_{\text{vis}}^{\text{MAX}}$, the average value is used. Since both $\mu_{\text{vis}}^{\text{MAX}}$ and $\Sigma_x \Sigma_y$ are measured directly in the vdM scan, the only remaining unknown quantity in Eq. (124) is the bunch population product $n_1 n_2$. It is measured by a set of current transformers located around the LHC ring. Finally, it is worth noting that Eq. (124) remains valid also in the case that the crossing-angle of the beams is non-zero. The derivation of the formula does, however, become considerably more involved [63].

6.4 Algorithms

This section describes how LUCID measures μ_{vis} using the counting algorithms listed in Section 5.3.2. Let ϵ_{OR} denote the efficiency of the OR condition. This is the probability for a single pp interaction to result in a hit in at least one of the LUCID tubes. Similarly, let ϵ_{ORA} , ϵ_{ORC} and ϵ_{AND} denote the

efficiencies of the ORA, ORC and AND conditions. The efficiencies are *inclusive*, which means that a given pp interaction can fulfill any number (zero or more) of the conditions. It is convenient to also define four *exclusive* conditions, such that a pp interaction must fulfill exactly one of them. Let $P_{00}(n)$ denote the probability that a bunch crossing with n pp interactions does not cause a hit in any tube. Similarly, let $P_{10}(n)$ denote the probability to register at least one hit on side A but *not* on side C, $P_{01}(n)$ the probability to register at least one hit on side C but *not* on side A, and $P_{11}(n)$ the probability to register at least one hit both on side A and side C. Equations (125) to (128) show how the exclusive detection probabilities are related to the efficiencies, and Figure 41 illustrates the same relations with Venn diagrams.

$$P_{00}(1) = 1 - \epsilon_{\text{OR}} \quad (125)$$

$$P_{10}(1) = \epsilon_{\text{ORA}} - \epsilon_{\text{AND}} \quad (126)$$

$$P_{01}(1) = \epsilon_{\text{ORC}} - \epsilon_{\text{AND}} \quad (127)$$

$$P_{11}(1) = \epsilon_{\text{AND}} \quad (128)$$

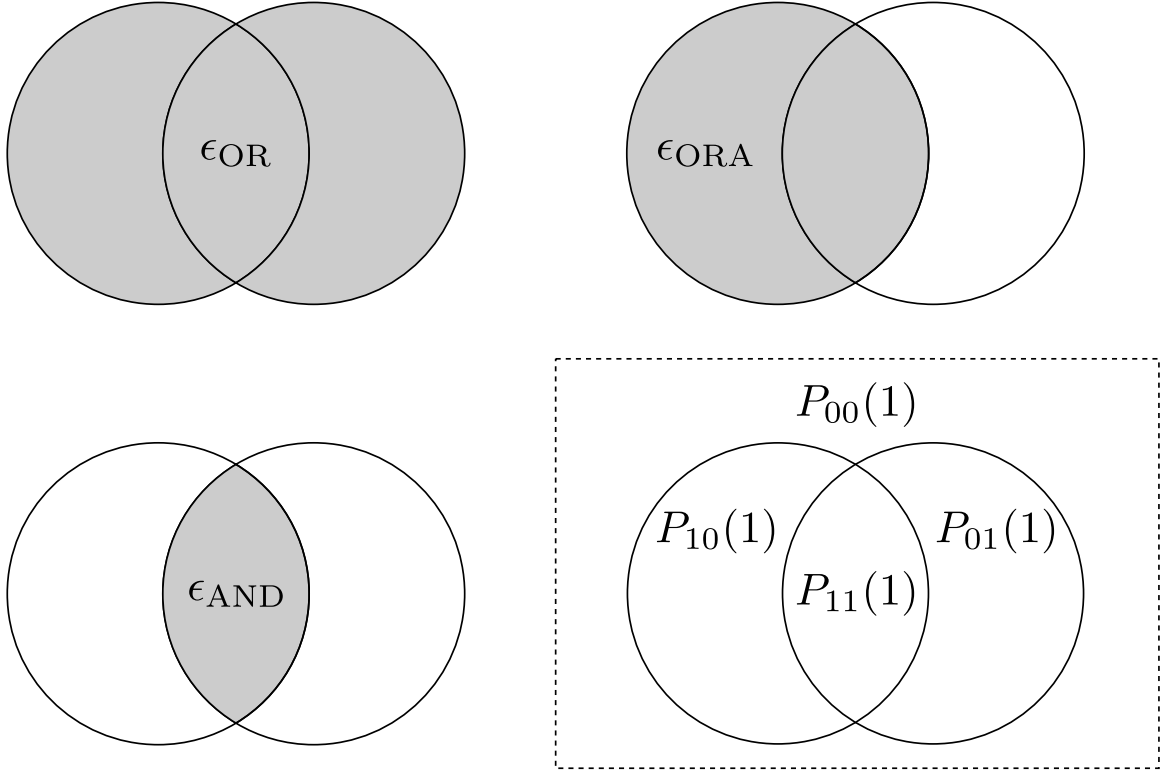


Figure 41: Venn diagrams illustrating the meaning of the LUCID efficiencies and their relation to the exclusive detection probabilities. The leftmost circle in each Venn diagram represents a hit in the side A readout module, while the rightmost circle represents a hit in the side C readout module.

6.4.1 EventOR

The EventOR algorithm counts the number of bunch crossings that fulfill the OR condition. Since the total number of bunch crossings is known, an equivalent way of counting the number of events that

fulfill the OR condition is to count the number of events that do *not* fulfill the OR condition. If all pp interactions in a bunch crossing are independent, the probability to observe no hits in a bunch crossing with n interactions is given by:

$$P_{00}(n) = P_{00}(1)^n \quad (129)$$

The event counting is carried out over many bunch crossings. Since n is a Poisson-distributed variable with a mean of μ , the average probability to observe no hits in LUCID is given by:

$$P_{00}(\mu) = \sum_{n=0}^{\infty} P_{00}(n) \frac{e^{-\mu} \mu^n}{n!} = \sum_{n=0}^{\infty} P_{00}(1)^n \frac{e^{-\mu} \mu^n}{n!} = e^{-\mu} \sum_{n=0}^{\infty} \frac{(P_{00}(1) \cdot \mu)^n}{n!} \quad (130)$$

Equation (130) can be simplified by noting that it contains the Maclaurin expansion for an exponential.

$$e^x = \sum_{n=0}^{\infty} \frac{x^n}{n!} \quad (131)$$

Inserting Eqs. (131) and (125) into Eq. (130) yields

$$P_{00}(\mu) = e^{-\mu} \cdot e^{P_{00}(1) \cdot \mu} = e^{-\mu(1-P_{00}(1))} = e^{-\epsilon_{\text{OR}} \mu} = e^{-\mu_{\text{vis}}^{\text{OR}}} \quad (132)$$

which can be solved for $\mu_{\text{vis}}^{\text{OR}}$ according to:

$$\mu_{\text{vis}}^{\text{OR}} = -\ln(P_{00}(\mu)) \quad (133)$$

LUCID measures the probability $P_{00}(\mu)$ by counting the fraction of bunch crossings that do not fulfill the OR condition. This can be expressed as

$$P_{00}(\mu) = 1 - \frac{N_{\text{EventOR}}}{N_{\text{BC}}} \quad (134)$$

where N_{EventOR} is the number of events that *do* fulfill the OR condition over the course of N_{BC} bunch crossings. Similarly, $\mu_{\text{vis}}^{\text{ORA}}$ and $\mu_{\text{vis}}^{\text{ORC}}$ are measured by counting the number of events N_{EventORA} and N_{EventORC} that fulfill the ORA and ORC conditions, respectively. The final formulas for measuring μ_{vis} using the EventOR, EventORA and EventORC algorithms are given by:

$$\mu_{\text{vis}}^{\text{OR}} = -\ln\left(1 - \frac{N_{\text{EventOR}}}{N_{\text{BC}}}\right) \quad (135)$$

$$\mu_{\text{vis}}^{\text{ORA}} = -\ln\left(1 - \frac{N_{\text{EventORA}}}{N_{\text{BC}}}\right) \quad (136)$$

$$\mu_{\text{vis}}^{\text{ORC}} = -\ln\left(1 - \frac{N_{\text{EventORC}}}{N_{\text{BC}}}\right) \quad (137)$$

In the derivation of Eqs. (135) to (137), it was assumed that all pp interactions in a bunch crossing are independent. More specifically, the assumption states that the probability for a given pp interaction to cause a hit in LUCID does not depend on how many other pp interactions take place in the same bunch crossing. This is a good assumption at low or moderate luminosities, but as explained in Section 7.4, it does not necessarily hold when the luminosity is very high.

6.4.2 EventAND

The EventAND algorithm counts the number of events that fulfill the AND condition. The formula for measuring $\mu_{\text{vis}}^{\text{AND}}$ using this algorithm is derived following a procedure similar to the one in Section 6.4.1. It starts by noting that the probability for a bunch crossing with n pp interactions to *not* fulfill the AND condition is given by:

$$P_{00,10,01}(n) = P_{00}(n) + P_{10}(n) + P_{01}(n) \quad (138)$$

If all pp interactions in a bunch crossing are independent, $P_{00}(n)$ is given by Eq. (129). The other probabilities are slightly more complicated to calculate. All permutations of k interactions producing a hit on side A with the remaining $n - k$ interactions producing no hits contribute to $P_{10}(n)$. Similarly, all permutations that produce hits only on side C contribute to $P_{01}(n)$.

$$P_{10}(n) = \sum_{k=1}^n P_{10}(1)^k P_{00}(1)^{n-k} \binom{n}{k} = (P_{10}(1) + P_{00}(1))^n - P_{00}(1)^n \quad (139)$$

$$P_{01}(n) = \sum_{k=1}^n P_{01}(1)^k P_{00}(1)^{n-k} \binom{n}{k} = (P_{01}(1) + P_{00}(1))^n - P_{00}(1)^n \quad (140)$$

Inserting Eqs. (129), (139) and (140) into Eq. (138) yields:

$$P_{00,10,01}(n) = (P_{10}(1) + P_{00}(1))^n + (P_{01}(1) + P_{00}(1))^n - P_{00}(1)^n \quad (141)$$

Since the average number of pp interactions per bunch crossing is μ , the average probability for a bunch crossing to not fulfill the AND condition is given by

$$P_{00,10,01}(\mu) = \sum_{n=0}^{\infty} (P_{10}(1) + P_{00}(1))^n \frac{e^{-\mu} \mu^n}{n!} + \sum_{n=0}^{\infty} (P_{01}(1) + P_{00}(1))^n \frac{e^{-\mu} \mu^n}{n!} - \sum_{n=0}^{\infty} P_{00}(1)^n \frac{e^{-\mu} \mu^n}{n!} \quad (142)$$

which, using the Maclaurin expansion in Eq. (131), can be simplified to:

$$P_{00,10,01}(\mu) = e^{-\mu(1-P_{10}(1)-P_{00}(1))} + e^{-\mu(1-P_{01}(1)-P_{00}(1))} - e^{-\mu(1-P_{00}(1))} \quad (143)$$

Inserting Eqs. (125) to (127) and noting that $\epsilon_{\text{OR}} + \epsilon_{\text{AND}} = \epsilon_{\text{ORA}} + \epsilon_{\text{ORC}}$ further simplifies the equation.

$$P_{00,10,01}(\mu) = e^{-\mu\epsilon_{\text{ORA}}} + e^{-\mu\epsilon_{\text{ORC}}} - e^{-\mu\epsilon_{\text{OR}}} \quad (144)$$

Rewriting Eq. (144) in terms of $\mu_{\text{vis}}^{\text{AND}} = \mu\epsilon_{\text{AND}}$ then yields

$$\begin{aligned} P_{00,10,01}(\mu) &= e^{-\mu_{\text{vis}}^{\text{AND}} \frac{\epsilon_{\text{ORA}}}{\epsilon_{\text{AND}}}} + e^{-\mu_{\text{vis}}^{\text{AND}} \frac{\epsilon_{\text{ORC}}}{\epsilon_{\text{AND}}}} - e^{-\mu_{\text{vis}}^{\text{AND}} \frac{\epsilon_{\text{OR}}}{\epsilon_{\text{AND}}}} \\ &= e^{-\mu_{\text{vis}}^{\text{AND}} \frac{\epsilon_{\text{ORA}}}{\epsilon_{\text{AND}}}} + e^{-\mu_{\text{vis}}^{\text{AND}} \frac{\epsilon_{\text{ORC}}}{\epsilon_{\text{AND}}}} - e^{-\mu_{\text{vis}}^{\text{AND}} \left(\frac{\epsilon_{\text{ORA}}}{\epsilon_{\text{AND}}} + \frac{\epsilon_{\text{ORC}}}{\epsilon_{\text{AND}}} - 1 \right)} \end{aligned} \quad (145)$$

which can also be expressed in terms of visible cross-sections.

$$P_{00,10,01}(\mu) = e^{-\mu_{\text{vis}}^{\text{AND}} \frac{\sigma_{\text{vis}}^{\text{ORA}}}{\sigma_{\text{vis}}^{\text{AND}}}} + e^{-\mu_{\text{vis}}^{\text{AND}} \frac{\sigma_{\text{vis}}^{\text{ORC}}}{\sigma_{\text{vis}}^{\text{AND}}}} - e^{-\mu_{\text{vis}}^{\text{AND}} \left(\frac{\sigma_{\text{vis}}^{\text{ORA}}}{\sigma_{\text{vis}}^{\text{AND}}} + \frac{\sigma_{\text{vis}}^{\text{ORC}}}{\sigma_{\text{vis}}^{\text{AND}}} - 1 \right)} \quad (146)$$

Equation (146) can not be solved analytically for μ_{vis} . In practice, it is instead inverted numerically using a lookup table. The probability $P_{00,10,01}(\mu)$ is measured by counting the fraction of bunch crossings that do not fulfill the AND condition.

$$P_{00,10,01}(\mu) = 1 - \frac{N_{\text{EventAND}}}{N_{\text{BC}}} \quad (147)$$

6.4.3 Hit counting

The HitOR algorithm counts the number of hits in bunch crossings that fulfill the OR condition. The formula used to calculate μ_{vis} with the HitOR algorithm has the same functional form as that of the EventOR algorithm. However, rather than simply counting the fraction of *events* per bunch crossing, the quantity of interest is the average number of *hits* per bunch crossing such that

$$\mu_{\text{vis}}^{\text{HitOR}} = -\ln\left(1 - \frac{N_{\text{HitOR}}}{N_{\text{BC}}N_{\text{CH}}}\right) \quad (148)$$

where N_{HitOR} is the number of hits accumulated over the course of N_{BC} bunch crossings using N_{CH} readout channels. Whereas event counting algorithms saturate once all bunch crossings give rise to at least one hit, the HitOR algorithm does not saturate until all bunch crossings give rise to a hit in every readout channel. It is worth noting that in case a detector has only one readout channel, hit counting and event counting are identical.

The HitAND algorithm counts the number of hits in bunch crossings that fulfill the AND condition. It is not possible to analytically derive a formula that allows μ_{vis} to be measured using this algorithm. However, as shown in Section 7, data-driven methods can still be used.

6.5 The Online Luminosity Calculator

All subdetectors in ATLAS are able to share data over a TDAQ infrastructure called the *information service* (IS). The luminosity detectors periodically publish luminosity information to IS. In the case of LUCID, the LUMAT board publishes the number of counts recorded by each algorithm as well as the number of orbits over which the counts were accumulated. This information is then picked up by a program called the Online Luminosity Calculator (OLC). Whenever the CTP signals the end of a Luminosity Block, the OLC polls the IS publications of each luminosity detector and uses the published information to calculate $\mathcal{L}$ according to a set of predefined formulas. For example, to calculate the luminosity based on the EventOR algorithm, the OLC first calculates $\mu_{\text{vis}}^{\text{OR}}$ based on Eq. (135) and then obtains $\mathcal{L}$ using Eqs. (116) and (114). The number of colliding bunch pairs n_b is obtained from IS based on the publications of other ATLAS subdetectors, while the visible cross-section for the EventOR algorithm (which is just a calibration constant) is read from an offline database.

In addition to calculating $\mathcal{L}$, the OLC performs a series of data quality checks to ensure that the luminosity infrastructure is working as it should. In case problems are detected, a warning is issued to the ATLAS control room. The OLC also republishes select pieces of information to IS for use in various trending plots. Such information is used for monitoring purposes and to quickly diagnose detector problems. For example, if two luminosity monitors disagree, a trending plot showing luminosity measurements from several detectors could be used to determine which monitor is in error and when the problem started. Finally, all the relevant information from all the relevant luminosity detectors is saved to disk. This is handled by a separate application called OLC2COOL. It is named after the COOL software suite, which is a framework used to manage the database where the information is stored. The data stored in the COOL database is then analyzed offline in order to calculate the official ATLAS luminosity. The offline luminosity calculation includes background corrections that are difficult to apply online. It is therefore more precise than the online calculation.

7 The LUCID pileup method

7.1 Introduction

The LUCID detector (see Section 5) is a luminosity monitor that measures the luminosity of the LHC. This is done by counting *hits* and *events* using a variety of algorithms as discussed in Section 6.4. The average probability $P(\mu)$ for a bunch crossing to be counted by LUCID is given by

$$P(\mu) = \sum_{n=0}^{\infty} P(n) \frac{e^{-\mu} \mu^n}{n!} \quad (149)$$

where $P(n)$ is the probability to count a bunch crossing containing exactly n pp interactions and μ is the average number of pp interactions per bunch crossing. The counting algorithms measure $P(\mu)$, and Eq. (149) is solved for μ by calculating $P(n)$. This can be done analytically, but only under the assumption that the probability for a given pp interaction to cause a hit in LUCID does not depend on how many other pp interactions take place in the same bunch crossing. Another restriction is that no analytic solution exists when measuring $P(\mu)$ based on the HitAND algorithm. This section describes a data-driven calibration method that can be applied to any counting algorithm and that makes no assumptions about how the number of pp interactions within a bunch crossing affects the detection probability. The method is referred to as the *pileup method*. Its implementation and performance in the LUCID detector are discussed. It should be noted that all data in this section comes from standard tubes, and any mention of PMTs refers to standard PMTs rather than fiber PMTs.

7.2 Method

The goal of the data-driven method is to measure $P(n)$ explicitly up to some $n_{\max}$. The infinite series in Eq. (149) is then truncated, and it is assumed that contributions from omitted terms are negligible.

$$P(\mu) = \sum_{n=0}^{n_{\max}} P(n) \frac{e^{-\mu} \mu^n}{n!} \quad (150)$$

Once $P(n)$ is known for every n , Eq. (150) can be solved numerically. Since $P(n)$ is the probability for LUCID to count a bunch crossing that contains n pp interactions, one way of measuring it is to construct a data sample consisting purely of bunch crossings with exactly n pp interactions. This is done for every n up to $n_{\max}$. The probabilities $P(n)$ are then obtained by counting the number of hits or events in each sample.

7.2.1 Obtaining a single-interaction sample

The first step is to obtain a data sample consisting purely of bunch crossings with $n = 1$. Since n is really a random variable that follows a Poisson distribution, real data contains an admixture of bunch crossings with various different n . However, it is possible to obtain a data sample that consists almost entirely of bunch crossings with $n = 0$ or $n = 1$ by running at a very low luminosity. A quick numerical exercise shows that when running with $\mu = 0.01$, approximately 99% of all bunch crossings have $n = 0$. The remaining bunch crossings are 200 times more likely to have $n = 1$ than $n > 1$.

Figure 42 shows a charge spectrum obtained from a single PMT during a run taken at low μ . The data was obtained by triggering randomly on a specific BCID. Three signal categories can be defined. The first category is made up of signals above the discriminator threshold, i.e. signals that produced a hit. If such a signal is present in any PMT, it is almost certain that the bunch crossing contained a pp interaction. The second category is made up of signals below the discriminator threshold but far from

the pedestal. If such a signal is present in any PMT, it is likely that the bunch crossing contained a pp interaction. The final category is made up of signals close to or within the pedestal distribution. If all PMTs contain such signals, it is impossible to determine whether or not the bunch crossing contained any pp interactions.

Let N_{EventOR} and N_{EventAND} denote the number of bunch crossings that fulfill the OR and AND conditions, respectively. A pure single-interaction sample containing N_{Single} bunch crossings must fulfill

$$N_{\text{Single}} \cdot \varepsilon_{\text{OR}} = N_{\text{EventOR}} \quad (151)$$

$$N_{\text{Single}} \cdot \varepsilon_{\text{AND}} = N_{\text{EventAND}} \quad (152)$$

where ε_{OR} and ε_{AND} are the (known) detector efficiencies. In case the discriminator thresholds are not known exactly, they can be determined using Eqs. (151) and (152). Offline cuts corresponding to the approximate threshold levels are applied to each PMT. These cuts are then varied until

$$N_{\text{EventOR}} \cdot \frac{\varepsilon_{\text{AND}}}{\varepsilon_{\text{OR}}} = N_{\text{EventAND}} \quad (153)$$

is satisfied. The final cut positions give the discriminator thresholds.

The low- μ data sample can be transformed into a good approximation of a single-interaction sample by discarding bunch crossings until only N_{Single} bunch crossings remain. The bunch crossings to be discarded are chosen at random from the pedestal category, i.e. from bunch crossings in which all PMTs are below the pedestal boundary. The final sample, by construction, is guaranteed to fulfill:

$$\frac{N_{\text{EventOR}}}{N_{\text{Single}}} = \varepsilon_{\text{OR}} = P(1)_{\text{EventOR}} \quad (154)$$

$$\frac{N_{\text{EventAND}}}{N_{\text{Single}}} = \varepsilon_{\text{AND}} = P(1)_{\text{EventAND}} \quad (155)$$

The position of the pedestal boundary (see Figure 42) is somewhat arbitrary. It is only constrained by the fact that all signals above the boundary are kept for the final sample. If the position is such that the number of kept events is greater than N_{Single} , Eqs. (151) and (152) can not be satisfied. In practice, this only happens if the boundary is placed inside of the pedestal distribution. Figure 43 shows the charge spectrum after transforming the low- μ sample into a pure single-interaction sample. Only about 1% of the events below the boundary remain. However, the shape of the spectrum outside of the pedestal is almost unchanged because most of those events were above the boundary in other PMTs. Therefore, the exact position of the boundary has a negligible effect on the final result as long as it is outside of the pedestal distribution.

7.2.2 Pileup

After producing a sample consisting purely of bunch crossings with $n = 1$, the next step is to construct bunch crossings with $n = 2$, $n = 3$ and so on. This is achieved through a pileup process where n pulses from single-interaction events are summed up to simulate the pulse shape produced by one bunch crossing with n pp interactions. The principle is illustrated in Figure 44. In order to correct for timing effects, the signals are shifted such that their maxima coincide before being summed.

There is no need to save the complete pulse shape of the piled-up events. The only relevant information is the number of hits, or the *hit multiplicity* as it is also called. Figure 45 shows the hit multiplicity in LUCID for events with $n = 1$ and $n = 15$ after imposing the OR and AND conditions. These histograms are called reference distributions. The OR reference distributions can be used to calculate the reference

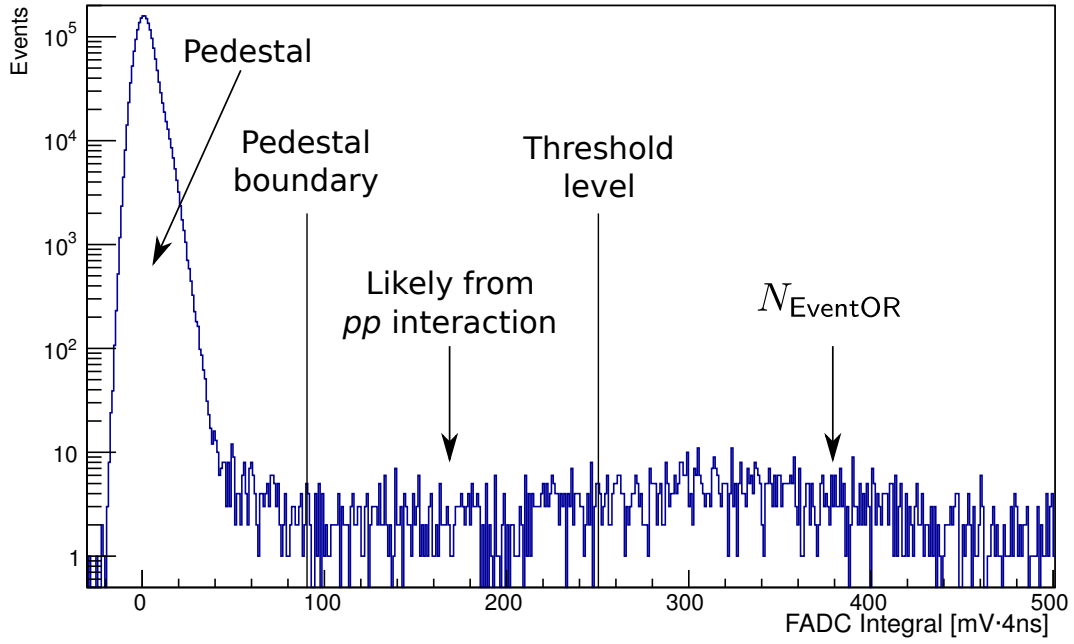


Figure 42: Charge spectrum from a low- μ run. Signals far from the pedestal are likely due to single-interaction bunch crossings. The boundary separating such signals from the pedestal is indicated. The discriminator threshold level is indicated as well.

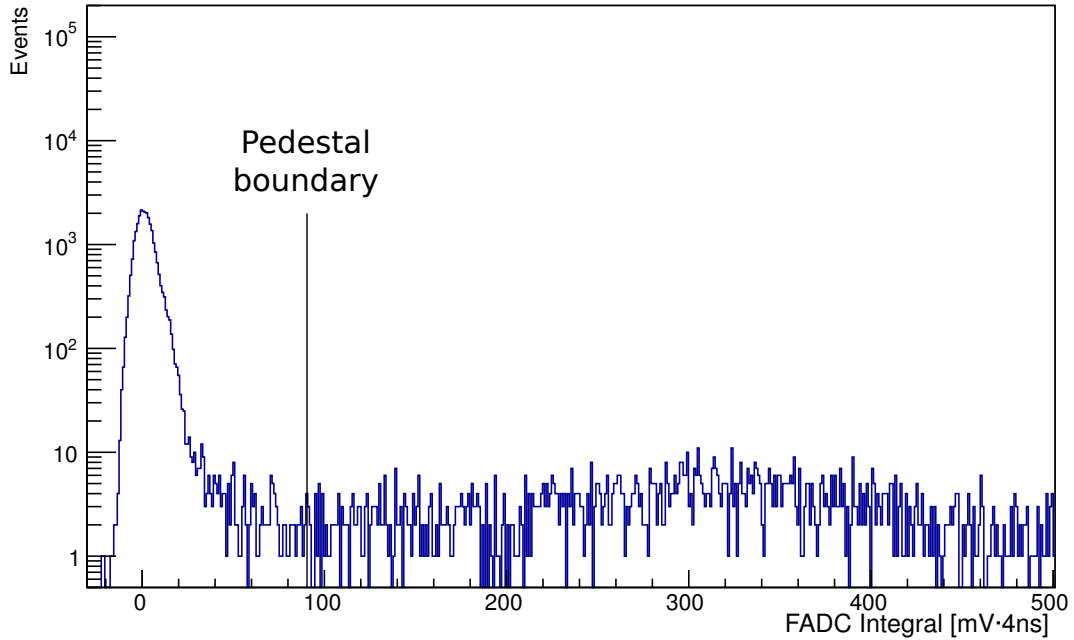


Figure 43: Charge spectrum after transforming the low- μ run into a pure single-interaction sample.

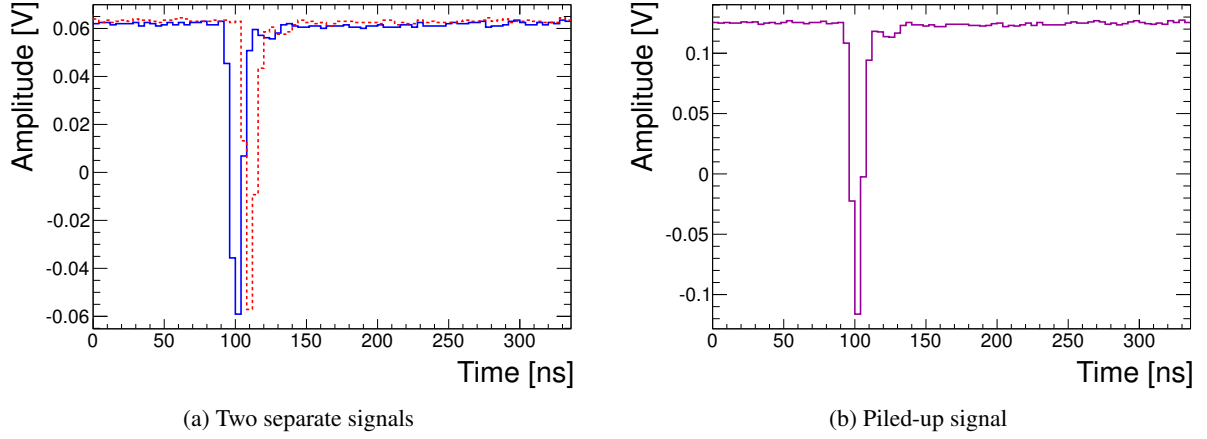


Figure 44: The two pulses shown in (a) are shifted so that their maxima coincide. They are then summed to produce the piled-up pulse shown in (b).

probabilities $P(n)_{\text{EventOR}}$ and $P(n)_{\text{HitOR}}$. Similarly, the AND reference distributions are used to obtain $P(n)_{\text{EventAND}}$ and $P(n)_{\text{HitAND}}$. Distributions for other conditions, such as ORA and ORC, can be obtained as well. The goal of the pileup procedure is to build reference distributions up to some sufficiently high n .

The maximum required n is dictated by the running conditions of the LHC. Recall that the purpose of obtaining the reference probabilities $P(n)$ is to solve Eq. (150). In theory, $P(\mu)$ receives contributions from every $P(n)$ up to infinity. In practice, $P(n)$ only needs to be calculated up to the point where the corresponding Poisson weight becomes negligible. A reasonable stopping condition is to round μ up to the nearest integer and require the Poisson weight for n_{max} to be less than a permil of the Poisson weight for μ . This can be written as

$$\frac{\mathcal{P}(n_{\text{max}}|\mu)}{\mathcal{P}(\mu|\mu)} < 0.001 \quad (156)$$

where

$$\mathcal{P}(n|\mu) = \frac{e^{-\mu} \mu^n}{n!}. \quad (157)$$

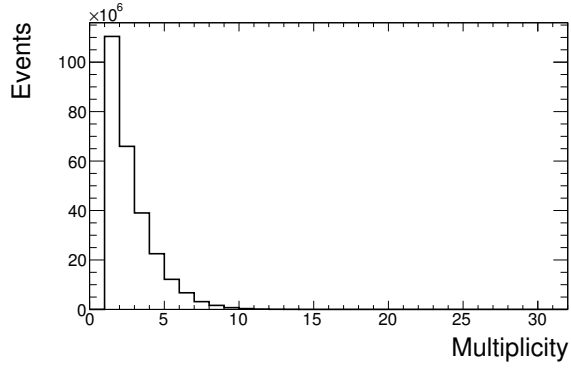
The pileup events are drawn at random from the single-interaction sample. In the simplest case where the same event can be chosen multiple times, the number of possible ways to construct an n -interaction bunch crossing is the number of n -multicombinations in a set of size N_{Single} .

$$N_{\text{Comb}} = \left(\binom{N_{\text{Single}}}{n} \right) \quad (158)$$

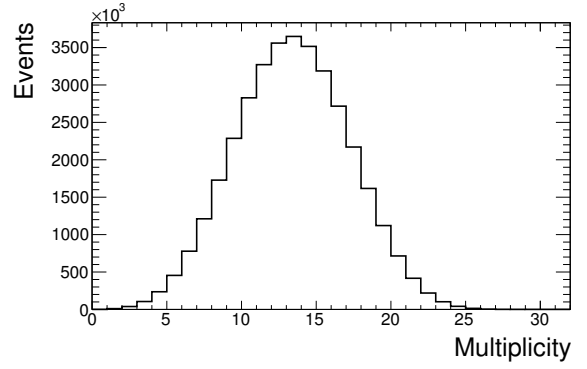
7.3 Data

This analysis is based on data from two different runs. The threshold uncertainty, which is discussed in Section 7.6.1, is evaluated using $\mu \approx 0.007$ data taken from ATLAS run 182204. The single-interaction sample size is $N_{\text{Single}} = 13209$. All other studies are done using $\mu \approx 0.012$ data taken from ATLAS run 200805. The single-interaction sample size is $N_{\text{Single}} = 52083$.

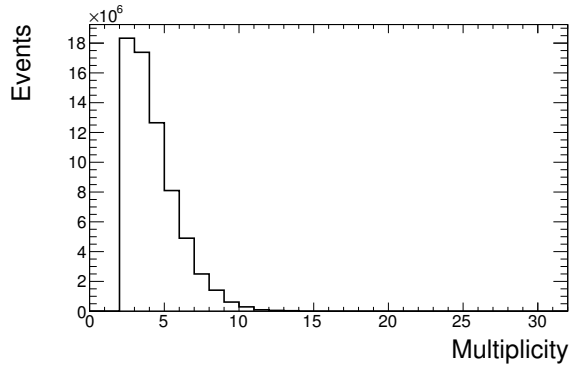
Using $N_{\text{Single}} = 13209$ in Eq. (158) yields approximately $N_{\text{Comb}} = 10^8$ for $n = 2$ and $N_{\text{Comb}} = 10^{11}$ for $n = 3$. This is clearly a sufficient sample size for generating unique events at high n .



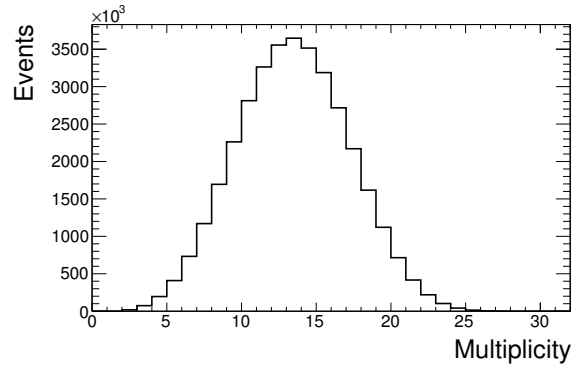
(a) OR distribution with $n = 1$



(b) OR distribution with $n = 15$



(c) AND distribution with $n = 1$



(d) AND distribution with $n = 15$

Figure 45: Reference distributions obtained after imposing the OR condition for $n = 1$ (a) and $n = 15$ (b) are shown together with reference distributions obtained after imposing the AND condition for $n = 1$ (c) and $n = 15$ (d). At high n , the OR and AND distributions are almost identical because nearly all events fulfill the AND condition.

7.4 Migration effects

The analytic formulas derived in Section 6.4 assume that all pp interactions in a bunch crossing are independent. In other words, they assume that the probability for a given pp interaction to cause a hit in LUCID does not depend on the total number of pp interactions in that bunch crossing. This assumption is not necessarily good at high μ . As far as LUCID is concerned, the particles that result from a pp interaction can be divided into two categories. *Primary* particles come straight from the interaction point and produce large signals in LUCID. *Secondary* particles are created when primaries interact with dead material in the ATLAS detector or the beam pipe. They have less energy than primaries, can be incident from any angle, and typically produce smaller signals in LUCID.

When running at high μ , it is possible for small signals from many secondaries to add up in order to cause hits above threshold. This phenomenon is called *migration*, as the signals are said to migrate away from the low end of the charge spectrum. It is a nonlinear effect that normally leads to an overestimation of the luminosity, because the standard (analytic) algorithms are calibrated at low μ where migration is negligible. The pileup method, on the other hand, takes migration into account by design. The impact of migration effects on the nominal μ -measurement can therefore be estimated by comparing the pileup method to the standard algorithms.

Figure 46 shows the ratio of the μ -values obtained from the standard EventOR and HitOR algorithms and the pileup method. The difference is negligible for EventOR, while for HitOR the migration effects result in a significant μ -dependence. During 2012, the official LUCID luminosity measurement was based on the pileup method for the HitOR and HitAND algorithms. All other algorithms used the standard analytic formulas.

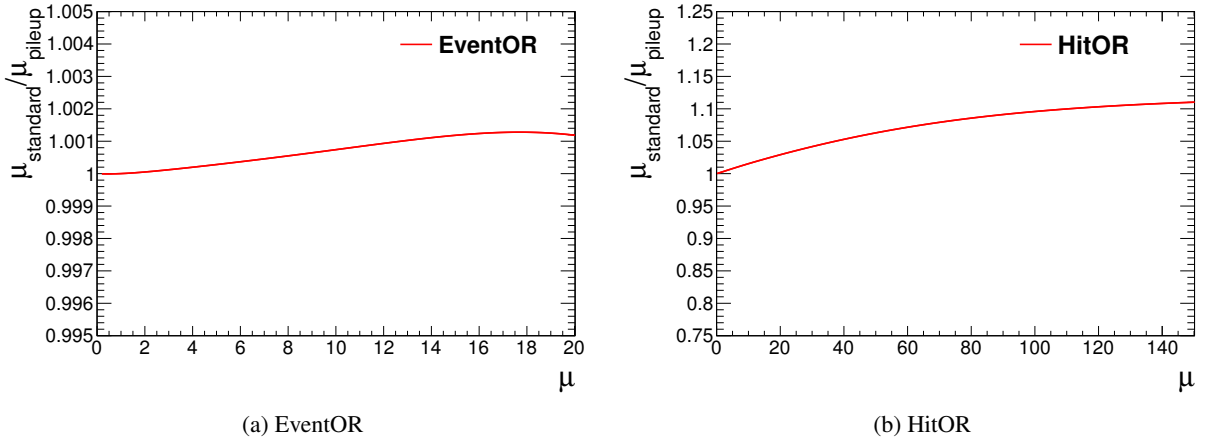


Figure 46: The ratio of the μ -values obtained from the standard algorithms and the pileup method. The EventOR (a) and HitOR (b) algorithms are shown. Differences between the standard algorithms and the pileup method are due to migration.

7.5 Statistical uncertainty

The first step towards obtaining the statistical uncertainty on μ is to estimate the statistical uncertainty of each reference probability $P(n)$. Once they are known, the resulting statistical uncertainty of $P(\mu)$ can be calculated. Finally, this is translated into a statistical uncertainty on μ by numerically inverting and differentiating Eq. (150).

The only source of statistical uncertainty on $P(n)$ is the number of recorded hits or events. The number of bunch crossings and readout channels in LUCID can be measured exactly. In the case of

event counting, the uncertainty is given by

$$\sigma_{P(n)\text{Event}} = \frac{\partial P(n)\text{Event}}{\partial N_{\text{Event}}} \sigma_{N_{\text{Event}}} = \frac{1}{N_{\text{BC}}} \sqrt{N_{\text{Event}}} = \sqrt{\frac{P(n)\text{Event}}{N_{\text{BC}}}} = \sqrt{\frac{nP(n)\text{Event}}{N_{\text{Piled}}}} \quad (159)$$

where N_{Piled} is the total number of events that were piled up in order to build the n :th reference distribution. The statistical uncertainty of $P(\mu)$ can be calculated according to:

$$\sigma_{P(\mu)} = \sum_{n=0}^{\infty} \frac{\partial P(\mu)}{\partial P(n)} \sigma_{P(n)} = \sum_{n=0}^{\infty} \frac{e^{-\mu} \mu^n}{n!} \sigma_{P(n)} \quad (160)$$

Plugging in $\sigma_{P(n)\text{Event}}$ yields:

$$\sigma_{P(\mu)\text{Event}} = \sum_{n=0}^{\infty} \frac{e^{-\mu} \mu^n}{n!} \sqrt{\frac{nP(n)\text{Event}}{N_{\text{Piled}}}} \leq \sqrt{\frac{\mu P(\mu)\text{Event}}{N_{\text{Piled}}}} \quad (161)$$

The statistical uncertainty of $P(\mu)$ in the hit counting case is derived analogously and has the same form except for an extra factor N_{CH} . It is given by:

$$\sigma_{P(\mu)\text{Hit}} = \sum_{n=0}^{\infty} \frac{e^{-\mu} \mu^n}{n!} \sqrt{\frac{nP(n)\text{Hit}}{N_{\text{Piled}} N_{\text{CH}}}} \leq \sqrt{\frac{\mu P(\mu)\text{Hit}}{N_{\text{Piled}} N_{\text{CH}}}} \quad (162)$$

The last step is to invert and differentiate Eq. (150) numerically. The statistical uncertainty on μ is then given by:

$$\sigma_{\mu} = \frac{\partial \mu}{\partial P(\mu)} \sigma_{P(\mu)} \quad (163)$$

Figure 47 shows the relative statistical uncertainty $\sigma_{\text{stat}} = \sigma_{\mu}/\mu$ for run 200805. The number of pileup events is $N_{\text{Piled}} = 5.4 \cdot 10^8$. For the event counting algorithms, the uncertainty quickly approaches unity as they saturate at high μ . This is an intrinsic limitation of the algorithms that can not be improved by increasing N_{Piled} . For the hit counting algorithms, the statistical uncertainty is negligible in the μ -range used for physics running.

7.6 Systematic uncertainty

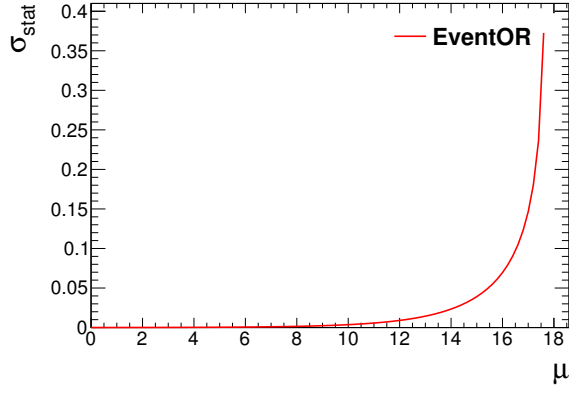
7.6.1 Thresholds

The threshold values are not known exactly. In order to estimate the impact of this uncertainty on the final μ -measurement, the entire pileup procedure is redone after scaling the thresholds T up and down according to

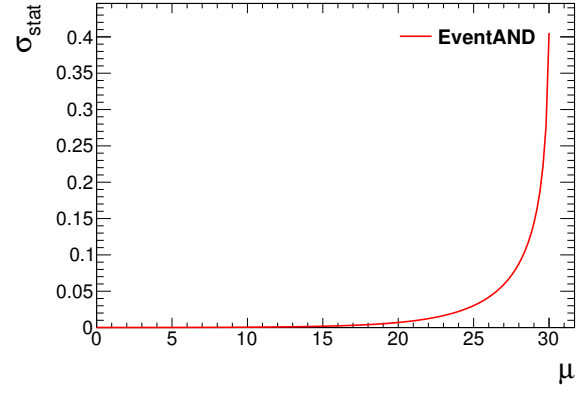
$$T_{\text{scaled}} = F \cdot T_{\text{nominal}} \quad (164)$$

where F is a scaling factor. The scaled result is compared to the nominal one, and the difference is taken to be a measure of the threshold systematic uncertainty σ_{thresh} . Scaling the thresholds means that Eq. (153) can no longer be satisfied, and therefore N_{single} is not uniquely determined. Equation (151) has been used to calculate N_{single} for the scaled samples.

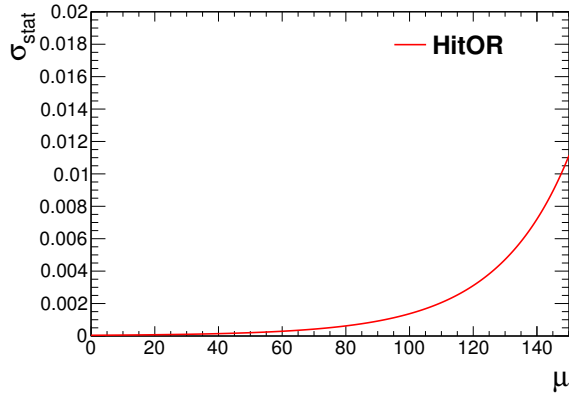
Figure 48 shows the effect of scaling the thresholds up and down by a conservative 5%. The impact on the EventOR algorithm is negligible because Eq. (154) still holds by construction. For EventAND, the difference stabilizes at 0.3%. This is taken to be the uncertainty also for EventOR. For both HitOR and HitAND, the difference stabilizes at 1.0%.



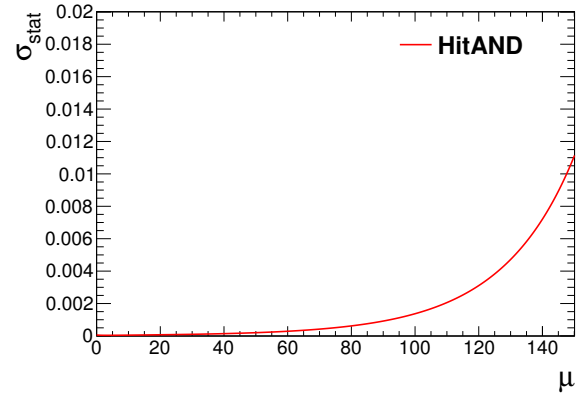
(a) EventOR



(b) EventAND

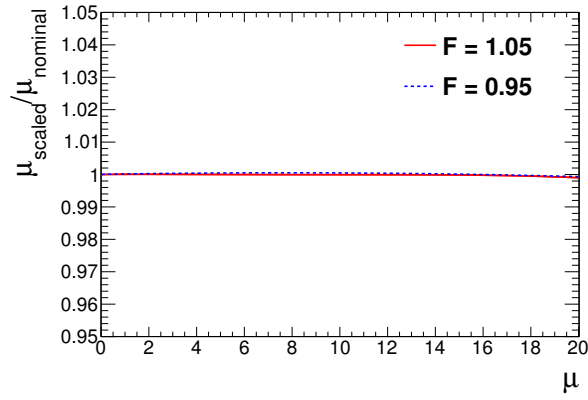


(c) HitOR

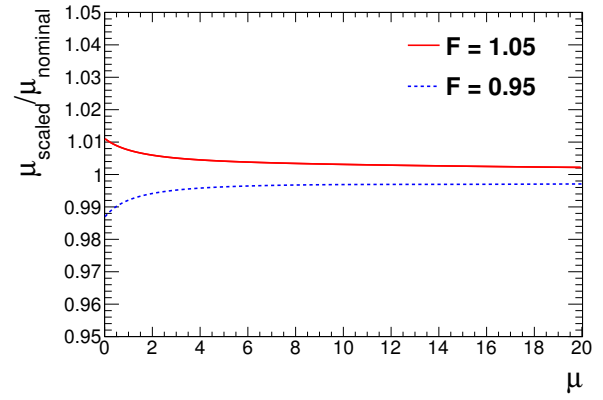


(d) HitAND

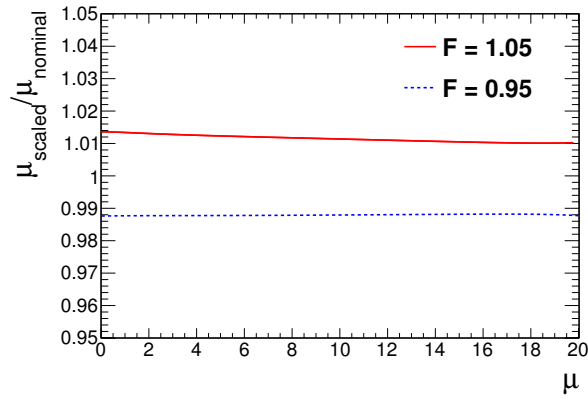
Figure 47: The relative statistical uncertainty σ_{stat} of the EventOR (a), EventAND (b), HitOR (c) and HitAND (d) algorithms.



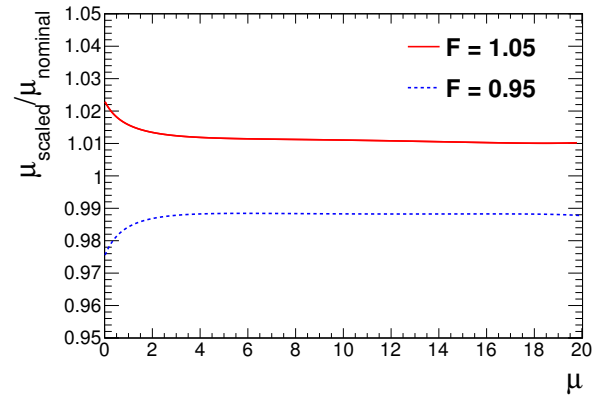
(a) EventOR



(b) EventAND



(c) HitOR



(d) HitAND

Figure 48: The ratio of the μ -values obtained with scaled and nominal thresholds. The EventOR (a), EventAND (b), HitOR (c) and HitAND (d) algorithms are shown.

7.6.2 Consistency

The uncertainty of the official ATLAS luminosity measurement involves a comparison between several different detectors and algorithms [1]. This section will show that the pileup method produces results that are consistent with other methods. The pileup method is therefore usable in the comparisons on which the official luminosity uncertainty is based, and a luminosity measured by the pileup method should not include any special uncertainties due to limited consistency. Unless specified otherwise, all hit algorithms shown in this section are based on the pileup method and all event algorithms are based on the standard formula.

Figure 49 compares the standard HitOR algorithm to the pileup HitOR and HitAND algorithms. All measurements are normalized to TileCal in such a way that they agree in the middle of the range by construction. The plot therefore shows the relative μ -dependence of the algorithms. The standard HitOR algorithm has a μ -dependence of 3% in the range $7 < \mu < 27$ while the pileup algorithms have no μ -dependence at all. This confirms the prediction made in Section 7.4.

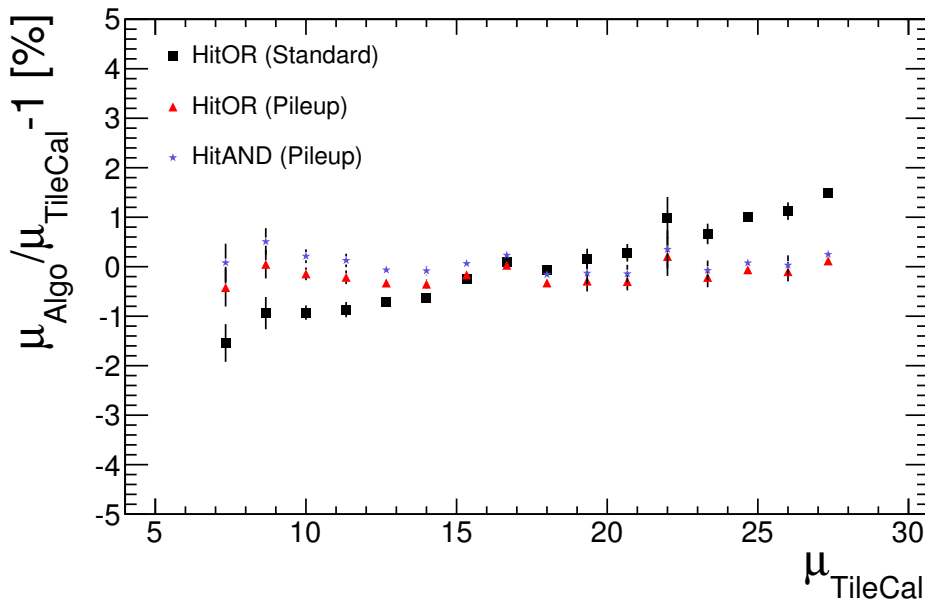


Figure 49: The standard HitOR algorithm compared to the pileup HitOR and HitAND algorithms. All measurements are normalized to TileCal. The data comes from run 206818, which had 36 colliding bunch pairs.

Figure 50 (a) shows various algorithms in a run with a high number of colliding bunches. All measurements are normalized to TileCal. The slope is caused by a nonlinear response from the LUCID PMTs when exposed to high currents. This effect is compensated for by applying a correction according to

$$\mu_{\text{corr}} = \frac{\mu}{a + b \cdot I + c \cdot I^2} \quad (165)$$

where I is the current drawn by the PMTs. Equation (165) is referred to as the *current correction*. It is used for runs where the number of colliding bunches is high. The constants a , b and c are determined by calibrating against another detector. This calibration is done only once using a specific run. Figure 50 (b) shows the current-corrected measurements. The pileup-based hit algorithms agree perfectly with the standard event algorithms. Figure 51 shows various algorithms from a run with a low number of colliding bunches. It is clear that the pileup algorithms agree with the standard ones also for runs in which no current correction is applied.

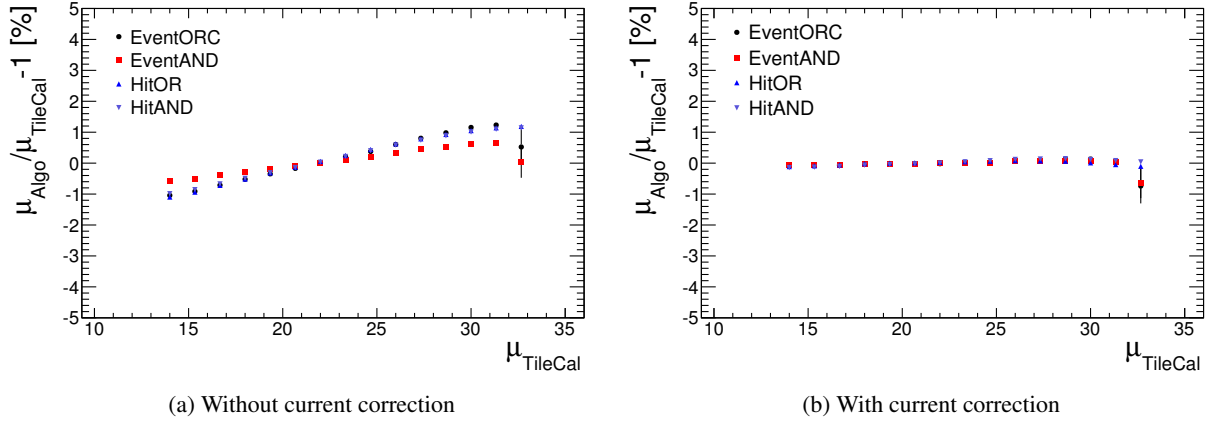


Figure 50: The figure shows various algorithms before (a) and after (b) applying the current correction. All measurements are normalized to TileCal. The data comes from run 206818, which had 1368 colliding bunch pairs.

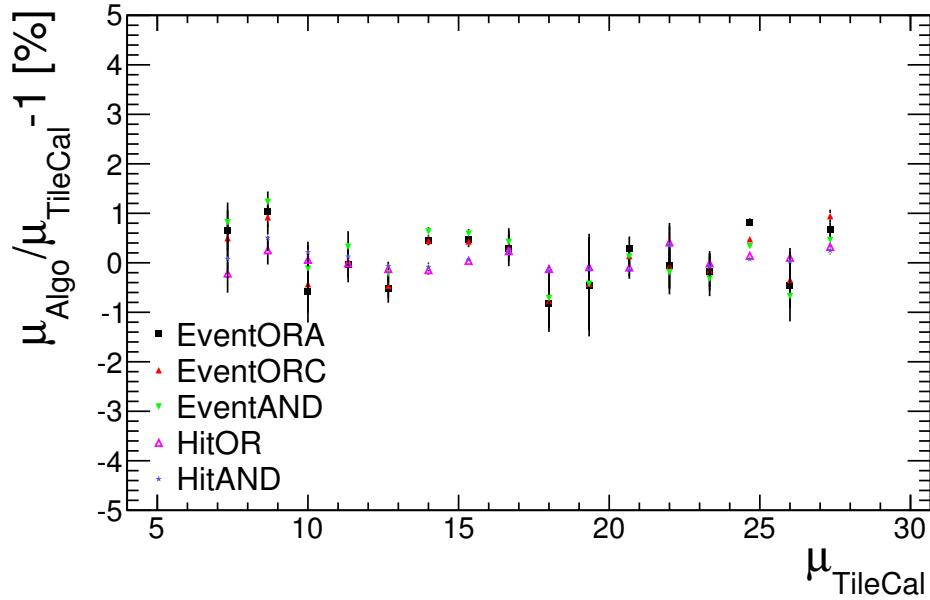


Figure 51: The figure shows various algorithms in run 206818. This run had only 36 colliding bunch pairs, so no current correction was applied. All measurements are normalized to TileCal.

Figure 52 shows the stability of various algorithms over time. Each data point was obtained by picking a particular algorithm in a particular run and calculating the average deviation per LB with respect to BCM. This shows that the time stability of the pileup method is consistent with the standard algorithms. Figure 52 also illustrates the manner in which the pileup method contributes to the uncertainty of the official ATLAS luminosity measurement. This measurement takes many different sources of uncertainty into account, one of which is the limited long-term consistency of the ATLAS luminosity determination. The impact of long-term inconsistencies is evaluated by comparing various different algorithms from various different detectors over time. The spread of the algorithms is taken to be a measure of the uncertainty. Figure 52 shows how LUCID would contribute to such an uncertainty. Since the spread of HitOR and HitAND is contained within the remaining algorithms, the pileup method does not drive the uncertainty of the official ATLAS luminosity measurement.

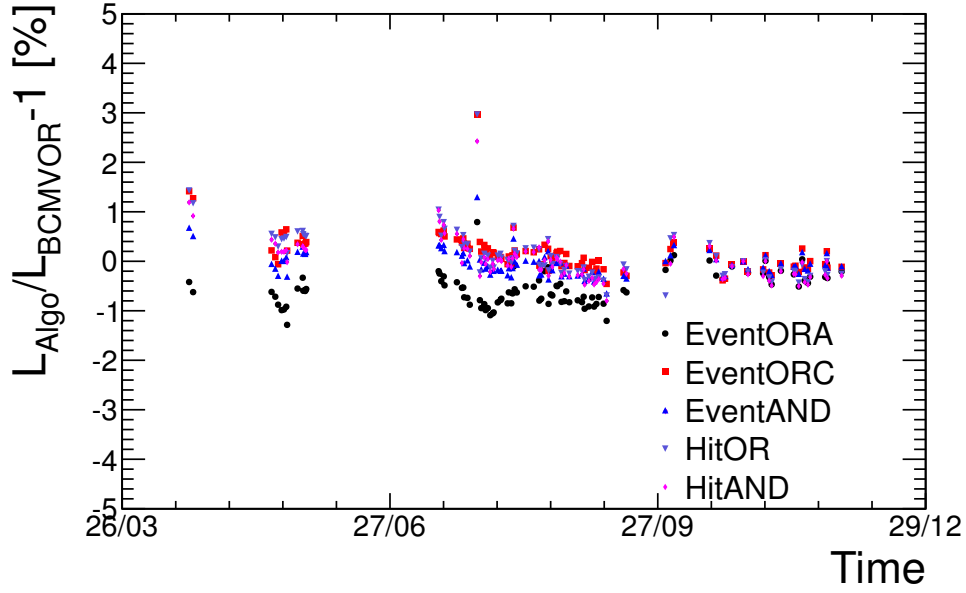


Figure 52: The figure shows the average deviation per LB of various algorithms with respect to BCM. Each point represents one run and algorithm. The x-axis shows the date in 2012.

7.7 Conclusions

A data-driven luminosity calibration method, called the pileup method, has been presented and evaluated using the LUCID detector. The method works with any counting algorithm, including those where no analytic solution exists, and corrects for migration effects by design. Luminosity measurements based on the pileup method have a method-specific uncertainty that should be added in quadrature to any other uncertainties associated with the measurement. This method-specific uncertainty is 0.3% for event counting algorithms and 1.0% for hit counting algorithms. The pileup method was used to calibrate the HitOR and HitAND algorithms during 2012. Use of the pileup method does not increase the uncertainty of the official ATLAS luminosity determination.

8 The LUCID particle counting method

8.1 Introduction

As discussed in Section 6, any quantity that is proportional to $\mathcal{L}$ could in principle be used to measure the luminosity. One example of such a quantity is the visible interaction rate μ_{vis} . The LUCID detector (see Section 5) measures μ_{vis} using various hit counting and event counting algorithms. This approach suffers from a few limitations. Firstly, counting algorithms can in general not be derived analytically. Solutions exist only for a few special cases (see Section 6.4). While a general solution can be obtained numerically (see Section 7), the method is reliant on input from dedicated low- μ runs that take time away from regular physics running. Secondly, counting algorithms work by imposing a threshold condition on the signal. This leads to migration effects as explained in Section 7.4. Finally, the algorithms inevitably saturate at high $\mathcal{L}$ when too many signals are above threshold.

An alternative to counting signals above threshold is to measure the total charge per bunch crossing produced by the LUCID photomultipliers. This method is referred to as *particle counting* in order to distinguish it from hit counting and event counting. Like μ_{vis} , the total signal charge is directly proportional to $\mathcal{L}$ and can therefore be used as a basis for luminosity determination. Unlike hit counting and event counting, particle counting works without imposing a threshold condition on the signal. All signals, regardless of their size, are allowed to add to the total charge. The method is therefore immune to migration. The absence of a threshold also means that there is no upper limit to how much charge an individual signal is allowed to contribute. This means that particle counting does not saturate in the same way that hit counting and event counting algorithms do. This section describes how particle counting can be implemented in LUCID and evaluates the performance of the method based on 2012 data. During run II, the upgraded LUCID detector will use particle counting alongside the standard algorithms.

8.2 Charge measurement

The charge measurement is carried out using data from the LUCID local stream. As explained in Section 5.3, the signal is recorded by a flash analog to digital converter (FADC). When the bunch spacing is 50 ns, the range of the FADC is such that it can record signals in four BCIDs after the trigger and one BCID before the trigger. The range before the triggering BCID is usually empty, however, as the triggering BCID is nominally chosen to be the first bunch in a bunch train. Figure 53 shows a PMT signal as recorded by the FADC. The charge of the signal is obtained by integrating over the pulse. This can be done on a per-BCID basis by simply adding up the baseline subtracted contents of each bin contained within the BCID of interest. However, finding the baseline position is not trivial. It varies from event to event and can even vary from BCID to BCID due to afterglow effects. *Afterglow* is a process caused by highly energetic primary particles. The primaries cause hadronic showers as they interact with the ATLAS detector material. The showers lead to nuclear excitations, and the photons emitted by the nuclei as they de-excite show up as a signal in the PMTs. Three different methods for determining the baseline are discussed in this chapter. They are referred to as the Static, PreTrain and Dynamic methods.

8.2.1 Static baselines

Static baselines are determined on a run-by-run basis. To get the baseline for a given run, one starts by selecting a subset of the run for which there was no beam in the LHC. For each readout channel, a starting value for the baseline is assumed. Using this value, the charge for every event in the subset is measured. The resulting charge spectrum should contain only a pedestal distribution as shown in Figure 54. The position of the pedestal distribution is then determined, for example using a Gaussian fit around the peak.

There are two disadvantages of using the Static baseline method. Firstly, because the same baseline

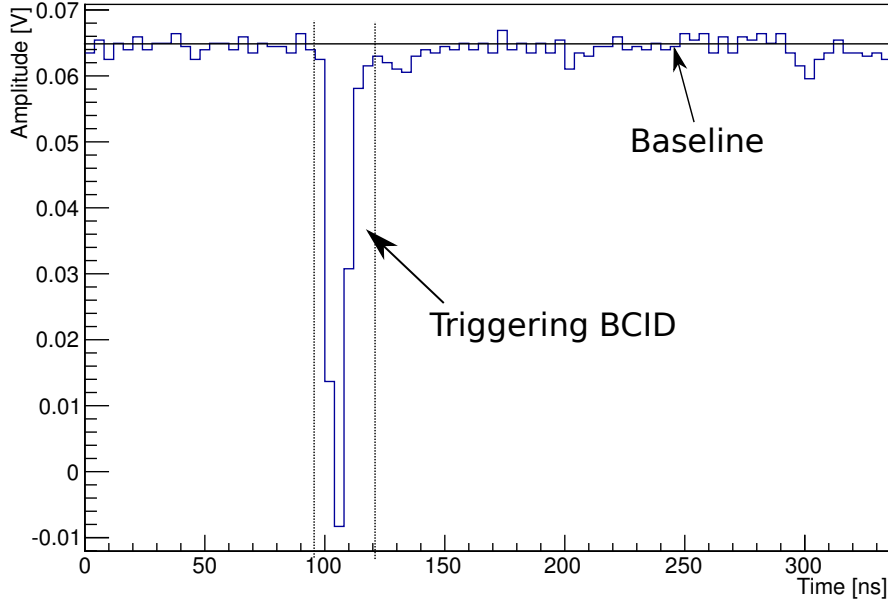


Figure 53: A flash ADC readout from one of the PMTs. The triggering BCID contains a pulse. The approximate baseline position and BCID width are indicated.

is used for the entire run, it is only correct on average. Event-to-event variations are not accounted for. Secondly, the baseline can not be determined until the run is already over. This makes Static baseline determination an inherently offline method. A baseline method that can be implemented online must be used in order to produce luminosity measurements in real time.

8.2.2 PreTrain baselines

PreTrain baselines are determined on an event-by-event basis. The principle is illustrated in Figure 55. Each bunch train is preceded by a range of empty BCIDs that is suitable for baseline calculation. Once the baseline has been determined, the same baseline is used for every bunch in the train. The PreTrain method can be implemented online, making it preferable to the Static method.

8.2.3 Dynamic baselines

Dynamic baselines are determined on an event-by-event basis and calculated individually for each bunch in a train. The principle is illustrated in Figure 56. When running with a 50 ns bunch spacing, there are gaps between the filled bunches that can be used for baseline determination. Each such gap is used to calculate the baseline for the bunch that comes immediately after it. However, the method fails when running with a 25 ns bunch spacing because there is insufficient space left between the bunches. This makes the PreTrain method a more realistic alternative as it always works.

8.3 Calibration

In practice, the charge measurements from all standard PMT readout channels are summed up to form a total PMT charge. In the same way, the charge measurements from all fiber readout channels are summed

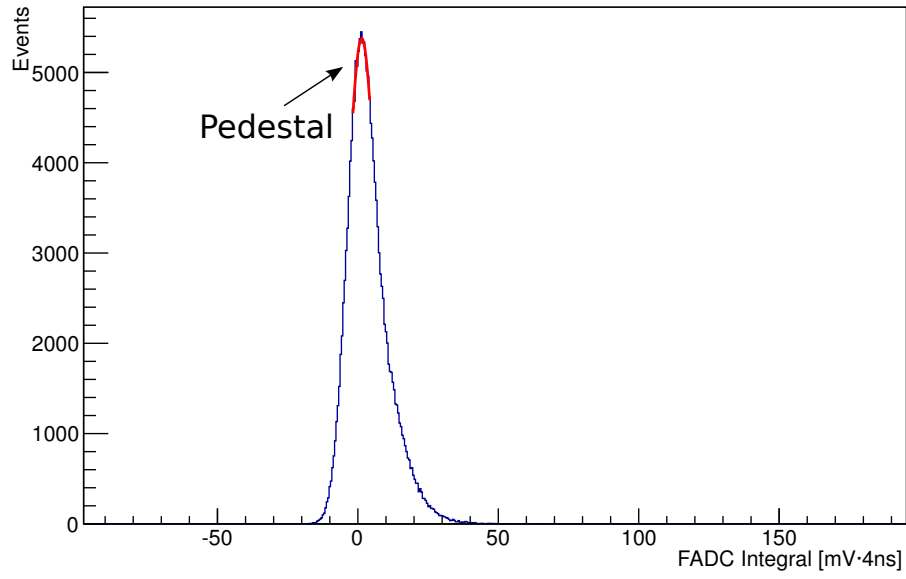


Figure 54: A charge spectrum. The events were recorded when there was no beam in the LHC, so only a pedestal distribution is visible. The same baseline was used for every event.

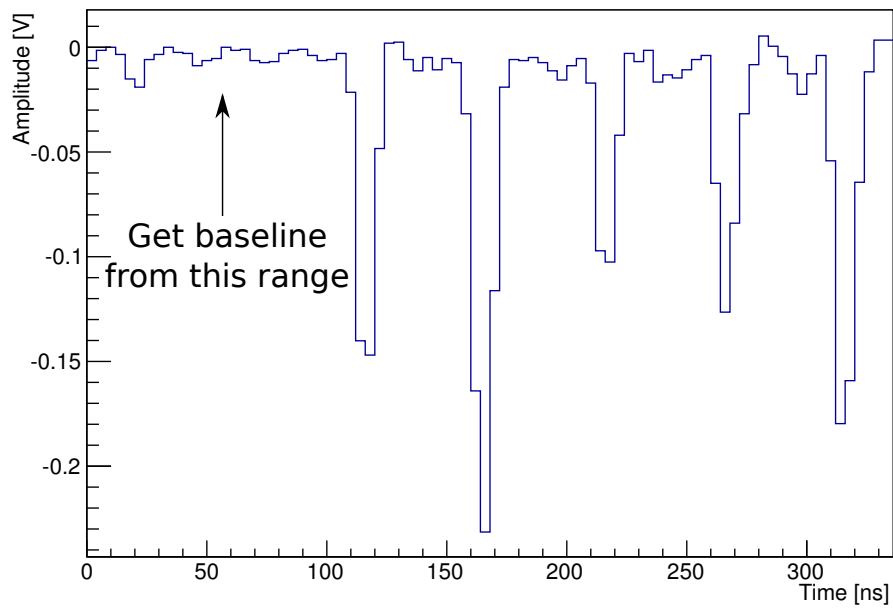


Figure 55: When using PreTrain baselines, the range before the first bunch in the train is used for the baseline calculation. The same baseline is then used for all bunches in the train.

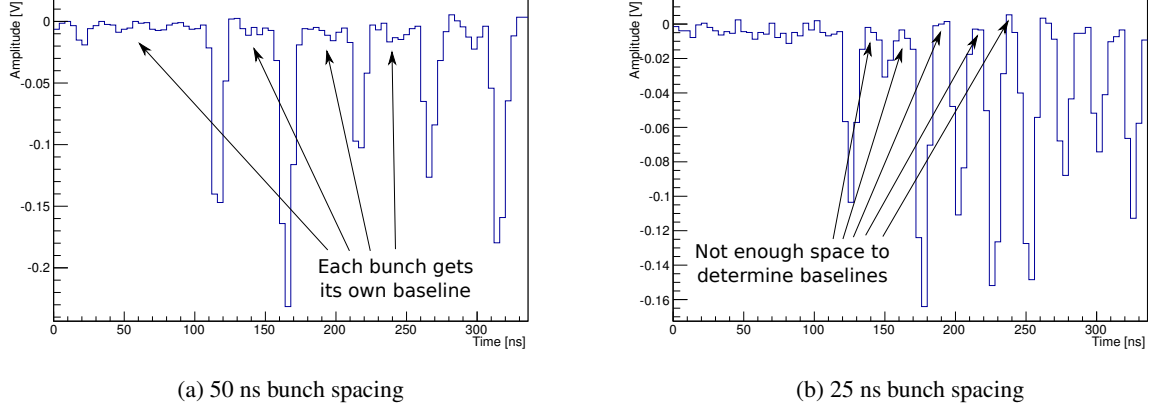


Figure 56: When using the Dynamic method, the baselines are calculated individually for each bunch. This is possible with a 50 ns bunch spacing (a), but the method fails with a 25 ns bunch spacing (b).

up to form a total fiber charge. The charge is integrated over the course of a lumiblock (LB) and then divided by the number of events in that LB. This yields the average charge per event $\langle Q \rangle_{\text{LB}}$, which is a per-LB quantity that is proportional to μ according to:

$$\mu = k \cdot \langle Q \rangle_{\text{LB}} \quad (166)$$

The constant of proportionality k is obtained in this study by calibrating against BCMV_EventOR, which was the ATLAS-preferred algorithm during 2012. Figure 57 shows a measurement of $\langle Q \rangle_{\text{LB}}$ for a single bunch using fibers. The BCM μ -measurement for the same bunch is shown alongside of it. It is clear that Eq. (166) holds.

The most accurate luminosity measurement is obtained by calibrating each bunch individually. A simpler option is to use a bunch-averaged calibration $\langle k \rangle_{\text{bunch}}$. This is done by calculating k for all n bunches within the FADC range and simply taking the average.

$$\langle k \rangle_{\text{bunch}} = \frac{\sum_{i=0}^n k_i}{n} \quad (167)$$

8.4 Systematic uncertainties

The systematic uncertainty associated with the charge measurement, denoted σ_{charge} , is evaluated bunch by bunch by fitting μ_{LUCID} as obtained from Eq. (166) to $\mu_{\text{BCM}_V\text{EventOR}}$. An example of such a fit is shown in Figure 58 (a). The bottom plot contains the Data/Fit ratio, which is projected onto the y-axis to yield the distribution shown in Figure 58 (b). The standard deviation of this distribution is taken to be a measurement of σ_{charge} . It should be noted that σ_{charge} does not depend on k .

There is also a systematic uncertainty associated with the run-to-run stability of the calibration. Essentially, different values of k are obtained depending on which run is used to calibrate. The calibration uncertainty, denoted σ_{calib} , is evaluated bunch by bunch by calculating k for a series of runs with identical running conditions. This yields a mean value $\langle k \rangle_{\text{run}}$ and a standard deviation $\sigma_{k,\text{run}}$. The uncertainty is then taken to be:

$$\sigma_{\text{calib}} = \frac{\sigma_{k,\text{run}}}{\langle k \rangle_{\text{run}}} \quad (168)$$

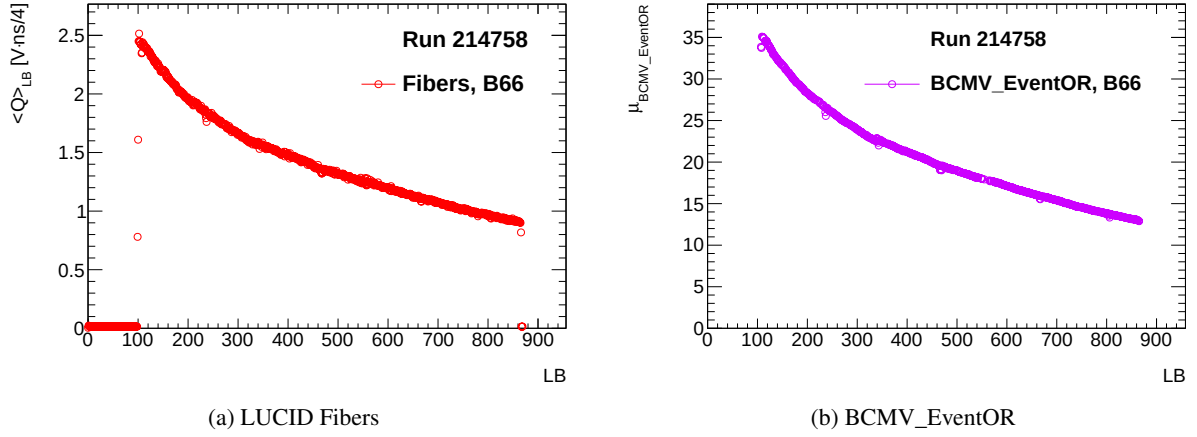


Figure 57: The charge measurement from the fiber readout channels (a) is compared to a luminosity measurement from BCM (b) as a function of LB. The data comes from bunch 66 in run 214758.

The longest available series of runs during which the LUCID running conditions did not change is listed in Section 8.6. It was recorded over the course of ten days. For each bunch, the slope of the k -values as a function of run number is statistically compatible with zero. Therefore, it is assumed that k does not change over time, and no systematic uncertainty is associated with the time stability of the calibration.

In case a bunch-averaged calibration is used, an additional systematic uncertainty σ_{bunch} associated with the difference between bunches must be introduced. Calculating k for each bunch within a run yields a mean value $\langle k \rangle_{\text{bunch}}$ and a standard deviation $\sigma_{k,\text{bunch}}$. The uncertainty for a given run is then taken to be:

$$\sigma_{\text{bunch}} = \frac{\sigma_{k,\text{bunch}}}{\langle k \rangle_{\text{bunch}}} \quad (169)$$

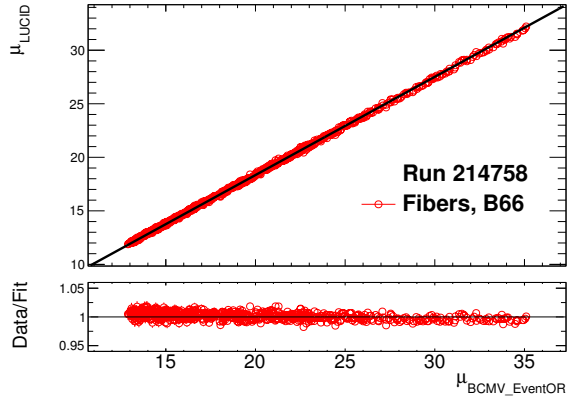
It is also possible to estimate the sum of σ_{charge} and σ_{bunch} by producing a plot equivalent to the one shown in Figure 58 (b), but instead of plotting only a single bunch, all the bunches are included. An example is shown in Figure 59. Bunches far from the average show up as separate peaks. The sum of σ_{charge} and σ_{bunch} is given by the standard deviation of the distribution. In general, σ_{bunch} is found to be much larger than both σ_{charge} and σ_{calib} . In runs with a 25 ns bunch spacing, it is too large for the method to be competitive with other luminosity measurements. Therefore, each bunch will be calibrated individually and σ_{bunch} will not be discussed further.

The charge and calibration systematic uncertainties are independent. If they are evaluated bunch by bunch, the total uncertainty σ_{tot} for the bunch in question is calculated by adding the two sources of uncertainty in quadrature.

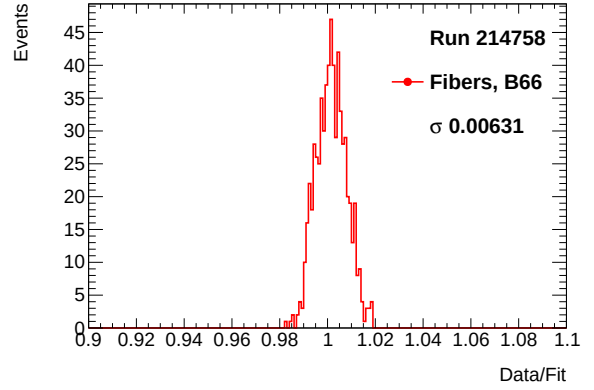
$$\sigma_{\text{tot}} = \sqrt{\sigma_{\text{charge}}^2 + \sigma_{\text{calib}}^2} \quad (170)$$

8.5 Runs with low HV

Figure 60 (a) shows a fit of $\langle Q \rangle_{\text{LB}}$ to BCM as obtained from standard PMTs under nominal conditions. The points in the Data/Fit ratio follow a slight downward arc. This is a saturation effect that grows as a function of μ . There is a limit to the amount of charge that the photomultipliers can deliver, and when this limit is exceeded, the PMT response ceases to be linear. Saturation can be avoided by lowering the high voltage supplied to the PMTs, thereby decreasing the signal amplification. Figure 60 (b) shows a run that was taken with a lowered HV. There is no sign of saturation.

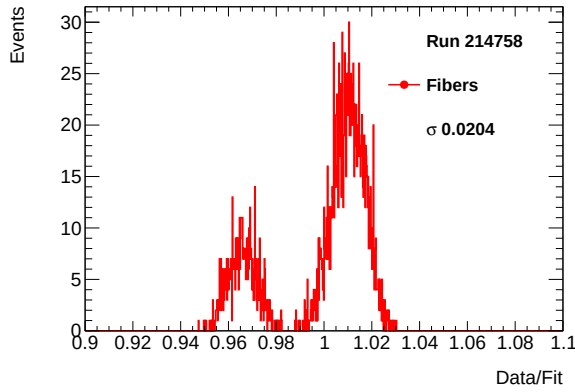


(a) LUCID fit to BCM

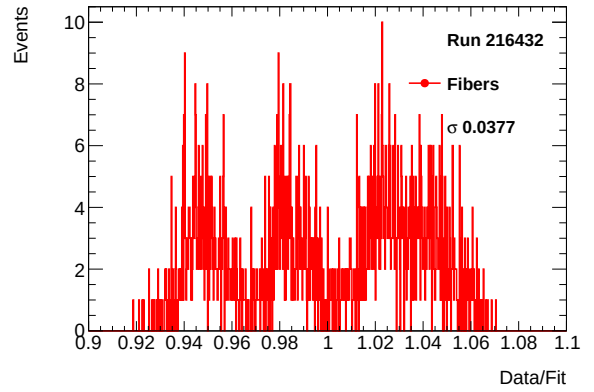


(b) Projection of the Data/Fit ratio

Figure 58: To evaluate the uncertainty of the charge measurement, the calibrated μ is fit to BCM (a). The Data/Fit ratio is then projected onto the y-axis, and the standard deviation of the resulting distribution (b) is taken to be a measurement of σ_{charge} . The data comes from bunch 66 in run 214758.



(a) Run with 50 ns bunch spacing



(b) Run with 25 ns bunch spacing

Figure 59: When using an average calibration, bunches far from the average show up as a separate peaks in the projection of the Data/Fit ratio. The standard deviation of the distribution corresponds to the sum of σ_{charge} and σ_{bunch} .

Event counting and hit counting algorithms are not affected by PMT saturation as long as the threshold level is placed in a range where the detector response is linear. Particle counting algorithms, on the other hand, have no concept of a threshold level. Saturation violates the linearity of Eq. (166) and leads to the wrong result. The HV should therefore be lowered until the detector response is linear over the entire μ -range of interest. The tradeoff is that lowering the HV leads to a smaller integrated charge and consequently a higher statistical uncertainty on the charge measurement. This shows up as an increased σ_{charge} . Finally, it should be noted that saturation is a problem only for standard PMTs. Fibers operate far from saturation because less Cherenkov light makes it to the photocathode.

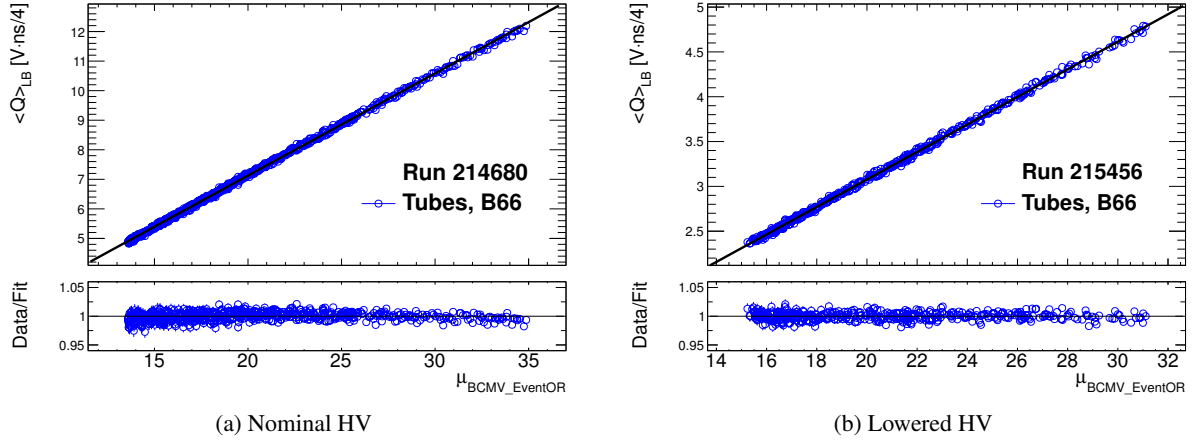


Figure 60: When running at nominal HV (a), the integrated PMT charge shows signs of saturation. The saturation effect disappears when running at a lowered HV (b).

8.6 Results

The following list of run numbers contains a series of runs with 50 ns bunch spacing during which LUCID did not change its running conditions. This list has been used to evaluate σ_{charge} and σ_{calib} . The local stream was set to trigger randomly on BCID 66.

- 214651
- 214680
- 214721
- 214758
- 214777
- 215027
- 215061
- 215063
- 215091

Tables 14 and 15 show σ_{charge} for the different baseline methods in fibers and standard PMTs, respectively. These numbers were obtained by taking the worst-case scenario out of all the runs considered. Tables 16 and 17 show σ_{calib} in fibers and standard PMTs, respectively.

Table 14: The table shows the worst-case σ_{charge} for BCID 66, 68, 70 and 72 in fibers as obtained using Static, PreTrain and Dynamic baselines in 50 ns runs.

σ_{charge}	B66	B68	B70	B72
Static	0.76%	0.57%	0.61%	0.63%
PreTrain	0.77%	0.70%	0.95%	0.67%
Dynamic	0.77%	0.76%	0.70%	0.70%

Table 16: The table shows σ_{calib} for BCID 66, 68, 70 and 72 in fibers as obtained using Static, PreTrain and Dynamic baselines in 50 ns runs.

σ_{calib}	B66	B68	B70	B72
Static	0.60%	0.63%	0.60%	0.70%
PreTrain	1.32%	0.91%	0.93%	0.45%
Dynamic	1.32%	0.56%	0.61%	0.36%

Table 15: The table shows the worst-case σ_{charge} for BCID 66, 68, 70 and 72 in standard PMTs as obtained using Static, PreTrain and Dynamic baselines in 50 ns runs.

σ_{charge}	B66	B68	B70	B72
Static	1.87%	1.61%	1.06%	1.50%
PreTrain	1.65%	1.04%	0.83%	1.05%
Dynamic	1.65%	1.23%	1.10%	1.35%

Table 17: The table shows σ_{calib} for BCID 66, 68, 70 and 72 in standard PMTs as obtained using Static, PreTrain and Dynamic baselines in 50 ns runs.

σ_{calib}	B66	B68	B70	B72
Static	1.50%	1.26%	0.84%	1.22%
PreTrain	1.28%	0.81%	0.76%	0.77%
Dynamic	1.28%	1.25%	1.52%	1.12%

The following two runs with 25 ns bunch spacing are available to cross-check the 50 ns systematics. The local stream was set to trigger randomly on BCID 301. Only every second BCID within the FADC range was used for the cross-check, retaining the same BCIDs relative to the triggering one that were populated also in the 50 ns runs.

- 216416
- 216432

Note that Dynamic baselines are not usable in this case. In the same manner as above, Tables 18 and 19 show σ_{charge} in fibers and standard PMTs, respectively. Tables 20 and 21 show σ_{calib} in fibers and standard PMTs, respectively. It should be noted that the value of σ_{calib} obtained for runs with a 25 ns bunch spacing is an unreliable estimate of the true uncertainty because of the small number of available runs. Nevertheless, it is consistent with the 50 ns result. The total uncertainty σ_{tot} for each baseline method is therefore calculated using the worst-case σ_{charge} and σ_{calib} out of the 50 ns and 25 ns runs. Tables 22 and 23 show σ_{tot} for fibers and standard PMTs, respectively. Since the local stream was set to trigger on different BCIDs during the 50 ns and 25 ns runs, σ_{tot} is quoted for bunch positions relative to the triggering BCID rather than absolute BCID numbers. Detailed data for each individual run can be found in Appendix A.

8.7 Conclusions

A study of particle counting using the LUCID detector has been presented. The method relies on measuring the total charge deposited by particles passing through the detector. This charge was shown to be directly proportional to the luminosity. The total uncertainty of the luminosity measurement is better than 1.6% for fiber readout channels and better than 2.1% for standard PMT readout channels using an algorithm that can be implemented online in 25 ns runs. During run II, LUCID will use particle counting alongside the standard algorithms.

Table 18: The table shows the worst-case σ_{charge} for BCID 301, 303, 305 and 307 in fibers as obtained using Static and PreTrain baselines in 25 ns runs.

σ_{charge}	B66	B68	B70	B72
Static	0.96%	0.86%	0.85%	0.86%
PreTrain	0.95%	0.90%	0.99%	0.86%

Table 19: The table shows the worst-case σ_{charge} for BCID 301, 303, 305 and 307 in standard PMTs as obtained using Static and PreTrain baselines in 25 ns runs.

σ_{charge}	B66	B68	B70	B72
Static	2.01%	1.28%	1.20%	1.37%
PreTrain	1.18%	1.02%	1.03%	0.97%

Table 20: The table shows σ_{calib} for BCID 301, 303, 305 and 307 in fibers as obtained using Static and PreTrain baselines in 25 ns runs.

σ_{calib}	B66	B68	B70	B72
Static	0.60%	0.18%	0.17%	0.56%
PreTrain	0.77%	0.29%	0.35%	0.67%

Table 21: The table shows σ_{calib} for BCID 301, 303, 305 and 307 in standard PMTs as obtained using Static and PreTrain baselines in 25 ns runs.

σ_{calib}	B66	B68	B70	B72
Static	1.63%	1.40%	1.03%	1.66%
PreTrain	0.08%	0.36%	0.12%	0.65%

Table 22: The table shows σ_{tot} for every second BCID after the triggering one in fibers as obtained using Static, PreTrain and Dynamic baselines.

σ_{tot}	Trig	Trig+2	Trig+4	Trig+6
Static	1.13%	1.07%	1.04%	1.11%
PreTrain	1.63%	1.28%	1.36%	0.95%
Dynamic	1.53%	0.94%	0.93%	0.79%

Table 23: The table shows σ_{tot} for every second BCID after the triggering one in standard PMTs as obtained using Static, PreTrain and Dynamic baselines.

σ_{tot}	Trig	Trig+2	Trig+4	Trig+6
Static	2.59%	2.13%	1.58%	2.24%
PreTrain	2.09%	1.32%	1.28%	1.30%
Dynamic	2.09%	1.75%	1.88%	1.75%

9 A search for sleptons and gauginos in final states with two leptons

9.1 Introduction

Early searches for supersymmetry at the LHC were designed to find gluinos and squarks. The reason for this choice is illustrated by Figure 61. The figure shows the cross-section of various SUSY production processes in 8 TeV pp collisions as a function of sparticle mass. For a given mass, the production cross-section of colored SUSY particles ($\tilde{g}$ and $\tilde{q}$) exceeds that of electroweak SUSY particles ($\tilde{\chi}^\pm$, $\tilde{\chi}^0$, $\tilde{\ell}$ and $\tilde{\nu}$) by several orders of magnitude. The ideal SUSY scenario is therefore one in which the lightest sparticles are colored. However, the early searches found no sign of new physics and instead put strong lower limits on the masses of gluinos and squarks. At the end of run I, gluinos below 1 TeV and squarks below 0.5 TeV were in general excluded [64].

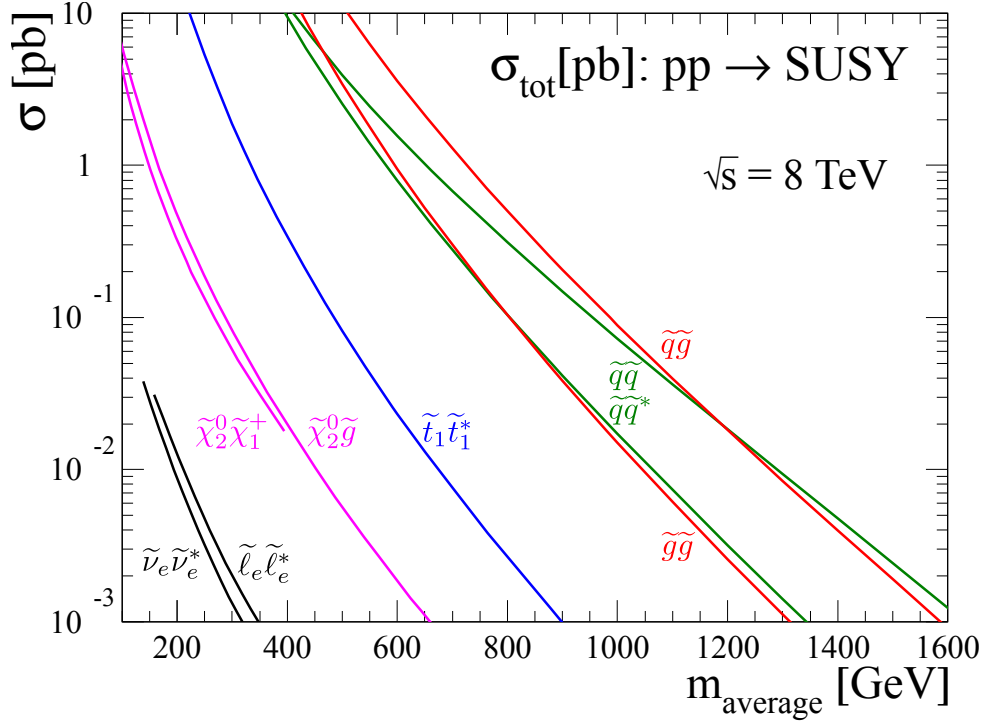


Figure 61: The figure shows the production cross-section of various pairs of sparticles in 8 TeV pp collisions as a function of the average mass of the produced sparticles.

The fact that colored SUSY particles must be heavy motivates searches for alternative scenarios in which the lightest sparticles are electroweak. This turns out to be the case in many GMSB scenarios [65–70]. The signature of electroweak supersymmetry production is missing energy from the LSP, little hadronic activity due to the absence of colored particles, and the presence of one or more leptons in the final state. This section describes a search for sleptons and gauginos in final states with exactly two oppositely charged leptons (where lepton means either electron or muon). It is based on 20.3 fb^{-1} of proton–proton collision data recorded with the ATLAS detector in 2012 at a center-of-mass energy of $\sqrt{s} = 8 \text{ TeV}$. The analysis has been published and is available in Ref. [3]. All figures in Sections 9.7 and 9.8 marked *ATLAS* are part of the official publication.

9.2 Signal models

The processes targeted by the analysis are shown in Figure 62. It is assumed that R-parity is conserved and that all sparticles except the LSP decay promptly. Slepton pair production leads to two-lepton final states because sleptons decay according to $\tilde{\ell} \rightarrow \ell \tilde{\chi}_1^0$ as illustrated in Figure 62 (a). Chargino pair production also leads to two-lepton final states. If sleptons are lighter than charginos, then charginos decay preferentially to sleptons which subsequently decay to leptons as in Figure 62 (b). If sleptons are heavier than charginos, then charginos are forced to decay via Standard Model W bosons as in Figure 62 (c). Both scenarios are considered in this analysis.

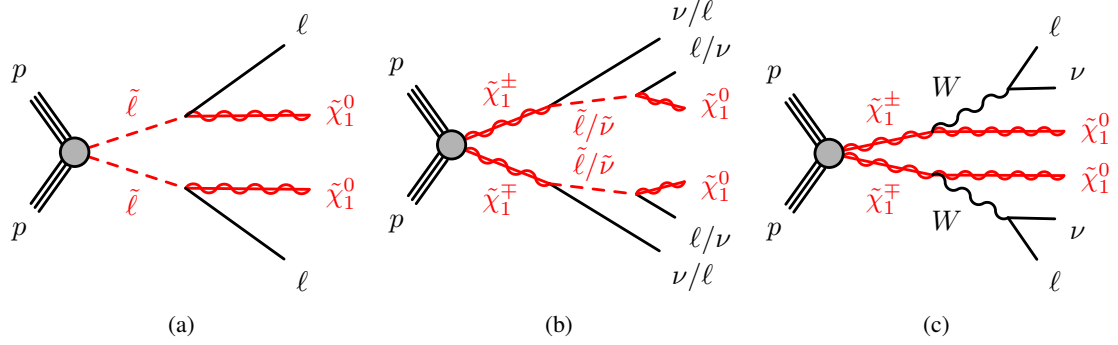


Figure 62: The search targets slepton pair production (a), chargino pair production followed by decay via sleptons (b) and chargino pair production followed by decay via Standard Model W bosons (c).

One way to search for SUSY is to make assumptions that reduce the number of free parameters in the symmetry breaking part of the MSSM Lagrangian and then scan the remaining parameter space. This is the approach taken in models such as mSUGRA, GMSB and the pMSSM. Another approach is to pick a single process, such as the production and subsequent decay of a pair of sleptons, and then try to set limits on the cross-section times branching ratio ($\sigma \times \text{BR}$) of that process. This can be done using so-called *simplified models*. The results of this analysis are interpreted both in terms of simplified models and within the pMSSM.

9.2.1 Simplified models

A simplified model is a type of minimalistic extension of the SM used to set limits on the $\sigma \times \text{BR}$ of a particular process. The only new particles that are added to the model are those that participate in the process in question. The only free parameters of the model are the masses of those particles. For example, the process shown in Figure 62 (a) can be described by a simplified model that, in addition to the SM, contains only sleptons $\tilde{\ell}$ and the LSP $\tilde{\chi}_1^0$. The $\sigma \times \text{BR}$ of the process is determined by the masses and mixing properties of the particles involved. Some simplifying assumptions are made about the mixing properties such that the only free parameters left are the masses $m_{\tilde{\ell}}$ and $m_{\tilde{\chi}_1^0}$. Limits set using the resulting simplified model can then be used to constrain any theory that contains the process in question. The drawback of using simplified models is that the limits are not completely generic. The assumptions made when constructing a simplified model may lead to the exclusion of a particular combination of sparticle masses that would be allowed by a more complex model.

9.2.1.1 Slepton pair production

This simplified model describes the process shown in Figure 62 (a). The LSP is assumed to be bino-like. Only selectron and smuon production is considered ($\tilde{\ell} = \tilde{e}, \tilde{\mu}$), with the selectron and smuon assumed to

be mass-degenerate. The reason for not including stau production is that the MSSM typically predicts $m_{\tilde{\tau}}$ to be significantly different from $m_{\tilde{\ell}}$, and any assumption about the position of the stau in the mass hierarchy would increase the model dependence of the limits. Excluding stau production lowers the $\sigma \times \text{BR}$ to two-lepton final states, ensuring that the limits remain conservative. The reason for not considering sneutrino production is that the analysis has been unable to demonstrate any sensitivity to such scenarios.

Only the case of a slepton NLSP is considered. This ensures that the sleptons decay according to $\tilde{\ell} \rightarrow \ell \tilde{\chi}_1^0$. It is worth noting that this process always leads to same-flavor final states, meaning ee and $\mu\mu$ but not $e\mu$. Another possible scenario is one in which the NLSP is a chargino, enabling the decay $\tilde{\ell} \rightarrow \nu \tilde{\chi}_1^\pm$. This signature, however, is already covered by the search for chargino pair production (which has a much higher cross-section, see Figure 61).

Production of both left-handed ($\tilde{\ell}_L^+ \tilde{\ell}_L^-$) and right-handed ($\tilde{\ell}_R^+ \tilde{\ell}_R^-$) sleptons is considered. Limits are set under the assumption that they are mass-degenerate. Since this assumption does not necessarily hold, limits are also set on $\tilde{\ell}_L^+ \tilde{\ell}_L^-$ and $\tilde{\ell}_R^+ \tilde{\ell}_R^-$ production separately. For illustrative purposes, Figure 63 shows the signal grid used to set limits on slepton pair production. A *signal grid* is a set of points in the plane spanned by the free parameters of the model in question. In the simplified model describing slepton pair production, the free parameters are the slepton and LSP masses. Limits are therefore set in the $m_{\tilde{\ell}}-m_{\tilde{\chi}_1^0}$ plane. The density of points is greater in the region where those limits are expected to be.

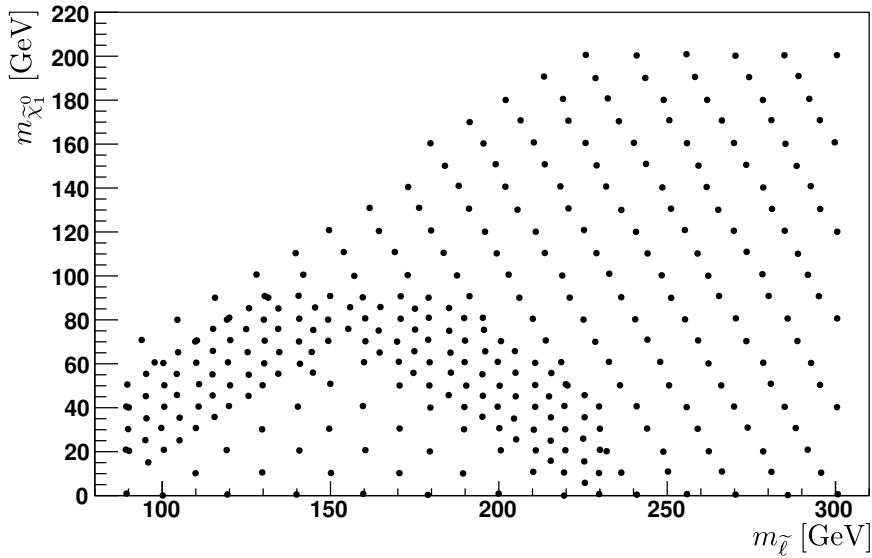


Figure 63: The signal grid used to set limits on slepton pair production.

9.2.1.2 Chargino pair production with light sleptons

This simplified model describes the process shown in Figure 62 (b). The charged sleptons and sneutrinos are assumed to be mass-degenerate. Their mass is placed halfway between the masses of the lightest chargino $\tilde{\chi}_1^\pm$ and the LSP $\tilde{\chi}_1^0$.

$$m_{\tilde{\ell}} = m_{\tilde{\nu}} = \frac{m_{\tilde{\chi}_1^\pm} + m_{\tilde{\chi}_1^0}}{2} \quad (171)$$

By fixing the slepton mass in this way, the model can be parametrized in terms of just $m_{\tilde{\chi}_1^0}$ and $m_{\tilde{\chi}_1^\pm}$. The chargino is assumed to be pure wino so that it decays only into left-handed sleptons. The branching ratio is 50/50 to $\tilde{\ell}$ and $\tilde{\nu}$. Each slepton flavor (including stau) occurs with equal probability. Since

the sleptons are not correlated in flavor, both same-flavor (SF) and different-flavor (DF) final states are equally probable.

9.2.1.3 Chargino pair production without light sleptons

This simplified model describes the process shown in Figure 62 (c). It is identical to the model described in Section 9.2.1.2, except light sleptons are not included. This forces the charginos to decay to Standard Model W bosons with a 100% branching ratio. Each W boson can then decay either leptonically or hadronically (i.e. to quarks). The branching ratio to two-lepton final states is approximately 10.5%, which makes this scenario considerably more difficult to discover than the one depicted in Figure 62 (b).

9.2.2 pMSSM

The pMSSM has 19 free parameters (see Table 10). Nine parameters mainly affect colored SUSY. They are the gluino mass M_3 , the six squark masses and the trilinear couplings A_t and A_b . In this search, those parameters are set such that colored sparticles become very heavy with masses > 2 TeV. This causes the colored SUSY sector to become mostly decoupled from the physics reachable at the LHC. The mass of the stop is still important, however, as it contributes significantly to the mass of the lightest Higgs boson. This is exploited in order to allow the other pMSSM parameters to float freely — For a given set of parameters, the stop sector can always be tuned in order produce a Higgs mass close to 125 GeV.

The A^0 is decoupled from the rest of the phenomenology by setting $M_A = 500$ GeV. Two scenarios are then considered, similarly to what was done for the simplified models. The first scenario contains light sleptons. The five pMSSM parameters relevant to the slepton sector are the two left-handed slepton masses, the two right-handed slepton masses and the trilinear coupling A_τ . Left-handed sleptons are assumed to be beyond reach of LHC physics, while the right-handed sleptons are all assumed to be mass-degenerate with a mass midway between the masses of the two lightest neutralinos.

$$m_{\tilde{\ell}_R} = \frac{m_{\tilde{\chi}_1^0} + m_{\tilde{\chi}_2^0}}{2} \quad (172)$$

Finally, A_τ is set to zero. These assumptions leave no free parameters in the slepton sector of the pMSSM. The remaining parameters are $\tan\beta$, μ , M_1 and M_2 . They govern the gaugino sector of the pMSSM. In this search, $\tan\beta$ is set to 6. The remaining free parameters are then scanned by choosing three values of M_1 and generating a signal grid in the μ - M_2 plane for each choice. The M_1 values 100 GeV, 140 GeV and 250 GeV are used.

Choosing the left-handed sleptons to be heavy is done for the following reasons. Left-handed sleptons carry both weak isospin and weak hypercharge while right-handed sleptons carry only weak hypercharge. The production cross-section for left-handed sleptons is therefore higher. By only including right-handed sleptons in the model, the slepton production cross-section is minimized and the resulting limits become maximally conservative. Secondly, the wino component of the gauginos couples only to left-handed sleptons. By making the left-handed sleptons heavy, this decay channel is no longer available, which causes the gauginos to decay via Standard Model gauge bosons instead. As explained in Section 9.2.1.3, this reduces the branching ratio to two-lepton final states, which again results in limits that are maximally conservative.

Finally, one should keep in mind that the assumptions made about the mass of the sleptons in Eq. (172) are not necessarily true. When in competition with other decays, lighter sleptons would enjoy an increased branching ratio which, in turn, would increase the branching ratio to two-lepton final states. Heavier sleptons would lead to the opposite result.

In the second scenario, all sleptons are made heavy. This forces all gauginos to decay via Standard Model gauge bosons. The values of $\tan\beta$ and M_1 are set to 10 and 50 GeV, respectively, in order to keep

the relic dark matter density [71] below cosmological bounds across the entire signal grid. The grid is then generated by scanning the μ – M_2 plane.

9.3 Event reconstruction and selection

All events are subject to a set of basic selection criteria. They must fall within a Luminosity Block that has been flagged as being good for physics, they must be free from cosmic muon candidates, and none of the relevant subdetectors must indicate problems with data quality. The events must also contain a primary vertex, which is defined as a proton–proton collision vertex that is associated with at least five tracks with transverse momentum $p_T > 400$ MeV. If an event contains several vertices that satisfy this condition, the one with the greatest $\sum p_T^2$ is considered to be the primary one. Events that pass this basic selection are then subject to additional cuts based on the number of leptons and jets, overlap removal, trigger requirements and other variables that are defined in this section.

9.3.1 Object definitions

The ATLAS detector measures various properties of the particles created during pp collisions and identifies electron, muon, tau and jet candidates. The detector performance groups then define sets of criteria that are used to select which objects should be considered in physics analyses. Each set is given a name, such as `medium` or `tight`, depending on its stringency. Since the criterion definitions are often quite technical, it is customary to refer to the performance group documentation rather than quoting them outright.

This section describes which criteria the object candidates should satisfy in order to be used in the analysis. Each candidate is subject to two selections. If it passes a looser set of requirements, it is called a *baseline* object. Such objects are used for overlap removal (see Section 9.3.2) and in the Matrix Method (see Section 9.5.4). A baseline object that in addition passes a stricter set of requirements is called a *signal* object. Such objects are used in the definition of the signal regions (see Section 9.4).

9.3.1.1 Electrons

Electron candidates are reconstructed by matching tracks in the inner detector to energy deposits in the calorimeters. Baseline electrons are required to pass the `medium` criteria defined in Ref. [72], have $p_T > 10$ GeV and satisfy $|\eta| < 2.47$.

Signal electrons are required to pass the `tight` criteria defined in Ref. [72]. They must also fulfill two isolation requirements. The first isolation requirement is defined by identifying all tracks associated with the primary vertex that have $p_T > 400$ MeV and reside within $\Delta R < 0.3$ of the electron. The scalar sum of their transverse momenta is required to be less than 16% of the electron p_T . The second isolation requirement is defined by identifying all calorimeter clusters within $\Delta R < 0.3$ of the electron (regardless of their energy or which primary vertex they are associated with). The scalar sum of their transverse energies is required to be less than 18% of the electron p_T . Finally, the longitudinal impact parameter must satisfy $z_0 \sin \theta < 0.4$ mm and the transverse impact parameter d_0 must be within five standard deviations of the primary vertex. This helps reject electrons originating from the decay of long-lived particles such as b -quarks.

9.3.1.2 Muons

Muon candidates are reconstructed by matching tracks in the inner detector to tracks in the muon spectrometer [73]. Baseline muons are required to pass loose selection criteria, have $p_T > 10$ GeV and satisfy $|\eta| < 2.50$. For signal muons, the cut on η is tightened to $|\eta| < 2.40$ due to limited trigger coverage. Signal muons must also fulfill an isolation requirement similar to the one applied for electrons. All

tracks associated with the primary vertex that have $p_T > 1$ GeV and reside within $\Delta R < 0.3$ of the muon are identified. The scalar sum of their transverse momenta is required to be less than 12% of the muon p_T . Finally, the longitudinal impact parameter must satisfy $z_0 \sin \theta < 1$ mm and the transverse impact parameter d_0 must be within three standard deviations of the primary vertex.

9.3.1.3 Taus

A tau can decay either leptonically to an electron or muon (in which case it is treated as an electron or muon rather than a tau) or hadronically to quarks. Since this analysis only considers ee , $\mu\mu$ and $e\mu$ final states, it does not explicitly look for hadronic taus. However, other SUSY searches within ATLAS do [74, 75]. Events containing hadronic taus are therefore rejected in order to remain orthogonal to such analyses. This ensures that the results of the analyses can be statistically combined.

Hadronic tau candidates are identified using a Boosted Decision Tree [76]. Baseline taus are required to pass the loose criteria defined in Ref. [76], have $p_T > 20$ GeV and satisfy $|\eta| < 2.47$. Signal taus must in addition pass the medium criteria.

9.3.1.4 Jets

Jet candidates are reconstructed from calorimeter clusters using the anti- k_t algorithm [77, 78] with a distance parameter of $R = 0.4$. Baseline jets must have $p_T > 20$ GeV and satisfy $|\eta| < 4.5$. Three different categories of signal jets are then defined. The definitions are summarized in Table 24. Jets with $2.4 < |\eta| < 4.5$ are required to have $p_T > 30$ GeV. A jet satisfying these requirements is placed in the *forward* category. Jets with $|\eta| < 2.4$ are checked against a multivariate b -jet identification algorithm called MV1 [79]. The identification efficiency of the algorithm is 80% for true b -jets, while the misidentification rate (i.e. the probability of incorrectly identifying a jet as a b -jet) is 30% for c -jets and 4% for u -jets, d -jets, s -jets and gluon jets. If a jet is identified as a b -jet, it is placed in the b -jet category. Jets that are *not* identified as b -jets are placed in the *central* category. Central jets with $p_T < 50$ GeV are additionally required to have a nonzero jet vertex fraction (JVF). The JVF is obtained by identifying all tracks within the jet that originate from the primary vertex and dividing the sum of their momenta by the total momentum of the jet. When running at high luminosity, stray particles from many different pp interaction vertices are sometimes classified as jets. The JVF cut is designed to reject such jets. The b -tagging algorithm rejects them in a similar manner, so no JVF cut is required in this case. The reason no JVF cut is applied for forward jets, as well as the reason why forward jets are not checked against the b -tagging algorithm, is that there is insufficient tracker coverage in the region $|\eta| > 2.4$.

Table 24: Summary of the three signal jet category definitions.

	Forward jets	b -jets	Central jets
p_T	> 30 GeV	> 20 GeV	> 20 GeV
$ \eta $	$[2.4, 4.5]$	< 2.4	< 2.4
b -tagged	–	Yes	No
JVF	–	–	> 0 if $p_T < 50$ GeV

9.3.2 Overlap removal

The ATLAS classification of tracks is not unambiguous. For example, it is possible for a track to be classified both as an electron candidate and as a jet candidate. An overlap removal procedure is applied in order to resolve such ambiguities. The procedure also takes care of rejecting individual objects that

are very close to each other, as this causes them to be badly reconstructed. Pairs of tracks with small ΔR are identified. One or both tracks are then removed from the event as described in Table 25. Finally, the procedure identifies pairs of same-flavor opposite-sign (SFOS) electrons and muons with an invariant mass $m_{\ell\ell} < 12$ GeV and removes them from the event. This excludes particles from quarkonia resonances and photon conversions. Only events that contain exactly two leptons with $m_{\ell\ell} > 20$ GeV after overlap removal are kept for analysis, where the final $m_{\ell\ell}$ cut is applied in order to reject uninteresting physics from low-mass resonances.

Table 25: Description of the overlap removal procedure. Objects that are removed are not considered in subsequent steps. All objects that satisfy the baseline definition are used. The one exception to this rule is step 10, which uses signal taus.

Step	Pair identified	Object removed	Comment
1	$\Delta R(e, e) < 0.05$	Electron with lowest E_T	Removes cases where one track is matched to many clusters in the calorimeter.
2	$\Delta R(j, e) < 0.2$	Jet	Removes electrons duplicated as jets.
3	$\Delta R(\tau, e) < 0.2$	Tau	Removes electrons duplicated as taus.
4	$\Delta R(\tau, \mu) < 0.2$	Tau	Removes muons duplicated as taus.
5	$\Delta R(j, e) < 0.4$	Electron	Removes electrons within jets.
6	$\Delta R(j, \mu) < 0.4$	Muon	Removes muons within jets.
7	$\Delta R(\mu, e) < 0.01$	Electron and muon	Muons can emit bremsstrahlung photons that are misidentified as electrons. Both objects are likely to be badly reconstructed.
8	$\Delta R(\mu, \mu) < 0.05$	Both muons	Both muons are likely to be badly reconstructed.
9	$m_{\ell\ell}^{\text{SFOS}} < 12$ GeV	Both leptons	Removes quarkonia and photon conversions.
10	$\Delta R(j, \tau_{\text{signal}}) < 0.2$	Jet	Removes taus duplicated as jets.

9.3.3 Triggers

As explained in Section 4.8, ATLAS uses a set of triggers to identify events that contain interesting physics. This analysis selects events from the ATLAS dataset by requiring that certain di-lepton triggers have fired. A di-lepton trigger is a trigger that detects the presence of at least two leptons. At the trigger stage, the analysis is divided into three different *channels* depending on the flavor of leptons present in the final state. The ee channel consists of events with exactly two signal electrons. In the same way, the $\mu\mu$ channel consists of events with exactly two signal muons. Finally, the $e\mu$ channel consists of events with exactly one signal electron and exactly one signal muon. This section motivates the choice of triggers for each channel.

There are two types of di-lepton triggers. *Symmetric* triggers require the presence of two leptons with p_T above a given threshold. The same p_T threshold is imposed on both leptons. *Asymmetric* triggers require the presence of one lepton with p_T above a higher threshold and another lepton with p_T above a lower threshold. The efficiency of a trigger is in general a function of the p_T of the leptons. The efficiency is low for leptons just above threshold, but then rises rapidly with increasing lepton p_T (this is called the *turn-on* region) until it stabilizes at some constant value a few GeV above the threshold (this is the *plateau* region). Triggers are therefore considered reliable only for leptons with p_T somewhat above threshold. Finally, triggers can be *prescaled* in which case only a subset of the events that pass the trigger are actually saved. In order to retain as much data as possible, the strategy used in this analysis is to select events based on the lowest- p_T unprescaled symmetric and asymmetric triggers available for

each channel.

Table 26 lists the lowest unprescaled triggers available for each channel. Also listed are the *Level 1 seeds* of each trigger. The seed is the Level 1 trigger that initiates the trigger decision. For example, only events that pass the L1_2MU10 trigger (which checks for the presence of two muons with $p_T > 10$ GeV at Level 1) are considered for the EF_2mu13 trigger (which checks for the presence of two muons with $p_T > 13$ GeV at the Event Filter). To save the reader some confusion, it is perhaps worth taking an example to explain what the (rather cryptic) trigger names mean. EF_2e12Tvh_loose1 is an Event Filter trigger that requires two electrons with $p_T > 12$ GeV. The T signifies that the p_T threshold of the trigger is within the turn-on of the seed, so lower efficiencies are to be expected close to threshold. The vh denotes that the trigger imposes a veto on the *hadronic leakage* of the triggering electrons. Hadronic leakage is a measure of the amount of energy deposited by the electrons in the hadronic calorimeter. Lastly, loose1 means that the electrons must satisfy some variety of "loose" tracking quality criteria. Some muon triggers contain EFFS, which means Event Filter Full Scan. This signifies that the Event Filter performs a full pattern recognition using all the information available from the muon spectrometer rather than reusing information from trigger decisions upstream [80].

Table 26: List of the lowest unprescaled triggers available for each channel. Also shown are the corresponding Level 1 seeds and the p_T thresholds above which the triggers are considered to be reliable. The leading (greater p_T) and subleading electrons are denoted e_1 and e_2 , with a similar notation for muons. Since no symmetric triggers exist for the $e\mu$ channel, two asymmetric ones are used instead.

Trigger name	Level 1 seed	Region of reliability
EF_2e12Tvh_loose1	L1_2EM10VH	$p_T^{e_1} > 14$ GeV, $p_T^{e_2} > 14$ GeV
EF_e24vh_medium1_e7_medium1	L1_EM18VH	$p_T^{e_1} > 25$ GeV, $p_T^{e_2} > 8$ GeV
EF_2mu13	L1_2MU10	$p_T^{\mu_1} > 14$ GeV, $p_T^{\mu_2} > 14$ GeV
EF_mu18_tight_mu8_EFFS	L1_MU15	$p_T^{\mu_1} > 18$ GeV, $p_T^{\mu_2} > 8$ GeV
EF_e12Tvh_medium1_mu8	L1_EM10VH_MU6	$p_T^e > 14$ GeV, $p_T^\mu > 8$ GeV
EF_mu18_tight_e7_medium1	L1_MU15	$p_T^e > 8$ GeV, $p_T^\mu > 18$ GeV

Figure 64 shows the extent of the p_T phase space that is covered by the triggers. In regions where the coverage of two triggers overlap, there are two options. The first option is to select events based on whichever trigger is more efficient. The second option is to use the logical OR of the two triggers. This option is only motivated if the gain in efficiency is significant, as it complicates the calculation of trigger weights (see Section 9.3.3.2). Table 27 shows the choice of trigger for each trigger region. The choice can be understood as follows. In the $\mu\mu$ channel, region A corresponds to both muons being above the high- p_T leg of the asymmetric trigger. The asymmetric trigger is more efficient than the symmetric one in this region, because the asymmetric trigger only requires one of the muons to fire the Level 1 seed. Either muon can fulfill this requirement, while for the symmetric trigger both muons are required to fire the seed. In region B, the efficiency of the asymmetric trigger drops as only the leading muon has sufficient p_T to match the mu18_tight part of the trigger. The OR of the symmetric and asymmetric triggers is used in this region. In the ee channel, one might expect the asymmetric trigger to be more efficient in the overlap region for the same reasons as in the $\mu\mu$ channel. However, the asymmetric di-electron trigger works differently from the di-muon one. Only the leading electron is allowed to fire the e24vh_medium1 part of the trigger at the Event Filter, regardless of the p_T of the subleading electron. This, combined with the stricter track quality criteria of the asymmetric trigger, makes the symmetric trigger more efficient. In the $e\mu$ channel overlap region, the trigger whose seed has lower p_T thresholds

is used. All other regions of phase space are free from overlaps in trigger coverage.

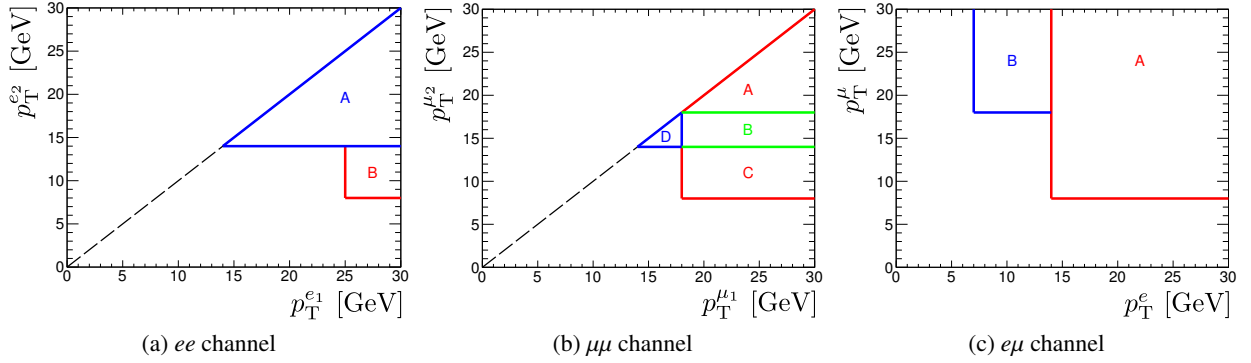


Figure 64: The figure shows the trigger regions used in the ee (a), $\mu\mu$ (b) and $e\mu$ (c) channels.

Table 27: Triggers used in the regions shown in Figure 64.

Region	Trigger used
ee channel A	EF_2e12Tvh_loose1
ee channel B	EF_e24vh_medium1_e7_medium1
$\mu\mu$ channel A	EF_mu18_tight_mu8_EFFS
$\mu\mu$ channel B	EF_mu18_tight_mu8_EFFS EF_2mu13
$\mu\mu$ channel C	EF_mu18_tight_mu8_EFFS
$\mu\mu$ channel D	EF_2mu13
$e\mu$ channel A	EF_e12Tvh_medium1_mu8
$e\mu$ channel B	EF_mu18_tight_e7_medium1

9.3.3.1 Trigger matching

Every object in an event is reconstructed twice. The first reconstruction is done online by the Event Filter. The resulting *online objects* are used to check if any trigger conditions are satisfied. If at least one trigger is fired, the event is saved to disk. The second reconstruction is done offline using more sophisticated reconstruction algorithms. Since the offline reconstruction is more accurate than the online one, all physics analyses are based on the resulting *offline objects*. In order to make sure that the Event Filter did not make any mistakes, and to make sure that the objects used by the analyses are indeed the ones that fired the trigger, the offline objects are *matched* against the online objects. This procedure is called trigger matching. It is applied only to real ATLAS data (as opposed to data from Monte Carlo simulation).

The offline objects are taken to be the signal objects defined in Section 9.3.1 after overlap removal. Matching is done by calculating the ΔR between those objects and the online objects that fired the trigger. Two objects are considered matched if they are within $\Delta R < 0.15$ of each other. The general trigger matching problem can be expressed in terms of graph theory. The offline and online objects make up a bipartite graph with edges connecting matched objects. A match of n tracks is considered successful if it is possible to find a subgraph of n offline objects and n online objects that contains a *perfect matching*. This is illustrated in Figure 65. If $n \leq 2$, as is always the case in this analysis, there exists a simple criterion for determining whether or not a perfect matching can be found. The match is considered

successful if at least n of the available offline objects can be matched to an online object and at least n of the online objects can be matched to an offline object.

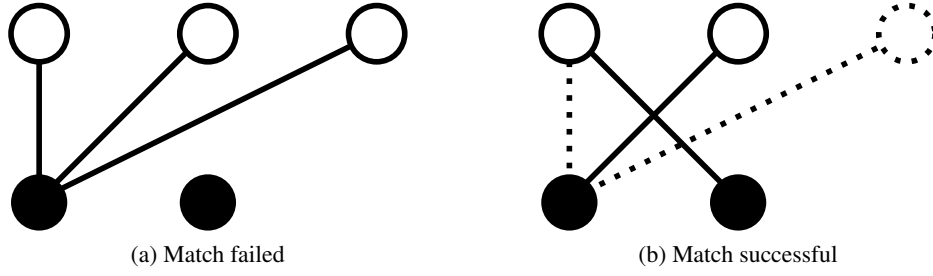


Figure 65: The offline objects (filled circles) and online objects (hollow circles) can be represented as a bipartite graph with edges connecting matched objects. A typical problem is the matching of two offline leptons to a di-lepton trigger. Graph (a) fails the trigger matching since only one of the offline objects could be matched. It does not matter that it was matched to three different online objects. In graph (b), the second offline lepton is also matched. This graph passes the trigger matching. An example of a subgraph containing a perfect matching is indicated. Dashed nodes and edges are excluded from the subgraph.

Matching to a symmetric di-lepton trigger is implemented as a match of two leptons to the corresponding single-lepton trigger. For example, a trigger match to `EF_2mu13` is done by matching two leptons to `EF_mu13`. This is a technical detail, as matching two leptons to the di-lepton trigger would be equivalent. Matching to an asymmetric di-electron or di-muon trigger is a two-step process. First, a match of two leptons is done to the full di-lepton trigger. Then, a match of one lepton is done to the single-lepton trigger that corresponds to the high- p_T leg of the di-lepton trigger. For example, a trigger match to `EF_e24vh_medium1_e7_medium1` is done by first matching two electrons to `EF_e24vh_medium1_e7_medium1` and then matching one electron to `EF_e24vh_medium1`. The second step is necessary because it is possible for both leptons to fire the low- p_T leg of the trigger without necessarily firing the high- p_T leg. Matching the single-lepton trigger separately guarantees that the reconstructed objects are the same as those that fired the trigger at the Event Filter. Finally, matching to an electron-muon trigger is done by matching the electron to the corresponding single-electron trigger and the muon to the corresponding single-muon trigger. For example, a trigger match to `EF_e12Tvh_medium1_mu8` is done by matching the electron to `EF_e12Tvh_medium1` and the muon to `EF_mu8`.

9.3.3.2 Trigger reweighting

The triggers are not 100% efficient. An event that should in principle fire a certain trigger might, due to misreconstructed objects, fail to fire that trigger in practice. This trigger inefficiency must be accounted for when producing events using Monte Carlo simulation. There are two principal approaches to the problem. The first approach is to rely on the detector simulation to properly model the response of the triggers. Monte Carlo events that fail the trigger simulation are then rejected just as they would be in real data. The second approach is to measure the trigger efficiencies in data. Each Monte Carlo event is then weighted by the appropriate trigger efficiency. This method is called trigger reweighting. Since a data-driven approach is expected to yield more accurate results than an approach based purely on Monte Carlo simulation, trigger reweighting is typically the preferred method. Since no Monte Carlo events are rejected, it also results in smaller statistical uncertainties. This analysis therefore uses trigger reweighting to account for the inefficiencies of the triggers.

The trigger reweighting procedure is based on the assumption that a di-lepton trigger efficiency can be factorized in terms of single-lepton trigger efficiencies. The simplest situation arises in case of symmetric triggers or asymmetric triggers for which each lepton is checked against only one of the trigger legs. The asymmetric triggers that fall into this category are the $e\mu$ triggers, the asymmetric ee trigger (recall that the high- p_T leg of this trigger is only checked against the leading electron regardless of the subleading electron p_T) and the asymmetric trigger used in $\mu\mu$ region C (because in this region, only the leading muon has sufficient p_T to fire the high- p_T leg of the trigger). The probability $P(D)$ to fire these di-lepton triggers is given by

$$P(D) = P(T_1)P(T_2) \quad (173)$$

where $P(T_1)$ and $P(T_2)$ are the probabilities for the leading and subleading leptons to fire their respective single-lepton triggers.

The situation is more complicated in case of asymmetric triggers for which each lepton is allowed to fire either part of the trigger. The asymmetric trigger used in $\mu\mu$ region A is the only trigger that falls into this category. The probability $P(D)$ is now given by

$$\begin{aligned} P(D) = & P(H_1)P(H_2) \cdot [P(L_1|H_1) + [1 - P(L_1|H_1)] \cdot P(L_2|H_2)] \\ & + P(H_1) \cdot [1 - P(H_2)] \cdot P(L_2|!H_2) \\ & + P(H_2) \cdot [1 - P(H_1)] \cdot P(L_1|!H_1) \end{aligned} \quad (174)$$

where $P(H_1)$ is the probability for the leading muon to fire the high- p_T leg of the trigger, $P(L_1)$ is the probability for the leading muon to fire the low- p_T leg of the trigger, $P(L_1|H_1)$ is the conditional probability for the leading muon to fire the low- p_T leg of the trigger given that it also fired the high- p_T leg of the trigger and $P(L_1|!H_1)$ is the conditional probability for the leading muon to fire the low- p_T leg of the trigger given that it did not fire the high- p_T leg of the trigger. An analogous notation is used for the subleading muon, with $P(H_2)$ denoting the probability for the subleading muon to fire the high- p_T leg of the trigger and so on. The first line of Eq. (174) represents the case where both muons fire the high- p_T leg of the trigger. In this case, either muon can fire the low- p_T leg. The second line represents the case where the leading muon fires the high- p_T leg and the subleading muon fires the low- p_T leg. Finally, the third line represents the case where the subleading muon fires the high- p_T leg and the leading muon fires the low- p_T leg.

The last remaining case is $\mu\mu$ region B, where the logical OR of the symmetric and asymmetric triggers is used. The probability $P(D)$ is given by

$$\begin{aligned} P(D) = & P(S_1)P(S_2) \\ & + P(H_1) \cdot [1 - P(S_2)] \cdot P(L_2|!S_2) \end{aligned} \quad (175)$$

where $P(S_1)$ and $P(S_2)$ are the probabilities for the leading and subleading muons to fire the legs of the symmetric trigger, $P(H_1)$ is the probability for the leading muon to fire the high- p_T leg of the asymmetric trigger and $P(L_2|!S_2)$ is the conditional probability for the subleading muon to fire the low- p_T leg of the asymmetric trigger given that it did not fire the symmetric trigger. Equation (174) takes advantage of several simplifying circumstances to eliminate terms that would otherwise appear in the expression. As shown in Figure 64 (b), the phase space of region B is such that the subleading muon has $p_T < 18$ GeV and therefore can not fire the high- p_T leg of the asymmetric trigger. Only the leading muon can do this. Furthermore, the asymmetric trigger imposes both a higher p_T cut and tighter track quality criteria on the leading muon than the symmetric trigger does. This implies that it is not possible for the leading muon to pass the high- p_T leg of the asymmetric trigger while failing the symmetric one. Consequently, the only events recovered by adding the asymmetric trigger to the symmetric one are those in which the leading muon passes the high- p_T leg of the asymmetric trigger and the subleading muon passes the low- p_T leg

of the asymmetric trigger while failing the symmetric one. This case is represented by the second line of Eq. (174). The first line represents the case where both muons fire the symmetric trigger.

All the probabilities discussed in this section are measured in data using the *Tag & Probe* method. In order to measure the single-electron efficiencies, a sample of $Z \rightarrow ee$ events in which at least one of the electrons was matched the lowest unprescaled single-electron trigger is selected from the ATLAS dataset. The electron that fired the trigger is called the *tag*. The other electron in the event is called the *probe*. Since the electrons originate from a Z boson, they both have high p_T and are therefore expected to fire all the relevant triggers. The efficiency of a particular trigger is given by

$$\epsilon_{\text{Trig}} = \frac{N_{\text{Trig}}^{\text{Pass}}}{N_{\text{Trig}}^{\text{Tot}}} \quad (176)$$

where $N_{\text{Trig}}^{\text{Pass}}$ is the number of probes that were matched to the trigger and $N_{\text{Trig}}^{\text{Tot}}$ is the total number of probes. Conditional probabilities are measured by considering only the subset of probes that did or did not pass the relevant trigger. Each event in the sample is processed twice, so that both electrons are given a chance to be the tag. An analogous method is then used to measure the single-muon trigger efficiencies based on a sample of $Z \rightarrow \mu\mu$ events. The electron trigger efficiencies are parametrized in terms of the electron p_T and η . Muon trigger efficiencies are found to depend only weakly on p_T and are instead parametrized in terms of η and ϕ . All trigger efficiencies are measured separately in the barrel and end-cap regions of ATLAS. Table 28 shows the approximate efficiencies of the single-lepton triggers used in the analysis.

Table 28: The table shows the efficiencies of various single-lepton triggers for leptons with p_T above the specified threshold. The values are for illustrative purposes only. They were evaluated based on data from Period A (meaning early 2012 running). Other run periods may have other trigger efficiencies.

Trigger name	Threshold	ϵ_{barrel}	$\epsilon_{\text{end-cap}}$
EF_e12Tvh_loose1	$p_T > 14 \text{ GeV}$	0.99	0.97
EF_e24vh_medium1	$p_T > 25 \text{ GeV}$	0.96	0.91
EF_e7_medium1	$p_T > 10 \text{ GeV}$	0.97	0.93
EF_mu13	$p_T > 14 \text{ GeV}$	0.72	0.88
EF_mu18_tight	$p_T > 18 \text{ GeV}$	0.67	0.87
EF_mu8	$p_T > 8 \text{ GeV}$	0.73	0.88
EF_mu8_EFFS	$p_T > 8 \text{ GeV}$	0.98	0.98

Strictly speaking, the *Tag & Probe* method only yields the trigger efficiencies for $Z \rightarrow \ell\ell$ events. Other processes have different event kinematics, which might in turn result in different trigger efficiencies. The results obtained with the *Tag & Probe* method are therefore validated using a closure test. This test has two steps. In the first step, the *Tag & Probe* is applied to a sample of $Z \rightarrow \ell\ell$ events obtained from Monte Carlo simulation. In the second step, a separate Monte Carlo sample containing a variety of processes is used to compare the resulting trigger weights to the trigger simulation. A systematic uncertainty is then assigned to the trigger weights to cover the difference, which is typically $< 1\%$.

9.3.4 Relative missing transverse energy

The missing transverse momentum $\mathbf{p}_T^{\text{miss}}$ of an event is defined as the negative of the vector sum of all electron, muon and photon candidates with $p_T > 10 \text{ GeV}$, all jet candidates with $p_T > 20 \text{ GeV}$ as well as all calorimeter energy deposits with $|\eta| < 4.9$ that are not associated with any such objects. The missing transverse energy E_T^{miss} is the magnitude of the missing transverse momentum. SUSY events are expected

to give rise to large amounts of missing transverse energy since the LSP escapes undetected. This means that E_T^{miss} is a good way to discriminate between SUSY events and Standard Model background.

There is a second source of E_T^{miss} , however. If the energy of an object is not properly measured by ATLAS, then the discrepancy shows up as missing energy in the event. Typically, the energy is underestimated such that the fake E_T^{miss} becomes aligned with the misreconstructed object. The *relative* missing transverse energy is defined as

$$E_T^{\text{miss,rel}} = \begin{cases} E_T^{\text{miss}} & \text{if } \Delta\phi_{\ell,j} \geq \pi/2 \\ E_T^{\text{miss}} \cdot \sin \Delta\phi_{\ell,j} & \text{if } \Delta\phi_{\ell,j} < \pi/2 \end{cases} \quad (177)$$

where $\Delta\phi_{\ell,j}$ is the azimuthal angle between $\mathbf{p}_T^{\text{miss}}$ and the closest signal electron, signal muon, b -jet or central jet. If an event contains significant fake E_T^{miss} due to an object that was misreconstructed, $E_T^{\text{miss,rel}}$ considers only the E_T^{miss} component perpendicular to that object. Using $E_T^{\text{miss,rel}}$ rather than E_T^{miss} therefore helps reject events where the missing energy is due to misreconstructed objects rather than invisible particles that escape detection.

9.3.5 Stransverse mass

The stransverse mass [81, 82] is a variable that can be defined in events with two visible and two or more invisible particles in the final state. In particular, it can be calculated for events where two massive particles both decay semi-invisibly. An example of such an event is the pair production of two W bosons that decay leptonically according to $WW \rightarrow \ell\nu\ell\nu$. The neutrinos are invisible and escape detection.

In order to explain stransverse mass, it is instructive to first introduce the transverse mass m_T . Consider a W boson decaying according to $W \rightarrow \ell\nu$. If it were possible to measure the four-momenta of both the lepton and the neutrino, the mass of the W could be calculated according to

$$m_W^2 = (p_\ell + p_\nu)^2 = m_\ell^2 + m_\nu^2 + 2(E_T^\ell E_T^\nu \cosh \Delta\varphi - \mathbf{p}_T^\ell \mathbf{p}_T^\nu) \quad (178)$$

where p_ℓ and p_ν are the four-momenta of the lepton and the neutrino, m_ℓ and m_ν are their masses, E_T^ℓ and E_T^ν are their transverse energies, $\mathbf{p}_T^\ell$ and $\mathbf{p}_T^\nu$ are their transverse momenta and $\Delta\varphi$ is the difference in rapidity between them. However, since the neutrino is an invisible particle, its momentum cannot be fully measured. Only its transverse component can be inferred from the missing transverse momentum in the event. The longitudinal component remains unknown, which means that the difference in rapidity $\Delta\varphi$ between the lepton and the neutrino can not be calculated. This leads to the introduction of transverse mass.

$$m_T^2(\mathbf{p}_T^\ell, \mathbf{p}_T^\nu) = m_\ell^2 + m_\nu^2 + 2(E_T^\ell E_T^\nu - \mathbf{p}_T^\ell \mathbf{p}_T^\nu) \quad (179)$$

The transverse mass is always calculable, and since $\cosh(x) \geq 1$ it holds that:

$$m_W^2 \geq m_T^2 \quad (180)$$

If the transverse mass is measured for many events, the resulting distribution has an endpoint at the W mass. This allows m_W to be measured even though the neutrino momentum cannot be fully reconstructed.

Now consider an event with a pair of W bosons each decaying leptonically. If the transverse momenta of the individual neutrinos were measurable, it would be possible to calculate the transverse mass for each W boson separately. The W mass would then be constrained by the greater of the two transverse masses, so that

$$m_W^2 \geq \max(m_T^2(\mathbf{p}_T^{\ell_a}, \mathbf{p}_T^{\nu_a}), m_T^2(\mathbf{p}_T^{\ell_b}, \mathbf{p}_T^{\nu_b})) \quad (181)$$

where a and b denote the first and the second particle, respectively. In reality, the transverse momenta of the two neutrinos can not be measured separately. The only information about $\mathbf{p}_T^{\nu_a}$ and $\mathbf{p}_T^{\nu_b}$ comes from

the total missing transverse momentum $\mathbf{p}_T^{\text{miss}}$, which is equal to their sum.

$$\mathbf{p}_T^{v_a} + \mathbf{p}_T^{v_b} = \mathbf{p}_T^{\text{miss}} \quad (182)$$

Since the true momentum partition is not known, *all* partitions that satisfy Eq. (182) must be considered. The only statement that can be made with certainty is that m_W is greater than the minimum transverse mass that can be constructed in this way. This is the definition of the transverse mass m_{T2} .

$$m_W^2 \geq m_{T2}^2 = \min_{\mathbf{p}_T^{\text{miss}}} \left[\max \left(m_T^2(\mathbf{p}_T^{\ell_a}, \mathbf{p}_T^{v_a}), m_T^2(\mathbf{p}_T^{\ell_b}, \mathbf{p}_T^{v_b}) \right) \right] \quad (183)$$

The Standard Model process that yields the highest values of m_{T2} in the two-lepton final state is pair production of W bosons. This means that the m_{T2} distribution for any on-shell Standard Model process of interest in this analysis is constrained to have an endpoint below m_W . The curious reader may want to skip ahead to Figure 82, which shows m_{T2} distributions for various Standard Model processes.

The exercise can be repeated for SUSY events, where the NLSP (typically $\tilde{\chi}_1^\pm$ or $\tilde{\ell}$) plays the role of the W boson and the LSP (typically $\tilde{\chi}_1^0$) plays the role of the invisible particle that escapes detection. For example, Figure 62 (a) describes a SUSY scenario with a slepton NLSP. The m_{T2} distribution then satisfies

$$m_\ell^2 \geq m_{T2}^2 = \min_{\mathbf{p}_T^{\text{miss}}} \left[\max \left(m_T^2(\mathbf{p}_T^{\ell_a}, \mathbf{p}_T^{\tilde{\chi}_{1a}^0}), m_T^2(\mathbf{p}_T^{\ell_b}, \mathbf{p}_T^{\tilde{\chi}_{1b}^0}) \right) \right] \quad (184)$$

which has an endpoint beyond m_W in models where $m_W < m_{\tilde{\ell}}$. However, calculation of m_{T2} requires that the mass of the invisible particle is known. This is obviously not the case for SUSY. The solution is to simply *assume* that the invisible particle has a mass χ and carry out the calculations anyway. The assumed mass may be different from the true mass $m_{\tilde{\chi}_1^0}$. It can be shown [83] that choosing $\chi \neq m_{\tilde{\chi}_1^0}$ offsets the m_{T2} endpoint by a term proportional to the squared mass difference $\chi^2 - m_{\tilde{\chi}_1^0}^2$. The usual approach is to choose $\chi = 0$. The transverse mass then satisfies

$$\left(m_\ell^2 - m_{\tilde{\chi}_1^0}^2 \right) \left(1 - \frac{m_{\tilde{\chi}_1^0}^2}{2m_\ell^2} \right) \geq m_{T2}^2 \quad (185)$$

so that the endpoint of the distribution is moved towards lower masses. The m_{T2} endpoint for the processes shown in Figure 62 (b) and (c) is slightly more complicated to calculate because there are multiple invisible particles in each leg. The solution is to replace all invisible particles of a given leg with one dummy particle whose mass equals the sum of the masses that were replaced. Since neutrinos are approximately massless, the endpoint fulfills

$$\left(m_{\tilde{\chi}_1^\pm}^2 - m_{\tilde{\chi}_1^0}^2 \right) \left(1 - \frac{m_{\tilde{\chi}_1^0}^2}{2m_{\tilde{\chi}_1^\pm}^2} \right) \geq m_{T2}^2 \quad (186)$$

to a very good approximation.

The m_{T2} distribution for SUSY models in which the mass splitting between the NLSP and the LSP is large stretches far beyond that of the Standard Model. In such cases, the m_{T2} variable is an excellent way to discriminate between SUSY signal and Standard Model background. For more compressed models (i.e. models with smaller mass splittings), the discriminating power is lost. Finally, it is interesting to note that m_{T2} and $E_T^{\text{miss,rel}}$ are strongly correlated. This is because both quantities are proportional to the mass difference between the NLSP and the LSP.

9.4 Signal regions

A signal region (SR) is a region within phase space optimized to maximally accept signal events while rejecting background. This analysis defines six signal regions. Three of them are designed to target decays that proceed via sleptons. This corresponds to the processes shown in Figure 62 (a) and (b). These signal regions, which use the m_{T2} variable as their main discriminant, are collectively referred to as SR- m_{T2} . The remaining three signal regions target decays that proceed via Standard Model W bosons. This corresponds to the process shown in Figure 62 (c). These signal regions use either $E_T^{\text{miss,rel}}$ or m_{T2} as their primary means to discriminate between signal and background. They are collectively referred to as SR- WW . Table 29 summarizes the signal region definitions. The notation and motivation for the chosen cuts are discussed in Sections 9.4.2 and 9.4.3.

At the end of the analysis, the observed number of events in each signal region is compared to the expected number of events from Standard Model background processes, and a statistical test is performed in order to determine whether or not there is any significant excess. If not, another test is done in order to determine which of the signal models under consideration can be excluded.

Table 29: Signal region definitions.

SR	m_{T2}^{90}	m_{T2}^{120}	m_{T2}^{150}	WWa	WWb	WWc
Lepton charge	Opposite-sign			Opposite-sign		
Forward jets	0	0	0	0	0	0
b -jets	0	0	0	0	0	0
Central jets	0	0	0	0	0	0
$ m_{\ell\ell} - m_Z $ [GeV]	> 10 for SF channels			> 10 for SF channels		
$p_T^{\ell 1}, p_T^{\ell 2}$ [GeV]	> 35, > 20			> 35, > 20		
$m_{\ell\ell}$ [GeV]	—	—	—	< 120	< 170	—
$E_T^{\text{miss,rel}}$ [GeV]	—	—	—	> 80	—	—
$p_T^{\ell\ell}$ [GeV]	—	—	—	> 80	—	—
m_{T2} [GeV]	> 90	> 120	> 150	—	> 90	> 100

9.4.1 Optimization technique

A signal region is defined by finding the set of cuts that maximizes the significance with respect to a set of SUSY models. During optimization, all background estimates are based purely on Monte Carlo simulation with the exception of the non-prompt background, which is obtained using the Matrix Method (see Section 9.5.4). The Z_n variable is used as a measure of significance. It is defined as

$$Z_n = \sqrt{2} \operatorname{erf}^{-1}(1 - 2p) \quad (187)$$

where p is the probability that the expected Standard Model background fluctuates to at least the number of events observed in data. The probability p is given by

$$p = \int_0^\infty \left(\mathcal{G}(x|N_b, \delta b N_b) \sum_{n=N_d}^\infty \mathcal{P}(n|x) \right) dx \quad (188)$$

where $\mathcal{G}$ denotes the Gaussian function, $\mathcal{P}$ denotes the Poisson function, N_b denotes a background expectation with uncertainty δb , and N_d is the number of events observed in data. Equations (187) and (188) look complicated, but their interpretation is simple. Given a background and a signal, Z_n translates the

excess into an approximation of the equivalent number of Gaussian standard deviations. A model is considered excludable if a signal region can be found that yields $Z_n > 1.64$. This corresponds to exclusion at a 95% confidence level (CL) for a one-sided test. Hypothesis testing is discussed further in Appendix C.

9.4.2 SR- m_{T2}

These signal regions target the processes shown in Figure 62 (a) and (b). Some cuts follow directly from the processes in question. The sparticles are pair produced and therefore have opposite electric charge. This leads to final states with opposite-sign (OS) leptons, and events with same-sign (SS) leptons can therefore be rejected. The diagrams do not contain any quarks, which means that all events that contain jets can be rejected. Such a jet veto is especially efficient at removing background from top processes (see Section 9.5.3). In the ee and $\mu\mu$ channels, events where the invariant mass $m_{\ell\ell}$ of the di-lepton system is compatible with a Z boson are rejected by requiring $|m_{\ell\ell} - m_Z| > 10$ GeV. This removes any process involving the decay of a Z . Cuts on the p_T of the leading and subleading leptons are then imposed in order to bring the signal region definition closer to that of SR- WW . They are placed at $p_T^{\ell_1} > 35$ GeV and $p_T^{\ell_2} > 20$ GeV. These cuts are found to have no impact on the significance, but they simplify the analysis by allowing SR- WW and SR- m_{T2} to use shared control regions for the estimation of certain backgrounds (see Section 9.5). Finally, a cut on m_{T2} is imposed. This cut mainly removes WW production.

The performance of a given m_{T2} cut is evaluated by producing a plot of Z_n for each point on the relevant signal grid. The m_{T2} cut is then scanned with the goal of maximizing the exclusion region. A flat 30% systematic uncertainty is assumed for all background processes except WW , whose uncertainty is evaluated using the nominal method described in Section 9.6.2. It is worth pointing out that placing the cut at $m_{T2} > m_W$ is not necessarily optimal. Off-shell processes and events with misreconstructed objects still contribute some Standard Model background above this cut.

The scan found that lower m_{T2} cuts perform better for models with a small mass splitting between the NLSP and LSP, while higher m_{T2} cuts perform better for models with a large mass splitting. Three nested signal regions are therefore defined, each with a successively tighter m_{T2} cut. Placing the cuts at $m_{T2} > 90$ GeV, $m_{T2} > 120$ GeV and $m_{T2} > 150$ GeV is found to be optimal. For illustrative purposes, Figure 66 shows the resulting exclusion regions in the chargino pair production grid with light sleptons. The same cuts were found to be optimal for slepton pair production. Only same-flavor events were used in this case. The signal regions corresponding to the three m_{T2} cuts are referred to as SR- m_{T2}^{90} , SR- m_{T2}^{120} and SR- m_{T2}^{150} .

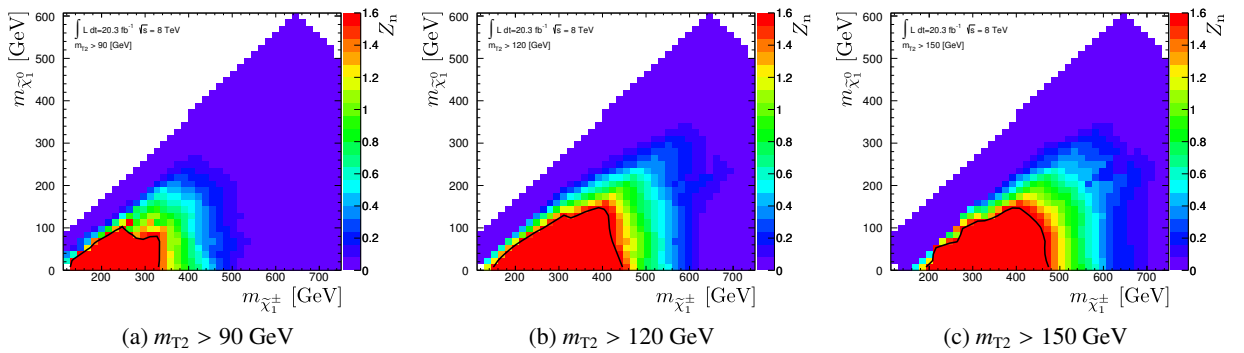


Figure 66: The plots show Z_n as well as the expected exclusion regions in the chargino pair production grid with light sleptons for $m_{T2} > 90$ GeV (a), $m_{T2} > 120$ GeV (b) and $m_{T2} > 150$ GeV (c).

9.4.3 SR-WW

These signal regions target the process shown in Figure 62 (c). Following the same motivation as for SR- m_{T2} , opposite-sign leptons are required and a jet veto is imposed. In the ee and $\mu\mu$ channels, the leptons are required to fulfill $|m_{\ell\ell} - m_Z| > 10$ GeV. Cut optimization studies show that models with small $\tilde{\chi}_1^\pm - \tilde{\chi}_1^0$ mass splittings are experimentally challenging. This is because they produce leptons with small p_T (the kinetic energy of the leptons is proportional to the mass splitting), which is also the case for most Standard Model processes. Effort is instead directed towards model points that have larger mass splittings and consequently yield leptons with higher p_T . Cuts on the leading and subleading lepton p_T are placed at $p_T^{\ell_1} > 35$ GeV and $p_T^{\ell_2} > 20$ GeV. These cuts are also found to reduce the background from non-prompt leptons (see Section 9.5.4).

Cut optimization is done similarly to SR- m_{T2} , with a flat 15% systematic uncertainty assumed for all background processes based on experience from previous iterations of the analysis [2]. Model points with small $m_{\tilde{\chi}_1^\pm}$ are targeted with a signal region based on $E_T^{\text{miss,rel}}$. Cutting at $E_T^{\text{miss,rel}} > 80$ GeV is found to be optimal. A cut on the transverse momentum $p_T^{\ell\ell}$ of the di-lepton system is then placed at $p_T^{\ell\ell} > 80$ GeV in order to remove background from Standard Model WW production (which has significant $E_T^{\text{miss,rel}}$ from neutrinos). Finally, an $m_{\ell\ell} < 120$ GeV cut is found to improve Z_n . The resulting signal region is called SR-WWa. For model points with larger $m_{\tilde{\chi}_1^\pm}$, the m_{T2} variable is again a powerful discriminant between SUSY and Standard Model processes. The signal region SR-WWb is defined by cutting at $m_{T2} > 90$ GeV. An additional $m_{\ell\ell} < 170$ GeV cut is then found to improve Z_n . The signal region SR-WWc is defined by cutting at $m_{T2} > 110$ GeV with no further cuts. The reason the optimal m_{T2} cuts are looser than for SR- m_{T2} is that chargino decays that proceed via W bosons have smaller branching ratios to two-lepton final states. This results in fewer events populating the tail of the m_{T2} distribution.

9.5 Background estimation

A background is a Standard Model process with a signature similar to that of the SUSY process under study. Events from such background processes can therefore make it into the signal regions. The purpose of the background estimation procedure is to provide a prediction of the expected number of Standard Model events in a given signal region as well as the uncertainty of that number. Any excess is then attributed to SUSY.

The dominant backgrounds are due to Standard Model processes that produce two isolated leptons in the final state. Their contribution is estimated using a semi-data-driven technique. The technique works as follows. For each background process, a special region of phase space called a *control region* (CR) is defined. The control region should be as close as possible to the signal region while still being dominated by the background process in question. At the same time, the control region should be orthogonal to the signal region so that their phase spaces do not overlap. Let N_d^{CR} denote the number of events in the control region as measured in real ATLAS data. Most of them are from the background process under study, but some may be from other Standard Model processes or even a SUSY process. Let $N_{\text{NonProc}}^{\text{CR}}$ denote the number of events in the control region from any process other than the background in question, as predicted by Monte Carlo simulation. The difference $N_d^{\text{CR}} - N_{\text{NonProc}}^{\text{CR}}$ is then a data-driven measure of the number of events in the control region from the background under study. This number is extrapolated to the signal region based on Monte Carlo simulation. Let $N_{\text{Proc}}^{\text{CR}}$ and $N_{\text{Proc}}^{\text{SR}}$ denote the number of events from the background process in question in the control region and signal region, respectively. A semi-data-driven prediction of the number of background events in the signal region for the desired process is then given by:

$$N_d^{\text{SR}} = \left(N_d^{\text{CR}} - N_{\text{NonProc}}^{\text{CR}} \right) \cdot \frac{N_{\text{Proc}}^{\text{SR}}}{N_{\text{Proc}}^{\text{CR}}} \quad (189)$$

It is often convenient to talk about certain parts of Eq. (189) by themselves, and they are therefore given

special names. A ratio between the number of events observed in data and predicted by Monte Carlo simulation is called a *scale factor* $\mathcal{S}$. This can be thought of as the factor by which a Monte Carlo prediction should be multiplied to match data.

$$\mathcal{S} = \frac{(N_d^{\text{CR}} - N_{\text{NonProc}}^{\text{CR}})}{N_{\text{Proc}}^{\text{CR}}} \quad (190)$$

A ratio between the number of events in a signal region and control region as predicted by Monte Carlo simulation is called a *transfer factor* $\mathcal{T}$. This is the factor used to extrapolate from a control region to a signal region.

$$\mathcal{T} = \frac{N_{\text{Proc}}^{\text{SR}}}{N_{\text{Proc}}^{\text{CR}}} \quad (191)$$

Table 30 summarizes the control region definitions. SR- m_{T2} , SR- WWb and SR- WWc are similar enough that they can share control regions, while SR- WWa requires its own set. The notation and motivation for the chosen cuts are discussed in Sections 9.5.1 to 9.5.3. The background contribution from non-prompt processes is discussed in Section 9.5.4, while any remaining subdominant backgrounds are covered in Section 9.5.5. Finally, Section 9.5.6 discusses how to fit multiple backgrounds at the same time. A description of the Monte Carlo samples used for the background estimation can be found in Appendix B.

Table 30: Control region definitions.

SR	m_{T2} and WWb/c			WWa		
CR	WW	ZV	Top	WW	ZV	Top
Lepton flavor	DF	SF	DF	DF	SF	DF
Lepton charge	Opposite-sign			Opposite-sign		
Forward jets	0	0	0	0	0	0
b -jets	0	0	≥ 1	0	0	≥ 1
Central jets	0	0	0	0	0	0
$ m_{\ell\ell} - m_Z $ [GeV]	—	< 10	—	—	< 10	—
$p_T^{\ell 1}, p_T^{\ell 2}$ [GeV]	$> 35, > 20$			$> 35, > 20$		
$m_{\ell\ell}$ [GeV]	—	—	—	< 120	—	< 120
$E_T^{\text{miss,rel}}$ [GeV]	—	—	—	$[60, 80]$	> 80	> 80
$p_T^{\ell\ell}$ [GeV]	—	—	—	> 40	> 80	> 80
m_{T2} [GeV]	$[50, 90]$	> 90	> 70	—	—	—

The kinematic plots shown in this section include a hatched region that represents the uncertainty on the Monte Carlo simulation. For technical reasons, the treatment of uncertainties is slightly simplified when creating the plots. More precisely, the total uncertainty is obtained by adding the statistical uncertainty and the systematic uncertainties discussed in Section 9.6.1 together in quadrature. Any correlations between the uncertainties are neglected. The simultaneous fit, which is described in Section 9.5.6, uses the full list of uncertainties and properly accounts for correlations when present. The final results of the analysis are all based on the simultaneous fit.

9.5.1 WW

Pair production of W bosons is the main background for this analysis. It leads to two-lepton final states when both W bosons decay leptonically. The corresponding Feynman diagram is shown in Figure 67. The m_{T2} variable is an effective means to remove WW background, but only for models in which the

mass splitting between the NLSP and LSP is large. Even for models with moderate mass splitting, WW remains a dominant background due to off-shell decays and events with misreconstructed objects. Ultimately, what makes the WW signature so similar to SUSY is that they both contain real missing energy in the final state. In the case of WW , the missing energy comes from neutrinos while in the SUSY case it comes from the LSP.

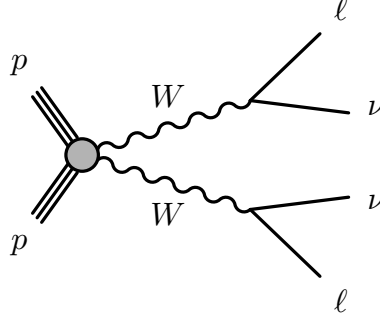


Figure 67: Feynman diagram showing WW production with subsequent decay into a two-lepton final state.

A region pure in WW is difficult to find while still retaining high statistics. The best approach is to require events with high $E_T^{\text{miss,rel}}$. The WW control region for SR- WWa is therefore defined by selecting a slice of $E_T^{\text{miss,rel}}$ in the range $60 < E_T^{\text{miss,rel}} < 80$ GeV, which is right next to the signal region. The $p_T^{\ell\ell}$ cut is then relaxed to $p_T^{\ell\ell} > 40$ GeV in order to increase statistics. The WW control region for SR- m_{T2} and SR- WWb/c is defined by selecting a slice of m_{T2} in the range $50 < m_{T2} < 90$ GeV. Scale factors are then calculated using data from the $e\mu$ channel. This channel is more pure than ee and $\mu\mu$, which suffer from contamination by Z bosons. Although all channels produce statistically compatible scale factors, the one obtained using $e\mu$ data is considered the most reliable. Figures 68 and 69 show some relevant kinematic distributions from the WW control regions for SR- WWa and the remaining SRs, respectively.

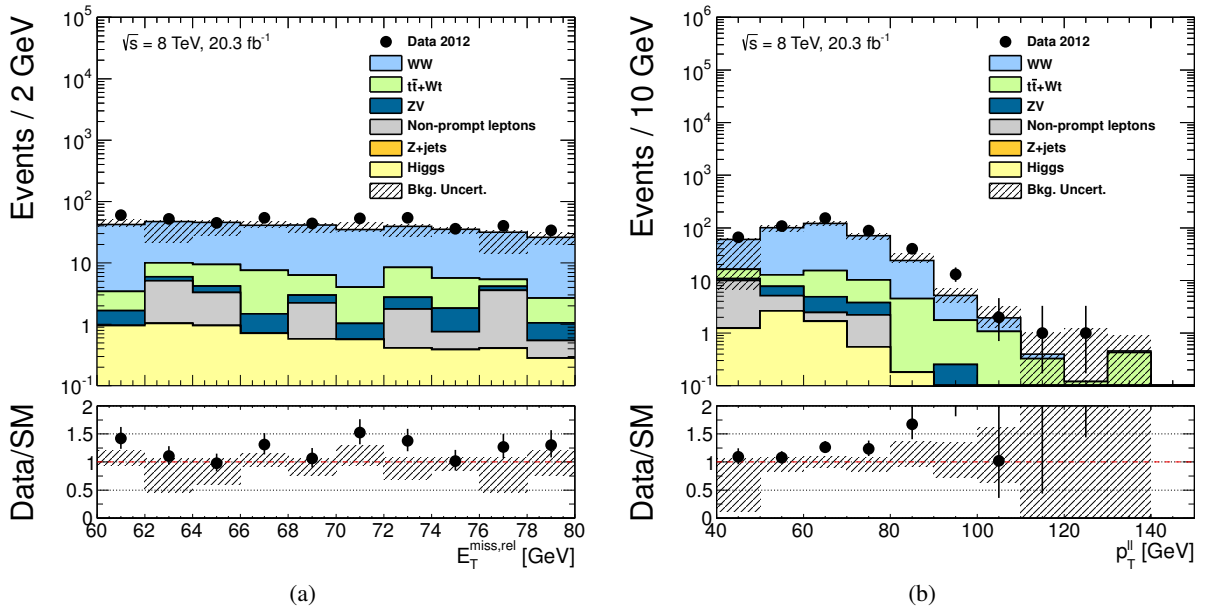


Figure 68: Distributions of $E_T^{\text{miss,rel}}$ (a) and $p_T^{\ell\ell}$ (b) in the WW control region for SR- WWa .

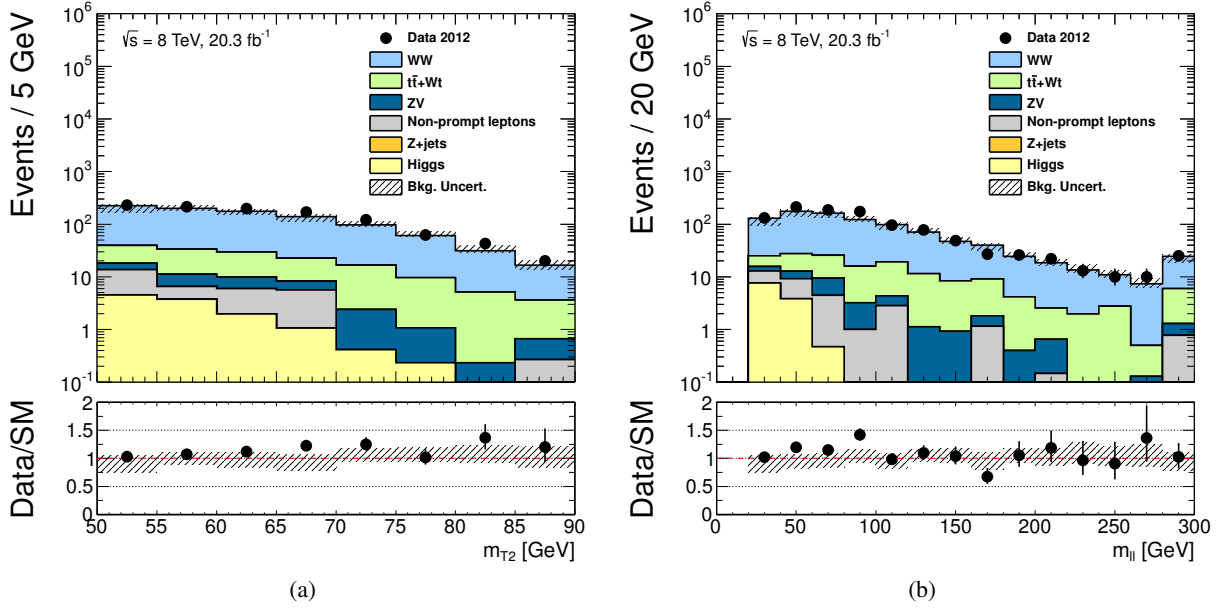


Figure 69: Distributions of m_{T2} (a) and $m_{\ell\ell}$ (b) in the WW control region for $SR\text{-}m_{T2}$ and $SR\text{-}WWb/c$.

Because the signature of SUSY is similar to that of WW , some SUSY model points could significantly contaminate the WW control regions. This would lead to an overestimation of the scale factor, which in turn would lead to an overestimation of the Standard Model background. The result is a decreased exclusion power for the analysis overall. SUSY contamination in the WW control regions is only a concern for the model points with the lowest particle masses — Both because they are the most Standard Model-like and because they have the highest production cross-sections. For higher-mass points, the contamination is negligible. The scenario that produces the highest contamination levels is chargino pair production with light sleptons. It contributes up to 52% of the events in the WW control region for $SR\text{-}m_{T2}$ and $SR\text{-}WWb/c$ as shown in Figure 70. This is tolerable for two reasons. Firstly, because the model points with the highest contamination also have the greatest signal yields, they all manage to remain excludable. Secondly, all the points with high contamination levels have already been excluded by a previous iteration of this analysis that used an independent 7 TeV dataset from 2011 [2]. Signal contamination is therefore not a deciding factor in the definition of the control regions. Contamination from direct slepton production is at most 14%. Finally, contamination from chargino production with heavy sleptons is at most 13% in the control region for $SR\text{-}WWa$. All model points with high contamination remain excludable.

The WW scale factors are shown in Table 31. It is worth noting that the scale factor for $SR\text{-}WWa$ is significantly greater than unity. This result is in agreement with the official ATLAS measurement of the WW production cross-section [84]. The experimentally measured cross-section is approximately a factor 1.22 ± 0.11 above the Standard Model prediction.

Table 31: The table shows the purity and scale factor obtained for each WW control region.

Control region	Channel	SM contamination	Scale factor
WWa	$e\mu$	15%	1.22 ± 0.07
m_{T2} and WWb/c	$e\mu$	17%	1.15 ± 0.13

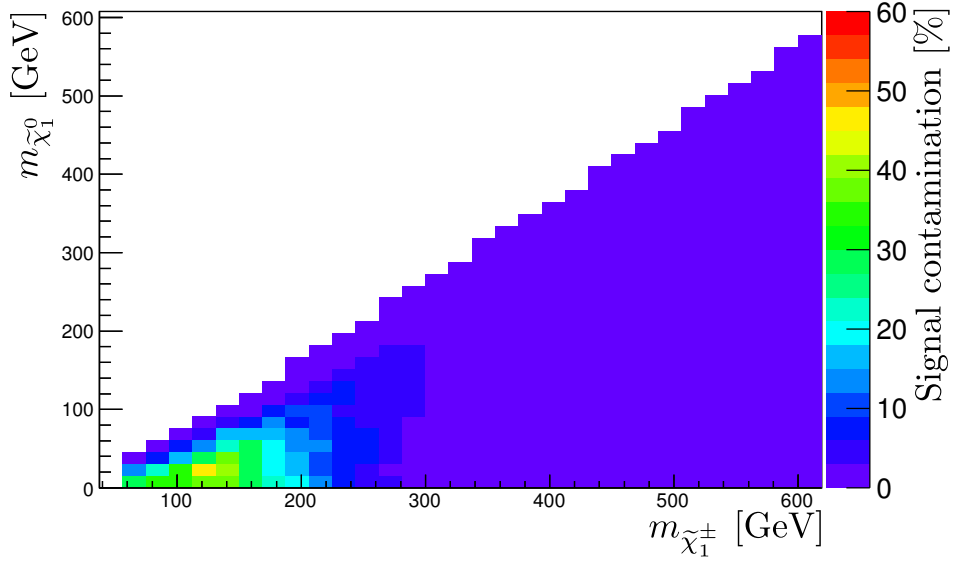


Figure 70: The contamination in the WW control region for m_{T2} and SR- WWb/c from chargino pair production models with light sleptons. The contamination level peaks at 52% for low sparticle masses, and then falls off as the sparticle masses increase.

9.5.2 ZV

The second most important background comes from WZ and ZZ production. The processes are collectively referred to as ZV (meaning Z in association with another vector boson). Figure 71 shows some examples of processes that lead to two-lepton final states. The $WZ \rightarrow qq\ell\ell$ and $ZZ \rightarrow qq\ell\ell$ processes are heavily suppressed by the jet veto. The $WZ \rightarrow \ell\nu\ell\ell$ process leads to a three-lepton final state, but it still contributes to the signal regions due to events where the third lepton is not reconstructed. The $ZZ \rightarrow \nu\nu\ell\ell$ process is the main contributor to the ZV background since it leads to two leptons and missing energy. It is similar to the WW background in this regard.

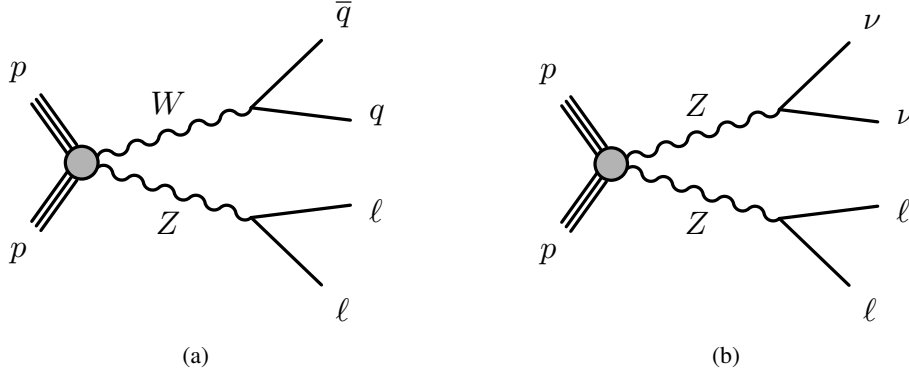


Figure 71: Feynman diagrams showing WZ (a) and ZZ (b) production with subsequent decay into two-lepton final states.

The ZV background always contains two leptons that originate from the decay of a Z boson. The control region for SR- WWa is therefore defined by inverting the Z -veto and requiring $|m_{\ell\ell} - m_Z| < 10$ GeV. In the case of SR- m_{T2} and SR- WWb/c , the control region is defined by requiring $|m_{\ell\ell} - m_Z| < 10$ GeV

while keeping the m_{T2} cut at the lowest common denominator, meaning $m_{T2} > 90$ GeV. Scale factors are evaluated only in the ee and $\mu\mu$ channels because the ZV contribution to the $e\mu$ channel is small. Figures 72 and 73 show some relevant kinematic distributions from the ZV control regions for SR- WWa and the remaining SRs, respectively. The ZV scale factors are shown in Table 32.

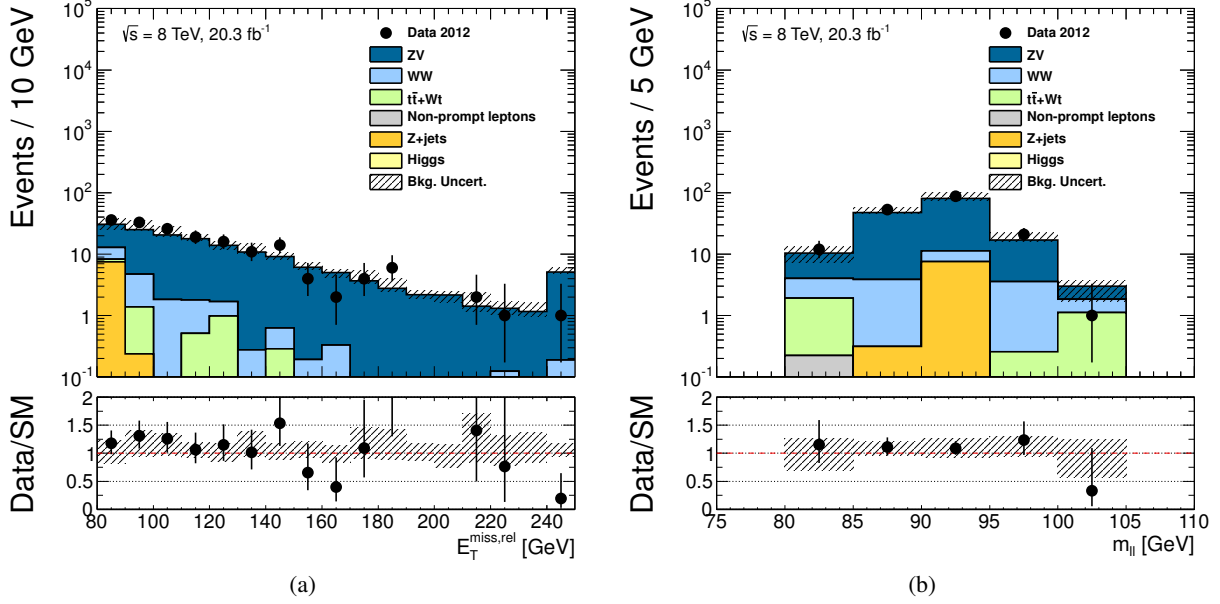


Figure 72: Distributions of $E_T^{\text{miss,rel}}$ (a) and $m_{\ell\ell}$ (b) in the ZV control region for SR- WWa .

Table 32: The table shows the purity and scale factor obtained for each ZV control region.

Control region	Channel	SM contamination	Scale factor
WWa	ee	19%	1.14 ± 0.16
WWa	$\mu\mu$	15%	1.12 ± 0.13
m_{T2} and WWb/c	ee	5%	1.03 ± 0.11
m_{T2} and WWb/c	$\mu\mu$	3%	1.15 ± 0.09

Similarly to the WW background, the scenario with the highest signal contamination is chargino pair production with light sleptons. Contamination levels of up to 20% are observed for low-mass points. But going to higher masses, the purity increases rapidly. For the model points not already excluded by the previous version of the analysis [2], all ZV control regions have a signal contamination less than 4% regardless of SUSY scenario.

9.5.3 Top

Top pair production, here denoted $t\bar{t}$, and associated production of a top quark and a W boson lead to two-lepton final states when the top quarks decay leptonically. The corresponding Feynman diagrams are shown in Figure 74. These two processes are collectively referred to as the *top background*. Although considerably smaller than WW and ZV , top is still the third most important background in the analysis. Since the final state always contains at least one b -quark, the top background is heavily suppressed by the jet veto. However, it has a high production cross-section and can therefore contribute in the signal

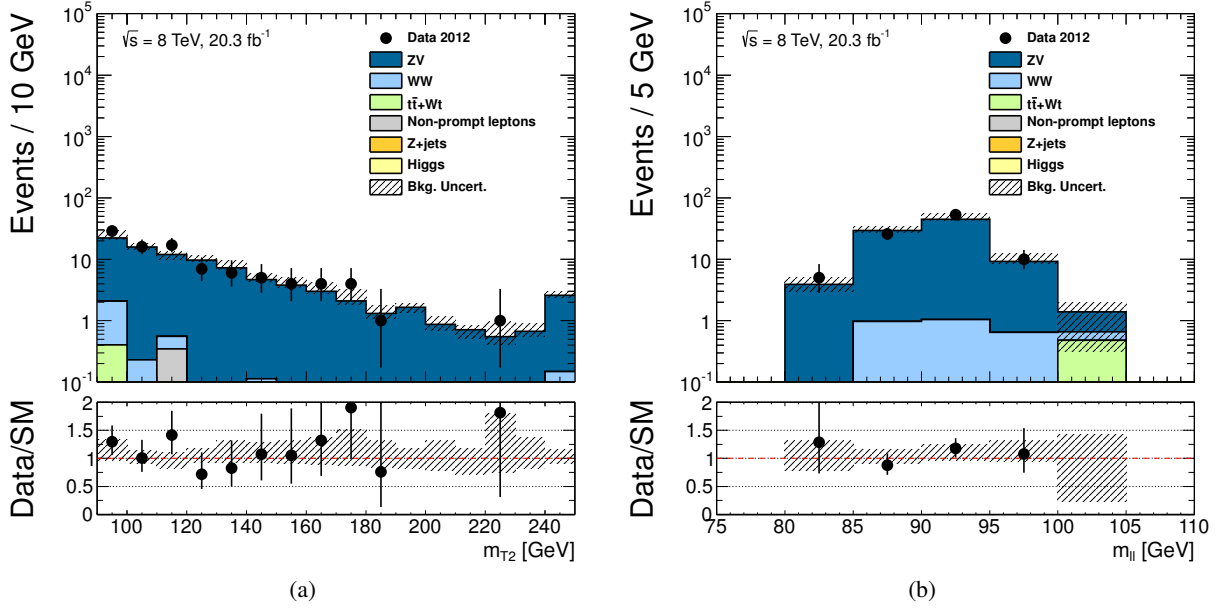


Figure 73: Distributions of m_{T2} (a) and $m_{\ell\ell}$ (b) in the ZV control region for SR- m_{T2} and SR-WWb/c.

regions owing to the occasional events where the b -jets are not reconstructed. The top control region for SR-WWa is obtained by inverting the veto on b -jets. For SR- m_{T2} and SR-WWb/c, it is obtained by both inverting the b -jet veto and relaxing the m_{T2} cut to $m_{T2} > 70$ GeV in order to increase statistics. As for the WW background, the top scale factors are calculated using $e\mu$ data only. The scale factors obtained from all channels are statistically compatible, but the ee and $\mu\mu$ ones are considered less reliable because the channels suffer from contamination by Z boson production.

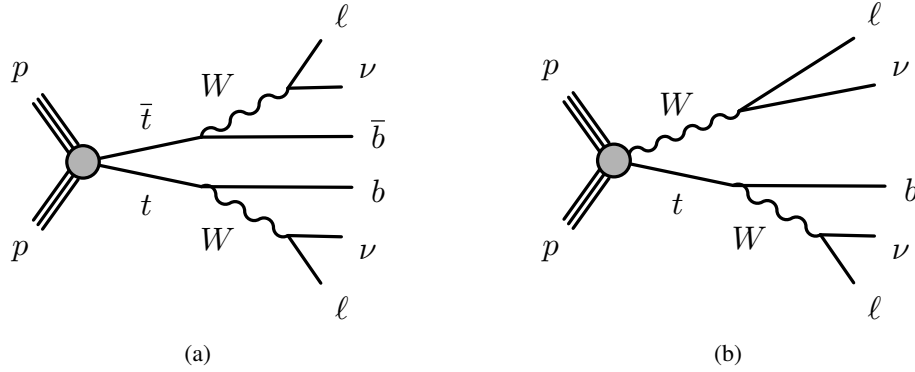


Figure 74: Feynman diagrams showing $t\bar{t}$ (a) and Wt (b) production with subsequent decay into two-lepton final states.

Figure 75 shows some kinematic distributions from the top control regions. The agreement between data and Monte Carlo simulation is good. In particular, Figure 75 (b) shows agreement in the range $70 < m_{T2} < 90$ GeV, which means that the top scale factor is not biased by relaxing the m_{T2} cut. The top purity is clearly very high in both control regions, and signal contamination is found to be negligible. The top scale factors are shown in Table 33.

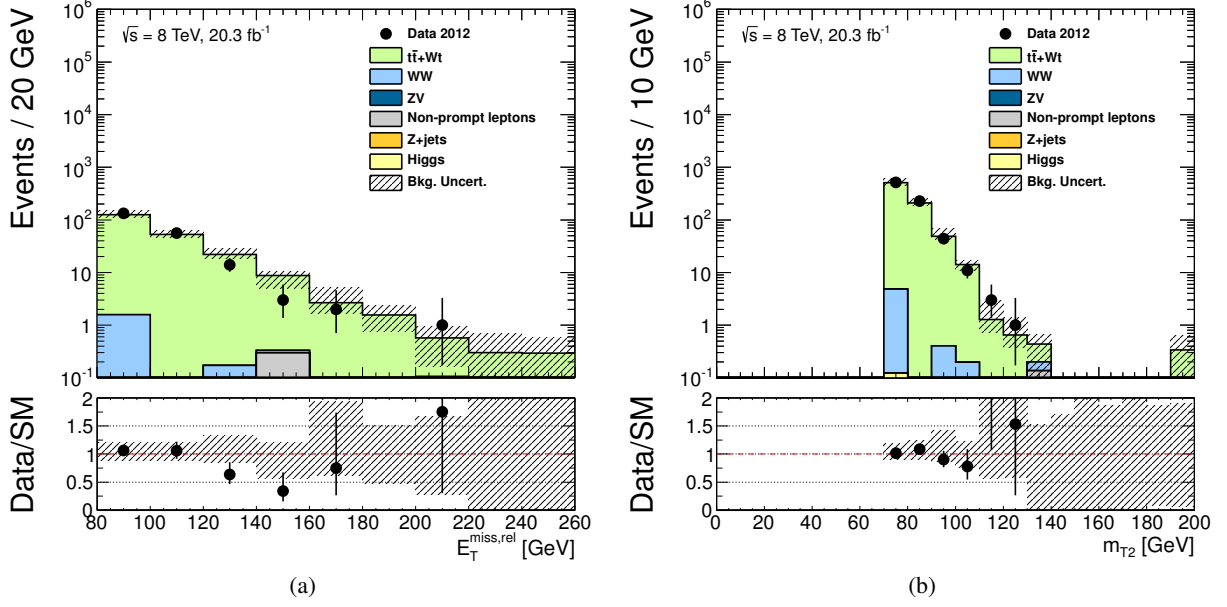


Figure 75: Distributions of $E_T^{\text{miss,rel}}$ in the top control region for SR- WWa (a) and m_{T2} in the top control region for the remaining SRs (b).

Table 33: The table shows the purity and scale factor obtained for each top control region.

Control region	Channel	SM contamination	Scale factor
WWa	$e\mu$	1%	1.00 ± 0.08
m_{T2} and WWb/c	$e\mu$	1%	1.03 ± 0.18

9.5.4 Non-prompt

The ATLAS detector sometimes misidentifies objects and reconstructs them as leptons even though they are not. Such misidentified objects are called *fake* leptons. They can cause events that have fewer than two leptons in the final state to contribute to the signal regions. For example, single-top production normally leads to a one-lepton final state as shown in Figure 76 (a). But if the b -jet is misidentified as a lepton, the process can contribute to the background. Another type of background arises from leptons that are *non-prompt*. That is, they are decay products of long-lived particles. Figure 76 (b) shows a $W\gamma$ event where the photon converts asymmetrically into a high- p_T lepton that is reconstructed and a low- p_T lepton that escapes detection. This event has two real leptons in the final state, but one of them does not come directly from the primary vertex. A perfect detector would have reconstructed the photon instead. Although there is a conceptual difference between *fake backgrounds* and *non-prompt backgrounds*, the two are often lumped together and the expressions are used analogously. This is because both backgrounds are estimated in the same way, using a data-driven technique called the *Matrix Method*. Other examples of fake backgrounds that are covered by the Matrix Method are W +jets and $b\bar{b}$ production.

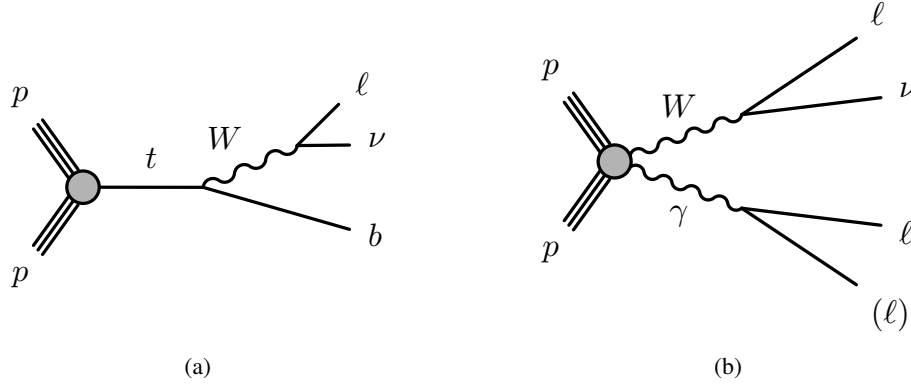


Figure 76: Feynman diagrams showing processes with fake and non-prompt leptons. Single-top production (a) leads to a two-lepton final state when the b -jet is misidentified as a lepton, and $W\gamma$ production (b) can lead to a two-lepton final state when the photon undergoes asymmetric pair production.

As explained in Section 9.3.1, every lepton is subject to two sets of selection criteria. Leptons that satisfy the baseline criteria are called *loose*, while leptons that in addition satisfy the signal criteria are called *tight*. The Matrix Method is based on two quantities called the *real efficiency* and the *fake rate*. The real efficiency r is the probability for a real, prompt lepton that passes the loose selection criteria to also pass the tight ones. The fake rate f is the probability for a fake (or non-prompt) lepton that passes the loose selection criteria to also pass the tight ones. Given these definitions, the number of events with any combination of tight and loose leptons can be expressed as a function of the number of events with real and fake leptons in the same region of phase space. For example, the number of events with two tight leptons N_{TT} is given by

$$N_{\text{TT}} = r_1 r_2 \cdot N_{\text{RR}} + r_1 f_2 \cdot N_{\text{RF}} + f_1 r_2 \cdot N_{\text{FR}} + f_1 f_2 \cdot N_{\text{FF}} \quad (192)$$

where r_1 and r_2 are the real efficiencies for leading and subleading leptons, f_1 and f_2 are the fake efficiencies for leading and subleading leptons, N_{RR} is the number of events with two real leptons, N_{RF} is the number of events with a leading real lepton and a subleading fake lepton, N_{FR} is the number of events with a leading fake lepton and a subleading real lepton, and N_{FF} is the number of events with two fake leptons. The number of events with other combinations of tight and loose leptons are calculated

analogously. This leads to a system of four equations that can be written in matrix form:

$$\begin{pmatrix} N_{TT} \\ N_{TL} \\ N_{LT} \\ N_{LL} \end{pmatrix} = \begin{pmatrix} r_1 r_2 & r_1 f_2 & f_1 r_2 & f_1 f_2 \\ r_1(1-r_2) & r_1(1-f_2) & f_1(1-r_2) & f_1(1-f_2) \\ (1-r_1)r_2 & (1-r_1)f_2 & (1-f_1)r_2 & (1-f_1)f_2 \\ (1-r_1)(1-r_2) & (1-r_1)(1-f_2) & (1-f_1)(1-r_2) & (1-f_1)(1-f_2) \end{pmatrix} \cdot \begin{pmatrix} N_{RR} \\ N_{RF} \\ N_{FR} \\ N_{FF} \end{pmatrix} \quad (193)$$

where N_{TT} is the number of events with two tight leptons, N_{TL} is the number of events with a leading tight lepton and a subleading loose lepton, N_{LT} is the number of events with a leading loose lepton and a subleading tight lepton, and N_{LL} is the number of events with two loose leptons. The quantities N_{TT} , N_{TL} , N_{LT} and N_{LL} are given by the event selection while the real efficiencies and the fake rates are measured as described in Section 9.5.4.1. Equation (193) is therefore a system of four equations and four unknowns. It can be inverted and solved to yield N_{RR} , N_{RF} , N_{FR} and N_{FF} . The fake background is the sum of N_{RF} , N_{FR} and N_{FF} , meaning all events that contain at least one fake lepton.

In order to validate the Matrix Method, the agreement between data and Monte Carlo simulation is checked in a region dominated by the fake background. This validation region is defined by requiring exactly two same-sign signal leptons outside the Z peak and $E_T^{\text{miss,rel}} > 40$ GeV. Figure 77 shows the resulting $E_T^{\text{miss,rel}}$ distributions.

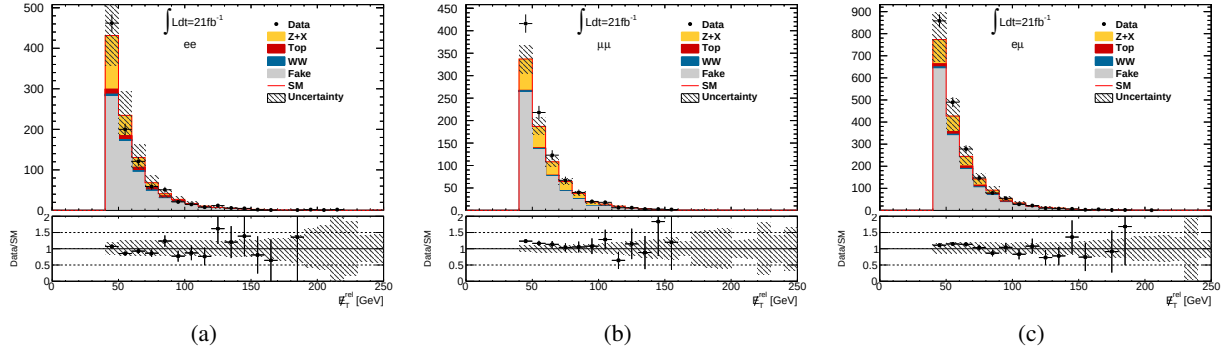


Figure 77: Distributions of $E_T^{\text{miss,rel}}$ in the fake validation region. The ee (a), $\mu\mu$ (b) and $e\mu$ (c) channels are shown separately.

9.5.4.1 Real efficiencies and fake rates

The real efficiencies and fake rates are functions of p_T . For example, a jet could never fake a muon unless it had enough energy to leak through the hadronic calorimeter and into the muon spectrometer. There is also a process dependence. The characteristic b -jets of a $t\bar{t}$ decay are less likely to fake a muon than the lighter jets of a Z decay. It is therefore necessary to measure r and f individually for each process that contributes to the fake background. Once they are known, the real efficiencies and fake rates appropriate for a particular region of phase space are obtained by weighting the process-specific r and f by the fraction of events that each process is expected to contribute to that region. Since it is difficult to know with certainty what kind of process gave rise to a given event in data, both r and f are measured in Monte Carlo samples by making use of generator-level (also called *truth*) information. They are then corrected using data-driven scale factors that are assumed to be process-independent.

Put into an equation, the real efficiency r_ℓ^{XR} for a lepton of type ℓ (electron or muon) in a region of phase space XR is given by

$$r_\ell^{\text{XR}}(p_T) = \sum_i \left(r_i(p_T) \cdot R_i^{\text{XR}} \cdot S^r \right) \quad (194)$$

where $r_i(p_T)$ is the real efficiency for process category i , R_i^{XR} is the fraction of events from process category i in the region XR and $\mathcal{S}^r$ is a scale factor that corrects for differences between data and Monte Carlo simulation. The real efficiency $r_i(p_T)$ for a given process is simply given by

$$r_i(p_T) = \frac{N_i^T(p_T)}{N_i^L(p_T)} \quad (195)$$

where $N_i^T(p_T)$ and $N_i^L(p_T)$ are the number of tight and loose leptons in a given p_T range as measured using truth information from the appropriate Monte Carlo sample. All events with exactly two baseline leptons and $E_T^{\text{miss,rel}} > 20$ GeV are included in the calculation. The fraction R_i^{XR} is calculated as

$$R_i^{\text{XR}} = \frac{N_i^{\text{XR}}}{\sum_k N_k^{\text{XR}}} \quad (196)$$

where N_i^{XR} is the number of events in the region XR from process i and $\sum_k N_k^{\text{XR}}$ is the total number of events in XR from all processes. The real efficiency scale factor $\mathcal{S}^r$ is obtained using a *Tag & Probe* method. A sample of $Z \rightarrow ll$ events is selected from the ATLAS dataset by requiring exactly two opposite-sign baseline leptons with $|m_{\ell\ell} - m_Z| < 10$ GeV. The contribution from other processes is assumed to be negligible such that all leptons in the sample are real and prompt. Events are then selected in which at least one of the leptons (the tag) passes the tight selection. The other lepton in the event is used as the probe. The data-driven real efficiency r_{Data} is given by

$$r_{\text{Data}} = \frac{N_T}{N_L} \quad (197)$$

where N_T and N_L are the number of tight and loose probes. The scale factor is then calculated as

$$\mathcal{S}^r = \frac{r_{\text{Data}}}{r_{\text{MC}}} \quad (198)$$

where r_{MC} is the real efficiency obtained by repeating the exact same *Tag & Probe* method using Monte Carlo truth information. The real efficiencies obtained from Monte Carlo simulation are thus corrected by a data-driven scale factor that is measured for $Z \rightarrow ll$ events but assumed to be equally valid for all processes. Finally, recall that leading leptons do not necessarily have the same real efficiency as subleading leptons and electrons do not necessarily have the same real efficiency as muons. The calculation is therefore repeated once for each case, yielding a total of four real efficiencies.

The fake rates are calculated analogously to the real efficiencies, except they are also parametrized depending on whether the background is due to jets or photon conversions. The sum therefore runs over an additional index j that denotes the fake type:

$$f_{\ell}^{\text{XR}}(p_T) = \sum_{i,j} (f_{ij}(p_T) \cdot R_{ij}^{\text{XR}} \cdot \mathcal{S}_j^f) \quad (199)$$

where $f_{ij}(p_T)$ is the fake rate for process category i and type j , R_{ij}^{XR} is the fraction of events from process category i and type j in the region XR and $\mathcal{S}_j^f$ is a scale factor that corrects for differences between data and Monte Carlo simulation for fakes of type j . Since there are two fake types (jets and photon conversions), two fake rate scale factors are required. They are calculated using a *Tag & Probe* method. Aside from the criteria used to select tags and probes, the procedure is identical to the real efficiency *Tag & Probe* method described above.

For the jet fake rate, a sample dominated by QCD events (i.e. jets) is selected from the ATLAS dataset. All leptons in this sample are assumed to be either fake or non-prompt. One characteristic source

of fakes is jets that punch through the hadronic calorimeter and end up in the muon spectrometer, thereby being identified as muons. All dilepton events that contain a muon matched to a b -jet are therefore selected from the QCD sample. The matched muon is used as the tag. The other lepton in the event (which can be either an electron or a muon) is used as the probe.

The photon conversion fake rate is a special case. Because the cross-section for conversion into muons is negligible compared to the cross-section for conversion into electrons, the conversion fake rate need only be measured for electrons. This is done by selecting a sample of $Z \rightarrow \mu\mu e$ events from the ATLAS dataset. The muons should be opposite-sign and the events should fulfill $|m_{\mu\mu e} - m_Z| < 10$ GeV. All electrons in the sample are then assumed to be due to conversions. In other words, the sample is really made up of $Z \rightarrow \mu\mu\gamma$ events where the photon has converted into an electron. The muon pair is used as the tag and the electron is used as the probe.

9.5.5 Others

A process with high cross-section is the production of a Z/γ^* in association with zero or more jets. This background is referred to as Z +jets. It is heavily suppressed by the requirement of significant missing energy in the final state. Another process, which can lead to two-lepton final states and missing energy but has a very low cross-section, is Higgs production with subsequent decay into W or Z bosons. Examples of Z +jets and Higgs processes are shown in Figure 78. Since they make only subdominant contributions to the signal regions, the background predictions are based purely on Monte Carlo simulation. No control regions are defined.

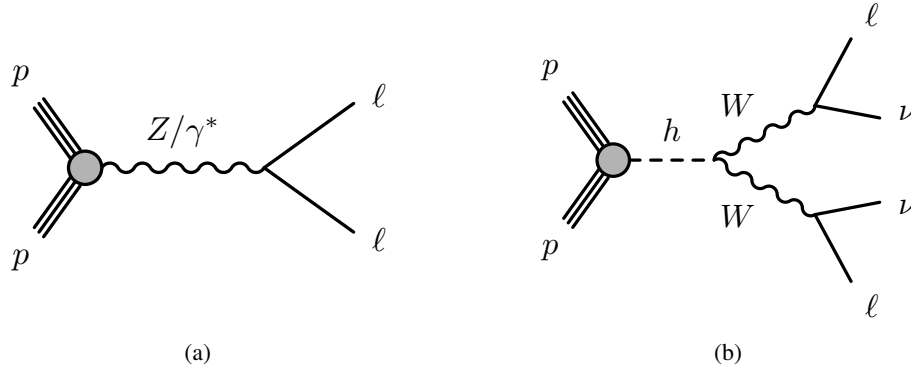


Figure 78: Feynman diagrams showing Z/γ^* (a) and Higgs (b) production with subsequent decay into two-lepton final states.

9.5.6 Simultaneous fit

When the semi-data-driven method described in Section 9.5 is used to obtain background predictions for more than one process, the different predictions become interdependent. This is because the control region of each background receives contributions also from the other processes. If the method determines that the scale factor for a given background is different from unity, then the contribution from this background to the control regions of all the other backgrounds should in principle be rescaled. If this interdependence is taken into account, the naive approach of evaluating one background after the other results in a situation where the final result may depend on the order in which the backgrounds are evaluated. This analysis therefore performs a fit to all backgrounds simultaneously. This also ensures that systematic uncertainties that are correlated between the control regions are properly accounted for.

The fitting procedure is described in detail in Appendix C. While the full fit is used to set limits on the free parameters of various SUSY models, it is also possible to use a simplified version of the fit in order to set limits on the background normalization. This is called a *background fit*. In a background fit, all information about the signal regions are ignored and the background normalization parameters are treated as the parameters of interest. Table 34 shows the result of the background fit. The normalizations should be compared to the scale factors presented in Sections 9.5.1 to 9.5.3. Good agreement is observed.

Table 34: The table shows the number of observed and predicted events, the normalization factor and the composition of each control region as obtained from the simultaneous fit.

SR	m_{T2} and WWb/c			WWa		
CR	WW	ZV	Top	WW	ZV	Top
Observed events	1061	94	804	472	175	209
MC prediction	947	91	789	385	162	215
Normalization	1.14	1.04	1.02	1.27	1.08	0.97
Statistical uncertainty	0.05	0.12	0.04	0.08	0.12	0.08
Composition						
WW	84.6%	5.0%	1.4%	86.8%	10.5%	1.7%
ZV	2.0%	94.9%	0.1%	1.9%	82.9%	< 0.1%
Top	10.4%	< 0.1%	98.5%	7.3%	2.8%	98.2%
Non-prompt	1.9%	< 0.1%	< 0.1%	2.7%	< 0.1%	< 0.1%
Other	1.1%	0.1%	< 0.1%	1.3%	3.8%	< 0.1%

9.6 Background uncertainty

The uncertainty of the background estimation determines the exclusion power of the analysis with respect to the models discussed in Section 9.2. If SUSY exists, it also determines the significance of a signal excess. This section discusses the various sources of uncertainty that are taken into account in the analysis. The impact of each source of uncertainty is first evaluated individually. This information is then used as input for a profile likelihood fit (see Appendix C) which determines the impact of all uncertainties simultaneously. The fit accounts for any correlations between the uncertainties.

Table 35 shows the impact of the various sources of uncertainty on the final background estimate as obtained from the profile likelihood fit. The CR statistics uncertainty is due to the limited number of data events in the control regions. The MC statistics uncertainty is due to the limited number of Monte Carlo events in both the control regions and the signal regions. All remaining uncertainties are systematic rather than statistical. The experimental uncertainties are described in Section 9.6.1. The modelling uncertainties of the three dominant backgrounds are then discussed in Sections 9.6.2 to 9.6.4. Special attention is given to WW, which is the main background in the analysis overall and therefore represents the dominant source of uncertainty. The ZV modelling uncertainty is discussed in moderate detail, and a short summary is presented for top.

9.6.1 Experimental uncertainties

This category covers instrumental sources of error such as the limited resolution of a calorimeter or the precision of a tracker. It also covers errors in the methods used to obtain the background estimate, such as the uncertainty of the trigger weights or the Matrix Method fake rates. In order to evaluate the impact of a given source of uncertainty, the background estimation is redone after varying the uncertainty by one standard deviation. This typically requires re-processing the relevant Monte Carlo samples in order to

Table 35: Breakdown of the various sources of uncertainty considered in the analysis and their impact in % on the final background estimation. The same-flavor and different-flavor channels are shown separately. For asymmetric uncertainties, the average variation is shown. Note that, because of correlations between the uncertainties, the total is not necessarily the quadratic sum of the sources.

SR	m_{T2}^{90}		m_{T2}^{120}		m_{T2}^{150}		WWa		WWb		WWc	
	SF	DF	SF	DF	SF	DF	SF	DF	SF	DF	SF	DF
CR statistics	5	3	6	4	8	4	5	5	5	3	6	4
MC statistics	5	7	7	12	10	23	3	4	5	8	6	10
Jet	4	1	2	1	5	7	3	6	4	2	4	3
Lepton	1	2	1	1	4	1	1	3	2	3	1	8
Soft-term	3	4	1	1	2	8	< 1	2	3	5	1	6
b -tagging	1	2	<1	<1	<1	<1	1	1	1	2	<1	1
Non-prompt	<1	1	<1	<1	1	<1	1	1	1	2	<1	1
Luminosity	<1	<1	<1	<1	<1	<1	<1	<1	<1	<1	<1	<1
Modelling	11	13	21	31	18	40	6	6	8	10	15	19
Total	13	16	24	34	23	47	9	11	12	14	17	24

redo the event selection. The recipes for applying the instrumental uncertainties are supplied via software packages by the ATLAS performance groups.

A distinction can be made between symmetric and asymmetric uncertainties. Symmetric uncertainties are equally likely to increase or decrease a given background. For example, the energy resolution uncertainty of a calorimeter is modelled as a Gaussian smearing of the measured energies. To evaluate its impact, the Monte Carlo samples are re-processed with the smearing applied. The change in the background estimation is then symmetrized — If the result was a 3% decrease of the background, then a $\pm 3\%$ uncertainty is assigned. Asymmetric uncertainties are those that can be varied either *up* or *down*. They require the Monte Carlo samples to be re-processed twice. For example, the up variation of the energy scale uncertainty of a calorimeter is applied by increasing the measured energy of all the relevant objects in the analysis. This might result in a 4% background increase. The down variation is then applied by decreasing all the measured energies. This might result in a 2% background decrease. An uncertainty of $^{+4\%}_{-2\%}$ is then assigned.

9.6.1.1 Jet

This category covers the jet energy resolution [85] and jet energy scale [86] uncertainties. They are derived from a combination of simulation, test-beam data and in-situ measurements. The energy scale uncertainty is provided as a function of p_T , while the energy resolution uncertainty is provided as a function of both p_T and η . The uncertainties also depend on whether a jet originated from a gluon or a quark, whether it was a light or heavy quark, and the number of other pp interactions in the same bunch crossing. These effects are accounted for.

9.6.1.2 Lepton

This category covers the uncertainty of the muon momentum measurement, which is provided as a function of p_T [87], as well as the electron energy resolution and scale uncertainties, which are provided as functions of p_T and η [88]. It also includes the uncertainty on the electron and muon efficiency weights. These are small corrections that are applied to events from Monte Carlo simulation in order to compensate for the fact that ATLAS sometimes fails to identify or reconstruct a lepton. Finally, the

lepton category includes the uncertainty associated with the trigger weights (see Section 9.3.3.2). The uncertainty is estimated by re-calculating the weights after varying some of the parameters used in the calculation. This includes the size of the Z-window, the ΔR cut used to match tag leptons against the single-lepton trigger and the bin sizes used for the p_T , η and ϕ parametrizations. The μ -dependence and any dependence on the run period during which the ATLAS data was recorded are also taken into account. These uncertainties are added in quadrature to the uncertainty obtained from the closure test.

It should be noted that this category does not include any uncertainties associated with taus (as opposed to electrons and muons). All tau uncertainties are found to have a negligible impact on the analysis.

9.6.1.3 Soft-term

The missing transverse energy (see Section 9.3.4) is obtained by adding up the transverse momenta of all objects in an event as well as all calorimeter energy deposits that are not associated with any objects. Those energy deposits are referred to as the *soft term* of the E_T^{miss} calculation. This category covers the uncertainty associated with calculating the soft term [89]. It is worth pointing out that the application of any experimental uncertainty that would affect the missing energy, such as calorimeter scales, calorimeter resolutions and the soft term, is correctly propagated to the E_T^{miss} calculation when the Monte Carlo samples are re-processed.

9.6.1.4 b -tagging

This category covers any uncertainty associated with the MV1 b -tagging algorithm. This includes uncertainties on the b -jet identification efficiency as well as uncertainties on the misidentification rate of lighter jets and gluons [90]. All uncertainties are provided as functions of p_T and η .

9.6.1.5 Non-prompt

This category covers any uncertainty associated with the Matrix Method. The dominant uncertainties are associated with the calculation of the fake rates. They are due to η dependence, limited statistics in the *Tag & Probe* scale factor calculation and differences in fake rate between light and heavy jets.

9.6.1.6 Luminosity

The uncertainty of the integrated luminosity measurement (see Section 6.3) affects all backgrounds equally. It is estimated to be 2.8% following the methodology detailed in Ref. [1]. The dominant uncertainty of the luminosity calibration is associated with the determination of the bunch population product. It should be noted that the impact of the luminosity uncertainty on the final background estimation is always less than 2.8%. This is because the main backgrounds are estimated using the semi-data-driven method described in Section 9.5. The luminosity uncertainty is fully correlated between the signal regions and the control regions. It therefore cancels out of the transfer factor.

9.6.2 WW modelling uncertainty

This section describes four separate and statistically independent sources of theoretical uncertainty that affect the WW background estimation. They are due to the choice of parton showering scheme, the choice of Monte Carlo generator, the choice of QCD scale and the evaluation of the parton distribution functions. The estimation of the uncertainties is based purely on Monte Carlo simulation. All studies are carried out at generator level so that they are independent of any reconstruction effects introduced by the ATLAS detector. In general, uncertainties are obtained either by comparing two different Monte Carlo

generators to each other or by comparing a generator to itself after varying some parameter of the event generation. The uncertainty due to differences between a nominal sample A and a varied sample B is given by

$$\sigma = \frac{\mathcal{T}_A - \mathcal{T}_B}{\mathcal{T}_A} \quad (200)$$

where $\mathcal{T}_A$ and $\mathcal{T}_B$ are the transfer factors obtained from sample A and B , respectively. The uncertainty is thus estimated individually for each signal region by using the appropriate transfer factors as defined in Eq. (191).

Table 36 shows the total WW modelling uncertainty. It is obtained by adding the uncertainties of all four sources together in quadrature. Note that all uncertainties are quoted with error bars. They reflect the limited statistics available in the evaluation of the transfer factors. Only events from the $e\mu$ channel are used in the uncertainty calculation. This is because the WW scale factor (see Section 9.5.1) is based on $e\mu$ channel data only. The ee , $\mu\mu$ and $e\mu$ channel uncertainties have been cross-checked and are all found to be statistically compatible.

Table 36: Total modelling uncertainty in % on the WW background.

WW	$SR-m_{T2}^{90}$	$SR-m_{T2}^{120}$	$SR-m_{T2}^{150}$	$SR-WW_a$	$SR-WW_b$	$SR-WW_c$
Total	18 ± 3	34 ± 5	43 ± 8	6 ± 2	13 ± 4	24 ± 4

9.6.2.1 Hard scattering

The first step in the generation of a Monte Carlo event at the LHC is to simulate precisely what processes take place when the incident protons interact. This step is called the *hard scattering* (HS). The second step is to simulate the hadronization that happens afterwards (see Section 1.4.3). This is the *parton shower* (PS) step. While most Monte Carlo generators are capable of performing both steps, it is also possible to interface two different generators with each other such that one takes care of the hard scattering and the other simulates the parton shower. This feature is exploited in order to evaluate the associated uncertainties. A pair of Monte Carlo generators is denoted by `Generator1+Generator2`. For example, `Powheg+Herwig` signifies that the `Powheg` generator was used to simulate the hard scattering while the `Herwig` generator was used for the parton shower step.

The uncertainty associated with the hard scattering simulation is estimated by comparing `Powheg+Herwig` to `aMC@NLO+Herwig`. The outcome is then cross-checked by comparing `Powheg+Pythia6` to `aMC@NLO+Pythia6`. Table 37 shows the results. Figure 79 shows some relevant kinematic distributions.

Table 37: Hard scattering uncertainty in % on the WW background as obtained by comparing `Powheg+Herwig` to `aMC@NLO+Herwig` and `Powheg+Pythia6` to `aMC@NLO+Pythia6`.

WW	$SR-m_{T2}^{90}$	$SR-m_{T2}^{120}$	$SR-m_{T2}^{150}$	$SR-WW_a$	$SR-WW_b$	$SR-WW_c$
HS (with Herwig PS)	17 ± 3	29 ± 5	32 ± 8	-1 ± 2	10 ± 4	23 ± 4
HS (with Pythia6 PS)	14 ± 6	40 ± 8	33 ± 16	-5 ± 4	11 ± 8	24 ± 7

9.6.2.2 Parton shower

The parton shower uncertainty is estimated by comparing `Powheg+Herwig` to `Powheg+Pythia8`. The outcome is then cross-checked by comparing `aMC@NLO+Herwig` to `aMC@NLO+Pythia6`. Table 38 shows

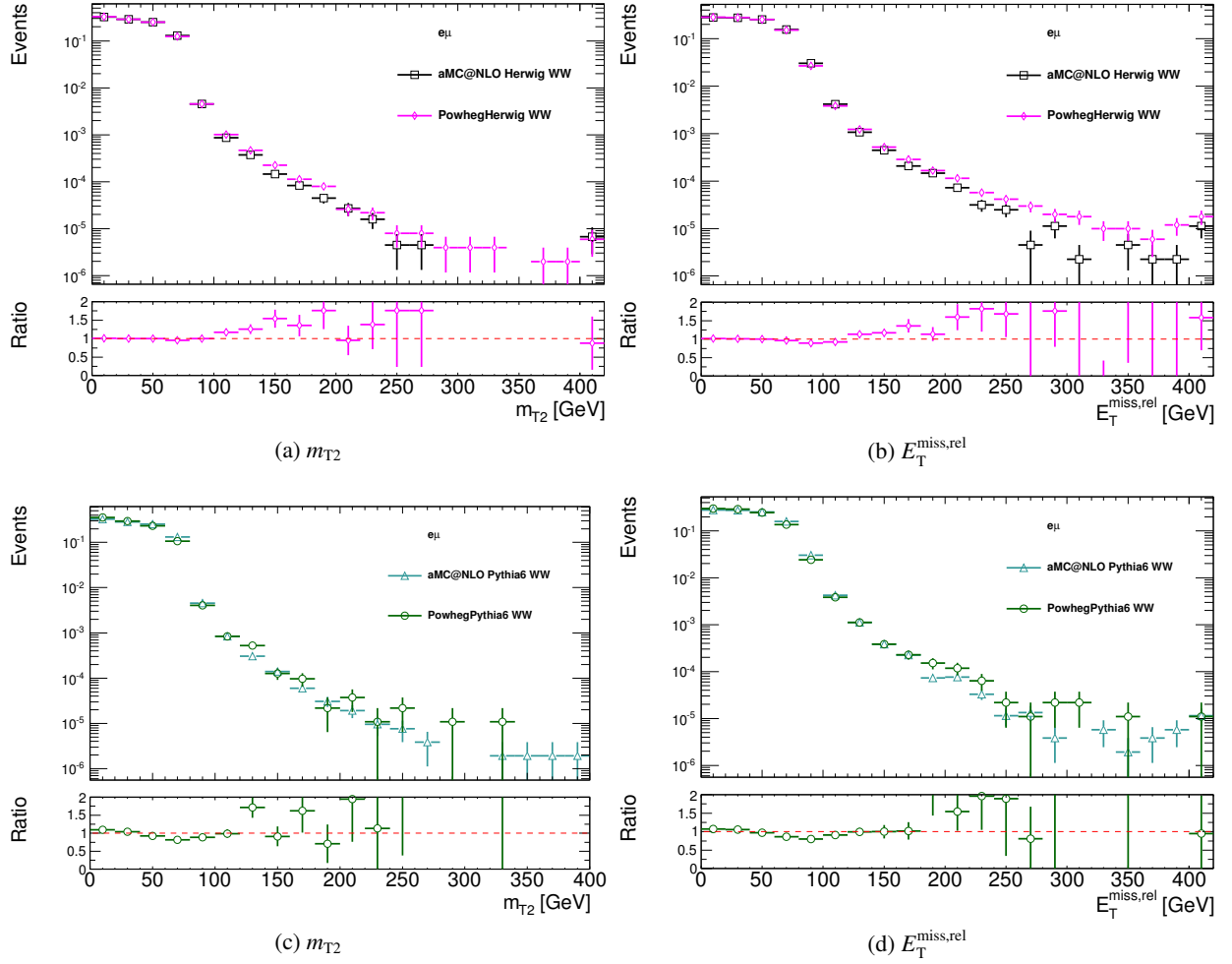


Figure 79: The upper plots show differences in m_{T2} (a) and $E_T^{\text{miss,rel}}$ (b) when comparing Powheg+Herwig to aMC@NLO+Herwig in a region corresponding to SR- m_{T2} without the m_{T2} cut. The lower plots show differences in m_{T2} (c) and $E_T^{\text{miss,rel}}$ (d) when comparing Powheg+Pythia6 to aMC@NLO+Pythia6 in the same region. All plots have been normalized to unity, and the errors are statistical only.

the results, and Figure 80 shows some relevant kinematic distributions.

Table 38: Parton shower uncertainty in % on the WW background as obtained by comparing Powheg+Herwig to Powheg+Pythia8 and aMC@NLO+Herwig to aMC@NLO+Pythia6.

WW	$SR-m_{T2}^{90}$	$SR-m_{T2}^{120}$	$SR-m_{T2}^{150}$	$SR-WW_a$	$SR-WW_b$	$SR-WW_c$
PS (with Powheg HS)	-1 ± 4	4 ± 6	-11 ± 12	-4 ± 2	-4 ± 4	0 ± 4
PS (with aMC@NLO HS)	4 ± 4	17 ± 7	19 ± 11	1 ± 2	5 ± 4	9 ± 5

9.6.2.3 QCD scale

Cross-sections at the LHC are calculated by use of perturbative QCD. When evaluated at a finite order, the equations generally contain singularities that tend to infinity. Only in the limit where all orders are included do the calculations yield a finite result. However, the *factorization theorem* [91] states that a calculation can always be separated into two parts. The first part is process-specific and calculable. The second part contains the singularity. While not calculable, it is experimentally *measurable* (the necessary input being the parton distribution functions).

The factorization process leads to the introduction of a new parameter called the factorization scale μ_F . Its value is in some sense arbitrary as it is not set by the theory. Physically speaking, it corresponds to the energy scale at which the proton is being probed. The calculation further depends on the strong coupling constant introduced in Eq. (41), whose value is again a function of the energy scale at which the process takes place. This energy scale is referred to as the renormalization scale μ_R . It is identical to $|q|$ in Eq. (41). When evaluating a cross-section, it is conventional to set both scales to a common value $\mu_R = \mu_F = \mu$. The value of μ is chosen such that it is close to the typical interaction energy. The uncertainty of the cross-section calculation is then estimated by varying μ_R and μ_F up and down independently. Technically, this is implemented by generating a sample of aMC@NLO+Herwig events at the nominal value of μ as well as four additional samples with the variations $0.5\mu_R$, $2\mu_R$, $0.5\mu_F$ and $2\mu_F$. Each varied sample is compared to the nominal one, and an uncertainty is calculated according to Eq. (200). The final uncertainty is then obtained according to

$$\sigma_{\text{tot}} = \sqrt{\frac{\sum_i \sigma_i^2}{2}} \quad (201)$$

where i runs over the four variations. The factor one half appears because the *up* and *down* variations of each scale are averaged. Table 39 shows the result.

Table 39: QCD scale uncertainty in % on the WW background.

WW	$SR-m_{T2}^{90}$	$SR-m_{T2}^{120}$	$SR-m_{T2}^{150}$	$SR-WW_a$	$SR-WW_b$	$SR-WW_c$
Scale	4 ± 3	16 ± 5	26 ± 8	4 ± 2	6 ± 3	6 ± 4

9.6.2.4 PDF

All the main backgrounds (WW , ZV and top) are simulated using the CT10 PDF set [92]. This set has 26 free parameters which are fit based on data available from previous particle physics experiments. Provided together with the main PDF set is a collection of error sets in which one of the free parameters has been varied either up and down within its 90% confidence interval. This yields 52 varied Monte

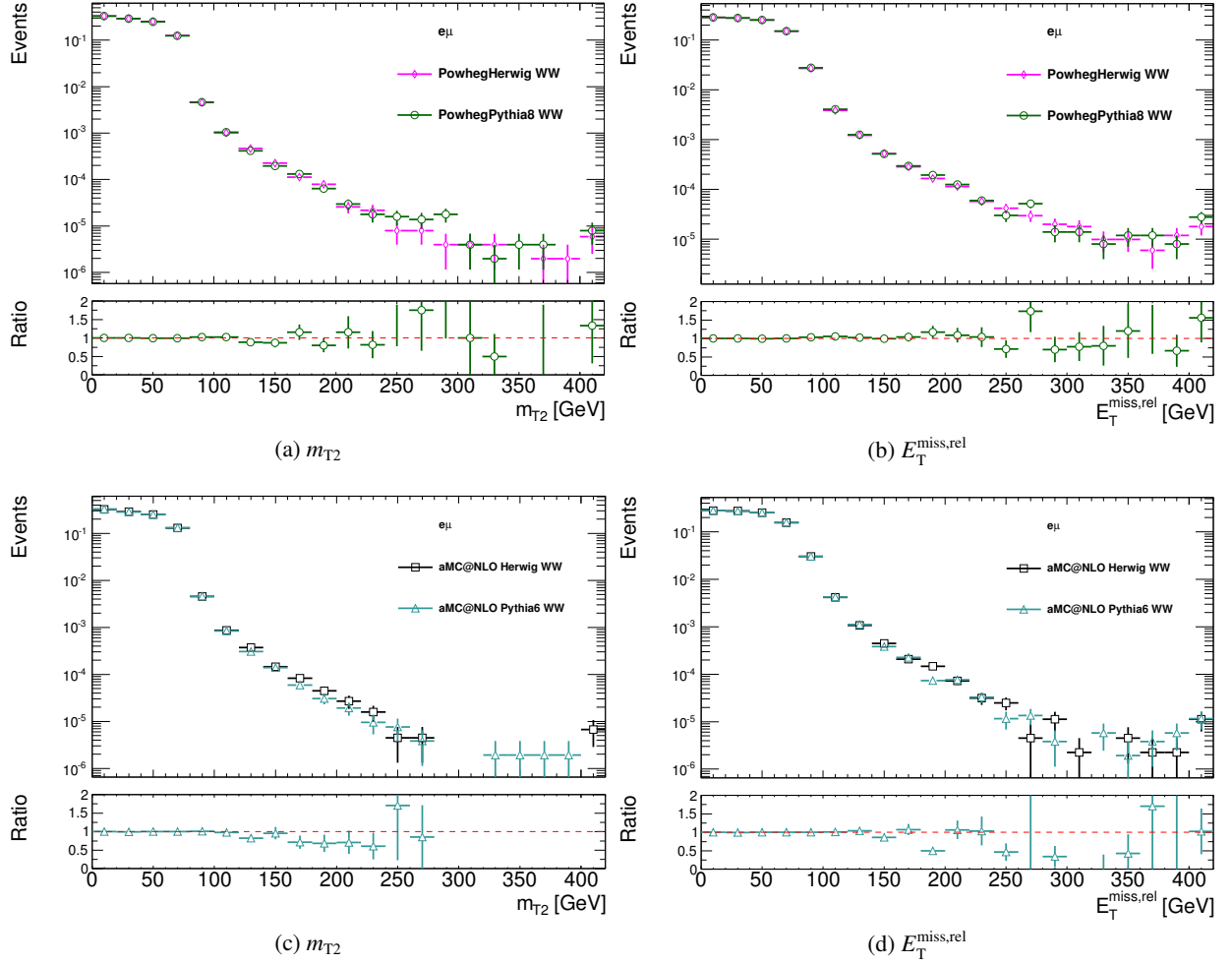


Figure 80: The upper plots show differences in m_{T2} (a) and $E_T^{\text{miss,rel}}$ (b) when comparing Powheg+Herwig to Powheg+Pythia8 in a region corresponding to SR- m_{T2} without the m_{T2} cut. The lower plots show differences in m_{T2} (c) and $E_T^{\text{miss,rel}}$ (d) when comparing aMC@NLO+Herwig to aMC@NLO+Pythia6 in the same region. All plots have been normalized to unity, and the errors are statistical only.

Carlo samples in addition to the nominal one. An uncertainty is calculated for each variation using Eq. (200). Those uncertainties are then summed using Eq. (201). Table 40 shows the result. It is based on Powheg+Pythia8.

Table 40: PDF uncertainty in % on the WW background.

WW	$SR-m_{T2}^{90}$	$SR-m_{T2}^{120}$	$SR-m_{T2}^{150}$	$SR-WW_a$	$SR-WW_b$	$SR-WW_c$
Pdf	2 ± 0	3 ± 0	5 ± 0	1 ± 0	2 ± 0	3 ± 0

The estimation of the PDF uncertainty is arguably more conservative than most other uncertainties because the error sets correspond to 90% rather than 68% confidence intervals. Another way to estimate the PDF uncertainty is to compare the results obtained with CT10 to other PDF sets. Inflating the errors ensures that any such differences are covered [93], thereby obviating the need for comparisons to other sets. This, however, remains a technical detail as the PDF uncertainties are negligible compared to the modelling uncertainty overall.

9.6.2.5 Simulation of W bosons in aMC@NLO

A discrepancy is observed when comparing the W mass distribution as obtained from aMC@NLO to mass distributions from other generators. It turns out that aMC@NLO generates $WW \rightarrow \ell\nu\ell\nu$ events in different ways depending on the masses of the W bosons involved. Only W bosons that are on-shell appear in the event record. For example, an event with two on-shell W bosons would be generated as $pp \rightarrow WW \rightarrow \ell\nu\ell\nu$ while an event with one off-shell W would be generated as a three-body process according to $pp \rightarrow W\ell\nu \rightarrow \ell\nu\ell\nu$. This behavior is illustrated in Figure 81 (a), which plots the mass m_{V1} of the heaviest W boson in WW events. If the W mass peak is constructed from the W bosons that are present in the event record, it appears to be truncated at roughly 80 ± 30 GeV in aMC@NLO. On the other hand, if the W bosons are reconstructed from their respective decay products (ℓ and ν), the shape of the mass peak matches that of Powheg. Note that the W bosons can always be reconstructed unambiguously because all data comes from the $e\mu$ channel. The same check could not have been done in the ee or $\mu\mu$ channels where each event contains two neutrinos of the same type.

As a cross-check, Figure 81 (b) shows the correlation between m_{T2} and the mass of the heaviest W boson in aMC@NLO. The theoretical boundary at $m_{T2} < m_{V1}$ is always satisfied, including when the W is off-shell. The behavior of aMC@NLO is therefore understood, and the samples are safe to use when evaluating systematic uncertainties based on the shape of the m_{T2} distribution.

9.6.3 ZV modelling uncertainty

The method used to estimate the ZV modelling uncertainty is almost identical to the one used for WW (see Section 9.6.2). The differences are that other samples are used for the comparisons, and all calculations are based on the sum of the ee and $\mu\mu$ channels rather than the $e\mu$ channel. This is because the scale factors used for the ZV background estimate are evaluated in the ee and $\mu\mu$ channels only. Table 41 shows the total uncertainties.

Table 41: Total modelling uncertainty in % on the ZV background.

ZV	$SR-m_{T2}^{90}$	$SR-m_{T2}^{120}$	$SR-m_{T2}^{150}$	$SR-WW_a$	$SR-WW_b$	$SR-WW_c$
Total	12 ± 7	30 ± 20	17 ± 16	18 ± 8	12 ± 7	20 ± 9

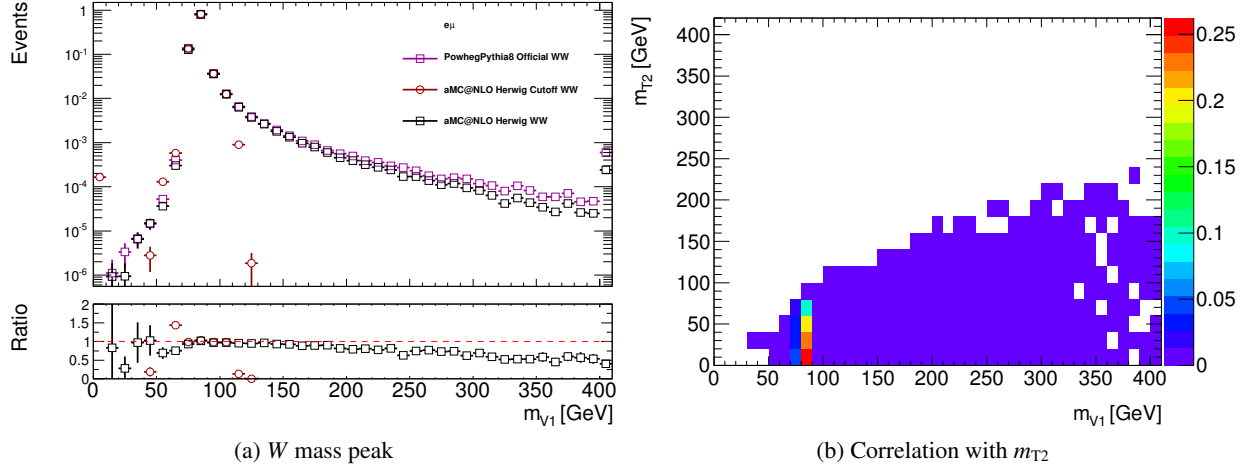


Figure 81: The mass m_{V1} of the heaviest W boson is shown for opposite-sign $e\mu$ events (a). If the mass is taken from the W bosons present in the event record, the peak becomes truncated at 80 ± 30 GeV in aMC@NLO. If the W bosons are instead reconstructed from the final state leptons, the mass distribution matches that of Powheg. Also shown (b) is m_{T2} as a function of m_{V1} as obtained from aMC@NLO in a region corresponding to SR- m_{T2} without the m_{T2} cut. The theoretical boundary at $m_{T2} < m_{V1}$ is satisfied. All distributions have been normalized to unity.

9.6.3.1 Hard scattering and parton shower

The availability of Monte Carlo samples suitable for the calculation of ZV modelling uncertainties is smaller than that of WW . In particular, no combination of samples exists that would allow for a separate estimation of hard scattering and parton shower uncertainties. The two uncertainties are therefore evaluated simultaneously by comparing Powheg+Pythia8 to Sherpa (with Sherpa simulating both the HS and PS steps.) Table 41 shows the result.

Table 42: The sum of the hard scattering and parton shower uncertainties in % on the ZV background as obtained by comparing Powheg+Pythia8 to Sherpa.

ZV	SR- m_{T2}^{90}	SR- m_{T2}^{120}	SR- m_{T2}^{150}	SR- WWa	SR- WWb	SR- WWc
HS + PS	10 ± 4	10 ± 6	12 ± 8	11 ± 4	10 ± 4	9 ± 5

9.6.3.2 QCD scale

The uncertainty associated with the choice of renormalization and factorization scale is estimated using aMC@NLO+Herwig samples with scale variations. See Section 9.6.2.3 for a description of the method. Table 43 shows the result.

Table 43: QCD scale uncertainty in % on the ZV background.

ZV	SR- m_{T2}^{90}	SR- m_{T2}^{120}	SR- m_{T2}^{150}	SR- WWa	SR- WWb	SR- WWc
Scale	7 ± 11	28 ± 21	12 ± 22	14 ± 10	7 ± 11	17 ± 10

9.6.3.3 PDF

The PDF uncertainty is estimated by comparing 52 varied Powheg+Pythia8 samples to a nominal one. See Section 9.6.2.4 for a description of the method. Table 44 shows the result.

Table 44: PDF uncertainty in % on the ZV background.

ZV	$SR-m_{T2}^{90}$	$SR-m_{T2}^{120}$	$SR-m_{T2}^{150}$	$SR-WW_a$	$SR-WW_b$	$SR-WW_c$
Pdf	1 ± 0	2 ± 0	2 ± 0	1 ± 0	1 ± 0	1 ± 0

9.6.4 Top modelling uncertainty

The method used to estimate the top modelling uncertainty is analogous to the one used for WW (see Section 9.6.2), except one additional source of uncertainty is taken into account. This source is the modelling of the initial and final state radiation (ISR and FSR). To evaluate its impact, samples with the amount of radiation varied up and down are generated using AcerMC. The total uncertainty is then obtained using Eq. (201). Similarly to WW , the top uncertainty is calculated using only events from the $e\mu$ channel. Table 45 shows the total uncertainties. No number is reported for $SR-m_{T2}^{150}$ because the top background in this region is zero.

Table 45: Total modelling uncertainty in % on the top background.

Top	$SR-m_{T2}^{90}$	$SR-m_{T2}^{120}$	$SR-m_{T2}^{150}$	$SR-WW_a$	$SR-WW_b$	$SR-WW_c$
Total	24 ± 16	48 ± 23	—	24 ± 6	23 ± 14	32 ± 19

9.7 Results

The scale factors used to estimate the background in each signal region are provided by a profile likelihood fit as described Appendix C. The fit also provides the final uncertainty on the background. In case the number of events observed in data is greater than the background estimation, the significance of the excess is given by the one-sided probability p_0 of the background fluctuating to at least the number of events observed in data. The CL_s technique [94, 95] is used to set an upper limit at a 95% confidence level (CL) on the number of non-SM events in each signal region. This number is then translated into an upper limit on the visible cross-section for physics beyond the Standard Model.

9.7.1 $SR-m_{T2}$

Table 46 shows the results obtained in $SR-m_{T2}$. The error bars on the background estimation include all sources of uncertainty. Any correlations between the different sources, backgrounds and channels are properly accounted for. No significant excess over the Standard Model prediction is observed. Figure 82 shows the m_{T2} distribution that is used to define $SR-m_{T2}$. All background sources have been rescaled using the scale factors obtained from the simultaneous fit, and the error band includes all uncertainties. Both Table 46 and Figure 82 include the signal expectation for two representative SUSY models. These models are clearly incompatible with the observed data. As shown in Section 9.8, they are part of the excluded parameter space.

Table 46: The table shows the number of observed and expected events in SR- m_{T2} together with the signal significance and 95% CL upper limits on the visible cross-section for physics beyond the Standard Model. The same-flavor and different-flavor channels are presented separately. Also shown is the number of signal events predicted for some representative SUSY models.

SR	m_{T2}^{90}		m_{T2}^{120}		m_{T2}^{150}	
	SF	DF	SF	DF	SF	DF
Expected background						
WW	22.1 ± 4.3	16.2 ± 3.2	3.5 ± 1.3	3.3 ± 1.2	1.0 ± 0.5	0.9 ± 0.5
ZV	12.9 ± 2.2	0.8 ± 0.2	4.9 ± 1.6	0.2 ± 0.1	2.2 ± 0.5	< 0.1
Top	3.0 ± 1.8	5.5 ± 1.9	$0.3^{+0.4}_{-0.3}$	< 0.1	< 0.1	< 0.1
Others	0.3 ± 0.3	0.8 ± 0.6	$0.1^{+0.4}_{-0.1}$	0.1 ± 0.1	$0.1^{+0.4}_{-0.1}$	$0.0^{+0.4}_{-0.0}$
Total	38.2 ± 5.1	23.3 ± 3.7	8.9 ± 2.1	3.6 ± 1.2	3.2 ± 0.7	1.0 ± 0.5
Observed events	33	21	5	5	3	2
Predicted signal						
$(m_{\tilde{\chi}_1^+}, m_{\tilde{\chi}_1^0}) = (350, 0)$	24.2 ± 2.5	19.1 ± 2.1	18.1 ± 1.8	14.7 ± 1.7	12.0 ± 1.3	10.1 ± 1.3
$(m_{\tilde{t}}, m_{\tilde{\chi}_1^0}) = (251, 10)$	24.0 ± 2.7		19.1 ± 2.5		14.3 ± 1.7	
p_0	0.50	0.50	0.50	0.27	0.50	0.21
Observed σ_{vis}^{95} [fb]	0.63	0.55	0.26	0.36	0.24	0.26
Expected σ_{vis}^{95} [fb]	$0.78^{+0.32}_{-0.23}$	$0.62^{+0.26}_{-0.18}$	$0.37^{+0.17}_{-0.11}$	$0.30^{+0.13}_{-0.09}$	$0.24^{+0.13}_{-0.08}$	$0.19^{+0.10}_{-0.06}$

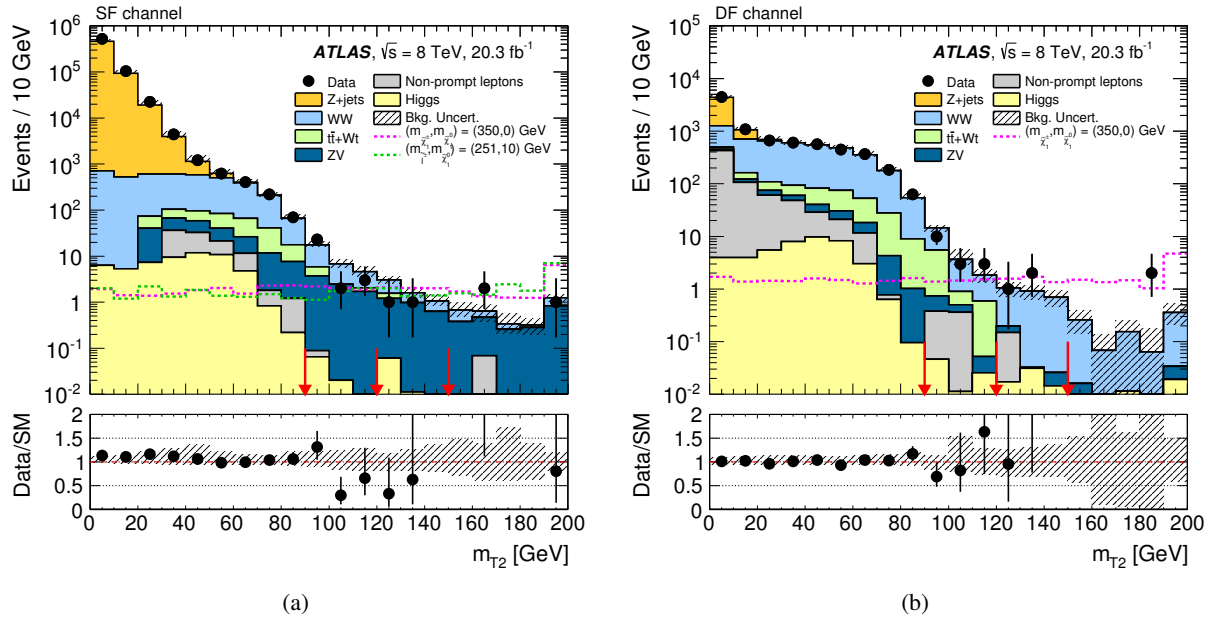


Figure 82: The plots show the m_{T2} distribution for SF (a) and DF (b) events in a region corresponding to SR- m_{T2} without the m_{T2} cut. The signals expected from two benchmark SUSY models are overlaid. Arrows indicate the position of the cuts that define SR- m_{T2}^{90} , SR- m_{T2}^{120} and SR- m_{T2}^{150} .

9.7.2 SR-WW

Table 47 shows the results obtained in SR-WW. The error bars on the background estimation include all sources of uncertainty. Any correlations between the different sources, backgrounds and channels are properly accounted for. No significant excess over the Standard Model prediction is observed. The deficit observed in the SR-WWc same-flavor channel could be due to either a statistical fluctuation or an overestimation of the background. Use of the CL_s technique ensures that the limits remain conservative. Figure 83 shows the $m_{\ell\ell}$ and $E_T^{\text{miss,rel}}$ distributions that are used to define SR-WWa. All background sources have been rescaled using the scale factors obtained from the simultaneous fit, and the error band includes all uncertainties. Both Table 47 and Figure 83 include the signal expectation for a representative SUSY model. It is not obvious from the figure whether or not the model is compatible with the observed data. As shown in Section 9.8, the model is in fact barely excluded.

Table 47: The table shows the number of observed and expected events in SR-WW together with the signal significance and 95% CL upper limits on the visible cross-section for physics beyond the Standard Model. The same-flavor and different-flavor channels are presented separately. Also shown is the number of signal events predicted for some representative SUSY models.

SR	WWa		WWb		WWc	
	SF	DF	SF	DF	SF	DF
Expected background						
WW	57.8 ± 5.5	58.2 ± 6.0	16.4 ± 2.5	12.3 ± 2.0	10.4 ± 2.7	7.3 ± 1.9
ZV	16.3 ± 3.5	1.8 ± 0.5	10.9 ± 1.9	0.6 ± 0.2	9.2 ± 2.1	0.4 ± 0.2
Top	9.2 ± 3.5	11.6 ± 4.3	2.4 ± 1.7	4.3 ± 1.6	$0.6^{+1.2}_{-0.6}$	0.9 ± 0.8
Others	3.3 ± 1.5	2.0 ± 1.1	0.5 ± 0.4	0.9 ± 0.6	$0.1^{+0.5}_{-0.1}$	0.4 ± 0.3
Total	86.5 ± 7.4	73.6 ± 7.9	30.2 ± 3.5	18.1 ± 2.6	20.3 ± 3.5	9.0 ± 2.2
Observed events	73	70	26	17	10	11
Predicted signal						
$(m_{\tilde{\chi}_1^+}, m_{\tilde{\chi}_1^0}) = (100, 0)$	25.6 ± 3.3	24.4 ± 2.2				
$(m_{\tilde{\chi}_1^+}, m_{\tilde{\chi}_1^0}) = (140, 20)$			8.3 ± 0.8	7.2 ± 0.8		
$(m_{\tilde{\chi}_1^+}, m_{\tilde{\chi}_1^0}) = (200, 0)$					5.2 ± 0.5	4.6 ± 0.4
p_0	0.50	0.50	0.50	0.50	0.50	0.31
Observed σ_{vis}^{95} [fb]	0.78	1.00	0.54	0.49	0.29	0.50
Expected σ_{vis}^{95} [fb]	$1.13^{+0.46}_{-0.32}$	$1.11^{+0.44}_{-0.31}$	$0.66^{+0.28}_{-0.20}$	$0.53^{+0.23}_{-0.16}$	$0.52^{+0.23}_{-0.16}$	$0.41^{+0.19}_{-0.12}$

9.8 Interpretation

The results are interpreted in terms of the SUSY models described in Section 9.2. For each model point, an exclusion fit (see Appendix C) is performed using all available signal regions. This fit properly takes into account any signal contamination in the control regions as predicted by the model under study. Upper limits on the number of events from physics beyond the Standard Model are set using the CL_s prescription. The signal region with the best expected limit is then used to test the model for exclusion. A model point is considered excluded if it is incompatible with the observed data at a 95% CL. The result of the limit setting procedure is a contour plot in the mass plane (see e.g. Figure 63) of each SUSY scenario under consideration. The contour separates the excluded model points from the unexcluded ones.

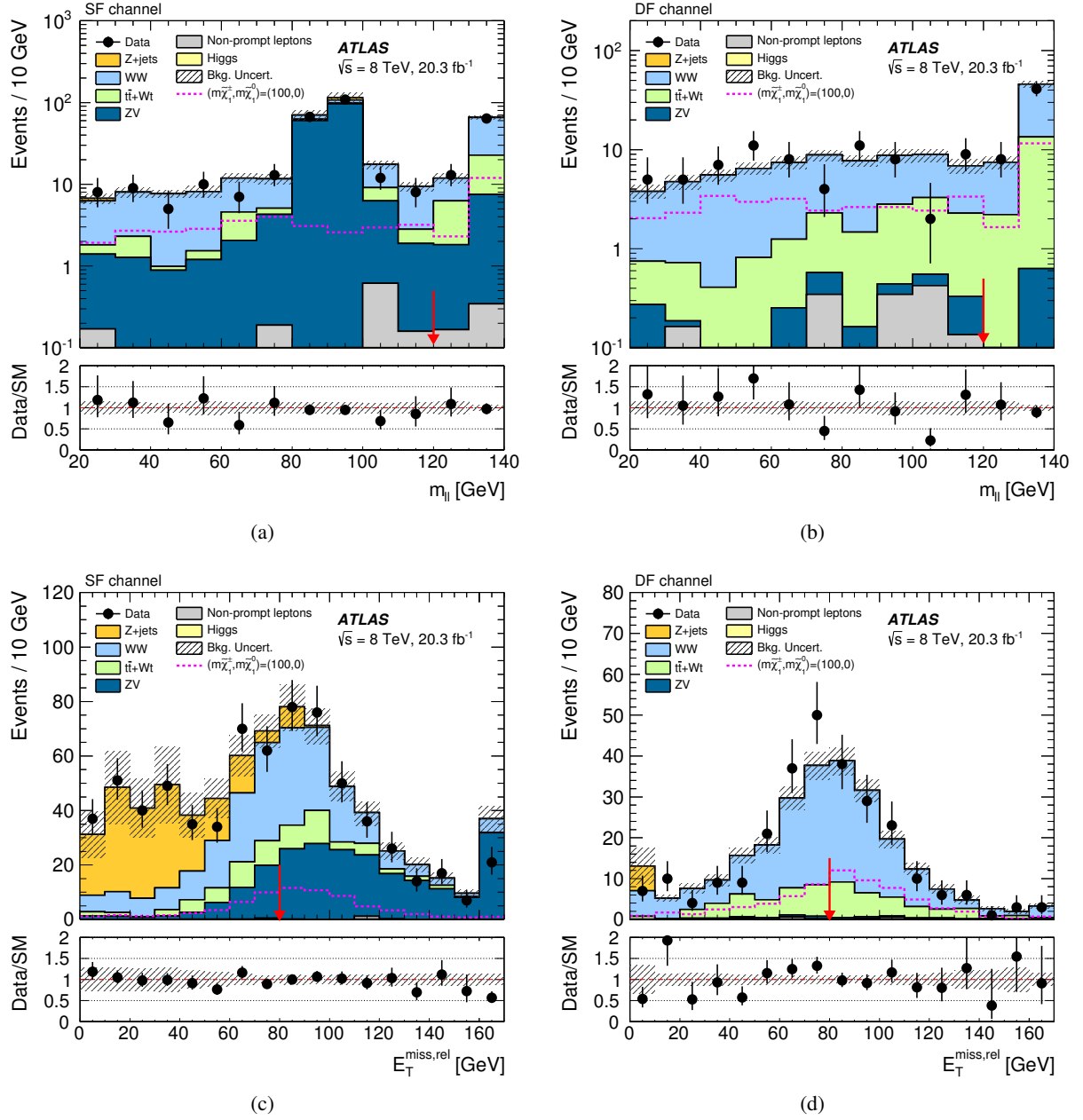


Figure 83: The upper plots show the $m_{\ell\ell}$ distribution for SF (a) and DF (b) events in a region corresponding to SR-WWa without the $m_{\ell\ell}$ and $E_T^{\text{miss,rel}}$ cuts. The lower plots show $E_T^{\text{miss,rel}}$ for SF (c) and DF (d) events in the same region. The signal expected from a benchmark SUSY model is overlaid. Arrows indicate the position of the cuts that define SR-WWa.

9.8.1 Simplified models

9.8.1.1 Slepton pair production

The slepton pair production model is described in Section 9.2.1.1. Limits on the parameters of this model are set using SR- m_{T2} . Figure 84 displays the corresponding exclusion contours. Contours are obtained separately for models that include only right-handed sleptons (a), only left-handed sleptons (b) and both right- and left-handed sleptons (c). In each plot, the dashed line shows the *expected* limit. This is the limit that the analysis expects to set under the Standard Model hypothesis. The solid band around the dashed line represents the expected limit when the background estimation is varied by $\pm 1\sigma$. The solid line shows the *observed* limit. This is the limit that the analysis can set given the data recorded by ATLAS. Due to statistical fluctuations, it is sometimes above and sometimes below the expected limit. The dotted lines surrounding the solid line represent the observed limits when the theoretical signal cross-section is varied by $\pm 1\sigma$. Finally, the filled area to the bottom left shows previous limits from the LEP collider [96].

In order to stay conservative when quoting limits, it is customary to use the values corresponding to the -1σ variation of the theoretical signal cross-section. Sleptons are thus excluded in the range 90-325 GeV for a massless $\tilde{\chi}_1^0$. It is also possible to set limits on selectron and smuon production separately by only using data from the ee and $\mu\mu$ channels, respectively. The selectron exclusion contours are shown in Figure 85. The smuon contours are shown in Figure 86.

9.8.1.2 Chargino pair production with light sleptons

The chargino pair production model with light sleptons is described in Section 9.2.1.2. Limits on the parameters of this model are set using SR- m_{T2} . Figure 87 shows the corresponding exclusion contours. Chargino masses in the range 140-465 GeV are excluded for a massless $\tilde{\chi}_1^0$.

9.8.1.3 Chargino pair production without light sleptons

The chargino pair production model with heavy sleptons is described in Section 9.2.1.3. Limits on the parameters of this model are set using SR- WW . Figure 88 (a) shows the corresponding exclusion contours. As expected, the limits are much weaker than in the case where sleptons are light. Chargino masses in the ranges 100-105 GeV, 120-135 GeV and 145-160 GeV are excluded for a massless $\tilde{\chi}_1^0$. Figure 88 (b) shows the 95% CL limit on the signal cross-section for a massless $\tilde{\chi}_1^0$. The limit has been normalized to the simplified model prediction such that any points < 1 , as marked by the dashed line, are considered excluded.

9.8.2 pMSSM

All signal regions are considered when setting limits in the pMSSM. For each model point, the SR with the best expected sensitivity is used. As explained in Section 9.2.2, two pMSSM scenarios are tested. Figure 89 shows the exclusion contours from the scenario with light sleptons for $M_1 = 100$ GeV (a), $M_1 = 140$ GeV (b) and $M_1 = 250$ GeV (c). Figure 90 shows the exclusion contours from the scenario with heavy sleptons. To help guide the eye, the excluded area enclosed by the expected limit uncertainty band has been filled with solid green.

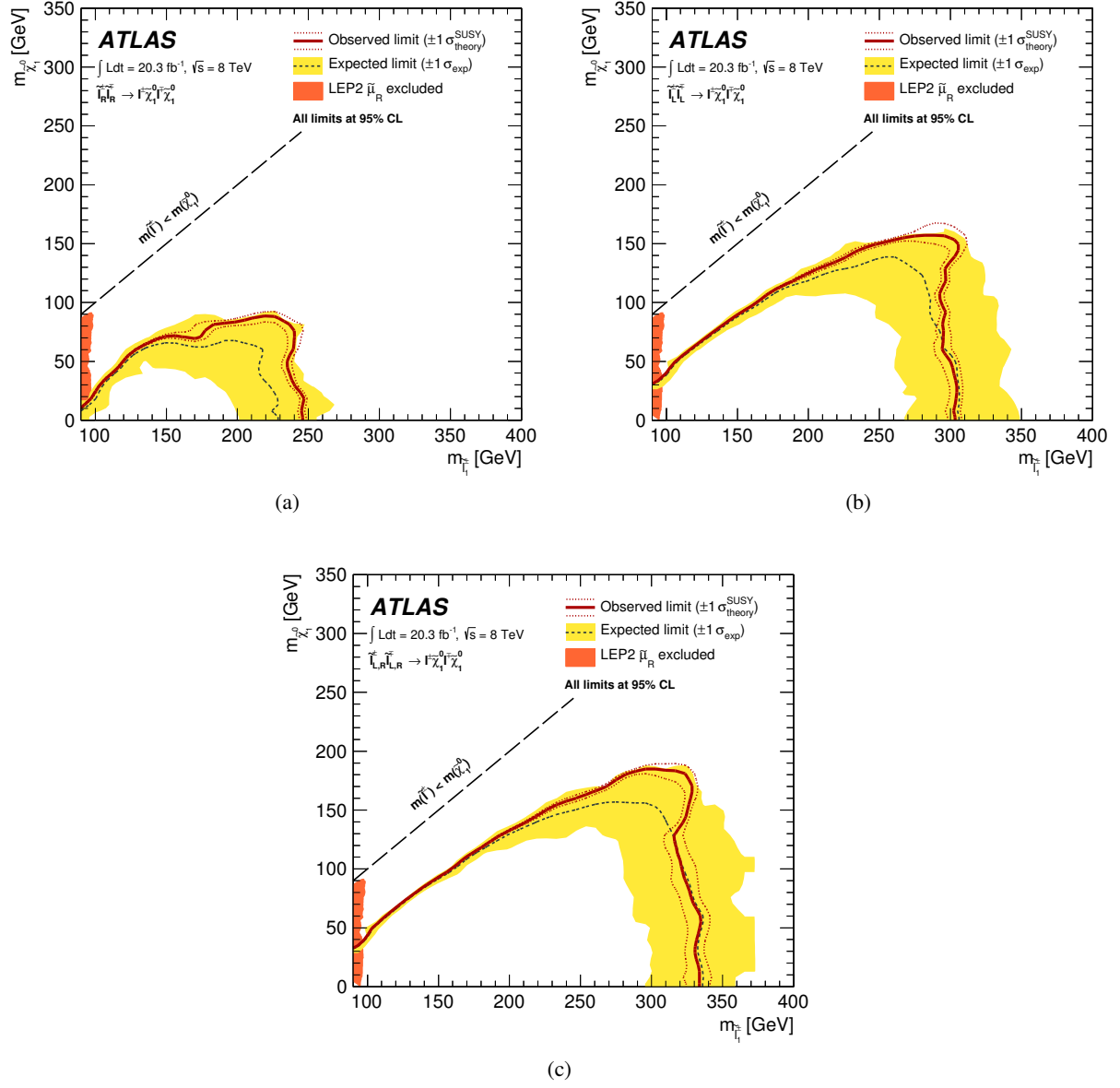


Figure 84: The figure shows 95% CL exclusion contours for slepton pair production in models with only right-handed (a), only left-handed (b) and both right- and left-handed (c) sleptons. The limits were obtained using SR- m_{T2} . LEP limits on the mass of the right-handed smuon are also included.

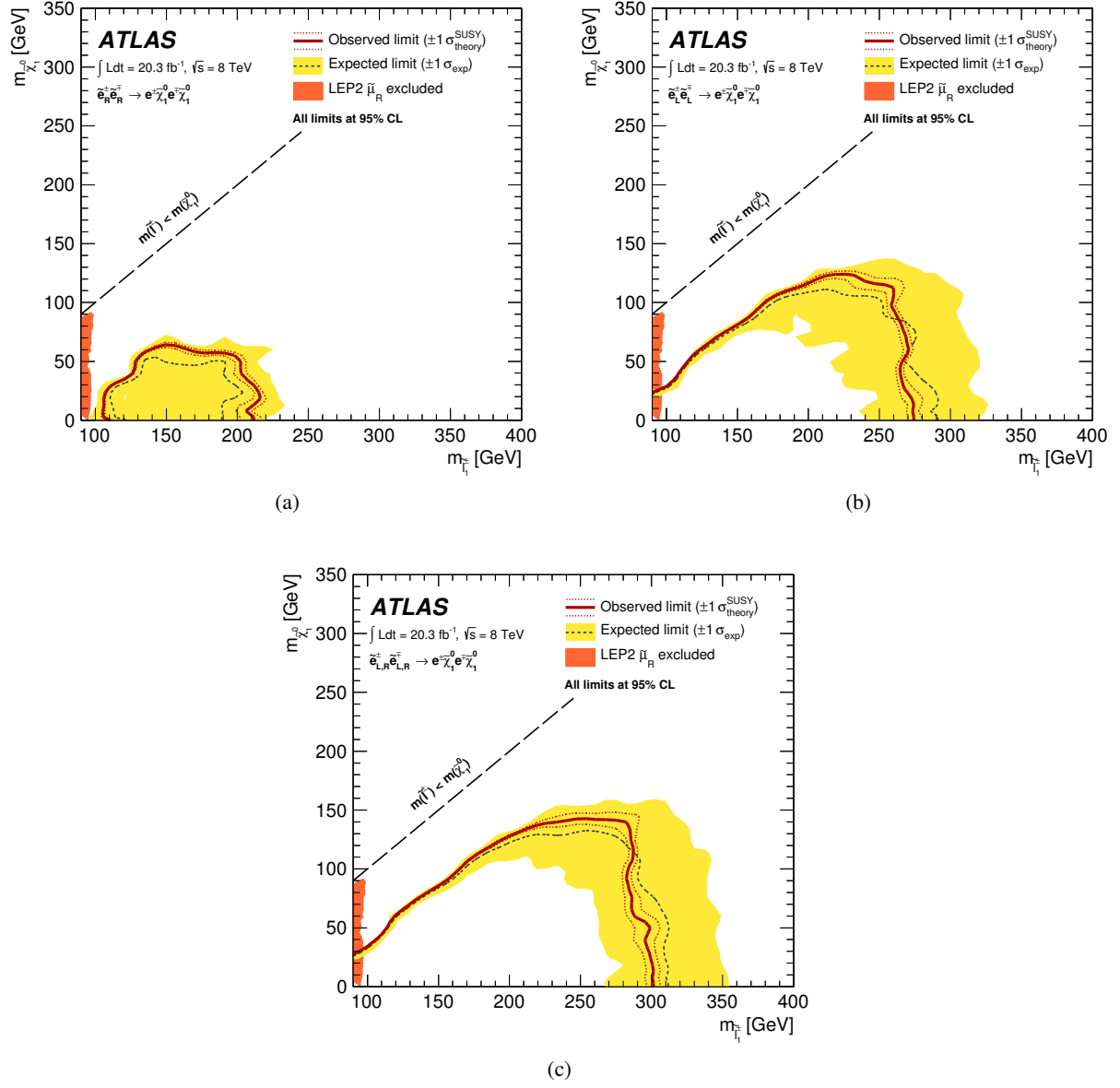


Figure 85: The figure shows 95% CL exclusion contours for selectron pair production in models with only right-handed (a), only left-handed (b) and both right- and left-handed (c) selectrons. The limits were obtained using SR- m_{T2} . LEP limits on the mass of the right-handed smuon are also included.

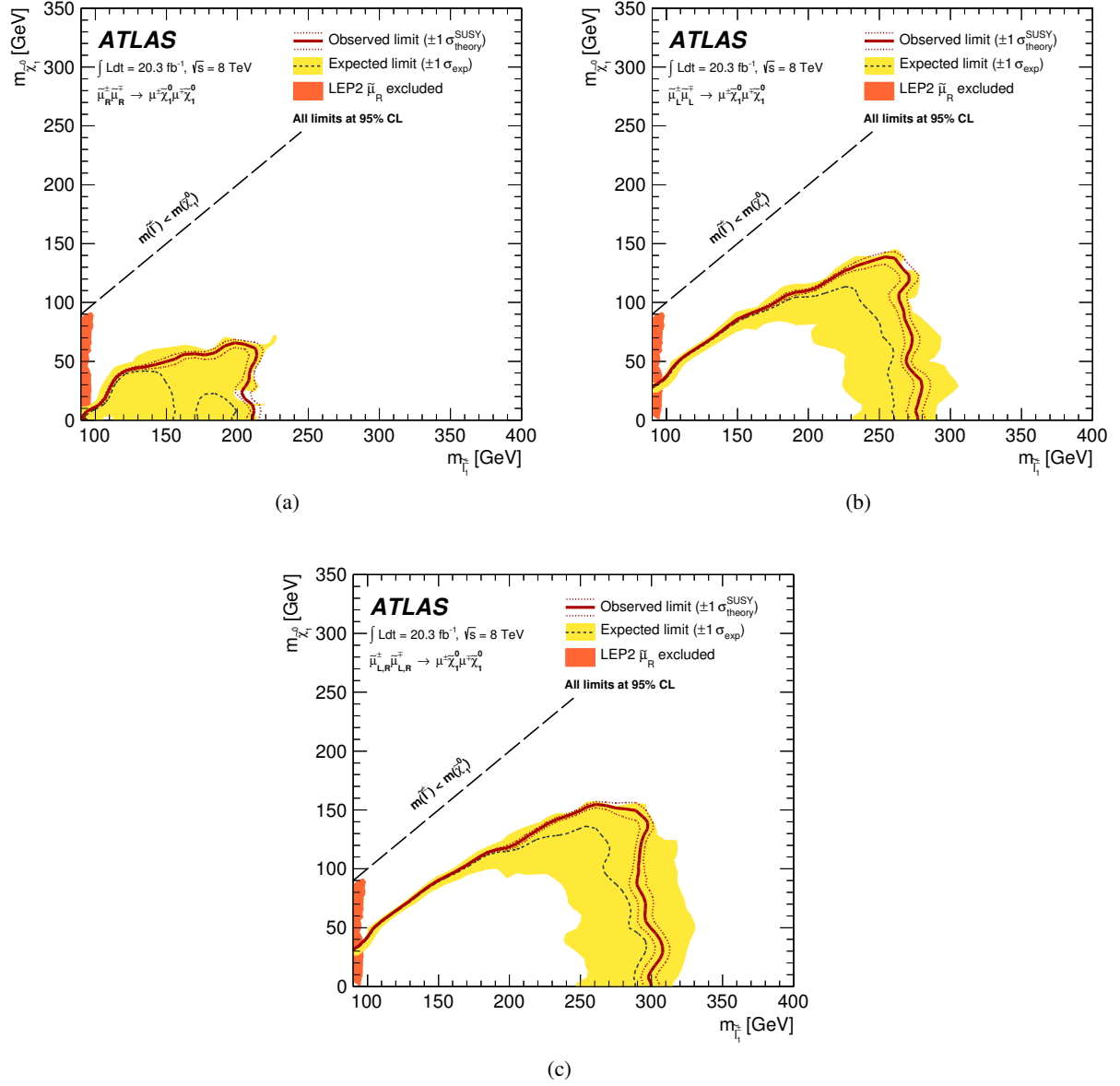


Figure 86: The figure shows 95% CL exclusion contours for smuon pair production in models with only right-handed (a), only left-handed (b) and both right- and left-handed (c) smuons. The limits were obtained using SR- m_{T2} . LEP limits on the mass of the right-handed smuon are also included.

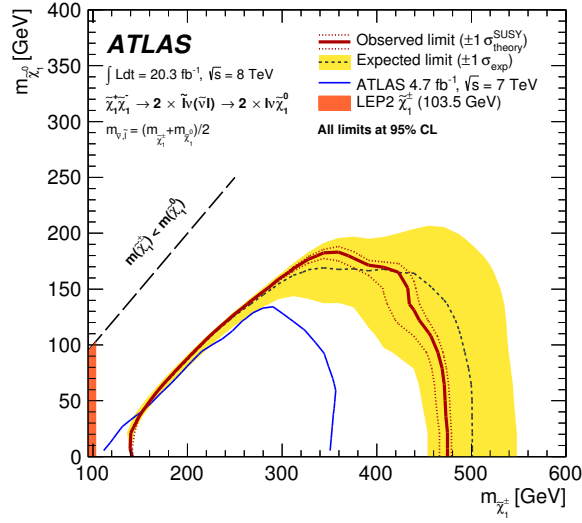


Figure 87: The figure shows 95% CL exclusion contours for chargino pair production in models with light sleptons. The thin solid line shows the exclusion contour set by a previous iteration of the analysis [2]. The limits were obtained using SR- m_{T2} . LEP limits on the mass of the chargino are also included.

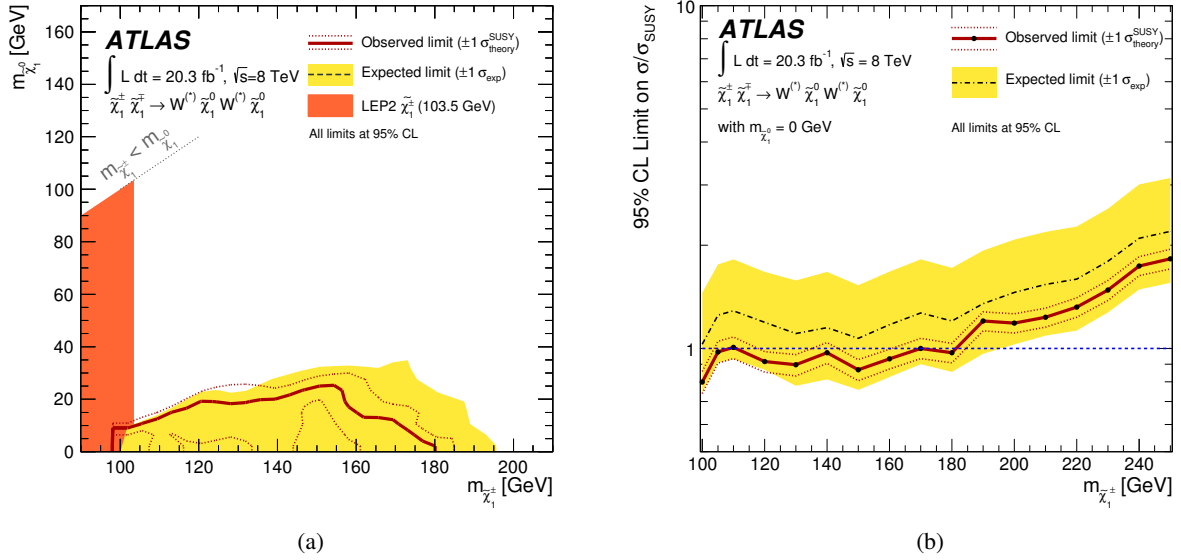
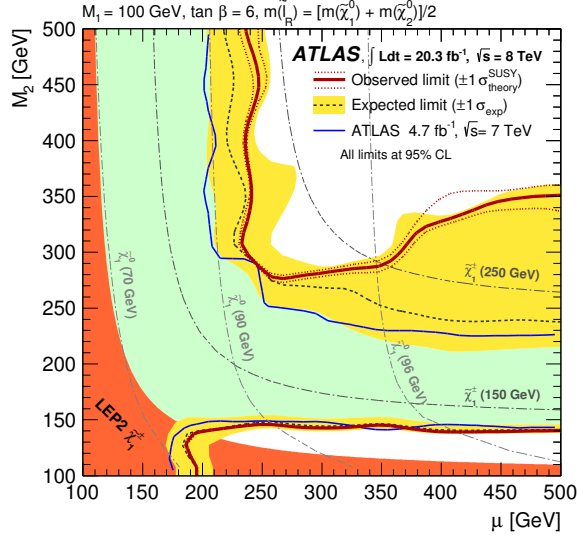
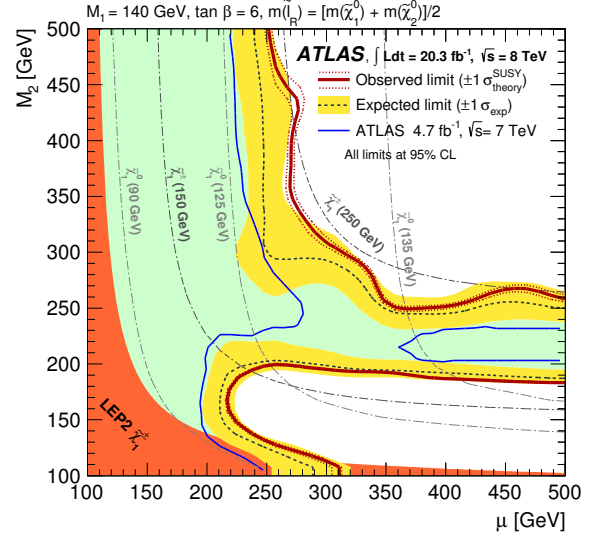


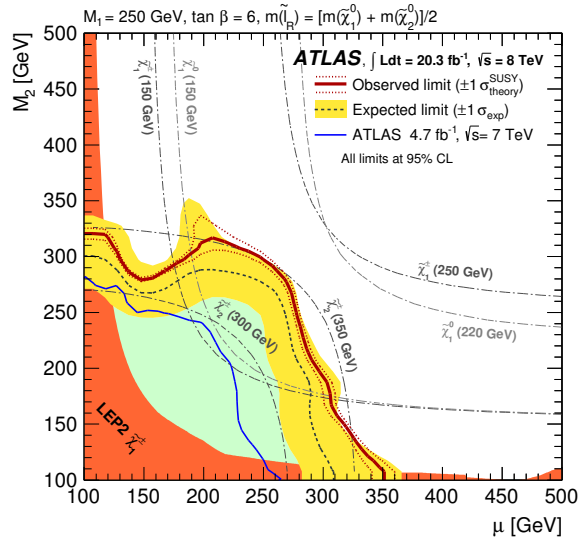
Figure 88: The figure shows 95% CL exclusion contours for chargino pair production in models with heavy sleptons (a) as well as the 95% CL limit on the signal cross-section normalized to the simplified model prediction in the case of a massless $\tilde{\chi}_1^0$ (b). The limits were obtained using SR-WW. LEP limits on the mass of the chargino are also included.



(a)



(b)



(c)

Figure 89: The figure shows 95% CL exclusion contours in the pMSSM scenario with light sleptons for $M_1 = 100$ GeV (a), $M_1 = 140$ GeV (b) and $M_1 = 250$ GeV (c). The thin solid line shows the exclusion contour set by a previous iteration of the analysis [2]. LEP limits based on the mass of the chargino are also included.

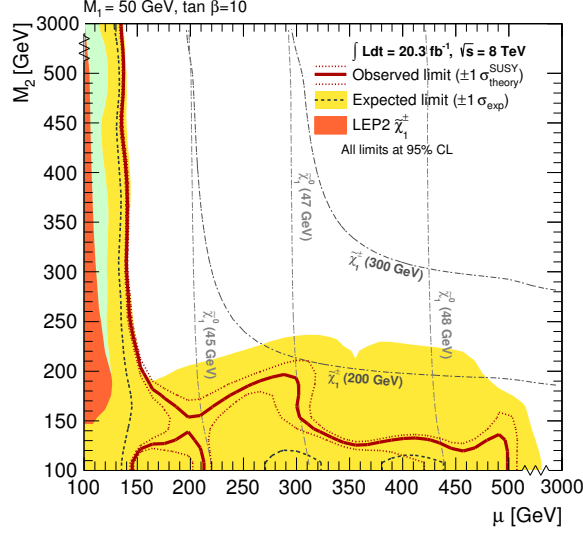


Figure 90: The figure shows 95% CL exclusion contours in the pMSSM scenario with heavy sleptons. LEP limits based on the mass of the chargino are also included.

9.9 Conclusions

A search for the electroweak production of sleptons and gauginos in final states with exactly two oppositely charged leptons has been presented. The search is based on 20.3 fb^{-1} of proton–proton collision data recorded with the ATLAS detector in 2012 at a center-of-mass energy of $\sqrt{s} = 8 \text{ TeV}$. No significant excess over the Standard Model expectation is observed. Limits on the masses of sleptons and gauginos are set both in simplified models and within the pMSSM. In the simplified model describing slepton pair production with subsequent decay to $\tilde{\chi}_1^0$, common masses for right- and left-handed sleptons are excluded in the range 90–325 GeV for a massless $\tilde{\chi}_1^0$. In the simplified model describing chargino pair production with subsequent decay to $\tilde{\chi}_1^\pm$ via sleptons with mass halfway between $\tilde{\chi}_1^\pm$ and $\tilde{\chi}_1^0$, chargino masses in the range 140–465 GeV are excluded for a massless $\tilde{\chi}_1^0$. In the simplified model describing chargino pair production with subsequent decay to $\tilde{\chi}_1^\pm$ via Standard Model gauge bosons, chargino masses in the ranges 100–105 GeV, 120–135 GeV and 145–160 GeV are excluded for a massless $\tilde{\chi}_1^0$.

9.10 Outlook

In order for supersymmetry to solve the hierarchy problem, it must be found at or below the TeV scale. Current limits generally exclude gluinos below 1 TeV and squarks below 0.5 TeV [64]. The electroweak sector is much less constrained. Section 9.9 summarizes the limits on slepton and chargino pair production. A search for $\tilde{\chi}_1^\pm \tilde{\chi}_2^0$ production in three-lepton final states is available in Ref. [74]. In the best-case scenario, degenerate masses for $\tilde{\chi}_1^\pm$ and $\tilde{\chi}_2^0$ are excluded below 700 GeV. In less favorable scenarios, the constraints are significantly weaker. Unlike the limits on colored SUSY, the best limits on electroweak supersymmetry rely on many assumptions and are therefore rather non-general.

Figures 84 to 88 highlight another challenge for the electroweak searches. Compressed models, i.e. models in which the mass splitting between the NLSP and the LSP is small, lead to low- p_T particles in the final state. This makes the SUSY events difficult to distinguish from Standard Model background. Consequently, the excluded region in the limit plots never gets close to the diagonal. Several analyses are underway to probe this region. One approach is to look for same-sign leptons in the final state. This greatly suppresses most Standard Model processes, but it also increases the impact of instrumental uncertainties since fake leptons become the dominant background. Another approach is to make use of

new kinematic variables, such as the *super-razor* [97], that are specifically designed to be sensitive to compressed models.

Although many new analysis designs are in the pipeline, the greatest and most immediate gain in sensitivity will be due to another reason entirely. The next step of LHC operation, run II, is about to commence at a center-of-mass energy of $\sqrt{s} = 13$ TeV. This will lead to higher production cross-sections for SUSY particles across the board, and the corresponding limits will improve as a result. Even if supersymmetry is not found at the LHC, some argue [98] that large parts of the relevant phase space will still be available to explore at future colliders such as the ILC. Others argue that the first hints of SUSY have already shown up in the form of discrepancies in various Standard Model precision tests. Both the ATLAS [84] and CMS [99, 100] experiments measure a WW production cross-section that is higher than the Standard Model prediction. The excess could be interpreted as being due to contributions from supersymmetric processes [101, 102].

Finally, it should be noted that most searches, both at the LHC and at previous colliders, have focused on the *Minimal* Supersymmetric Standard Model. It is possible to come up with plenty of more exotic models [103–106] that would thus far have escaped detection. Killing supersymmetry entirely is a very difficult thing to do.[†] But regardless of whether or not SUSY exists, *some* manner of physics beyond the Standard Model has to be out there. Because of that, the future of particle physics is sure to be an interesting one.

[†]Opinions differ on whether or not this is a good property for a model to have.

A Additional particle counting tables

Tables 48, 49 and 50 show σ_{charge} for each bunch within the FADC range using fibers with Static, Pre-Train and Dynamic baselines respectively. The numbers have been obtained from runs with a 50 ns bunch spacing. Tables 51, 52 and 53 show the same numbers for standard PMTs, again in runs with a 50 ns bunch spacing. Tables 54 and 55 show the same numbers for fibers in runs with a 25 ns bunch spacing, and Tables 56 and 57 show them for standard PMTs in runs with a 25 ns bunch spacing. Note that the Dynamic baselines are not usable in 25 ns runs, which is why the corresponding tables are not included. The worst-case scenario is considered when calculating σ_{tot} . This number is highlighted in each table.

Tables 58, 59 and 60 contain k , σ_{calib} and σ_{bunch} for each bunch within the FADC range using fibers with Static, PreTrain and Dynamic baselines respectively. The numbers have been obtained from runs with a 50 ns bunch spacing. Tables 61, 62 and 63 show the same numbers for standard PMTs, again in runs with a 50 ns bunch spacing. Tables 64 and 65 show the same numbers for fibers in runs with a 25 ns bunch spacing, and Tables 66 and 67 show them for standard PMTs in runs with a 25 ns bunch spacing. The alert reader will notice that the values of k listed in these tables are close to unity even though Figure 57 suggests that they should be significantly larger. When producing the tables, all $\langle Q \rangle_{\text{LB}}$ were multiplied by a calibration factor before calculating k . The factor was such that k would be approximately equal to unity, making the numbers easier to compare. Note that this does not change the value of σ_{calib} .

Table 48: The table contains σ_{charge} for BCID 66, 68, 70 and 72 in 50 ns runs with identical running conditions using fibers with Static baselines. The worst-case scenario is highlighted.

Run	B66	B68	B70	B72
214651	0.00498	0.00389	0.00526	0.00483
214680	0.00646	0.00574	0.00546	0.00577
214721	0.00506	0.00513	0.00498	0.00497
214758	0.00759	0.00553	0.00597	0.00574
214777	0.00613	0.00502	0.00539	0.00536
215027	0.00637	0.00559	0.00603	0.00609
215061	0.00518	0.00389	0.00438	0.00451
215063	0.006	0.00583	0.0061	0.00626
215091	0.00584	0.00574	0.00588	0.00584
Mean	0.00596	0.00515	0.00549	0.00548

Table 49: The table contains σ_{charge} for BCID 66, 68, 70 and 72 in 50 ns runs with identical running conditions using fibers with PreTrain baselines. The worst-case scenario is highlighted.

Run	B66	B68	B70	B72
214651	0.00509	0.0041	0.0056	0.00502
214680	0.00651	0.00549	0.00555	0.00564
214721	0.00477	0.00501	0.00496	0.00485
214758	0.00631	0.00583	0.0067	0.00628
214777	0.00618	0.00498	0.00547	0.00535
215027	0.00769	0.00698	0.00948	0.00668
215061	0.00507	0.00403	0.00462	0.00452
215063	0.00675	0.0066	0.00787	0.00659
215091	0.00559	0.00575	0.00583	0.00581
Mean	0.00599	0.00542	0.00623	0.00564

Table 50: The table contains σ_{charge} for BCID 66, 68, 70 and 72 in 50 ns runs with identical running conditions using fibers with Dynamic baselines. The worst-case scenario is highlighted.

Run	B66	B68	B70	B72
214651	0.00509	0.00418	0.00568	0.00508
214680	0.00651	0.00757	0.00623	0.00684
214721	0.00477	0.00526	0.00531	0.00533
214758	0.00631	0.00638	0.00657	0.00703
214777	0.00618	0.00613	0.00621	0.00549
215027	0.00769	0.00695	0.00633	0.00658
215061	0.00507	0.00406	0.00494	0.0048
215063	0.00675	0.00699	0.00704	0.00703
215091	0.00559	0.00624	0.0064	0.00633
Mean	0.00599	0.00597	0.00608	0.00606

Table 51: The table contains σ_{charge} for BCID 66, 68, 70 and 72 in 50 ns runs with identical running conditions using standard PMTs with Static baselines. The worst-case scenario is highlighted.

Run	B66	B68	B70	B72
214651	0.00808	0.00903	0.00669	0.00875
214680	0.0124	0.0127	0.00884	0.0122
214721	0.00894	0.0127	0.00698	0.0102
214758	0.0187	0.0161	0.0106	0.0142
214777	0.0103	0.0145	0.00695	0.0136
215027	0.0118	0.0151	0.00919	0.015
215061	0.00673	0.00815	0.006	0.00791
215063	0.0105	0.0134	0.00891	0.0133
215091	0.00978	0.0125	0.00787	0.0106
Mean	0.0108	0.0127	0.00801	0.0118

Table 52: The table contains σ_{charge} for BCID 66, 68, 70 and 72 in 50 ns runs with identical running conditions using standard PMTs with PreTrain baselines. The worst-case scenario is highlighted.

Run	B66	B68	B70	B72
214651	0.00661	0.00499	0.00574	0.00573
214680	0.00954	0.00829	0.00712	0.0084
214721	0.00806	0.00796	0.00625	0.00718
214758	0.0165	0.0103	0.00782	0.00945
214777	0.01	0.0104	0.00686	0.0105
215027	0.00912	0.00847	0.00757	0.0089
215061	0.0064	0.00605	0.00555	0.00626
215063	0.0102	0.00938	0.00821	0.00935
215091	0.0105	0.00999	0.00827	0.00934
Mean	0.00966	0.00842	0.00704	0.00835

Table 53: The table contains σ_{charge} for BCID 66, 68, 70 and 72 in 50 ns runs with identical running conditions using standard PMTs with Dynamic baselines. The worst-case scenario is highlighted.

Run	B66	B68	B70	B72
214651	0.00661	0.00531	0.00741	0.00668
214680	0.00954	0.00818	0.00851	0.00872
214721	0.00806	0.00651	0.00865	0.0069
214758	0.0165	0.00809	0.00995	0.00879
214777	0.01	0.00868	0.00826	0.00892
215027	0.00912	0.0123	0.0096	0.0135
215061	0.0064	0.00617	0.00697	0.00685
215063	0.0102	0.00882	0.011	0.00944
215091	0.0105	0.00844	0.0102	0.00933
Mean	0.00966	0.00805	0.00895	0.00879

Table 54: The table contains σ_{charge} for BCID 301, 303, 305 and 307 in 25 ns runs with identical running conditions using fibers with Static baselines. The worst-case scenario is highlighted.

Run	B301	B303	B305	B307
216416	0.00754	0.00689	0.00774	0.00615
216432	0.0096	0.00912	0.00991	0.00866
Mean	0.00857	0.008	0.00882	0.0074

Table 55: The table contains σ_{charge} for BCID 301, 303, 305 and 307 in 25 ns runs with identical running conditions using fibers with PreTrain baselines. The worst-case scenario is highlighted.

Run	B301	B303	B305	B307
216416	0.00756	0.00686	0.0078	0.00608
216432	0.00954	0.00901	0.00988	0.00862
Mean	0.00855	0.00794	0.00884	0.00735

Table 56: The table contains σ_{charge} for BCID 301, 303, 305 and 307 in 25 ns runs with identical running conditions using standard PMTs with Static baselines. The worst-case scenario is highlighted.

Run	B301	B303	B305	B307
216416	0.0103	0.00862	0.00683	0.00762
216432	0.0201	0.0128	0.012	0.0137
Mean	0.0152	0.0107	0.00943	0.0107

Table 57: The table contains σ_{charge} for BCID 301, 303, 305 and 307 in 25 ns runs with identical running conditions using standard PMTs with PreTrain baselines. The worst-case scenario is highlighted.

Run	B301	B303	B305	B307
216416	0.00963	0.00897	0.00788	0.00801
216432	0.0118	0.0109	0.012	0.00999
Mean	0.0107	0.00996	0.00994	0.009

Table 58: The table contains k for BCID 66, 68, 70 and 72 in 50 ns runs with identical running conditions using fibers with Static baselines. The resulting σ_{calib} and σ_{bunch} are also shown.

Run	B66	B68	B70	B72	σ_{bunch}
214651	0.949	0.999	0.997	1	2.21%
214680	0.965	1	1.01	0.999	1.77%
214721	0.957	0.997	0.995	0.988	1.61%
214758	0.952	0.991	0.996	0.997	1.89%
214777	0.947	0.98	1.01	0.983	2.18%
215027	0.958	0.999	1.01	1	2.02%
215061	0.956	0.99	1.01	1	2.04%
215063	0.948	0.997	1.01	1.01	2.44%
215091	0.957	0.992	1	0.996	1.82%
$\langle k \rangle_{\text{run}}$	0.954	0.994	1	0.997	
$\sigma_{k,\text{run}}$	0.00568	0.00627	0.00601	0.00702	
σ_{calib}	0.596%	0.63%	0.599%	0.704%	

Table 59: The table contains k for BCID 66, 68, 70 and 72 in 50 ns runs with identical running conditions using fibers with PreTrain baselines. The resulting σ_{calib} and σ_{bunch} are also shown.

Run	B66	B68	B70	B72	σ_{bunch}
214651	0.917	0.971	0.969	0.971	2.39%
214680	0.94	0.978	0.989	0.975	1.9%
214721	0.937	0.979	0.976	0.969	1.71%
214758	0.917	0.958	0.965	0.963	2.08%
214777	0.926	0.962	0.988	0.965	2.3%
215027	0.906	0.955	0.962	0.96	2.42%
215061	0.912	0.955	0.972	0.964	2.42%
215063	0.904	0.958	0.967	0.966	2.75%
215091	0.93	0.968	0.98	0.97	1.97%
$\langle k \rangle_{\text{run}}$	0.921	0.965	0.974	0.967	
$\sigma_{k,\text{run}}$	0.0122	0.00878	0.00903	0.00435	
σ_{calib}	1.32%	0.91%	0.927%	0.45%	

Table 60: The table contains k for BCID 66, 68, 70 and 72 in 50 ns runs with identical running conditions using fibers with Dynamic baselines. The resulting σ_{calib} and σ_{bunch} are also shown.

Run	B66	B68	B70	B72	σ_{bunch}
214651	0.917	0.929	0.886	0.895	1.88%
214680	0.94	0.923	0.897	0.887	2.27%
214721	0.937	0.931	0.88	0.886	2.84%
214758	0.917	0.916	0.883	0.889	1.73%
214777	0.926	0.915	0.891	0.886	1.83%
215027	0.906	0.92	0.895	0.891	1.24%
215061	0.912	0.923	0.885	0.893	1.68%
215063	0.904	0.918	0.889	0.892	1.27%
215091	0.93	0.926	0.884	0.892	2.23%
$\langle k \rangle_{\text{run}}$	0.921	0.922	0.888	0.89	
$\sigma_{k,\text{run}}$	0.0122	0.00517	0.00539	0.00323	
σ_{calib}	1.32%	0.561%	0.607%	0.363%	

Table 61: The table contains k for BCID 66, 68, 70 and 72 in 50 ns runs with identical running conditions using standard PMTs with Static baselines. The resulting σ_{calib} and σ_{bunch} are also shown.

Run	B66	B68	B70	B72	σ_{bunch}
214651	0.887	0.917	0.935	0.953	2.62%
214680	0.901	0.927	0.948	0.967	2.64%
214721	0.892	0.919	0.945	0.961	2.81%
214758	0.915	0.933	0.941	0.974	2.24%
214777	0.875	0.902	0.939	0.943	3.06%
215027	0.911	0.925	0.958	0.964	2.35%
215061	0.905	0.905	0.946	0.95	2.32%
215063	0.919	0.939	0.958	0.977	2.29%
215091	0.908	0.931	0.956	0.979	2.82%
$\langle k \rangle_{\text{run}}$	0.901	0.922	0.947	0.963	
σ	0.0135	0.0116	0.00796	0.0117	
σ_{calib}	1.5%	1.26%	0.84%	1.22%	

Table 62: The table contains k for BCID 66, 68, 70 and 72 in 50 ns runs with identical running conditions using standard PMTs with PreTrain baselines. The resulting σ_{calib} and σ_{bunch} are also shown.

Run	B66	B68	B70	B72	σ_{bunch}
214651	0.708	0.753	0.773	0.784	3.86%
214680	0.717	0.759	0.78	0.791	3.72%
214721	0.705	0.75	0.773	0.783	4.01%
214758	0.724	0.755	0.772	0.79	3.18%
214777	0.694	0.746	0.774	0.782	4.61%
215027	0.72	0.763	0.79	0.798	3.95%
215061	0.709	0.748	0.778	0.785	3.99%
215063	0.722	0.763	0.785	0.796	3.7%
215091	0.717	0.761	0.784	0.798	4.03%
$\langle k \rangle_{\text{run}}$	0.713	0.755	0.779	0.79	
σ	0.00915	0.00611	0.00594	0.00604	
σ_{calib}	1.28%	0.809%	0.763%	0.765%	

Table 63: The table contains k for BCID 66, 68, 70 and 72 in 50 ns runs with identical running conditions using standard PMTs with Dynamic baselines. The resulting σ_{calib} and σ_{bunch} are also shown.

Run	B66	B68	B70	B72	σ_{bunch}
214651	0.708	0.653	0.636	0.628	4.77%
214680	0.717	0.658	0.631	0.633	5.25%
214721	0.705	0.647	0.63	0.623	4.97%
214758	0.724	0.643	0.637	0.616	6.3%
214777	0.694	0.635	0.612	0.614	5.23%
215027	0.72	0.652	0.636	0.626	5.57%
215061	0.709	0.658	0.631	0.63	4.86%
215063	0.722	0.661	0.65	0.633	5.02%
215091	0.717	0.66	0.64	0.633	4.98%
$\langle k \rangle_{\text{run}}$	0.713	0.652	0.633	0.626	
$\sigma_{k,\text{run}}$	0.00915	0.00818	0.00966	0.00699	
σ_{calib}	1.28%	1.25%	1.52%	1.12%	

Table 64: The table contains k for BCID 301, 303, 305 and 307 in 25 ns runs with identical running conditions using fibers with Static baselines. The resulting σ_{calib} and σ_{bunch} are also shown.

Run	B301	B303	B305	B307	σ_{bunch}
216416	1.86	1.95	2.07	2.01	3.93%
216432	1.88	1.96	2.08	2.03	3.73%
$\langle k \rangle_{\text{run}}$	1.87	1.95	2.07	2.02	
$\sigma_{k,\text{run}}$	0.0112	0.00347	0.00359	0.0113	
σ_{calib}	0.598%	0.178%	0.173%	0.56%	

Table 65: The table contains k for BCID 301, 303, 305 and 307 in 25 ns runs with identical running conditions using fibers with PreTrain baselines. The resulting σ_{calib} and σ_{bunch} are also shown.

Run	B301	B303	B305	B307	σ_{bunch}
216416	1.85	1.94	2.06	2	3.93%
216432	1.88	1.95	2.07	2.02	3.72%
$\langle k \rangle_{\text{run}}$	1.86	1.94	2.06	2.01	
$\sigma_{k,\text{run}}$	0.0143	0.00571	0.00721	0.0134	
σ_{calib}	0.766%	0.294%	0.349%	0.666%	

Table 66: The table contains k for BCID 301, 303, 305 and 307 in 25 ns runs with identical running conditions using standard PMTs with Static baselines. The resulting σ_{calib} and σ_{bunch} are also shown.

Run	B301	B303	B305	B307	σ_{bunch}
216416	0.267	0.29	0.32	0.314	7.08%
216432	0.276	0.298	0.327	0.324	6.83%
$\langle k \rangle_{\text{run}}$	0.271	0.294	0.324	0.319	
$\sigma_{k,\text{run}}$	0.00443	0.00412	0.00335	0.0053	
σ_{calib}	1.63%	1.4%	1.03%	1.66%	

Table 67: The table contains k for BCID 301, 303, 305 and 307 in 25 ns runs with identical running conditions using standard PMTs with PreTrain baselines. The resulting σ_{calib} and σ_{bunch} are also shown.

Run	B301	B303	B305	B307	σ_{bunch}
216416	0.253	0.279	0.305	0.301	7.28%
216432	0.254	0.281	0.305	0.305	7.37%
$\langle k \rangle_{\text{run}}$	0.253	0.28	0.305	0.303	
$\sigma_{k,\text{run}}$	0.000191	0.000997	0.000371	0.00198	
σ_{calib}	0.0753%	0.356%	0.122%	0.653%	

B Monte Carlo generators

This section lists the processes that are included in each background category as well as the Monte Carlo generators used to model them. It also summarizes the Monte Carlo generator combinations used to estimate the modelling uncertainty. Tables 68 to 73 show the processes included in the WW , WZ , top, non-prompt, Z +jets and Higgs backgrounds, respectively. Tables 74 to 76 show the generators used to evaluate the WW , ZV and top modelling uncertainties, respectively.

All **AlpGen** v2.14 [107] and **Madgraph5** v1.3.33 [108] samples use **Pythia** v6.426 [109] to model the parton shower. The **Powheg-Box** v1.0 [110] samples instead use **Pythia** v8.865 [109]. In the case of the **MC@NLO** v4.06 [111–113] samples, those used for the background estimation employ **Pythia** v6.426 for the parton shower. The rest of the **MC@NLO** samples are used to evaluate the parton shower modelling uncertainty, with the varied samples being interfaced with **Herwig** v6.520 [114] instead. All leading order generators use the CTEQ6L1 PDF set [115] while next-to-leading order generators use CT10 [92]. In almost all cases, the detector simulation is implemented fully in **GEANT4** [116, 117]. The exceptions, which are noted in Table 72, use the faster **AFII** simulation [118, 119] to model the calorimeters and **GEANT4** elsewhere.

The effect of having multiple pp interactions per bunch crossing (that is, $\mu > 1$) is simulated by overlaying minimum-bias events onto the hard scattering. This part of the simulation is called the *underlying event*. In case the parton shower is handled by **Herwig**, the underlying event is modelled by **Jimmy** v4.31 [120]. In all other cases, the underlying event is instead modelled by **Pythia** using ATLAS tune AUET2B [121]. Since most Monte Carlo samples are generated prior to data taking, the distribution of μ in Monte Carlo simulation does not necessarily match the one in real ATLAS data. This discrepancy is corrected for by comparing the μ -distributions in data and Monte Carlo simulation and applying a weight ω_μ to every Monte Carlo event such that $\omega_\mu > 1$ for events with a value of μ that is underrepresented in Monte Carlo simulation and $\omega_\mu < 1$ for events that are overrepresented.

Finally, although they are not shown in any tables, the SUSY signal samples were generated using **Herwig++** v2.5.2 [122] with the CTEQ6L1 PDF set. All signal cross-sections were calculated using **Prospino2.1** [123].

Table 68: The table shows all processes included in the WW background as well as the generators used to model each process.

Process	Generator	Comments
WW	Powheg-Box v1.0	
WW	gg2WW v3.1.2 [124]	Via gluon-gluon fusion
$WWjj$	Sherpa v1.4.1 [125]	
$W^\pm W^\pm$	Sherpa v1.4.1	Same-sign production
$W^\pm W^\pm jj$	Sherpa v1.4.1	Same-sign production
$VV \rightarrow \ell\nu qq$	Sherpa v1.4.1	Includes both WW and WZ
WWW^*	Madgraph5 v1.3.33	

Table 69: The table shows all processes included in the ZV background as well as the generators used to model each process.

Process	Generator	Comments
WZ	Powheg-Box v1.0	
ZZ	Powheg-Box v1.0	
ZZ	gg2ZZ v3.1.2 [126]	Via gluon-gluon fusion
$VV \rightarrow \ell\ell qq$	Sherpa v1.4.1	Includes both WZ and ZZ
ZWW*	Madgraph5 v1.3.33	
ZZZ*	Madgraph5 v1.3.33	
WZjj	Sherpa v1.4.1	

Table 70: The table shows all processes included in the top background as well as the generators used to model each process. A top mass of 172.5 GeV is assumed.

Process	Generator	Comments
$t\bar{t}$	MC@NLO v4.06	
Wt	MC@NLO v4.06	
Single-top	MC@NLO v4.06	s-channel production
Single-top	AcerMC v3.8 [127]	t-channel production
$t\bar{t}V$	Madgraph5 v1.3.33	

Table 71: The table shows various processes that are considered by the Matrix Method when estimating the non-prompt background as well as the generators used to model each process.

Process	Generator	Comments
W+jets	Alpgen v2.14	
Wbb	Alpgen v2.14	
Wcc	Alpgen v2.14	
Wc	Alpgen v2.14	
$W + \gamma$	Sherpa v1.4.1	
$b\bar{b}$	Pythia v8.165	
$c\bar{c}$	Pythia v8.165	

Table 72: The table shows all processes included in the Z+jets background as well as the generators used to model each process. The Drell-Yan $Z \rightarrow e\bar{e}$ and $Z \rightarrow \mu\mu$ samples are simulated using AFII. Each sample is used only in the mass range indicated.

Process	Generator	Comments
Drell-Yan	Sherpa v1.4.1	$m_{\ell\ell} < 40$ GeV
Z	Sherpa v1.4.1	$40 < m_{\ell\ell} < 60$ GeV
Z	Alpgen v2.14	$m_{\ell\ell} > 60$ GeV
Zcc	Alpgen v2.14	$m_{\ell\ell} > 60$ GeV
Zbb	Alpgen v2.14	$m_{\ell\ell} > 60$ GeV

Table 73: The table shows all processes included in the Higgs background as well as the generators used to model each process.

Process	Generator	Comments
WH	Pythia v8.165	
ZH	Pythia v8.165	
Vector boson fusion	Powheg-Box v1.0	
Gluon-gluon fusion	Powheg-Box v1.0	

Table 74: The table shows all the generators used in the estimation of the WW modelling uncertainty as well as the sources of uncertainty they are used to estimate.

Generator	Uncertainty
Powheg+Pythia8	PDF
aMC@NLO+Herwig	HS, PS, QCD scale
aMC@NLO+Pythia6	HS, PS
Powheg+Herwig	HS, PS
Powheg+Pythia6	HS, PS

Table 75: The table shows all the generators used in the estimation of the ZV modelling uncertainty as well as the sources of uncertainty they are used to estimate.

Generator	Uncertainty
Powheg+Pythia8	HS + PS, PDF
Sherpa	HS + PS
aMC@NLO+Herwig	QCD Scale

Table 76: The table shows all the generators used in the estimation of the top modelling uncertainty as well as the sources of uncertainty they are used to estimate.

Generator	Uncertainty
MC@NLO+Jimmy	HS, PDF
Powheg+Jimmy	HS, PS
Powheg+Pythia6	PS
Powheg+Pythia6	QCD Scale
AcerMC	ISR+FSR

C Limit setting techniques

The ultimate goal of the analysis described in Section 9 is to say something about the existence, or non-existence, of supersymmetry. Such statements can be formulated in terms of hypothesis tests. Hypotheses in particle physics are typically represented by Lagrangians. In order to figure out which Lagrangian is the null hypothesis and which Lagrangian is the alternative, it is important to first understand the purpose of the test. A test that tries to establish the existence of supersymmetry is called a test for *discovery*. In this case, the null hypothesis is that nature is described by the Standard Model Lagrangian (see Section 1). The alternative hypothesis is that nature is *not* described by the SM, in which case it is hopefully described by some manner of supersymmetric Lagrangian (see Section 2). Such a test either rejects or accepts the Standard Model. A test that tries to establish the non-existence of a particular SUSY model is called a test for *exclusion*. In this case, the null hypothesis is that nature is described by the model in question and the alternative hypothesis is that the Standard Model is correct. Exclusion tests can be used to set upper limits, like the ones presented in Section 9.7, on the cross-section for physics beyond the Standard Model. They are also used to produce exclusion contours like the ones presented in Section 9.8.

Regardless of whether the goal is discovery or exclusion, the hypothesis test itself works in the same way. The analysis defines a special function, called a *test-statistic*, that is designed to take on different values under the null hypothesis compared to the alternative hypothesis. The greater the difference, the easier it is to distinguish between the hypotheses, and the more powerful the analysis. This section describes how the test-statistic is defined and interpreted in order to set limits on supersymmetry.

C.1 Hypothesis testing

A hypothesis is tested by conducting an experiment. If the outcome of that experiment is sufficiently unlikely under the assumption that the hypothesis is correct, then the hypothesis is rejected. The degree of confidence in a given hypothesis is usually quantified by calculating its *p*-value. A *p*-value is defined as the probability to obtain an outcome that is equally or more extreme if the experiment were to be repeated. Whether or not an outcome is *extreme* depends on the definition of the test-statistic. Some test-statistics are defined in such a way that very large values indicate poor agreement with the hypothesis, but very small values do not. The *p*-value is then given by the *one-sided* probability to observe an equally or more extreme outcome. This is illustrated in Figure 91 (a). Other test-statistics are defined in such a way that very large values and very small values both indicate poor agreement with the hypothesis. The *p*-value is then given by the *two-sided* probability to observe an equally or more extreme outcome. This is illustrated in Figure 91 (b).

Rejecting a hypothesis if the outcome is *sufficiently unlikely* is obviously not a well-defined criterion. Physicists therefore agree on precisely how unlikely the outcome needs to be *before* carrying out the experiment. They also agree in advance on whether a one-sided or two-sided *p*-value should be used. A typical criterion when testing a model for exclusion is to reject the null hypothesis if $p < 5\%$. The model is then said to be rejected at a *confidence level* (CL) of 95%. When testing for discovery, the null hypothesis is the Standard Model. The criterion for rejection is therefore much more stringent. A typical choice is $p < 2.87 \cdot 10^{-7}$. This number, while it may look somewhat arbitrary, has a well-defined origin. It corresponds to the one-sided probability for a Gaussian-distributed random variable to fluctuate five standard deviations above its mean. In particle physics, *p*-values are often expressed as the equivalent number of Gaussian standard deviations *Z*. Equation (202) can be used to convert between *p* and *Z*.

$$Z = \begin{cases} \sqrt{2} \operatorname{erf}^{-1}(1 - 2p) & \text{if } p \text{ is one-sided} \\ \sqrt{2} \operatorname{erf}^{-1}(1 - p) & \text{if } p \text{ is two-sided} \end{cases} \quad (202)$$

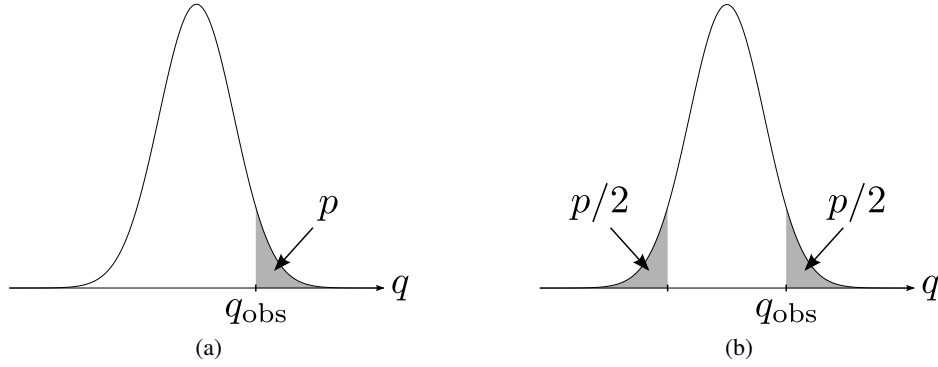


Figure 91: The figure shows the probability density function (pdf) of a test-statistic q under some hypothesis. An experiment has yielded the outcome q_{obs} . The one-sided (a) and two-sided (b) p -values are given by the shaded areas under the curves. They correspond to the probability of obtaining an outcome that is equally or more extreme if the experiment were to be repeated.

C.2 Constructing the test-statistic

Performing an experiment ultimately results in one or more observables, collectively referred to as the *outcome* of the experiment. A typical example of an observable is the number of events observed in some signal region. Since this is a random variable, running the experiment repeatedly would not yield the same outcome every time. A test-statistic q is a function whose purpose is to order the outcomes from least to most extreme. There are two kinds of test-statistics. A *simple* test-statistic is a function with no free parameters. Hypothesis testing with a simple test-statistic is done by building a Monte Carlo model of the experiment and running it many times in order to obtain a probability density function for q . The real experiment is then run once, and the observed value of the test-statistic q_{obs} is compared to the expected distribution as in Figure 91. The hypothesis can then be either rejected or accepted. A *composite* test-statistic is a function with free parameters. It can be used not only to reject or accept hypotheses but also to obtain confidence intervals for the parameters in question. This section discusses composite test-statistics.

The first step towards constructing the test-statistics used for the analysis in Section 9 is to define the *likelihood function* $\mathcal{L}$. A likelihood is a function that models the outcome of an experiment. The input to the likelihood function is the values of the parameters used in the model, and the output is the probability to observe a given outcome. For example, consider a simple experiment with a single signal region whose purpose is to exclude a specific SUSY model. The expected signal and background have been estimated to s and b events, respectively. The total number of events expected in the signal region is therefore given by

$$\lambda_s(\mu) = \mu s + b \quad (203)$$

where the *signal strength parameter* μ is a convenient way to express the strength of the signal in a standardized manner. A signal strength $\mu = 0$ corresponds to the background-only hypothesis while $\mu = 1$ corresponds to the nominal signal hypothesis. Since the number of observed events is a Poisson-distributed variable, the likelihood function becomes

$$\mathcal{L}(n_s|\mu) = \mathcal{P}(n_s|\lambda_s(\mu)) = \frac{e^{-(\mu s + b)} (\mu s + b)^{n_s}}{n_s!} \quad (204)$$

where n_s is the number of events that were observed in the signal region after running the experiment. The likelihood function serves two purposes, both of which intuitively connect it to its role in the construction

of the test-statistic. Firstly, $\mathcal{L}$ can be used to determine how compatible an observed outcome is with a given signal hypothesis (i.e. a given value of μ). Secondly, $\mathcal{L}$ can be used to determine what set of parameters best describe the outcome that was observed. This is done by finding the set of parameters that maximize the likelihood.

In a real analysis, a more realistic model is used. Such a model should account for the interplay between different backgrounds by using a simultaneous fit as motivated in Section 9.5.6. It should also include the impact of systematic uncertainties. This is accomplished by adding new factors and additional free parameters to the likelihood function. For the sake of terminology, a distinction is made between μ and the other parameters. The signal strength μ , which encodes all the information about the existence or non-existence of supersymmetry, is called the *parameter of interest*. The other parameters, while necessary for the model to work, do not contain information about supersymmetry. They are therefore referred to as *nuisance parameters*. There is no technical difference between the parameter of interest and the nuisance parameters — The grouping is simply a matter of what information the physicists are trying to extract from the experiment.

Each background process that is included in the simultaneous fit comes with a factor

$$\mathcal{L}(n_i|\mu, \mathbf{b}) = \mathcal{P}(n_i|\lambda_i(\mu, \mathbf{b})) \quad (205)$$

where the subscript i enumerates the control regions, such that n_i and λ_i are the observed and expected number of events in the i :th control region, respectively. The expected number of events λ_i depends on μ (because of contamination from SUSY processes) and on the normalization of each background. There are three background normalization parameters b_{WW} , b_{WZ} and b_{Top} . They are collectively denoted $\mathbf{b}$.

The impact of systematic uncertainties is modelled by including yet more information in the likelihood function. Each source of uncertainty comes with a nuisance parameter θ_{Syst} whose effect on the analysis has been determined by an auxiliary measurement (see Section 9.6). All event counts are parametrized in terms of θ_{Syst} so that they shift as appropriate with the systematic. It is convenient to normalize the nuisance parameters such that their nominal value is $\theta_0 = 0 \pm 1$. That is, the nuisance parameters are defined such that they have a nominal value of zero and a standard deviation of unity. For example, if an auxiliary measurement has determined the impact of the jet uncertainty to be $\pm 5\%$, then setting $\theta_{\text{Jet}} = 0.2$ in the likelihood function would shift the background up by 1%. The value of θ_{Jet} is then constrained by adding an additional factor to the likelihood.

$$\mathcal{L}(0|\theta_{\text{Jet}}) = \mathcal{G}(0|\theta_{\text{Jet}}, 1) \quad (206)$$

The factor $\mathcal{G}(0|\theta_{\text{Jet}}, 1)$ is the probability to obtain zero from a Gaussian distribution with a mean of θ_{Jet} and a standard deviation of unity. When dealing with multiple sources of uncertainty, the corresponding nuisance parameters are collectively denoted $\boldsymbol{\theta}$.

The full likelihood function is given by

$$\mathcal{L}(\mathbf{n}|\mu, \mathbf{b}, \boldsymbol{\theta}) = \mathcal{P}(n_s|\lambda_s(\mu, \mathbf{b}, \boldsymbol{\theta})) \prod_{i \in \text{CR}} \mathcal{P}(n_i|\lambda_i(\mu, \mathbf{b}, \boldsymbol{\theta})) \prod_{j \in \text{Syst}} \mathcal{G}(0|\theta_j, 1) \quad (207)$$

where $\mathbf{n}$ denotes the set of observed event counts n_s and n_i in the signal and control regions. The aim is to use the likelihood function in order to construct test-statistics from which limits on μ can be derived. Let $\hat{\mu}$, $\hat{\mathbf{b}}$ and $\hat{\boldsymbol{\theta}}$ denote the values of μ , $\mathbf{b}$ and $\boldsymbol{\theta}$ that maximize the likelihood given an outcome $\mathbf{n}$. These are the parameters that are most likely to be true given the observed data.

$$\mathcal{L}(\mathbf{n}|\hat{\mu}, \hat{\mathbf{b}}, \hat{\boldsymbol{\theta}}) \geq \mathcal{L}(\mathbf{n}|\mu, \mathbf{b}, \boldsymbol{\theta}) \quad \forall \quad \mu, \mathbf{b}, \boldsymbol{\theta} \quad (208)$$

Intuitively, a confidence interval for μ is constructed by determining how much it can be varied from its most likely value before the model becomes incompatible (at a given CL) with the observed data. Let $\hat{\mathbf{b}}$

and $\hat{\theta}$ denote the values of $\mathbf{b}$ and θ that maximize the likelihood given a choice of $\mu \neq \hat{\mu}$.

$$\mathcal{L}(n|\mu, \hat{\mathbf{b}}, \hat{\theta}) \geq \mathcal{L}(n|\mu, \mathbf{b}, \theta) \quad \forall \quad \mathbf{b}, \theta \quad (209)$$

Constraining the nuisance parameters in this way is sometimes referred to as *profiling*. Once the above likelihood functions are defined, the Neyman-Pearson lemma [128] provides the recipe for constructing the test-statistics. The lemma states that the most powerful test-statistic available for distinguishing between two hypotheses is a *likelihood ratio*. In ATLAS, the following likelihood ratio is used:

$$\text{LR} = \frac{\mathcal{L}(n|\mu, \hat{\mathbf{b}}, \hat{\theta})}{\mathcal{L}(n|\hat{\mu}, \hat{\mathbf{b}}, \hat{\theta})} \quad (210)$$

The Neyman-Pearson lemma technically holds only for simple test-statistics, but likelihood ratios are commonly used for composite test-statistics as well. Equivalently, a log-likelihood ratio can be used.

$$\text{LLR} = -2 \log \left(\frac{\mathcal{L}(n|\mu, \hat{\mathbf{b}}, \hat{\theta})}{\mathcal{L}(n|\hat{\mu}, \hat{\mathbf{b}}, \hat{\theta})} \right) \quad (211)$$

Defining a test-statistic to be a log-likelihood ratio has several advantages over using a plain likelihood ratio. The information content of the LR and the LLR are, of course, exactly the same. The advantages are purely technical. Firstly, a log-likelihood ratio is a sum whereas a plain likelihood ratio is a product. Maximizing a sum is computationally less demanding than maximizing a product. Secondly, minimization algorithms are generally more efficient than maximization algorithms. The LLR is therefore defined as a *negative* log. Thirdly, it can be shown [129, 130] that $-2 \log(\text{LR})$ follows a χ^2 -distribution with one degree of freedom in the limit $n \rightarrow \infty$. The negative log is therefore multiplied by a factor two to make it more convenient to use in a χ^2 -test.

Arriving at the definition of the test-statistics used for the analysis in Section 9 requires jumping through a few more hoops. It is clear that supersymmetric processes should only be able to *add* events to the signal regions. They should never be able to reduce the number of events. The following definitions therefore assume that only models with $\mu \geq 0$ will be tested.

When testing for discovery, a large excess of events should be interpreted as evidence for supersymmetry but a large deficit should not. The quantity

$$q_{\text{CI}} = \begin{cases} -2 \log \frac{\mathcal{L}(n|\mu, \hat{\mathbf{b}}, \hat{\theta})}{\mathcal{L}(n|0, \hat{\mathbf{b}}, \hat{\theta})} & \text{if } \hat{\mu} < 0 \\ -2 \log \frac{\mathcal{L}(n|\mu, \hat{\mathbf{b}}, \hat{\theta})}{\mathcal{L}(n|\hat{\mu}, \hat{\mathbf{b}}, \hat{\theta})} & \text{if } \hat{\mu} \geq 0 \end{cases} \quad (212)$$

declares any outcome that would be best described by a negative signal strength to be maximally compatible with the hypothesis $\mu = 0$. The test-statistic used when testing for discovery is obtained by setting $\mu = 0$ in Eq. (212). This definition guarantees that the Standard Model can never be rejected in favor of supersymmetry as a result of a deficit of events.

When testing a model for exclusion, the aim is to set an upper limit on μ . Appendix C.3 describes the technique used to set such limits. In general, q_{CI} would be an appropriate test-statistic to use if the aim were to find a confidence interval for μ . A test-statistic appropriate for upper limits can be obtained by, in addition, requiring $\mu \geq \hat{\mu}$.

$$q_{\text{UL}} = \begin{cases} -2 \log \frac{\mathcal{L}(n|\mu, \hat{\mathbf{b}}, \hat{\theta})}{\mathcal{L}(n|0, \hat{\mathbf{b}}, \hat{\theta})} & \text{if } \hat{\mu} < 0 \\ -2 \log \frac{\mathcal{L}(n|\mu, \hat{\mathbf{b}}, \hat{\theta})}{\mathcal{L}(n|\hat{\mu}, \hat{\mathbf{b}}, \hat{\theta})} & \text{if } \mu \geq \hat{\mu} \geq 0 \\ 0 & \text{if } \mu < \hat{\mu} \end{cases} \quad (213)$$

The test-statistic used when testing for exclusion is precisely Eq. (213). Hypotheses with $\mu < \hat{\mu}$ are impossible to exclude by definition, and the limits on μ are consequently one-sided.

C.3 Setting the limits

In order to understand how limits are set on the parameters in the likelihood function, it is instructive to begin by considering the simple case of a parameter of interest μ together with single nuisance parameter θ . As explained in Appendix C.2, μ is the only free parameter in the likelihood function. Nuisance parameters are not free. Instead, they are *profiled* such that they are always equal to their maximum likelihood estimators $\hat{\theta}$. Suppose that an experiment has been conducted and yielded some outcome. The parameters that maximize the likelihood of that outcome are $\hat{\mu}$ and $\hat{\theta}$. Other signal hypotheses (i.e. other values of μ) are less favored by the experiment. Some hypotheses might be so unlikely that they can be excluded. Indeed, *limits* are formally defined as inverted hypothesis tests. Consider the test-statistic q_{CI} as defined by Eq. (212). This test-statistic is non-zero both for $\mu < \hat{\mu}$ and $\mu > \hat{\mu}$. It can therefore be used to set both upper and lower limits on μ . In contrast, the test-statistic q_{UL} , which is defined in Eq. (213), is identically zero for $\mu < \hat{\mu}$. It can therefore be used only to set upper limits.

Figure 92 (a) illustrates how confidence intervals can be obtained using q_{CI} . The line shows how $\hat{\theta}$ depends on μ . The point $(\hat{\mu}, \hat{\theta})$ maximizes the likelihood and therefore corresponds to $q = 0$. All other values of μ , being less compatible with the observed outcome, yield $q > 0$. To calculate the p -value of a given point along the line, it is necessary to know the probability density function of q under the corresponding signal hypothesis. The p -value is then calculated according to

$$p = \int_{q^{\text{point}}}^{\infty} f(q|\mu) dq \quad (214)$$

where $f(q|\mu)$ is the pdf of q given the signal hypothesis μ and q^{point} is the value of q in the point of interest. Suppose that the dotted segment of the line covers all points that satisfy $p > 0.32$ and the dashed segment covers all points that satisfy $p > 0.05$. Any value of μ outside of the dotted region can then be excluded at a 68% CL given the observed data. The resulting confidence interval would be approximately $\mu = 0.42^{+0.11}_{-0.14}$. In the same way, any value of μ outside of the dashed region can be excluded at a 95% CL.

Figure 92 (b) illustrates how upper limits are obtained using q_{UL} . The procedure is analogous to that of confidence intervals obtained with q_{CI} . A p -value can be calculated for any point along the line using Eq. (214). Suppose that the dotted segment of the line covers all points that satisfy $p > 0.32$ and the dashed segment covers all points that satisfy $p > 0.05$. Any value of μ above the dotted region can then be excluded at a 68% CL given the observed data. The result is an upper limit of approximately $\mu < 0.51$. In the same way, any value of μ above the dashed region can be excluded at a 95% CL. The hashed region represents points that satisfy $\mu < \hat{\mu}$. They can not be excluded since they correspond to $q = 0$ by definition. Removing these points from the tail of the pdf allows additional points with $\mu > \hat{\mu}$ to be excluded instead. This is what makes an upper limit obtained using q_{UL} stronger than an upper limit obtained using q_{CI} (at the cost of not getting any lower limit).

Two different methods are commonly used in order to find the probability density function $f(q|\mu)$. The first method is to simulate a large number of Monte Carlo experiments under the signal hypothesis μ . By calculating q for each experiment, a probability density function is obtained. This method is robust in the sense that $f(q|\mu)$ can be determined to an arbitrary precision in the limit of infinite Monte Carlo experiments. The tradeoff is that generating many Monte Carlo experiments can be computationally demanding. The second method is to use the Wilks [129] and Wald [130] approximations. Wilks theorem states that $f(q|\hat{\mu})$, in the limit of infinite statistics, follows a χ^2 -distribution with one degree of freedom as shown in Figure 93. Wald generalized this result by showing that $f(q|\mu)$, with $\mu \neq \hat{\mu}$, follows a *noncentral* χ^2 -distribution with one degree of freedom. The noncentrality parameter can be obtained numerically, yielding a fully defined pdf. Using this approximation is much quicker than using Monte Carlo experiments. Particle physicists often refer to it as the *Asimov approximation* or the *asymptotic formula*. In practice, 95% CL limits set with the asymptotic formula are astonishingly accurate even with

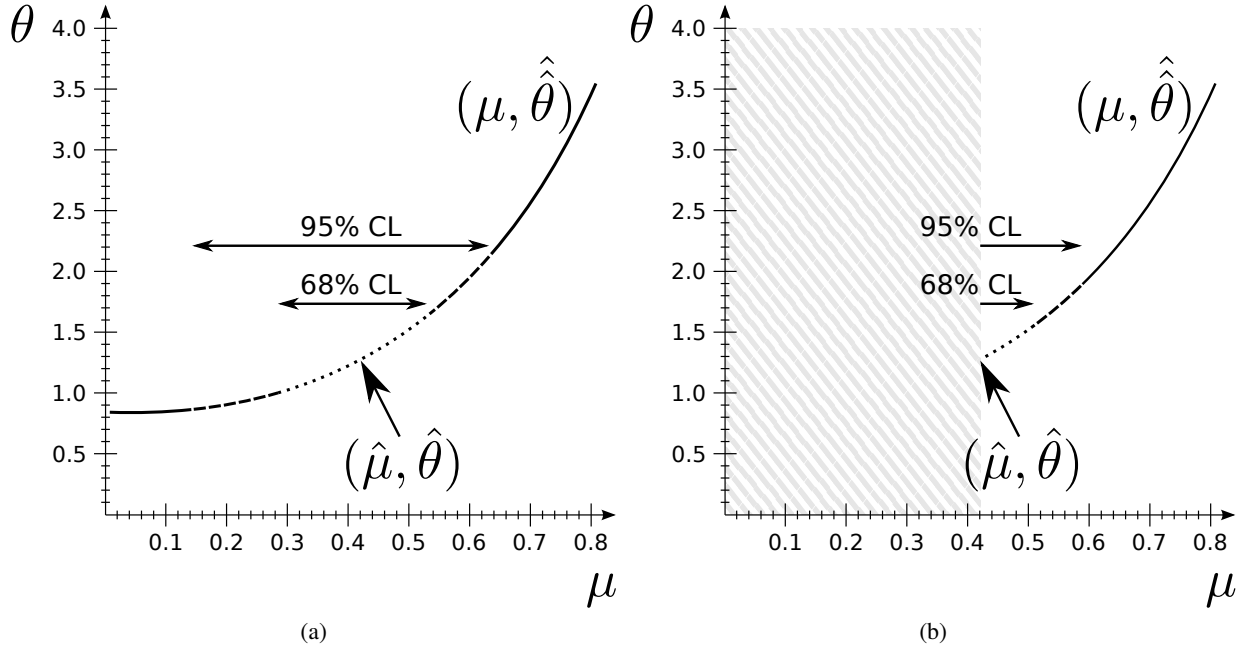


Figure 92: The figure illustrates how a confidence interval (a) and an upper limit (b) on the parameter of interest can be obtained in the case of a single nuisance parameter.

very few events in a signal region [131]. This accuracy drops off if a higher CL is required. The limits quoted in Section 9 are obtained with the asymptotic formula and cross-checked using Monte Carlo simulations.

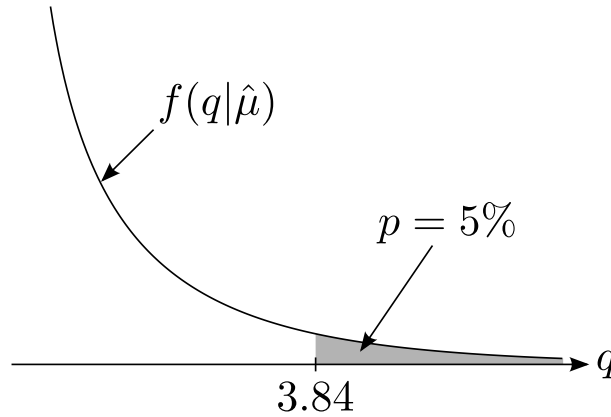


Figure 93: In the asymptotic limit, $f(q|\hat{\mu})$ follows χ^2 -distribution with one degree of freedom. Once a pdf is explicitly known, finding the value of q that corresponds to a given p -value is simple. For example, if q follows a χ^2 -distribution with one degree of freedom as in the figure, $p = 5\%$ corresponds to $q = 3.84$.

Profiling with multiple nuisance parameters follows the same principle as profiling with one nuisance parameter. Simply add one new dimension per nuisance parameter to the plots in Figure 92. Intuition dictates that increasing the number of nuisance parameters (i.e. the number of systematics) should necessarily weaken the limits on μ . This is indeed the case. Each new nuisance parameter adds additional flexibility to the model, thereby increasing the range of μ -values that can be made compatible with the observed data.

Finally, recall that all nuisance parameters are normalized such that they have a nominal value $\theta_0 = 0 \pm 1$. The maximum likelihood fit might find that the values of θ that best describe the observed data are different from the nominal values. This difference is referred to as the *pull*. In general, pull is defined as

$$g = \frac{\theta_0 - \hat{\theta}}{\sigma_{\theta_0}} \quad (215)$$

which, given the normalization of θ_0 , reduces to just $\hat{\theta}$. The uncertainty on the pull is evaluated in the following way. First, the value of μ is fixed such that $\mu = \hat{\mu}$. The nuisance parameter θ is then treated as if it were a parameter of interest. Limits are obtained by varying θ while profiling the remaining nuisance parameters, analogously to how limits on μ were obtained in Figure 92. This procedure is then repeated for each nuisance parameter in the likelihood function. The numbers quoted in Table 35 reflect the impact of varying each nuisance parameter within its fitted uncertainty.

C.4 The CL_s technique

An observed upper limit depends on statistical fluctuations in data. If the background happens to fluctuate downwards in a particular signal region, the observed limit will be stronger than expected. Given a large enough downward fluctuation, it is in principle possible to exclude arbitrarily small signals. Particle physicists find this property of the limits to be undesirable. Statistical outliers are bound to show up every now and then, and they should not be interpreted as evidence against new physics. Doing so, the particle physicist argues, would be no different from throwing dice. The ATLAS collaboration has therefore adopted the use of the so-called *modified frequentist* or CL_s technique [94, 95]. This is a method designed to make models harder to exclude if they are not well-resolved from the background-only hypothesis.

Suppose that a signal hypothesis $\mu \neq 0$ is to be tested for exclusion using the CL_s method. The first step is to obtain $f(q|\mu)$, which is the probability density function of q under the signal-plus-background hypothesis. This can be done by calculating q for a large number of Monte Carlo experiments simulated using a signal strength parameter of μ . The next step is to obtain $f(q|0)$, which is the probability density function of q under the background-only hypothesis. This can be done by calculating q for a large number of Monte Carlo experiments simulated using a signal strength parameter of zero. It is worth pointing out that μ is set to zero only when simulating the number of events observed by the Monte Carlo experiments. When computing the value of the test-statistic q , a signal strength of μ is still used.

Normally, the signal-plus-background hypothesis would be rejected if it had a sufficiently small p -value. In the context of the CL_s method, the p -value of the signal-plus-background hypothesis is called CL_{s+b} .

$$\text{CL}_{s+b} = \int_{q^{\text{point}}}^{\infty} f(q|\mu) dq \quad (216)$$

The idea of CL_s is that the signal-plus-background hypothesis should not be excluded if the outcome of the experiment was extreme also under the background-only hypothesis. Let $1 - \text{CL}_b$ denote the p -value of the background-only hypothesis.

$$1 - \text{CL}_b = \int_{q^{\text{point}}}^{\infty} f(q|0) dq \quad (217)$$

When using the CL_s technique, the p -value of the signal-plus-background hypothesis is penalized by normalizing it to the p -value of the background-only hypothesis. Let CL_s denote the ratio of the two p -values.

$$\text{CL}_s = \frac{\text{CL}_{s+b}}{1 - \text{CL}_b} \quad (218)$$

Hypothesis tests are then done by replacing any p -value with its corresponding CL_s value. A hypothesis is rejected at a 95% CL if it satisfies $CL_s < 0.05$. Small values of CL_s are obtained only if the outcome of the experiment is sufficiently extreme under the signal-plus-background hypothesis compared to the background-only hypothesis. A large downward fluctuation of the background, which is unlikely both under the signal-plus-background and background-only hypotheses, will not result in a small value of CL_s . This makes hypothesis tests based on CL_s robust against the problem of experiments excluding signals to which they are not really sensitive. The tradeoff is that limits set using CL_s are weaker than limits set using CL_{s+b} . Formally, the limits are said to *overcover*, which means that they cover the true value of the parameter more often than the stated confidence level. Finally, it is worth pointing out that CL_s as presented here is only used when setting *upper* limits. A test based on CL_s can not exclude the background-only hypothesis by definition.

References

- [1] ATLAS Collaboration, *Improved luminosity determination in pp collisions at $\sqrt{s} = 7$ TeV using the ATLAS detector at the LHC*, *Eur.Phys.J.* **C73** (2013) 2518, [arXiv:1302.4393 \[hep-ex\]](#).
- [2] ATLAS Collaboration, *Search for direct slepton and gaugino production in final states with two leptons and missing transverse momentum with the ATLAS detector in pp collisions at $\sqrt{s} = 7$ TeV*, *Phys.Lett.* **B718** (2013) 879–901, [arXiv:1208.2884 \[hep-ex\]](#).
- [3] ATLAS Collaboration, *Search for direct production of charginos, neutralinos and sleptons in final states with two leptons and missing transverse momentum in pp collisions at $\sqrt{s} = 8$ TeV with the ATLAS detector*, *JHEP* **1405** (2014) 071, [arXiv:1403.5294 \[hep-ex\]](#).
- [4] A. Floderus, *Searches for direct supersymmetric gaugino and slepton pair production in final states with leptons with the ATLAS detector*, in *HCP 2012 - Hadron Collider Physics Symposium 2012*, vol. 49. EDP Sciences, 2013. <http://dx.doi.org/10.1051/epjconf/20134915008>.
- [5] D. J. Griffiths, *Introduction to elementary particles*. Wiley, New York, NY, 2008. [cds:111880](#).
- [6] J. F. Donoghue, E. Golowich, and B. R. Holstein, *Dynamics of the Standard Model*. Cambridge Univ. Press, Cambridge, 1992. [cds:238727](#).
- [7] G. L. Kane, *Modern elementary particle physics: the fundamental particles and forces*. Addison-Wesley, Reading, MA, 1993. [cds:364760](#).
- [8] J. Beringer et al. (Particle Data Group), *Review of Particle Physics*, *Phys. Rev. D* **86** (2012) 010001.
- [9] H. Goldstein, C. Poole, and J. Safko, *Classical Mechanics*. Addison-Wesley, San Francisco, CA, 2002. [cds:572817](#).
- [10] M. Gell-Mann, *A Schematic Model of Baryons and Mesons*, *Phys.Lett.* **8** (1964) 214–215.
- [11] O. W. Greenberg, *Spin and Unitary-Spin Independence in a Paraquark Model of Baryons and Mesons*, *Phys. Rev. Lett.* **13** (1964) 598–602.
- [12] C. C. Chevalley, *Theory of Lie groups*. Princeton Univ. Press, Princeton, NJ, 1946. [cds:112010](#).
- [13] C. N. Yang and R. L. Mills, *Conservation of Isotopic Spin and Isotopic Gauge Invariance*, *Phys. Rev.* **96** (1954) 191–195.
- [14] M. Kobayashi and T. Maskawa, *CP-Violation in the renormalizable theory of weak interaction*, *Prog. Theor. Phys.* **49** (1973) 652–657, [cds:1117037](#).
- [15] J. Goldstone, K. A. Johnson, A. Salam, and S. Weinberg, *Broken symmetries*, *Phys. Rev.* **127** no. 3, (1962) 965–970, [cds:405863](#).
- [16] ATLAS Collaboration, *Observation of a new particle in the search for the Standard Model Higgs boson with the ATLAS detector at the LHC*, *Phys. Lett.* **B716** (2012) 1–29, [arXiv:1207.7214 \[hep-ex\]](#).
- [17] CMS Collaboration, *Observation of a new boson at a mass of 125 GeV with the CMS experiment at the LHC*, *Phys. Lett.* **B716** (2012) 30–61, [arXiv:1207.7235 \[hep-ex\]](#).

- [18] V. Rubin, N. Thonnard, and W. K. Ford Jr, *Rotational Properties of 21 Sc Galaxies with a Large Range of Luminosities and Radii from NGC 4605 ($R=4kpc$) to UGC 2885 ($R=122kpc$)*, *Astrophysical Journal* **238** (1980) 471.
- [19] D. Close et al., *A direct empirical proof of the existence of dark matter*, *Astrophysical Journal* **648** (2) (2006) L109–L113, [arXiv:astro-ph/0608407](#) [[astro-ph](#)].
- [20] Planck Collaboration, *Planck 2013 results. I. Overview of products and scientific results*, [arXiv:1303.5062](#) [[astro-ph.CO](#)].
- [21] G. G. Ross, *Grand Unified Theories*. Frontiers in Physics. Westview Press, Reading, MA, 1985. [cds:1107830](#).
- [22] S. P. Martin, *A Supersymmetry Primer*, [arXiv:hep-ph/9709356](#) [[hep-ph](#)].
- [23] M. C. Gonzalez-Garcia and M. Maltoni, *Phenomenology with Massive Neutrinos*, *Phys. Rept.* **460** (2008) 1–129, [arXiv:0704.1800](#) [[hep-ph](#)].
- [24] H. Miyazawa, *Baryon number changing currents*, *Prog. Theor. Phys.* **36** no. 6, (1966) 1266–1276, [cds:478673](#).
- [25] P. Ramond, *Dual theory for free fermions*, *Phys. Rev. D* **3** no. 10, (1971) 2415–2418, [cds:432382](#).
- [26] Y. A. Gelfand and E. P. Likhtman, *Extension of the algebra of Poincare group generators and violation of P invariance*, *JETP Lett.* **13** no. 8, (1971) 323–325, [cds:433516](#).
- [27] N. A. and J. H. Schwarz, *Factorizable dual model of pions*, *Nucl. Phys.* **B31** (1971) 86–112.
- [28] N. A. and J. H. Schwarz, *Quark Model of Dual Pions*, *Phys. Rev.* **D4** (1971) 1109–1111.
- [29] J.-L. Gervais and B. Sakita, *Field Theory Interpretation of Supergauges in Dual Models*, *Nucl.Phys.* **B34** (1971) 632–639.
- [30] D. V. Volkov and V. P. Akulov, *Is the Neutrino a Goldstone Particle?*, *Phys. Lett.* **B46** (1973) 109–110.
- [31] J. Wess and B. Zumino, *A Lagrangian Model Invariant Under Supergauge Transformations*, *Phys. Lett.* **B49** (1974) 52.
- [32] J. Wess and B. Zumino, *Supergauge Transformations in Four-Dimensions*, *Nucl. Phys.* **B70** (1974) 39–50.
- [33] G. 't Hooft, *Symmetry breaking through Bell-Jackiw anomalies*, *Phys. Rev. Lett.* **37** no. 1, (1976) 8–11. 11 p, [cds:428662](#).
- [34] H. Goldberg, *Constraint on the photino mass from cosmology*, *Phys. Rev. Lett.* **50** no. NUB-2592-REV, (1983) 1419. 10 p, [cds:143045](#).
- [35] J. R. Ellis, J. S. Hagelin, D. V. Nanopoulos, K. A. Olive, and M. Srednicki, *Supersymmetric relics from the big bang*, *Nucl. Phys. B* **238** no. SLAC-PUB-3171, (1983) 453–476. 38 p, [cds:782076](#).
- [36] A. Djouadi et al., *The Minimal Supersymmetric Standard Model: Group Summary Report*, [arXiv:hep-ph/9901246](#) [[hep-ph](#)].

- [37] O. S. Bruning et al., *LHC Design Report*, vol. I. CERN, Geneva, 2004. [cds:782076](#).
- [38] M. Benedikt et al., *LHC Design Report*, vol. III. CERN, Geneva, 2004. [cds:823808](#).
- [39] ATLAS Collaboration, *ATLAS: technical proposal for a general-purpose pp experiment at the Large Hadron Collider at CERN*. LHC Tech. Proposal. CERN, Geneva, 1994. [cds:290968](#).
- [40] CMS Collaboration, *Technical proposal*. LHC Tech. Proposal. CERN, Geneva, 1994. [cds:290969](#).
- [41] LHCb Collaboration, *LHCb : Technical Proposal*. Tech. Proposal. CERN, Geneva, 1998. [cds:622031](#).
- [42] ALICE Collaboration, *ALICE: Technical proposal for a Large Ion collider Experiment at the CERN LHC*. LHC Tech. Proposal. CERN, Geneva, 1995. [cds:293391](#).
- [43] ATLAS Collaboration, *ATLAS detector and physics performance: Technical Design Report, Volume 1*. Technical Design Report ATLAS. CERN, Geneva, 1999. [cds:391176](#).
- [44] ATLAS Collaboration, *ATLAS detector and physics performance: Technical Design Report, Volume 2*. Technical Design Report ATLAS. CERN, Geneva, 1999. [cds:391177](#).
- [45] ATLAS Collaboration, *The ATLAS Experiment at the CERN Large Hadron Collider*, [JINST 3 \(2008\) S08003](#).
- [46] ATLAS Collaboration, *ATLAS magnet system: Technical Design Report*. Technical Design Report ATLAS. CERN, Geneva, 1997. [cds:338080](#).
- [47] ATLAS Collaboration, *ATLAS inner detector: Technical Design Report, Volume 1*. Technical Design Report ATLAS. CERN, Geneva, 1997. [cds:331063](#).
- [48] ATLAS Collaboration, *ATLAS inner detector: Technical Design Report, Volume 2*. Technical Design Report ATLAS. CERN, Geneva, 1997. [cds:331064](#).
- [49] ATLAS Collaboration, *ATLAS pixel detector: Technical Design Report*. Technical Design Report ATLAS. CERN, Geneva, 1998. [cds:381263](#).
- [50] ATLAS Collaboration, *ATLAS calorimeter performance: Technical Design Report*. CERN, Geneva, 1996. [cds:331059](#).
- [51] ATLAS Collaboration, *ATLAS liquid-argon calorimeter: Technical Design Report*. CERN, Geneva, 1996. [cds:331061](#).
- [52] ATLAS Collaboration, *ATLAS tile calorimeter: Technical Design Report*. Technical Design Report ATLAS. CERN, Geneva, 1996. [cds:331062](#).
- [53] ATLAS Collaboration, *ATLAS muon spectrometer: Technical Design Report*. CERN, Geneva, 1997. [cds:331068](#).
- [54] L. Tompkins, *Performance of the ATLAS Minimum Bias Trigger in pp collisions at the LHC*, Tech. Rep. ATL-DAQ-PROC-2010-033, CERN, Geneva, Sep, 2010. [cds:1295946](#).
- [55] M. L. Miller, K. Reygers, S. J. Sanders, and P. Steinberg, *Glauber modeling in high energy nuclear collisions*, [Ann.Rev.Nucl.Part.Sci. 57 \(2007\) 205–243](#), [arXiv:nucl-ex/0701025 \[nucl-ex\]](#).

- [56] J. D. Jackson, *Classical electrodynamics*. Wiley, New York, NY, 1999. [cds:490457](#).
- [57] ATLAS Collaboration, *ATLAS level-1 trigger: Technical Design Report*. Technical Design Report ATLAS. CERN, Geneva, 1998. [cds:381429](#).
- [58] P. Jenni and M. Nessi, *ATLAS Forward Detectors for Luminosity Measurement and Monitoring*, Tech. Rep. CERN-LHCC-2004-010. LHCC-I-014, CERN, Geneva, Mar, 2004. [cds:721908](#).
- [59] M. Villa, *The Luminosity Monitor of the ATLAS Experiment*, Tech. Rep. ATL-LUM-PROC-2009-004, CERN, Geneva, Nov, 2009. [cds:1222513](#).
- [60] <https://twiki.cern.ch/twiki/bin/view/AtlasPublic/LuminosityPublicResults>.
- [61] S. van der Meer, *Calibration of the effective beam height in the ISR*, Tech. Rep. CERN-ISR-PO-68-31. ISR-PO-68-31, CERN, Geneva, 1968. [cds:296752](#).
- [62] C. Rubbia, *Measurement of the luminosity of $p\bar{p}$ collider with a (generalized) Van der Meer Method*, Tech. Rep. CERN- $p\bar{p}$ -Note-38, CERN, Geneva, Nov, 1977. [cds:1025746](#).
- [63] W. Herr and B. Muratori, *Concept of luminosity*, [cds:941318](#).
- [64] I. Melzer-Pellmann and P. Pralavorio, *Lessons for SUSY from the LHC after the first run*, *Eur.Phys.J. C* **74** (2014) 2801, [arXiv:1404.7191 \[hep-ex\]](#).
- [65] M. Dine and W. Fischler, *A Phenomenological Model of Particle Physics Based on Supersymmetry*, *Phys. Lett. B* **110** (1982) 227.
- [66] L. Alvarez-Gaume, M. Claudson, and M. B. Wise, *Low-Energy Supersymmetry*, *Nucl. Phys. B* **207** (1982) 96.
- [67] C. R. Nappi and B. A. Ovrut, *Supersymmetric Extension of the $SU(3) \times SU(2) \times U(1)$ Model*, *Phys. Lett. B* **113** (1982) 175.
- [68] M. Dine and A. E. Nelson, *Dynamical supersymmetry breaking at low-energies*, *Phys. Rev. D* **48** (1993) 1277, [arXiv:hep-ph/9303230](#).
- [69] M. Dine, A. E. Nelson, and Y. Shirman, *Low-energy dynamical supersymmetry breaking simplified*, *Phys. Rev. D* **51** (1995) 1362, [arXiv:hep-ph/9408384](#).
- [70] M. Dine, A. E. Nelson, Y. Nir, and Y. Shirman, *New tools for low-energy dynamical supersymmetry breaking*, *Phys. Rev. D* **53** (1996) 2658, [arXiv:hep-ph/9507378](#).
- [71] G. Jungman, M. Kamionkowski, and K. Griest, *Supersymmetric dark matter*, *Phys.Rept.* **267** (1996) 195–373, [arXiv:hep-ph/9506380 \[hep-ph\]](#).
- [72] ATLAS Collaboration, *Electron performance measurements with the ATLAS detector using the 2010 LHC proton-proton collision data*, *Eur. Phys. J. C* **72** (2012) 1909, [arXiv:1110.3174 \[hep-ex\]](#).
- [73] ATLAS Collaboration, *Expected Performance of the ATLAS Experiment - Detector, Trigger and Physics*, [arXiv:0901.0512 \[hep-ex\]](#).
- [74] ATLAS Collaboration, *Search for direct production of charginos and neutralinos in events with three leptons and missing transverse momentum in $\sqrt{s} = 8\text{TeV}$ pp collisions with the ATLAS detector*, *JHEP* **1404** (2014) 169, [arXiv:1402.7029 \[hep-ex\]](#).

- [75] ATLAS Collaboration, *Search for electroweak production of supersymmetric particles in final states with at least two hadronically decaying taus and missing transverse momentum with the ATLAS detector in proton-proton collisions at $\sqrt{s} = 8$ TeV*, Tech. Rep. ATLAS-CONF-2013-028, CERN, Geneva, Mar, 2013. [cds:1525889](#).
- [76] ATLAS Collaboration, *Identification of the Hadronic Decays of Tau Leptons in 2012 Data with the ATLAS Detector*, Tech. Rep. ATLAS-CONF-2013-064, CERN, Geneva, Jul, 2013. [cds:1562839](#).
- [77] M. Cacciari, G. P. Salam, and G. Soyez, *The Anti- k_t jet clustering algorithm*, **JHEP** **0804** (2008) 063, [arXiv:0802.1189 \[hep-ph\]](#).
- [78] M. Cacciari and G. P. Salam, *Dispelling the N^3 myth for the k_t jet-finder*, **Phys.Lett.** **B641** (2006) 57–61, [arXiv:hep-ph/0512210 \[hep-ph\]](#).
- [79] ATLAS Collaboration, *Commissioning of the ATLAS high-performance b -tagging algorithms in the 7 TeV collision data*, Tech. Rep. ATLAS-CONF-2011-102, CERN, Geneva, Jul, 2011. [cds:1369219](#).
- [80] M. Biglietti et al., *Muon Event Filter Software for the ATLAS Experiment at LHC*, [cds:820783](#).
- [81] C. Lester and D. Summers, *Measuring masses of semi-invisibly decaying particles pair produced at hadron colliders*, **Phys.Lett.** **B463** (1999) 99–103, [arXiv:hep-ph/9906349 \[hep-ph\]](#).
- [82] A. Barr, C. Lester, and P. Stephens, *m_{T2} : The truth behind the glamour*, **J.Phys.** **G29** (2003) 2343–2363, [arXiv:hep-ph/0304226 \[hep-ph\]](#).
- [83] A. J. Barr, B. Gripaios, and C. G. Lester, *Weighing Wimps with Kinks at Colliders: Invisible Particle Mass Measurements from Endpoints*, **JHEP** **0802** (2008) 014, [arXiv:0711.4008 \[hep-ph\]](#).
- [84] ATLAS Collaboration, *Measurement of the W^+W^- production cross section in proton-proton collisions at $\sqrt{s} = 8$ TeV with the ATLAS detector*, Tech. Rep. ATLAS-CONF-2014-033, CERN, Geneva, Jul, 2014. [cds:1728248](#).
- [85] ATLAS Collaboration, *Jet energy resolution in proton-proton collisions at $\sqrt{s} = 7$ TeV recorded in 2010 with the ATLAS detector*, **Eur.Phys.J.** **C73** (2013) 2306, [arXiv:1210.6210 \[hep-ex\]](#).
- [86] ATLAS Collaboration, *Jet energy measurement with the ATLAS detector in proton-proton collisions at $\sqrt{s} = 7$ TeV*, **Eur.Phys.J.** **C73** (2013) 2304, [arXiv:1112.6426 \[hep-ex\]](#).
- [87] ATLAS Collaboration, *Preliminary results on the muon reconstruction efficiency, momentum resolution, and momentum scale in ATLAS 2012 pp collision data*, Tech. Rep. ATLAS-CONF-2013-088, CERN, Geneva, Aug, 2013. [cds:1580207](#).
- [88] ATLAS Collaboration, *Electron performance measurements with the ATLAS detector using the 2010 LHC proton-proton collision data*, **Eur.Phys.J.** **C72** (2012) 1909, [arXiv:1110.3174 \[hep-ex\]](#).
- [89] ATLAS Collaboration, *Performance of Missing Transverse Momentum Reconstruction in Proton-Proton Collisions at 7 TeV with ATLAS*, **Eur.Phys.J.** **C72** (2012) 1844, [arXiv:1108.5602 \[hep-ex\]](#).

- [90] ATLAS Collaboration, *Measuring the b -tag efficiency in a top-pair sample with 4.7 fb^{-1} of data from the ATLAS detector*, Tech. Rep. ATLAS-CONF-2012-097, CERN, Geneva, Jul, 2012. [cds:1460443](#).
- [91] J. C. Collins, D. E. Soper, and G. F. Sterman, *Factorization of Hard Processes in QCD*, Adv.Ser.Direct.High Energy Phys. **5** (1988) 1–91, [arXiv:hep-ph/0409313](#) [hep-ph].
- [92] H.-L. Lai, M. Guzzi, J. Huston, Z. Li, P. M. Nadolsky, et al., *New parton distributions for collider physics*, Phys.Rev. **D82** (2010) 074024, [arXiv:1007.2241](#) [hep-ph].
- [93] LHC Higgs Cross Section Working Group Collaboration, S. Heinemeyer et al., *Handbook of LHC Higgs Cross Sections: 3. Higgs Properties*, [arXiv:1307.1347](#) [hep-ph].
- [94] A. L. Read, *Presentation of search results: the CL_s technique*, J. Phys. G **28** no. 10, (2002) 2693–704, [cds:722145](#).
- [95] A. L. Read, *Modified frequentist analysis of search results (the CL_s method)*, [cds:451614](#).
- [96] LEP SUSY Working Group (ALEPH, DELPHI, L3, OPAL). Notes LEPSUSYWG/01-03.1, 04-01.1, <http://lepsusy.web.cern.ch/lepsusy/Welcome.html>.
- [97] M. R. Buckley, J. D. Lykken, C. Rogan, and M. Spiropulu, *Super-Razor and Searches for Staletons and Charginos at the LHC*, Phys.Rev. **D89** (2014) 055020, [arXiv:1310.4827](#) [hep-ph].
- [98] A. Bharucha, S. Heinemeyer, and F. von der Pahlen, *Does the LHC exclude SUSY Particles at the ILC?*, [arXiv:1404.0365](#) [hep-ph].
- [99] CMS Collaboration, *Measurement of WW production rate*, Tech. Rep. CMS-PAS-SMP-12-005, CERN, Geneva, 2012. [cds:1440234](#).
- [100] CMS Collaboration, *Measurement of $W+W^-$ and ZZ production cross sections in pp collisions at $\sqrt{s} = 8\text{ TeV}$* , Phys.Lett. **B721** (2013) 190–211, [arXiv:1301.4698](#) [hep-ex].
- [101] K. Rolbiecki and K. Sakurai, *Light stops emerging in WW cross section measurements?*, JHEP **1309** (2013) 004, [arXiv:1303.5696](#) [hep-ph].
- [102] D. Curtin, P. Meade, and P.-J. Tien, *Natural SUSY in Plain Sight*, [arXiv:1406.0848](#) [hep-ph].
- [103] U. Ellwanger, C. Hugonie, and A. M. Teixeira, *The Next-to-Minimal Supersymmetric Standard Model*, Physics Reports **496** no. 1–2, (2010) 1 – 77.
- [104] G. D. Kribs, E. Poppitz, and N. Weiner, *Flavor in supersymmetry with an extended R-symmetry*, Phys.Rev. **D78** (2008) 055010, [arXiv:0712.2039](#) [hep-ph].
- [105] S. Choi, M. Drees, A. Freitas, and P. Zerwas, *Testing the Majorana Nature of Gluinos and Neutralinos*, Phys.Rev. **D78** (2008) 095007, [arXiv:0808.2410](#) [hep-ph].
- [106] G. D. Kribs and A. Martin, *Supersoft Supersymmetry is Super-Safe*, Phys.Rev. **D85** (2012) 115014, [arXiv:1203.4821](#) [hep-ph].
- [107] M. L. Mangano, M. Moretti, F. Piccinini, R. Pittau, and A. D. Polosa, *ALPGEN, a generator for hard multiparton processes in hadronic collisions*, JHEP **0307** (2003) 001, [arXiv:hep-ph/0206293](#) [hep-ph].

- [108] F. Campanario, V. Hankele, C. Oleari, S. Prestel, and D. Zeppenfeld, *QCD corrections to charged triple vector boson production with leptonic decay*, *Phys.Rev.* **D78** (2008) 094012, [arXiv:0809.0790 \[hep-ph\]](#).
- [109] T. Sjostrand, S. Mrenna, and P. Z. Skands, *PYTHIA 6.4 Physics and Manual*, *JHEP* **0605** (2006) 026, [arXiv:hep-ph/0603175 \[hep-ph\]](#).
- [110] S. Frixione, P. Nason, and C. Oleari, *Matching NLO QCD computations with Parton Shower simulations: the POWHEG method*, *JHEP* **0711** (2007) 070, [arXiv:0709.2092 \[hep-ph\]](#).
- [111] S. Frixione and B. R. Webber, *Matching NLO QCD computations and parton shower simulations*, *JHEP* **0206** (2002) 029, [arXiv:hep-ph/0204244 \[hep-ph\]](#).
- [112] S. Frixione, P. Nason, and B. R. Webber, *Matching NLO QCD and parton showers in heavy flavor production*, *JHEP* **0308** (2003) 007, [arXiv:hep-ph/0305252 \[hep-ph\]](#).
- [113] S. Frixione, E. Laenen, P. Motylinski, and B. R. Webber, *Single-top production in MC@NLO*, *JHEP* **0603** (2006) 092, [arXiv:hep-ph/0512250 \[hep-ph\]](#).
- [114] G. Corcella, I. Knowles, G. Marchesini, S. Moretti, K. Odagiri, et al., *HERWIG 6: An Event generator for hadron emission reactions with interfering gluons (including supersymmetric processes)*, *JHEP* **0101** (2001) 010, [arXiv:hep-ph/0011363 \[hep-ph\]](#).
- [115] J. Pumplin, D. Stump, J. Huston, H. Lai, P. M. Nadolsky, et al., *New generation of parton distributions with uncertainties from global QCD analysis*, *JHEP* **0207** (2002) 012, [arXiv:hep-ph/0201195 \[hep-ph\]](#).
- [116] GEANT4 Collaboration, S. Agostinelli et al., *GEANT4: A Simulation toolkit*, *Nucl.Instrum.Meth.* **A506** (2003) 250–303.
- [117] ATLAS Collaboration, *The ATLAS Simulation Infrastructure*, *Eur. Phys. J.* **C70** (2010) 823, [arXiv:1005.4568 \[physics.ins-det\]](#).
- [118] E. Richter-Was, D. Froidevaux, and L. Poggioli, *ATLFAST 2.0 a fast simulation package for ATLAS*, Tech. Rep. ATL-PHYS-98-131, CERN, Geneva, Nov, 1998. [cds:683751](#).
- [119] ATLAS Collaboration, *The simulation principle and performance of the ATLAS fast calorimeter simulation FastCaloSim*, Tech. Rep. ATL-PHYS-PUB-2010-013, CERN, Geneva, Oct, 2010. [cds:1300517](#).
- [120] J. Butterworth, J. R. Forshaw, and M. Seymour, *Multiparton interactions in photoproduction at HERA*, *Z.Phys.* **C72** (1996) 637–646, [arXiv:hep-ph/9601371 \[hep-ph\]](#).
- [121] ATLAS Collaboration, *ATLAS tunes of PYTHIA 6 and Pythia 8 for MC11*, Tech. Rep. ATL-PHYS-PUB-2011-009, CERN, Geneva, Jul, 2011. [cds:1363300](#).
- [122] M. Bahr, S. Gieseke, M. Gigg, D. Grellscheid, K. Hamilton, et al., *Herwig++ Physics and Manual*, *Eur.Phys.J.* **C58** (2008) 639–707, [arXiv:0803.0883 \[hep-ph\]](#).
- [123] W. Beenakker, M. Klasen, M. Kramer, T. Plehn, M. Spira, et al., *The Production of charginos / neutralinos and sleptons at hadron colliders*, *Phys.Rev.Lett.* **83** (1999) 3780–3783, [arXiv:hep-ph/9906298 \[hep-ph\]](#).
- [124] T. Binoth, M. Ciccolini, N. Kauer, and M. Kramer, *Gluon-induced W-boson pair production at the LHC*, *JHEP* **0612** (2006) 046, [arXiv:hep-ph/0611170 \[hep-ph\]](#).

- [125] T. Gleisberg, S. Hoeche, F. Krauss, M. Schonherr, S. Schumann, et al., *Event generation with SHERPA 1.1*, **JHEP** **0902** (2009) 007, [arXiv:0811.4622](#) [[hep-ph](#)].
- [126] T. Binoth, N. Kauer, and P. Mertsch, *Gluon-induced QCD corrections to $pp \rightarrow ZZ \rightarrow \ell\bar{\ell}\ell'\bar{\ell}'$* , [arXiv:0807.0024](#) [[hep-ph](#)].
- [127] B. P. Kersevan and E. Richter-Was, *The Monte Carlo event generator AcerMC versions 2.0 to 3.8 with interfaces to PYTHIA 6.4, HERWIG 6.5 and ARIADNE 4.1*, **Comput.Phys.Commun.** **184** (2013) 919–985, [arXiv:hep-ph/0405247](#) [[hep-ph](#)].
- [128] J. Neyman and E. S. Pearson, *On the Problem of the Most Efficient Tests of Statistical Hypotheses*, **Philosophical Transactions of the Royal Society of London. Series A, Containing Papers of a Mathematical or Physical Character** **231** (1933) pp. 289–337.
- [129] S. S. Wilks, *The Large-Sample Distribution of the Likelihood Ratio for Testing Composite Hypotheses*, **The Annals of Mathematical Statistics** **9** no. 1, (1938) 60–62.
- [130] A. Wald, *Tests of Statistical Hypotheses Concerning Several Parameters When the Number of Observations is Large*, **Transactions of the American Mathematical Society** **54** no. 3, (1943) pp. 426–482.
- [131] G. Cowan, K. Cranmer, E. Gross, and O. Vitells, *Asymptotic formulae for likelihood-based tests of new physics*, **European Physical Journal C** **71** (2011) 1554, [arXiv:1007.1727](#) [[physics.data-an](#)].