



Performance Studies of the ATLAS Transition Radiation Tracker Barrel using SR1 Cosmics Data

Richard Wall

March 8, 2007

Abstract

The ATLAS experiment at the Large Hadron Collider (LHC) is designed to measure Nature at the energy scale often associated with electroweak symmetry breaking. When it comes online in 2008, the LHC and ATLAS will work to discover, among other things, the Higgs boson and any other signatures for physics beyond the Standard Model. As part of the ATLAS Inner Detector, the Transition Radiation Tracker will be an important part of ATLAS's ability to make precise measurements of particle properties. This paper summarizes work done to study and categorize the performance of the TRT, using a combination of cosmic ray test data from the SR1 facility and Monte Carlo. In general, it was found that the TRT is working well, with module-level efficiencies around 90 % and module-level noise just above 2 %. Reasonably good agreement was observed with Monte Carlo, though there are some apparently pathological differences between the two that deserve further attention.

Contents

1	Acknowledgements	4
2	Introduction	5
2.1	CERN and the Large Hadron Collider	5
2.2	ATLAS, the Inner Detector, and the Transition Radiation Tracker	7
3	Methodology	10
3.1	Efficiency Defined	10
3.2	Noise Defined	11
3.3	Minor Complications	12
4	Efficiency Studies	13
4.1	Module Level Efficiency	14
4.1.1	Stability of Efficiency Measurements	14
4.1.2	Efficiency as a Function of Distance	15
4.1.3	Efficiency as a Function of Track Length Inside the Straw	16
4.1.4	Efficiency as a Function of Track Quality	20
4.1.5	Efficiency as a Function of Alignment	23
4.2	Straw Level Efficiency	24
4.2.1	Efficiency Per Straw within the Module	24
4.3	Efficiency Conclusion	28
5	Noise Studies	29
5.1	Module Level Noise	29
5.1.1	Noise Levels per Module	29
5.1.2	Noise as a Function of Track Quality	31
5.2	Straw Level Noise	31
5.2.1	Noise as a Function of Time over Threshold	32

5.2.2	Noise per Straw within a Module	33
5.2.3	Preliminary Measurements of Cross-Talk	34
5.3	Noise Conclusions	37

List of Figures

1	The ATLAS Detector	8
2	Cross-Sectional View of the TRT	9
3	Module 1.6.0 Efficiency as a Function of Run Number	15
4	Basic Efficiency vs Distance Profile	16
5	Efficiency vs Distance in Data and Monte Carlo by Layer	17
6	Efficiency as a Function of Track Length by Layer	19
7	Chisq/ndof Spectra	21
8	Efficiency vs Distance with Refined Track Samples by Layer	22
9	Average Efficiency per Layer vs. Track Chisq/Ndof	23
10	Efficiency vs Distance with Aligned and Unaligned Geometries	23
11	Efficiency per Straw per Module in Data	26
12	Efficiency per Straw per Module in Data and Monte Carlo	26
13	Comparison of Straw Occupancy in Layer 2	27
14	Distribution of Straw Efficiency in Layer 2	27
15	Average Noise Levels by Layer	30
16	Average Noise Levels by Phi-Sector	30
17	Noise Levels vs. Track Quality	31
18	Time over Threshold Distributions	33
19	Straw-Level Noise Occupancy	33
20	Noise Level per Straw per Module in Data	35
21	Noise Level per Straw per Module in Monte Carlo	35
22	Noise Level per Straw per Module in Data, ToT ≥ 1 bin	36
23	Noise Level per Straw per Module in Monte Carlo, ToT ≥ 1 bin	36

1 Acknowledgements

Over the past year, I have had the honor to work as part of the ATLAS collaboration as it prepares to push back the energy frontier and reveal an as yet unknown side of Nature. Standing, as I am, at the end of my time as a part of the Duke University Physics Department, I should take this time to thank those who have helped me get to where I am today.

First and foremost among the people I should thank is my thesis advisor, Mark Kruse. If he hadn't taken a chance and given me the opportunity to work in his group almost a year ago, this thesis certainly would not be here, and I may well not have pursued a career in physics at all. Throughout my time in the Duke ATLAS group, he has given me enough freedom to work and explore on my own while not leaving me lost in the woods and has always been there to provide a sanity check in case I've lost sight the bigger picture. I also have to thank Andrea Bocci who helped me get started at CERN and taught me essentially everything I know about ATHENA. Whether the mistake was a missed comma or a pathologically broken piece of code, Andrea has always been more than willing to sit down with me and figure it out, and I only wish I could have learnt more from him. In the same vein, I owe a lot of this thesis and indeed my whole understanding of the TRT to Wouter Hulsbergen. His expansive expertise with straw-tube detectors made Wouter an irreplaceable asset, and his welcoming, easy-going demeanor made him truly a pleasure to work with. Aside from all the advice he gave regarding the direction of this project, Wouter also was the original author of much of the ATHENA code used in this paper. Fianlly, I would like to thank the many members of the Duke faculty (many from outside the world of high-energy physics) who have fostered my interest and ability in physics. The warm, collegial environment they created made my undergraduate years truly something I will cherish for the rest of my life.

On a more personal note, I would like to thank my friends and family for putting up with me throughout this process. Its tempting, sometimes, to lose yourself in your work, even more so when the work is as amazing and wonderful as high-energy physics is for me, and so it is especially comforting to have a group of people around you who can bring you out of the grind, if only for a drink here or a movie there. I owe you all so much, I cannot possibly express it in words.

2 Introduction

When it comes online, the ATLAS experiment will need a variety of pieces to work in exact concert if reasonable measurements are to be obtained and new physics discovered. One key area where such measurements will be important is, of course, in particle identification. As part of the ATLAS Inner Detector (ID), the Transition Radiation Tracker (TRT) will provide precision measurements for particle trajectories as they leave the interaction point. These measurements, when combined with magnetic field information, will shed light on the particle's momentum and thus aid in general particle identification. Further, the TRT will aid uniquely in the identification of electrons when they interact with the radiator material located in the space between straws.

All of the functions of the TRT mentioned above, however, depend intimately on how well the straws that comprise the TRT perform in real conditions. In this paper, we study the performance of the TRT as it was during the SR1 cosmic ray test runs, recently conducted on site at CERN. This study has two principal goals; to study the efficiency of and noise levels in the TRT as seen in SR1 cosmics data and compare these results to available Monte Carlo simulation. While the former goal is clearly of interest, there are a number of reasons that making absolute statements about the TRT's performance is difficult in the present context. As such, we shall focus on comparing the available data to Monte Carlo, with the hope that better insight into the latter's ability to reproduce data can be gained.

2.1 CERN and the Large Hadron Collider

Founded 1954, the European Council for Nuclear Research (CERN) has had a rich and interesting history as one of the premier research institutions for experimental particle physics. Located on the Swiss-French border near Geneva, CERN has a uniquely international flavor not found in many other major research laboratories in the world, and throughout its history it has clearly benefited from this international nature. Funded mostly by the 20 European member states, CERN is also one of the most important joint ventures currently being undertaken by the newly unified Europe.

As to its contribution to science in general, CERN has been very productive over its 50 year history. Some of its major discoveries include the first evidence for the W and Z bosons in 1983 and the first experimental verification of CP-violation in 2001. Today, CERN is ramping up for its next major project, the Large Hadron Collider (LHC), set to come online in early 2008. Designed to collide protons at a center of mass collision energy of 14 TeV, the LHC will be the most powerful particle accelerator in the world, with the ability to probe Nature on energy scales not seen since the very early universe. Specifically, the LHC is designed to study the energy regime in which the electroweak symmetry is thought to break down.

Within the Standard Model of particle physics, the spontaneous breakdown of the electroweak symmetry occurs via a process called the Higg's mechanism. It turns out that this same process leads to the generation of mass for all fundamental particles (quarks, leptons, and the W and Z bosons). Specifically, the coupling of the Higgs field (whose strength is everywhere a non-zero constant) to a fundamental particle generates a "drag" on the particle that we perceive as its rest mass. Though some indirect suggestions for the Higg's mechanism have already been found, direct experimental evidence for it has yet to be seen. The most likely candidate for such verification will be an observation of the Higg's boson, the massive spin 0 particle associated with the Higg's field itself. Precision measures of various electroweak phenomena at the Tevatron at Fermilab have constrained the Higg's mass to a very well-defined range that will be easily within the reach of the LHC when it reaches its design specifications. Interestingly, however, some of these same measurements indicate the Higg's may be accessible by the Tevatron itself (which operates a center-of-mass collision energy of only 2 TeV), in which case the Higg's may be discovered before the LHC ever turns on.

Like past experiments, the LHC will operate multiple detectors, each with their own areas of specialization. One (ALICE) will study heavy-ion collisions, but the remaining four (ATLAS, CMS, LHCb, and TOTEM) will study the proton-proton collisions mentioned above. Of these, LHCb has been specialized to study b-quark events [6] and will thus be less sensitive to new physics. Similarly, the TOTEM detector has been specially designed to measure the total cross-

section, elastic scattering, and diffractive properties of the LHC [7] and is thus not expected to contribute to searches for new physics. The two remaining detectors, ATLAS and CMS, are designed to measure a larger range of phenomena beyond the Standard Model, from the Higgs mechanism to the existence of extra dimensions and even low-energy signatures of supersymmetry.

2.2 ATLAS, the Inner Detector, and the Transition Radiation Tracker

Within the larger framework of CERN and the LHC, the ATLAS collaboration is very important, since many of the LHC's most visible physics goals are tied to discoveries at either ATLAS or CMS. As with most large particle detectors, ATLAS consists of three main components, the Inner Detector (ID), the Calorimeters, and the Muon Chambers. As its name suggests, the Calorimeters are designed to measure the energy content of out-going particles. Since muons rarely deposit a significant amount of their energy in standard calorimeters, they tend to pass right through the entire detector. With this in mind, muons are primarily measured when they leave tracks in the specially designed Muon Chambers, which reside outside the calorimeters. Figure 1 shows the ATLAS detector in its entirety [5].

The ID, by contrast, is designed to measure the momentum, rather than the energy, of out-going particles. It accomplishes this goal largely through an accurate reconstruction of the particle's trajectory. When this information is combined with local magnetic field readings, an accurate momentum measurement can be made. The ATLAS ID has three main components, the Silicon Pixel Tracking System, the Semiconductor Tracker (SCT), and the Transition Radiation Tracker (TRT). The first two components use high-precision silicon strips to detect out going particles and are responsible for the majority of the ID's precision measurement capability. By contrast, the TRT is built with wire straw chambers in which a central wire (held at a positive potential) is suspended inside straw containing a gaseous mixture of xenon (70%), carbon dioxide (27%), and oxygen (3%) [1]. Particles leaving the interaction point will ionize molecules in the gas, and the electrons created in this process will drift towards the wire with a characteristic drift velocity. Once they impact the wire, higher end electronics will signal a hit, which can be correlated to the known

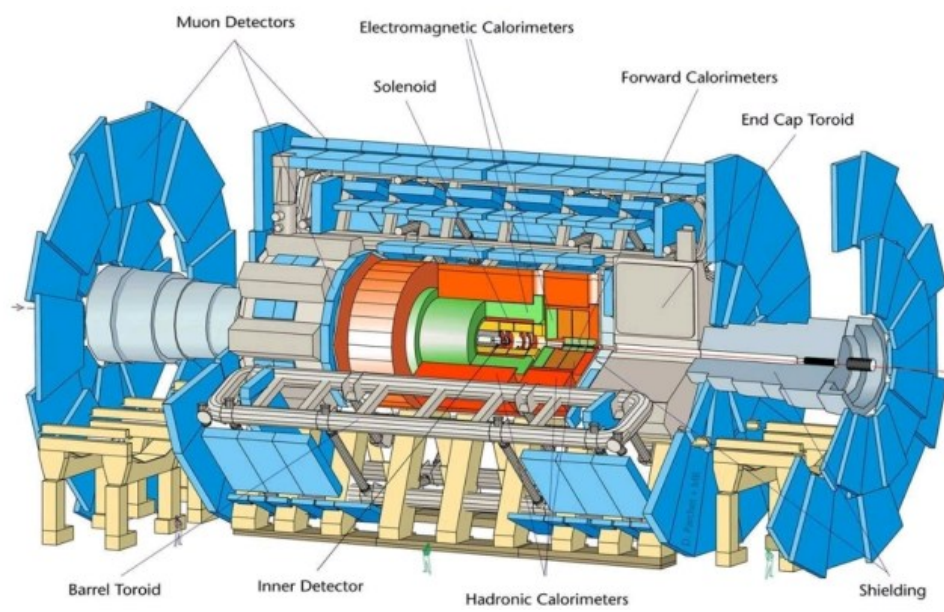


Figure 1: The ATLAS Detector

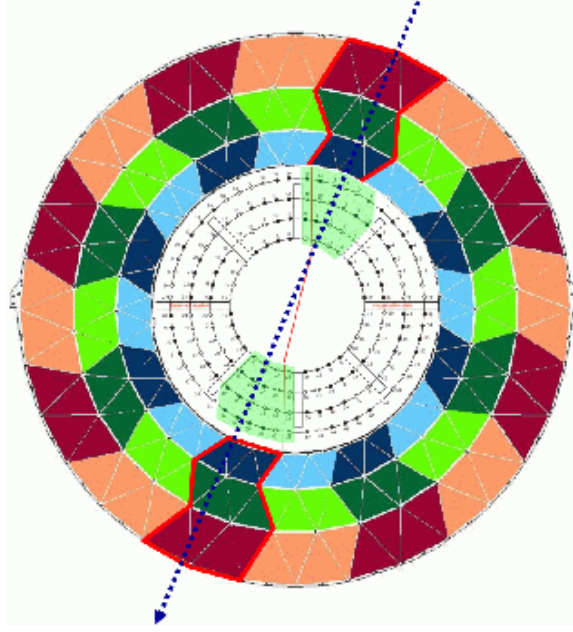


Figure 2: Cross-Sectional View of the TRT

position of the wire to deduce the particle's location at that moment. By using many such wires, an accurate picture of the particle's trajectory can be obtained. See Figure 2 for a cartoon make-up of a cross-sectional view of the TRT with a track passing through it [4].

Organizationally, the TRT itself has three main components, the Barrel (which is the focus of this paper) and two endcaps, which are designed to measure particles with a higher pseudo-rapidity than is accessible to the barrel. The barrel of the TRT itself is further subdivided into modules that are identified by layer (a local radial coordinate), ϕ -sector (an angular coordinate), and z-coordinate (either +1 or -1 relative to the interaction point) which lies along the beam line. The number of straws per module depends (as one might expect) on layer number, with 793 straws over 30 strawlayers in layer 2, 520 straws over 24 strawlayers in layer 1, and 329 straws over 19 strawlayers in layer 0 [1]. Within a given module, strawlayer and strawnumber are used to uniquely identify straws in the r - ϕ plane.

3 Methodology

Since these experiments are being conducted with real data on the nearly complete TRT, it is important to properly choose how one defines quantities like efficiency and noise. In other settings, it would be possible to conduct a more controlled series of experiments to determine the absolute efficiency and noise levels within the detector, and indeed such detailed analyses have been carried out elsewhere. The goal here is primarily to conduct a series of explorations of the TRT's performance in data and compare these findings with Monte Carlo, with the belief that any systematic bias in methodology will affect both samples equally. As such, the following definitions for efficiency and noise are used.

3.1 Efficiency Defined

As one might expect, efficiency in the TRT is designed to be a measure of how often the detector responds when a particle passes by it. In terms of the straw drift chambers themselves, how often does a wire register a hit when a particle passes within the straw radius? Given this intuitive picture, the efficiency (ϵ) is defined as the ratio of observed hits to expected hits:

$$\epsilon \equiv \frac{hits_{observed}}{hits_{expected}}. \quad (1)$$

Clearly, the number of observed hits simply a count of the number of times a wire sends a signal to the electronics. Expected hits, by contrast, measure the number of reconstructed tracks that pass within some distance R of the wire. Ideally, expected hits would count the number of times an actual particle passes within some distance R , but such information is not available.

One very important but subtle consequence of this definition is that there is really no way to define the absolute efficiency of a detector element, since all such measurements are made with respect to this distance parameter R . That being said, the straw is expected to be relatively uniform in its ability to register hits, and thus a reasonable estimate of detector efficiency can be obtained by integrating the efficiency over this distance parameter for a suitable range of R values. Exactly what

constitutes "a suitable range of R values" is up for discussion, but the general idea is that efficiency begins to drop rather sharply near the straw radius. This is due to the fact that the probability of a charged particle ionizing a molecule in the gas mixture is proportional to the particle's track length within the gas [2]. Thus, particles with a small track length inside a straw (which occur when R is nearly the straw radius) will not generate hits as well as particles that pass through the center of the straw. It is, however, misleading to include these edge effects, since they do not really reflect any inherent problem with the straw. As such, the upper bound of R is chosen such that these edge effects are negligible.

In general, the efficiency of the TRT is studied on two levels. The first, and most obvious, level is the straw itself. In this context, the efficiency of individual straws are studied, specifically with the goal of identifying straws with particularly low efficiency. These may represent broken or malfunctioning straws that should not be considered in further measurements. The next and larger level at which the efficiency can be studied is the module level. In these instances, the efficiency of individual straws is integrated over entire modules (or layers or ϕ -sectors), allowing insight into how the efficiency varies as a function of other more globally-relevant parameters, like track quality. Such explorations would be difficult at the straw level due to lowered statistics. In both cases, there is ample opportunity for comparison to Monte Carlo, though these comparisons are easier to categorize on the module level due to access to "reasonable" measures of detector performance like an average efficiency.

3.2 Noise Defined

In many ways, a study of noise is as simple as looking for efficiency in the wrong places. For the purpose of these studies, noise is taken to be the efficiency of a detector element for an R value outside the straw radius. Since poor track reconstruction or misalignments within the detector can blur distances at or very near the straw radius itself, it is more instructive to look well outside the straw radius for noise. Such hits, it can reasonably be assumed, could not be due to a real particle that is victim to a poor reconstruction. More will be said about an appropriate choice for what is

considered “well outside the straw radius” in later sections, however. As with the efficiency inside of the straw radius, this question is explored both at the straw level, to identify noisy straws, and at the module level, to gain a picture of ambient noise and see if these levels are affected by any parameters like location within the detector or track quality. Further, noise can be studied in terms of the time an individual hit over which a hit persisted in the detector, since this can be a valuable indicator of random, electronic noise.

3.3 Minor Complications

Let us now discuss a few potential complications that have arisen with the cosmics samples. All of these factors have been taken into account and will be discussed in more detail in later sections, but we still outline at a broad stroke these difficulties here. Data from cosmic ray test runs at the SR1 facility have been compiled over the months leading up to the installment of the TRT barrel along the beamline. In these studies, high energy particles (mostly muons) enter the detector and leave tracks in much the same way as particles leaving the interaction site would under real operating conditions. Since these studies were done prior to the TRT’s installation, there is no magnetic field present and hence all tracks follow a straight line trajectory.

This introduces the complication that momentum information about the incoming muons is not available. Since cosmic muons have a wide momentum spectrum, tracks in the TRT will experience various degrees of Multiple Coulomb Scattering (MCS), a phenomenon that drastically affects the quality of the reconstructed track. Simply put, if a track has numerous “kinks” in it due to MCS, reconstruction software will “smooth over” these kinks and make a line of best fit. In light of the definitions of efficiency and noise, it is clear that such “bad tracks” could certainly produce both low efficiencies and high noise levels, even if the detector is working perfectly. In an ideal world, tracks with low momentum (frequent MCS) would be removed from the data samples, but absent a measurement of track momentum, such cuts cannot be made. Thankfully, the contribution of such low momentum tracks is small, and there are indeed ways (to be discussed in later sections) of getting the benefits of a cut on track momentum even in the absence of such information.

Another important issue that deserves mention is the fact that the TRT was not completely finished when these studies were undertaken. Specifically, only the upper sector modules ($\phi = 6, 7$) were cabled in both + and - z directions. As a note, the Monte Carlo used throughout this paper reflects this rather atypical situation. In principle, this should not effect the efficiency of the detector at all, but as has been hinted at, the methods used to measure the TRT's efficiency and noise levels rely heavily on the quality of track reconstruction. Since only half of the detector is cabled fully in z, it is expected that there will be two kinds of reconstructed tracks, ones with hits in both the upper and lower sectors of the detector and those with hits only in the upper sector. This latter group of tracks are expected to be of lesser quality because they, in effect, "miss out on" the lower sector modules ($\phi = 22, 23$) and thus are less well-determined than tracks in the other group, which have approximately twice the number of hits.

4 Efficiency Studies

The first goal of these studies is to understand the efficiency of the TRT inside the straw radius. Obviously, this will give an idea of how faithful the TRT is at recording out-going particle tracks, but it is also important to compare the observed efficiencies in the data with those obtained from Monte Carlo. This comparison is not exact however, since the Monte Carlo's input momentum spectrum for cosmic muons is believed not to be precisely correct. While it is hard to categorize such a difference, the effects of it (as will be seen shortly) are generally small, since the agreement between data and simulation is usually quite good. That being said, there are some obvious but surprising differences that deserve further attention. The following discussions are divided first by the scope (module or straw level) of the study and then by the quantity under direct consideration (e.g. track quality, R value, etc.).

4.1 Module Level Efficiency

To begin studying the performance of the TRT, large scale efficiencies, that is to say the efficiency summed over all straws in a given part of the detector, is studied. The goal here is not to diagnose specific hardware problems with individual straws but rather to get a general sense for the TRT's performance. Direct comparison with Monte Carlo is also well-defined at this level, since it allows for access to tangible quantities like the average efficiency of a given module. In the following sections, we explore the stability (in time) of the efficiency and explore efficiency as a function (roughly) of distance from the wire, track quality, and detector alignment.

4.1.1 Stability of Efficiency Measurements

In order to remove statistical uncertainties from cosmic ray data samples, it would be nice to chain numerous runs together, since no single run has more than about ten thousand tracks. To do this, however, it must be seen whether the efficiency has a significant dependence on time (alternately the run number). If the efficiency appears stable, it is safe to chain multiple runs together and decrease the statistical uncertainties associated with all following measurements. To demonstrate the stability of the efficiency measurements, the average efficiency of a single module (1.6.0) as a function of run number is shown in Figure 3. The error associated with each point simply reflects the number of tracks in each run, so runs with more tracks have smaller errors.

The data in Figure 3 represents essentially all of the good data sets that came out of the SR1 facility during its cosmic ray test runs in June 2006 [3]. The precise times of each run have not been included, but the essential message is that the efficiency measurements appear stable over the timescale of the entire run of data taking. Though no plots are included for the other modules, similar results are obtained elsewhere in the TRT, and thus it is safe to conclude the efficiency of the TRT is indeed stable (at least over the few days during which data was being taken). Thus, all results presented from here on out will include a data from all good cosmic ray runs. Explicitly, this amounts to a sample of over thirty thousand tracks, taken from three data separate runs.

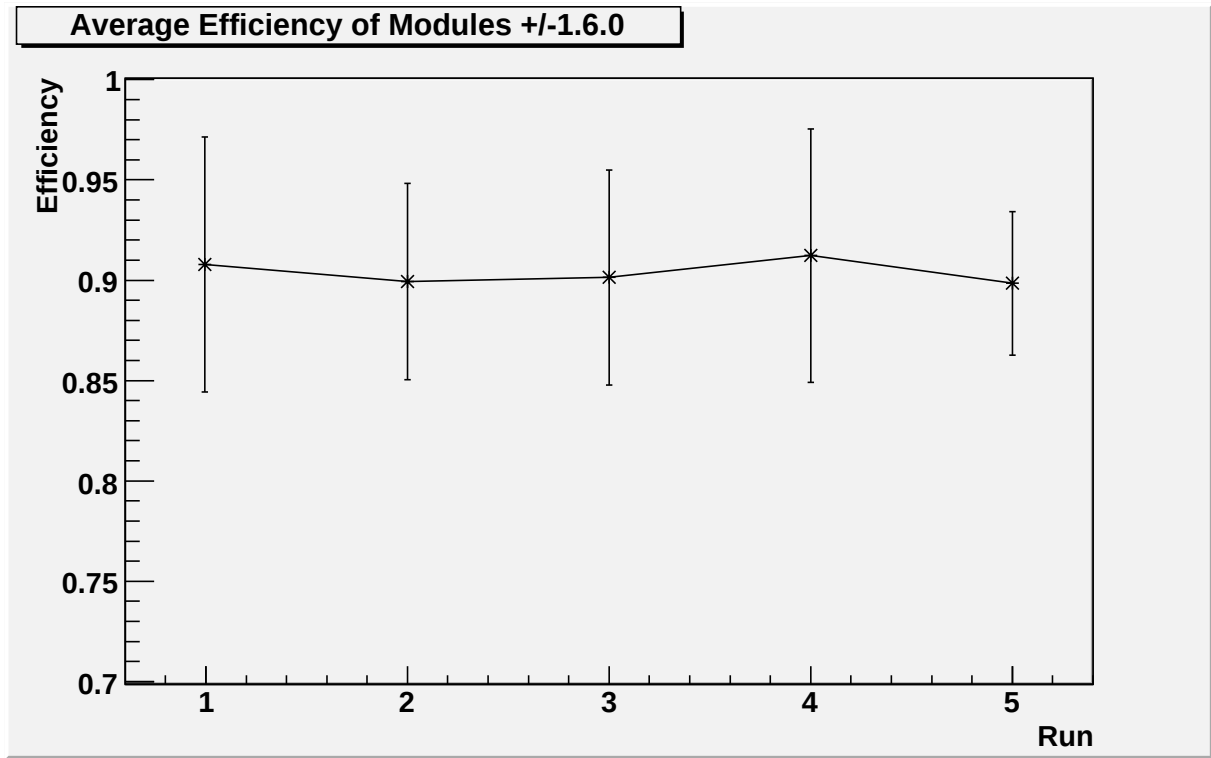


Figure 3: Module 1.6.0 Efficiency as a Function of Run Number

4.1.2 Efficiency as a Function of Distance

The first and most basic result from an efficiency study is a profile of efficiency as a function of distance. In general terms, the profile is hoped to be a square-wave centered at 0 with height 1 and a width of two times the straw radius. This would represent a detector with perfect efficiency and zero noise. In practice, however, this is not observed. Since the probability of a particle ionizing a molecule in the gas mixture is proportional to this track length, it is expected that the efficiency will drop near the straw radius and the profile will generally become rounded. The details of this process are left to the next section.

In 4, a standard efficiency profile as a function of distance to the wire is displayed. This plot is generated using hits from the entire TRT, with data and Monte Carlo distributions shown next to each other. Clearly, there is very good agreement, but this is somewhat misleading, since there are a number of factors that make this agreement "too good to be true." Shown in Figure 5 are a series of similar plots, but with each layer's efficiency separated out.

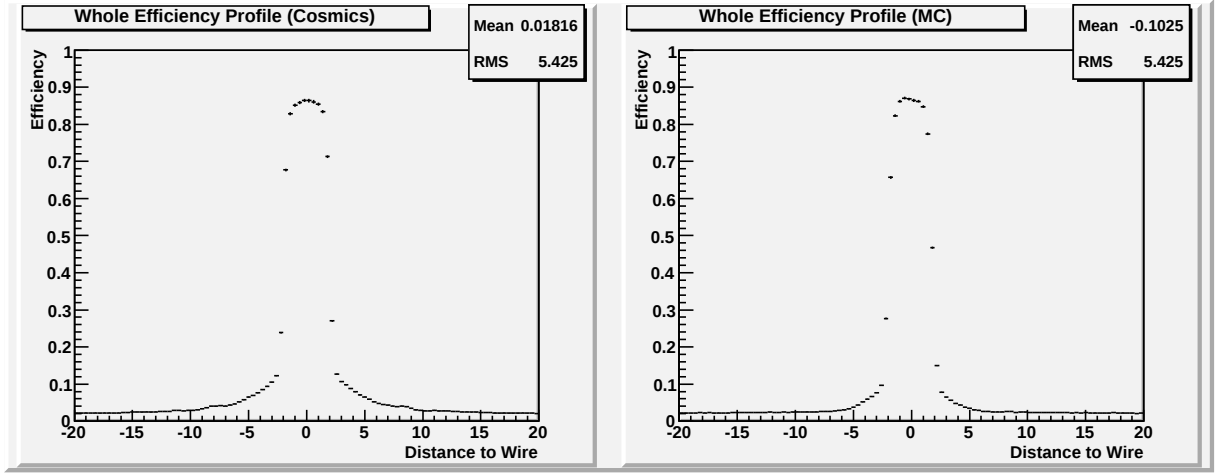


Figure 4: Basic Efficiency vs Distance Profile

From this, two important results become apparent. While it may be difficult to see from Figure 5 directly, the Monte Carlo efficiencies are consistently lower than the ones from data by roughly 2-3 %. In future sections, this will be commented on in more detail. The other surprising result that comes from this is that there is a clear layer dependence to the efficiency, with layer 0 being the most efficient and layer 2 being the least. This phenomenon is best explained by the fact that tracks are constructed such that they are better defined closer to the SCT and Pixels (since these components have a higher inherent precision than the TRT). As such, track segments in the outer layers are less well-determined and thus it is more likely a reconstructed track will "miss" a straw. It should be emphasized, however, that this does not mean the straws in layer 2 are somehow different than the others. Merely, this reflects one of the limitations of our definition of efficiency. Later we discuss a method to remove this effect.

4.1.3 Efficiency as a Function of Track Length Inside the Straw

As mentioned above, it is very important to understand the particulars of the ionization process within the straw to get an understanding of how well the detector is working. Along these lines, we study how the efficiency depends on track length within the straw, a quantity intimately related to the distance variable discussed above. Naively, we can relate the track length in a straw to the

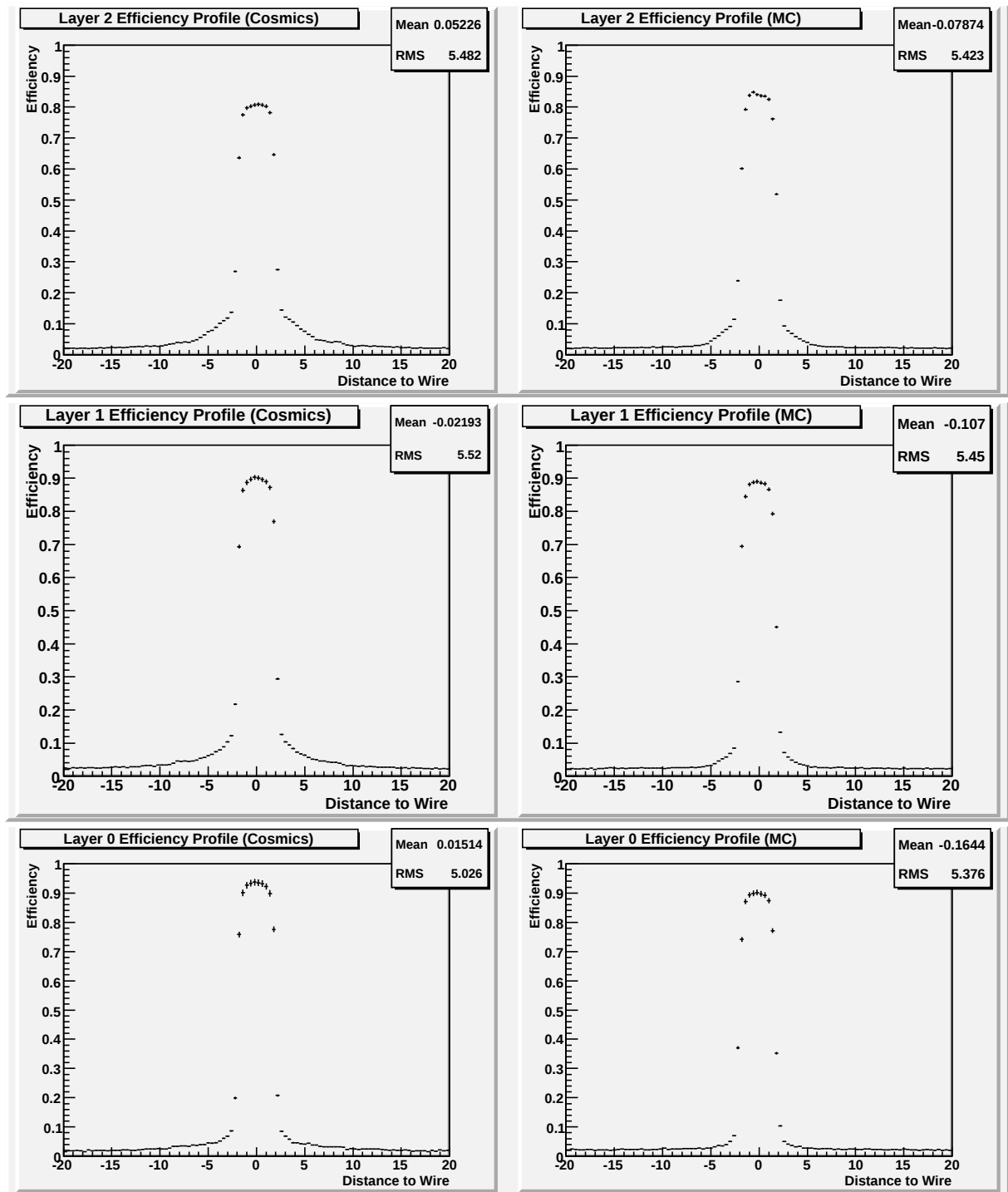


Figure 5: Efficiency vs Distance in Data and Monte Carlo by Layer

track's R value with the following formula;

$$tracklength = 2(2 - R) \quad (2)$$

At this point, however, we have added nothing to our understanding of the efficiency since we have simply changed variables. It is quite easy, however, to derive from theory what a distribution of efficiency as a function of track length should look like and relate this distribution to some parameters of the detector [2]. Specifically, it can be shown that the general shape of this distribution has the form

$$\epsilon = 1 - e^{-\alpha L} \quad (3)$$

where L is taken to be the track length inside the straw, and α is a constant related to the ionization cluster density of the gas mixture. While this picture is not precisely correct, it certainly captures many important features of the observed distributions. Taking this into account, we have plotted below in Figure 6 the efficiency of modules in ϕ -sector 6 as a function of track length for data and Monte Carlo and fit each to a theoretical curve based on the above considerations. From each fit, a value for α is obtained (indicated in Figure 6 as $p0$), allowing for a clean comparison between data and Monte Carlo.

Comparing the obtained fit values for α , we see that the agreement between data and Monte Carlo is inconsistent at best. There is a clear layer dependence, however, and it is readily observed that agreement decreases with layer number. In a way, this explains the lowered efficiency observed above, but more than anything, it simply begs the question why this is so. The gas mixture is the same through out, so there ought not be a noticeable difference in α over this range. Further, in both cases, it is observed that the efficiency does not go to 0 at the straw radius as predicted by our theoretical model, though this is expected to some extent based on the "smearing" of reconstructed tracks in this region.

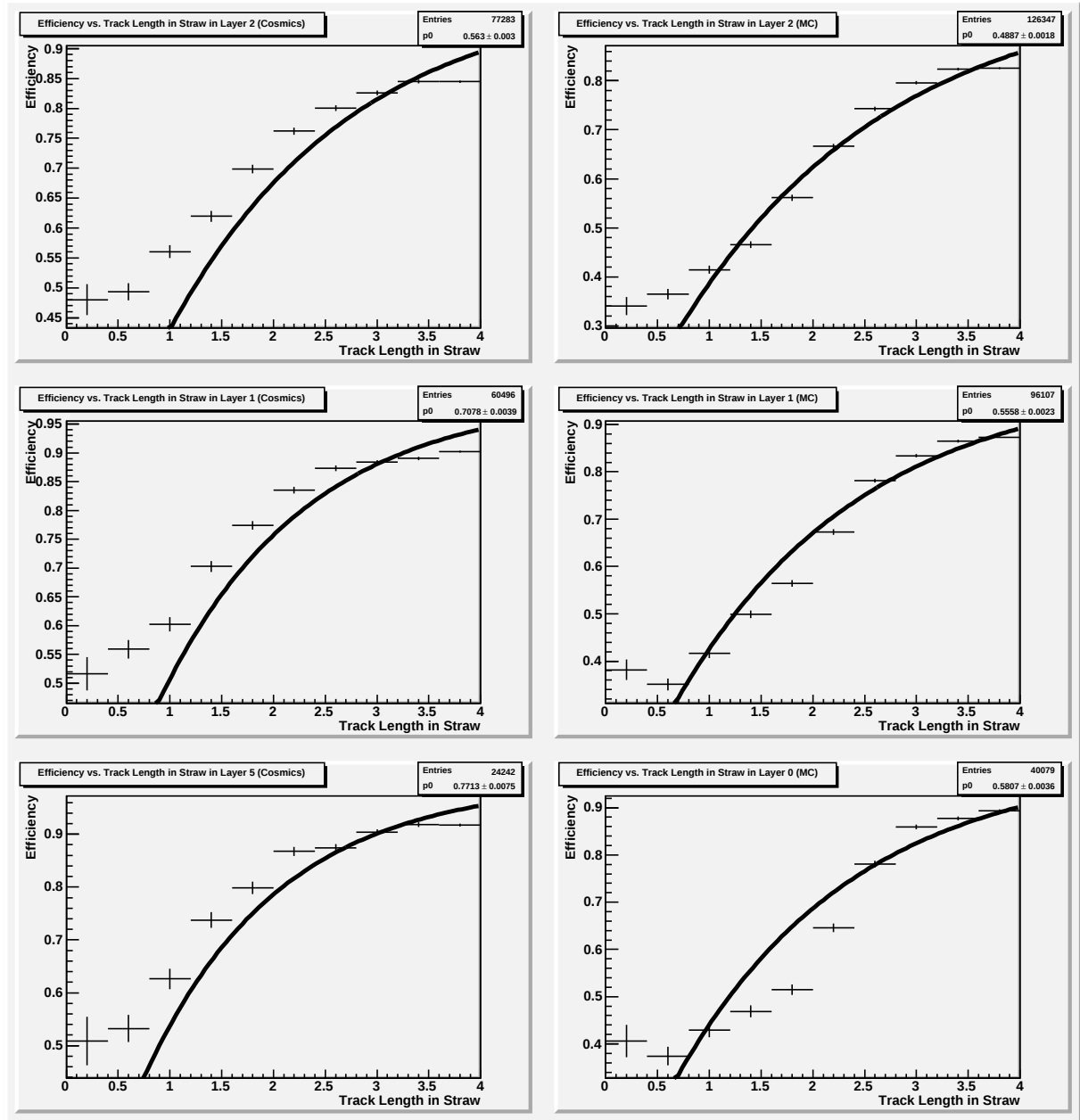


Figure 6: Efficiency as a Function of Track Length by Layer

4.1.4 Efficiency as a Function of Track Quality

The next major variable under consideration is the quality of tracks used in the analysis. It has been noted in numerous places that our efficiency calculations rely strongly on good track reconstruction. For particles that experience lots of MCS, such a reconstruction is not always possible, since the default reconstruction algorithms (for SR1 data sets) fit tracks only to a straight line. While there is no way to ignore such tracks in the initial analysis, it stands to reason that such tracks, once reconstructed, will have a higher chisq/ndof than "good" tracks which experience little to no MCS. In this way, we can refine the data and Monte Carlo samples to minimize the effects of poor track reconstruction on the observed efficiency. In many ways, this more accurately reflects how the detector will perform under real circumstances, since in these cases, momentum information will be available, and it will be possible to select out only high- p_T tracks.

Before we perform a cut on track chisq/ndof , however, it is important to explore and compare the spectra in data and Monte Carlo. If, for instance, the spectra are markedly different, a cut at a given value will produce different results in data and Monte Carlo. Obviously, this has potential to color an comparison between the two. Presented in Figure 7 is the chisq/ndof spectra obtained for data and Monte Carlo. While agreement is relatively good, there is a noticeable shift in the data towards higher values. The differences are, however, small enough that the cuts we will be considering should not have markedly different effects on data and Monte Carlo.

Since it is apparent that data and Monte Carlo will react similarly to a cut at a given value, we now consider how the efficiency of a given detector element changes with track quality. The first result we present is a simple distance profile, where a cut excludes tracks with a $\text{chisq}/\text{ndof} \geq 6$. Cut and uncut samples of layers 0, 1 and 2 are superimposed to allow for comparison to previous results. As can be seen in Figure 8, the effects are most dramatic, in both data and Monte Carlo, in layer 2 and least dramatic in layer 0. Happily, this is consistent with our earlier hypothesis that track segments in the outer layers are less well-determined and thus more susceptible to errors caused by tracks with lots of MCS.

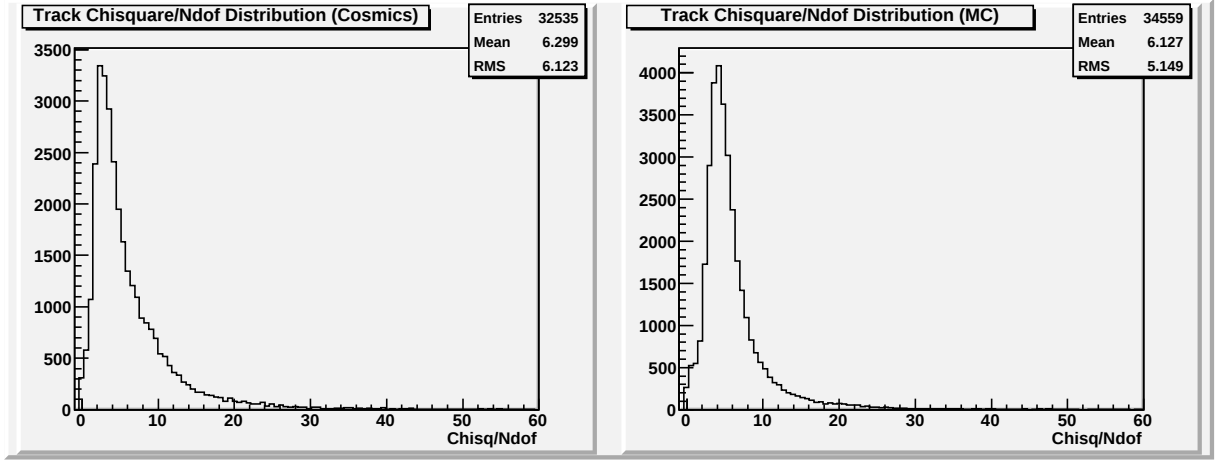


Figure 7: Chisq/ndof Spectra

Two results are here worth noting. First and most importantly, the efficiencies appear to lose some of the layer dependence they had in the uncut samples. This allows us to speak more clearly about an "average" efficiency for the TRT, and is thus a very pleasing result to have. Secondly, though, the efficiency observed inside the straw is becoming increasingly uniform. That is to say that an increase in track quality flattens out the distribution of efficiency as a function of track length inside the straw. In this way, the detector approaches the ideal "guess" we made earlier regarding what an efficiency profile should look like. In short, the detector responds more nearly as expected when presented with the type of tracks the reconstruction algorithm is designed for.

To get a better sense for the quantitative change in average efficiency observed as track quality becomes increasingly good, we have also plotted the average efficiency over an entire layer as a function of track quality. As can be seen in Figure 9, the average efficiency of the detector does indeed converge nicely to a single, layer-independent value (up to uncertainties) in both data and Monte Carlo. The pattern of convergence is, however, atypical in the cosmics sample, since layer 2 becomes more efficient than either layers 0 or 1 once the strongest cuts are performed. As yet, there is no positive explanation for this phenomenon.

As a final note for this exploration, it is worth pointing out that the p_T spectrum of the Monte Carlo is almost certainly not the same as is observed in Nature. While efforts have been made to understand this point better, no real conclusion has been reached as of this writing, and it is likely,

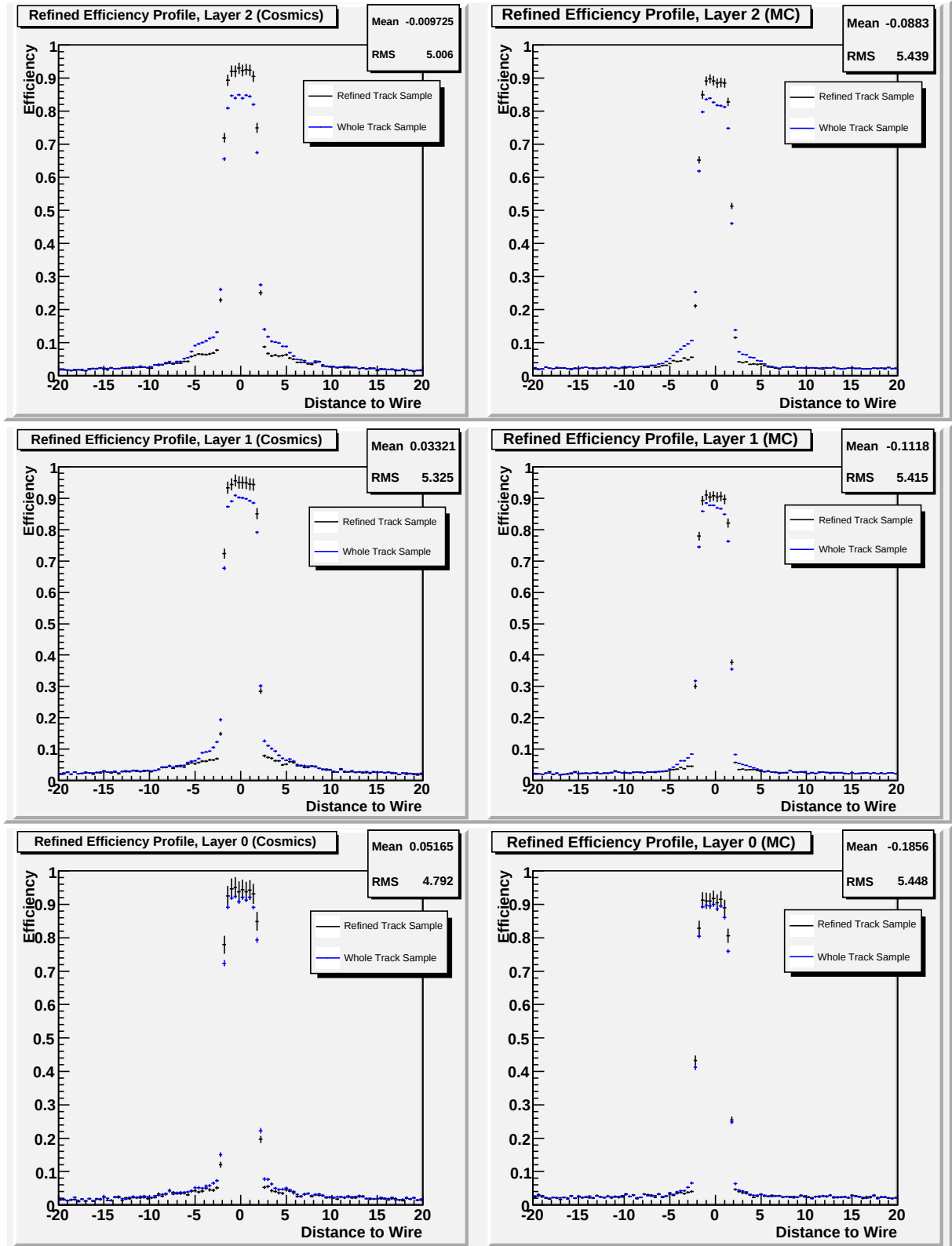


Figure 8: Efficiency vs Distance with Refined Track Samples by Layer

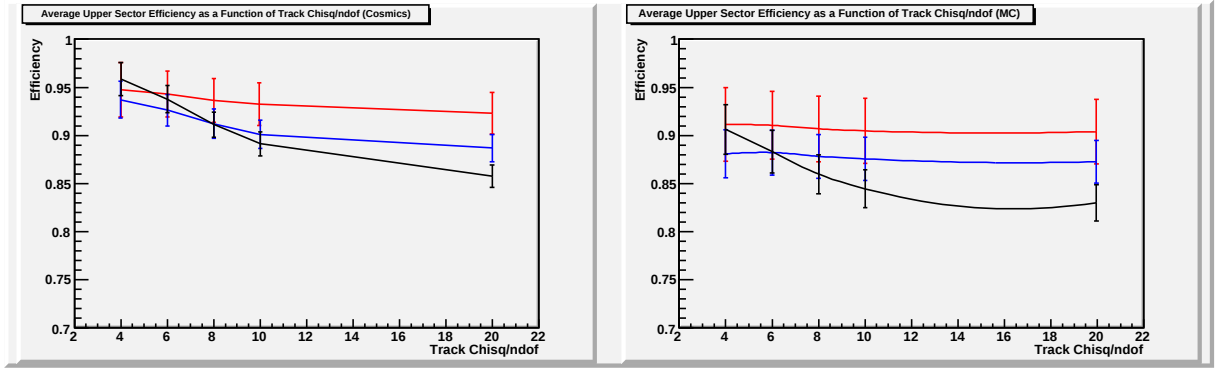


Figure 9: Average Efficiency per Layer vs. Track Chisq/Ndof

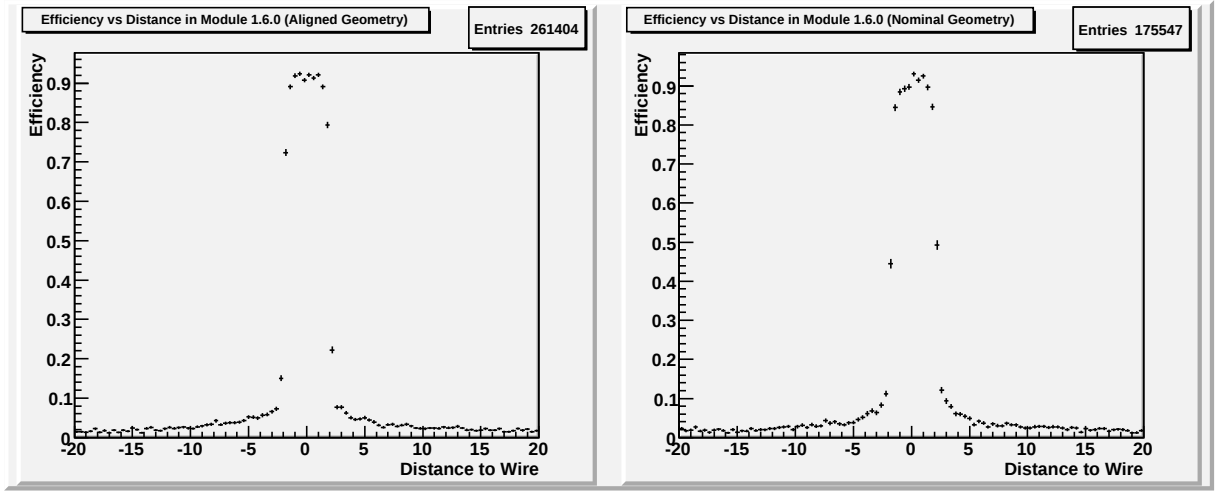


Figure 10: Efficiency vs Distance with Aligned and Unaligned Geometries

with the installment of the barrel on the beam line, that no such conclusion will be made before the LHC turns on. Though one cannot draw an exact correlation between the p_T spectrum and the observed chisq/ndof spectrum, previous discussions have hinted that some connection exists.

4.1.5 Efficiency as a Function of Alignment

The plots presented up to this point were generated using alignment constants developed by other members of the TRT working group that compensate for minor misalignments present in the real detector. Since the method employed here to study efficiency depends intimately on the location of the straws and their relation to the tracks, it is interesting to see how minor misalignments in the detector affect its performance.

Rather than manually introducing arbitrary misalignments, it suffices to use the detector as it appears in the absence of any alignment in the first place. Given this, the efficiency of the upper sector is studied in each layer with both aligned and unaligned detector geometries. As can be seen in Figure 10, the effect of alignment on the detector's performance is actually quite visible. It is apparent that efficiencies in the unaligned detector are visibly asymmetric about the wire and thus actually takes on a lower average value if the efficiency of the straw is integrated over the straw itself. Both of these results are consistent with what one would expect concerning a misaligned detector, so for the rest of this paper, only the properly aligned detector will be considered.

4.2 Straw Level Efficiency

The next major step in these studies is to look at the efficiency of the TRT in terms of the individual straws in each module. While there will inevitably be statistical issues with such studies, it is still worth while to explore the performance of the TRT at this level in order that potentially dead straws can be identified and properly treated. To reduce statistical uncertainties, only straws with more than 10 expected hits are considered for any of these plots. Finally, comparisons with Monte Carlo are more difficult on this level, since all elements of the detector function as they should in the simulations. One of the major findings from these studies was the presence of large (apparent) inefficiencies in layer 2 of the lower sector.

4.2.1 Efficiency Per Straw within the Module

In this section, we present results obtained from the study of straw-level efficiency within each module. In order to assign "an efficiency" to each straw, it is required that we pick some distance from the straw R and integrate the observed efficiency up to that point. For reasons described above, this value of R was chosen to be 1.5 mm. Figure 12 shows the preliminary results for the efficiency per straw distributions of modules 1.6.0, 1.6.2, 1.22.0, and 1.22.2.

As can be seen in Figure 12, lower-right half of module 1.22.2 appears to be missing in the sense that efficiencies in this region appear to be 0. This same affect was observed for $\phi = 23$. In essence, it appears that part of the lower sector contains mostly dead straws. However, this is not, in any likelihood, the case. If the straws in this region are "dead" in the sense of not being able to mechanically conduct a signal, then we would expect them to be flagged by online monitoring systems. Since no such flagging occurs, it stands to reason that something else must be at work. To investigate this, occupancy diagrams with the total number of hits a given straw registers were made. If the occupancy distribution of straws is the same in the upper and lower sectors, a miscabling of straws in the lower sector is the most likely explanation for the dead region. This, however, would imply a higher noise rate in the lower sector as well, something that is not in fact observed. Further, the occupancy diagrams (Figure 13) themselves are not observed to be very similar at all.

Clearly, there is a marked difference between the total occupancy in the upper and lower sectors of layer 2. Further, these occupancy diagrams replicate very closely the efficiency plots shown above (especially with the large region of low efficiency and occupancy in the lower half of the lower sector). This makes it unlikely that a portion of the lower sector has simply been miscabled. As of this writing, the best explanation for this phenomena is simply that these regions of the detector fall outside the scintillator acceptance range for some unknown reason and thus get do not get hit with that many incoming muons.

Finally, it is interesting to study how many straws in a module possess a given efficiency and see whether this can yield any insight into the lowered efficiency of the lower sector. As before, only straws with at least 10 expected hits are included for statistical reasons, but all modules in layer 2 of each ϕ -sector are considered. It is hoped that the distribution in the upper sector will appear similar to the one from the lower sector (at least if the straws in the dead region are excluded). In Figure 14, such results are presented, and as expected, the distributions appear remarkably similar, even down to the proportion of straws with the peak value for the efficiency. Thus, it is safe to conclude that there is nothing really pathologically wrong with the straws in layer 2 of the lower sector, though some other factors must clearly be involved.

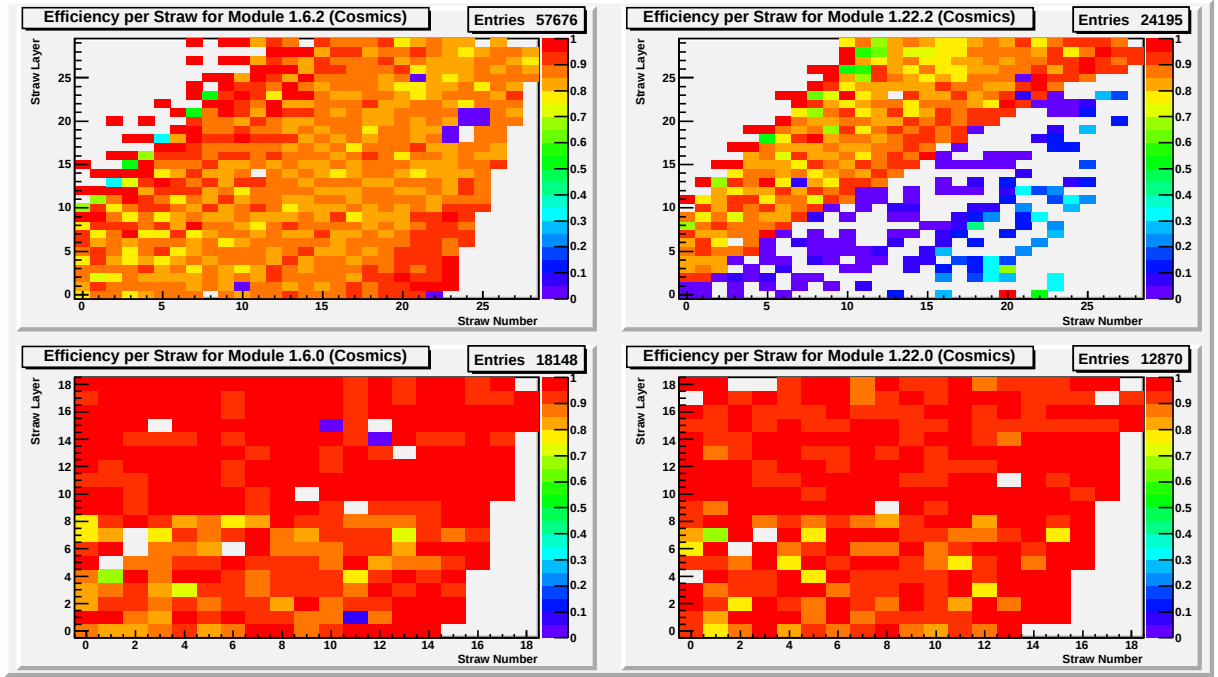


Figure 11: Efficiency per Straw per Module in Data

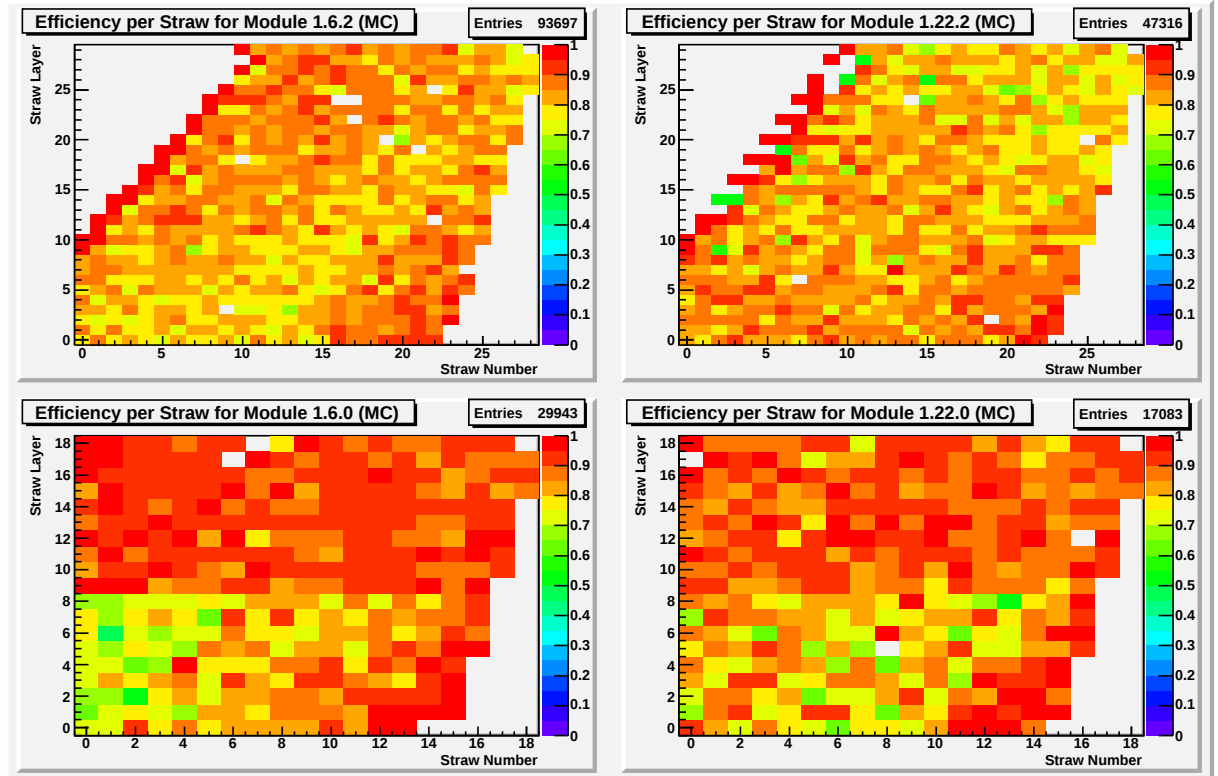


Figure 12: Efficiency per Straw per Module in Data and Monte Carlo

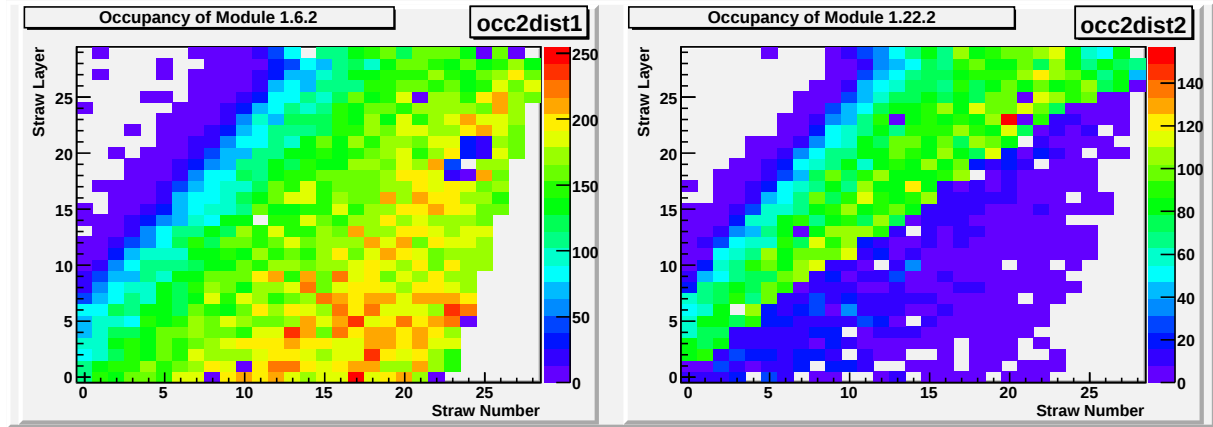


Figure 13: Comparison of Straw Occupancy in Layer 2

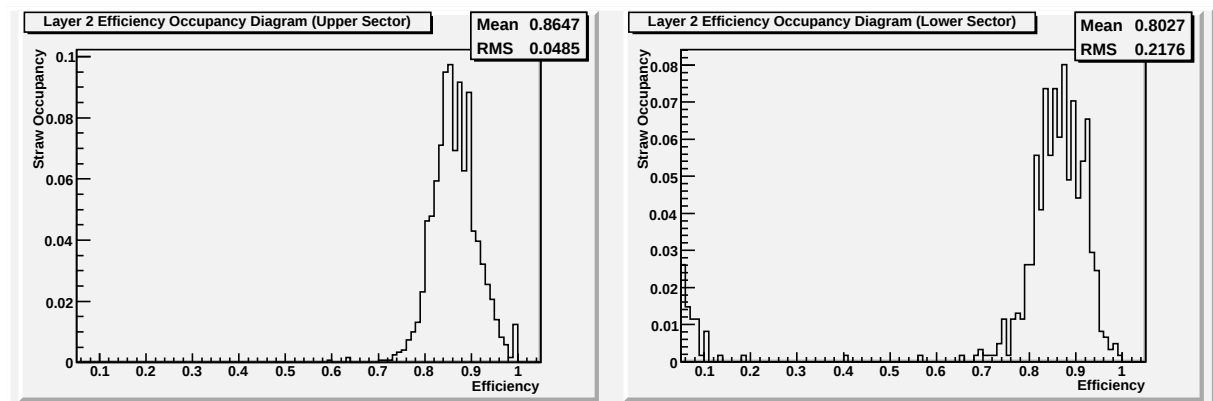


Figure 14: Distribution of Straw Efficiency in Layer 2

One obvious difference between the two distributions, however, is the larger fraction of genuinely inefficient straws in the lower sector. This is a somewhat puzzling finding, since truly dead straws would register an efficiency of exactly zero, while truly miscabled straws might have non-zero efficiency but a regular occupancy. The fact that neither of these expectations are realized points to something deeper, though the precise nature of this difference is not known.

4.3 Efficiency Conclusion

In total, the TRT appears to be operating with a tolerable efficiency. Given the difficulties of our definition, there is no way to pin down exactly how efficient the detector is going to be when presented with real data. Rather, the correspondence (or lack thereof) between data and Monte Carlo can, if the same definition of efficiency is used, shed light on how well the simulations are reproducing the internal state of the detector and help guide future refinements of the latter.

That being said, the TRT appears to be operating at around 92 % efficiency, with little to no layer dependence observed once even modest track quality cuts are made. This clearly bodes well for the overall performance of the detector, since most particles can be expected to produce TRT hits and that this expectation holds relatively well across the entire detector. More importantly, though, very good agreement between data and Monte Carlo was observed in many qualitative features of the efficiency. A few key differences were present though.

Most significantly, the Monte Carlo was observed to be less efficient by about 2 - 3 %, even in the face of track quality cuts. Naively, we would of course expect that the Monte Carlo would be more efficiency (if anything) than the data, so this result is indeed somewhat surprising. As of this writing, no concrete reason can be given for this phenomenon, though it is strongly suspected that the problems encountered in simulating the p_T spectrum of cosmic muons mentioned above lie at the heart of this difference. Assuming the p_T spectra are different, it would follow that data and Monte Carlo would contain different amounts of MCS, even once track quality cuts have been made. As discussed above, this could certainly affect the overall efficiency of the detector.

The final key result from these studies is the apparent dead zone in the lower sector of layer

2. As discussed above, there is as yet no definitive explanation of this result either. It is observed, however, that a large fraction of straws in layer 2 of the lower sector have exactly 0 (or at best, nearly 0) efficiency. These same straws, further, have drastically lower occupancies than those that are functioning properly, so it is unlikely a simple miscabling is to blame.

5 Noise Studies

By definition, noise is a straw registering a hit when there appears to be no track nearby. Since all we have to go on are reconstructed tracks, however, "noise" becomes a rather subtle quantity, since a poorly reconstructed track can miss some straws that legitimately had a particle pass by them and thus create the appearance of noise within the detector. In this section, we outline how noise is studied here and present some results analogous to those shown above for the detector's efficiency.

5.1 Module Level Noise

At a broadstroke, it is desirable to have some sense of how noisy the detector is as a whole. This is exactly analogous to the module level efficiency studies conducted above, however now we are interested in the efficiency outside of the straw radius rather than within it. This is, in fact, exactly the manner in which these studies are conducted.

5.1.1 Noise Levels per Module

The goal of this section of study is to determine the average noise levels in the detector. More rigorously, the average efficiency of the TRT for distances well outside the straw radius can be plotted to determine the "average noise" of a module. For these studies, the interval of 10 to 20 mm was chosen, primarily to agree with average noise levels as measured by other methods.

Specifically, the noise levels per layer and per ϕ -sector are presented to determine if there is any significant geometric dependence. In Figure 15, the average noise level per layer (integrated over all ϕ -sectors) is presented. A slightly lower noise level is observed in layer 0 than in the other

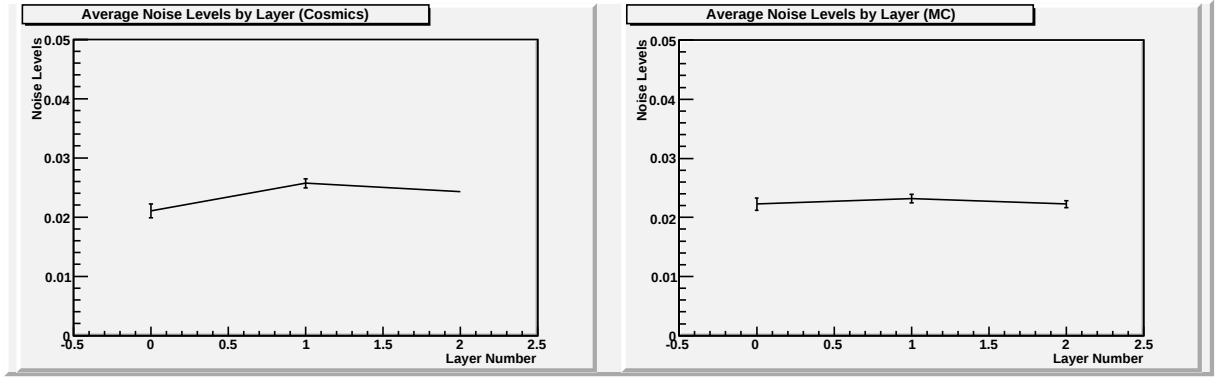


Figure 15: Average Noise Levels by Layer

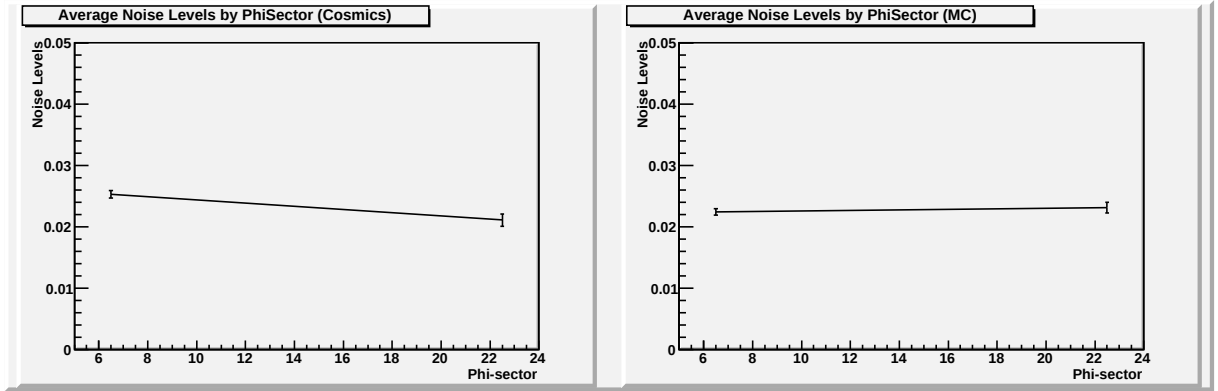


Figure 16: Average Noise Levels by Phi-Sector

two, but this effect is generally small.

In like manner, it is possible to consider the average noise level as a function of ϕ -sector, since one could guess that the lower-than-expected efficiencies in layer 2 of the lower sector will “reappear” as entirely noise. As such, the noise level is presented for the upper and lower ϕ -sectors separately (integrated over all layers). These results are presented below in Figure 16.

As hoped, the noise levels in the TRT are fairly uniform and not terribly high. Further, it is observed that the Monte Carlo is consistently less noisy, a result that is well expected. Again, it should be noted that these results include the contribution of poor track refitting, so this is not strictly a measure of a wire’s mechanical propensity to fire randomly. This is important, since the Monte Carlo is designed to minimize the presence of random noise, so in principle its noise levels should be even lower.

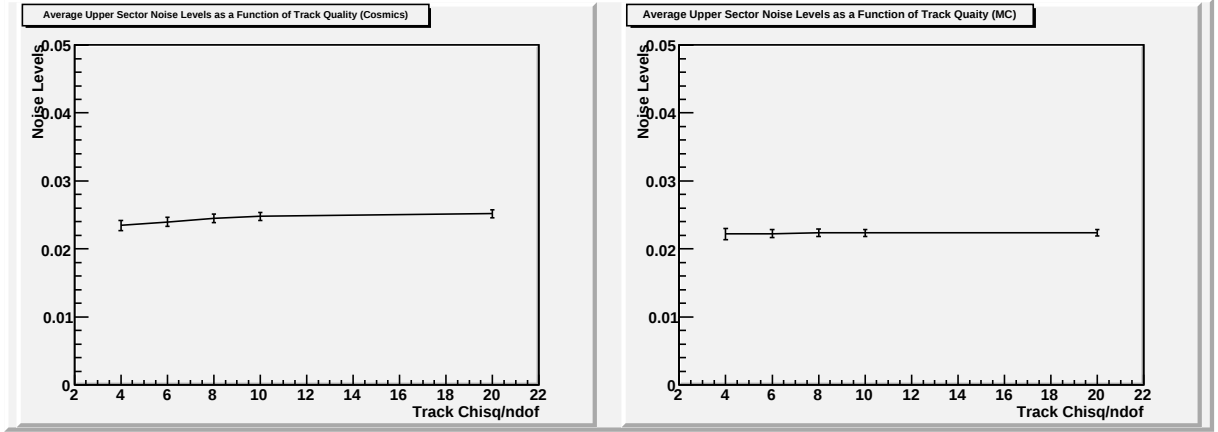


Figure 17: Noise Levels vs. Track Quality

5.1.2 Noise as a Function of Track Quality

Following the same arguments used above with regards to the efficiency measurements, it seems prudent to see if track quality has any effect on the noise levels in the detector. It is hoped that a cut on track quality will remove some of the noisy hits that are due only to a poor track fitting and not the mechanical properties of the wires themselves. Since the noise levels appear to be fairly constant over different layers, they are now studied only for the entire detector.

As can be seen in Figure 17, there is a noticeable difference in the noise levels observed in data as track quality increases. This is expected based on the arguments given above. There is no similar change, however, for the simulated data. This is somewhat unexpected since the same tracking errors that are present in the data should survive in the Monte Carlo, but this does not pose a major problem, since noise levels in both data sets are already so low and in such good agreement.

5.2 Straw Level Noise

As a final point of interest, noise levels per straw will be examined in direct analogy to the efficiency studies in section 4.2.1. The goal of these studies is primarily to confirm the uniformity of the noise levels discussed above. While modest in scope, these explorations can also serve as a way to verify the existence of potentially noisy straws and determine the extent of cross-talk that occurs between

wires. Unfortunately, many of these more interesting studies on the TRT's noise levels are not presented here, though there is hope they will be in the near future.

5.2.1 Noise as a Function of Time over Threshold

When a wire registers a hit, the total time the wire is over the electronic's threshold voltage is recorded. As one could imagine, this time is catalogued as the hit's time over threshold (or simply ToT). Since the electronics read out each straw approximately once every $3\ \mu\text{s}$ (citation from the TDR??), ToT measurements are made in units of this read-out time and are referred to here as bins. Since the ionisation event due to a passing particle within a straw is not an instantaneous process, it is expected that any hit due to such an event will record a ToT of about two or three bins ($6 - 9\ \mu\text{s}$). A hit purely due to electronic noise will, however, only tend to have a ToT of one bin, since there is nothing actually driving the wire above threshold.

As such, one way to eliminate purely electronic noise from the detector is to require a ToT in excess of one bin. Aside from cutting out unnecessary hits in the detector, this process can give further insight into the extent which the noise levels currently observed in the detector are due to tracking problems. A visible dependence of the noise level on track quality has already been observed, so it would be interesting to see (or if) this changes if ignore hits that essentially have to be noise.

As with the track quality cuts made above, it is informative to compare the ToT distributions seen in data and Monte Carlo, if only to make sure something isn't going terribly wrong in one or the other. The results of such a comparison can be seen in Figure 18, and it is clear from this there is little effective difference between the two.

With the relative tameness of the hit ToT distribution established, it is possible to compare cut and uncut noise levels. In Monte Carlo, it is hoped such cuts will dramatically reduce noise levels, while in data the result is much less clear, though some reduction is still expected to occur. In Figure 19, a distribution of straw-level noise occupancy (i.e. the number of straws in the detector with a given noise level) is generated with a raw data sample and then again for a sample that

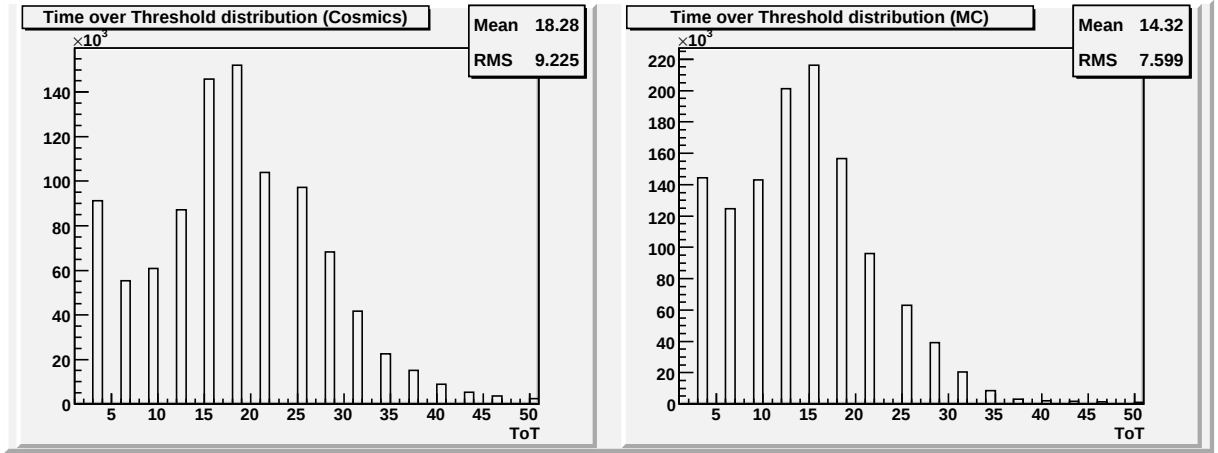


Figure 18: Time over Threshold Distributions

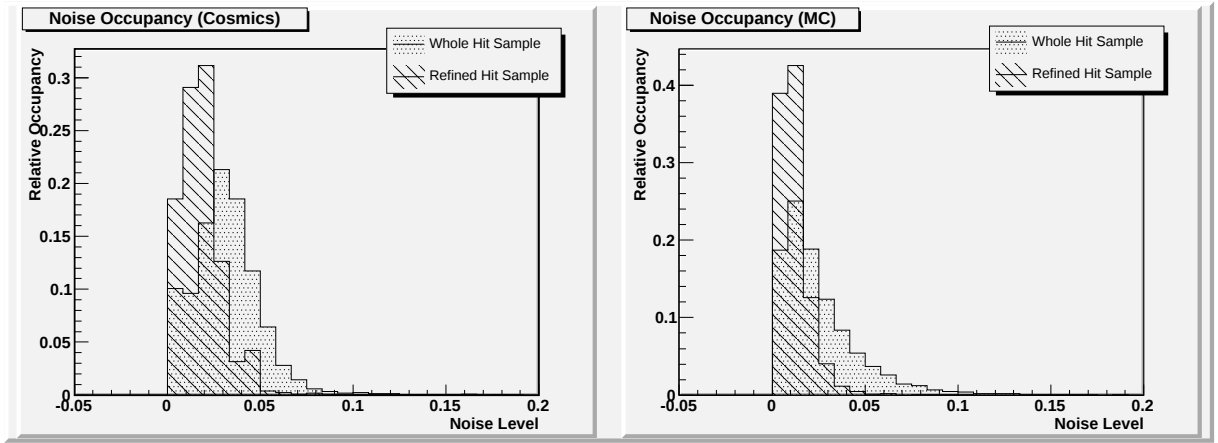


Figure 19: Straw-Level Noise Occupancy

excludes hits with a ToT of just one bin.

As hoped, the average noise level in both data and Monte Carlo drops steeply once this cut on time over threshold is made. While the relative agreement between the two is still not perfect, it is encouraging that both respond positively to such a cut.

5.2.2 Noise per Straw within a Module

In direct analogy with the efficiency plots from above, it is possible to study how much noise is present in each straw within a given module by essentially considering the "efficiency" of that straw for distances far outside the straw radius. To be consistent with earlier measurements, we

consider the region 10 to 20 mm away from the wire. In Figure 21, the plots of the integrated noise levels in modules 1.6.0, 1.6.2, 1.22.0, and 1.22.2 are presented.

As can be seen in Figure 21, the noise levels appear rather uniform, save for a few particularly noisy straws. Happily, this agrees with the intuitive picture that noise should be due more to intrinsic electronic uncertainties than anything geometrical like the straw's position relative to the scintillators. Further, the high noise levels observed in these straws probably arise from a very small number of expected and observed hits. Near the limits of the detector's geometric acceptance, fewer tracks are recorded and so fewer statistics are expected. In this environment, even a single noisy hit would appear significant, even though it would likely behave completely normally if it were placed in another part of the detector.

In Figure 23, noise levels per straw are presented when hits with a ToT of just 1 bin are excluded. As expected, the noise levels visibly decrease once this cut is applied, though the same generally uniform pattern is observed.

5.2.3 Preliminary Measurements of Cross-Talk

Cross-talk between wires occurs when a track passes through one wire but a hit is registered somewhere else. In principle, this can occur because of some physical leakage of drift electrons from one straw to another (in which case cross-talk occurs only between geometrically adjacent straws), but this is not expected to be a major contributor to the observed levels of cross-talk, since the straw coverings themselves are fairly durable. To investigate this kind of cross-talk, one can look for higher-than-expected noise levels up to one straw radius away from a given wire.

More than likely, the majority of cross-talk occurs one step removed from the straws themselves, in the readout electronics. Since groupings of eight wires are read out by the same ADS chip, it is possible (and likely) that some mixing of signals occurs at this level, leading to the rather counter-intuitive picture of cross-talk between wires with the same read out chip. To gain an estimate of the amount of cross-talk generated by this process, it is necessary to look for elevated noise levels near the straw, but farther out than just one straw radius.

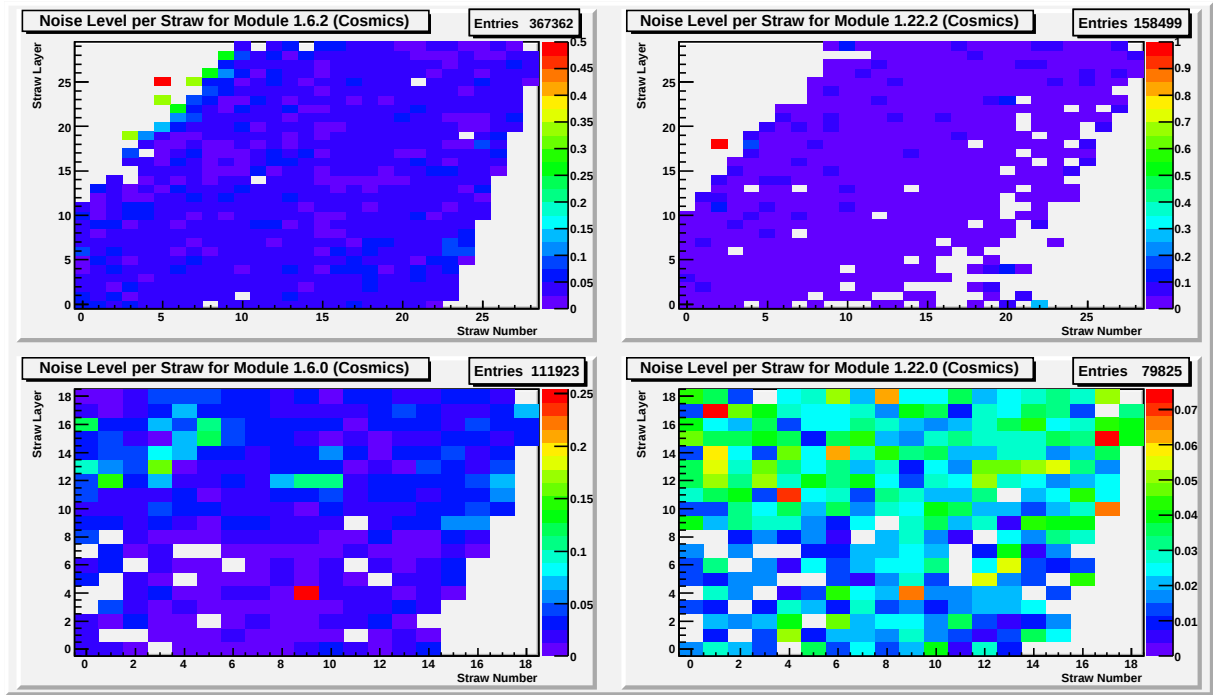


Figure 20: Noise Level per Straw per Module in Data

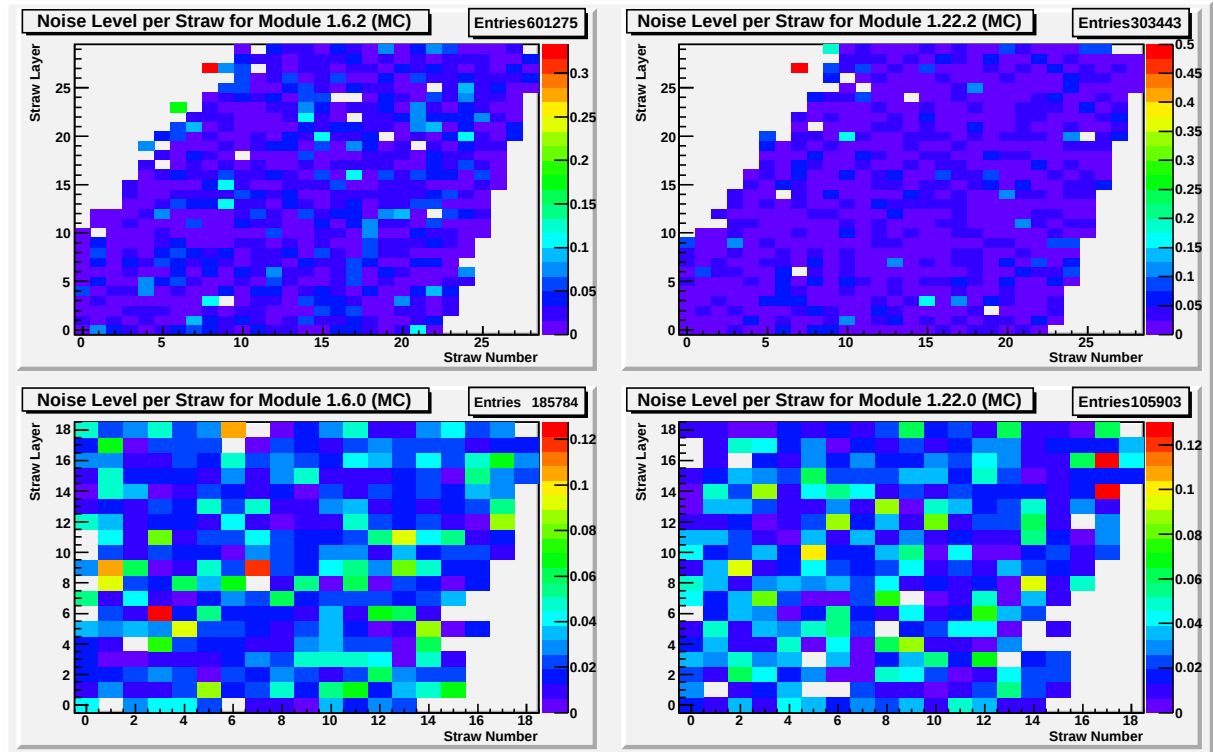


Figure 21: Noise Level per Straw per Module in Monte Carlo

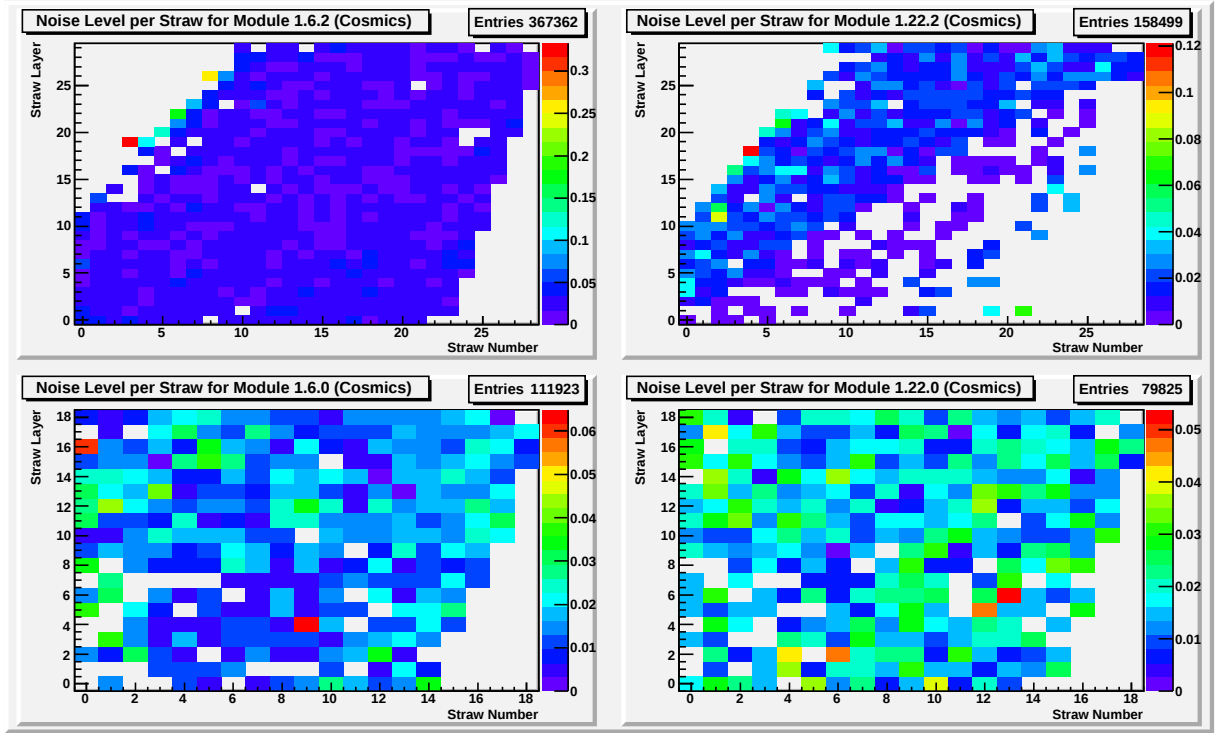


Figure 22: Noise Level per Straw per Module in Data, ToT ≥ 1 bin

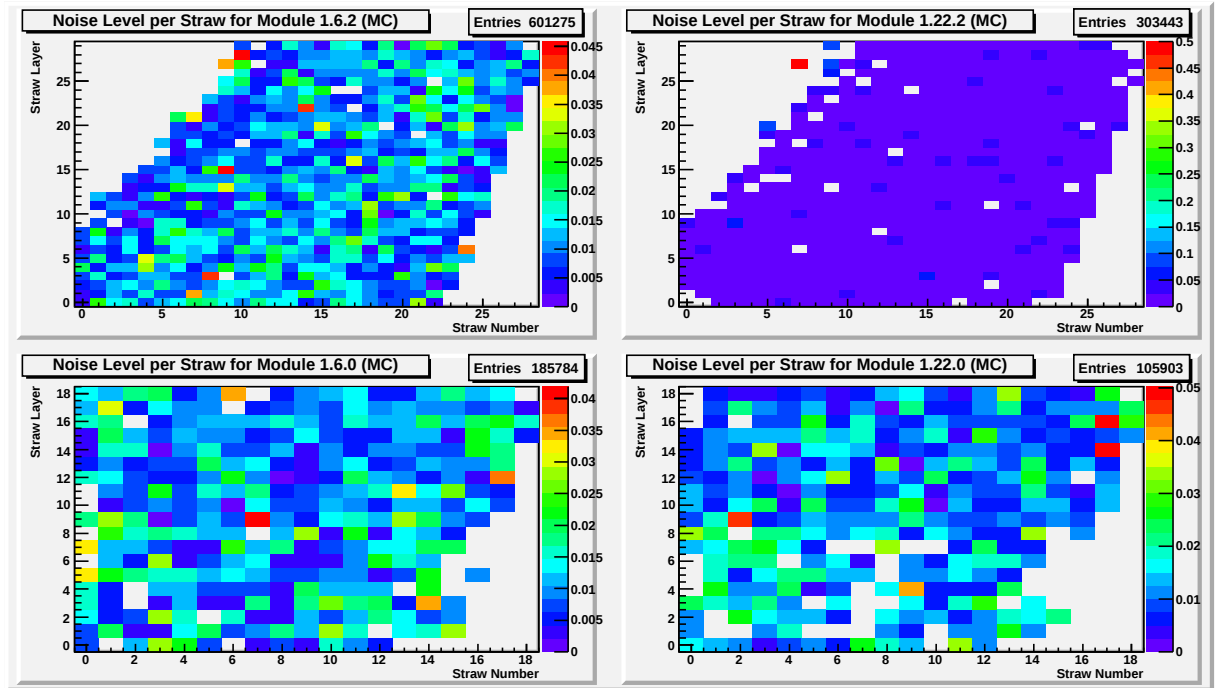


Figure 23: Noise Level per Straw per Module in Monte Carlo, ToT ≥ 1 bin

Given these considerations, the noise level observed 5 - 10 mm away from the wire are recorded and compared with the noise levels seen 10 to 20 mm away. By subtracting out this later background, an estimate for the amount of cross-talk can be attained. The interval 5 - 10 mm is chosen to allow ample room for minor tracking errors to smear the efficiency of a wire out to adjacent elements (which is not cross-talk) while still capturing a region near enough for chip-level cross-talk to occur. While neither of these processes are a rigorous measure of cross-talk, they do serve to gain an approximate cross-talk level, in preparation for more in depth studies to be conducted in the near future.

When such studies are done, it is observed that the noise level near a wire is 0.042549 ± 0.000943854 in data and $0.0262334 \pm 0.000623844$ in Monte Carlo. By contrast, when regions very far from the wire are considered, noise levels drop to $0.0242602 \pm 0.000514959$ in data and $0.0225845 \pm 0.000421213$ in Monte Carlo. If the Monte Carlo results are taken to offer a reasonable estimate for the level of tracking error which contributes to the noise levels in the region nearer the straw, then cross-talk in the cosmics samples still accounts for around 1.5 % of the observed hits in the detector. While this is by no means rigorous, it certainly offers motivation for further study.

5.3 Noise Conclusions

As can be seen from these analyses, the noise levels in the TRT appear to be around 2 - 3 % of the total number of hits registered by a given straw, with even lower values if the hit's time over threshold is taken into consideration. Unfortunately, these numbers are not to be taken too seriously, since the method for their computation relies more on the tracking system than would be ideal for a real study of intrinsic noise. That being said, these values agree reasonably well with those found through other analysis methods, especially so when track quality cuts are made.

Finally, there is little utility in comparing the results obtained in this section to Monte Carlo, since the simulations are supposed to have no intrinsic noise. It is possible that the "noise" as measured by the Monte Carlo would give an indication of how many noisy hits can be properly

attributed to tracking issues (since this is essentially the only reason for there to be noise in the Monte Carlo), but such an estimate would be approximate at best, so we neglect to make it.

References

- [1] T. Akesson, F. Anghinolfi, E. Arik, et al.
Status of design and construction of the Transition Radiation Tracker (TRT) for the ATLAS experiment at the LHC.
Nuclear Instruments and Methods in Physics Research A 522 (2004).
- [2] Hulsbergen, Wouter.
A Study of Track Reconstruction and Massive Dielectron Production in Hera-B
University of Amsterdam (2002).
- [3] DAQ Online Logbook at SR1
<http://pcphsctr02.cern.ch/cgi-bin/logbook.cgi?locn=SR1logbookIndex=12physics=1tselected=0>
April 2007
- [4] <http://hausch.home.cern.ch/hausch/ATLAS/SR1/>
April 2007
- [5] <http://www.hep.physik.uni-siegen.de/pics/>
April 2007
- [6] LHCb Technical Design Report
<http://lhcb.web.cern.ch/lhcb/TDR/>
April 2007
- [7] TOTEM Technical Design Report
<http://totem.web.cern.ch/Totem/>
April 2007