

Measurement of the charged current
quasi-elastic cross-section for electron
neutrinos on a hydrocarbon target

by

Jeremy Wolcott

Submitted in Partial Fulfillment

of the

Requirements for the Degree

Doctor of Philosophy

Supervised by

Professor Steven Manly

Department of Physics and Astronomy

Arts, Sciences and Engineering

School of Arts and Sciences

University of Rochester

Rochester, New York

2016

To my father Alan, who taught me the joy of learning for learning's sake.

Biographical Sketch

The author was born in Highland Park, IL, and grew up in Richmond, VT. He attended the University of Rochester and earned a Bachelor of Science degree in Physics, as well as a Bachelor of Arts in Mathematics, during his studies there, graduating Magna Cum Laude in 2007. His research experience in particle physics began during this time when he participated in a Research Experience for Undergraduates program at Fermilab, searching for novel bottom quark resonances under the direction of Arie Bodek.

The author's graduate career at the University of Rochester began in 2008 when he enrolled in the Ph.D. program in Physics. After obtaining his Master of Arts degree from the University, he joined the MINER ν A experiment in 2009, working under the supervision of Steven Manly. He has been an active participant in a number of facets of the experiment including data acquisition software engineering and operations, reconstruction software development, software architecture management, and the physics result detailed in this work.

Acknowledgments

My first and my most heartfelt gratitude go (of course!) to Kate, who so far has exhibited no regret at deciding to share my name and my life shortly before I embarked on my graduate school journey. Her cheerful patience with “seven” (-ish) years of a graduate student’s... umm... “flexible” work schedule and fickle social life has been unwavering and perpetually encouraging. Without her dedication to a regular routine (not to mention paycheck!) and her boundless determination to keep me and our little zoo fed, exercised, and healthy, this trek would’ve come out a whole lot differently.

Thanks must also go to my mom (and dad, *in abstentia*): first, for so many years of entertaining my ceaseless “why?” questions, and then beyond that, for their commitment to teaching me the structure and discipline that ultimately enabled an impulsive boy with attention problems to undergo an improbable transformation into a professional scientist. And not only they but also my sister, my brother, and the rest of my extended family deserve my gratitude for asking about what I do, and then for humoring me when I answered as if it was actually interesting; for visiting, and writing, and calling, even when I got distracted and forgot to reciprocate; and for never pressuring me to hurry up and get on with life. Their forbearance and supportiveness has been admirable, and I’m grateful. In a very similar way, Doug Jackson’s constant encouragement, perspective, and genuine interest in the progress of my own work and the general advancements in the field has been inspiring and reassuring, and I can’t imagine doing without it. And all of the fine folks at Browncroft Community Church have provided a vibrant community for discussion, recreation, and personal growth that has been essential to my sanity and perseverance.

I am of course indebted both to Steve Manly and to Kevin McFarland for their outstanding scientific and professional advice, and their many questions and suggestions, and fruitful discussion, over the years. It has been wonderful to work with a group so clearly committed to the success of the students in it, and I am exceptionally grateful for the opportunity! I am equally grateful to Phil Rodrigues for the countless hours he patiently spent with me deliberating analysis technique, software arcana, and the minutiae of particle physics' daily grind; his fingerprints are all over this analysis, to its great benefit. Thanks as well to Chris Marshall for great office conversation and some very timely analysis suggestions; to my other colleagues at Rochester (team-building exercises and all); to Gabe Perdue, who, maybe unwisely, gave a bright-eyed first-year student with more technical ability than common sense the keys to the MINERvA DAQ; and to my many remaining colleagues on MINERvA who have contributed to my success in innumerable ways big and small.

I must also hasten to thank the wonderful administrative staff I have had the privilege of working with over the years. Connie Jones stands out for her perpetually flawless negotiation of the quagmire of travel funding and restrictions that sometimes threaten to get in the way of actually doing science; she deserves special thanks for what is a very thankless job. And I am personally indebted to Laura Blumkin for going above and beyond the call of duty in helping me to finally get this degree program finished.

Finally, above all, while his name will not appear elsewhere in the pages of this document, his world and his design thereof are the foundation, inspiration, subject, and sustenance of all of this work. Soli Deo gloria.

Abstract

Appearance-type neutrino oscillation experiments, which observe the transition from muon neutrinos to electron neutrinos, promise to help answer some of the fundamental questions surrounding physics in the post-Standard-Model era. Because they wish to observe the interactions of electron neutrinos in their detectors, and because the power of current results is typically limited by their systematic uncertainties, these experiments require precise estimates of the cross-section for electron neutrino interactions. Of particular interest is the charged-current quasi-elastic (CCQE) process, which figures significantly in the composition of the reactions observed at the far detector. However, no experimental measurements of this cross-section currently exist for electron neutrinos; instead, current experiments typically work from the abundance of muon neutrino CCQE cross-section data and apply corrections from theoretical arguments to obtain a prediction for electron neutrinos. Verification of these predictions is challenging due to the difficulty of constructing an electron neutrino beam, but the advent of modern high-intensity muon neutrino beams—together with the percent-level electron neutrino impurity inherent in these beams—finally presents the opportunity to make such a measurement.

We report herein the first-ever measurement of a cross-section for an exclusive state in electron neutrino scattering, which was made using the MINER ν A detector in the NuMI neutrino beam at Fermilab. We present the electron neutrino CCQE differential cross-sections, which are averaged over neutrinos of energies 1-10 GeV (with mean energy of about 3 GeV), in terms of various kinematic variables: final-state electron angle, final-state electron energy, and the square of the four-momentum transferred to the nucleus by the neutrino, Q^2 . We also provide a total cross-section vs. neutrino energy.

While our measurement of this process is found to be in agreement with the predictions of the GENIE event generator, we also report on an unpredicted photon-like process we observe in a similar kinematic regime. The absence of this process from models for neutrino interactions is a potential stumbling block for future on-axis neutrino oscillation experiments. We include kinematic and particle species identification characterizations which can be used in building models to help address this shortcoming.

Contributors and Funding Sources

The defense committee for this dissertation is comprised of Professor Steven Manly of the Department of Physics and Astronomy (the student’s advisor), Professors Kevin McFarland and Lynne Orr (also of the Department of Physics and Astronomy), Professor James Farrar, of the Department of Chemistry, and committee chair Professor David McCamant, of the Department of Chemistry.

The work presented in this dissertation was performed under the auspices of the MINER ν A collaboration, the members of which are listed in appendix D. In addition to fulfilling detector monitoring shifts, the author was heavily involved in a number of detector operations efforts, including construction of the data acquisition software, software used for monitoring data quality while running, and the software environment used for performing shifts from universities away from Fermilab. He performed the *in situ* characterization of photomultiplier tube cross-talk laid out in ch. 5 and implemented the resulting model that was added into the experiment’s detector simulation, in addition to a number of other contributions to the detector simulation. He also served as a member of the team charged with supervising the development and deployment of the experiment’s analysis software stack. Finally, in the course of performing the physics analysis documented in ch. 7, the author additionally was responsible for the development of the particle identification methods discussed in ch. 6 and app. B.

Funding for the work described in this dissertation was provided by the U.S. Department of Energy under grants DE-FG02-91ER40685 and DE-FG02-12ER41823. FNAL is operated by Fermi Research Alliance, LLC under contract no. DE-AC02-07CH11359 with the United States Department of Energy.

Table of Contents

Biographical Sketch	iii
Acknowledgments	iv
Abstract	vi
Contributors and Funding Sources	viii
List of Tables	xii
List of Figures	xiii
 1. Introduction and motivation	 1
1.1. The Standard Model	4
1.2. Beyond the Standard Model: neutrino oscillations	21
1.3. Neutrino interactions and scattering	26
 2. Creating a high-intensity neutrino beam	 39
2.1. Proton acceleration	40
2.2. Protons to neutrinos	41
 3. Characterizing the neutrino flux	 46
3.1. g4numi: GEANT4 simulation	46
3.2. External constraints from hadron production experiments	48
3.3. Flux uncertainties	52

3.4. <i>In situ</i> constraint: neutrino-electron elastic scattering	54
3.5. Final flux prediction	57
4. The MINERνA detector	59
4.1. Subdetector arrangement and function	59
4.2. Active plane composition and instrumentation	63
4.3. Readout electronics and data acquisition system	65
5. Data collection and simulation	71
5.1. Data acquisition and unpacking	71
5.2. Data simulation via Monte Carlo	75
5.3. Data calibration	88
6. ν_e CCQE event reconstruction	101
6.1. Generic reconstruction	101
6.2. ν_e CCQE-specific reconstruction	106
7. Measuring the differential cross-section $\frac{d\sigma}{dQ^2}$ for the ν_e CCQE process	126
7.1. Signal and fiducial volume definitions	127
7.2. Selection of events	128
7.3. Selected event sample	138
7.4. Background model and constraints	140
7.5. Unfolding in observables	148
7.6. Correction for inefficiency	156
7.7. Normalization by flux and target number	156
7.8. Measured cross-sections	158

7.9. Ratios to ν_μ CCQE	160
7.10. Systematic uncertainties	163
8. Conclusions	176
Bibliography	177
A. Uncertainties due to Feynman scaling	183
A.1. Comparing NA61 and NA49	183
A.2. Estimating a systematic uncertainty for NuMI	186
A.3. Effect on the neutrino flux prediction	193
B. Photon-like data excess	198
B.1. Characterization of the excess	198
B.2. Candidate explanations for the excess	206
B.3. <i>Ad hoc</i> models for extrapolating into the signal region	219
C. Alternate analysis used for ratio to ν_μ CCQE	228
D. List of MINERνA collaborators	232

List of Tables

1.1. Lagrangians in field theory	7
1.2. Particles in the Standard Model	17
4.1. Composition of scintillator strips and constructed planes	64
5.1. Progression of electron counts in dynode cross-talk model	83
6.1. Calorimetric reconstruction constants	110
7.1. Fitted background scale factors	145
7.2. Uncertainties on GENIE model parameters. Nominal values for for the FSI parameters in GENIE are not published anywhere and so are not given here.	168
A.1. Applied Gaussian blur parameters	193

List of Figures

1.1.	Feynman diagram for inverse muon decay	29
1.2.	Feynman diagram for ν_e CCQE scattering	32
2.1.	Plan view of Fermilab accelerator complex	39
2.2.	Overview of the NuMI beam line	42
2.3.	NuMI target design	43
2.4.	NuMI horn system	44
3.1.	Predicted NuMI neutrino flux	47
3.2.	Beam attenuation correction	50
3.3.	Reweighted NuMI neutrino flux prediction	52
3.4.	Final <i>a priori</i> $\nu_e + \bar{\nu}_e$ flux	54
3.5.	Effect of flux constraint on $\nu + e$ scattering	56
3.6.	Effect of flux constraint on $\nu_e + \bar{\nu}_e$ flux	57
3.7.	Final flux prediction	58
4.1.	Engineering drawing of a MINER ν A tracker module	61
4.2.	Layout of the modules in MINER ν A	62
4.3.	The MINER ν A coordinate system	63
4.4.	Arrangement of strips within a plane	64
4.5.	Partially assembled PMT box	67
5.1.	Sample pedestal measurement	75

5.2. Measured cross-talk fractions	81
5.3. Reconstructed vs. true cluster energies	94
5.4. Fits to plane displacement and rotation	96
5.5. Cross-talk PMT face to scintillator strip mapping	99
5.6. Cross-talk p -values	100
6.1. True energy fraction captured by various cone angles	108
6.2. Electron energy reconstruction residuals	111
6.3. Endpoint energy fraction variable energy dependence	113
6.4. Endpoint energy fraction data-simulation comparison	114
6.5. Mean dE/dx variable energy dependence	115
6.6. Mean dE/dx data-simulation comparison	116
6.7. Median transverse width variable energy dependence	118
6.8. Median transverse width data-simulation comparison	119
6.9. Separation of PID variables vs. energy	120
6.10. kNN classifier response	122
7.1. Non-MIP-cluster fraction variable distribution	129
7.2. Transverse spread score variable distribution	131
7.3. Position of minimum front dE/dx in particle gun	133
7.4. Optimization of mean front dE/dx cut	134
7.5. Selected cut on mean front dE/dx	135
7.6. Optimization of extra energy fraction variable	136
7.7. Selected cut on extra energy fraction	137
7.8. Selected event distribution	139
7.9. Matching GENIE's coherent π production to MINER ν A	141

7.10. Observables in the extra energy sideband	142
7.11. Observables in the Michel-match sideband	142
7.12. π^+ multiplicity in signal and sideband regions	143
7.13. Fitted signal scale factors	144
7.14. Observables after background fitting	146
7.15. Background-subtracted distributions	147
7.16. Uncertainties on background-subtracted distributions	148
7.17. Migration matrices in the cross-section variables	149
7.18. Progression of residuals during unfolding	150
7.19. Unfolded distributions in the cross-section variables	152
7.20. Electron angle weighting function	153
7.21. Pull distributions (θ_e, E_e)	154
7.21. Pull distributions (E_ν, Q_{QE}^2)	155
7.22. Predicted efficiencies in the cross-section variables	156
7.23. Measured cross-sections	158
7.24. Closure tests in the cross-section variables	159
7.25. Quasi-elastic-likeness breakdown of ν_μ sample	161
7.26. The ratio $\frac{d\sigma_{\nu e}}{dQ_{QE}^2} / \frac{d\sigma_{\nu \mu}}{dQ_{QE}^2}$	162
7.27. Fractional uncertainties on the cross-sections	164
7.28. Absolute uncertainties on the cross-sections	165
7.29. Neutrino interaction model uncertainties on the cross-sections . . .	166
7.30. Weights applied for effective nuclear radius variations	169
7.31. Detector model uncertainties on the cross-sections	171
7.32. Mean cross-talk fractions for PMTs	173

7.33. Pulse heights of hits identified as cross-talk	174
A.1. Bin centers of NA49, NA61, and Barton data	187
A.2. NA49-NA61 fractional residuals	188
A.3. Simulated MINER ν A neutrino x_F, p_T distributions	189
A.4. FLUKA correction factors	191
A.5. Smoothed NA49-NA61 fractional residuals	194
A.6. Smoothed FLUKA correction factors	195
A.7. NA49-NA61 scaling systematic uncertainty	196
A.8. Effect of NA49-NA61 scaling on NuMI flux	197
B.1. Observables in the two-particle sideband	199
B.2. Shape of electron angle in the two-particle sideband	200
B.3. Shape of electron angle in the two-particle sideband	200
B.4. Vertex-attached energy in the two-particle sideband	201
B.5. Upstream inline energy in the two-particle sideband	202
B.6. Extra energy in the two-particle sideband	203
B.7. Vertex R^2 in the two-particle sideband	204
B.8. Vertex z in the two-particle sideband	205
B.9. Shower axis projection in the two-particle sideband	208
B.10. Electron shower mismodeling	209
B.11. Front dE/dx in Michel-match sideband	210
B.12. Front dE/dx with simulated extra protons	212
B.13. Extra energy fraction in the two-particle sideband	213
B.14. $E\theta^2$ in the two-particle sideband	214
B.15. Comparisons of excess kinematics to diffractive model	217

B.16. Upstream inline energy shape, excess vs. diffractive	218
B.17. $E\theta^2$ shape, excess vs. diffractive	218
B.18. Photon <i>ad hoc</i> sample reweight functions	221
B.19. Kinematics of reweighted <i>ad hoc</i> samples	221
B.20. Front dE/dx distributions with fitted <i>ad hoc</i> samples	222
B.21. Separation via extra energy ratio	224
B.22. Separation via median plane shower width	225
B.23. Separation via transverse asymmetry	227
C.1. Extra energy sideband for the alternate analysis	229
C.2. Michel-match sideband for the alternate analysis	229
C.3. Efficiency predictions in the alternate analysis	230
C.4. Measured cross-sections in the alternate analysis	231

1 Introduction and motivation

Particle physics in the twenty-first century is inseparably tied to the Standard Model (“SM”). This *tour de force* composed of insights about elementary particles and the forces by which they interact, assembled from the collective experimental evidence obtained over the previous several centuries, is remarkably successful. Besides synthesizing data from thousands of disparate measurements into a cohesive system, the SM’s predictive ability is unmatched elsewhere in science. Its prediction for the anomalous magnetic moment of the electron is perhaps its most striking accomplishment: agreement between the best calculations and the best measurements is on the order of 10 significant figures, better than any other theory in the history of the scientific method [1]. And multiple times the SM has predicted a new particle (the Z boson and the Higgs boson, for instance) that occupied parameter space beyond the limits of detection, only to be later confirmed when detector technology gained sufficient ground to probe it [2]. Any new data collected is, therefore, stacked up against the predictions of the SM as a matter of course; any new theory is likewise expected to justify its divergence in a compelling fashion.

Yet for all of its predictive power, the Standard Model is known to be incomplete. In broad strokes, major blemishes include: the SM’s inability to unify with the theory of general relativity (itself incredibly successful) in a consistent fashion; the SM’s silence on the large-scale properties of our universe observed in cosmology (such as the evidence giving rise to the postulation of dark matter and dark energy); and theoretical-experimental discrepancies in a few very precise predictions (such as the atomic radius of muonic hydrogen [3]). In addition, to synthesize major discoveries over the past several decades, theorists have been forced to turn to

ad hoc extensions to the model, particularly when confronting the enormous range of masses measured for the family of fermions. Together with the cosmetic distastefulness of the need to fine-tune the roughly 20 free parameters that result so as to align with the data, all of these features are troubling: and as a result, searches for “beyond the Standard Model” phenomena—and the hope that a discovery will lend clues to what shape a more fundamental system might take—are numerous, diverse, and vigorous.

As the lightest and, at least initially, the most apparently featureless, of the fermions described by the SM, neutrinos were long nothing more than a curiosity whose role was relegated to the conservation of energy in weak-force interactions. But the discovery that at least two of the three known flavors bear mass, in stark contrast to the construction of the classic SM, suddenly catapulted the neutrino into the forefront of “beyond the Standard Model” research. Modern, large-scale investigation into the phenomenon of neutrino oscillations (the original litmus test which first demonstrated conclusively that neutrinos indeed do have mass) is now a two-decades-old affair, and though much of the relevant parameter space has been nailed down, major fundamental questions concerning the nature of neutrinos still remain. By virtue of the same weakness of neutrino interactions that originally relegated the neutrino to an anonymous role in particle physics—Pauli, the man who first proposed its existence, famously lamented that he had proposed “a particle that cannot be detected”—neutrino oscillation experiments require extremely large detectors that are made of complex materials. Yet they must simultaneously work at high precision to obtain results that have any hope of addressing the most pressing questions that remain in the field (like, for example, the ordering of the neutrino masses, and whether neutrinos and their antimatter counterparts, antineutrinos, are exactly mirror images of one another or not). These competing demands necessitate input data of very high quality to ensure that imprecisions are not propagated through to the final measurements.

Chief among the measurements that are assumed by neutrino oscillation experiments is a characterization of the interaction properties of neutrinos on the materials used as targets within detectors. Such is the importance of these neutrino-nucleus cross-sections, in fact, that the development of cross-section measurement programs has recently been thrust into the limelight in neutrino oscillations research. But improving the quality of the cross-section measurements used in oscillation results is only one arm of the dedicated neutrino scattering experiment program. Such experiments are also valuable contributors to the world of nuclear physics as well, where they can provide an electromagnetically neutral probe into the complex physics of particles bound in nuclei. In so doing, they have uncovered several complications to the simple “Fermi gas” model of the nucleus usually taken for granted by oscillation predictions, complementing the picture first advanced by electron scattering experiments.

The MINER ν A experiment at the Fermi National Accelerator Laboratory in Batavia, IL, operates in this intersection of neutrino oscillation and nuclear physics. While the vast majority of results from MINER ν A will focus on the interaction of muon neutrinos (which are the primary content of NuMI’s beam) with the various nuclei in the detector, the electron neutrino cross-section is equally important. Appearance-type oscillation measurements—which usually begin with a nearly pure muon neutrino beam, and measure what fraction of them oscillate to electron neutrinos at a far detector—depend on knowing the ν_e cross-section to make an accurate prediction of their expected signal event rate. Moreover, since accelerator-made muon neutrino beams are not completely pure, but typically contain an impurity of electron neutrinos on the order of 1%, an accurate background rate prediction also depends on knowledge of σ_{ν_e} .

However, there exist only two direct measurements of σ_{ν_e} anywhere near the energy range of interest to oscillation experiments [4, 5]; and these measurements’

applicability to oscillation experiments is challenged by their small statistics (and, in the former case, its markedly different target material, heavy freon). In lieu of using direct measurements, oscillation calculations typically instead work from the muon neutrino cross-section and apply corrections from theoretical arguments to obtain a prediction for the electron neutrino cross-section. The foundational nature of this parameter in the prediction, however, occasions a better measurement. The work presented in this thesis intends to fill that vacuum. Because the intensity of the NuMI beam and the length of MINER ν A’s data taking period were such that the electron neutrino impurity of the beam produced several thousand recorded ν_e events in the detector, there is opportunity to make a significant improvement upon the existing data. We therefore describe herein the first-ever direct measurement of the cross-section of the charged-current quasi-elastic (CCQE) interaction of electron neutrinos on a hydrocarbon target in the few-GeV energy region.

1.1. The Standard Model¹

1.1.1. Fields; gauge theories

The common currency of the Standard Model is the *field*. Though the insights governing the postulates on which the edifice of the SM is built are much deeper, in essence, the SM effectively attempts to answer the questions: “What fields govern nature, how do they evolve in time and space, and how do they interact?” So we begin with the notion of a field: a physical quantity—whether directly observable or not—which is defined everywhere in space and time; we represent it by the notation $\psi(x)$ (or represent one component of it as $\psi^\mu(x)$, or $\psi^{\mu\nu}(x)$, etc. if it is a tensorial quantity with more than a single component), where x is a

¹Throughout this section we will employ natural units, $\hbar = c = 1$, as well as the Einstein summation convention.

four-component Lorentz vector identifying a point in space-time.²

One can construct from a field (or multiple fields) and its (their) derivatives a quantity which encodes its (their) contribution(s) to the energy structure of space-time; this quantity is called the “Lagrangian density,” $\mathcal{L}(\psi, \partial\psi)$. To determine the dynamics of the system modeled by this Lagrangian³, the SM then applies the principle of stationary action (also known as the principle of least action, or Hamilton’s Principle [6]) via the Euler-Lagrange equations, where we will use the shorthand ∂_μ to refer to $\partial/\partial x^\mu$:

$$\partial_\mu \left(\frac{\partial \mathcal{L}}{\partial (\partial_\mu \psi)} \right) - \frac{\partial \mathcal{L}}{\partial \psi} = 0 \quad (1.1)$$

The result is the system’s equations of motion. Historically, Lagrangians have been selected such that their Euler-Lagrange equations reproduce the known equations of motion for specific fundamental forces (for example, electrodynamics). This, in fact, was the starting place for quantum electrodynamics (QED), which, molded by a number of brilliant theorists between roughly 1930 and 1980, eventually grew into the shape of the Standard Model as we know it today. Stripped of its historical trappings, the result for electrodynamics is as follows: writing the electric and magnetic potentials as the four-vector $A = (V, \vec{A})$ and the charge and current sources as $J = (\rho, \vec{J})$ let us consider the Lagrangian

$$\mathcal{L} = -\frac{1}{16\pi}(\partial^\mu A^\nu - \partial^\nu A^\mu)(\partial_\mu A_\nu - \partial_\nu A_\mu) + J^\mu A_\mu \quad (1.2)$$

²As is customary, we will let x^μ refer to the μ th component of this (contravariant) four-vector; one then uses the Minkowski metric $g_{\mu\nu}$ to obtain the covariant four-vector $x_\mu = x^\mu g_{\mu\nu}$ from it.

³It is commonplace in quantum field theory to refer to Lagrangian densities as simply “Lagrangians,” even though strictly speaking the Lagrangian corresponding to the Lagrangian density $\mathcal{L}$ is $\int d^3x \mathcal{L}$. Since in quantum field theory one always works with the Lagrangian density, we will follow this tradition here.

The first term of the Euler-Lagrange equations (eq. 1.1) results in

$$\frac{\partial \mathcal{L}}{\partial(\partial_\mu A_\nu)} = -\frac{1}{4\pi}(\partial^\mu A^\nu - \partial^\nu A^\mu), \quad (1.3)$$

while the second term yields

$$\frac{\partial \mathcal{L}}{\partial A_\mu} = J^\mu, \quad (1.4)$$

Thus the equations of motion are

$$\partial^\mu A^\nu - \partial^\nu A^\mu = 4\pi J^\mu, \quad (1.5)$$

which can be shown to be equivalent to the more canonical form of Maxwell's equations. [2] There are a handful of other well-known sets of equations of motion in quantum field theory; they and the Lagrangians that result in them are listed in table 1.1.2.

It was known long before the advent of quantum mechanics that the laws of electrodynamics are invariant under a so-called *gauge transformation*: an operation whereby the electric and magnetic potentials are modified by the addition of the gradient of a scalar field (that is, $A'_\mu = A_\mu + \partial_\mu \lambda$ for a scalar field $\lambda(x)$). The same principle can be seen (as it must) when viewed through the lens of field theory; notice that

$$\partial^\mu A^{\nu'} - \partial^\nu A^{\mu'} = \partial^\mu (A_\nu + \partial_\nu \lambda) - \partial^\nu (A_\mu + \partial_\mu \lambda) \quad (1.6)$$

$$= \partial^\mu A^\nu + \partial^\mu \partial_\nu \lambda - \partial^\nu A^\mu - \partial^\nu \partial_\mu \lambda \quad (1.7)$$

$$= \partial^\mu A^\nu - \partial^\nu A^\mu, \quad (1.8)$$

leaving eqn. 1.5 invariant. The other fields noted in table 1.1.2 admit similar gauge transformations.

Name	Field	Lagrangian	Equations of motion
Klein-Gordon	Scalar field ϕ	$\mathcal{L} = \frac{1}{2}(\partial_\mu \phi)(\partial^\mu \phi) - \frac{1}{2}m^2 \phi^2$	$\partial_\mu \partial^\mu \phi + m^2 \phi = 0$
Dirac	Spinor field ψ^i	$\mathcal{L} = i\bar{\psi}\gamma^\mu \partial_\mu \psi - m\bar{\psi}\psi$	$i\gamma^\mu \partial_\mu \psi - m\psi = 0$
Proca	Vector field A^μ , with $F^{\mu\nu} = \partial^\mu A^\nu - \partial^\nu A^\mu$	$\mathcal{L} = -\frac{1}{16\pi}F^{\mu\nu}F_{\mu\nu} + \frac{1}{8\pi}m^2 A^\nu A_\nu$	$\partial_\mu F^{\mu\nu} + m^2 A^\nu = 0$

Table 1.1: Well-known Lagrangians in field theory and their resulting equations of motion. What is meant by “spinor” and “vector” fields is discussed in sec. 1.1.2.

1.1.2. Particles in the SM ⁴

So far we have nothing more than a mathematical model of free fields; no mention (aside from the oblique introduction of the parameter m into the Lagrangians in tab. 1.1.2) has been made of anything resembling a particle. Since the Lagrangian density in fact represents an energy density, the quanta of the field a Lagrangian operator represents can (in light of the energy-mass equivalence principle Einstein made famous) very well be interpreted as “particles;” but the problem remains that so far our fields—and therefore particles—do not interact. It was Chen Ning Yang and Robert Mills who finally established a way to rectify this problem in a theoretically consistent way.

We must begin by clarifying some notation that we will use extensively below. It was noted in tab. 1.1.2, for example, that the Dirac Lagrangian corresponds to a “spinor field.” A (Dirac) spinor Ψ , in this case, is a collection of *four* fields Ψ_i which transform under a Lorentz transformation such that the quantity $\bar{\Psi}\Psi = |\Psi_0|^2 + |\Psi_1|^2 - |\Psi_2|^2 - |\Psi_3|^2$ is invariant. In fact, introducing the so-called “Dirac matrices” γ^μ ,

$$\gamma^0 = \begin{bmatrix} \mathbb{1} & 0 \\ 0 & \mathbb{1} \end{bmatrix} \tag{1.9}$$

$$\gamma^i = \begin{bmatrix} 0 & \sigma^i \\ -\sigma^i & 0 \end{bmatrix} \tag{1.10}$$

$$\tag{1.11}$$

(where 0 and $\mathbb{1}$ represent 2×2 additive and multiplicative identity matrices, respectively, and σ^i are the 2×2 Pauli matrices), we can write out the definition

⁴The treatment in this section was largely inspired by ref. [2].

for $\bar{\Psi}$ more straightforwardly:

$$\bar{\Psi} = \Psi^\dagger \gamma^0 \quad (1.12)$$

Thus the Dirac Lagrangian is in actuality an expression which sums over the four components of Ψ (multiplied by the appropriate components of γ^μ and $\partial/\partial x^\mu$).

With those preliminaries in mind, let us consider two Dirac fields, Ψ_1 and Ψ_2 , simultaneously. With the tools we have thus far, all we can do is write down their (non-interacting) sum:

$$\mathcal{L} = i\bar{\Psi}_1 \gamma^\mu \partial_\mu \Psi_1 - m_1 \bar{\Psi}_1 \Psi_1 + i\bar{\Psi}_2 \gamma^\mu \partial_\mu \Psi_2 - m_2 \bar{\Psi}_2 \Psi_2 \quad (1.13)$$

If we perform the same sleight-of-hand with Ψ_1 and Ψ_2 that we have already done with the components of each of the Ψ themselves—that is, bundle them together into an object like a column vector—we can in fact condense this back to something that looks exactly like what we started with:

$$\mathcal{L} = i\bar{\Psi} (\gamma^\mu \partial_\mu \mathbb{1}) \Psi - \bar{\Psi} M \Psi \quad (1.14)$$

if only we take

$$\Psi = \begin{bmatrix} \Psi_1 \\ \Psi_2 \end{bmatrix} \quad (1.15)$$

$$\bar{\Psi} = \begin{bmatrix} \bar{\Psi}_1 & \bar{\Psi}_2 \end{bmatrix} \quad (1.16)$$

and

$$M = \begin{bmatrix} m_1 & 0 \\ 0 & m_2 \end{bmatrix} \quad (1.17)$$

Now, this Lagrangian has the peculiar property that it remains invariant if the field Ψ is transformed by applying some unitary 2×2 matrix U (that is, a U where $U^\dagger U = \mathbb{1}$) to it: that is, if $\Psi \rightarrow U\Psi$ (and thus $\bar{\Psi} \rightarrow \bar{\Psi}U^\dagger$), we get

$$\mathcal{L} = i\bar{\Psi}U^\dagger\gamma^\mu\partial_\mu U\Psi - \bar{\Psi}U^\dagger MU\Psi \quad (1.18)$$

If we choose the masses are the same, such that $M = m\mathbb{1}$, then we can commute the U in the second term past it; furthermore, since all that separates the $U^\dagger$ and the U in the first term is the identity matrix⁵, we can do the same there:

$$\begin{aligned} \mathcal{L} &= i\bar{\Psi}U^\dagger U\gamma^\mu\partial_\mu\Psi - \bar{\Psi}U^\dagger U m\Psi \\ &= i\bar{\Psi}(\gamma^\mu\partial_\mu)\Psi - m\bar{\Psi}\Psi \end{aligned} \quad (1.19)$$

which is the same as eq. 1.14.

It is well-known that an arbitrary 2×2 unitary matrix can be written alternatively: $U = e^{iH}$, for some Hermitian matrix H ($H = H^\dagger$). (Here, the exponentiation of the matrix means that one should expand it in the usual power series: $\mathbb{1} + iU + \frac{i^2}{2!}U^2 + \dots$) And in fact, a Hermitian matrix can be further decomposed: every H could be broken up into $H = \lambda_0\mathbb{1} + \lambda_1\tau_1 + \lambda_2\tau_2 + \lambda_3\tau_3$, where

⁵We are being somewhat elusive with the notation here. Technically, the γ^μ are already 4×4 matrices in their own right, but when we combined the two spinors into a larger one, we implicitly created an 8×8 block-diagonal matrix for the γ^μ with the γ^μ in both upper and lower blocks. Similarly, if we were to explicitly write out all the components of the U matrix, it would actually be an 8×8 block matrix with U_{11} in the upper-left block; etc. We will henceforth dispense with the explicit identity matrix $\mathbb{1}$ in eq. 1.14 with the understanding that it is implied where necessary.

the λ_i are real numbers and the τ_i are the Pauli matrices. Thus, we may rewrite:

$$U = e^{i(\lambda_0 \mathbb{1} + \vec{\lambda} \cdot \vec{\tau})} \quad (1.20)$$

using the dot-product as a convenient shortcut for the sum of the component-wise products. By virtue of eqs. 1.14 and 1.19, then, we have shown that the joint Dirac Lagrangian is invariant under transformations of the form $\Psi \rightarrow e^{i(\lambda_0 \mathbb{1} + \vec{\lambda} \cdot \vec{\tau})} \Psi$. The space of 2×2 unitary matrices corresponds to the group known as $U(2)$ in group theory, so we have established that the joint Lagrangian is invariant under a “global” gauge transformation (that is, one that is independent of the space-time coordinate of the fields) in the group $U(2)$.

But what happens if the transformation is *not* independent of space-time? In other words, what if we make a transformation $\Psi \rightarrow S\Psi = e^{i(\lambda_0(x)\mathbb{1} + \vec{\lambda}(x) \cdot \vec{\tau})} \Psi$? The answer was suggested by work of Hermann Weyl, who laid the groundwork using a single field Ψ ; the major insight of Yang and Mills, and the reason the theory is named after them, was how to handle the non-commutative (non-Abelian) nature of the matrices involved with multiple fields. The basic strategy is to offset the extra term resulting from the differentiation of the exponential by “re-defining” the derivative so that it is cancelled exactly. Specifically, when eq. 1.18 is rewritten with S instead of U , the space-time derivative ∂_μ applies to S as well, now:

$$\begin{aligned} \partial_\mu(S\Psi) &= (\partial_\mu S)\Psi + S\partial_\mu\Psi \\ &= i((\partial_0\lambda_0)\mathbb{1} + (\partial_1\lambda_1)\tau_1 + (\partial_2\lambda_2)\tau_2 + (\partial_3\lambda_3)\tau_3) S\Psi + S\partial_\mu\Psi \end{aligned} \quad (1.21)$$

If we wanted to counteract this action, so that the Lagrangian was invariant under this sort of transformation, we can imagine replacing the derivative with something

that *also* transforms under the transformation; we want

$$\mathcal{D}_\mu \Psi \rightarrow S(\mathcal{D}_\mu \Psi) \quad (1.22)$$

It is not at all straightforward to derive how this can be done, but Yang and Mills showed that it can. The trick is to define this mysterious new derivative—usually called the “covariant derivative”—such that

$$\mathcal{D}_\mu = \partial_\mu - i \left(A_\mu^0 \mathbb{1} + \vec{\tau} \cdot \vec{A}_\mu \right) \quad (1.23)$$

where the four functions A_μ^i (which, because each has four space-time components, remember, means a total of 16 functions...) transform in just the right way to cancel the extra pieces in eq. 1.21. The result, after all the hard work is done, is that our Lagrangian has a new term:

$$\begin{aligned} \mathcal{L} &= i\bar{\Psi}\gamma^\mu \mathcal{D}_\mu \Psi - m\bar{\Psi}\Psi \\ &= i\bar{\Psi}\gamma^\mu \partial_\mu \Psi - m\bar{\Psi}\Psi - \left[\bar{\Psi}\gamma^\mu \Psi A_\mu^0 + (\bar{\Psi}\gamma^\mu \vec{\tau} \Psi) \cdot \vec{A}_\mu \right] \end{aligned} \quad (1.24)$$

and the new fields A_μ^i transform such that

$$\begin{aligned} A_\mu^0 &\rightarrow A_\mu^0 + \partial_\mu \lambda_0 \\ \vec{A}_\mu &\rightarrow \vec{A}_\mu + \partial_\mu \vec{\lambda} + 2(\vec{\lambda} \times \vec{A}_\mu) \end{aligned} \quad (1.25)$$

(to first order in $|\vec{\lambda}|$). The process of exchanging the usual derivatives in this way, and introducing the fields necessary to make it gauge invariant, is often referred to as “minimal coupling.”

Now, since we have introduced some new fields, to be consistent, we ought to include free Lagrangians for them as well. Since each of them is a four-vector field, the Proca Lagrangian (see tab. 1.1.2) is the appropriate one. The only hiccup is

that the Proca Lagrangian includes two terms. The first, that proportional to $F^{\mu\nu}$, is perfectly gauge invariant when we make the substitutions in eq. 1.25, since the mixed partial derivatives of λ are the same. But the term proportional to $A^\nu A_\nu$ would destroy the beautiful gauge invariance we have worked so hard to maintain. The only obvious way out of the quandary—and the one that, with the benefit of hindsight, can be identified as the best one—is to set the mass of the gauge field to 0. To extend this to our set of four vector fields, in addition to the scalar $F^{\mu\nu}$ in tab. 1.1.2, we define

$$\vec{F}^{\mu\nu} = \partial^\mu \vec{A}^\nu - \partial^\nu \vec{A}^\mu - 2(\vec{A}^\mu \times \vec{A}^\nu) \quad (1.26)$$

so that the final Lagrangian, assembling all the pieces, is

$$\mathcal{L} = i\bar{\Psi}\gamma^\mu\partial_\mu\Psi - m\bar{\Psi}\Psi - \left[\bar{\Psi}\gamma^\mu\Psi A_\mu^0 + (\bar{\Psi}\gamma^\mu\vec{\tau}\Psi) \cdot \vec{A}_\mu\right] - \frac{1}{16\pi} \left(F^{\mu\nu}F_{\mu\nu} + \vec{F}^{\mu\nu} \cdot \vec{F}_{\mu\nu}\right) \quad (1.27)$$

Remarkably, we have reproduced something that looks much like electromagnetism simply by starting from two independent spin- $\frac{1}{2}$ Lagrangians and joining them by demanding a local gauge invariance under $U(2)$ transformations. In fact, this Lagrangian explicitly *contains* the Maxwell one; simply return to a single Ψ field, and the necessary vector fields reduce to a *single* field A^μ , requiring invariance under a phase transformation, $\Psi \rightarrow e^{i\lambda(x)}\Psi$. (This theory goes by the name of quantum electrodynamics, QED.) One can equally well do this procedure for quantum chromodynamics (QCD), the theory of the strong force, starting instead from *three* Dirac fields and thus using the symmetry group $SU(3)$.

The weak force, however—the force that underlies neutrino interactions—is much subtler. The difference lies in the mass of the gauge bosons generated by the coupling of the fields. In the Yang-Mills case above, we noted that we were

forced to set the gauge boson mass to zero, in order to respect the symmetry of the Lagrangian under the gauge transformation. (The same is true in QCD.) For QED and QCD, this poses no problems, as the gauge bosons (the photon and the gluon, respectively) appear empirically to be massless. But for the weak force, this is abundantly untrue: the $W^\pm$ (the two charged mediators of the weak force) have masses of roughly 80 times that of the proton, and the Z (their neutral partner) is even heavier than that. To make this machinery work for the weak force, we need one more trick up our sleeve.

So far we have always performed the Yang-Mills local gauge invariance procedure on Dirac fields. But it proved worthwhile to ponder what would happen if one began instead with fields representing particles of different spin: for example, what about *scalar* particles (that is, those having no spin at all)? If we were to follow the Yang-Mills prescription with two copies of the Klein-Gordon Lagrangian from tab. 1.1.2, we would again wind up with four gauge fields (since we would be once again working with $U(2)$). But the expression that results after the introduction of the covariant derivative is [7], instead,

$$\mathcal{L} = (\mathcal{D}_\mu \Phi^\dagger)(\mathcal{D}^\mu \Phi) - \mu^2 \Phi^\dagger \Phi - \lambda(\Phi^\dagger \Phi)^2 \quad (1.28)$$

(where μ and λ simply reparameterize the constants introduced by the minimal coupling procedure).

Now, in all the previous results, the potential part of the Lagrangian (here corresponding to the latter two terms) was minimized for null states; interpreted in quantum field theory, this means that the vacuum state had zero potential energy: $\langle 0 | \Psi^\dagger \Psi | 0 \rangle = 0$. However, *this* potential is instead minimized for $\Phi^\dagger \Phi = v^2/2$; since in the theory there are multiple states which would satisfy this relation, we are left with some freedom in how we wish to write the form of the ground state, and thus the other (excited) states which are expanded around this minimum. An astute

choice is to write the expectation value of the ground state as

$$\langle 0|\Phi|0\rangle = \frac{1}{\sqrt{2}} \begin{bmatrix} 0 \\ v \end{bmatrix} \quad (1.29)$$

with

$$v = \sqrt{-\frac{\mu^2}{\lambda}}; \quad (1.30)$$

then, when expanded around this ground state, we can rewrite the field doublet as

$$\Phi(x) = \frac{1}{\sqrt{2}} \exp\left(\frac{i}{2v} \vec{\xi} \cdot \vec{\tau}\right) \begin{bmatrix} 0 \\ v + H(x) \end{bmatrix} \quad (1.31)$$

(where the three $\xi_i(x)$ and the $H(x)$ fields are real-valued and thus reorganize the four degrees of freedom of the two complex-valued fields of Φ). But, we have invested much effort in rendering these Lagrangians invariant under local gauge transformations, so if we exploit that freedom by choosing a transformation with

$$\Phi \rightarrow e^{-\frac{i}{2v} \vec{\xi} \cdot \vec{\tau}} \Psi \quad (1.32)$$

then the entire exponential in eq. 1.31 cancels out. (This choice of gauge is termed the *unitary* gauge.) From there, we proceed to apply the covariant derivative to Ψ , as before:

$$\mathcal{D}_\mu(x)\Phi(x) = \left[\partial_\mu + A_\mu^0 + \vec{A}_{mu} \cdot \vec{\tau} \right] \Phi(x) \quad (1.33)$$

To obtain a version of the Lagrangian which makes the physical content clear, it

will be necessary to reparameterize the A_μ^i as well:

$$W^\mu = -ig \frac{1}{\sqrt{2}} (A_1^\mu - iA_2^\mu) \quad (1.34)$$

$$A^\mu = -i \frac{1}{\sqrt{2}} \left(g \sin(\theta_W) A_3^\mu + \frac{g'}{2} \cos(\theta_W) A_0^\mu \right) \quad (1.35)$$

$$Z^\mu = -i \frac{1}{\sqrt{2}} \left(g \cos(\theta_W) A_3^\mu - \frac{g'}{2} \sin(\theta_W) A_0^\mu \right) \quad (1.36)$$

with the famous “weak mixing angle” θ_W . (The g and g' are constants that were previously absorbed in the fields; we write them explicitly here because they are known to be related by $\tan \theta_W = g'/g$.) If these are inserted into the covariant derivative in eq. 1.33, and the algebra is worked all the way through, at last, one arrives at the Higgs Lagrangian [7]:

$$\begin{aligned} \mathcal{L} = & \frac{1}{2}(\partial H)^2 - \lambda v^2 H^2 - \frac{g^2 v^2}{4} W_\mu^\dagger W^\mu + \frac{g^2 v^2}{8 \cos^2 \theta_W} Z_\mu Z^\mu \\ & + (\text{higher-order terms in } H \text{ and products of } H \text{ with } W \text{ and } Z) \end{aligned} \quad (1.37)$$

These would be added to the free terms corresponding to $F^{\mu\nu} F_{\mu\nu}$ for each of the gauge fields.

The triumph of this approach is that despite explicitly setting the masses in the Proca field Lagrangian to zero for all the gauge fields, masses still result from combining the Yang-Mills technique with the non-zero vacuum expectation value for the scalar field. In particular, the second term in eq. 1.37 corresponds to a mass term for the H field (now known usually as the Higgs field); the third and fourth correspond to mass terms for the W and Z fields (the vector bosons that mediate the weak force). When we combine the Higgs Lagrangian with what results from the Yang-Mills procedure applied to the fermion fields, the sum prescribes fields for all the fundamental fermions in the standard model, all of the gauge bosons that create interactions between them, and a Higgs field from which the mass of

Particle	electric charge (e)	Spin ($\hbar/2$)	Mass (MeV/c^2)
Higgs boson H	0	0	$\sim 125,000$
Electron e	-1	1/2	~ 0.5
Muon μ	-1	1/2	~ 100
τ	-1	1/2	~ 1750
Electron neutrino ν_e	0	1/2	< 0.00001
Muon neutrino ν_μ	0	1/2	< 0.002
Tau neutrino ν_τ	0	1/2	< 20
u quark	2/3	1/2	< 5
d quark	-1/3	1/2	< 10
c quark	2/3	1/2	~ 1200
s quark	-1/3	1/2	~ 100
t quark	2/3	1/2	$\sim 175,000$
b quark	-1/3	1/2	~ 4250
photon γ	0	1	0
gluon g	0	1	0
$W^\pm$ boson	± 1	1	$\sim 80,000$
Z boson	0	1	$\sim 90,000$

Table 1.2: Particles whose fields are described in the Standard Model. They are further subdivided into quarks (constituents of baryons, including protons and neutrons) and leptons (do not form larger structures); each of the fermions (spin- $1/2$ particles) also has an anti-particle partner with opposite charge but the same spin and mass. Charges are given in terms of the elementary charge e ; spins are terms of the Planck's constant $\hbar$.

the weak vector bosons emerges. This was a truly spectacular accomplishment, and its synthesis resulted in several Nobel prizes. The plethora of particles whose fields the Standard Model describes (along with some of their quantum numbers, which—excluding their masses, on which more in sec. 1.1.3—are consequences of the field structure, even though we did not discuss them here) is shown in tab. 1.2.

1.1.3. Fermion masses in the SM and why the neutrino was massless

To this point we have worried only about the masses of the gauge bosons. Most of the fermions—of which the neutrinos are a subclass—empirically do have nonzero masses, as noted in tab. 1.2; but, as we saw in sec. 1.1.2, the masses (parameter m) that we thought to carry around starting from the Dirac Lagrangian in tab. 1.1.2 cannot be responsible for them. Instead, the Standard Model endows the fermions with mass in the same way it does the vector bosons: through spontaneous symmetry breaking.

In eq. 1.24, we saw that the introduction of the covariant derivative produces extra terms quadratic in the fermion field and linear in the gauge fields. Inspired by this phenomenon, we can find terms similar to them that respect the gauge symmetry, where, instead of the gauge field, we use the Higgs scalar field Φ . Such terms resemble (taking, for the moment, the simplification where we have only field singlets ψ and ϕ):

$$\mathcal{L}_Y = y\bar{\psi}\phi\psi \quad (1.38)$$

where y is a constant characterizing the strength of the interaction between the fermions and the scalar field. When the scalar field ϕ is re-expressed around its vacuum expectation value as in eq. 1.31, this component of the Lagrangian becomes

$$\mathcal{L}_Y = -\frac{1}{\sqrt{2}}y\bar{\psi}(v + H(x))\psi \quad (1.39)$$

and while the second term is still an interaction term with the H field, the first term—the one proportional to v —is now a mass term for the field Ψ .

As expressed here, any fermion field could be used for ψ , so there is no reason

to expect that the neutrinos would be any different from the other, massive leptons. However, there is one further complication in the Standard Model that rendered it attractive to conceptualize the neutrinos, particularly, as having no mass: and it has to do with the nature of the chirality of the weak force.

In general, Dirac fields can be decomposed into two (orthogonal) components using the chirality projection operators:

$$P_{R,L} = \frac{\mathbb{1} \pm \gamma^5}{2} \quad (1.40)$$

with

$$\gamma^5 = i\gamma^0\gamma^1\gamma^2\gamma^3 \quad (1.41)$$

so that

$$\begin{aligned} \psi_L &= P_L\psi = \frac{\mathbb{1} - \gamma^5}{2}\psi \\ \psi_R &= P_R\psi = \frac{\mathbb{1} + \gamma^5}{2}\psi \end{aligned} \quad (1.42)$$

Now, empirically, the charged-current interactions of the weak force are always left-handed, and the gauge theory used to describe it must reflect that. The way this requirement is implemented in the Standard Model is by decomposing the $U(2)$ gauge group into $SU(2)$ (the space of all unitary 2×2 matrices whose determinant is 1) and $U(1)$ (rotations by a phase angle), and requiring that the $SU(2)$ piece—the interactions resulting from the vector bosons in the gauge theory—act only on the *left-handed* component of the fermion fields. The right-handed projection, therefore, is unaffected by the $SU(2)$ component altogether; it is thus frequently said that ψ_R is a *singlet* under these $SU(2)$ operations. To denote this situation, the gauge group for weak interactions is usually written as $SU(2)_L \times U(1)$.

While it is irrelevant for those fermions whose masses can be measured more or less directly, an interesting theoretical point can be raised about these chiral decompositions. Notice, first of all, that since γ^5 is symmetric and real, $(\gamma^5)^\dagger = \gamma^5$, and so $P_{L,R}^\dagger = P_{L,R}$. Then, because γ^0 anticommutes with γ^5 :

$$\overline{\psi_R} = \overline{(P_R\psi)} = (P_R\psi)^\dagger \gamma^0 = \psi^\dagger \gamma^0 P_L = \overline{\psi} P_L \quad (1.43)$$

and in the same way, $\overline{\psi_L} = \overline{\psi} P_R$. Furthermore, written in terms of the chiral decomposition, the Dirac Lagrangian in tab. 1.1.2 becomes

$$\mathcal{L}_{\text{chiral}} = i\overline{\psi_R}\partial_\mu\gamma^\mu\psi_R + i\overline{\psi_L}\partial_\mu\gamma^\mu\psi_L - m(\overline{\psi_R}\psi_L + \overline{\psi_L}\psi_R) \quad (1.44)$$

$$= i\overline{\psi_R}\partial_\mu\gamma^\mu\psi_R + i\overline{\psi_L}\partial_\mu\gamma^\mu\psi_L \quad (1.45)$$

if the field is massless. That is to say: for a massless field, the chiral components of a state completely decouple in the Lagrangian, and therefore, their time-evolution is not related. In other words, a massless field can exist indefinitely in a purely chiral state, whereas a massive field cannot. Even more suggestively, for a massless particle, its eigenvalue under the chirality operator P_R is identical to its helicity $\vec{\sigma} \cdot \vec{p}/|\vec{p}|$: that is, helicity and chirality coincide. Taking the fact that neutrinos that interact via the charged current experimentally always have left-handed helicity, and combining it with the fact that prior to the resolution of the solar neutrino conundrum (on which see below) there was no convincing evidence that neutrinos did have mass, many physicists found the idea of a massless neutrino extremely compelling.

1.2. Beyond the Standard Model: neutrino oscillations

1.2.1. Massive neutrinos

The conception of the neutrino as a massless particle made an extremely convenient fit into the mathematical picture offered by the Standard Model, and was ultimately only challenged—in the way that seems to be customary for science—by an experiment that was seeking to establish something else. In the 1960s, models of the mechanism underlying the Sun’s enormous power output were converging on a complicated fusion cycle, which was predicted to produce a copious flux of neutrinos. Confirmation of this “solar neutrino” hypothesis was expected to provide important evidence in the race to understand the Sun’s core. So, in what has become a seminal experiment for the history of neutrino physics, Ray Davis and collaborators constructed a detector which would capture the by-products of neutrinos interacting with tetrachloroethylene (C_2Cl_4) via the reaction $\nu + {}^{37}Cl \rightarrow {}^{37}Ar + e^-$. (The resultant radioactive argon atoms were flushed out of the enormous tank of tetrachloroethylene using a helium gas current forced through the volume, and their decays observed in a proportional counter.) Their measurement, that the rate of argon atoms they collected was roughly a third of what was expected based on the best predictions [8], gave rise to what has usually been termed the “solar neutrino problem.”

Many subsequent experiments further confirmed the deficit observed by Davis, eventually precluding the dismissal of the result as experimental error. Parallel refinements of the calculations used to produce the predicted rates also ruled out egregious errors in the model. Eventually, the clever idea proposed by Bruno Pontecorvo in 1957 [9]—positing that there could be multiple “flavors” of neutrinos which could “oscillate” back and forth into one another (based on the previously

known model of $K^0 \leftrightarrow \overline{K}^0$ oscillations)—rose to the forefront of the discussion. Crucially, though, such a mechanism required at least one of the neutrino flavors to have a nonzero mass – medicine that many theorists originally found difficult to swallow.

Decades of successive experiments attempting to substantiate or disprove the theory of neutrino oscillations finally culminated in the definitive measurement by the Subary Neutrino Observatory (SNO) collaboration that the neutrinos detected by Davis make up only about 35% of the total flux of neutrinos originating from the Sun. The rest, they found, were consistent with having undergone precisely the transition predicted by the neutrino oscillation model. [10] This was seen as the convincing proof that at least two of the three flavors that are now known (known by their charged lepton partners, electron, muon, and tau) are not massless. A number of other experiments have since confirmed this conclusion, measuring the disappearance (or new appearance) of neutrinos creating by solar, cosmic-ray, and terrestrial (accelerator) sources subject to propagation across various distances.

1.2.2. Formalism

The mathematical formulation of neutrino oscillations sees the three neutrinos that were previously known by the lepton they produce in a charged-current reaction— ν_e , ν_μ , and ν_τ —as quantum-mechanical states that do not coincide exactly with the basis of states corresponding to their masses. That is, using α to run over the flavors, and i to denote the various masses,

$$|\nu_\alpha\rangle = \sum_{i=1}^3 U_{\alpha i}^* |\nu_i\rangle \tag{1.46}$$

where $U_{\alpha i}^*$ are the entries in a unitary matrix that is not diagonal.⁶ (We assume here that both the flavor states $|\nu_\alpha\rangle$ and the mass states $|\nu_i\rangle$ are unit normalized, and that each is an orthogonal basis set; these assumptions can be shown to be reasonable, though we will not do so here. [7]) Because in quantum mechanics the stationary states of a system correspond to states of fixed energy, the stationary neutrino states are those of definite mass (which, by Einstein's equation $E^2 = \vec{p}^2 + m^2$, means that at zero momentum, energy and mass are interchangeable). Such states can be expanded as plane wave solutions to the time-dependent Schrodinger equation:

$$|\nu_i(t)\rangle = e^{-i(p_i \cdot x_i)} |\nu_i(t=0)\rangle \quad (1.47)$$

(with p and x the customary momentum and position four-vectors).

Now, since neutrinos' masses are very small (whatever their precise values may be), they propagate in the ultrarelativistic regime ($E^2 \gg m^2$), where Einstein's equation can be simplified via Taylor expansion:

$$|\vec{p}| = \sqrt{E^2 - m^2} \quad (1.48)$$

$$= E \sqrt{1 - \frac{m^2}{E^2}} \quad (1.49)$$

$$\approx E \left(1 - \frac{m^2}{2E^2} \right) \quad (1.50)$$

Then, assuming we align our coordinate system so that the plane wave is propagating exclusively along the z -direction (so that $\vec{p} \cdot \vec{x} = p_z z$), and note that for an

⁶It would be equally sensible to use this matrix's Hermitian conjugate here; traditionally, however, the un-starred version is reserved for the matrix that decomposes the mass states in terms of the flavor ones. To avoid confusion, we will retain that tradition here.

ultrarelativistic particle traveling along the z -axis, $t \approx z$, we find

$$e^{-i(p_i \cdot x_i)} = e^{-i(E_i t - \vec{p} \cdot \vec{x})} \quad (1.51)$$

$$= e^{-i(E_i - p_z)z} \quad (1.52)$$

$$\approx e^{-i\left(E_i - E_i + \frac{m^2}{2E_i}\right)z} \quad (1.53)$$

$$= e^{-i\left(\frac{m^2}{2E_i}z\right)} \quad (1.54)$$

so that the propagation of a neutrino mass eigenstate obeys

$$|\nu_i(z)\rangle = e^{-im_i^2 z/2E_i} |\nu_i(0)\rangle \quad (1.55)$$

The probability of observing a neutrino of flavor β after a neutrino initially in state α has propagated a distance z is then, according to quantum mechanics,

$$P_{\alpha \rightarrow \beta} = |\langle \nu_\beta | \nu_\alpha(z) \rangle|^2 \quad (1.56)$$

If we expand both states in terms of the mass eigenstates per eq. 1.46, and then apply eq. 1.55, we get

$$P_{\alpha \rightarrow \beta}(z) = \left| \left(\sum_i U_{\beta i} \langle \nu_i | \right) \left(\sum_j U_{\alpha j}^* |\nu_j(z)\rangle \right) \right|^2 \quad (1.57)$$

$$= \left| \left(\sum_i U_{\beta i} \langle \nu_i | \right) \left(\sum_j U_{\alpha j}^* e^{-im_j^2 z/2E} |\nu_j\rangle \right) \right|^2 \quad (1.58)$$

$$= \left| \sum_i U_{\alpha i}^* U_{\beta i} e^{-im_i^2 z/2E} \right|^2 \quad (1.59)$$

where we have dropped the subscript on the energy since the flavor state's total energy is fixed. It is traditional to work this expression out somewhat further (and

also to use L instead of z); after some algebra, one can show [7]

$$\begin{aligned}
P_{\alpha \rightarrow \beta} = & \delta_{\alpha\beta} - 4 \sum_{k>j} \Re(U_{\alpha k}^* U_{\beta k} U_{\alpha j} U_{\beta j}^*) \sin^2\left(\frac{\Delta m_{kj}^2 L}{4E}\right) \\
& + 2 \sum_{k>j} \Im(U_{\alpha k}^* U_{\beta k} U_{\alpha j} U_{\beta j}^*) \sin\left(\frac{\Delta m_{kj}^2 L}{2E}\right)
\end{aligned} \tag{1.60}$$

(with $\delta_{\alpha\beta}$ being the Kronecker delta, and $\Delta_{kj}^2 = m_k^2 - m_j^2$). That is: oscillation occurs if and only if at least one pair of states has nonzero Δm^2 between them, which requires at least one of the neutrinos to be massive.

So far we have said nothing about the oscillation parameters $U_{\alpha i}$. This matrix should be unitary by virtue of the orthogonality and completeness of the sets of states $|\nu_\alpha\rangle$ and $|\nu_i\rangle$.⁷ It is known that the N^2 degrees of freedom of a $N \times N$ unitary matrix can be decomposed into real parameters by a suitable parameterization; for a 3×3 matrix, this yields three “mixing angles” and six phase factors. Arguments based on the invariance of the weak current under phase transformations reduce this total for the physical case of neutrinos to three mixing angles and one phase factor. [7] Traditionally, the mixing angles are denoted θ_{ij} (corresponding to mixing between $|\nu_i\rangle$ and $|\nu_j\rangle$) and the phase factor as δ (or δ_{CP} , since it is related to the breaking of the symmetry of the Lagrangian under charge-parity operations). If we further introduce the notation $c_{ij} = \cos(\theta_{ij})$ and $s_{ij} = \sin(\theta_{ij})$, then the matrix can be written out as [11]:

$$U = \begin{bmatrix} c_{12}c_{13} & s_{12}c_{13} & s_{13}e^{-i\delta} \\ -s_{12}c_{23} - c_{12}s_{23}s_{13}e^{i\delta} & c_{12}c_{23} - s_{12}s_{23}s_{13}e^{i\delta} & s_{23}c_{13} \\ s_{12}s_{23} - c_{12}c_{23}s_{13}e^{i\delta} & -c_{12}s_{23} - s_{12}c_{23}s_{13}e^{i\delta} & c_{23}c_{13} \end{bmatrix} \tag{1.61}$$

⁷The extent to which this is exactly true is directly related to the search for so-called “sterile neutrinos,” about which we will remain silent for the remainder of this document. Further reading can be obtained in, e.g., ref. [11].

This matrix is called the “PMNS” matrix (after Bruno Pontecorvo, Ziro Maki, Masami Nakagawa, and Shoichi Sakata).

1.3. Neutrino interactions and scattering

1.3.1. Oscillations to cross-sections

In the decade-and-a-half since the SNO result, the focus of neutrino physics research has drifted from the question of *whether* neutrinos have mass to the question of *how much*, and to a certain extent, *why*. Fully determining the parameters that govern neutrino oscillation would constitute significant progress in this campaign. Among the most tantalizing prospects is the measurement of the parameter δ_{CP} discussed above, a nonzero value of which would indicate that neutrinos and their antiparticles undergo oscillation differently. (Such a discovery would have, improbably, significant bearing on the fundamental question of why our universe appears to be filled with matter, rather than equivalent amounts of matter and anti-matter, as our original models of cosmogenesis predicted.) Other lingering fundamental questions include the *signs* of the various Δm^2 (notice that the leading term in oscillation probability depends only on the *square* of the sine of these parameters) and the absolute values of the masses themselves.

The effort to measure the parameters of neutrino oscillation catapults us squarely into the realm of precision measurement. For instance, the measurement of δ_{CP} in experiments using accelerators (where beams of muon neutrinos are generated, and the observation is of their oscillation into electron neutrinos) relies on the difference in sign between neutrinos and antineutrinos in the following term in the oscillation probability [12]:

$$\frac{\sin 2\theta_{12} \sin 2\theta_{23}}{2 \sin \theta_{13}} \sin \frac{\Delta m_{21}^2 L}{4E} \sin^2 2\theta_{13} \sin^2 \frac{\Delta m_{31}^2 L}{4E} \sin \delta_{CP} \quad (1.62)$$

The size of this quantity is strongly influenced by the factors involving θ_{13} ; since θ_{13} is small (roughly 9°), this term is small. Furthermore, as it is at most about 25% of the leading term in the oscillation probability, one must search for a small variation on a larger trend.

Searching for small features in a potentially noisy dataset requires sufficient statistical power to resolve them. And, superimposed on the general problem of the weakness of neutrino interactions in the first place, the desire for precision neutrino physics results has driven the experimental community to build large detectors from massive nuclei, in order to maximize the neutrino target. But here increased statistical power comes at a heavy price: while we are able to predict the interaction properties of neutrinos with single nucleons with reasonable success, the added complications of the nuclear environment introduce significant uncertainties into any result that requires a model of the rate or outgoing kinematics of particles in a neutrino reaction. As neutrino oscillation campaigns typically reverse-engineer the oscillation probability by comparing a predicted event rate with what is actually observed in their detectors, oscillation experiments cannot escape this curse.

The mathematical framework used to quantify interaction probabilities in the Standard Model is that of the *cross-section*. From the relevant particle flux Φ , one can then compute the event rate as a function of any parameter of interest (call it ξ , so that the event rate is $dN/d\xi$), so long as the cross-section is available as a function of that parameter ($d\sigma/d\xi$):⁸

$$\frac{dN}{d\xi} = \Phi \frac{d\sigma}{d\xi} \tag{1.63}$$

⁸Once again, we must follow the traditional notation here, even though the objects which appear to be derivatives in this equation are strictly not. The differential notation is typically used because often we use a cross-section divided into bins in some parameter which is not an intrinsic property of the interacting particle(s)—say, for instance, the energy of an *outgoing* particle—and we can obtain the total event count by integrating eq. 1.63 across that parameter so long as the “differential” cross-section is correct.

So, then, the success of a neutrino oscillation measurement—where a precise prediction of $N(E_\nu) = \Phi(E_\nu)\sigma(E_\nu)$ is crucial to the determination of the result—is vitally tied to the quality of its cross-section model.

1.3.2. Cross-sections in the SM

The Standard Model prescribes a specific set of rules which can be used to compute the cross-section for any process where a Feynman diagram can be drawn of it using only elementary quarks, leptons, and gauge bosons. In fact, the formula known as Fermi’s Golden Rule specifies that for a scattering process where non-identical particles numbered 1 (with four-momenta p_1 and mass m_1) and 2 (etc.) interact to yield non-identical particles numbered 3 through n [2],⁹

$$\sigma = \frac{1}{4\sqrt{(p_1 \cdot p_2)^2 - (m_1 m_2)^2}} \int |\mathcal{M}|^2 (2\pi)^4 \delta^4(p_1 + p_2 - (p_3 + p_4 + \dots + p_n)) \times \prod_{j=3}^n \frac{1}{2\sqrt{\vec{p}_j^2 + m_j^2}} \frac{d^3 \vec{p}_j}{(2\pi)^3}; \quad (1.64)$$

here $\mathcal{M}$ is what is known as the *matrix element* or *amplitude* for the process in question, and it is calculated using Feynman rules (derived ultimately from the Lagrangian of the governing fundamental force) applied to the Feynman diagram of the process.

One of the simplest weak-force processes for which the Standard Model cross-section can be computed exactly¹⁰ is what is called “inverse muon decay,” in which a muon neutrino interacts with an atomic electron: $\nu_\mu + e^- \rightarrow \mu^- + \nu_e$. This process can be represented by the diagram shown in fig. 1.1.

One can deduce from the Lagrangian a set of rules that correspond to the lines in the diagram representing incoming and outgoing and internally exchanged

⁹Identical particles in either the initial or final state reduce the result by a factor of $1/s!$ for each set of s identical particles.

¹⁰Computed exactly to first order in perturbation theory, anyway.

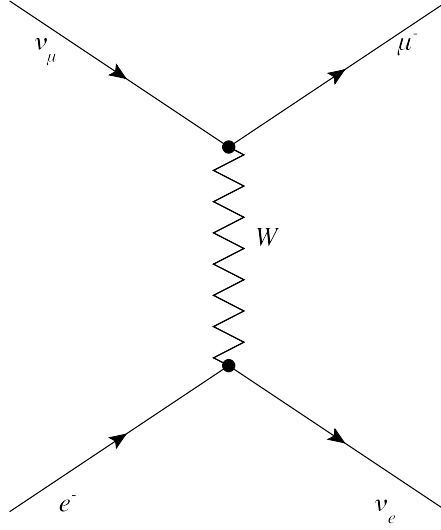


Figure 1.1: Feynman diagram for inverse muon decay. Reading from left to right, this diagram indicates a muon neutrino (top left) and an electron (bottom left) interacting by exchanging a W boson (center) to produce a muon (top right) and an electron neutrino (bottom right).

particles as well as the vertices that join them. [1] The amplitude for the inverse muon decay process, when the momentum transferred to the electron is much less than the W mass (about 80 GeV), works out to be [2]:

$$\mathcal{M} = \frac{g_w^2}{8M_W^2} [\bar{u}(3)\gamma^\mu(1 - \gamma^5)u(1)][\bar{u}(4)\gamma_\mu(1 - \gamma^5)u(2)] \quad (1.65)$$

where in this expression g_w is the weak coupling constant, M_W is the mass of the W boson, and the $u(i)$ are particular Dirac wavefunction states that solve the Dirac equation.

It is known that this expression can be simplified substantially if we wish to average over the polarization of the initial state fermions. Since, for our purposes, we do not care what orientation the incoming or outgoing particles had—we only want the cross-section—we can use a technique attributed [2] to Casimir, followed by the application of some theorems involving the traces of the resulting matrices, to obtain the spin-averaged amplitude squared $\langle |\mathcal{M}|^2 \rangle$. This will be inserted into Fermi's Golden Rule to obtain the differential cross-section vs. the solid angle Ω .

The result is:

$$\langle |\mathcal{M}|^2 \rangle = 8 \left(\frac{g_w E}{M_W} \right)^4 \left[1 - \left(\frac{m_\mu}{2E} \right)^2 \right] \quad (1.66)$$

$$\frac{d\sigma}{d\Omega} = \frac{1}{2} \left(\frac{g_w^2 E}{4\pi M_W^2} \right)^2 \left[1 - \left(\frac{m_\mu}{2E} \right)^2 \right]^2 \quad (1.67)$$

Here, the kinematics (energy E) are given in the center-of-momentum reference frame, where the neutrino and electron have equal and opposite momenta before the collision: to compare them to those measured in an experiment, they would need to be transformed back into the laboratory coordinate system (where the electron is essentially at rest and the neutrino carries all the momentum). But, regardless, the cross-section is a function only of the initial-state kinematics and constants that are either known from the Standard Model or can be measured experimentally. Thus, the cross-section can be predicted with very good precision.

The situation becomes much more difficult when hadrons—particles composed of bound states of quarks—are involved. In particular, this difficulty extends to the nucleons that comprise nearly all of the mass in ordinary matter. Though QCD does provide a framework describing the interactions between quarks and the gluons that mediate the strong force, the perturbative mechanism it uses for calculations does not apply to the especially strong binding that exists between quarks inside hadrons. For instance, the charged-current quasi-elastic (CCQE) process we will consider in the next section is the analog of inverse muon decay where a neutrino interacts with a nucleon instead of an electron. In that process, however, we cannot write down a closed-form expression for the amplitude like was done in 1.65. Instead, as we will see shortly, we will be forced to make do with parameterizing our ignorance in terms of *form factors*, replacing the second factor in eq. 1.65 with one in which the vector-minus-axial-vector tensorial structure $\gamma_\mu(1 - \gamma^5)$ is exchanged for a parameterization in terms of every possible type of

tensor [7]:

$$\mathcal{M} \propto \bar{u}_p(p_p) \left[\gamma^\mu F_1(Q^2) + \frac{i}{2m_N} \sigma^{\mu\nu} q_\nu F_2(Q^2) - \gamma^\mu \gamma^5 F_A(Q^2) - \frac{q^\mu}{m_N} \gamma^5 F_P(Q^2) \right] u_n(p_n) \quad (1.68)$$

These various F are the form factors, and will play heavily into the exposition below.

1.3.3. CCQE

Around 1 GeV, the most common charged-current interaction that neutrinos experience when interacting with ordinary matter is charged-current quasi-elastic (CCQE) scattering: the neutrino exchanges a (charged) W boson with a nucleon, converting the neutrino into a charged lepton, and the nucleon into its isospin partner. For appearance-type neutrino oscillation experiments, this process is typically a large fraction of the predicted event rate, and it is thus essential to have an accurate characterization of its cross-section. In those experiments, usually the initial neutrino beam is of muon neutrinos, which oscillate into electron neutrinos; at the far end, where the electron neutrino interacts, the process may be written as follows:

$$\begin{aligned} \nu_e + n &\rightarrow e^- + p^+ \\ \bar{\nu}_e + p^+ &\rightarrow e^+ + n \end{aligned} \quad (1.69)$$

This reaction is illustrated in the Feynman diagram in fig. 1.2.

The formalism for the derivation of the cross-section explored below works equally well for any flavor, even though the outgoing lepton has different mass depending on its flavor. We will apply it in general, and note where the lepton's mass becomes a relevant quantity.

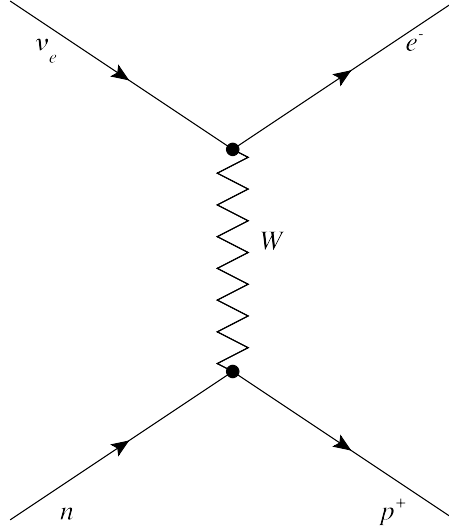


Figure 1.2: Feynman diagram for ν_e CCQE scattering. (The antineutrino version would exchange $\nu_e \rightarrow \bar{\nu}_e$, $e^- \rightarrow e^+$, $n \leftrightarrow p^+$.)

CCQE with free nucleons

V – A formula The full calculation for a differential cross-section of the CCQE process originates from the late 1960s. C. H. Llewellyn Smith’s review [13] is the derivation most commonly followed in the literature; it is written in terms of the Lorentz invariants called Mandelstam variables: $Q^2 = -q^\mu q_\mu$, where q represents the four-momentum transferred from the neutrino to the nucleon it interacts with; $s = (p_\nu + p_N)^2$; and $u = (p_l - p_N)^2$. (Sometimes the Mandelstam variable $t = -Q^2$ is also used.) The lengthy computation results in the cross-section (after being recast in somewhat more modern notation [7]):

$$\frac{d\sigma}{dQ^2} = \frac{M^2 g_W^2 \cos^2 \theta_c}{8\pi E_\nu^2} \left(A(Q^2) \pm B(Q^2) \frac{s-u}{M^2} + C(Q^2) \frac{(s-u)^2}{M^4} \right) \quad (1.70)$$

where the A , B , and C are functions of various form factors of various tensorial natures (F_1 and F_2 are vector; F_P is pseudo-scalar; and F_A is axial vector):

$$\begin{aligned}
A &= \frac{m^2 + Q^2}{M^2} \left[(1 + \tau)F_A^2 - (1 - \tau)(F_1^2 - \tau F_2^2) + 4\tau F_1 F_2 \right. \\
&\quad \left. - \frac{m^2}{4M^2} \left((F_1 + F_2)^2 + (F_A + 2F_P)^2 - \frac{1}{4}(1 + \tau)F_P^2 \right) \right] \\
B &= 4\tau F_A(F_1 + F_2) \\
C &= \frac{1}{4}(F_A^2 + F_1^2 + \tau F_2^2)
\end{aligned} \tag{1.71}$$

In these expressions, m is the lepton mass; M is the initial nucleon mass; and $\tau = \frac{Q^2}{4M^2}$. (The $\pm$ in the cross-section is taken to mean that it is positive for the neutrino-incident process, and negative for the antineutrino one.) It is worth observing here that for electron and muon neutrino scattering, the fraction $\frac{m^2}{M^2}$ is small, and so any terms proportional to it are negligible; this simplifies the expression for A , in particular, enormously. Finally, it should also be mentioned that originally, Llewellyn Smith considered so-called “second-class” currents, of higher-order in the Lagrangian; on theoretical grounds these are commonly neglected (and, in any case, have never been observed). The fixing of the remaining form factors is discussed in the following sections.

Vector form factors; CVC The vector form factors F_1 and F_2 are inherently electromagnetic, as one might expect from the vector nature of the electromagnetic interaction. At $Q^2 = 0$, in fact, they reduce to (respectively) the difference between the proton and neutron electric charges and the difference between their anomalous magnetic moments [7]. But to move away from $Q^2 = 0$ —since, of course, in CCQE, some four-momentum must always be exchanged with the nucleon, if for no other reason than that the nucleon changes mass—we must turn to some kind of scattering measurement.

The same argument that results in the simplifications at $Q^2 = 0$ establishes a relationship between F_1 and F_2 and another set of form factors used in the scattering of *charged* leptons from nucleons; the fundamental relationship is called the Conserved Vector Current (CVC) hypothesis, which is a consequence of the weak isospin invariance of the QCD Lagrangian [7]. Those form factors can be straightforwardly measured in electron scattering experiments as functions of Q^2 (since the initial and final state kinematics are available); the data used in neutrino models (the BBBA05 form factors [14]) is usually parameterized in the form used by Kelly [15]:

$$G(Q^2) \propto \frac{\sum_{k=0}^n a_k \tau^k}{1 + \sum_{k=1}^{n+2} b_k \tau^k} \quad (1.72)$$

Notice that when Q^2 (and therefore τ) is small, the Kelly form factor reduces to what is known as “dipole” form:

$$G(Q^2) \xrightarrow{Q^2 \ll M^2} \frac{1}{1 + \tau} \quad (1.73)$$

This form will also serve usefully in some of the other form factors.

Pseudoscalar form factor; PCAC Another conservation hypothesis allows the relation of the pseudoscalar form factor to the axial one (discussed shortly), again at $Q^2 = 0$. In this case, one employs the Partially-Conserved Axial Current (PCAC) hypothesis, which supposes that the Klein-Gordon pion field is related to the axial part of the hadronic current from weak hadronic interactions. This provides for the expression of the hadronic F_P in terms the F_A and an experimental constant:

$$F_P = \frac{2M^2 F_A(Q^2)}{Q^2 + m_\pi^2} \quad (1.74)$$

(in this case using m_π for the charged pion mass).

Axial vector form factor The axial form factor F_A is the only parameter in the cross-section that cannot be satisfactorily fixed using data from non-neutrino experiments. The current best estimates of this form factor, in fact, arise from fits to ν_μ CCQE data sets. [16] For these fits, the most common parameterization of F_A is again a dipole:

$$F_A(Q^2) = \frac{g_A}{(1 + \frac{Q^2}{M_A^2})^2} \quad (1.75)$$

with $g_A \sim 1.26$ measured in free neutron decay [17]. Though values fit for the “axial mass” M_A have varied somewhat [18, 19], MINER ν A’s own measurements favor $M_A = 0.99$ [20, 21], and we will employ that value for the analysis in this work.

Nuclear/initial state effects

The formalism above was all derived assuming that the target nucleon is free (i.e., does not experience any potential) and at rest. Given that modern neutrino detectors are nearly all constructed from materials made of nuclei more complex than hydrogen, these are questionable assumptions. The effect of the initial state nucleon momentum is particularly notable.

Because the nucleus is a bound state of nucleons, the nucleons within it exist in states of negative potential energy. However, since nucleons are all fermions, they obey the Pauli exclusion principle, which requires that no more than a single particle in a system can occupy a state with a given set of quantum numbers. This means that, for example, no two protons can occupy the same bound state energy level. Thus, when a neutrino interacts with a nucleon, if the energy transferred to the nucleon is not enough to eject it from the nucleus altogether, the nucleon

would be left co-existing in an energy level with the nucleon already occupying it. Because this is prohibited, the interaction is suppressed. As a result, neutrino interactions which transfer less than the energy required to elevate a nucleon above the so-called Fermi energy (maximum energy within the nucleus) are suppressed, and a minimum energy transfer threshold for the process is observed.

Models for how these energy levels within the nucleus are distributed vary widely. The simplest—and the one currently used as a basis in most neutrino interaction generators—is the Relativistic Fermi Gas (RFG), first proposed by Smith and Moniz [22]. In this model, nucleons are seen as non-interacting, and confined in a simple gas within the nucleus. In the initial calculation, as well as in early implementations, the threshold effect due to the Pauli blocking mentioned above was enforced by setting the cross-section for momentum transfers of less than the Fermi energy to 0. (An average binding energy E_b is also typically subtracted from the prediction of the final-state energy to account for the energy lost in ejecting the nucleon from its place in the potential well.) Because there is good evidence that pairs of nucleons often form stronger bound states within the nucleus, which can allow them to have opposite momenta of magnitude larger than the Fermi energy allows, more modern implementations, like the GENIE model discussed in sec. 5.2.1, often modify the step function somewhat. In these models, the struck nucleon has a decaying but nonzero probability of having more momentum than that corresponding to the Fermi energy (a “Bodek-Ritchie tail,” after the parameterization in ref. [23]).

Other models attempt to further improve upon the characterization of the interaction of nucleons within the nucleus. Spectral function models, for example, replace the step function of the RFG with a probability distribution derived from a nuclear shell model [24, 25]. Meson exchange current (MEC) models attempt to remove the non-interacting nature of the other models altogether by calculating

amplitudes for the exchange of mesons (typically pions) between baryons in the nucleus [26, 27]; some of these models can even predict the kinematics of the simultaneous ejection of multiple nucleons that existed in correlated states prior to the neutrino interaction [28]. (The amplitudes of MEC processes’ final states can interfere with those of the traditional CCQE, and therefore affect the final observable distributions.) There are also empirical attempts based on electron scattering data to parameterize the effect of nuclear initial state effects on the final cross-section [29].

FSI

The nuclear medium has another significant consequence for neutrino interactions: any final-state particles composed of quarks can interact via the strong force with the remnant nucleus as they exit. These interactions are often lumped together under the label “final-state interactions,” or FSI. Final-state interactions can severely influence the energy and angle of particles as they traverse the nucleus, smearing out estimators of those quantities; or, worse, particles can be absorbed, or new ones created. In a search for CCQE events, the most serious of these effects is the potential for the absorption of pions that were created by the decay of a nucleon resonance (excited by the neutrino interaction): for instance, $\nu_\mu + n \rightarrow \mu^- \Delta^+ \rightarrow \mu^- p^+ \pi^0$. If the pion in this reaction is absorbed by the nucleus before it can be observed, the final-state particles are identical to those from the CCQE interaction of eq. 1.69.

Modelers take several different approaches to FSI. The most popular is to adopt a cascade-type model like those used for hadron formation in higher-energy situations (like particle colliders). These step final-state particles through the nuclear medium subject to interaction probabilities from various models that run as functions of the path traveled. Simulations used by theorists [30] as well as

neutrino experiments [31] employ this approach. Another method is the more phenomenological one used in GENIE [32] (discussed further in sec. 5.2.1), in which dedicated hadron-nucleus scattering data is used to predict the kinematics and multiplicities of hadrons exiting the nucleus.

1.3.4. Previous attempts to measure σ_{ν_e}

The prediction of the Standard Model that ν_μ and ν_e interactions should follow the same behavior apart from the effects of the different final-state lepton masses—known generally as “lepton universality”—has been extensively investigated. (The impact of the differing masses is extensively studied in ref. [33].) In addition to verification using charged lepton probes, experiments using wide-band neutrino beams, with energies typically ranging from 0 to several hundred GeV, made a number of measurements of the ratios of the interaction rates or cross-sections of muon and electron neutrinos primarily in the 1970s and 1980s. [34, 35, 36, 37, 5, 38, 39] Of these experiments, only the Gargamelle collaboration [5, 38] published any cross-sections for the ν_e interactions they observed; they produced a total charged-current cross-section σ_{ν_e} on heavy freon using neutrinos of energies between 0 and 10 GeV. No significant deviations from the universality predicted by the Standard Model were found in any of these results.

More recently, the T2K collaboration also measured cross-sections for inclusive charged-current electron neutrino interactions, this time on a carbon target. [4] T2K’s result is for much lower energy neutrinos (roughly 0-3 GeV), but made use of a much more intense beam, such that they obtained sufficient statistics to also produce cross-sections differential in the electron angle and energy, and the reconstructed Q^2 of the interaction, in addition to a total cross-section. Again, their results were consistent with the predictions of lepton universality.

2 Creating a high-intensity neutrino beam

The Fermi National Accelerator Laboratory’s NuMI (“Neutrinos at the Main Injector”) high-intensity neutrino beam provided the neutrinos used for the measurement described in this work. The creation of a high-intensity neutrino beam is a complex process involving a number of stages, all of which are described herein. The various components described below (as well as a few others used in other experiments) are depicted in figure 2.1.

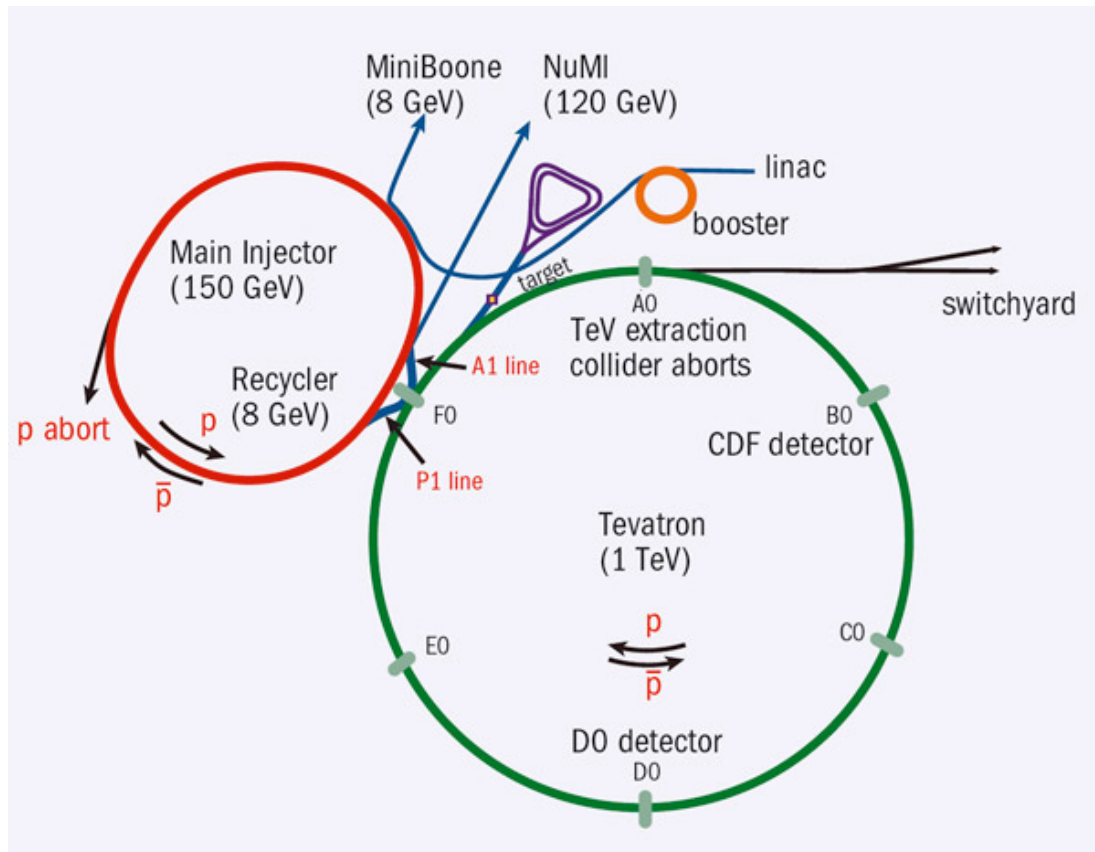


Figure 2.1: A plan view of the Fermilab accelerator complex. From ref. [40].

2.1. Proton acceleration

Neutrinos in NuMI arise from the decay of particles produced by the collisions of high energy protons ($E_p = 120 \text{ GeV}$) with a target. The creation of the high-intensity proton beam used in this process employs much of the Fermilab accelerator complex; each stage will be described in sequence below.

Proton beams at Fermilab begin their life with a negative-Hydrogen (H^-) ion source, which creates H^- ions from diatomic hydrogen gas in a magnetized sputtering chamber. [41] H^- ions emerge from the source at roughly 18 KeV and are accelerated to 750 KeV by a Cockroft-Walton accelerator, after which they are fed into a linear accelerator (“Linac”) which accelerates them to about 116 MeV via an Alvarez drift-tube system. [41] A second linear accelerator (a newer side-coupled system, which is more powerful and more efficient) then accelerates the H^- to 400 MeV. [41] Since the Linac principle is based on RF power, they produce particles in “bunches” that are synchronized to the RF phase frequency.

Upon exiting the so-called “pre-accelerator” chain (the Cockroft-Walton and Linacs), the H^- are “debunched” (since the Linac bunches are not compatible with the Booster’s RF magnets) and then fed into the Booster (a synchrotron accelerator with a diameter of 150 meters) for acceleration to 8 GeV. During injection, each successive bunch from the Linac is combined with a previously-injected bunch already circulating in the Booster, after which they are passed through a carbon foil, which strips the electrons from the H^- ions. This merged collection of protons then can be further merged with successive bunches; once all the Linac bunches that simultaneously fit into the accelerator reach full intensity of about 3×10^{12} protons, they are accelerated through multiple orbits around the ring until they reach the Booster’s target energy of 8 GeV. [41, 42]

Once the beam is fully accelerated to 8 GeV, it is passed to the Main Injector, which was the immediate preaccelerator for the 1-TeV Tevatron accelerator until the

latter was decommissioned in 2011. Extraction and transfer to the Main Injector from the Booster is conceptually (if not technically) much simpler than extraction from the Linac: so-called “kicker” quadrupole magnets (which are synchronized to the proton bunch time) steer the beam out of the plane of the ring, and further magnets direct beam along a path into the Main Injector field.

While the Main Injector serves a number of other purposes even after the decommissioning of the Tevatron, we will pass over all aspects of it except those that are relevant to the creation of the NuMI beam. As with the Booster, all the bunches from the preceding step are loaded before beginning acceleration, though in this case, the bunch structure from the Booster is retained. Acceleration is to 120 GeV, after which protons can be fast extracted using another set of kicker magnets into the line that delivers them to NuMI. [43]

2.2. Protons to neutrinos

Several steps which do not involve the acceleration of protons remain before the neutrino beam itself results. An overview of the various components of the NuMI beam line is presented in fig. 2.2; the individual stages in the process are detailed further below.

First, as was noted above, bunches of protons are extracted from the Main Injector using kicker magnets. These are directed downwards into the earth at an angle to the horizon of 58 milliradians (this angle being chosen for the purposes of the MINOS experiment, the user for whom the beam line was originally constructed). [45] These beam “spills” (as they are commonly known) nominally occur at 0.45 Hz, with a design intensity of roughly 35×10^{12} protons per pulse. (For reasons that will become clear shortly, the number of protons in a spill is typically characterized by the number of protons ultimately colliding with a target; thus the beam intensity is measured in “protons on target,” or P.O.T.) This

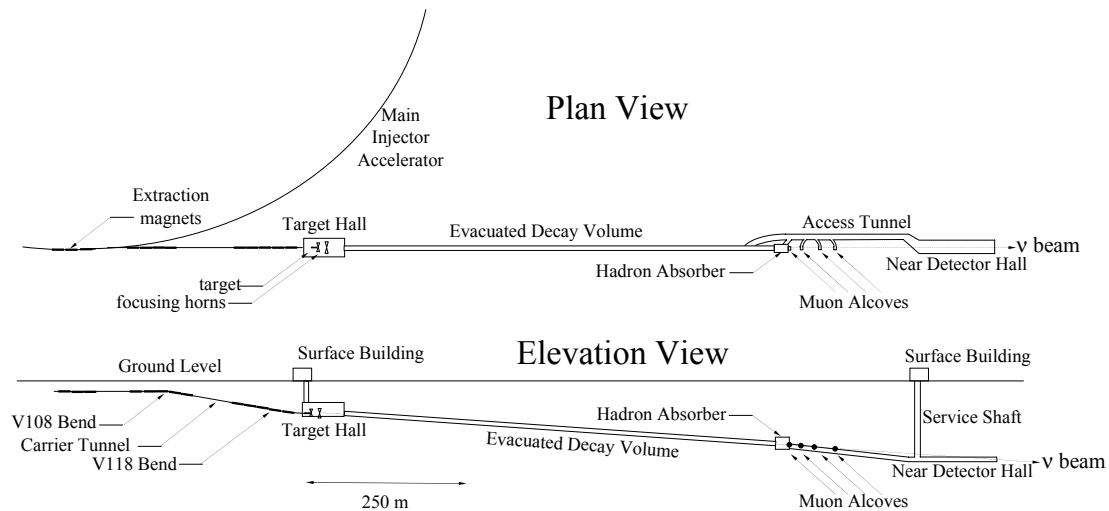


Figure 2.2: An overview of the NuMI beam line. From ref. [44].

translates to delivery of approximately 4×10^{20} P.O.T. per year under normal operating circumstances. [46]

The beam line magnets then direct protons through a graphite baffle (for the purposes of collimation and beamline equipment downstream from large mis-steerings of protons [44]) before they are forced to collide with the meson production target. NuMI production targets for the run period containing this analysis’s data are long (97 cm), narrow ($1.5 \text{ cm} \times 0.64 \text{ cm}$ cross-sectional area) graphite targets segmented into 47 “fins,” each of which is approximately 2 cm long, and which are spaced by 0.3 cm longitudinally. The target structure is affixed to hardware for stabilization, cooling, and environmental insulation. The target and associated support structure are illustrated in figure 2.3.

The reaction products of the proton beam-target interaction are extraordinarily diverse. Not only are a large number of final-state particles possible from a single interaction of a proton with the graphite, but furthermore, any particle produced in the proton-target interaction can itself interact again with the target or its carrier, or other beam line components¹, which can result in neutrino-producing decays of

¹Other examples include the magnetic horns discussed below, the decay tunnel walls, and the beam dump at the end of the tunnel.

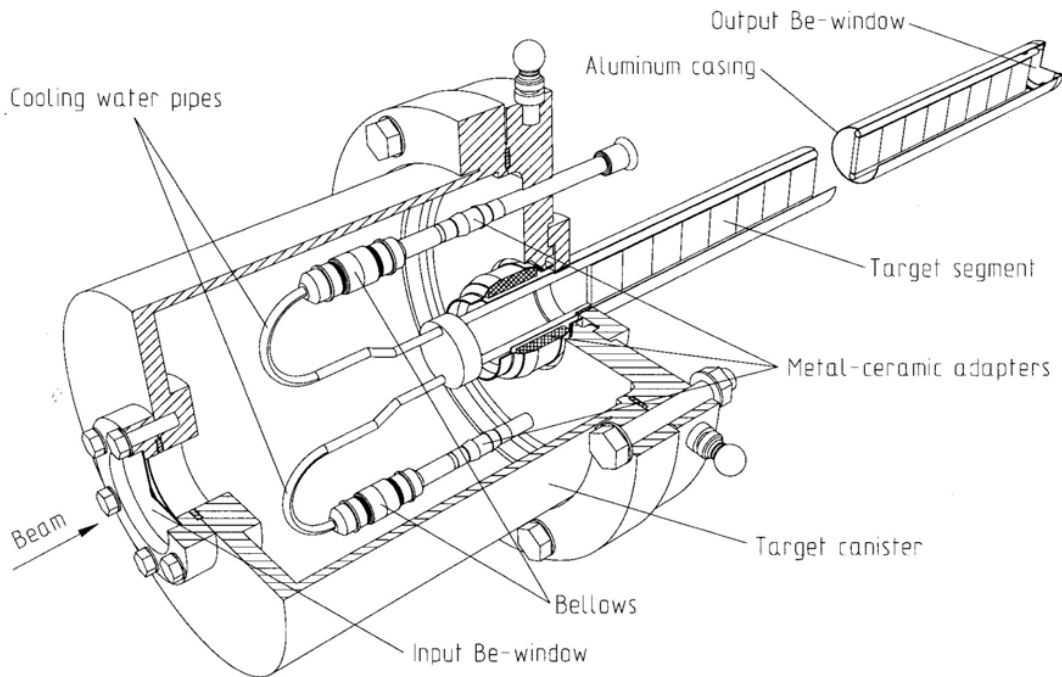


Figure 2.3: The NuMI meson production target design. From ref. [45].

different energy spectra. This mechanism—which is often given the label “tertiary neutrino production,” as the most recent ancestor undergoing an interaction was itself a secondary or later product—is one of the largest sources of systematic uncertainty in the final neutrino flux prediction. This will be discussed in more detail in chapter 3.

By design, the makeup of the particle spray which exits the target after interaction is heavily dominated by charged pions (and, to a lesser extent, charged kaons). Now, the goal of NuMI is to provide, at high intensity, as monochromatic and “mono-helitic”² a neutrino beam as possible, where the beam energy and helicity are tunable parameters. Therefore, these secondary (resulting from a primary proton’s interaction with the target) and tertiary (resulting from a secondary particle’s interaction) mesons, whose decay products will include the neutrinos for the beam, must be filtered based on both the sign of their electric charge and on their momentum. The NuMI technique, which is shared with other high-intensity

²That is, containing only a single helicity: either ν or $\bar{\nu}$, but not a mixture.

neutrino beams worldwide, is to do the selection via magnetic “horns:” large, thin aluminum conductors with roughly parabolic cavities inside. These are carefully designed to act as a magnetic lens system: when the appropriate current is pulsed peripherally around the horns, pions with kinematics that will result in the desired neutrino beam configuration are focused down the beam line by the resulting toroidal magnetic field, while pions with undesirable kinematics (as well as other particle species generally) are largely deflected aside. The current in the horn can be modulated both in amplitude and direction to achieve the desired hadron spectrum and sign. Fig. 2.4 illustrates the horn system with a schematic.³

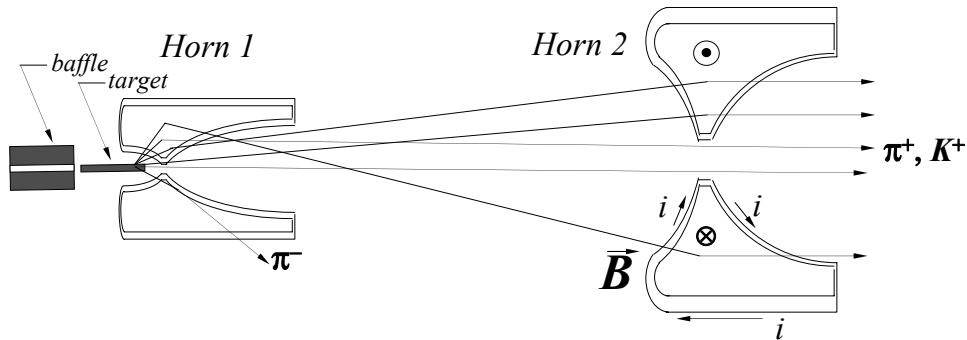


Figure 2.4: A schematic of the NuMI horn system in cross-section. From ref. [44].

Along the beam axis then follows a 675 m long, absorber-clad decay pipe, where most of the surviving particles (i.e., those not directed into the side walls) decay in flight. Here positively-charged pions and kaons quickly decay to anti-muons and muon neutrinos (the charges and helicities reversed for negatively-charged mesons)—the latter of which, because of their minimal mass, are strongly boosted along the initial focusing direction and travel mostly along the beam axis. The electron neutrinos and muon antineutrinos resulting from the decay of anti-muons (again, reversing the helicities for negative mesons) are not as strongly boosted due to the three-body nature of the muon decay, but some still pass through the

³While it is possible to further improve the chromatic spread of the beam by placing the detector slightly displaced from the main axis of the beam (see, e.g., ref. [47]), one must sacrifice much of the beam intensity by doing so. NuMI trades monochromaticity for intensity.

detector, resulting in a mild “contamination” of the muon neutrino beam with electron neutrinos and wrong-helicity muon anti-neutrinos.⁴ (Much more will be made of this in ch. 3.) Finally, a beam dump at the end of the volume absorbs leftover proton beam and any other undecayed hadrons. Muons that have not decayed in the decay volume typically stop in the 270 m of rock between the decay volume and the detector hall (which is approximately 100 m underground). Figure 2.2 gives an overview schematic of the beam line setup. [46]

Though originally the decay pipe was designed such that it would sustain a medium vacuum of around 1 Torr, safety concerns prompted the later introduction of a system to maintain an atmosphere of approximately 0.9 atm of helium gas in the decay pipe. While less than ideal (secondary mesons can now be deflected from the beam axis by interactions with the helium gas), this configuration has proven to introduce only manageable errors into the flux prediction. [48] All of MINER ν A’s data were acquired with a beam configuration containing helium in the decay volume.

The data presented in this analysis were collected in the so-called “low-energy, forward horn current” (LE FHC) beam line configuration, where the target is partially inserted into the more upstream of the two horns, leaving 10 cm separation between the target end and the narrowest part of the horn. The horn current direction in this configuration is chosen to select positively-charged mesons and its magnitude is fixed at 185 kA.⁵ As chapter 3 will demonstrate, this configuration is projected to yield a neutrino beam illuminating MINER ν A which consists of 94.13% ν_μ , 5.14% $\bar{\nu}_\mu$, 0.69% ν_e , and 0.04% $\bar{\nu}_e$.

⁴Additional electron neutrinos and antineutrinos are also produced directly in the decays of kaons, typically at higher energies than we will be concerned with for this analysis. The kaon-parent flux is roughly 10% of the electron neutrino flux and is treated together with it in ch. 3.

⁵This target-horn configuration pairing is often assigned the code “le010z185i”.

3 Characterizing the neutrino flux

Because neutrinos rarely interact with matter, it is difficult to make a direct measurement of the energy spectrum of a neutrino beam. In addition, analytic calculations of this neutrino flux are notoriously difficult: first, because the strong-force interactions in the NuMI target are poorly modeled by current theory; and second, because the target system is comprised of many different components in which beam-line particles can undergo re-interactions. MINER ν A therefore relies on a detailed simulation program—based on the GEANT4 simulation package [49]—significantly augmented by internal and external data constraints to provide a prediction. The final flux prediction has uncertainties of order 8-10%.

3.1. g4numi: GEANT4 simulation

We begin with a GEANT4-based [49] simulation package known as g4numi, which models the interaction of 120 GeV protons with the graphite neutrino target and the resultant decay product interactions with the target and horns (as described above in section 2.2).¹ Because GEANT4 is modular, it allows the run-time selection of various interaction models bundled with the package. Based on its superior agreement [50] with the NA49 data set mentioned in section 3.2.1 below, we elected to use the newer “FTFP-BERT” (FRITIOF Precompound + Bertini cascade) physics list rather than the default “QGSP-BERT” model [51]. After propagation through the target and horn system, g4numi yields the neutrino fluxes depicted in fig. 3.1.

¹Our version of g4numi is built upon GEANT4 version 9.2p03.

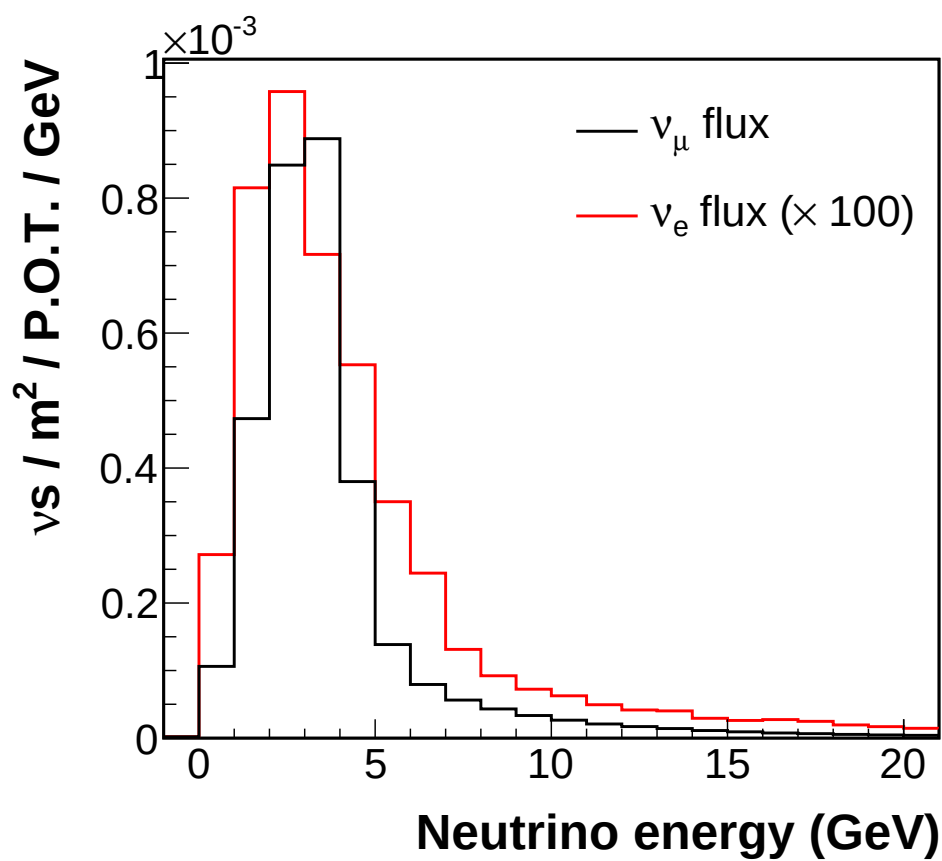


Figure 3.1: The predicted NuMI beam flux from g4numi for the le010z185i configuration described in the text.

3.2. External constraints from hadron production experiments

The dominant uncertainties in the simulation stem from an underlying inability to predict the products of the strong interactions of the beam protons with the graphite NuMI target. Because we cannot predict with certainty the particle composition of the by-products and their energies, the neutrino beam composition and spectrum are also consequently uncertain. We therefore turn to experimental data to fill in the gaps in our predictive machinery.

Detailed measurements of the particle production species and spectra for reactions such as those occurring in the NuMI target are present in the literature. But, mainly because NuMI’s target is thick (multiple interaction lengths), until very recently there was no single published work presenting results that apply to NuMI directly. (The MIPP experiment at Fermilab collected data on a replica NuMI target, at proton energies of 120 GeV, but their conclusions [52] were not available to be incorporated into our prediction in time to be included with this analysis.) Instead, for an external flux constraint, we are forced to choose between data sets that were collected on much thinner targets or at different energies (the data from CERN’s NA49 experiment [53], for example, and some of that from NA61 [54]).

For the results discussed in this work, we elected to use the data from NA49 [53] as our primary external constraint, supplemented by data from Barton *et al.* [55] in specific cases. NA49 collected data from 377,000 proton interactions on a thin (1.5% interaction length) graphite target at proton energies of 158 GeV (the CERN SPS beam) and measured the resulting charged particle spectrum using a proportional counter detector. Barton’s group used a 100-GeV proton beam (M6E at Fermilab) and similar target and detector technology to NA49, though their

statistics were much poorer. We use these datasets to re-normalize the interaction rates predicted by the GEANT4 simulation of the proton-target collisions in NuMI (sec. 3.1) as described below in section 3.2.1. The validity of the scaling procedure used to translate from NA49 and Barton energies to NuMI energies is discussed in appendix A.

3.2.1. Central-value reweighting of NuMI simulation

Beam attenuation correction

Since the neutrino production target is multiple interaction lengths long, as the proton beam passes through the target, progressively more protons interact, and the beam is therefore progressively attenuated. Though the beam simulation attempts to model this phenomenon, it does so using the cross-sections from the FTFP-BERT model, which we found we needed to adjust using NA49’s measurements (on which see the detail further below). We therefore correct the attenuation calculation itself based on the difference between GEANT4’s FTFP-BERT cross-section and the one measured by NA49. To do so, we warp the neutrino spectrum by applying an additional weight to generated neutrino interactions (sec. 5.2.1) according to the following formula:

$$w_{\text{atten}} = \exp(-x(\sigma_{\text{data}} - \sigma_{\text{sim}})\rho) \quad (3.1)$$

where x is the axial distance into the target the interaction occurred; the σ s are the production cross-sections measured by NA49 and used in the simulation, respectively; and ρ is the average target density (graphite plus a small admixture of helium due to the spaces between the target fins and the presence of helium in the decay pipe (sec. 2.2)). The effect of this correction is small, but not negligible; it is illustrated in fig. 3.2.

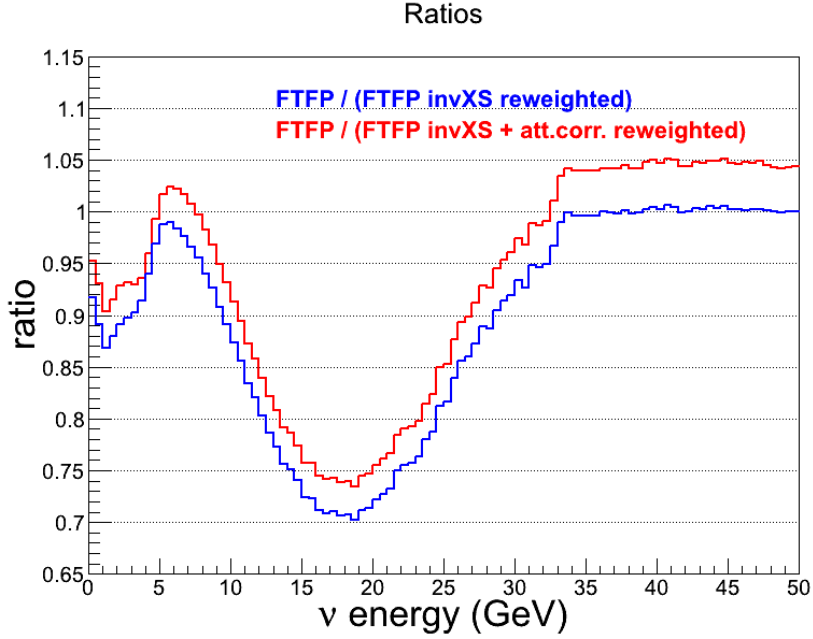


Figure 3.2: The effect of the beam attenuation correction as a ratio to the nominal FTFP-BERT flux. The attenuation correction is the difference between the blue and the red curves. Courtesy L. Aliaga (MINERνA).

Production cross-section correction

Though one can describe the data from experiments like NA49 and Barton in terms of the outgoing particle kinematics directly (outgoing momentum and angle, for example), these quantities are bound in an inflexible way to the initial conditions of the experiment (particularly the incoming proton momentum). Thus, it is traditional to report results instead in terms of quantities that scale sensibly to different beams. In particular, the resultant distributions of production cross-sections are usually given in terms of the particle’s transverse momentum (the component transverse to the proton beam), p_T , and the scaling parameter introduced by Feynman [56]:

$$x_F = \frac{2p_L}{\sqrt{s}}$$

where p_L is the component of the particle’s momentum that is parallel to the beam in the center-of-momentum frame, and s is the square of the total center-

of-momentum energy of the reaction. According to Feynman’s postulate [56] (well-supported in energy regimes above ~ 10 GeV by accelerator data; see, e.g., ref. [57]), the invariant cross-section (“invariant” because it is a Lorentz invariant [11]) of an inclusive particle production process,

$$f = E \frac{d^3\sigma}{dp^3} \quad (3.2)$$

depends on x_F and p_T , but not s .

We use this so-called Feynman scaling to apply the NA49 and Barton measured pion production cross-sections to the NuMI data by assuming that the Feynman scaling (along with measured corrections and violations) implemented in the FLUKA simulation package [58, 59] is correct. (The uncertainty in the flux due to this procedure is negligible, as discussed in app. A.) Under this hypothesis, we scale (for example) the NA49 measured cross-section by the ratio of FLUKA’s predicted cross-section at the target energy to FLUKA’s prediction at the NA49 energy. That is:

$$f(E_{\text{NuMI}}) = f_{\text{NA49}}(158 \text{ GeV}) \times \frac{f_{\text{FLUKA}}(E_{\text{NuMI}})}{f_{\text{FLUKA}}(158 \text{ GeV})}$$

Because NuMI’s target is much thicker than that used by NA49 and Barton, we can only apply reweighting factors from these experiments to events in NuMI which are comparable. We term particles from interactions where the original beam proton interacted just once in the target, and the exiting charged particle did not re-interact, “NA49-like” particles, and reweight only neutrino events with neutrinos stemming from “NA49-like” particles. With this caveat, then, we reweight each event exiting the NuMI simulation whose neutrino ancestry contains an “NA49-like” particle whose kinematics are covered by the data from NA49 (or Barton) by the ratio of the measured NA49 (Barton) cross-section to the cross-section used by

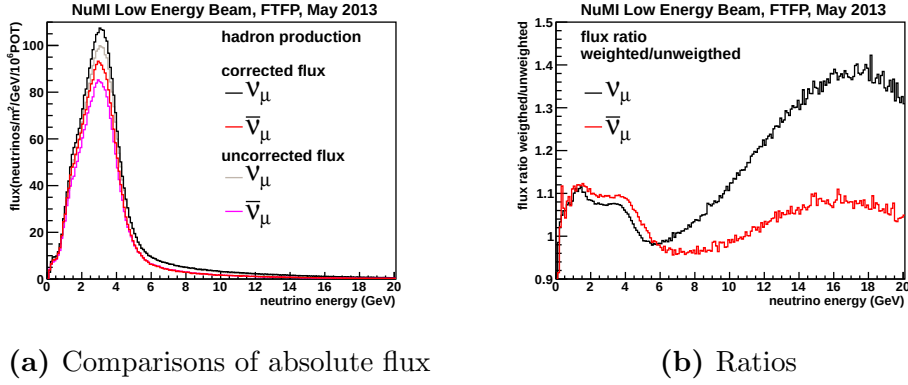


Figure 3.3: The predicted NuMI le010z185i beam flux after reweighting. Courtesy L. Aliaga, M. Kordosky, A. Norrick (MINERνA).

GEANT. More specifically, pion weights are calculated from the NA49 data when $x_F < 0.5$ and from the Barton dataset when $x_F > 0.5$; we apply weights to events where a kaon is produced if $x_F < 0.2$, based on a different study from NA49 [60]; and we apply weights calculated from NA49 data for events producing a proton with $x_F < 0.95$ from a third study by NA49 [61]. The weight applied to a simulated neutrino event is then simply the ratio of the scaled NA49 (or Barton) cross-section to the GEANT cross-section at the NuMI energy (120 GeV):

$$w_{\text{NA49}} = \frac{f_{\text{NA49}}(E_{\text{NuMI}})}{f_{\text{GEANT}}(E_{\text{NuMI}})} \quad (3.3)$$

Plots of the modifications that this procedure makes to GEANT’s prediction of the NuMI beam flux are in fig. 3.3. Note that the alterations are, in general, quite large (up to 40%).

3.3. Flux uncertainties

The lion’s share of the uncertainties in the *a priori* flux prediction that we obtain from the GEANT4 prediction of sec. 3.1 are due to re-interactions of hadrons within

the target material. When the simulation throws an interaction with outgoing kinematics that are governed by the external hadron production data of sec. 3.2, we are able to use that data as a constraint, as explained there. However, for the remaining interactions—of which there are, unfortunately, many—we are forced to choose some method for evaluating the uncertainty in the final flux prediction due to the uncertainties in the underlying models used in GEANT4.

The uncertainties in these parameters are often correlated in such a way that they are difficult to propagate through calculations analytically. To cope with this fact, we have adopted a sampling-style approach in which we generate a large number (usually roughly 1000) of variations on the desired distribution in which the various parameters in question have been sampled based on their central values and known uncertainties. We then compute the uncertainty on the final distribution—here, the flux—from the spread of these varied distributions in each bin.

The various parameters which must be varied in GEANT according to their uncertainties are grouped together into discrete “tunes,” which, in the GEANT jargon, are usually called “physics models” or “physics lists.” We use the difference between the hadron kinematics predicted by these physics lists, which model the interactions of hadrons with matter in different ways, to determine the mean and variance of the distribution from which the relevant parameters are sampled in the variations. In particular, for this analysis we compared the predicted outgoing kinematic distributions of pions and kaons (which decay in flight to neutrinos) from the following models to determine the uncertainties on non-NA49-constrained neutrinos:

QGSP “quark-gluon string” model with “pre-compound” model back-end

QGSP-BERT above QGSP coupled to Bertini cascade model

QGSC-BERT like QGSP except with a “chiral-invariant phase space” model

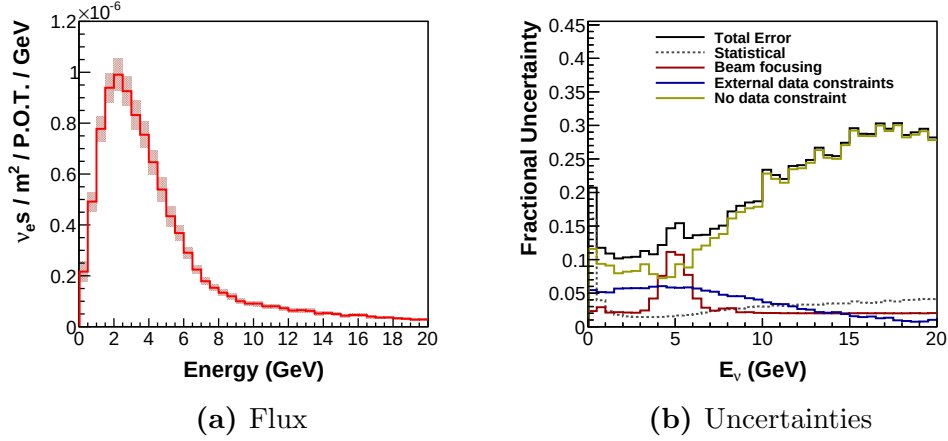


Figure 3.4: The final *a priori* $\nu_e + \bar{\nu}_e$ flux prediction and its uncertainties.

back-end

FTFP-BERT Fritiof-like string pre-compound model with Bertini cascade

FTF-BIC Fritiof-like string model with a binary-cascade model

All of these models are described in further detail in ref. [51].

3.4. *In situ* constraint: neutrino-electron elastic scattering

Because of the difficulties associated with the *a priori* flux prediction and uncertainties that we obtain from secs. 3.1-3.3, we furthermore employ an *in situ* measurement technique to arrive at our best estimate of the flux. This method relies on a particular reaction channel which is predicted very precisely in the standard model: the neutral-current, elastic scattering of neutrinos from atomic electrons within the detector medium, $\nu + e \rightarrow \nu + e$. (The charged-current scattering of electron neutrinos on electrons yields an identical final state, and therefore interferes constructively with the neutral-current process; it is similarly well predicted by the SM and will be lumped together with the neutral-current

process for the remainder of this section.) The isolation and extraction of the rate of such events is a challenging process, and will not be considered in detail here; it is described in depth in ref. [62]. Instead we will focus here on the application of the constraint it provides to the flux prediction.

To constrain the flux prediction using $\nu + e$ events, we consider the distribution of incoming neutrino energies from the *a priori* prediction of secs. 3.1-3.3 given in fig. 3.4a. The uncertainties illustrated in fig. 3.4b are computed, as noted in sec. 3.3, by varying parameters in the underlying models, and computing the spread of the resulting distributions. Since each of these variations in the flux can be propagated through our simulation to yield a prediction for the number and spectrum of $\nu + e$ events, we begin by comparing the measured distribution in the data to the predicted distribution of $\nu + e$ events in the variation and computing a χ^2 statistic between them:

$$\chi^2 = \sum_{i=1}^{N_{\text{bins}}} \frac{(D_i - V_i)^2}{\sigma_i^2} \quad (3.4)$$

where D_i is the content of the i th bin of the measured distribution in data (with its uncertainty σ_i), and V_i is the content of the i th bin in the simulated variation. We assign the value of the probability computed in the standard way from the χ^2 statistic [11] to each variation as its weight. After all the weights have been assigned, we renormalize them such their mean is 1. To determine the predicted uncertainty from these variations, we then compute a *weighted* variance from their spread (instead of the unweighted version we used previously) in each bin:

$$\sigma^2 = \frac{1}{\sum w_i} \sum w_i (x_i - \mu')^2 \quad (3.5)$$

where w_i is the weight for variation i , and x_i and μ' are the value of the variation and the weighted mean of all of the variation distributions in bin i , respectively.

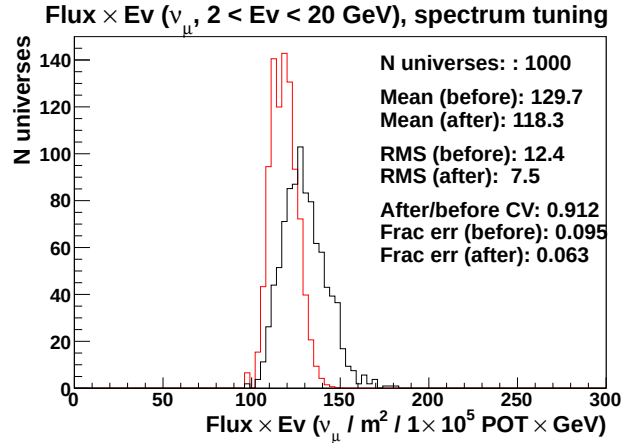
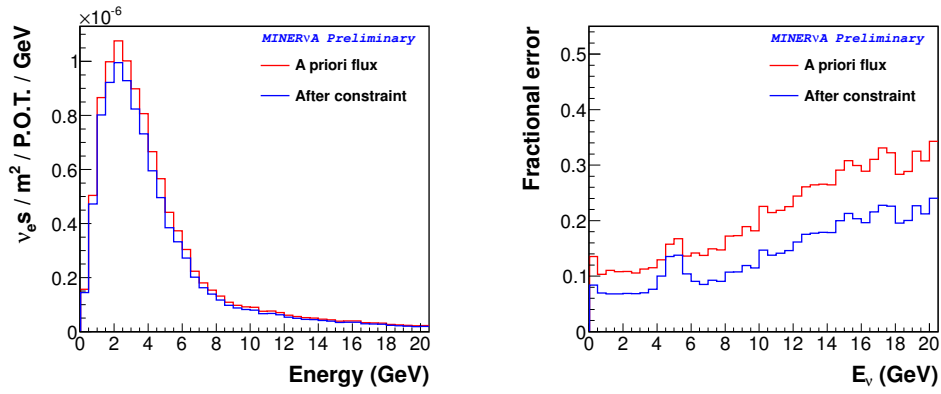


Figure 3.5: Effect of the $\nu + e$ scattering constraint on the predicted number of $\nu + e$ events. Courtesy J. Park (MINER ν A).

The effect that the constraint has on an example variable—the number of $\nu + e$ events, which is a single-bin variable, chosen for simplicity—is presented in fig. 3.5.

Applying the weights to the flux-parameter variations in this fashion reduces the spread of the variations, and thus the uncertainty, but it does not adjust the central-value prediction. However, we can also make use of the information contained in the weights to improve the central value as well. In this case, we examine the mean of the distribution of variations in each bin of the variable in question. The ratio of the weighted mean (that is, after the constraint is applied) to the unweighted mean is multiplied (bin by bin) against the central value distribution to yield the corrected central value distribution. Finally, any other variations associated with the variable in question that are not variations in the flux parameters are translated by the same amount as the difference in the central value, so as to render the fractional variance invariant. The effect on the ν_e flux is illustrated in fig. 3.6.

Because this approach uses the predicted values of the variable of interest as they are fluctuated by varying the underlying flux model parameters, we can use it



(a) The ν_e flux prediction before and (b) Uncertainties on the ν_e flux before and after constraint.

Figure 3.6: Effect of the flux constraint described in the text on the $\nu_e + \bar{\nu}_e$ flux prediction and its uncertainties.

to predict the effect of the flux constraint on *any* distribution, not just the flux itself. This feature will be heavily exercised in sec. 7.

3.5. Final flux prediction

The final flux prediction (the *a priori* prediction with the constraint of sec. 3.4 applied) is shown in fig. 3.7.

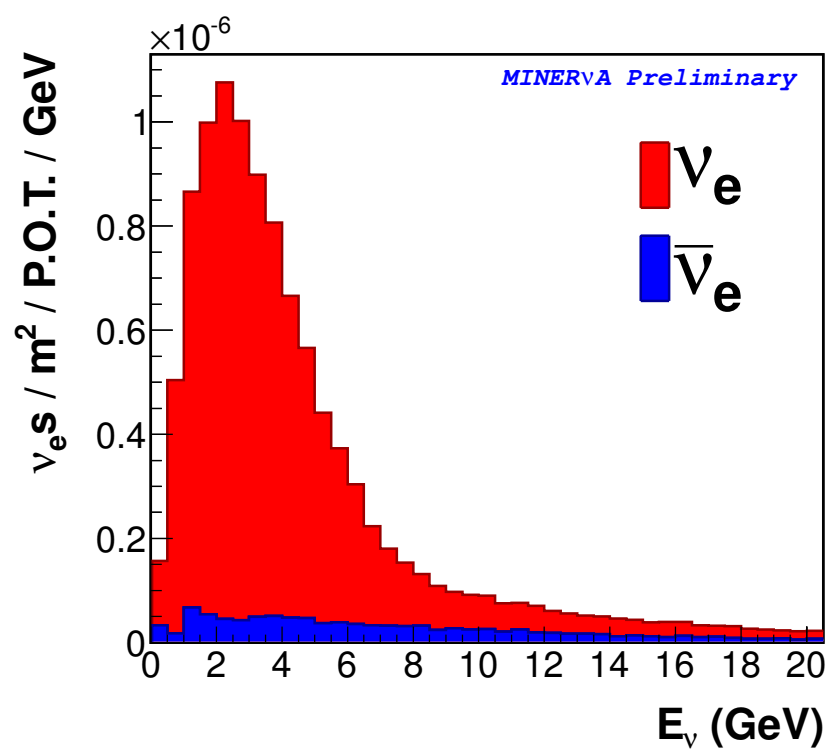


Figure 3.7: The final flux prediction used in ch. 7.

4 The MINER ν A detector

MINER ν A is a finely segmented, scintillator-based tracking detector and sampling calorimeter. Like many other particle physics detectors, it can be subdivided into a handful of subdetectors which serve different purposes. It is instructive to consider the detector in a cylindrical coordinate system, in which case we can subdivide along two different axes: radially, the detector breaks into an Inner Detector (“ID”), designed primarily for precise particle tracking, and an Outer Detector (“OD”), which primarily serves as a sampling calorimeter; axially, from upstream to downstream along the beam line, MINER ν A is built of a “nuclear targets” region (passive targets of different materials interleaved with scintillator), a “tracker” region (pure scintillator), an electromagnetic calorimeter (ECAL) region, and a hadronic calorimeter (HCAL) region. The active parts of the detector transmit their light into a network of photomultiplier tubes (PMTs) via optical fibers, and custom electronics read out the response of the PMTs. The nature of each subsection of the detector, as well as the technology used to instrument it, is enumerated in detail below; further descriptions, as well as detailed performance metrics, can also be found in ref. [63].

4.1. Subdetector arrangement and function

MINER ν A consists of 120 welded steel hexagonal “modules” that hang in a manner akin to file folders on a steel support structure attached to the floor of the underground experiment hall. Though conceptually the ID and OD are different subdetectors, the steel frame of the OD serves also as the support structure for

each module, and therefore each module contributes to both subdetectors. As shown in fig. 4.1, moving radially inward from the outer edge of the OD region of each module, one first encounters four so-called “stories,” layers of two rectangular strips of scintillator, interleaved with the steel. Contained further inward, one finds one or two inner detector hexagonal planes (each half the axial depth of the module in the latter case), the composition of which varies depending on the module type: nuclear target-type modules are single planes built from sheets of graphite, lead, and/or steel, edge-welded together in various shapes to make a hexagon; tracker-type modules are made of two planes of edge-glued scintillator strips (further described in sec. 4.2); ECAL-type modules are identical to scintillator modules except that a 0.2 cm layer of lead is affixed to the front of the more upstream plane; HCAL-type planes are composed of a plane of solid steel followed by a plane of iron. Furthermore, a “collar” of lead (shaped like a hexagonal annulus; again see fig. 4.1) covers the outermost portion of scintillator in the tracker-type planes. The only exception to the general pattern described here is the water target, residing in the middle of the targets region, which (due to design considerations) has no OD frame and is composed of Kevlar sheets fixed to a circular steel frame and is filled with distilled water.

Each longitudinal region—nuclear targets, tracker, ECAL, HCAL—is then composed of an assortment of modules of the appropriate types; this is depicted in fig. 4.2. At the most upstream end, the nuclear target region alternates each of the five nuclear target-type modules with four successive tracker-type modules, for a total of 28 modules. (As the nuclear target region is not used for the analysis described in this work, it will be neglected from here on.) This is followed by the tracker region (62 modules), which occupies the bulk of the detector volume (though still a minority of the mass), made entirely of tracker-type modules. The downstream end of the detector contains the ECAL (10 modules; ~ 8 radiation lengths) and HCAL (20 modules; ~ 3 nuclear interaction lengths), in that order,

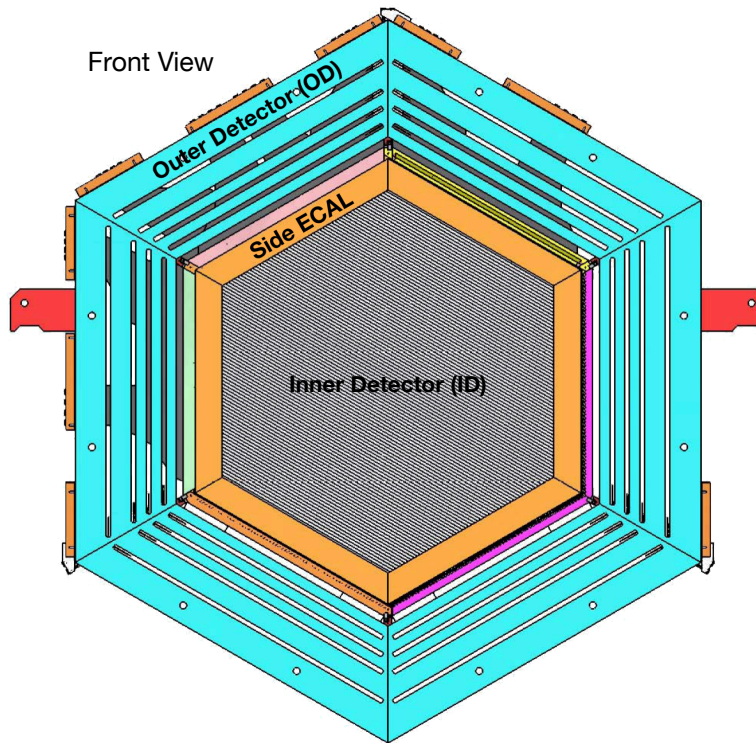


Figure 4.1: An engineering drawing of a MINERνA tracker-type module, as seen from the front (looking downstream along the beam line). The “ears” from which it hangs on the support structure are shown in red; optical fiber couplings are the outermost yellow elements; the OD steel is shown in cyan; the lead collar is yellow; and scintillator strips are gray.

which are built from the corresponding type of modules. (A “transition” module, which is like an ECAL-type module except that the lead sheet is attached to the back of the module instead of the front, sits in place of a normal tracker module at the end of the tracker region.)

One additional complication to the geometry of the detector is the plane orientation. Because each strip of scintillator runs the full width of the ID portion a plane, any given plane can yield only two-dimensional information about the position of a particle passing through it: the coordinate transverse to the axis of the strip, and the coordinate along the beam direction (the latter being shared by all strips within the same plane). Therefore, to fully reconstruct a three-dimensional path through the detector, it is necessary to orient the planes along different axes—termed “views”—in an alternating fashion. MINERνA uses three plane

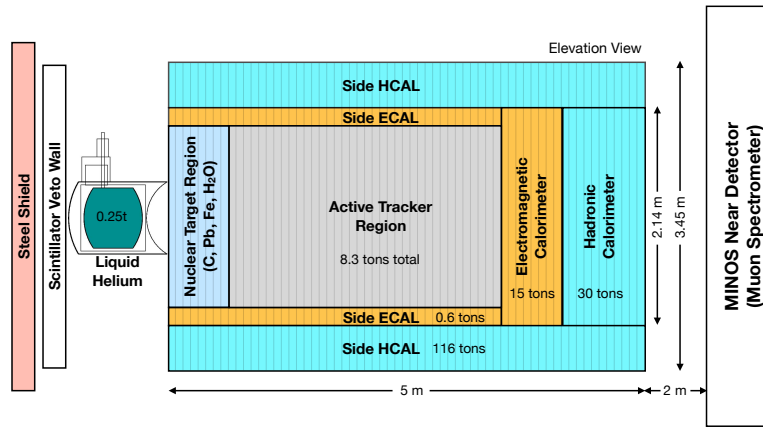


Figure 4.2: The layout of the modules in MINERνA. Beam is incident from the left. The upstream veto system, liquid helium target, and MINOS near detector are not mentioned further in the text as they are unused in this analysis.

views: “X,” in which the strip axis points along the vertical (as defined by gravity); and “V” and “U,” where the axis is rotated by $\pm 60^\circ$ (respectively) as compared to X. Successive modules group these varying views in a regular pattern: the first module in a sequence is composed of a U plane and an X plane (in that order); it is followed by a module containing a V plane an X plane (again in that order); and then the pattern is repeated.

This choice of plane views informs the coordinate system chosen for the detector geometry. The $+\hat{y}$ -direction is the the usual “up,” antiparallel to the direction of gravitation. The $+\hat{z}$ -axis is chosen to lie parallel to both the cavern floor and the downstream direction of the beam (which, since $\hat{z}$ must be perpendicular to $\hat{y}$, and since the beam points slightly downward, means that $\hat{z}$ forms an angle of about 3° to the beam). To complete a right-handed coordinate system, then, the $+\hat{x}$ -direction must point to the left when viewing the detector from the most upstream end and looking downstream. $(x, y) = (0, 0)$ represents the center of an ID plane, while $z = 12000$ mm is anchored to the front of the MINOS near detector (which was installed in the detector hall many years prior to MINERνA). With these choices, the unit vectors corresponding to the positive strip directions in the U and V planes, respectively, become $\frac{1}{2}(\hat{x} \pm \sqrt{3}\hat{y})$. The coordinate system is better

described by an illustration, given in fig. 4.3.

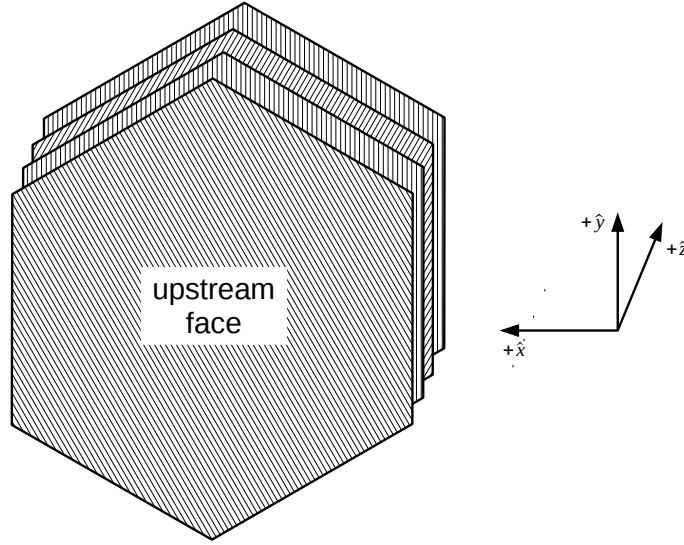


Figure 4.3: The MINERνA coordinate system, shown with representative schematics of detector planes for reference. This view is from upstream toward downstream, where the $+\hat{z}$ -direction is the direction in which the planes are stacked. From front to back, the planes shown are in the UXVX ordering described in the text.

4.2. Active plane composition and instrumentation

Despite the various orientations they are assembled into, planes whose ID region are composed of active scintillator (and not of some passive material) are all constructed from the same materials and share the same basic structure. The ID region of such a plane is composed of 127 triangular strips of scintillator material made from a co-extrusion of doped polystyrene and a reflective titanium dioxide outer coating. Each of these is, cross-sectionally, an isosceles triangle with width 3.3 cm and height 1.7 cm; strips are cut longitudinally such that, when assembled, they fit into a regular hexagon of apothem 1.07 m. (Note that this implies strip lengths which vary from 1.2 m on the plane edges to about 2.5 m at the plane center.)

Component	H	C	O	Al	Si	Cl	Ti
Strip	7.59%	91.9%	0.51%	-	-	-	0.77%
Plane	7.42%	87.6%	3.18%	0.26%	0.27%	0.55%	0.69%

Table 4.1: Chemical composition of scintillator strips and constructed planes, by mass percentage. From ref. [63].

Glued (using an optically clear epoxy) into a 2.6 mm hole running lengthwise along the strip is so-called “wavelength-shifting” (“WLS”) fiber, a proprietary product of Kuraray Co., which transports scintillation light out of the strip. In so doing, light is downshifted from the ultraviolet spectrum of the scintillator response into the visible spectrum where the photomultiplier tubes (described further below) are most sensitive. This results in typical attenuations of roughly 40% for the longest strips. WLS fibers are further coupled to clear fibers, which carries light the remaining distance to the photomultipliers (where electronic readout begins, as in sec. 4.3), further attenuating the light by a factor ranging from about 0.85 (for the shortest fibers) to 0.45 (for the longest ones).

Strips are glued edgewise (using a different, translucent epoxy) in an alternating fashion so as to make a consistent edge. (See figure 4.4.) In order to ensure that ambient light in the detector hall cannot leak into the strip and overwhelm the comparatively dim signal from particle interactions, each plane face is then covered with sheets of Lexan (glued on with an opaque gray epoxy), and further secured with black PVC electrical tape. Planes assembled as such have an average areal density of $2.02 \pm 0.03 \text{ g/cm}^2$, of which the chemical composition is listed in table 4.1.

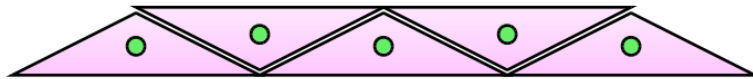


Figure 4.4: Illustration of the arrangement of strips within a plane, viewed along the axis of the strips. The green circle represents the WLS fiber.

4.3. Readout electronics and data acquisition system

The electronic readout system converts actual scintillator light to an electronic representation suitable for storage and analysis. As with the detector design, each subsystem is explained below, with significantly more detailed specifications available in ref. [64].

4.3.1. Photomultiplier tubes

Light incoming via the clear fibers of sec. 4.2 is coupled into light-tight capsules containing 64-channel multi-anode photomultiplier tubes (PMTs) using commercial connectors. Inside these capsules, known better as “PMT boxes,” optical fibers on the receiving end of the connectors are woven into a special arrangement which maps the light from neighboring strips in a plane in the detector onto non-neighboring PMT pixels, thereby mitigating the effect of the phenomenon known as “cross-talk” (discussed in further detail below). Each fiber is mated to a single PMT channel using a specially machined connector.

Photons passing into the PMT strike a cathode, and, by the photoelectric effect, liberate typically zero or one electron per photon (the quantum efficiency of the sensor being roughly 20%). Electrons are then accelerated through twelve “dynodes” in turn, each of which typically bears a higher fraction of the total PMT potential of several hundred volts and each of which is made of a material which has a large secondary electron emission probability. When struck by an incoming electron, therefore, more electrons are created from a dynode, resulting in a gain in the number of electrons of a factor of 10^5 - 10^6 by the end of the chain. [65] A custom printed circuit board attaches to the base of the PMT to facilitate collection of the charge from the 64 independent anodes. A photograph of a partially assembled

PMT box is given in fig. 4.5.

Multi-anode PMTs experience two notable side effects which are undesirable and must be taken into account. First, because the pixels on the PMT face (and the corresponding dynode chains within the PMT itself) are relatively small and border each other directly, photons intended for one pixel (or photoelectrons intended for one dynode chain) are able to wander into a neighboring pixel (or dynode chain). This phenomenon is known as “cross-talk;” the first variety is called “optical cross-talk,” and the second will be referred to as “dynode-chain cross-talk.” Based on bench measurements, we both simulate and attempt to subtract this cross-talk as a calibration. (The measurements and procedures used are described in sections 5.2.3 and 5.3.7, respectively.) Second, PMTs are susceptible to what is known as “afterpulsing,” in which residual gasses which have leaked into the PMT through the glaze face are ionized by electrons progressing down the chain. Because the resulting ions are positively charged, they accelerate the wrong way through the dynodes, eventually colliding with a dynode or the photocathode and producing a secondary (lesser) electron cascade. This results in a delayed signal. As discussed below in sec. 4.3.2, afterpulsing plays a role in so-called “dead-time” creation, forcing a compensatory procedure of overlaying data onto our simulation. (Further discussion is below.)

4.3.2. Front-end electronics

The current exiting the anodes of one PMT is collected by a custom printed circuit board called a “front-end board” (“FEB”), which mounts directly on to the PMT box. Besides serving as the high-voltage power supply for the PMT via the onboard Cockroft-Walton generator, the FEB also digitizes the incoming charge in three gain ranges using one of six onboard “TriP-t” (i.e., Trigger and Pipeline with timing) integrated circuits. Charge is collected during a 16-microsecond window

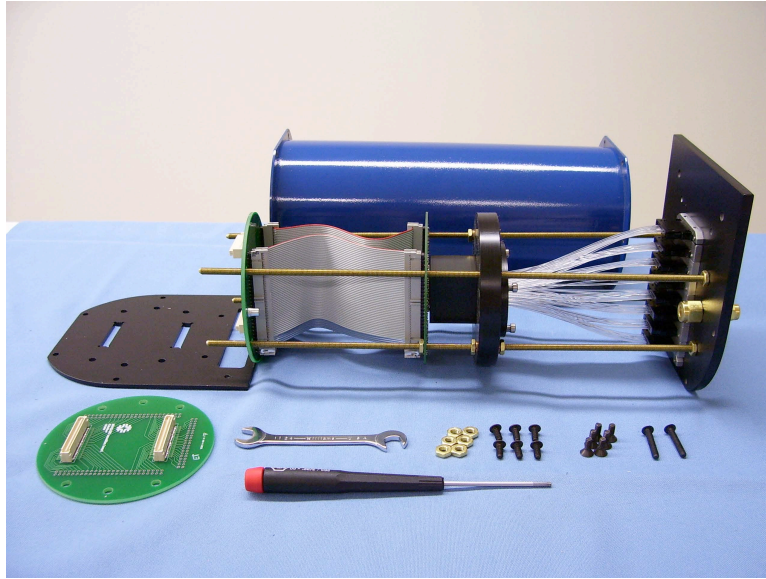


Figure 4.5: A partially assembled PMT box. From right, visible are the fibers, the coupling fixture (black cylinder), the PMT itself (black rectangular solid), PMT base boards with ribbon cable, and end cap. A spare PMT base board is front left. From ref. [63].

called a “gate,” designed to capture the entire breadth of a roughly 10 microsecond Main Injector beam spill plus a large fraction of slower particle decays.

The digitization scheme is roughly as follows. First, analog charge in each gain range is continuously fed into a discriminator circuit, which latches when a pre-set threshold is reached. Then, following a discriminator latch, charge in each gain range is pushed into time-stamped buckets in a charge pipeline for a pre-set period (which, during the run period of this analysis, was equivalent to 16 9.4ns-long system clock ticks, or 150.4ns). Finally, charge is pushed out of the pipeline and digitized and stored, and the discriminator capacitor is cleared, during a 20 clock tick (188ns) window during which no new incoming charge is recorded. After this the board is ready to accept charge once again. The pipelines contain enough register space to store seven of these time-stamped segments plus an eighth, untimed readout (which happens at the very end of the 16-microsecond gate).

The system clock on an FEB, which governs the clock ticks noted above, is

regulated by a 53 MHz crystal oscillator built into a field-programmable gate array (FPGA) chip integrated onto the FEB. Although the crystal's fundamental resonance controls the system clock tick length and provides the resolution quoted above (9.4 ns), when the FEB is time-stamping activity in the charge pipeline, it is further able to compare the phase of the oscillator to a reference and thus subdivide a clock tick to into quarters. Therefore, the best resolution on the time stamp of each pipeline bucket is one quarter-tick, that is, 2.35 ns.

Since neutrino interactions during this running period are relatively rare, occurring only a few times per beam spill, very little of the detector is illuminated at any given time. However, the design of the FEBs contains one feature which, unfortunately, tends to amplify any effects from the pile-up of multiple particle interactions in one gate when such a thing does transpire. There are only six TriP-t chips for any given FEB, as observed above, while there are 64 incoming channels whose charge is being distributed among them; this results in a number of channels being ganged together so that they can be serviced by the same TriP-t. (The high- and medium- gain ranges of the channels are split between four of the TriP-ts, meaning that those TriP-ts handle 16 channels each; meanwhile, the low-gain ranges are divided between the remaining two, making them responsible for 32 channels each.) Consequently, when the discriminator latches for any particular channel, all of the channels in its TriP-t group experience the same digitization and reset periods. Should any new activity be incoming on *any* of the channels in this group during the reset time, it will be lost. Typically we refer to this phenomenon as “discriminator dead time,” and it necessitates a modification technique to the simulation in order to simulate events that are not observed as a result (as discussed in sec. 5.2.3). Furthermore, in an event with many neutrino interactions or other related activity (e.g., PMT afterpulsing, explained in sec. 4.3.1), the readout buffers can be saturated, such that real physics activity is relegated to the untimed eighth bucket in the charge pipeline. The overlay technique mentioned above is a

partial simulation of this effect, as well, but cannot compensate for it completely since we do not accurately simulate afterpulsing from activity originating in the simulation itself. We disregard any activity in the eighth charge bucket, which is a negligible correction for this analysis.

4.3.3. Rack-mounted electronics

The FEBs are synchronized to each other and to the Main Injector time by daisy-chaining them together and operating the chain as a slave to a device called a chain read-out controller (“CROC”). Synchronization and clock signals from the CROC (whose provenance will be further discussed shortly) are transmitted through the FEB loop, adding a pre-programmed time offset (measured in an *in situ* timing calibration (see sec. 5.3.6) stored in each FEB’s FPGA chip (sec. 4.3.2) to account for the delay corresponding to the FEB’s position in the loop. The maximum number of FEBs thus chained together is governed by the degradation of the serialization lock and timing signals through the ethernet cable connecting them together; this constraint imposes a practical limit of a maximum of 10 FEBs per chain. Each CROC supports four chains; in total MINERνA uses 15 CROCs to service the 507 FEBs used on the detector (more than the minimum of 13 because of design considerations arising from the PMT mounting locations).

The 15 CROCs are VMEbus (“VERSAmodule Eurocard”) devices [66], and are therefore used exclusively from within a readout device known as VME “crate” (which, through a “crate controller” module also within the crate, ultimately provides their connection to the data acquisition computers). Other devices also call the VME crate their home, including four “CROC interface modules” (CRIMs), bridging up to four CROCs each to the crate controller, and a “MINERνA Timing Module” (MTM), which is directly connected to the Main Injector’s timing distribution system. Timing signals thus arrive via the MTM, are distributed to

the CRIMs, fan out from them to the CROCs, and ultimately run through the FEB loop. Since a crate can service up to eight CROCs in addition to the crate controller and two CRIMs, MINER ν A requires two crates to address all of the front-end electronics. [64]

In addition to distributing timing and programming information, the FEB daisy-chain system also serves as the conduit for the extraction of pulse height data collected during data taking gates. At the end of a gate, each CROC polls all of the FEBs in each of its chains sequentially. The pipelined and buffered data (as described above in section 4.3.2) are aggregated by the CROC and then transmitted to the data acquisition computer responsible for that crate (see next section).

4.3.4. Data acquisition computers

MINER ν A employed a cluster of three data acquisition computers during the running period corresponding to the data used in this analysis: two “slave” machines, each of which was responsible for communicating directly with one VME crate; and one “master” machine, which was responsible for synchronizing readouts, assembling the data from both slaves together into a single data stream, and serving as the interface to the user via run control software. In addition, a fourth “slave” machine was used to duplicate the data stream for the purposes of online monitoring. All of these computers run custom DAQ software based on top of the Event Transfer package (version 9.0), based on CODA [67]; this software originates from Jefferson Laboratory.

5 Data collection and simulation

The data used in this analysis were collected and calibrated using a number of techniques designed to maximize detector live-time and efficiency. The models to which the result are to be compared were used to generate simulated data using a Monte Carlo technique. Though the simulated data are used in a manner as similar as possible to that collected using the real detector, several steps differ; the steps unique to the real data collection and simulation are described in secs. 5.1 and 5.2, respectively, and the calibrations that are applied to both, irrespective of their origin, are then elaborated in sec. 5.3.

5.1. Data acquisition and unpacking

5.1.1. DAQ, run control, operator procedures

During the run period corresponding to this analysis, calibration read-outs—high-voltage pedestal measurements and light from calibrated LEDs directly injected into the clear fibers—were interspersed in the time between Main Injector spills. (The calibration done with these measurements will be discussed further in sec. 5.3.) Gates were recorded in blocks of order 1000 consisting of the same trigger pattern (NuMI spills alternating with pedestals, NuMI spills alternating with light injection, or NuMI spills only), called “sub-runs;” sub-runs of varying trigger pattern but identical detector conditions were further grouped into “runs” lasting up to a maximum of roughly 24 hours. When no errors requiring operator or expert intervention were encountered in the read-out hardware or network apparatus, the

supervisory run control software maintained continuous operation with an average live-time of over 99% compared to the time the accelerator was delivering protons to the target.

5.1.2. On-line monitoring

Whenever the accelerator is delivering beam, an operator oversees the collection of data by monitoring certain fundamental and reconstructed quantities calculated from the secondary on-line data stream noted in sec. 4.3.4. Technically, the information the operator monitors is not exactly real-time because, as discussed in sec. 5.1.3, a significant amount of unpacking and, in some cases, preliminary calibration must be applied to the data read out by the DAQ before it is in a form digestible by a human. This processing delays arrival of on-line information from a few minutes to several hours depending on the nature of the information in question.

The general data flow is as follows. The raw data stream from the DAQ is first duplicated to a subsidiary server also housed in the detector hall. There, a process unpacks it as it is received (see sec. 5.1.3) and performs the very lightweight pedestal suppression process (sec. 5.1.4), followed by the generation of a handful of low-level histograms (charge distributions read off the PMTs for each gate, distributions of hit times relative to the gate start, the deviations of PMT voltages from their setpoints, the time the detector spent reading out after each gate, etc.) which are then fed back to the operator through a shared disk system. Because the unpacking procedure is too slow to keep pace with the rate of data acquisition, a pre-scaling is applied such that only 20-30% of incoming gates are represented in these histograms. Typical end-to-end delay for this chain between data acquisition and histogram update is a few minutes.

Further processing begins after the data acquisition for an entire sub-run has

completed (sec. 5.1.1). The compressed data file is distributed to a server farm consisting of four workers where, in addition to unpacking every gate with no prescaling, the pedestal suppression and full calibration chain (sec. 5.3) are applied. Following this is basic event reconstruction (sec. 6.1). Finally, a number of higher-level diagnostic histograms are generated (e.g., number of candidate rock muon tracks per proton on target, number of time bunches in each gate, etc.); these are relayed back to the operator. Normally this process requires several hours altogether for each sub-run. Operators review these plots and make initial judgments about a sub-run’s data quality, which inform the choice of canonical data set for analysis later on.

5.1.3. Unpacking

Data arrives from the DAQ in a continuous stream of frames which correspond to the FPGA register settings, discriminator content, and digitized ADC counts in each channel. (Note that the FPGA registers are read out only once per gate, rather than once per discriminator push, since the information in them does not change during the gate.) Unpacking them involves first separating the stream into separate gates; then, within each gate, the discriminator and ADC information are collected together to form so-called “raw hits,” where each raw hit corresponds to a single charge bucket on a single channel of one FEB’s digitized charge.¹ The unpacking procedure opens a new data stream, divided into gates this time, each of which is composed of the FPGA information and a collection of raw digits.

¹Hereafter, a “hit” will always refer conceptually to an energy deposition in a scintillator strip. Types of hits are distinguished by the amount of calibration to which they have been subjected (“raw” vs. “calibrated”), their location in the detector (“ID” vs “OD”), and their provenance (“MC hits” thus hailing from the Monte Carlo simulation).

5.1.4. Pedestal suppression

The analog-to-digital converters (ADCs) described in sec. 4.3.2 report a modestly stable nonzero number of “ADC counts” (the output unit of the ADC) even when no current is incoming off the PMT. This offset is known as the “ADC pedestal.” And even though the discriminator circuits described in sec. 4.3.2 are tuned to ensure that very few hits are made of zero-content readouts, the ADC counter can occasionally fluctuate to large enough values that a “pedestal hit” is created.

The *post facto* removal of such pedestal activity from the data stream requires dedicated measurements of the pedestal offsets for each gain range of each channel, since the pedestal value is a function of environmental factors like temperature and can therefore drift over time. As was presaged in sec. 5.1.1, dedicated readouts in between neutrino gates during normal running—as well as dedicated subruns when the neutrino beam is not operating—are collected for this purpose. Since real activity not induced by the neutrino beam (e.g., cosmic rays, residual radioactivity in the detector materials or detector hall, etc.) can be captured by the detector during the pedestal readout, we apply a method based on Peirce’s Criterion [68] to remove outlier gates from the distribution. A sample outlier pedestal measurement, consisting of the measured pedestal values for a single gate, is shown in fig. 5.1; notice the inconsistent channel roughly 100 ADC counts above the pedestal level.

After scrubbing the measurements of outlier gates, tables of characteristic pedestal values (consisting of the mean pedestal value and the standard deviation of the distribution) of each gain range for each channel are created (for intervals of roughly one day) and stored in a database. These tables are then retrieved during the offline calibration procedure, and hits whose high gain ADC value is either below the pedestal mean value or within 3 standard deviations of the mean are removed. (The medium- and low-gain ranges are not considered during suppression

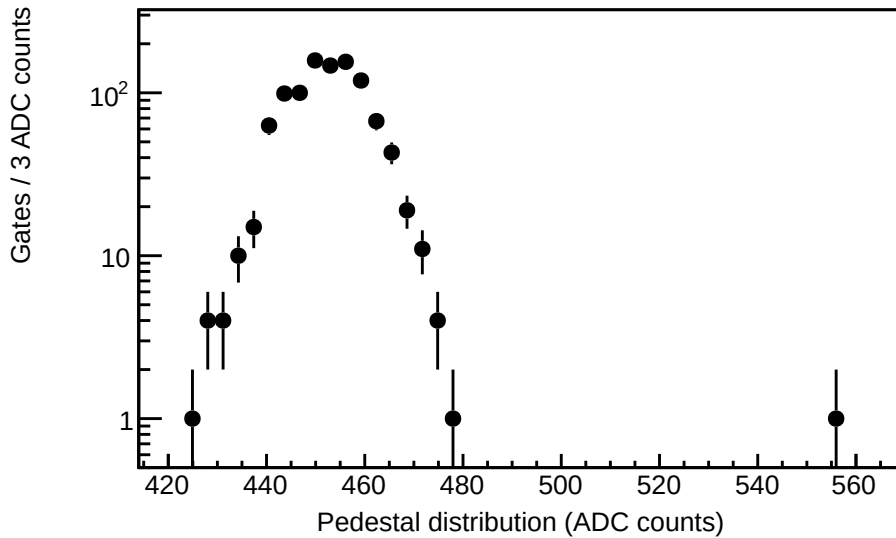


Figure 5.1: A sample pedestal measurement distribution from a gate containing an outlier. From ref. [63].

because any hit from real activity whose medium- or low-gain pipeline contains activity above threshold will also have above-threshold activity in the high-gain range.) The medium- and low-gain channel pedestal values are used separately in a quality-assurance check of the pedestal tables performed before they are inserted into the database.

5.2. Data simulation via Monte Carlo

5.2.1. Neutrino interaction simulation (GENIE)

The simulation chain begins with the GENIE neutrino interaction generator [32], of which we use version 2.6.2 for this analysis. Based on the characterization of the incoming neutrino flux described in 3.1, a model of the detector based on the description in sec. 4, and its internal neutrino interaction model, GENIE produces an ensemble of neutrino interaction “stubs” appropriate to a requested exposure of

protons on target. Loosely, each of these outgoing “stubs” (thrown randomly from GENIE’s library of available interaction types according to the distributions of their various cross-sections multiplied by the MINERvA flux prediction described in ch. 3) captures the incoming and outgoing particle content and kinematics of a reaction in which a neutrino consistent with the flux interacted with one of the particles in the detector.

GENIE’s collection of interaction models span most—but, significantly, not all—of the cross-section models popularly used by the neutrino physics community. The description of the signal reaction, the charged-current quasi-elastic reaction of electron neutrinos with carbon nuclei, is, as presented in the introduction, nearly the same as that for ν_μ CCQE on carbon. For this, GENIE uses (in the formalism of Lewellyn Smith, discussed in sec. 1.3.3) a dipole-form axial form factor with the parameter $m_A = 0.99 \text{ GeV}/c$. The vector form factors are fixed via CVC to electron scattering results, and the pseudo-scalar form factor used is that suggested by the PCAC hypothesis (on both of which again see sec. 1.3.3). The cross-section expressed by this model is then modified to adjust for the smaller cutoff in transferred momentum due to the much smaller mass of the final state lepton ($m_e \sim 0.5 \text{ MeV}$ vs. $m_\mu \sim 106 \text{ MeV}$), as discussed in ref. [33]. (A much more detailed treatment of the ν_μ CCQE model in GENIE is given in ref. [32].)

GENIE’s model of the nuclear medium in version 2.6.2 treats nucleons as quasi-independent entities under the assumption that they exist in a relativistic Fermi gas (RFG) inside the nucleus. Therefore, as presented in sec. 1.3.3, the nucleon momenta are distributed in energy levels appropriate to the nuclear binding potential up to the cutoff at the Fermi momentum p_F . GENIE’s RFG model modifies the upper end of this spectrum according to the short-range correlation effects parameterized by Bodek and Ritchie [23] (the so-called “Bodek-Ritchie tail”), which allow for momenta larger than the Fermi momentum. GENIE furthermore

attempts to model the interactions of final-state particles (as of version 2.6.2, only pions and nucleons) with other spectator nucleons as the former exit the nucleus using a third-party library called INTRANUKE (see op. cit. in ref [32]).

There are numerous background processes relevant to this analysis also modeled by GENIE. Among them are two major classes: first, non-CCQE ν_e interaction types where additional hadrons are absorbed in the nucleus or lost in the ensuing electron shower (chiefly resonant and coherent single charged pion production, though deep-inelastic scattering—DIS—contributes as well); and second, ν_μ interactions which produce a neutral pion that subsequently decays to two photons, one of which is lost or is incorrectly reconstructed (again typically resonant or coherent production). Resonant and coherent pion production (again, computed for muon neutrinos and extended via theoretical corrections to electron neutrinos) are both treated by the Rein-Sehgal model [69], while DIS is simulated according to the Bodek-Yang model [70]. The implementation of these and other processes is further elucidated in ref. [32].

5.2.2. Particle propagation simulation in MINER ν A (GEANT4)

For each event, GENIE produces only a list of outgoing particles and their associated four-momenta. The next necessary step is to determine the detector’s response to each of these particles, which we do by employing the GEANT4 simulation package noted first in sec. 3.1. GEANT performs this by stepping particles through the geometry model 0.1 mm at a time until they exit the simulated detector volume, decay, interact inelastically in such a way that they are destroyed, or have exhausted their momentum in non-destructive interactions. The interactions, decays, and energy loss experienced by a particle are governed by modular collections of physics models, denoted “physics lists.” MINER ν A’s choice of physics list for the

detector geometry simulation (as distinct from that used in g4numi, sec. 3.1) is QGSP_BERT.

In addition to the above, GEANT also records the amount of energy deposited in each geometry volume designated “active.” In the MINER ν A geometry used for this analysis, these volumes comprise the inner- and outer-detector scintillator bars, which fluoresce when ionizing particles pass through them. A mapping from active volumes to energy deposited in them is passed to the next stage of the simulation, described in sec. 5.2.3.

5.2.3. Detector read-out simulation

The final step in the simulation chain is to simulate the response of the detector to energy deposited within it. There are a handful of models which are used in series to accomplish this.

Optical model

MINER ν A’s optical model is responsible for predicting the number of photons that arrive at the photomultiplier tube for each channel based on the predicted energy deposited in each active volume of the detector from GEANT4.

As noted in sec. 5.2.2, the dopant fluor material in the detector’s scintillator strips responds to charged particle passage by emitting ultraviolet radiation according to a formula known as Birks’s Law [11, 71], which accounts for a known saturation effect at higher energy depositions. (We use the value of Birks’s constant used by the MINOS collaboration for their detectors, which use the same scintillator material: $k_B = 0.133 \pm 0.040$ mm/MeV [72]; the effect of uncertainty in this parameter is addressed in a systematic uncertainty in sec. 7.10.) The optical model begins, therefore, by calculating, for each energy deposition in an active geometry volume, the mean number of photons created by such a deposition. A

photon pulse object is then constructed to match the appropriate photo-statistics for the detector's optical response. Specifically, we record a bunch of photons with size drawn from a Poisson distribution; this distribution has mean equal to the number of photons given by Birks's Law, modified by the light yield factor computed in the global energy calibration (sec. 5.3.4), further multiplied by a value drawn from a Gaussian with mean 0 and standard deviation 0.0557 (our measured smearing). Though it is physically inaccurate to do so, we assume that all photons in the pulse were created at the same moment (that is, when GEANT reported the parent particle to have traversed the active volume); we compensate for the actual spread due to the non-negligible fluor decay times in a later step of the simulation.

Pulses so created are first smeared according to a scale factor keyed to the strip in which they were created. This scale factor is drawn from a Gaussian of mean 0 whose width was chosen to obtain agreement between the simulated and data distributions of photoelectrons after all other calibrations were applied; in this analysis, the smearing is roughly 6% ($\sigma = 0.0557$). The scale factor for a channel is chosen randomly at the beginning of a simulation job and is fixed for the remainder of the job; as one job typically corresponds to 10^{17} P.O.T., each channel samples approximately 18,900 independently-chosen scale factors throughout the entire simulation sample.

Once scaled, the pulse is then reduced in magnitude by the values measured in various calibration efforts to better model the physical behavior of the detector: the mean attenuation lengths in the wavelength-shifting fiber within the strip and the clear fibers which connect to the PMT; and the strip response factor measured in the calibration described in sec. 5.3. Thus two new average numbers of photoelectrons are constructed (one modeling a cohort of photons traveling directly from the interaction point towards the photomultiplier, and another modeling a second cohort traveling to the mirrored end of the fiber and reflected back), and are used as

the means of Poisson distribution from which new integer numbers of photons are drawn. The pulse henceforth is regarded as being the sum of these two. Similarly, we correct the pulse time for the time spent propagating in the fibers.

PMT model

The PMT model simulates the response of the photomultipliers to the light arriving at the end of the optical chain and produces a prediction of the outgoing current on their anodes.

We begin by attempting to model the phenomenon of optical cross-talk (first described in sec. 4.3.1). The simulation of optical cross-talk is based on bench measurements inherited from MINOS (described in detail in ref. [73]). The model employed in the measurement assumed that optical cross-talk photons are produced according to a Poisson distribution whose mean the tests purported to measure. Accordingly, in the MINER ν A simulation, we introduce additional pulses in each cross-talking channel in keeping with the number of optical cross-talk photons produced according to the model. That is, for each neighbor pixel, given the charge (due to physics activity) simulated in a particular illuminated pixel, we randomly sample from a Poisson distribution with mean given by the MINOS measurement's measured cross-talk strength to that neighbor. However, we find that this is not sufficient to model the cross-talk, so we use the *in situ* calibration results from sec. 5.3.7 to further tune the simulation. In this procedure, we examine the muon data on a channel-by-channel basis and compute the ratio between the cross-talk fraction as measured in the data to what was measured in the unscaled simulation. The resulting scale factor is then applied to the Poisson mean for the channel pair as returned by the MINOS measurement. This does not completely account for the difference between the distributions, as shown in fig. 5.2, but the disagreement in the mean is largely resolved (they differ by 1.5% instead of the nearly 5% before

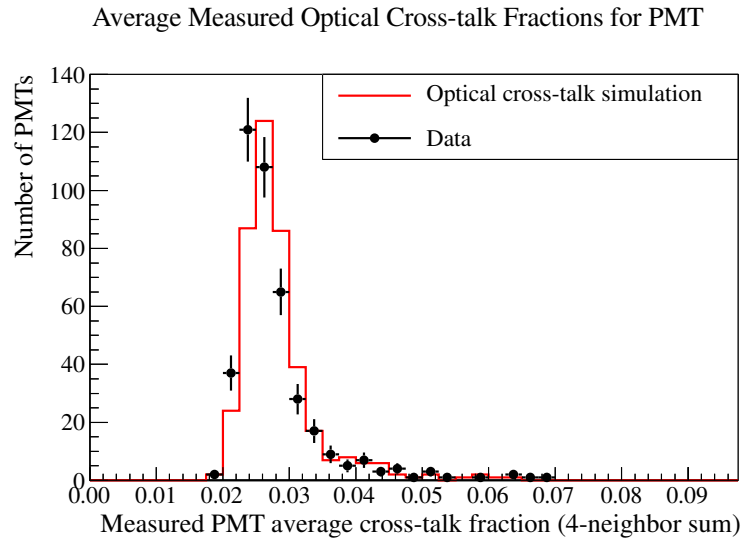
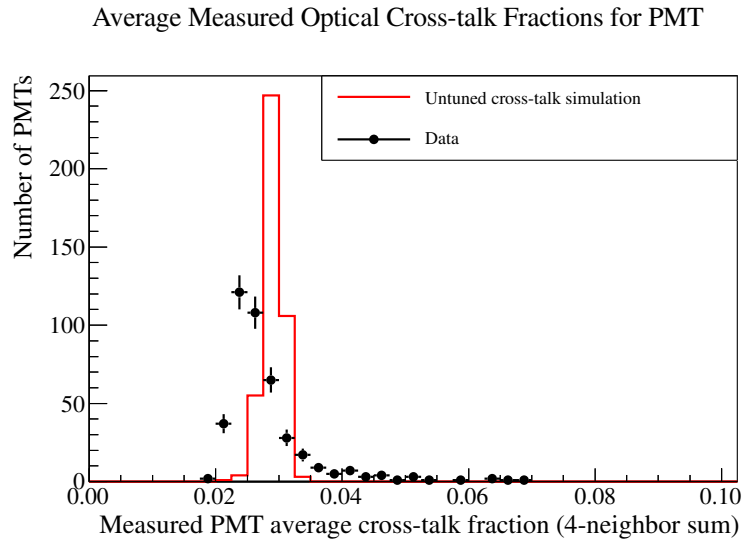


Figure 5.2: Cross-talk fractions measured via the algorithm described in sec. 5.3.7, before and after applying the calibration noted in that section.

the second tuning), and the shape matches much better in the high-side tail.

Each pulse (either from physics or from optical cross-talk) then undergoes a simulation of the amplification process. Because the details of the electron shower development in the first two stages of the dynode chain dominate the result, we simulate them individually. This is done by drawing the number of electrons in

the subsequent stage from a Poisson distribution with mean equal to the number at the current stage multiplied by a factor g_n modeling the gain due to the voltage ratios between dynodes, according to Rademacker's model [74]:

$$g_n = g_0 \left(\frac{r_n}{r_0} \right)^\alpha \quad (5.1)$$

with

$$g_0 = G^{1/N_D} \left(\frac{1}{r_0^{N_D-1}} \prod_{i=1}^{N_D-1} r_i \right)^{-\alpha/N} \quad (5.2)$$

where the r_i are the voltage ratios on the dynodes (3:2:2:1:1:1:1:1:1:2), G is the overall channel gain measured in the calibration of sec. 5.3.2, N_D is the total number of dynodes (in MINERνA, 12), and α is a value provided by the PMT manufacturer (0.75).

Once these two stages of the acceleration chain have been simulated in their entirety, the remainder are simulated together using a Gaussian approximation where the final number of electrons is drawn from a Gaussian with mean and standard deviation given by the following:

$$\mu = n_2 \prod_{i=3}^{N_D} g_i \quad (5.3)$$

$$\sigma = \mu \sqrt{n_2 \sum_{i=3}^{N_D} \left(\prod_{j=3}^i g_j \right)^{-1}} \quad (5.4)$$

(where n_2 is the number of electrons in the channel at stage 2, and the other parameters are as before).

Within the simulation of the dynodes, we also account for dynode-chain cross-talk (again see sec. 4.3.1) using the following model. Once again we begin with mean leakage fractions measured in MINOS PMT bench testing [73]. Their model

Dynode stage	original channel	cross-talk channel
0 (input)	P	0
1	$g_1 P$	$f P$
2	$g_1 g_2 P$	$g_2(f P) + f(g_1 P) = (g_1 + g_2)f P$
3	$g_1 g_2 g_3 P$	$g_3((g_1 + g_2)f P) + f(g_1 g_2 P) = (g_1 g_3 + g_2 g_3 + g_2 g_3)f P$
$\vdots$	$\vdots$	$\vdots$
N	$P \prod_{i=1}^N g_i$	$(g_1 g_2 \cdot \dots \cdot g_N + g_1 g_3 \cdot \dots \cdot g_N + \dots)f P$

Table 5.1: Number of electrons in each stage of the dynode-chain cross-talk model used in the simulation. P is the original number of photoelectrons illuminating the channel, the g_i are the gain ratios for the dynodes, and f is the constant leakage fraction that is to be determined according to eq. 5.6.

assumes that during the acceleration of electrons from dynode to dynode, secondary electrons created at one dynode may be ejected in such a direction that they end up in a neighboring channel. The result is then a smaller cascade in the non-illuminated channel and therefore a weak charge signal on the corresponding anode. To implement such a model within the simulation, we make two further assumptions: first, that the leakage fraction is constant from dynode to dynode²; second, that electrons leaked to an adjacent channel do not create a large enough pulse to generate significant cross-talk in the originally illuminated channel³. Under these assumptions, the number of electrons in the original channel and a cross-talk channel as the simulation progresses are illustrated in tab. 5.1.

The measured total cross-talk fraction for the channel F (from the MINOS measurements) is then the ratio of the cross-talk electrons to the original channel

²This assumption is suspect, particularly because the voltage ratios between successive dynodes are not constant. However, dynode-chain cross-talk is a small enough effect that the model used is ultimately somewhat unimportant. This one was chosen because it is simple to work with.

³Given the magnitudes of the measured cross-talk fractions, *this* assumption is certainly true within the tolerances of the measurement.

electrons at the end of the simulation, or

$$F = f \sum_{i=1}^N (g_i)^{-1} \quad (5.5)$$

which makes

$$f = F \left(\sum_{i=1}^N (g_i)^{-1} \right)^{-1}. \quad (5.6)$$

The calculated f is then used to simulate the number of cross-talk electrons for each dynode stage, for each channel.

Finally, after all the amplification stages and cross-talk simulation are complete, the resultant number of electrons is converted to a charge, at which point it is passed to the next phase of simulation.

FEB model

Simulation of the FEBs' charge digitization scheme described in sec. 4.3.2 takes place in four separate phases. In the first, simulation begins with the amplification of charge incoming from each PMT anode according to the three gain ranges of the FEB system (which amplify by factors of 1, 4, and 16 from “low” to “high” gain). After amplification, extra charge from a rudimentary simulation of cross-talk between the channels within the FEB is then inserted into the channels. Because the strength of FEB cross-talk is very low, we simulate it only in nearest- and next-to-nearest-neighbor channels; moreover, the measurements on which its scale is based were of very limited statistics. Fortunately, in nearly every case, FEB cross-talk does not ultimately produce charges of any real consequence.

Once amplified, pulse objects are next funneled into a simulation of the discriminator circuits, where simulated charge is divided into the “buckets” mentioned in sec. 4.3.2. To perform the segregation, we employ a scanning algorithm which

recursively subdivides the charges in the high gain range into “windows” in time. We begin by creating a discriminator window object which corresponds to the entire gate. Following this, we iteratively examine simulated PMT anode charge objects, sequentially in time order, until all of them are dispatched, using the following procedure. For each charge, we first add it to whichever discriminator window contains it in time. Then we ask if this charge, added together with all the other charge contained in the window, crosses the discriminator threshold; if it does, then the window is marked as corresponding to a discriminator having “fired,” and all this charge together is considered to have occurred at the time when the discriminator threshold was first surpassed (which mirrors the behavior of the hardware). In this case the end of the discriminator window is moved to occur at the end of the “push time” (the discriminator latch time + 16 system ticks, as described in 4.3.2), a window for “dead time” is created, and a new readout window is created spanning from the end of the dead time window until the beginning of the next window (or the end of the gate if there are no more windows). If, on the other hand, the charge falls in a window which already has its “discriminator fired” flag set, the charge is merely added to the window, though in this case its full timing information is retained (since it occurred during a dedicated readout window). Finally, charge that is assigned to a “dead time” window is masked from further propagation through the simulation chain.

The next stage in the FEB model is for these now-discretized charge buckets to be collected together and summed. Thus, each channel’s discriminator window objects and their charge are considered, and summed pulse objects are formed with time information based on the discriminator flag status: in windows for which the discriminator fired, the hit time is equal to the discriminator fire time (or precise later time within the discriminator window); and in windows for which the discriminator did not fire, the hit time is equal to the end of the discriminator window (which, in the case of the last discriminator window, is the end of the

gate). Once again this is consistent with the means by which times are assigned to hits in the real detector data stream. The resultant containers correspond to the analog quantity of charge pushed out of the FEB pipelines.

Finally, “hit” objects (in the sense of sec. 5.1.3) are formed by simulating the analog-to-digital converters of the FEBs. As onboard a real FEB, the incoming analog charge is “amplified” once more—though in this case, the gain factor actually *reduces* the charge by a factor of $1/16$ —and then added to a value drawn from a Gaussian which models the ADC’s baseline pedestal (see sec. 5.1.4). (The parameters for this Gaussian are chosen by examining the mean and RMS of the measured pedestal distributions in data; typical choices are $\mu = 440.0$ and $\sigma = 8.0$.) Hereafter, however, the model diverges somewhat from the actual workings of the hardware: instead of creating an ADC-digitized charge for every channel served by a discriminator that latched (and then deferring pedestal suppression to a later step), as is done with data (on which again see sec. 5.1.4), in the simulation we immediately remove all charges except those which, when digitized, lie above a so-called “sparsification threshold” of $3\sigma_{\text{ADC}}$ designed to mimic the pedestal suppression process. This method significantly reduces the size of the output data stream. The timing information propagated from the discriminator simulation is then attached, and the charge is assigned a “hit number” corresponding to the discriminator pipeline push sequence. All of this information together—digitized charge and discriminator status, along with timing information—comprises so-called “raw hits” which correspond to the data read out of the DAQ on the real detector. Subsequently, all of the calibrations described in sec. 5.3 below are applied to the raw hits to produce calibrated hits which can be directly compared to those from the data.

Overlay of data onto simulation

Thus far we have adopted a more or less “first principles” approach to the simulation of interactions in the detector. But there are a handful of phenomena remaining which impact the data stream in ways that make them difficult to simulate with this *a priori* technique. First and foremost among these is the nature of the beam as containing bunches of neutrinos which must traverse several hundred meters of material between production and the detector itself; the resultant copious flux of particles which enter the detector from the outside—as opposed to the products of a neutrino interaction within the detector—requires vast computing resources to simulate in sufficient detail. But the effect of externally-produced particles (particularly muons produced in the rock upstream of the detector) cannot be ignored, because they are a primary source of the dead time described in sec. 4.3.2. Other incidental activity like radioactive decays within the detector materials or the occasional cosmic ray passage can be mostly negligible but may provide backgrounds for certain classes of events. To cope with these challenges in a time- and labor-effective manner, we have chosen to augment our simulation by overlaying each simulated neutrino event with a gate drawn from the data.

Implementation of this overlaying procedure is fairly straightforward. After simulation of a neutrino event completes, we select fully calibrated data gates from the same running period as the simulation is intended to model. To conserve disk space, we choose hits from the data with time no less than 50 ns before and no more than 200 ns after the extent of the hits from the simulated event. After inserting the data hits into the stream, we retroactively mask hits which fall into a dead-time window (induced either by the simulation or the overlaid data as appropriate), irrespective of their origin in data or simulation, using the same principle as in the simulation of dead time described above.

At this point a simulated event is ready for reconstruction.

5.3. Data calibration

A series of calibrations is necessary to convert the digitized measurement of charge read off the anode of a channel on a PMT to a measurement of energy actually deposited in a scintillator strip. This section describes each calibration in the sequence, along with the procedure for obtaining the relevant constants. The full calibration chain primarily serves to convert ADC counts from the FEB (for digits surviving the pedestal suppression step of sec. 5.1.4) into an estimate of the energy deposited in a scintillator strip, according to the following rule:

$$E = Q_i(ADC) \times \frac{1}{g_i(t)} \times \exp\left(-\frac{l_i}{\lambda}\right) \times \eta_i \times S_i(t) \times C(t) \quad (5.7)$$

with:

- Q_i The FEB's ADC response function, described in "FEB calibration";
- $g_i(t)$ The value of channel i 's gain measured closest to the time of the event t , as described in "PMT gain calibration";
- l_i Channel i 's clear optical fiber length;
- λ The clear optical fiber attenuation length (measured *ex situ* [63]);
- η_i The attenuation factor within the scintillator strip, described in "strip response calibrations";
- $S_i(t)$ The value of strip i 's relative response measured closest to event time t , as described in "strip response calibrations";
- $C(t)$ The value of the overall energy scale constant closest to event time t , described in "global energy scale calibration".

Two other calibrations apply to the raw data (a calibration of hit times, and a calibration for the cross-talk subtraction procedure), as well as a calibration which

determines the actual installed hardware alignment, which are also described here. Though the FEB and PMT gain calibrations are applied only to the data (for technical reasons discussed in sec. 5.2.3 above), all of the remaining calibrations apply equally to the data and the simulation samples.

5.3.1. FEB calibration

The individual responses of each FEB’s ADC chip to incoming analog charge from the discriminator pipelines vary somewhat, thus requiring a calibration to anchor ADC response (which is what is actually reported by the FEB) to the anode charge output by the PMT. Therefore, before installing them on the detector, we perform an *ex situ* calibration on each FEB in which a known amount of charge is injected into the FEB and the response in each of the three gain channels is fit to a three-piece linear spline function (called a “trilinear function” in MINER ν A jargon); the result is Q_i in eq. 5.7. We store the output of these fits (whose free parameters include the slopes and intercepts of the three line segments and the ADC values where they intersect) in a database for use during offline reconstruction of data.

5.3.2. PMT gain calibration

The gain function of a photomultiplier is influenced by environmental factors such as temperature in addition to evolving over time as the PMT ages. Therefore, MINER ν A collects dedicated calibration data on the PMT gains on a daily basis. (Because temperature in the detector hall is well controlled, and the time dependence of the gains on age is characterized by a scale of weeks to months, this is a sufficient granularity to correct for the effects we observe.)

We employ an *in situ* technique to measure the gains of the PMTs installed on the detector. Light from 23 blue (472 nm) LEDs is fed into an optical distribution

system which delivers light from two independent fibers to each PMT box. Because pulses and the subsequent readout are short—less than 1s—they can be interleaved between beam spills, as noted in sec. 5.1.1; this allows us to collect as many LED flashes as we require for the calibration to be successful. Now, the light injected into the optical fibers that couple to the LEDs is not uniform from fiber to fiber, so we require a method of calibration that is insensitive to the number of photons collected during calibration readouts. Rademacker’s model for PMT gains [74] admits such a procedure, assuming that the rate of photon incidence is constant in time: the gain g is given by

$$g = \frac{\sigma_Q^2 - \sigma_p^2}{\overline{Q}e(1 + w(g)^2)} \quad (5.8)$$

where $\overline{Q}$ and σ_Q are the mean and standard deviations of the measured anode charges, respectively, σ_p is the standard deviation of the ADC pedestal, e is the electron charge, and the function $w(g)$ is defined such that

$$w^2 = \sum_{j=1}^{n_{\text{dynodes}}} \left(\prod_{i=1}^j \frac{1}{g_i} \right) \quad (5.9)$$

(with g_i from eq. 5.1). (This procedure is explained in further detail in ref. [63].) The values of g so calculated for each channel i on a daily basis yield the $g_i(t)$ in eq. (5.7).

5.3.3. Strip response calibrations

There are two calibrations applied to correct the response seen at the photodetector to an ideal absolute response value directly proportional to the energy deposited

in the strip. The first is a correction for the attenuation of the scintillator light as it propagates down the strip and then the optical fibers which couple the strip to the PMT. Owing to their varying distances from the PMTs, different strips are connected to different length fibers, and thus the correction for the clear fiber, $\exp\left(-\frac{l_i}{\lambda}\right)$ in eq. (5.7), is a function of the channel i . For the attenuation in the strip, by contrast, the information available in the initial calibration phase is not sufficient to determine the longitudinal distance from the strip end to the location where the energy was deposited, so we apply the same attenuation correction (half strip length) to each hit. After the tracks of sec. 6.1.3 have been reconstructed, we make further corrections to the actual point of traversal (using the 3D information from the track fit) for those hits that were able to be identified with a track, which is why the correction η is written as a function of channel, η_i , in eq. (5.7).

The second calibration is intended to correct for the differing response of the detector strips to a fixed amount of energy deposited in them. In particular, manufacturing and assembly effects like voids (bubbles) in the epoxy used to glue the optical fibers into the strips, minor inconsistencies in the material composition of the scintillator material, and couplings in the optical systems result in light yields that vary from strip to strip. Therefore, we construct correction factors for each strip which adjust each strip's response such that the mean correction factor is 1.0.

As the relative strip response calibration is an *in situ* calibration, we require a mechanism by which as many strips as possible are illuminated with an identical amount of light (which must furthermore be within the tolerances of the amplification system). An ideal source for this is the flux of external particles which originate from interactions of the neutrino beam with the rock upstream of the detector hall. Because muons are produced in every charged-current muon neutrino interaction, and are comparatively long-lived and interact only electroweakly, they

are the dominant externally-incident particle, and are, in fact, abundantly present in the data stream (we typically observed several per neutrino bunch during the running period of this analysis). Moreover, at energies typical to those created by the NuMI beam (several GeV), muons interact with matter primarily by ionization of atomic electrons, leaving behind energy deposits very near the minimum of their well-known energy-loss distribution (e.g., Fig. 30.1 in ref. [11]). (This behavior makes them the prototype of what is known as a “minimum-ionizing particle,” or “MIP”). Thus we can use the minimum energy deposited (per unit areal density) by these “rock muons” as they traverse the detector as the standard candle we need for the strip response calibration.

We thus begin with a sample of rock muons, which we select by examining events where the long track reconstruction of sec. 6.1.3 has found a track, and further require that the track entered from the front of the detector. Obtaining sufficient statistics in every single channel of the inner detector region typically requires of order 10^5 rock muons within the sample, which corresponds to roughly 2-4 weeks’ worth of data taking. Once this sample is assembled, the energy deposited in each strip must be corrected for the estimated length of the muon’s path through that strip, because rock muons can pass through MINER ν A at different angles, and the energy deposited by a MIP depends directly on the amount of matter with which it interacts. (This requires that procedurally, we must finalize the alignment correction of sec. 5.3.5 first, to ensure that the path length estimate is reliable.) The reconstructed track parameters are used to perform this correction. The peak of the resultant energy distribution for each channel is then fitted with an iterative truncated mean technique, and the final constant for channel i , S_i in eq. (5.7), is calculated from the following formula:

$$S_i = \frac{(\overline{E}_i)^{-1}}{\frac{1}{N} \sum_j (\overline{E}_j)^{-1}} \quad (5.10)$$

where $\overline{E}_i$ is the peak (truncated mean) energy observed in strip i . This procedure is repeated roughly monthly to follow the time dependence of the constants, yielding $S_i(t)$.

As before, much more detail on this procedure is given in ref. [63].

5.3.4. Global energy scale calibration

The final step in the energy calibration of eq. 5.7 is the determination of $C(t)$, the global energy scale, which relates the best estimate of light in a channel (after all the previous corrections) to an absolute energy in physical units.

For this calibration we once again make use of the rock muon sample described in sec. 5.3.3. This time, however, we also require a corresponding simulation sample for comparison. The latter is generated by using only muons from the data rock muon sample which pass into the MINOS detector, where they can be charge- and momentum-analyzed. After adding back an estimate of the energy lost within the MINER ν A detector during their traversal, the momentum vectors of this sample of rock muons can be used as a seed for the simulation, resulting in a set of simulated muons with similar kinematics to the rock muon sample.

We apply the basic reconstruction techniques described in sec. 6.1 to both samples, obtaining spatially contiguous clusters of illuminated strips (sec. 6.1.2) which can be fit to straight-line tracks (sec. 6.1.3). For muons in which the tracking succeeds (which is to say, muons for which exactly one track object was reconstructed, and for which the track spans the detector and matches to a track reconstructed in the MINOS detector), we proceed to compare the energy distributions of clusters calculated with a trial C . In particular, we fit both the peak regions of each distribution with a fifth-order polynomial. (Disagreements between these quantities are postulated to be due to the optical response of the scintillator—the “light level”—and are corrected by the tuning of a light yield

parameter which is fed back into the simulation such that the fitted peak of both distributions agree.) The value of the simulated reconstructed cluster energy distribution is plotted against the distribution of the true energy deposited in the illuminated strips, and fitted to a line; the slope of this line is the global energy scale C (see fig. 5.3).

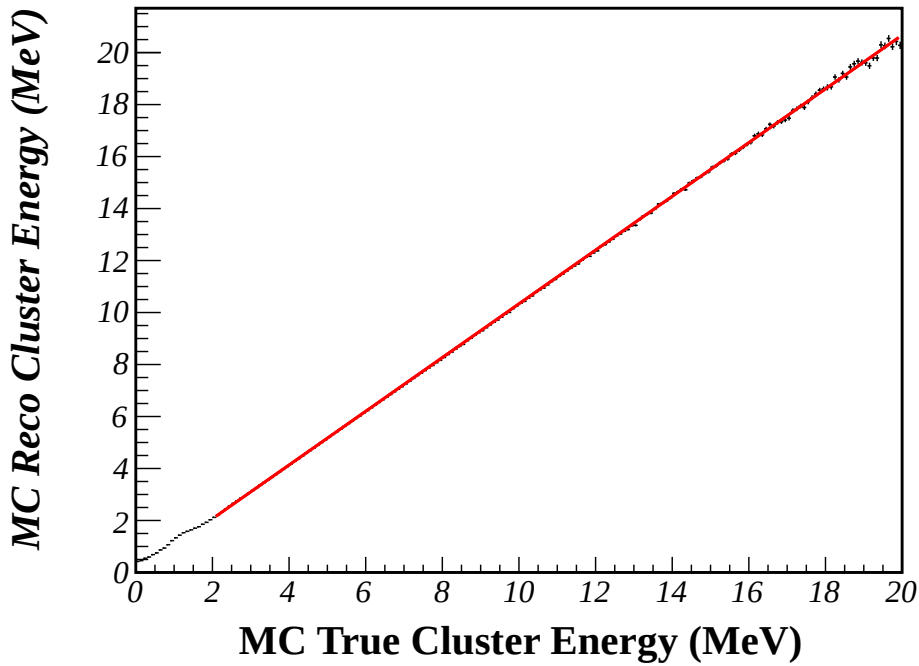


Figure 5.3: Reconstructed cluster energies vs. true cluster energies, fitted to a line. Courtesy A. Mislivec (MINER ν A).

Because the light output of scintillator is known to degrade as a function of time [75], we must repeat the global energy scale calibration periodically to account for it, making C actually $C(t)$.

As previously, further detail on this calibration beyond what is given in this section is described in full detail in ref. [63].

5.3.5. Alignment calibration

Variations in the positioning of the detector elements relative to the detector axis are inevitably introduced during the installation process. We expect the largest variance to be between the positioning of the individual planes, which are pre-assembled and then installed using a crane. We apply a calibration procedure to correct for two types of plane misalignments: transverse variation (perpendicular to a plane's strip direction) and rotation. Between these we can completely characterize any misorientation of a detector plane in the $x - y$ coordinate space.

To compute the corrections necessary, we once again turn to the rock muon sample of sec. 5.3.3. By plotting the energy observed in a strip as a function of its reconstructed transverse position, we are able to fit for the actual location of the triangle peak. The displacement of this fitted peak from its expected position is identically the transverse correction. Similarly, the rotation correction is deduced by performing the above correction in bins along the longitudinal strip direction. Since a strip that is perfectly aligned with the expectation will produce an identical shift in all the bins, a rotation can be determined by fitting the calculated shift as a function of longitudinal position. We average these over all the strips in a plane to obtain the corrections for the plane. Illustrations of both of these fits are given in fig. 5.4. As usual, further details are given in ref. [63].

5.3.6. Timing calibration

Light generated in a scintillator strip in the detector must traverse several meters of optical fiber (both within the strip and without) before it reaches a photosensor. Since the amount of optical fiber between the strip end and the corresponding PMT varies from channel to channel, we must account for this delay as part of the calibration chain. In addition, the daisy-chain system used to connect the detector FEBs in serial (described in sec. 4.3.3) introduces a measurable offset due to the

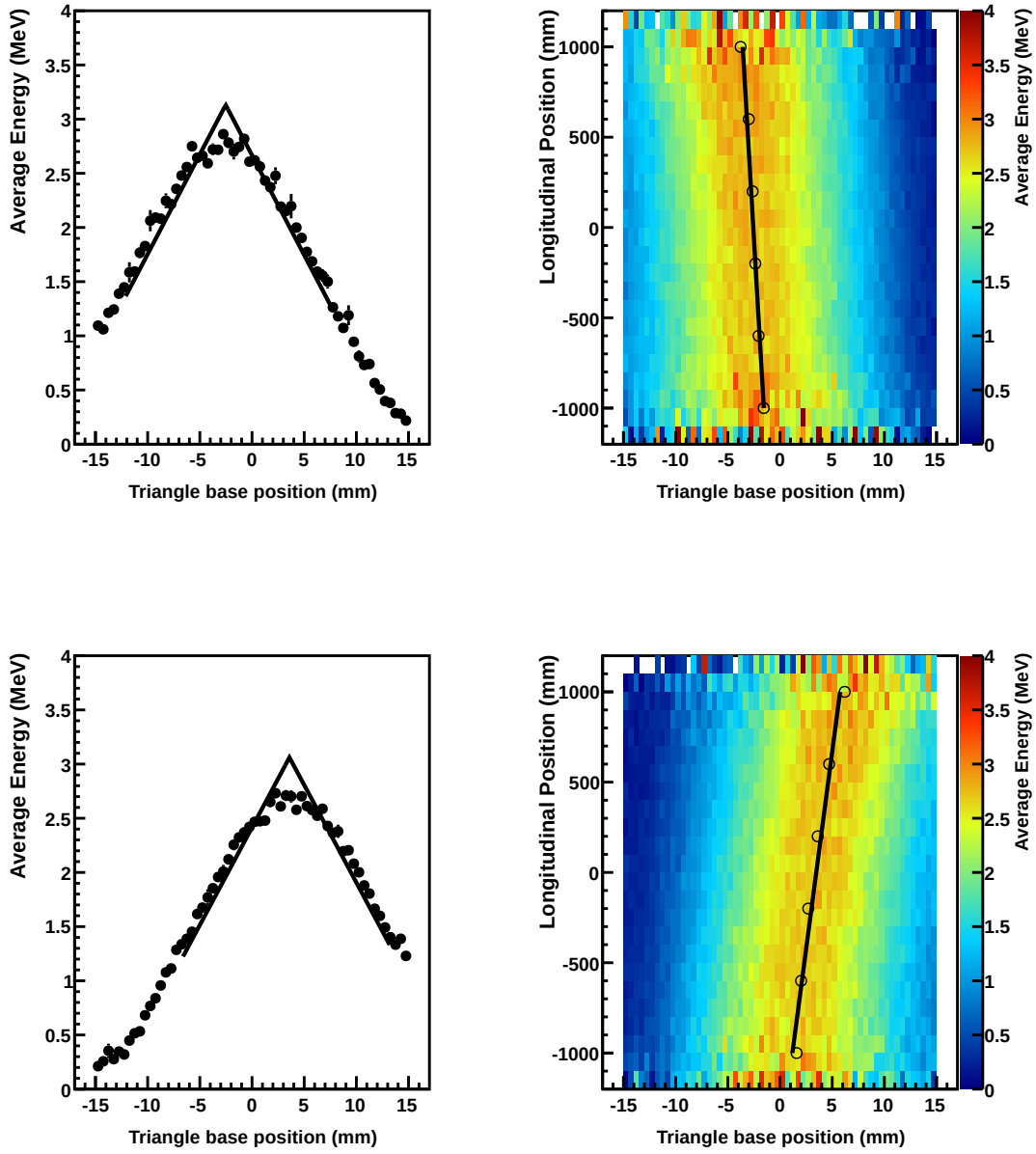


Figure 5.4: Sample fits to the plane displacement and rotation for two different planes: at left, reconstructed energy as a function of displacement along expected triangle base position (averaged over many muons); at right, reconstructed energy as a function of both triangle base position and longitudinal position. In both cases, fits are illustrated with solid lines. From ref. [63].

transport time of information through the electronics. Finally, the relaxation time of the fluoror compound used to dope the scintillator strip polystyrene (making

it scintillate in the presence of ionizing radiation as particles pass through) is of the order of nanoseconds, which is comparable to the timing resolution of the electronics, so again, we must compensate for it. (The non-linear time response of the scintillator as a function of incident light is often known as “time slewing,” jargon which will be adopted throughout the remainder of this section.)

Transport time in the fiber is corrected based on the measured index of refraction of the fibers (from which the speed of light therein can be calculated) and the measured lengths of the fibers. Taking this adjustment into account, we compute the remaining corrections using an iterative procedure which takes as input the distributions of the calibrated number of photoelectrons estimated in strips (see secs. 5.3.2-5.3.4) and the corresponding times in rock muon events. In these events, the distribution of hit times (corrected for muon time-of-flight and fiber transport) for hits occurring in a particular FEB—effectively a distribution of the time slewing as a function of pulse height—is fitted to a third-order polynomial in $1/\sqrt{\text{hit PE}}$. These distributions are then compared and provisional offsets are calculated to make them agree. In successive iterations, the offsets are applied as corrections and fine-tuned. The calibrated offsets necessary are in some cases as large as 30 ns.

5.3.7. Cross-talk calibration

A number of processes can cause a signal in one PMT channel to produce a signal in another channel. These are collectively known as “cross-talk”. While these can, in principle, be differentiated by tests on the bench, once the detector components are assembled and installed it is virtually impossible to separate them from one another with any significant confidence, particularly at large pulse-heights. The dominant types of cross-talk in MINERνA, as noted in sec. 4.3.1, are optical (fiber-to-PMT coupling) and PMT internal (dynode chain).

The ideal probe for a measurement of either type of cross-talk in the detector

would be one in which individual pixels were illuminated with a light pulse of well-known energy. Unfortunately, once a PMT is mounted on the detector there is no system available to MINER ν A that can accomplish this goal. (In particular, the light injection system used for the gain calibration of sec. 5.3.2 cannot be used because it illuminates all of the pixels of a PMT simultaneously.) The next best option is to use data generated by neutrino interactions. For this measurement, we once again sample from the rock muon data set.

As in the preceding sections, we rely on basic reconstruction of rock muons into time slices and spatial tracks. Hits falling within a rock muon’s time slice are then classified as signal or noise based on whether or not they have been associated to the track by the track reconstruction software, since the fiber weave described in sec. 4.3.1 displaces cross-talk hits far enough away from the track that they are not associated to it. Cross-talk hits can be further distinguished from other noise by assuming that such hits occur in the same PMT as on-track activity; we then assume each cross-talk candidate hit to be associated with the on-track hit that is nearest to it on the PMT pixel grid. A sketch of how this process would work for a typical muon event is shown in fig. 5.5.

Once hits have been identified as signal or cross-talk and the rest discarded, an average cross-talk fraction f_{xt} for the PMT is defined as the ratio of the summed energy of cross-talk hits associated to that PMT to the energy of the on-track hits. Various permutations of this metric are also used (most notably the “nearest-neighbor” pixel cross-talk $f_{xt,NN}$ average for each PMT, because the strongest cross-talking pixels are generally nearest-neighbors).

We use the calculated values of $f_{xt,NN}$ to provide the tuning for the simulation of cross-talk described in sec. 5.2.3. Furthermore, we attempt to subtract cross-talk from both simulated and real data events based on a model where the null hypothesis corresponds to cross-talk for a given channel based on $f_{xt,NN}$. The

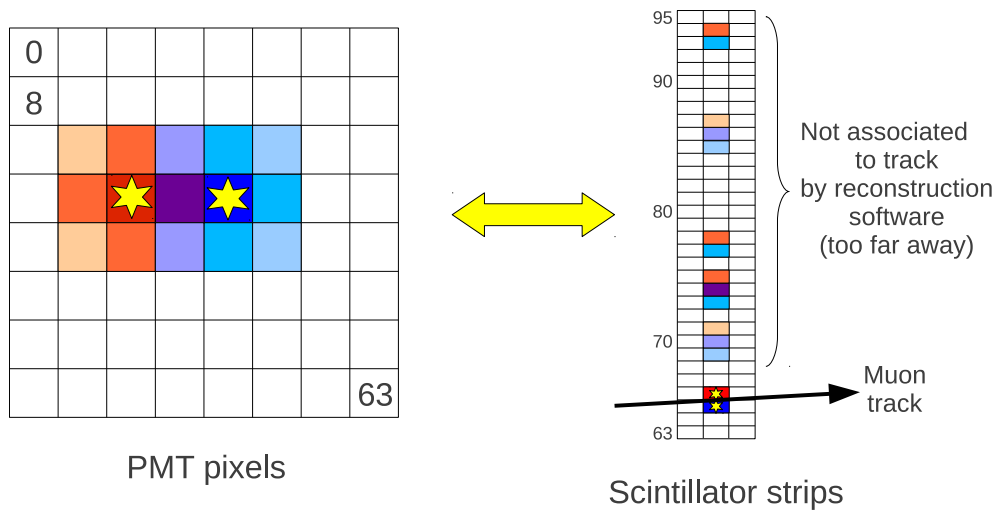


Figure 5.5: Schematic depiction of how cross-talk on the PMT face maps to scintillator strips. The darkest blue and red (stars on the PMT diagram; strips 65 and 66 in the scintillator sketch) are the original signal from a muon track; the cross-talk energy is colored according to which original signal hit it will be associated with by the algorithm described in the text (darker means stronger cross-talk). Purple represents cross-talk that will be associated to either hit at random since it is ambiguous.

p -values calculated for hits in 2500 simulated neutrino events as a function of the true cross-talk fraction of those hits is shown in fig. 5.6. We tag any hits with $p > 0.0001$ (i.e., $\log_{10}(p) > -4$) as cross-talk candidates; they are not used by the reconstruction algorithms of ch. 6.

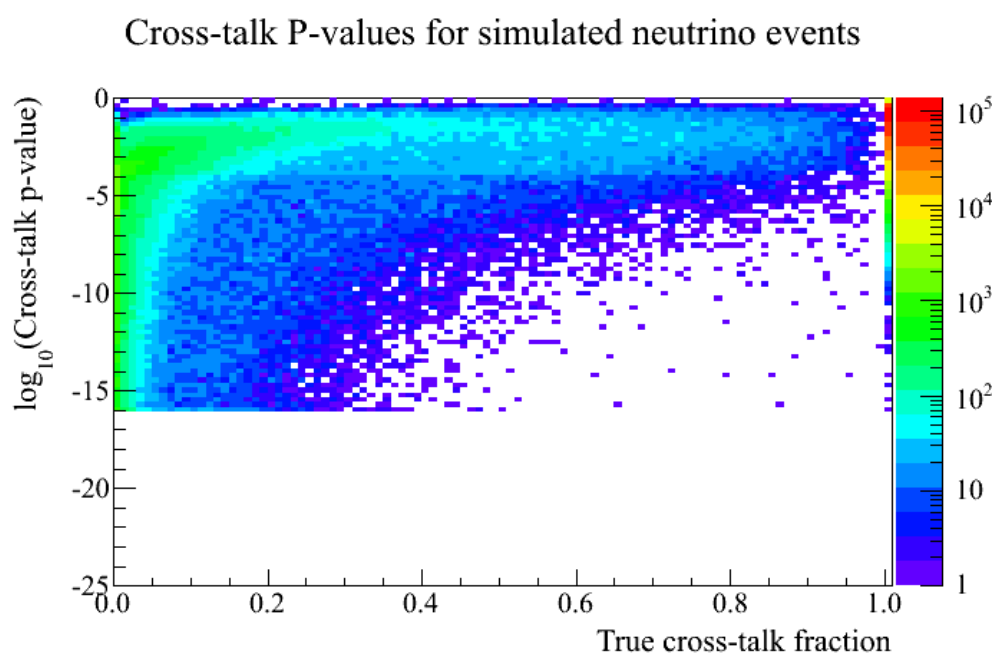


Figure 5.6: Cross-talk p -values for hits in simulated neutrino events as a function of true cross-talk fraction, where the p -value is calculated as described in the text.

6 ν_e CCQE event reconstruction

Reconstruction of neutrino events begins with the fully calibrated (strip, energy, time) multiplets which are the output of the lengthy calibration procedure of sec. 5.3. These are commonly referred to as “Digits” (short for “digitizations”), jargon which we will employ below. Digits are grouped first by time, then further by spatial proximity, and finally, into higher-level reconstructed objects like particle-trajectory tracks.

6.1. Generic reconstruction¹

Event reconstruction begins with a number of algorithms that perform generic reconstruction tasks designed to help isolate neutrino interactions irrespective of the reaction channel.

6.1.1. Time bunch separation

Because the MINER ν A detector’s time resolution is very good (~ 3 ns) compared to the mean time between neutrino interactions in MINER ν A in this data set (typically hundreds of ns), we begin isolating neutrino interactions by grouping Digits into so-called “Time Slices” according to their reconstructed time. The peak-finding algorithm we employ to do the grouping is based on a scanning method in which Digits are first sorted by time and then examined in sequential order. In particular, we require Time Slices to contain a minimum of roughly

¹In this discussion, the names of specific reconstructed objects—e.g., “Cluster,” “Track,” etc.—will be capitalized so as to distinguish them from the generic usage of the terms.

2 MeV worth of energy (on the order of 70% of the expected energy deposited by a minimum-ionizing particle in one scintillator plane of MINER ν A) by searching for 80 ns “windows” of Digits which cumulatively surpass this energy threshold. If any such windows overlap, they are then merged, subject to the constraint that the first 30 ns and last 50 ns of the merged distribution are excluded from the resultant Time Slice (which prevents overlap of close Time Slices with a true valley between them).

Digits whose discriminators did not latch are excluded from the creation phase of Time Slice construction since their times are unreliable (see sec. 4.3.2). However, because such a Digit can stem from a read-out cycle induced by activity in another channel within the same TRiP-t whose discriminator *did* latch, when we find an undiscriminated Digit that falls within the same TRiP-t and read-out cycle as a discriminated Digit, we attach such Digits to the Time Slice of the Digit which exceeded the threshold.

6.1.2. Spatial clustering

Once collected into Time Slices, we proceed to group Digits sharing a slice into groupings known as “Clusters,” collecting them by spatial proximity. Clusters so formed are first of all required to be contained entirely within a single detector plane (in the inner detector region) or module (in the outer detector region). Further, Digits within a Cluster always must be immediately adjacent; no transverse gaps between Clusters are allowed. In the OD, this requirement implies that Clusters always consist of either a single Digit or the two neighboring scintillator bars in a single story. For ID Clusters, we additionally classify Clusters by the following patterns:

Cross-talk candidate Cluster A single Digit which was identified by the cross-talk rejection scheme (see sec. 5.3.7) as likely cross-talk. (Cross-talk Digits are

therefore disallowed from being used in any other types of Clusters.)

Low-activity Cluster Any collection of a nonzero number of Digits whose energies sum to less than 1 MeV.

Trackable Cluster A collection of Digits whose energies sum to between 1 and 12 MeV, so long as:

- (a) there are four or fewer Digits;
- (b) at least one Digit contains 0.5 MeV or more; and
- (c) any Digits with more than 0.5 MeV of energy are contiguous.

Heavily-ionizing Cluster Any otherwise Trackable Cluster which exceeds the 12 MeV cutoff for total Digit energy.

Super-Cluster Any Cluster containing five or more Digits; also, any Cluster which is not Low-activity and does not satisfy the requirements for Trackable or Heavily-ionizing Clusters.

The clustering procedure’s primary purpose is to ease the recognition of trajectories, and these definitions reflect that aim. In particular, Trackable Clusters are intended to correspond primarily to the energy deposition pattern of minimum-ionizing particles as they traverse the inner detector, while Heavily-ionizing Clusters tend to collect the energy from stopping particles which do not initiate cascades. Super-Clusters and Low-activity Clusters are “catch-all” designations created to collect other activity that does not fit these patterns; most showering activity (for example) produces some of each type.

6.1.3. Track formation

We attempt to reconstruct single-trajectory objects called Tracks in the inner-detector region of the detector using a multi-stage line segment fitting-and-merging algorithm.

Track seeding

The initial stage of the track finding algorithm consists of forming track “seeds” out of triplets of Trackable Clusters which share a Time Slice. In this phase, we search for groups of three Clusters lying within three consecutive planes of the same detector view (X, U, or V; see sec. 4.1). If the Clusters within such a grouping can be fit to a line segment, we consider them a potential piece of a track, and therefore classify them together as a track “seed.” As we consider all possible triplets satisfying these criteria, a single Cluster is allowed to contribute to multiple seeds. (This does not result in double-counting since seeds will be merged together in the next step of the algorithm.)

It should be noted that the threefold combination of (first) the requirement of a straight-line fit, (second) the finite width of scintillator strips and finite plane pitch in MINER ν A, and (finally) the insistence on Trackable Clusters containing no more than four Digits, all taken together, result in an upper limit of about 70° for the angle between the detector z -axis and the seed, which propagates through to a similar constraint for the resulting Track object. Though this is a potential limitation of the tracking technique, the kinematics of NuMI neutrinos impinging upon the MINER ν A detector result in the vast majority of neutrino interaction products traveling within reasonably shallow angles to $\hat{z}$. The marginal impact that this weakness has on the analysis will be discussed in sec. 6.2.

Seed merging to form track candidates

Once we have obtained a set of seeds in each of the three detector views, we next attempt to combine them into longer trajectories. Seeds within the same view are examined pairwise for consistency in slope and overlap in z , and if a pair should satisfy those criteria, a “Track Candidate” is formed from the union of their Clusters. This Track Candidate is then compared pairwise to all the remaining

seeds within its view, absorbing any that are consistent according to the test noted above. This procedure is repeated with any unused seeds until no candidates are formed or no seeds remain.

After the seed merging step, we also attempt to merge Track Candidates together so as to provide the minimum number of inputs to the final step. Candidates are combined with very similar criteria to those for seeds; notably, however, we do not require Candidates to overlap to be merged. This allows the resulting tracks to span portions of the detector containing dead strips.

Candidate merging to form three-dimensional tracks

So far we have three independent sets of two-dimensional Track Candidates separated by detector view. At this point, we sequentially apply two separate strategies to combine as many of them as possible into three-dimensional trajectories in the detector.

The first strategy searches through the Candidates to find triplets which contain one from each detector view and which substantially overlap along the detector z -axis. Next, a Kalman filter-based fitter [76] attempts to fit them to a three-dimensional line, allowing for mild scattering along the track. If the fitting procedure converges, a reconstructed trajectory known as a Track is constructed out of line segments between a series of nodes, each of which corresponds to the fitted position at one of the input Clusters. (Triplets of Candidates corresponding to a subdominant topology where a significantly longer Candidate matches to two pairs of shorter Candidates in other views—corresponding to the case where a particle undergoes a scattering parallel to the orientation of one view—are found by this search as well. In their case, the longer Candidate is broken such that two sets of three shorter Candidates result. These Candidates are then formed into tracks by the procedure described above.) Because it requires at least one

Track Candidate—and therefore one track seed—in each view, and because the arrangement of views (on which see sec. 4.1) is such that there are twice as many X-view Clusters as U-view or V-view Clusters, this approach requires a minimum of 11 sequentially illuminated planes to obtain a Track.

Once all of the Tracks that can be found by the above technique have been exhausted, any Track Candidates that have not been thus used are passed to the second Track creation algorithm. In this case, we search for doublets (instead of triplets) of Candidates in differing views that substantially overlap in z . For each doublet, we then search in the remaining detector view for Clusters not part of any Track Candidate that are consistent in three-dimensional position with the three-dimensional trajectory formed from the Track Candidates. If sufficiently many are found, a Track object is made from the combination of the Track Candidate pair and these Clusters. As with the three-Candidate version, Tracks found this way are also subjected to a fitting procedure based on a Kalman filter, which is required to converge in order for the Track to be retained.

After both of these algorithms have formed as many Tracks as they can, any remaining Track Candidates are discarded and their constituent Clusters are released for use in other reconstruction algorithms.

6.2. ν_e CCQE-specific reconstruction

Broadly speaking, events which will pass the signal selection described in more detail in sec. 7.1.1 can be characterized as containing an electron or positron, no muons, and “not much else” (the precise definition of which is elaborated later). Therefore, the task of the reconstruction detailed in the following sections is primarily to locate events containing electrons and no muons. Further requirements spelled out in ch. 7 will be used to enforce the remainder of the signal selection.

6.2.1. Data reduction

We begin our reconstruction by attempting to reject events that are obvious background events, since our signal is a rare process in a beam dominated by muon neutrinos. We demand that candidate events contain at least one reconstructed Track (as constructed in sec. 6.1.3) and that no Tracks exit the back of the MINER ν A detector as muons are much more likely to do. At this stage, we remove any activity from candidate events consistent with cross-talk by eliminating any cross-talk candidate Clusters (cf. sec. 6.1.2) present in the event.

6.2.2. Electron candidate reconstruction

Cone construction

Once we have identified an event containing at least one Track (with no exiting Tracks), we attempt to determine whether any of its activity is of electromagnetic origin. The basic philosophy of this step is to construct a cone around each Track in succession and apply particle identification techniques to the collection of Clusters contained within it (including the Track).

We iterate through the Tracks and consider each in turn. To construct a cone about a Track, we need an estimate for its vertex and its main direction, which will become the cone vertex and axis. Here, we identify the vertex and axis of a Cone object with the location and direction of the fitted position and direction of the first node of the Track. We studied the effect of using various opening angles for the cone; because above 7.5° the fraction of electron energy contained in the cone did not improve substantially (fig. 6.1), we chose a cone opening angle of 7.5° .

Using the Cone just constructed, we examine all other reconstructed objects in the event and decide whether they are contained within the Cone. Any objects that are deemed inside the Cone are consolidated together with it. After every

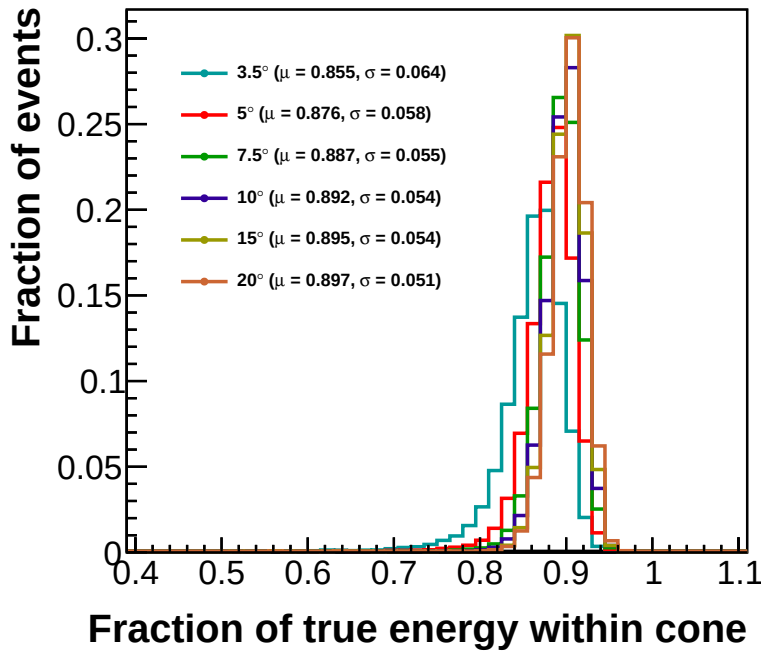


Figure 6.1: Fraction of true electron energy retained by cones of various opening angles as measured in (simulated) particle cannon samples.

other unused reconstructed object has been considered for addition to the Cone, we enforce a *post facto* limitation on the separation between energy along the Cone's (longitudinal) axis, since we do not expect electromagnetic cascades to propagate very far in the detector without depositing any energy. To do this, we first sort the Clusters within the Cone in order of their projection onto the Cone's axis. (This projection is performed in three dimensions for any Clusters which have been refit during the Tracking process, and which therefore have a three-dimensional position; for all other clusters, the projection is two-dimensional, onto the projection of the Cone's 3D axis into the Cluster's detector view.) We then calculate the distance between neighboring Clusters along the projection, in units of radiation length. If anywhere there is a gap of 3 radiation lengths or more, any Clusters beyond the gap are removed from the Cone.

At this point, we have determined the final set of Clusters which will be

considered as part of this Cone for the purposes of particle identification. Because we may have added a substantial number of clusters to the Cone since its original creation as a Track, our initial guess for the Cone axis may no longer be the best estimate for the energy we now have reconstructed. Therefore, we attempt to re-fit the Cone axis. To do so, we create a collection of single-use Clusters, each of which corresponds to all the Digits within a single plane inside the Cone; from each of these we create a track Node. Then, we use the same Kalman filter used for Track correction in the common reconstruction (sec. 6.1.3) to estimate the best Track through these Clusters. If the Kalman filter’s fit succeeds, we update the Cone’s axis vector and vertex. Furthermore, we also use the three-dimensional information from the new Cone axis to determine the longitudinal position of all the (original) Clusters within the Cone along the strips in the detector. This information is used to provide the best estimate of energy attenuation within the strips for these Clusters.

Energy reconstruction

An estimate of the incident energy of the particle represented by the Cone object of the previous section is an observable that contributes significantly to the particle identification (and will furthermore feature centrally in the cross-section analysis of ch. 7). We assign each Cone object an energy reconstructed according to the hypothesis that the Cone originated from an electromagnetic cascade.

The reconstructed energy is assembled by summing the energies of the Clusters various regions of the detector, then multiplying them by a set of scale factors that compensate in an average way for energy lost in the passive materials therein.²

²The various multiplicative factors applying to the different ECAL regions owe to the difference in geometry between the downstream ECAL region (where sheets of lead completely interleave with sheets of scintillator) and the tracker region (where a hexagonal “collar” of lead obscures parts of some strips and all of others). Their precise derivation is given in ref. [62].

Parameter	Description	value
α_{scale}	Overall scale (compensates for passive fraction of scintillator)	1.326
α_E	ECAL constant (extra passive material: lead)	2.205
α_H	HCAL constant (extra passive material: steel)	9.540

Table 6.1: Calibration constants used for calorimetric energy reconstruction.

The particulars of this calculation are laid out in eq. 6.1:

$$E_{reco}^{EM} = \alpha_{scale} (E_T + \alpha_E E_E + (2\alpha_E - 1)E_{SE}^{X \text{ view}} + (4\alpha_E - 1)E_{SE}^{UV \text{ view}} + \alpha_H E_H) \quad (6.1)$$

Here, the aggregate calibrated visible energies in the Tracker, ECAL, Side ECAL (X-view Clusters), Side ECAL (U- and V-view Clusters), and HCAL are represented by E_T , E_E , $E_{SE}^{X \text{ view}}$, $E_{SE}^{UV \text{ view}}$, and E_H , respectively. The values of the calibration constants α are shown in table 6.1. These constants were computed by minimizing the residuals between the reconstructed and true energies in a simulation of electrons in the detector volume, as shown in fig. 6.2.

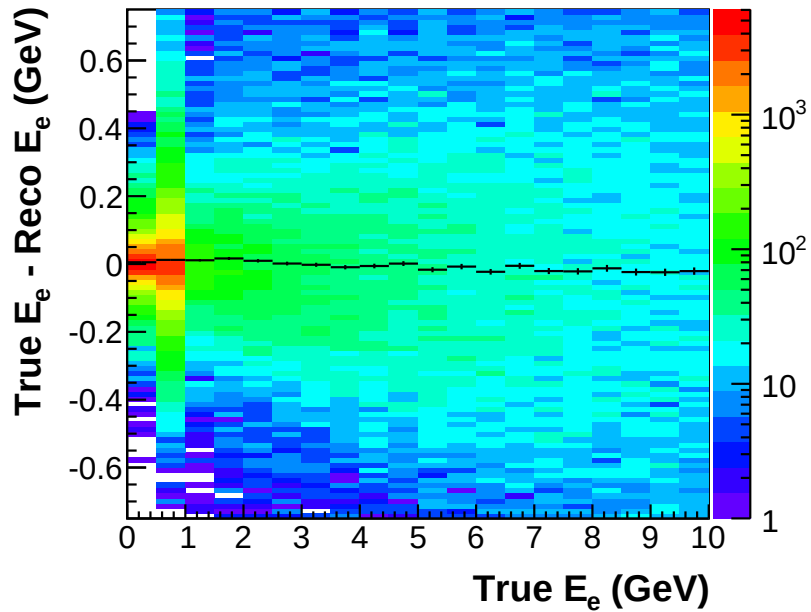


Figure 6.2: Residuals between reconstructed (eq. 6.1) and true energy for electrons in simulation. The calibration constants used to build the reconstructed energy are those given in fig. 6.1.

Electromagnetic cascade identification

The judgment of whether the energy collected in a Cone represents an electromagnetic cascade is performed by examining a multivariate analysis (MVA) classifier that combines information from three measures of its energy deposition pattern. Before describing them, however, a preliminary word about lengths is appropriate. Because MINER ν A is a heterogeneous detector, it is challenging to devise variables which compare in an unbiased way when particle tracks span multiple material types. To help (though not completely alleviate) this difficulty, where a particle identification variable uses a length as a parameter, we will usually use one of two alternatives depending on which is the more natural scale based on the problem at hand: the radiation length (on which see ref. [11]) or integrated columnar density (i.e., $\int \rho dx$ along the path), the latter of which will usually be denoted as “range.”

Secondly, we will consider the performance of these variables by examining their

separation of different particle types as simulated in the MINER ν A detector in single-particle samples (which we will refer to using the jargon “particle cannon”). We use samples of electrons and photons (both considered signal for the purposes of this section), as well as pions, protons, and muons (considered backgrounds). All particles were generated with initial momenta of 0-10 GeV, in the central part of the MINER ν A Tracker region, with angles to the z -axis of less than 45° .

Finally, we will evaluate the accuracy of the simulation to predict the response of the variables we choose by comparing distributions selected by various techniques used to select signals in other analyses within MINER ν A. For muons, we will compare rock muons to simulated front-entering muons (see sec. 5.3.3); for pions, we will use those selected by a charged pion cross-section analysis [77]; for protons, we compare protons selected in a ν_μ CCQE analysis [78]; for photons, we examine the decay products of neutral pions in a charged-current neutral pion analysis [79].

Endpoint energy fraction The first of the three variables incorporated into the MVA examines the fraction of the Cone’s energy deposited closest to the end (along the Cone axis). When most particles stop in MINER ν A primarily due to ionization-type energy loss, the characteristic signature is a fractionally large deposition of energy at the end of the particle track. Muons, protons, and pions often manifest themselves this way. An electromagnetic cascade, by contrast, is different: typically its longitudinal energy profile rises to a maximum some distance from either end of the cascade and then decays away more slowly as the cascade dies out. Therefore, we construct a measure f_{endpoint} examining how much of the Cone’s energy is deposited at the end:

1. Project the Cone’s energy into bins of range (size 10 g/cm^2) along the Cone axis.
2. Remove any Clusters off the end off the Cone in bins with less than 2 MeV.
(This reduces the algorithm sensitivity to detector noise.)

3. Determine the median bin energy (excluding the last bin) E_{median} .
4. Compute $f_{\text{endpoint}} = E_N/E_{\text{median}}$ (where E_N is the last bin energy).

The performance of the endpoint energy fraction variable is shown in fig. 6.3. Comparisons of simulation and data are shown in fig. 6.4; agreement is very good except for in the muon sample, but because there are very few muons in the signal sample, this is not significant.

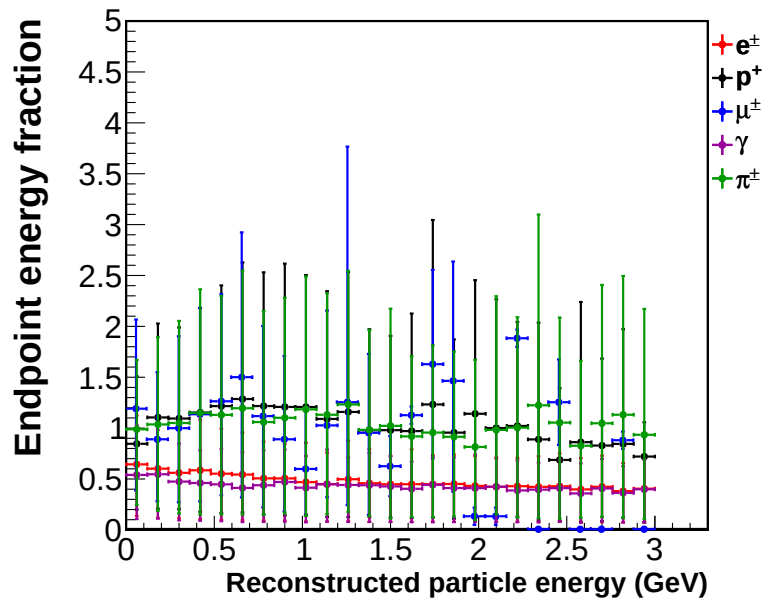


Figure 6.3: Performance of the endpoint energy fraction variable as a function of calorimetrically reconstructed energy. Error bars on each point represent the smallest interval containing 68% of the distribution in that bin.

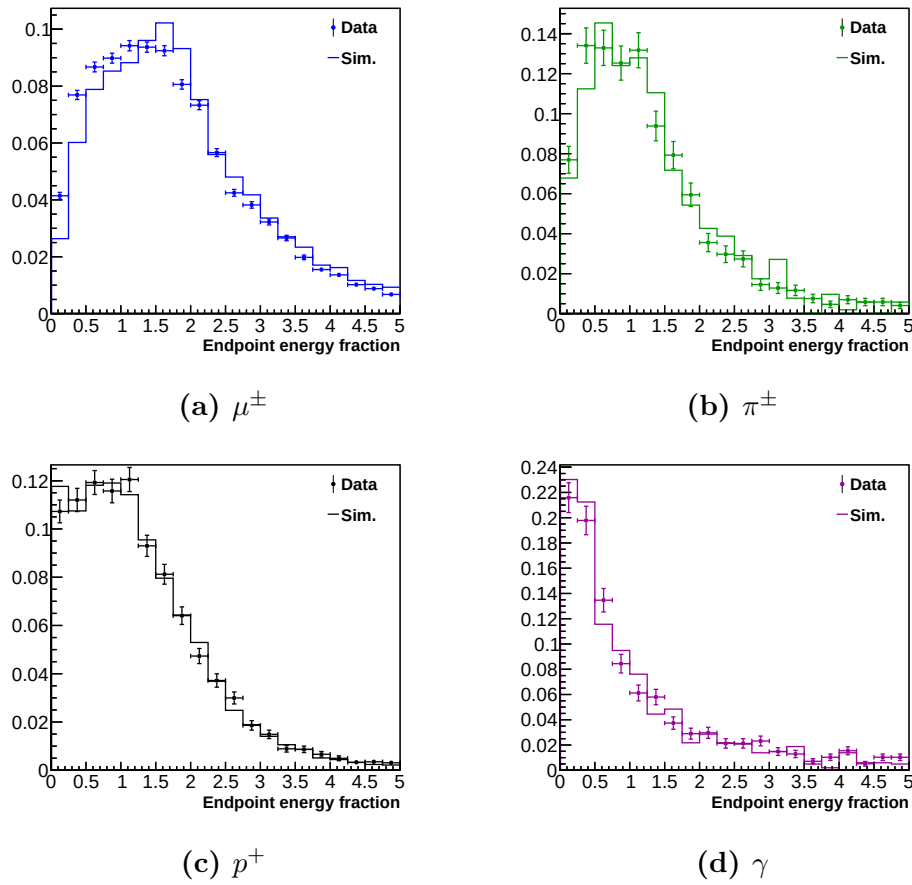


Figure 6.4: Data-simulation comparisons of the endpoint energy fraction variable for various particle types. All plots are unit normalized.

Mean dE/dx Another way to exploit the different longitudinal energy profile between particles that stop in MINER ν A due to ionization energy loss and those that initiate electromagnetic cascades (as discussed above) is to consider the amount of energy deposited per unit range, dE/dx . We use the mean dE/dx (total reconstructed energy divided by total integrated density traversed along Cone axis) to reduce sensitivity to fluctuations in the plane-by-plane energy loss.

The performance of the mean dE/dx variable is shown in fig. 6.5. Comparisons of simulation and data are shown in fig. 6.6. We find the agreement to be exceptional.

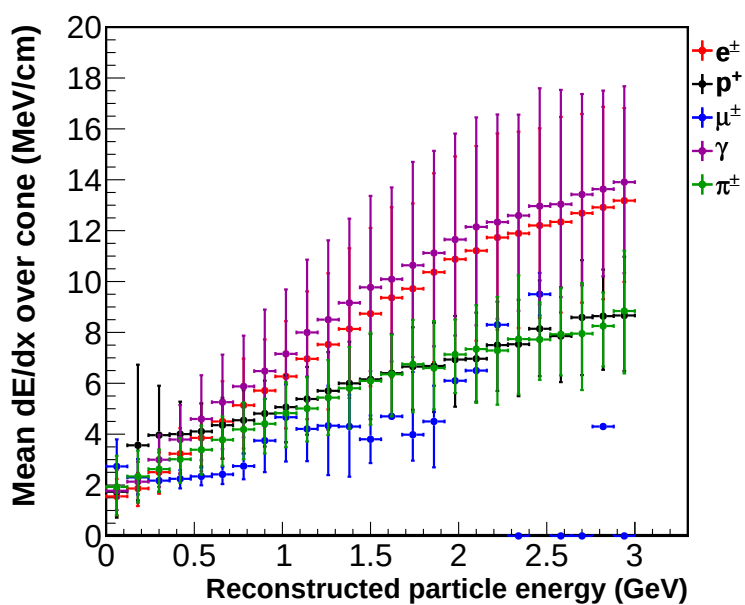


Figure 6.5: Performance of the mean dE/dx variable as a function of calorimetrically reconstructed energy. Error bars on each point represent the smallest interval containing 68% of the distribution in that bin.

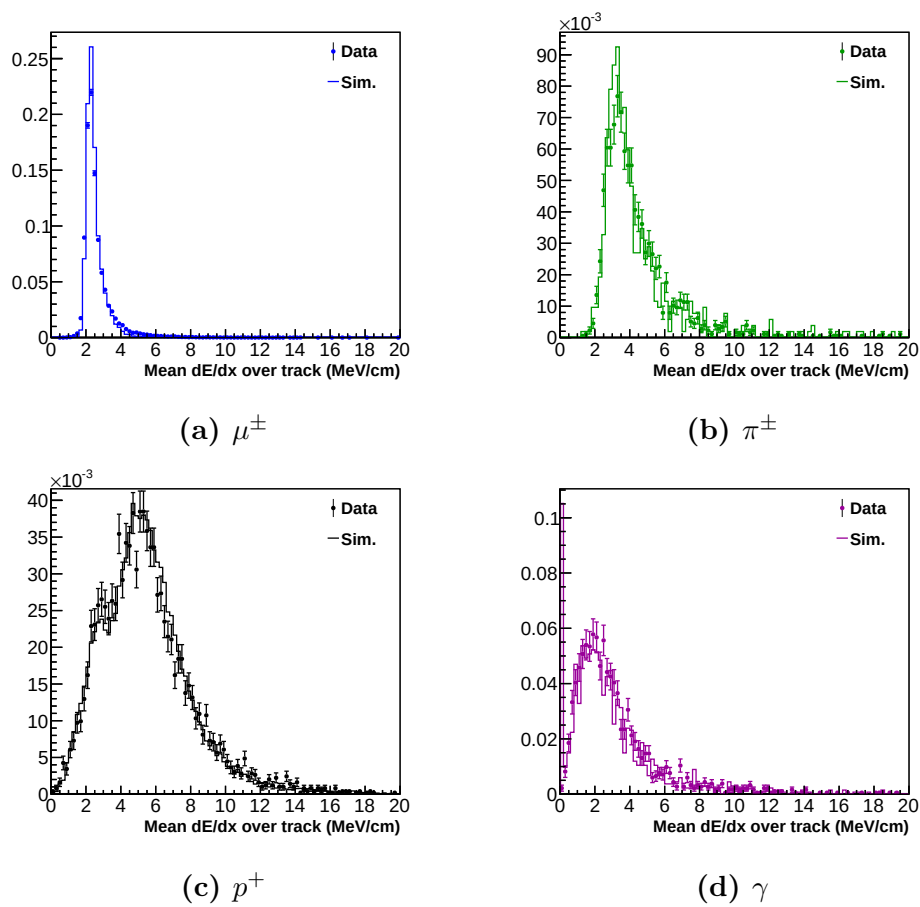


Figure 6.6: Data-simulation comparisons of the mean dE/dx variable for various particle types. All plots are unit normalized.

Median transverse width One of the most recognizable features of an electromagnetic cascade is the extent to which it spreads transverse to its direction of propagation as the cascade progresses. (This behavior arises from the increasing number of particles of which the cascade is composed, none of which carry exactly the same longitudinal and transverse components of the original particle’s momentum.) Our attempt to characterize this “width” is by computing a variable we call “median transverse width”:

1. Sort Digits by plane.
2. Search for pairs of neighboring Digits in each plane. If the two Digits with the largest energy in a plane are neighbors, merge them into one pseudo-Digit. (This significantly improves the differentiation between cascade- and single-particle-like energy profiles.)
3. For each plane, compute the weighted standard deviation of strip numbers, using Digits’ energy as weights.
4. The output score is the median of the planes’ standard deviations from step 3.

The performance of the median transverse width variable is shown in fig. 6.7. Comparisons of simulation and data are shown in fig. 6.8. Here the agreement is tolerable.

We further examine the performance of the particle identification variables via a metric known as their “separation,” $\langle S^2 \rangle$, which, as a function of the particle’s reconstructed energy, is computed as follows [80]:

$$\langle S_\xi^2 \rangle = \frac{1}{2} \int \frac{(s_\xi(E) - b_\xi(E))^2}{s_\xi(E) + b_\xi(E)} dE \quad (6.2)$$

where $s_\xi(E)$ and $b_\xi(E)$ are the probability distribution functions of variable ξ as a function of energy for signal and background, respectively. The separation of the

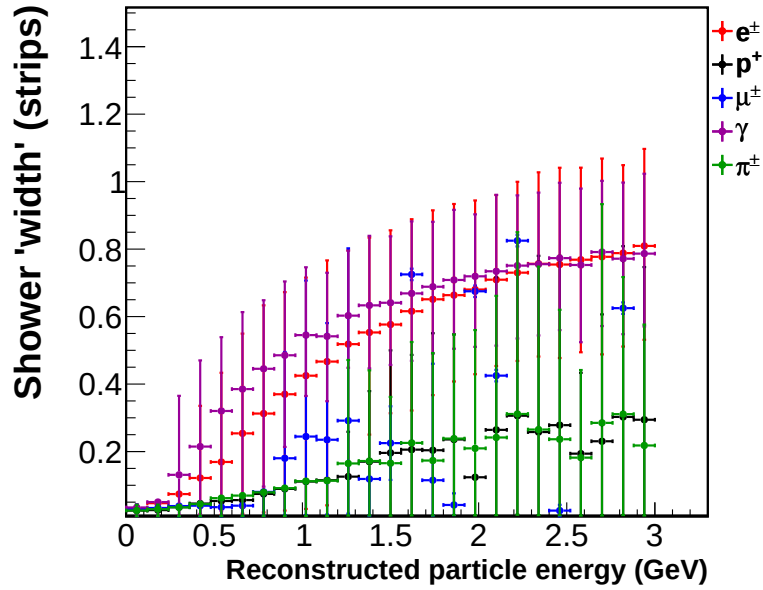


Figure 6.7: Performance of the median transverse width variable as a function of calorimetrically reconstructed energy. Error bars on each point represent the smallest interval containing 68% of the distribution in that bin.

three variables described above is shown in fig. 6.9. Significantly, the separation of two of the three variables (the mean dE/dx and the median transverse width) is negligible at low energy and grows substantially with energy, while the third (endpoint energy fraction) garners only moderate separation at all energies. This results from the lower probability of electrons and positrons to initiate cascades at lower energies, and makes it fundamentally more difficult to distinguish electrons from other particle types in that regime. This difficulty will present one of the central challenges in the analysis (ch. 7).

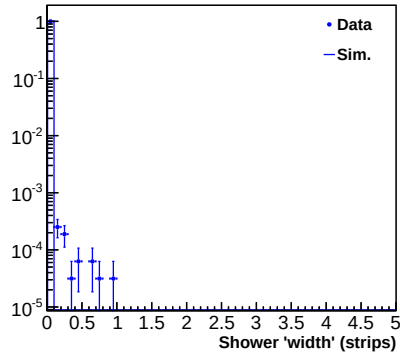
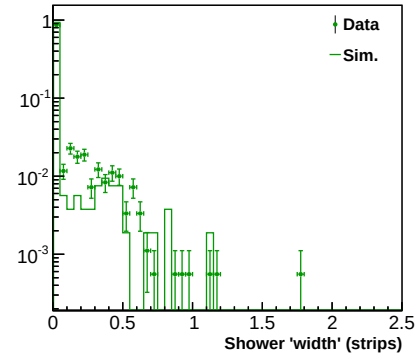
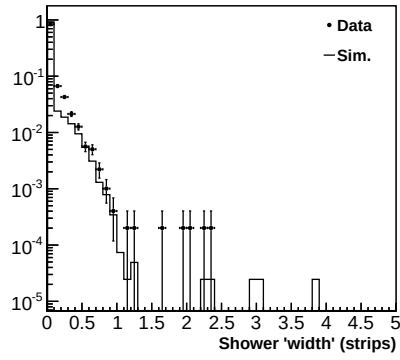
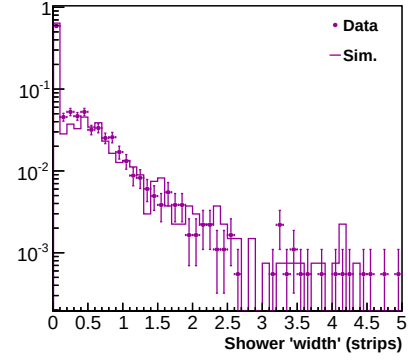
(a) $\mu^\pm$ (b) $\pi^\pm$ (c) p^+ (d) γ

Figure 6.8: Data-simulation comparisons of the median transverse width variable for various particle types. All plots are unit normalized.

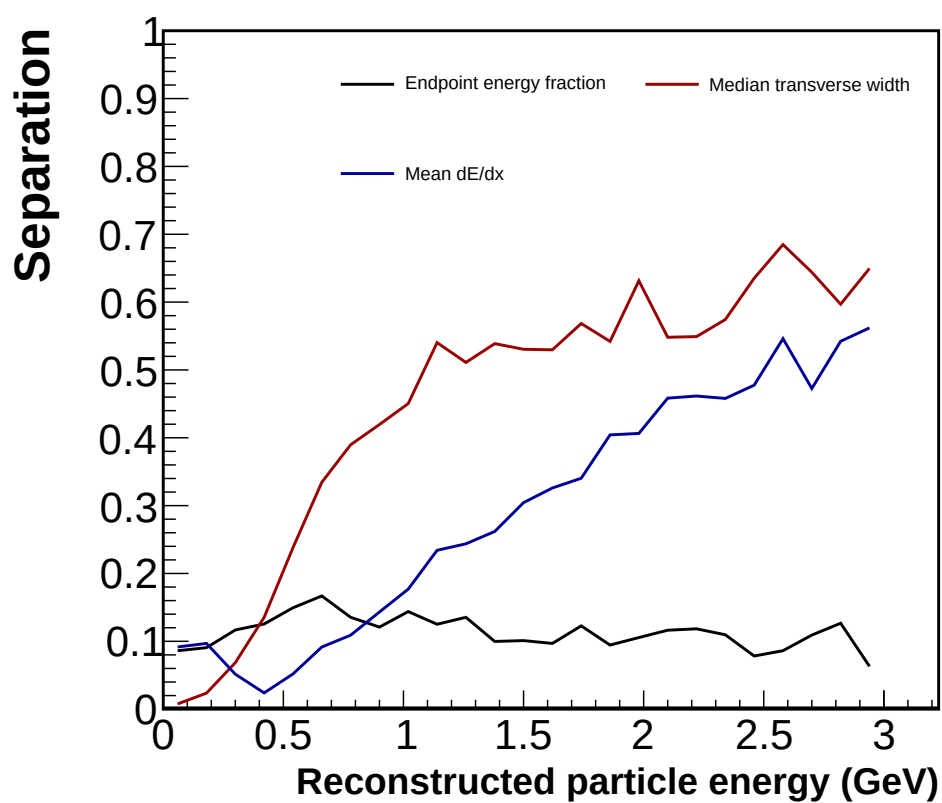


Figure 6.9: Separation (as defined in the text) of the three chosen PID variables, as a function of calorimetrically reconstructed energy.

Constructing the multivariate classifier To avail ourselves of the maximum discrimination power inherent the variables chosen in the above sections, we employ a k -nearest-neighbors discriminator [81] (kNN) trained on the input for the distributions shown above using the TMVA package distributed with ROOT [80]. The training of such a discriminator (when used to investigate N training variables simultaneously) begins by creating an N -dimensional space populated by the training signal and background points. From this, weights are constructed for each region of that space which code for what fraction of the k nearest points (using the Euclidean distance between points in N -space) are true signal events. These weights can then be used to classify a candidate event; we will refer to the output of the kNN algorithm as ζ_{kNN} below.

To actually employ the kNN output, we must choose a threshold in ζ_{kNN} (that is, a threshold in the fraction of neighboring training events were signal) above which events are retained. We began with the product of the efficiency ϵ and purity π as our figure of merit. To ascertain the right value of ζ_{kNN} to cut on, we studied $\epsilon \times \pi$ in the sample selected by the cuts of sec. 7.2 with the requirement of electromagnetic-likeness (which uses this classifier) removed; this sample (when scaled to the full data P.O.T. exposure) contained 504 events where the electron candidate was of electromagnetic cascade in origin and 6884 events where it was not. The distribution of EM-like and non-EM-like events in ζ_{kNN} , and the corresponding $\epsilon \times \pi$ curve for cuts at various values of ζ_{kNN} , are shown in fig. 6.10.

Though the optimum cut for this sample suggested by the $\epsilon \times \pi$ metric is about $\zeta_{\text{kNN}} = 0.9$, we relaxed it to $\zeta_{\text{kNN}} = 0.7$ in order to retain more events at lower candidate electron energies where, as noted in the discussion in the preceding section, the separation of the PID variables is much poorer.

Thus, if a Cone has $\zeta_{\text{kNN}} \geq 0.7$, it is deemed to be an “EM-like” Cone, and it is retained as reconstructed. If it does not, then it is dissolved, and any reconstructed

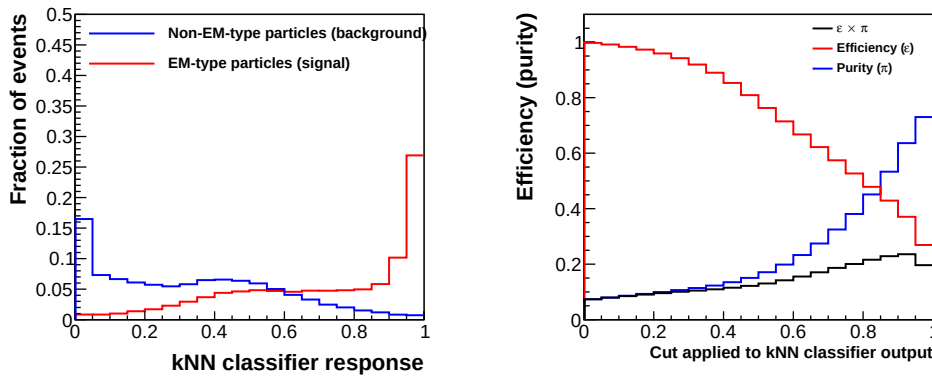


Figure 6.10: Response to the kNN classifier (left) and efficiency-purity curve (right).

objects that were added into it during the Cone expansion step noted above are freed to be used in expansion of the next candidate Cone. Every Track that exists in an event is considered for expansion in this way.

6.2.3. Event-level reconstruction

Further event cleanup and reconstruction quality cuts

Signal events should have exactly one EM-like object. Therefore, at this stage, we now require that candidate events contain exactly one Cone judged EM-like by the PID above; this is the electron candidate. Events which meet this criterion can then be subjected to further cleaning to aid in better reconstruction of the event kinematics. We eliminate all reconstructed objects from the event whose times do not lie in a $[-20, +35]$ ns window around the electron candidate's time.

We also attempt to remove events from the sample where any non-electron particles are obviously not nucleons or the event is obviously not CCQE-like. We eliminate any events with more than two Tracks emanating from the electron candidate vertex. Any events where non-electron-candidate Tracks exit the inner detector or penetrate more than halfway through the downstream HCAL are rejected.

Finally, we absorb energy not part of the Cone but near the electron candidate vertex into a separate group such that is excluded from calculations of the nuclear recoil energy. Any activity within 30cm of the vertex is grouped together and set aside.

Computation of variables used in selection cuts

There are three variables used in signal selection which can only be calculated after reconstruction is fully complete.

1. **Fiducial vertex:** Whether the electron candidate's source vertex is within the fiducial volume prescribed in sec. 7.1.2 is determined at this juncture.
2. **First-fire fraction:** PMT afterpulsing in the detector (see sec. 4.3.1) can cause phantom signal events composed mostly or entirely of afterpulse hits. (For events that would otherwise pass the signal selection, the inducing event is typically either a signal event or bremsstrahlung from a particularly energetic muon.) Since there was no simulation of afterpulsing in the simulation chain used for this analysis, it is imperative that these events be cut out from the data sample. To measure what fraction of a candidate electron Cone was likely due to afterpulsing, we examine all the Digits contained in the Cone, and determine how many of them follow another hit (above pedestal suppression threshold) in the same channel within the same gate. We call the fraction of these that were the first in their channel in the gate the "first-fire fraction;" the cut we make on it is detailed in sec. 7.2.
3. **Extra energy fraction Ψ :** As sec. 7.2 will demonstrate, the most powerful selection criterion we have in our collection is one made on the amount of energy in the recoil system. For CCQE-like events, we expect very little energy not contained in the electron candidate Cone or the associated vertex energy grouping. Therefore, we sum up all other energy in the event and

apply the calorimetric estimator described in sec. 6.2.2 to it; call this E_{other} . Because the leakage of energy outside of the cone grows with increasing shower energy, however, we do not cut on E_{other} directly. Instead, we form the ratio of E_{other} to the calorimetrically reconstructed energy inside the electron candidate (E_e) to form the “extra energy ratio” Ψ , which is stable with increasing energy:

$$\Psi = \frac{E_{\text{other}}}{E_e} \quad (6.3)$$

Once again, the cut we make on this variable is elaborated in sec. 7.2.

Energy and neutrino kinematics reconstruction

As we will see shortly, the most important observables in this analysis are the electron angle θ_e and electron energy E_e . θ_e is taken directly from the refitted electron Cone axis described in sec. 6.2.2. Similarly, the estimate of E_e was detailed above (sec. 6.2.2).

For the reconstruction of the neutrino interaction kinematics, we make use of the CCQE assumptions frequently used in the measurement of the ν_μ CCQE process [20], in which the initial nucleon is assumed to be stationary within a binding potential. Under those conditions, the neutrino energy can be estimated from the electron’s measured kinematics noted above and a few constants (the particle masses, $m_{e,n,p}$, and a binding energy E_b) alone:

$$E_\nu^{QE} = \frac{m_n^2 - (m_p - E_b)^2 - m_e^2 + 2(m_p - E_b E_e)}{2(m_p - E_b - E_e + p_e \cos \theta_e)} \quad (6.4)$$

We can furthermore use E_ν^{QE} and the observables to infer the square of the four-

momentum transferred to the nucleus, Q_{QE}^2 :

$$Q_{QE}^2 = 2E_\nu^{QE} (E_e - p_e \cos \theta_e) - m_e^2 \quad (6.5)$$

For this analysis, we use the same value of the binding energy, $E_b = 34$ MeV, as in the published ν_μ CCQE result from MINER ν A [21].

Plausibility of true event reconstruction (simulation only)

Finally, we are obligated to make one more concession to the nature of our simulation. Since we overlay data information onto the simulation to approximate the effect of multiple neutrinos interacting simultaneously (sec. 5.2.3), it is possible for a signal candidate from the *data* to be reconstructed instead of the interaction from the simulation. We wish to avoid events where this has occurred because they will erroneously be classified as the background process from the simulation; therefore, we exclude from consideration any simulation events where less than 50% of the electron candidate energy truly originated from the simulation. The effect on the sample from this requirement is minimal; 99.4% of signal events are retained.

7 Measuring the differential cross-section $\frac{d\sigma}{dQ^2}$ for the ν_e CCQE process

We calculate differential cross-sections in three variables (E_e , θ_e , and Q_{QE}^2) and a total cross-section vs. E_ν^{QE} using neutrinos of energies 0-10 GeV.¹ Generally, the formula for computing a differential cross-section from an observable event distribution N_{events} in some variable ξ is:

$$\left(\frac{d\sigma}{d\xi}\right)_i = \frac{1}{\epsilon_i \Phi T_n(\Delta_i)} \times \sum_j U_{ij} \left(N_j^{\text{data}} - N_j^{\text{bknd pred}}\right) \quad (7.1)$$

Here, the index i refers to bins in ξ , ϵ represents signal acceptance, Φ is the flux integrated over the energy range of the measurement, T_n corresponds to the number of targets (CH molecules) in the target area (fiducial region), Δ_i is the width of bin i , and U_{ij} represents a matrix correcting for detector smearing in the variable of interest. (The formula for the total cross-section, which is computed as a function of neutrino energy, differs only in that the flux is integrated only over the energy of bin i , rather than the whole energy range, and that we therefore do not need to divide by the bin width Δ_i .) We will examine how each factor of this equation is estimated in detail in the sections below.

¹We use this upper threshold to remain in a regime where the uncertainties on the neutrino flux are tolerable; see, e.g., fig. 3.6 and note the behavior of the uncertainties above 10 GeV.

7.1. Signal and fiducial volume definitions

7.1.1. Signal definition

As was noted in sec. 1.3.3, a number of effects interfere with the basic CCQE process to make the identification of precisely those events which interacted as CCQE within a nucleus very difficult. Furthermore, the non-magnetization of the MINER ν A detector renders it impossible to distinguish on an event-by-event basis between electrons and positrons, so that the ν_e and $\bar{\nu}_e$ processes' final states are identical. Because of these challenges, we elected to search for what we will term a “ ν_e CCQE-like” process, which is defined by the (observable) final state particles. In particular, we require the final state to consist of:

- Exactly one lepton, either an electron or a positron
- Any number of nucleons (protons or neutrons)
- Exactly zero other hadrons
- Exactly zero photons with energies above 15 MeV (photons below this threshold stem chiefly from nuclear de-excitation and are allowed)
- Any number of nuclear fragments

7.1.2. Fiducial volume

Our fiducial volume consists of the majority of the central plastic Tracker region, which corresponds to a fiducial mass of approximately 5.57 tons. We leave the following buffers of scintillator between the surrounding volumes:

- 25 cm in z between the upstream end of the fiducial volume and the most downstream passive target;
- 26 cm in z between the upstream end of the fiducial volume and the most upstream lead sheet in the downstream ECAL;

- 6 cm radially between the outer edge of the fiducial volume and the innermost edge of the side ECAL shielding (and 17 cm further of side ECAL before reaching the outer edge of the ID).

In detector coordinates (see sec. 4.1) the set of points (x, y, z) enclosed by the fiducial volume can be expressed by the following conditions:

$$5990.0 \text{ mm} \leq z \leq 8340.0 \text{ mm}$$

$$a \leq 850 \text{ mm}$$

(where a is the apothem of the smallest regular hexagon contained in the $x - y$ plane, centered at $(0, 0, z)$, and containing (x, y, z)).

7.2. Selection of events

Broadly, the event selection can be broken down into three phases: electromagnetic cascade selection, electron-photon discrimination, and CCQE selection. The first of these was discussed extensively in sec. 6.2.2, where it played a central role in event reconstruction, but we apply further selections here to isolate an optimally pure sample. We will also treat the others in full detail below.

7.2.1. Electromagnetic cascade selection

We make two selection cuts beyond the particle identification requirements of sec. 6.2.2, both of which attempt to further purify the sample at low reconstructed electron candidate energies.

Fraction of energy in non-MIP-like-clusters

Low-energy electromagnetic cascades tend to burn themselves out before spreading out significantly in the transverse direction. But because they still tend to consist of a plurality of particles, fluctuations in the longitudinal energy profile of such cascades are usually more frequent and more pronounced than in the profiles of single particles traversing the detector (the primary background). This distinction is roughly captured by the Cluster classification system used in the reconstruction, which was described in sec. 6.1.2. Since “Trackable” clusters roughly correspond to the deposited energy of single particles traversing a plane of the detector, we construct a “non-MIP-cluster energy fraction” variable by dividing the total energy of non-Trackable clusters within the electron Cone by the total Cone energy to attempt a measure of these fluctuations.

The distribution in this variable and the cut threshold selected are shown in fig. 7.2.

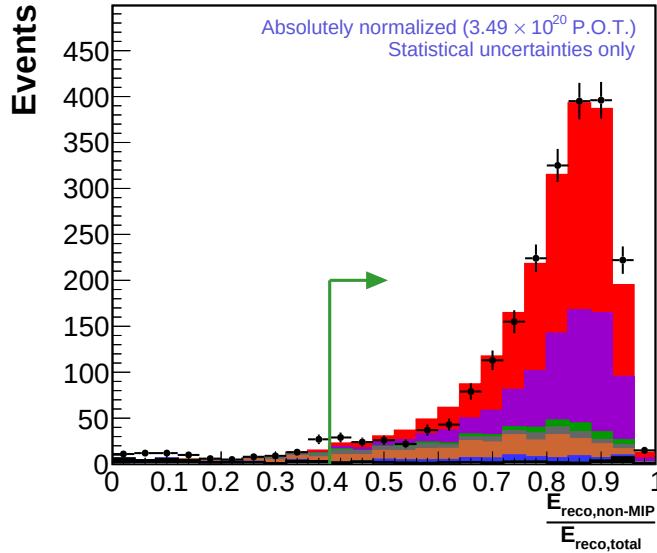


Figure 7.1: Distribution of “non-MIP-cluster fraction” (with all other selection cuts applied). The cut made on this variable is indicated by the dashed line and arrow.

Transverse Spread Score

The second variable used to help reduce contamination at low electron candidate energies is a so-called “transverse spread score.” The primary backgrounds at low electron candidate energy are single particles, usually from neutral current interactions, and they tend to leave very narrow tracks of energy. In this variable, such energy patterns produce substantially lower scores than signal-type events. The basic principle is to sum energy-weighted strip numbers (which are a proxy for the transverse coordinate). In detail:

1. Sort Digits (n.b. *not* Clusters) by plane and then by strip number.
2. Determine the energy-weighted mean strip number for each plane j ,

$$\bar{N}_j = \frac{1}{(\sum_i E_i)} \sum_i N_i E_i$$

3. Determine the median energy deposit for each plane j , $\tilde{E}_j$.
4. The score for each plane is summed over each digit i :

$$S_j = \left(\sum_i E_{j,i} \right)^{-1} \sum_i \left| \bar{N}_j - N_{j,i} \right| \left| \tilde{E}_j - E_{j,i} \right| \quad (7.2)$$

5. The final score $S = 1/N_{planes} \sum_j^{N_{planes}} S_j$.

The distribution in this variable and the cut we make are shown in fig. 7.2. The underprediction at low values of transverse spread score is largely due to activity leaking in from the outer detector; the simulation sample used for this analysis was generated using only the inner detector as the neutrino target, due to technical constraints, and therefore does not model events whose interaction vertex is in the outer detector but whose products can spill into the fiducial region. Spill-over that would otherwise be accepted by this analysis is typically low in visible energy (less than 500 MeV), at high angles, and very narrow in trajectory, somewhat like the

neutral current events predicted by the simulation. The cut shown here is chosen in a region where the agreement between the model and the data is reasonable.

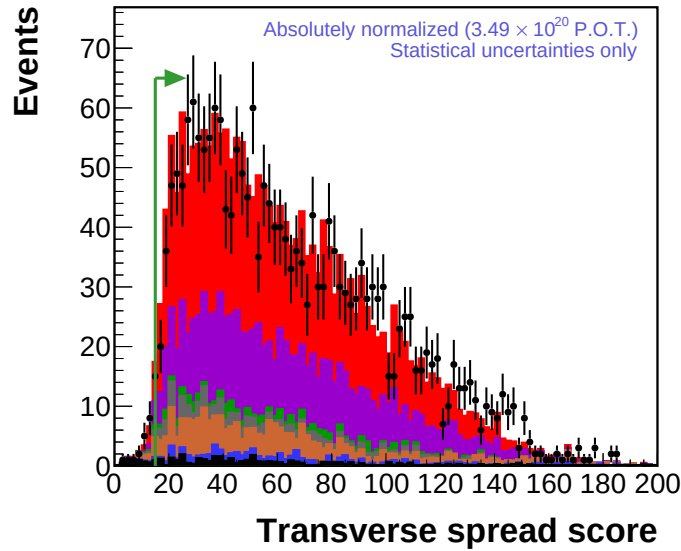


Figure 7.2: Distribution of “transverse spread score” (with all other selection cuts discussed in this section applied). The cut made on this variable is indicated by the dashed line and arrow.

7.2.2. Electron-photon discrimination

At typical MINER ν A energies ranging from several hundred MeV to several GeV, photons interact with the detector primarily through pair production in the nuclear field. [11] Because of their boost in the direction of travel, the resultant electron and positron typically overlap, and therefore, the cascade that they produce begins with a Track whose energy deposition rate is approximately twice that of a single particle. Therefore, to reduce backgrounds from events in which a photon mimics an electron or positron signal, we inspect the energy observed in the front the electron candidate object.

We expect that in some signal events (as well as many background events where an electron or positron is the primary particle), other charged particles exiting

the nucleus may overlap in one or more views with the electron candidate. In such a case, if we were to measure strictly the energy observed at the very front of the electron candidate Cone, we could measure a dE/dx consistent with two (or more) particles, and therefore erroneously classify the event as a photon-like rather than an electron-like one. Therefore, we instead sequentially compute the dE/dx of 100 mm blocks along the cone axis, progressing in 25 mm increments, until the further end of the block reaches 500 mm from the reconstructed event vertex (or the end of the Tracker region, whichever is first). We then select the *minimum* dE/dx from this collection as our classifier. Using the minimum in this way allows the algorithm to step over any overlapping nuclear activity towards the front while still using the most upstream part of the cascade available to attempt to differentiate between electrons and photons. Note that because of the stochastic nature of shower development, even if there is no overlap from another particle, the minimum dE/dx may occur some distance from the initial interaction point (see fig. 7.3, but note that the radiation length in the Tracker material is roughly 400 mm).

Our choice for the cut threshold is optimized by using an $\epsilon \times \pi$ figure of merit, as was done for the particle identification (sec. 6.2.2). In this case, we use the superior statistics of dedicated particle cannon simulations of electrons and photons (but reconstructed with the same techniques as described in sec. 6.2) to perform the optimization, which is illustrated in fig. 7.4. (The particle cannon sample used for the tuning is described in more detail in section 6.2.2.) Our simulation indicates that the search-based dE/dx method described above compares favorably to one in which a fixed portion of the electron cone is used for the purposes of separating electrons from photons. This is illustrated in fig. 7.4. The distribution of the minimum front dE/dx variable in the beam data and simulation is in fig. 7.5.

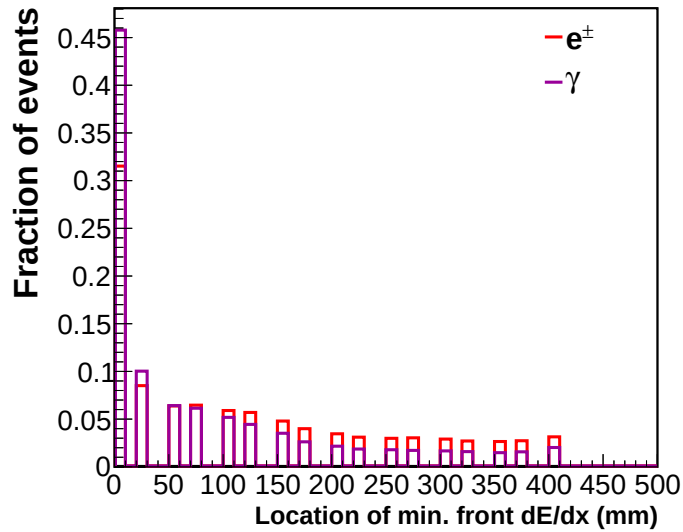


Figure 7.3: Distribution of the location of the minimum front dE/dx along the shower axis (as described in the text) for simulated electrons and photons (sample described in the text).

The disagreement between the simulation and the data above 2 MeV/cm in fig. 7.5 suggests that there is substantially more photon-like background than is predicted by the simulation. Candidate explanations for the origin of this data excess are explored in app. B; our attempt to constrain and estimate the effect it may have on this result is detailed in sec. 7.4.

7.2.3. CCQE selection

No Michel electron candidates

A substantial fraction of the predicted background events contain a charged pion in them. While π^- tend to capture on any of the detector materials, π^+ , by contrast, in some cases will stop without an inelastic interaction, in which case the positron from $\pi^+ \rightarrow \mu^+ X \rightarrow e^+ X$ can be observed after a short ($\tau \sim 2 \mu s$) delay. We search for these Michel electrons (positrons) near the beginning and end vertices of every track in the event.

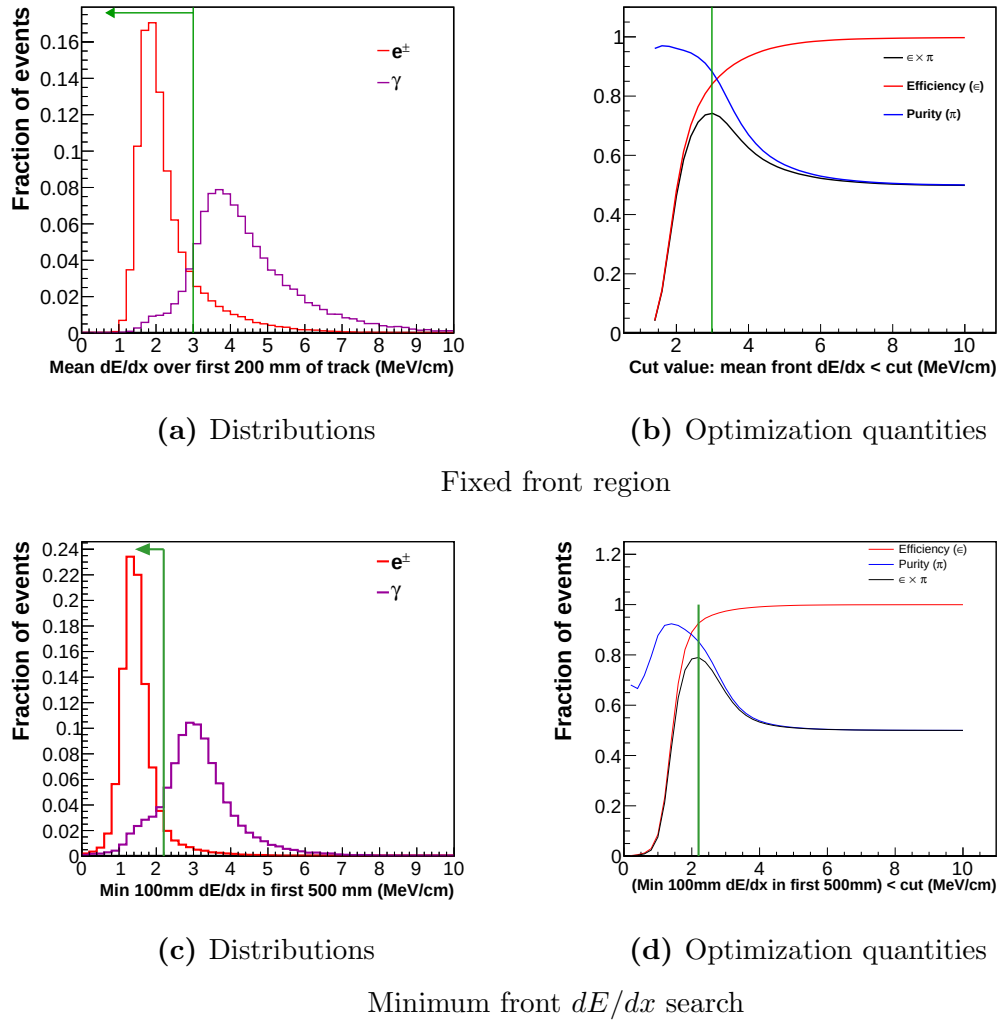


Figure 7.4: Distributions used to tune the mean front dE/dx cut, constructed from the particle cannon sample noted in the text. The green lines indicate the optimum cut chosen.

Requiring the absence of a Michel electron candidate for signal events results in an increase in the predicted final sample purity of roughly 6 percentage points. The events removed from the sample are precisely those contained in the Michel-match sideband discussed in sec. 7.4; more detailed discussion of the event content can be found there.

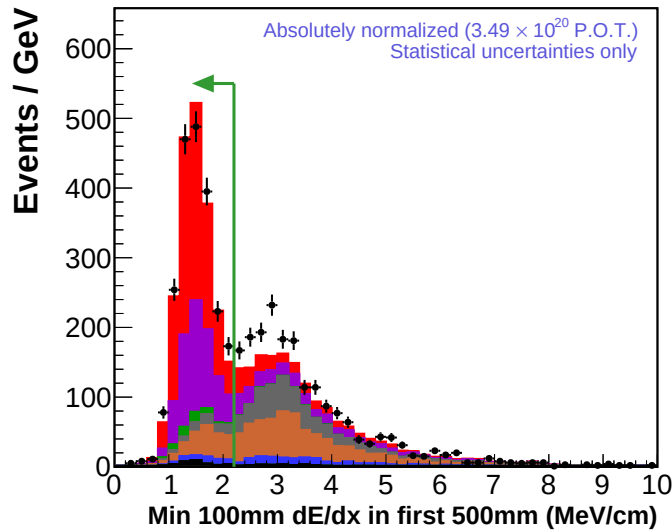


Figure 7.5: Distribution of minimum 100 mm dE/dx over first 500 mm in events (with all other selection cuts applied). The cut made on this variable is indicated by the line and arrow.

Extra energy fraction

Our most powerful selection to separate CCQE-like events from other backgrounds is a cut on the extra energy fraction Ψ described in sec. 6.2.3 as a function of the visible energy in the event. This cut is optimized in the following bins of E_{vis} (in GeV): $(0, 0.5)$, $(0.5, 2)$, $(2, 3)$, $(3, 5)$, $(5, \infty)$. We optimize in each of these bins by examining the efficiency (ϵ) and purity (π) of a cut at each bin edge in Ψ and find the cut which maximizes $\epsilon \times \pi$; this procedure is depicted for a sample bin in fig. 7.6. The distribution of Ψ vs. E_{vis} , along with the cuts optimized over the bins, are shown in fig. 7.7. The final tunings (the black curve in fig. 7.7) are as follows:

E_{vis} range (GeV)	Ψ upper bound
(0, 0.5)	0.10
(0.5, 2)	0.10
(2, 3)	0.06
(3, 5)	0.03
(5, ∞)	0.03

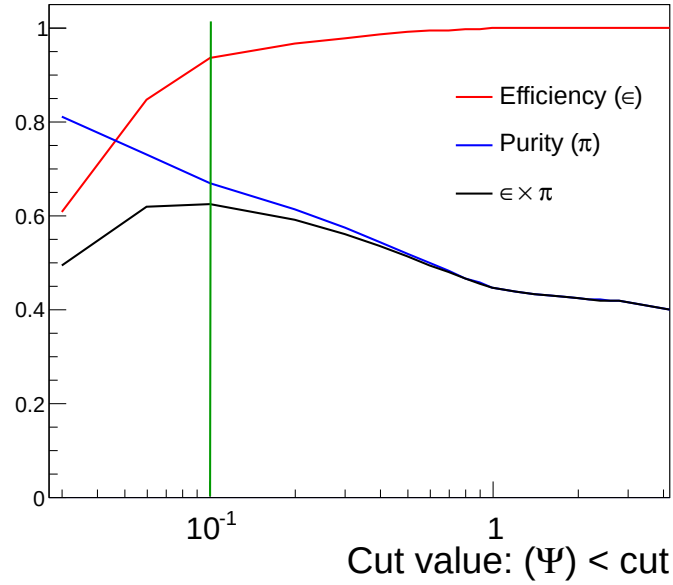


Figure 7.6: Efficiency, purity, and their product for various cuts on Ψ in the $0.5 \text{ GeV} \leq E_{\text{vis}} < 2 \text{ GeV}$ bin. The optimum cut, indicated by the green line, is for $\Psi < 0.1$.

Final cleanup

The remainder of the selection cuts are intended to guard against catastrophically mis-reconstructed events that somehow managed to slip through all of the preceding requirements. The first cut requires a positive reconstructed neutrino energy, E_{ν}^{QE} , as defined in sec. 6.2.3.

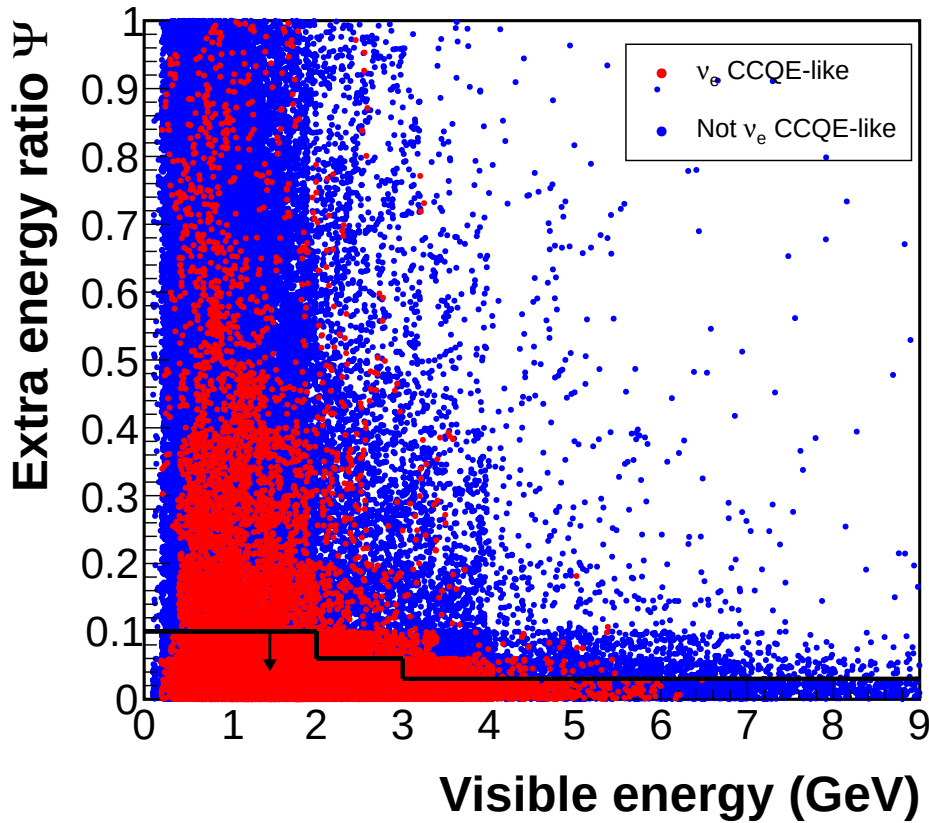


Figure 7.7: Distribution of Ψ vs. visible energy in events (with all other selection cuts discussed in this section applied). Events below the black curve are retained as signal candidates.

The final two cuts eliminate event misreconstruction due to known detector effects; they do not affect the simulation. The first is a cut on a quantity which quantifies whether activity upstream of a track has been obscured by the discriminator dead time induced by previous activity (cf. sec. 4.3.2). (The quantity itself is calculated by projecting the electron candidate track into the two modules upstream of the vertex and counting the number of discriminator pairs that should be experiencing deadtime in the strips intersected by the projection or their immediate neighbors, based on recorded activity immediately preceding the time of the electron candidate.) The primary purpose of this cut for this analysis is to eliminate events in which an upstream part of the seed Track has been lost, which

would bias the vertex position reconstruction as well as possibly the reconstruction of the energy.

The second detector effects cut we make is intended to remove fake signal events that arise from afterpulsing. We make use of the “first-fire fraction” variable calculated as described in sec. 6.2.3, and require $f_{\text{first-fire}} < 0.25$. (This cut was tuned by examining the distribution of $f_{\text{first-fire}}$ in candidate signal events which were determined to be from afterpulsing using a hand scan of event displays.)

7.3. Selected event sample

The samples selected in the data and simulation by these cuts are shown as a function of the two observables (electron angle and energy) and the two variables inferred from quasi-elastic kinematics (neutrino energy and Q^2) in fig. 7.8. Comments about each of the observables are in order:

Electron angle The electron angle used here is the angle with respect to the neutrino beam, not the polar angle with respect to the detector z -axis. (The latter is the θ calculated in detector coordinates and differs from the θ reported here by $\sim 3.4^\circ$.)

Electron energy The reconstructed electron energies for the simulated events are multiplied by 1.05 here to account for a measured 5% electromagnetic energy scale discrepancy between data and simulation in the π^0 invariant mass spectrum constructed for ref. [79].

We classify events as follows:

ν_e CCQE-like Signal as defined in sec. 7.1.1, with $E_\nu < 10$ GeV.

$\nu + e$ elastic Elastic scattering of any neutrino flavor from atomic electrons.

Other CC ν_e Charged-current ν_e interactions that are not quasielastic-like as per sec. 7.1.1.

NC Coh Coherent neutral-current interactions (which produce a single π^0 in the final state).

Other NC π^0 Incoherent neutral-current scattering which produces a π^0 and possibly other particles in the final state.

CC $\nu_\mu \pi^0$ Any muon neutrino charged-current process which produces a π^0 in the final state.

Other Anything else. (Includes, in particular, signal events that have $E_\nu > 10$ GeV.)

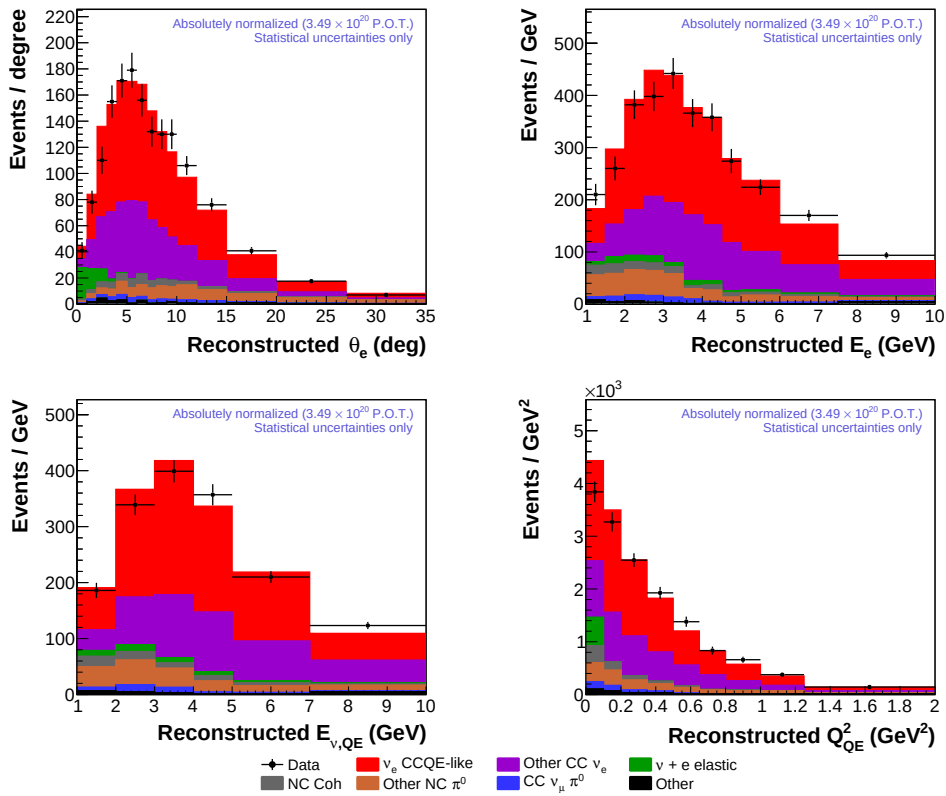


Figure 7.8: Selected events after all cuts described in the text. (The background categories are also described in the text.)

7.4. Background model and constraints

From fig. 7.8 it is plain that a sizeable fraction of the remaining events in the working samples actually result from background processes and must be subtracted before any analysis on the signal distribution can be made. Moreover, fig. 7.5 suggests that the simulation either underestimates or is entirely missing one or more processes that lie immediately adjacent to the signal region in front dE/dx ; and this may influence the background estimate in the signal region. Since the sizes and shapes of the backgrounds subtracted will manifestly affect the final results in a significant way, we require constraints from the data on the predicted backgrounds before we subtract them from the data distribution.

The flux constraint of sec. 3.4 is in fact a data constraint for the $\nu + e$ scattering background (green in fig. 7.8), since it consists of weights that were obtained by matching the prediction in that channel to a direct measurement. Therefore, once we have applied the flux constraint to that background category in each distribution, we consider that class to be constrained by data. (Applying the flux constraint results in a roughly 10% reduction in the normalization, as well as minor changes in the shapes of various kinematic distributions.)

MINER ν A's own measurement of charged-current coherent pion production found that the predicted normalization in GENIE is roughly correct [82], and the same model (from Rein and Sehgal [83]) produces the neutral-current coherent sample here. As shown in fig. 7.9, MINER ν A found that the agreement between the model and the measurement can be substantially improved by simply reducing the cross-section for coherent pion production by 50% for $E_\pi < 450$ MeV. We apply that reweighting to the coherent events in our sample (both CC and NC) and afterwards (along with the introduction of the uncertainties described in sec. 7.10.1) consider that process to be fully constrained.

The impact that the process responsible for the excess noted in sec. 7.2.2 has

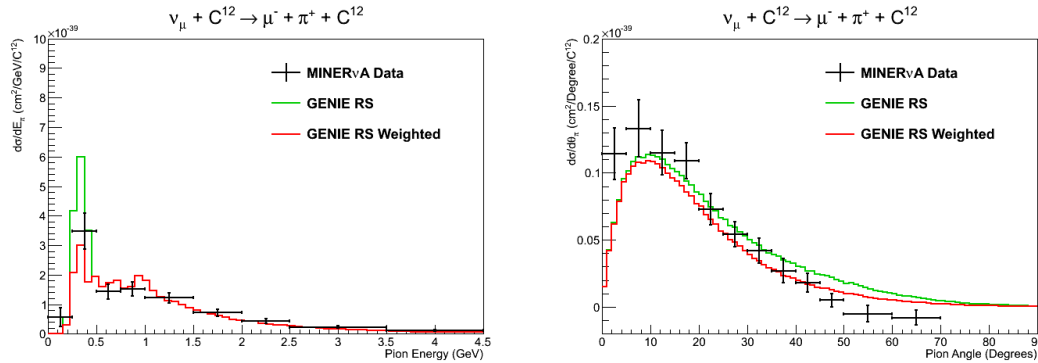


Figure 7.9: The effect of reweighting GENIE’s coherent model down by 50% for pions with $E_\pi < 450$ MeV. The data-simulation agreement is significantly improved. Courtesy A. Mislivec (MINERνA).

in the signal region is estimated using an *ad hoc* sample of π^0 s whose kinematic distributions are fitted to match those observed in a dE/dx region immediately neighboring the signal region. (The selection of this model for the excess is the subject of app. B.) The excess process is predicted to have very little effect on the signal region, as can be seen in fig. B.20b. Nevertheless, we will subtract off the minimal contribution along with the backgrounds predicted by GENIE after the latter are constrained, as detailed below.

7.4.1. Construction of sidebands

We attempt to construct sidebands for the purposes of constraining the normalizations of three of the primary background classes: ν_e CC inelastics and incoherent π^0 production (both charged- and neutral-current). Both sidebands result from the modification or inversion of a single selection cut:

“Extra energy” sideband For this sideband, we select events which do not pass the cut from sec. 7.2.3, but otherwise have $\Psi < 0.225$, in addition to the remainder of the signal selection cuts. (Here compare to the standard cut in fig.

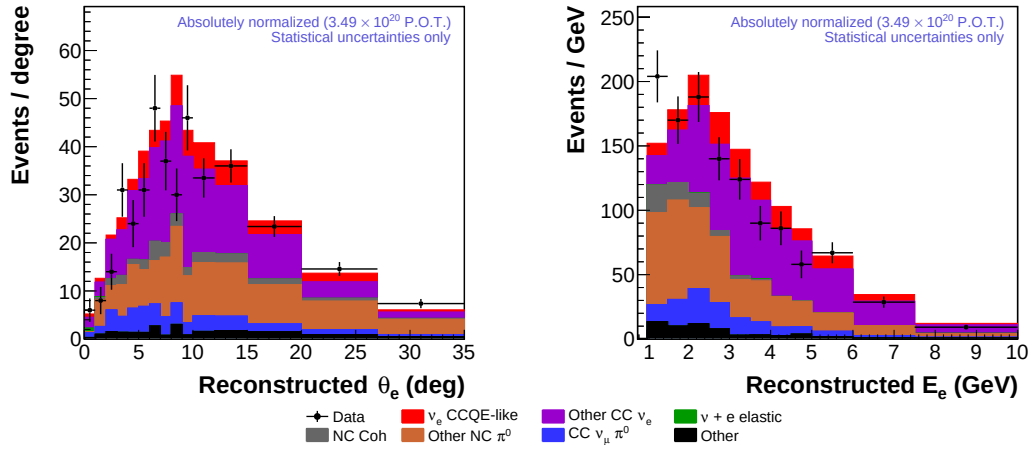


Figure 7.10: Distributions in the observables from events in the extra energy sideband.

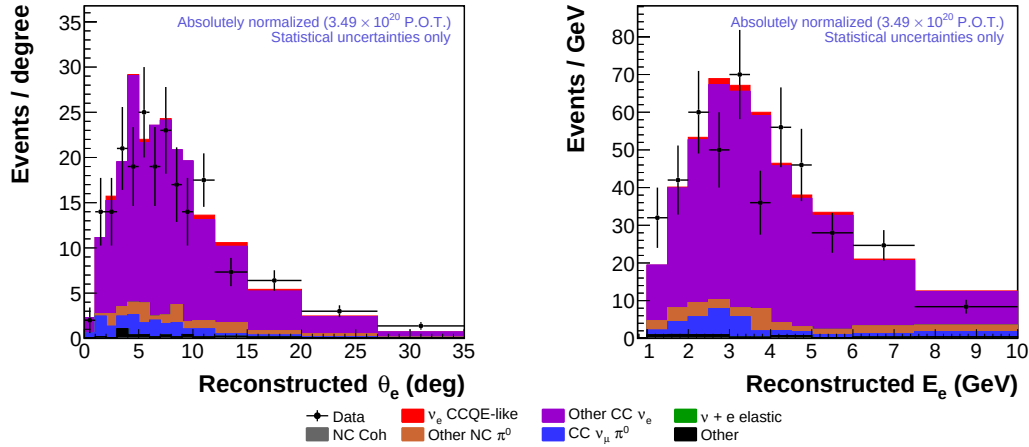


Figure 7.11: Distributions in the observables from events in the Michel-match sideband.

7.7.) This selection prefers events that are somewhat more inelastic than the signal, and thus creates a sideband rich in the inelastic backgrounds. Distributions of the observables in this sideband are shown in fig. 7.10.

Michel electron candidate sideband The final sideband's criteria consist of the usual signal selection criteria with the Michel electron rejection cut (sec. 7.2.3) reversed. This sideband is almost entirely comprised of events from the ν_e inelastic class, as seen in the distributions of the observables in fig. 7.11.

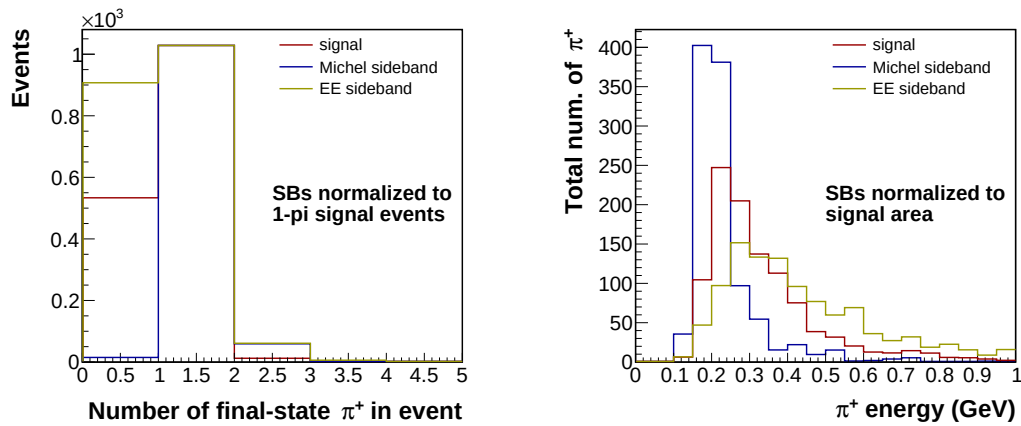


Figure 7.12: True π^+ multiplicity (left) and energy (right) in the signal and sideband regions.

It is potentially worth considering whether the events in each background category in the sidebands are compatible with the background events present in the signal region. (If they are substantially different, then any constraints derived in the sidebands may not be fully applicable to the signal region.) We examined the multiplicity and energies of π^+ in events falling into the “other ν_e ” background category in the signal region and Michel-match and extra energy sideband regions. Since, as shown in fig. 7.12, the signal straddles the two sidebands’ distributions in π^+ energy, we conclude that our best chance to model the kinematics present in the signal region is to use both sidebands together.

7.4.2. Signal process constraint

From the plots in figs. 7.10-7.11, it is evident that there is a non-negligible presence of signal events in the extra energy sideband region. Flaws in the prediction of the signal process may therefore influence any attempts to correct the background distributions based on the sidebands, and uncertainties in the signal process will also contribute to uncertainty in the background prediction. To combat this problem, we apply constraints to the signal predictions—for the purpose of the

background constraint calculation only—which are computed by simultaneously fitting the normalizations of the four predicted cross-sections to the measured data cross-sections shown in sec. 7.8. (We fit using exactly the same machinery used for the backgrounds, described below in sec. 7.4.3.) Because the result of this normalization constraint will affect the background prediction and thus the cross-section measured in the end, we perform this process iteratively until the scale factors for the signal process stabilize. As shown in fig. 7.13, this required only three iterations.

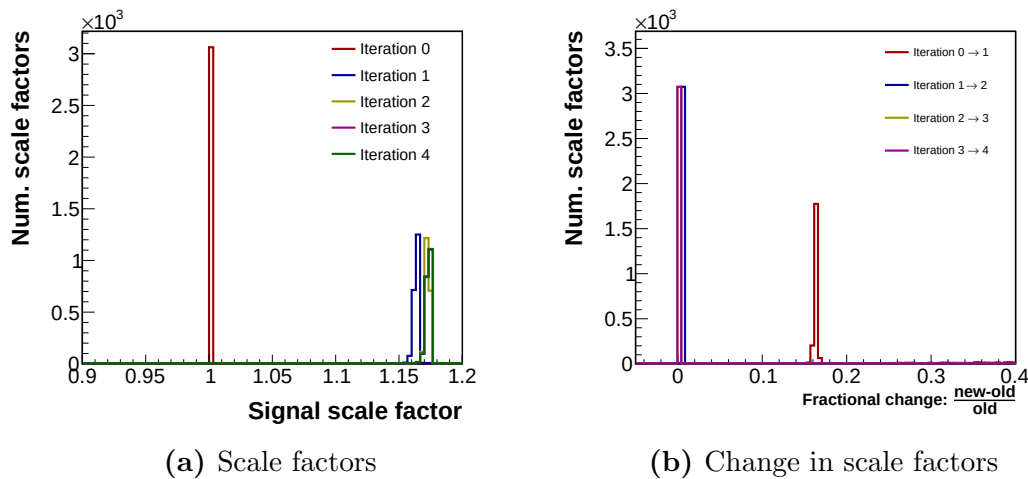


Figure 7.13: Fitted scale factors, and their fractional change from iteration to iteration, for the signal prediction for all the error band universes. Since there was virtually no change from iteration 3 to iteration 4, we conclude that 3 iterations is sufficient.

We wish to stress again that these scale factors are used *only* during the calculation of background scales in sec. 7.4.3, and are not propagated any further or used anywhere else in the analysis.

7.4.3. Calculation of background normalizations

As was noted at the beginning of sec. 7.4.1, we intend to constrain the two most important background classes for which we have reason to suspect the

Background class	scale factor
Other CC ν_e	0.89 ± 0.08
CC & NC incoherent π^0	1.06 ± 0.12

Table 7.1: Fitted background scale factors. The uncertainties shown here include an extra factor of $\sqrt{2}$ to account for the correlations between the distributions, as described in the text.

normalizations may be incorrectly predicted: ν_e CC inelastics, and incoherent π^0 production. (In this scheme, we will constrain the CC and NC incoherent π^0 categories together, because we have only two independent sets of events arising from the two sidebands. In this way we have two unknowns and two constraints.) To obtain our best estimate of the normalization of these backgrounds, we fit the four histograms in figs. 7.10-7.11 simultaneously using a single free parameter for each of the background classes we wish to constrain. (Since we apply the signal constraint from sec. 7.4.2, the flux constraint of sec. 3.4 provides a constraint for the $\nu + e$ scattering background, and as was pointed out in the introductory material to sec. 7.4, MINER ν A's measurement of CC coherent production provides an indirect constraint for NC coherent π^0 , the only background classes that will remained unconstrained after this procedure are negligible.) Using the Minuit2 fitter within ROOT [84], we minimize the total χ^2 between all four histograms' data and MC distributions (where index n refers to the histogram, i to the bin number, and α to the background class), universe by universe, to fit the best set of scale factors $\{f^\alpha\}$:

$$M = \sum_{n=1}^4 \sum_{i=1}^{N_{bins,n}} \frac{(H_{data,n}(i) - \sum_{\alpha} f^{\alpha} H_n^{\alpha}(i))^2}{\sigma_{data,n}^2(i)} \quad (7.3)$$

The fitted scale factors in the central value universe are noted in tab. 7.1. Because we use the same events in each pair of E_e , θ_e distributions in the same sideband, the four distributions are not all statistically independent; this means

that the statistical uncertainties on the final parameters reported by the minimizer are *underestimated*. However, we repeated the fitting procedure using all four of the independent combinations of only two distributions (E_e in the extra energy sideband and E_e in the Michel-match sideband; E_e in the extra energy sideband and θ_e in the Michel-match sideband; and so on), where the uncertainties returned by the minimizer are correct. In all these cases, the uncertainties were roughly $\sqrt{2}$ larger. We therefore have included an extra factor of $\sqrt{2}$ in the uncertainties reported in tab. 7.1 to account for the effect of the correlations on the uncertainties.

The distributions after the best fit are given in fig. 7.14.

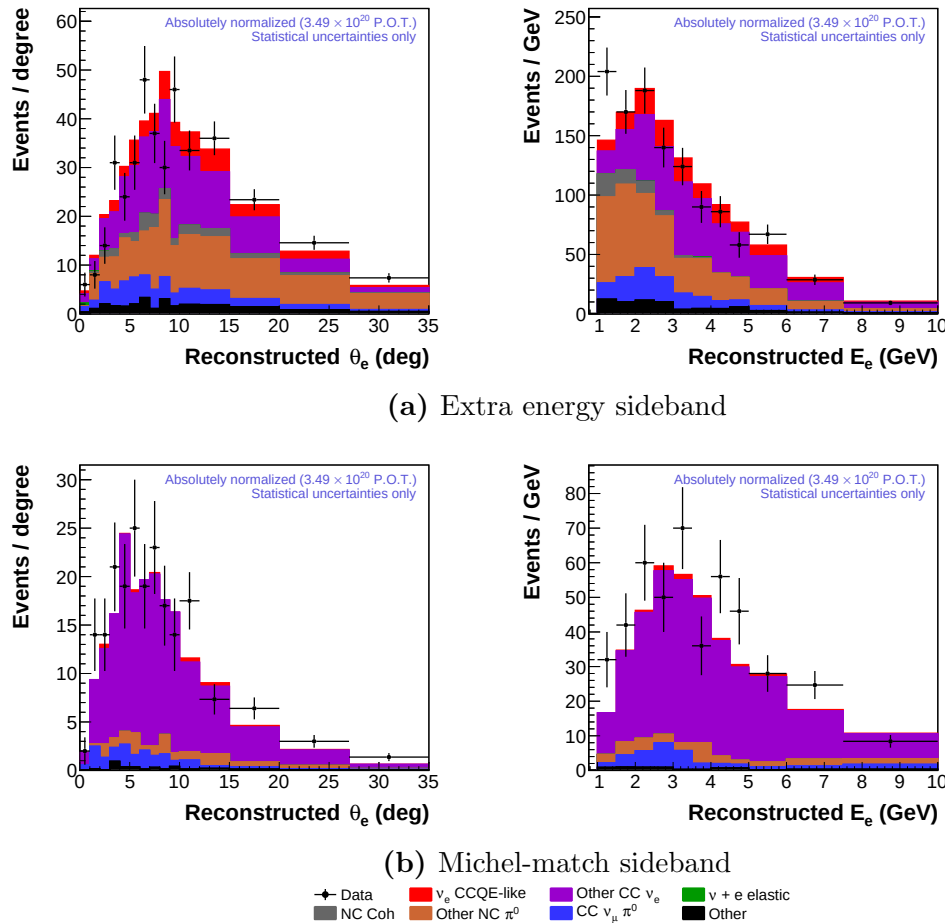


Figure 7.14: Observable distributions in the sidebands after the background fitting of sec. 7.4.3. These distributions can be compared to figs. ??-7.11

The f^α are then applied to the backgrounds in the signal region, and the

estimated backgrounds subtracted from the data. The background-subtracted data is compared to the signal prediction in fig. 7.15.

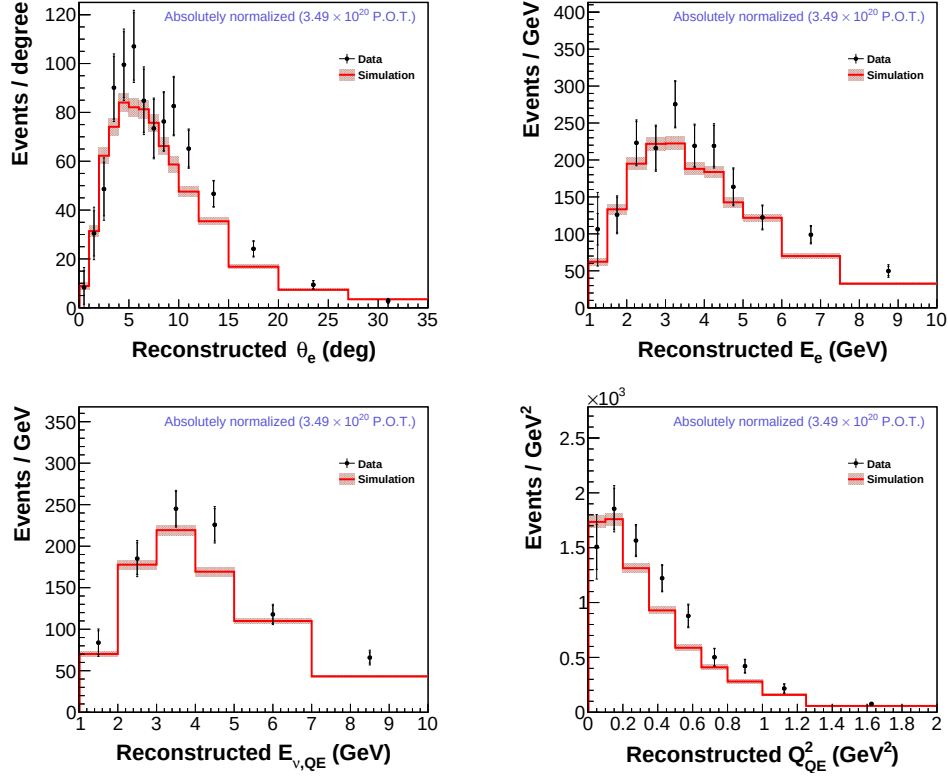


Figure 7.15: Background-subtracted distributions compared to the signal prediction. Data uncertainties are elaborated in fig. 7.16; MC uncertainties are statistical only.

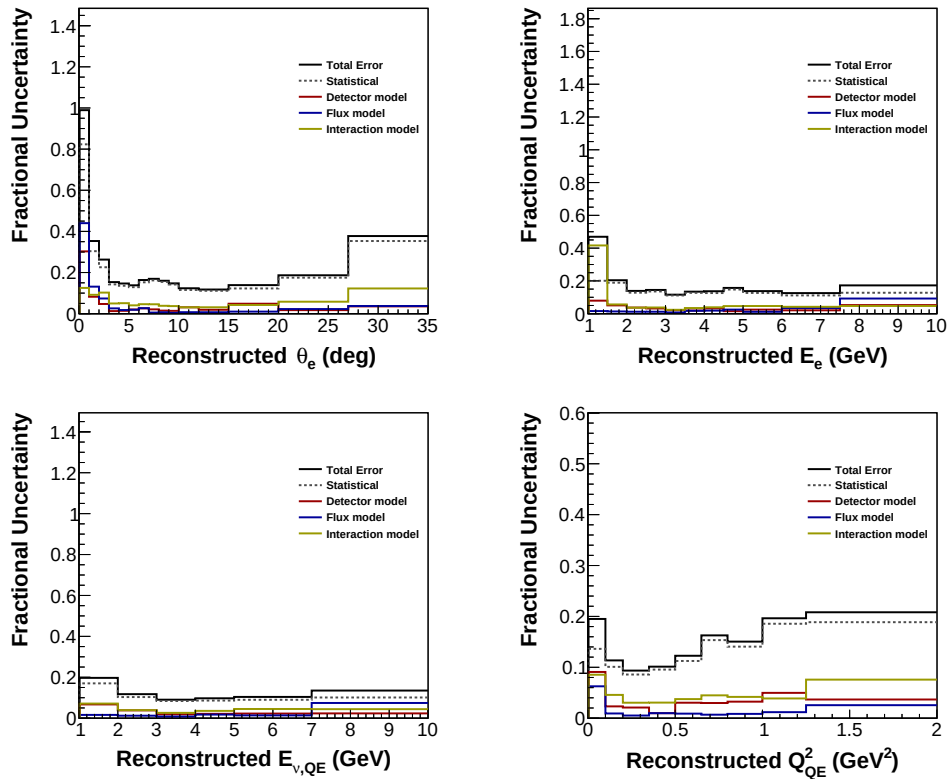


Figure 7.16: Uncertainties on the background-subtracted data distributions. Further detail on the uncertainties is given in sec. 7.10.

7.5. Unfolding in observables

After background subtraction is complete, we are left with our best estimate of the signal rate in reconstructed quantities. Because we measure efficiency as a function of true variables (as not every true signal event can even be reconstructed), however, we attempt to correct the reconstructed values to their true values before correcting for inefficiency. We do this using an unfolding procedure that relies on the simulation to predict the migration of events from their true to their reconstructed values due to detector effects like resolution.

We wish to avoid introducing dependence on the physics models of the neutrino interactions used by GENIE through the unfolding. Therefore, we unfold only the observable quantities E_e and θ_e (and since the inferred quantities E_ν^{QE} and Q_{QE}^2

can both be calculated from these alone, we unfold E_ν^{QE} and Q_{QE}^2 to the true E_ν^{QE} and Q_{QE}^2 calculated from true E_e and θ_e instead of to the true E_ν and Q^2 used by the generator). The migration matrices as predicted by the simulation are shown in fig. 7.17. Though it appears that there is a bias in the reconstructed electron energy (systematically too large), this is due to the 5% energy scale correction applied to the simulation to make it agree with the data (introduced in sec. 7.3).

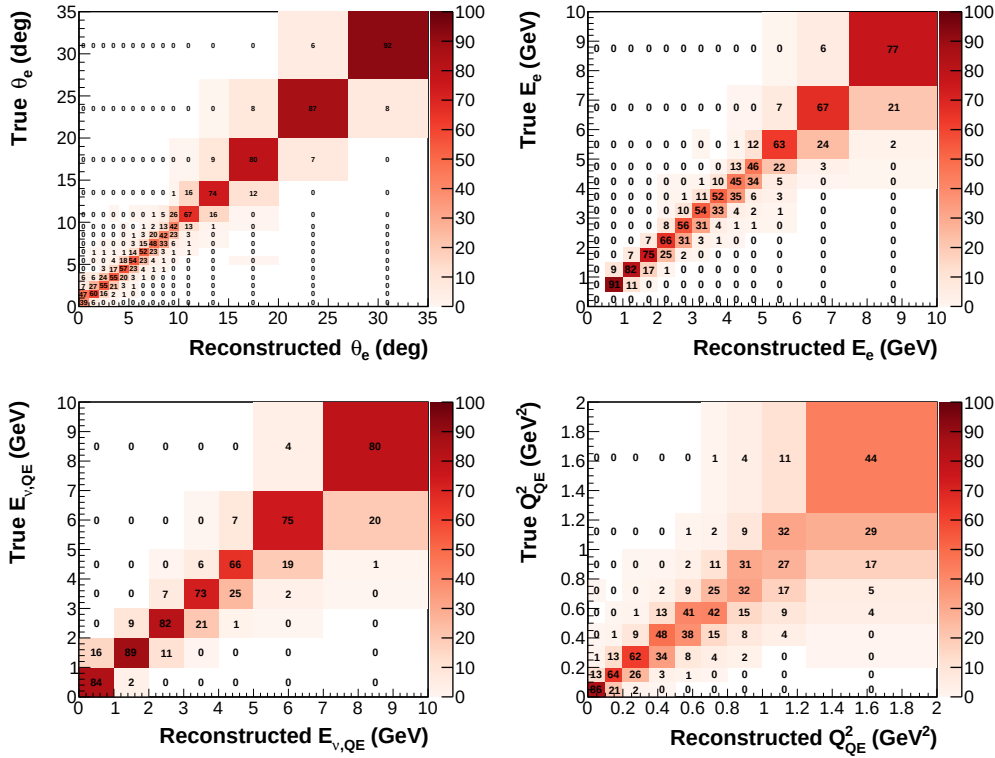


Figure 7.17: Migration matrices predicted by the simulation. Entries are normalized by column (for a given reconstructed event, what are the probabilities of it having a particular true energy?); values are in percent.

We unfold using the Bayesian unfolding method first proposed by d’Agostini [85] and implemented in RooUnfold [86]. To determine the optimum number of iterations of unfolding we should use, we examined the progression of the residuals from the true value in the simulation as a function of unfolding iteration, and compared them to the progression of statistical errors; these are shown in fig. 7.18. (Note that the wild fluctuations in bin 1 of the electron angle distributions arise

from the very small number of events in that bin; the uncertainties on this quantity account for this behavior, as discussed further below.) Because the statistical uncertainties and residuals both increase after a single iteration in most bins of most distributions, we conclude that a single iteration is ideal. (As RooUnfold uses the true distribution as its prior for the first iteration [86], this is not completely unexpected.) The distributions resulting from one iteration of Bayesian unfolding are shown in fig. 7.19.

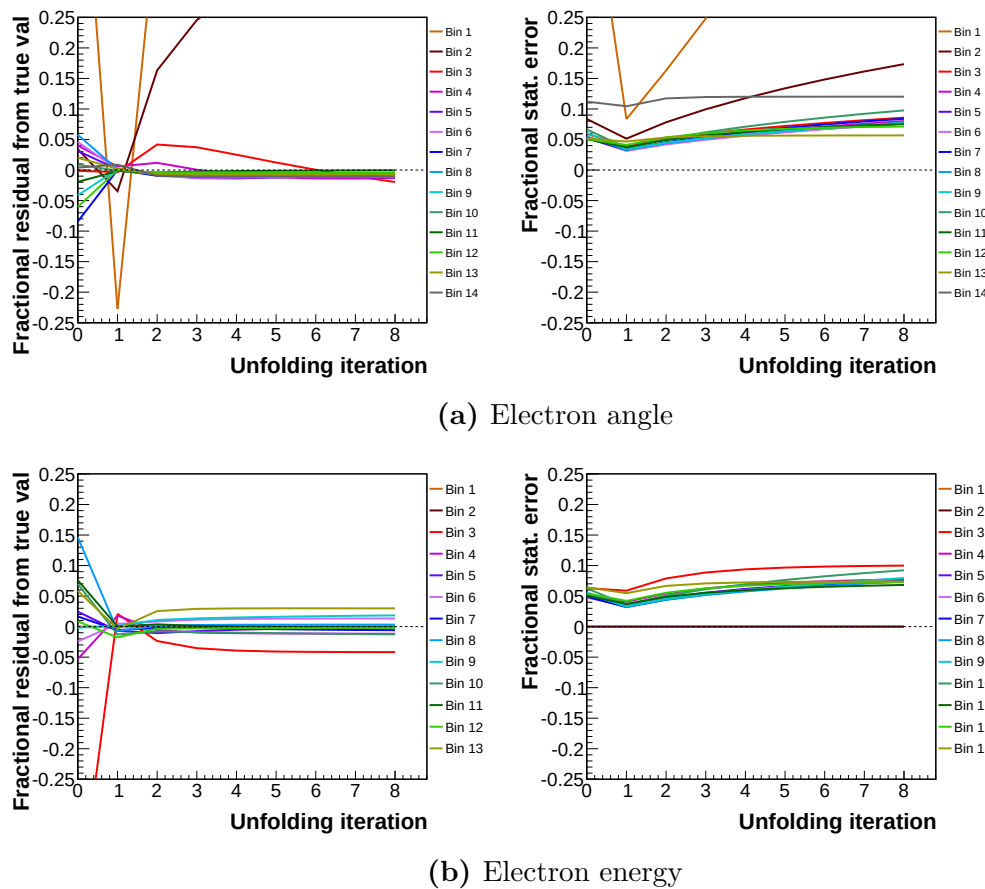
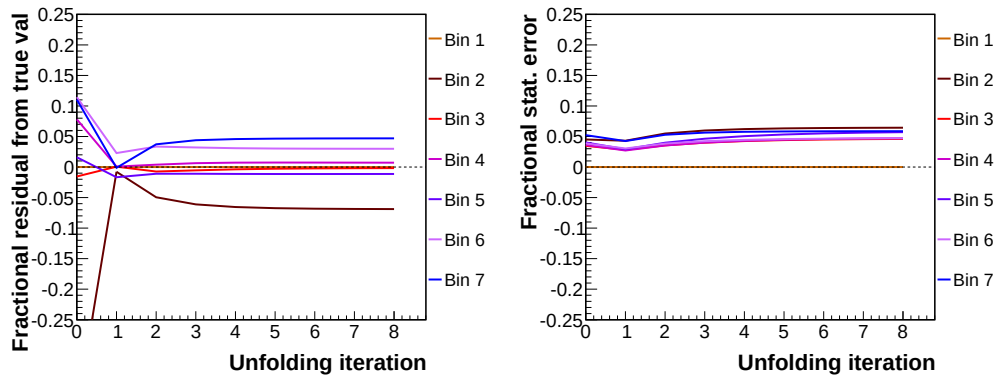


Figure 7.18: Residuals from true value (left) and statistical uncertainties (right) as a function of number of iterations of unfolding.

To further establish that 1 iteration is a sensible choice, we performed a pseudodata study in which the simulated event distributions were warped to produce a different reconstructed distribution. Comparing the result of the unfolding on the



(c) Neutrino energy (QE hypothesis)

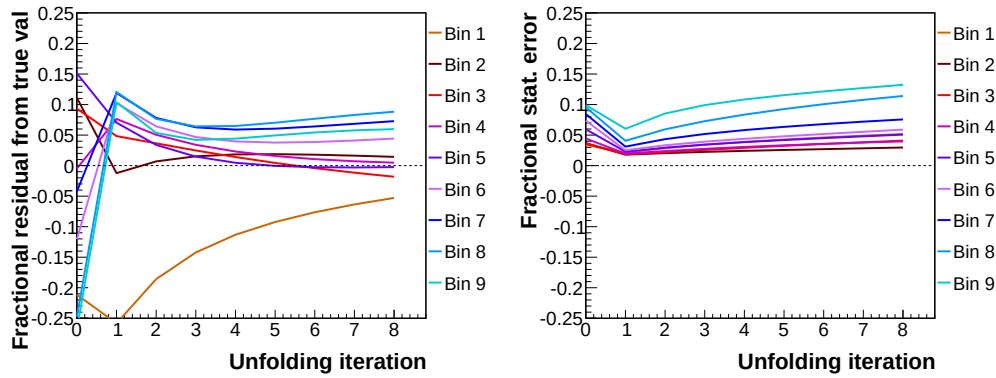
(d) Q^2 (QE hypothesis)

Figure 7.18: Unfolding residuals (left) and statistical uncertainties (right) as a function of number of iterations of unfolding.

warped distribution to the underlying true distribution then provides a measure of any bias introduced by the unfolding (though we expect that any such effect would be unimportant anyway since the simulated distributions are very similar to the data ones even before unfolding). We chose the reweighting function depicted in fig. 7.20 based on an earlier version of the analysis in which there was a significant difference in shape in electron angle. We then threw 1000 Poisson variations from the resultant distributions, where the weighted central value was taken as the Poisson mean in each bin. From these variations we calculated the pull g in each

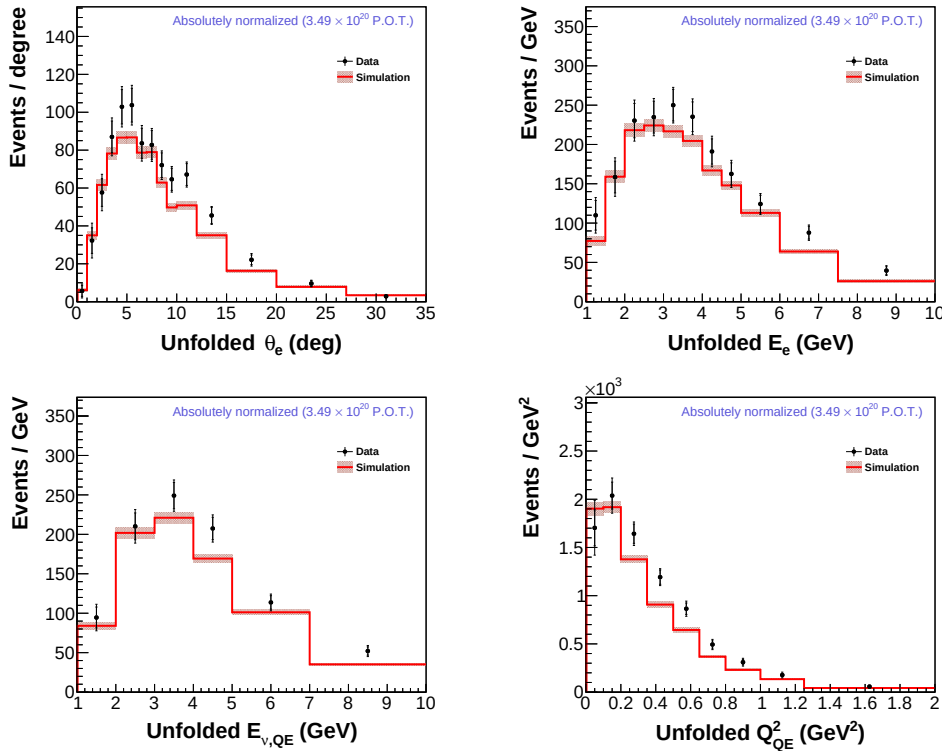


Figure 7.19: Distributions in the cross-section variables after unfolding.

bin i of each variable ξ :

$$g_i = \frac{\xi_i^{\text{unfolded}} - \xi_i^{\text{true}}}{\sigma_i^\xi} \quad (7.4)$$

where σ_i^ξ is the uncertainty on the unfolded value in bin i of variable ξ . If the unfolding is unbiased, then the mean of the pull distribution in every bin should be 0; if the (statistical) errors are correctly estimated by the unfolding procedure, then the standard deviation of the pull distribution in every bin should be 1. Both the pulls themselves and the moments of the pull distributions are plotted in fig. 7.21. Based on these plots, we conclude that the performance of the unfolding is not significantly biased and that the errors are reasonable.

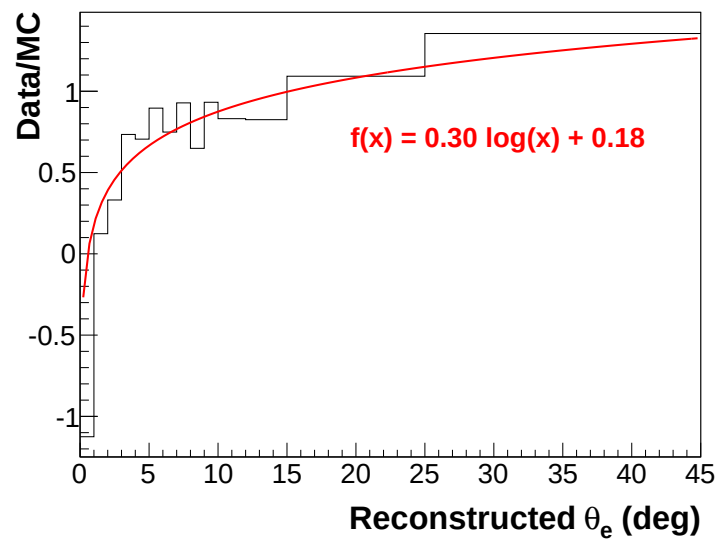
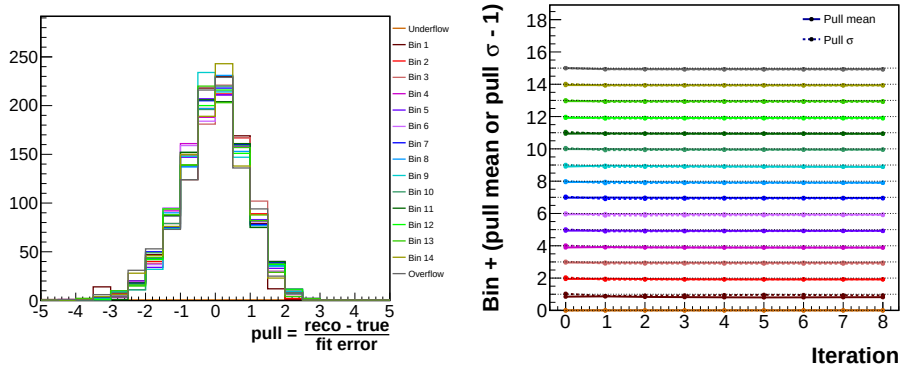
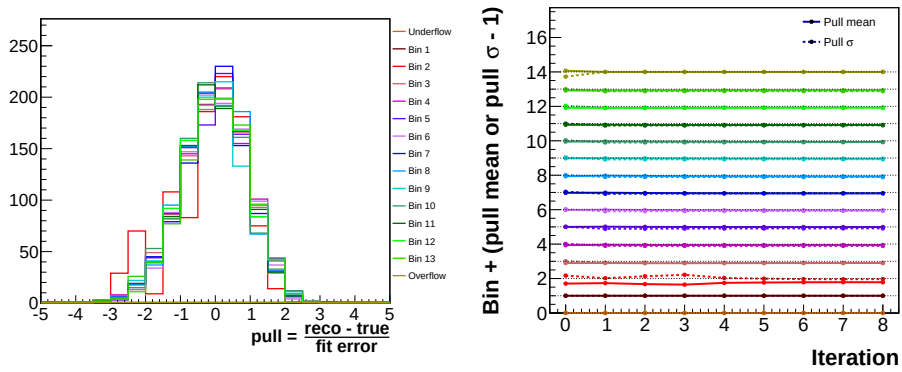


Figure 7.20: Reweighting function used to create the warped distributions for pseudo-data studies.

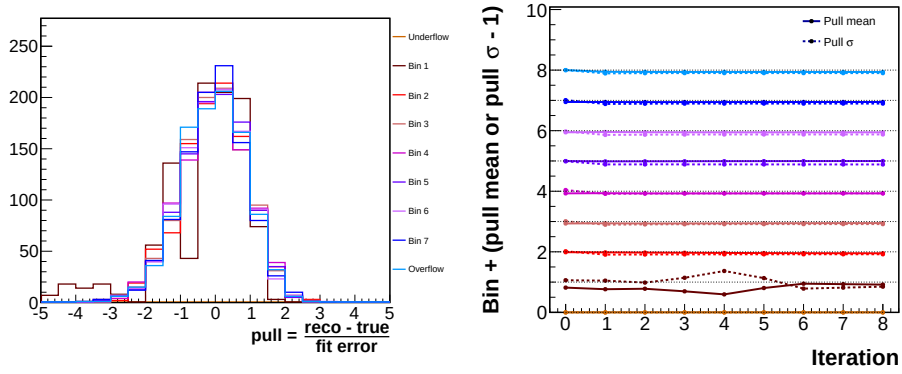


(a) Electron angle



(b) Electron energy

Figure 7.21: Pull distributions after one iteration (left); mean and standard deviation of pulls as a function of iteration (right). In the right plots, the mean (solid curve) and standard deviation (dashed curve) will both make a straight line at the bin number if the unfolding is unbiased and the uncertainties are correctly estimated.



(c) Neutrino energy (QE hypothesis)

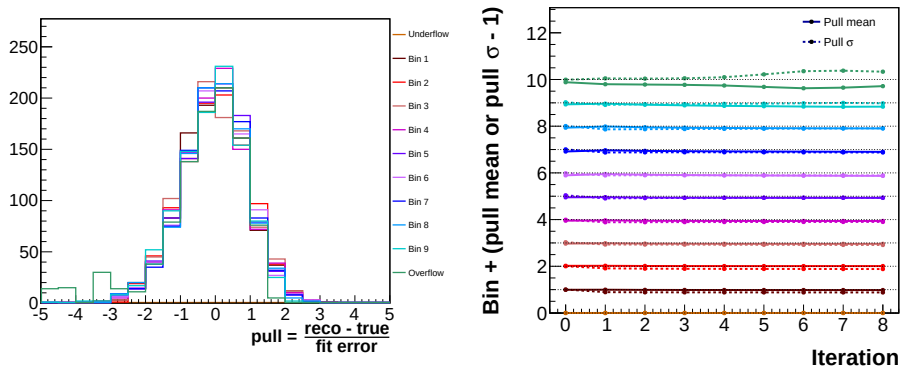
(d) Q^2 (QE hypothesis)

Figure 7.21: Pull distributions after one iteration (left); mean and standard deviation of pulls as a function of iteration (right). In the right plots, the mean (solid curve) and standard deviation (dashed curve) will both make a straight line at the bin number if the unfolding is unbiased and the uncertainties are correctly estimated.

7.6. Correction for inefficiency

Once the distributions have been unfolded to our best estimate of their values, we then use the simulation to correct for events which were truly signal events but lost by the reconstruction or selection cuts. We measure the efficiency in some variable ξ , as predicted by the simulation, by dividing the predicted distribution of selected true signal events in ξ by the predicted distribution of all true signal events in ξ . The nominal efficiencies calculated in this way are shown in fig. 7.22.

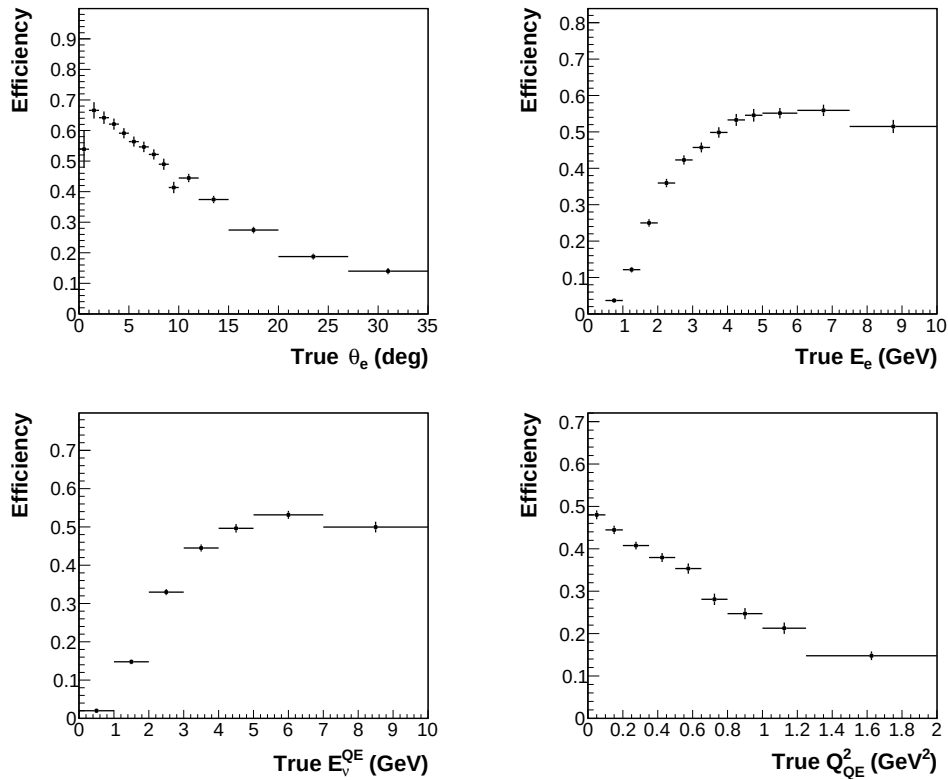


Figure 7.22: Efficiency predictions in the cross-section variables.

7.7. Normalization by flux and target number

The final steps in the calculation of the cross-section are normalization (the remaining factors are in the denominator in eq. 7.1). First, we divide by an

estimate of the number of scattering targets (N_t) available to neutrinos in the fiducial volume of sec. 7.1.2. We will define our results to be given per CH molecule in the fiducial volume. Therefore, we wish to know the number of CH molecules inside a volume shaped like a hexagonal prism with apothem 850 mm and z-extent from $z = 5990$ mm to $z = 8340$ mm; since these z -positions do not correspond exactly to the edges of detector planes ($z = 5990$ mm contains only the latter 42% of module 27, plane 1, and $z = 8340$ mm contains only 31% of module 79, plane 1), this corresponds to 103.73 inner detector planes.² This results in 2.306×10^{29} Carbon atoms (and therefore CH molecules) in the simulated detector volume, and 2.308×10^{29} Carbon atoms in the actual detector volume. (The detector geometry used in the simulation is slightly different than the current best estimate for the actual detector, and we must be consistent with what was used to generate the neutrino interactions.)

We have already examined the flux prediction, Φ , in detail in sec. 3. What remains is to establish its manner of use. For the differential cross-sections, we numerically integrate the flux histogram in fig. 3.7 across the entire energy range of the measurement, that is, 0-10 GeV. The total cross-section, by contrast, requires the flux in each bin of neutrino energy of the cross-section; so, we integrate fig. 3.7 in the region of each output bin, using a spline to interpolate when the bin edges do not coincide.

Finally, each bin i of the differential cross-sections are normalized by their bin widths Δ_i (which effectively causes the flux we just computed to become the average flux across 0-10 GeV). (As mentioned in the introductory material to ch. 7, we do not do this for the total cross-section since both numerator and the flux in the denominator are measured as a function of neutrino energy.)

²Though it may seem unusual to choose a fiducial volume definition with partial planes in it, this is the definition used by preceding MINER ν A analyses, and we have retained it to be consistent.

7.8. Measured cross-sections

The cross-sections we obtain after all of the preceding steps are shown in fig. 7.23.

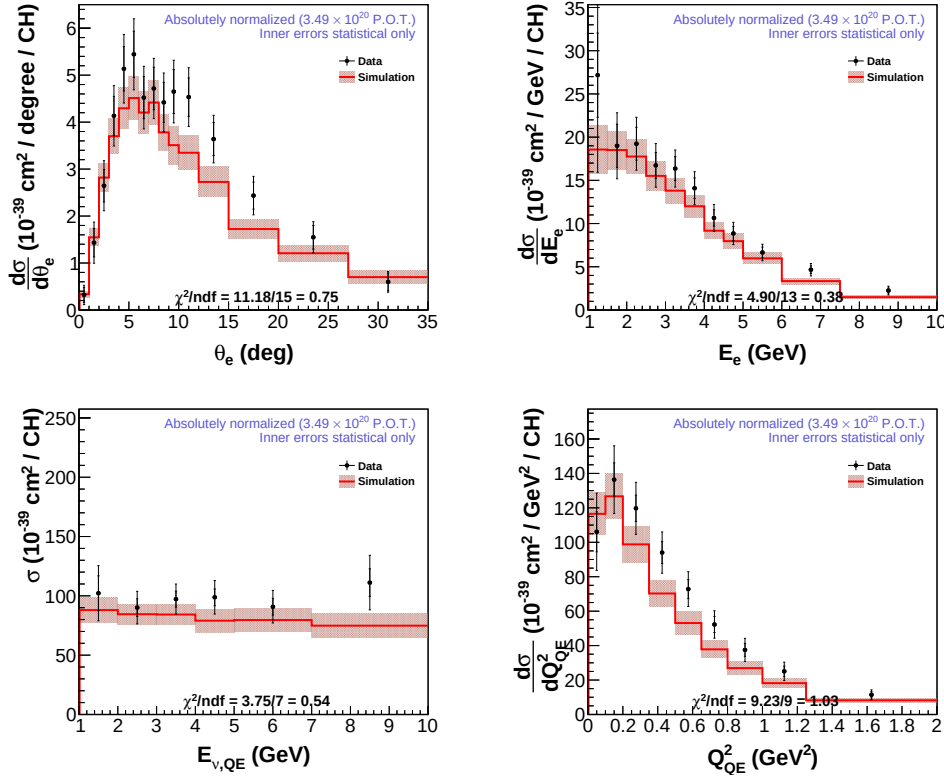


Figure 7.23: Measured cross-sections.

7.8.1. Closure test

To ensure that the calculations described in secs. 7.1-7.7 were carried out correctly, we performed a closure test in which we compared our measured cross-section from the simulation to the cross-section that the generator used when generating events. Because the $\nu + \bar{\nu}$ CCQE-like cross-section is not a process for which the generator has a closed-form cross-section, however, we used computed the effective cross-section for this process by extracting all the true signal events generated by the generator and normalizing by the flux and effective exposure.

One important modification to the analysis was necessary to compare the two cross-sections. The application of the flux constraint (sec. 3.4) is designed in such a way that it can—and indeed, should—alter the shape of the distributions to which it is applied. We therefore must not use it, either for the event rate distributions or for the flux prediction, when comparing to the extracted GENIE cross-section.

Ratios of our calculated cross-section to that of GENIE are shown in fig. 7.24. As is clear from the figure, we were not able to obtain perfect closure; we suspect that there is some minor residual difference in the event selections that we were not able to find. Since any potential error here is substantially smaller than the uncertainties we will be quoting on the result (see sec. 7.10), we choose not to pursue it any further.

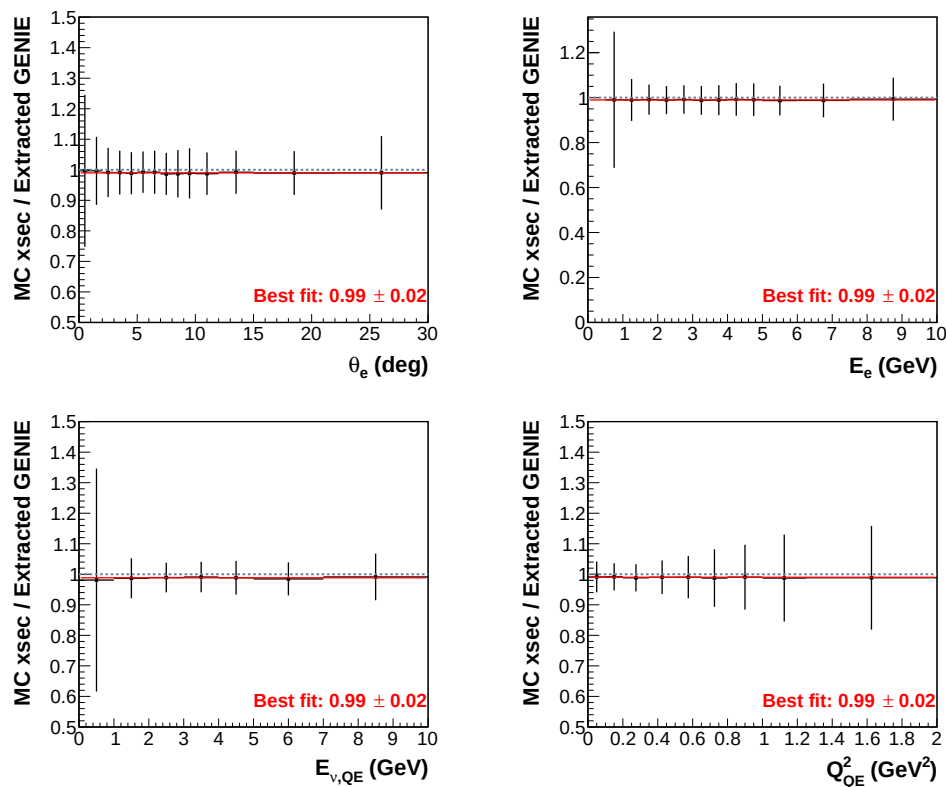


Figure 7.24: Closure tests comparing measured cross-section to the GENIE cross-section the simulation began with. The error bars are not meaningful and serve mainly to demarcate the location of the points.

7.9. Ratios to ν_μ CCQE

We have presented the ν_e CCQE cross-section measurement as being of critical importance because of its centrality in the neutrino oscillation arena. However, as we noted in ch. 1, because of the strong expectation of lepton universality, current experiments typically use the corresponding ν_μ cross-section as a placeholder. Therefore, the ratio between the two cross-sections serves as a useful indication of the potential effect that the substitution that is usually made could have on oscillation results. Furthermore, if we can make a comparison to a measurement that was made using the same detector and flux, we can potentially cancel some of the systematic uncertainties discussed in sec. 7.10 that are common to the measurements.

7.9.1. Ensuring consistency with ν_μ cross-section

We will compare to the previously-published measurement of ν_μ CCQE from MINER ν A [21]. To do so, we will need to make some alterations to the analysis to ensure maximum consistency between the two results.

Signal definition

The prior ν_μ result was computed using the GENIE generator’s identification of events as CCQE for its definition of signal. As a result, it contains no events that originated as non-quasi-elastic events in GENIE (e.g., baryon resonances that decayed to a final state including pions) and are phenomenologically identical to CCQE because of the effect of final-state interactions. This has the effect of reducing the signal component of the selected sample by about 10% from that which would be expected for a “CCQE-like” definition; see fig. 7.25.

For the ratio to be meaningful, we must employ a ν_e cross-section using the

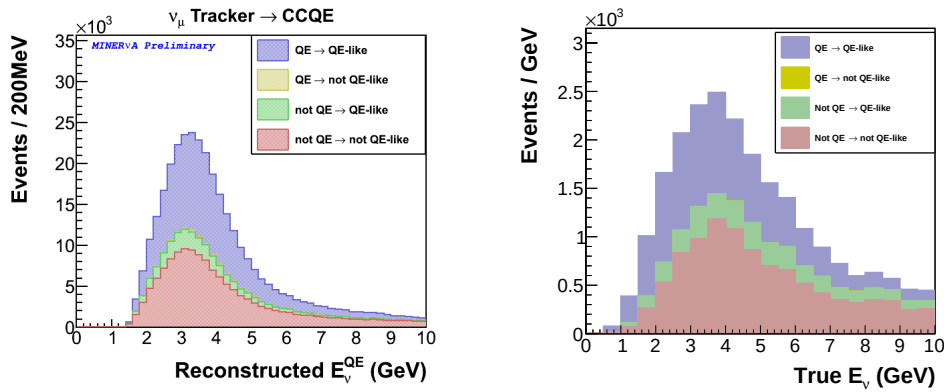


Figure 7.25: The breakdown of events in the selected ν_μ (left) and ν_e (right) CCQE samples, separated by their “quasielastic-likeness.” ν_μ version courtesy C. Patrick (MINERνA).

same definition of CCQE in the numerator. Therefore, we re-perform the analysis with this in mind; the background scale factors (secs. 7.4 and 7.4.3), the efficiency corrections (sec. 7.6), and of course the end cross-section all change as a result. Please see the plots in Appendix C (which can be compared to those in the sections mentioned above).

Flux constraint and uncertainties

The flux constraint described in sec. 3.4 was not used for the ν_μ result. Because we would like to benefit from the partial correlation of the ν_μ and ν_e fluxes, and thus the partial cancellation of their uncertainties, we cannot apply the flux constraint to the ν_e cross-sections either. Therefore we prepare a version of the result without it.

Ratio and discussion

The measured ratio in $\frac{d\sigma}{dQ^2}$ between the ν_e and ν_μ results, with the uncertainties described in sec. 7.10 applied to the ν_e cross-section and those mentioned in ref. [21] applied to the ν_μ cross-section, is shown in fig. 7.26.

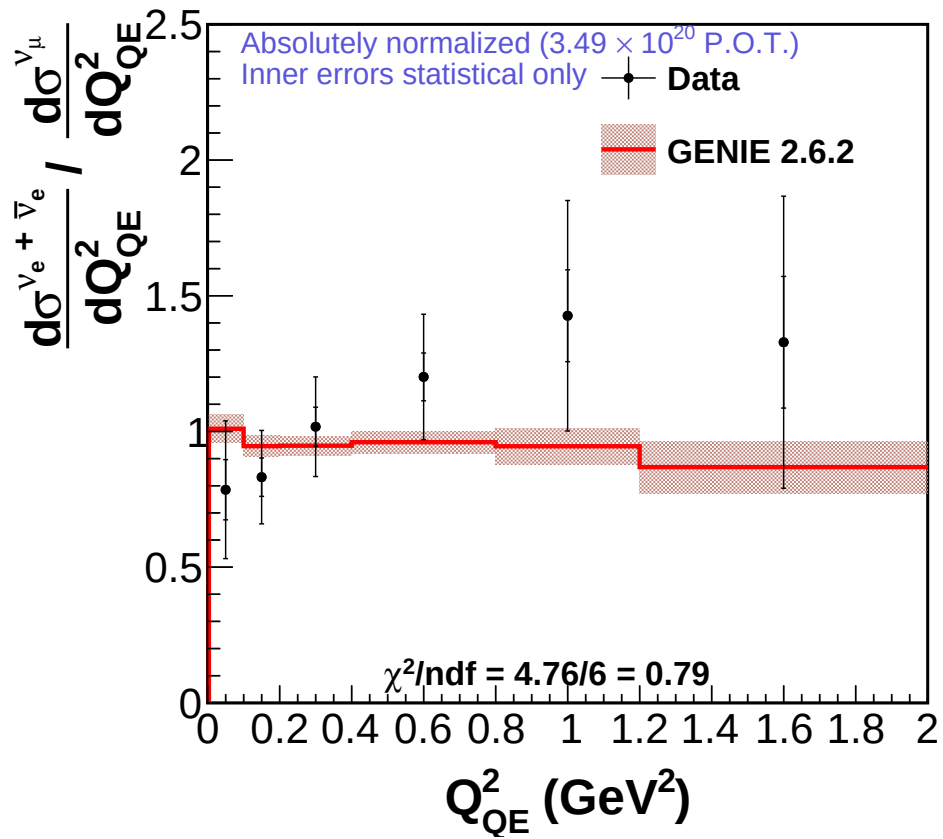


Figure 7.26: The ratio $\frac{d\sigma_{\nu_e}}{dQ_{QE}^2} / \frac{d\sigma_{\nu_\mu}}{dQ_{QE}^2}$. The red line indicates the GENIE prediction, while points represent the MINER ν A measurement (inner errors statistical; outer are statistical summed in quadrature with the systematic uncertainties enumerated in sec. 7.10).

7.10. Systematic uncertainties

The uncertainties we estimate for this result can be grouped along three broad categories: uncertainties in the primary neutrino interaction model, uncertainties in the flux model, and uncertainties in our model of the detector’s response to particle activity.

One special technique is used occasionally in the estimation of systematic uncertainties in this analysis and warrants special attention here. We will term it “many universes”; this strategy simultaneously varies multiple parameters $\{p_\alpha\}$ by throwing a large number (typically 1000) of collections of their values from Gaussian distributions $g(\mu_\alpha, \sigma_\alpha)$ using pre-determined estimates for the means μ_α and uncertainties σ_α of each parameter p_α . Each throw of the parameters $\{p_\alpha\}$ is then used to compute the values of a set of fundamental variables $\{\xi_\alpha\}$ (for example, a cross-section used at generation time), and events generated by those parameters are assigned weights corresponding to (for example) the ratio of the $\{\xi_\alpha\}$ to the nominal value. We can use the weights for each event to construct alternate histograms for each derived quantity in which we are interested, and the covariance of these many histograms provides our uncertainty estimate.

The manner in which the various systematics were evaluated is detailed below. Summaries of their sizes (with respect to the final result) are shown in figs. 7.27 and 7.28.

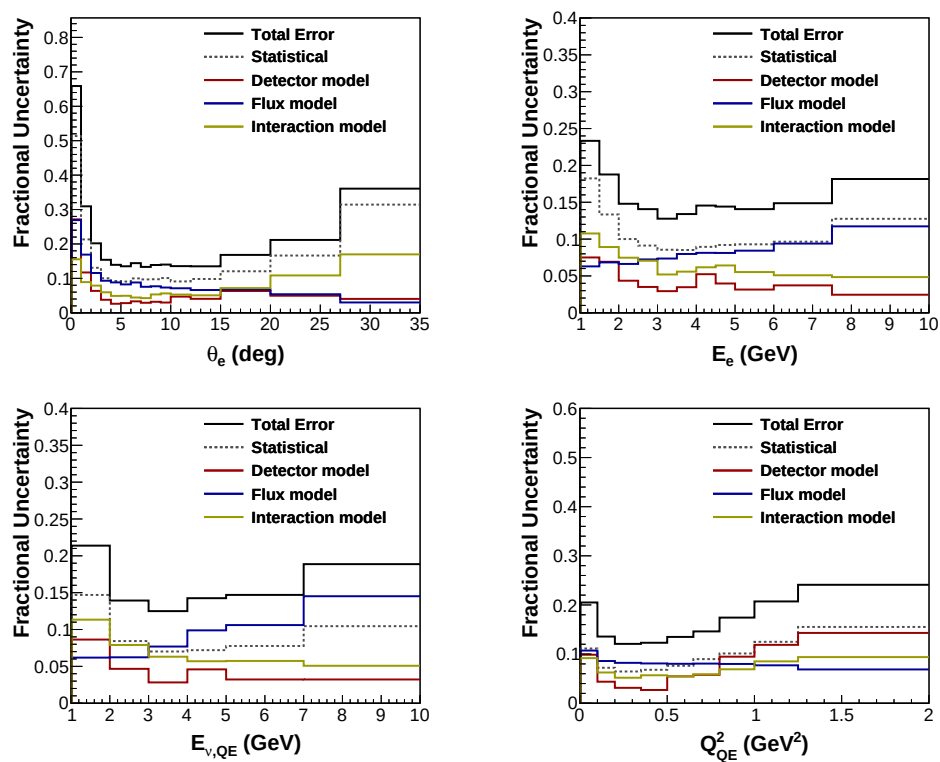


Figure 7.27: Uncertainties on the cross-section measurements given as fractions of the central value bin content.

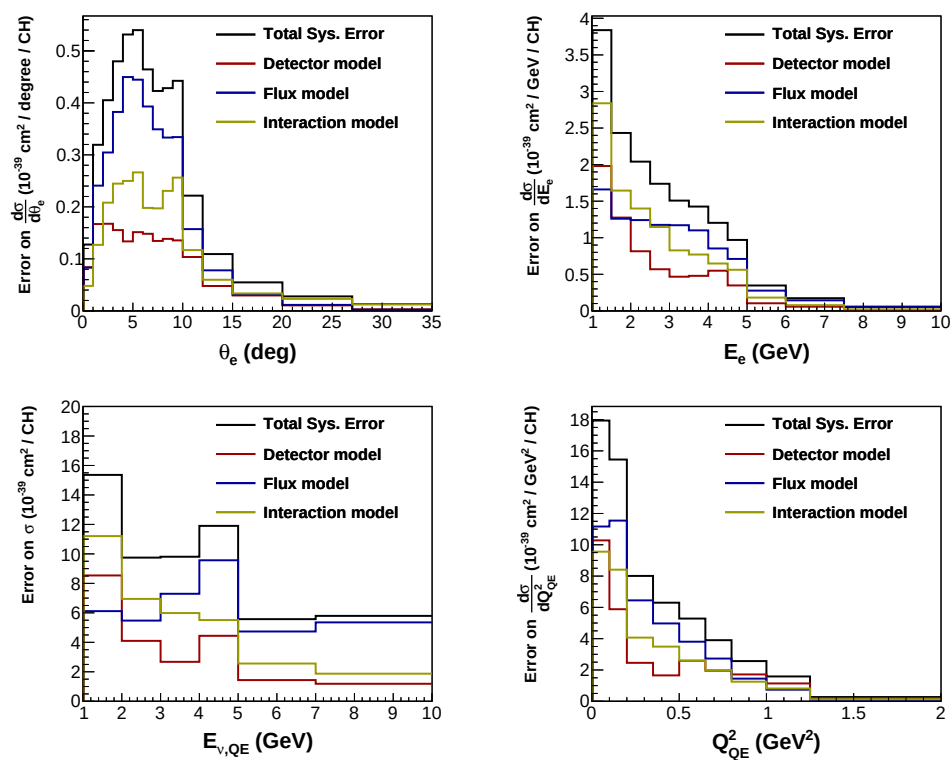


Figure 7.28: Uncertainties on the cross-section measurements.

7.10.1. Interaction model

A breakdown the uncertainties in the interaction model group is given in fig. 7.29; further detail on how each band is computed is given in the sections below.

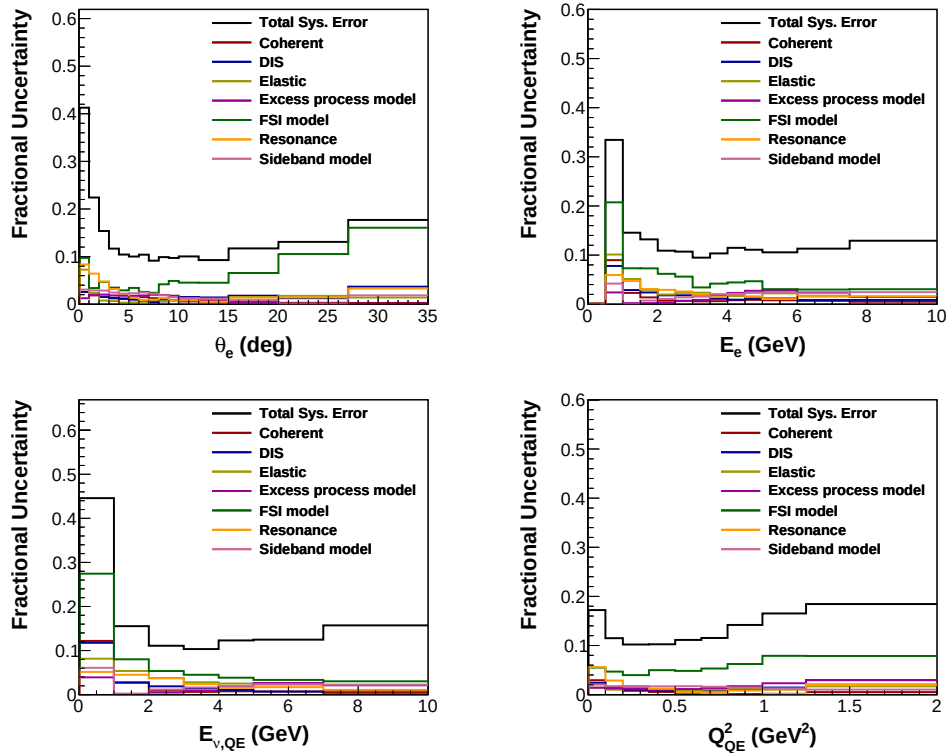


Figure 7.29: Uncertainties on the cross-section measurements due to the neutrino interaction model given as fractions of the central value bin content. (Note that the “total” curve includes systematic uncertainties from all groups, not just the interaction model.)

Primary interaction model: GENIE

Uncertainties in the GENIE generator’s model of neutrino interactions enter primarily through the background prediction detailed in sec. 7.4 and the efficiency correction of sec. 7.6. To estimate most of them, we use event-by-event weights generated for $\pm 1\sigma$ shifts for a number of uncertainties estimated by the GENIE collaboration. The parameters being varied include shape and normalization knobs for elastic and resonance productions (both charged- and neutral-current); nuclear

model parameters which principally affect deep inelastic scattering; and parameters which control the strength and behavior of final-state interactions. A full list of the parameters, their nominal values (where available), and the uncertainty used for them is given in tab. 7.2. (Fuller discussion of each of these knobs is available in the GENIE physics manual [87].)

Parameter	Nominal value	Uncertainty
<i>NC elastic model uncertainties</i>		
Axial mass M_A (GeV/c^2)	0.99	$\pm 25\%$
Strange axial FF parameter η	0.12	$\pm 30\%$
<i>CCQE model uncertainties</i>		
Axial mass M_A (GeV/c^2)	0.99	+25% / -15%
Vector form factor model	—	Changes model to dipole from BBBA parameterization
Pauli-blocking momentum cutoff		$\pm 30\%$
<i>Resonant production model uncertainties</i>		
Axial mass M_A^{res} (GeV/c^2)	1.12	$\pm 20\%$
Vector mass M_V^{res} (GeV/c^2)	0.84	$\pm 10\%$
<i>Non-resonant pion production / DIS model uncertainties</i>		
νn or $\bar{\nu} p \rightarrow 1\pi + X$ s.f.	0.3	$\pm 50\%$
νn or $\bar{\nu} p \rightarrow 2\pi + X$ s.f.	1.0	$\pm 50\%$
$\bar{\nu} n$ or $\nu p \rightarrow 1\pi + X$ s.f.	0.1	$\pm 50\%$
$\bar{\nu} n$ or $\nu p \rightarrow 2\pi + X$ s.f.	1.0	$\pm 50\%$
Bodek-Yang parameter A_{HT}	0.538	$\pm 25\%$
Bodek-Yang parameter B_{HT}	0.305	$\pm 25\%$
Bodek-Yang parameter C_{V1u}	0.291	$\pm 30\%$
Bodek-Yang parameter C_{V2u}	0.189	$\pm 30\%$
<i>Hadronic system / FSI model uncertainties</i>		
Nucleon mean free path		$\pm 20\%$
Nucleon absorption probability		$\pm 20\%$
Nucleon charge-exchange probability		$\pm 50\%$
Nucleon elastic scattering probability		$\pm 30\%$

Parameter	Nominal value	Uncertainty
Nucleon inelastic scattering probability		$\pm 30 \%$
Pion mean free path		$\pm 20 \%$
Pion absorption probability		$\pm 30 \%$
Pion charge-exchange probability		$\pm 50 \%$
Pion elastic scattering probability		$\pm 10 \%$
Pion inelastic scattering probability		$\pm 10 \%$
Pion production probability from pions		$\pm 20 \%$
Pion production probability from nucleons		$\pm 20 \%$
Δ decay isotropy	—	Change to non-isotropic
Δ photonic decay branching ratio		$\pm 50 \%$
x_F distribution in AGKY hadronization model		$\pm 20\%$

Table 7.2: Uncertainties on GENIE model parameters. Nominal values for the FSI parameters in GENIE are not published anywhere and so are not given here.

We also estimated the uncertainty due to a GENIE parameter that has no associated event weights: the “effective nuclear radius,” which controls the distance over which newly-formed hadrons traveling out from the neutrino interaction vertex can undergo final-state interactions within that nucleus. To perform the estimate, we generated a very large sample of GENIE events using alternate values of this parameter. (The nominal values are 0.9 fm for nucleons and 0.6 fm for pions; we used 0.4 fm and 0.0 fm for a downward variation, respectively, and 1.6 fm and 1.1 fm for our upwards variation.) We then compared the true kinematics of outgoing hadrons in the central value simulation to those in the modified simulation, and generated reweighting functions from them; these reweighting functions are illustrated in fig. 7.30. (The same smoothing technique discussed in appendix A.2.2 was applied to the two-dimensional plots here as well.) We applied these weights to

the selected samples of sec. 7.3 to determine the effect on the cross-section result.

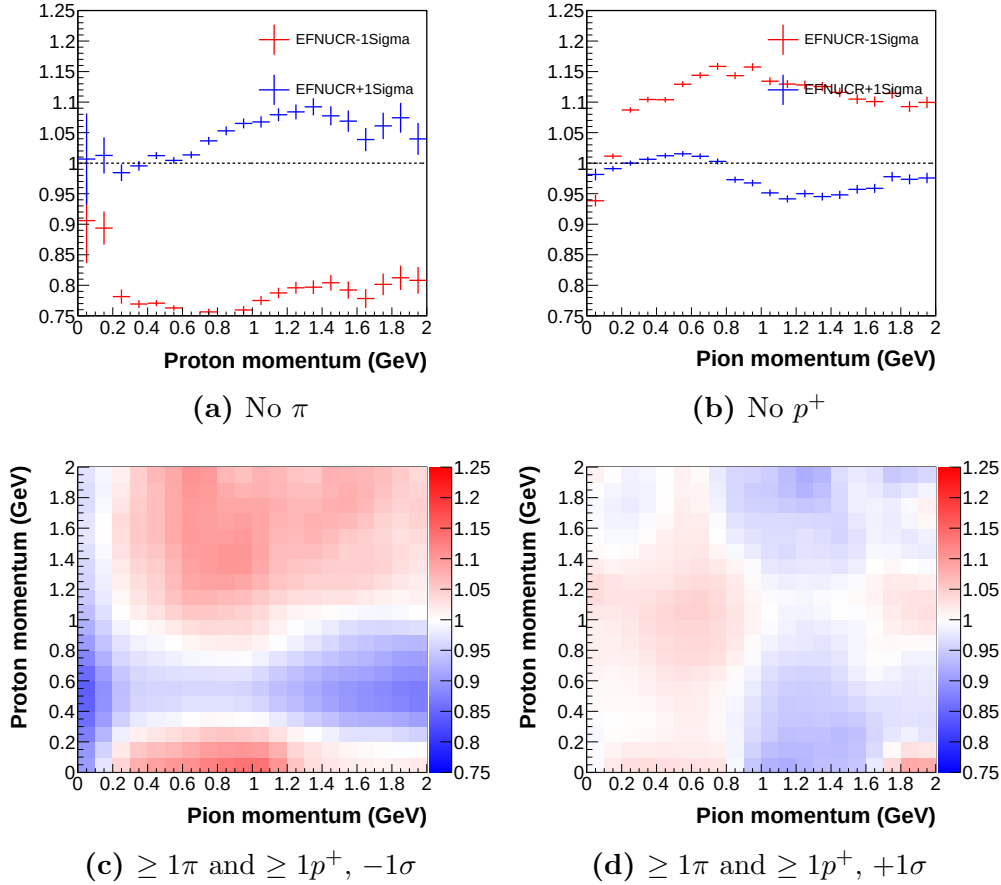


Figure 7.30: Weights applied for the effective nuclear radius variations.

It is important to appreciate that all of the *a priori* GENIE uncertainties computed above are significantly reduced by applying the background constraints of sec. 7.4.

Background fitting (sideband model)

Our uncertainty budget must account for the statistical uncertainty associated with the scale factors computed in sec. 7.4.3. To do so, we consider variations of the background prediction where the scale factors are shifted relative to their nominal values consistent with the uncertainties shown in tab. 7.1. To correctly account

for the correlations between the scale factors, however, we obtain the variations in the parameters from the eigenvectors of the statistical covariance matrix returned by the fitter, multiplied by the eigenvalues and the factor of $\sqrt{2}$ explained in sec. 7.4.3.

Coherent neutrino interaction model

Based on the measurement uncertainties and the residual disagreement between the data and the model after the corrections described in sec. 7.4, we assume a flat 20% uncertainty on the NC coherent distributions predicted in the signal and sideband regions. (Note that without this uncertainty, the coherent process would not be ascribed *any* uncertainty, since GENIE’s uncertainty knobs, described in sec. 7.10, do not vary the coherent cross-section at all.)

Excess process model

We selected the neutral pion model from sec. B.3 to model the excess process in the central value prediction in sec. 7.4. Therefore, we consider the single photon model as a variation to generate the uncertainty on the background due to the excess. (As noted there, coherent η production is predicted not to exist; since it is furthermore disfavored by the particle identification variables, we concluded that it is not a sensible variation to use here.)

7.10.2. Flux model

Uncertainties in the flux model are estimated using the “many universes” method; this is discussed in more detail in sec. 3.3. As noted above in sec. 3.4, these uncertainties are somewhat reduced by the application of the flux constraint.

7.10.3. Detector response

Uncertainties on the response of the detector to the passage of charged particles can be further divided into a handful of categories. A summary is shown in fig. 7.31, and more detail on each individual band is given below.

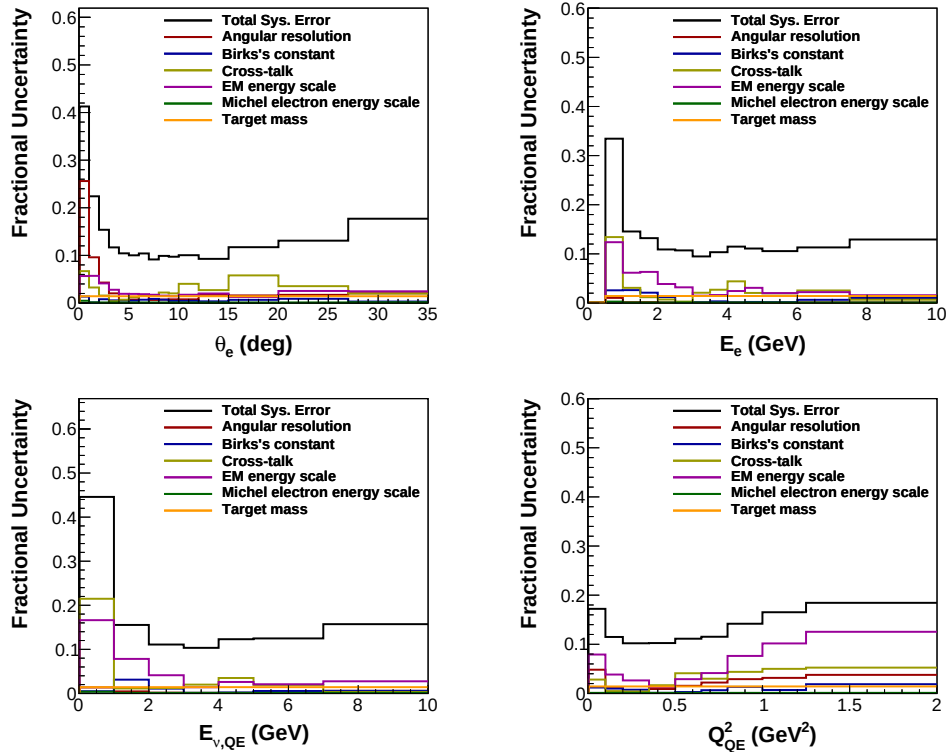


Figure 7.31: Uncertainties on the cross-section measurements due to the detector model given as fractions of the central value bin content. (Note that the “total” curve includes systematic uncertainties from all groups, not just the detector model.)

Angular resolution modeling

Because the angle of the electron is identified with the angle of the reconstructed Track that was first reconstructed, we associate any uncertainty in the modeling of the resolution of the angle with the angular resolution of Tracks. This was estimated to be 1 mrad [63]. To evaluate the impact of this uncertainty, we construct 1000 universes in which the electron angle is shifted by an amount drawn from a

Gaussian centered at 0 and of width 0.057 deg (= 1 mrad)). (This accounts for changes to E_ν^{QE} and Q_{QE}^2 resulting from the difference in the angle.)

Birks's constant

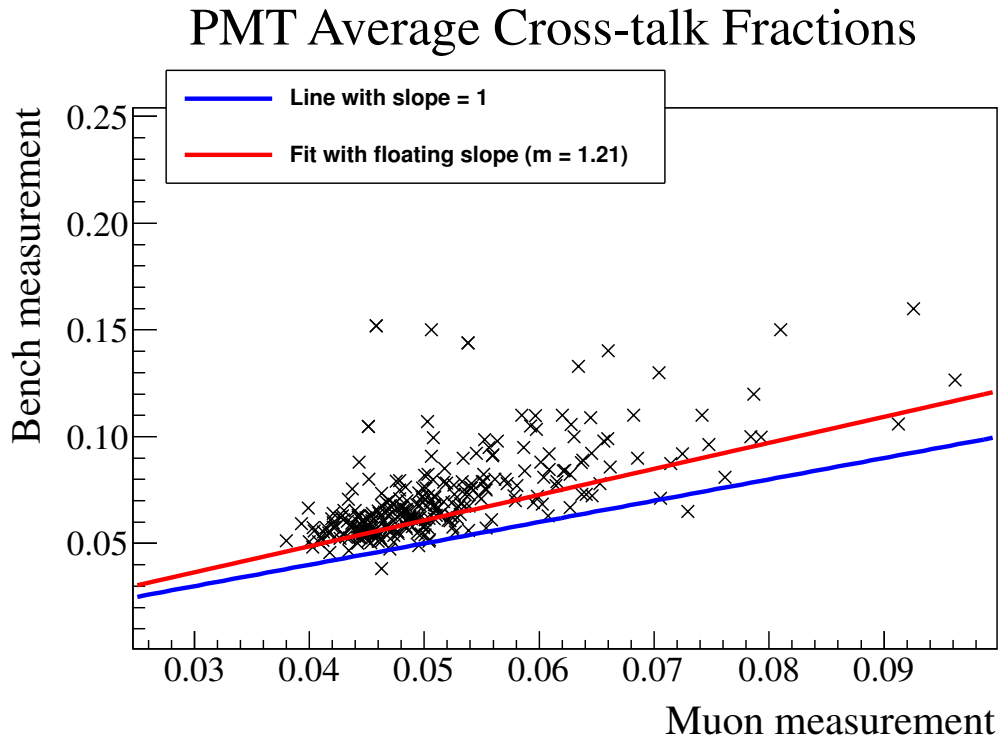
There was some uncertainty associated with the value of the independent parameter used in the simulation of Birks's Law (discussed further in sec. 5.2.3). We used a central value of $k_B = 0.133$ m/GeV, obtained from MINOS [72]; the uncertainty was estimated to be 30%. To examine the effect of this uncertainty, we generated varied samples where the energy depositions of particles not interacting via electromagnetic cascades were varied $\pm 30\%$. (The uncertainty in the energy scale constraint of sec. 7.10.3 includes any variation in Birks's constant; therefore, we do not apply the Birks's constant uncertainty to EM activity here so as not to double-count.)

Cross-talk

Because cross-talk energy is dominated by optical cross-talk (see sec. 5.2.3) and because assembly variations affect optical cross-talk most strongly, we assume that the only significant cross-talk uncertainties are due to optical cross-talk. We factorize our estimate into two separate components.

First, we have a large set of PMTs that are both mounted on the detector and which underwent cross-talk bench characterization prior to installation. We compared the measurements in fig. 7.32. As noted in the figure, the best-fit slope to this distribution differs from 1 at the 20% level; therefore, we conclude that a 20% uncertainty on the overall (optical) cross-talk scale is appropriate.

Secondly, careful inspection of fig. 7.32 reveals a handful of PMTs whose muon-based measured muon cross-talk fraction substantially exceeds that from the bench test. We inspected the difference between the nominal cross-talk fraction in the simulation for a channel and what was actually seen in the *in situ* measurement,



1

Figure 7.32: Mean cross-talk fractions for PMTs as measured by the bench tests vs. the measurement from the muon measurement. The 20% difference between the best-fit slope and 1 implies that a 20% uncertainty on the overall cross-talk scale is reasonable.

and constructed a scale factor from each difference. Then, during simulation, we can randomly sample a smearing factor from the distribution of these scale factors and apply it, in addition to the nominal scale factor, to simulate a pair whose scaling factor was incorrectly measured too low. It was estimated that there is an upper limit of about 2% of channels which could be mis-tuned in this fashion.

The combined effect of this re-smearing of 2% of the channels, in tandem with the global scale shift, is illustrated in figure 7.33.

Detector mass

The detector mass and composition breakdowns are estimated to be uncertain at the level of 1.4%. [63] We create an error band corresponding to this uncertainty by varying the number of targets used in the denominator of eq. 7.1. (This means

Pulse Height of Hits Identified as Cross-talk

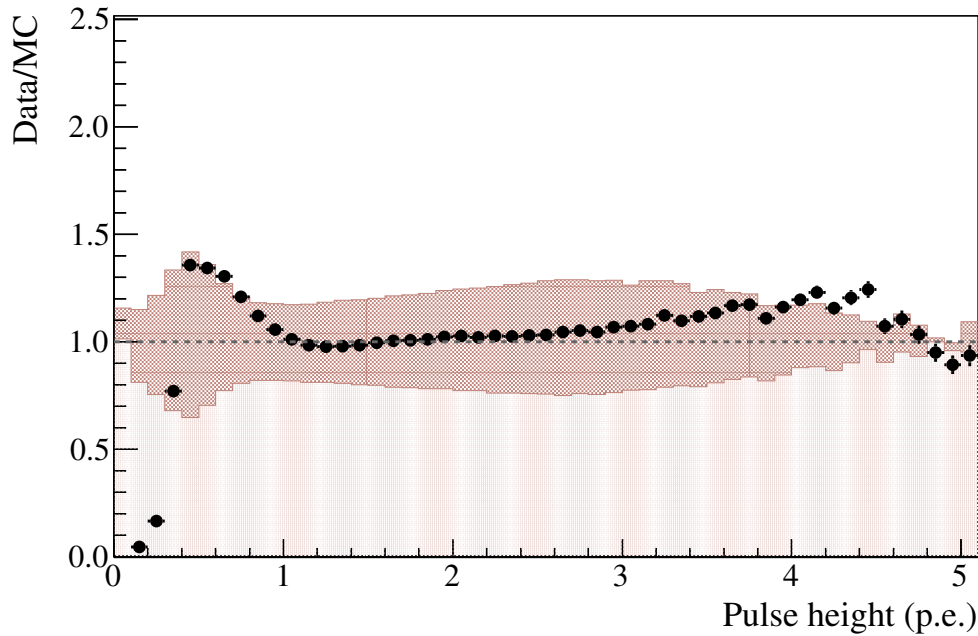


Figure 7.33: Distribution of pulse heights of hits identified as cross-talk, with a band illustrating the variations described in the text.

that the uncertainty is fully correlated across all the bins of each cross-section measurement.)

Electromagnetic energy scale

There is a 1.8% uncertainty associated with the extra 5% EM energy scaling we apply to the simulation (sec. 7.3). We therefore construct alternate versions of the result in which the electron energy is shifted by 3.2% and 6.8% on an event-by-event basis and compute the difference in the result to estimate the effect on the cross-sections.

Michel electron energy scale

The cuts made to eliminate any events with obvious Michel electrons (sec. 7.2.3) depend on the reconstructed energy of Michel electron candidates. While the data

energy scale for Michel electrons was determined to be about 3% larger than that in the simulation [63], no formal uncertainty on this parameter has been published. We assumed it to be uncertain at the level of 50% of itself; therefore, we evaluated the difference in the results when the energies of candidate Michel electrons in the simulation were multiplied by 1.015 and 1.045 to estimate the effect on the cross-sections.

8 Conclusions

We have presented here the first-ever cross-section measurements made exclusively of electron neutrino charged-current quasielastic scattering. A comparison of our measurement to a cross-section measurement of muon neutrino CCQE scattering vs. Q^2 made in the same detector using the same beam finds them to be statistically consistent, supporting the principle of lepton universality. We also compared the ν_e cross-sections themselves to the predictions (based on the Standard Model) in the GENIE generator, and found satisfactory agreement there as well. Given the community's reliance upon this correspondence in the neutrino oscillation campaign, the agreement we report in this channel is significant. Other studies considering different exclusive channels in electron neutrino scattering with MINER ν A could also be contemplated in order to bolster the verification of the cross-section models.

We have also observed an unpredicted π^0 -like process with similar characteristics to the electron neutrino quasi-elastic process. While it does not substantially affect this measurement, detectors for neutrino oscillation measurements which are unable to differentiate between photon- and electron-initiated electromagnetic showers (such as water Cherenkov detectors) could be susceptible to confusion from such a reaction. Further study aimed at fully characterizing and identifying the process in question in our event sample is ongoing.

Bibliography

- [1] M. E. Peskin and D. V. Schroeder, *An Introduction to Quantum Field Theory* (Perseus Books, Cambridge, MA, 1995).
- [2] D. J. Griffiths, *Introduction to Elementary Particles*, second ed. (Wiley-VCH, KGaA, Weinheim, Germany, 2008).
- [3] R. Pohl, R. Gilman, G. A. Miller, and K. Pachucki, *Ann.Rev.Nucl.Part.Sci.* **63**, 175 (2013), 1301.0905.
- [4] T2K Collaboration, K. Abe *et al.*, *Phys.Rev.Lett.* **113**, 241803 (2014), 1407.7389.
- [5] J. Blietschau *et al.*, *Nuclear Physics B* **133**, 205 (1978).
- [6] S. T. Thornton and J. B. Marion, *Classical Dynamics of Particles and Systems* (Thomson Learning, Belmont, CA, 2004).
- [7] C. Giunti and C. W. Kim, *Fundamentals of Neutrino Physics and Astrophysics* (Oxford University Press, New York, New York, 2007).
- [8] B. Cleveland *et al.*, *Astrophys.J.* **496**, 505 (1998).
- [9] B. Pontecorvo, *Sov.Phys.JETP* **7**, 172 (1958).
- [10] SNO Collaboration, Q. Ahmad *et al.*, *Phys.Rev.Lett.* **89**, 011301 (2002), nucl-ex/0204008.
- [11] J. Beringer *et al.*, *Phys. Rev. D* **86**, 010001 (2012).
- [12] T2K Collaboration, K. Abe *et al.*, *Phys.Rev.Lett.* **112**, 061802 (2014), 1311.4750.
- [13] C. Llewellyn Smith, *Phys.Rept.* **3**, 261 (1972).
- [14] R. Bradford, A. Bodek, H. S. Budd, and J. Arrington, *Nucl.Phys.Proc.Suppl.* **159**, 127 (2006), hep-ex/0602017.

- [15] J. Kelly, Phys.Rev. **C70**, 068202 (2004).
- [16] V. Bernard, L. Elouadrhiri, and U. Meissner, J.Phys. **G28**, R1 (2002), hep-ph/0107088.
- [17] H. Abele, Nucl.Instrum.Meth. **A440**, 499 (2000).
- [18] MiniBooNE Collaboration, A. Aguilar-Arevalo *et al.*, Phys.Rev.Lett. **100**, 032301 (2008), 0706.0926.
- [19] NOMAD Collaboration, V. Lyubushkin *et al.*, Eur.Phys.J. **C63**, 355 (2009), 0812.4543.
- [20] MINERvA Collaboration, L. Fields *et al.*, Phys.Rev.Lett. **111**, 022501 (2013), 1305.2234.
- [21] MINERvA Collaboration, G. Fiorentini *et al.*, Phys.Rev.Lett. **111**, 022502 (2013), 1305.2243.
- [22] R. Smith and E. Moniz, Nucl.Phys. **B43**, 605 (1972).
- [23] A. Bodek and J. L. Ritchie, Phys. Rev. D **24**, 1400 (1981).
- [24] A. M. Ankowski and J. T. Sobczyk, Phys.Rev. **C77**, 044311 (2008), 0711.2031.
- [25] O. Benhar, Nucl.Phys.Proc.Suppl. **229-232**, 174 (2012), 1012.2032.
- [26] M. Martini, M. Ericson, G. Chanfray, and J. Marteau, Phys.Rev. **C81**, 045502 (2010), 1002.4538.
- [27] J. Nieves, I. Ruiz Simo, and M. Vicente Vacas, Phys.Rev. **C83**, 045501 (2011), 1102.2777.
- [28] J. T. Sobczyk, Phys.Rev. **C86**, 015504 (2012), 1201.3673.
- [29] A. Bodek, H. Budd, and M. Christy, Eur.Phys.J. **C71**, 1726 (2011), 1106.0340.
- [30] J. Sobczyk, PoS **NUFACT08**, 141 (2008).
- [31] Y. Hayato, Acta Phys.Polon. **B40**, 2477 (2009).
- [32] GENIE collaboration, C. Andreopoulos *et al.*, Nucl. Instrum. Meth. **A614**, 87 (2010), 0905.2517.

- [33] M. Day and K. S. McFarland, Phys. Rev. D **86**, 053003 (2012).
- [34] CHARM Collaboration, J. Allaby *et al.*, Phys.Lett. **B179**, 301 (1986).
- [35] V. Ammosov *et al.*, Z.Phys. **C40**, 487 (1988).
- [36] H. Ballagh *et al.*, Phys.Lett. **B79**, 320 (1978).
- [37] C. Baltay *et al.*, Phys.Rev. **D41**, 2653 (1990).
- [38] T. Eichten *et al.*, Phys.Lett. **B46**, 281 (1973).
- [39] Bari-Birmingham-Brussels-Ecole Poly-Rutherford-Saclay-London Collaboration, O. Erriquez *et al.*, Phys.Lett. **B102**, 73 (1981).
- [40] R. L. Dixon, Cern Courier **14** (2011).
- [41] J. Thompson, FNAL Report No. FERMILAB-TM-1909, 1994 (unpublished).
- [42] B. Worthel, Booster rookie book, http://www-bdnew.fnal.gov/operations/rookie_books/Booster_V4.1.pdf, 2009.
- [43] B. Baller, NuMI rookie book, http://www-bdnew.fnal.gov/operations/rookie_books/NuMI_v1.pdf, 2005.
- [44] R. Zwaska *et al.*, Nuclear Instruments and Methods in Physics Research Section A: Accelerators, Spectrometers, Detectors and Associated Equipment **568**, 548 (2006).
- [45] NuMI technical design handbook, http://www-numi.fnal.gov/numwork/tdh/tdh_index.html, 2002.
- [46] S. Kopp, Fermilab-Conf-05-093-AD (2005), arXiv:physics/0508001v1 [physics.acc-ph].
- [47] K. T. McDonald, (2001), hep-ex/0111033.
- [48] A. I. Himmel, *Antineutrino Oscillations in the Atmospheric Sector*, PhD thesis, California Institute of Technology, Pasadena, California, 2011.
- [49] S. Agostinelli *et al.*, Nucl. Instr. Meth. A **506**, 250 (2003).

- [50] L. Aliaga, MINERA flux: Current uncertainties and future plans, <https://www.jlab.org/indico/contributionDisplay.py?contribId=51&confId=0>, 2012, Presentation at NuINT2012.
- [51] J. Apostolakis *et al.*, Nuclear Science Symposium Conference Record , 833 (2008).
- [52] MIPP Collaboration, J. Paley *et al.*, Phys.Rev. **D90**, 032001 (2014), 1404.5882.
- [53] C. Alt *et al.*, Eur. Phys. J. C **49**, 897 (2007).
- [54] NA61/SHINE Collaboration, N. Abgrall *et al.*, Phys. Rev. C **84**, 034604 (2011).
- [55] D. S. Barton *et al.*, Phys. Rev. D **27**, 2580 (1983).
- [56] R. P. Feynman, Phys. Rev. Lett. **23**, 1415 (1969).
- [57] H. Boggild, K. Hansen, and M. Suk, Nucl.Phys. **B27**, 1 (1971).
- [58] <http://www.fluka.org/fluka.php>.
- [59] A. Fassò *et al.*, The physics models of FLUKA: status and recent developments, CHEP03, 2003, arXiv:hep-th/0306267.
- [60] NA49 Collaboration, T. Anticic *et al.*, Eur. Phys. J. **C68**, 1 (2010), 1004.1889.
- [61] NA49 Collaboration, T. Anticic *et al.*, Eur. Phys. J. **C65**, 9 (2010), 0904.2708.
- [62] J. Park, *Neutrino-electron scattering in MINERvA*, PhD thesis, University of Rochester, 2013.
- [63] MINERvA Collaboration, L. Aliaga *et al.*, Nucl. Instr. Meth. A **743**, 130 (2014).
- [64] MINERvA Collaboration, G. Perdue *et al.*, Nucl. Instr. Meth. A **694**, 179 (2012).
- [65] http://www.hamamatsu.com/resources/pdf/etd/H8804_TPMH1333E01.pdf.
- [66] IEEE standard 1014-1987 (1987).

- [67] http://coda.jlab.org/wiki/index.php/Main_Page, 2011.
- [68] S. Ross, J. Eng. Tech. **20**, 38 (2003).
- [69] D. Rein and L. M. Sehgal, Annals of Physics **133**, 79 (1981).
- [70] A. Bodek and U. K. Yang, Journal of Physics G: Nuclear and Particle Physics **29**, 1899 (2003).
- [71] J. B. Birks, Proc. Phys. Soc. **A64**, 874 (1951).
- [72] D. A. Petyt, *A Study of parameter measurement in a long baseline neutrino oscillation experiment*, PhD thesis, University of Oxford, Oxford, United Kingdom, 1998.
- [73] A. Cabrera, A. D. Santo, P. S. Miyagawa, N. Tagg, and A. Weber, Oxford Report No. NuMI-NOTE-SCINT-934, 2003 (unpublished), MINOS Doc-DB 934.
- [74] J. Rademacker, Nucl. Instrum. Meth. A **484**, 432 (2001).
- [75] MINOS Collaboration, M. Mathis, J.Phys.Conf.Ser. **404**, 012039 (2012).
- [76] R. Fruhwirth, Nucl.Instrum.Meth. **A262**, 444 (1987).
- [77] MINERvA Collaboration, B. Eberly *et al.*, (2014), 1406.6415.
- [78] MINERvA Collaboration, T. Walton *et al.*, Phys. Rev. **D91**, 071301 (2015), 1409.4497.
- [79] MINERvA Collaboration, T. Le *et al.*, Phys. Lett. **B749**, 130 (2015), 1503.02107.
- [80] A. Hoecker *et al.*, PoS **ACAT**, 040 (2007), physics/0703039.
- [81] T. Hastie, R. Tibshirani, and J. Friedman, *The Elements of Statistical Learning*, second ed. (Springer, New York, 2009).
- [82] MINERvA Collaboration, A. Higuera *et al.*, Phys. Rev. Lett. **113**, 261802 (2014), 1409.3835.
- [83] D. Rein and L. M. Sehgal, Nucl. Phys. **B223**, 29 (1983).

- [84] R. Brun and F. Rademakers, Nucl.Instrum.Meth. **A389**, 81 (1997).
- [85] G. D'Agostini, Nucl.Instrum.Meth. **A362**, 487 (1995).
- [86] T. Adye, p. 313 (2011), 1105.1160.
- [87] Genie physics and user manual, http://genie.hepforge.org/manuals/GENIE_PhysicsAndUserManual_v2.10.00a.pdf, Retrieved 28 May 2015.
- [88] UA5 Collaboration, G. Alner *et al.*, Z.Phys. **C33**, 1 (1986).
- [89] K. S. Lackner, Nucl. Phys. **B153**, 526 (1979).
- [90] MiniBooNE Collaboration, A. A. Aguilar-Arevalo *et al.*, Phys. Lett. **B664**, 41 (2008), 0803.3423.
- [91] C. Baltay *et al.*, Phys. Rev. Lett. **57**, 2629 (1986).
- [92] CHARM Collaboration, F. Bergsma *et al.*, Phys. Lett. **B157**, 469 (1985).
- [93] H. Faissner *et al.*, Phys. Lett. **B125**, 230 (1983).
- [94] SKAT Collaboration, H. J. Grabosch *et al.*, Z. Phys. **C31**, 203 (1986).
- [95] NOMAD Collaboration, C. T. Kullenberg *et al.*, Phys. Lett. **B682**, 177 (2009), 0910.0062.
- [96] E. Isiksal, D. Rein, and J. G. Morfin, Phys. Rev. Lett. **52**, 1096 (1984).
- [97] D. Rein, Nucl. Phys. **B278**, 61 (1986).
- [98] B. Z. Kopeliovich and P. Marage, Int. J. Mod. Phys. **A8**, 1513 (1993).

A Uncertainties due to Feynman scaling

Though Feynman’s postulate is reasonably well substantiated in the literature, it is also known to be at least partially violated at the high energy extreme [88]; this implies that we need to quantify how reliable it is for our purposes here. We have chosen to bracket the uncertainty on the correction to the invariant cross-section as follows: we employ the scaling procedure to scale from the NA49 proton energy ($E_p = 158 \text{ GeV}$) down to the energy of the NA61 experiment ($E_p = 31 \text{ GeV}$)—which also measured $p + C$ pion production [54]—and measure the difference, then partially invert the scaling correction to return the error estimate to the NuMI proton energy ($E_p = 120 \text{ GeV}$). Once we have a systematic uncertainty estimate, we insert it into the neutrino flux prediction and examine its effect.

A.1. Comparing NA61 and NA49

To quantify our systematic uncertainty on the scaling procedure, we wish to perform it to scale NA49’s results for f ($E_p = 158 \text{ GeV}$) to match the energy of NA61’s results ($E_p = 31 \text{ GeV}$), and examine the difference:

$$\Delta f_{\text{NA49} \rightarrow \text{NA61}} = f_{\text{NA61}}(31 \text{ GeV}) - f_{\text{NA49}}(158 \text{ GeV}) \times \frac{f_{\text{FLUKA}}(31 \text{ GeV})}{f_{\text{FLUKA}}(158 \text{ GeV})}. \quad (\text{A.1})$$

A.1.1. Translating $\frac{1}{p}\sigma(p, \theta)$ to $f(x_F, p_T)$

Unfortunately, while NA49 reports f in two-dimensional bins of x_F, p_T (ideal for our purposes here, since this makes the results straightforward to use at any energy assuming Feynman scaling), NA61 does not report f . Instead, they give $\frac{d\sigma}{dp}$ in bins of p integrated over various ranges of θ , which we cannot apply directly (the p, θ distributions are not invariant, nor is $\frac{d\sigma}{dp}$)¹. Therefore, we are obligated to convert NA61's results into invariant cross-sections to compare them with NA49.

To do this, we begin by noting that, by definition (applying a Lorentz transformation to the center-of-momentum frame),

$$x_F = \frac{2p_L^*}{\sqrt{s}} = \frac{2\gamma_{CM}(p_{\pi,z} - \beta_{CM}E_\pi)}{\sqrt{2M_p^2 + M_p E_p}} \quad (\text{A.2})$$

$$p_T = p \sin \theta \quad (\text{A.3})$$

where $p_{\pi,z}$ refers to the z -component of the outgoing pion's momentum in the lab frame, γ_{CM} and β_{CM} are the Lorentz transformation parameters to the center-of-momentum frame, and M_p and E_p are the proton mass and incoming energy in the lab frame, respectively.

Then, also by definition, NA61's result (denoted hereafter as $g([p_{low}, p_{high}], [\theta_{low}, \theta_{high}]))$

¹This parameterization of the results *is* convenient for T2K, NA61's primary data consumer, however, who use a replica target in a basically identical beam.

can be rewritten as follows:

$$g([p_{low}, p_{high}], [\theta_{low}, \theta_{high}]) = \frac{1}{\Delta p} \sigma([p_{low}, p_{high}], [\theta_{low}, \theta_{high}]) \quad (\text{A.4})$$

$$= \frac{1}{\Delta p} \iiint \frac{d^3 \sigma}{dp^3} d^3 p \quad (\text{A.5})$$

$$= \frac{1}{\Delta p} \iiint \left(\frac{1}{E} \right) \left(E \frac{d^3 \sigma}{dp^3} \right) d^3 p \quad (\text{A.6})$$

$$= \frac{1}{\Delta p} \int_{p_{low}}^{p_{high}} \int_{\theta_{low}}^{\theta_{high}} \int_0^{2\pi} \frac{1}{E} f(x_F, p_T) p^2 \sin \theta dp d\theta d\phi \quad (\text{A.7})$$

Now we can find a bin-averaged (over the p, θ bin) f_{NA61} in terms of g . The bin average of $\frac{1}{E} f$ is:

$$\left\langle \frac{1}{E} f(x_F, p_T) \right\rangle = \frac{2\pi \int_{p_{low}}^{p_{high}} \int_{\theta_{low}}^{\theta_{high}} \frac{1}{E} f(x_F, p_T) p^2 \sin \theta dp d\theta}{2\pi \int_{p_{low}}^{p_{high}} \int_{\theta_{low}}^{\theta_{high}} p^2 \sin \theta dp d\theta} \quad (\text{A.8})$$

(where the denominator is the phase space volume element, necessary to get the average); noting that the numerator is just $g\Delta p$ (see eq. A.7),

$$\left\langle \frac{1}{E} f(x_F, p_T) \right\rangle = \frac{g\Delta p}{2\pi \int_{p_{low}}^{p_{high}} \int_{\theta_{low}}^{\theta_{high}} p^2 \sin \theta dp d\theta} \quad (\text{A.9})$$

$$= \frac{g\Delta p}{2\pi \Delta(\cos \theta) \Delta(p^3/3)} \quad (\text{A.10})$$

after doing the integrals.

Now, if we assume that the covariance of f and $1/E$ isn't extreme over the kinematic bin (since the bins are relatively small), then

$$\left\langle \frac{1}{E} f(x_F, p_T) \right\rangle \approx \left\langle \frac{1}{E} \right\rangle \langle f(x_F, p_T) \rangle, \quad (\text{A.11})$$

which means

$$\langle f_{\text{NA61}}(x_F, p_T) \rangle \approx \frac{\Delta p}{\langle \frac{1}{E} \rangle} \frac{1}{2\pi \Delta(\cos \theta) \Delta(p^3/3)} g(\Delta p, \Delta \theta)$$

with $x_F = x_F(\langle p \rangle, \langle \theta \rangle)$ and $p_T = p_T(\langle p \rangle, \langle \theta \rangle)$.

A.1.2. Phase space coverage

The region of kinematic (x_F, p_T) space sampled by NA49 and NA61 differs significantly: as illustrated in figure A.1, NA49's bin choices span a broader range of the available coordinate area, while NA61's data points partition a smaller region much more finely. We can thus only report the difference $\Delta f_{\text{NA49} \rightarrow \text{NA61}}$ (eq. A.1) in regions where the data from the two experiments overlap.

With this in mind, figure A.2 presents the fractional residuals between scaled NA49 and translated NA61 ($\frac{\Delta f_{\text{NA49} \rightarrow \text{NA61}}}{f_{\text{NA61}}}$, with the numerator from eq. A.1). A comparison plot of the expected kinematic distributions for neutrinos in two energy bins from π^+ parents in the le010z185i (forward horn current, low-energy horn and target settings) configuration mentioned previously is in figure A.3. Notice that the regions of strongest disagreement between NA49 and NA61 do not in general overlap with the region sampled by the majority of NuMI pions.

A.2. Estimating a systematic uncertainty for NuMI

To compute an uncertainty appropriate for the pions which create neutrinos in NuMI, we must remove the extra contribution in figure A.2 due to the scaling between NuMI energy and the NA61 lower bound. (Recall that for NuMI, $E_p =$

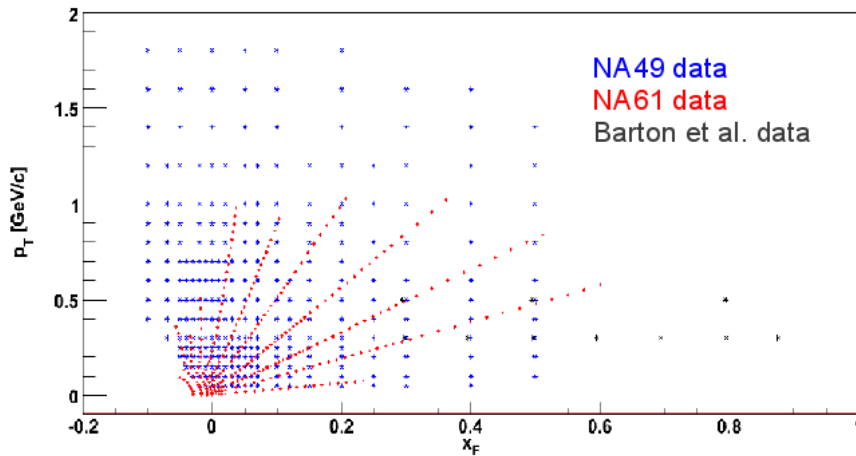
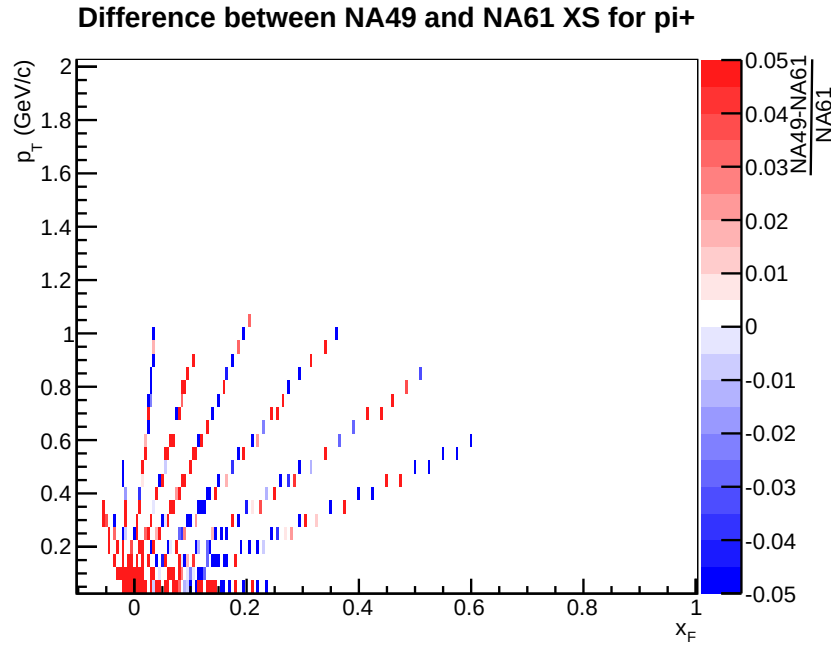


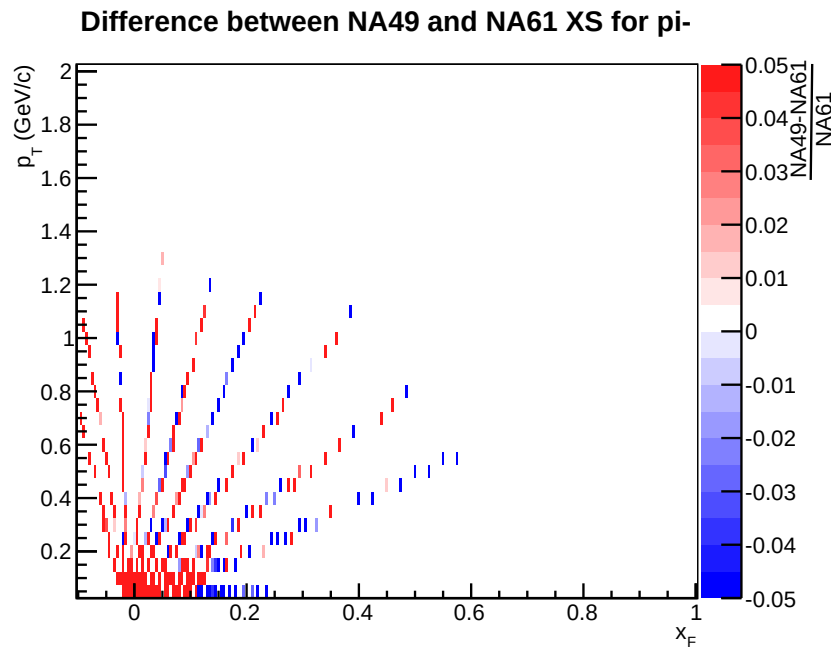
Figure A.1: Bin centers of the data points from NA49’s pion results [53], NA61 [54], and Barton et al. [55].

120 GeV, which is not nearly as far from NA49’s $E_p = 158$ GeV as NA61’s $E_p = 31$ GeV is; this means the uncertainty should be smaller than fig. A.2 would imply.)

Once we have a result that has been corrected to the NuMI energy, we must correct for two further flaws: first, the data we have been working with to this point covers only a limited subset of the range of (x_F, p_T) that pions producing neutrinos in NuMI occupy (compare figs. A.2 and A.3); second, the fluctuations between adjacent bins in fig. A.2 imply that uncertainties in the NA49 and NA61 measurements themselves (which have not been controlled for here) are contributing to the differences in a non-insignificant way.



(a) $p + C \rightarrow \pi^+ X$



(b) $p + C \rightarrow \pi^- X$

Figure A.2: Fractional residuals in the invariant cross-section ($\frac{\Delta f_{\text{NA49} \rightarrow \text{NA61}}}{f_{\text{NA61}}}$, where the numerator is defined in eq. A.1) for $p + C \rightarrow \pi^+ X$ (left) and $p + C \rightarrow \pi^- X$ (right).

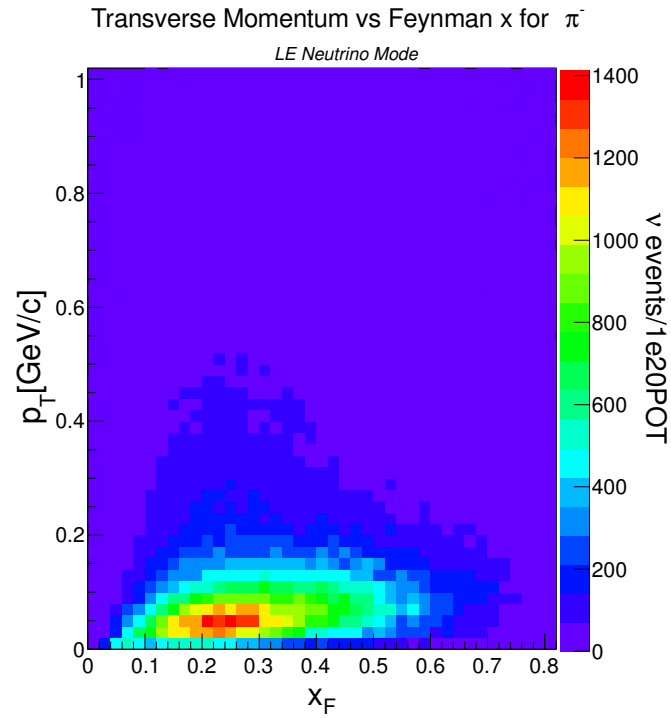
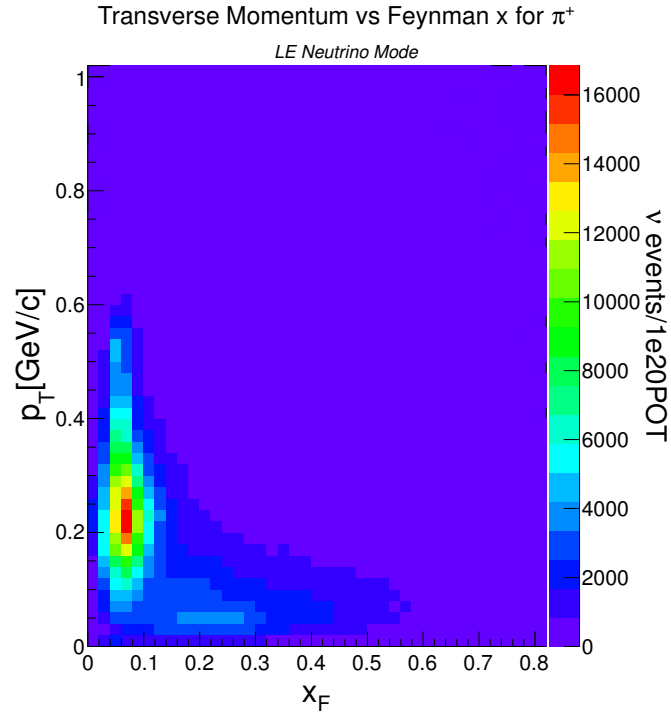


Figure A.3: Simulated x_F, p_T distributions for neutrinos from pions creating neutrinos in NuMI which strike the MINER ν A detector. Courtesy L. Aliaga (MINER ν A).

A.2.1. “Unscaling” back to NuMI energy

To “undo” the extra scaling between 120 and 31 GeV, we begin by assuming that the fractional scaling error $\Delta f_{E_i \rightarrow E_f} / f_{E_i \rightarrow E_f}$ is linear with the scaling factor $S_{E_i \rightarrow E_f}$:

$$\frac{\Delta f_{E_i \rightarrow E_f}}{f_{E_i \rightarrow E_f}} \sim k S_{E_i \rightarrow E_f} + b \quad (\text{A.12})$$

$$= k \left(\frac{f_{\text{FLUKA}, E_f}}{f_{\text{FLUKA}, E_i}} \right) + b \quad (\text{A.13})$$

$$= k \left(1 - \frac{f_{\text{FLUKA}, E_f}}{f_{\text{FLUKA}, E_i}} \right) \quad (\text{A.14})$$

(where in the last step the intercept was chosen such that for zero scaling, where $S_{E_i \rightarrow E_i} = 1$, the scaling error $\Delta f_{E_i \rightarrow E_i} = 0$).

Now, since the maps in figure A.2 correspond to $\frac{\Delta f_{158 \text{ GeV} \rightarrow 31 \text{ GeV}}}{f_{\text{NA61}}}$ for π^+ and π^- , for each case we can use it as the constraint needed to fix k , and then compute $\Delta f_{158 \text{ GeV} \rightarrow 120 \text{ GeV}}$:

$$\frac{\Delta f_{158 \text{ GeV} \rightarrow 120 \text{ GeV}}}{f_{158 \text{ GeV} \rightarrow 120 \text{ GeV}}} = \frac{\Delta f_{158 \text{ GeV} \rightarrow 31 \text{ GeV}}}{f_{158 \text{ GeV} \rightarrow 31 \text{ GeV}}} \left(\frac{1}{1 - \frac{f_{\text{FLUKA}, 31 \text{ GeV}}}{f_{\text{FLUKA}, 158 \text{ GeV}}}} \right) \left(1 - \frac{f_{\text{FLUKA}, 120 \text{ GeV}}}{f_{\text{FLUKA}, 158 \text{ GeV}}} \right) \quad (\text{A.15})$$

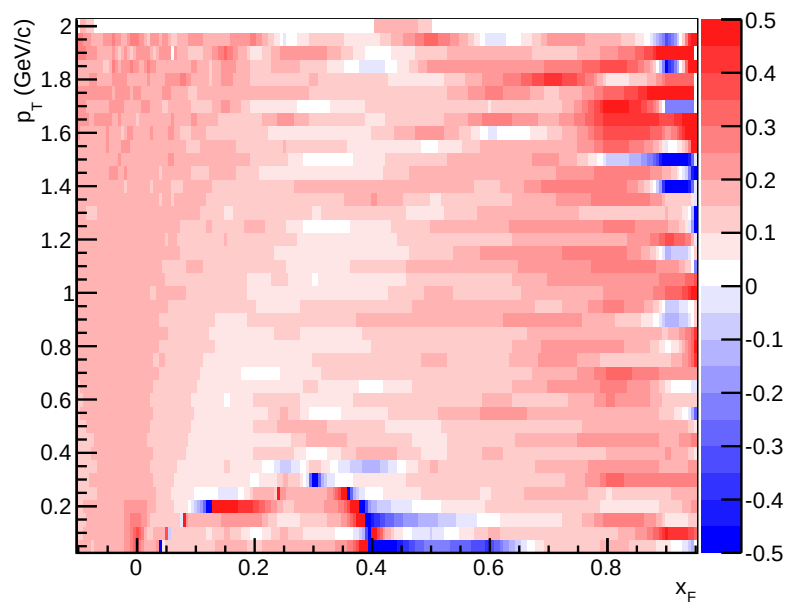
$$\frac{\Delta f_{\text{NA49} \rightarrow \text{NuMI}}}{f_{\text{NA49} \rightarrow \text{NuMI}}} = \frac{\Delta f_{\text{NA49} \rightarrow \text{NA61}}}{f_{\text{NA49} \rightarrow \text{NA61}}} \left(\frac{f_{\text{FLUKA}, 158 \text{ GeV}} - f_{\text{FLUKA}, 120 \text{ GeV}}}{f_{\text{FLUKA}, 158 \text{ GeV}} - f_{\text{FLUKA}, 31 \text{ GeV}}} \right) \quad (\text{A.16})$$

$$= \alpha(x_F, p_T) \frac{\Delta f_{\text{NA49} \rightarrow \text{NA61}}}{f_{\text{NA49} \rightarrow \text{NA61}}} \quad (\text{A.17})$$

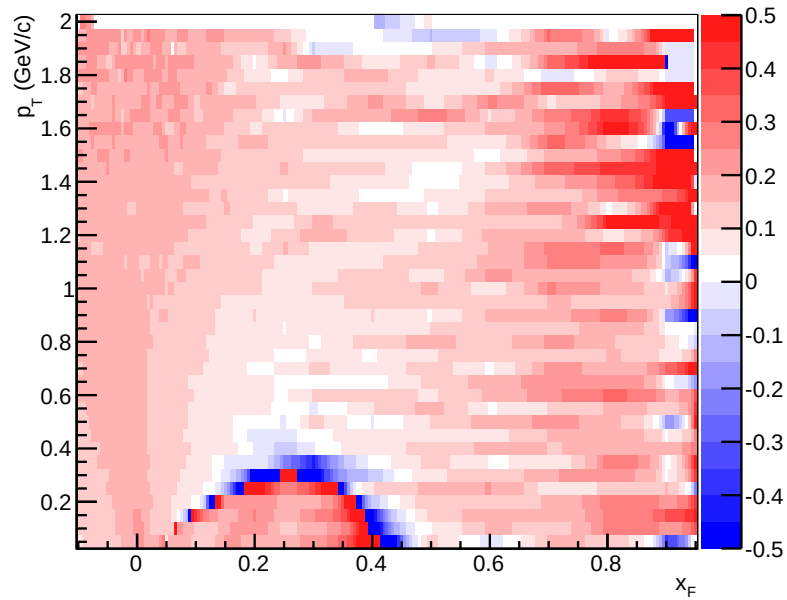
with the additional definition of

$$\alpha(x_F, p_T) \equiv \left(\frac{f_{\text{FLUKA}, 158 \text{ GeV}} - f_{\text{FLUKA}, 120 \text{ GeV}}}{f_{\text{FLUKA}, 158 \text{ GeV}} - f_{\text{FLUKA}, 31 \text{ GeV}}} \right) \quad (\text{A.18})$$

for notational convenience. The values of α as a function of (x_F, p_T) are plotted in figure A.4.



(a) $\alpha(x_F, p_T)$ for π^+



(b) $\alpha(x_F, p_T)$ for π^-

Figure A.4: The FLUKA correction factors (α in eq. A.18) for π^+ (left) and π^- (right) plotted as a function of (x_F, p_T) .

A.2.2. Smoothing

We have chosen to deal with the fluctuations in the measurements by smoothing the histograms in a systematic way: the so-called “Gaussian blur” technique, which is widely used in image processing. Gaussian blurring uses a weighted average to combine the signals from neighboring pixels in an image. In particular, if $f(x, y)$ is some function that is plotted as a function of pixels numbered (x, y) relative to a given origin, then the transformation applied by blurring is:

$$f(x, y) \rightarrow f'(x, y) = \frac{\sum_{i,j} f(i, j) G(x, y; i, j)}{\sum_{i,j} G(x, y; i, j)} \quad (\text{A.19})$$

where the Gaussian weights are defined in terms of the “Cartesian distance” between the pixel in question (x, y) and its neighbor pixels (i, j) , $r = \sqrt{(x - i)^2 + (y - j)^2}$:

$$G(x, y; i, j) = N \exp \left(-\frac{r^2}{2\sigma^2} \right). \quad (\text{A.20})$$

In this formulation, σ is a tunable parameter (which then controls the ‘width’ of the blur filter). To ensure that the computation does not become inordinately expensive for little gain, we also truncate the summations in eq. A.19 such that they are performed only over pixels (i, j) within a tunable distance r_{max} of (x, y) . Furthermore, the blurring transformations were applied iteratively to achieve a satisfactory level of smearing.

Blurring is applied both to the distributions in fig. A.2, both to account for the neighboring-bin fluctuations and as a cheap means of interpolating between the bins data in them. The parameters σ , r_{max} , and N (number of iterations) were chosen by trial-and-error, attempting to preserve as much of the structure

Variable	Fig. number	r_{max}	σ	N
$\frac{\Delta f_{NA49 \rightarrow NA61}}{f_{NA61}}$	A.2	4.0	2.0	10
α	A.4	5.0	3.0	10

Table A.1: Gaussian blur parameters (explained in the text) applied to the distributions feeding into a systematic uncertainty result.

as possible while simultaneously flatten out the large, unphysical ridges in the distributions. The values ultimately selected for each of the distributions are listed in table A.1.

The results of applying the blurring procedure with the parameters listed in table A.1 are in figures A.5 (for the cross-section residuals) and A.6 (for α).

A.2.3. The systematic uncertainty result

The distributions of the systematic uncertainty, which are the product of the distributions in figures A.5 and A.6 according to the particle species, are presented in figure A.7. It is plain that there are some residual minor peaks and troughs that were not smoothed out in this result; further tuning of the parameters in table A.1 could be applied to minimize these. However, in view of the negligible effect the systematic uncertainty has on the neutrino flux prediction (as described below), we judged that these distributions were sufficient.

A.3. Effect on the neutrino flux prediction

To evaluate the effect that the systematic uncertainty calculated above has on the neutrino flux prediction, the uncertainties from fig. A.7 are used to generate 300 universes in which the weights applied by the NA49 correction are multiplied by a random number drawn from a Gaussian with standard deviation equal to

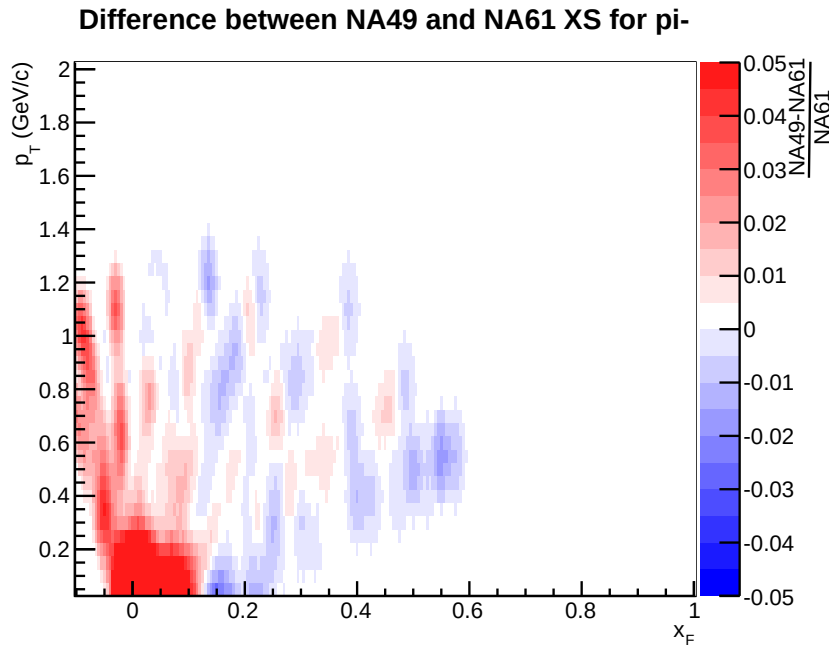
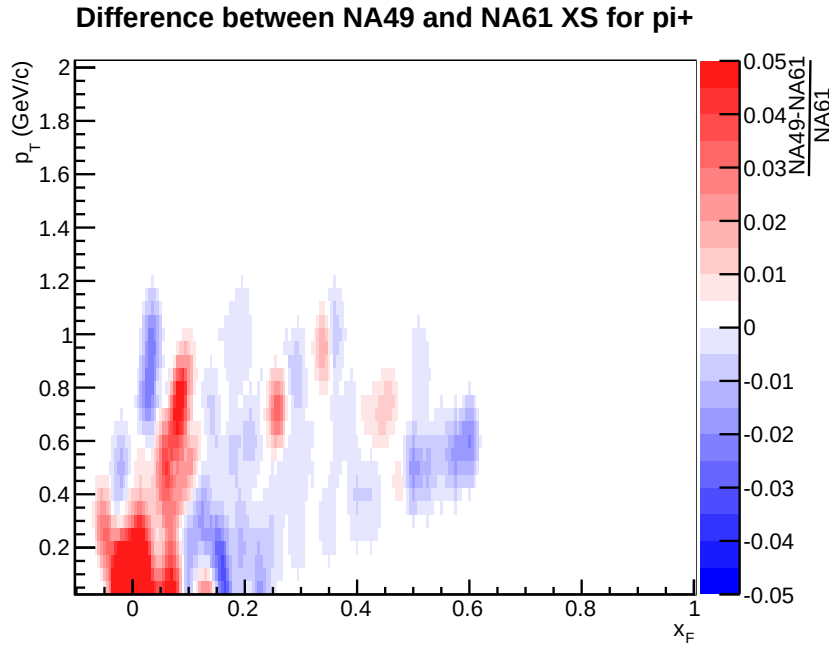
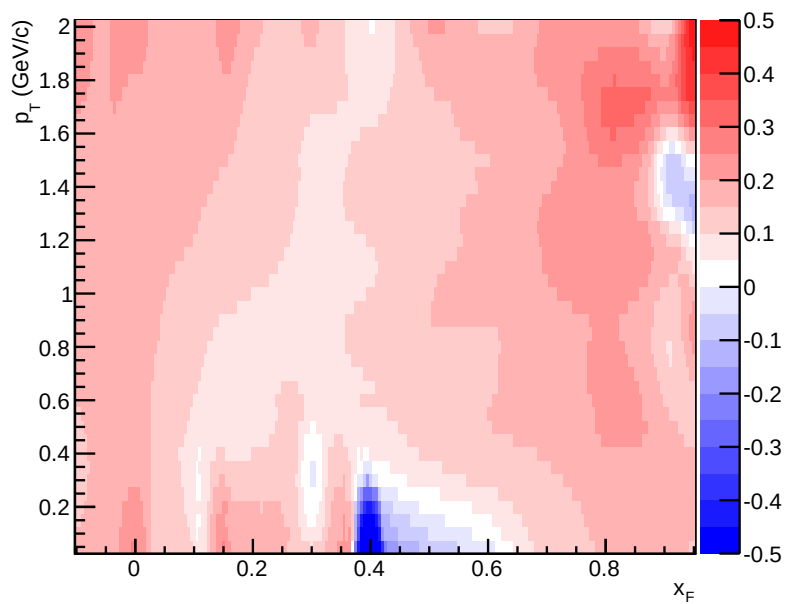
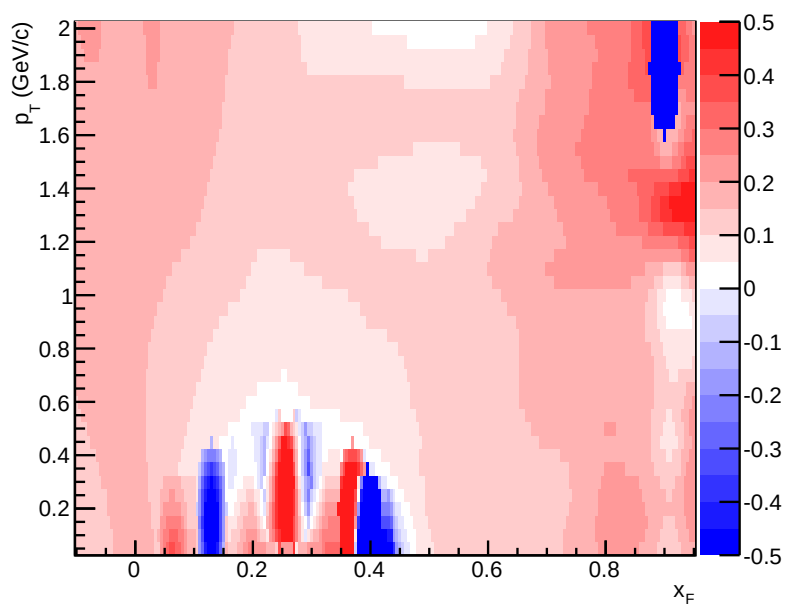


Figure A.5: Distributions of $\frac{\Delta f_{\text{NA49} \rightarrow \text{NA61}}}{f_{\text{NA61}}}$ after the blurring described in the text has been performed. Compare to fig. A.2.



(a) $\alpha(x_F, p_T)$ for $p + C \rightarrow \pi^+ X$



(b) $\alpha(x_F, p_T)$ for $p + C \rightarrow \pi^- X$

Figure A.6: Distributions of α after the blurring described in the text has been performed. Compare to fig. A.4.

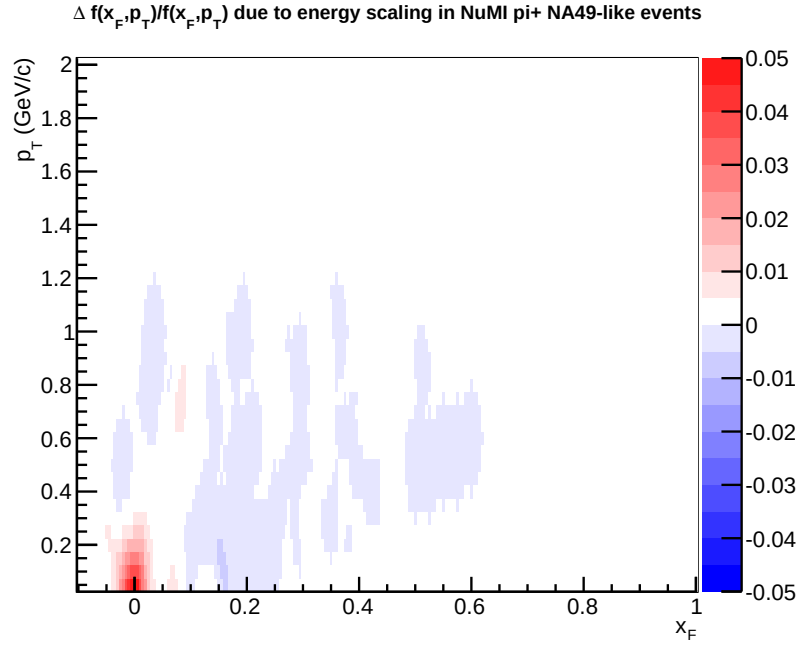
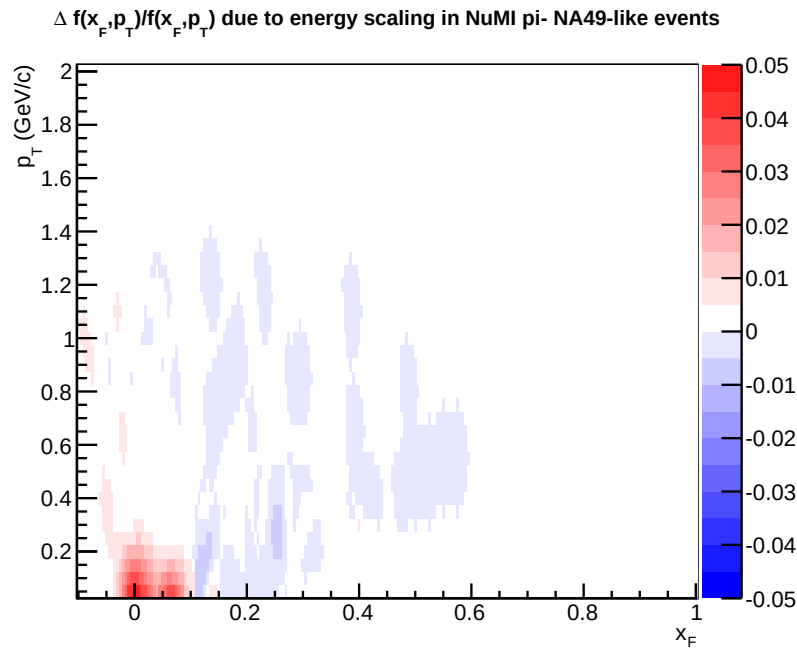
(a) $p + C \rightarrow \pi^+ X$ (b) $p + C \rightarrow \pi^- X$

Figure A.7: Distributions of the systematic uncertainty estimate, $\frac{\Delta f_{\text{NA49} \rightarrow \text{NA61}}}{f_{\text{NA61}}}$, at the NuMI proton energy.

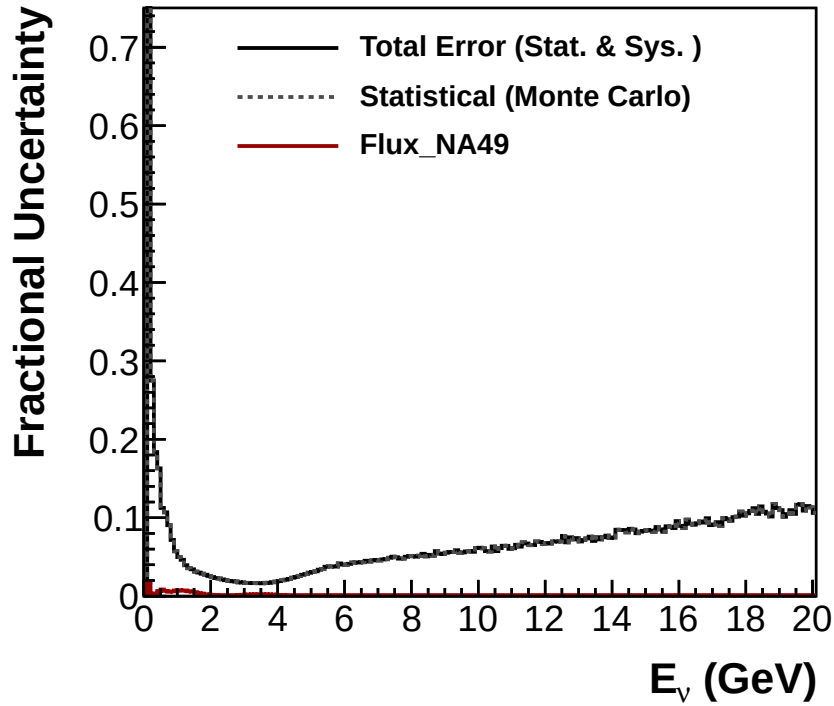


Figure A.8: Effect of the systematic uncertainties in fig. A.7 on the flux prediction for le010z185i mode beam (red). The gray dashed line indicates the statistical uncertainties for the sample used (10^{20} P.O.T. equivalent) for comparison.

the uncertainty. The systematic change to the le010z185i (low-energy, ν rich beam) flux prediction is shown in figure A.8. Since the largest contribution to the uncertainty budget is about 1%, and this at energies well outside the focusing peak region (which in le010z185i is roughly 2-5 GeV), we conclude that the effect on the neutrino flux prediction is negligible.

B Photon-like data excess

Fig. 7.5 indicates that the prediction from GENIE in the two-particle minimum front dE/dx region (that is, $2.2 < dE/dx < 3.4 \text{ MeV/cm}$) falls significantly short of what is observed in the data. (As observed in sec. 7.2.2, in this sample, that region is predominantly composed of events where the showers of one or more photons together form the electron candidate.) While the identity of the process responsible for the missing events in data is not essential to this analysis *per se*, processes that dominate in the two-particle region can be expected to also contribute at a lower level in the signal region as well, as fig. 7.5 makes clear. Therefore, an accurate estimate of the characteristics of any such events falling in the signal region is vital to the background prediction used in sec. 7.4. With that in mind, we attempt to characterize the properties of the excess in sec. B.1. We then consider several possible hypotheses for the origin of the excess in sec. B.2. Finally, in sec. ?? we summarize our findings and present the model selected for use in the background estimate of sec. 7.4.

B.1. Characterization of the excess

In the same manner as sec. 7.4.1, for the purposes of characterizing the excess, we form a “two-particle” sideband by selecting events that would otherwise pass all of the selection cuts of sec. 7.2 but have minimum front dE/dx within the two-particle range noted above: $2.2 < dE/dx < 3.4 \text{ MeV/cm}$. We may then form comparisons like those in the other sidebands (figs. 7.10 and 7.11). As illustrated in fig. B.1, while it is not especially notable vs. electron candidate angle, the

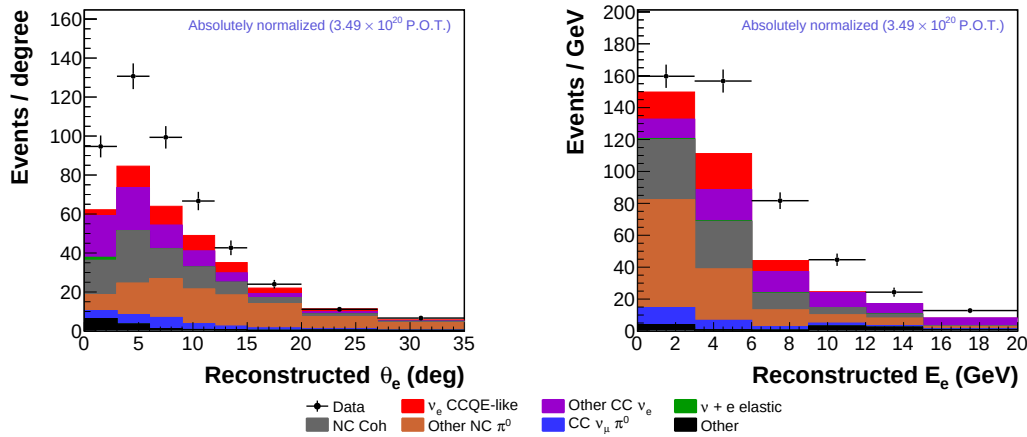


Figure B.1: Distributions in the observables from events in the two-particle sideband.

excess tends to be dominated by events of much higher electron candidate energy than the bulk of the predictions from GENIE.

Plots in the style of fig. B.1 do not directly yield information about the shape of the excess itself, which is important in the evaluation we wish to do. To better characterize the shapes, we subtracted the prediction from the data after having applied the flux constraint of sec. 3.4 and the background constraints from sec. 7.4.3. The shape of the resulting distribution may then be compared to shapes of the various background classes predicted by GENIE in any variable of interest. The shapes in the observables, for instance, are shown in figs. B.2-B.3. (These plots reinforce the inferences about the excess made above.)

In addition to the kinematics, a full characterization of the excess region will include information about the amount and location of non-electron-candidate energy in these events. Based on the comparisons in figs. B.4-B.5, it is apparent that typical excess events have little non-shower activity, either near the vertex of the shower (fig. B.4) or upstream of and in-line with the shower (fig. B.5). Energy elsewhere in these events is also consistent with the shower being the only energy (fig. B.6). All of these together imply a neutral-current process, though its particular identity remains unknown (see the further discussion of the implications

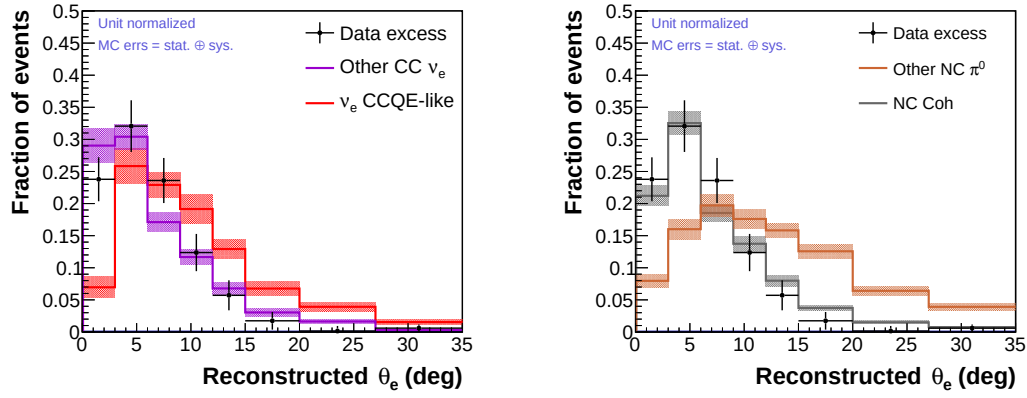


Figure B.2: Shape of the excess compared to MC categories in the electron candidate angle from events in the two-particle sideband.

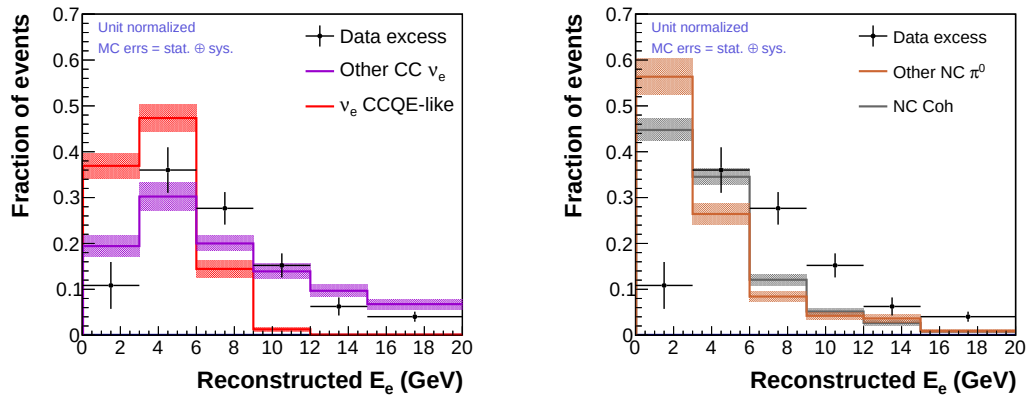
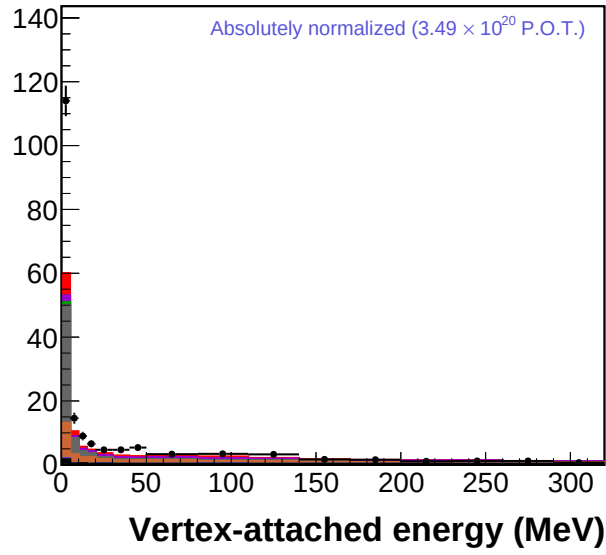
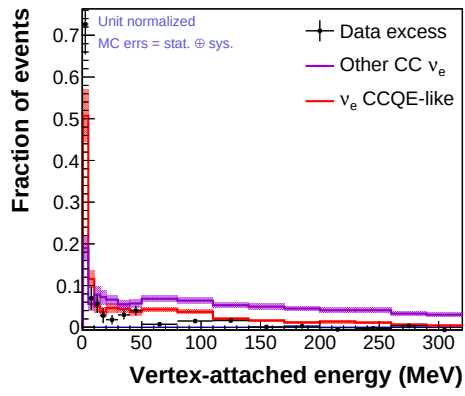
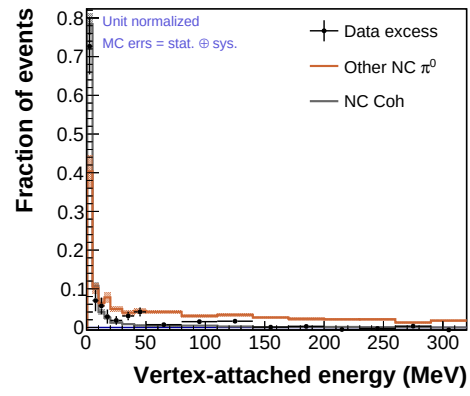


Figure B.3: Shape of the excess compared to MC categories in the electron candidate energy from events in the two-particle sideband.



(a) Data vs. total sim.

(b) ν_e categories(c) NC π^0 categories**Figure B.4:** Vertex-attached energy in the two-particle sideband.

of these distributions in sec. B.2).

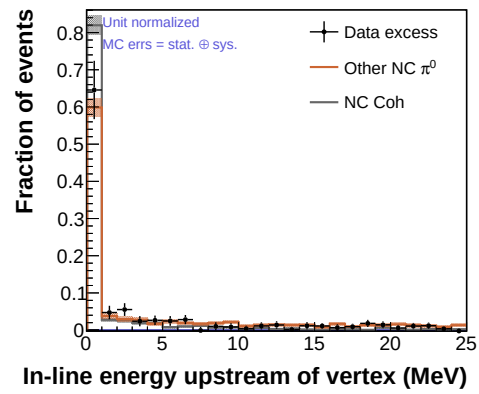
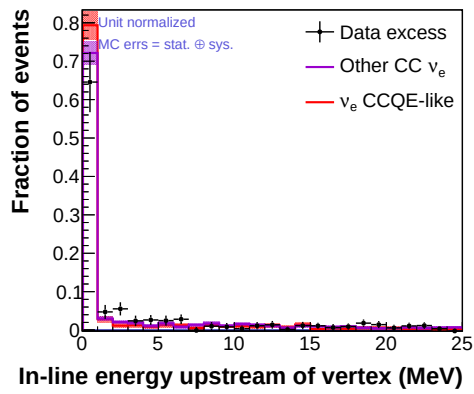
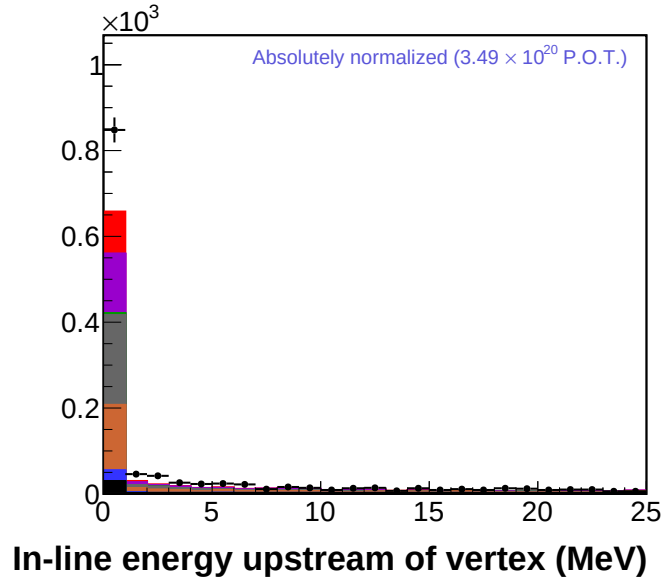
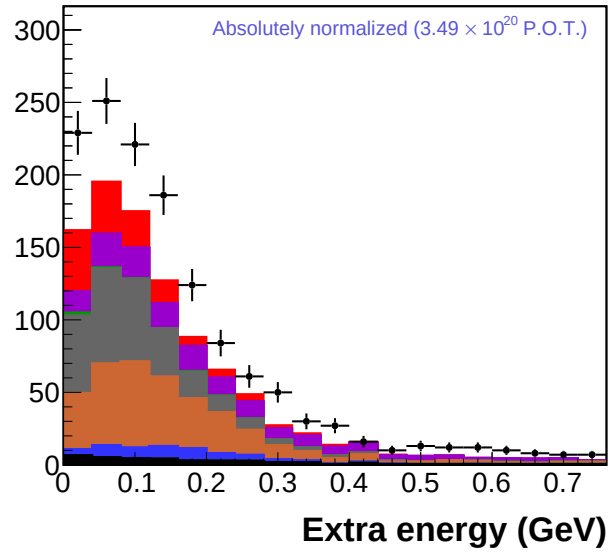


Figure B.5: Energy inside a 7.5° cone along the shower axis but pointing upstream, in the two-particle sideband.



(a) Data vs. total sim.

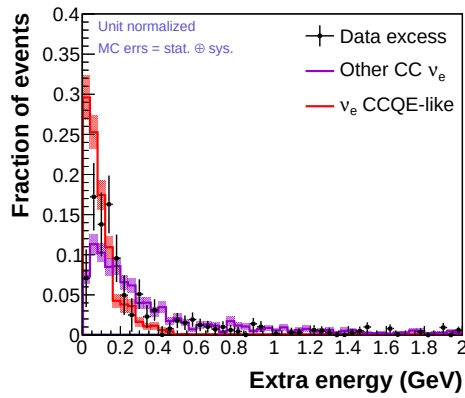
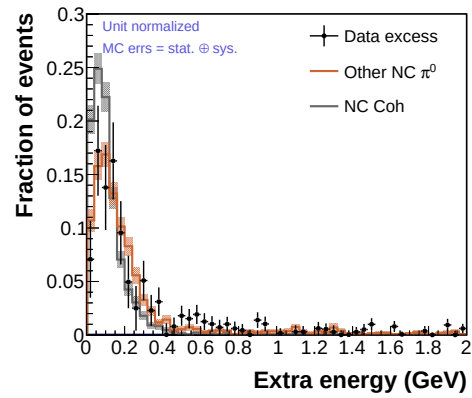
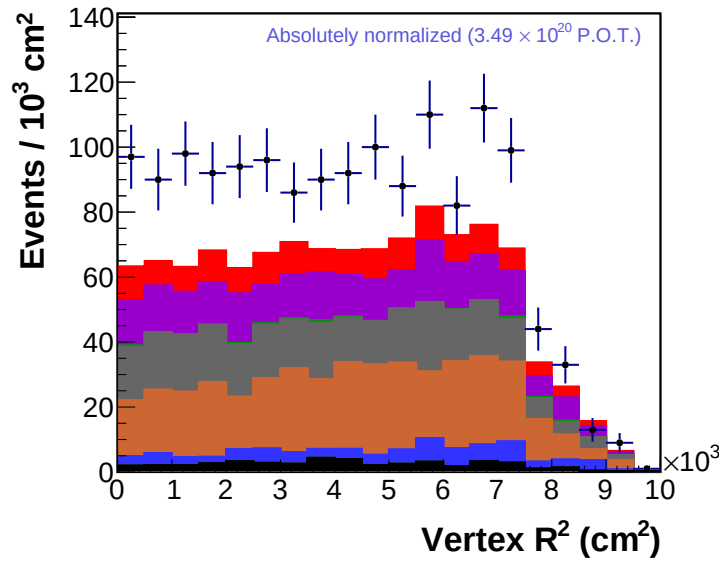
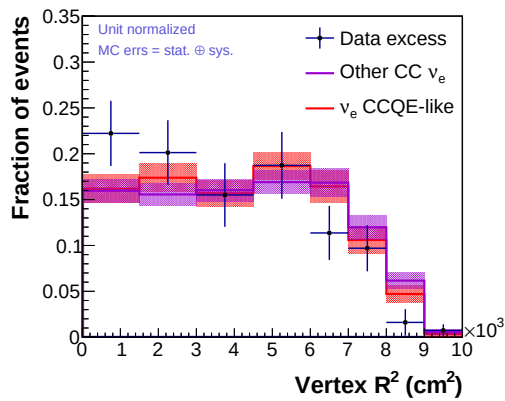
(b) ν_e categories(c) NC π^0 categories

Figure B.6: Energy outside the electron cone and vertex region, in the two-particle sideband.

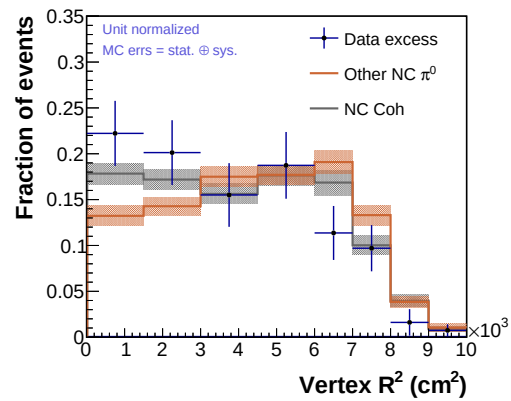
Finally, in order to ensure that some pathology in the reconstruction or simulation is not responsible for the excess, we also examine where the events are reconstructed as originating. Figs. B.7-B.9 demonstrate that the reconstructed vertex position and shower direction are consistent with real events originating inside the detector fiducial volume. (This point will be further elaborated in sec. B.2.)



(a) Data vs. total sim.

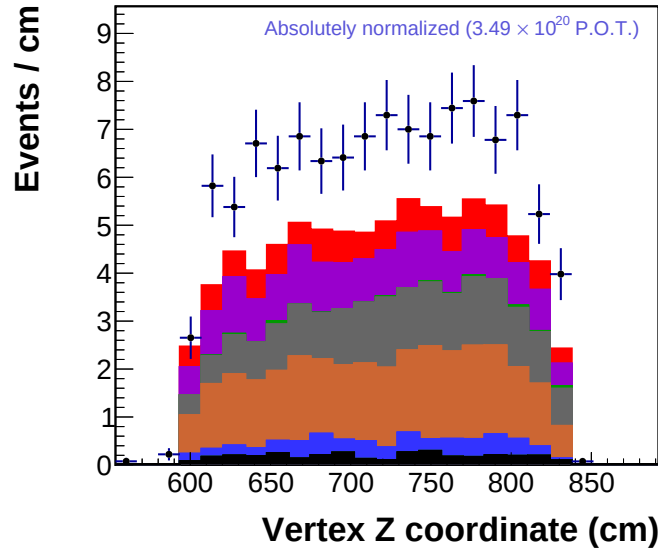


(b) ν_e categories

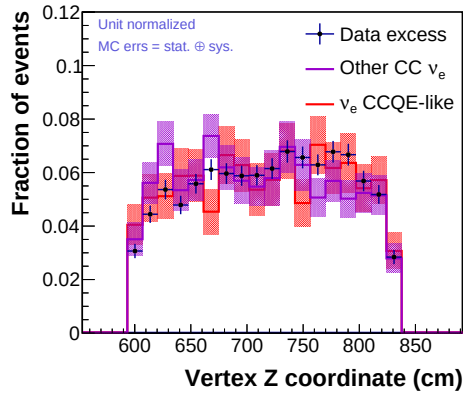
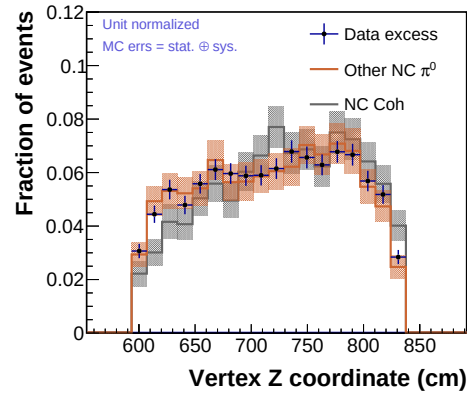


(c) NC π^0 categories

Figure B.7: Vertex R^2 ($= x^2 + y^2$) of events in the two-particle sideband.



(a) Data vs. total sim.

(b) ν_e categories(c) NC π^0 categories**Figure B.8:** Vertex z position of events in the two-particle sideband.

B.2. Candidate explanations for the excess

We have considered several hypotheses for the origin of the excess. They are each treated in turn below.

B.2.1. Non-fiducial/externally-entering particles

We first sought to understand whether the excess could be caused by neutrino interactions outside of our chosen fiducial volume. Although we expect that the data overlay procedure described in sec. 5.2.3 should result in externally-entering events in the simulation, the plausibility requirement of sec. 6.2.3 will remove any selected event where the electron candidate originates from data. Since we simulate only neutrino interactions within the inner detector and rely on our selection cuts to remove other interactions from the data sample, therefore, outside-detector neutrino interactions constitute (*prima facie*) an important candidate for the excess.

However, when we examine the distributions of the reconstructed vertex positions in the excess region, the evidence disfavors this hypothesis. The z position of the reconstructed vertices, shown in fig. B.8, is consistent with the simulated processes that arise inside the fiducial volume; anything arising upstream would develop earlier in z , and we see no such behavior. Furthermore, events entering the side of the detector from the outside should preferentially interact towards the outer edges of the detector, resulting in the excess being concentrated towards larger reconstructed $R^2 = x^2 + y^2$; but we observe in fig. B.7 that the excess is distributed uniformly in R^2 .

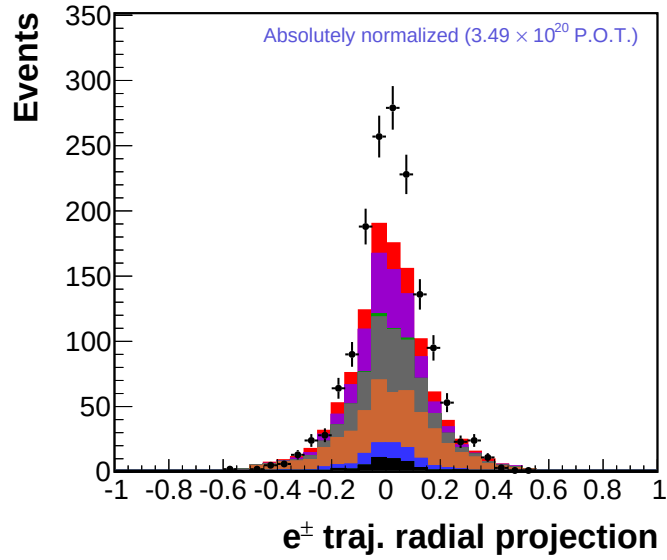
To further establish that the excess process cannot be explained by side-entering events, we examined one other characteristic of the direction of events in the excess region. If we were to superimpose a cylindrical coordinate system on the detector,

where the axis of the cylinder coincides with the z -axis of the detector, then the radial coordinate $\hat{\rho}$ points directly outward in the x - y plane. The projection of the electron candidate trajectory r onto $\hat{\rho}$, that is, the quantity

$$\pi_{\hat{\rho}} = \frac{\vec{r} \cdot \hat{\rho}}{||\vec{r}||} \quad (\text{B.1})$$

should differentiate between events entering from the exterior (which will always point inward, and thus have negative projection onto $\hat{\rho}$) and events originating inside the detector (which should be distributed on both sides of 0 in $\pi_{\hat{\rho}}$). As shown in fig. B.9, the excess is consistent with events inside the detector.

Based on this evidence, we rule out the hypothesis that the excess arises from externally-entering events.



(a) Data vs. total sim.

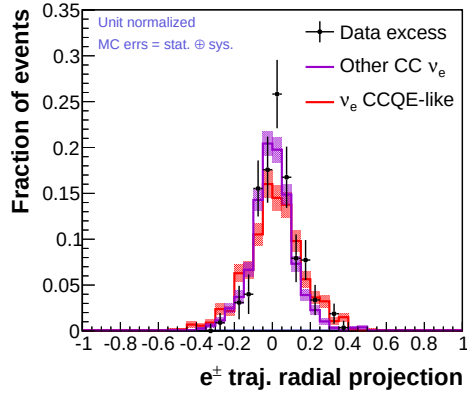
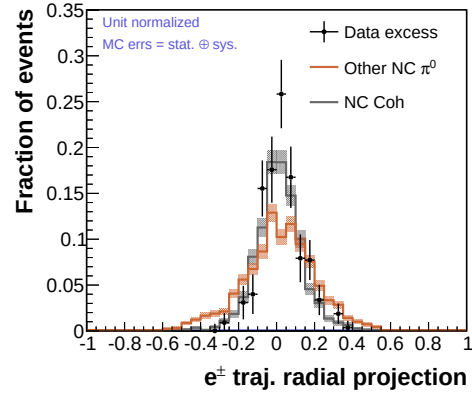
(b) ν_e categories(c) NC π^0 categories

Figure B.9: Projection of shower axis onto radial coordinate for events in the two-particle sideband. (Projection of +1 means pointing directly outward from z -axis; -1 is directly inward.)

B.2.2. Charged-current interaction model

A second potential explanation for the origin of the excess posits that the detector simulation significantly mismodels the front dE/dx distribution of electron processes, and that the excess events are really electron events in which the shower begins to develop earlier than GEANT predicts (c.f. the cartoon in fig. B.10). (The same effect would apply to positrons with no modification.)

In this scenario, both the signal process and the “other CC ν_e ” background should be equally affected because both yield electrons (or positrons) as the primary showering particle in the final state, and this hypothesis supposes that the effect is due entirely to a process affecting only the final state. Therefore, we must consider them together when considering whether this conjecture is consistent with the excess. However, we observe (fig. B.4) that the excess prefers very little shower vertex-attached energy, in contrast to the ν_e events (which have nucleons and frequently other charged particles originating from the shower vertex). Moreover, there is no corresponding excess observed in the Michel-match sideband (fig. B.11), which is a sample strongly dominated by ν_e events. We will also notice in passing that it is difficult to understand how the photon backgrounds would not also be

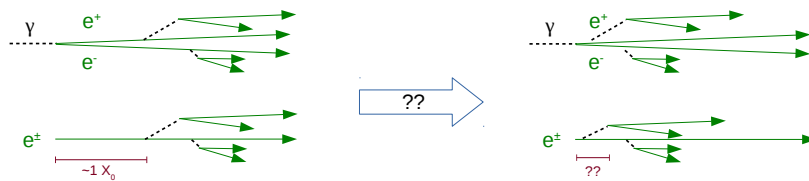


Figure B.10: Schematic illustration of how mismodeling of electron showers could result in the excess process. At left: the model in the simulation, where the first Bremsstrahlung photon (initiating the cascade) is radiated (on average) after about 1 radiation length ($1 X_0$) of travel, permitting measurement of the initial dE/dx over that range. Right: postulated true behavior, in which the initial radiation is emitted sooner than expected; this results in a dE/dx corresponding to more than one particle for some of the measurement range, and could therefore (on average) lie in the two-particle region.

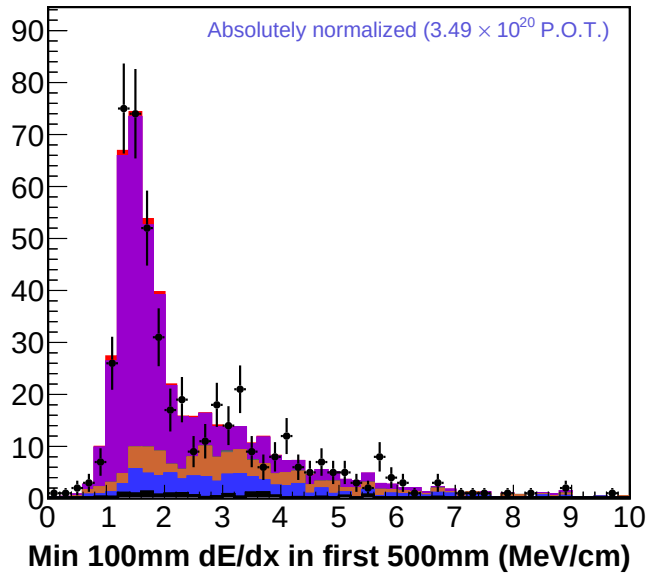


Figure B.11: Data-simulation comparison in Michel-match sideband.

affected by the kind of effect considered here, and there does not appear to be any of the broadening of the data “two-particle” peak in fig. 7.5 relative to that of the simulation that we would expect if a photon-initiated cascade developed more quickly than the simulation predicts.

Therefore, we conclude that the excess does not arise from a catastrophic mismodeling of the electron shower or electron final-state processes.

Additionally, the argument regarding vertex-attached energy from nucleons made above applies also to charged-current neutral pion production, and the latter process is therefore ruled out equally well.

An alternative way in which the charged-current model could be responsible for the excess is via the non-shower activity associated with an event. In particular, MINER ν A has found that in muon neutrino scattering, GENIE’s model systematically underestimates the energy observed near the muon vertex [21, 20], and, by lepton universality, it is expected that a similar phenomenon should happen for electron neutrino CCQE. When this activity overlaps with the beginning of

the electron shower, as it is bound to do at least some fraction of the time, it could bias the dE/dx measurement. The sliding nature of the way the front dE/dx is measured (see sec. 7.2.2) is specifically intended to minimize the impact that unsimulated nucleons could have on the front dE/dx , but nevertheless, we sought to quantify its potential contribution. To do so, we added an extra proton of kinetic energy sampled uniformly from 0-225 MeV to 25% of our simulated events, consistent with the findings in muon neutrino scattering [21]. Because those results provided no insight into the angular distribution of the extra proton, we performed the study twice: once with the protons distributed uniformly in angle with respect to the electron, and one in which the proton was constrained to lie in the same hemisphere as the electron. We compared the shape of the nominal dE/dx distribution (with all other selection cuts applied) to that of the samples with extra protons, shown in fig. B.12. As can be readily seen there, the effect is minimal, and cannot be responsible for an excess of the magnitude we observe.

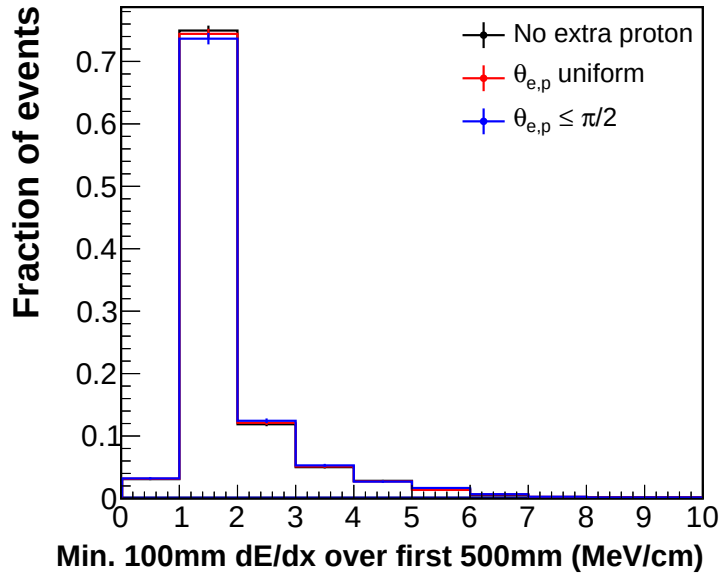


Figure B.12: Shape of front dE/dx distribution with: nominal simulated sample (black); extra protons distributed uniformly in angle w.r.t. the electron (red); extra protons constrained to lie in the same hemisphere as the electron (blue).

B.2.3. Neutral-current π^0 model

Due to the extremely different spectra — particularly in shower energy (fig. B.3), but also in angle (fig. B.2) — we already expect that the excess is very unlikely to be due to a mismodeling of the NC π^0 models GENIE offers. To confirm or disprove this hypothesis, however, we wish to examine other properties of these classes of events which depend less on the neutrino interaction model than the kinematics do, as we did in the sections above. As noted in sec. B.1 above, the consistency of the vertex-attached energy (fig. B.4) and in-line upstream energy (fig. B.5) criteria with both NC event classes indicate that the excess process is very likely neutral-current.

However, we expect that the amount of non-shower, non-vertex energy in the event should differentiate between NC coherent events (in which no other particle is produced in the final state besides the π^0) and NC incoherent events (which,

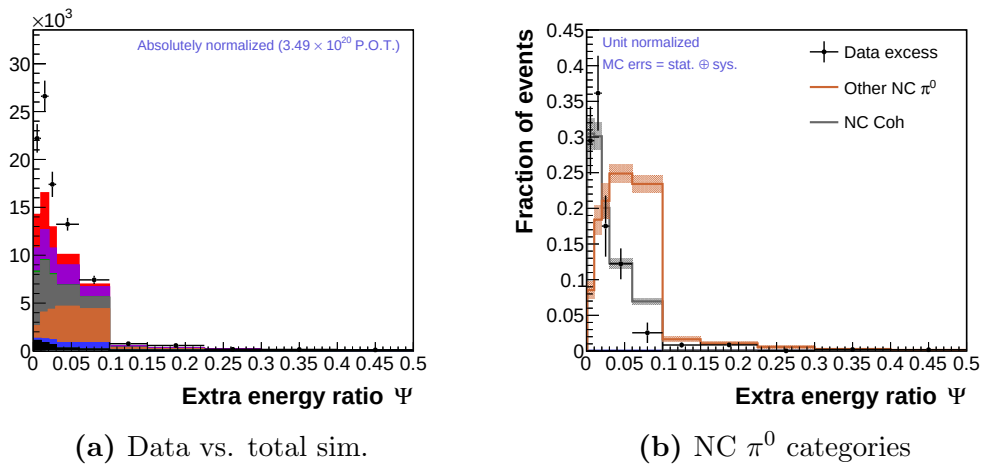


Figure B.13: Extra energy fraction (as defined in sec. 6.2.3), in the two-particle sideband.

in general, will have other particle activity). Fig. B.6, at first glance, seems to indicate that the excess is more like the incoherent process than the coherent one. However, as was noted in sec. 6.2.3, the energy leaked outside of the reconstructed cone by EM showers increases with energy. Since the excess occupies a much higher range of shower energy, on average, than the NC processes simulated by GENIE, the extra energy distributions are not directly comparable. On the other hand, the Ψ variable of sec. 6.2.3 was designed to combat precisely this problem. Examining fig. B.13, we in fact come to the opposite conclusion: the excess process appears to be coherent-like.

To further evaluate the consistency of the excess with a coherent process, we investigated one other quantity. Based on the work of Lackner [89] and Rein and Sehgal [83] on neutral-current coherent π^0 production, experiments attempting to observe NC coherent production traditionally work in the variable $E_{\pi^0}(1 - \cos \theta_{\pi^0}) \approx E_{\pi^0}\theta_{\pi^0}^2$ (refs. [90, 91, 92, 93, 94, 95]; others in narrow neutrino energy ranges sometimes use just $\theta_{\pi^0}^2$, e.g., ref. [96]). Because it is constrained by the requirement that the wavelength corresponding to the momentum transfer t be greater than a constant r_0 (roughly the inverse of the nuclear size), this variable

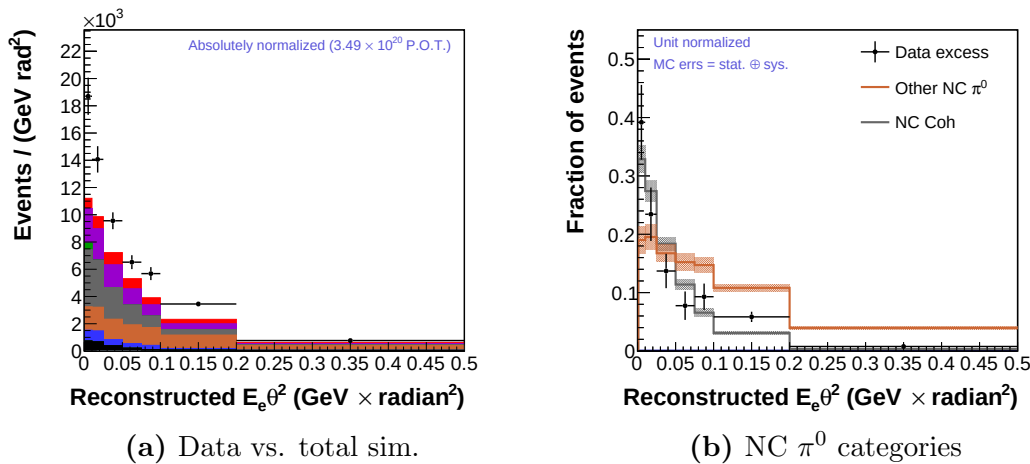


Figure B.14: $E\theta^2_{\text{shower}}$, in the two-particle sideband.

obeys the relation $E(1 - \cos \theta) \lesssim 1/r_0$ [89]. Thus in this variable, the coherent process tends towards much smaller values than the incoherent (resonant, DIS) ones, where it does not obey the same condition. It also has the virtue of being an approximately energy-independent measure of the coherence of the reaction, a vital characteristic in the current analysis (where the excess occupies a much different energy regime than the predicted coherent process and is thus difficult to compare directly). Accordingly, we compared the shape of the excess in $E\theta^2_{\text{shower}}$ to the processes simulated by GENIE. As shown in fig. B.14, the process is again most consistent with a coherent process.

All this evidence appears to indicate a coherent neutral-current process, producing a final state with at least one (and possibly more) electromagnetic shower(s), which has a much harder shower energy spectrum than the neutral-current π^0 production predicted by the Rein-Sehgal model as implemented in GENIE. However, we feel that it is very unlikely our excess arises from that process since, as argued in sec. 7.4, this same model predicts charged-current coherent production of pions as well, and the excellent MINER ν A measurement of this process [82] agrees very well with the model at pion energies above 0.5 GeV. Furthermore, to approach the magnitude of disagreement evidenced in fig. 7.5, the normalization of the NC

coherent π^0 process would need to be incorrect at the level of a factor of two, and such a disagreement would be in strong tension with the charged-current results. Thus, we conclude that none of the models incorporated into the simulation are good candidates for the origin of the excess.

B.2.4. Diffractive NC scattering from H

Since MINER ν A's tracker material is a CH-based hydrocarbon, it contains approximately-free protons in abundance. Therefore, we expect that neutral pions originating from a diffractive process occurring on the hydrogen (analogous to the coherent process on carbon and the other heavy nuclei in the detector) should be present in our data sample. However, given that in diffractive scattering, the proton absorbing the momentum transferred from the neutrino is much less massive than an entire nucleus, it often recoils substantially. This has an important consequence: the proton will often receive enough kinetic energy to make it visible in MINER ν A, resulting in visible energy away from the shower vertex.

Some quantitative evaluation of diffractive production as a candidate for the excess is clearly desirable. GENIE in fact contains an implementation of a model for charged-current diffractive scattering from D. Rein [97], but it is disabled by default since it has not been subjected to the detailed benchmarking against data that GENIE's default models are required to undergo. We nevertheless attempted to use it to estimate the characteristics of diffractive production by modifying it such that in its output events, charged pions were replaced with neutral pions and the outgoing charged lepton with a neutrino. We also multiplied the cross-section by 0.5 in accordance with the expectation for the CC/NC ratio from coherent production [98]. (These changes result in slightly incorrect kinematics for values of true Q^2 near and below the sum of the muon and pion masses, which is the threshold for the charged-current process; but since we observe showers with an

order of magnitude more energy than that, should this be the correct explanation for the excess, we expect that the threshold effects will be unimportant.) We then generated events using this model in the predicted NuMI flux (see ch. 3) with the MINER ν A detector model (see ch. 4).

The spectra in front dE/dx , electron candidate angle, and electron candidate energy, where the diffractive prediction is added to the constrained nominal prediction, are compared to the data spectra in figs. B.15. Though it fills in the excess region in dE/dx nicely, the combined prediction with the diffractive piece included significantly overpredicts the rate at high angles and low energies.

Regardless of the particulars of the kinematics, as noted above, there is one feature that will be present regardless of the implementation, and that is the recoiling proton. We examined the energy contained in a cone beginning at the reconstructed shower vertex and pointing directly backwards along the shower axis. The comparison of the shape in this variable between the diffractive events and the data excess is presented in fig. B.16. Though qualitatively the behavior is correct—both the model and the excess have a large number of events where the visible energy upstream is consistent with zero, with a long tail up to tens of MeV—the diffractive model consistently produces more of it.

Thus, while we cannot reject the diffractive hypothesis out of hand (given the qualitative similarities to the excess, including in the “coherence” variable $E\theta^2$ in fig. B.17), we are forced to conclude that the model we have available cannot simulate the excess well enough to use it in a subtraction in the signal region.

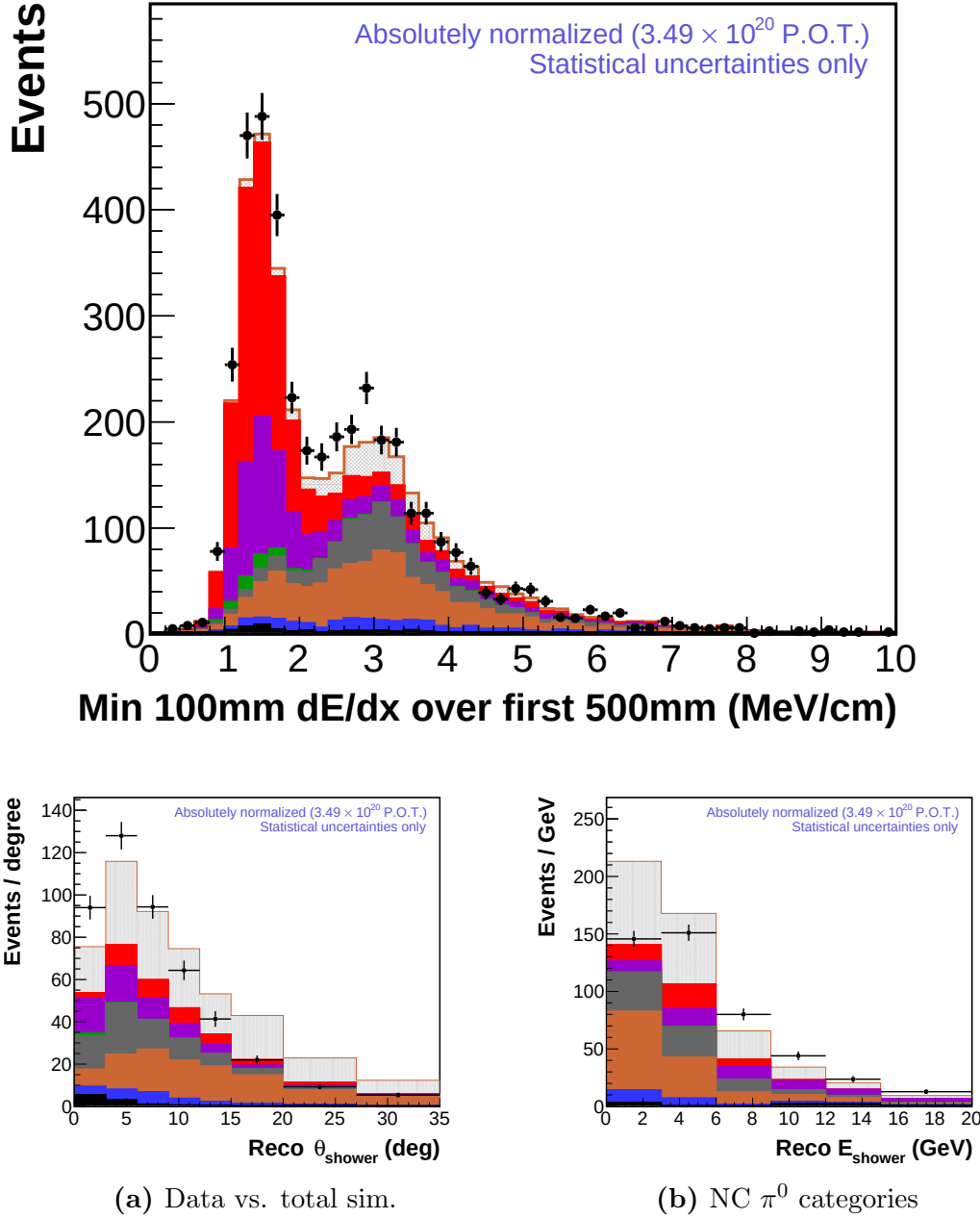


Figure B.15: Comparisons of the properties of the diffractive model against the excess. Top: front dE/dx distribution with all cuts except the dE/dx cut applied (diffractive contribution shown with solid orange line and gray shading). Bottom: kinematic variables with the diffractive model added to the constrained nominal prediction, compared to the data.

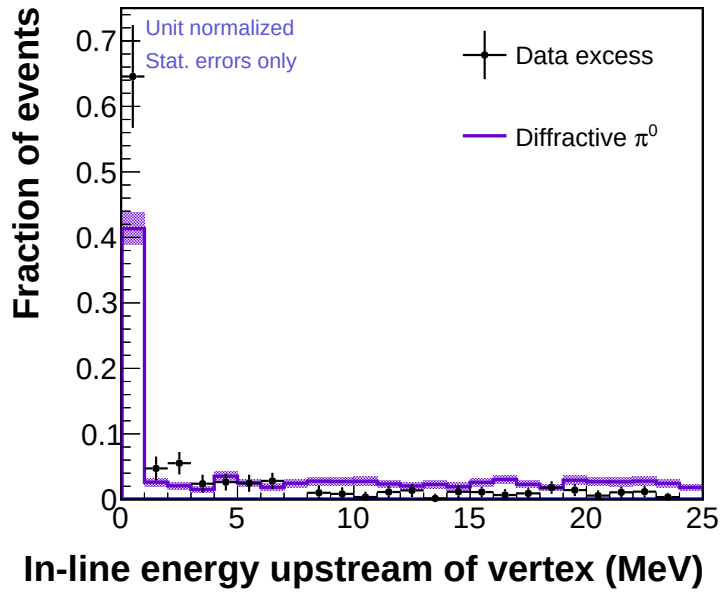


Figure B.16: Shape comparison in upstream inline energy (defined in the text), diffractive model vs. data excess.

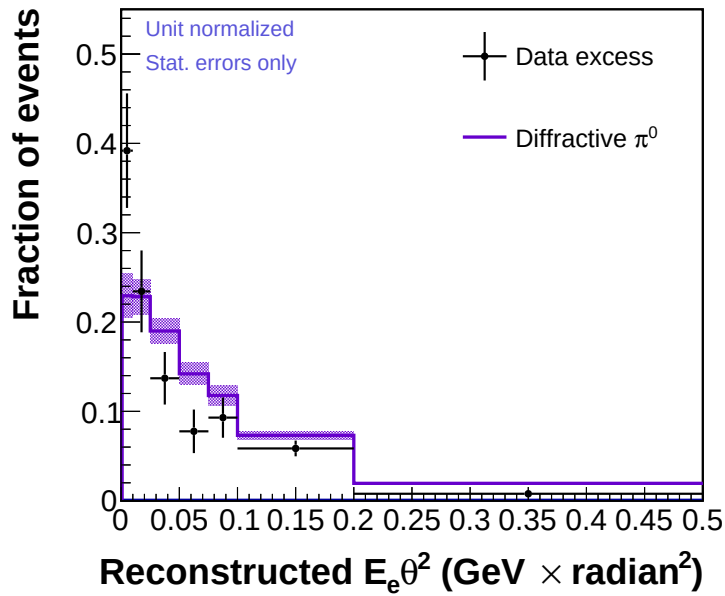


Figure B.17: Shape comparison in $E\theta^2$, diffractive model vs. data excess.

B.3. *Ad hoc* models for extrapolating into the signal region

B.3.1. Construction of samples

As detailed above, we were unable to find any candidate process among those included in our simulation which appears to correctly predict the spectra of the excess. Therefore, in order to estimate the effect that the excess process may have in the signal (*i.e.*, 1-particle dE/dx) region, we are forced to turn to an *ad hoc* model engineered explicitly to match the measured final-state kinematics and content. Because the evidence as presented thus far does not conclusively differentiate between single- and multiple-photon showers, we constructed several alternatives: single photons, neutral pions (which decay to two photons 99% of the time with a very short lifetime), and etas (which decay into two-photon or three- π^0 states about 70% of the time, also very quickly).¹ As is explained in the summary below, we chose from among them to form the predicted contribution of the excess process in the signal region in sec. 7.4.

All of our *ad hoc* models were constructed in the same fashion. We began with simulations of 500,000 single particles of each type in our detector, where the initial spectra were very broad: kinetic energies were sampled uniformly from 0 to 20 GeV, angles with respect to the beam axis were sampled uniformly from 0 to $\pi/2$, and the azimuthal angle ϕ was unconstrained (uniformly sampled from

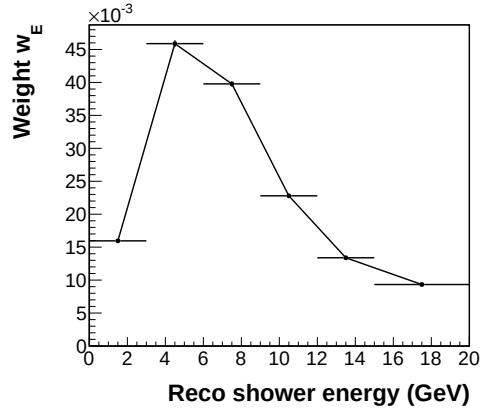
¹These choices were inspired in part by the particles expected to be produced in coherent neutrino interactions and in part by the requirement that the particle should quickly decay to a final state involving only photons. We were unable to find any higher-mass predicted coherent processes that match our observations: ρ^0 and $\rho^\pm$ decays nearly always involve a charged pion (ruled out by the extra energy this would produce), and three-pion coherent production mediated by the a_1 [98] occurs only via the charged current, which would result in a charged lepton in the final state (ruled out by the significant non-vertex energy this would entail, regardless of which particle was chosen as the shower). Though coherent η production is predicted not to exist [92], we included the η as an example of a heavier mass state decaying to photons.

0 to 2π). These samples were fed into the reconstruction and analysis chain of chs. 6 and 7. We then formed sets of weights for the events reconstructed into the two-particle sideband such that the weighted reconstructed shower energy and angle distributions for the simulation matched the distributions observed in the data. Because the correlations between the angle and energy are important (presumably only certain combinations are allowed based on the allowed kinematics of the reaction, like in the interactions we have models for), but the statistics of the excess in the data are not sufficient for a simple bin-by-bin reweighting in the two-dimensional energy-angle space, we constructed a bin-by-bin weight function in the energy and a set of one bin-by-bin function in angle for each energy bin:

$$w(E_0, \theta_0) = C_i w_E(E_0) w_\theta^i(\theta) \quad (\text{B.2})$$

(where the index i corresponds to the bin of the w_E function that energy E_0 falls into, and the C_i are normalization constants). The functions interpolate linearly between the ratios calculated in the bin centers. Example weight functions (drawn from the photon sample) are shown in fig. B.18. To verify that kinematics of the resulting reweighted samples match the excess, we compare them in fig. B.19. The front dE/dx distributions with the various *ad hoc* models (after fitting to the excess) included are presented in fig. B.20.

With these samples in hand, we can make comparisons of a similar form to those considered in sec. B.1. We will exploit this capability repeatedly in the next section.



(a) Energy weight function

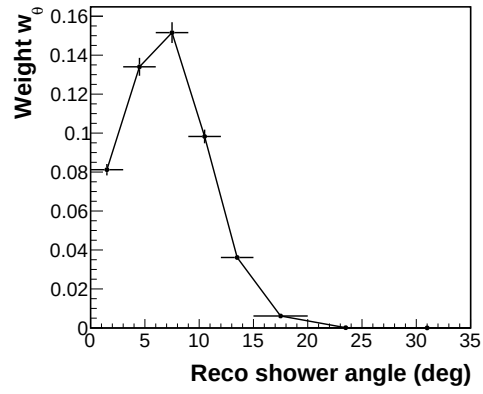
(b) Angle weight function for $3 \text{ GeV} < E_{\text{shower}} < 6 \text{ GeV}$.

Figure B.18: Sample weight functions from among those used to achieve agreement between the photon *ad hoc* single-particle sample and the data excess.

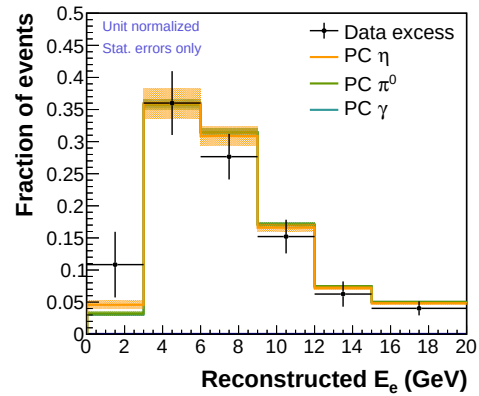
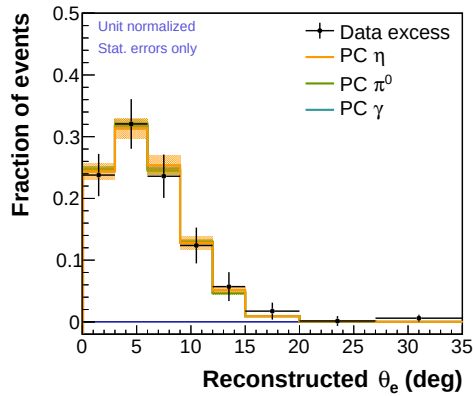


Figure B.19: Comparisons of the shapes of the reweighted *ad hoc* samples with the measured excess.

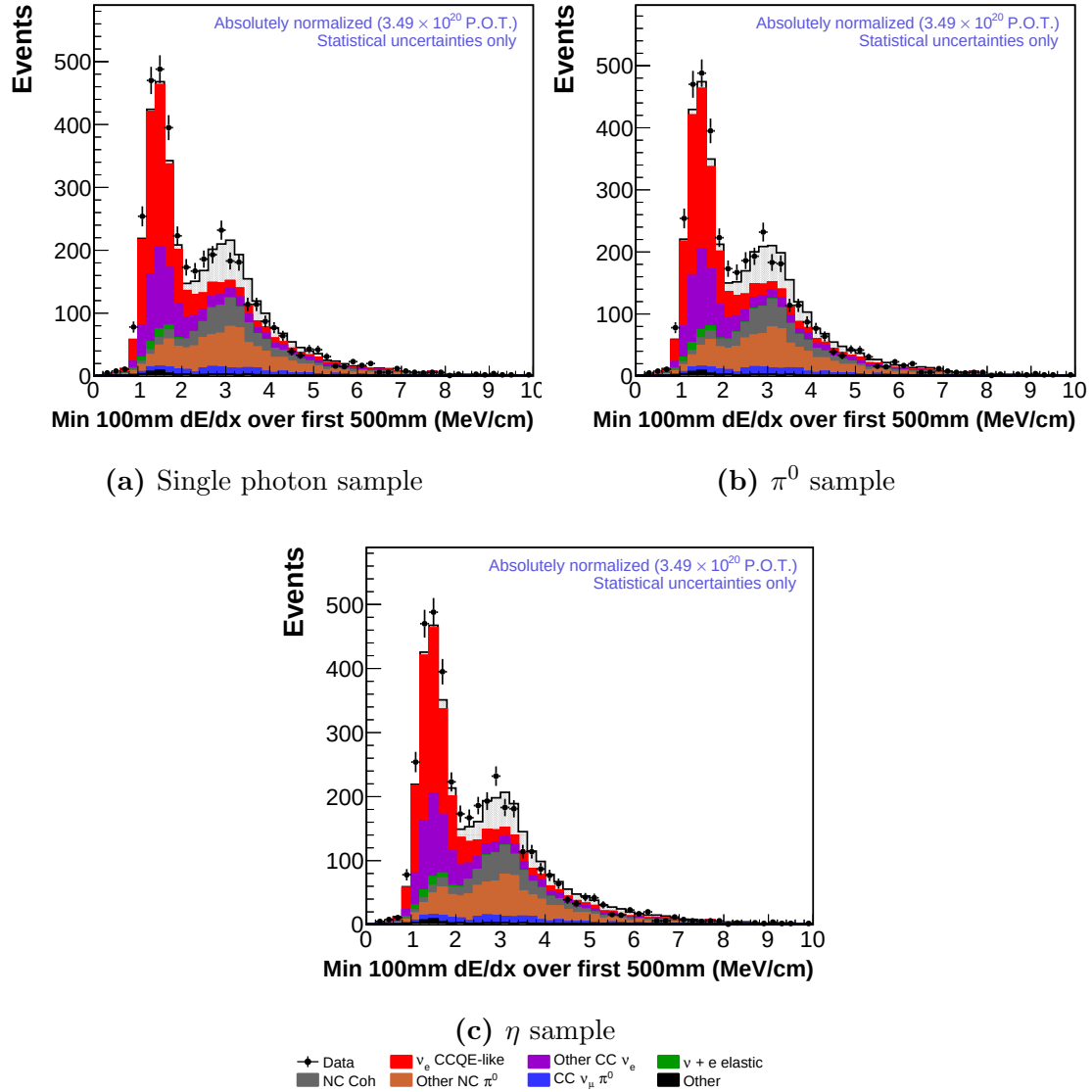


Figure B.20: Front dE/dx distributions comparing the fitted *ad hoc* samples added to the constrained GENIE prediction with the data. The topmost (shaded gray) histogram is the *ad hoc* model in each case.

B.3.2. Differentiating between particle hypotheses

Particle identification variables

Our primary objective is to select a single candidate model we can use to extrapolate into the ν_e CCQE signal region. Selecting from the *ad hoc* models we constructed in sec. B.3.1 requires us to differentiate between single-photon showers and multi-photon particle decays where the photons are significantly boosted in the forward direction and the resulting showers have substantial overlap. This is a challenging proposition, even in a fine-grained detector like MINER ν A. Nevertheless, we have managed to find a few variables which provide some degree of separation. Two of them will be familiar from sec. 6.2.2.

Extra energy When a particle decaying to multiple photons does so at low to moderate kinetic energies (i.e., it has relativistic $\beta \lesssim 25$), the photon showers will frequently receive little enough boost to be fully separated. In that case we may hope to observe the photon that was not chosen as the reconstructed shower as “extra energy” (in the sense of sec. 6.2.2); we can contrast this behavior with that of single photons. Fig. B.21 illustrates this point (using the flat-spectrum original sample that the reweighted *ad hoc* samples were constructed from): note the separation in extra energy ratio between photons and neutral pions in the lower energy range, and the separation from etas for much of the energies shown. Intriguingly, though fig. B.1 appears to indicate that although the excess events prefer an energy range that is in large part above the separation power of this variable, the excess is nevertheless more consistent with single photons or neutral pions than etas.

Shower width Above the range in kinetic energy where separate photon showers begin to merge, there is still some range of energy where it is possible to distinguish

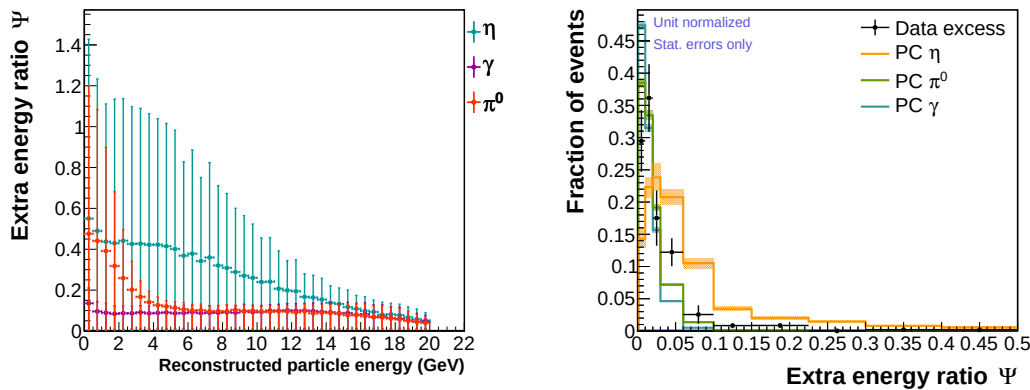


Figure B.21: Left: profile of extra energy ratio Ψ (defined in sec. 6.2.3) for various particle types in single-particle events. Each point represents the mean of the range of Ψ for that energy, while the error bars correspond to the standard deviation. Right: shapes of distributions of *ad hoc* samples in Ψ compared to that of the excess.

between single- and multi-photon showers based on their transverse size. That is, because a multi-photon shower in this “transition regime” in energy will have partial but not total overlap, the shower will appear to be broader in its transverse dimension. While we enforced a minimum shower width via the median plane shower width criterion of sec. 6.2.2 to find electromagnetic showers in the main selection, we can also contemplate examining its upper bound. Fig. B.22 this time indicates that the excess is more consistent with a multi-photon origin than a single-photon one, though here we are unable to differentiate between the neutral pion and eta hypotheses.

Transverse asymmetry The transverse energy distribution of multi-photon showers has one additional feature besides its width that can potentially be used to discriminate between the various particle types. Because the conversion of photons into electron-positron pairs is stochastic, and the radiation length in the MINER ν A plastic is large ($X_0 \sim 40$ cm) compared to the longitudinal resolution of the detector ($\sigma_z \sim 2$ cm), an event containing multiple photons will frequently see them convert far enough apart longitudinally that the track reconstruction (sec.

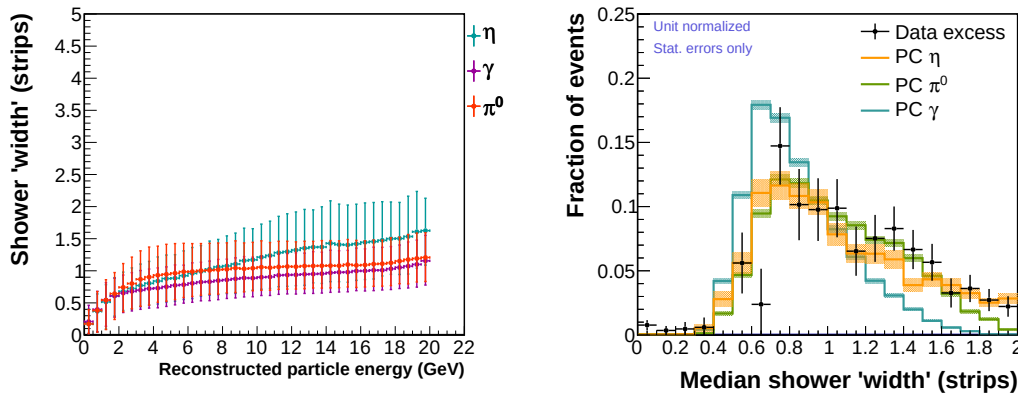


Figure B.22: Left: profile of median plane shower width (defined in sec. 6.2.2) for various particle types in single-particle events. Right: shapes of distributions of *ad hoc* samples in shower width compared to that of the excess.

6.1.3) effectively follows the earlier conversion. And since the direction of the track determines the reconstructed Cone axis (sec. 6.2.2), if the opening angle between the two photons is more than a degree or two, then the energy will (on average) be distributed asymmetrically around the cone axis. It is this asymmetry that we seek to measure.

We calculate the asymmetry of the event around the reconstructed axis as follows. The hits inside the reconstructed Cone object are first sorted according to the detector view (sec. 4.1) they appear in. The reconstructed Cone axis is then projected into each plane where a strip was illuminated and the strip it intersects recorded. Computing the energy-weighted mean of the difference between the illuminated strips x_i and the projected intersection point of the axis x_p yields an asymmetry score for the view (where i runs over the strips in the view):

$$A_{X,U,V} = \frac{1}{\sum_i E_i} \sum_i E_i (x_i - x_p) \quad (\text{B.3})$$

Now, though we have three values here, they are not all independent: we have summed across the detector $\hat{z}$ -coordinate, leaving us with only two degrees of freedom. To make this explicit, we recall from sec. 4.1 that the unit vectors for

the U and V directions are, respectively, $\frac{1}{2}(\hat{x} \pm \sqrt{3}\hat{y})$. Thus, we can recover the Cartesian coordinates with the combinations:

$$\hat{x} = \hat{u} + \hat{v} \quad (\text{B.4})$$

$$\hat{y} = \frac{1}{\sqrt{3}}(\hat{u} - \hat{v}) \quad (\text{B.5})$$

So we may as well represent the two measurements of asymmetry that we have represented here as an x -coordinate asymmetry and a y -coordinate asymmetry. However, to combine additive functions F_X , F_U , and F_V into an x -coordinate one $F_{\hat{x}}$, we should take into account the fact that the X planes sample this coordinate twice as frequently as the combination (U+V), and therefore the estimator from the X view should have error $\frac{1}{\sqrt{2}}$ of that of the estimator from (U+V). So, using an error-weighted mean, we have

$$F_{\hat{x}} = \frac{1}{\left(\frac{1}{2}\right)^{-2} + 1^{-2}} \left(F_X \left(\frac{1}{\sqrt{2}} \right)^{-2} + (F_U + F_V)(1)^{-2} \right) \quad (\text{B.6})$$

$$= \frac{1}{3} (2F_X + (F_U + F_V)) \quad (\text{B.7})$$

$$F_{\hat{y}} = \frac{1}{\sqrt{3}} (F_U - F_V) \quad (\text{B.8})$$

Therefore, we can construct the x - and y -coordinate asymmetries as:

$$A_{\hat{x}} = \frac{1}{3 \sum_{\text{all strips}} E_i} \left(2 \sum_{\text{X strips}} E_i(x_i - x_p) + \left(\sum_{\text{U strips}} E_i(x_i - x_p) + \sum_{\text{V strips}} E_i(x_i - x_p) \right) \right) \quad (\text{B.9})$$

$$A_{\hat{y}} = \frac{1}{\sqrt{3} \sum_{\text{U,V strips}} E_i} \left(\sum_{\text{U strips}} E_i(x_i - x_p) - \sum_{\text{V strips}} E_i(x_i - x_p) \right) \quad (\text{B.10})$$

To obtain a single metric describing the asymmetry, we can add $A_{\hat{x}}$ and $A_{\hat{y}}$ as we

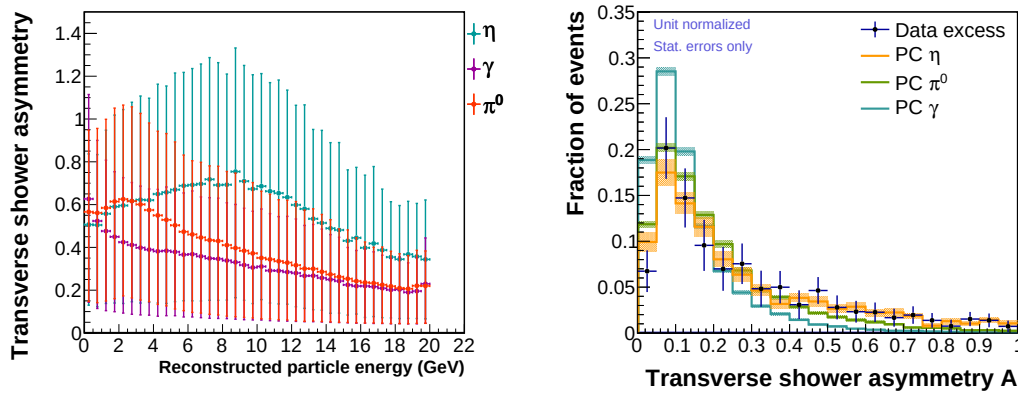


Figure B.23: Left: profile of transverse shower asymmetry (defined in the text) for various particle types in single-particle events. Right: shapes of distributions of *ad hoc* samples in transverse shower asymmetry compared to that of the excess.

would to get a distance in Cartesian coordinates:

$$A = \sqrt{(A_{\hat{x}})^2 + (A_{\hat{y}})^2} \quad (\text{B.11})$$

The values of A for different particle species, as well as a comparison of their shapes to the excess, are given in fig. B.23. Once again, we observe a preference for the multi-photon states over the single-photon.

Conclusions

Because the distributions in extra energy ratio rule out the η hypothesis, and the distributions in shower width and transverse asymmetry rule out the γ hypothesis, we conclude that the π^0 hypothesis is the most appropriate for modeling of the excess.

C Alternate analysis used for ratio to ν_μ CCQE

In order to compare the ν_e CCQE cross-section measured in ch. 7 to the existing MINER ν A measurement of ν_μ CCQE [21] in a consistent way, it was necessary to construct a version of the cross-section measurement which used an alternate signal definition (events which were CCQE according to the event generator), which used a smaller range of neutrino energy acceptance ($1.5 \leq E_{nu} \leq 10$ GeV), and which did not use the flux constraint of sec. 3.4. (The reasoning is further discussed in sec. 7.9.)

The observable distributions in the sidebands (figs. C.1-C.2), the efficiency corrections (fig. C.3), and the final result for this alternate analysis (fig. C.4) are shown below. (Though they contain different content because the distributions just listed are different, the background-subtracted and unfolded distributions are not shown here because the methodology for producing them is the same, and the differing signal distribution does not have a direct effect on them.)

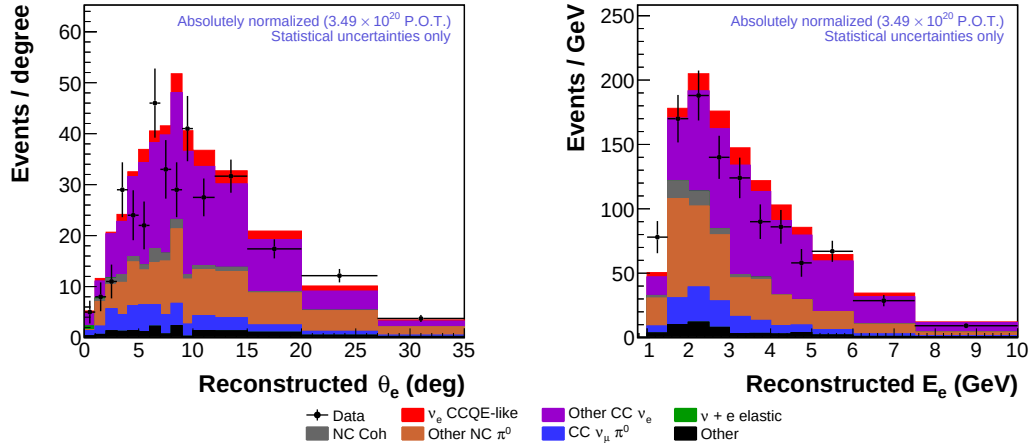


Figure C.1: Distributions in the observables from events in the extra energy sideband for the alternate analysis (app. C).

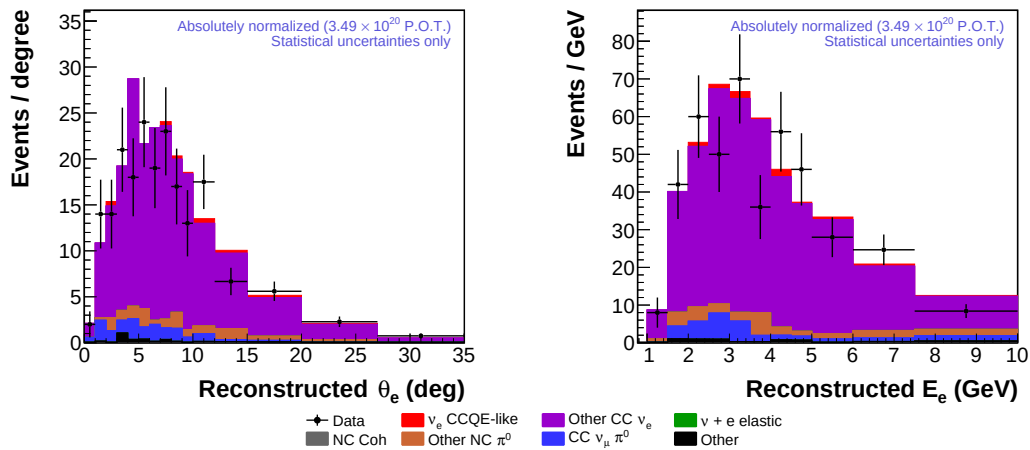


Figure C.2: Distributions in the observables from events in the Michel-match sideband for the alternate analysis (app. C).

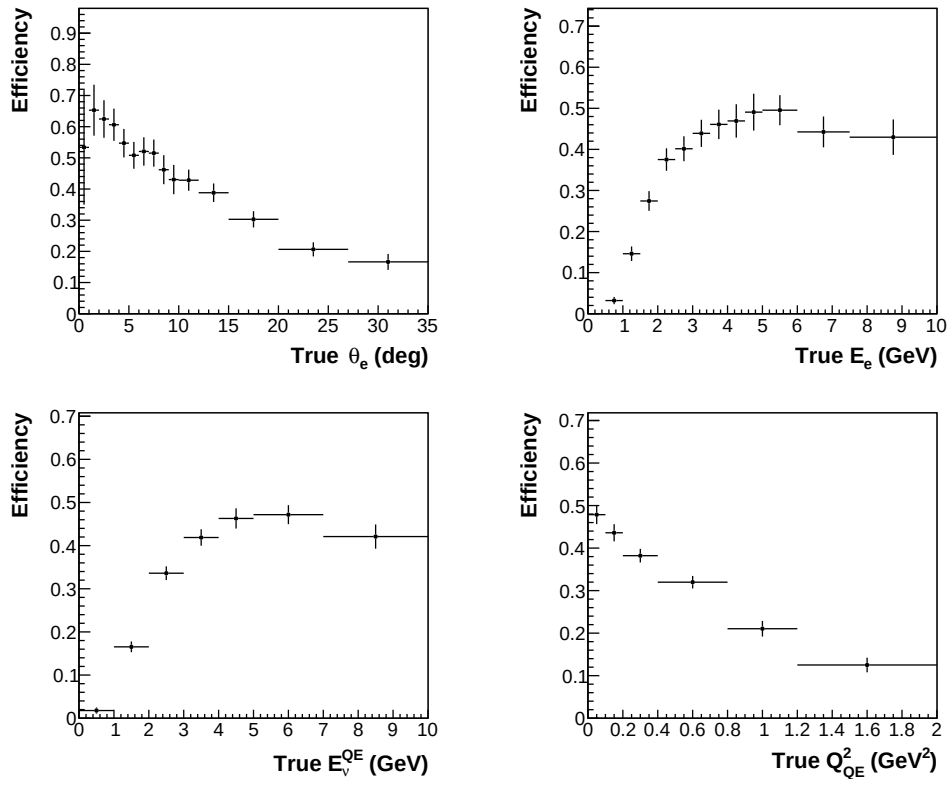


Figure C.3: Efficiency predictions in the cross-section variables, after correcting for the shape of the data, in the alternate analysis (app. C).

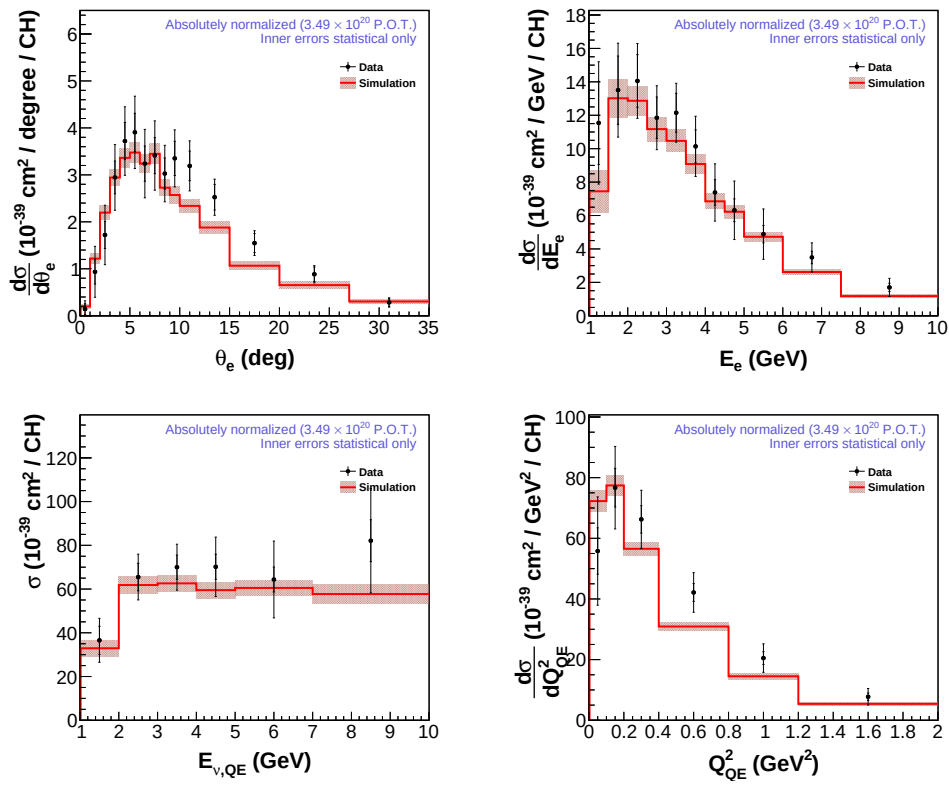


Figure C.4: Measured cross-sections in the alternate analysis (app. C).

D List of MINER ν A collaborators

G. Tzanakos

University of Athens

J. Cravens, M. Jerkins, S. Kopp, J. Ratchford

University of Texas at Austin

D.A. Martinez Caicedo, M.F. Carneiro, C.M. Castromonte, H. da Motta, G.A.

Fiorentini, A. Ghosh, K. Hurtado, J.L. Palomino, H. Yepes-Ramirez

Centro Brasileiro de Pesquisas Físicas

D.W. Schmitz

University of Chicago

J. Mousseau, B. Osmanov, H. Ray, D. Rimal, M. Wospakrik

University of Florida

L. Bellantoni, M. Betancourt, D. Boehnlein, R. DeMaat, N. Grossman, D.A.

Harris, M. Kiveni, J.G. Morfin, J. Osta, E. A. Paschos, G.N. Perdue, L.

Rakotondravohitra, P. Rubinov, F.D. Snider, J.T. Sobczyk, R. Stefanski

Fermilab

A. Bravar, Y. Karadzhov

University of Geneva

Victor M. Castillo-Vallejo, J. Felix, Z. Urrutia, E. Valencia, G. Zavala

Universidad de Guanajuato

M.E. Christy, M. Datta, C.E. Keppel, W. Tan, T. Walton, L. Y. Zhu

Hampton University

A. Butkevich, S.A. Kulagin

Inst. Nucl. Reas. Moscow

G. Niculescu, I. Niculescu

James Madison University

E. Maher

Mass. Col. Lib. Arts

L. Fields, B. Gobbi, J.A. Hobbs, C.E. Patrick, L. Patrick, H. Schellman

Northwestern University

N. Tagg

Otterbein University

S. Boyd, S.A. Dytman, B. Eberly, Z. Isvan, C.L. McGivern, D. Naples, V. Paolone,

L. Ren

University of Pittsburgh

M. J. Bustamante , E. Endress, A.M. Gago, J.P. Velásquez

Pontificia Universidad Catolica del Peru

R. Napora

Purdue University Calumet

S. Avvakumov, A. Bercellie, A. Bodek, R. Bradford, H. Budd, J. Chvojka, M. Day,
G.A. Díaz , R. Fine, T. Golan, A. Gomez, A. Higuera, J. Kleykamp, H. Lee, L.
Loiacono, S. Manly, C.M. Marshall, K.S. McFarland, A.M. McGowan, A. Mislivec,
J. Park, P.A. Rodrigues, D. Ruterbories, J. Wolcott

University of Rochester

T. Le, R.D. Ransome, E.C. Schulte, B.G. Tice

Rutgers University

O. Altinok, H. Gallagher, T. Kafka, W.A. Mann, W. P. Oliver

Tufts University

C. Simon, B.P.Ziemer

University of California at Irvine

J. Demgen, R. Gran, M. Lanari, E. Miltenberger

University of Minnesota at Duluth

G. Salazar, C.J. Solano Salinas, A. Zegarra

Universidad Nacional de Ingeniería

W.K. Brooks, R.Galindo, G. Maggi, J. Miller, C. Peña

Universidad Técnica Federico Santa María

L. Aliaga, J. Devan, M. Kordosky, J.K. Nelson, A. Norrick, J. Walding, D. Zhang

College of William and Mary