

Determination of neutron background for the studies of in-beam gamma ray spectroscopy using GEANT4

JACQUELINE B. KURODA

2025 NSF/REU Program
Department of Physics and Astronomy
University of Notre Dame

ADVISOR(S): Ani Aprahamian, Wanpeng Tan

Abstract

Cross section data of neutron induced reactions is critical for applications in nuclear physics, astrophysics, and nuclear technology. In nuclear physics, it is important to simulate and properly attribute potential contributions to measurements from incidental background reactions. The University of Notre Dame Nuclear Science Laboratory (NSL) has constructed and implemented a neutron beam of $10^8 \text{ cm}^{-2} \text{ s}^{-1}$ using the ${}^7\text{Li}(p, n){}^7\text{Be}$ reaction. The flux and energy distribution of neutrons is well characterized. This project is focused on determining the gamma ray background resulting from the neutrons in the target room. The intent is to determine the feasibility for in-beam inelastic neutron scattering experiments with backgrounds low enough to enable the measurements of gamma rays of interest. GEANT4 simulations were used to gain a better understanding of the neutron reactions and resulting gamma rays in the room. The setup was optimized to reduce the neutron background seen by the detector.

1 Background

The measurement of cross section data of neutron induced reactions are relevant in many disciplines such as nuclear physics, astrophysics and applications in nuclear technology. Cross section data is important because it helps predict neutron interactions. In reactor technology it is important to understand the interactions of neutrons to produce energy safely and efficiently. These cross section measurements are also fundamental for gaining an understanding of nucleosynthesis, the process by which new atomic nuclei are created in stellar environments through the slow and rapid neutron capture processes.

2 Introduction

To study neutron induced reactions, the University of Notre Dame Nuclear Science Laboratory (NSL) has constructed the Neutron Irradiation Station (NIS), a neutron beam with a high neutron flux and well characterized energy distribution [1]. NIS can generate nearly monoenergetic neutrons within a wide range of energies using the ${}^7\text{Li}(p, n){}^7\text{Be}$ reaction. NIS performs this reaction using monoenergetic protons delivered to the beam line by the 5U accelerator at the NSL, a 5 MV single end pelletron built by the National Electrostatics corporation. The protons are incident on a metallic Lithium target held in place by a copper target holder [1].

NIS is intended to be used for in-beam inelastic neutron scattering experiments, a type of neutron induced reaction. In this approach a beam of neutrons is incident on a target whose material properties are to be investigated. The incoming neutrons will excite the target nuclei and as these excited states decay they will emit characteristic gamma rays. With these experiments the intent is to perform in-beam gamma ray spectroscopy during the neutron irradiation, enabling the detection of decay products in real time. This method would allow the detection of short-lived states, providing valuable insight into the structure and the properties of the target nucleus.

As the nucleus of interest is being investigated, fast neutrons produced by the source are emitted around the room. These fast neutrons will undergo collisions in the room, slowing them down; this process is called moderation. Neutron moderation occurs until the neutrons reach thermal equilibrium with their surroundings, at which point they are referred to as thermal neutrons. The thermalized and fast neutrons in the room will interact with the detector creating background gamma rays. These gamma rays are produced by thermal neutrons undergoing neutron capture and fast neutrons undergoing inelastic scattering off detector nuclei. These gamma rays can hide the actual signal of interest, so to understand the signals seen by the detector it is important to simulate and properly attribute potential contributions to measurements from incidental background reactions. As well as these neutrons hiding the actual signal of interest they could also cause radiation damage to the detector, harming the detector.

3 Methods

To evaluate the background and to optimize the shielding around the detector for experiments using NIS, simulations play a critical role. Monte Carlo methods are widely used in the simulation of physical processes; they are computational techniques that use random sampling to model the probability of different outcomes in a statistical system. In physics, many different softwares use these methods to simulate the interactions of particles in matter, and one of the most widely used is GEANT4.

3.1 GEANT4

GEANT4 is a software toolkit developed using C++ for the simulation of the passage of particles in matter using Monte Carlo methods [2]. Developed by CERN for use in high-energy physics, it was later extended to other fields with applications including particle physics, nuclear physics, and accelerator design. To simulate the passage of particles through matter, GEANT4 uses a variety of physics and theoretical models; these models describe the electromagnetic and hadronic interactions of particles as they propagate through matter.

Properly simulating hadronic interactions is important in the evaluation of the background seen by the detector. In this work the QGSP_BIC_AllHP list was chosen to simulate some of the hadronic physics as it is the recommended package for simulating hadronic interactions with protons with matter below 20 MeV [3]. The QGSP_BIC_AllHP list is a high precision neutron model made from a combination of other GEANT4 high precision models; the list accurately simulates various physical processes, such as inelastic neutron interactions and neutron capture processes. The list is built from a set of cross sections and interaction models utilizing G4TENDL, the GEANT4 version of the TENDL nuclear library, which derives its data from calculations using the TALYS nuclear model code. The QGSP_BIC_AllHP list cannot be used directly because the list is still under development and is not available for default use. As a result, the list must be configured manually by having a user environment variable pointing to the data directory. The G4HadronElasticPhysicsHP list was also used for the elastic scattering of hadrons, providing high precision models for neutrons below 20 MeV. For simulating the other physical processes GEANT4 standard physics lists were used.

4 Simulations

4.1 Energy and Angular Dependence of the ${}^7\text{Li}(p, n){}^7\text{Be}$ Reaction

Before developing the full model of the shielding for the fast neutron source, the energy and angular dependence of ${}^7\text{Li}(p, n){}^7\text{Be}$ reaction was simulated to gain an understanding of GEANT4 to ensure

that the neutron production reaction is treated properly in the simulation. The geometry for this simulation was all set up using GEANT4. A 2 μm thick disk of lithium-7 acted as the lithium foil used to produce neutrons in NIS. The detector was a cylinder measuring 2 inches in diameter and 2 inches tall, that recorded the neutron's kinetic energy when they entered the detector volume. The angle of the detector was varied from 0-180° in 10° steps in order to test the angular dependence of the ${}^7\text{Li}(p, n){}^7\text{Be}$ reaction. To produce neutrons, the GEANT4 particle gun class was used to fire monogenetic protons that varied from 2.2-4.2 MeV at the Li disk in order to test the energy dependence of the reaction.

The angular dependence of ${}^7\text{Li}(p, n){}^7\text{Be}$ reaction was simulated using 1.9 MeV protons created by GEANT4's particle gun class and varying the detector angle from 0-180° in 10° steps. The mean neutron energies were found by fitting a Gaussian to the simulated spectrum. This was then compared to the calculated values of neutrons produced at this energy for the various angles (Figure 1) as calculated by [4].

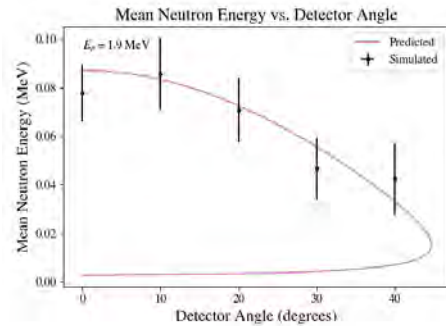


Figure 1: Shows the mean neutron energy produced at various angles by the ${}^7\text{Li}(p, n){}^7\text{Be}$ reaction from 1.9 MeV protons compared to the calculated values [4]

The energy dependence of the ${}^7\text{Li}(p, n){}^7\text{Be}$ reaction was simulated using monogenetic neutrons using GEANT4's particle gun class with proton energies ranging from 2.2-4.2 MeV in 0.1 MeV steps. The mean neutron energies were again found by fitting a Gaussian to the simulated spectrum. This was then compared to the calculated values of neutrons produced at this angle for the various energies (Figure 2) as calculated by [4].

The simulated data for both the angular and energy dependence of ${}^7\text{Li}(p, n){}^7\text{Be}$ reaction fit the

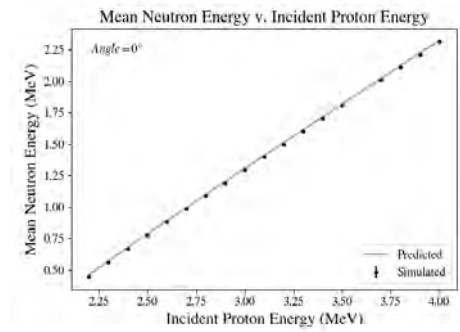


Figure 2: Shows the mean neutron energy produced at various energies by the ${}^7\text{Li}(p,n){}^7\text{Be}$ reaction from 2.2-4.2 MeV protons, with the detector at 0° compared to the calculated values [4]

data well verifying that the simulation is accurately reproducing the neutron production.

4.2 Varying Lithium Thickness to Study Neutron Energy Distribution

In NIS, the Li foil used for neutron production is created by the evaporation of lithium metal onto a copper target backing. The perimeter of this is water cooled in order to minimize the heating of the foil [1]. The water cooling runs around the perimeter to minimize the moderation of produced neutrons. The design of this was made in CAD a 3D modeling software and imported to GEANT by converting the CAD file to a gdmf file using FASTRAD, a 3D radiation analysis and shielding optimization software (Figure 3).

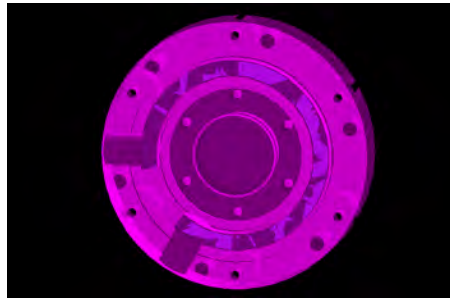


Figure 3: The foil holder geometry created in CAD imported to the GEANT4 simulation

To understand the neutron energies that would be produced in the room for different target thickness, the thickness of foil was varied and with the imported foil holder geometry in place. The neutrons were detected in the same way as previously done in the energy and angular dependence of ${}^7\text{Li}(p,n){}^7\text{Be}$ reaction, with the detector at a 0° and using 2.2 MeV protons.

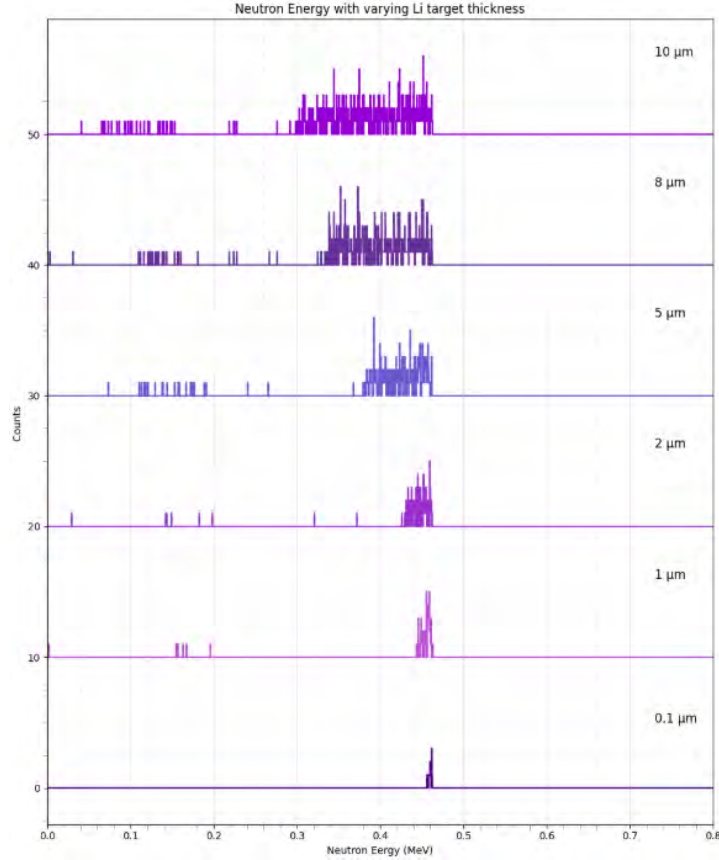


Figure 4: The neutron energy detected with the detector at 0 ° angle for, 10μm, 8μm, 5μm, 2μm, 1μm, 0.1μm thick lithium foils

The thicker the lithium foil used in the GEANT4 simulation the more neutrons were produced and the spread in neutron energy was larger. More neutrons are produced in the thicker foils as the protons have a longer path to travel to go through the foil so there is more chances for them to undergo the ${}^7\text{Li}(p, n){}^7\text{Be}$ reaction. This also leads to the larger spread of neutron energy from the thicker foils as the protons will lose more energy as they have more electromagnetic interactions with the electrons in the lithium slowing them down. These electromagnetic interactions happen more in the thicker foils as the protons have more distance to travel so more electromagnetic interactions happen leading to a larger spread in proton energies producing the neutrons, causing a larger spread in neutron energy.

4.3 Optimization of Shielding with Simulations

To effectively shield neutrons created by the ${}^7\text{Li}(p, n){}^7\text{Be}$ reaction borated-polyethylene was chosen to be the shielding material. The polyethylene acts as a neutron moderator slowing down neutrons through collisions with the hydrogen atoms. The boron effectively absorbs neutrons as boron-10 has a high neutron capture cross section. With these combined in the borated-polyethylene, the polyethylene can slow down fast neutrons from the source and the boron can absorb the slow thermalized neutrons from the room and the ones that have been slowed down in the polyethylene. The borated-polyethylene was made using G4Material, G4Element and G4Isotope, these are GEANT4 classes that allow for the creation of custom made materials.

To optimize the shielding around the detector, three shielding set ups were tested to see how the neutron energy spectrum seen by the detector changed. The detector was a clover germanium detector consisting of four germanium crystals and an aluminum shell and it was placed in a 8x8x6 m room with 1.2 m thick concrete walls to get the interactions of fast neutrons with the room, producing thermal neutrons in the simulation. The first set up was the detector with no shielding to set a baseline (Setup A), then shielding was placed around the detector and no shielding by the neutron source (Setup B), and the last set up was shielding around the detector and the neutron source (Setup C). The shielding around the detector was made in CAD and converted to a gdml file using FASTRAD. The detector was placed at two angles relative to the target position 60° (Figure 5.a) and 90° (Figure 5.b). It went around the target position and not the neutron source to test shielding as in an actual experiment the neutron beam will be irradiated a target of interest whose decay products are to be tested.

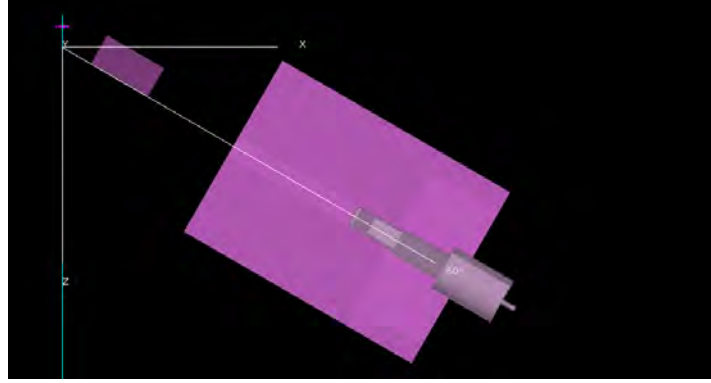


Figure 5: Shows the configurations with shielding around the Ge detector and the neutron source with the detector at 60°

To test the effectiveness of the shielding setups, the neutron energy seen by the crystals was recorded. 0.15 MeV neutrons were produced emitted from the source in a solid angle around the detector, using the GEANT4 particle gun class. They were not emitted isotropically from the source to save on computational time.

Tests of shielding configurations showed that the shielding round the detector reduced the background the neutrons create in the energy spectra of the detector especially at higher energies (figure 6). The test also showed that the shielding placed around the neutron source showed limited additional attenuation, suggesting that most of the reduction in background created by the neutrons is achieved by shielding around the detector and not the neutron source itself (table 1).

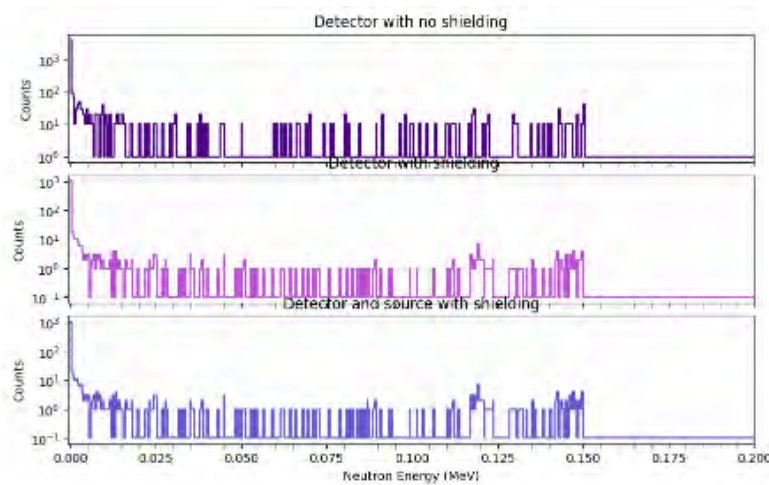


Figure 6: Shows the neutron energy spectrum seen by the Ge detector by the interactions of the neutrons with the detector

Table 1: Detected neutrons for three shielding setups with detector angles 60° and 90°. Each simulation used 1,000,000 events.

Detector Angle	Setup A	Setup B	Setup C
60°	584	140	134
90°	757	226	227

5 Conclusion

This work focused on using GEANT4 to model the ${}^7\text{Li}(p,n){}^7\text{Be}$ reaction and to investigate the shielding effects on neutrons produced from this reaction. The angular and energy dependence of the neutron production via the ${}^7\text{Li}(p,n){}^7\text{Be}$ reaction was tested and showed agreement with expected behavior. Simulations were also performed with varying lithium foil thicknesses, revealing that thicker foils result in both a broader neutron energy distribution and a higher neutron yield. Tests of different shielding configurations demonstrated that shielding placed around the detector significantly reduced neutron-induced background in the energy spectra, particularly at higher energies. In contrast, shielding near the neutron source provided minimal additional attenuation.

References

- [1] M. Matney, A. Simon, D. Robertson, K. Manukyan, A. Aprahamian, D. Bardayan, R. deBoer, E. Stech, W. Tan, and M. Wiescher, Nuclear Instruments and Methods in Physics Research Section A **1075**, 170406 (2025).
- [2] S. Agostinelli, J. Allison, K. Amako, J. Apostolakis, H. Araujo, P. Arce, *et al.*, Nuclear Instruments and Methods in Physics Research Section A **506**, 250 (2003).
- [3] Y. Lu, Z. Xu, L. Zhang, Z. Wang, T. Li, and M. S. Khan, Nuclear Instruments and Methods in Physics Research Section B **506**, 8 (2021).
- [4] S. Sjeue, Relativistic kinematics calculator for two-body nuclear reactions, skisickness.com (2020), online tool and plotter, last modified 22Feb2020.

Electronic Transport in Ti_4MnBi_2

NOAH LAPE

2025 NSF/REU Program
Department of Physics and Astronomy
University of Notre Dame

ADVISOR: Prof. Kateryna Foyevtsova

Abstract

We aim to model the electronic transport properties of Ti_4MnBi_2 near its 1D manganese spin chains by developing a tight-binding model. We start with ab initio magnetic Density Functional Theory (DFT) calculations of the electronic band structure of Ti_4MnBi_2 and an idealized case of the 1D manganese spin chains isolated from the rest of the material. We then develop a six-orbital basis Hamiltonian for our tight-binding model using physical intuition gained from 3D Wannier function visualizations and orbital projections onto DFT bands. Using a gradient descent machine learning approach, we find the optimal hopping integrals and onsite energy values for our tight-binding model. Next, we downfolded our Hamiltonian to three orbitals to minimize complexity while still capturing the structure of our bands near the Fermi surface. We find that our model makes use of only near-neighbor hoppings and that our Wannier functions are well localized in space. Our model produces a flat band at the Fermi surface. We find that the notable hoppings in our model produce a special 1D geometry that is remarkably similar to the well-studied sawtooth lattice known to host nontrivial flat bands. We show that this model can produce perfectly flat bands as a result of destructive interference and propose a potential localized molecular orbital to explain the magnetic moment localization in Ti_4MnBi_2 .

1 Introduction

Ti_4MnBi_2 is a worthwhile material to investigate, as it is one of a select few materials that house one-dimensional spin chains that are strongly coupled to a sea of conduction electrons.[1] In contrast, most materials that have been discovered to house one-dimensional spin chains are insulators. Additionally, it has been observed that this spin chain consists of spin $\frac{1}{2}$ moments when manganese atoms most often host spin $\frac{5}{2}$ moments. Experimental results indicate exotic physics may be found in this system encompassing magnetic frustration that arises from competing antiferromagnetic and ferromagnetic phases, kondo physics, strong correlation, and heavy fermions. But in order to better understand the rich physics of this system, a kinetic model for electronic transport must be developed. Later, many body techniques may be deployed to add to our model and develop a stronger understanding of how electron-electron interactions effect the richness of our material's physics.

By investigating our system from the perspective of a tight-binding model we are able to develop one such kinetic model for electronic transport. A tight-binding model is a method used to describe how electrons move in a crystal starting with the perspective that they are mainly strongly

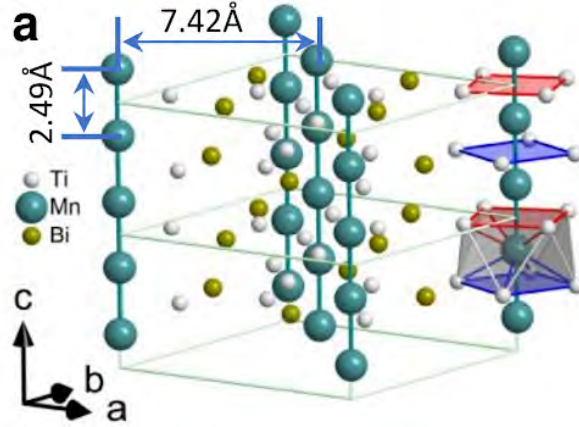


Figure 1: Diagram of the Crystal Structure of Ti_4MnBi_2 .

bound to individual atoms but have the ability to "hop" to others nearby in space. In this article we walk through the development of the tight-binding model and using our results make preliminary efforts to explain the source of experimentally observed electron localization and spin $\frac{1}{2}$ moments in Ti_4MnBi_2 .

2 Developing a Tight-Binding Model

2.1 Magnetic Density Functional Theory

In order to develop a tight-binding model for Ti_4MnBi_2 we first need to know its band structure. We can get this computationally by using a technique called Density Functional Theory (DFT). Naively, if you want to solve for the band structure of a material you would need to solve a many-electron Schrodinger equation. However, this involves taking into account electron-electron interactions for every electron in a system which is computationally impossible. DFT is a technique that approximates the solution to the many-electron Schrodinger equation by looking at the density of electrons as opposed to individual electrons. This transforms a many variable problem into one that only requires three spatial variables to describe the electron density. Because the magnetic moments are very important to this systems behavior it is necessary to include them as well. This is done by using a refinement of DFT simply called magnetic Density Functional Theory that takes into

account how electron spin effects the band structure.

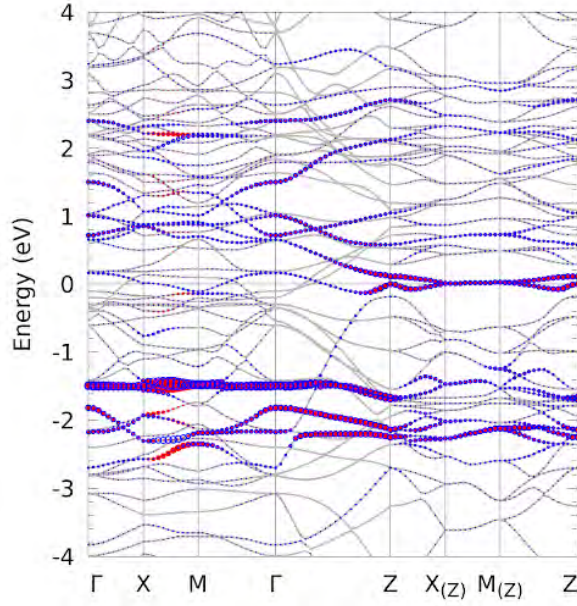


Figure 2: DFT calculated band structure of Ti_4MnBi_2 .

Using this approach we can create an accurate graph of the electronic band structure. However, due to frequent band crossings, splittings, and hybridization the band structure is difficult to work with. In order to simplify it we can calculate the band structure of a hypothetical material that only consisted of a manganese spin $\frac{1}{2}$ chain and the closest titanium atoms to it. This is the minimal model required to capture potentially novel physics regarding the manganese spin $\frac{1}{2}$ chain. We can be fairly sure we are not missing any important features by looking at orbital contributions to each band in the band structure. We find that predominantly near the Fermi surface (0 eV) and in the region surrounding our manganese chain only orbitals from the manganese and nearby titanium atoms contribute to the kinetic behavior of electrons. Thus, we are only concerned with these bands that have this orbital character denoted by the purple and green color in future analysis.

This band structure while simplified is still difficult to work with for many of the same reasons that the original band structure was difficult to work with. Due to interactions with unimportant orbitals the shape of our bands near the fermi level that have strong contributions from our orbitals of interest is left vague. Due to this, traditional methods can not be used to develop a tight-binding

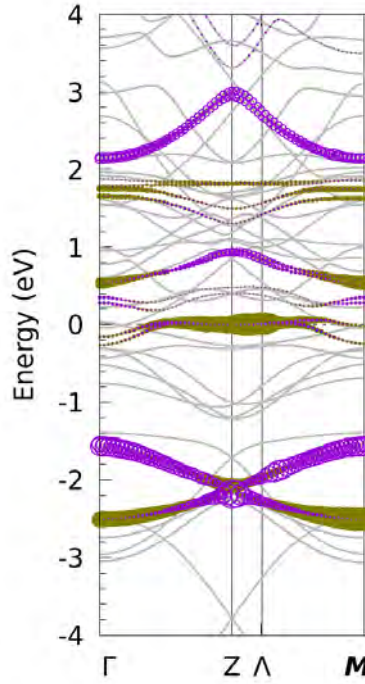


Figure 3: DFT calculated band structure of manganese spin $\frac{1}{2}$ chain.

model. We will discuss how we overcame this obstacle in section 2.3.

2.2 Creating a Hamiltonian

First we have to create a Hamiltonian which in the context of the tight-binding model is a mathematical object that describes the total energy of an electron if it is alone moving through our crystal field. Our Hamiltonian has two types of parameters which are onsite energies and hopping integrals. Onsite energy refers to the potential energy and electron has if it was placed on an atom in our crystal. Hopping integral refers to how easily an electron can move from one orbital to another. Using these two ideas we can encode both potential and kinetic energy into our Hamiltonian which we represent as a matrix. The columns and rows represent distinct orbitals and entries represent how those orbitals are connected to each other.

By looking at Ti_4MnBi_2 's crystal structure we can identify orbitals (or linear combinations of orbitals) and their position in space that might be relevant to recreating the bands we see. We made assumptions about which orbitals could hop to which mainly based on their proximity to others in

space. In this process it is okay if we do not include an orbital or hopping that may be important as during the optimization process the effective hopping of other degrees of freedom, our hopping integrals and onsite energies, will change to compensate. In fact, the real space representation of our orbitals in a sense is not "real" because it is not uniquely defined or measurable. This is because in a solid an electron does not actually hop from one atomic orbital to another but instead moves as a delocalized wavefunction. The more important concept to capture is the Bloch wave solution given to us by our band structure which is uniquely defined and measurable.

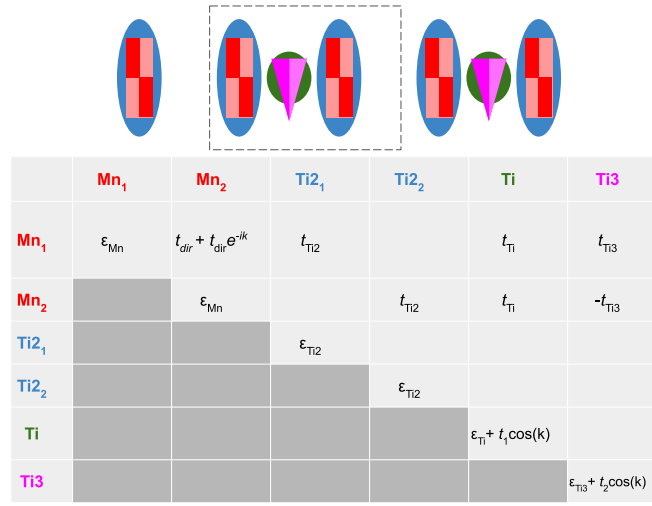


Figure 4: Cartoon and Hamiltonian of 6-orbital model of Ti_4MnBi_2 .

Finally, it is important that the Hamiltonian be a 6x6 matrix. This is because we want six solution branches so that we can match the six solution branches in the magnetic DFT band structure that have strong contribution from the orbitals we are interested in. Eventually we will want to consolidate our model into a 3x3 Hamiltonian so that we only capture the necessary physics at the Fermi level however due to the poor definition of our bands of interest at the Fermi level we need to incorporate other bands at first so we can be reasonably certain we get a good fit.

2.3 Fitting our Model Parameters

In the last step we constructed a 6x6 Hamiltonian for our system however we still do not know the values for the onsite energies or hopping integrals in that model. We now have to obtain them by fitting our model to our magnetic DFT band structure. By using these three additional bands that are far from the Fermi level as a guide rail we can be reasonably confident that if our fit matches those well then it also accurately describes the true nature of the bands near the Fermi level.

In order to get this fit we utilize a gradient descent algorithm. First we choose a few k-points spread across our x-axis and consisting of many points of interest throughout or Brillouin Zone. Then we find the energy of each band at that k-point. Next, we choose random starting conditions for our parameters. We then go through each parameter, one by one, and make it a bit smaller and a bit larger. As we do this, one change will make our curves fit to the predicted DFT energy values slightly better. After a few thousand iterations and slight improvements we end with a set of parameters that match the DFT band structure very well.

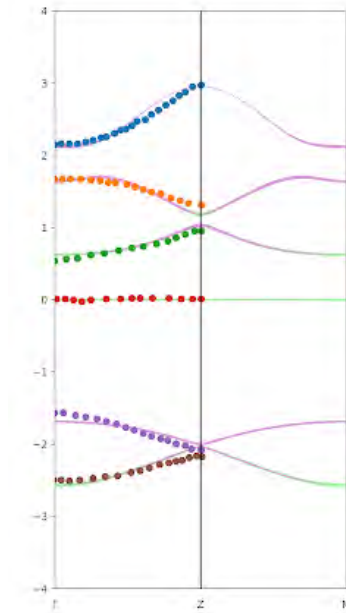


Figure 5: Best fit tight-binding model with DFT data overlaid.

Along with fitting to the shape of the band structure we also incorporated a mechanism that would fit to the orbital contributions of the band structure as well. We did this because in order to have an accurate tight-binding model it not only needs to fit the shape of the magnetic DFT band structure but also predict which orbitals are responsible for which bands.

We found that in this system there is a flat band directly on the Fermi surface. a flat band is the horizontal line and it represents that at that energy electrons have zero velocity or are stationary. Due to this, electron-electron interactions are more important than normal and a plethora of exotic phenomena could occur.

2.4 Downfolding

Now that we have a well fit 6x6 tight-binding model that captures the structure at the fermi level we want to get rid of the extra information. This helps us interpret the model and also makes it easier for many-body physicists to add electron-electron interactions into our model. We can achieve this through a computational process called downfolding which utilizes a Fourier transform to reduce our 6x6 Hamiltonian into a 3x3 Hamiltonian by removing the solution branches we no longer want, the ones far from the Fermi surface, and perfectly replicating the ones we are interested in.

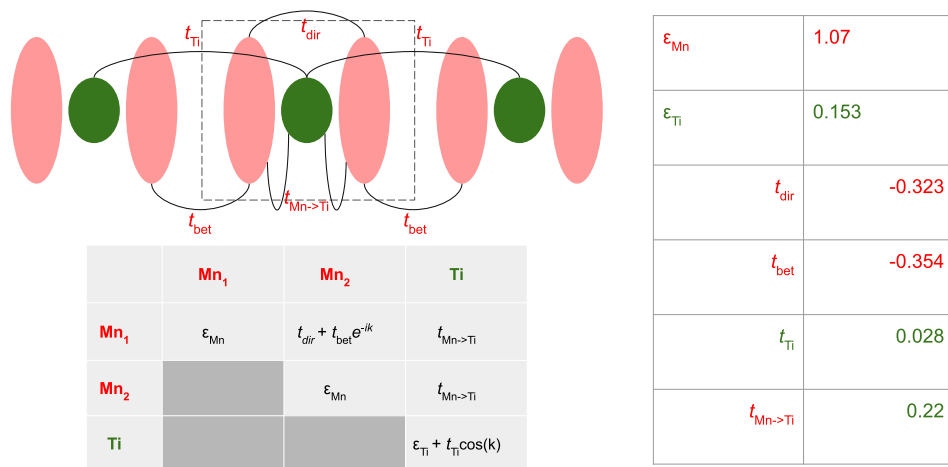


Figure 6: Cartoon and Hamiltonian of 3-orbital model of Ti₄MnBi₂.

Our results are promising. In order to replicate our bands to a high degree of accuracy we did not need to include high harmonics. In addition real space Wannier functions were constructed from this simplified Hamiltonian and found to be highly local. This indicates our original 6x6 Hamiltonian and fitting procedure worked adequately. Our model remains physical and useful with only close range interactions.

3 Investigating the Flat Band

We developed a solid tight binding model that reproduces the important DFT band structure and verified it is physical. Not only that, but we also found that there is a flat band at the Fermi surface in this model. We want to answer the question of where does the flat band come from? A flat band is always the result of an electron being localized but this can be done trivially or nontrivially. Trivially, an electron can be localized on an atom with no ability to "hop" to another atom in our model resulting in a flat band. Nontrivially, an electron can be free to move with strong couplings between all the atoms however the special geometry or topology of our system, results in destructive interference at some energy level which causes the electron to be localized to some group of orbitals.

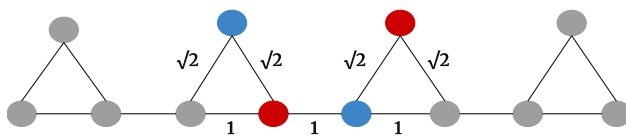


Figure 7: Depiction of our special geometry that allows for perfectly flat bands. Red and blue nodes represent anti-phase relationships and all together the colored nodes show the localized molecular orbital.

There is good evidence that our system hosts this type of nontrivial flat band and localized molecular orbital. This is because the geometry of our tight-binding model when the parameters are tweaked to be appropriate can host mathematically perfect nontrivial flat bands as shown. A

close match to our system called the sawtooth lattice has been studied extensively in research concerning one dimensional flat band systems.[2] The interesting finding about this interpretation is that our localized molecular orbital that is responsible for the flat band takes up the space of two unit cells. This predicts that magnetic moments are localized in a slightly broader region than previous theory had predicted.[1] This prediction also aligns with ARPES data better making it an exciting discovery. While our system does not quite have the parameters to be a perfectly flat band it does sit close in parameter space to the solution surface of parameter combinations that do.

4 Conclusion

We have successfully created a tight-binding model for the material Ti_4MnBi_2 to model electronic transport near the manganese spin $\frac{1}{2}$ chains. We then discovered a flat band at the Fermi surface and used our model and this finding to propose an explanation for the localized spin and magnetic moments found in the material that is in better agreement with ARPES data than leading theories. We also showed the geometry of this system most likely hosts non-trivial topological properties.

Further work should be done on verifying the nature of this localized molecular orbital and investigating the non-trivial topology of our system. Many-body physics should be added to our models including coulomb repulsion and other electron-electron interactions to see if any exotic physics emerges.

References

- [1] Li, X.Y., Nocera, A., Foyevtsova, K. et al. Frustrated spin-1/2 chains in a correlated metal. Nat. Mater. 24, 716–721 (2025). <https://doi.org/10.1038/s41563-025-02192-z>
- [2] Flach, S., Leykam, D., Bodyfelt, J. D., Matthies, P., & Desyatnikov, A. S. Detangling flat bands into Fano lattices. Europhysics Letters, 105(3), 30001 (2014). <https://doi.org/10.1209/0295-5075/105/30001>

Optimizing the Computational Solution of the SRG Flow Equation in *Ab Initio* Nuclear Many-Body Theory

ABBASS H. MOURTADA

2025 NSF/REU Program
Department of Physics and Astronomy
University of Notre Dame

ADVISOR(S): Prof. Ragnar Stroberg

Abstract

Solving the nuclear many-body problem from first principles (ab initio) remains one of the central challenges in theoretical nuclear physics, due to the complexity of inter-nucleon forces. Accurate yet efficient computational methods are needed to make progress in modeling medium-mass and heavy nuclei. This study addresses the need for improved numerical techniques to solve the similarity renormalization group (SRG) flow equation. We aim to optimize the computational performance of SRG solvers by introducing and evaluating strategies that enhance efficiency while maintaining accuracy. Specifically, we target issues of numerical stiffness, computational cost, and flow stability in evolving large-scale Hamiltonians by implementing commutator norm thresholding within SRG solvers, and evaluate their performance on realistic nuclear structures such as ^{14}C . Our findings demonstrate that these optimizations reduce computational time by up to

These results highlight the value of threshold-based strategies in scaling SRG techniques for broader many-body applications. This work contributes to the development of scalable, structure-preserving solvers in many-body theory, laying a foundation for future extensions to higher-body interactions and time-dependent systems.

1 Introduction

The nuclear many-body problem lies at the heart of theoretical nuclear physics. It involves describing the structure and dynamics of atomic nuclei starting from fundamental interactions between nucleons. The increasing complexity of nuclear forces with the number of nucleons requires powerful computational methods [1].

One such method is the In-Medium Similarity Renormalization Group (IMSRG), which utilizes a continuous unitary transformation to decouple low-energy states from high-energy ones. This transformation is governed by flow equations that evolve the Hamiltonian as a function of a flow parameter s [2]. The goal of the SRG transformation is to drive the Hamiltonian toward a band- or diagonal-dominated form, making the resulting many-body problem more manageable both analytically and numerically.

A distinguishing feature of SRG is its adaptability: depending on the choice of generator $\eta(s)$, the flow can be tailored to optimize convergence and computational simplicity. Because the equation changes rapidly, the system size grows very fast with more particles, and the calculated commu-

tators are complex and repetitive, solving them requires careful and efficient computational choices. [3; 4].

This work focuses on optimizing the computational aspects of solving the SRG flow equation via systematic code-based improvements, empirical benchmark testing, and theoretical error scaling analysis.

2 Background

The SRG framework evolves the Hamiltonian $H(s)$ via the flow equation:

$$\frac{dH(s)}{ds} = [\eta(s), H(s)], \quad (1)$$

where $\eta(s)$ is an anti-Hermitian generator responsible for suppressing the off-diagonal elements of $H(s)$. The canonical choice $\eta = [H_d, H]$, with H_d being the diagonal part of H , is widely used due to its proven convergence properties and computational feasibility [3]. However, alternative generator choices such as White's generator, which scales off-diagonal terms by their energy differences causing rapid decoupling, and the arc-tangent generator, which smoothly suppresses large off-diagonal elements by applying an arc-tangent function of energy denominators, have been explored in the literature to enhance convergence and reduce numerical instabilities [1].

In the IMSRG framework, the Hamiltonian is expressed in normal-ordered form and truncated at the two-body level, to become what is known as IMSRG(2) [4]. This truncation, although an approximation, captures most of the relevant physics for medium-mass nuclei and allows for practical computations.

The primary benefit of IMSRG is its ability to evolve operators consistently with the Hamiltonian, ensuring that observables are calculated in a transformed basis.

To further understand SRG's utility, we examine its application in a model system: a quantum particle confined in a one-dimensional infinite square well with a random potential. This example is conceptually simpler than full nuclear many-body problems but reveals enough about SRG's

mathematical and computational behavior.

The disordered system lacks analytical solutions, and therefore the finite-difference method is used to discretize the Schrödinger equation and build a Hamiltonian consisting of a kinetic matrix and a symmetric random potential matrix. The SRG flow equation used is

$$\frac{dH}{ds} = [[T, H], H], \quad (2)$$

where T is the kinetic energy operator and also serves as the generator for the equation. This double commutator formulation is known to drive the Hamiltonian toward a diagonal form while preserving its eigenvalues due to the unitary nature of the SRG transformations [5].

The results of this SRG flow are multifaceted. The wavefunctions retain their structure and symmetry after the evolution, verifying that the content of the system is preserved.

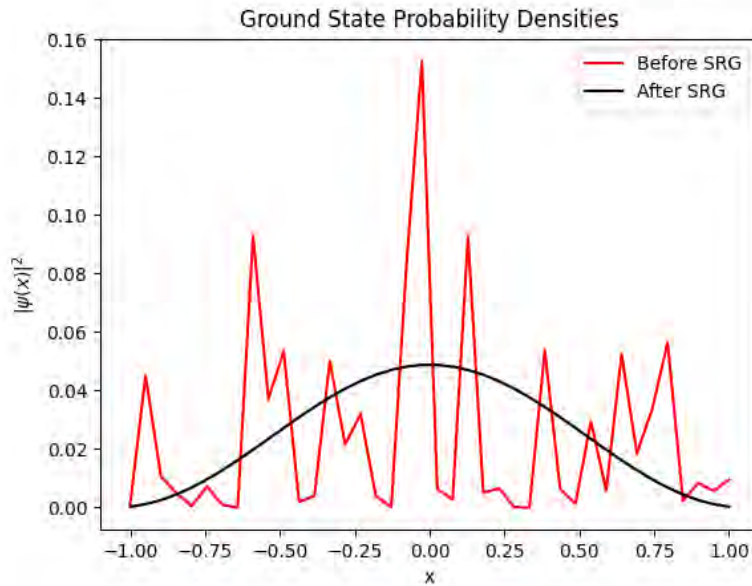


Figure 1: Ground state probability densities before and after SRG flow.

Energy eigenvalues remain conserved, confirming that the SRG flow is unitary.

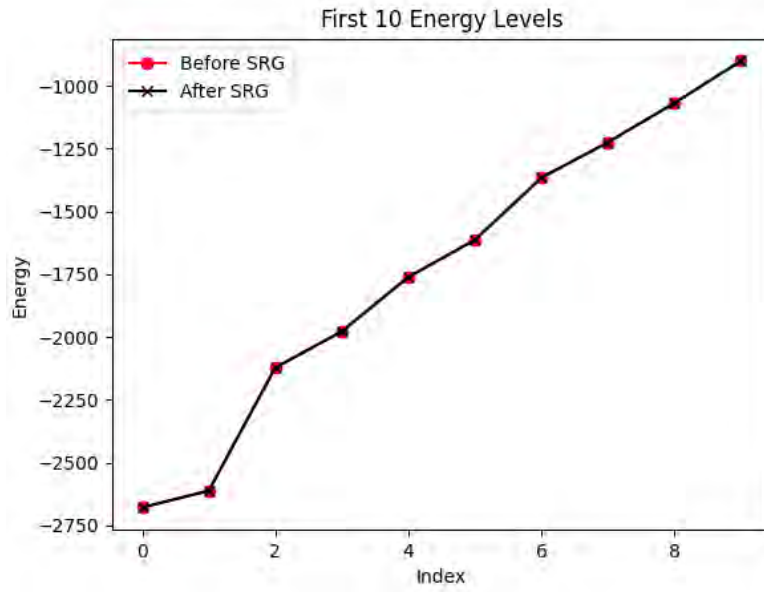
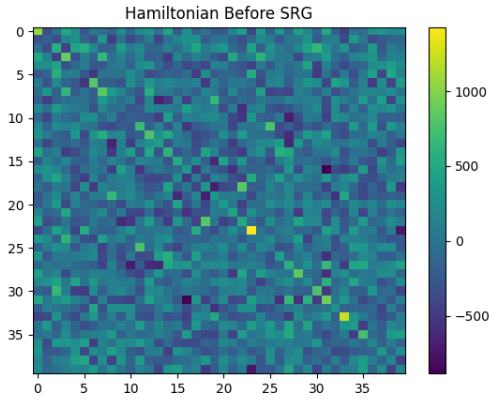
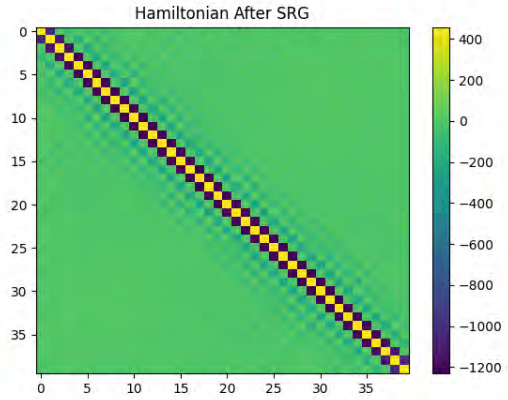


Figure 2: Energy eigenvalues before and after SRG flow.

The Hamiltonian matrix evolves from a densely entangled form to a nearly diagonal matrix.



(a) Hamiltonian before SRG flow.



(b) Hamiltonian after SRG flow.

Figure 3: Comparison of Hamiltonian matrices before and after SRG flow, showing suppression of off-diagonal elements.

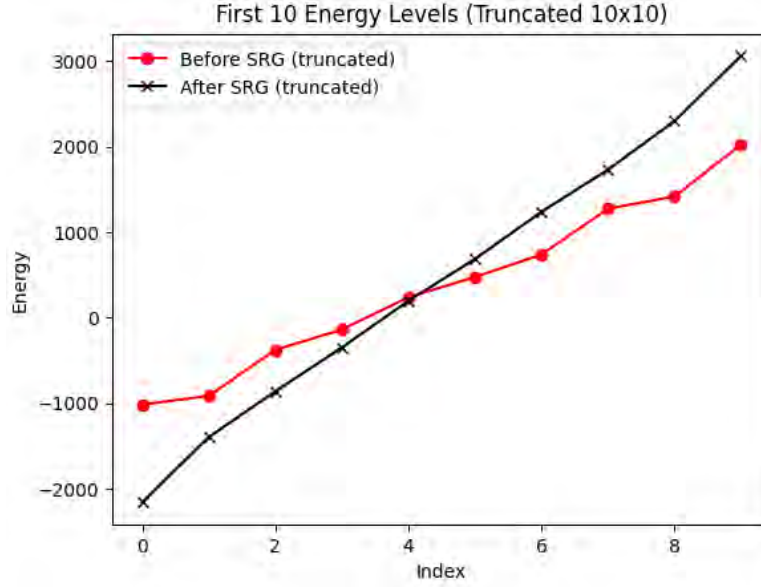


Figure 4: Truncated spectrum comparison after SRG. Some discrepancies arise due to loss of higher energy contributions.

This numerical application reinforces the utility of SRG techniques developed in nuclear physics and highlights their broader relevance to general quantum systems with complex interactions.

However, the computational cost of solving the flow equations increases rapidly with the size of the single-particle basis, motivating the need for optimization techniques discussed in the following section.

3 Methods

3.1 Computational Environment

All simulations were executed on the Center for Research Computing (CRC) cluster at the University of Notre Dame. The IMSRG codebase was compiled using Intel compilers with optimization flags for vectorization and parallelization. Dependencies included Python 3.x, NumPy, and CMake for build configuration. Bash scripts were used to automate job submission and data aggregation across flow parameter values. Intermediate data, including Hamiltonian norms and flow diagnostics, were stored in HDF5 format for post-processing.

3.2 Commutator Threshold Truncation

The SRG flow equation involves recursive nested commutators, with each derivative requiring the evaluation of increasingly complex expressions of the form:

$$\frac{dH}{ds} = [\eta, H], \quad \frac{d^2H}{ds^2} = [\eta, [\eta, H]], \quad \text{and so on.} \quad (3)$$

These nested commutators become increasingly computationally expensive, and their repeated evaluation is a major bottleneck in large-scale IMSRG calculations. As the flow progresses and off-diagonal elements are suppressed, many of the higher-order commutators contribute negligibly to the evolution of $H(s)$ as seen in Figure 5 where we ran the calculation with seven different thresholds.

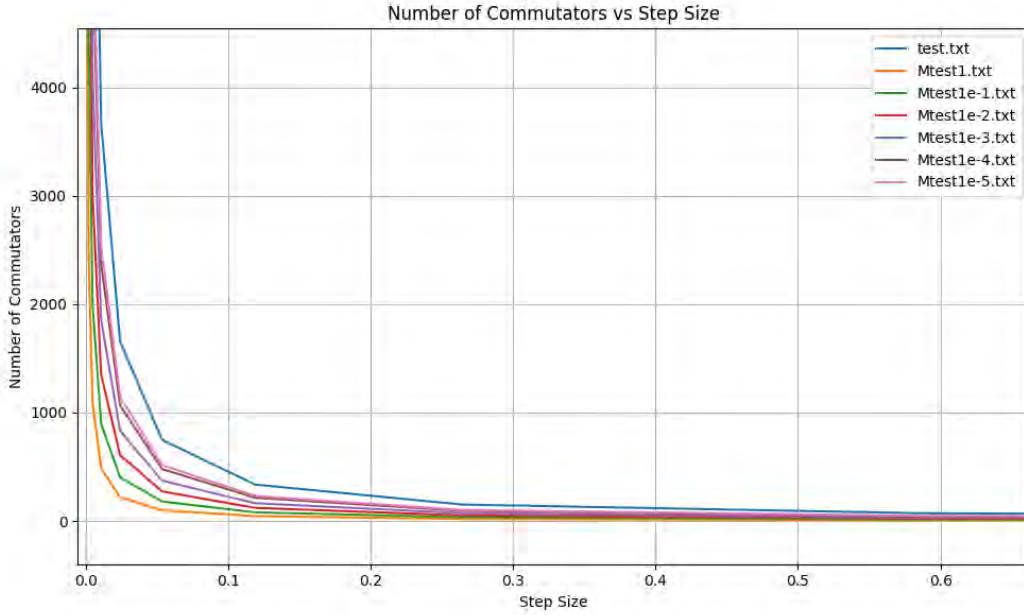


Figure 5: Number of Commutators vs Step Size for thresholds $[1, 10^{-5}]$.

However, the downside would be the resultant error which can be observed in the ground state energy calculation.

To exploit this, we introduced a thresholding mechanism that conditionally skips the evaluation of higher-order commutators based on their estimated magnitude. Specifically, each commutator

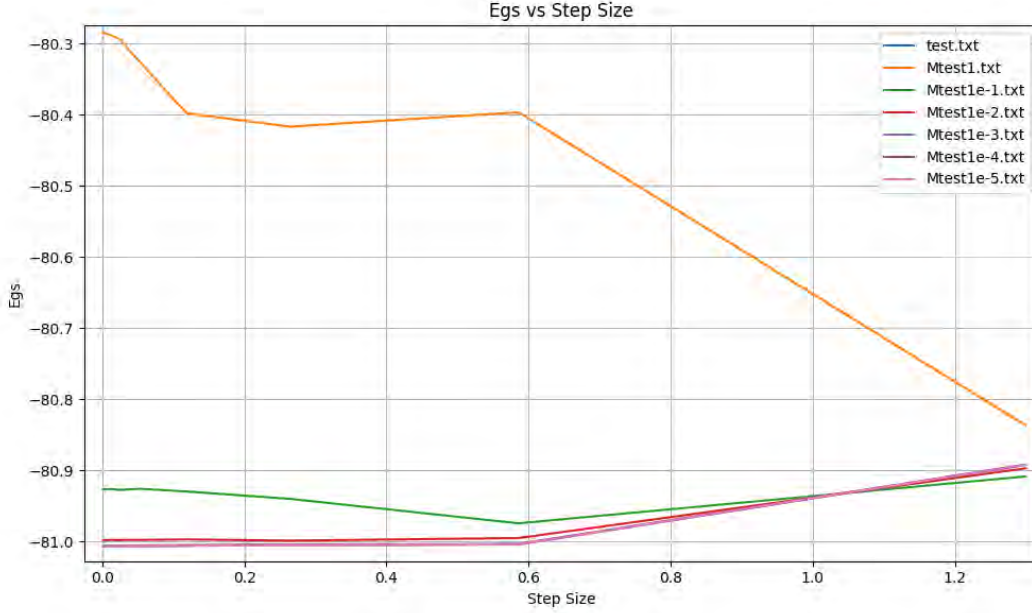


Figure 6: Grounds State Energy vs Step Size for thresholds $[1, 10^{-5}]$.

in the recursive chain,

$$C_n = [\eta, C_{n-1}], \quad (4)$$

was computed only if its norm satisfied the inequality:

$$\|C_n\| > \epsilon \cdot \|\eta\| \cdot \|C_{n-1}\|, \quad (5)$$

where ϵ is the threshold parameter (10^{-4} - 10^{-5}). This estimate is motivated by bounding the contribution of C_n using submultiplicativity of the operator norm. If the estimate falls below the threshold, we terminate the expansion at C_{n-1} , effectively truncating higher-order terms that are unlikely to impact the result significantly.

This method stems from our use of a Taylor expansion of the flow operator up to the fourth order:

$$H(s + \Delta s) \approx H(s) + \Delta s \cdot C_1 + \frac{(\Delta s)^2}{2!} C_2 + \frac{(\Delta s)^3}{3!} C_3 + \frac{(\Delta s)^4}{4!} C_4 + \dots, \quad (6)$$

where $C_n = \text{ad}_\eta^n(H)$ is the n^{th} nested commutator with $\text{ad}_\eta(H) := [\eta, H]$. Given that higher-order terms contribute only at $\mathcal{O}(\Delta s^n)$, their neglect becomes justifiable once their norms fall below the

set threshold.

4 Results

4.1 Efficiency Gains

Thresholding of commutators had a pronounced effect on performance, particularly in the latter stages of the SRG flow where the many nested commutator contributions fall below significance. BWe avoided evaluating terms that were unlikely to affect physical observables, thereby reducing redundant computation without sacrificing accuracy.

To quantify the performance benefits, a systematic benchmark was conducted. We ran 20 calculations with step sizes logarithmically spaced between 10^{-3} and 0.7, both without, with the thresholding and then halting conditions enabled. The results showed that enabling these optimizations reduced the cumulative runtime by approximately 80 seconds across the full range of simulations.

Figure 7 illustrates the runtime differences for each step size. The effect of thresholding is particularly prominent at small step sizes, where the increased number of integration steps would have otherwise led to a quadratic or cubic growth in commutator overhead.

# dsmax	N_Commutators	Egs	# dsmax	N_Commutators	Egs	# dsmax	N_Commutators	Egs
1.0000000e-03	40004	-81.0066884934	1.0000000e-03	27750	-81.0066884934	1.0000000e-03	33077	-81.006688493
1.41169866e-03	28336	-81.0066840189	1.41169866e-03	19655	-81.0066840189	1.41169866e-03	23428	-81.006684018
1.99289310e-03	20072	-81.0066778159	1.99289310e-03	13922	-81.0066778159	1.99289310e-03	16595	-81.006677815
2.81336452e-03	14220	-81.0066693106	2.81336452e-03	9862	-81.0066693105	2.81336452e-03	11755	-81.006669310
3.97162292e-03	10072	-81.0066569051	3.97162292e-03	6985	-81.0066569050	3.97162292e-03	8326	-81.0066569051
5.60673475e-03	7136	-81.0066400204	5.60673475e-03	4948	-81.0066400203	5.60673475e-03	5897	-81.0066400204
7.91501993e-03	5056	-81.0066162737	7.91501993e-03	3505	-81.0066162734	7.91501993e-03	4177	-81.0066162737
1.11736230e-02	3580	-81.0065818806	1.11736230e-02	2481	-81.0065818800	1.11736230e-02	2958	-81.0065818806
1.57737886e-02	2536	-81.0065350237	1.57737886e-02	1757	-81.0065350225	1.57737886e-02	2094	-81.0065350237
2.22678362e-02	1800	-81.0064734173	2.22678362e-02	1246	-81.0064734150	2.22678362e-02	1484	-81.0064734173
3.14354746e-02	1276	-81.0063856732	3.14354746e-02	882	-81.0063856683	3.14354746e-02	1051	-81.0063856732
4.43774173e-02	904	-81.0062656255	4.43774173e-02	625	-81.0062656163	4.43774173e-02	744	-81.0062656255
6.26475405e-02	640	-81.0061053626	6.26475405e-02	442	-81.0061053440	6.26475405e-02	526	-81.0061053626
8.84394480e-02	456	-81.0059080231	8.84394480e-02	313	-81.0059087858	8.84394480e-02	372	-81.0059080231
1.24849851e-01	324	-81.0056607301	1.24849851e-01	222	-81.0056606555	1.24849851e-01	264	-81.0056607301
1.76250360e-01	228	-81.0053617699	1.76250360e-01	155	-81.0053615954	1.76250360e-01	185	-81.0053617698
2.48812408e-01	164	-81.0050678694	2.48812408e-01	110	-81.0050675006	2.48812408e-01	131	-81.0050678693
3.51248142e-01	116	-81.0047248108	3.51248142e-01	78	-81.0047241361	3.51248142e-01	92	-81.0047248105
4.95856531e-01	84	-81.0042460268	4.95856531e-01	55	-81.0042454432	4.95856531e-01	65	-81.0042460264
7.00000000e-01	60	-81.0030259425	7.00000000e-01	40	-81.0030253463	7.00000000e-01	45	-81.0030259426
Runtime: 462.25 seconds			455.51 seconds			382.15 seconds		

Figure 7: Runtime comparison for 20 logarithmically spaced step sizes from 10^{-3} to 0.7.

5 Conclusion

This research demonstrated that SRG solvers benefit significantly from adaptive commutator thresholding and numerical halting. These tools maintain accuracy while dramatically improving efficiency. Implementation of these optimizations provides practical tools for extending IMSRG calculations to larger nuclei and more sophisticated observables. Future work will involve three-body interactions, machine-learning-guided thresholding, and real-time evolution of SRG-driven observables [1; 4].

References

- [1] H. Hergert, S. K. Bogner, T. D. Morris, A. Schwenk, and K. Tsukiyama, *Physics Reports* **621**, 165 (2016).
- [2] S. K. Bogner, R. J. Furnstahl, and R. J. Perry, *Physical Review C* **75**, 061001 (2007).
- [3] F. Wegner, *Annalen der Physik* **506**, 77 (1994).
- [4] K. Tsukiyama, S. K. Bogner, and A. Schwenk, *Physical Review Letters* **106**, 222502 (2011).
- [5] S. Kehrein, *The Flow Equation Approach to Many-Particle Systems*, Springer Tracts in Modern Physics, Vol. 217 (Springer, 2006).

Investigation of $^{11}\text{C}(\text{p}, \text{p})^{11}\text{C}$ Resonant Scattering in the TriSol System

MICHAEL POTTS

2025 NSF/REU Program
Department of Physics and Astronomy
University of Notre Dame

ADVISOR(S): Prof. Dan Bardayan, Dr. Patrick O'Malley

Abstract

The resonances of ^{12}N are of significant astrophysical interest, as they are related to the $^{11}\text{C}(\text{p}, \gamma)^{12}\text{N}$ reaction rate, which can significantly impact the evolution of metal-deficient massive stars (Population III stars) and inform us about carbon production in the early universe through the $^7\text{Be}(\alpha, \gamma)^{11}\text{C}(\text{p}, \gamma)^{12}\text{N}(\beta^+, \nu)^{12}\text{C}$ chain. This cycle could bypass the triple- α process in the production of ^{12}C (and CNO cycle elements) through the $^{12}\text{N}(\beta^+, \nu)^{12}\text{C}$ decay. Due to the experimental difficulty of measuring the $^{11}\text{C}(\text{p}, \gamma)^{12}\text{N}$ reaction directly, detailed information on the energies, widths, and spin-parity assignments of these resonances is crucial. Such information is obtained from measuring the $^{11}\text{C}(\text{p}, \text{p})^{11}\text{C}$ system. We therefore have performed a measurement of the $^{11}\text{C}(\text{p}, \text{p})^{11}\text{C}$ system, aiming to study the resonances in ^{12}N above the $^{11}\text{C} + \text{p}$ decay threshold.

1 Background

1.1 Significance to Nuclear Astrophysics

Primordial (Population III) stars are thought to have been very massive and short lived. These stars initially consisted of very light ingredients such as hydrogen and helium that were synthesized into heavier elements such as carbon, nitrogen, and oxygen. In these stars, we see that some of them had an enhancement of these three elements but are poor in heavier ones [1]. As the primordial stars are thought to be more massive than second-generation stars, they will burn via the hot pp-chain and the triple- α process [1]. The α -capture reactions on ^7Be and proton-capture reactions on ^{11}C are identified as being rate determining steps and thus are important in understanding reactions in the primordial stars, but radioactive beams are required for the study of these reactions and data for these are sparse and the isotopes of interest are unstable.

One possible important reaction chain for carbon production is the $^{11}\text{C}(\text{p}, \gamma)^{12}\text{N}(\beta^+, \nu)^{12}\text{C}$ chain [2]. It is possible that the cross section for the $^{11}\text{C}(\text{p}, \gamma)^{12}\text{N}$ reaction may be large as indicated by transfer and knock-out reaction studies. We plan to use TriSol to produce a high-quality ^{11}C beam at low energy to measure the $^{11}\text{C}(\text{p}, \text{p})^{11}\text{C}$ reaction, which will be an important step for constraining the $^{11}\text{C}(\text{p}, \gamma)^{12}\text{N}$ reaction.

1.2 An Overview of TriSol

TriSol is the next-generation iteration of the TwinSol system, a beamline setup originally consisting of two superconducting solenoids. TwinSol has been used for producing, collecting, transporting, focusing, and analyzing both stable and radioactive ion beams for nuclear physics experiments for year. Between 2020 and 2022, TwinSol underwent a major upgrade to become TriSol [4].

The TriSol upgrade introduced a third superconducting solenoid and an additional dipole magnet. This significantly enhanced the system's selectivity and focusing capabilities. These additions improved beam transport and reduced beam spot sizes at the target position, enabling more precise measurements [4]. The incorporation of vertical and horizontal slits further increased the purity of secondary beams by allowing for tighter control over the beam's spatial and momentum distribution. The way in which TriSol manipulates a beam is shown in Fig. 1. When we use TriSol to filter and focus a beam, the particles that make it to the target chamber all have a comparable magnetic rigidity: the particle's resistance to deflection in a magnetic field as determined by its charge and momentum.

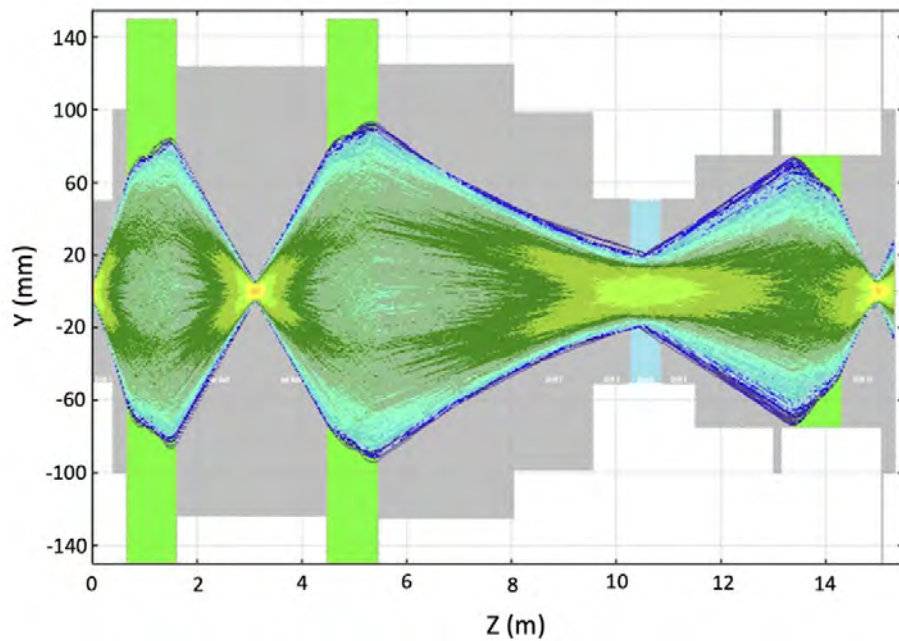


Figure 1: A simulation of a particle beam in the TriSol Line.[5]

1.3 An Overview of Our Experiment

Before going more into the methods, it is worthwhile to give a broad-scale overview of the experiment as well as our goals. We want to perform a measurement of proton scattering in ^{11}C . However, ^{11}C is unstable. It has a half-life of approximately 20 minutes [3]. Therefore, we cannot maintain a sample of it for an extended period. This is an issue since the experiment takes place over several days. Instead, we need TriSol and other NSL facilities to produce and focus a beam of ^{11}C that can be used for studying scattering.

For our measurement we first need to produce a beam of ^{11}C with a high enough intensity. We start with a beam of $^{10}\text{B}^{-1}$, a stable isotope that won't spontaneously decay, that is produced using a sputter source. It is then sent to the FN Tandem Accelerator. Once the beam reaches the terminal inside the FN, it hits a stripper foil that removes electrons from the ^{10}B and leaving it positively charged. We are particularly interested in 5+ charge state since it can be accelerated to a high enough energy for ^{11}C production. After being accelerated and exiting the FN, the $^{10}\text{B}^{5+}$ hits a deuterium gas target where they can undergo a transfer reaction to produce ^{11}C . The beam then enters the TriSol line, where the three superconducting solenoids and dipole magnet are used to filter and focus the beam before sending it into the target chamber. The beam hits a polyethylene target and stops, protons through scattering. The resulting protons are then detected by an array of silicon detectors.

2 Data Collection and Processing

2.1 Correlation of Events

When we produce the beam of ^{11}C , we also unavoidably produce many other things. For example, when the $^{10}\text{B}^{-1}$ beam hits the stripper foil, we are looking to produce the 5+ charge state. However, we also produce some in the 3+ and 4+ charge states. Then, when the beam hits the deuterium gas cell, we produce many different isotopes such as ^6Li , ^7Li , ^4He , and ^9Be while also spreading out the momentum distribution. Furthermore, a portion of the ^{10}B simply passes through

the gas without undergoing any reaction at all. This necessitates a way to separate the events caused by ^{11}C from those caused by everything else. Thankfully, the energy loss of a particle in matter is strongly dependent on its initial energy as well as its charge. We can use this fact to create a plot like in Fig. 2(a) of to get an idea of what we are producing in the beam. From that, we can construct a spectra like Fig. 2(b). It should be noted that at the moment the experiment was only recently completed, so the data is not all calibrated, and the plots in Fig. 2 are pulled directly from the DAQ.

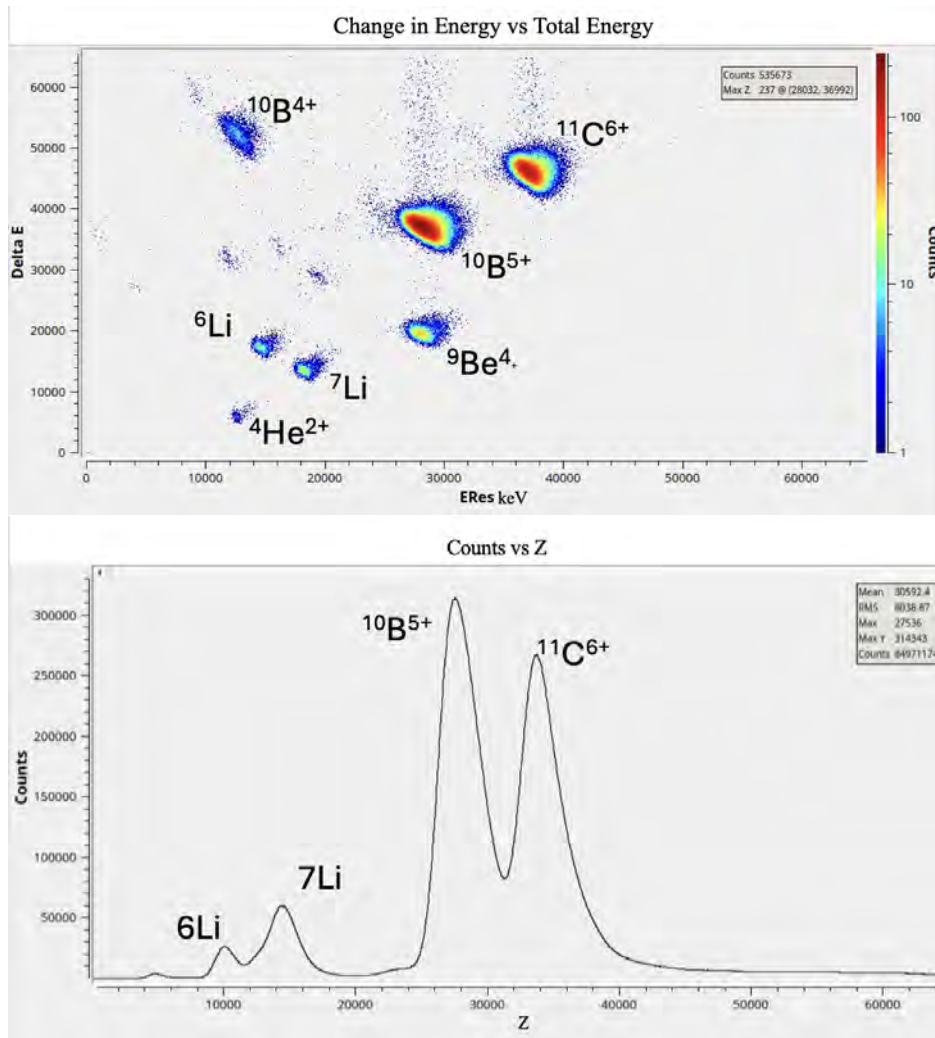


Figure 2: (a) The plot of change in energy vs total energy of the particles, with specific particles identified. (b) Spectra of different nuclei observed.

2.2 Target Chamber and Detector Set Up

The interactions of the beam in the target chamber are measured using two stacked circular arrays of silicon detectors positioned to capture angular distributions, as well as a detector positioned at 0 degrees to the beam axis. The two segments on the bottom did not have detectors due to location being inconvenient to attach ribbon cables. The top layer of detectors is for measuring a change in energy, and the bottom layer is for stopping and measuring the remaining energy of the protons ejected from the polyethylene targets. As such, the top layer is thinner than the bottom layer. It should be mentioned that we did not have enough thick detectors (about $1000\ \mu\text{m}$) for the bottom layer, so for one section we used two $500\ \mu\text{m}$ detectors.

To use the detector mount with the new target chamber, modifications had to be made so it could be properly aligned inside. It was also modified to improve stability and prevent rocking. Initially, we planned to also adapt a target ladder for the experiment. A modified design was developed using CAD software. However, during setup, an existing ladder was discovered that met our requirements, so there was no need for the modification. Pictures of the chamber and detector mount are included in Fig. 3.



Figure 3: (a) Detector mount with Si Detectors (b) Target Chamber (c) Detector mount without detectors

2.3 Thick Target Experiments and Polyethylene Targets

Polyethylene (CH_2) is widely used in nuclear physics experiments, particularly in studies of transfer and knock-out reactions, because of its high hydrogen content and relatively simple chemical composition. The abundance of hydrogen nuclei (protons) in polyethylene makes it an ideal target for inverse kinematics experiments involving proton scattering. Furthermore, since polyethy-

lene contains only carbon and hydrogen, it simplifies the analysis by reducing background signals from unwanted nuclei.

In this experiment, we performed a thick target measurement. The targets are thick enough so that the incident particle beam is completely stopped within the target material. We probe a wide range of beam energies in a single run, and during this process some reactions knock out protons from the polyethylene target. These ejected protons can then be detected and analyzed to determine the energies, widths, and spin-parity assignments of the ^{12}N resonances above the $^{11}\text{C} + \text{p}$ threshold. However, since polyethylene contains both carbon and hydrogen, it is not straightforward to determine which protons originated from which nucleus. To resolve this, we also performed a separate run using a pure carbon target and subtracted the resulting spectrum as a background to isolate the signal from hydrogen-induced reactions.

The polyethylene foils used in the experiment were sourced from Goodfellow, a well-known supplier of materials for scientific use. The target thicknesses were provided by the manufacturer with a precision of approximately half a micron, and in this case, they were approximately $120\ \mu\text{m}$ thick.

3 Conclusion

At this time, the experiment was just recently completed, and the data is pending analysis. When this experiment was run previously, they were unable to obtain sufficient data due to losing data in the 0 degree detector and the background caused by the carbon in the polyethylene was not accounted for since no background run was done. This time, we had a significant improvement in the data and should be able to determine more information about the resonances in ^{12}N above the $^{11}\text{C} + \text{p}$ level. This in turn will allow us to better understand the reactions that powered Population III stars as well as how carbon was produced in the early universe.

References

- [1] P.D. O'Malley, T. Ahn, D.W. Bardayan, M. Brodeur, S. Coil, J.J. Kolata, Nuclear Inst. and Methods in Physics Research, A 1047 (2023) 167784, <https://www.sciencedirect.com/science/article/abs/pii/S0168900222010762>
- [2] P.D. O'Malley, 2025, $^{11}\text{C}(\text{p},\text{p})$ resonant scattering with the TriSol system, [PowerPoint slides]
- [3] National Nuclear Data Center, information extracted from the NuDat 3.0 database, <https://www.nndc.bnl.gov/nudat/>
- [4] P.D. O'Malley, 2022, TriSol – Cleaning up RIBS at the NSL, [PowerPoint slides]
- [5] T. Ahn, D.W. Bardayan, D. Blankstein, C. Boomershine, M. Brodeur, S. Carmichael, S. Coil, J.J. Kolata, P.D. O'Malley, W. Porter, J.S. Randhawa, F. Rivero, J. Rufino, W.W. von Seeger, R. Zite, Nuclear Instruments and Methods in Physics Research B 541 (2023) 216–220, <https://www.sciencedirect.com/science/article/abs/pii/S0168583X23002033>

Methodological Development for Investigating Low-Energy Electron Induced Bond Rupture in DNA

KENTON REEH

2025 NSF/REU Program
Department of Physics and Astronomy
University of Notre Dame

ADVISOR(S): Prof. Sylwia Ptasinska

Abstract

Even though low energy electrons (LEEs) are known to contribute significantly to DNA damage, standardized methodology to quantify this damage remains lacking. In this study, thin films of DNA were prepared with <80% structural integrity through treatment in atmospheric, lyophilization, and ultra high vacuum environments. Data indicates that usage of a beam shutter, a stainless steel sample holder, LEE beam stabilization time, and beam centering techniques allow more precise observation of energy-dependent bond rupture in the DNA backbone. Detailed methodology enhances the understanding of LEE-induced DNA damage in the cellular environment, with implications for uncovering biological effects of high energy radiation used in radiation therapy and diagnostic imaging.

1 Introduction

High-energy radiation common in radiation therapy and diagnostic imaging primarily induces biological damage via the formation of low-energy electron (LEE) cascades between 3-20eV [1]. Because the average secondary electron energy is beneath the ionization threshold of water, the mechanism for biological damage obeys the quantum process of dissociative electron attachment (DEA) [1–3]. Observing DEA in LEE-irradiated DNA molecules requires strict methodology due to damage from DNA's structural fragility, sample preparation, and transportation between atmospheric and vacuum conditions. Therefore, ensuring damage resulted from specifically LEE beam irradiation involves several parameters which may be impossible to observe without risking beam integrity or without proper detector systems and picoammeters.

1.1 Formation of Secondary LEEs from High Energy Radiation

High-energy photons interacting with biological media generate cascades of LEEs through primarily ionization and absorption processes [3]. Initially, high energy photons undergo Compton scattering, which transfers some of the incident photon's energy to an outer shell electron, ejecting it. Both the resultant lower energy photon and ejected Compton electron have a likelihood to ionize media: the Compton electron by colliding with water and the photon by undergoing absorption [3].

If the resultant photon is absorbed by an inner shell electron, Resonant Auger decay and inter-

molecular Coulombic decay (ICD) generate more LEEs. In resonant Auger decay, a core electron is excited to an outer shell by a photon, and the subsequent filling of the core hole releases enough energy to eject a low energy outer shell electron [4]. In ICD, an inner-valence ionized molecule relaxes by transferring its excess energy to a weakly coupled neighboring molecule, causing emission of a LEE [3; 4]. Through ionization and absorption, about 5×10^4 secondary LEEs - the most abundant secondary species made - are formed per MeV of incident radiation [1].

1.2 The DEA mechanism

While LEEs scatter inelastically, there exists a likelihood for them to temporarily localize to a molecule if the electron energy matches that of a vacant antibonding orbital (either σ^* or π^*) [5]. Attachment occurs on a timescale such that nuclei can't reposition themselves immediately, causing them to become vibrationally excited; these are called transient negative ions (TNIs) [3].

Because TNIs exist for only 10^{-15} to 10^{-12} seconds [6], two competing decay pathways occur: if auto-detachment occurs faster than bond elongation leading to cleavage, the electron is ejected and will probabilistically scatter or attach to other molecules. If bond elongation is faster, dissociation occurs [3]. This results in the formation of a stable anionic fragment and neutral co-fragment [3]:

Representative Reaction Pathway



AB indicates a diatomic molecule, # indicates vibrational excitation, and $\bullet$ indicates a free radical.

If molecule AB is DNA, DEA primarily cleaves a strong C-O σ bond in the DNA backbone. Two main processes govern this: firstly, attachment to the lowest-energy π^* orbital in the P=O bond can break a 3' or 5' C-O bond [5; 7]. Alternatively, an electron may attach to a nucleobase via a π^* orbital, forming a metastable state. This results in competition between auto-detachment and dissociation; to dissociate, the electron can adiabatically migrate to the sugar-phosphate 3' C-O σ^* orbital, rupturing it [7].

Broken bonds in the DNA backbone are called strand breaks. Without irradiation, they manifest

as nicked circles (single strand break; SSB) and full-length linear strands (double strand break; DSB) while undamaged DNA is supercoiled (SC) [1]. Of these, DSBs are the most debilitating as only one is sufficient to trigger apoptosis or disturb genomic integrity [8]. DNA also may deform differently when exposed to radiation. SSBs, DSBs, and SC DNA can dimerize [9], and SSBs can form complexes from multiple plasmids [10]. When exposed to radiation and 45°C heat, DNA may denature, forming single-stranded DNA (ssDNA) [11].

2 Methods

Before LEEs are introduced, DNA molecules suspended in pure water must be isolated. This is because DNA binds up to 33 water molecules per nucleotide which form radicals when irradiated with LEEs [12], and these radicals react with DNA within picoseconds [3]. Therefore it would be impossible to prove DEA to DNA resulted in strand breaks.

2.1 Sample Handling Development

DNA (pUC18; 2686bp) suspended in deionized water at 0.5 µg/µL was diluted to a concentration of 0.1 µg/µL. Three stainless steel substrates each received 5.00 µL of DNA solution, with each to be dried differently. One was placed in a lyophilization chamber at -50°C and 4×10^{-4} mbar, the second was placed in a ~ 150 mbar vacuum at room temperature, and the third was left on the counter to air dry. Once each was observably dry, three “wash steps” were performed to incorporate the DNA back into solution. Each one consisted of pipetting 5.00 µL of deionized water onto the substrate and waiting three minutes before collecting the resulting solution without touching the substrate to ensure the most DNA was collected. After analyzing each via gel electrophoresis, all three were over 65% damaged, while the original diluted solution was 6.0% damaged.

The large percentage of damaged DNA indicates the need for testing the preparation technique. It was discovered that dried samples became damaged from wait time in atmospheric conditions, and between 0 and 60 minutes samples became 0.87% more damaged after each minute (n=4, $R^2=0.9984$). Wait time during wash steps incurred more damage: a sample with 0-minute wait

was 31% damaged while waiting three minutes incurred an additional 35% damage. Damage was also induced through pipetting vigorously: spending 3 seconds allowing the solution to exit the pipette tip resulted in no more damage than the control (3.9% damage). Samples more vigorously pipetted incurred an additional 1.3% damage. Lastly, neither centrifuge usage nor gently touching dried DNA with the pipette tip damaged samples further.

Incorporating these results into the water extraction procedure resulted in less damage from each drying technique. DNA dried in the lyophilization chamber was 5.4% damaged, a sample that wasn't dried was 5.8% damaged, the air dried sample was 6.5% damaged, and the sample dried with the room temperature vacuum was 8.1% damaged. Because the lyophilization chamber was most effective at preserving DNA integrity, it became the drying method for future experiments.

When samples finished drying in the lyophilization chamber, they were fastened into a sample holder. The sample holder was inserted into an ultra high vacuum (UHV) chamber, where samples remained for an hour while the vacuum pump reached 4×10^{-2} mbar and then for the UHV pump to reach 10^{-7} mbar. Sample integrity was retained while keeping dried samples under UHV conditions for 16 hours.

Samples were irradiated with a 10eV electron beam (0.3eV energy resolution), 1550V focusing cup, 100V 1st anode, 9V grid, and 15 μ A emission. After one sample received LEE treatment, the grid voltage was increased to ensure no electron emission while the next sample was aligned with the beam. Once alignment was finished, the grid voltage would slowly be decreased to allow emission of more electrons, but the beam control system doesn't provide how many electrons reach the sample as the grid voltage is decreasing. A form of shutter would be required to avoid ambiguity in electron number since the beam can remain on between trials.

Once all samples received treatment, wash steps were done and the concentration of each resultant DNA solution was determined via UV-vis spectrophotometry. Solutions were diluted with different volumes of water so the final concentrations of DNA were equal for all samples. Thermo Scientific 6X DNA loading dye was added at a volume of 20% of the resultant solution, and each tube was centrifuged for ~ 10 s to incorporate the dye.

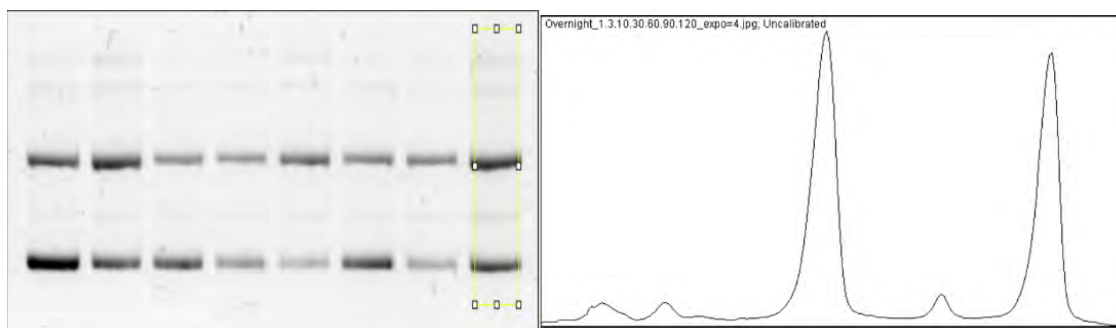


Figure 1: ImageJ averages the greyscale pixel intensity per row of pixels in a highlighted region. Therefore the highlighted lane from top to bottom produces the peaks from left to right on imageJ.

Dyed, irradiated samples underwent agarose gel electrophoresis for analysis. The agarose gel was dyed with SYBR Safe (10,000X in DMSO) and suspended in a 10X dilution of TBE buffer. 10.00 μ L of each centrifuged solution was slowly transferred onto the gel. Electrophoresis was ran for 55 minutes at 400mA and 80V. The gel is then removed and placed in a MICROTEK Bio-1000F gel scanner, and images were analyzed using imageJ software. ImageJ allows the user to highlight each lane of a gel, where it assigns an intensity level from the picture's grayscale for each row of pixels inside a highlighted rectangle (figure 1). This also removes noise because a speck of dust averaged with a row of white pixels becomes negligible. The relative intensity of each band is calculated to determine the percentage of strand breaks.

Each band on the gel is formed because electrophoresis takes advantage of the negative charge from DNA's phosphate groups, allowing different conformations of DNA to traverse a gel matrix at different rates when a voltage is applied. The bottom 3 bands in figure 1, from top to bottom, correspond to SSB, DSB, and SC

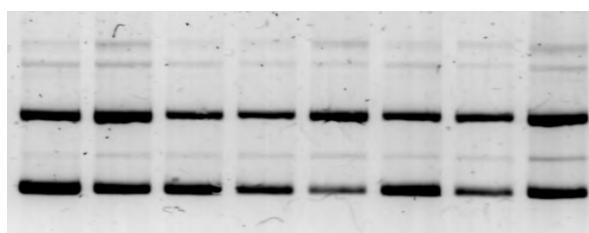


Figure 2: A sixth band becomes visible but not analyzable with imageJ when exposure time in the gel scanner is increased.

DNA. This is because SC DNA moves fastest through the gel as it's more compact than the linear DSBs. SSBs are circular (the least compact of the three): they get impaled by dangling fibers in the gel matrix, making them move slowest [13]. The top two bands are unknown. SSB dimers (a circle of DNA twice as large) [9], SSB complex forms[10], and denatured ssDNA [11] may all

move slower through the gel than SSBs. In fact, the same gel from figure 1 has a sixth band that becomes visible with greater exposure time in the gel scanner (see figure 2), although imageJ can't analyze it because the grayscale values become too large.

3 Results and Discussion

To observe DEA, DNA is irradiated with different electron energies. For each trial, a graph showing the trends of DNA damage per incident electron is formed and the slope of the first four data points that form the linear region of a larger decay curve indicate electron induced DNA damage effectiveness for that energy level [1]. This linear region is used to indicate strand breaks per incident electron because linear proportionality implies the interaction between one electron and one DNA molecule. As a result, to demonstrate the existence of DEA this effectiveness must be greater at specific energy levels, as this implies the energy overlap between the electron and anti-bonding orbital of DNA.

Figure 3 attempts to find the linear region for a 10eV electron beam. It demonstrates samples treated with 0 incident electrons being <80% undamaged, indicating sample preparation technique that makes distinguishing between sample preparation induced damage and radiation induced damage possible. Although there is a trend of increasing damage with increasing incident electrons, too much error is inherent to depict a linear region of a larger decay curve. The linear region might be evaluated if more samples were irradiated with fewer elec-

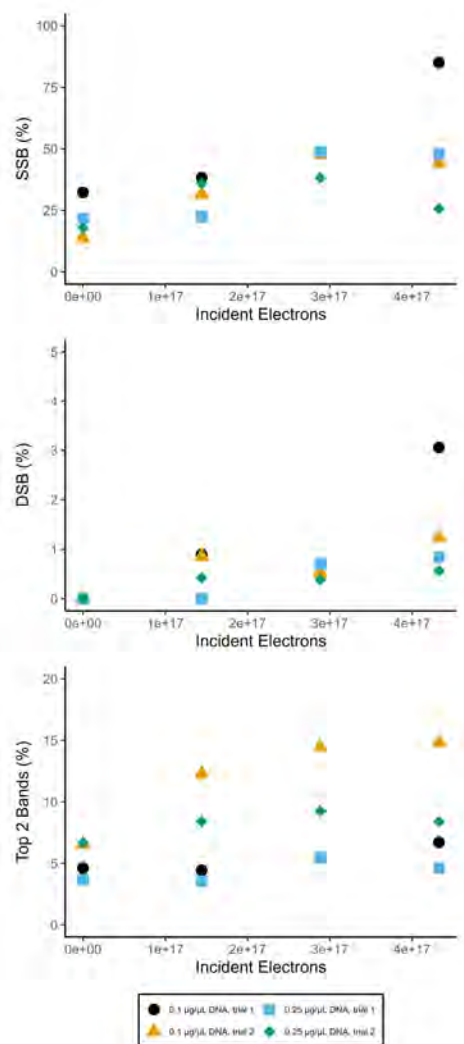


Figure 3: DNA damage from 10eV LEE irradiation. Two trials were performed at each of two DNA concentrations. The 0.1 µg/µL trial 1 was given a larger dose of incident electrons.

trons, as more samples may yield a smoother curve and fewer electrons may increase the likelihood of an interaction happening between one electron and one DNA molecule.

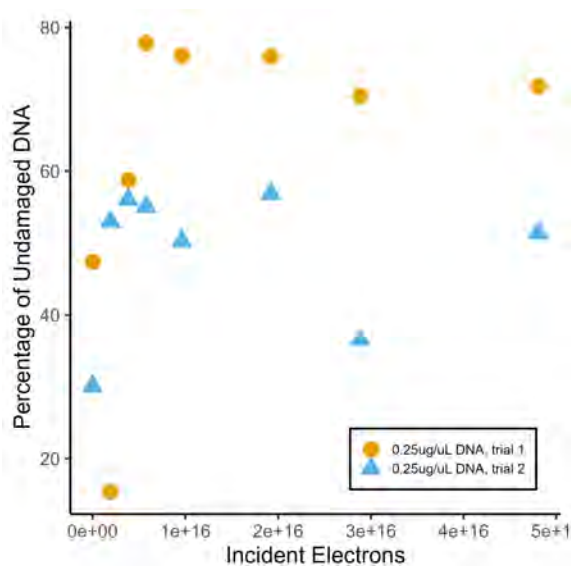


Figure 4: DNA damage from 10eV LEE irradiation with lower current. In both trials, the control sample (0 incident electrons) is about twice as damaged as the least damaged sample.

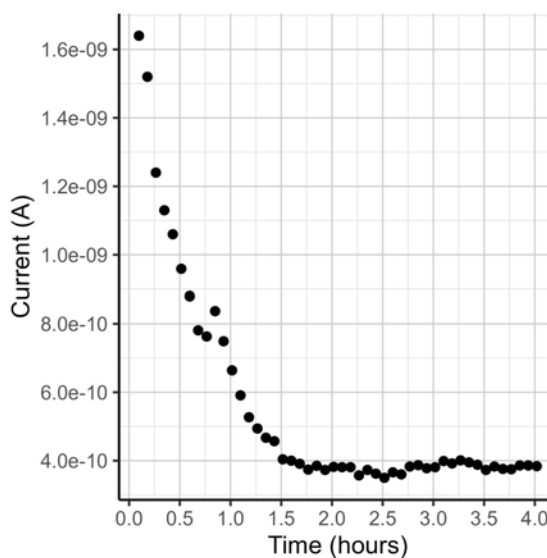


Figure 5: LEE beam current recorded by the Faraday cup at the end of the beam path.

In both of the trials in figure 4, the number of incident electrons was decreased by an order of magnitude and both control samples were more damaged than the samples irradiated with $\sim 5 \times 10^{15}$ incident electrons. Because irradiated samples were less damaged than the control but the bands resolved on the gels, the other data points demonstrate good sample preparation, and this trend is apparent in both trials, there are likely further experimental design problems.

By observing that samples treated with less than $\sim 5 \times 10^{15}$ electrons show a decrease in DNA damage with increasing incident electrons, the beam may need time to stabilize after it's turned on. As a result, a Keithley 6485 picoammeter was wired to a computer and the Faraday cup, which collects current at the end of the electron beam path. A script using Python and SCPI was made to automate current measurements, and figure 5 found the beam requires about two hours to stabilize.

The other trends in figure 4 that indicate experimental design problems are the variability between samples (also observed in figure 3) and control sample damage. To account for these, the beam may be bending or misaligned with DNA samples.

Beam bending may occur as the sample holder used was made of ceramic with a metal baseplate, meaning neither the baseplate nor the stainless steel sample holders are grounded. This also means the ceramic surrounding the samples becomes charged if the beam misses or while electrons are emitted during sample alignment with the beam. A stainless steel sample holder was modeled using Fusion360 (see figure 6) to ensure the substrates and baseplate could be grounded easily, allowing a greater uniformity of the electron beam.



Figure 6: Stainless steel sample holder model. The center is dyed white where ceramic is in the current design.

The beam is likely misaligned with the sample holder because the beam diameter when it reaches the samples isn't measurable without proper detector systems. The lowest energy beam measurable without a detector was 500eV, using carbon tape; it was <1mm in diameter. Alternatively, the DNA film is within a 4mm diameter of the center of the substrate. Not knowing the 10eV beam size means risking charging the ceramic sample holder if the beam

is too wide and risking missing part or all of the DNA sample if the beam is too narrow. Even if the beam size was large enough to irradiate the whole sample without charging the sample holder, the beam must be widened to account for error in the beam alignment process.

4 Conclusion

Developing methodology that demonstrates DEA for the fragile DNA molecule continues to be the most significant pitfall to understanding which LEE energies cause biological damage. Because these methods demonstrate the capability for <80% of DNA to remain undamaged throughout the irradiation procedure, radiation damage can be appropriately quantified but questions remain regarding centering the beam on the samples, employing a shutter, and grounding the sample holder to minimize sample variability. Answering these may provide a more nuanced picture of DEA

mechanism differences for solutions more closely mirroring the cellular environment, which is necessary to improve radiation therapy and diagnostic imaging.

References

- [1] B. Boudaïffa, P. Cloutier, D. Hunting, M. A. Huels, and L. Sanche, *Science* **287**, 1658 (2000).
- [2] E. Alizadeh, T. M. Orlando, and L. Sanche, *International Journal of Mass Spectrometry* **377**, 146 (2015).
- [3] S. Ptasinska, *Physical Chemistry Chemical Physics* **15**, 14636 (2013).
- [4] F. Trinter, M. S. Schöffler, H.-K. Kim, F. P. Sturm, K. Cole, N. Neumann, A. Vredenburg, J. Williams, I. Bocharova, R. Guillemin, M. Simon, A. Belkacem, A. L. Landers, T. Weber, H. Schmidt-Böcking, R. Dörner, and T. Jahnke, *Nature* **505**, 664 (2014).
- [5] L. Sanche, *European Physical Journal D* **35**, 367 (2005).
- [6] A. D. Bass and L. Sanche, *Radiation Physics and Chemistry* **68**, 3 (2003).
- [7] J. Simons, *Accounts of Chemical Research* **39**, 772 (2006).
- [8] T. Bohgaki, M. Bohgaki, and R. Hakem, *Genome Integrity* **1**, 15 (2010).
- [9] M. Schleef, R. Baier, W. Walther, M.-L. Michel, and M. Schmeer, *BioProcess International* , 48 (2006).
- [10] C. Herskind, *Int J Radiat Biol Relat Stud Phys Chem Med.* **52**, 565 (1987).
- [11] A. Sebastian, D. Lipa, and S. Ptasinska, *ACS Omega* **8**, 3984 (2023).
- [12] Y. Gao, X. Wang, P. Cloutier, Y. Zheng, and L. Sanche, *Molecules* **29**, 6033 (2024), accessed: 2025-07-16.
- [13] N. C. Stellwagen and E. Stellwagen, *Journal of Chromatography A* **1216**, 1917 (2009).

Explodability of Core-Collapse Supernovae with a Transition to Quark Matter

OLIVER RUPP

2025 NSF/REU Program
Department of Physics and Astronomy
University of Notre Dame

ADVISOR(S): Prof. Grant Mathews

Abstract

The phase transition to quark matter as a shock revival method for core-collapse supernovae offers a new mechanism to produce successful supernova explosions. Past studies have explored this transition with various EOSs and codes, but few have examined it with turbulence from neutrino-heated convection, which adds pressure support and aids explosions. Thus, we study the effect of the hadron-quark phase transition (HQPT) on supernova explodability in the context of Supernova Turbulence in Reduced-Dimensionality (STIR) using the general relativistic hydrodynamics code GR1D. We use two hybrid hadron-quark equations of state (EOSs) and a single solar metallicity progenitor with zero-age main sequence mass $25 M_{\odot}$ (s25). We find that some models with turbulence explode before reaching the phase transition, while runs with and without turbulence exhibit signs of a phase transition but crash soon after due to instabilities in neutrino transport computations. Future work should address the neutrino treatment in the quark matter regime to fully explore the dynamics of the phase transition.

1 Introduction

Core collapse supernovae (CCSNe) are some of the most energetic events in the universe, releasing $\sim 10^{51}$ ergs, but their explosion mechanism is still uncertain. Current models of supernovae often predict failed explosions, and require a careful treatment of all the physics involved. During core collapse, a hydrodynamical shock wave forms due to infalling matter rebounding off the dense iron core. The shock wave then stalls, losing energy as it travels through the outer shells of the star. In order to revive the stalled shock wave and produce a successful explosion, a shock revival method must be considered. The standard neutrino heating revival method [1] generates a shock revival from neutrino interactions, which provide the stalled shock the energy needed to explode the star. We, however, study a newer shock revival method that offers a promising alternative to the standard neutrino heating mechanism: the hadron-quark phase transition (HQPT).

Shock revival may be driven by the phase transition from hadronic to quark matter [2]. During core collapse after the first bounce, falling mass accretes onto the proto-neutron star (PNS) causing the density of the inner core to rise. When the central density reaches values of $2.5\rho_{sat}$ ($\sim 6 \times 10^{14} \text{ g cm}^{-3}$), where ρ_{sat} is nuclear saturation density, quarks become deconfined from the strong force and matter is composed of a mix of free quarks and nucleons [3]. In this mixed phase, the

polytropic index decreases and causes a sudden drop in pressure, which can cause the PNS to collapse for a second time. During this second collapse, densities sharply rise to $\sim 1 \times 10^{15} \text{ g cm}^{-3}$, which is sufficiently high for pure quark matter formation. Once pure quark matter has been reached, the polytropic index sharply increases. This leads to an increase in pressure and a second bounce of the infalling matter from the second collapse [4]. The second bounce and resulting shockwave can then interact with the first stalled shock wave and provide it the energy needed to break through the outer layers of the star for a successful explosion.

In this work, we extend the work of previous studies by considering the transition to quark matter in CCSN driven by neutrino-heated convection in the context of relativistic supernova turbulence in reduced dimensionality (STIR) [5].

2 Background

2.1 GR1D

We use GR1D, a one-dimensional general relativistic spherically symmetric hydrodynamical code for stellar collapse [6]. GR1D maps out a computational grid and solves the hydrodynamic equations of the conserved variables, evolving them with Runge-Kutta integrators (see Eqs. 7-10 of [6]). The center cell values are reconstructed at the interfaces between cells and flux is calculated.

2.2 Neutrino Interactions

Multiple schemes have been used for modeling neutrino interactions in stellar evolution codes such as GR1D. The simplest of these is the leakage scheme, which approximates neutrino emission rates and energies using an estimate of the local optical depth. A more rigorous approach for modeling neutrino transport is to approximate the relativistic neutrino distribution function [7] as a series of moments. We use one such scheme, GR1D's M1 moment scheme for neutrino transport [8]. The zeroth and first moments (energy and momentum density) are updated in each time step from the flux and source terms. The system is then closed by approximating the next highest moment as a

function of the first two.

In order to provide the source terms for neutrino interactions, the absorption opacities, scattering opacities, and emissivities are calculated using Nulib [8], which reads in an EOS and calculates the neutrino rates from the particle mass fractions and chemical potentials. In CCSN simulations, three species are typically considered: electron neutrinos (ν_e), antielectron neutrinos ($\bar{\nu}_e$), and an average of the muon and tau neutrinos (ν_x).

2.3 Hadron-Quark Hybrid EOS

To close the system of hydrodynamical equations, an equation of state (EOS) is needed to provide energy, pressure, and other thermodynamic variables as a function of density, temperature, and electron fraction. A typical hadron-quark hybrid EOS is constructed with a first order phase transition from hadronic to quark matter with an intermediate mixed phase containing separate domains of quark and hadronic matter.

2.3.1 DD2F-SF EOS

We use the hadron-quark hybrid DD2F-SF EOS of [9]. The EOS is calculated for the hadronic phase with the relativistic mean field (RMF) theory of [10], using the DD2 parameterization of [11] and a flow constraint. RMF models nucleon interactions with mean-field potentials instead of direct two-body interactions. Nuclear clusters are modeled in the nuclear statistical equilibrium (NSE) of [10], which assumes nuclear reactions occur rapidly enough to maintain equilibrium among species.

The quark phase is calculated with the string-flip model introduced in [12], which considers only up and down quarks. The phase transition between the two models uses the Gibbs conditions for phase equilibrium, which require that the two regions be in chemical, thermal, and mechanical equilibrium during the transition.

2.3.2 STOS-B145 EOS

For the hadronic phase, the STOS-B145 is constructed of the RMF EOS of [13]. The quark phase is calculated with the MIT bag model of [14]. In the bag model of quark matter, a gas of up, down, and strange quarks is assumed to be confined in a "bag" where they can move around freely. The model is defined by the bag constant B : the positive potential per unit volume. We choose a bag model with $B = 145$ MeV. This value is in the midrange of the available EOSs and has been used in previous studies ([4], [15]), allowing for a direct comparison of our results.

3 Methods

We used GR1D for all simulations, its M1 transport to model neutrino propagation, and Nulib to generate the neutrino opacity tables. We use a spatial grid of 700 cells, with the inner 20 km at a constant spacing of 0.2 km, and then logarithmic spacing out to an outer boundary where the density is at $2 \times 10^3 \text{ g cm}^{-3}$. We use a Courant factor of 0.5, which is reduced at both the first and second bounce to 0.2. We use a single solar metallicity progenitor with zero-age main sequence mass $25 M_{\odot}$ (s25)

We use two hadron-quark hybrid EOSs: DD2F-SF and B145. Both of the tables are from CompOSE [16] and converted to a format compatible with GR1D. In [9], a number of different models of the DD2F-SF EOS were provided, but we chose to use DD2F-SF 1.7 as it has one of the smallest onset densities for quark matter formation and thus reduces simulation time. We will refer to DD2F-SF 1.7 as simply DD2F-SF for the remainder of the report. We will also refer to the STOS-B145 EOS as simply B145.

We use two other purely hadronic EOSs in this study as a benchmark: the LS220 [17] and DD2 ([11], [10]) EOSs. The LS220 EOS uses a Skyrme model for nuclear interactions and the single-nucleus approximation (SNA), a compressible liquid-drop model. They are coupled to a network of 3335 nuclides in NSE at low densities. The DD2 EOS is a similar model as the hadronic part of the DD2F-SF EOS. The LS220 EOS was generated by the Schneider-Roberts-Ott (SRO) code of

[18], while the DD2 EOS is available from [19].

For neutrinos, we use 18 logarithmically spaced energy groups, centered from 1 MeV to ~ 300 MeV. We use the neutrino opacities in [20], and also include the weak magnetism and recoil corrections from [21], as well as a virial correction to neutrino scattering from [22]. We also include inelastic electron-neutrino scattering.

Our main addition to the study of the effect of phase transition to quark matter on supernovae is studying it in the context of neutrino heating turbulence using the STIR of [23], which was first studied with GR1D in [5]. Although the effects of turbulent fluid flow are inherently multi-dimensional, an approximation to the additional energy and pressure generated by turbulence is given by STIR, which adds a term to GR1D’s hydrodynamical equations to evolve the turbulent energy. The STIR model is characterized by the mixing length parameter α_{MLT} which determines the characteristic length a fluid element travels before mixing with the surrounding medium. Following [5], we adopt $\alpha_{\text{MLT}} = 1.48$ for all runs with turbulence.

4 Results

For the majority of simulation time, densities are below saturation density, and thus the dynamics are not affected by the HQPT. When temperatures and densities are high enough for a phase transition to pure quark matter, GR1D crashes. Nearly all of these crashes are because the reconstructed neutrino flux is unable to converge to the expected value within the allowed tolerance. This is likely due to the neutrino opacity table or the EOS table, where the sudden change in opacities or thermodynamic values in the quark matter regime is either poorly constructed or too discontinuous for GR1D to handle without modifications. We do, however, find results from just before the crash and from runs without a phase transition.

In Fig. 1(a, b), we show the shock radius of the s25 progenitor for different EOSs with and without turbulence. All models run without turbulence do not explode, although future work should confirm if the B145 and DD2F-SF EOSs crash due to the HQPT or black hole formation. For example, the shock radius B145 run with no turbulence sharply drops off, implying black hole

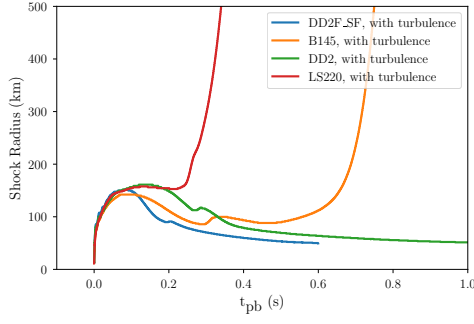
formation. However, the central lapse function values decrease to only ~ 0.55 before increasing again, where runs with black hole formation typically reach central lapse function values of ~ 0.2 before crashing. This suggests that although black hole formation may eventually occur, it is not the cause of the run crashing.

The inclusion of turbulence leads to successful explosions for the LS220 and B145 EOSs. The successful run of the B145 EOS is likely because the star explodes before a phase transition to quark matter occurs. This means the energy from turbulence and neutrino heating is sufficient to revive the shock wave and explode the star before central densities are high enough to form quark matter and crash the run. The shock radius of the s25 progenitor run with the LS220 EOS and turbulence is in good agreement with [5].

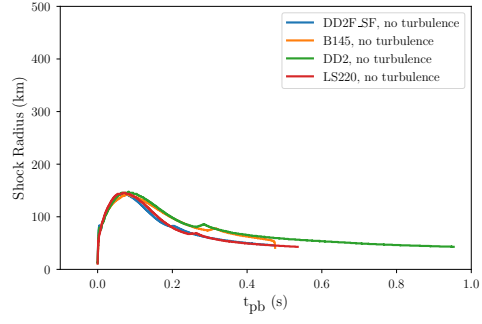
Recent studies ([24], [25]) have shown that the second collapse induced by the HQPT is indicated by a spike in central density as the PNS loses pressure support immediately before pure quark matter is formed. Central density then levels out when pure quark matter is reached as the PNS stabilizes and a second shock is generated. Central temperatures typically increase slightly from the compression of the PNS and then drop from the latent heat of the HQPT.

In Fig. 1(c, d), we show the central temperature and density of the B145 EOS run with and without turbulence. The run without turbulence qualitatively agrees with previous studies ([24], [25]), although we are unable to assess how far the temperature drops. The run with turbulence is more interesting. Although the central temperature drops as it would with a phase transition, the central density levels out instead of spiking. This suggests that the inner core reached the mixed phase, but the additional pressure support from turbulence was sufficient to prevent a second collapse and the formation of pure quark matter.

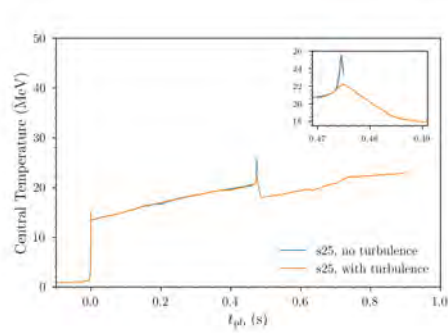
In Fig. 1(e), we show the shock evolution of the B145 EOS run without turbulence. The large dip is the first stalled shock wave. The velocities in the inner core show another collapse and then a new outgoing shock wave which moves towards the stalled shock right before the run crashes. This suggests that the interaction between the two shock waves is delicate and must be treated carefully to ensure numerical instabilities do not crash the code.



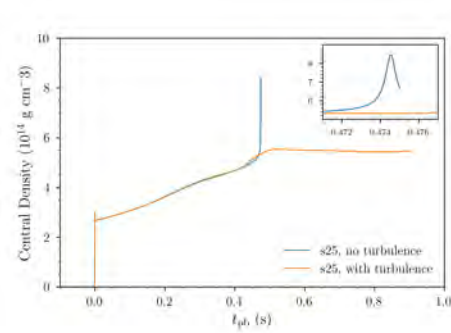
(a) Shock radius with turbulence



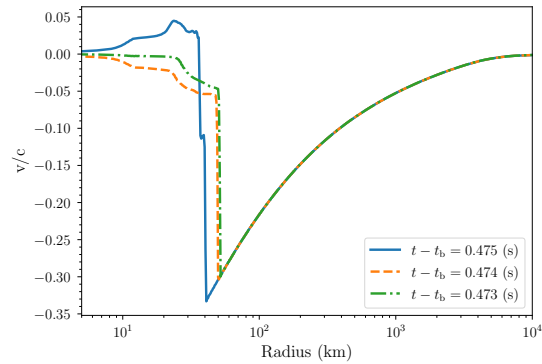
(b) Shock radius without turbulence



(c) Central temperature



(d) Central density



(e) Velocity vs. radius at various post-bounce times showing shock evolution without turbulence for the B145 EOS.

Figure 1: GR1D simulation results, comparing effects of turbulence on shock evolution and central thermodynamic variables. Times are shown relative to bounce, i.e., time post bounce t_{pb} .

5 Conclusion

We have investigated the effect of a hadron-quark phase transition (HQPT) on the explodability of core-collapse supernovae using the GR1D code with two hybrid hadron-quark EOSs: DD2F-SF

and B145. Our simulations include the impact of turbulence through the STIR model to assess its role in aiding explosions. Our results demonstrate that the inclusion of turbulence increases the likelihood of successful explosions. In particular, the B145 EOS with turbulence produces an explosion before pure quark matter formation, while the same EOS without turbulence undergoes a second collapse, consistent with a HQPT, before crashing. While some models exhibit behavior consistent with the HQPT triggering a second shock wave, none of our hadron-quark simulations successfully evolve past the pure quark matter regime due to numerical instabilities in neutrino transport.

Future work should improve the stability of the code at the phase transition and interaction between the two shock waves, addressing discontinuities or unphysical assumptions in the neutrino opacity table and EOS in the quark matter region. Additionally, a range of progenitors should be tested to more thoroughly study the interaction of turbulence and the HQPT.

References

- [1] H.-T. Janka, *Annual Review of Nuclear and Particle Science* **62**, 407 (2012).
- [2] I. Sagert, T. Fischer, M. Hempel, G. Pagliara, J. Schaffner-Bielich, A. Mezzacappa, F.-K. Thielemann, and M. Liebendörfer, *Phys. Rev. Lett.* **102**, 081101 (2009).
- [3] S. Zha and E. O’Connor, *Physical Review D* **106** (2022).
- [4] P. Jakobus, B. Müller, A. Heger, A. Motornenko, J. Steinheimer, and H. Stoecker, *MNRAS* **516**, 2554 (2022).
- [5] L. Boccioli, G. J. Mathews, I.-S. Suh, and E. P. O’Connor, *The Astrophysical Journal* **926**, 147 (2022).
- [6] E. O’Connor and C. D. Ott, *Classical and Quantum Gravity* **27**, 114103 (2010).
- [7] R. W. Lindquist, *Annals of Physics* **37**, 487 (1966).

- [8] E. O'Connor, The Astrophysical Journal Supplement Series **219**, 24 (2015).
- [9] N.-U. F. Bastian, Phys. Rev. D **103**, 023001 (2021).
- [10] M. Hempel and J. Schaffner-Bielich, Nuclear Physics A **837**, 210 (2010).
- [11] S. Typel, G. Röpke, T. Klähn, D. Blaschke, and H. H. Wolter, Phys. Rev. C **81**, 015803 (2010).
- [12] M. A. R. Kaltenborn, N.-U. F. Bastian, and D. B. Blaschke, Phys. Rev. D **96**, 056024 (2017).
- [13] H. Shen, H. Toki, K. Oyamatsu, and K. Sumiyoshi, PTP **100**, 1013 (1998).
- [14] E. Farhi and R. L. Jaffe, Phys. Rev. D **30**, 2379 (1984).
- [15] S. Zha, E. P. O'Connor, and A. da Silva Schneider, The Astrophysical Journal **911**, 74 (2021).
- [16] CompOSE Collaboration, CompOSE: CompStar Online Supernovae Equations of State.
- [17] J. M. Lattimer and F. Douglas Swesty, Nuclear Physics A **535**, 331 (1991).
- [18] A. S. Schneider, L. F. Roberts, and C. D. Ott, Phys. Rev. C **96**, 065802 (2017).
- [19] E. P. O'Connor, <https://ttt.astro.su.se/eoco/eos.html>.
- [20] S. W. Bruenn, The Astrophysical Journal **58**, 771 (1985).
- [21] C. J. Horowitz, Phys. Rev. D **65**, 043001 (2002), arXiv:astro-ph/0109209 [astro-ph] .
- [22] C. J. Horowitz, O. L. Caballero, Z. Lin, E. O'Connor, and A. Schwenk, Phys. Rev. C **95**, 025801 (2017), arXiv:1611.05140 [nucl-th] .
- [23] S. M. Couch, M. L. Warren, and E. P. O'Connor, The Astrophysical Journal **890**, 127 (2020).
- [24] N. Khosravi Largani, T. Fischer, and N.-U. F. Bastian, The Astrophysical Journal **964**, 143 (2024).
- [25] T. Fischer, The European Physical Journal A **57**, <https://doi.org/10.1140/epja/s10050-021-00571-z> (2021).

Using $^{23}\text{Na} + p$ Reaction Analysis to Characterize $^{12}\text{C} + ^{12}\text{C}$

JORDAN SHEGOG

2025 NSF/REU Program
Department of Physics and Astronomy
University of Notre Dame

ADVISOR(S): Dr. Ani Aprahamian, Dr. Shelly Leshner, Dr. Wanpang Tan

Abstract

The $^{12}\text{C} + ^{12}\text{C}$ fusion reaction is a fundamental process in nuclear astrophysics, crucial for understanding the late stages of burning in massive stars. However, direct measurement of this reaction at low energies remains challenging due to its complexity and the highly suppressed cross-section resulting from the Coulomb barrier. To overcome these challenges, we use an alternative approach using the inverse reaction of $^{23}\text{Na} + \text{p}$, which shares some resonances in their compound nucleus: $^{24}\text{Mg}^*$. The application of R-Matrix formalism is essential for accurately modeling narrow, isolated resonances present in low-energy fusion reactions. The results in this paper reveal insights into the levels of the $^{24}\text{Mg}^*$ nucleus, which are directly linked to the $^{12}\text{C} + ^{12}\text{C}$ fusion reaction through the analysis of the higher excitation proton channels at run energies 4.1-4.74 MeV. This characterization contributes to our understanding of nuclear astrophysics by offering a more comprehensive and comprehensible analysis of this fundamental process.

1 Introduction

In nuclear astrophysics, the $^{12}\text{C} + ^{12}\text{C}$ fusion reaction is a major focus of study. This reaction takes place in the cores of massive stars, where lighter elements undergo fusion to form heavier nuclei. Carbon fusion plays a critical role in the late-stage evolution of stars and has direct implications for stellar nucleosynthesis and the onset of supernovae[1].

Despite its importance, detailed information about the reaction mechanisms—particularly the reaction rate—remains limited. This is largely due to the extremely low cross-section associated with the process at astrophysical energies. The cross-section, which represents the probability of an interaction between two particles, becomes vanishingly small at low energies due to the Coulomb barrier, making direct measurement in laboratory settings exceedingly difficult.

To avoid these challenges, the reaction was studied through the $^{23}\text{Na} + \text{p}$ channel, which serves as one of the decay pathways of the intermediate compound nucleus, $^{24}\text{Mg}^*$. This exit channel exhibits a comparatively higher cross-section, making it more experimentally accessible. By examining this channel using inverse kinematics, we can extract information about the resonant states of $^{24}\text{Mg}^*$ and, in turn, gain insights into the original $^{12}\text{C} + ^{12}\text{C}$ fusion process.

2 Background

2.1 Reaction Mechanisms

When a $^{12}\text{C} + ^{12}\text{C}$ fusion reaction occurs, the two carbon nuclei combine to form an excited magnesium nucleus, $^{24}\text{Mg}^*$. This compound nucleus is highly unstable and promptly decays into either the $^{23}\text{Na} + \text{p}$ or $^{20}\text{Ne} + \alpha$ exit channels, as illustrated in Figure 1.

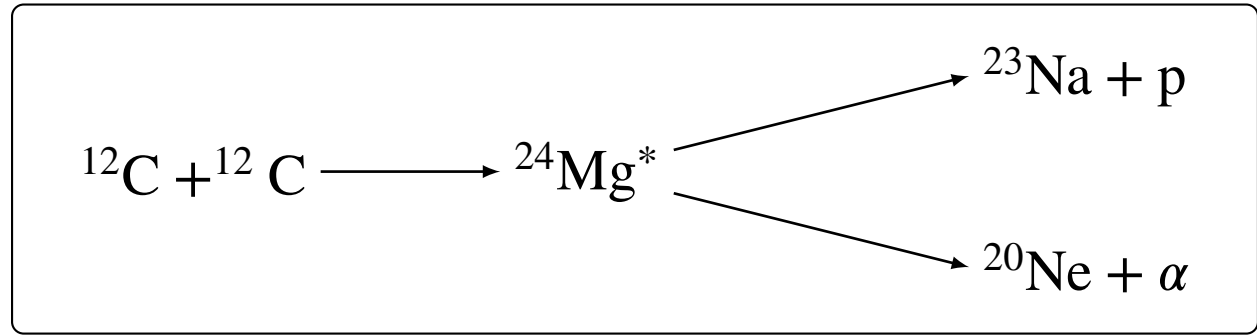


Figure 1: Reaction Diagram of $^{12}\text{C} + ^{12}\text{C}$ forming $^{24}\text{Mg}^*$, with decay into either $^{23}\text{Na} + \text{p}$ or $^{20}\text{Ne} + \alpha$.

By treating the $^{23}\text{Na} + \text{p}$ exit channel as an entrance channel in a time-reversed reaction, we obtain a different but related diagram, shown in Figure 2. The resonant structures observed in the $^{23}\text{Na} + \text{p}$ system provide valuable information about the properties of the compound nucleus. Since both reactions share the same intermediate state, $^{24}\text{Mg}^*$, a detailed understanding of this nucleus offers insight into the original $^{12}\text{C} + ^{12}\text{C}$ fusion process.

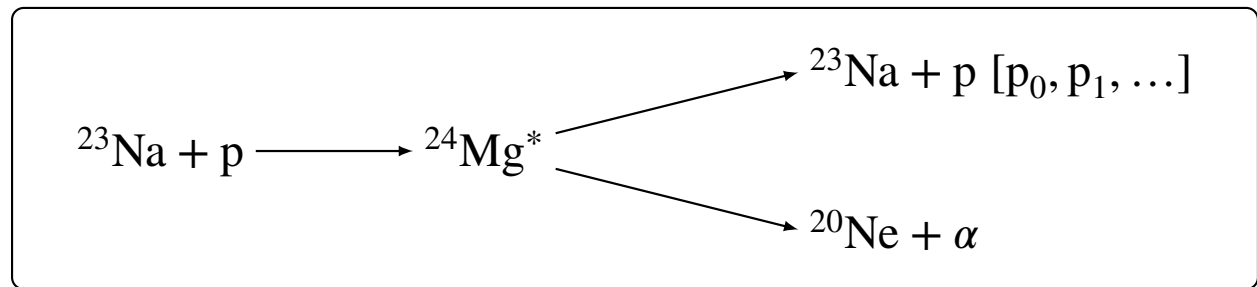


Figure 2: Reaction diagram using $^{23}\text{Na} + \text{p}$ as the input channel.

The experiment conducted at the University of Notre Dame is aimed at improving our understanding of the excited state $^{24}\text{Mg}^*$ and the underlying dynamics of the $^{12}\text{C} + ^{12}\text{C}$ fusion reaction.

This is achieved by analyzing the reaction mechanisms and cross-sections observed in the $^{23}\text{Na} + \text{p}$ channel.

2.2 R-Matrix Theory

In nuclear physics, especially in low-energy reactions relevant to astrophysics, the behavior of particles is often dominated by resonances. These are short-lived excited states of the compound nucleus formed during a reaction. Because these states affect the cross-section it is important to model them accurately. The R-Matrix formalism provides a way to do this by separating the reaction into two regions: an internal region where the interacting nuclei are close enough for the nuclear force to dominate, and an external region where the particles are far apart and primarily influenced by long-range forces like the Coulomb interaction.

The boundary between these regions is defined by a channel radius a , typically chosen large enough to contain all nuclear interactions[2]. Inside this boundary, the wavefunction of the system is expanded in terms of basis functions $\phi_\lambda(r)$, which are eigenfunctions of the internal Hamiltonian. Each state λ has an associated energy E_λ and reduced width amplitude $\gamma_{\lambda c}$, which describes the coupling of the internal state to channel c . The R-matrix is then defined as:

$$R_{cc'}(E) = \sum_{\lambda} \frac{\gamma_{\lambda c} \gamma_{\lambda c'}}{E_{\lambda} - E}$$

where $R_{cc'}(E)$ represents the interaction between channels c and c' , and E is the center-of-mass energy. This matrix acts as a boundary condition that connects the internal solution to the external scattering wavefunctions. From it, physical quantities like the scattering matrix S , phase shifts δ , and cross-sections σ can be calculated. This formalism aids in the analysis of reactions with narrow, isolated resonances, like those seen in low-energy fusion or proton capture reactions present in this research.

3 Methods

The experiment, referred to as R2D2, was carried out over three days at the University of Notre Dame using the FN tandem Van de Graaff accelerator. The experimental configuration is shown in Figure 3, featuring a central $\text{Na}_2\text{WO}_4/\text{C}$ target, with silicon photo-diode detectors placed at a range of forward and backward scattering angles to capture emitted particles.

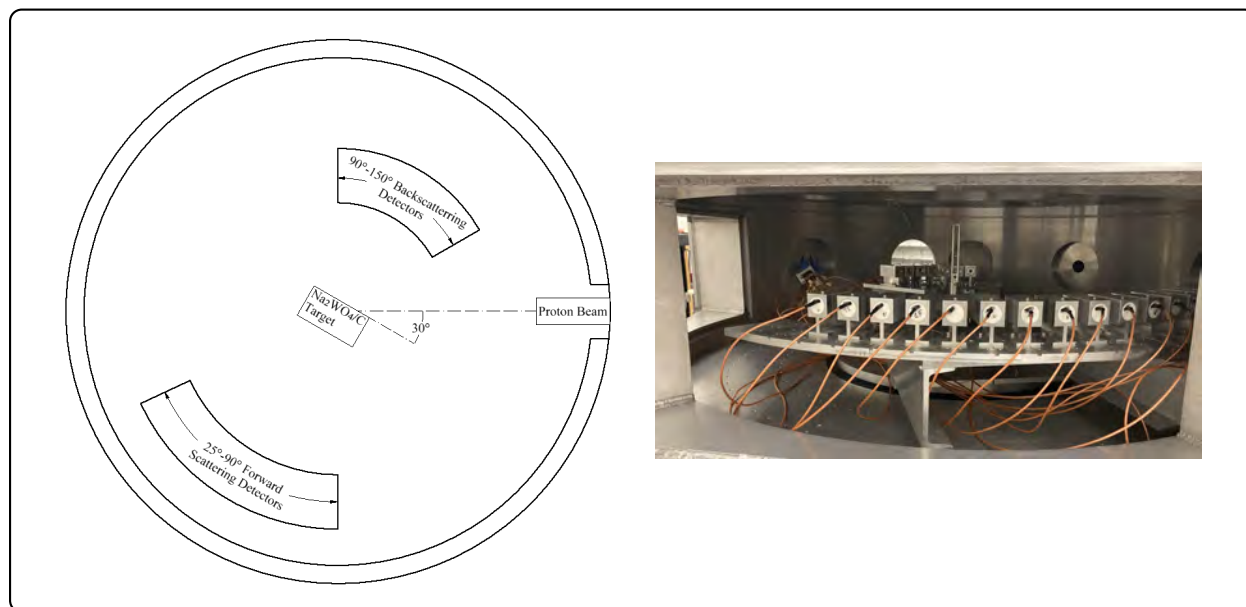


Figure 3: Schematic of Detector and Target

The target was bombarded with protons, inducing the $^{23}\text{Na} + \text{p}$ reaction to occur, and the products of those reactions were then scattered and noted in angle and energy by the detectors. Figure 4 is an example of the data that is acquired from an accelerator run of this experiment. It shows the number of particles detected of a certain energy at a given detector angle. A total of 206 runs were conducted, ranging from 4.1-6.1 MeV. With the energy profiles of the different scattered particles, we can focus on specific exit channels for further analysis.

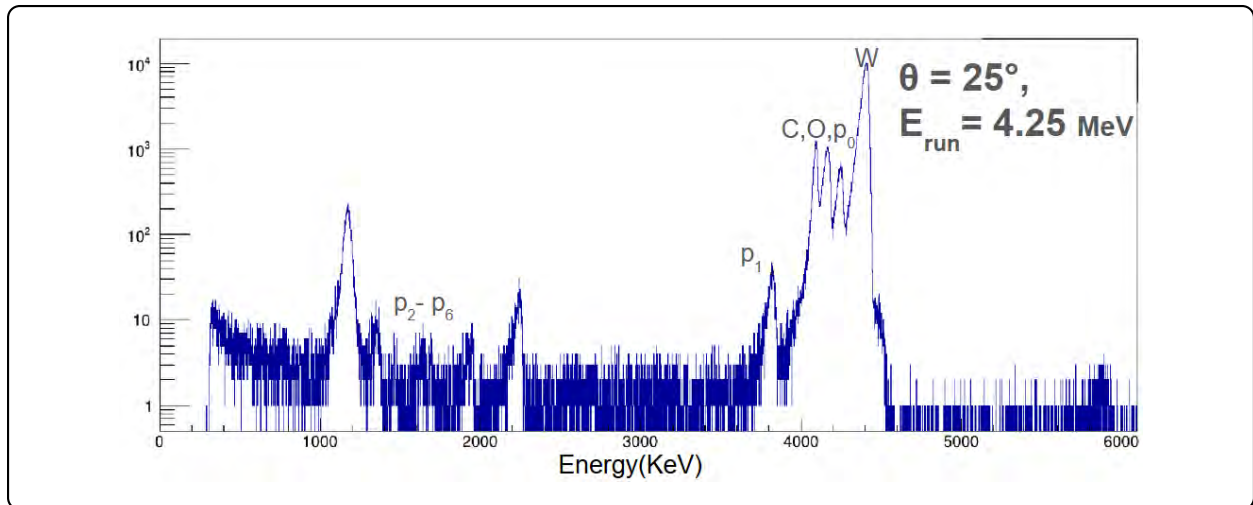


Figure 4: Direct Detector Energy Output from Accelerator Run with Resonance Channels Labeled

With the graph generated, processing and analysis can begin. A Gaussian fit is applied to the resonance of interest to better filter out noise and other particles' resonances. For this paper, the higher proton channels ($p_2 - p_6$) of the graph will be focused on. Figure 5 gives an example of a fit generated for the p_2 . The characteristics of the fit created allow for the extraction of cross-section and error data from the resonance. With that, the R-Matrix Formalism can be used to obtain the cross-section vs. energy plot needed for understanding the compound nucleus.

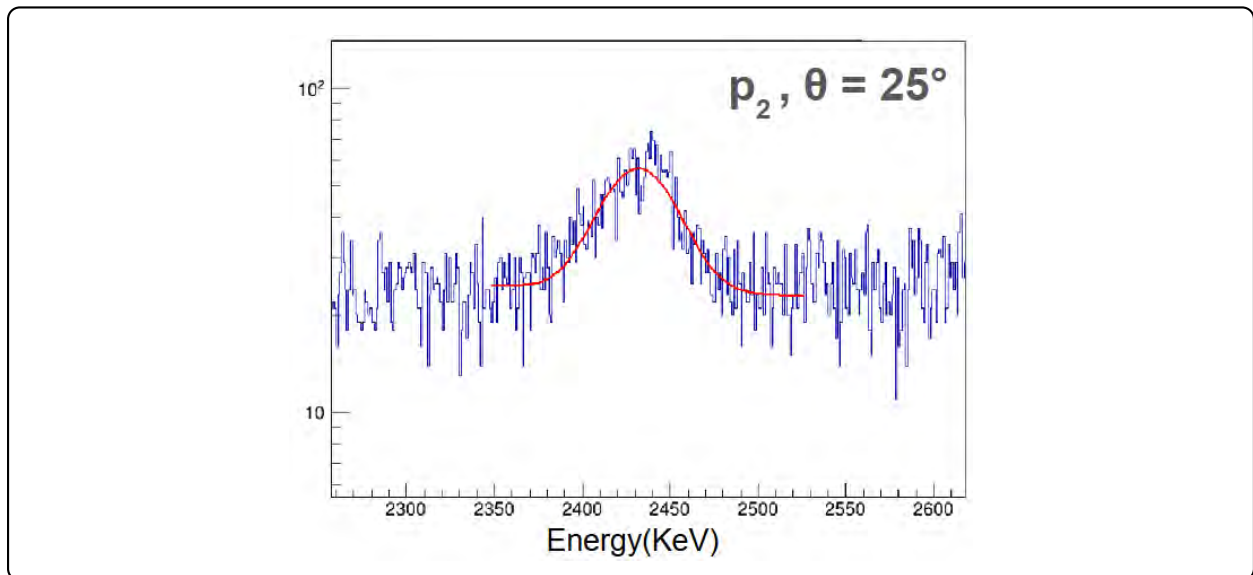


Figure 5: Plot of a Gaussian fit applied to p_2 Peak at $\theta = 25^\circ$

4 Results

A large portion of this analysis was spent on data wrangling for the Gaussian fits and getting the fit data ready for implementation in the R-Matrix software, AZURE2. This is a software that allows for the calculation of the different resonances present in the compound nucleus of a given reaction.

The data acquired from the Gaussian fits are loaded into the program, and resonances can be fit to the cross-section vs. energy created for different angles. At the end of the resonance fitting, Figure 6 was produced.

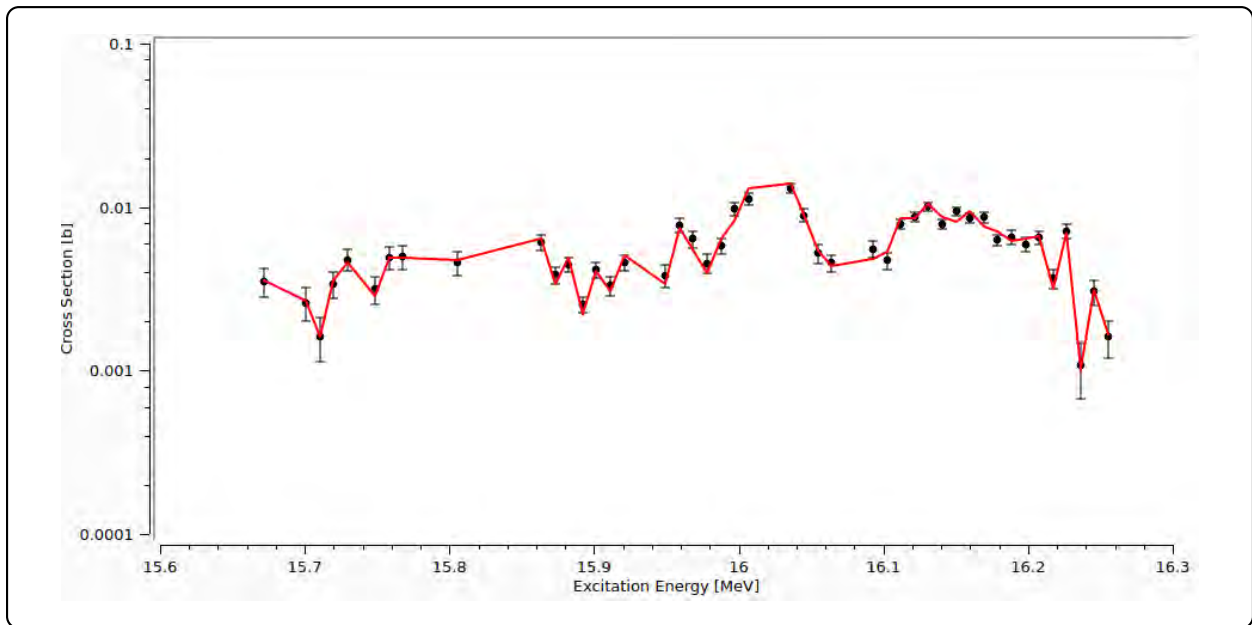


Figure 6: Plot of the resonance structure of $^{24}\text{Mg}^*$ at $\theta = 25^\circ$ using the p_2 exit channel

It is important to note that the cross-section is angularly dependent, so much work will need to be done in order to ensure that every angle of each proton channel is fit correctly. In theory, when a perfect characterization of the compound nucleus is achieved, one set of fit parameters should be able to accurately approximate the compound nucleus at every angle.

Discussion and Conclusion

The results from the initial R-matrix analysis using AZURE2 is another step toward understanding the resonance structure of the compound nucleus $^{24}\text{Mg}^*$. The ability to extract resonant behavior from the $^{23}\text{Na} + \text{p}$ channel offers an indirect method for probing the $^{12}\text{C} + ^{12}\text{C}$ fusion reaction.

One of the central challenges in this analysis lies in the angular dependence of the differential cross-section. Because of this, different scattering angles can emphasize different resonance contributions. As such, the goal of creating a single consistent set of resonance parameters that reproduces cross-section data at all angles is nontrivial. Discrepancies between fits at various angles may suggest incomplete level schemes, misassigned spin-parity combinations, or limitations in the experimental resolution and statistics.

Despite these challenges, the preliminary fits show that a number of the key features in the cross-section data can be reproduced within the R-matrix formalism. This confirms the suitability of the $^{23}\text{Na} + \text{p}$ channel as a tool for exploring the intermediate states populated in $^{12}\text{C} + ^{12}\text{C}$ fusion.

Further work will focus on expanding the fit across all measured angles, improving the precision of the Gaussian fitting process, and exploring different energies and channels of the $^{23}\text{Na} + \text{p}$. Ultimately, a full multi-channel R-matrix analysis will enable a more complete and reliable characterization of the $^{24}\text{Mg}^*$ nucleus, deepening our understanding of carbon fusion in astrophysical environments.

Acknowledgments

This work is supported in part by NSF Grant No. PHY-2310059 and NSF Grant No. 2411543. I would like to thank Dr. Aprahamian, Dr. Wanpeng, and Dr. Leshner for advising me throughout this REU. I would also like to thank all the members of the Nuclear Science Lab who helped me over the summer.

References

- [1] W. Nan, Y. Wang, J. Su, Y. Sheng, R. J. deBoer, Y. Zhang, L. Song, F. Cao, C. Chen, C. Dong, and et al., Determination of resonances in gamow window for $^{12}\text{C}+^{12}\text{C}$ fusion reaction via thick-target inverse kinematics method (2025).
- [2] A. M. Lane and R. G. Thomas, R-matrix theory of nuclear reactions (1958).

Simulation Integration and Dielectric Modeling for the St. Benedict Experiment

BROOKS SIMS

2025 NSF/REU Program
Department of Physics and Astronomy
University of Notre Dame

ADVISOR(S): Prof. Maxime Brodeur

Abstract

Recent theoretical corrections to the Cabibbo-Kobayashi-Maskawa (CKM) matrix have altered the value of its largest element, V_{ud} , leading to a 3σ tension with unitarity. This has prompted the need to extract a more accurate value of V_{ud} experimentally, a goal that the Superallowed Transition BEta-NEutrino Decay Ion Coincidence Trap (St. Benedict) experiment at the University of Notre Dame is focused on. To achieve this, radioactive ions are guided to a Paul trap via a series of ion optical components. By analyzing the time-of-flight spectra of the recoiling daughter nuclei confined in the Paul trap one can determine V_{ud} . This work focuses on connecting and improving simulations of two components of St. Benedict in SimIon, an ion optics simulation software. Previously, there was no method for smoothly transitioning ions across the radio frequency (RF) carpet and RF quadrupole ion guide simulations. To address this, an analysis of the potential maps of DC and RF electrodes in both simulations is done to create a new tool that allows for ions to be transitioned across simulations without having them experience a discontinuity in the electric field. This allows one to study the beam across various segments and scan St. Benedict for various experimental parameters. Additionally, PEEK and Kapton dielectrics that are part of the experimental RF carpet setup in the lab but not in simulation were added in SimIon. By implementing these dielectrics, the RF carpet simulation is expected to more closely resemble what is observed experimentally.

1 Introduction

The Standard Model (SM) is the most precise description of nature that is currently available, but it is still incomplete. Some of its major faults include lacks of: a gravitational interaction between particles, an explanation for the asymmetry between matter and antimatter, and an explanation for the observed dark matter and dark energy. Because of its deficiencies, many experiments are underway that attempt to precisely probe characteristics of the SM to look for physics Beyond the Standard Model. One such test is the unitarity of the CKM matrix. This is a 3×3 unitary matrix ($V_{CKM}^\dagger V_{CKM} = I$) that contains information on the flavor changing weak-interaction between quarks. The 9 elements of the CKM matrix are fundamental to the SM, so experimentally determining them and testing unitarity is a way to discover new physics and confirm existing physics.

1.1 St. Benedict

The St. Benedict experiment at Notre Dame aims to improve the accuracy on V_{ud} via a series of beta-neutrino correlation measurements obtained from an analysis of the time-of-flight (ToF) spectra of recoiling daughter nuclei from β -decaying ions in a Paul trap. Prior to their trapping, the ions must be guided through a series of beam manipulation components including a radio frequency (RF) carpet and a RF quadrupole ion guide (RFQ IG). This work focuses on improving our understanding of the transport of ions between these two components.

The TwinSol facility uses a double superconducting solenoid system to generate a low-energy radioactive ion beam (RIB). Upon arriving at St. Benedict, this beam loses energy in thin layers of aluminum before entering the gas catcher that is filled with helium gas at a pressure of around 30-40 Torr. By undergoing multiple collisions with the helium atoms, the radioactive ions eventually fully stop. These thermalized ions are then drifted towards the nozzle of the gas catcher using a combination of constant (DC) + RF voltages. A DC voltage gradient is applied to cause the ions to drift in the correct direction, and the RF signals are applied to repel the ions from the walls of the chamber. Near the nozzle the ions experience a large helium flow that drags them to the differentially pumped RF carpet chamber that sits at a pressure of around 2 Torr. On the wall opposite to the nozzle rests a so-called RF carpet that has a combination of DC + RF + LF (low frequency) signals to guide the ions towards another small opening after which lies the RFQ IG. LF signals are used here to create an electromagnetic traveling wave aimed towards the opening hole of the RFQ IG that the ion can “surf” on [1]. In the RFQ IG chamber, there are 4 cylindrical electrodes in the same positions as the plates on a baseball field (birds-eye view). These rods have alternating RF signals (+, -, +, -) that focuses the ions towards the center axis of the 4 rods (the pitcher’s plate). These ions then go into the Cooler-Buncher, which is used to form discrete packets of ions out of the continuous ion beam. These bunches are finally injected into the Paul trap where the measurement of interest is done. Further details about St. Benedict can be found in [2].

1.2 Simulating St. Benedict in SimIon

Each component of St. Benedict has its own individual simulation in a software called SimIon, an ion optics simulator. SimIon enables one to test various voltages, pressures, and geometries virtually, which allows one to assess the performance of the system. Essentially, one can better understand where ion transport losses occur and find ways to mediate them. SimIon calculates the potential of individual electrodes by discretizing space and then solving Laplace's/Poisson's equation based on boundary conditions. Then it calculates ion trajectories in these potential maps. With user programs, it can also handle time-varying signals (RF) and gas collisions.

In an ideal world, one would start ions in the gas catcher and have one simulation all the way to the end of the Paul trap. This, however, is not possible in SimIon due to chambers having different pressures and some compartments requiring a more fine discretization of simulation space (or different symmetry) than others. This limits the possibility of testing things like the transport efficiency of a beam across compartments at certain voltages since one cannot simulate ions continuously throughout all chambers of St. Benedict.

In this work, a method to transition ions across simulations is explored. First, a Python script is created to locate where ions should be transitioned at various voltages. Then we used this new transition tool to scan the various electrode potentials of the RF carpet and the RFQ IG compartments to study properties of the beam. Finally, we explored adding dielectrics in SimIon to further improve the physical representation of the simulation setup of the St. Benedict RF carpet.

2 Transitioning Ions Across Simulations

There are 2 simulations of interest: the RF carpet simulation and the RFQ IG simulation. The first will be referred to as CS (Carpet Simulation) and the latter as IGS (Ion Guide Simulation). Because these segments are connected sequentially, the end of the CS is the beginning of the IGS. The region of overlap between simulations will be referred to as the transition region since the goal is to find the x-value of the yz-plane in this region that transitions ions the smoothest. There are two rules

to follow to transition ions smoothly: no teleportation and no abrupt change in the electric field the ions experience.

The first rule implies that the ions must have the same (x, y, z) across simulations, and the second implies that it is transported to a place that has roughly the same V and $\vec{E} = -\nabla V$. The CS has 8 electrodes and the IGS has 10 electrodes. There are three electrodes in each simulation that are close enough to the transition region to have an impact. These electrodes are the RF carpet, the RF carpet aperture, and the upstream section of the ion guide. Due to the small dimension of the RF carpet electrodes ($160 \mu m$), the CS is very finely discretized but benefits from a cylindrical symmetry. On the other hand, the IGS does not have cylindrical symmetry but benefits from a more coarse discretization. Because of the cylindrical symmetry of the CS, the upstream section of the ion guide had to be mimicked by a hollow cylinder. Likewise, because of the coarse discretization of the IGS, the RF carpet had to be mimic as a solid electrode. As such, no RF from these mimic electrodes will be perceived in the simulations. Hence, in the next section we explored what is the effect of ignoring these RFs.

2.1 Effect of the RF in the transition region

To find the greatest impact of the electrodes with RF, we applied the amplitude of the signal as potential (75V for the RF carpet and 292V for the ion guide). Next, a potential map for this electrode (x, y, z, V) is generated by spawning uncharged particles at each point of interest then immediately killing them at the first time step. There are 961 sample points per yz -plane. To limit computation time, we limited ourselves to a maximum of 1000 ions and 961 is the closest square to 1000 ($31^2 = 961$). These ions are evenly spaced on a square yz -plane with $y, z \in [-0.55, 0.55]$ mm, and to fill the 3D space these planes repeat every $\Delta x = 0.05$ mm. These planes will run from $x \in [0, 5]$ mm, where 0 mm is inside the upstream section of the ion guide and 5 mm is just below the RF carpet.

Figure 1 shows a 2D interpolated contour plot of the potential map at $y = 0$ for the CS and the IGS respectively. In the IGS simulation, the RF amplitude is always less than 1V in this region, which means that it is small enough to be considered negligible. In the CS, the amplitude stays

small until $x \rightarrow 5$ mm, which is just below the RF carpet. From both of these results, one can conclude that it is safe to ignore time-varying signals for $x \in [0, 4.5]$ mm since the potentials here are fairly small. Now that RF can be ignored, the problem has turned into finding for which value of x the electrostatic potential between simulations is the most alike.

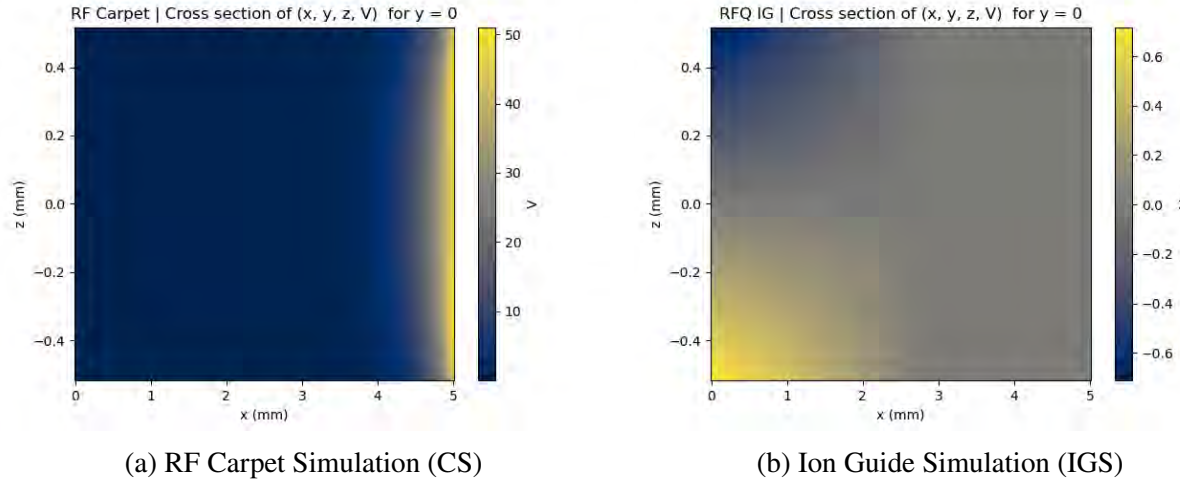


Figure 1: Comparison of the RF amplitude potential maps at $y = 0$ for the CS (left) and IGS (right).

2.2 Electrostatics potential matching

As mentioned previously, both simulations have 3 DC electrodes that have an impact on (x, y, z, V) near the transition region. Because of the linearity of Laplace's equation, one is allowed to work with one electrode at a time and then scale and sum up at the end to find the resulting (x, y, z, V) . The process of finding (x, y, z, V_{total}) is the following:

1. Set the first electrode to $V_1 = 1V$ and the other two to 0V.
2. Find (x, y, z, V_1) for all 961 points on all yz -planes for $x \in [0, 0.05, \dots, 5.0]$ mm.
3. Adjust the voltage of this electrode by multiplying V_1 by s_1 (s_1 is the actual voltage of the first electrode).
4. Repeat this for all three electrodes and sum to generate $(x, y, z, V_{total}) = (x, y, z, s_1 V_1 + s_2 V_2 + s_3 V_3)$.

If one does this process for both simulations, $(x, y, z, V_{total,CS})$ and $(x, y, z, V_{total,IGS})$ can be generated. The goal is to find where the difference between both simulations is the smallest, so the next step is to create (x, y, z, V_{diff}) , where $V_{diff} = |V_{total,CS} - V_{total,IGS}|$. Once (x, y, z, V_{diff}) is created, all that is left to do is sort through this space and find the x-value of the plane that gives the lowest difference between simulations.

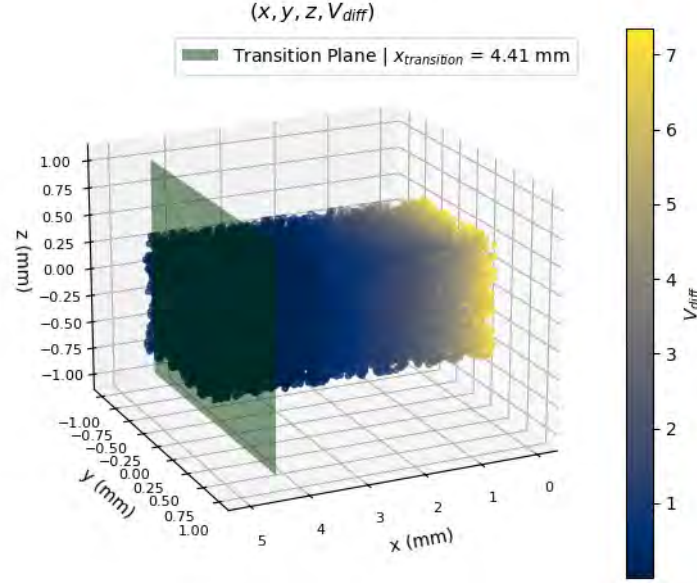


Figure 2: Plane of Transition Found by Sorting Through Potential Map

To accomplish this, a Python script was created. The data consists of 961 data points for each yz-plane with x-values $\in [0, 0.05, \dots, 5.0]$ mm. If one sums V_{diff} for all points on a plane and then divides by 961, $V_{diff,avg}$ has been found for this plane. This process can be repeated for all planes, and $x_{transition}$ will be the x-value of the yz-plane with the smallest $V_{diff,avg}$. The Python script picks a plane, finds the average difference on it, and stores it to an array. This is repeated for all yz-planes. There are now 2 arrays: $x_{plane} = [0, 0.05, \dots, 5.0]$ mm and $V_{diff} = [V_{diff,plane1}, V_{diff,plane2}, \dots]$. Using SciPy, these 2 arrays can be interpolated to find V_{diff} for all x-values. All that is left to do is to find the minimum of this interpolated function, and the minimize function from SciPy accomplishes just this. Figure 2 shows the result of this entire process. The plane that cuts through this 3D space at $x = 4.41$ mm is the plane that will be used to transition ions between simulations. It can be visually observed that V_{diff} for all points on this plane are close to 0V from the color map.

With this new tool, ions can be sent across simulations for any voltages desired.

3 Scanning DC Voltages to Test Transport Efficiency and ToF

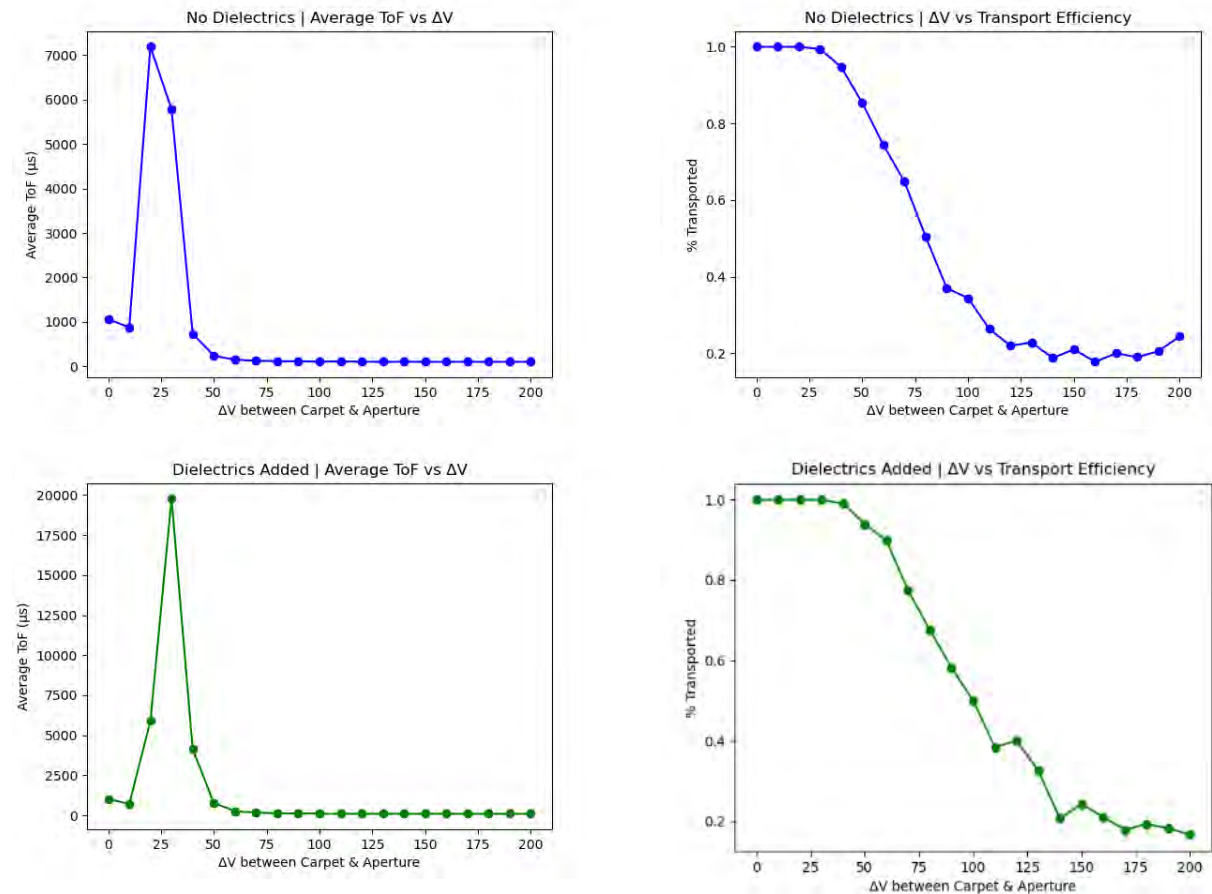
The next goal is to use this tool to scan potentials of the St. Benedict RF carpet and ion guide in SimIon. In particular, we are interested in answering the question of what potentials gives the best transport efficiency and a reasonable time-of-flight. The three DC electrodes discussed so far are: carpet electrode (CE), aperture electrode (AE), and upstream ion guide electrode (UIGE). These abbreviations will be used for brevity.

The CE is held at -75V always, the AE varies from [-75, -85, ..., -275]V and the UIGE is always 70V below the AE for all scans. The key behind these scans is automating everything in the Lua language: the language that SimIon user-programs run on. For each voltage triplet, the user-program must tell SimIon where to transition ions, check if an ion reached the transition plane, and write ion end states to a .txt for each set of voltages. The method for doing this was to hard code an array into the Lua script that has the $x_{transition}$ value for all voltage triplets for this scan. This has the limitation of not being dynamic, but it still automates the scans. The Lua script also writes the end state of ions that reached the transition plane to a .ion file (files that are used to initialize particles) so that they can be transitioned to the IGS.

As indicated in Figure 3, when ΔV between the CE and AE is relatively small, the transport efficiency is the highest and the time-of-flight is low. This indicates that we should have a low potential difference between CE and AE. When ΔV increases to around 20-30V, however, the ions get hung up and take a very long time to find their way to the next segment. Additionally, at high voltages ions were going between adjacent electrodes on the carpet which was unusual behavior. This could be due to the large negative potential of the AE penetrating between the CE. Both of these predicaments led to the idea that dielectrics should be added to the RF carpet simulation. In the lab there is an appreciable amount of dielectrics on the RF carpet which might affect the simulation results.

4 Adding dielectrics to the RF Carpet Simulation

SimIon 8.1 allows one to include dielectrics via its Poisson's equation solver. To do this, Lua scripts + Windows batch scripts were written to turn the geometry files of the dielectrics into potential arrays, and then Poisson's equation was solved with the permittivity depending on the dielectric potential array.



(a) Average ToF vs. ΔV with and without dielectrics

(b) Transport efficiency vs. ΔV with and without dielectrics

Figure 3: Comparison of ToF and transport efficiency vs. ΔV for simulations with and without dielectrics

Figure 3 shows the comparison of the ToF and transport efficiency results for the CS with and without dielectrics. Clearly, the dielectrics cause the ions to hang up at the same voltages as they did without them, but their hang-up time is increased dramatically. Additionally, the transport efficiency with and without dielectrics is nearly identical for all voltages. This result indicates that

the dielectrics certainly did have an impact on the simulation, but it is unfortunately an unwanted impact. Future work must be done on understanding why ions are getting hung-up at these voltages and how to improve these simulations to remove hang-up spots and more closely resemble what's in the lab.

5 Conclusion

The lack of continuity between the RF carpet and RFQ IG simulations has been remedied by finding an electrically similar transition plane between both simulations. In both simulations near the transition region, RF signals were found to be much weaker than DC signals which simplified the problem to an electrostatics one. To solve this electrostatics problem, one must sort through (x, y, z, V) in the transition region for both simulations and find the x -value of the yz -plane that minimizes the potential difference between simulations. To accomplish this, a Python script was made to analyze the potential map at all discrete points in space. For any triplet of voltages, the Python script will output the x -value of where one should transition ions across simulations. With this tool, scans were done for 21 triplets of voltages and it was observed that ions were having unexpected behavior. To remedy this, dielectrics were added to the RF carpet simulation to improve its accuracy to what is observed experimentally. The same problems persisted even after dielectrics were incorporated, so future work needs to be done on understanding why there are problem voltages and how to further match SimIon simulations to what is observed while running St. Benedict.

References

- [1] G. Bollen, International Journal of Mass Spectrometry **299**, 131 (2011).
- [2] M. Brodeur, T. Ahn, D. Bardayan, D. Burdette, J. Clark, A. Gallant, J. Kolata, B. Liu, P. O'Malley, W. Porter, R. Ringle, F. Rivero, G. Savard, A. Valverde, and R. Zite, Nuclear Instruments and Methods in Physics Research Section B: Beam Interactions with Materials and Atoms **541**, 79 (2023).

Analyzing Over 500 Stars in Local Group Galaxies Using Chemfit

AIDAN STARRS

2025 NSF/REU Program
Department of Physics and Astronomy
University of Notre Dame

ADVISOR(S): Lauren Henderson, Prof. Evan Kirby

Abstract

The metallicity ($[M/H]$) and alpha element abundances ($[\alpha/H]$) of Local Group dwarf spheroidal galaxies (dSphs) provide insight into their star formation histories. Learning about the star formation of these galaxies can give information about galactic evolution in the Milky Way and the universe as a whole. We worked to determine the metallicity distribution and alpha element abundances of over 500 stars in the dSphs Sculptor and Leo I using a newly developed software called Chemfit that measures stellar parameters and abundances from spectra. In comparison to previously established values, our results for Sculptor corroborate the metallicity distribution and the evolution of $[\alpha/M]$ over metallicity. In addition, our study of Leo I showed an increase in the detail of the $[\alpha/M]$ - $[M/H]$ relationship as a result of using bluer spectra than previous work. However, some discrepancy between results from red and blue spectra remain to be studied. The agreement between our values and previously accepted determinations emphasizes the successful capabilities of Chemfit and opens numerous avenues for further investigation. For example, Chemfit can be used to analyze individual chemical abundances, in addition to being used for other dSphs to further increase our understanding of the Local Group.

1 Introduction

The Local Group dwarf galaxies, specifically dwarf spheroidal galaxies (dSphs) such as Sculptor and Leo I, are key contributors to understanding chemical evolution. Their old ages and relatively simple abundance patterns allow for studying galactic evolution at early times in our universe [1]. The stars within these galaxies provide insight into the creation of elements and the formation history of the galaxy as a whole.

The α element (O, Ne, Mg, Si, S, Ca, Ar, Ti) abundances ($[\alpha/H]$) are often used in the determination of the star formation history. Core-collapse supernovae (CCSNe) are primarily responsible for the release of the α elements and thus provide information about the approximate chemical evolution of the galaxies. These occur when a massive star exhausts its supply of nuclear fuel and can no longer withstand the force of gravity, resulting in the implosion, core bounce, and subsequent explosion of the star.

This study uses the software Chemfit to analyze absorption spectra and determine stellar parameters and abundances of over 500 stars in the Local Group dSphs Sculptor and Leo I. Moreover, this study compares the results of the newly developed Chemfit to the established literature values

of Kirby *et al.* [2].

In Section 2, the literature used for the comparison is described. In Section 3, I describe the process for writing a Python script in order to properly run Chemfit for the various files of data and the methods used in determining the quality of data and which parts to mask. Throughout Section 4 I discuss the significance of these results specifically in comparison to the literature [2; 3]. Section 5 is a summary of our conclusions and their importance.

2 Background

2.1 Established Literature of Sculptor

Sculptor is a very well studied dwarf spheroidal galaxy and is often referred to as a “textbook dSph.” Hill *et al.* [4] defines it as so in their analysis of 99 high-resolution spectra in Sculptor. They detail how abundance ratios evolve with metallicity and time very closely to predictions from basic nucleosynthesis. Specifically, they highlight a steady decrease in $[\alpha/\text{Fe}]$ after a knee starting at the Galactic halo plateau value for low $[\text{Fe}/\text{H}]$.

Kirby *et al.* [2] furthers the understanding of Sculptor by analyzing 388 medium-resolution spectra. They produced an unbiased metallicity distribution function that exhibits an asymmetric shape with a metal-poor tail. Furthermore, they confirmed the decrease in $[\alpha/\text{Fe}]$ with $[\text{Fe}/\text{H}]$. Their results also indicate the ratio of $[\alpha/\text{Fe}]$ becoming negative (subsolar) at higher values of $[\text{Fe}/\text{H}]$ as predicted by previous models.

2.2 Established Literature of Leo I

Leo I is not as frequently studied as Sculptor but is just as helpful in understanding the origins of galaxies. Kirby *et al.* [3] derives the star formation history of Leo I from its alpha abundance pattern. Additionally, they found a moderate correlation between $[\alpha/\text{Fe}]$ and $[\text{Fe}/\text{H}]$ particularly with lower metallicity stars displaying higher $[\alpha/\text{Fe}]$ than more metal-rich stars.

Another study by Gullieuszik *et al.* [5]. studied the metallicity of just over 50 stars in Leo I and

found the metallicity distribution function to be symmetric and narrow. Their results combined with existing photometric data provide a age-metallicity relation. The age-metallicity relation suggests a rapid initial enrichment. Additionally, the chemical history of the galaxy is not well constrained at early epochs because it only contains a minor old stellar component.

3 Methods

To conduct this study we used data from the DEIMOS instrument at Keck Observatory in Hawaii. It obtains visible-wavelength spectra using masks with slits that are aligned with the target stars in the sky. The blue spectra range from approximately 4000Å to 6800Å and the red from approximately 6300Å to 9100Å.

In order to determine the stellar parameters like temperature, surface gravity, $[\alpha/H]$, $[M/H]$, $[C/H]$, and the redshift of these Local Group stars, we created a Python script that runs Chemfit for any reduced DEIMOS spectrum. Chemfit was used to determine the stellar parameters of each star by fitting the observed spectrum to a model spectrum. To do this, Chemfit accesses a large grid of synthetic spectra with varying stellar parameters and interpolates between them to retrieve the parameters that best fit the observed spectrum. Chemfit requires the input of various data accessed from the input data files, such as the wavelength, flux, inverse variance, and photometry. These are used for measuring temperature, surface gravity, alpha-element abundance, total metallicity, carbon abundance, and redshift of all of the stars.

Although the code will run without the photometry information being included, it is highly important as it creates a “prior” for the model. By including the photometry, we are telling Chemfit the color of the star, which in turn restricts the temperature values of the star because a star’s color is directly related to its temperature. This prevents a gross over- or under-estimation of the temperature.

In detail, Chemfit computes synthetic photometry from the model spectra. It convolves the filter transmission function with the flux-calibrated model spectrum. The result of this convolution is converted to a photometric magnitude. The difference between the magnitude of two filters is

a color. The color is then compared with the observed photometric color. The color is treated like another “pixel” in the χ^2 goodness-of-fit function, except that this pixel has significantly more weight than any single spectral pixel.

We performed a quality check on all of the stars, filtering out any star for which a radial velocity is not consistent with the dwarf galaxy in which it is located or could not be reliably measured. Moreover, for each star, we masked out some wavelength ranges due to significant noise. For instance, we masked the ends of the spectra due to the heavy noise present. In addition, we discarded any points that showed a zero flux and a few of the surrounding indices so as not to influence the model and parameter determination. Moreover, every spectrum from DEIMOS contains a CCD chip gap that must be masked out because the spectral pixels near the chip gap are often poorly calibrated.

We began with Sculptor as our primary target because it is a frequent source of research (it is sometimes called a “textbook dwarf galaxy”), and thus there was already a significant amount of data to compare with. Since it is so well studied, most of the stars also contained two spectra, a bluer and redder spectrum. One of the main purposes of the script I wrote was to filter through all of the stars between two data files and plot their absorption spectra depending on whether they had just blue, both blue and red, or just red spectra. It is important to note that the script we wrote for Chemfit can run any of these three combinations.

Likewise, we ran Chemfit for the stars of Leo I. We ran separate analyses of the blue, red, and both blue and red spectra, respectively to test how the measurements depend on wavelength coverage.

4 Results

4.1 Literature Comparison of Sculptor

In analyzing our Chemfit results, we created a metallicity distribution as well as a plot of the α abundance against the total metallicity ($[\alpha/\text{H}]$ vs. $[\text{M}/\text{H}]$). Figure 1 depicts the metallicity distri-

bution as determined from Chemfit. It displays a highly metal-poor tail and overall asymmetric distribution that is skewed left. Additionally, there is a peak at approximately $[M/H] = -1.6$.

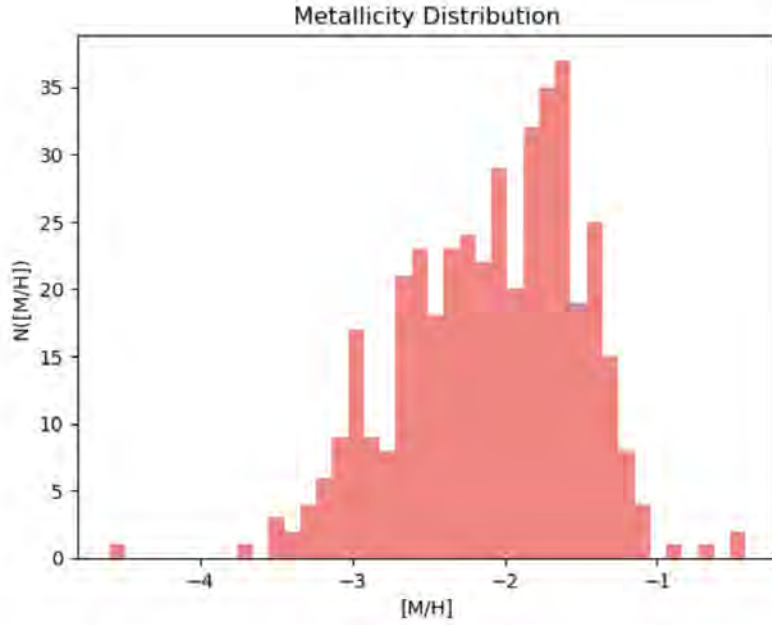


Figure 1: Sculptor’s metallicity distribution as determined in this study.

A metallicity distribution is valuable for understanding the star formation history of a galaxy. Kirby *et al.* [2] established the metallicity distribution of Sculptor using an alternative method to Chemfit. They observed a very asymmetric distribution skewed left with an extremely metal-poor tail. As seen in Figure 1, there is a strong similarity in the shape and distributions of both our data and that of Kirby. Furthermore, Kirby observes a peak in the distribution at approximately $[Fe/H] = -1.4$. The most likely reason for the differences in peaks is the fact that Kirby uses specifically the abundance of iron ($[Fe/H]$) on the x-axis rather than the bulk metallicity, M . However, this remains a safe comparison to make because the metallicity is primarily determined through the iron abundance in stars.

Figure 2 depicts the abundances of α elements versus the metallicity of the stars. It shows a plateau in the alpha abundance until approximately $[M/H] = -2.0$ in which a knee forms and $[\alpha/M]$ begins to decrease.

This is very similar to the trend that is predicted for relatively simple stellar populations like

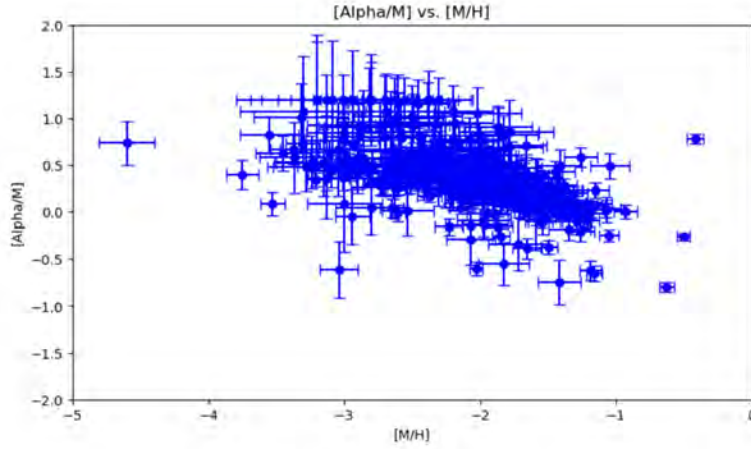


Figure 2: Sculptor's α element abundance with respect to the metallicity.

Sculptor. Metallicity increases monotonically over time [6]. Different galactic events like core-collapse SNe and Type Ia SNe have different delay periods, creating the shape of the classic knee. CCSNe have a relatively short delay time, creating α elements at early times, as indicated by the plateau at low metallicity. In other words, at low metallicities massive stars are dying (CCSNe) and creating α elements like Mg. But as time progresses, the lower mass stars begin dying as well (Type Ia SNe) which increases the metallicity, through the creation of elements like iron, and thus decreases their ratio and results in a negative slope.

These observations are corroborated once again by Kirby's discoveries as seen in their $[\alpha/\text{H}]$ - $[\text{Fe}/\text{H}]$ plot [2]. It is important to note that the differences in the location of the knee may be accounted for by Kirby's use of $[\text{Fe}/\text{H}]$ rather than total metallicity. However, once again it is still a safe comparison to make as the general shape does not change significantly.

In addition, we also created two color magnitude diagrams (CMD) to confirm our measurements are acting as predicted as shown in Figure 3. The CMD on the left, colored by temperature, depicts a rainbow gradient as the stars progress up toward the top right. This shows the evolution of red giant stars up the red giant branch. As the stars get older and cool down, they begin to expand and thus become brighter, hence their higher position on the plot. It is important to note that the vertical axes of these plots are inverted, and a lower I value on the y-axis corresponds to a brighter star. Likewise, the CMD colored by $[\text{M}/\text{H}]$ is also as predicted as seen by the rainbow aligned with the

curvature of the red giant branch. The metals in the stars atmosphere absorb more light and thus make the star more opaque. More metals in the atmosphere of the star make the star appear redder.

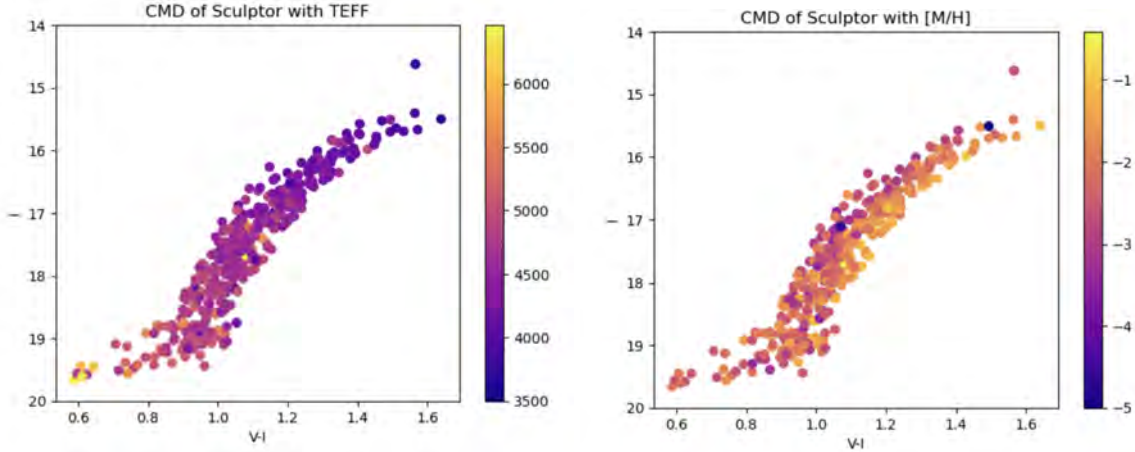


Figure 3: The plot on the left is a color magnitude diagram (CMD) with respect to the temperature of the Sculptor stars. The plot on the right is a CMD of the same Sculptor stars with respect to metallicity.

4.2 Literature Comparison of Leo I

Similar to Sculptor, we made comparisons to previous literature for our results from Leo I. We created a $[\alpha/\text{H}]$ vs. $[\text{M}/\text{H}]$ plot and compared to a previously determined $[\text{Mg}/\text{Fe}]$ vs. $[\text{Fe}/\text{H}]$ plot. Mg is the main contributor in the α element abundance and therefore is a safe comparison to make.

Figure 4 shows the alpha abundances determined using just blue, just red, and blue and red spectra, respectively. There is similarities between the strictly blue spectra points and the points with both spectra seen by the plateau until around $[\text{M}/\text{H}] = -1.5$, where a knee begins to form. However, it is important to note that the red spectra consistently yielded higher metallicities and a higher spread of data than the other points. These discrepancies remain to be studied.

Our results differ from that of Kirby *et al.* [3], where they do not define a clearly visible knee or plateau. Additionally, our data exhibits less spread than Kirby, with almost all the $[\alpha/\text{H}]$ values

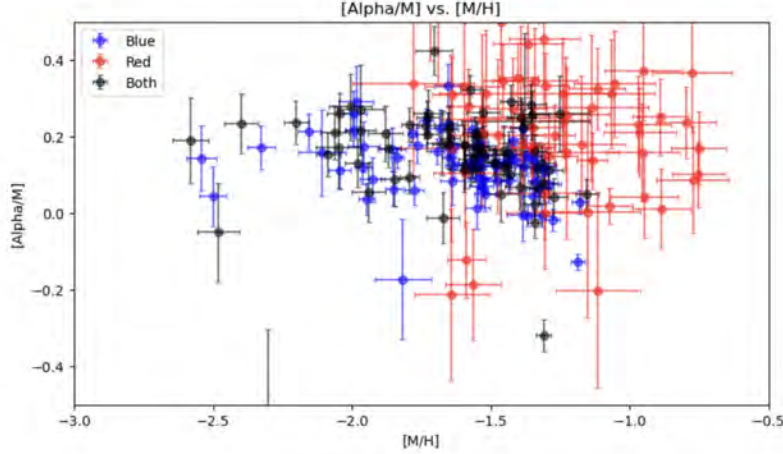


Figure 4: Leo I's α element abundance with respect to the metallicity. The red points indicate $[\alpha/M]$ determined from only red spectra. The blue points indicate $[\alpha/M]$ determined from only blue spectra. The black points indicate $[\alpha/M]$ determined from a synthesis of the blue and red spectra.

falling between $-0.4 < [\alpha/H] < 0.4$, whereas in Kirby they fall between $-0.5 < [\text{Mg/Fe}] < 1$. But, Kirby's analysis uses strictly red spectra whereas ours includes blue spectra as well. Compared with just our red data they follow the same general trend, where lower metallicity stars tend to have higher α abundances.

5 Conclusion

In this project we analyzed stars from the dSphs Sculptor and Leo I and determined their stellar parameters and abundances using Chemfit. We then compared our results to the previously accepted literature values. Primarily, we found that Chemfit using blue spectra gives more precise measurements than previous results using red spectra, though our results are in agreement with that of the literature from Kirby. Additionally, the $[\alpha/M]$ differences in Leo I between the strictly red and strictly blue spectra are a result that requires further investigation. Furthermore, these stellar populations can give knowledge into topics like dwarf galaxy star formation more broadly.

The success of using Chemift opens many avenues for future research. For instance, it can be used to determine the individual elemental abundances to more accurately determine the star formation histories and nucleosynthesis of the dSphs. Also, Chemfit can be applied to other dSphs,

or other Local Group galaxies to increase the sample size to offer understanding of our own Milky Way.

References

- [1] L. E. Henderson, E. N. Kirby, M. A. C. de los Reyes, R. Gerasimov, and V. Manwadkar, *ApJ* **983**, 117 (2025), arXiv:2502.16447 [astro-ph.GA] .
- [2] E. N. Kirby, P. Guhathakurta, M. Bolte, C. Sneden, and M. C. Geha, *ApJ* **705**, 328 (2009), arXiv:0909.3092 [astro-ph.GA] .
- [3] E. N. Kirby, J. G. Cohen, G. H. Smith, S. R. Majewski, S. T. Sohn, and P. Guhathakurta, *ApJ* **727**, 79 (2011), arXiv:1011.5221 [astro-ph.GA] .
- [4] V. Hill, Á. Skúladóttir, E. Tolstoy, K. A. Venn, M. D. Shetrone, P. Jablonka, F. Primas, G. Battaglia, T. J. L. de Boer, P. François, A. Helmi, A. Kaufer, B. Letarte, E. Starkenburg, and M. Spite, *A&A* **626**, A15 (2019), arXiv:1812.01486 [astro-ph.GA] .
- [5] M. Gullieuszik, E. V. Held, I. Saviane, and L. Rizzi, *A&A* **500**, 735 (2009), arXiv:0904.2718 [astro-ph.CO] .
- [6] C. Kobayashi, A. I. Karakas, and M. Lugaro, *ApJ* **900**, 179 (2020), arXiv:2008.04660 [astro-ph.GA] .

Allometry-Informed Inference of Species Interactions from Equilibrium Data in Stochastic Lotka–Volterra Ecosystems

JOSEPH TEMPLE

2025 NSF/REU Program
Department of Physics and Astronomy
University of Notre Dame

ADVISOR(S): Prof. Derviş Can Vural

Abstract

Ecological network reconstruction from limited observational data remains a fundamental challenge in understanding ecosystem dynamics and stability. We present a novel method to infer the interaction matrix of a stochastic Lotka–Volterra ecosystem using only noisy steady-state population densities and allometric scaling laws. Our algorithm, framed as a Bayesian maximum a posteriori (MAP) estimation with biologically-informed priors defined by unknown hyperparameters, combines global optimization via Differential Evolution with local refinement via L-BFGS-B. We evaluate the method’s accuracy on synthetically generated ecosystems, varying key system parameters: ecosystem size, deviation from allometric laws, and noise in equilibrium observations. Our *in silico* experiments demonstrate reliable inference of interaction matrices in low-noise regimes, with performance declining predictably as noise and prior deviation increase.

1 Introduction

The behavior and stability of an ecosystem are governed by interactions among its constituent species. As such, the network structure of any ecosystem can be codified in an *interaction matrix*, $\mathbf{A}$, which quantifies both the type and strength of each pairwise species interaction. One such interaction is known as an *interaction coefficient*, A_{ij} . Accurately determining these coefficients provides ecologists with powerful mathematical tools to predict community stability and responses to environmental changes, as well as to guide targeted conservation efforts.

Direct measurement of interaction coefficients has proven experimentally challenging due to high sensitivity with respect to environmental conditions and the complexity of real interactions [1]. As such, significant research is focused on developing *inference* methods, particularly with time-series data [2]. Our work addresses a gap in this literature; we approach inference by combining noisy equilibrium population data with allometry (the scaling of biological traits with body mass).

This inverse problem is underdetermined: equilibrium measurements alone do not carry enough information to uniquely determine an ecosystem structure. To overcome this, we cast inference as a Bayesian maximum a posteriori (MAP) problem, introducing allometric priors. Our resulting algorithm represents a combination of two competing goals: (i) accurate recovery of the interaction matrix, with (ii) minimal assumptions on the structure of that matrix.

2 Background

2.1 The Generalized Lotka–Volterra Model

The generalized Lotka–Volterra model (see Box 1) is a system of ordinary differential equations, Equation (1), that describe the time evolution of population densities for a set of interacting species. The equilibrium point of this system, Equation (2), is determined solely by the interaction matrix $\mathbf{A}_0$ and the intrinsic growth rate vector $\vec{r}$.

Box 1: Deterministic Lotka–Volterra Model	
$\frac{d\vec{x}}{dt} = \vec{x} \circ (\vec{r} + \mathbf{A}_0 \vec{x}) \quad (1)$	$\vec{x}$: Population densities of each species.
$\frac{d\vec{x}}{dt} = \vec{0} \rightarrow \vec{x} = -\mathbf{A}_0^{-1} \vec{r} \quad (2)$	$\vec{r}$: Intrinsic growth rates of each species.
$\vec{x}, \vec{r} \in \mathbb{R}^N; \mathbf{A}_0 \in \mathbb{R}^{N \times N}$	$\mathbf{A}_0$: Interaction matrix; A_{ij} is the effect of species j on population of species i .
	N : Number of species in the ecosystem.
	$\circ$: Hadamard (element-wise) vector product.

We introduce stochasticity (see Box 2) into the elements of the interaction matrix—more accurately modeling the biological reality of fluctuating interactions—leading to several stable equilibria from the same underlying structure. That is, stochastic perturbations in the interaction matrix result in stochastic perturbations in the observed equilibrium.

Box 2: Stochastic Lotka–Volterra Model	
$\frac{d\vec{y}}{dt} = \vec{0} \rightarrow \vec{y} = -\mathbf{A}^{-1} \vec{r} \quad (3)$	$\vec{y}$: Population densities of each species under stochastic interaction matrix.
$\mathbf{A} = \mathbf{A}_0 + \mathbf{W} \quad (4)$	$\mathbf{A}$: Interaction matrix with stochastically drawn elements.
$W_{ij} \sim \mathcal{N}(0, \sigma_{ij}^2)$	σ : Interaction deviation matrix, element σ_{ij} is the standard deviation of A_{ij} .
$\vec{y} \in \mathbb{R}^N; \mathbf{A}, \mathbf{W}, \sigma \in \mathbb{R}^{N \times N}$	$\mathbf{W}$: Noise matrix, each W_{ij} sampled from a zero-mean Gaussian with variance σ_{ij}^2 .

After substituting Equation (4) into Equation (3), taking a first-order Neumann expansion (a matrix analog of the Taylor series expansion for $(1+x)^{-1}$) provides us with the linear approximation $\vec{y} \approx [\mathbf{I} - \mathbf{A}_0^{-1} \mathbf{W}] \vec{x}$. From this approximation, the distribution of stochastic equilibria $\vec{y}$ can be reasonably approximated as a multivariate Gaussian whose mean is $\vec{x}$ and whose covariance matrix is defined $\Sigma \equiv \mathbf{A}^{-1} \mathbf{D} \mathbf{A}^{-\top}$, where $\mathbf{D}$ is a diagonal matrix with entries $D_{pp} = \sum_k \sigma_{pk}^2 x_k^2$.

2.2 Allometry: Constraining Inference with Biology

Allometry is the study of how different biological functions scale with respect to body mass. Kleiber’s law is an empirically derived relationship stating that the basal metabolic rate (the rate of energy expenditure at rest) scales as $\text{BMR} \propto m^{3/4}$; this holds remarkably true across all scales of organism [3]. Similarly, the intrinsic growth rate, the rate at which a population grows in the absence of interactions with other species, scales as $r_i \propto m_i^{-1/4}$, according to Damuth’s law.

2.2.1 Carrying Capacity: Diagonal Elements

Diagonal elements of the interaction matrix represent self-regulation (the effect of species i on itself), in the absence of other species. Comparing the common logistic model for limited population growth to the Lotka–Volterra model for $i = j$, we can define $A_{ii} = \frac{-r_i}{K_i}$, where K_i is the carrying capacity of the species. By Damuth’s law, carrying capacity per area scales as $K_i \propto m_i^{-3/4}$ [4].

2.2.2 Sight and Speed: Upper-Triangle

Following the approach of Rall et al., we interpret the magnitude of inter-species interaction as the attack rate in a Holling Type II functional response [5]. This leads to the scaling relationship $|A_{ij}| \propto m_i^\alpha m_j^\beta$, where m_i and m_j are the body masses of species i and j . The exponents α and β depend on ecological traits of the “attacking” species i , such as its typical speed, perceptual range, and whether it operates in a two- or three-dimensional environment, all of which scale as power laws with respect to mass. For ease of modeling, we assume each species has access to 2.5 dimensions—neither restricted to 2D nor unbound in 3D—resulting in $\alpha = 0.75$ and $\beta = 0.5$ [5].

2.2.3 Predator-Prey and Resource-Mediated Interactions: Lower-Triangle

For predator-prey interactions (where A_{ij} and A_{ji} have opposite signs), assimilation efficiency c quantifies the fraction of prey biomass that gets converted into predator biomass, giving $A_{ji} = -c \left(\frac{m_i}{m_j} \right) A_{ij}$, with predator j and prey i [6]. With no specified predator-prey relationship, we take a ratio of the relationship between A_{ij} and A_{ji} from Section 2.2.2, finding that $A_{ji} = \left(\frac{m_j}{m_i} \right)^{\alpha-\beta} A_{ij}$.

2.3 Problem Formulation via Bayes' Theorem

Say we have a set of M experimental equilibrium population density observations, $\mathcal{Y} = \{\vec{y}^{(k)}\}_{k=1}^M$. The question we ask is “given the dataset $\mathcal{Y}$, what interaction matrix $\mathbf{A}$ is most likely to have generated that data?” Bayes’ theorem, Equation (5), allows us to answer exactly this question.

$$P(\mathbf{A} \mid \mathcal{Y}) \propto P(\mathcal{Y} \mid \mathbf{A})P(\mathbf{A}) \quad (5)$$

In our work, the likelihood $P(\mathcal{Y} \mid \mathbf{A})$ is the probability density function of the distribution discussed in Section 2.1, while the prior $P(\mathbf{A})$ is a product of Gaussians with means μ_{ij} determined by allometry and standard deviations $\tau_{ij} = f|\mu_{ij}|$ being some fraction f of the mean’s magnitude. Our goal is to find the $\mathbf{A}$ matrix for which the posterior $P(\mathbf{A}|\mathcal{Y})$ is greatest, the *maximum a posteriori* (MAP) estimate. To stabilize numerical optimization and leverage existing optimizers, we do this by *minimizing* the negative log posterior, $\mathcal{L} \equiv -\ln [P(\mathcal{Y}|\mathbf{A})P(\mathbf{A})]$, serving as a loss function.

3 Methods

Inferring the interaction matrix $\mathbf{A}$ from noisy equilibrium observations is a high-dimensional, non-linear optimization problem. Because equilibrium data alone does not uniquely determine $\mathbf{A}$, we regularize inference with allometric priors derived from biological scaling laws. However, due to the scarcity of experimental datasets containing both steady-state population data and allometric parameterization, we instead generate synthetic data to benchmark the accuracy of our method.

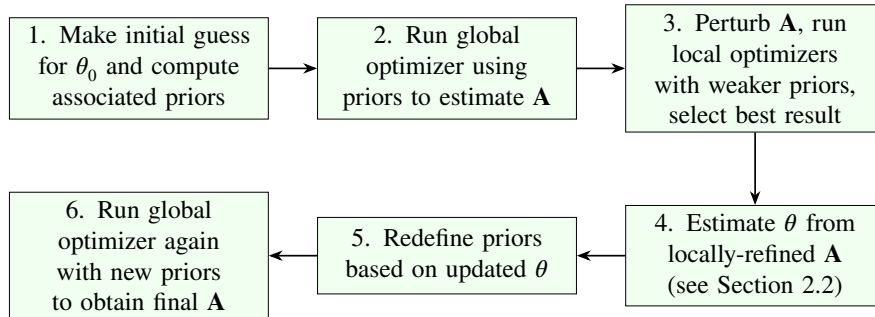


Figure 1: Optimization pipeline flowchart, combining global and local solves to estimate interaction matrix $\mathbf{A}$, inferring allometric hyperparameters θ from data in the process.

3.1 Synthetic Data Generation

3.1.1 Interaction Matrix, $\mathbf{A}$

Our interaction matrices are generated according to the allometric laws discussed in Section 2.2. We must choose values for the proportionality constants— K_0 for the carrying capacity density, a_0 for the attack rate, and c for assimilation efficiency—the set of which we call the *allometric hyperparameters*, θ . With these, each element of the synthetic interaction matrix is sampled from Gaussian distributions similar in structure to those defining the priors.

Interaction types are assigned by labeling each species with a “trophic level” from the set $\{1, 2, 3\}$, corresponding to the classes of producers, primary consumers, and secondary consumers. The sign of each interaction is then defined according to the difference in trophic level, ΔT . If $\Delta T = 0$, the interaction coefficients A_{ij} and A_{ji} will have the same sign, representing mutualism (positive) and competition (negative) respectively. If $\Delta T = 1$, we assign A_{ij} and A_{ji} to have opposite signs, modeling a predator-prey relationship where predator i on prey j results in $A_{ij} > 0$ and $A_{ji} < 0$. Finally, if $\Delta T = 2$, we assume the species have no interaction such that $A_{ij} = A_{ji} = 0$.

3.1.2 Equilibrium Population Densities, $\mathcal{Y}$

Each realization of the steady-state for an ecosystem is a solution to Equation (3). Thus, to generate $\mathcal{Y}$ we solve this equation according to M random samples of W_{ij} for a given interaction deviation matrix σ , where each σ_{ij} is fixed as some fraction g of $|A_{ij}|$.

3.1.3 Ecosystem Validity Criteria

For multiple equilibrium measurements to be taken, an ecosystem must be *stable*, a condition which is guaranteed only when all eigenvalues of the interaction matrix have negative real parts [7]; if a synthetic $\mathbf{A}$ fails this test, we sample each element again. It is also possible for these fluctuated matrices to result in $\vec{y}$ with one or more negative entries, violating the condition of *feasibility*, as a population density cannot be negative. As such, any $\vec{y}$ generated where $y_i \leq 0$ for any i is discarded.

3.2 Global Search: Differential Evolution

Differential Evolution (DE) is a population-based global optimization algorithm. The population consists of candidate solutions, flattened vector representations $\vec{A}_x$ of interaction matrices. At each generation, new candidates $\vec{A}_{\text{new}}$ are proposed by adding $\vec{A}_x$ to the randomly selected scaled difference $\eta(\vec{A}_y - \vec{A}_z)$, where η is chosen to balance exploration of the surface with local convergence (commonly ≈ 0.8). If $\mathcal{L}(\vec{A}_{\text{new}}) < \mathcal{L}(\vec{A}_x)$, it replaces $\vec{A}_x$ in the population. The scaled difference vector both codifies useful surface geometry information and shrinks as candidates get closer to one another, leading DE to efficiently locate minima on a high-dimensional nonlinear surface [8].

3.3 Local Search: L-BFGS-B

We apply L-BFGS-B, a gradient-based local optimization algorithm, to refine the DE minimum. It seeks to efficiently navigate the local curvature of the loss surface by approximating the Hessian matrix. We apply this method to several perturbations of the DE output, using relaxed priors to avoid overfitting to θ_0 . We estimate the allometric hyperparameters from the least loss $\mathbf{A}$ by rearranging the relationships given in Section 2.2, averaged across all interactions of a given type.

4 Results

Each of our results is a plot of recovery error versus some property of the ecosystem or the optimizer. This recovery error is defined as the magnitude of the percent error between each matrix element recovered by our method and the corresponding element of the synthetically-generated true matrix, averaged across the entire matrix. This is regularized by a denominator of the form $A_{ij}^{\text{true}} + \epsilon$, where $\epsilon = 0.01$, in order to avoid excessive penalization when the true value is near-zero.

$$\text{Recovery Error} = \sum_{ij} \frac{|A_{ij}^{\text{true}} - A_{ij}^{\text{recovered}}|}{A_{ij}^{\text{true}} + \epsilon}$$

For each parameter value tested, figures report the median, the interquartile range (25th–75th percentile), and a broader distribution range extending from the 10th to 90th percentiles.

4.1 Recovery Error vs. Ecosystem Size

We performed multiple experiments for each $N \in \{2, 3, 4, 5, 6\}$, running the optimization algorithm on different synthetic interaction matrices. The result is a dataset containing recovery errors for each matrix, for each N .

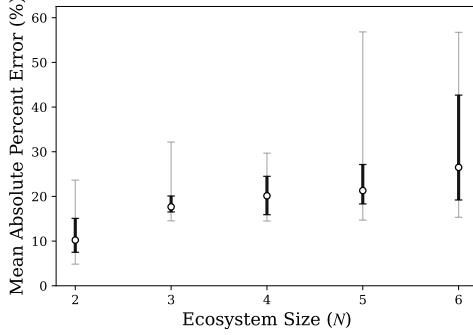


Figure 2: Recovery error as a function of ecosystem size, N . Twenty independent matrices were recovered per N . Error bars show interquartile range in bold, with fainter 10th–90th percentile range.

4.2 Recovery Error vs. Deviation from Allometric Laws

Next we investigate algorithm performance as the synthetic matrix deviates from theoretical allometric behavior. We sweep f from 0 to 0.5, where f scales the standard deviation of the normal distribution from which elements of $\mathbf{A}$ are sampled, as discussed in Section 3.1.1.

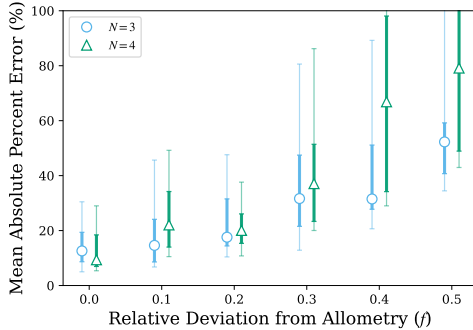


Figure 3: Recovery error as a function of relative deviation f from allometric laws. Twenty-five independent matrices were optimized per f . Y-axis capped at 100% for visual clarity. See Section 4.4 for performance details. Error bars show interquartile range in bold, with fainter 10th–90th percentile range.

4.3 Recovery Error vs. Equilibrium Noise

Finally, we test algorithm performance as elements of the noise matrix σ grow in size relative to the magnitude of elements of $\mathbf{A}$. That is, test g from 0.05 to 0.5, where g scales the noise in realizations of $\vec{y}$. This step requires a minor adjustment to the optimization pipeline as shown in Figure 1: after updating the priors, we perform a sweep over g , allowing the optimizer to select the value that minimizes the loss.

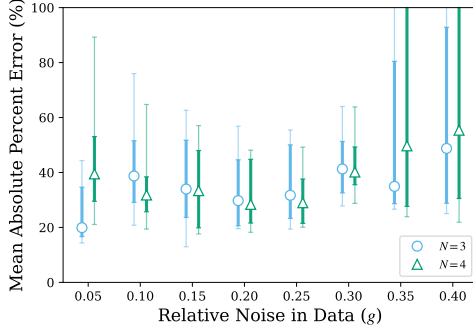


Figure 4: Recovery error as a function of relative noise g in observed equilibria $\vec{y}$. Twenty independent matrices were optimized per g . Y-axis capped at 100% for visual clarity. See Section 4.4 for performance details. Error bars show interquartile range in bold, with fainter 10th–90th percentile range.

4.4 Discussion

Our method consistently achieves recovery errors at or below 20% for small ecosystems ($N \leq 4$), with both accuracy and consistency declining as system size increases. Incorporating allometric priors provides strong regularization: when priors are reasonably accurate ($f < 0.3$), recovery errors remain at or under 40% with low variability. Low noise in equilibrium measurements ($g \leq 0.3$) also results in lower recovery error, similarly remaining below 40%. However, the interquartile range is wider in this case, indicating reduced reliability. We attribute this inconsistency to the additional step of inferring g during the optimization. Beyond the data shown in the axis-limited figures, the worst-case trials sweeping over f had recovery errors greater than 150% for $f = 0.4$ and $f = 0.5$. For the experiments over g , worst-case recovery errors exceeded 200%.

A key limitation of this method is its reliance on both repeated equilibrium measurements and prior knowledge of interaction signs. The former are difficult to come by experimentally, while the latter infringes upon our initial goal (ii) of making minimal assumptions. Similarly, while the initial guesses for allometric hyperparameters can be estimated from experiments or ecological literature, we currently lack a principled method to derive or calibrate them directly from data.

For future work, we propose jointly optimizing the negative log posterior over both the interaction matrix and the allometric hyperparameters, $\mathcal{L}(\mathbf{A}, \theta)$, rather than treating the two as separate sub-problems. Incorporating posterior sampling via MCMC could also help to quantify uncertainty in recovered matrices. Finally, inferring the interaction signs, rather than taking them as known, would increase the method’s applicability to systems with sparser experimental data.

5 Conclusion

This work demonstrates that species interaction networks can be reliably inferred from noisy equilibrium data so long as ecological information is properly integrated into the optimization process. By combining stochastic Lotka–Volterra modeling with allometric priors in a Bayesian MAP framework, we infer interaction matrices with high accuracy in low-dimensional systems, even under moderate observational noise. Across a range of controlled experiments, we found that recovery error remained below 20% for small ecosystems ($N \leq 4$), improved significantly when prior structure aligned with true interactions ($f \leq 0.3$), and remained stable for noise levels up to $g = 0.3$. These results demonstrate the potential of steady-state data, when paired with biologically informed constraints, to serve as a reliable foundation for ecological inference, with potentially wide-ranging benefits in conservation and field ecology.

References

- [1] J. T. Wootton and M. Emmerson, *Annual Review of Ecology, Evolution, and Systematics* **36**, 419 (2005).
- [2] F. Carrara, A. Giometto, M. Seymour, A. Rinaldo, and F. Altermatt, *Methods in Ecology and Evolution* **6**, 895 (2015).
- [3] G. B. West, J. H. Brown, and B. J. Enquist, *Science* **276**, 122 (1997).
- [4] J. Damuth, *Nature* **290**, 699 (1981).
- [5] B. C. Rall, U. Brose, M. Hartvig, G. Kalinkat, F. Schwarzmüller, O. Vucic-Pestic, and O. L. Petchey, *Philosophical Transactions of the Royal Society B* **367**, 2923 (2012).
- [6] J. C. D. Terry, M. B. Bonsall, and R. J. Morris, *Oikos* **129**, 147 (2019).
- [7] R. M. May, *Nature* **238**, 413 (1972).
- [8] R. Storn and K. Price, *Journal of Global Optimization* **11**, 341 (1997).

Preparation and Characterization of Lanthanide Oxides through Solution Combustion Synthesis (SCS) for Nuclear Science Applications

MARIEMIKA THEGENUS

2025 NSF/REU Program
Department of Physics and Astronomy
University of Notre Dame

ADVISOR(S): Prof. Ani Aprahamian & Prof. Leshar

Abstract

Lanthanide oxides are utilized in nuclear science applications. This paper explores the preparation of lanthanide oxides using solution combustion synthesis (SCS), a process that synthesizes fine powders. In this process, metal nitrates are the oxidizers mixed with fuels and solvents such as acetylacetone, ethanol, and 2-methoxyethanol and ignited under controlled conditions, with an close attention on the process and parameters used. The attention that is placed on understanding the role of different parameters, including the fuel-to-oxidizer ratio, different solvents, temperature, and the purity of the product. The resulting reaction is an oxide powder. The synthesized materials were characterized using techniques such as X-ray diffraction (XRD). The results indicate successful synthesis of pure lanthanide oxides with phase composition and crystal structure. This work highlights SCS for lanthanide-based materials.

1 Introduction

This study focuses on the preparation and characterization of lanthanide oxides, specifically dysprosium (Dy), samarium (Sm), europium (Eu), cerium (Ce), and neodymium (Nd) via SCS. These lanthanides are commonly used as surrogates for highly radioactive actinide elements such as californium (Cf), plutonium (Pu), americium (Am), thorium (Th), and uranium (U) because of their similar oxidation states, making them useful for research of nuclear fuel and waste without the complications of contamination. This work is part of a larger project to develop innovative SCS methods for preparing actinide oxides for nuclear science measurements.

Over a 9 week period, more than 100 combustion reactions were performed to evaluate the effects of solvent type and fuel to oxidizer ratio (ϕ) on combustion factors and the final product. Ethanol and 2-methoxyethanol were tested as solvents. Acetylacetone was tested as the fuel. In ethanol based solutions, combustion events in the range of $\phi = 0.5\text{--}1.3$ rarely produced a flame, just gas release. After Jordan's paper was reviewed, the study shifted to 2-methoxyethanol. Therefore, reactions involving 2-methoxyethanol consistently produced intense flames.

To understand the thermal characteristics of these reactions, thermocouples were used to mea-

sure both the solution and flame temperatures. A second thermocouple was added to capture flame temperatures better, but measuring the flame temperature proved to be difficult due to the nature of the flame. Combustion events were visually recorded.

Actinides such as americium, plutonium, uranium, thorium, and californium showed issues for experimental work because of their radioactivity and cost. Lanthanides such as samarium (Sm) and europium (Eu) offer a safer and more cost-effective alternative for studying SCS and thermo behavior.

This work emphasizes the role of solvent choice and stoichiometric ratio in determining the combustion behavior and the overall product of the SCS process. The synthesized powders were characterized using X-ray diffraction (XRD) to assess their phase purity and crystal structure. The results provide insights into improving SCS parameters for the valid synthesis of lanthanide oxide surrogates important to nuclear materials research.

Here is how to calculate the correct amount of acetylacetone with different phi ratio.

Let:

- $[N]$ = metal nitrate
- $[F]$ = fuel
- ϕ = fuel-to-oxidizer molar ratio
- $\rho_{[F]}$ = density of the fuel in g mL^{-1}

$$\text{Volume of } [F] \text{ (mL)} = \text{Mass of } [N] \text{ (g)} \left(\frac{1 \text{ mol of } [N]}{\text{Molar mass of } [N] \text{ (g)}} \right) \times \left(\frac{\phi \text{ mol } [F]}{1 \text{ mol } [N]} \right) \times \left(\frac{\text{Molar mass of } [F]}{1 \text{ mol } [F]} \right) \times \left(\frac{1 \text{ mL}}{\rho_{[F]} \text{ (g mL}^{-1}\text{)}} \right) \quad (1)$$

This formula calculates the fuel volume needed to achieve a desired ratio ϕ , taking into account the fuel density.

2 Background

Lanthanides are essential in many nuclear science experiments as chemical surrogates for radioactive actinides. Their oxidation state and similar chemical behavior make them ideal surrogates for actinides under experimental conditions that do not require handling of radioactive material (RAM). Their similar oxidation states allow researchers to investigate combustion behavior and material characteristics.

Solution combustion synthesis (SCS) has emerged as a method for preparing oxide powders and thin films for nuclear science measurements. This project specifically investigated Dy, Sm, Eu, Ce, and Nd. SCS is an effective method for the synthesis of lanthanide oxides. It involves a reaction between a metal nitrate oxidizer and a fuel such as acetylacetone. Key parameters of the reaction include the fuel-to-oxidizer ratio (ϕ), the solvent combustion, and environmental conditions such as containment of the final product. Although ethanol is widely used, it does not perform well in combustion. 2-methoxyethanol has been shown in the literature to produce more consistent reactions.

Normally, actinide targets have been prepared by traditional methods such as vacuum evaporation and electro spraying. However, these methods often involve complex infrastructure, result in poor material utilization, and exhibit issues such as incompatibility with certain backings and stoichiometry. These different challenges have motivated scientists to explore alternative methods.

3 Methods

Lanthanide nitrate - Dy, Sm, Eu, Ce, and Nd were used as oxidizers. Each was dissolved in either ethanol or 2-methoxyethanol, with values ranging from 0.5 to 1.3. Reactions were carried out in ceramic vessel placed on a hot plate.

Temperature measurements were recorded using:

- **Bottom thermocouple:** tracked the solution temperature before and after ignition.
- **Top thermocouple:** added later to attempt capturing flame temperature.

Visual observations were documented with videos and photos. After combustion, powders were characterized using X-ray diffraction (XRD). Diffraction patterns were analyzed with DIFFRAC.EVA to compare the final product to database references. Below are images of combustion and the XRD patterns. The primary goal was to investigate the effect of solvent and fuel-to-oxidizer ratio on



(a) Gas release of samarium oxide with ethanol as a solvent (b) Full combustion of samarium oxide with 2-methoxyethanol as a solvent

Figure 1: Two combustion of samarium oxide with a phi ratio of 0.625

combustion efficiency and product formation using SCS. The organic fuel used in all reactions was acetylacetone. Solvents tested included ethanol and 2-methoxyethanol.

The study began with 0.5 mL solution batches but switched midway to 1 mL batches for all lanthanides to improve product yield. The mixed solutions were placed in quartz combustion vessels and positioned on a hot plate set to 425°C. Combustion initiation times were recorded: ethanol-based reactions ignited after approximately 2–3 minutes, while those with 2-methoxyethanol required 6–7 minutes before ignition. After each reaction, the vessels were allowed to cool to room temperature. The resulting oxide product was then scraped from the quartz and transferred to vials for storage and analysis.

4 Results

A solvent mixture study was conducted using samarium nitrate hexahydrate to evaluate how varying ratios of ethanol and 2-methoxyethanol affect combustion behavior. Many solvent compositions were tested: 95:5, 80:20, 40:60, and 2:98 (2-methoxyethanol:ethanol, vis-versa). Surprisingly, even with only 2% 2-methoxyethanol and 98% ethanol, a highly visible combustion occurred, demonstrating that very small amounts of 2-methoxyethanol can enhance combustion efficiency. These results support the idea that solvent choice is a critical parameter in SCS.

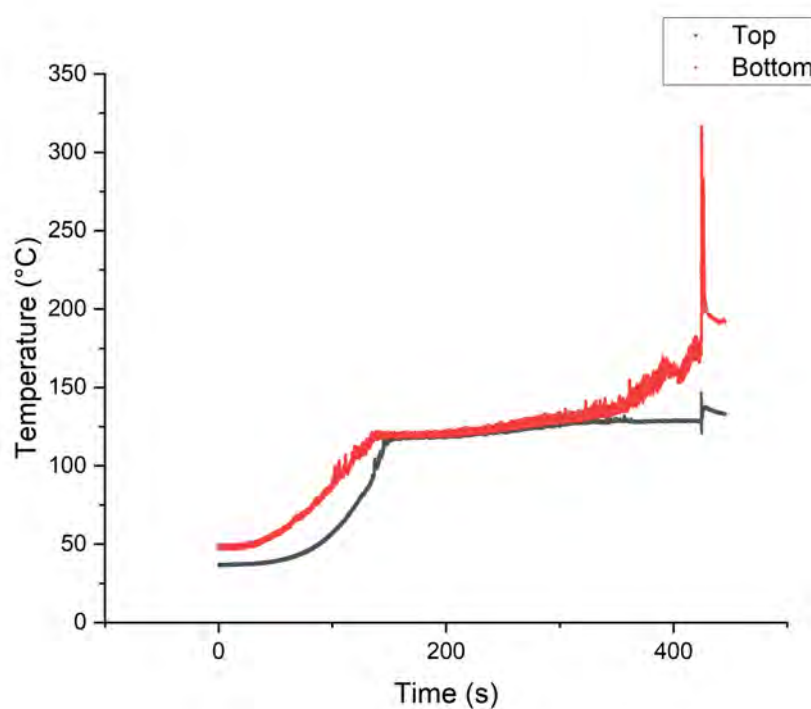


Figure 2: Time vs. Temperature of a samarium oxide combustion. The thermocouple caught the flame temperature.

The products of selected reactions were analyzed using X-ray diffraction (XRD). While the XRD patterns did not show 100% purity, a significant number of peaks aligned with reference patterns for the expected lanthanide oxide phases.

5 Conclusion

This study explored the synthesis of lanthanide oxides using solution combustion synthesis (SCS), with a focus on understanding the impact of solvent choice and fuel-to-oxidizer ratio (ϕ) on combustion efficiency and product quality. The lanthanides selected: Dy, Sm, Eu, Ce, and Nd, served as safe and chemically appropriate surrogates for radioactive actinides. Over 100 SCS reactions were performed using acetylacetone as the fuel and either ethanol or 2-methoxyethanol as the solvent.

Ethanol-based reactions typically produced only gas release across a wide range of ϕ values, while 2-methoxyethanol consistently produced complete combustion. A solvent mixture study with samarium nitrate revealed that even small amounts of 2-methoxyethanol (as little as 2%) drastically improved combustion, highlighting its catalytic behavior in SCS reactions. Thermal measurements and visual documentation provided insight into the combustion process, while XRD analysis confirmed that the resulting oxides were largely phase-pure, though not 100%.

These results demonstrate that SCS is an efficient method for producing lanthanide oxide powders with applications in nuclear materials research. The findings help lay the groundwork for applying similar techniques to the preparation of thin film targets composed of radioactive actinides such as Pu and Am.

References

M. R. Buchmeiser, M. G. Ehrlich, M. Zupančič, A. R. T. Nugroho, P. Erhart, S. Harder, and C. Scheurer, *The role of chemical bonding in structure formation: New perspectives from experimental and theoretical solid-state chemistry*, Journal of Solid State Chemistry, vol. 324, 124258, 2023. <https://doi.org/10.1016/j.jssc.2023.124258>

P. Jain, N. S. Dalal, B. H. Toby, H. W. Kroto, and A. K. Cheetham, *Recent Advances in the Chemistry of Metal–Organic Frameworks*, Chemical Reviews, vol. 116, no. 19, pp. 10512–10539, 2016. <https://doi.org/10.1021/acs.chemrev.6b00279>

A. R. West, *Powder Diffraction: An Overview*, Course notes, University of Trieste, accessed July 2025.

https://moodle2.units.it/pluginfile.php/374253/mod_resource/content/1/Powder-3.pdf

Cabanas, J. M., Epping, B., Qiao, J., Zhang, X., Smith, M. R., Wood, M. D., et al. (2025). *ThO₂ and Th_{1-x}UxO₂ nanoscale materials and thin films for nuclear science applications*. Journal of Nuclear Materials, 585, 154480.

<https://doi.org/10.1016/j.jnucmat.2024.154480>

Automating IMSRG(3) expression generation with Python

VICTOR VAIDA

2025 NSF/REU Program
Department of Physics and Astronomy
University of Notre Dame

ADVISOR(S): Prof. Ragnar S. Stroberg

Abstract

IMSRG is a powerful theoretical method for studying medium-mass nuclei. However, manually deriving the expressions needed for IMSRG calculations is time-consuming and error-prone. In this work, we develop a fully automated Python framework to generate and evaluate IMSRG expressions. We benchmark our code against the existing IMSRG codebase with promising accuracy. This project allows for the future exploration of new nested commutator structures. Ultimately, we hope to improve the accuracy of existing nuclear calculations.

1 Introduction

In nuclear physics, theoretical models are essential for confirming and predicting results. When theory and experiment misalign, it can hint at new physics beyond the Standard Model. We also need theoretical nuclear predictions for systems we can't access experimentally (e.g., neutrinoless double-beta decay, neutron-rich stars). However, developing theoretical machinery is tedious. For methods like coupled cluster [1] or the in-medium similarity renormalization group, it can take weeks to draw diagrams, months to write the code, and days to compute a result.

In this paper, we focus on automating the IMSRG expression generation process. Our main motivation is to explore specific nested three-body commutators, which would otherwise require writing hundreds of separate C++ routines by hand. We build one general routine using Python.

This paper is organized as follows. Section 2 provides IMSRG background and discusses notation. Section 3 outlines the rules behind making IMSRG diagrams and expressions. Section 4 benchmarks the accuracy of this work's code against the existing `imsrg` code.

2 Background

The IMSRG (in-medium similarity renormalization group) is an *ab initio* method for nuclear physics. *Ab initio* methods predict nuclear properties by considering the fundamental interactions between the nucleons, rather than using empirical data from the nucleus being studied.

2.1 Schrödinger equation

The overall goal for IMSRG is to solve the many-body Schrödinger equation for a nucleus.

$$\hat{H}\Psi = E\Psi \quad (1)$$

For almost all nuclei, it is impossible to get an exact solution to the Schrödinger equation. To approximate solutions, IMSRG transforms the Hamiltonian H into a more manageable form.

First, we perform a unitary transformation on H . For IMSRG, it is convenient to choose the unitary operator $U = e^\Omega$, where $\Omega = -\Omega^\dagger$ is the anti-Hermitian Magnus operator [2]. After the transformation, the resulting Hamiltonian $\tilde{H}$ is more block-diagonal than H , which makes the Schrödinger equation easier to solve [3].

$$\tilde{H} := U H U^\dagger = e^\Omega H e^{-\Omega} \quad (2)$$

By choosing our unitary operator to be an operator exponential ($U = e^\Omega$), we can conveniently employ the Baker-Campbell-Hausdorff (BCH) expansion [3]. This turns our unitary transformation into an infinite sum of nested commutators. Our problem now reduces to evaluating these nested commutators.

$$\tilde{H} = H + [\Omega, H] + \frac{1}{2!}[\Omega, [\Omega, H]] + \frac{1}{3!}[\Omega, [\Omega, [\Omega, H]]] + \dots \quad (3)$$

2.2 Commutator notation

The commutator between two operators is defined as $[\hat{B}, \hat{A}] = \hat{B}\hat{A} - \hat{A}\hat{B}$. A nested commutator like $[\hat{C}, [\hat{B}, \hat{A}]]$ is expanded from the inside out:

$$[\hat{C}, [\hat{B}, \hat{A}]] = \hat{C}\hat{B}\hat{A} - \hat{C}\hat{A}\hat{B} - \hat{B}\hat{A}\hat{C} + \hat{A}\hat{B}\hat{C} \quad (4)$$

In this paper, we often talk about the " n -body" part of a given operator $\hat{A}$. $\hat{A}_0$ is the zero-body

part (a scalar), $\hat{A}_1$ is the one-body part (acts on single particles), $\hat{A}_2$ is the two-body part (acts on pairs of particles), and so on.

$$\hat{A} = \hat{A}_0 + \hat{A}_1 + \hat{A}_2 + \dots \quad (5)$$

In addition to *full* operators, we can also take a commutator of the *n-body parts* of operators. Writing $[\hat{B}_3, \hat{A}_1]$ denotes the commutator of the three-body part of $\hat{B}$ with the one-body part of $\hat{A}$. Adding a subscript 2 to this commutator makes this $[\hat{B}_3, \hat{A}_1]_2$, meaning we compute the commutator but extract only its two-body part.

In this paper, we write such commutators $[\hat{B}_3, \hat{A}_1]_2$ as $[3, 1]_2$, for brevity. A triple-nested commutator like $[\hat{D}_2, [\hat{C}_3, [\hat{B}_2, \hat{A}_2]_3]_4]_2$ turns into $[2, [3, [2, 2]_3]_4]_2$.

2.3 Current IMSRG scope

The standard framework is IMSRG(2), which only includes the nested commutators that involve two-body operators at every intermediate step. This leaves out many three-body terms (e.g., $[2, 2]_3$, $[2, [2, 3]_3]_2$) from its calculations, leading to error in calculations. IMSRG(3) includes all the missing three-body terms, but many are computationally intensive and have limited impact on the final result. By automating the IMSRG expression generation process, we can explore which higher-body commutators are worth computing.

3 Diagram and expression theory

This section provides an overview of how to evaluate nested commutators.

3.1 Commutator $\rightarrow$ Diagram

We will use the commutator $[2, [2, [3, 1]_2]_2]_2$ as an example. We first draw four operator nodes labeled $\hat{D}, \hat{C}, \hat{B}, \hat{A}$, top to bottom. Each node has incoming and outgoing fermion lines, which represent the particles the operator acts on. The operator rank determines the number of fermion lines. For example, since $\hat{B}_3$ has rank 3, node $\hat{B}$ gets 3 incoming and 3 outgoing fermion lines.

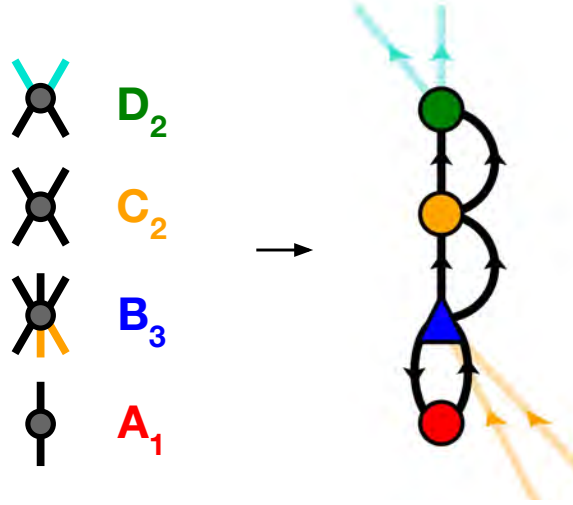


Figure 1: Making one possible Hugenholtz diagram for $[2, [2, [3, 1]_2]_2]_2$.

We connect the operators via fermion lines from the inside-out, using the nested commutator as a guide. Our commutator is in the form $[D, [C, [B, A]_p]_q]_r$. The innermost commutator is $[B, A]_p = [3, 1]_2$. We connect these two operators in such a way that their resulting $\hat{B}\hat{A}$ operator has $p = 2$ lines that come in and out. Then we draw the next innermost commutator, $[C, p]_q = [2, 2]_2$. We connect node $\hat{C}$ and the $\hat{B}\hat{A}$ node combination to form a two-body operator. The last step is to draw the commutator $[D, q]_r = [2, 2]_2$, where we connect $\hat{D}$ to the $\hat{C}\hat{B}\hat{A}$ combination in a way that forms a final two-body operator. The resulting diagram is a "Hugenholtz diagram," abbreviated as "HH diagram" in this paper.

An HH diagram is a directed graph, so it can be uniquely represented by an adjacency matrix [4]. Our code automatically generates HH diagrams by generating adjacency matrices based on a set of rules. We can use `matplotlib` and `feynman`¹ to convert the matrix to a diagram (like Fig. 1).

In general, there is usually more than one way to turn a commutator into an HH diagram. Luckily, it is feasible to draw every diagram by hand for single or double-nested commutators (e.g., $[2, 2]_2$ or $[2, [2, 2]_3]_2$). However, for triple-nested commutators like $[2, [2, [2, 2]_3]_3]_2$, there can be hundreds of diagrams per commutator. Thankfully, the code developed in this paper automatically generates these triple-nested HH diagrams within seconds.

¹<https://gkantoniuss.github.io/feynman/index.html>

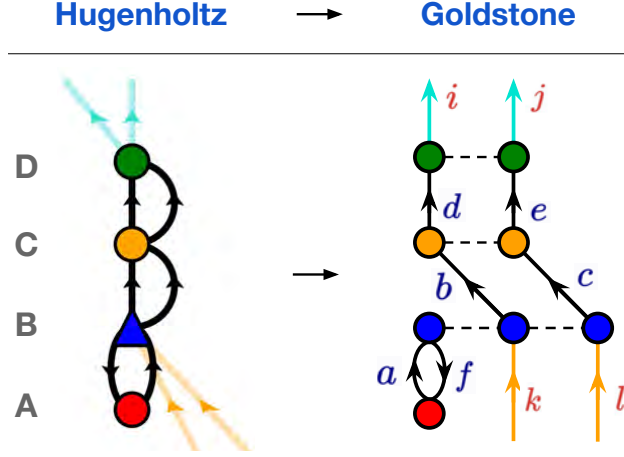


Figure 2: Hugenholtz (HH) diagram turning into a Goldstone (GS) diagram

3.2 Diagram → Expression

The next step is to convert these Hugenholtz diagrams to algebraic expressions. (These expressions are what HH diagrams symbolically represent.)

We first must turn the Hugenholtz (HH) diagram into a Goldstone (GS) diagram (shown in Fig. 2). GS nodes only have 1 incoming and outgoing fermion line. For each HH node, we split it into multiple GS nodes and connect them with a dashed line. We then assign a letter to each fermion line. Traditionally, we label external lines with i, j, k, l , and internal lines with a, b, c, d , etc.

$$\underbrace{(-1)^{\ell+h}}_{\text{Sign}} \underbrace{\frac{1}{(2!)^4}}_{\text{Equivalency factor}} \sum_{abcdef} \underbrace{(1 - P_{ij})(1 - P_{kl})}_{\text{Permutation factors}} \underbrace{(\bar{n}_a \bar{n}_b \bar{n}_c \bar{n}_d \bar{n}_e n_f)}_{\text{Occupation factors}} \underbrace{D_{ijde} C_{debc} B_{fbcakl} A_{af}}_{\text{Operator matrix elements}}$$

Figure 3: Annotated expression of the GS diagram in Fig. 2

Now we build the expression (shown in Fig. 3). We put a matrix element (e.g., C_{debc}) for each operator ($\hat{C}$), where the leftmost indices are outgoing lines (d, e) and the rightmost indices are incoming lines (b, c). We sum over all internal line indices ($\sum_{abcdef}$), meaning only external line indices ($ijkl$) show up in the result tensor. In this paper, we only deal with one or two-body resulting tensors (T_{ij} or T_{ijkl}).

All lines that go up are "particle lines," and all lines that go down are "hole lines." For each internal line, we include an occupation factor. The occupation factor is either n_a if the line a is a hole, or $\bar{n}_a$ if a is a particle, where $\bar{n}_a = 1 - n_a$.

If the resulting tensor is a two-body operator, we must include the term $(1 - P_{ij})(1 - P_{kl})$ ² in the expression to preserve anti-symmetry [5]. P_{ij} is a permutation operator that swaps indices i and j in the expression to its right. [2]

If there are lines connecting the same two HH nodes in the same direction, we call them "equivalent" [5]. Each group of n equivalent lines contributes a symmetry factor of $1/n!$. For example, a group of 3 equivalent lines and a group of 2 equivalent lines contribute a factor of $\frac{1}{2! \cdot 3!}$.

The sign of the expression is $(-1)^{h+\ell}$ where h is the number of hole lines and ℓ is the number of loops in the diagram. A loop is a path that starts and ends at the same GS node.³

At this point, our expression (Fig. 3) only corresponds to the $\hat{D}\hat{C}\hat{B}\hat{A}$ piece of the full commutator $[\hat{D}, [\hat{C}, [\hat{B}, \hat{A}]]]$, which when expanded yields seven other terms:

$$\hat{D}\hat{C}\hat{B}\hat{A} - \hat{D}\hat{C}\hat{A}\hat{B} - \hat{D}\hat{B}\hat{A}\hat{C} + \hat{D}\hat{A}\hat{B}\hat{C} - \hat{C}\hat{B}\hat{A}\hat{D} + \hat{C}\hat{A}\hat{B}\hat{D} + \hat{B}\hat{A}\hat{C}\hat{D} - \hat{A}\hat{B}\hat{C}\hat{D}. \quad (6)$$

Other works [2] account for this by adding duplicate expressions with operator labels swapped, which re-indexes the matrix elements (e.g., $B_{ia}A_{aj} \rightarrow A_{ia}B_{aj}$). In this paper, we kept the operator labels fixed and repositioned the operators diagrammatically (e.g., sliding $\hat{A}$ above $\hat{B}$ without detaching lines). This changes the occupation factors instead of the operator indices (particle lines become holes).⁴

Our complete expression now consists of 2–8 "mini-expressions." We use Python's `einsum`⁵ to evaluate each mini-expression, contracting it to a numpy tensor. Adding all the mini-einsum tensors gives the numerical result corresponding to our original HH diagram.

²This is because $(1 - P_{ij})(1 - P_{kl})$ expands to $1 - P_{ij} - P_{kl} + P_{ij}P_{kl}$, which represents all four permutations: swap none, swap ij , swap kl , or swap both.

³Previously, the sign also depended on fermion line crossings, which was difficult to implement in code. However, we found a simpler rule that only depends on holes and loops. We connect our outgoing external lines to our internal external lines with GS nodes. In this closed diagram, adding a crossing now corresponds to subtracting a loop.

⁴While both methods work, admittedly the first method may have been simpler in hindsight.

⁵We use `einsum`'s 'greedy' setting.

Previously, we would use Prof. Ragnar Stroberg’s C++-based IMSRG(3) codebase (named `imsrg`, found in ref. [6]) to evaluate these expressions. While C++ offers faster runtime performance than Python, every diagram demanded its own hand-coded contraction routine, which was labor-intensive. With our method, every expression is evaluated with the same framework (`einsum`), making it significantly easier to explore new commutator topologies.

4 Results

In this section, we benchmark this code against the well-established `imsrg` codebase.

Commutator	$e_{\max} = 1$	$e_{\max} = 2$	$e_{\max} = 3$	$e_{\max} = 4$
$[1, 1]_0$	0 (0.00 s)	✓ (0.01 s)	✓ (0.01 s)	✓ (0.03 s)
$[2, 2]_0$	✓ (0.00 s)	✓ (5.04 s)	✓ (1.14 min)	✓ (10.34 min)
$[1, 1]_1$	0 (0.00 s)	✓ (0.00 s)	✓ (0.01 s)	✓ (0.04 s)
$[1, 2]_1$	0 (0.13 s)	✓ (4.94 s)	✓ (1.09 min)	✓ (10.27 min)
$[2, 2]_1$	✓ (0.13 s)	✓ (4.59 s)	✓ (1.38 min)	✓ (17.04 min)
$[1, 2]_2$	✓ (0.19 s)	✓ (6.70 s)	✓ (2.50 min)	✓ (18.24 min)
$[2, 2]_2$	✓ (0.20 s)	✓ (7.44 s)	✓ (1.78 min)	✓ (46.96 min)
$[3, 3]_0$	0 (36.29 s)	—	—	—
$[2, 3]_1$	✓ (36.35 s)	—	—	—
$[3, 3]_1$	✓ (38.62 s)	—	—	—
$[1, 3]_2$	0 (38.32 s)	—	—	—
$[2, 3]_2$	✓ (45.98 s)	—	—	—
$[3, 3]_2$	× (2.36 min)	—	—	—

Table 1: Single commutator verification benchmark, results and runtimes. Checking agreement between this paper’s code and the `imsrg` code across several $e_{\max}$ values. ✓ means matching results, and × means non-matching results. 0 means the norms of both results are zero.

Table 1 above shows the list of all single commutators $[B, A]_p$ that the `imsrg` code can evaluate. To compare our code with the `imsrg` code, we first generate two random operators $\hat{B}$ and $\hat{A}$ ⁶. We then use both our code and the `imsrg` code to evaluate each commutator $[B, A]_p$, and we compare the two resulting tensors element-wise (up to a certain tolerance).

$e_{\max}$ is a cutoff that controls how many single-particle states we include in the calculation. Increasing $e_{\max}$ makes the calculation more accurate but also more expensive. Due to computational limitations, we are unable to run three-body operator calculations at $e_{\max} = 2$. We chose to stop at $e_{\max} = 4$ because of computation time.

The $[3, 3]_2$ commutator test in Table 1 failed, indicating that further debugging is required. However, it is promising that all other tests pass up to $e_{\max} = 4$. It is unfortunate we cannot evaluate the three-body commutators beyond $e_{\max} = 1$, because the failure with $[3, 3]_2$ suggests a possibility of a bug involving three-body operators. We require $e_{\max}$ levels higher than 1 to investigate this.

It is clear that the runtimes of these tests increase exponentially with $e_{\max}$. We also observe a strong dependence on operator rank: one-body tests run almost instantly, two-body tests are moderately expensive, and three-body tests are significantly more costly.

Commutator	$e_{\max} = 1$	$e_{\max} = 2$	$e_{\max} = 3$
$[2, [2, [2, 2]_2]_2]_2$	✓ (2.16 s)	✓ (3.89 min)	✓ (3.97 hr)

Table 2: Benchmark results and runtimes for $[2, [2, [2, 2]_2]_2]_2$ across several $e_{\max}$ values.

We have not conducted a *complete* analysis on double or triple-nested commutators, but Table 2 shows a small comparison test for the triple-nested commutator $[2, [2, [2, 2]_2]_2]_2$. While this commutator is a simple example, it is encouraging that our expressions match `imsrg` for multiple $e_{\max}$ values. However, the 4-hour runtime for $e_{\max} = 3$ further emphasizes the need for optimization.

Once we have confidence in this code’s accuracy in evaluating single and double-nested commutators, we will conduct a thorough analysis on triple-nested commutators, which is a novel area. We are most interested in the commutators: $[2, [2, [2, 2]_3]_3]_1$, $[2, [2, [2, 2]_3]_3]_2$, and $[2, [2, [2, 2]_3]_4]_2$.

⁶Both operators are anti-symmetric. We choose $\hat{B}$ to be anti-Hermitian and $\hat{A}$ to be Hermitian. This is to mirror $[\Omega, H]$, where Ω is anti-Hermitian and H is Hermitian. The random seed is set to 1, which is the default.

5 Conclusion

We find that our implementation is accurate for all commutators in the `imsrg` code up to $e_{\max} = 4$, except for the commutator $[3, 3]_2$ which fails at $e_{\max} = 1$. We also find our implementation is accurate for the triple-nested commutator $[2, [2, [2, 2]_2]_2]_2$ up to $e_{\max} = 3$.

The main priority for future work is to optimize this code for higher $e_{\max}$ levels. There are many ways to do this. We can reduce the number of "mini-expressions" through algebraic simplification, which could reduce the computation time by 2-16 times. We can also improve our contraction paths using packages like `opt_einsum`.

6 Acknowledgements

A big thanks to Prof. Ragnar Stroberg for his glowing support and financial assistance toward this project. Thank you also to Kristen Amsler and Prof. Umesh Garg for setting up an amazing REU.

References

- [1] J. Brandejs, J. Pototschnig, and T. Saue, arXiv preprint arXiv:2409.06759 (2025).
- [2] B. He and S. Stroberg, arXiv preprint arXiv:2405.19594 (2024).
- [3] Wikipedia contributors, Unitary transformation (quantum mechanics) (2025), accessed: 2025-07-30.
- [4] Wikipedia contributors, Adjacency matrix (2025), accessed: 2025-07-30.
- [5] D. Sahan, Diagrammatic techniques for many-body quantum theory (2010), accessed: 2025-07-30.
- [6] S. Stroberg, `imsrg++`, <https://github.com/ragnarstroberg/imsrg> (2018), accessed: 2025-07-30.

Ion Mobility Effects on Gas Catcher Simulation Efficiency

MATTHEW WATSON

2025 NSF/REU Program
Department of Physics and Astronomy
University of Notre Dame

ADVISOR(S): Prof. Maxime Brodeur, Fabio Rivero

Abstract

The Cabibbo-Kobayashi-Maskawa (CKM) quark mixing matrix is a component of the Standard Model which describes the weak interactions between quarks. If the Standard Model contains all existing quarks, then this matrix must be unitary. The goal of the Superaligned Transition Beta-Neutrino Decay Ion Coincidence Trap (St. Benedict) experiment is to improve the accuracy of the V_{ud} element of this matrix, or the coupling strength of the up and down quarks, to test the unitarity of the matrix. St. Benedict was constructed with multiple key components designed to rapidly stop, extract, and transport a low energy radioactive ion beam (RIB) to an ion trap for decay measurements. The first key component is a large-volume gas catcher. It uses ultra-high purity helium gas to stop the RIB, and a combination of RF and DC fields to extract the RIB species from the gas for low-energy delivery to the rest of St. Benedict. Preceding the gas catcher is TwinSol, a set of twin solenoids designed to focus the RIB into St. Benedict. Efficient transport is critical for performing faster measurements, and using ion optics simulations, lab time and resources can be saved by tuning ahead of experiment. Simulations in SimIon were performed to model the beamline and gas catcher to run efficiency tests for various parameters using an updated ion mobility value for high pressures to match offline commissioning results.

1 Introduction

The Standard Model (SM) describes the behavior and nature of the most fundamental particles and forces in the universe, including particles like quarks and the weak nuclear force. A cornerstone of the SM is the Cabibbo-Kobayashi-Maskawa (CKM) matrix which rotates the regular mass eigenstates of the quarks to their eigenstates under the weak interaction. Because this matrix represents a rotation of quark eigenstates it must be unitary, meaning its conjugate transpose multiplied by the matrix itself should yield the identity matrix. Recent theoretical corrections used in calculations of V_{ud} have yielded a value which disagrees with unitarity by 2.4σ [1]. In light of this, a more high-precision accurate determination of V_{ud} must be made in order to shed light on this discrepancy.

St. Benedict aims to determine the V_{ud} element from superallowed beta decay between mirror nuclei with the goal of shedding more light on the present tension with unitarity when the superallowed pure Fermi data set is used. St. Benedict consists of four main components: a gas catcher, the RF carpet and ion guide, the cooler-buncher, and the Paul trap. The gas catcher is designed to thermalize incoming RIBs and transport ions into the RF carpet and ion guide which guide ions

through a differentially pumped section of St. Benedict. Then, the RIB encounters the cooler-buncher which cools and bunches the ions before they enter into the Paul trap which keeps the ions in one place as they decay so that high precision measurements can be taken.

These RIBs are produced with the TwinSol magnetic recoil separator at the Nuclear Science Lab (NSL) at the University of Notre Dame. A stable ion-beam is accelerated to 20-100 MeV of energy, using the 10 MV FN Tandem Van de Graaf accelerator, and impinged on a gas target (typically deuterium). A “cocktail-beam” composed of various reaction products is then refocused through a small aperture which allows for selection of the species of interest by the magnetic rigidity, or charge-to-mass ratio. Then, the RIB passes through a variable angle aluminum degrader that takes a large part of the energy away before passing through an aluminum entrance window which accomplishes the same effect at a fixed thickness. The RIB then enters the gas catcher, the aforementioned first component of St. Benedict, where it encounters high pressure helium gas and loses energy via elastic collisions with the He atoms, causing it to quickly thermalize. The ions, now stopped, can be manipulated by applying RF and DC fields to several electrodes inside the gas catcher. A potential gradient across the body of the gas catcher creates an electric field that drags the thermalized ions through that first part of the gas catcher. This gradient is produced by applying two different potentials, "body high" and "body low" on the outermost electrode of that section. All intermediate electrodes are connected by resistors resulting in voltage divide that creates the potential gradient. Then, the cone high and cone low do the same as ions reach the end of the gas catcher, in the cone. A RF signal with opposite polarity on adjacent electrodes is applied to repel ions away from the walls. In the offline commissioning, a potassium ion source was used. See Fig. 1 for more details.

To get the best possible results, St. Benedict requires the best possible efficiency of ions from TwinSol making it all the way to the Paul trap. Simulations play a critical role in this. By having an accurate model of the beam through each part, optimizations can be made. Simulations using the SimIon ion optics software work best for running bulk ions and determining precise ion movement paths as well as adjusting electrode potentials to find the most efficient setup for transport. However,

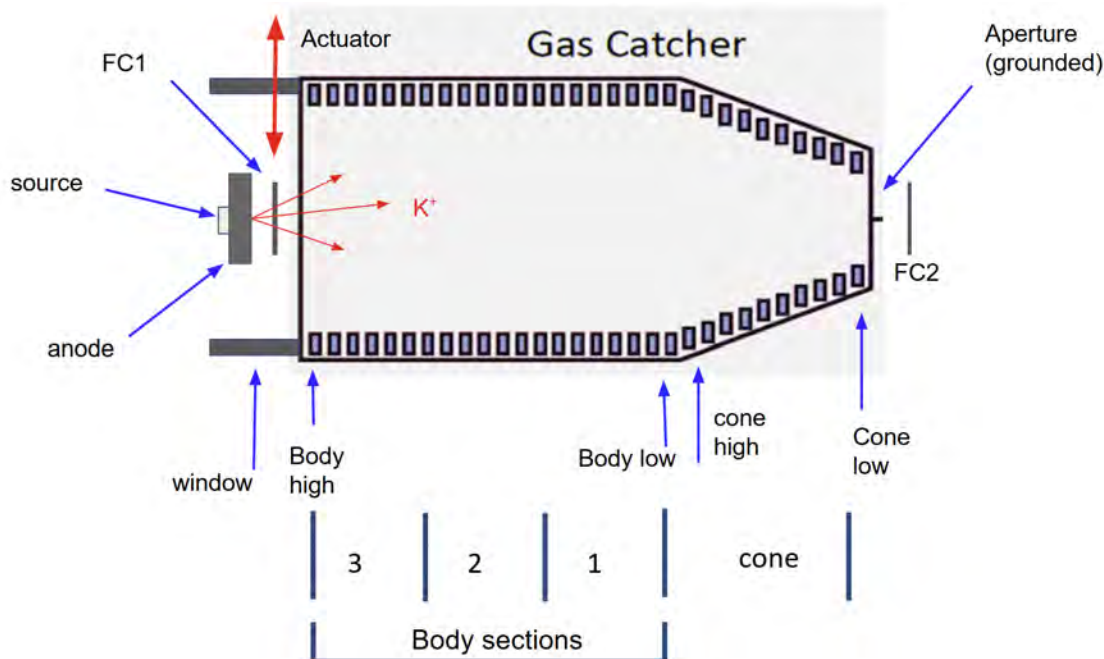


Figure 1: Diagram of the gas catcher for its offline commissioning.

running SimIon with a hard-sphere collision model at the relatively high pressures of 25 to 75 Torr is computationally expensive and so statistical models must be used instead. These rely on a value referred to as the ion mobility parameter, K . This value determines the drift velocity that ions will come to in a given electric field based on a number of variables. Primarily of interest here is checking and correcting for the ion mobility parameter which dictates the drift velocity of ions moving in a field at a specific temperature and pressure.

2 Methods

To accurately model the gas collisions at this pressure, a special lua code is called within SimIon to run statistical models on ions as they travel. This code requires a value for the ion mobility to be entered, and without it, it will make an attempt to guess at what this value is. This guess has proven inaccurate and fails to give results congruent with the offline commissioning, which used a thermionic potassium source. A new ion mobility was used from previously done ab initio calculations, yielding a value of $21.24 \text{ cm}^2/\text{Vs}$ [2]. These results provide better dynamics, but to

test whether or not this value is accurately reproducing physical effects, a simpler SimIon geometry file was created that would allow testing of the time of flight to determine whether or not ions had the expected drift velocity for the given ion mobility. This geometry file consists of two 1 mm thick, 250 mm radius electrodes separated by a distance of 500 mm with a potential difference of 10 V. See Fig. 2. Potassium ions were spawned 10 mm away from the initial electrode to prevent gas collisions from slamming them into the wrong electrode. Spawned ions had no initial energy and all originated at the center of the plates. There was no bulk gas velocity component.

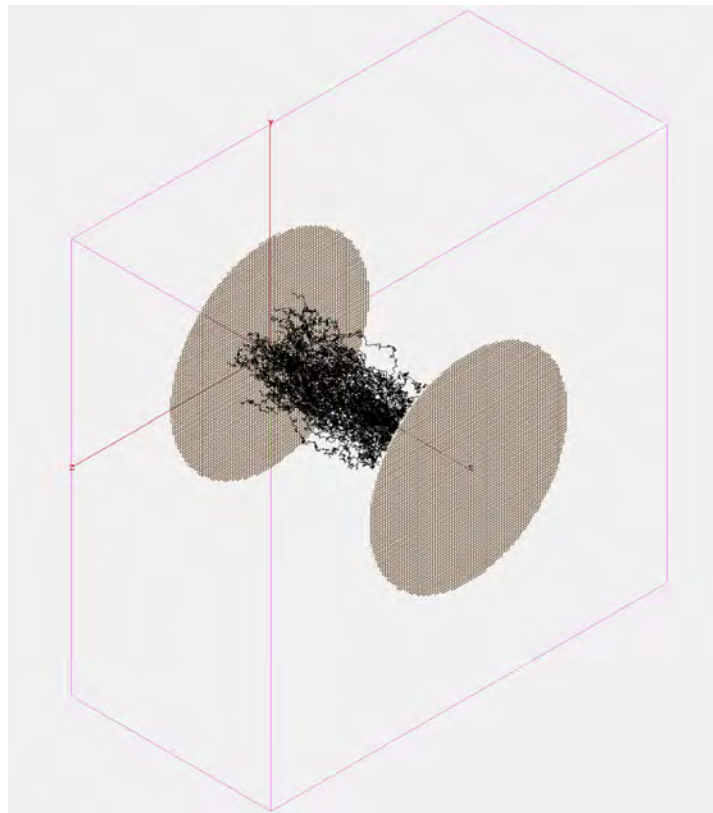


Figure 2: Isometric view of the electrodes in SimIon.

Timing and position data was recorded for 1000 ions and then processed using AWK, a linux based programming language. When the ions passed the halfway mark at 250 mm, the time of flight was recorded. This was to give the ions ample time to accelerate and reach a close to constant velocity. When the ions struck the end electrode, their time of flight was recorded again. The average velocity of each particle was found, and these velocities were averaged to give a drift velocity for

the potassium ions in a simple, controlled simulation environment under a constant electric field.

These velocities were then compared with theoretical outcomes dictated by the following equation where V_d is the drift velocity, E is the electric field, K_0 is the reduced ion mobility constant, and P is pressure in Torr.

$$V_d = E * K_0 * (760/P) \quad (1)$$

These theoretical results were calculated for three different pressure values, 33, 66, and 100 mbar. Then, the ion mobility parameter was adjusted in SimIon until the drift velocity as calculated by AWK matched the theoretical results. This adjusted ion mobility should help account for any discrepancies between theory and simulation.

With the new ion mobility, the gas catcher was run at the three previously mentioned pressures over a range of electrode potentials. The three electrodes in question were, as labelled, the RF, the body low (BL), and the body high (BH). The RF was run from a range of 0 V to 150 V, and the BH was run from 45 V to 110 V both with voltage steps of 5 V simulating 300 ions for each step. The values for the RF and BH when they were not being swept were 100 V and 80 V, respectively. The BL was kept at a constant 50 V. This value was subtracted from BH when plotting to visualize the BH-BL. Sweeping over the potential values, histograms were generated of the amount of particles that were successfully reaching a defined “transported zone” within the gas catcher where the RF carpet could catch them. Then, the transport efficiency was calculated for each voltage.

3 Results

For each of the pressure values, an ion mobility was found that gave an ion mobility which closely agreed with the theoretical expectations as well as the other pressure values. This new ion mobility was found to be 19.3 cm²V/s. Using this value, new efficiency versus voltage plots were made to compare to offline efficiency test results and the efficiency results using the old SimIon estimate for the ion mobility of 11.5 cm²V/s.

From the offline results, efficiency should sharply increase for the lower voltage steps before plateauing quickly. While a general increase in efficiency with voltage and slight concavity is observed in the old ion mobility simulation, it fails to adequately match what was seen in offline testing at all pressures. In the updated ion mobility simulation, the same thing can be observed. There is only a general increase that lacks agreement with experimental data. The BH-BL plot for the new ion mobility fits experiment worse than the old mobility did, losing some concavity that experiment at 100 mbar possessed.

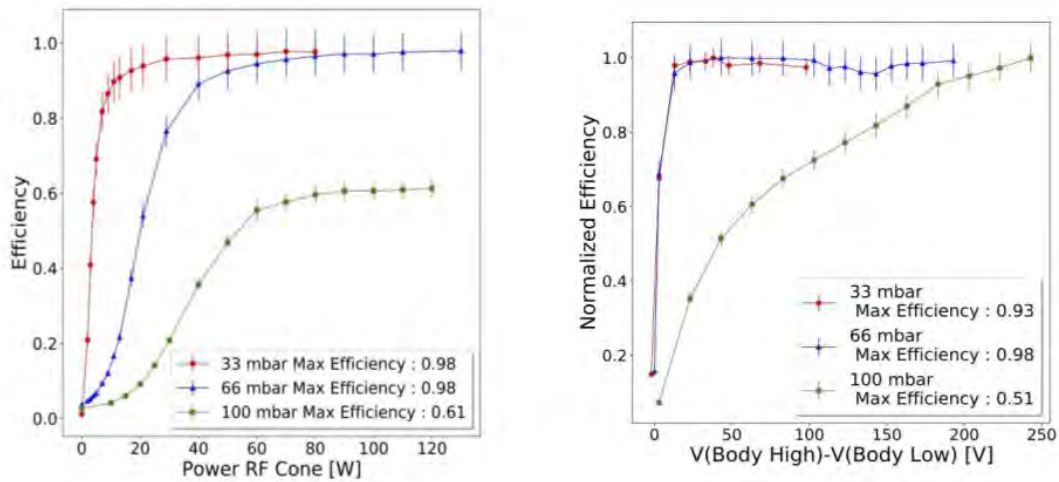


Figure 3: Offline commissioning efficiency results.

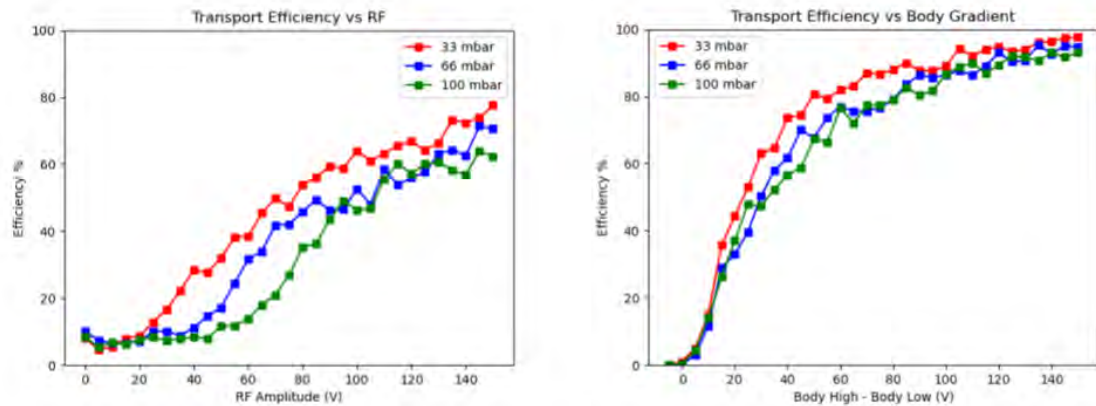


Figure 4: 11.5 cm²V/s simulation efficiency results.

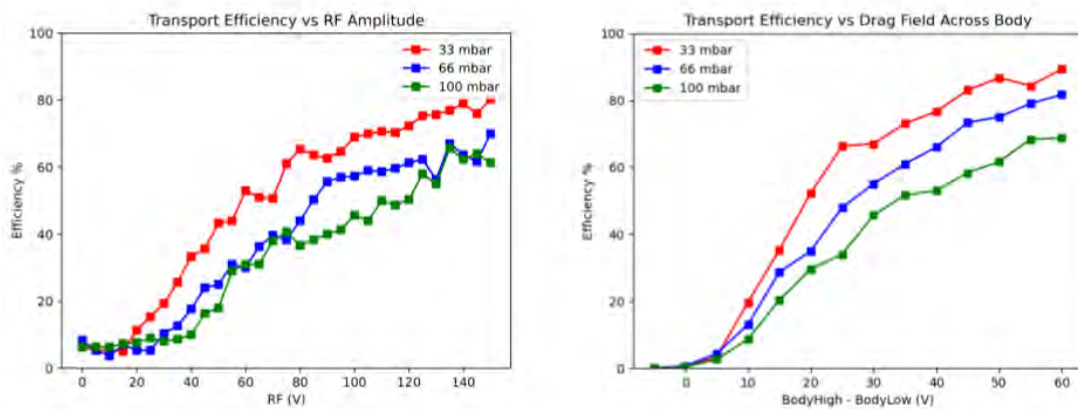


Figure 5: 19.3 cm²V/s simulation efficiency results.

4 Conclusion

The plots did not change significantly with the introduction of a more accurate ion mobility parameter. This indicates that the ion mobility is not a leading order factor in the gas catcher simulation. From here, more parameters can be investigated to look for leading order factors and make simulation plots as close as possible to the offline testing plots. Potential candidates for leading order factors include ion-ion collisions and gas flows.

References

- [1] W. Porter, D. Bardayan, M. Brodeur, D. Burdette, J. Clark, A. Gallant, A. Houff, J. Kolata, B. Liu, P. O'Malley, and et al., *Atoms* **11**, 10.3390/atoms11100129 (2023).
- [2] L. Viehland, *International Journal for Ion Mobility Spectrometry* **15**, 15:21–29 (2011).

Build a $P(E)$ Model Based on the Josephson Junction in MATLAB Language

PENGLIN XIN

2025 NSF/REU Program
Department of Physics and Astronomy
University of Notre Dame

ADVISOR(S): Prof. Xiaolong Liu

Abstract

This work presents the construction and simulation of a model for a Josephson junction embedded in a scanning tunneling microscope (STM) setup using MATLAB. The $P(E)$ theory provides a framework for describing inelastic Cooper pair tunneling processes, where energy is exchanged with the electromagnetic environment. We incorporate a complex transmission line impedance model characterized by a fundamental resonance frequency $\nu_0 = 30\text{GHz}$, dissipation factor $\alpha=0.75$, and dispersion parameter $\kappa=0.005$. Using the iterative algorithm proposed by Ingold and Grabert, we calculate the energy exchange probability distributions $P(E)$, and obtain the resulting I – V characteristics. Our model successfully captures key features such as supercurrent peaks and spectral resonances observed in low-temperature STM experiments.

1 Introduction

The critical Josephson current I_0 serves as a fundamental parameter reflecting the coupling strength between two superconducting condensates across a tunnel junction. It provides direct insight into the superconducting ground state, making it a powerful probe in the study of superconducting phenomena. Integrating the Josephson effect with scanning tunneling microscopy (STM), a technique sometimes referred to as SJTM (Superconducting Josephson Tunneling Microscopy), enables such investigations with atomic-scale resolution.[1]

To date, SJTM experiments have largely been conducted at relatively elevated temperatures, where thermal fluctuations dominate over the Josephson coupling energy. In this regime, unambiguous measurement of the critical current becomes challenging. Achieving direct access to I_0 typically requires either careful pre-engineering of the junction parameters—an approach feasible in lithographically defined planar junctions—or a precise theoretical framework for modeling the system based on known parameters. In STM-based setups, where design flexibility is limited, this approach is suitable. However, not all theoretical models are valid in the STM regime. For instance, the Ivanchenko-Zilberman (IZ) model assumes large junction capacitance and a frequency-independent environment impedance $Z(0)$, conditions that are incompatible with STM junctions, which possess small capacitances and frequency-dependent environmental responses.

In some of the existing STM measurements, we find that the Josephson coupling energy E_J is

much smaller than the charging energy E_C , placing our system firmly in the dynamical Coulomb blockade (DCB) regime. Under such conditions, the conventional DC Josephson effect at zero bias voltage is suppressed. Instead, Cooper pairs tunnel at finite bias voltages $\Delta V \neq 0$, giving rise to a distinct supercurrent peak, as illustrated in Fig. 1.[2] Since a Cooper pair consists of two electrons with opposite spin and momentum, it carries neither net spin nor kinetic energy.[3] Consequently, tunneling in this regime becomes an inelastic process, with the bias energy $2eV_J = h\nu$ released into the environment in the form of photons. Recent studies have analyzed these emitted photon spectra in the context of non-linear quantum electrodynamics. Moreover, if the junction environment supports resonant electromagnetic modes—such as standing waves at frequency ν —these can enhance inelastic Cooper pair tunneling at voltages matching the resonance condition. Such interactions manifest as additional features in the current-voltage characteristics, which can be observed experimentally (see Fig. 1).

To quantitatively describe these inelastic tunneling processes, we employ the $P(E)$ theory, which provides a framework for calculating the probability $P(E)$ that a tunneling Cooper pair emits or absorbs energy E through interactions with its electromagnetic environment. From this probability distribution, the resulting inelastic tunneling current $I(V_J)$ can be directly computed.[4]

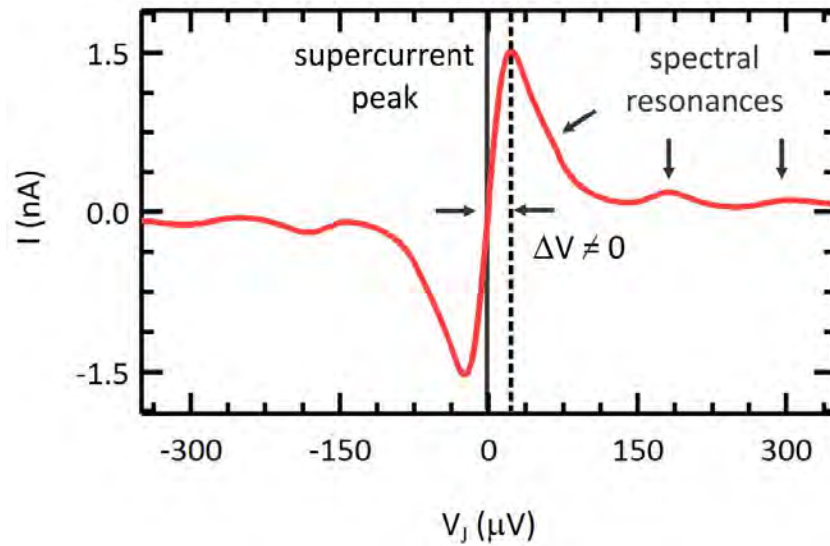


Figure 1: Typical $I(V_J)$ tunneling characteristics measured at $T = 15mK$ [3].

2 $P(E)$ Theory

2.1 Formulation of $P(E)$

By decoupling the Josephson junction's intrinsic dynamics from the environmental influence, this tunneling rate can be reformulated in terms of the probability function $P(E)$:

$$\tilde{\Gamma}(V_J) = \frac{\pi E_J^2}{2\hbar} P(2eV_J), P(E) = \frac{1}{2\pi\hbar} \int_{-\infty}^{\infty} dt \exp \left[J(t) + \frac{i}{\hbar} Et \right]. \quad (1)$$

Here, $J(t)$ is a phase correlation function primarily governed by the frequency-dependent total environmental impedance $Z_t(\nu)$, which is normalized to the Cooper pair resistance quantum $R_Q = h/(2e)^2$. This effective impedance includes contributions from both the junction capacitance C_J and the external circuit environment, and is described as:

$$Z_t(\nu) = \frac{1}{i2\pi\nu C_J + Z^{-1}(\nu)}. \quad (2)$$

We approximate the impedance $Z_T(\nu)$ of the STM tip using the model of an open-ended transmission line. This approximation holds because such a line exhibits similar impedance features to those of a monopole antenna — namely, a strong mismatch at the end and a matched base, giving rise to comparable resonance behavior. It should be noted, however, that this model does not account for potential grounding effects. The full input impedance seen by the junction is expressed as:

$$Z_T(2\pi\nu) = Z(0) \frac{1 + \frac{i}{\alpha} \tan \left(\frac{\pi\nu}{2\nu_0} \right)}{\left[1 - \bar{\kappa} \frac{\nu}{\nu_0} \tan \left(\frac{\pi\nu}{2\nu_0} \right) \right] + i\alpha \left[\kappa \frac{\nu}{\nu_0} + \tan \left(\frac{\pi\nu}{2\nu_0} \right) \right]}. \quad (3)$$

In this formulation, ν_0 denotes the fundamental resonance frequency, α quantifies dissipation in the transmission line, and κ characterizes its frequency dispersion. The function $P(E)$, captures the spectral probability that a tunneling Cooper pair exchanges energy $E = 2eV_J$ with its electromagnetic environment when a voltage bias V_J is applied. As $P(E)$ directly determines the inelastic tunneling rate, a measurable current only occurs when the environment can absorb or supply the

corresponding energy $h\nu = 2eV_J$. Since $P(E)$ represents a probability distribution, it must be normalized over the full energy domain:

$$1 = \int_{-\infty}^{\infty} P(E) dE. \quad (4)$$

Accordingly, the net current through the junction is given by the difference in tunneling probabilities for energy emission and absorption:

$$I(V_J) = \frac{\pi e E_J^2}{\hbar} [P(2eV) - P(-2eV)]. \quad (5)$$

Beyond interactions with the electromagnetic environment, other processes such as thermal fluctuations across the junction capacitance also contribute to inelastic tunneling. A complete model must account for all such energy exchange channels. Therefore, further extensions to the $P(E)$ framework are necessary and described in the following subsection.

2.2 Extensions of the $P(E)$ Framework

To numerically evaluate the energy exchange distribution $P_Z(E)$, which incorporates an arbitrary environmental impedance $Z_t(\nu)$, one could in theory compute Eq. (2) and (3) directly. However, due to their computational complexity, Ingold and Grabert proposed a more tractable integral equation approach:[3]

$$P_Z(E) = I(E) + \int_{-E_0}^{E_0} d\nu K(E, \nu) P_Z(E - h\nu), \quad (6)$$

which can be solved iteratively to obtain the full $P_Z(E)$ distribution. Here, $I(E)$ is a Lorentzian function that serves as the inhomogeneous part of the equation:

$$I(E) = \frac{1}{\pi h} \frac{D}{D^2 + E^2/h^2}, \quad (7)$$

$$D = \frac{\pi k_B T}{h} \frac{\Re[Z_t(0)]}{R_Q}.$$

This inhomogeneity reflects the minimum width of $P_Z(E)$ due to thermal voltage noise originating from a resistive element $R_{\text{DC}} = \Re[Z_T(0)]$ in the circuit. The kernel $K(E, \nu)$ encompasses the full frequency-dependent interaction with the environment and is expressed as:

$$\begin{aligned}
K(E, \nu) &= \frac{E/h}{D^2 + (E/h)^2} k(\nu) + \frac{D}{D^2 + (E/h)^2} \tilde{k}(\nu), \\
k(\nu) &= \frac{1}{1 - \exp[h\nu/(k_B T)]} \frac{\Re[Z_t(\nu)]}{R_Q} - \frac{k_B T}{h\nu} \frac{\Re[Z_t(0)]}{R_Q}, \\
\tilde{k}(\nu) &= \frac{1}{1 - \exp[h\nu/(k_B T)]} \frac{\Im[Z_t(\nu)]}{R_Q} - \frac{2k_B T}{h} \sum_{n=1}^{\infty} \frac{\omega_n}{\omega_n^2 + \nu_n^2} \frac{Z_t(-i\omega_n)}{R_Q}.
\end{aligned} \tag{8}$$

Here, $\omega_n = 2\pi n k_B T$ are the Matsubara frequencies. These equations allow for a self-consistent numerical solution of $P_Z(E)$, using $I(E)$ as the initial guess $P_{m=0}(E)$. Additionally, thermal fluctuations across the junction capacitor C_J introduce voltage noise u , which provides an additional energy exchange channel $E = 2e(V_J + \bar{u})$. The resulting energy distribution, denoted $P_C(E)$, follows a Gaussian form:

$$P_C(E) = \sqrt{\frac{1}{4\pi E_C k_B T}} \exp\left[\frac{-E^2}{4E_C k_B T}\right]. \tag{9}$$

In our case, Coulomb blockade-induced shifts in $P_C(E)$ are neglected, since the low DC impedance $R_{\text{DC}} \ll R_Q$ strongly suppresses such effects in the observed $I(V_J)$ data. For the STM experiments considered here, thermal voltage noise and impedance interaction represent the two primary inelastic processes. Because these mechanisms are statistically independent, the overall energy exchange probability is obtained via convolution:

$$P(E) = \int_{E_0}^{-E_0} dE' P_Z(E') P_C(E - E'). \tag{10}$$

3 Results

Construct the required $P(E)$ model using MATLAB. Under the conditions of $\nu_0 = 30$ GHz, $\alpha = 0.75$, and $\kappa = 0.005$, the resulting frequency-dependent behavior of the complex impedance of a transmission line is shown in Fig.2. Among them, ν_0 characterizes fundamental resonance of the

circuit, α characterizes the dissipative response or loss characteristics of the transmission line, and κ describes the frequency dispersion of the transmission line.

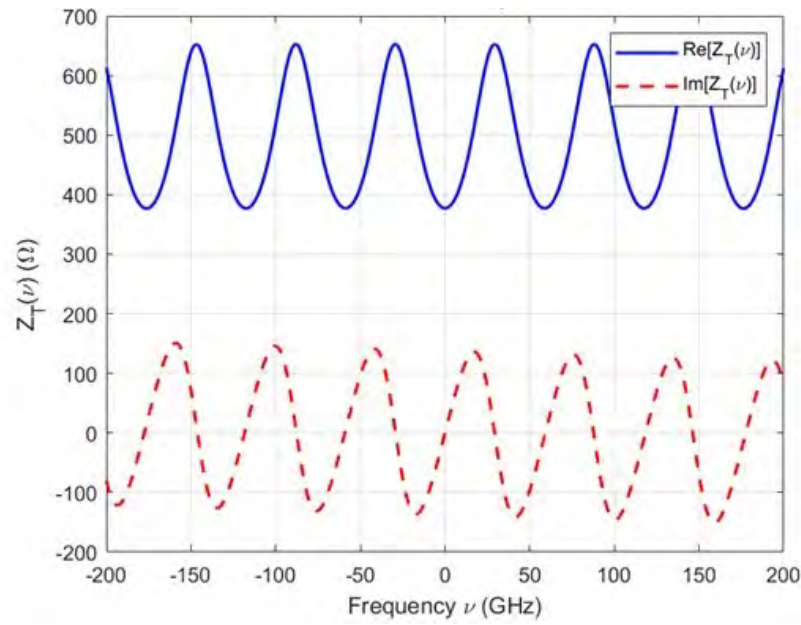


Figure 2: Z_T versus ν

Using the iterative algorithm described in Section 2.2, the calculated $P_Z(E)$ and $P(E)$ curves are shown in Fig.3.

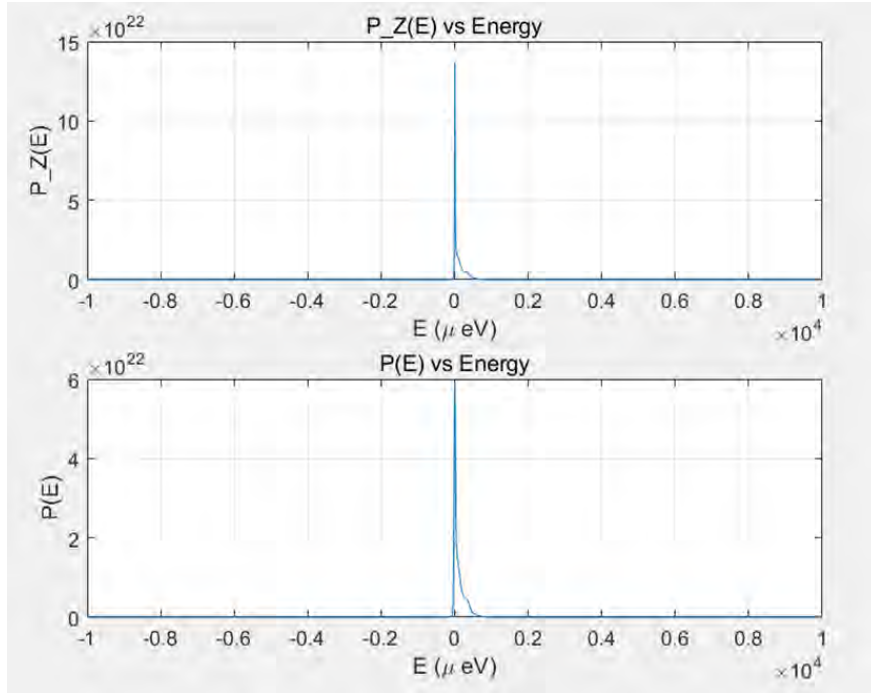


Figure 3: P_Z vs E and P vs E

The final constructed I – V curve is shown in Fig.4.

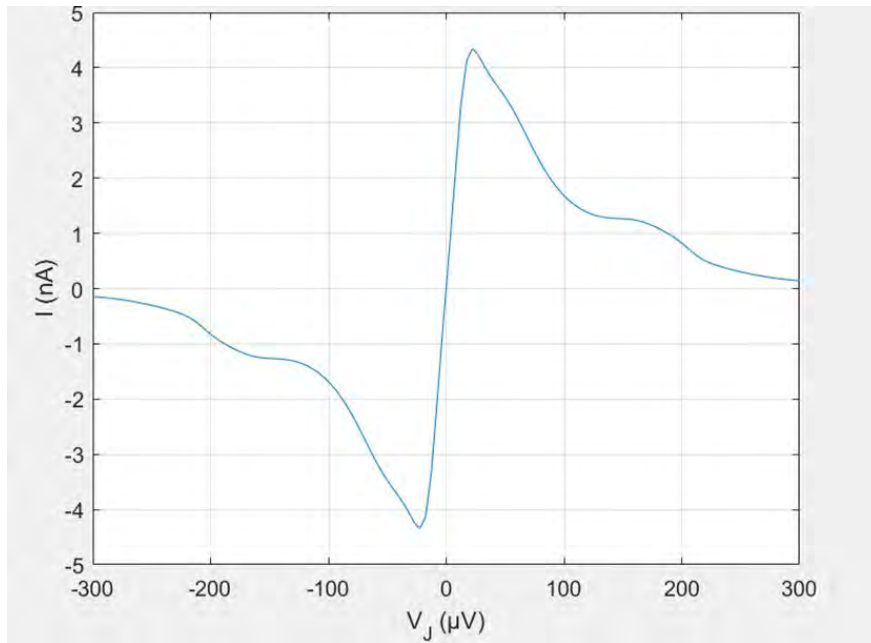


Figure 4: I – V Curve

4 Conclusion

In this work, we established $P(E)$ simulations in MATLAB to describe inelastic Cooper pair tunneling in a Josephson junction operating in the dynamical Coulomb blockade (DCB) regime. The environmental impedance was modeled as a frequency-dependent complex function based on a transmission line approximation. Using this model, we computed the energy exchange probability distributions $P_Z(E)$ and $P_C(E)$, and obtained the total tunneling probability $P(E)$ via convolution. The resulting I - V characteristics successfully reproduce key experimental features such as finite-voltage supercurrent peaks and resonant structures. Our results confirm that the $P(E)$ theory, extended to include both environmental impedance and thermal voltage noise, provides a powerful framework for interpreting Josephson dynamics in STM-based setups at millikelvin temperatures.

5 Acknowledgements

I would like to express my sincere gratitude to my advisor, Prof. Xiaolong Liu, for his invaluable guidance, encouragement, and support throughout the course of this project. His expertise and insights have been instrumental to my understanding and progress. I am also deeply thankful to the members of the laboratory — Dr. Fangjun for his helpful discussions and technical advice, and PhD students Matt, Nileema, and James for their kind assistance and collaboration during experiments. I had a really good summer here.

References

- [1] J. Shmakov, I. Martin, and A. V. Balatsky, Phys. Rev. B **64**, 212506 (2001).
- [2] A. Steinbach, P. Joyez, A. Cottet, D. Esteve, M. H. Devoret, M. E. Huber, and J. M. Martinis, Phys. Rev. Lett. **87**, 137003 (2001).
- [3] B. Jäck, *Josephson Tunneling at the Atomic Scale: The Josephson Effect as a Local Probe*,

Ph.D. thesis, École Polytechnique Fédérale de Lausanne (EPFL), Stuttgart: Max Planck Institute for Solid State Research (2015), PhD thesis.

- [4] B. Jäck, M. Eltschka, M. Assig, M. Etzkorn, C. R. Ast, and K. Kern, Phys. Rev. B **93**, 020504 (2016).

Band Structure of Twisted Bilayer Graphene and WTe_2 via BM Model

LIYUAN XU

2025 NSF/REU Program
Department of Physics and Astronomy
University of Notre Dame

ADVISOR(S): Prof. Yi-Ting Hsu

Abstract

The Bistritzer-McDonald (BM) model has played a pivotal role in understanding the emergence of flat bands and correlated electronic phenomena in twisted bilayer graphene (TBG). In this report, we first reconstruct the theoretical framework developed by Bistritzer and McDonald, providing a detailed derivation of the continuum model. Second, we analyze and comment on the electronic band structure of monolayer WTe, highlighting its key features and symmetries. Finally, we explore the applicability of the BM model to twisted bilayer WTe and attempt to compute and visualize its low-energy band structure near the Dirac point.

1 Introduction

Since its introduction, the Bistritzer-McDonald (BM) model has become a cornerstone in the theoretical study of twisted bilayer graphene (TBG), offering a minimal yet powerful continuum framework to describe the emergence of flat electronic bands at small twist angles. These flat bands have been linked to a variety of strongly correlated phenomena, including unconventional superconductivity and correlated insulating states. The success of the BM model has inspired broader efforts to extend similar approaches to other two-dimensional (2D) materials, in order to investigate whether comparable moiré-induced electronic phenomena might arise.

In this work, we begin by revisiting the BM model[1] in detail. We reconstruct the derivation of the continuum Hamiltonian for TBG, emphasizing the role of interlayer coupling and momentum-space truncation in shaping the low-energy band structure. This forms the theoretical foundation for exploring moiré physics beyond graphene.

Next, we turn our attention to monolayer WTe₂, a material of significant current interest due to its strong spin-orbit coupling, topological characteristics, and anisotropic electronic structure. We analyze the features of its band structure using first-principles and symmetry-based insights, laying the groundwork for a bilayer extension.

Finally, we attempt to generalize the BM framework to twisted bilayer WTe. While the material presents unique challenges compared to graphene including lack of rotational symmetry and a more complex orbital structure we make a first attempt at constructing a continuum model and computing the resulting band structure near the Dirac point. Our results provide a preliminary

step toward understanding moiré effects in transition metal dichalcogenides (TMDs) and suggest possible directions for future refinement of the model.

2 Construction of the BZ Model Hamiltonian

2.1 diagonal elements

The low-energy physics we concern is near the $K_{\pm}$ Dirac points showing linear dispersion, so Dirac fermion is used as the layer Hamiltonian.

$$\langle \mathbf{k}'_l, \alpha' | \hat{H}_0 | \mathbf{k}_l, \alpha \rangle = \hbar v_F [(\mathbf{k}_l - \mathbf{K}_l) \cdot \boldsymbol{\sigma}]_{\alpha' \alpha} \delta_{\mathbf{k}'_l, \mathbf{k}_l} \quad (1)$$

Fold it into a mini Brillouin zone and we obtain

$$\langle \mathbf{k}'_l, \mathbf{Q}, \alpha' | \hat{H}_0 | \mathbf{k}_l, \mathbf{Q}', \alpha \rangle = \hbar v_F [(\mathbf{k}_l - \mathbf{K}_l + \mathbf{Q}) \cdot \boldsymbol{\sigma}]_{\alpha' \alpha} \delta_{\mathbf{k}'_l, \mathbf{k}_l} \delta_{\mathbf{Q}, \mathbf{Q}'} \quad (2)$$

Where $\mathbf{k}_l \in mBZ$ and $\mathbf{Q}$ is the linear combination of the basis of the mBZ . $l = \pm$ denotes the operation of rotating $\pm\theta/2$.

2.2 off-diagonal elements

Here we apply the two-center approximation, ignoring the influence of nearby atoms, assuming that the interlayer hopping only depends on the vector connecting the two atoms, not on the surrounding lattice. Moreover, we assume the homogeneity of the $t(\mathbf{q})$, i.e. it's only a function of the modulus of $\mathbf{q}$.

$$t(\mathbf{R}_i, \mathbf{R}_j) = t(\Delta\mathbf{R}) \quad (3)$$

Also, the result obtained from π -band tight-binding model shows that $t_{\mathbf{q}}$, the fourier transformation of $t(\mathbf{r})$ decline rapidly for big $\mathbf{q}$. This enables us to only keep the leading-order terms of the off-diagonal elements.

In an untwisted frame, the momentum before and after the tunneling is $\mathbf{k}$ and $\mathbf{k}'$. The sum we later operate will be under the twisted frame, so $\pm$ /notation is still used.

In a π -band tight-binding model the projection of the Bloch function of the two layers to a given sublattice is

$$|\phi_{\alpha,\mathbf{k}}\rangle = \frac{1}{\sqrt{N}} \sum e^{i\mathbf{k}(\mathbf{R}+\delta^\alpha)} |\mathbf{R} + \delta^\alpha\rangle \quad (4)$$

Substitute this into the matrix element to obtain

$$\langle \phi_{\beta,\mathbf{k}'_l} | \hat{H} | \phi_{\alpha,\mathbf{k}_l} \rangle = \frac{1}{N} \sum_{\mathbf{R}'_l \in \Lambda_{-l}} \sum_{\mathbf{R}_l \in \Lambda_l} e^{-i\mathbf{k}_{-l} \cdot (\mathbf{R}'_l + \delta_{-l}^\beta)} e^{i\mathbf{k}_l \cdot (\mathbf{R}_l + \delta_l^\alpha)} t(\mathbf{R}'_l + \delta_{-l}^\beta - \mathbf{R}_l - \delta_l^\alpha) \quad (5)$$

Because of the finite volume(area) of the crystal, the reciprocal boundary condition leads to the discrete form of the fourier transformation of $t(\mathbf{R})$. From Bloch Theorem we get

$$e^{i\mathbf{k} \cdot (N_1 \mathbf{a}_1 + N_2 \mathbf{a}_2)} = 1 \quad (6)$$

implying that the momentum $\mathbf{k}$ only takes its value on a mesh containing $N = N_1 N_2$ lattices. Then the fourier transformation has the form of

$$t(\Delta \mathbf{R}) = \frac{1}{(2\pi)^2} \int_{\mathbf{R}^2} d^2 \mathbf{q} \cdot \hat{t}(\mathbf{q}) e^{i\mathbf{q} \cdot \Delta \mathbf{R}} = \frac{1}{N A_{u.c.}} \sum t_{\mathbf{q}} \cdot e^{i\mathbf{q} \cdot \Delta \mathbf{R}} \quad (7)$$

Substitute in equation (5) and we get

$$\begin{aligned} \langle \phi_{\beta,\mathbf{k}'_l} | \hat{H} | \phi_{\alpha,\mathbf{k}_l} \rangle = & \sum_{\mathbf{q} \in \Delta q^2 \mathbb{Z}^2 \text{ in BZ}} \frac{t(\mathbf{q})}{A_{u.c.}} \left(\frac{1}{N} \sum_{\mathbf{R}_{-l} \in \Lambda_{-l}} \exp [i(\mathbf{q} - \mathbf{k}_{-l}) \cdot \mathbf{R}_{-l}] \right) \cdot \\ & \left(\frac{1}{N} \sum_{\mathbf{R}_l \in \Lambda_l} \exp [-i(\mathbf{q} - \mathbf{k}_l) \cdot \mathbf{R}_l] \right) \\ & \times \exp [i(\mathbf{q} - \mathbf{k}_l) \cdot \delta_{-l}^\beta - i(\mathbf{q} - \mathbf{k}_{-l}) \cdot \delta_l^\alpha] \end{aligned} \quad (8)$$

Using the orthogonality relationship to obtain

$$\langle \phi_{\beta, \mathbf{k}'_{-l}} | \hat{H} | \phi_{\alpha, \mathbf{k}_l} \rangle = \sum_{\mathbf{G}_{1(-l)} \in BZ^-, \mathbf{G}_{2(l)} \in BZ^+} \frac{t(\mathbf{k}_l + \mathbf{G}_{2(l)})}{A_{\text{u.c.}}} \exp [i\mathbf{G}_{1(-l)} \cdot \delta_{-l}^\beta - i\mathbf{G}_{2(l)} \cdot \delta_l^\alpha] \cdot \delta_{\mathbf{k}'_{-l} + \mathbf{G}_{1(-l)}, \mathbf{k}_l + \mathbf{G}_{2(-l)}} \quad (9)$$

As the rotation of frames keeps the inner product unchanged, the equation above can be rewritten as

$$\langle \phi_{\beta, \mathbf{k}'_{-l}} | \hat{H} | \phi_{\alpha, \mathbf{k}_l} \rangle = \sum_{\mathbf{G}_1, \mathbf{G}_2 \in \text{untwisted } BZ} \frac{t(\mathbf{k} + \mathbf{G}_2)}{A_{\text{u.c.}}} \exp [i\mathbf{G}_1 \cdot \delta^\beta - i\mathbf{G}_2 \cdot \delta^\alpha] \cdot \delta_{\mathbf{k}'_{-l} + \mathbf{G}_{1(-l)}, \mathbf{k}_l + \mathbf{G}_{2(-l)}} \cdot \delta_{\mathbf{k}'_{-l} + \mathbf{G}_{1(-l)}, \mathbf{k}_l + \mathbf{G}_{2(-l)}} \quad (10)$$

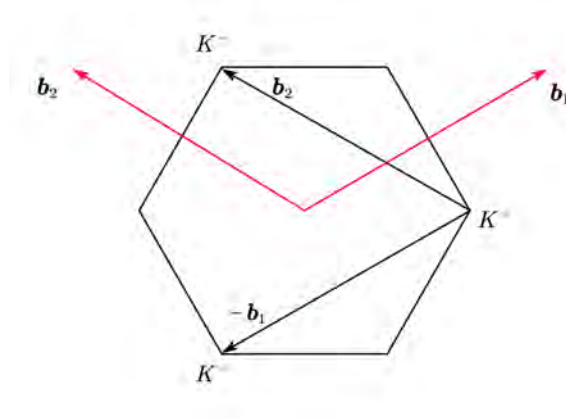


Figure 1: equivalent dirac points

From this we can interpret the quasi momentum conservation of the tunneling process

$$\mathbf{k}'_{-l} + \mathbf{G}_{1(-l)} = \mathbf{k}_l + \mathbf{G}_{2(-l)} \quad (11)$$

We mainly focus on the low-energy physics in the neighborhood of the Dirac points, and expand the off-diagonal matrix element about the difference between the tunneling momentum and Dirac momentum $\Delta \mathbf{k} \equiv \mathbf{k} - \mathbf{K}$. We could separately deal with the two inequivalent Dirac points for they are completely uncoupled, and the ultimate result can be written as the direct sum of the two valley hamiltonian. Given the rapid declining of $t(\mathbf{q})$, the 3 leading-order terms correspond to the three equivalent Dirac points making the minimum $|\mathbf{q}| \equiv |\mathbf{K}^+ + \mathbf{G}_2|$ as shown in fig 1:

In these terms, $\mathbf{G}_2 = 0, -\mathbf{b}_1, \mathbf{b}_2$, or written in full form

$$\langle \phi_{\beta, \mathbf{k}'_l} | \hat{H} | \phi_{\alpha, \mathbf{k}_l} \rangle = \sum_i \frac{t(|\mathbf{K}|)}{A_{u.c.}} \exp[i\mathbf{G}_i \cdot (\delta^\beta - \delta^\alpha)] \delta_{\Delta \mathbf{k}'_l, \Delta \mathbf{k}_l + \mathbf{q}_i} \quad (12)$$

The matrix form of the 3 terms is

$$T_1 = \begin{pmatrix} w_0 & w_1 \\ w_1 & w_0 \end{pmatrix}, \quad T_2 = \begin{pmatrix} w_0 & w_1 e^{-i\frac{2\pi}{3}} \\ w_1 e^{i\frac{2\pi}{3}} & w_0 \end{pmatrix}, \quad T_3 = \begin{pmatrix} w_0 & w_1 e^{i\frac{2\pi}{3}} \\ w_1 e^{-i\frac{2\pi}{3}} & w_0 \end{pmatrix}$$

If not taken the structural relaxation into account, the coefficients w_0 and w_1 are equal.

This corresponds to the "tripod" model introduced in the Bistritzer paper, which describes the interlayer coupling between nearest-neighbor sites in the moiré lattice. Technically, the BM Hamiltonian constructed via this procedure results in a large matrix whose dimension scales with the number of moiré lattice sites contained within the (untwisted) Brillouin zone, posing significant computational challenges during diagonalization. To overcome this difficulty, a truncation scheme is typically employed, in which only moiré lattice sites within a certain nearest-neighbor range, shown in equation 13, are included in the calculation.

$$\mathbf{k} = \mathbf{K}^\pm + n_1 \mathbf{G}_{M1} + n_2 \mathbf{G}_{M2}, n_1, n_2 \in [-N, N] \quad (13)$$

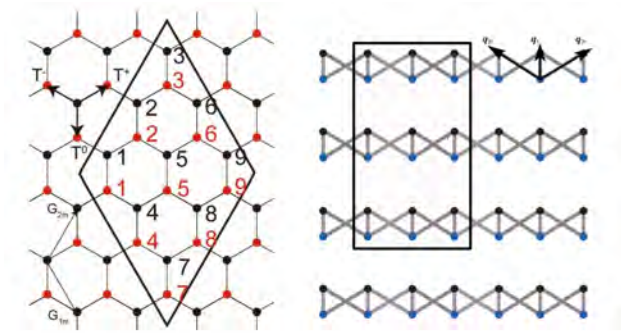


Figure 2: truncations of grapene and WTe₂ at N=1

3 Monolayer and Twisted Bilayer WTe₂

3.1 Monolayer WTe₂

The monolayer 1T'-WTe₂ exhibits a unique electronic band structure characterized by band inversion near the Fermi level. This inversion arises from strong spin-orbit coupling and the particular lattice distortion of the 1T' phase, where the W atoms form zigzag chains. As a result, the conduction and valence bands near the Γ point are inverted compared to a trivial insulator. This inverted band ordering leads to a nontrivial topological invariant, classifying monolayer WTe as a two-dimensional topological insulator (2D TI), or quantum spin Hall (QSH) insulator. In this phase, the bulk remains insulating while conducting helical edge states emerge, protected by time-reversal symmetry. These properties make WTe a promising platform for exploring topological quantum phenomena in two dimensions.

The layer Hamiltonian that we'll later use as the diagonal element of the bilayer Hamiltonian is constructed via a low-energy effective model, which is tight-binding model involving two d orbitals from W (namely $d_{x^2-y^2}$ and d_{xy}) and two p orbitals from Te (p_x and p_y)[2]. We draw the monolayer band structure according to the Hamiltonian in [1] as follows:

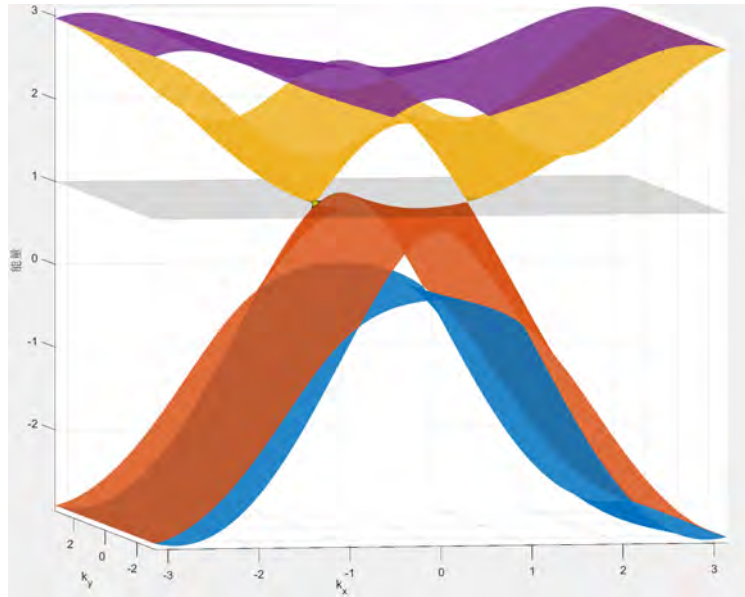


Figure 3: band structure of monolayer WTe₂

3.2 Twisted Bilayer WTe₂

3.2.1 Coupling

In the previous discussion, we assume that the different valleys, namely $\mathbf{K}^+$ and $\mathbf{K}^-$ have no coupling because of the distance between the two dirac points is significant. However, in the WTe₂ case, the two valleys are close to each other, giving rise to the considerable coupling. However in the first place we decide to omit the coupling between valleys to get a preliminary result.

3.3 Hamiltonian

For the interlayer elements, we only consider the coupling between the nearest valley of the same kind, and the coupling matrix is apparently

$$T = w * \begin{pmatrix} 1 & 1 & 1 & 1 \\ 1 & 1 & 1 & 1 \\ 1 & 1 & 1 & 1 \\ 1 & 1 & 1 & 1 \end{pmatrix} \quad (14)$$

where w is the coupling intensity, according to the graphene result.

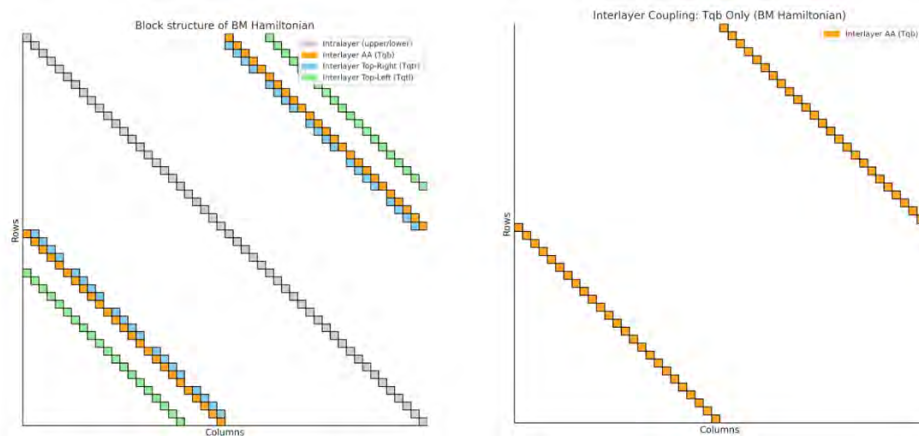


Figure 4: Hamiltonian scheme of bilayer graphene and WTe₂

The Hamiltonian schematic graphs is as above, which is a quite sparse matrix. We diagonalize

it and plot the band structure near the valley at $w = 0.1, \theta = 1.05^\circ$, as follows:

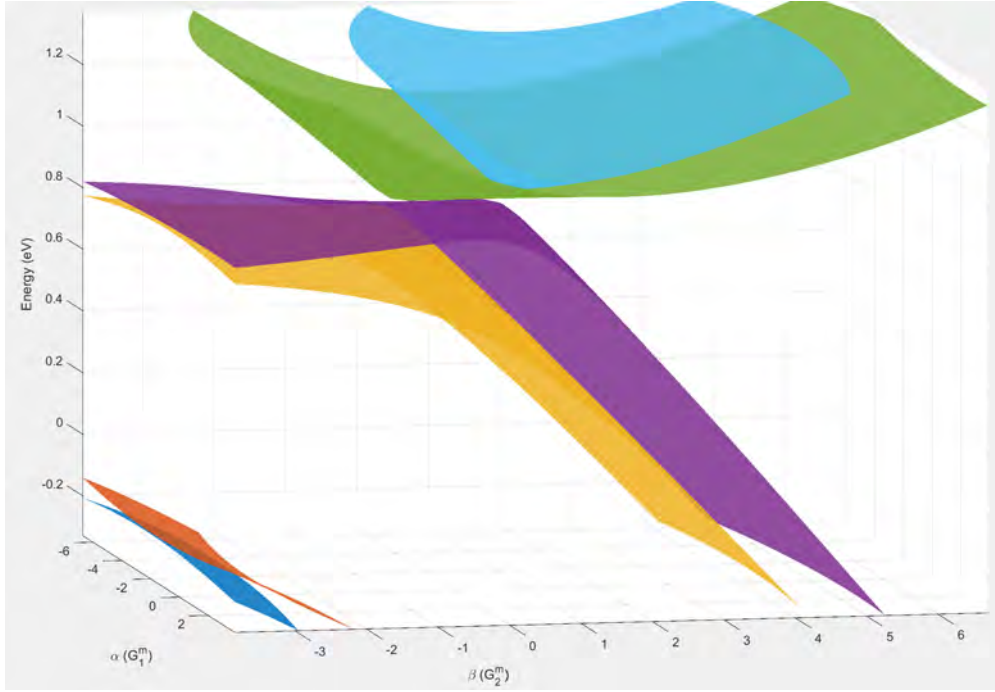


Figure 5: preliminary band structure of bilayer WTe_2

From this picture we can interpret similar behavior of the bands in the neighborhood of the valley to Fig1 e-g of [4] because of similar Dirac cone coupling mechanism.

4 Conclusion

In this study, we extended the Bistritzer-MacDonald (BM) model originally formulated for twisted bilayer graphene to twisted bilayer WTe_2 . By incorporating the material-specific band structure and interlayer coupling characteristics of WTe_2 , we get a result of a continuum Hamiltonian suitable for small twist angles. Preliminary band structure calculations reveal significant reconstruction of the low-energy electronic states due to the moiré potential.

5 References

- [1]Bistritzer, R., MacDonald, A. H. (2011). Moiré bands in twisted double-layer graphene. Proceedings of the National Academy of Sciences, 108(30), 12233-12237.
- [2]Muechler, L., Alexandradinata, A., Neupert, T., Car, R. (2016). Topological nonsymmorphic metals from band inversion. Physical Review X, 6(4), 041069.
- [3]Cano, J., Fang, S., Pixley, J.H., Wilson, J.H. (2021). A moiré superlattice on the surface of a topological insulator. Physical Review B, 103(15), 155157.
- [4]Cao, Y. et al, (2018). Correlated insulator behaviour at half-filling in magic-angle graphene superlattices. Nature, 556(7699), 80-84.

Characterizing Thermionic Emission of Tungsten Filament

ZIHENG XU

2025 NSF/REU Program
Department of Physics and Astronomy
University of Notre Dame

ADVISOR(S): Prof. Dafei Jin

Abstract

Electron charge qubits are compelling candidates for solid-state quantum computing because of their inherent simplicity in qubit design, fabrication, control, and readout. The Emergent Quantum Systems Lab (EQSL) at Notre Dame is developing a novel quantum processor based on electron-on-solid neon (eNe) qubits. One critical process of building this quantum system is the production of electron qubits. Thermionic emission is the simplest way to generate electrons. To prevent damage to the cryogenic system while ensuring efficient electron capture, it is essential to determine the optimal voltage to apply to the filament. This project aims to verify the emission characteristic of thermionic emission.

1 Introduction

Quantum computing's advancement hinges on developing robust qubit platforms that simultaneously achieve long coherence times, high-fidelity control, and scalability[1; 2]. Traditional electron qubits in semiconductor heterostructures face significant coherence limitations due to material defects and electromagnetic noise. In response, the electron-on-solid-neon (eNe) platform has emerged as a breakthrough, leveraging the ultrapure, nuclear spin-free environment of solid neon crystals to host trapped single-electron qubits[3].

Solid neon's crystalline structure provides a uniquely clean substrate for electron trapping. When cooled below its triple point (24.6 K and 0.43 bar), neon forms a rigid surface with positive electron affinity—a property absent in most other materials[4]. In this ultra clean state, solid neon can form a free surface to vacuum with no uncontrollable impurities, enabling stable electron trapping through a balance of quantum forces: a repulsive Pauli barrier (0.7 eV) from neon atom and an attractive polarization potential ($V(z) = -(\epsilon - 1)/[(\epsilon + 1)e^2/4z]$) from induced image charges, yielding a ground-state binding energy of $E_{z0} = -15.8\text{meV}$. The in-plane motion of electrons can be engineered by an external electro-static trapping potential and delivers state-of-the-art coherence times ($T_1 = 0.1\text{ms}$, $T_{2E} = 0.1\text{ms}$) for charge qubits.

A key requirement for this platform's operation is the reliable injection of a single electron onto the neon surface. Since different filament has it own electrons emission characteristic, it's important to exam the emission ability of the filament EQSL is using now. In the present experiment, we focus

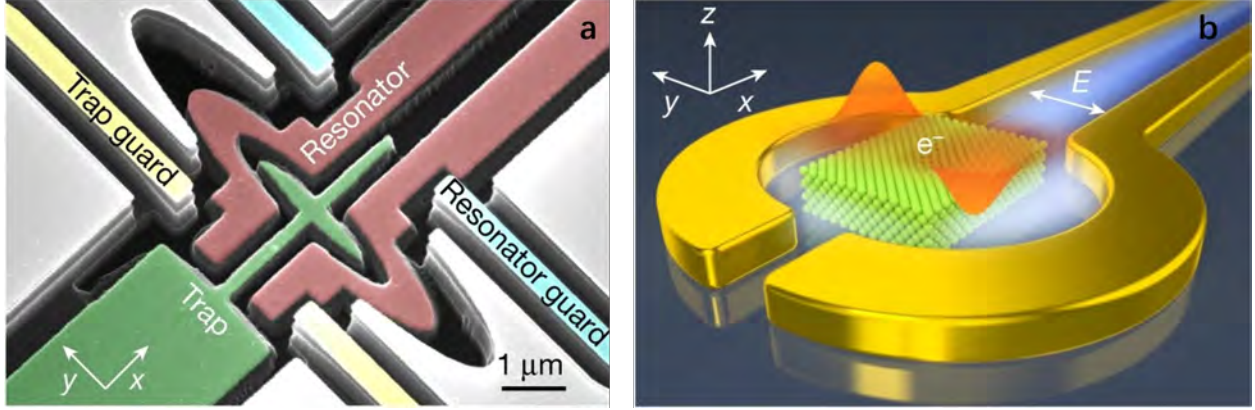


Figure 1: a. (False-color) scanning electron microscopy picture of the fabricated device around the electron trap and photon coupling region. The trap and two striplines reside inside an etched channel that is $3.5 \mu\text{m}$ wide and $1.5 \mu\text{m}$ deep. b. Conceptual illustration of a single electron trapped on a solid neon surface and interacting with microwave photons at the open end of a superconducting coplanar stripline resonator. The electric dipole transition and the electric field of microwave photons are aligned in the y direction. The Figure is reproduced from [3].

on quantifying steady-state thermionic emission current and the Joule heating is not considered under this context.

2 Background

Thermionic emission refers to the liberation of charged particles from a heated metal surface, whereby thermal energy provides sufficient kinetic energy for particles to overcome the material's surface potential barrier. The kinetic energies of these free electrons follow a Boltzmann distribution. At elevated temperatures, a fraction of electrons gains sufficient kinetic energy (normal to the surface) to surmount the work function barrier and escape into vacuum. The magnitude of thermionic emission is highly sensitive to temperature (T) and the work function (ϕ). In 1901, Owen Willans Richardson empirically demonstrated that thermionic current density scales approximately as $T^{1/2}e^{-b/T}$ [5], where b is an empirical parameter. Subsequent theoretical analysis led Richardson to derive the fundamental equation governing thermionic emission, now known as Richardson's Law [6]:

$$J = A_G T^2 \exp\left(-\frac{\phi}{k_B T}\right) \quad (1)$$

where J is the current density, A_G is the Richardson constant that is specific to different metals, and k_B is the Boltzmann constant. In our experimental setup, direct measurement of the filament's temperatures is unavailable. However, empirical evidence[7] establishes a strong correlation between the filament's temperatures and the applied electrical parameters (voltage or current).

3 Methods

Since the experimental objective is to measure electron emission characteristics, investigations must be conducted under ultra-high vacuum conditions. Otherwise, gas molecules would collide with the emitted electrons, preventing their detection at the anode. The mean free path of air at room temperature and atmospheric pressure is 10^{-5} m, which is much shorter than the distance between the filament and the anode. Such an ultra-high vacuum environment can be achieved through rigorous indium sealing combined with turbo pump evacuation or by placing the system in a cryogenic environment (Four9Design cryostat).

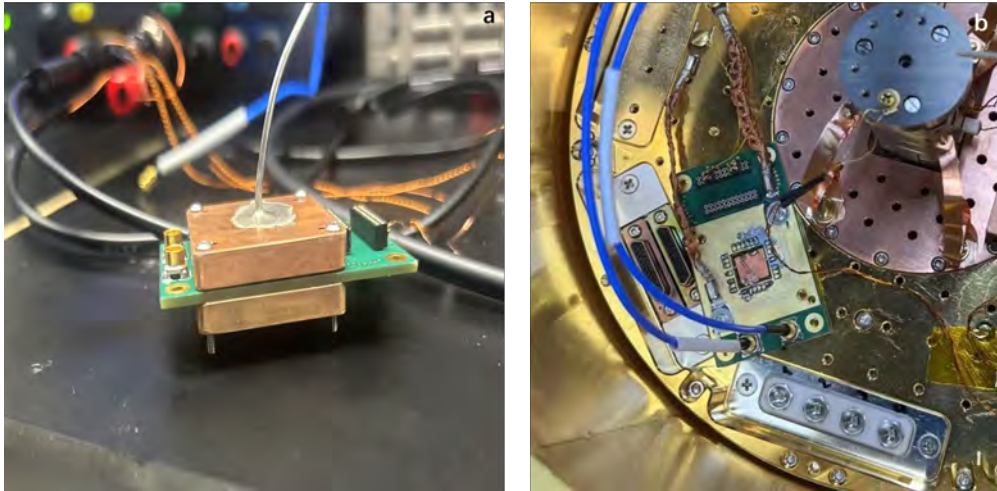


Figure 2: a. Initial version of the device. b. Final version of the device

The thermionic emission apparatus initially consisted of a printed circuit board (PCB) assembly enclosed within dual metallic shields and sealed with indium wire and high-vacuum grease. After

evacuation through a scroll pump, the system reached a base pressure of 10^{-2} mbar at room temperature, corresponding to a mean free path of 10^{-1} m—still comparable to the distance between the filament and the anode. Consequently, premature filament failure and undetectable anode signals ($< 10\text{pA}$) were observed, necessitating significant reconfiguration. In the second experimental run, we directly put the PCB assembly into a cryogenic system integrated with a turbo pump, the system achieves a base pressure of 10^{-8} mbar at low temperature and 10^{-5} mbar at room temperature, yielding a mean free path exceeding 1 m. This improvement resolved the signal-detection limitations and extended operational stability, simultaneously.

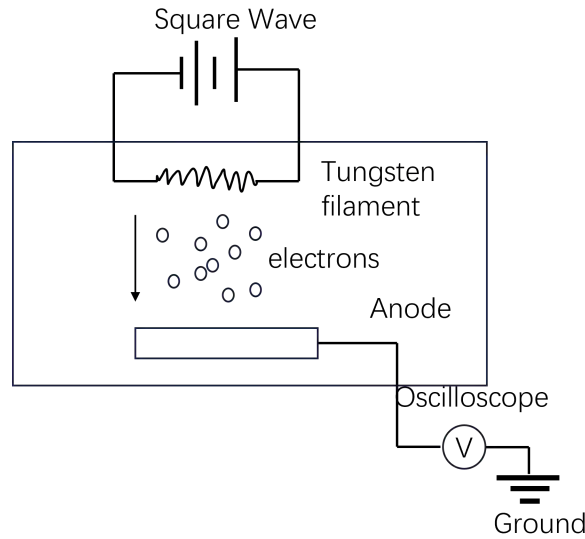


Figure 3: Schematic of device

As shown in Fig. 3, voltage pulse is applied to the filament using a Keysight E36312A DC power source. This induced electron emission toward the metal anode. Upon receiving these electrons, the anode generates a current signal that is transmitted to an oscilloscope for measurement. To synchronize data acquisition and capture images, a 2V binary trigger signal was simultaneously sent to the oscilloscope at the start of each pulse cycle.

4 Results

Despite recent studies indicating that abrupt filament resistance transitions occur exclusively in superfluid helium at cryogenic temperatures[8], we have observed this phenomenon in cryogenic vacuum environments. As illustrated in Fig. 4, the filament exhibits relatively low resistance when subjected to constant voltages below 0.1 V, while an abrupt resistance surge occurs at 0.75 V. The figure suggests favorable electron emission properties in the high-resistance region.

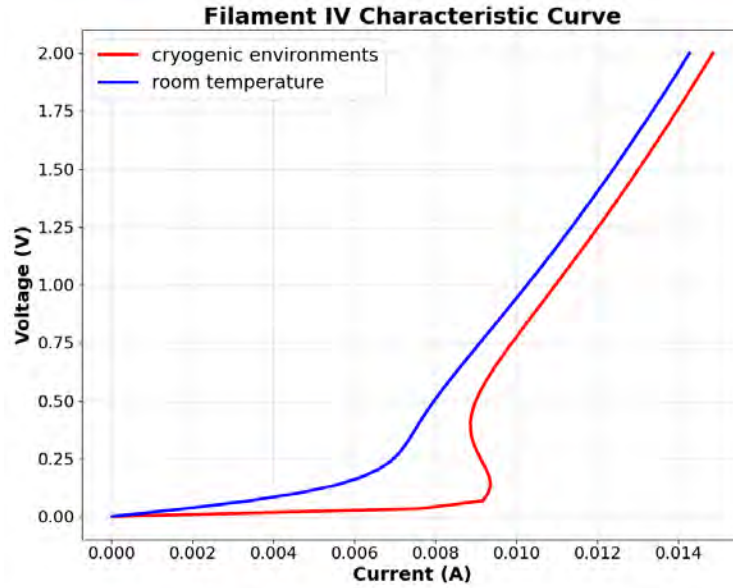


Figure 4: IV curve of filament in cryogenic environments and room temperature.

To characterize the fundamental electron emission properties, initial experiments employed a robust tungsten filament (larger diameter to prevent fracture) under cryogenic conditions. As shown in Fig. 5(a), the voltage applied to the filament triggered an abrupt current surge between the filament and the anode, followed by gradual recovery to a steady-state value. The termination of the applied voltage at $t=60$ s caused an immediate current collapse to near-zero levels, confirming the emission origin of the observed current. Fig. 5(b) presents comparative data acquired with the fine filament designed for eNe quantum systems, tested in both cryogenic and room temperature environments. The room temperature configuration demonstrates significantly enhanced emission characteristics: lower threshold voltages for electron emission and higher output currents at equivalent applied voltages compared to cryogenic operation.

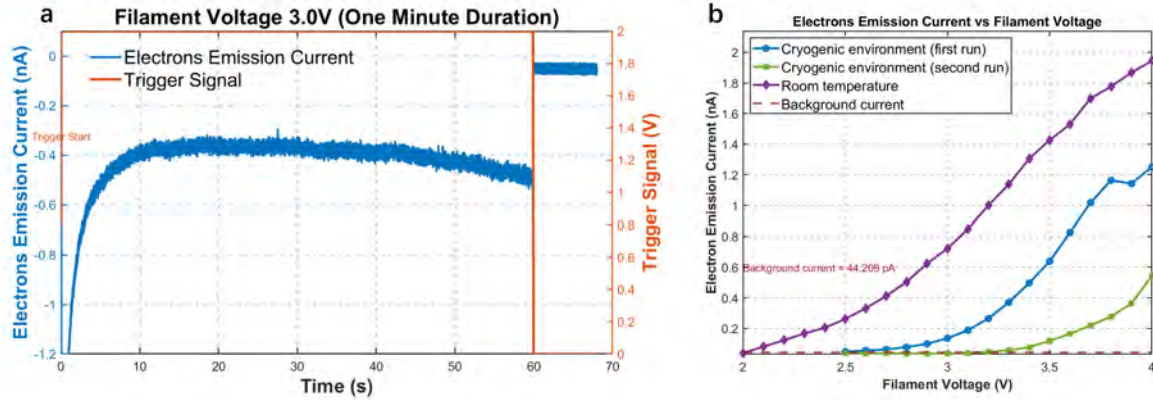


Figure 5: a. Emission current waveform. b. Measurement result in room temperature and cryogenic environments

A notable observation from Fig. 5(b) is the significantly higher threshold voltage reduced and output current during the second cryogenic measurement compared to the first run, indicating degraded electron emission capability. Since the intervals between voltage applications exceeded 30 s—longer than tungsten’s thermal relaxation time—this phenomenon cannot be attributed to incomplete filament cooling[9]. To investigate the underlying mechanism, we conducted experiments using a previously unheated filament, monitoring its resistance before each pulse. The resultant data in Fig. 6(a) suggest that chemical changes may have occurred in the filament, as reflected by the resistance variations. Optical microscopy images of the filament captured before and after the heating cycles are presented in Fig. 6(b). The distinctive color transition observed in the post-heating specimen provides empirical evidence supporting our hypothesis regarding material degradation. Additionally, previous studies have indicated that when filaments are fabricated from tungsten-thorium alloy, a portion of the thorium evaporates under high-temperature heating conditions. This behavior causes a reduction in the overall work function of the filament, subsequently diminishing its electron emission capability.[10; 11]

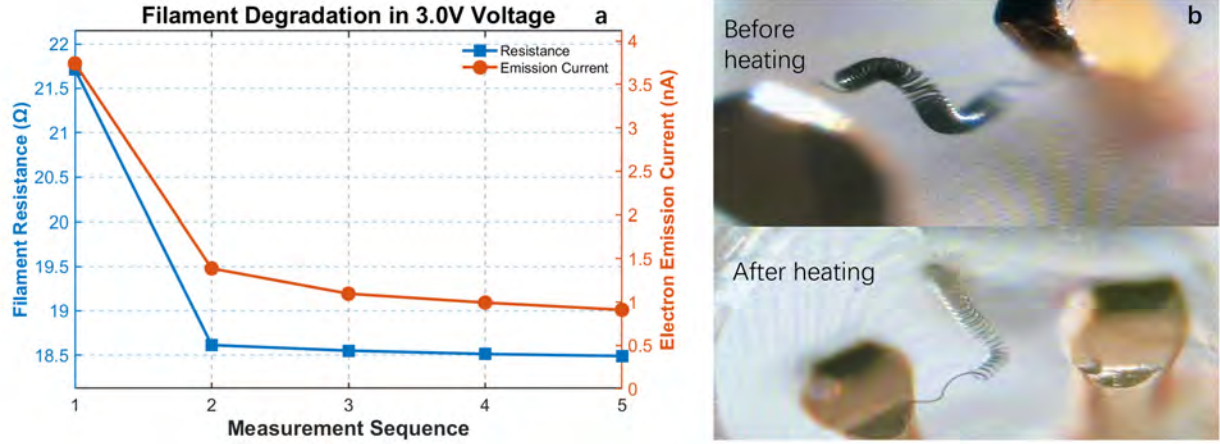


Figure 6: a. Resistance and emission current with measurement time. b. Photos of filament under microscopy.

5 Conclusion

This experiment demonstrates that tungsten filaments serve as effective thermionic electron sources, capable of supplying sufficient electrons for eNe quantum systems. However, the observed degradation in electron emission performance suggests that filaments in eNe configurations may require frequent replacement to maintain optimal output, as emission currents exhibit progressive diminution across successive heating cycles.

6 Outlook

Although thermionic emission from heated filaments has demonstrated promising results for electron delivery in eNe quantum systems at the EQSL Laboratory, the inherent energy dispersion of thermally emitted electrons presents fundamental limitations. A significant fraction of electrons may penetrate interstitial sites within the neon matrix rather than achieving precise localization on the neon surface, thereby degrading qubit coherence. Given the constant 0.7 eV repulsive barrier governing free-electron–neon atom interactions[3], photoelectron emission offers a compelling alternative. Through the photoelectric effect, the kinetic energy of emitted electrons can be precisely tuned by selecting the light frequency ($E_k = h\nu - W$), enabling controlled surface deposition. Following this REU program, the EQSL Laboratory will continue to investigate the feasibility of

implementing photoelectric emission techniques.

7 Acknowledgment

First, I express profound gratitude to Professor Dafei Jin for providing an exceptional opportunity for me in this summer's REU program. Special appreciation is extended to Postdoc Shan Zou, whose dedicated mentorship and considerable time investment in my project provided indispensable laboratory training and scientific direction. I am deeply grateful to other group members, Alec, Charles, Evgenii, Lauren, Juliana, and Yutian, for their generous assistance whenever needed. Finally, sincere thanks to Professor Garg and Ms. Kristen for their exceptional organization of the REU program and steadfast support throughout this two-month research journey.

References

- [1] G. Popkin, *Science* **354**, 1090 (2016).
- [2] N. P. de Leon, K. M. Itoh, D. Kim, K. K. Mehta, T. E. Northup, H. Paik, B. S. Palmer, N. Samarth, S. Sangtawesin, and D. W. Steuerman, *Science* **372**, eabb2823 (2021).
- [3] X. Zhou, G. Koolstra, and X. e. a. Zhang, *Nature* **605**, 46 (2022).
- [4] G. L. POLLACK, *Rev. Mod. Phys.* **36**, 748 (1964).
- [5] O. W. Richardson, *Proceedings of the Cambridge Philosophical Society*. **11**, 286 (1901).
- [6] O. W. Richardson, *The Emission of Electricity from Hot Bodies* (Longmans Green and Co., 1916) pp. 63–64.
- [7] J. quan Li, S. han Li, and P. Liu, *Vacuum* **210**, 111837 (2023).
- [8] I. F. Silvera and J. Tempere, *Physical Review Letters* **100**, 117602 (2008).
- [9] T. Durakiewicz and S. Halas, *Journal of Vacuum Science & Technology A* **17**, 1071 (1999).

[10] J. E. Polk, AIP Conference Proceedings **271**, 1435 (1993).

[11] I. Langmuir, Phys. Rev. **22**, 357 (1923).