

A SUPERCONDUCTING ION DETECTOR

By

Joseph Suttle

A DISSERTATION SUBMITTED IN PARTIAL FULFILLMENT OF THE
REQUIREMENTS FOR THE DEGREE OF

DOCTOR OF PHILOSOPHY

(PHYSICS)

at the

UNIVERSITY OF WISCONSIN – MADISON

2018

Date of final oral examination: May 15, 2018

The dissertation is approved by the following members of the Final Oral Committee:

Robert McDermott, Professor, Physics

Dan McCammon, Professor, Physics

Aki Hashimoto, Professor, Physics

Irena Knezevic, Professor, Electrical and Computer Engineering

Abstract

This thesis will explore the theory and design of a novel variant of the superconducting nanowire single photon detector (SNSPD). This new detector is dubbed the Superconducting Delay Line Detector (SCDLD) and is designed to enable the scaling of SNSPDs to unprecedented sizes while also adding the ability to discern the position of detected particles. The detection mechanism of the SCDLD is essentially identical to that of the SNSPD; a superconducting wire is biased near its critical current. When an incident particle's energy is absorbed, the local critical current density is suppressed, causing the wire to revert to its normal state, generating a measureable voltage pulse. In the SCDLD, this wire is designed to operate as a microstrip rather than an inductor. Now, an ion event generates a pair of voltage pulses of opposite polarity which travel to opposite sides of the microstrip. By measuring the relative delay between their arrivals, it is possible to calculate where in the detector the event occurred.

Because this detector architecture is substantially different from previous those studied in much of the SNSPD literature, careful attention will be paid to the dynamics of this detector system. In particular, the subject of latching is revisited, taking into account the new electrical system our detector comprises. Elements of the device physics which have not been thoroughly studied in the SNSPD literature will also be treated, such as the interaction of multiple interleaved detectors operating within close proximity to one another on a single die. In addition, the new capability of position discernment allows for unique methods of data collection and analysis compared to existing SNSPD work which will be studied extensively.

This new detector technology is developed with an eye towards applications in the field of materials' analysis. In particular, time of flight spectrometry of large biomolecules and atom probe tomography of myriad sample types are both interesting potential applications for these detectors. With these applications in mind, the work in this thesis is therefore geared towards the detection of ions which have energies of order 5 to 10 kiloelectronvolts.

Acknowledgements

It has been a long road to completing this thesis, and I have many people to thank for support along the way.

First off, I would like to thank my advisor, Robert McDermott. Throughout my graduate career, Robert has given me plenty of room to work and explore the various challenges associated with this project. He has also been a wellspring of expertise especially in the realms of superconductivity and electronics in general. Other members of the McDermott Group have been invaluable to my success over the years as well. The contributions made by the previous generation of graduate students and post-docs laid a good portion of the groundwork for this study, building up much of the infrastructure necessary to fabricate and measure the detectors discussed in this thesis. There is so much work that happens in the background of a research group like this just to keep labs operational, and I'd like to express my sincere thanks to all of my fellow groupmates who make this all run smoothly. In particular, Guilhem Ribeill and Ed Leonard have taken on many of these sometimes invisible tasks. I would like to thank Matt Beck for so many stimulating conversations over the years; there have been many times where we acting as sounding boards for one another's more outlandish ideas. That, and metal Thursdays. Matt, Ed, and I have all also collaborated extensively on fabrication processes and chased down issues like step coverage failure and etch recipe drifts. Umesh Patel was also an excellent resource; he helped me keep perspective on things, and was also a good source of ideas from a new perspective. In fact, he was the first to suggest to me that the reactive ion etch process we were using could have been causing my yield issues. Ted

Thorbeck has also been instrumental to making our group run efficiently, helping teach graduate students and instilling the group with a positive attitude. More recently, Alex Opremcak has been a great person to talk to. He asks interesting, thoughtful questions and pushes everyone around him to think more deeply about scientific questions. He's also a great source of levity when it all seems like too much to handle. Our newest postdoc, Brad Christensen has also been the source of many interesting discussions and is always willing to field my naive questions about quantum information experiments. Finally, I would like to extend a very sincere thanks to Francisco 'Paco' Schlenker who is instrumental to the general operations of our group and has been hugely helpful to this particular project on many occasions. Finally, I'd like to give a shout-out to the Executive Branch.

I would also like to thank several people at Cameca Instruments for their help in my time working on this project. Tom Kelly provided the original impetus for this project and has been an excellent advocate for this research. His enthusiasm for the applications is inspiring as well as his vision for the future of atom probe tomography. I'd also like to thank Ed Oltman who has been a helpful resource for understanding the nitty-gritty details of what is necessary to make a proper detector for atom probe microscopes. Discussions with him definitely helped me appreciate system level considerations which otherwise would've been buried by the minutiae of day to day experimental concerns. Yimeng Chen was always willing to take a couple of hours with me to electropolish fresh tungsten wires to use in our field ion microscope.

My time in Madison has been made so much more enjoyable by the friendships I've developed in my time here. Scott Mahoney, Justin Lin, and Carl Nelson have all been great friends, and our times playing board games, running around downtown playing a

broken augmented reality game, going to the Dells, and hitting up Geek-kon every year will all be cherished memories. Becky Lipsitz and Wes Jackson have also been excellent friends, providing a great outlet for just hanging out, talking about music, spending time with Bowser and Space Ghost, or watching incredibly bad movies.

I would like to thank my family as well; although they are so far away, my mom and dad have been extremely supportive throughout my entire educational career, including these past several years in graduate school. My mom in particular has always pushed me to be the very best I can be, and I really appreciate it.

Finally, I would like to thank my wife, Julia, who has been with me throughout graduate school and more, and has always been there for me even when things may have seemed bleak. Her support and caring have been integral not only to my success in graduate school, but to my overall well being these past 7 years. Thank you.

List of Figures

1.1	Time of Flight Spectrometry Concept	4
1.2	Atom Probe Concept	6
1.3	Example of 3D Atom Probe Data	7
1.4	Field Ion Microscope Data Example	7
1.5	Microchannel Plate coupled to a Delay Line Anode	9
2.1	Transition Edge Sensor Operating Regime	16
2.2	Transition Edge Sensor Equivalent Circuit	17
2.3	Illustration of SNSPD with Typical Readout Circuit	18
2.4	”Normal-Core” Photodetection Mechanism	20
2.5	Equivalent Circuit of a Standard SNSPD	21
2.6	Simulated Electrical Output from SNSPD	24
2.7	Frequency Multiplexed SNSPD Concept	28
2.8	Superconducting Delay Line Detector Concept	30
2.9	Superconducting Microstrip Geometry	31
2.10	Superconducting Delay Line Detector Avalanche Mode Equivalent Circuit	32
2.11	Simplified Equivalent Circuit Used in Hot Spot Dynamics Simulation . .	37
2.12	Hot Spot Stability Simulation Results	38
2.13	Hot Spot Recovery Simulation	39
2.14	Linear Mode Equivalent Circuit	43
2.15	Simplified Linear Mode Simulation Circuit	44

2.16	First Attempt at Linear Mode Simulation	45
2.17	Linear Mode Simulation	46
2.18	Helium in Silicon Nitride Stopping Power Example	48
2.19	5 keV Helium in Silicon Nitride Range	49
2.20	5 keV Helium in Niobium Damage Events	50
2.21	20 keV Tungsten in Niobium Damage Events	51
2.22	Ion Absorption and Phonon Emission in Encapsulation Layer	52
2.23	Example Transfer Function from Solid Angle to Detection Probability . .	55
2.24	Detection Probability versus Ion Impact Position (Solid Angle Model) . .	56
2.25	Pseudo-1D Model Fit to Solid Angle Detection Probability	57
2.26	Illustration of Ion Absorption Between Pitch Experiment Wires	58
2.27	Probability of Detecting an Ion with: $D_{max} = .9$, $w_{eff} = .75\mu\text{m}$, $\tau =$ $.1\mu\text{m}$, $s = .75\mu\text{m}$	59
2.28	Correlation Fraction versus spacing with the following parameters: $D_{max} =$ $.9$, $w_{eff} = .75\mu\text{m}$, $\tau = .1\mu\text{m}$	61
2.29	Pitch Experiment example layout. The length and number of meanders have been significantly reduced from their true values for ease of visual- ization.	62
2.30	Pitch Experiment Curve Fitting Monte Carlos Results	63
2.31	Illustration of Thermal Crosstalk from a Latched Detector	64
2.32	Illustration of Geometry Labels in Pitch Imager Derivation	65
2.33	Lag vs Gap, Parabolic Heat Model	66
2.34	Pitch Imager Concept	68

2.35	Timing Uncertainty due to Finite Signal to Noise Ratio in Threshold Detection	70
3.1	PURAIR Flow Hood in Lab	78
3.2	Aluminum Etch Stop Degradation	81
3.3	Lithography Overview	86
3.4	Lithography Defects Due to Suboptimal Developer Choice	87
3.5	Improved Lithographic Performance	88
3.6	Process Flow Overview	91
3.7	Detector Cross Sectional Image	97
3.8	Critical Current Density Test	101
3.9	Critical Current Density vs. Length Ion Mill	102
3.10	Critical Current Density vs. Length Ion Mill	103
3.11	Niobium Etch BCl ₃ +Cl ₂ +Ar Plasma Spectroscopy	104
4.1	Circuit used for measuring current-voltage curves	112
4.2	Circuit used for performing time domain reflectometry	113
4.3	Field Ion Microscope Ion Source	117
4.4	Field Ion Image Example	118
4.5	Field Ion Microscope Ion Source Picture	119
4.6	Wiring Diagram for Ion Detection Experiments	120
4.7	Basic Ion Detector Wire Bonded to Sample Mount	123
4.8	Event Rate vs FIM Voltage Bias Measured with MCP	126
4.9	Relative Delay Resolution Measurement Setup	127
4.10	Illustration of Pitch Experiment Concept	128

4.11 Pitch Experiment Sample Mount	129
4.12 Example of Simulated Waveforms for Time Tagging Tests	134
4.13 Example of Calculated Cross Correlation For Time Tagging	136
4.14 Comparison of Plain Cross Correlation Technique with Optimally Filtered Cross Correlation	137
5.1 Current Voltage Relation of Weakly Constricted SCDLD	141
5.2 Current Voltage Relation of Weakly Constricted SCDLD	142
5.3 Time Domain Reflectometry of Good Detector	144
5.4 Time Domain Transmissometry Equivalent Circuit	145
5.5 TDR Measurements of various SCDLD Failure Modes	147
5.6 Pulse Waveform Example	149
5.7 Pulse Waveform With Resistive Matching Network	151
5.8 Filtered Waveform Example	153
5.9 Measured Crosscorrelation Function	154
5.10 Relative Delay Histogram	156
5.11 Pulse Amplitude Distribution	157
5.12 Event Rate vs Detector Current Bias	158
5.13 Event Rate vs FIM Voltage Bias	159
5.14 Relative Delay Resolution Experiment Results	161
5.15 Relative Delay Resolution Experiment Results	163
5.16 Separately Biased Pitch Experiment Relative Delay Histograms	165
5.17 Pitch Experiment Relative Delay Correlations	166
5.18 Pitch Experiment Relative Delay Histograms	167

5.19 Pitch Experiment Correlated vs Uncorrelated Event Histogram	168
5.20 Correlation Fraction vs Pitch	169
5.21 Event Lag vs Pitch	169
5.22 Event Lag vs Pitch Model Fit	170
5.23 Pitch Imager Approximate Transfer Function	172
5.24 Pitch Imager FIM Tip Images, Set 1	173
5.25 Pitch Imager FIM Tip Images, Set 2	174
5.26 Pitch Imager FIM Tip Images, Set 3	175
5.27 Multilayer Pitch Experiment	177
A.1 Simplified Linear Mode Simulation Circuit	183
B.1 Typical Heat vs Time Curve	190

Contents

Abstract	i
Acknowledgements	iii
1 Introduction	1
1.1 Thesis Outline	1
1.2 Ion Detection Applications	3
1.2.1 Time of Flight Mass Spectrometry (TOF-MS)	3
1.2.2 Atom Probe Tomography (APT)	4
1.2.3 Field Ion Microscopy (FIM)	7
1.3 Ion Detector Technology	8
1.3.1 Microchannel Plates	8
1.3.2 Detector Performance Metrics	10
2 Detector Physics and Operation	14
2.1 Overview of Superconducting Detectors	14
2.2 Superconducting Nanowire Single Photon Detector Background	18
2.3 Superconducting Delay Line Detector Concept	29
2.3.1 Avalanche Mode	31
2.3.2 Linear Mode	42
2.4 Ion Interactions With Matter	47
2.4.1 Phonon Mediated Detection Scheme	51

2.5	Multi-Detector Experiments	53
2.5.1	Pitch Experiment Concept and Detection Efficiency Model	53
2.5.2	Pitch Imager Concept	64
2.6	Sources of Uncertainty in Timing Measurements	68
2.6.1	Uncertainty Due to Electrical Noise in Readout	69
2.6.2	Uncertainty Due to Velocity Noise in Detector	70
2.6.3	Time Varying, Spatially Homogeneous Inductance	72
2.6.4	Stationary, Spatially Varying Inductance	73
2.6.5	Time varying, spatially varying inductance analysis	74
3	Device Fabrication	76
3.1	Overview	76
3.2	Superconductor Deposition	76
3.2.1	Kurt-Lesker Sputter System	76
3.2.2	TLI Sputter System	77
3.2.3	PURAIR Flow Hood	78
3.3	Dielectric Deposition	79
3.3.1	PlasmaTherm 70	79
3.3.2	Denton RF Sputter System	82
3.4	Lithography	83
3.4.1	Karl Suss MA6 Contact Aligner	84
3.4.2	Nikon Body 8 i-Line Stepper	84
3.5	Dry Etching	89
3.5.1	Unaxis 790	89

3.5.2	Plasma-Therm 770	90
3.5.3	Kaufman & Robinson Ion Mill	90
3.6	Process Flow Overview	91
3.6.1	Niobium Ground Plane	92
4	SCDL Measurement	110
4.1	Measurement Setup Overview	110
4.2	Room Temperature Screening	111
4.3	Liquid Helium Screening	111
4.4	Detection Experiments	114
4.4.1	Ion Sources	115
4.4.2	Electronics	119
4.4.3	Basic Ion Detection Experiments	125
4.4.4	Relative Delay Resolution Measurement	126
4.4.5	Pitch Experiment	127
4.5	Data Analysis Procedures	129
4.5.1	Optimal Filtering	130
4.5.2	Time Tagging Algorithms	131
5	Superconducting Delay Line Detector Results	138
5.1	Liquid Helium Dip Results	138
5.1.1	Current Voltage Measurement Results	138
5.1.2	Time Domain Reflectometry Results	143
5.2	Single Wire Ion Detection Experimental Results	148
5.2.1	Time Domain Waveforms	148

5.2.2	Relative Delay Histogram	155
5.2.3	Pulse Amplitude Distribution	155
5.2.4	Event Rates	157
5.2.5	Relative Delay Resolution Results	160
5.3	Pitch Experiment Experimental Results	164
5.4	Pitch Imager Experimental Results	171
5.5	Conclusions and Outlook	176
A	Linear Mode Response	179
A.1	Voltage Due to time Varying Cooper Pair Density	179
A.2	Voltage Response of Linear Mode Equivalent Circuit	183
B	Parabolic Heat Equation between Two Closely Spaced Detectors	187
	Bibliography	191

Chapter 1

Introduction

1.1 Thesis Outline

Although there are several variations of experiments which will be presented in this thesis, the chapters are arranged functionally rather than temporally. The process of studying this new technology has not always been linear, but every effort will be made to reconstruct the lessons learned in a coherent, digestible order for the reader. Chapter 1 begins with an outline of the applications we are interested in developing superconducting detectors for. This section concludes with a description of the current state of the art technology for the detection of ions in these applications and its performance metrics. Following this, an overview of the performance parameters we are interested in addressing will be given.

Chapter 2 begins with a brief survey of superconducting detector technologies, followed by a more in-depth explanation of the operational principles of superconducting nanowire single photon detectors (SNSPDs). This leads into the explanation of how the superconducting delay line detector (SCDLD) varies from traditional nanowire detectors. From there, the physics of detection and reset will be explored for this device design. After covering these general theoretical considerations, the motivation and theoretical foundation for the experiments carried out on these detectors will be given. These

sections are meant to instill intuition in the reader to better digest the results presented in Chapter 5. This chapter concludes by considering possible sources of uncertainty in our identification of ion event position.

In Chapter 3, a thorough, step by step explanation of the fabrication process for making SCDLDs will be presented. The tools used in this process are described, followed by a general process guide. Throughout, special attention will be paid to the various pitfalls which can severely limit the success of this fabrication process. Other than the deposition of metal films and ion milling, all of this work was performed in the Wisconsin Center for Applied Microelectronics. This on campus cleanroom has been a critical resource in the success of this work.

In Chapter 4, the tools and setups used for measuring our detectors will be described. This chapter begins with the screening process used to quickly distinguish useful from deficient detectors. After this multistage screening procedure is detailed, the focus of this chapter will shift to characterization of SCDLD in ion detecting experiments. First, basic experiments involving a single detector will be explained. Finally, two experiments involving multi-detector designs are outlined with an explanation of the goal in each case.

In Chapter 5, the results of all of the experiments described in Chapter 4 will be presented and explained. We begin with some screening data, showing representative samples from both successful and unsuccessful detectors. The author's hope in doing so is that if future scientists/engineers would like to try and continue this work, some of the potential failure modes of these detectors will be known ahead of time, saving them both time and frustration. The results of characterizing single detectors with an ion source are then presented. Much of this characterization (and the methods for understanding it) will carry over into the following sections as well. In these later sections, we will

describe the results of the multi-detector experiments, dubbed the “Pitch Experiment” and “Pitch Imager”. Finally, we will provide an overview of these results and the outlook for future work in this technology.

1.2 Ion Detection Applications

1.2.1 Time of Flight Mass Spectrometry (TOF-MS)

The goal of Time of Flight Mass Spectrometry is to measure the mass to charge ratio of a substance. This is achieved by following the steps shown in Figure 1.1. First, samples are prepared and placed in a high vacuum environment. An electric field gradient is created by voltage biasing the sample container (on the left side of each panel) with respect to a drift tube (the cylinder in the center of each panel). At controlled times, bits of a sample are ionized (in the illustrated example, this is achieved by applying a short laser pulse). This marks the beginning of the flight time. The generated ions are then accelerated down the potential difference between the sample mount and the drift tube. Once they enter the drift tube, they continue traveling with a constant velocity. Once they exit the tube, they are detected by a particle detector, which then signals the end of the flight time. Taking the approximation that the time spent in the drift tube is much longer than the time taken to reach full velocity in the initial part of the path, the mass to charge ratio can be calculated from simple kinematics. This results in Equation 1.1.

$$\frac{m}{n} = 2e^{-}V_{tip}\left(\frac{t_{flight}}{f_{path}}\right)^2 \quad (1.1)$$

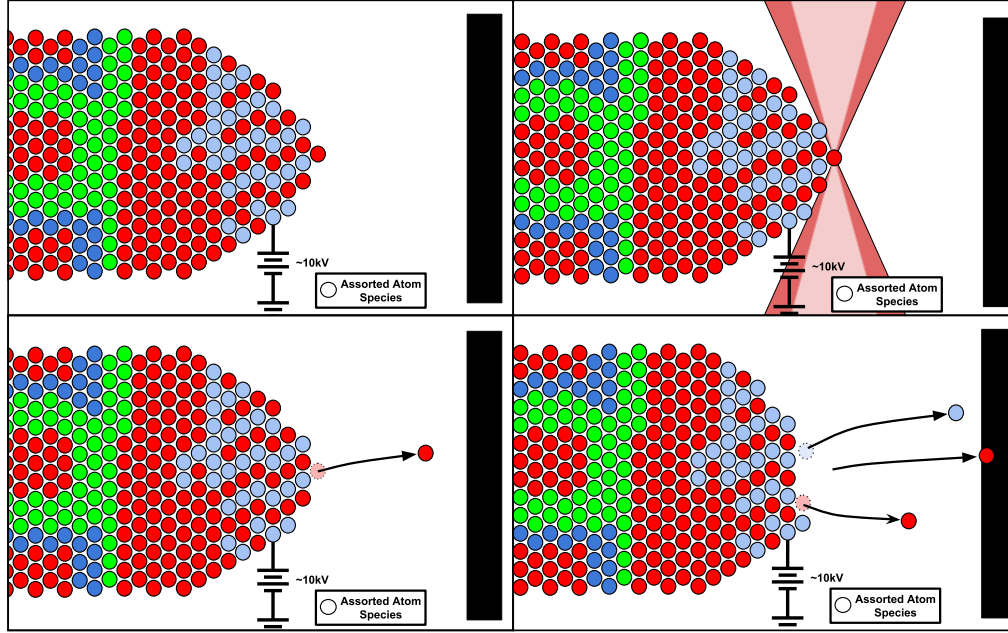


Figure 1.1: Time of Flight Spectrometry Concept

Where m is the mass of the ion, n is the charge state of the ion, e^- is the charge of an electron, V_{tip} is the voltage bias on the microtip, t_{flight} is the time of flight for the ion, and f_{path} is the length of the flight path from the sample to the detector.

1.2.2 Atom Probe Tomography (APT)

Atom Probe Tomography is a specialized, destructive form of microscopy[19]. The general idea is illustrated in Figure 1.2. First, a sample is prepared by shaping the original material or device into a long, narrow cylinder. This cylinder must then be sharpened to a microtip with a fine point (with a radius of order 50–200 nanometers). This sample is then loaded into an ultra high vacuum (UHV) environment and aligned in front of a detector. A large voltage bias is applied to the sample microtip such that the materials which comprise the tip are subject to an electric field just below the strength at which

they would field evaporate. This is illustrated in the upper right panel of Figure 1.2. The sample tip is then subjected to energetic perturbations at controlled times. This can be brought on by a laser pulse, or by pulsing the strength of the electric field temporarily. This perturbation provides enough energy that there is a chance of field evaporating an atom off of the tip. The goal is to only field evaporate at most a single atom every time; achieving this in practice is often difficult. This ionized atom is accelerated down the applied electric field and is collected by a detector, as shown in the bottom left panel of Figure 1.2. The detector needs to be capable of measuring the time from the energy perturbation to the detection of the ion as well as the position where the ion hit the detector. This process is repeated many times over, so that subsequent layers of atoms in the sample tip are exposed to the evaporating field. By collecting data about the location where the ions hit the detector, their trajectory can be calculated reconstructing where they came from on the sample tip. By keeping track of the order in which ions reach the detector and using information about their transverse location, an approximation of the depth in the sample they came from can be reconstructed. Using the time delay from the energy perturbation to when each ion was detected, the mass to charge ratio can be calculated for each ion can be calculated using Equation 1.1.

This collected data can then be used to build an atom by atom 3D reconstruction of the original sample, including the charge to mass ratio of each atom (which in many, but not all cases can be used to determine that atom's species). An example of such a 3D map is shown in Figure 1.3. This process can be used to analyze many different kinds of materials[18] and has found applications in geology[48], thin film metrology[15], and more. Although atom probe tomography has reached a point where it is being actively used for materials analysis, there are still many ways in which it could be improved[30].

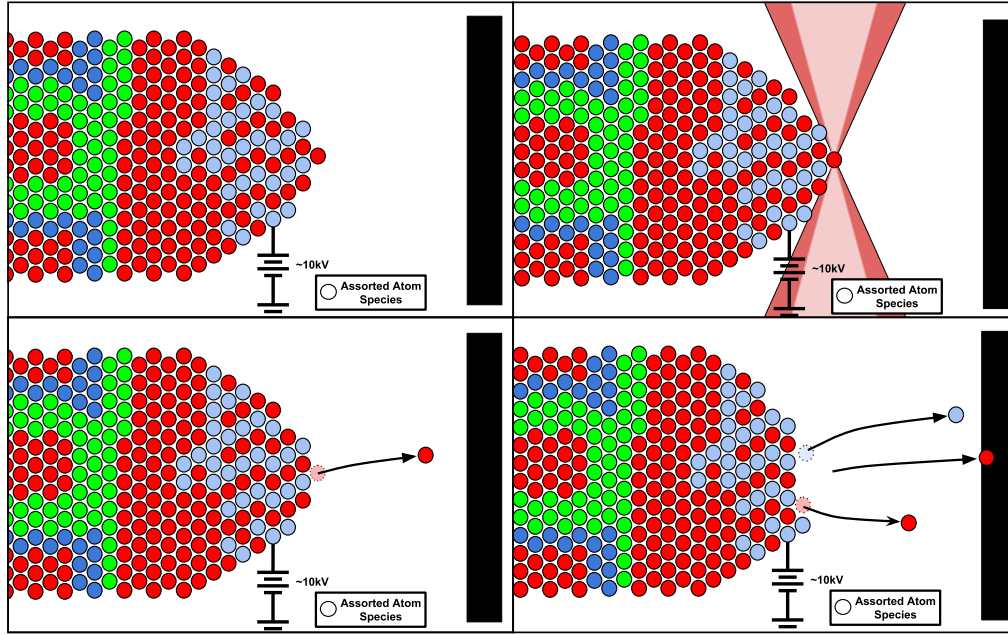


Figure 1.2: Atom Probe Concept

The prospect of using superconducting detectors to improve the performance of atom probes is not a new idea[30]; the focus of this thesis will be on addressing deficiencies related to the atom probe's use of microchannel plates as the main detection element.

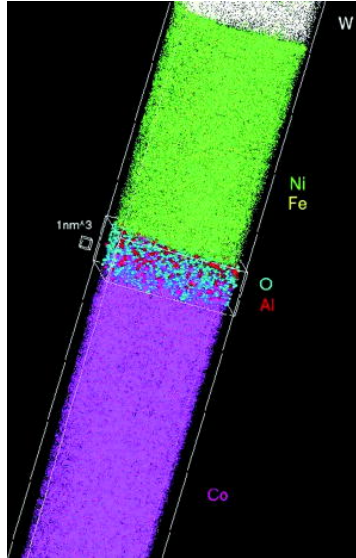


Figure 1.3: Example of 3D Atom Probe Data: Taken from Ultramicroscopy **101**[23]

1.2.3 Field Ion Microscopy (FIM)

Field Ion Microscopes can be thought of as the spiritual predecessors to atom probes. The original field ion microscope was constructed by Erwin Müller in 1951[32] and its usage marked the first direct observation of individual atoms[40]. The physical method involved in generating an image using FIM will be described in depth in Section 4.4.1.2. An example of data taken using FIM is shown in Figure 1.4.

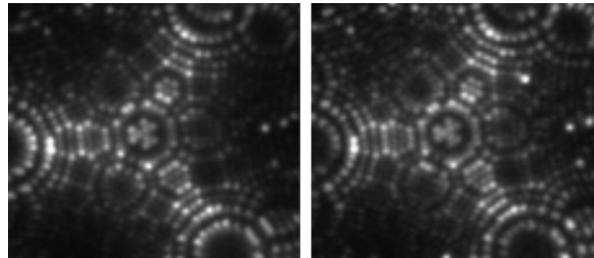


Figure 1.4: Field Ion Microscope Data Example: Taken from Ultramicroscopy **122**[47].

1.3 Ion Detector Technology

1.3.1 Microchannel Plates

The current state of the art technology for detecting ions over a large active area, with position sensitivity is the Microchannel Plate (MCP). In Figure 1.5 the standard detection configuration is shown where a microchannel plate is placed in front of a delay line anode. The microchannel plate itself is on top in this image; it consists of a plate of highly resistive material with a series of holes through it, spaced periodically across the face. A voltage bias is applied between the front and back faces of the plate. When an ion strikes the inner wall of one of the channels going through the plate, it releases secondary electrons. The number of secondary electrons generated follows a Poisson distribution[11] with a mean yield of λ_{se} which depends on the particle energy and species being detected. These secondary electrons are accelerated down the channel by the applied voltage bias. Then, each secondary electron collides with the channel walls, generating its own cascade of secondary electrons. This happens several times before the resulting cloud of electrons is ejected out the bottom of the microchannel plate, creating a charge multiplication effect whereby the incoming ion generates of order hundreds of thousands to tens of millions of secondary electrons (depending upon the design of the MCP).

This charge cloud is then collected by a pair of perpendicularly oriented delay line anodes, thus producing a pair of voltage pulses on each anode which propagate to either end of their respective delay line. Measuring the time of arrival of these pulses at either read out port of the anode, τ_1 and τ_2 , the ion arrival time can be calculated as:

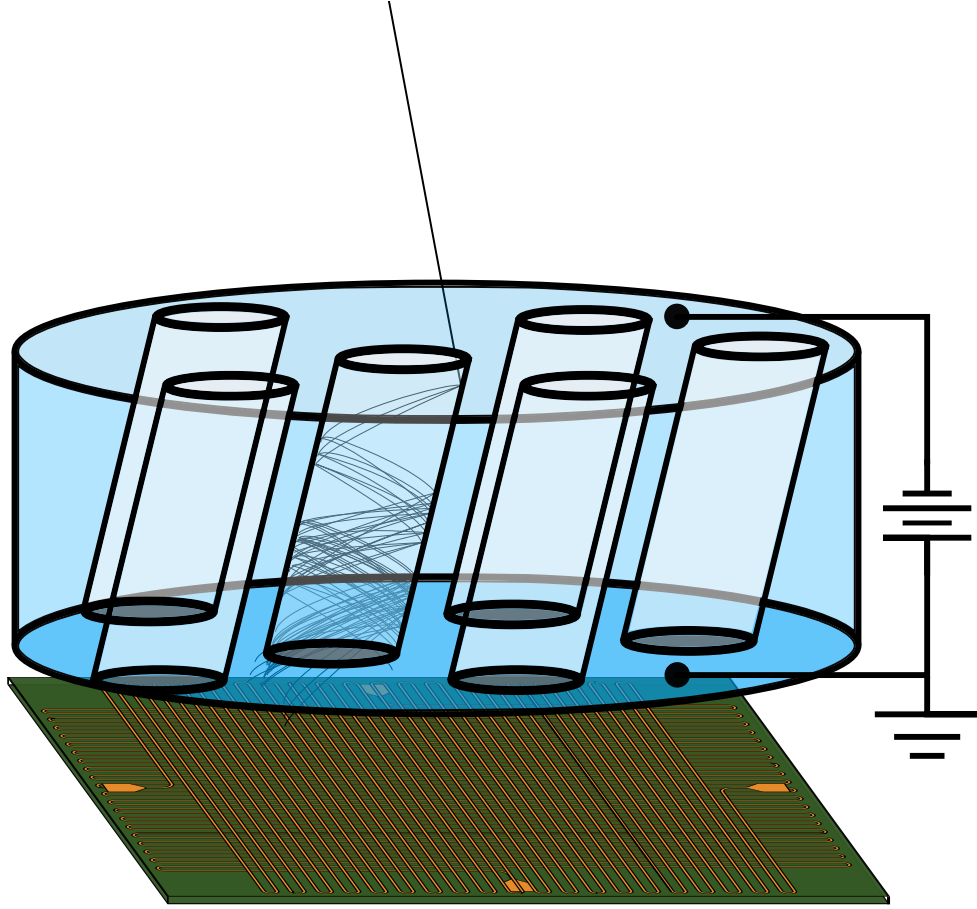


Figure 1.5: Microchannel Plate coupled to a Delay Line Anode

$$\tau_{ion} = \frac{\tau_1 + \tau_2}{2} - \tau_{delayline} \quad (1.2)$$

where $\tau_{delayline}$ is the total electrical length of the delay line. These times are measured relative to the time when the ion was generated. The position along the delay line where the charge cloud was collected, x_{ion} can be calculated as:

$$x_{ion} = l_{delayline} \frac{\tau_1 - \tau_2}{2\tau_{delayline}} \quad (1.3)$$

where x_{ion} is measured from the first read out port, and $l_{delayline}$ is the physical length

of the anode.

1.3.2 Detector Performance Metrics

1.3.2.1 Position Resolution

Position resolution is a measure of how accurately we can determine where on a detector an event occurred. The way this manifests in a particular detector depends on its design; for instance, if a detector is actually made up of an array of smaller detectors which do not have position sensitivity on their own, but can be separately read out, then the position resolution is given by the pitch of the detector pixels (that is, the spacing from pixel to pixel). In the case of both MCPs and the superconducting delay line detectors studied in this thesis, the position sensing mechanism is a mapping from timing information to positional. It is therefore necessary to calculate/measure the uncertainty in our assignment of a "time of arrival" of pulses coming out of the delay line and then map this to an effective uncertainty in position.

1.3.2.2 Timing Resolution

In addition to the position uncertainty induced by our uncertainty in the relative delay between pulses coming out of the delay line, there is also overall timing uncertainty in Equation 1.2. This is an important parameter, as the ability to accurately measure the time of flight of an ion is crucial for calculating the mass to charge ratio. This parameter is going to be related to the aforementioned position uncertainty; if errors in our timing estimation are correlated between the two readout channels of the delay line, the relative delay calculation will subtract out those errors, reducing the uncertainty in

that measurement compared to the time of flight measurement. If the errors between the two channels are uncorrelated, then the relative delay calculation will have essentially the same uncertainty as the time of flight calculation.

1.3.2.3 Detection Efficiency

There are two different types of detection efficiency to consider. The first is internal detection efficiency; this refers to the following situation: assuming a particle (or its energy) has been absorbed in the active portion of a detector, the internal detection efficiency is the probability that the detector generates a measurable response. The second is system detection efficiency. In this case, if a particle is generated as part of an experiment, the system detection efficiency is the probability that the detector generates a measurable response regardless of where the ion hits¹. Consider the case of the microchannel plate. If the active area is defined as the channel interior of the plate, the internal detection efficiency is given approximately by $D_{intern} = 1 - e^{-\lambda_{se}}$, where λ_{se} is the mean secondary electron yield for the particle being detected. Assuming the ion image being projected lands fully within the microchannel plate, the system detection efficiency (ignoring effects such as ion feedback) is approximately $D_{sys} = \frac{A_{channels}}{A_{ionimage}}(1 - e^{-\lambda_{se}})$, where the As are areas of the MCP channel openings and ion image.

¹Or even if it is deflected before reaching the detector; for instance, if an electrostatic grid is used to change the ion flight paths, some portion of the ions will collide with the grid and never reach the detector.

1.3.2.4 Multi-Hit Discrimination

This parameter deals with the situation where a detector is hit by two events which are closely spaced in time. For instance, if a microchannel plate is hit by two ions in two different channels, it is capable of multiplying both and creating two amplified charge clouds. However, the generated clouds have a reasonably large spatial extent, and the delay line anodes typically used with MCPs have somewhat poor high frequency properties. This means that the pulses generated by these clouds often end up overlapping so that it is impossible to distinguish that there were multiple events. This is exacerbated by the fact that the gain for each detection event is drawn from a relatively wide distribution (the standard deviation of this distribution is often about equal to the mean of the distribution). In an ideal detector, it would be possible to distinguish two events no matter how closely spaced they are in time and/or space.

1.3.2.5 Kinetic Energy Discrimination

In some cases, knowledge of the mass to charge ratio is not sufficient to distinguish between different species of ions. For instance, the mass to charge ratio of doubly ionized silicon is essentially equal to the mass to charge ratio of singly ionized nitrogen. If a detector is energy dispersive (that is, the output varies with particle energy) this degeneracy can be broken, as the silicon ion will have twice the kinetic energy of the nitrogen. It would be beneficial, then, to have a detector with at least modest kinetic energy discrimination. One figure of merit for energy resolution is the Energy Resolving Power which is defined as the mean energy of a peak divided by the full width at half maximum of an energy peak. In most atom probe experiments, we expect a maximum

ion charge state of approximately 4, and so having an energy resolving power of 4 or better would be extremely useful.

1.3.2.6 Parameter Summary and Targets

In the table below, the parameters discussed above will be listed along with the typical performance of existing microchannel plates and the target we would like to achieve with a superconducting detector.

Parameter	MCP Value	Target Value
Position Resolution (Pixels)	10^4 – 10^5	10^6
Active Area	40 cm	>9 cm
Detection Efficiency	67% (mass dependent)	100%
Maximum Count Rate	10-100 kHz	1 MHz
Energy Resolving Power	None	≈ 4
Timing Resolution	≈ 100 ps	≈ 20 ps

Table 1: Detector Figures of Merit

Chapter 2

Detector Physics and Operation

2.1 Overview of Superconducting Detectors

We seek to design a superconducting detector which is sensitive to ions in the kiloelectron volt energy range. There are numerous different strategies for detecting energetic particles/photons with superconductors; however, they can be loosely boiled down into two different categories: measurement of modulations in the device's inductance or resistance.

An example of a detector which primarily measures modulations in inductance¹ is microwave kinetic inductance detectors (MKIDs). MKIDs consist of a high-Q resonator made out of superconducting material; when a particle is absorbed in this resonator, non-equilibrium quasiparticles are formed which alter the kinetic inductance of the material. This causes the center frequency and quality factor of the resonator to shift which can easily be measured by taking a measurement of the scattering parameters of the resonator at or near its equilibrium resonance frequency. The magnitude of this shift is a smooth function of the amount of energy the particle imparts to the detector. MKIDs

¹More precisely, MKIDs measure the modulation of both the real and imaginary parts of a superconductor's conductivity. In this sense, they are actually sensitive to both variations in inductance and resistance at their resonance frequency. See [26] for more information.

can therefore be used as energy dispersive detectors (if properly designed).

One example of a detector architecture which primarily measures the modulation of a superconductor's resistance is the transition edge sensor (TES). TESs consist of a strip of superconducting material which is held in a state which is a mixture of superconducting and normal. This is achieved by operating very near the material's critical temperature (T_c) so that its measured resistivity is approximately half of its residual resistivity² as shown in Figure 2.1.

The detector is connected in series with a current sensitive element and in series with a shunt resistance as shown in Figure 2.2. A bias current is applied to the detector. When the temperature of the detector changes, its resistance shifts, causing some of the current going through the device to be shunted out through the shunt resistor, thus changing the current flowing through the ammeter. In a well designed device, the decrease in current going through the device will cause the device to cool back to its equilibrium operating temperature. In this way, the device generates a signal that is proportional to its temperature change (and therefore approximately proportional to the energy deposited by the particle/photon) and then self-resets. More detailed information about the operation of TESs can be found in the literature; for example, [27] and [46].

Although both TESs and MKIDs have the appealing property of being energy dispersive (and with extremely good energy resolution no less), there are issues with trying to use them in this application. Firstly, in both cases, it is difficult to imagine a detector arrangement which could avoid the same pitfalls associated with microchannel plates; that is, filling the entire area of the detector (or very nearly so) with active sensing

²In actual use, it is beneficial to work a bit below this halfway point in order to increase the dynamic range available without significantly degrading the linearity of the detector.

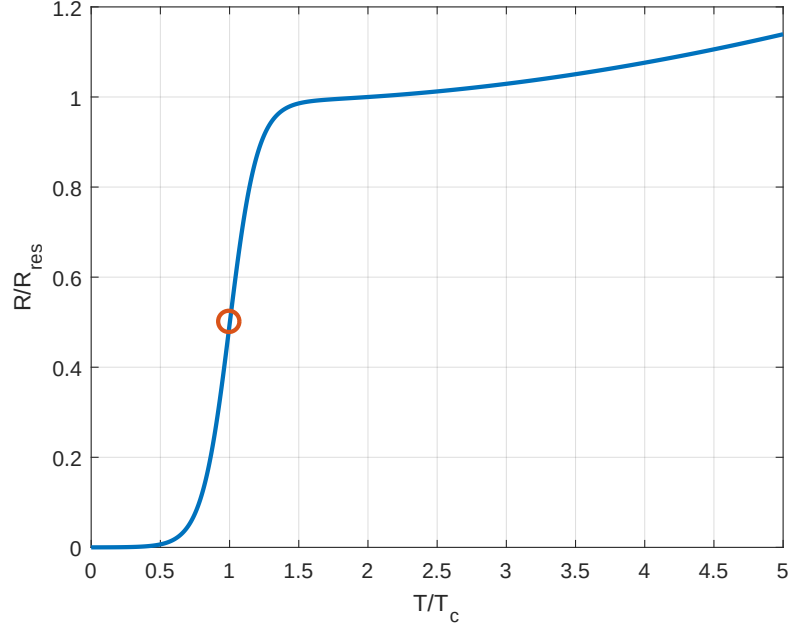


Figure 2.1: Shown above is a simple example of a superconductor's resistance vs temperature characteristic (normalized to the materials residual resistance and critical temperature, respectively). An orange circle marks the equilibrium temperature the detector should be held at in order to operate effectively as a transition edge sensor. At this point, even minute changes in temperature will cause the resistance of the device to change drastically.

elements. In principle, this could be overcome. However, another, more practical issue, is that these detectors must be operated at very low temperatures (typically in the range of 50-100 millikelvin). Reaching these operating temperatures require involved refrigeration technology such as adiabatic demagnetization refrigeration (ADR) or dilution refrigeration (DR). For the ion detecting applications of interest, we constrain our focus to technologies which can be implemented using only pulse tube coolers; that is, technologies that can operate at temperatures around 3 Kelvin or higher. Pulse tube coolers

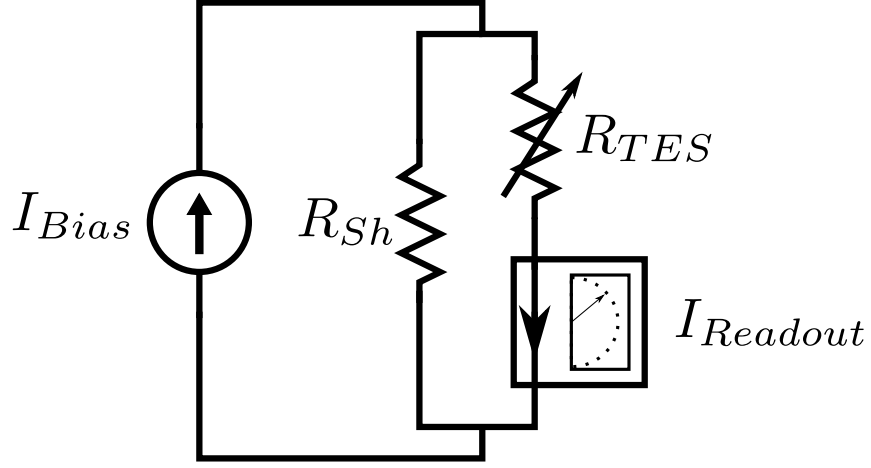


Figure 2.2: Simplified Transition Edge Sensor Equivalent Circuit

are have reached a maturity to where they are very simple to operate, require limited maintenance, and are closed cycle, so cryogenics do not need to be provided. They can also be operated continuously, unlike ADRs.

The final type of detector to consider is called a superconducting nanowire single photon detector (SNSPD). In contrast to the other two detectors, the response of an SNSPD is not energy dispersive; if an incident particle/photon has sufficient energy, and event will be triggered. Otherwise, the detector does not respond. It is possible, however, to vary the bias current in the SNSPD which shifts the minimum energy necessary to trigger an event[42][43], allowing the detector to make some statistical statements about the energy distribution of incident particles. The physics of this detector architecture will be explained in depth in Section 2.2.

2.2 Superconducting Nanowire Single Photon Detector Background

The development of our superconducting delay line detector (SCDLD) is based loosely on previous work on superconducting nanowire single photon detectors (SNSPDs or SSPDs). These detectors were an attractive starting point for several reasons; they are extremely fast, sensitive to single photons of relatively low energy [24], able to operate at high detection rates [20], and they exhibit very low timing jitter[22]. SNSPDs have been previously used to detect large, biological ions in the keV energy range[41][53][5]. Additionally, SNSPDs have been used to detect monoatomic species[6] and have been shown to have an internal detection efficiency of roughly 100%[37].

An SNSPD consists of a very narrow and thin superconducting wire ($\sim 10\text{nm} \times 100\text{nm}$) which is biased near its critical current connected to a read-out circuit as illustrated in Figure 2.3.

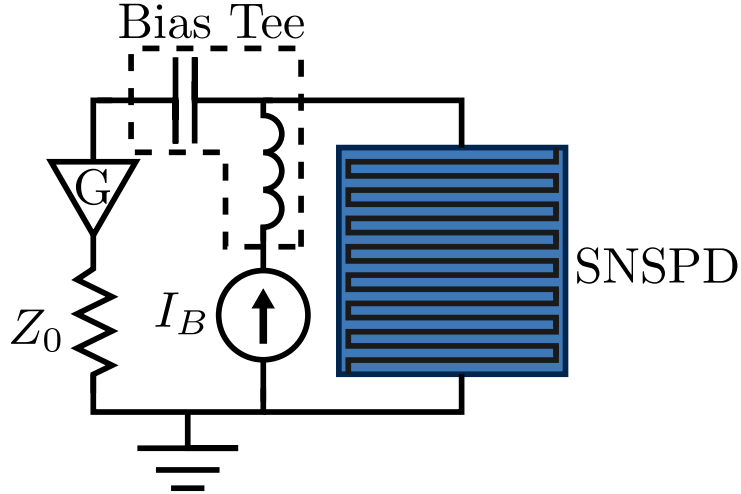


Figure 2.3: Illustration of SNSPD with Typical Readout Circuit

Incident photons are absorbed in the SNSPD and break up cooper pairs and form a non-equilibrium population of quasiparticles. There are several models for the steps following. In the simplest explanation, illustrated in Figure 2.4, the energy absorbed by the photon fully destroys superconductivity in a small region. This “normal-core” model is an excellent starting point for understanding the dynamics of photodetection in these devices; however, in recent years more sophisticated models have been developed [8] which can more robustly predict properties of SNSPDs such as dark-count rates, energy threshold variation with bias current, et cetera. These models are focused on the detection of relatively low energy photons, and thus are beyond the scope of this thesis. Furthermore, some work on the detection of ions in superconducting nanowires seems to confirm this normal-core model’s applicability[43].

There are, however, a number of aspects of the SNSPD design which are suboptimal when used for ion detection applications. In particular, the cross section of an SNSPD is typically of order 4 nanometers by 100 nanometers and it is meandered to cover an area of order $100\text{ }\mu\text{m}^2$ to $1000\text{ }\mu\text{m}^2$ [28]. We are interested in detecting particles which have several thousand times as much energy as the $1550\text{ }\mu\text{m}$ photons which SNSPDs are usually trying to detect. We therefore have the latitude to design detector with somewhat larger cross sectional areas (thus allowing us to cover a larger active area without sacrificing detection efficiency). Additionally, the flux of ions which are generated in a typical time of flight experiment (whether it be TOF-MS, APT, or FIM) is quite different from the flux of photons coming out of a single mode fiber (a common source of infrared photons for SNSPD experiments). For instance, the density of photons passing through an SNSPD is often many orders of magnitude higher than the density of ions

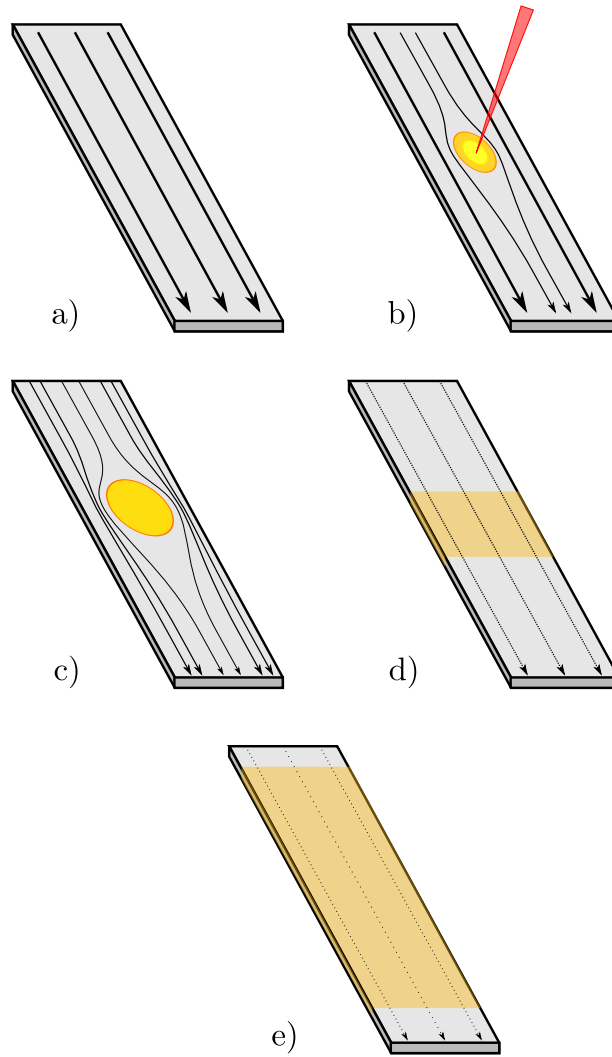


Figure 2.4: "Normal-Core" Photodetection Mechanism: a) A superconducting wire is biased with a constant current near its critical current. b) A photon is absorbed in the wire, breaking up cooper pairs and forming a central core where the wire is no longer superconducting. c) The bias current redistributes around this normal core. d) In the strips around the normal core, this redistributed current exceeds the critical current density, thus causing a normal zone to form that spans the entire cross section of the wire. The resulting increase in resistance causes the bias current flowing through the wire to decrease. e) The residual current that is flowing through the wire dissipates heat in the normal zone, causing it to grow to some equilibrium size.

we can produce. The number of ions produced initially is much lower, and it is difficult (and often counterproductive³) to try and focus them onto a tiny spot size like the one generated by a single mode fiber. All of this is to say that when one is interested in developing an ion detector, covering a large active area is absolutely necessary. For comparison, MCPs often have an active area of order 10 cm^2 or more.

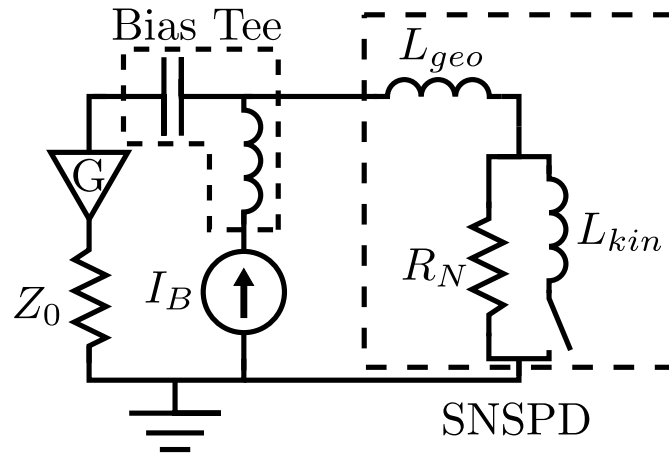


Figure 2.5: Equivalent Circuit of a Standard SNSPD

³Ion optics is an entire field of study and engineering unto itself; there are many different methods for focusing or defocusing ion trajectories, but they invariably induce aberrations in the image[7], and, in the case of electrostatic grids, can induce a loss in system detection efficiency because they have a non-unity transmissivity. In some applications, these drawbacks are essentially fatal— in APT, for instance, each ion is evaporated off of the original sample, so there is only one opportunity to detect it. If the ion instead collides with an electrostatic grid before reaching the detector, that information is lost forever. In contrast, if one is performing TOF-MS, there are many nominally identical copies of the ion of interest, and so if some are lost to such losses this is not so critical. Additionally, in something like TOF-MS, the location where an ion hits the detector (or equivalently, the image generated by the ions) is of no concern, so aberrations also do not impact the effectiveness of the measurement.

Although arrays of SNSPDs have been produced[29][51], there are several challenges involved in trying to scale up from $100\text{ }\mu\text{m}^2$ to 10 cm^2 . The first of these is a practical issue; SNSPDs, due to their narrow width, are typically fabricated using electron beam lithography (EBL). EBL typically writes features in a serial fashion; that is, an tightly focused beam of electrons is rastered across an exposure field, applying controlled doses of charge to the sample at each point. In some advanced EBL systems, there may be several beams which can work in parallel; however, compared to optical lithography (in which the entire sample is exposed at once, no rastering involved), this process is still incredibly slow. Therefore, when scaling to much larger active areas, the time to perform just the lithographic exposure becomes daunting. As mentioned above, we are trying to detect much higher energy particles in this thesis, so we are able to work with larger cross sectional areas– and it turns out that it is possible to create our detectors using only i-Line (365 nm) optical lithography to define our detector patterns.

The second issue with scaling the SNSPD design can be understood by looking at the equivalent circuit for the SNSPD as shown in Figure 2.5. The left half of the circuit is the bias and readout electronics which are extrinsic to the actual detector, which is modeled by the components on the right side. Before a detection event occurs, the switch which is in series with L_{kin} is closed so that the bias current (I_B) flows through the detector to ground with no resistance. In this simple model, a detection event is simulated by opening the switch so that all of the bias current is forced to flow through a resistor R_N which is a stand in for the resistivity of the newly formed hotspot/normal zone. In reality, the resistance of this normal zone will be a time varying property. We will simulate these more involved dynamics for our own detector design in section 2.3.1.1. For now, though, it is sufficient to consider only the timescales imposed by the

electrical parameters. When the switch is opened (a particle/photon is detected), current is shunted out of the detector and through the RF port of the bias tee, thus generating a measurable voltage in the read out circuit (marked as a resistor with an impedance of Z_0 , typically $50\ \Omega$). In this simplistic SPICE simulation, the rise time of our voltage signal is determined by the readout impedance, the normal state resistance, and the inductance of the detector (which is usually dominated by the kinetic inductance, L_{kin}); in the real world, the electrothermal dynamics previously alluded to play a large role by causing R_N to vary as a function of time. The fall time is set by the total inductance of the detector and the read out impedance[33]. These two time scales are as follows:

$$\tau_{rise} = \frac{L_{kin} + L_{geo}}{Z_0 + R_N} \quad (2.1)$$

$$\tau_{fall} = \frac{L_{kin} + L_{geo}}{Z_0} \quad (2.2)$$

At this point, there are a couple of clear issues. Both the rise and fall time of the detector pulse scale linearly with the total inductance of the detector⁴. Because the inductance scales with length, as a detector is lengthened to cover a larger active area, two things will happen: 1) The best achievable timing resolution (for some fixed noise level) will be degraded by the increased τ_{rise} and 2) The increased τ_{fall} could lead to issues with "pile-up"—that is, when detection events are occurring so often that their pulses overlap. However, this pile-up scenario is actually not the most crucial problem that arises from a growing τ_{fall} .

⁴By dividing a large detector into multiple, smaller detectors which are connected in parallel with each other, the total inductance of the entire structure can be maintained at a reasonable level[54].

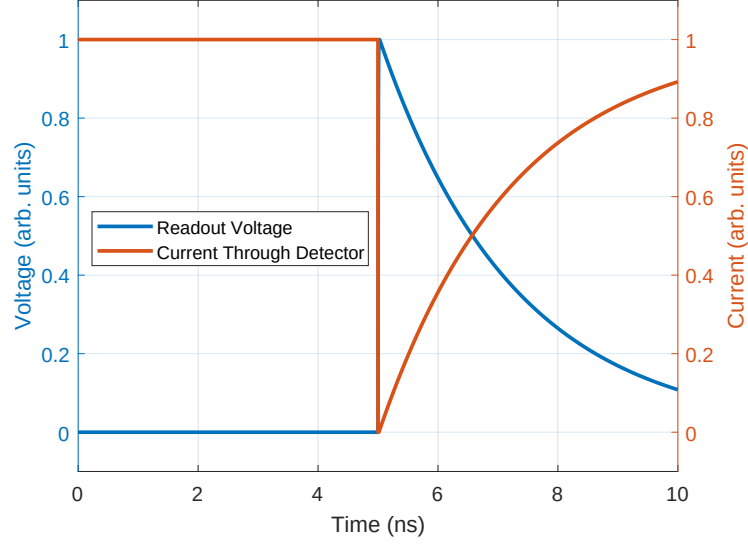


Figure 2.6: Simulated Electrical Output from SNSPD

The true significance of increasing τ_{fall} can be illuminated by calculating the current through the detector (I_{det}) after a detection event. For the following analysis, assume that the time scales involved are short enough to neglect the effects of the reactive components in the bias-tee. After the switch opens, the current through the detector branch falls exponentially towards a value $I_B \frac{Z_0}{Z_0 + R_N}$. The time constant for this fall is the same as the rise time of the voltage signal, given in equation 2.1. After some amount of time (which cannot be calculated from this simple SPICE simulation), the hotspot may recover. This process is approximated by closing the switch. The current through the detector (I_{det}) returns to its starting value of just I_B . The time scale over which the current is returned to the detector is given by equation 2.2. This means that the detector has some finite recovery time during which there is less than I_B flowing into the detector. As illustrated in Figure 2.4, the detection mechanism relies on there being sufficient bias current to generate a normal zone across the entire width of the superconducting strip.

That is, the current flowing through the strip needs to be roughly:

$$I_{det} > I_{crit} \frac{w - 2r_{normalcore}}{w} \quad (2.3)$$

where w is the width of the strip and $r_{normalcore}$ is the radius of the area where the particle/photon's energy completely destroyed superconductivity. Define a parameter β such that $I_B = \beta I_{crit}$. In a typical nanowire experiment, R_N is much greater than Z_0 so that we can take the initial value of I_{det} after the hotspot heals to be very small compared to I_B . We can now write an approximate expression for the dead time of the detector:

$$t_{dead} \approx \frac{L_{kin} + L_{geo}}{Z_0} \log\left(\frac{\beta w}{\beta w - w + 2r_{normalcore}}\right) \quad (2.4)$$

Again, this parameter scales linearly with inductance, and therefore with the length of the detector. Therefore, using the standard SNSPD design, the maximum achievable countrate will be inversely related to the length (and therefore active area) of the detector. For instance, the kinetic and geometric inductances of a nanowire can be estimated by[2]:

$$L_{kin} \approx \frac{\mu_0 \lambda_0^2 l_{wire}}{2A_{wire}} \quad (2.5)$$

$$L_{geo} \approx \frac{\mu_0}{2\pi} l_{wire} \left(\log\left(\frac{2l_{wire}}{r}\right) - 1 \right) \quad (2.6)$$

Using rough values for our envisioned device (1 μm by 40 nm cross section, made of Niobium with $\lambda_0 \approx 85\text{nm}$, this yields an inductivity of 2 $\mu\text{H m}^{-1}$ to 3 $\mu\text{H m}^{-1}$. Thus, if the design goal is to maintain the capability to detect ions at a rate of 1 MHz and

operate in a regime where we need to restore 95% of the bias current to be sensitive to these ions, the detector is limited to a length of approximately 6 m. With a fill factor of 50% (this is a very conservative estimate of the necessary fill factor), this translates to an active area of 0.12 cm per detector, thus requiring at least 150 pixels to cover the target 9 cm. At first blush, this number of pixels seems achievable. However, up to this point no consideration has been given to the added complication of position sensitivity— the target resolution for this detector is 1 megapixel, but an SNSPD conveys no information about where an event occurred within its bounds. This suggests that in order to obtain this megapixel of resolution, it would be necessary to multiplex 1 million pixels in some way.

One can envision many different strategies for multiplexing detectors of this type using strategies such as time domain multiplexing (only reading out $\mathcal{O}(1)$ pixel at a time, or using a more clever codeword scheme) or frequency domain multiplexing (embedding the detectors in a resonator or attaching bandpass filters with different frequencies to either end of the nanowire to cause each detector to ring with a distinct frequency). However, all of these techniques involve trade-offs among the following parameters: detector duty cycle⁵, readout complexity, detector layout complexity, and heatload to the cryogenic stage. We therefore seek a more scalable detector architecture which will allow us to cover a large active area while also maintaining the ability to resolve the position

⁵In time domain multiplexing, most of the detectors are off at any one time. If the readout system can chop between detectors fast enough, it is possible to take advantage of the finite length of the pulses coming out of each detector to catch more of the data being produced by the entire array. However, the benefit of SNSPDs is that they are fast; here that is a detriment because it requires the readout electronics to switch between channels more quickly, otherwise pulses are lost.

of detection events.

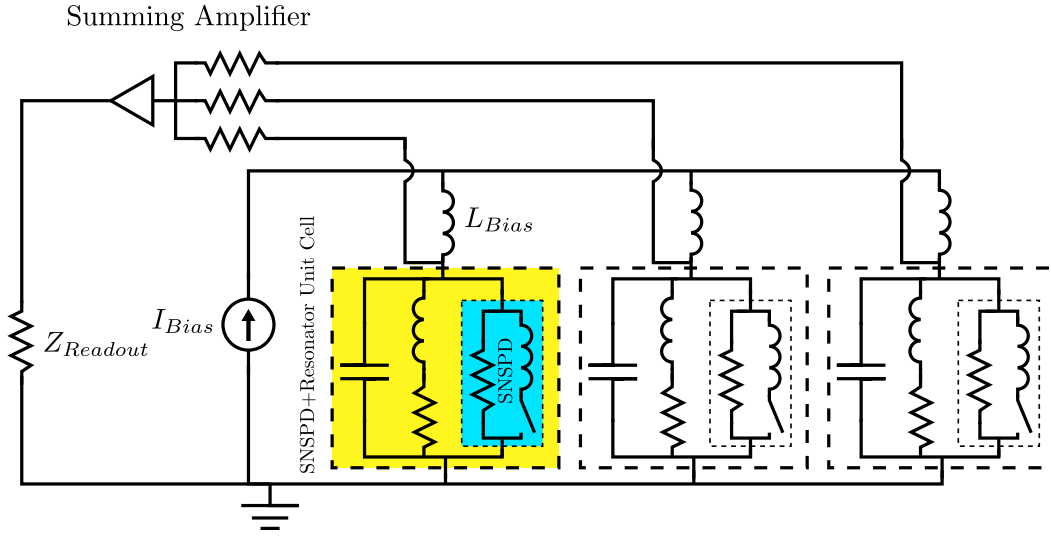


Figure 2.7: An example of a multiplexing scheme in which SNSPDs are connected in parallel with microwave resonators. When an SNSPD is tripped, current is shunted out through the resonator circuit which rings at its resonant frequency. The resistor in series with the inductor is used to control the Q of the resonator so that one can engineer a trade-off between how long it takes a detector to return to equilibrium and how tightly spaced each detector/resonator combo can be in frequency space. Each unit cell needs a relatively large bias inductor ($L_{Bias} \approx 10$ nH) to mitigate the impact of detection events in one cell on the bias current in other cells. The output of this scheme looks similar to the output from a single normal SNSPD, except the fall time is longer and oscillates at the center frequency of the resonator. Unlike more traditional frequency multiplexing schemes, one does not apply microwaves and measure the scattering parameters of the array; instead, the array generates its own microwaves which are readout through a summing amplifier. These can then be measured in the time-domain (possibly after mixing down to near base band depending on the bandwidth of the digitizer used) and then deconvolved in software.

2.3 Superconducting Delay Line Detector Concept

In order to solve this conundrum of creating a large, high resolution array of nanowire detectors, we drew inspiration from microchannel plate readout techniques. That is, instead of treating the nanowire detector as a lumped inductor which essentially transforms into a lumped resistor, it is modeled as an extended object. Using this model, the location where the normal zone forms affects the timing properties of the output pulse. Furthermore, because the detector isn't constrained to having an extremely small cross section (because the particles of interest here are much higher energy than in standard SNSPD experiments) and because the detector does not have to be embedded in a carefully tuned optical stack[31, 52], it can now be implemented as a microstrip transmission line⁶. In this implementation, voltage edges created by detection events travel with a well-defined propagation velocity— that is, the detector also acts as a delay line, much like the anode used with MCPs. This design, along with the readout and bias circuitry is shown in Figure 2.9. We call this new SNSPD variant the superconducting delay line detector (SCDLD).

In a superconducting microstrip with a width w , distance about the ground plane h , thickness t , relative dielectric constant ϵ_r , and London penetration depth $\lambda_L(T)$ [17], the characteristic impedance (Z_0) and propagation velocity (v_{prop}) are given by equations 2.7 and 2.8, respectively.

$$Z_0 = 120\pi \frac{h}{w\sqrt{\epsilon_r}} \sqrt{1 + 2 \frac{\lambda_L(T)}{h} \coth \frac{t}{\lambda_L(T)}} \quad (2.7)$$

⁶This strategy has been undertaken before[9], however, the method of reading out the detector in that prior work failed to take full advantage of this design.

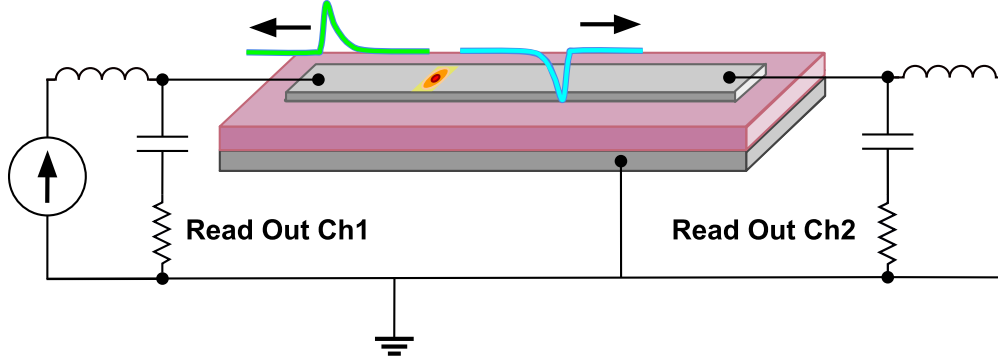


Figure 2.8: Illustration of a superconducting delay line detector. The yellow/orange band shows where a detection event occurred while the green and blue show hypothetical pulses created by the formation of the normal zone propagating to the upstream readout circuit (Ch1) and downstream readout (Ch2).

$$v_{prop} = \frac{C_{vacuum}}{\sqrt{\epsilon_r(1 + \frac{\lambda_L(T)}{h} + \frac{\lambda_L(T)}{h} \coth \frac{t}{\lambda_L(T)})}} \quad (2.8)$$

Where C_{vacuum} is the speed of light in a vacuum.

We also explicitly note that the London penetration depth is a function of temperature which can be calculated approximately as[44]:

$$\lambda_L(T) = \frac{\lambda_L(0)}{\sqrt{1 - (\frac{T}{T_{crit}})^4}} \quad (2.9)$$

There are two primary operating modes for the SCDLD. The first, the Avalanche Mode, is essentially the same as the detection method described for SNSPDs in Section 2.2. On the other hand, in the Linear Mode of operation the detector is biased with a lower constant current such that the ions of interest no longer have sufficient energy to completely destroy the superconducting state of the delay line. In this case, the output

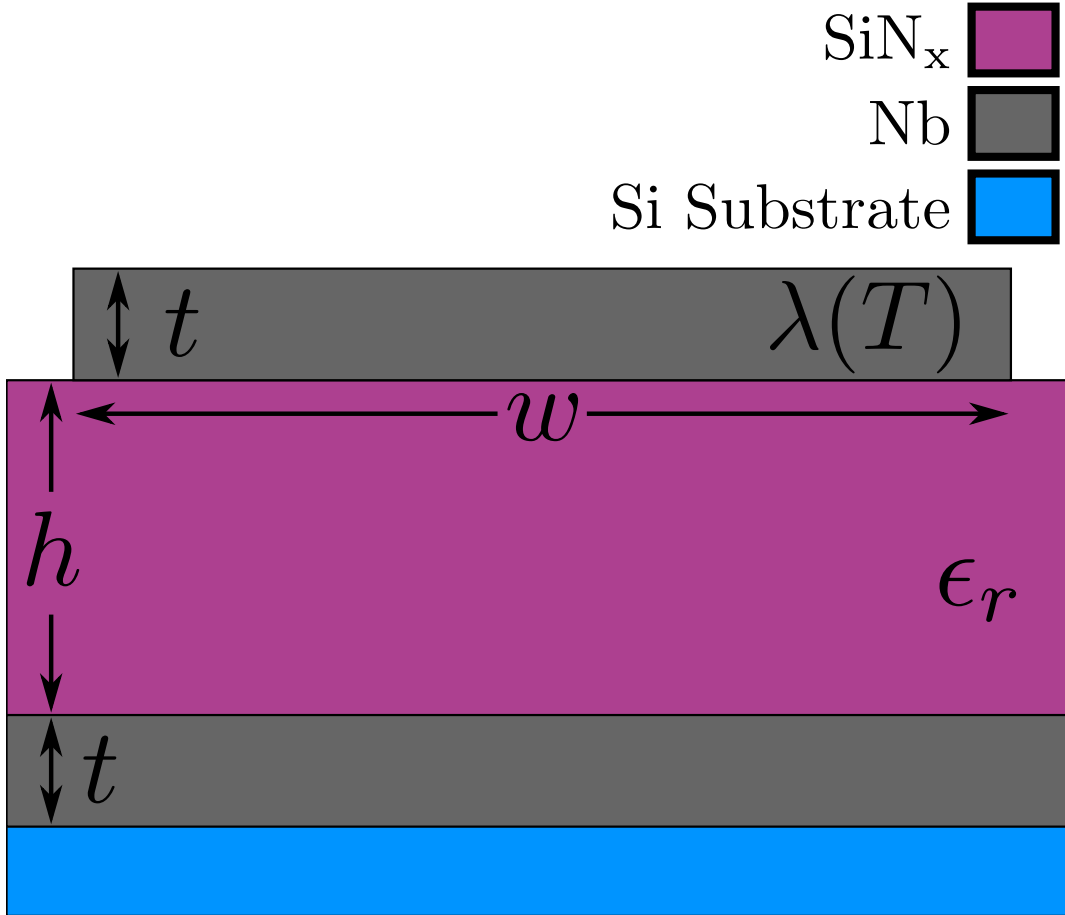


Figure 2.9: Superconducting Microstrip Geometry: In our devices, $w \gg h$ and the thickness of the ground plane is equal to the thickness of the signal trace (t).

of the detector is created by variations in the kinetic inductance of the detector and its amplitude will vary with input energy.

2.3.1 Avalanche Mode

In the avalanche mode, a bias current very near the detector's I_{crit} is sent through the detector. When energy from an ion is absorbed, superconductivity is entirely destroyed

and a normal zone is formed, generating a measurable voltage pulse following the same basic mechanisms shown in Figure 2.4. The output of the detector in this mode is expected to be the same regardless of the input energy (so long as this energy is sufficient to trigger the detector). An equivalent circuit is shown in Figure 2.10.

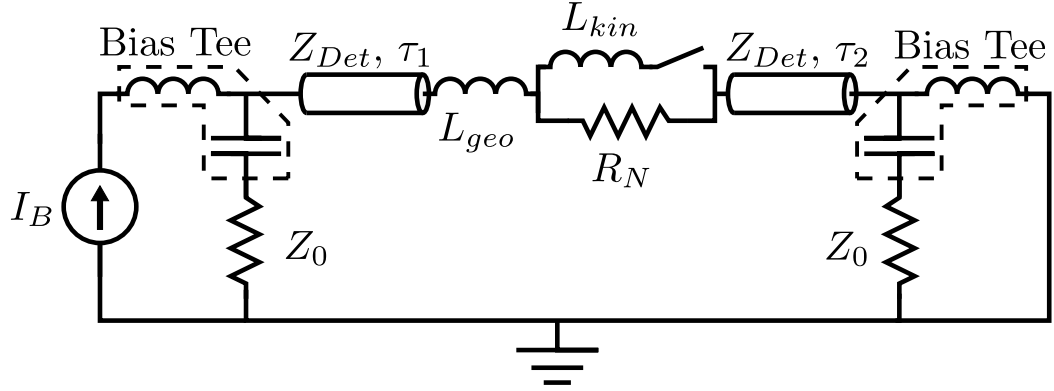


Figure 2.10: Equivalent circuit for an SCDLD operated in the avalanche mode. The central segment should look familiar; it is the same circuit we used to model an SNSPD previously. Now, though, large parts of the detector are unaffected by the detection event and maintain their properties as a transmission line. These segments are modeled by the two transmission line segments on either side of the normal zone. The resistors labeled as Z_0 are the read out channels which typically have an impedance of $50\ \Omega$.

In our conceptual illustration of the SCDLD in Figure 2.9 the detector response was depicted as a pair of pulses with a sharp rising edge, and relatively slow falling edge. However, this is based purely on previous measurements of SNSPD devices which generate such characteristic time traces. In reality, our detector "sees" a significantly different type of load than a traditional SNSPD, so it isn't immediately clear what will set the time scales previously defined in equations 2.1 and 2.2. In the equivalent

circuit, we labeled components with R_N, L_{kin} , and L_{geo} , but the actual value of these components will depend on the length of the hot spot at any one particular time. It is therefore necessary to model the electrothermal feedback which drives the growth and shrinkage of the hotspot in an SCDLD. This is done in the following section.

2.3.1.1 Hot Spot Dynamics Finite Difference Time Domain (FDTD) Simulation

The problem of electrothermal feedback in normal domains inside of superconductors is a subtle problem which involves a large number of material parameters and non-linear functions. This problem has previously been treated in a fairly general way[13][1], yielding some useful results we will employ in our analysis. The first of these is the definition of the Stekley parameter:

$$\alpha = \frac{\rho_r j_{crit}^2 d}{h(T_{crit} - T_0)} \quad (2.10)$$

In the expression above, ρ_r is the residual resistivity of the superconductor, j_{crit} is the critical current density of the superconductor, d is the ratio of the wire's volume to its surface area, and h is an effective thermal conductivity between the wire and the substrate the wire is grown on. This Stekley parameter characterizes the power dissipated in a normal zone versus the power conducted away into the substrate when the normal zone is just above the superconductor's T_{crit} and j_{crit} . If α is greater than 1, then stable normal zones can be formed above some minimum current density, j_p . If it is less than one, then the current that would be necessary to sustain a normal zone through joule heating is greater than j_{crit} of the device itself. Most of the material parameters in α are relatively easy to measure except for h . However, if $\alpha \gg 1$, its value can be

estimated as [13]:

$$\alpha \approx 2\left(\frac{j_{crit}}{j_p}\right)^2 \quad (2.11)$$

where j_p is the minimum current density at which a normal zone can be sustained with joule heating. This parameter can now be estimated by taking an IV curve measurement of an unshunted detector; j_{crit} is a material parameter we know ahead of time while j_p is approximately given by the retrapping current of the detector. The fact that there is a retrapping current which is less than I_{crit} already tells us that stable normal zones can be developed, and thus α must be greater than 1. Now, the velocity for the propagation of the domain boundary between the normal and superconducting regions can be calculated:

$$v_{ns}(I_{det}) = v_0 \frac{\alpha\left(\frac{I_{det}}{I_{crit}}\right)^2 - 2\frac{T-T_0}{T_{crit}-T_0}}{\sqrt{\alpha\left(\frac{I_{det}}{I_{crit}}\right)^2 - \frac{T-T_0}{T_{crit}-T_0}}\frac{T-T_0}{T_{crit}-T_0}} \quad (2.12)$$

However, the temperature in the superconductor should be a smooth function, so near the normal zone boundary, $T \approx T_{crit}$, thus simplifying equation 2.12 to

$$v_{ns}(I_{det}) = v_0 \frac{\alpha\left(\frac{I_{det}}{I_{crit}}\right)^2 - 2}{\sqrt{\alpha\left(\frac{I_{det}}{I_{crit}}\right)^2 - 1}} \quad (2.13)$$

as shown in [21]. Our analysis continues along the line of thinking laid out in [21] by identifying

$$R_{hs} = 2\rho_r v_{ns}(I_{det}) \quad (2.14)$$

The static R_N from our previous equivalent circuit has been replaced with this R_{hs} which is the resistance of the hotspot which can vary in time. This expression enables

the calculation of the resistance of the hot spot/normal zone which is, approximately, only dependent on the current in the detector⁷. This means it is possible to write a SPICE-type simulation to simulate these dynamics without having to explicitly treat the thermal system. Now, using this equation for a current dependent resistance, a finite difference, time domain (FDTD) simulation of the circuit shown in Figure 2.11 is implemented.

For this simulation, the detector/readout/bias system is split into three parts. The circuit block on the left of Figure 2.11 is the upstream bias and readout circuits along with the upstream port of the detector. The model has been simplified somewhat—we make the somewhat crude assumption that we can neglect the reactive components in our bias tee and just treat the bias as an ideal current source. The circuit block on the left side of Figure 2.11 is the downstream bias and readout circuits and the downstream port of the detector. Similarly, the reactive components in the bias tee are neglected. This implicitly assumes that there is always a current equal to I_{bias} going out through the DC port of the bias tee. This is roughly equivalent to assuming that the DC port has an infinite inductance so that its current can never change. The center circuit in Figure 2.11 is the portion of the detector where the detection event occurred, and so it contains the normal zone which has a resistance $R_{hs}(t)$. The inductance and capacitance per unit length that is present in the region of the normal zone have also been neglected because their impact on the overall timescales are miniscule compared

⁷This neglects the possibility of nonlinearity in the thermal conductivity to the substrate and assumes that the approximation of $T \approx T_{crit}$ is valid

to the time delays imposed by the transmission lines⁸. The three circuits are connected through voltage and current sources which implement the time delays associated with the long segments of transmission line between the hot spot and the ends of the detector. This implementation of the transmission lines assumes that our microstrips are very nearly ideal (no loss, dispersion, et cetera).

It has been previously shown[21] that for an SNSPD to recover without latching the total inductance of the nanowire (in their case, this is dominated by L_{kin}) must be large enough to make the electrical recovery time (which was approximated in equation 2.2) longer than a thermal recovery time. We expect to see a similar, but modified anti-stability criterion. Now, the specific location of the hotspot's creation is allowed to move around within the detector so that the apparent "inductance" will vary with the detection event's position. Furthermore, the hotspot is coupled to a transmission line rather than a lumped inductor. However, the total inductance of the transmission line between the hotspot and the upstream readout can be calculated as $L_{det} = Z_{det}\tau_1$ while the inductance between the hotspot and the downstream readout is $L_{det} = Z_{det}\tau_2$, so at long times compared to τ_1 and τ_2 there should be a correspondence between the original lumped inductor model and the transmission line model. Unsurprisingly, the short time dynamics do not play out in exactly the same way in the transmission line case; but, essentially, the expectation is that as the detection event location moves farther from the edges of the detector, it should be more likely to recover. Additionally, as the

⁸Simulating the full dynamics, we typically find a maximum hotspot length of approximately 1000 squares; that is, about 1mm long. This corresponds to a total geometric inductance of 1nH which has a characteristic time scale of roughly 20 picoseconds compared to the tens of nanoseconds of delay imposed by the transmission lines.

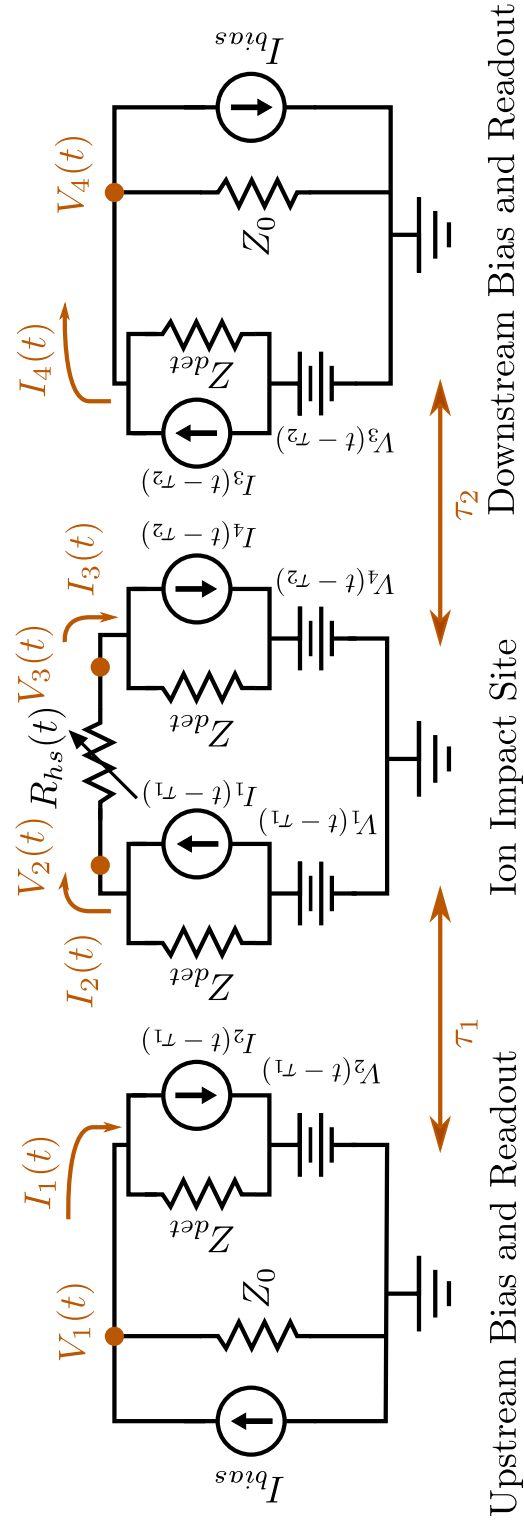


Figure 2.11: Simplified Equivalent Circuit Used in Hot Spot Dynamics Simulation

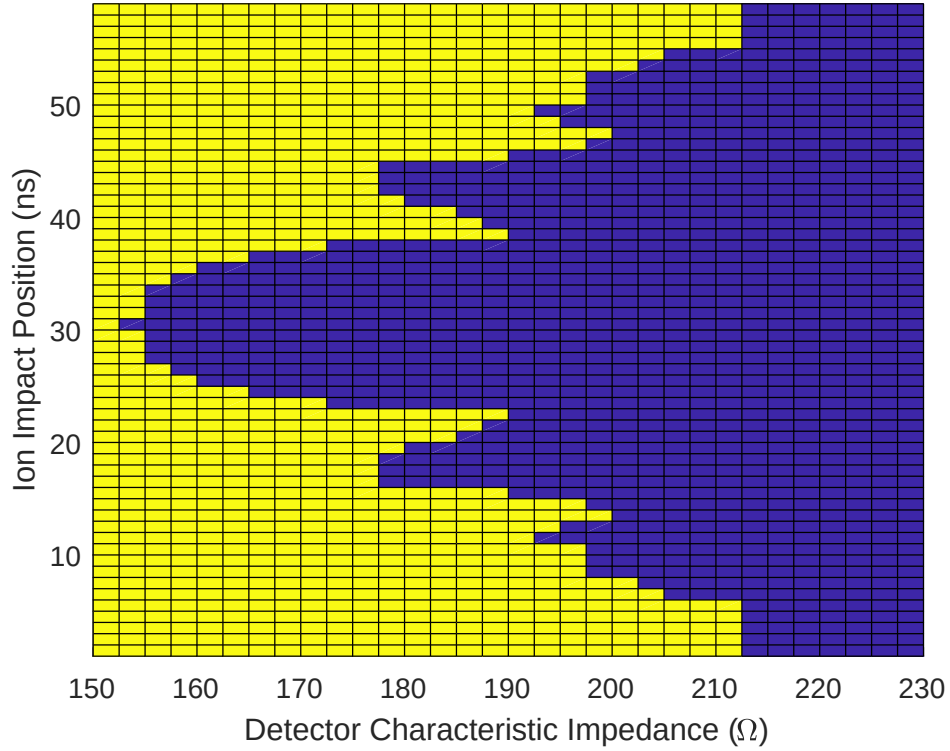


Figure 2.12: Results from our hot spot dynamics simulation. Cells which are colored blue indicate that the detector recovers superconductivity while yellow cells indicate that the detector latched into the normal state and would need to be reset by turning off the bias current in order to regain superconductivity.

characteristic impedance of the detector transmission line increase, it should become more likely to recover.

In order to assess the stability of hotspots in various detector designs, a fixed total electrical length⁹ is chosen, and the detection event position and the characteristic impedance of the detector are varied. The FDTD simulation is then run until either 1) the resistance of the hotspot stabilizes at some non-zero value or 2) the hotspot shrinks

⁹Which is approximately as long as the devices we make in the real world (60 ns)

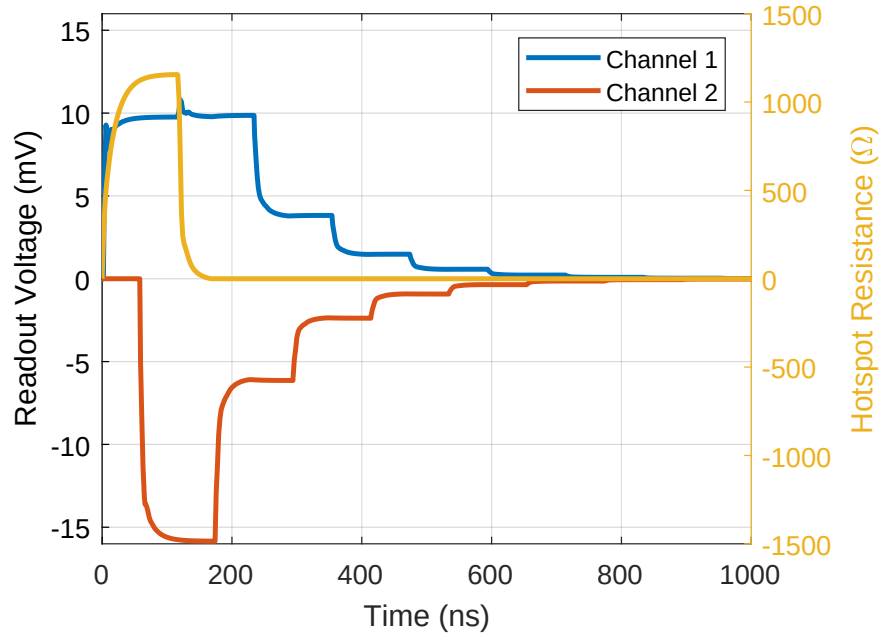


Figure 2.13: Time trace of the readout voltage from the upstream and downstream ports on the detector along with the time dependent hotspot resistance. In this simulation, the detector is designed to unconditionally recover; its characteristic impedance is $215\ \Omega$ and the detection event occurs 2 nsec from the upstream readout port.

down to a length of 0, at which point the normal zone is considered to be totally healed. These results are summarized in the plot shown in Figure 2.12.

As expected, increasing the detector’s characteristic impedance leads to a larger portion of the detector which will recover after an event. Also as expected, events near the middle are more likely to recover because of the larger “equivalent inductance” they see towards each of the read out lines. So, at first blush, it appears that we should just work with very high impedance detectors and call it a day. However, by examining the actual dynamics of the system, it becomes apparent that there is a significant trade off being made here. For example, in Figure 2.13, a detector with a characteristic

impedance of $215\,\Omega$ is simulated. Looking at the chart in 2.12, we see that this detector should unconditionally recover no matter where the ion impact occurs. Indeed, this is the case. Unfortunately, though, this comes at the cost of a lengthy recovery process, during which time the current flowing through the detector is suppressed from its initial value of I_{bias} . This is completely analogous to the situation in SNSPDs where there is a finite dead time after an event which scales with the length of the detector. We therefore do not gain any inherent benefit in τ_{fall} over a standard SNSPD. We do, however, find that τ_{rise} is independent on the length of the detector and the position where an event occurs. Now, the rise time is essentially completely dominated by the thermally driven growth of the hotspot (described by equation 2.13). This is because the hotspot essentially only sees the characteristic impedance of the detector as a resistive shunt to ground (at least at short times compared to τ_1 and τ_2 , which are the most critical for determining the rise time of the voltage edges) and this resistance has no dependence on the overall length of the detector nor the location of the detection event (except for extremely close to the edges where τ_1 or τ_2 are short compared to the rise time, which has been experimentally measured to be about 500 picoseconds). There is a tradeoff, though, in that lower detector characteristic impedances cause the current to redistribute away from the normal zone more quickly, so that τ_{rise} decreases with decreasing Z_{det} .

This analysis leaves presents the following options:

- Try to design a detector which will unconditionally recover, at the cost of having a long recovery time. Additionally, because the readout ports are coupled through the capacitive port of the bias tees, the waveform will actually consist of a series of reflected pulses after the initial event pulse, rather than the essentially flat voltage

edges seen in the FDTD simulation. This could confuse the analysis procedure somewhat if there is ambiguity between reflected pulses and actual event pulses.

- Try to design a detector which can detect ions below I_p so that it can operate at a low enough bias to where the detector will not latch regardless of its length or Z_{det} . This would be achieved by shrinking the cross section of the detector. This option is appealing, but we are limited by the lithography tools we have at our disposal, so it may not be possible to make detectors of an appropriate size. Additionally, as the width of our detector wire is reduced, the required length of detector we need to cover the same active area grows. Furthermore, the propagation velocity will also shrink, meaning that the total required electrical length grows faster than linearly with decreasing detector width. It is therefore advantageous to work at the largest width possible.
- Operate the detector in a mode where it is expected will latch. After an initial event, the detector will remain “dead” until the current flowing through the detector is reduced below I_p for some amount of time, allowing the hotspot to shrink and return to the superconducting state. This adds complexity to the system; however, due to the nature of time of flight spectrometry, it may not be a terrible idea. Because we need to measure the time of flight of the ions we are interested in, we will know when to expect them to be generated, and so before each repetition of the TOF experiment, the bias current can be lowered and then reset to prepare the detector for the next set of ions. We can minimize the time required to shut off the bias current and then restore it by matching the detector to the characteristic impedance of the readout/bias circuits (typically 50Ω). The reset time will then

be roughly equal to the time required for the hotspot to heal plus the total transit time of the detector ($\tau_1 + \tau_2$). Our FDTD simulation suggests that the time for the hotspot to heal is very short, so that the main limiting factor is just the transit time of the detector. This has the added benefit of suppressing the aforementioned pulse reflections, thus simplifying readout.

In the end, the decision is not completely in our hands; it turns out that for reasonable detector dimensions (set by our lithographic capabilities, deposition rates in our tools, and the energy of particles we are trying to detect), it is most convenient to work within the third strategy listed above, so we strive to match our detector to $50\ \Omega$.

2.3.2 Linear Mode

Another possible approach is to bias the detector with a current which is so low that a full normal zone never forms. This type of operation is called the “Linear Mode”. The idea is to use the energy of an incident particle to break up some portion of the cooper pairs in the detector, thus altering the kinetic inductance in a small area. Using Faraday’s law of induction,

$$\mathcal{E} = -\frac{d\Phi_B}{dt} \quad (2.15)$$

or, in terms of circuit parameters....

$$V = -\frac{d}{dt}(IL_{tot}) \quad (2.16)$$

This time varying inductance will generate a voltage which can then be detected. The equivalent circuit for the detector operated in this mode is shown in Figure 2.14.

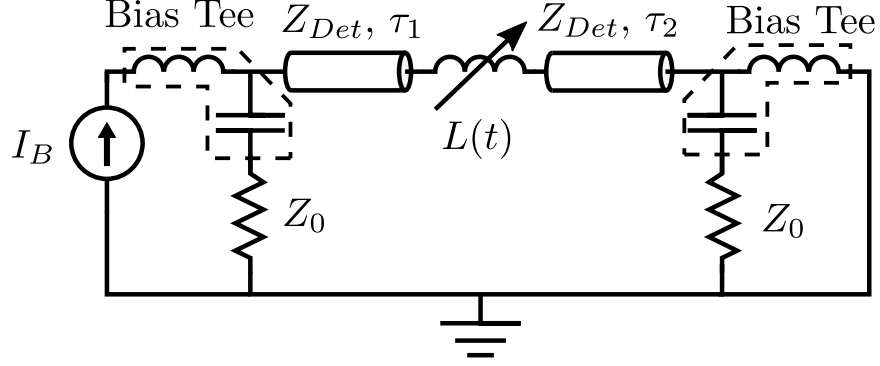


Figure 2.14: Linear Mode Equivalent Circuit

It has been suggested[2] that a time varying kinetic inductance should be treated differently from a time varying geometric inductance and that it may not actually enter the time derivative in Faraday's Law of Induction as shown below:

$$V = -\left(\frac{d}{dt}(IL_{geo}) + L_{kin}\frac{dI}{dt}\right) \quad (2.17)$$

In this case, changes in this parameter will not generate an electric field. This result is loosely derived by using the Drude model of conductivity. By taking a different approach, outlined in the Appendix A.1, it can be shown that the voltage generated by a time-varying cooper pair density is approximately given by

$$V = \frac{\dot{n}_s}{n_s}\left(A - \frac{\hbar}{q}\frac{\partial\Theta}{\partial x}\right)l \quad (2.18)$$

Which is equivalent to:

$$V = I\frac{\partial L_k}{\partial t} \quad (2.19)$$

This suggests that the kinetic inductance *does* in fact enter the time derivative of

Faraday's Law. Now, consider the equivalent circuit shown in Figure 2.15:

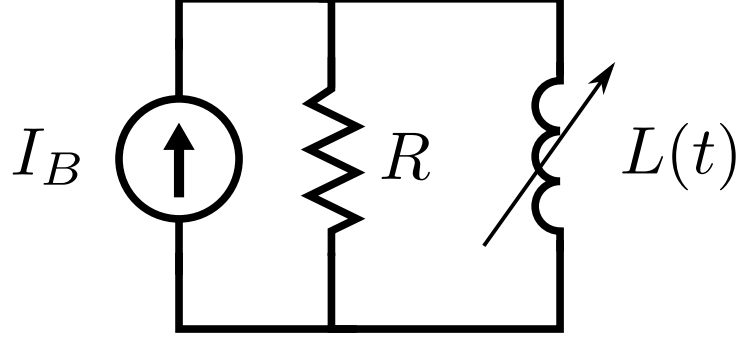


Figure 2.15: Simplified Linear Mode Simulation Circuit

It is shown in the Appendix A.2 that the voltage response generated by this circuit in the limit of a perturbation in L_k which is small compared to L_k and fast compared to L_k/R is approximately given by:

$$V_{out}(t_{final}) \approx I_B R \frac{\Delta L}{\Delta L + L_0} \quad (2.20)$$

We will run a simulation of this circuit in Section 2.3.2.1

2.3.2.1 Linear Mode FDTD Simulation

Equation A.3 suggests that operation in this mode should produce reasonably sized pulses. As a first estimate, consider the following parameters: $Z_0 = 50 \Omega$, $I_B = 1 \text{ mA}$, and $\Delta L/(L + \Delta L) \approx 1\%$. This should produce an output pulse with an amplitude of about $500 \mu\text{V}$. Simulating this circuit, again using FDTD techniques, produces a trace similar to Figure 2.16. The actual amplitude falls somewhat short of our estimation. This is due to the approximation of the perturbation being fasted compared to the timescale set by L/R in our model, which is not fully satisfied.

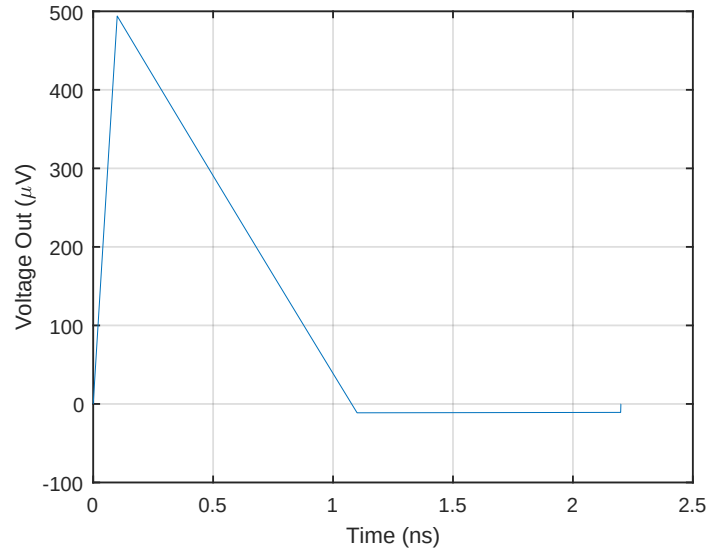


Figure 2.16: First Attempt at Linear Mode Simulation

This deficiency worsens when a more accurate estimation of the length scale over which the inductance varies in our detector is used in the simulation. As will be explained in Section 2.4.1, we plan on encapsulating the detector in a protective dielectric layer. This has the effect of spreading the energy of incoming ions over some finite range which of order the thickness of this encapsulation. This means that the inductance will be modulated over a stretch of transmission line which is of order 100s of nanometers long. Using the inductance of this length of transmission line yields a simulation result more like Figure 2.17. For the sake of clarity, plotted along with the raw output of the simulation is a trace which includes amplifier added white noise, a limited input bandwidth, and finite sampling rate which are all roughly matched to the electronics which are currently available to us.

Looking at these simulation results, it appears that reading out a detector in this mode would be extremely challenging due to the small signals generated, along with their

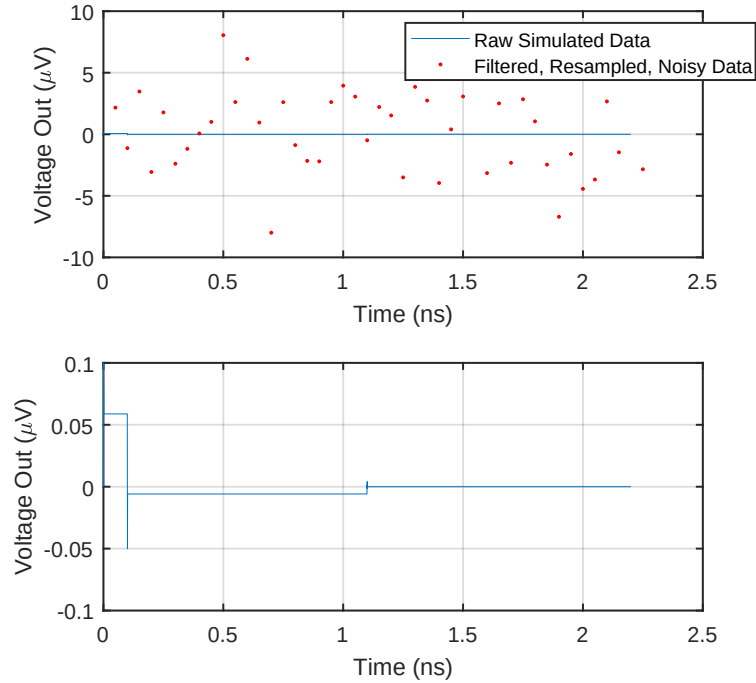


Figure 2.17: Linear Mode Simulation with more accurate estimation of L . The top panel includes both the raw simulated waveform along with data points which include the effects of finite oscilloscope bandwidth, finite oscilloscope sampling rate, and added noise from an amplifier. In this simulation, we take the bandwidth of the scope to be approximately 4 GHz, the sampling rate to be 20 GSa/s, the amplifier noise temperature to be 4K (very optimistic), and the amplifier bandwidth to be matched to the oscilloscope bandwidth. The bottom trace is just the raw simulated waveform. It is apparent from these plots that the actual signal we are looking for is totally swamped with noise.

wide bandwidth. It may be possible to improve the signal to noise ratio by employing such tricks as turning the detector into a resonator and probing it with a microwave drive; however, this is just a re-imagining of the MKID. The challenge with such kinetic

inductance based measurements in this experiment is that the target operation temperature is quite high. This leads to a reasonably large quasiparticle population which means that the lifetime of non-equilibrium quasiparticles is very short[16] (this is what leads to the extremely wide bandwidth response). It has long been recognized in the MKID[26] community that having quasiparticle lifetimes which are short compared to their resonator Qs is detrimental.

2.4 Ion Interactions With Matter

In order to understand the operation of our delay line detector, it is necessary to have at least a rudimentary understanding of the interaction of high energy ions in matter. The dynamics of ion scattering in matter has been thoroughly studied over the past century, and has reached such a mature stage that there are freeware simulations which can simulate these interactions with a high degree of accuracy. For the work in this thesis, we primarily utilized a program titled “SRIM/TRIM”; that is, the Stopping (Range) of Ions in Matter.

SRIM[56] can simulate the implantation of monoatomic ions of a wide range of energies into targets made of pure elements or simple compounds. It is worth noting that SRIM assumes the material is amorphous—phenomenon such as channeling are not included in the simulation. The general approach SRIM takes is to simulate ion interactions with a target material using Monte Carlos methods using the Binary Collision Approximation (BCA). BCA assumes that collisions between the incident ion and the target material atoms are rare so that each collision can be simulated separately; that is,

in any particular collision, exactly one target atom and the incident ion are involved¹⁰. The amorphous assumption means that SRIM does not need to create any kind of list of locations or overall structure for the atoms in the target. Instead, between each collision, a mean free path is calculated for the ion (which depends on the ion species and current energy, and the target density and composition). This is then used to calculate a distribution of likely free path lengths which the ion could traverse without experiencing a “significant” collision which would deflect it off its original path. Then a random number is drawn from this distribution, and the ion is then propagated that distance, where another new collision is calculated with a random impact parameter which will change the ion’s energy and direction.

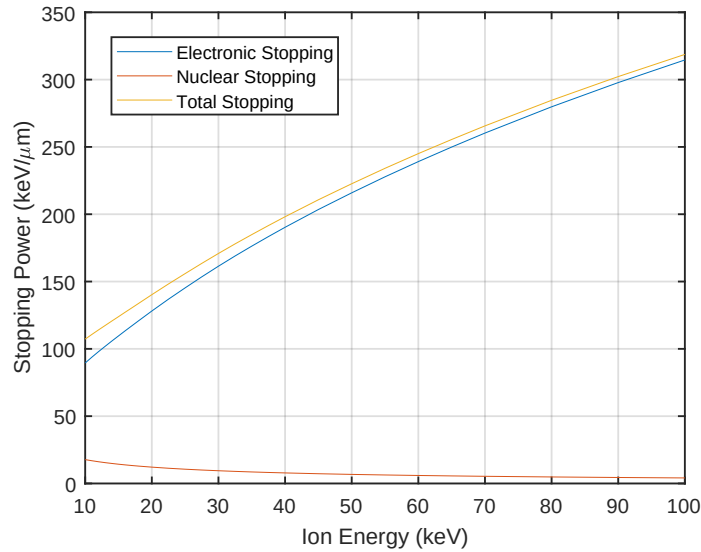


Figure 2.18: Helium in Silicon Nitride Stopping Power Example

There are two main types of energy loss for ions traveling through matter. The first is referred to as the nuclear stopping power. This is energy lost in collisions with nuclei

¹⁰If full cascades are calculated, it could be one target atom with a second target atom

in the target so that a portion of the ion's kinetic energy is transferred to the target atom. This energy can then either be dissipated as phonons in the target, or, if enough energy was transferred, the target atom can be displaced from its original location, creating defects in the target. The term nuclear stopping power specifically refers to the average energy lost by the ion to nuclear collisions (given a particular ion species, energy, and target material) per some unit length. The second type of energy loss is electronic stopping power. This refers to the energy lost by the ion to the electrons in the target material. Unlike nuclear stopping power which is treated in a kind of stochastic manner with well defined, rare collisions, electronic stopping power is more analogous to friction. Between collisions, as the ion loses some amount of energy per unit length (which is given by the electronic stopping power) which varies with the energy of the ion. Figure 2.18 shows an example of how these two stopping powers vary with energy for a given system (in this case, helium ions being implanted into silicon nitride).

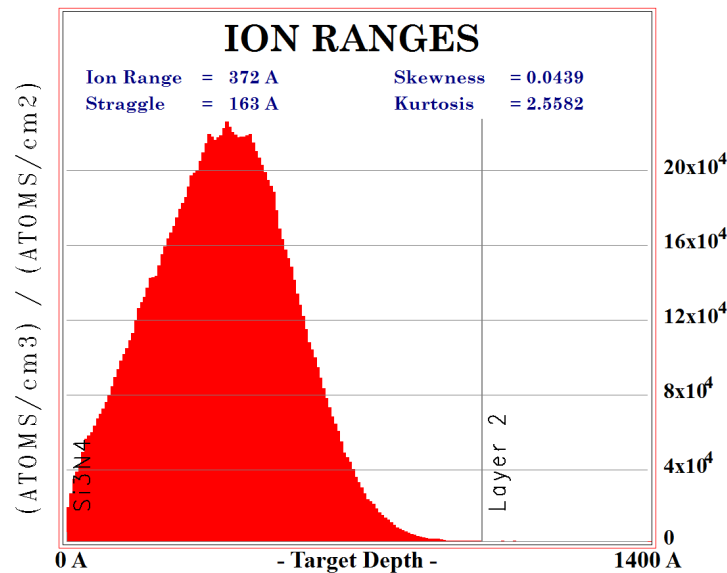


Figure 2.19: 5 keV Helium in Silicon Nitride Range

It is worth mentioning that ions which are stopped in the target material (in this case, our detector) remain there. Additionally, as mentioned above, if the incoming ion transfers enough energy to a nuclear collision, it can displace the target atom, creating defects in the target material. Although our niobium would best be described as polycrystalline, the disorder caused by long term exposure to ion damage could certainly suppress its critical temperature and critical current density over time. SRIM also includes the capability of making estimates of how much incoming ions will damage a target sample. In the work described in this thesis, we will mainly be testing our detectors with helium ions which are relatively light. However, in a real atom probe or time of flight spectrometer, the ions of interest could be of almost any species (and at higher energies than we can produce in our test setup), so it is worthwhile to look at the damage caused in these more pessimistic scenarios.

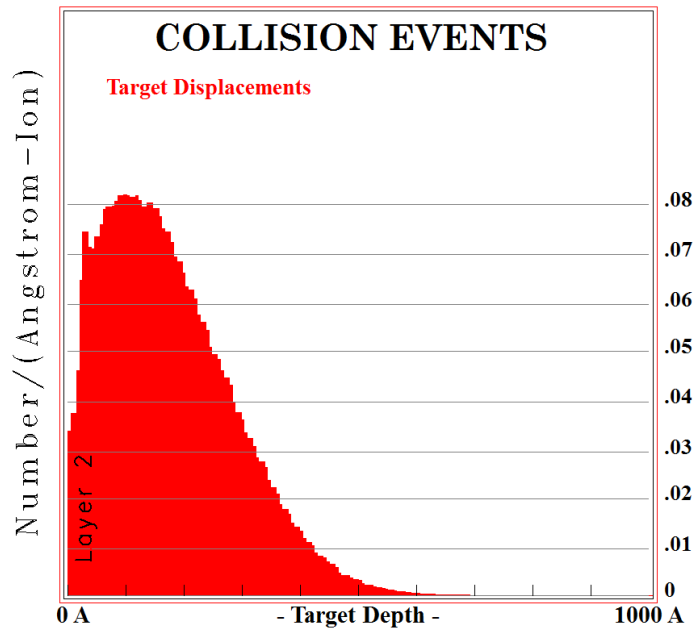


Figure 2.20: 5 keV Helium in Niobium Damage Events

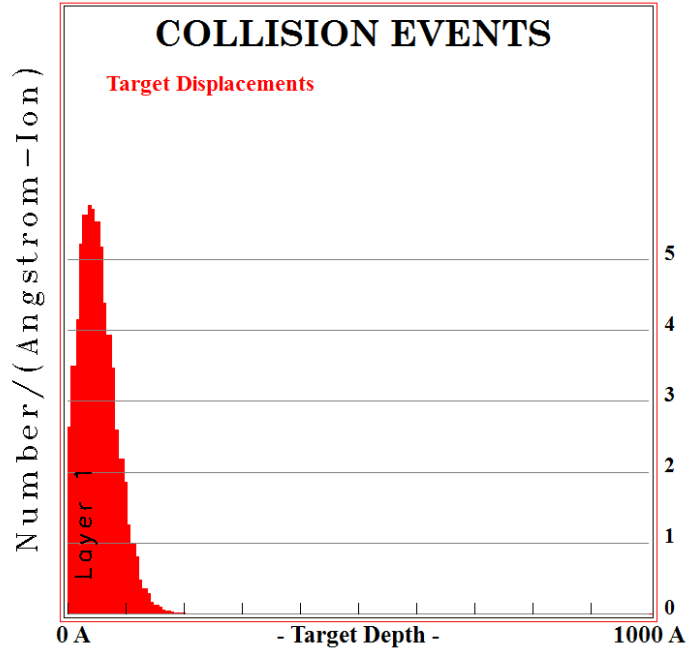


Figure 2.21: 20 keV Tungsten in Niobium Damage Events

As can be seen by comparing Figures 2.21 and 2.20, the damage per ion can change drastically with changing ion energy and mass. This begs the question of whether the detector can be designed in such a way as to protect it from damage from incoming ions.

2.4.1 Phonon Mediated Detection Scheme

In order to avoid accumulating ion induced damage in the detector itself, the detector is encapsulated with an amorphous layer of silicon nitride, deposited using plasma enhanced chemical vapor deposition. This also has the added benefit of protecting the detector from oxygen while stored in atmosphere, which can be detrimental to its performance[36].

With this encapsulation layer, energy from incident ions is no longer directly absorbed

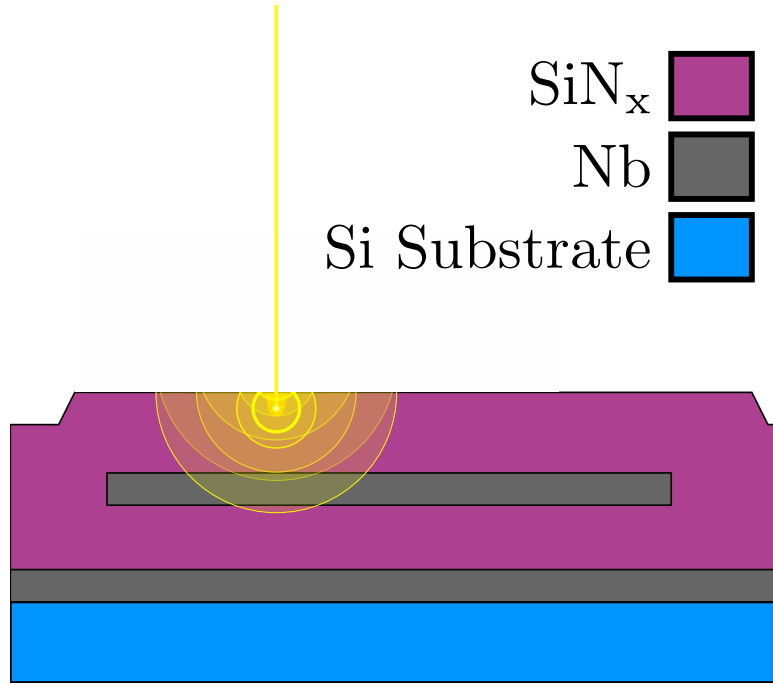


Figure 2.22: Ion Absorption and Phonon Emission in Encapsulation Layer

by the detector. Instead, the ion's energy is transferred into the electronic and phononic systems of the encapsulation layer, as described in Section 2.4. The energy absorbed as phonons (caused by nuclear scattering) is then transferred through the encapsulated layer and into the detector. Because there are essentially no free electrons in the silicon nitride, the thermal conductivity is dominated by phonon¹¹ transport[3]. Energy that is absorbed into the electronic system (via the electronic stopping power must be coupled out into the phononic system before it can be efficiently conducted to the detector.

¹¹Our silicon nitride is amorphous, and as such, does not truly have proper phonon modes due to its lack of crystal structure. However, measurements of the thermal conductivity of silicon nitride at cryogenic temperatures[57] indicate that pseudophononic modes exist which can have very long mean free paths, so we will continue to treat this transport as though it is phononic in nature.

2.5 Multi-Detector Experiments

After we fabricated and measured our first set of SCDLDs, it became clear there were some properties that we could not measure with these single detector experiments. For instance, in our test set up (which will be explained in Chapter 4), ions are generated with a non-uniform flux in space and time. In order to determine if the resulting “image” our detector generates is accurate, it would be useful to have a second witness detector very nearby. Additionally, we are unable to accurately calibrate the amount of ions being generated at our experiment at any particular time (much less know what kind of emission pattern they are being generated in). It is therefore difficult, if not impossible, to make a well calibrated estimate of the detection efficiency of the SCDLD. We therefore seek to implement a multi-detector experiment which will allow us to calculate this parameter. Finally, we are interested in potentially creating an array of these detectors in order to cover a very large active area and provide redundancy—creating small test patterns with more than one detector will allow us to study the interaction between multiple detectors operating on a single die in close proximity. In the following subsections, two experiments will be described which involve such dual detector devices.

2.5.1 Pitch Experiment Concept and Detection Efficiency Model

The first experiment is called the “Pitch Experiment”. The primary goal of this experiment is to obtain an estimate of the detection efficiency of the SCDLD as well as how this detection efficiency varies with ion impact position.

As described in Section 2.4.1, our detector is primarily detecting the phonons generated by the incident ion’s interaction with the silicon nitride encapsulation layer. At

cryogenic temperatures, phonons in even amorphous silicon nitride have very long mean free paths[57] compared to the lateral dimensions we are interested in (in a standard experiment, our detector width is 500–750 nm and the separation between detectors is 0.5 μm to 3 μm). We therefore usually work in the approximation of ballistic thermal transport— that is, we mainly concern ourselves with the solid angle subtended by the detector from the point of view of the ion’s impact point.

First, a simple model of detection where there is some threshold solid of angle of phonons the detector must absorb is considered. Above this value, the detection efficiency saturates to its maximum value. Below this value, the detection efficiency drops to essentially zero. The function used to model this is the sigmoid function:

$$P_{\text{detection}}(\Omega) = \frac{D_{\text{max}}}{1 + \exp(-\frac{\Omega - \Omega_{\text{thresh}}}{\sigma_{ph}})} \quad (2.21)$$

In Equation 2.21, D_{max} is the maximum detection efficiency, Ω is the solid angle, Ω_{thresh} is the threshold point at which the detection probability becomes appreciable, and σ_{ph} is a parameter that controls how quickly the detection probability changes in the vicinity of Ω_{thresh} (smaller values lead to a sharper turn on). An example of what this transfer function looks like is plotted in Figure 2.23.

The length of each detector is essentially infinite, but, from the normal core model (see Figure 2.4) it can be seen that collecting phonons along this long dimension is not so useful; the absorption of phonons in a narrow length near the ion impact site is more critical. It is therefore necessary to choose a somewhat arbitrary length scale to treat along this dimension. Given that the wavefunction of the superconductor can vary only over roughly the coherence length, this is chosen as the length scale to employ. Now, calculate the solid angle as a function of ion impact position[25](measured from the

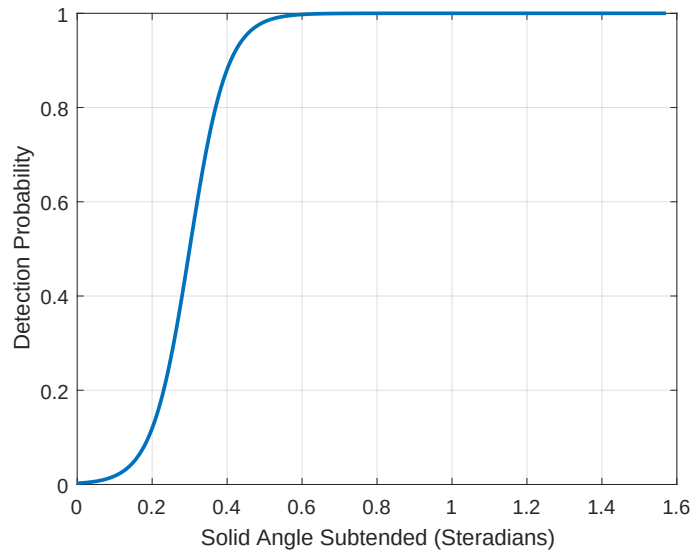


Figure 2.23: Example Transfer Function from Solid Angle to Detection Probability:

$$D_{max} = 1, \Omega_{thresh} = .3, \sigma_{ph} = .05$$

center of our detector), and then use Equation 2.21 to calculate the detection probability as a function of position.

An analytic form of the equation which connects ion position to detection probability can be written down in this model; however, it is quite complicated. Looking at Figure 2.24, the function looks very similar to the product of two sigmoid functions. In Figure 2.25 this similarity is illustrated by taking the previously calculated function for $P_{detection}(x_{ion})$ and fit it with the product of two sigmoid functions. This is a psuedo-1D model of detection probability.

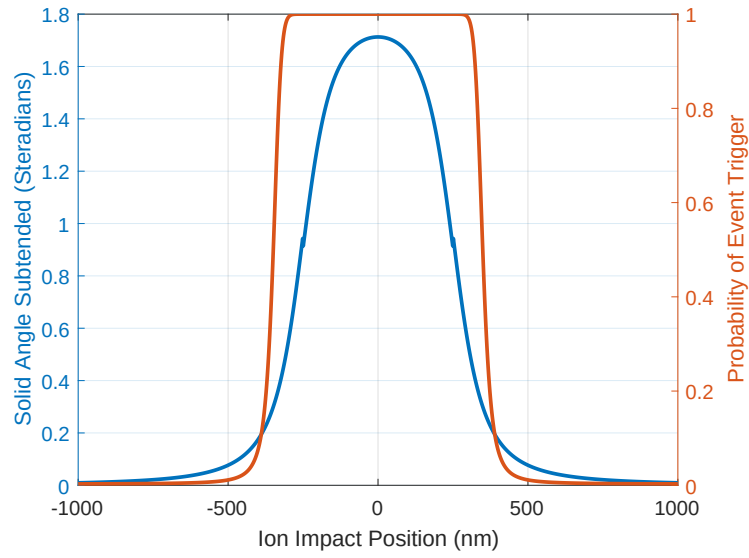


Figure 2.24: Detection Probability versus Ion Impact Position: $D_{max} = 1$, $\Omega_{thresh} = .3$, $\sigma_{ph} = .05$. The detector's geometry is as follows: 500 nm wide, 40 nm thick, encapsulated with 100 nm of silicon nitride. There is a small discontinuity in the calculated solid angle; this is due to the contribution from the sidewalls of the detector which only comes into play past a certain critical angle.

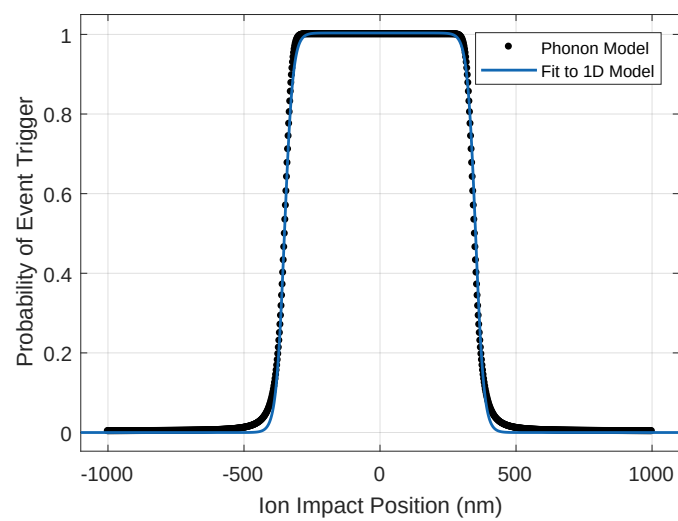


Figure 2.25: Psuedo-1D Model Fit to Soild Angle Detection Probability: The R-square value for this fit is .9993

We formally define our psuedo-1D model of the detector as follows:

$$P_{detection} = \frac{D_{max}}{(1 + e^{(x-x_c - \frac{w_{eff}}{2})/\tau})(1 + e^{(-x+x_c - \frac{w_{eff}}{2})/\tau})} \quad (2.22)$$

Where D_{max} is the peak detection efficiency, x_c is the center of the detector wire, x is the distance from the center of the wire to the ion impact location, w_{eff} is an effective width within which the wire is sensitive to ions, and τ is a parameter that controls how quickly the detection probability rolls off from D to zero.

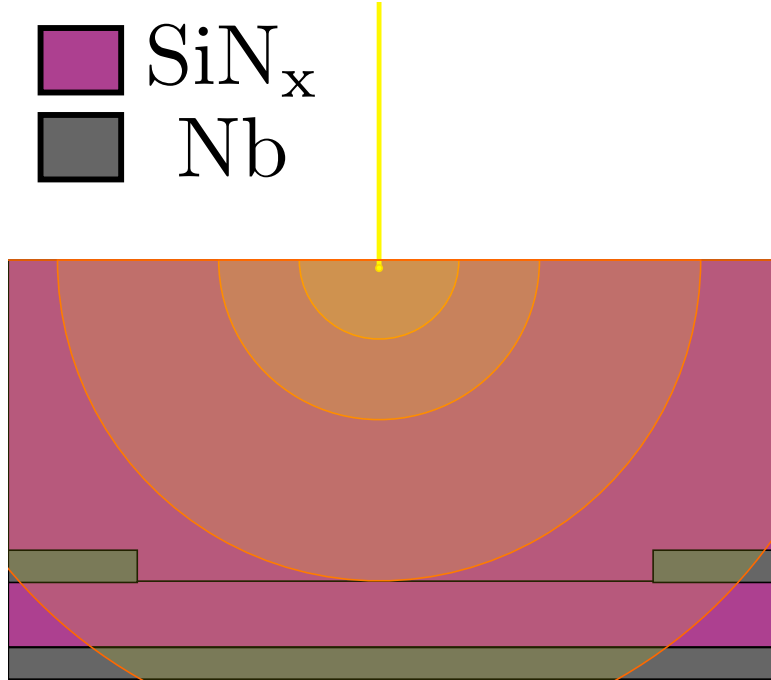


Figure 2.26: Ion Absorption Between Pitch Experiment Wires; In this illustration, the encapsulation is much thicker than the penetration depth of the ion. Phonons are released isotropically and travel ballistically for the majority of their journey.

Now, if two detectors are placed close to one another with a separation s , the joint probability that an ion will trigger both detectors can be calculated. Let P_1 be the

probability of wire 1 detecting an ion and P_2 be the probability that wire 2 detects an ion. Without loss of generality, we can center the two wire system around zero.

$$P_1(x) = \frac{D_{max}}{(1 + e^{(x - \frac{s+w_{eff}}{2})/\tau})(1 + e^{(-x - \frac{w_{eff}-s}{2})/\tau})}$$

$$P_2(x) = \frac{D_{max}}{(1 + e^{(x - \frac{w_{eff}-s}{2})/\tau})(1 + e^{(-x - \frac{w_{eff}+s}{2})/\tau})}$$

$$P_1(x)P_2(x) = \frac{D_{max}^2}{(1 + e^{(x - \frac{s+w_{eff}}{2})/\tau})(1 + e^{(-x - \frac{w_{eff}-s}{2})/\tau})(1 + e^{(x - \frac{w_{eff}-s}{2})/\tau})(1 + e^{(-x - \frac{w_{eff}+s}{2})/\tau})}$$

An example of these probability curves is illustrated in Figure 2.27 using reasonable parameters.

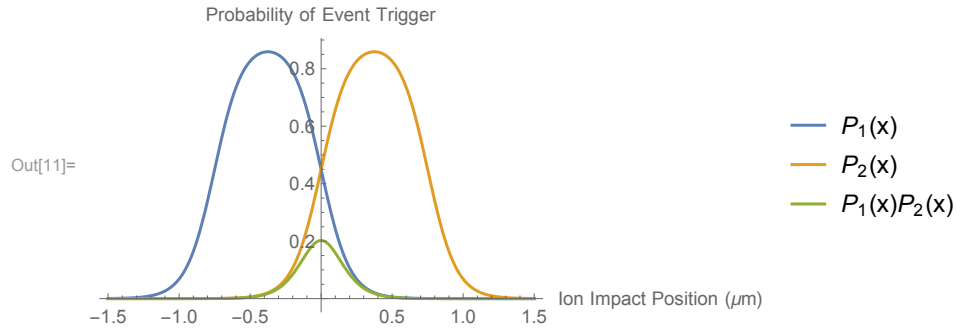


Figure 2.27: Probability of Detecting an Ion with: $D_{max} = .9$, $w_{eff} = .75\mu\text{m}$, $\tau = .1\mu\text{m}$, $s = .75\mu\text{m}$

These probabilities are not directly measurable in our experiment. However, it is possible to measure event rates (Γ) which are related to these probabilities as follows:

$$\begin{aligned}
\Gamma_1 &= \int_{-\infty}^{\infty} F_{ion}(x) P_1(x) dx \\
\Gamma_2 &= \int_{-\infty}^{\infty} F_{ion}(x) P_2(x) dx \\
\Gamma_{12} &= \int_{-\infty}^{\infty} F_{ion}(x) P_1(x) P_2(x) dx
\end{aligned}$$

In the above expressions, F_{ion} is the number of ions per unit time pass through the unit distance dx . In general, this ion flux will vary with x ; however, we have observed that in a typical experiment, the length scale over which F_{ion} varies is of order $100\mu\text{m}$. In comparison, a typical detector width would be 750nm and the spacing between detectors does not exceed $4\mu\text{m}$. It is thus reasonable to assume that F_{ion} does not vary considerably over the range where $P_1(x)$ and $P_2(x)$ are appreciable. We now calculate the "correlation fraction," Γ_{12}/Γ_1 .

Consider D , w_{eff} , and τ as fixed parameters and evaluate this expression as a function of the spacing, s . This yields:

$$\frac{\Gamma_{12}}{\Gamma_1}(s) = -\frac{D_{max}}{4w_{eff}} e^{\frac{w_{eff}}{2\tau}} \text{Csch}\left(\frac{s}{2\tau}\right) \left((w_{eff}-s) \text{Csch}\left(\frac{s-w_{eff}}{2\tau}\right) + (w_{eff}+s) \text{Csch}\left(\frac{s+w_{eff}}{2\tau}\right) \right) \quad (2.23)$$

An example plot of this function can be seen in Figure 2.28. Taking the limit of $\tau \rightarrow 0$ gives a sharp cutoff in $P_1(x)$ and yields a correlation fraction that, in the domain $0 < s < w_{eff}$ is linear in s .

$$\lim_{\tau \rightarrow 0} \frac{\Gamma_{12}}{\Gamma_1}(s) = D_{max} \left(1 - \frac{s}{w_{eff}}\right) \text{ in the range } 0 < s < w_{eff} \quad (2.24)$$

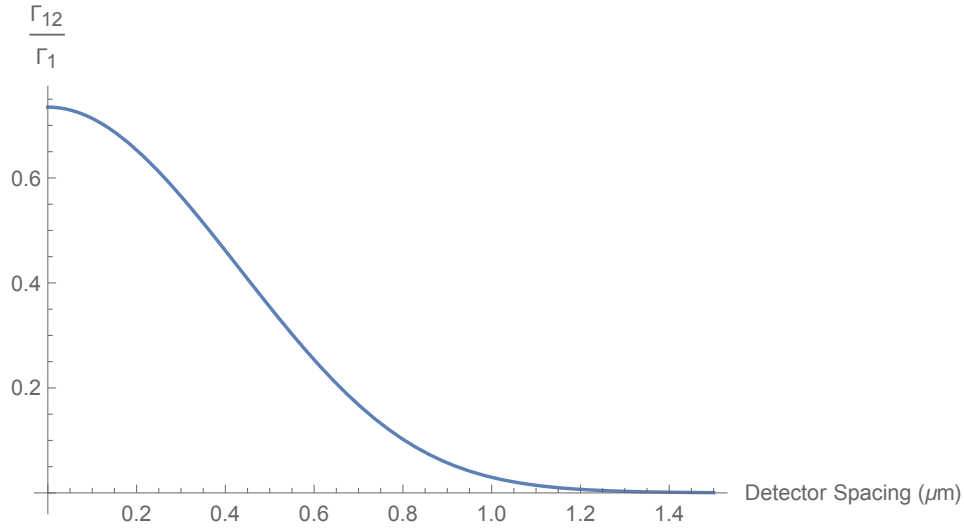


Figure 2.28: Correlation Fraction versus spacing with the following parameters: $D_{max} = .9$, $w_{eff} = .75\mu\text{m}$, $\tau = .1\mu\text{m}$

It is now possible to design an experiment that will measure the quantities D_{max} , w_{eff} , and τ . In Figure 2.29 an example of such an experiment is illustrated. In this design, two detectors are patterned on a single die and run parallel to one another. On one side of the die, they are very near each other, but with each meander they drift slowly apart. We chose a simple mapping such that the spacing varies linearly with distance along the detector.

The goal is to measure the correlation fraction as a function of s and fit it to Equation 2.23. In order to understand how well we can expect to fit this function, a Monte Carlo simulation of the experiment was written where ions are randomly generated across the detector face. The probability that an ion is detected in a particular wire is calculated in the phonon/solid angle model. This is used to generate a data set from which the correlation fraction can be calculated, and then fit to Equation 2.23. This fit is used to extract the parameters D_{max} , w_{eff} , and τ . Because we know what the original D ,

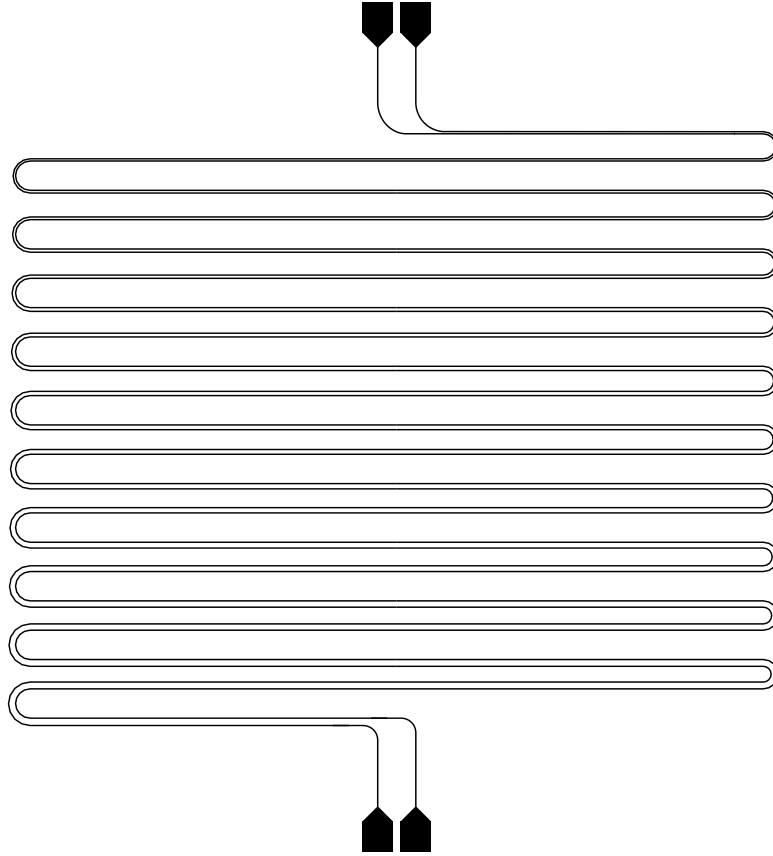


Figure 2.29: Pitch Experiment example layout. The length and number of meanders have been significantly reduced from their true values for ease of visualization.

w_{eff} , and τ used to generate the data was, the error in our parameter estimation can be calculated for each fit. This experiment is run many times over, and the root mean squared error of our parameter estimations is calculated. The actual detector parameters are also varied to see if the accuracy of our fits varies with the actual detector parameters.

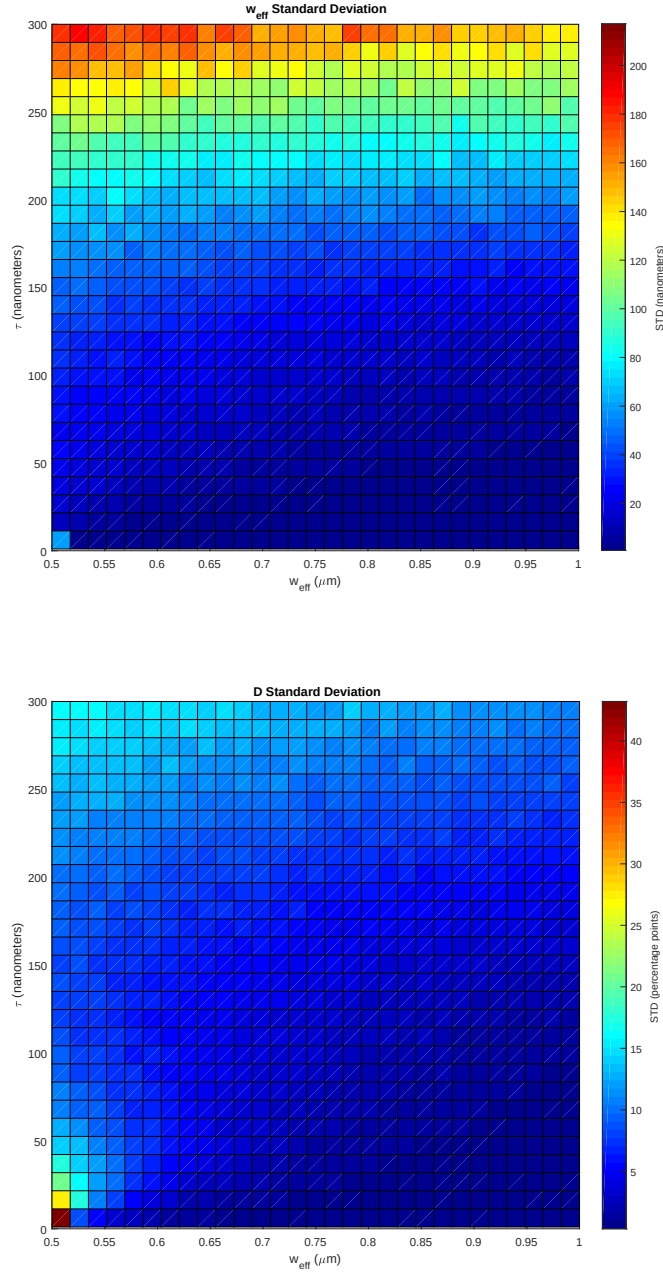


Figure 2.30: Pitch Experiment Curve Fitting Monte Carlos Results: In these simulations, we generate 1 million ions, roughly 12% of which are detected one or more wires. The pitch range we fit over is from $0.5\mu m$ to $1.25\mu m$ in 400 steps. We see that for most of the range of values simulated, we can place tight bounds on the parameters we are estimating. However, if w_{eff} is very small our data set has very few data points to fit to, and so the uncertainty becomes large.

2.5.2 Pitch Imager Concept

In Section 2.3.1.1 it was found that in most cases the detector will latch into the normal state after detecting an ion. In the devices fabricated for this thesis, this was certainly the case, and so it is worthwhile to think about what happens not only to the wire that latches, but other detectors which may be nearby. Recall that after the wire latches (but before the bias current is turned off), the hotspot grows to some equilibrium size where the electrical power dissipated in it is equal to the power it can dump into the surrounding substrate/encapsulation. This suggests that the latched wire will warm up nearby detectors, possibly causing them to trigger. This is illustrated in Figure 2.31.

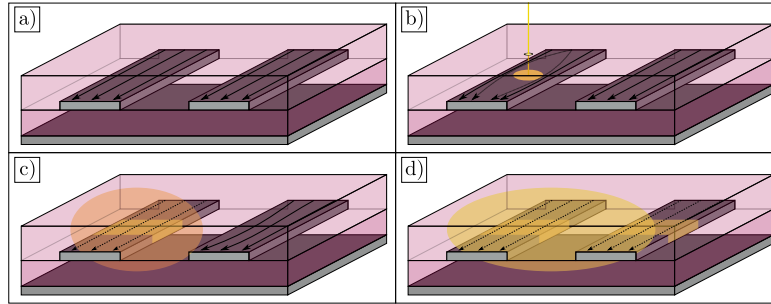


Figure 2.31: Illustration of Thermal Crosstalk from a Latched Detector

If this heating is sufficient to cause the second wire to trigger, it will be necessary to engineer around this in detector arrays with close spacing. This heating will happen on a relatively long time scale compared to phonon production by the original ion¹². We therefore opt to use the parabolic heat equation (or Fourier's Law of Heat Conduction)[12]:

¹²As an estimate for the phonon time, we consider the time it takes a phonon to traverse the gap from one detector to the other, which is approximately $1\mu\text{m}/10000\text{m/sec} \approx 100\text{ps}$.

$$\frac{\partial \phi}{\partial t} = D_\phi \nabla^2 \phi \quad (2.25)$$

Where Θ is heat and D_ϕ is the thermal conductivity of the material carrying that heat. The solution to this is derived in Appendix B as:

$$\phi(\vec{x}, t) = \frac{h_{joule}}{4D_\phi\pi s}(1 - \text{erf}(\frac{s}{\sqrt{4D_\phi t}})) + \frac{h_{joule}}{4D_\phi\pi\sqrt{s^2 + 4z_e^2}}(1 - \text{erf}(\frac{\sqrt{s^2 + 4z_e^2}}{\sqrt{4D_\phi t}})) \quad (2.26)$$

Where h_{joule} is the heat dissipated by the joule heating of the first wire, s is the center to center spacing between the two detectors, and z_e is the thickness of the silicon nitride encapsulation layer.

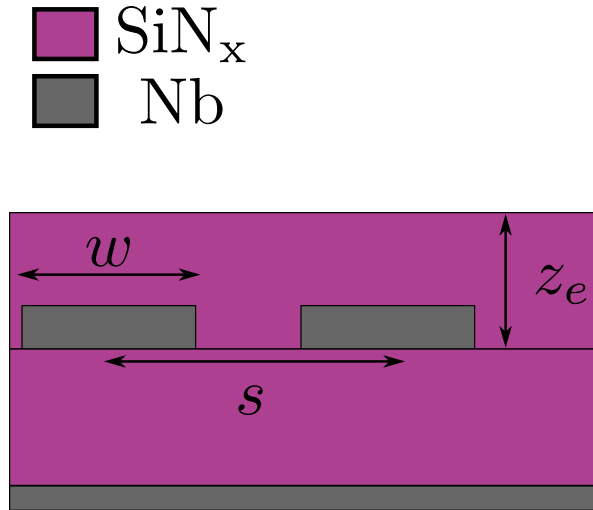


Figure 2.32: Illustration of Geometry Labels in Pitch Imager Derivation

It is enlightening to calculate quantities that we will be able to measure easily. In particular, making estimates of the thermal parameters for this system and then plotting the amount of time it takes the secondary wire to trigger after the first hotspot forms, yields a plot that looks like Figure 2.33

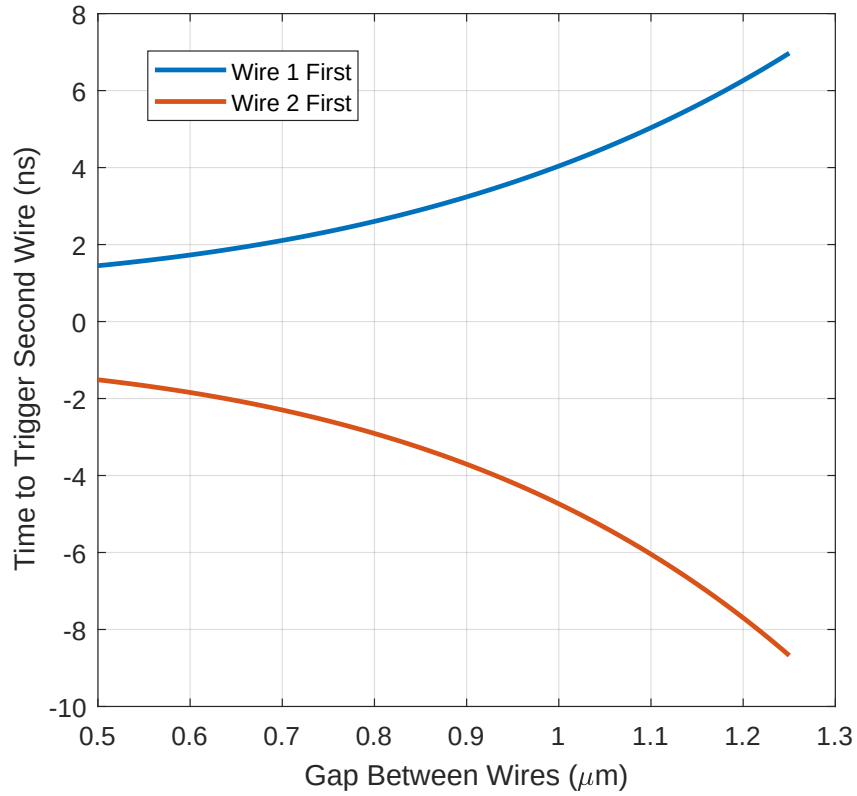


Figure 2.33: This is a plot of how long it takes from the first hotspot forming to reaching some critical heat flux which triggers the second wire. There are two branches which represent whether one wire triggers first or the second. We allow for the possibility that they may have different heat production rates and/or different critical heat levels—these factors cause them to have different curvature.

Curves like those shown in 2.33 can be directly measured in our Pitch Experiment samples to confirm this model’s validity. Additionally, these “Thermal Crosstalk” events can be utilized to improve our horizontal resolution (that is, along the length of each meander).

As we will discuss in Section 2.6, our timing measurement for both the average event

time and the relative delay time has some uncertainty. Regardless of how big or small that uncertainty is, our effective relative delay resolution in the vertical directional (perpendicular to each meander) will always be better than the horizontal timing resolution. The reason for this is because localizing an event along the vertical axis requires a measurement of the relative delay to within an uncertainty of about the time it takes to traverse an entire meander. The horizontal resolution, however, requires a measurement with an uncertainty which is fractions of this time¹³. This deficiency in the horizontal resolution is addressed by turning the Pitch Experiment on its side as shown in Figure 2.34.

Now, the pitch between the two detectors varies linearly as a function of horizontal position rather than as a function of vertical position. This allows us to measure the lag between thermal crosstalk events and then back out where along the horizontal dimension an event came from using the transfer function shown in Figure 2.33. A typical time scale for measured event lags is 5 ns to 100 ns for a pitch range of 1 μm to 4 μm .

¹³The obvious answer to this would be to make two meanders which are orthogonal to one another. Unfortunately, due to limitations related to our fabrication process, we have not been able to do this particular experiment.

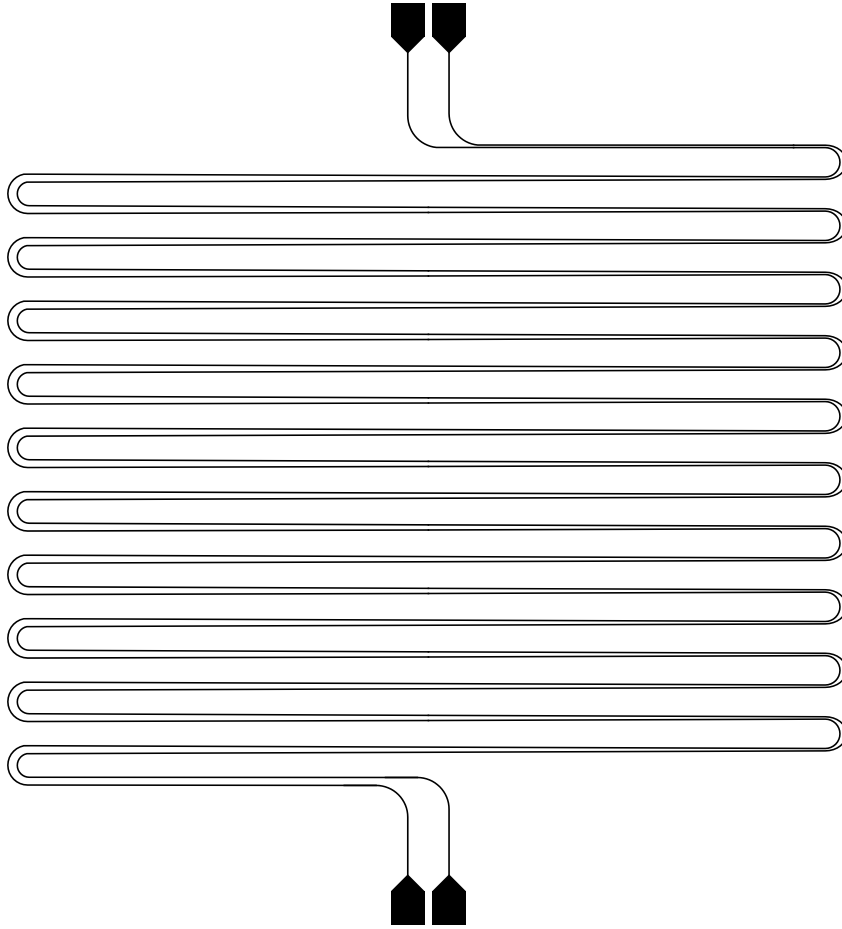


Figure 2.34: Pitch Imager Concept

2.6 Sources of Uncertainty in Timing Measurements

As we alluded to in Section 2.5.2, our measurements of pulse are not perfect. There are several possible sources of uncertainty in our timing measurements which will be discussed below.

2.6.1 Uncertainty Due to Electrical Noise in Readout

This algorithm is very simple; a threshold voltage, V_{thresh} , is set such that the time when the output from the detector, $V(t)$, first crosses this threshold is marked as t_{thresh} , and is considered the time when a pulse “happened”. In situations where there is a large scatter of pulse amplitudes (such as in the readout of MCPs) this technique fails and must be replaced with something more robust like a constant fraction discriminator. The pulse distribution coming out of our SCDLDs is quite narrow, though, so the basic threshold algorithm suffices. In order to understand the uncertainty in correctly identifying the time at which the threshold condition is satisfied, consider the noise present in our signal and how it affects our timing accuracy. If there is some noise voltage with an root mean square value of σ_V on our channel, then this will induce an uncertainty on our timing determination. To first order, we can say that $V(t) \approx V_0(t_{thresh}) + dV/dt|_{t_{thresh}}(t - t_{thresh})$, so that the uncertainty in our timing is then given simply by propagating the uncertainty in $V(t)$, yielding Equation 2.27.

$$\sigma_t = \sigma_V / \left(\frac{dV}{dt} \Big|_{t_{thresh}} \right) \quad (2.27)$$

The best resolution is achieved by setting the threshold time to the point where the derivative of $V(t)$ is at its largest. However, the measured pulses have an approximately exponential rising edge, so the time derivative monotonically decreases from the onset of the pulse. This means that the ideal threshold point would be some tiny voltage above 0; this is clearly not practical due to the presence of noise in the signal which would then cause false triggers at an unacceptably high rate. It is therefore necessary to balance the desire for improved timing resolution against rejection of noise when performing

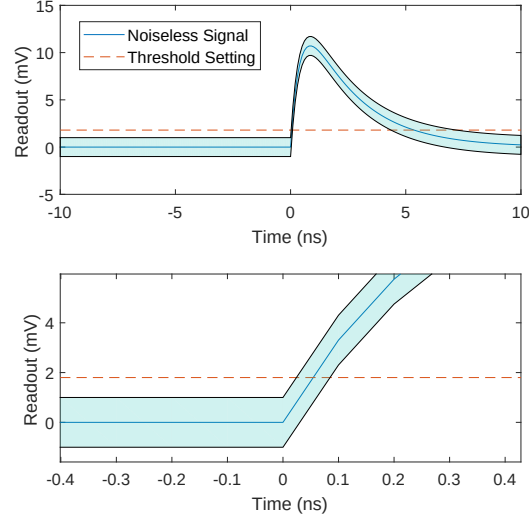


Figure 2.35: Timing Uncertainty due to Finite Signal to Noise Ratio in Threshold Detection

threshold detection. It is possible to improve the timing resolution separately from choosing the threshold voltage by filtering the signal to improve the signal to noise ratio. This process will be described in 4.5.2.

2.6.2 Uncertainty Due to Velocity Noise in Detector

Another possible source of timing uncertainty is if the timing of pulses coming out of the SCDLD has some intrinsic uncertainty. For instance, if the propagation velocity is varying inside the delay line. We consider a discretized version of our transmission line which has a randomly varying inductance component. We assume that the varying component of the inductance is relatively small (an assumption which is strongly supported by our experimental data). Let us divide the transmission line into N different components and label the inductance of each as L_n . We denote the mean of our inductances as L here,

and assume that the variations in L_n (denoted as δL_n) are normally distributed about zero with some variance we temporarily label as A .

$$L_n = L + \delta L_n$$

$$\delta L_n \sim \mathcal{N}(\mu = 0, \sigma = A)$$

We can now write down the time delay through each segment as a function of the transmission line parameters and then calculate the approximate variation of this delay for small perturbations in δL .

$$\tau_n = \frac{dx}{v_{prop_n}}$$

$$\tau_n = dx \sqrt{(L + \delta L_n)C}$$

$$\tau_0 \equiv \langle \tau_n \rangle = dx \sqrt{LC}$$

$$\tau_n \approx \tau_0 \left(1 + \frac{\delta L_n}{2L}\right) + \mathcal{O}(\delta L_n^2)$$

In the above, τ_0 is the propagation delay averaged over both temporal and spatial variations. We are interested in calculating the variance we expect to see in the relative delay for an event originating at any particular point. We label this point N_1 , and the relative delay is $RD(N_1)$.

$$RD(N_1) = \sum_{n=1}^{n=N_1} \tau_n - \sum_{n=N_1}^{n=N} \tau_n$$

$$\sigma_{RD}^2(N_1) = \langle RD(N_1)^2 \rangle - \langle RD(N_1) \rangle^2$$

The exact form of this final expression depends on how the inductance varies in space and time. We explore various possibilities in the following subsections.

2.6.3 Time Varying, Spatially Homogeneous Inductance

Consider the case where the inductance of the transmission line is varying in time, but this variation is happening everywhere in the detector equally. Note that τ_0 is the average propagation delay across variations in both space and time.

$$\begin{aligned}\delta L_n(t) &= \delta L(t) \\ \sigma_{RD}^2(N_1) &= \frac{\tau_0^2}{4L^2} \left(\sum_{n=1}^{n=N_1} \delta L_n - \sum_{n=N_1}^{n=N} \delta L_n \right)^2 \\ \sigma_{RD}^2(N_1) &= \frac{\tau_0^2 \text{var}^2(\delta L(t))}{4L^2} (N_1 - (N - N_1))^2 \\ \sigma_{RD}(N_1) &= \frac{\tau_0 A}{L} \left(N_1 - \frac{N}{2} \right)\end{aligned}$$

In this case, the uncertainty in relative delay depends on the position of the event being measured. The uncertainty reaches a minimum (which is 0) at the center of the detector. One physical mechanism which could create an uncertainty of this kind is variations in the temperature of the entire detector. Our detector is connected to the 3 K stage of a pulse tube cooler; typically, the temperature on this stage is about 3.2 Kelvin. However, this temperature oscillates at a rate of about 1 Hz, with a typical RMS amplitude of about 100 mK. Using Equation 2.8 we calculate an expected velocity variation of approximately .16%, which, in a 2 meter SCDLD translates to a worst case (ie, at either end of the delay line) uncertainty of 1.6 picoseconds.

2.6.4 Stationary, Spatially Varying Inductance

Consider the case of an spatially varying inductance which does not vary in time. Because the deviations in inductance are stationary, our 'average' relative delay term seems more complicated because we cannot just take a time average, thus eliminating the deviations; however, because we measure the same relative delay for any particular position everytime, the variance just goes to zero.

$$\begin{aligned}
 RD(N_1) &= \sum_{n=1}^{n=N_1} \tau_n - \sum_{n=N_1}^{n=N} \tau_n \\
 RD(N_1) &= \tau_0 \left(\sum_{n=1}^{n=N_1} \left(1 + \frac{\delta L_n}{2L}\right) - \sum_{n=N_1}^{n=N} \left(1 + \frac{\delta L_n}{2L}\right) \right) \\
 \sigma_{RD} &= 0
 \end{aligned}$$

The spatially varying propagation delays do not contribute to random noise in the relative delay measurement, but they do contribute to uncertainty in the transfer function connecting relative delay to position. Ideally, this function is perfectly linear; however, if the propagation delay is spatially varying, it will deviate slightly about this linear function. We can write an expression for the error (E_{map}) in our mapping between the two functions as follows.

$$\begin{aligned}
 E_{map}(N_1) &= \tau_0 \left(\sum_{n=1}^{n=N_1} \left(1 + \frac{\delta L_n}{2L}\right) - \sum_{n=N_1}^{n=N} \left(1 + \frac{\delta L_n}{2L}\right) - (2N_1 - N) \right) \\
 E_{map}(N_1) &= \tau_0 \left(\sum_{n=1}^{n=N_1} \left(\frac{\delta L_n}{2L}\right) - \sum_{n=N_1}^{n=N} \left(\frac{\delta L_n}{2L}\right) \right)
 \end{aligned}$$

One important thing to note is that no matter how large the deviations in δL_n are, the relative delay function must monotonically increase; having a negative propagation delay across a segment of transmission line is unphysical. Therefore, this type of inductance variation cannot "scramble" or blur information in the way we observe in our experimental measurements of relative delay timing resolution. This inductance variation instead warps and stretches the image we measure, but in a neat, one-to-one mapping.

2.6.5 Time varying, spatially varying inductance analysis

Finally, we now consider the case that the inductance of our transmission line varies in both space and time. We assume the fluctuations in inductance to be uncorrelated in different parts of the transmission line.

$$\begin{aligned}\sigma_{RD}^2 &= \left(\sum_{n=1}^{n=N_1} \tau_0 \left(1 + \frac{\delta L_n}{2L} \right) - \sum_{n=N_1}^{n=N} \tau_0 \left(1 + \frac{\delta L_n}{2L} \right) \right)^2 - \tau_0^2 (N_1 - (N - N_1))^2 \\ \sigma_{RD}^2 &= \frac{\tau_0^2}{4L^2} \left(\sum_{n=1}^{n=N_1} \delta L_n - \sum_{n=N_1}^{n=N} \delta L_n \right)^2\end{aligned}$$

Unsurprisingly, the mean value for the transit time to the upstream and downstream sides drop out. We are then left with the difference between two normal distributions, both of which have a mean of zero. Remembering that we are considering the case of L_n being independent from one another, and taking advantage of the fact that the distributions are symmetric about zero, we replace the subtraction with addition without loss of generality. This means that the uncertainty in the relative delay is equal to the uncertainty in total transit time, and is therefore totally independent of N_1 .

$$\begin{aligned}\sigma_{RD}^2(N_1) &= \frac{\tau_0^2}{4L^2} \left(\sum_{n=1}^{n=N_1} \delta L_n + \sum_{n=N_1}^{n=N} \delta L_n \right)^2 \\ \sigma_{RD}^2(N_1) &= \frac{\tau_0^2}{4L^2} \left(\sum_{n=1}^{n=N} \delta L_n \right)^2 \\ \sigma_{RD} &= \frac{\tau_0}{2L} A \sqrt{N}\end{aligned}$$

This kind of uncertainty will be difficult to distinguish from readout/electrical noise due to the fact that electrical noise will also impact the uncertainty of our measurement in a way that does not vary with position in the detector. However, if we compare the measured uncertainty to Equation 2.27, we should be able determine if our readout is the limiting factor or if it is velocity noise in the SCDLD.

Chapter 3

Device Fabrication

3.1 Overview

The detectors described in this thesis are very conceptually very simple; however, the process of actually realizing them is very challenging and time consuming. We designed and fabricated a number of different generations and variants, each with their own manufacturing complications. The overarching themes in this chapter will be an emphasis on strict procedures to ensure cleanliness and protecting the detectors from sources of oxidation. We will first outline the general procedure for creating individual layers of the detector. After this, we will describe the procedure for fabricating a device from the bottom up. The devices in this thesis were all fabricated by the author at the University of Wisconsin– Madison (UW). Most of this work was performed in the Wisconsin Center for Applied Microelectronics (WCAM), a shared cleanroom facility on the UW campus.

3.2 Superconductor Deposition

3.2.1 Kurt-Lesker Sputter System

Our lab contains two sputter deposition tools. The first of these, a Kurt-Lesker sputter deposition chamber, contains a Niobium and Aluminum target. Each of these is installed

on a magnetron sputter gun, and they share a 600 Watt DC power source which can be switched between them. This sputter tool is the primary workhorse of our group and is used to grow 'typical' films for our superconducting devices. This chamber is also outfitted with an ion gun which is used to remove native oxides from metals and to promote film adhesion between layers.

Samples are mounted on a circular, aluminum platen. This platen has two grooves milled into it where indium wire is inserted in order to improve thermal conduction from the wafer to the platen. One of these is meant for two inch wafers, while the other is for three inch wafers. The sample is then set on top of the indium wire and is held in place by an aluminum ring which overlaps with the edge of the wafer by approximately 2 millimeters. This ring is held to the platen with 6 screws, equally spaced along its circumference.

3.2.2 TLI Sputter System

The second sputter tool is used for more exotic film depositions. It contains the following targets: Niobium, Titanium, Tungsten Silicide, and Silicon Nitride. It is also plumbed with nitrogen gas to reactively sputter Niobium Nitride (NbN) and Titanium Nitride (TiN). The sample stage has an integrated heater which can reach temperatures of up to 900 degrees Celsius. There is also a Kaufman & Robinson Ion Source that was installed later on for removing oxidized metal on samples which have been exposed to atmosphere. This step ensures that newly sputtered metal layers will make good metallic contact with previously processed metal layers.

3.2.3 PURAIR Flow Hood

Originally, sample mounting was performed in open air in a completely uncontrolled, dirty environment. Most devices that are fabricated by our group have extremely small critical areas such that they are relatively insensitive to contamination with dust, hair, et cetera. These detectors, however, cover a large percentage of the wafer and any contamination at all will render them useless. It was therefore imperative that we provide a clean environment for this step. Ideally, the entire sputter tool and mounting area would be in a clean room; however, this simply wasn't practical and so we elected to invest in a 'flow hood' which creates a small, clean work space area on a bench. This is achieved by forcing HEPA¹ (or ULPA²) filtered air out through it's open front.



Figure 3.1: PURAIR Flow Hood in Lab

After installing this flow hood, we observed a drastic improvement in yield. When mounting in open air, our yield for detector widths $< 5\mu\text{m}$ was essentially zero. Mounting in the flow hood, this number increased to about 50% with no other changes. After

¹High Efficiency Particulate Air filter– rated to remove 99.97% of particles larger than 300nm

²Ultra Low Particulate Air filter– rated to remove 99.999% of particles larger than 100nm

implementing stricter procedures for cleaning all mounting tools and critical parts of the sputter chamber, and covering the user's face/head while mounting, this number further increased to $> 90\%$.

3.3 Dielectric Deposition

In order to form the microstrip detector, it is necessary to grow an insulating layer of dielectric to separate the ground plane from the detector layer(s). Although both amorphous silicon oxide (SiO_x) and amorphous silicon nitride (SiN_x) were experimented with over the course of this work, we found that using SiN_x led to better detector yield. Thus, we used SiN_x almost exclusively in the last several years of developing the detector fabrication process. There are several different methods and tools for growing dielectrics. The tools and methods we explored in this work will be summarized in the following subsections.

3.3.1 PlasmaTherm 70

Dielectrics grown in the PlasmaTherm 70 (PT-70) are formed via plasma enhanced chemical vapor deposition (PECVD). In this process, precursor gases flow into the reaction chamber which is continuously pumped through a throttle valve. This throttle valve is adjusted to reach a steady state pressure which is programmed by the user- typically it is on order of 100 mTorr. RF power is then capacitively coupled into the gases, igniting a plasma. This plasma allows the precursor gases to become ionized at much lower temperatures than would normally be possible, thus catalyzing reactions in the

chamber; it is worth mentioning, however, that the temperature still needs to be somewhat elevated- typical deposition temperatures are in the range of $150 - 350^{\circ}\text{C}$. These reactions deposit precipitates on the wafer. The wafer is also heated in order to increase the mobility of precipitates on the wafer's surface, thus improving the uniformity and smoothness of the deposited film.

There are several figures of merit for dielectrics grown in this way. The refractive index, density, loss tangent, surface roughness, and pinhole densities are of concern in many cases. These characteristics are sensitive to chamber conditions and recipe parameters which influence the stoichiometry and uniformity of the final films. For instance, even brief lapses in nitrogen flow (on order of 10 seconds) in the PT-70 can lead to a proliferation of pinholes in SiN_x films. When fabricating detectors, the most crucial characteristic of the dielectric is a lack of pinholes. The active area covered by the detector is huge, so even rare pinholes can present significant issues. If they are large, they allow the microstrip to short to ground. If they are small, they can act as small breaks in the line which can limit the critical current of the detector. Of secondary importance is the surface roughness. This number should be as low as possible- excessively rough surfaces can limit the critical current of devices. Typically, the surface roughness of the SiN_x as deposited has not been an issue in this work. Third (and significantly less important) is the loss tangent of the SiN_x . This parameter influences the spread of pulse amplitudes. In extreme cases, this could cause problems with detecting the rising edge of pulses- however, in practice, the loss tangent is never so bad as to cause a significant issue.

When developing the recipe for depositing dielectrics during PECVD, several factors

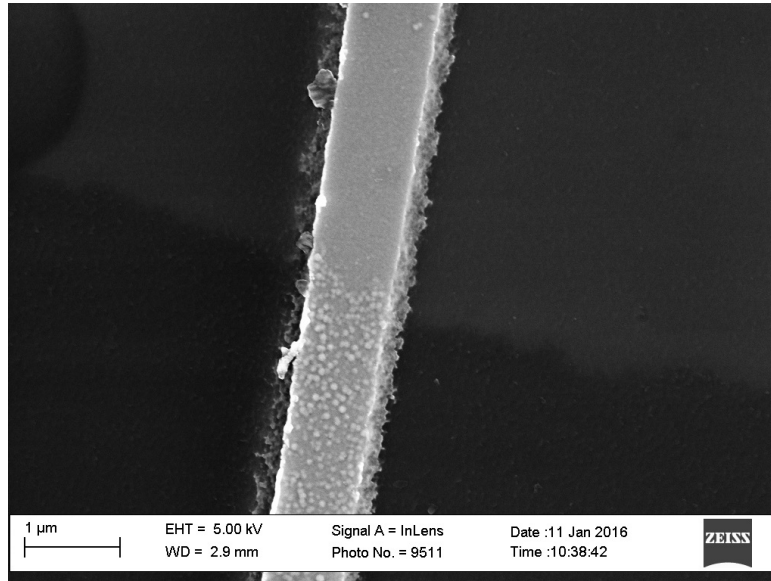


Figure 3.2: Aluminum Etch Stop Degradation: In the device pictured above, a niobium ground plane with an aluminum etch stop was sputtered, patterned, and etched. A portion of the aluminum etch stop was then wet etched off using TMAH (the top, lighter area of the image). Then, SiNx was deposited PECVD with the usual procedure. After this, a layer of niobium was deposited, patterned into a detector wire, and etched. It is difficult to see in the ground plane itself; however, where the detector wire crosses over parts of the ground plane which were capped with aluminum, the surface is clearly much rougher. It is unclear what the exact chemical mechanism is that causes the aluminum to develop this rough surface topology during PECVD. Measurements of niobium wires grown over these rough surfaces reveal that, perhaps unsurprisingly, the critical current density is severely suppressed by this surface topology.

needed to be considered. In addition to the previously mentioned metrics for the film itself, it is necessary to consider how the deposition process will affect the sample we are processing. For instance, we have found that the PECVD process can often lead

to degradation of aluminum films as shown in Figure 3.2. In this case, the solution we pursued was to eliminate the need for any aluminum films in our design rather than modifying the PECVD recipe. As another example, we found that oxygen has a tendency to suppress the critical current density of our films as well. It was therefore necessary to develop a SiNx recipe which contained no oxygen- therefore, we were constrained to only using SiN₄ and N₂ as precursor gases. This sensitivity to oxygen also meant that it was important to not allow unencapsulated detector wires to be heated above approximately 130°C while exposed to atmosphere. Because the PT-70 does not have a load lock, it was necessary to cool the sample platen to this lower temperature before loading a sample with exposed detector wires. Once the sample was pumped down to the base pressure of the tool, the platen temperature could be restored to the deposition temperature (250°C for our SiNx process).

3.3.2 Denton RF Sputter System

In addition to PECVD, we experimented with using RF sputtering to deposit dielectrics. When sputtering insulating materials, it is necessary to use an RF power source to avoid excessive charging of the dielectric sputter target. This also means it is necessary to use a matching network to make sure the RF power applied is actually being coupled into the plasma rather than being reflected back into the power supply.

We found that RF sputtering produced dielectric films which were high quality in some ways (high density, index of refraction much closer to the expected values for stoichiometric films) but failed with respect to our most important parameters. In particular, the pinhole density was higher than those grown by PECVD, thus limiting yield of layers

grown on top of the RF sputtered dielectric. From a practical standpoint, the procedure for RF sputtering proved to be far too slow compared to PECVD as well. The Denton sputter system has no loadlock and a very large deposition chamber, meaning that it takes well over 4 hours just to reach a modest vacuum of $\sim 10^{-6}$ Torr. Additionally, the procedure for sputtering dielectrics called for an extremely long ramp up and down procedure for the RF power to avoid thermally stressing the target. Due to the large pinhole density and the large time overhead for sputtering, we decided to focus our efforts on optimizing PECVD growth.

3.4 Lithography

Lithography is the process which allows us to define shapes/patterns which comprise devices.

As mentioned in Chapter 2, we are trying to make extremely large devices, and so ideally we will work with designs which can be patterned using optical lithography due to its extremely high throughput compared to electron beam lithography. In our fabrication facilities we have access to several different lithography tools, but they are all based on i-line light sources (that is, they use light with a wavelength, λ , of 365 nm). The minimum feature size (d_{min}) which can be made without using advanced deconvolutional/diffraction correction techniques is approximately given by:

$$d_{min} \propto \frac{\lambda}{NA} \quad (3.1)$$

where NA is the numerical aperture.

In practice, this minimum critical feature size depends on many factors and is often

specified for a particular lithography tool. This will limit the size of structures we can fabricate.

3.4.1 Karl Suss MA6 Contact Aligner

For coarse patterns which will only be shot on a small number of dies with no or only crude alignment, contact lithography is a quick, easy way to make patterns. In this method, a mask is held with a vacuum chuck above the surface of a wafer. The wafer is then coarsely aligned underneath the mask. Once it is roughly close to the location where the pattern is to be shot, the wafer stage lifts up until the wafer makes contact with the mask (hence the name). The stage then lowers slightly, allowing the user to make final, fine adjustments to the positioning using built in microscopes to check for alignment marks. This tool can be expected to achieve, at best, a resolution of $1\text{ }\mu\text{m}$. The overlay accuracy depends greatly on the skill of the user, but is essentially limited to $1\text{ }\mu\text{m}$ as well.

3.4.2 Nikon Body 8 i-Line Stepper

The Nikon Body 8 is a step and repeat aligner (“Stepper”) which uses i-line light and projection lithography to pattern multiple copies of the same pattern across a wafer. While designing a process, the user must decide what size of die will be used. This then constrains the size of devices that can be made and determines the number of dies which can fit on a wafer. In a Stepper, the lithographic masks only contain one copy of each pattern to be shot. For each lithographic step, the stepper then uses a series of blinds to isolate portions of the mask that the user wants to pattern with. It then shots that

pattern on each die on the wafer in turn, using stepper motors to position the wafer under the portion of the mask being used.

Unlike the MA6, alignment is performed automatically and in this tool the overlay accuracy can be as good as 100 nanometers. Another difference between this tool and the MA6 is the use of projection lithography. This means that there is no contact between the mask and the wafer which cuts down on contamination of the mask with both photoresist residues and particulate contamination which can cause future lithographic headaches. Finally, due to the design of the projection lithography system, there is a 5X size reduction from features written on the mask to the resulting features on the wafer.

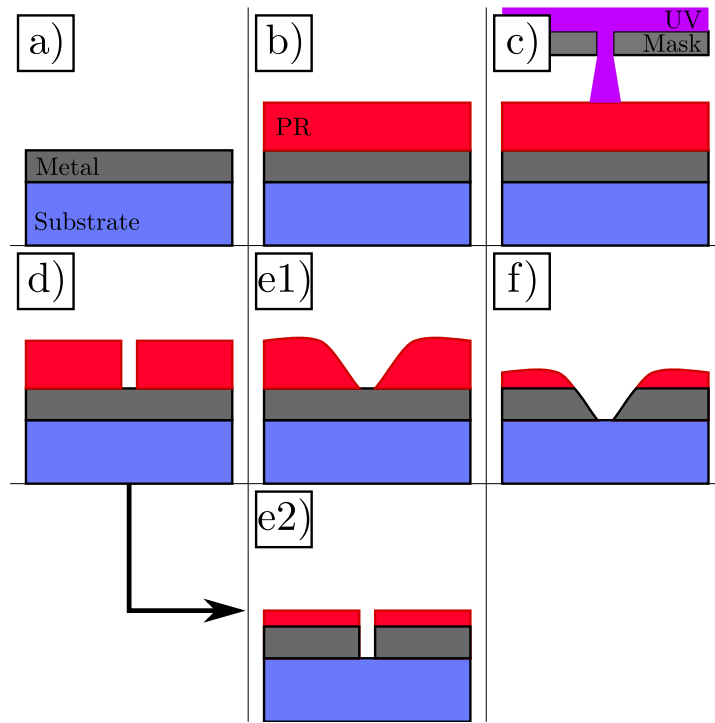


Figure 3.3: a) We start with a blanket film of material, which we want to remove some portion of (“Metal” in this case). b) We apply a coating of photoactive polymer called photoresist. This is typically then baked for a short amount of time. c) Ultraviolet light is used to expose the resist either where we want to remove it (for positive tone resist) or where we want it to remain (negative tone) d) The resist is then post baked for some amount of time and developed. In this case, we are illustrating positive tone resist, so the area which was exposed to UV is removed in the developing process e1) If we want to end up with sloped sidewalls in our etch pattern, we can do an additional bake step with a temperature near or slightly higher than the glass temperature of the photoresist. This causes it to liquify slightly and reflow, forming roughly 45° slopes. f) The underlying metal is then etched. If sloped sidewalls are desired, the etch recipe needs to be tuned to etch the photoresist at the same rate as the material underneath. e2) If straight walls are desired, the extra bake in e1 is skipped, and the material is etched with an anisotropic process, such as a reactive ion etch.

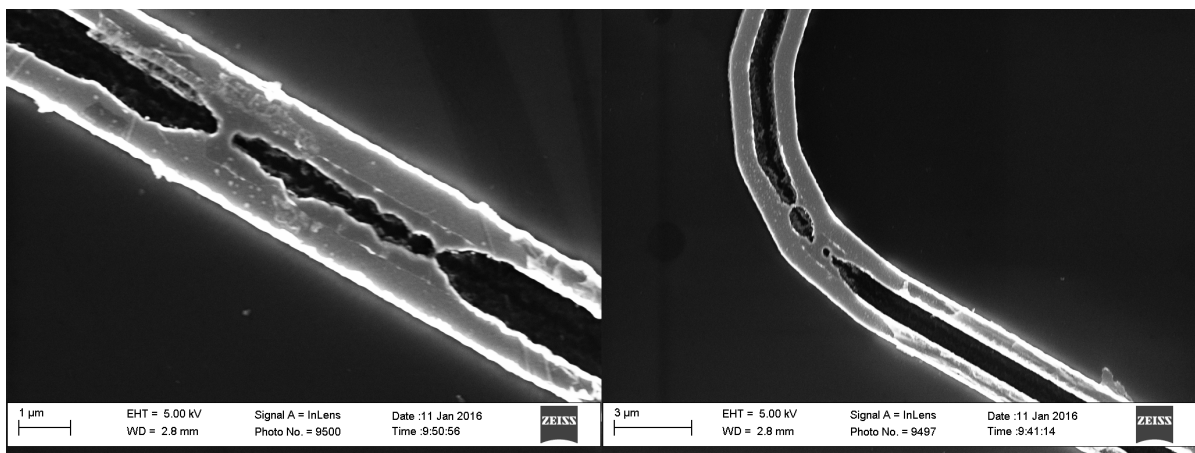


Figure 3.4: Lithography Defects Due to Suboptimal Developer Choice: This is an SEM image of two closely spaced (750 nanometer gap) detector wires which were patterned in SPR-955 and developed using MF-24A. In an SEM, it is clear that the photoresist does not properly clear between the two detectors- furthermore, this was the case even after varying both the exposure dose and focus offset during the lithography.

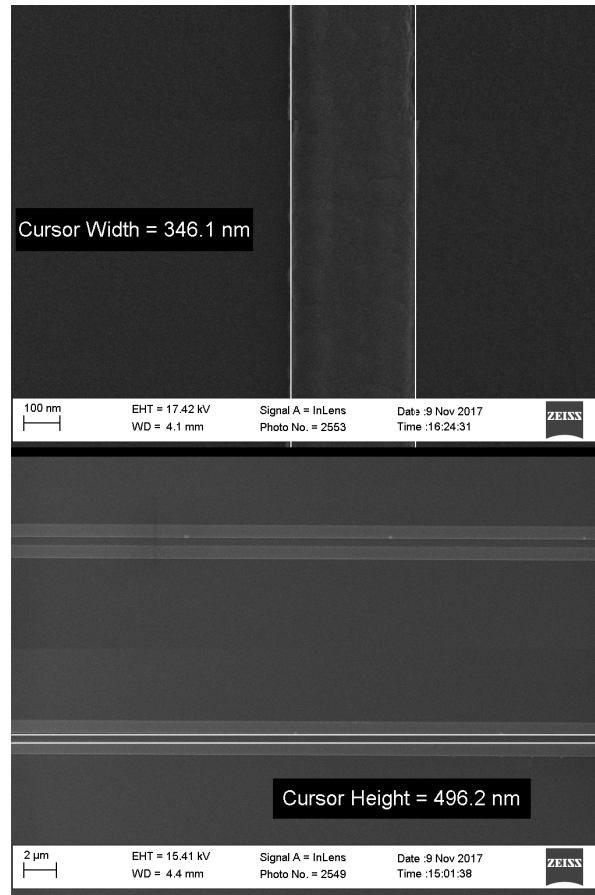


Figure 3.5: By switching to a different developer with a higher TMAH molality and no surfactents, we were able to greatly improve the resolution of our lithography and eliminate the issues shown in Figure 3.4. Smaller improvements were also made by adjusting the spin speed and bake procedures for this process.

3.5 Dry Etching

Dry etching is the primary method used in this work for turning lithographically defined patterns into actual patterns in metal or dielectric. The photoresist pattern left after performing lithography is used as an etch mask so that areas which are covered with photoresist are not etched. We then use various tools to generate plasmas or ion beams which then remove material from the surface of our sample, allowing us to remove unwanted metal/dielectric and create the desired pattern. The tools used to do this are described below.

3.5.1 Unaxis 790

The Unaxis 790 (790 for short) is a reactive ion etcher (RIE) used to etch both metal, dielectric and semiconducting materials in our cleanroom. The 790 is primarily used for flourine based processes, although it also has the ability to generate plasmas from argon, oxygen, and nitrogen. These reaction gases are introduced to the chamber through mass flow controllers, and then pumped out using a turbo pump which is throttled using a throttle gate valve. This allows a user to program recipes which involve different gas flow rates and process pressures independently. The 790 then generates a plasma in the resulting mixture of gases using an RF power supply. An automated matching network is used to impedance match the RF power supply to the plasma environment in the chamber. This creates a situation where the DC bias of the plasma is not directly controllable, and so for a particular process the plasma density and average energy of ions in the plasma are not separately controllable. In the process described in this thesis, the 790 is used primarily to etch the ground plane layer (niobium) and the final via layer

(silicon nitride). One important thing to note about this tool is that it does not have a loadlock and the base pressure of the system is typical in the 10^{-4} Torr range.

3.5.2 Plasma-Therm 770

The Plasma-Therm 770 is very similar to the 790, except that it is an inductively-coupled plasma (ICP) variant of RIE. This means that there are separate power supplies for generating the plasma and accelerating the plasma towards the sample. This allows the user to more carefully tailor the plasma density and ion bombardment energy. This tool also replaces the fluorine based recipes with chlorine, and it features a loadlock. This tool was used for performing reactive ion etches of our detector layers.

3.5.3 Kaufman & Robinson Ion Mill

Ion milling involves bombarding a sample with high energy ions which causes material on the surface to be sputtered off. This process is purely physical; the argon ions are not reactive, so there is no significant chemical component to the etching. The result is an extremely anisotropic etch. However, this can also lead to significant crosslinking in photoresist on the wafer which can lead to extreme difficulties in removing the resist after the etch is complete. Typically ion mills are used for removing a very thin layer of oxide (several nanometers) off of a metal before depositing more metal on top, thus ensuring good metal to metal contact. In this work, we also used the ion mill to completely etch through significant layers (40 nanometers) of metal.

3.6 Process Flow Overview

Now that we have covered the basic background of the tools involved in our process, we will step through a typical SCDLD process flow. Particular attention will be given to pitfalls which have been discovered in through process development and debugging. A birds-eye view of the coarse process steps is given in Figure 3.6.

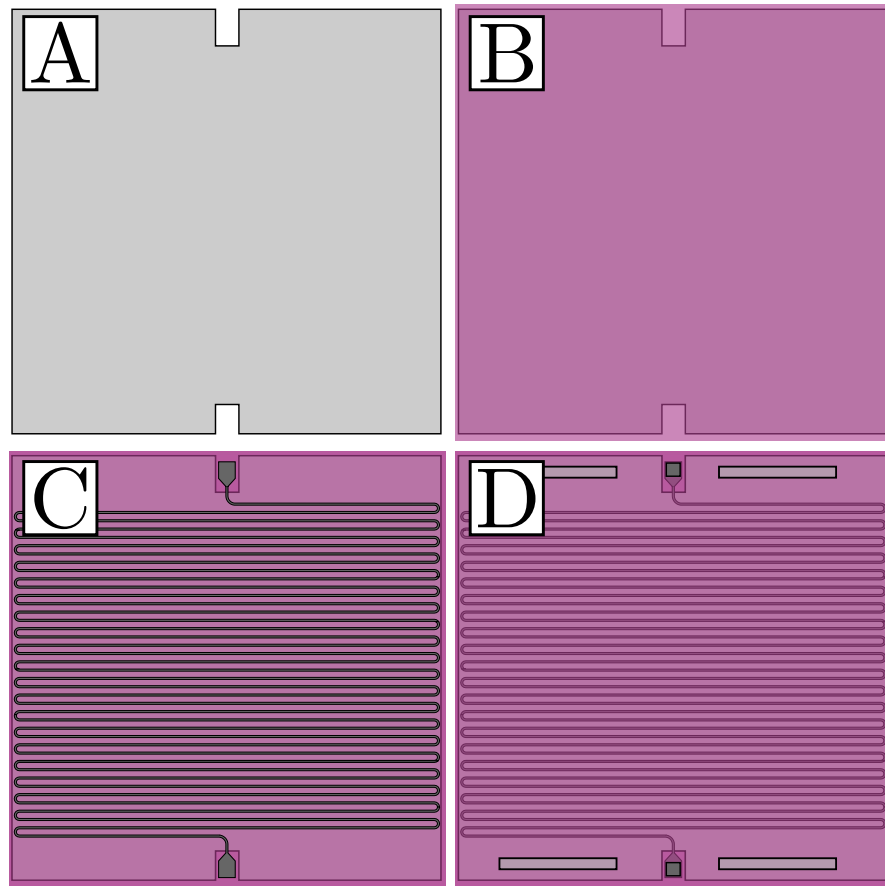


Figure 3.6: Process Flow Overview: A) Niobium Ground Layer B) Silicon Nitride Dielectric Layer C) Niobium Detector Layer D) Silicon Nitride Encapsulation Layer and VIA etch

3.6.1 Niobium Ground Plane

We begin with an oxidized silicon wafer. The manufacturer of these wafers produces these wafers in a clean environment, and they are double sealed to ensure they are not contaminated with particulates in transit. We have therefore found that the best strategy in our working environment is to do as little to the wafers as possible before depositing the first layer of metal, rather than trying to do a solvent rinse of any kind. When this thesis work was first started, standard operating procedure was to mount a wafer in an uncontrolled environment (as mentioned in Section 3.2.3) and then blow the surface of the wafer with compressed, dry nitrogen. It quickly became clear that this situation was poorly suited to fabricating devices with large active areas with any appreciable yield. With the installation of the flow hood, our yield improved significantly. Our yield was further improved after realizing the nitrogen blowing step was actually adding a significant amount of particulate contamination to the wafer surface, rather than clearing it off. In retrospect, it should have been obvious that a regularly maintained HEPA or ULPA filter would need to be installed at the output of the gun to prevent such a situation. Our deposition process proceeds as follows:

1. Put on a face mask, head covering, and eye covering (goggles or glasses).
2. Clean critical surfaces and tools inside the flow hood with isopropyl alcohol and a cleanroom wipe
3. Wipe off outer surfaces of sputter system near loadlock using isopropyl alcohol and a tekwipe
4. Vent sputter system load lock with dry nitrogen

5. Quickly transfer sample mount out of loadlock and into flow hood, then quickly close loadlock door to keep contamination out
6. Retrieve fresh oxidized silicon wafer from wafer boat
7. Center Wafer on sample mount, polished side up, and mount in place with aluminum ring
8. Holding the sample mount (and sample) upside down to minimize contamination, quickly transfer sample from flowhood into load lock of sputter system. Close loadlock and pump down as quickly as possible.
9. Follow standard pumpdown procedure and transfer sample into main chamber
10. Use argon ion mill to prepare substrate surface for deposition (improves metal adhesion to the oxide surface). Mill for approximately 15 seconds (removes roughly 7–10 nm of oxide).
11. Deposit 40 nm of niobium following the recipe shown in Table 2.
12. Transfer wafer into loadlock
13. Vent, following standard procedure
14. Remove wafer from mount, load into plastic wafer carrier. Double seal wafer carrier in plastic bags for transit to WCAM.

After transferring the wafer into the WCAM, it is time to pattern our ground plane features as follows:

Kurt Lesker Sputter System	
Argon Pressure	3.9–5 mTorr
Target Size	3 inch
Fixed Parameter	Power
Cleaning Power	400W
Deposition Power	500W
Deposition Rate	40–50 nm/min

Table 2: Niobium Deposition Parameters: The argon pressure is adjusted periodically to maintain a slight compressive stress in the niobium films[50]. The ideal pressure drifts over time, and changes based on how much material has been sputtered off of the niobium target.

1. Clean wafer of any contamination it collected while being handled in the non-clean environment of our labs by running it through a spin rinse dryer (SRD).
2. Spin coat the wafer with SPR-955 CM.7, spinning at 3000 rpm
3. Prebake the wafer at 100 °C for 90 seconds
4. Expose ground plane pattern in Body 8 stepper.
5. Post bake at 110 °C for 90 seconds
6. Develop in MF-CD26 developer for 60 seconds
7. Rinse developer off in DI water beaker for 60 seconds
8. Rinse wafer under flowing DI water for 30 seconds

9. Bake wafer at 125 °C for 180 seconds to reflow PR, creating sloped sidewalls
10. Examine wafer under microscope to verify pattern transferred correctly

After the lithography is complete, the ground plane needs to be etched. This process is performed in the Unaxis 790 RIE tool. The procedure for this process is as follows:

1. Preclean chamber with 10 minute, 500W O₂ Plasma
2. Preseed chamber with 5 minute run of plasma recipe shown in Table 3
3. Vent chamber, load sample, pump down
4. Run plasma recipe in Table 3 for 50 seconds
5. Vent chamber, retrieve sample, and check under microscope. Etched regions should appear blue.
6. Pump down sample chamber
7. Post clean chamber with 5 minute, 500W O₂ Plasma

After etching the ground plane, it is necessary to very thoroughly strip all of the leftover PR and any other organic byproducts which have been left behind by previous processing. This is generally done in the following way:

1. Preheat beaker full of Microposit Remover 1165³ to 75 °C
2. Perform coarse strip of PR residue with flowing acetone (this helps conserve 1165 by getting the bulk of the easily removable PR off ahead of time)

³Specialized resist stripper which mostly consists of n-methyl-pyrrolidinone (NMP)

Unaxis 790		
Chamber Pressure	40 mTorr	
Process Gas Flow	SF ₆	15 sccm
	O ₂	20 sccm
RF Power	150W	
Etch Rate	90 nm/min	

Table 3: Niobium Ground Etch Parameters

3. Without allowing the wafer to dry at all, submerge it in the beaker of hot 1165. Soak in hot 1165 for 30 minutes.
4. Quickly transfer wafer into beaker full of room temperature acetone without allowing it to dry at all in the process.
5. Sonicate wafer in beaker full of acetone for 15 minutes
6. Quickly transfer wafer into beaker full of de-ionized (DI) water, sonicate for 10 more minutes
7. Remove wafer, quickly place under ample supply of flowing DI water. Rinse in this way for 2 minutes
8. Quickly load wafer into SRD⁴
9. Unload, examine under microscope to check for any particulate contamination of solvent residue (which appears as discolored streaks on the wafer in an optical

⁴Spin rinse dryer

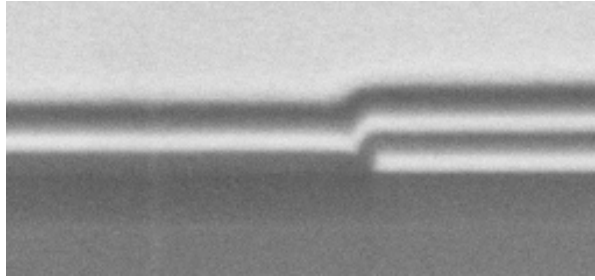


Figure 3.7: Cross Section of an SCDLD: The lowest bright layer is the ground plane of our detector. This layer is etched with the sloped sidewall process to improve the step-coverage of the detector layer which is the next bright color layer. The dark layers around these two are silicon nitride. On top of everything is a platinum layer which is only there to aid in imaging the device. Thanks to Edward Leonard who took this image on a Zeiss Focused Ion Beam-Scanning Electron Microscope.

microscope)

It is now time to deposit the dielectric layer which separates the ground plane from the detector wire(s). This is performed in the PT-70 as follows:

1. Heat the sample platen in the deposition chamber to the deposition temperature (250 °C)
2. Run plasma preclean process (O_2 and CF_4) for 10 minutes
3. Preseed by running deposition recipe (shown in Table 4) on empty chamber for 10 minutes
4. Vent deposition chamber. While waiting for chamber to finish venting, run wafer through SRD to clear any particulates that may have collected since its last cleaning

5. Load sample onto sample platen in deposition chamber, pump down chamber to base pressure
6. Run deposition recipe in Table 4 for 20 minutes
7. Vent, remove sample— set sample on a metal surface to cool before putting it back in plastic carrier
8. Run post clean plasma (same as preclean)
9. Meanwhile, examine wafer under microscope. In particular, make sure there are no obvious pinholes in the film. There should be no streaks in the color of the film either— if there is, solvent residue was trapped under the film and the yield for this wafer is likely to be very low.

PT-70			
Chamber Pressure		400 mTorr	
Sample Temperature		250 °C	
		N ₂	750 sccm
Process Gas Flow	2% Silane (in N ₂)	83.3 sccm	
	5% Ammonia (in N ₂)	200 sccm	
RF Power		100W	
Deposition Rate		7 nm/min	

Table 4: Silicon Nitride Deposition Parameters

We do not perform any lithography on this dielectric at this point. Doing so only adds unnecessary complexity to the process and more opportunities for solvent residue

or particulate contamination to cause issues later on. At this point, if a multi-wafer lot is being run, it is advisable to split the lot into multiple parts as the chances of success on a single run drop drastically from here on out. Typically, we begin a run with 4 wafers and at this point start processing them individually.

After finishing the silicon nitride deposition, the wafer should be run through the SRD, loaded into the wafer carrier, and then double bagged for transfer back to the deposition system in our lab. We then deposit another 40 nm of niobium following the procedure previously outlined. This is the detector layer. We then transfer the wafer back to the WCAM for more processing.

The lithography process for this layer is slightly modified:

1. Clean wafer of any contamination it collected while being handled in the non-clean environment of our labs by running it through a spin rinse dryer (SRD).
2. Spin coat a test wafer which also has a niobium layer on top with the photoresist (PR) SPR-955 CM.7, spinning at 5500 rpm
3. Prebake the test wafer at 100 °C for 90 seconds
4. Using a fine, clean cloth and isopropyl alcohol, wipe off the wafer stage in the stepper to remove particles which can be detrimental to the focus performance of the stepper.
5. Perform a focus exposure matrix test⁵ on the test wafer using the detector pattern in Body 8 stepper.

⁵This test automatically varies the focus offset and exposure dose from die to die across a wafer, allowing the user to determine which conditions are optimal at that moment for that particular pattern exposure (different parts of the mask often have different optimal focus offsets)

6. Post bake the test wafer at 110 °C for 90 seconds
7. Develop in MF-CD26 developer for 60 seconds
8. Rinse developer off in DI water beaker for 60 seconds
9. Rinse wafer under flowing DI water for 30 seconds
10. Carefully inspect the test wafer under a microscope and determine which exposure/focus offset combination yielded the best pattern.
11. Spin coat the real wafer with SPR-955 CM.7, spinning at 5500 rpm
12. Prebake the real wafer at 100 °C for 90 seconds
13. Expose detector pattern in the stepper, using the optimal focus and exposure parameters found previously.
14. Post bake at 110 °C for 90 seconds
15. Develop in MF-CD26 developer for 60 seconds
16. Rinse developer off in DI water beaker for 60 seconds
17. Rinse wafer under flowing DI water for 30 seconds
18. Examine wafer under microscope to verify pattern transferred correctly.

After this is complete, the detector pattern needs to be etched into the metal. This step is critical to the success or failure of a wafer. The recipe used in etching the niobium ground plane is a standard recipe used in many processes in our group for etching niobium. However, in this work, we found that etching wires with a small

cross section ($w < 2\text{ }\mu\text{m}$) this etch process severely degrades the measured critical current density of even moderately long wires (longer than approximately 1 mm). This was test by making single layer test wafers in pairs which had modified detector patterns as shown in Figure 3.8. These wafers were processed in parallell under identical conditions except for the detector etch. In one wafer, the detector was etched using the SF_6 process while the other was ion milled using the Kauffman source.

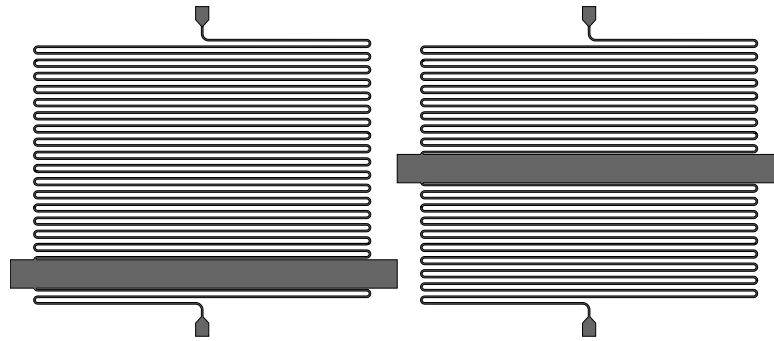


Figure 3.8: Critical Current Density Test: The wide bar in the middle is generated by modifying the stepper lithography program to purposefully not expose a band across each die. By moving this band around, we can split each detector into two shorter wires which are both shorted to ground on one side. This allows us to make measurements of the critical current density of many different wire lengths without having to purchase a special made mask with this type of test structure.

The results, summarized in Figures 3.9 and 3.10, clearly indicate that the SF_6 recipe is degrading the quality of our films. This etch recipe also has the unfortunate quality that it etches the silicon nitride under the detector at an extremely high rate (nearly twice as fast as the niobium itself). In response to these factors, we developed a chlorine based etch recipe, summarized in Table 5. The tool this etch is performed in is also

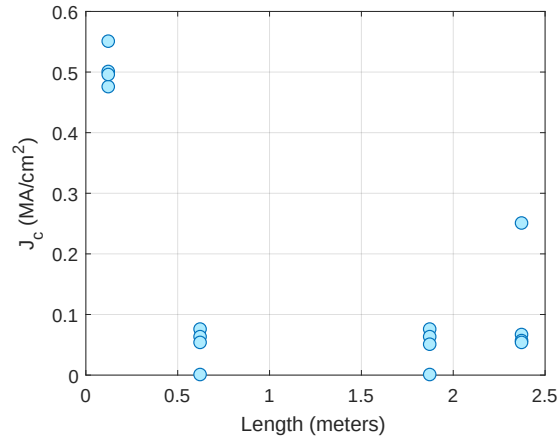


Figure 3.9: Critical Current Density vs. Length Ion Mill

loadlocked, which should cut down on exposure to ionized oxygen during the etch process.

The procedure for performing this etch is:

1. Preclean etch chamber and carrier wafer with 10 minute oxygen plasma
2. Preseed etch chamber and carrier wafer with 10 minute run of etch recipe with no sample
3. Mount sample wafer onto 6 inch carrier wafer using a tiny amount of santovac between them
4. Run etch recipe shown in Table 5
5. Remove carrier wafer, warm to 50 °C and then sample wafer from carrier wafer.
Rinse with water in case any volatile chlorine byproducts remain on the wafer

This etch recipe worked well sometimes, but occasionally the etch chamber would become contaminated by another user's process, and our measured critical current densities would be suppressed across every wafer we ran through the tool. We attempted to

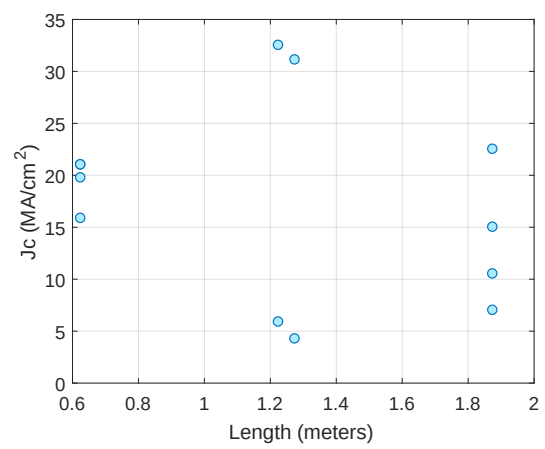


Figure 3.10: Critical Current Density vs. Length Ion Mill

PT-770		
Chamber Pressure	40 mTorr	
Process Gas Flow	BCl ₃	15 sccm
	Cl ₂	20 sccm
	Ar	20 sccm
RIE Power	50W	
ICP Power	300W	
Etch Rate	17 nm/min	

Table 5: Niobium Detector Chlorine Etch Parameters

use plasma spectroscopy to detect what contamination was present that was degrading our films, but when we compared the spectra of etches which yielded good devices versus bad we weren’t able to see any significant differences.

Due to the unreliability of this etch, we focused on just using the Kauffman Ion

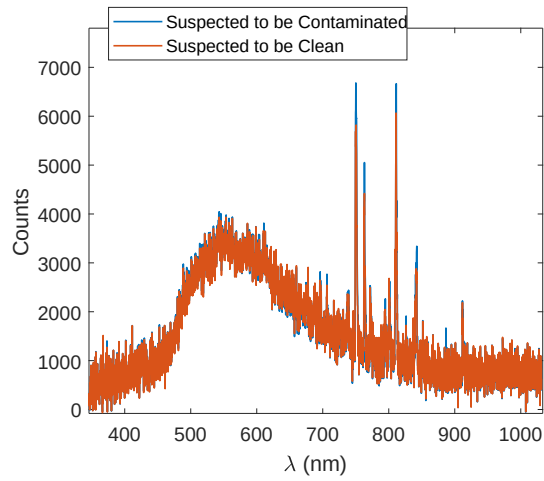


Figure 3.11: Niobium Etch $\text{BCl}_3+\text{Cl}_2+\text{Ar}$ Plasma Spectroscopy

Source to etch our patterns. The recipe for this etch is shown in Table 6.

Kauffman Ion Source	
Argon Pressure	.2 mTorr
Beam Voltage	600 V
Accelerator Voltage	120 V
Beam Current	46 mA
Neutralizer Current	51 mA
Etch Rate	24–30 nm/min

Table 6: Niobium Ion Mill Etch Parameters

After etching the detector layer, it is critical that as much of the remaining photoresist residue as possible is stripped. This can be especially difficult to do after ion milling. The procedure to do this is:

1. Preheat beaker full of Microposit Remover 1165⁶ to 80 °C
2. Perform coarse strip of PR residue with flowing acetone (this helps conserve 1165 by getting the bulk of the easily removable PR off ahead of time)
3. Without allowing the wafer to dry at all, submerge it in the beaker of hot 1165. Soak in hot 1165 for 3 hours. Meanwhile, preheat shared 1165 tank (which is built into a chemical bench) to maximum temperature of 75 °C.
4. Transfer wafer into shared 1165. Leave it there overnight.
5. The next day, preheat a beaker of acetone to 75 °C
6. Quickly transfer wafer into the beaker of acetone without allowing it to dry at all in the process.
7. Soak in the hot acetone for at least 30 minutes.
8. Sonicate wafer in beaker full of acetone for 30 minutes
9. Quickly transfer wafer into beaker full of de-ionized (DI) water, sonicate for 10 more minutes
10. Remove wafer, quickly place under ample supply of flowing DI water. Rinse in this way for 2 minutes
11. Quickly load wafer into SRD⁷

⁶Specialized resist stripper which mostly consists of n-methyl-pyrrolidinone (NMP)

⁷Spin rinse dryer

12. Unload, examine under microscope to check for any particulate contamination of solvent residue (which appears as discolored streaks on the wafer in an optical microscope).

After this heavily crosslinked resist is stripped, it is time to deposit the encapsulation layer. The procedure is similar to the procedure for depositing the wiring dielectric as described above, but slightly modified:

1. Heat the sample platen in the deposition chamber to the deposition temperature (250 °C)
2. Run plasma preclean process (O_2 and CF_4) for 10 minutes
3. Preseed by running deposition recipe (shown in Table 4) on empty chamber for 10 minutes
4. Vent deposition chamber. Turn off the sample platen heater. Allow the sample platen temperature to fall below at least 130 °C. When it is close to this temperature, run the wafer through the SRD one final time.
5. Load sample onto sample platen in deposition chamber, pump down chamber to base pressure
6. Run deposition recipe in Table 4 for 20 minutes
7. Vent, remove sample— set sample on a metal surface to cool before putting it back in plastic carrier
8. Run post clean plasma (same as preclean)

9. Meanwhile, examine wafer under microscope. In particular, make sure there are no obvious pinholes in the film. Solvent/PR residue isn't great at this point, but it shouldn't significantly affect the yield. It may, however, cause devices to have non-uniform propagation velocities.

In order to access the detector and ground plane, it is necessary to etch vias (holes) through both the encapsulation layer that was just deposited as well as the wiring dielectric layer which was deposited several steps ago. First, perform lithography as follows:

1. Spin coat the wafer with SPR-955 CM.7, spinning at 3000 rpm
2. Prebake the wafer at 100 °C for 90 seconds
3. Expose via etch pattern in Body 8 stepper.
4. Post bake at 110 °C for 90 seconds
5. Develop in MF-CD26 developer for 60 seconds
6. Rinse developer off in DI water beaker for 60 seconds
7. Rinse wafer under flowing DI water for 30 seconds
8. Examine wafer under microscope to verify pattern transferred correctly

This will generate a pattern with holes connecting to both the detector layer and the ground layer. Normally this would be inadvisable because there are different thicknesses of silicon nitride between the top surface and those two layers. However, our recipe for etching silicon nitride (shown in Table 7) has extremely good selectivity between silicon

nitride and niobium; that is, it etches silicon nitride at approximately 125 nm/min and niobium at only 5 nm/min. This means we can tolerate etching the detector part of the niobium slightly longer than necessary in order to etch the ground bond pads through the extra 100 nm of nitride. The procedure for this etch is given below:

1. Preclean chamber with 10 minute, 500W O₂ Plasma
2. Preseed chamber with 5 minute run of plasma recipe shown in Table 7
3. Vent chamber, load sample, pump down
4. Run plasma recipe in Table 7 for 120 seconds
5. Vent chamber, retrieve sample, and check under microscope. Etched regions should appear silver.
6. Pump down sample chamber
7. Post clean chamber with 5 minute, 500W O₂ Plasma
8. Probe bond pads which should be electrically connected. Verify that the measured resistance is correct (roughly 20 Ω for ground plane pads within the same die, roughly 10 M Ω for pads on either side of a detector)

At this point, we've completed all of the processing steps for the wafer. It is useful to probe all of the detectors on the wafer as described in Section 4.2. After this, the wafer can be diced into separate dies and fully characterized as described in Chapter 4.

PT-790		
Chamber Pressure	100 mTorr	
Process Gas Flow	CHF ₃	50 sccm
	O ₂	20 sccm
RF Power	150W	
Etch Rate	125 nm/min	

Table 7: Silicon Nitride Etch Parameters

Chapter 4

SCDLD Measurement

4.1 Measurement Setup Overview

In this chapter, the methods for characterizing superconducting delay line detectors (SCDLD) will be described. The yield of these devices has proven to be quite low, and so a multi-step screening process was developed which allowed the devices to be measured quickly and efficiently. The active area of each device is rather large, and even small defects can render the entire device useless. Because of this, manual inspection (using an optical microscope, SEM, et cetera) is only useful in the case of gross defects which are rare and easily distinguished. After this, room temperature DC measurements are performed to screen devices that meet the bare minimum requirements for usefulness; this will be described in section 4.2. In section 4.3, measurements performed in liquid helium will be described. These measurements provide enough information to determine which devices will be useful as detectors. The detectors that pass this screening test will then be tested by bombarding them with ions from various sources, as described in section 4.4. These experiments produce large amounts of data— the algorithms for processing this data will be presented in section 4.5.

4.2 Room Temperature Screening

As described in the Fabrication chapter, our devices are fabricated on 3 inch wafers which hold can well over 50 dies. The first thing that must be done after finishing a wafer is to screen the devices at room temperature. This is done using a Signatone S-1170 Probe Station. This setup allows us to measure the resistance of detectors, test structures, and check for shorts between the detector to ground before we even dice the wafer into dies. An entire wafer can be thoroughly probed in about one hour— performing individual measurements on dies in liquid helium would take several weeks worth of time. These measurements are used to screen out devices which are not worth measuring in liquid helium.

4.3 Liquid Helium Screening

Devices that passed the room temperature screening are then measured at 4.2 K in liquid helium. Individual dies are mounted on a sample test mount and are connected to test ports by wire bonding. The sample mounts we used have four RF test ports, so if each die has a single detector four can be dipped at once, or if they are Pitch Experiment or Pitch Imager dies, we can dip two at once. Each SCDLD has one side wire bonded to an RF port and the other side wire bonded to ground. There are also short ($\sim 200\text{ }\mu\text{m}$) test wires with the same cross sectional area as the SCDLD on each die which can be connected instead of the full detector for fabrication debugging purposes. The ground plane is then connected to the ground of the sample mount with numerous wire bonds to provide a clean reference plane. After the samples are mounted and bonded, the sample mount is connected to a dip probe via four rigid coaxial lines, and then surrounded with

a mu-metal shield.

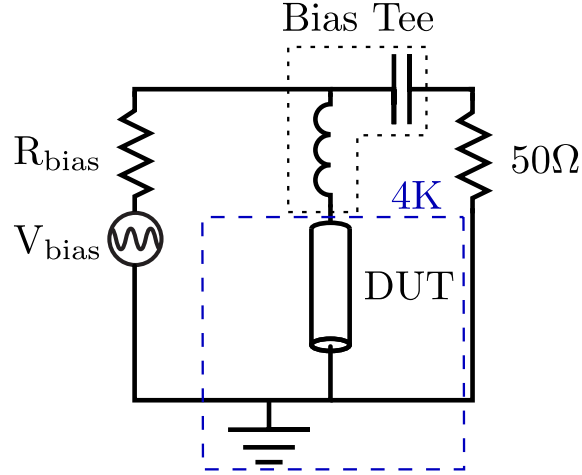


Figure 4.1: Circuit used for measuring current-voltage curves. The detector under test (DUT) is marked as a transmission line.

The first type of measurement we take in liquid helium is a current-voltage (“IV”) curve. This allows us to determine if the SCDLD or test structure transitions into the superconducting state, and what the effective critical current density (J_{ceff}) for the structure is. This number is essentially never equal to the theoretical critical current density (J_{crit}) due to inhomogeneities in the line which create local constrictions which suppress the measured critical current density for the entire line. The method for measuring these IV curves is shown in Figure 4.1. It is only necessary to perform a two wire measurement in this case because the residual resistance of an SCDLD is much larger ($\sim 3\text{ M}\Omega$) than the lead resistance in the rigid coaxial lines ($\sim 2\Omega$). A bias-tee is connected in series with the detector to act as a low pass filter to reject RF noise which can suppress the apparent critical current density of the device under test.

After measuring IV curves, we perform Time Domain Reflectometry (TDR) on each

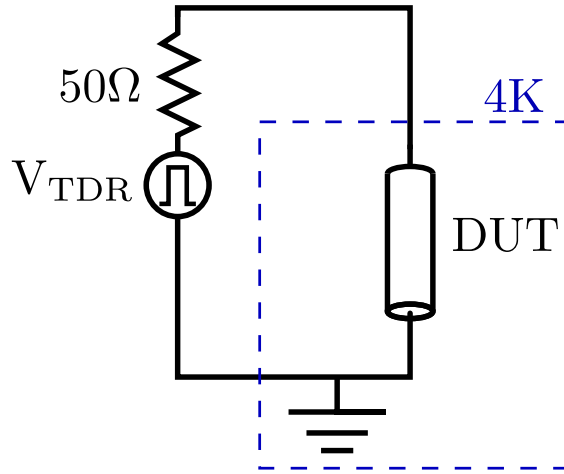


Figure 4.2: Circuit used for performing time domain reflectometry (TDR). The detector under test (DUT) is marked as a transmission line. A sampling oscilloscope, Tektronix 11801C is used to both generate the voltage edge and measure the reflected TDR pulse for this measurement.

sample. The circuit used for this measurement is illustrated in Figure 4.2. In this method, a voltage step with a very fast rise time (~ 100 ps) is generated by a TDR sampling head and transmitted down a transmission line (in this case, a set of 50Ω coaxial cables) and into the device under test. While this is happening, the TDR sampling head measures the sum of the voltage edge it generated and the reflected voltage coming out of the test circuit as a function of time. This measured voltage can then be used to calculate the apparent impedance as a function of time. With knowledge of the propagation velocity of the transmission line(s) under test, it is possible to then turn this into a plot of characteristic impedance as a function of distance into the test circuit¹.

¹Unless there are significant reactive impedances involved which complicate the deconvolution of impedance variations in time back into impedance variations in space

This type of measurement is extremely valuable in characterizing SCDLDs which are extremely long physically and electrically. In the case that the device we are measuring appeared to have a reasonable J_{eff} , TDR will allow us to measure the characteristic impedance and propagation velocity of the device. Additionally, there are several subtle defects that a DC measurement wouldn't reveal (such as a short between two adjacent wires or a short to ground far from the input port) which TDR will illuminate. On the other hand, if the measured critical current density is lower than expected, performing TDR can help narrow down what may be wrong with the device. In TDR, weak points in a detector (short segments of the line with suppressed critical current density) appear as essentially open circuits. For instance, if step coverage was an issue with a particular wafer, TDR will show almost every die as being an open circuit at the point when the voltage edge first encounters the detector. If, on the other hand, random defects of some kind are a limiting factor, TDR data for different dies will have apparent open circuits at different places. Furthermore, an estimate of defect density can be made by looking at the distribution of where in the line the first defect appears (it is often difficult or impossible to observe any defects beyond the first one). There is also the possibility of mask defects being a limiting factor, in which case the TDR data from different dies from the same wafer will all have defects at the same point.

4.4 Detection Experiments

After extensively screening SCDLD using the measurements discussed in Sections 4.2 and 4.3, the remaining devices are ready to be tested as ion detectors. First, the various ion sources used in this thesis work will be described. Then the electronics (both readout

and bias) which are used to perform these experiments will be described. After this, the method for collecting data will be discussed, followed by the software for analyzing that data.

4.4.1 Ion Sources

4.4.1.1 Radioactive Sources

In some cases it is desirable to generate ions with a uniform flux. This can be by using radioactive sources which generate α particles. ^{241}Am is a readily available source of α particles. It should be noted, however, that the α particles generated by ^{241}Am have energies of around 5.5 MeV which is a factor of 1000 higher than the helium ions made by field ionization (FIM) which will be discussed in Section 4.4.1.2. This can be useful for running tests where it is useful or informative to expose detectors to higher energies than are accessible with our FIM setup. It is also worth noting that not all of those 5 megaelectron volts actually end up in the detector. We can use SRIM (as described in 2.4) to simulate the amount of energy deposited in our detector by these high energy particles. A typical energy deposition for our geometry/stack would be roughly 45 keV. The energy deposited is almost always mostly in the electronic system of the detector stack. This is because the cross section for nuclear scattering falls sharply with energy in this case, so nuclear collisions which actually scatter a significant amount of energy are extremely rare.

4.4.1.2 Field Ionization Source

Although ^{241}Am was a useful starting point, the energy of the ions are not representative of the expected energy ranges involved in either TOF spectrometry or APT. In order to generate ions that are a more accurate representation of those created in an atom probe, we designed and implemented a small scale field ion microscope (FIM) in our cryostat. The mechanism for generating field ions in this way is shown in Figure 4.3.

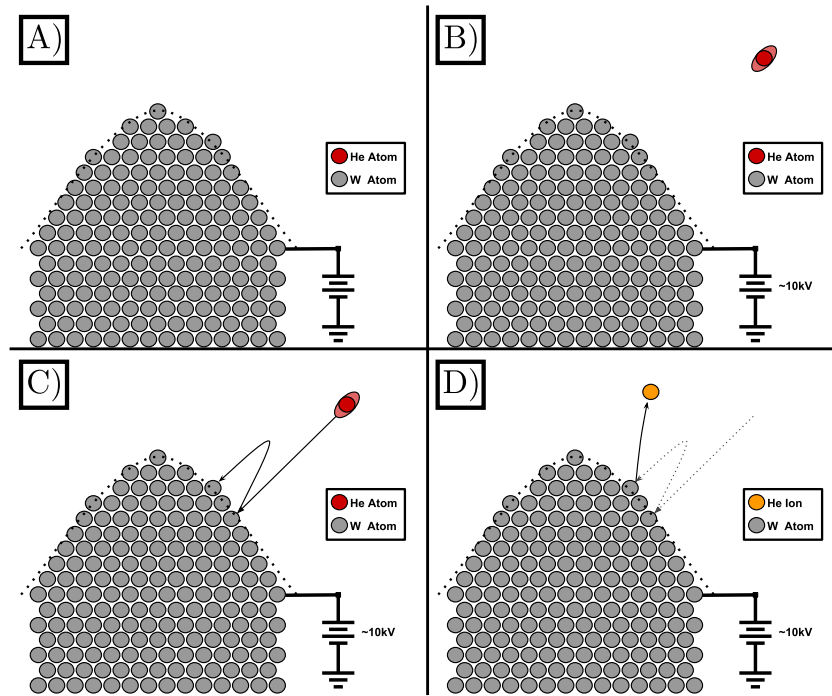


Figure 4.3: Field Ion Microscope Ion Source: A) The core element of the source is a tungsten wire which has been electropolished to a very sharp tip (tip radius of 10 nm to 100 nm) B) Imaging gas is bled into the chamber and a voltage bias is applied to the tungsten tip C) The electric field surrounding the tip polarizes imaging gas atoms which are then attracted towards the tip D) After colliding with the wire several times, an electron tunnels off the imaging gas atom, thus ionizing it. The resulting ion is repelled away from the tip, and then is collected in a detector (which would be above the tip in this drawing).

An important aspect of the ions generated in this way is that the ion flux generated by the tip is not uniform; the emission coming off the tip depends on the shape of the tungsten. The probability of ionizing an imaging gas atom at some point on the tip is dependent on the electric field strength at that point. Therefore, places where the

tungsten atoms' electron orbitals are protruding from the tip generate a larger flux of ions. The ion flux then creates an image of the surface of the emitter as shown in Figure 4.4.

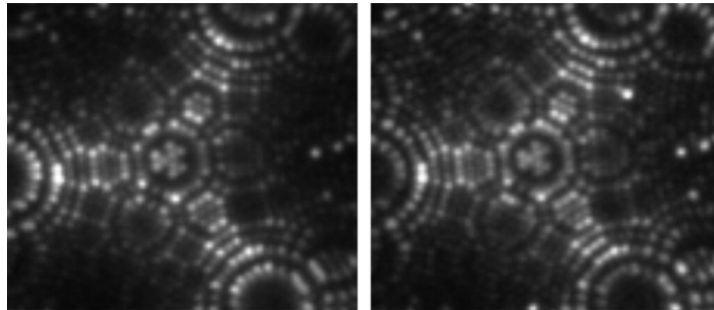


Figure 4.4: Field Ion Image Example: Generated using a tungsten tip and helium (left) and neon (right) imaging gas. Taken from Reference[47]

In our implementation, we include a gas line which can supply a constant, low flow of imaging gas. This is due to the fact that our cryostat is much colder than a usual FIM setup, and imaging gas can condense and/or adsorb onto surfaces. It has been observed that if the flow of imaging gas is stopped, the ion current drops off to zero over time. This imaging gas is plumbed in through a series of swagelok tubes and valves. Although this system was carefully assembled and leak checked, even trace amounts of gases which are inert under normal conditions (such as water, N_2 , et cetera) can be corrosive to the tip[47][10]— therefore, the imaging gas line which feeds down into our cryostat connects to a heat exchanger at the 50 K stage as illustrated in 4.6. This serves two purposes: it precools the imaging gas, and it freezes out any impurities which may have leaked into the imaging gas line.

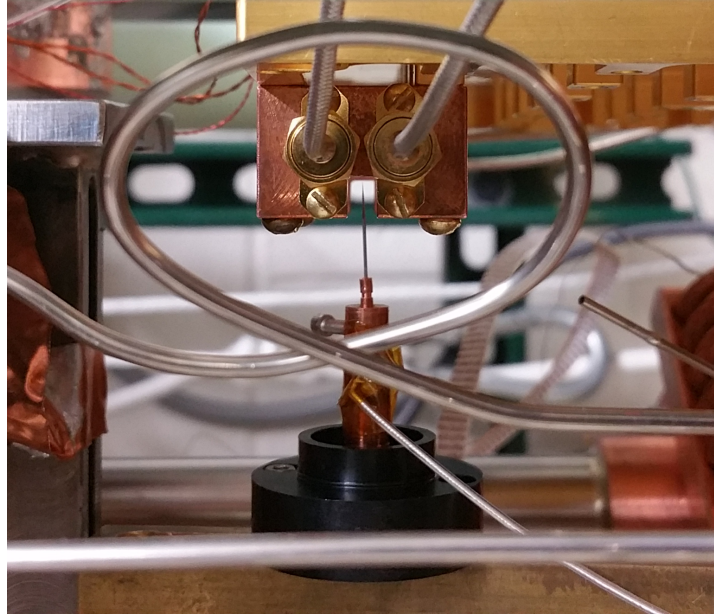


Figure 4.5: Picture of actual Field Ion Microscope setup in cryostat. In the center of the image is the tungsten emitter which is held in a copper holder which is screwed into an insulating mount made of PEEK. This is then connected to a pair of adjustable arms so that the emitter can be positioned over the detector. The detector is not directly shown; it is mounted in the copper sample mount near the top of the image with two SMA cables connected. These are the two wires connecting to the upstream ports of the detectors of a Pitch Experiment die.

4.4.2 Electronics

Part of the original design goal of the SCDLD was to create a detector which is relatively simple to operate and collect data from. As a result of this focus on simplicity, our electronics requirements are relatively simple, even for operating devices in the development phase where it is necessary to have the flexibility to store raw waveforms generated by the detector and be able to adjust bias current waveforms on the fly. Figure 4.6 shows an

overview of the main electrical systems necessary to operate our detector as well as the FIM ion source which we use to test it. In this section, the various subsystems involved in biasing and reading out our delay line detectors will be discussed.

Our experiments are carried out in between the two stages of a pulse-tube cooled cryostat. The experiments are located in an area that has access to a 50 K stage (where the FIM tip is mounted and the incoming imaging gas is cooled) and a 3 K stage (where the detector and its wiring is located).

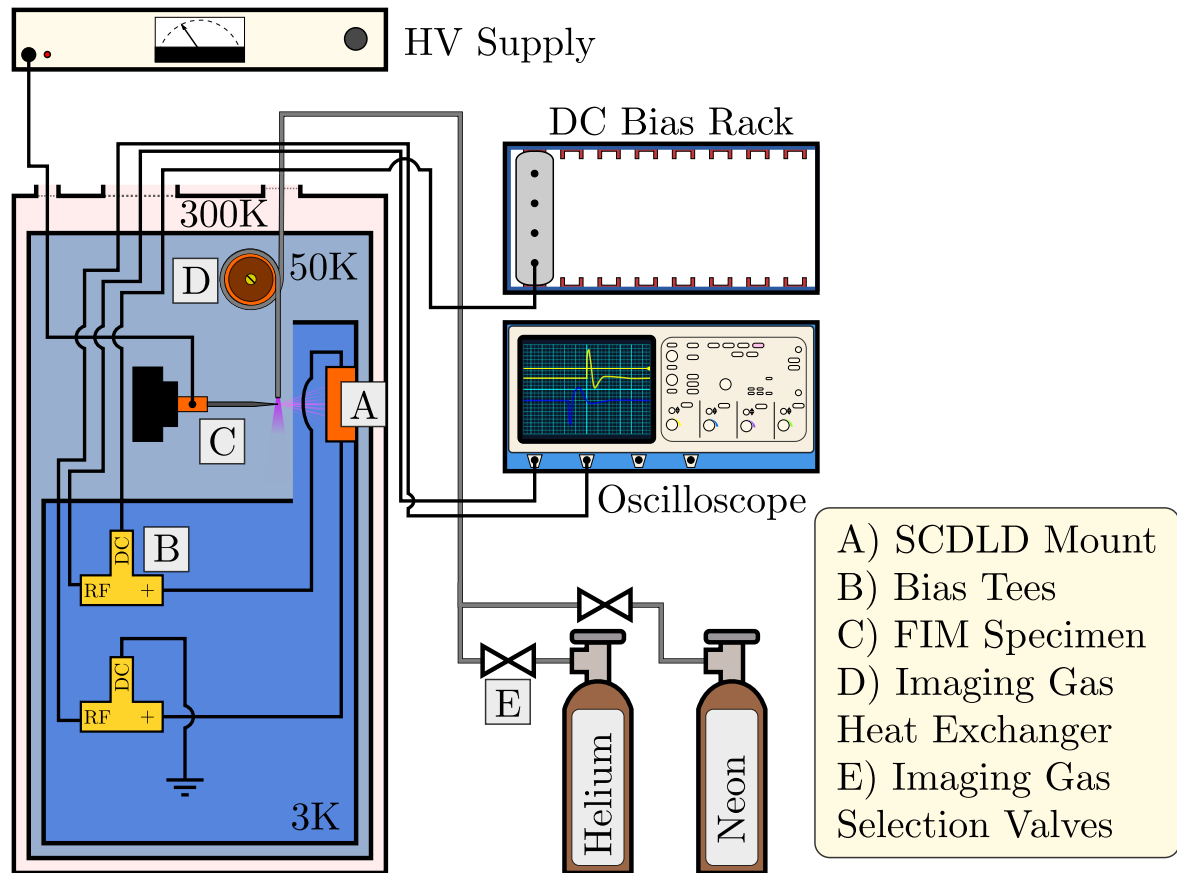


Figure 4.6: Wiring diagram for ion detection experiments. The necessary electronics and gas lines for generating ions as a FIM are also illustrated.

4.4.2.1 Cold Wiring

Our detectors only require a pair of bias tees in the cryostat for biasing and readout. Due to the fact that we are working at a relatively high temperature (3 K), rather than, say, the millikelvin stage of an ADR²), there is a significant amount of cooling power available (~ 1 W). This means that the thermal conductivity of our wiring is less of a concern than in millikelvin temperature experiments. Our wiring consists of the following:

1. Semi-rigid coaxial cable with a copper nickel ground, teflon dielectric, and silver plated copper nickel inner conductor (which improves their conductivity at microwave frequencies without adding excessive thermal conductivity/heat load to the fridge). These RF lines connect the feedthroughs in the vacuum chamber which contains the cryostat all the way to the 3 K stage of the pulse tube cooler. They are heat sunk with wide copper straps which are tightly bound around the coax (with a layer of vacuum grease in between for improved thermal conductivity) and mechanically anchored to the 50 K stage. They have similar heat sinking straps at the 3 K stage and also connect to bulkhead feedthroughs which provide additional thermalization channels. These RF lines have a loss coefficient of approximately 3 dB/m at 5 GHz (not including connector losses).
2. Handformable cables which provide connections between components within the 3 K stage. These cables consist of a silver plated copper inner conductor³, a teflon

²adiabatic demagnetization refrigerator

³This material has a much higher thermal conductivity than the cupronickel of the semi-rigid lines, hence why it is only used to connect components which are mounted at the same temperature stage of the cryostat

dielectric, and a copper outerbraid which is coated in tin. These cables, if connected properly, can achieve attenuations as low as .6 dB/m at 5 GHz.

3. Bias tees, manufactured by Anritsu which have proven to be robust through numerous cooldowns. These allow us to provide a quasi-DC bias current to the detector through a large inductance (designed to have a very high self-resonance frequency) while simultaneously reading out broad-bandwidth pulse coming out of the detector through a capacitively coupled port.
4. Transitions from coaxial cable to microstrip are made from patterned Duroid materials. These small sections of circuit board provide a clean, flat surface suitable for wire bonding. These wire bond connections go directly to the sample bondpads (both detector and ground plane). An image showing how this is implemented in the reality is shown in Figure 4.7.

4.4.2.2 Biasing Electronics

Early experiments performed in this thesis were performed using a DC voltage source to generate the bias current for our detectors. The voltage source used at that time was a SIM928 module, manufactured by Stanford Research Systems. However, after early measurements and simulations suggested that detecting ions with our SCDLDs would require us to operate in the latching regime, this DC source was replaced by a variety of different pulsed voltage sources. In the end, a specialized bias rack originally designed for performing fast biasing operations on qubits was employed due to its programmability, noise performance, and its lack of overshoot when generating voltage edges. This source is used to generate the high duty cycle square waves which are used to bias our detectors.

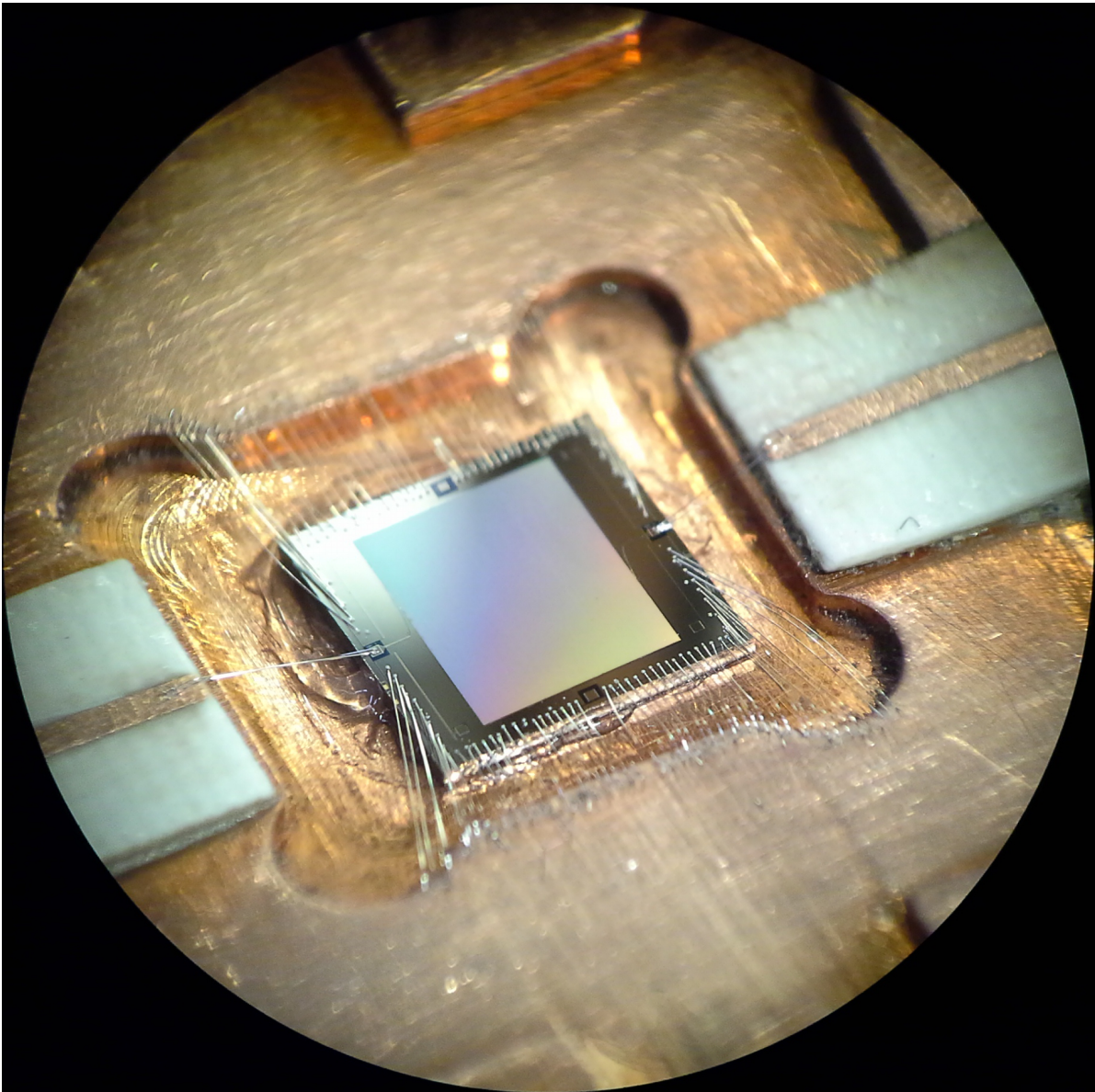


Figure 4.7: Basic Ion Detector Wire Bonded to Sample Mount: The entire copper mount is electrically grounded. The two microstrips on either side of the detector are silver pasted to the copper sample mount to ensure a solid ground connection all the way to the point where wire bonds are made. These microstrips are designed to be $50\ \Omega$.

Due to the fact that ions are generated stochastically in our current experiments, this waveform is not synchronized to anything; its repetition rate is set to be at least roughly 10 times the apparent event rate in the detector to minimize event loss due to the detector being in a latched state when an ion arrives.

4.4.2.3 Read Out Electronics

Very early on we experimented with reading detectors out using Nuclear Instrumentation Module (NIM) bin based electronics to measure event rates, relative delay distributions, et cetera. Procuring modules like these proved to be somewhat difficult, as many of them haven't been manufactured for years, and used modules proved to have reliability issues in some cases. Although analog electronics like these could be designed and used for simple read out of SCDLD in the future, we decided it was more suitable at this early stage of development to instead utilize high bandwidth oscilloscopes to collect raw waveforms from our detectors. This allows for more in depth analysis of the types of signals coming from our detectors, and also opens the possibility of storing raw data which can then be analyzed with different analysis procedures, even years later — something which could not be done with data generated hardware based time tagged data. This has proven to be an invaluable capability.

The primary oscilloscope for data collection used was a Tektronix CSA7404B which has a maximum (single channel) sampling rate of 20 GSa/s, an input bandwidth of 4 GHz, and maximum sensitivity of 2 mV/div. Another useful feature of this oscilloscope is a capability called “FastFrames”. This feature allows the user to capture a large number of separate waveforms which are all stored in one longer record, which can then be written to disk.

4.4.3 Basic Ion Detection Experiments

In the basic ion detection experiments a measurement can be made of the rate of ion events as functions of both bias current and FIM voltage. In the ideal case, the event rate should saturate as a function of bias current below J_{ceff} indicating that the detector has reached its peak detection efficiency. Furthermore, the event rate should be essentially zero when the FIM voltage is low, and then suddenly turn on at some minimum bias voltage. These rates are tracked in the measurement software by measuring how long it takes to fill a FastFrame record and then dividing by the number of waveforms stored in the record. Although this measurement method is not perfect due to the fact that there can be some latency at the beginning and end of filling a FastFrame record, it is possible to test the accuracy of this method. To do this, the amplitude of the current bias pulses should be set above the critical current of the detector. Then, the frequency of the current bias pulse train set to some known value to test. Then, the data collection program should start. Every time a bias pulse arrives at the detector, it will be driven normal and produce an “overdrive” event where the part of the wire with the lowest critical current density forms a normal zone and generates a pulse pair. Therefore, the event rate measured should be equal to the bias pulse frequency. We have performed these simple test and found that as long as the number of frames is approximately ten times larger than the number of events per second, the measured event rate is accurate to within about 5%. Above the minimum FIM turn on voltage, the event rate will continue increasing with increasing bias voltage. The pattern of emission will also change somewhat with FIM voltage as contaminants and protruding tungsten atoms can evaporate off as we increase the voltage, thus creating a new image.

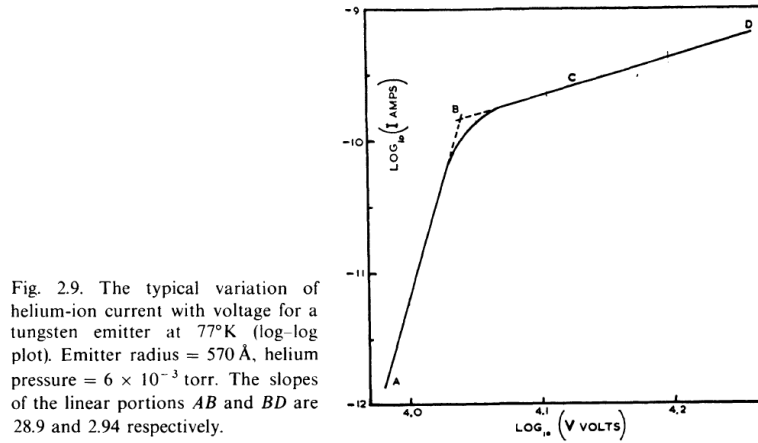


Figure 4.8: Event Rate vs FIM Voltage Bias Measured with Microchannel Plate: Taken from [40]. This plot shows the standard behavior of near zero emission below a threshold voltage followed by a rapid increase in emission current, followed by saturated behavior where the current only weakly depends on the bias voltage.

4.4.4 Relative Delay Resolution Measurement

The experimental set up for measuring our timing resolution with respect to relative delay is very similar to the basic ion detection experiments. The difference is that the FIM ion source is replaced with an ^{241}Am source (so that the ion flux generated will be uniform). This source is mounted on a block of copper which was designed to clamp onto our detector sample mount and has a very narrow and deep slot machined through it. The purpose of this mount is to block most of the α particles except for a narrow strip of illumination down the center of the detector. The layout of this experiment is shown in Figure 4.9.

Through some simple geometric calculations and by measuring the propagation velocity of the SCDLD under test, the ideal image/histogram which should be produced

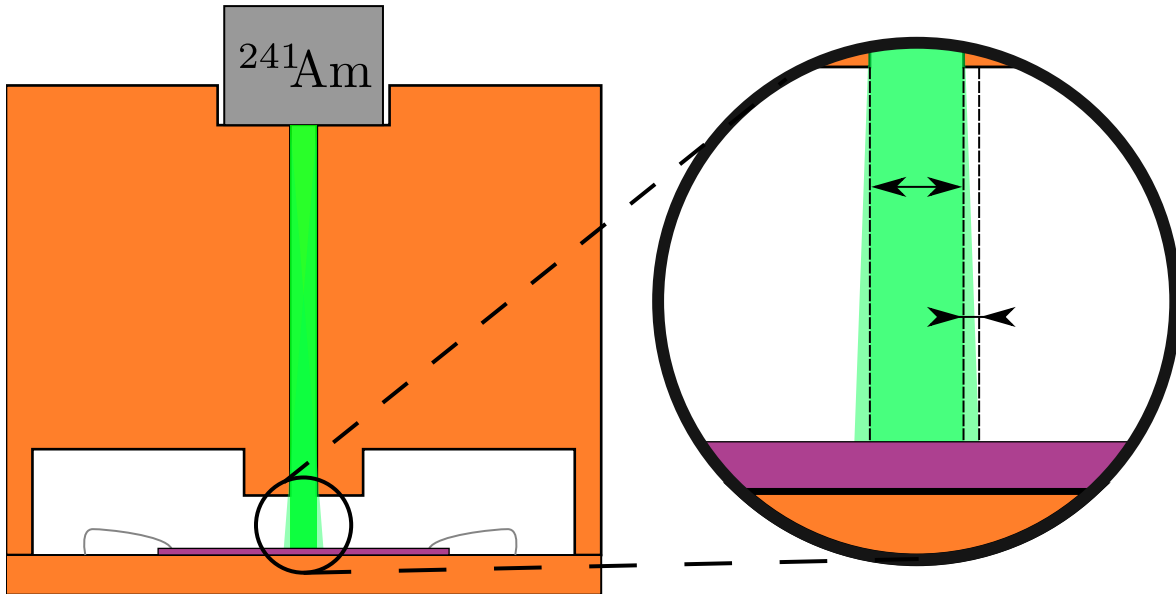


Figure 4.9: Relative Delay Resolution Measurement Setup: The ion source is ^{241}Am , held in a mount which also acts as a shadow mask, creating a narrow band of ion flux on the detector.

can be calculated. By comparing this to the measured histogram the measured timing resolution in relative delay can be calculated.

4.4.5 Pitch Experiment

The basic detection experiments described above provide a certain amount of information about a detector's performance; however, without careful calibration of the ion emission current from the FIM source and a tooling factor to account for pattern emission actually impinging on the detector, determination of the detection efficiency of the detector is impossible. It is possible, in principle, to calibrate for both of these factors. There is another method for measuring the detection efficiency which is simpler, and can provide

more detailed information on the detector's response to ion impacts.

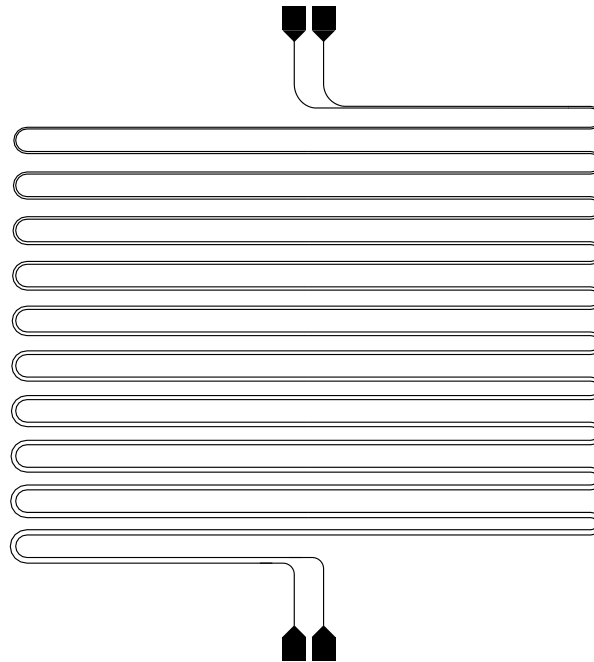


Figure 4.10: Illustration of Pitch Experiment Concept

This method calls for fabricating a pair of detectors on the same chip, laid in parallel. These two detectors are meandered together, with a gap that varies across the die. On one side of the chip, the two detectors are as close together as possible- on the other, they are farther apart (typically a few microns). The sample is then loaded into a cryostat, and run under essentially the same conditions as the Basic Ion Detection Experiments. The only difference now is that there are two sets of bias lines, and two sets of readouts. After the experiment is run, the pulse outputs can be analyzed in the normal way. The analysis is also augmented to check for correlated events between the separate wires which would indicate than an ion managed to trigger both lines. By calculating the ratio of correlated to uncorrelated events as a function of the separation between the wires, it is possible to calculate the detection efficiency of the detector as a function of

the ion's impact position relative to the detector. More details on the calculations of these parameters is presented in Section 2.5.1.

The experimental setup for the Pitch Experiment is very similar to the Basic Ion Detection experiment. The main differences are that now there are four bias tees (a pair for each detector on the chip), four channels used on the oscilloscope, and a new sample mount which can accommodate four RF connections. An illustration of this sample mount is shown in Figure 4.11. This redesigned sample mount also features a window which can be used for aligning FIM tips more simply. In Figure 4.5 this redesigned mount is being used to hold the detector sample and align the tungsten field source.

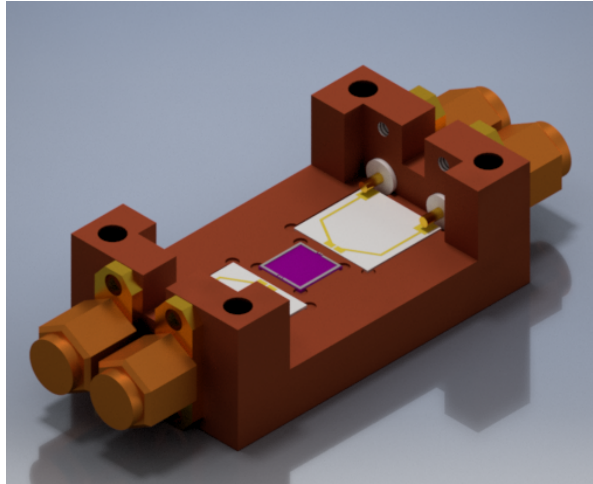


Figure 4.11: Pitch Experiment Sample Mount

4.5 Data Analysis Procedures

Ion detection experiments generate a large amount of raw data which needs to be processed into meaningful plots. In this section, the various methods used for turning a large collection of pulse waveforms into useful plots will be described.

4.5.1 Optimal Filtering

As alluded to in Section 2.6, the accuracy of our time tagging algorithms can be improved by first filtering our pulses. The strategy for doing so is to filter our signal such that the signal to noise ratio(SNR) of the waveform is optimized. An estimation of the ideal pulse edge's shape can be made based on theoretical grounds and from actual measurements. Furthermore, the noise present in the signal will be approximately white— that is, it has a constant spectral density at all frequencies of interest. Using these two pieces of information, it is possible to calculate a filter which optimizes the SNR in the frequency domain by setting the transfer function at any one frequency to be proportional to the ratio of the spectral density of the signal divided by the spectral density of the noise. In this case, doing so is very simple— because the noise is white, the transfer function is simply proportional to the spectral density of the signal at all frequencies.

$$H_{opt}(f) \propto \frac{S_{V_{ideal}}(f)}{S_{V_{noise}}(f)} \propto S_{V_{ideal}}(f) \quad (4.1)$$

This suggests that the impulse response of the matched filter is given by the time reversed version of our ideal waveform[45]. This means that by simply convolving the measured signal with a time reversed version of the idealized waveform results in an optimally filtered⁴ signal:

$$H_{opt}(f)S_{V_{sig}}(f) = \alpha V_{ideal}(-t)^* * V_{sig}(t) \quad (4.2)$$

However, if the the timing analysis is being performed via the threshold detection

⁴With respect to signal to noise ratio

mechanism (described in Section 4.5.2.2), the improvement in SNR achieved by performing this filtering procedure is somewhat offset by the suppression of the rise time of the signal. When using this time tagging method, the simple boxcar⁵ filter is actually optimal[39] for reducing noise while preserving the rise time of the time domain waveform. The response of the boxcar filter, $y(n)$, acting on the discrete time series, $x(n)$, is given by:

$$y(n) = \frac{1}{2p+1} \sum_{m=n-p}^{m=n+p} x(n-m) \quad (4.3)$$

4.5.2 Time Tagging Algorithms

Perhaps the most critical measurement taken in our ion detection experiments is the assignment of time tags to pulses coming out of the SCDLDs. There are a variety of different methods we have explored which will be explored below. In general, the more robust (accurate) methods have the trade-off of a longer computation time which can be limiting, especially in the case of large data sets.

4.5.2.1 Interpolation Method

It is often necessary to achieve a timing resolution which is in excess of the sampling rate of the oscilloscope we're using to collect data. It is possible to interpolate between points exactly given that the signal being interpolated is bandlimited below the Nyquist frequency (that is, half the sample rate)[38]. Typically, data is taken with the oscilloscope at its maximum possible sample rate (10 GSa/s for 2 channels, 5GSa/s for 4 channels). The rated analog bandwidth of the scope is 4 GHz, so this condition is weakly satisfied

⁵Also known as the moving average filter

just by the design of the oscilloscope. Additionally, the signals of interest to us are pulses with a rise time of approximately 500 ps and a fall time of roughly 5 ps. If we calculate the Fourier transform of such a pulse, at 2.5GHz the power spectral density (PSD) is reduced by more than 45 dB compared to the PSD at DC. This suggests that the criteria of using a signal that is bandwidth limited to 2.5 GHz should be no issue. In this case, we can exactly reconstruct the original, continuous signal[38] using Formula 4.4.

$$V(t) = \sum_{n=-\infty}^{\infty} (V_n \text{Sinc } F_s t - n) \quad (4.4)$$

In practice, this procedure can be carried out in many different ways. Observing that Equation 4.4 is essentially a convolution of our sampled signal with a Sinc function, it could be implemented by following this procedure:

1. Create new data and time vectors, $x'(n')$ and $t'(n')$, which are upsampled by a factor F_{re} from the original data ($x(n)$ and $t(n)$)
2. Load the original data into this vector as $x'(n') = \sum_{m=0}^{m=n_{max}} x(n) \delta(t'(n') - t(n))^6$
3. Perform a discrete convolution in the time domain or take the Fourier transform of $x'(n')$ and multiply it by the Fourier transform of the sinc function then Fourier transform it back to the time domain

On the other hand, it could also be implemented by carrying out the sum in Equation 4.4 directly. This crude method is certainly less computationally efficient than performing the convolution in the time domain as described above; however, it has one clear

⁶Note that this will leave most of the entries as $x'(n') = 0$

advantage in portions of $x'(n')$ can be calculated separately whereas in the previous method the entire $x'(n')$ vector is generated at the same time. In practice, the only portion of the upsampled vector we are interested in is the small region near either the peak of a crosscorrelation function or near the threshold crossing. It turns out that naively calculating the sum in this case is much faster than the more ‘clever’ procedure.

4.5.2.2 Threshold Detection

The simpler method (both to understand and to implement) for time tagging pulses is called threshold detection. The implementation is very simple; for a given discrete time series, $x(n)$, find the first sample which exceeds the threshold. Then, if a higher resolution estimation is desired, upsample the signal near this detected threshold cross using the methods described in Section 4.5.2.1. Then, find the point in this new series, $x'(n')$ which exceeds the threshold. Because the waveforms measured will invariably contain noise, it is worthwhile to estimate the uncertainty in our measurement of the time of arrival of a pulse using this method. This issue was discussed in Section 2.6.1, with the uncertainty calculated in Equation 2.27.

4.5.2.3 Cross-Correlation Time Tagging

A more robust method of time tagging waveforms is to use crosscorrelation functions. For instance, to measure the relative delay τ_{RD} between two sampled waveforms, $x_1(n)$ and $x_2(n)$, the crosscorrelation function (C_{12}) is first calculated⁷:

⁷Technically, this is only an approximation of C_{12} unless the sum extends from $-\infty$ to $+\infty$

$$C_{12}(m) = \sum_{n=-n_{max}}^{n=n_{max}} x_1(n)^* x_2(n+m) \quad (4.5)$$

This crosscorrelation function can be thought of as a measure of the similarity of the two waveforms $x_1(n)$ and $x_2(n)$ as a function of m , which is a parameter which shifts $x_2(n)$ relative to $x_1(n)$. The maximum of crosscorrelation is therefore located where m makes the two functions the most similar, and is therefore an estimate of the delay between the original signals. This method is more robust than the threshold algorithm due to the fact that it utilizes the entire waveforms for making the estimation rather than just a tiny portion of the waveforms around the point where they cross the threshold. This means that the effect of noise is mitigated somewhat.

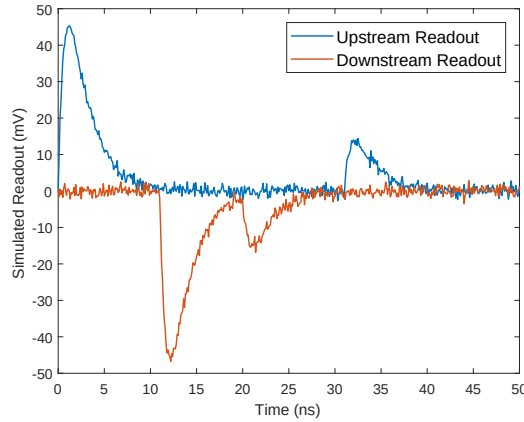


Figure 4.12: Example of Simulated Waveforms for Time Tagging Tests

Similar to the threshold algorithm, a coarse estimate of the location of the peak in C_{12} is made at the original sample rate of the data. Then, the crosscorrelation function is interpolated as described in Section 4.5.2.1 and a more refined estimate of the peak value is calculated. This process is shown in Figure 4.13.

Unlike the threshold algorithm, using crosscorrelation functions to time tag data does

not have a simple closed form equation for the uncertainty in estimating the parameter m (or equivalently, τ_{RD}). Therefore, we estimate this using a Monte Carlo simulation which generates pulse pairs with precisely known relative delays, adds random noise, and then calculates the estimated relative delay using the cross correlation method. This estimated delay can be compared to the actual delay, thus calculating the error in the method. This metric is plotted in Figure 4.14 as a function of the root mean square noise added to the signal and the rise time of the signal pulse. These simulations suggest that the uncertainty of this method is much lower than the uncertainty in the Threshold Algorithm.

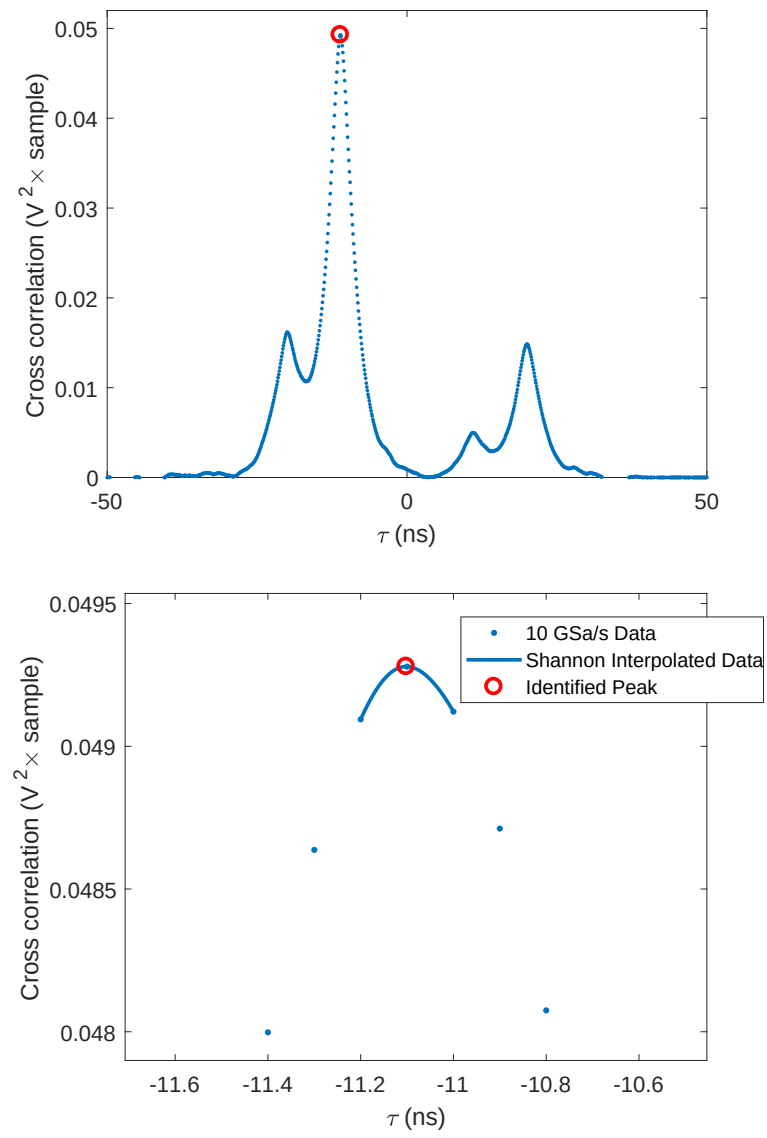


Figure 4.13: Example of Calculated Cross Correlation For Time Tagging

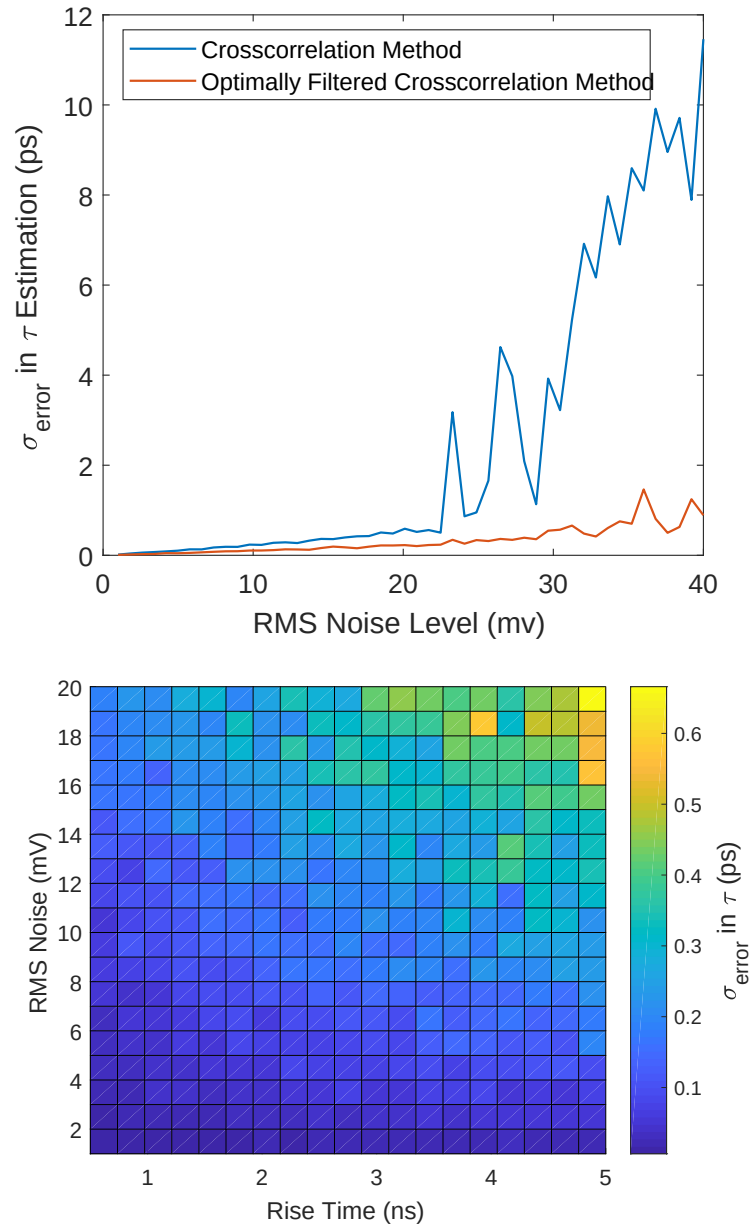


Figure 4.14: Comparison of Plain Cross Correlation Technique with Optimally Filtered Cross Correlation

Chapter 5

Superconducting Delay Line

Detector Results

5.1 Liquid Helium Dip Results

As previously described, devices are first screened by taking measurements in a dip probe in liquid helium. These measurements are performed at 4.2 Kelvin with all the electronics at room temperature. The first measurement taken is a current-voltage relationship (IV curve). After this, the device is characterized using Time Domain Reflectometry. The results of these measurements allow devices which will not work as ion detectors to be filtered out, saving valuable time. Measurements of both good and bad devices will be presented below with explanation of what kind of criteria we are looking at and why these criteria are important.

5.1.1 Current Voltage Measurement Results

Figure 5.1 shows an example IV curve for a “good” device which has an effective critical current density which is sufficiently high to detect ions in the several kiloelectron volt energy range. This curve is taken using the circuit in Figure 4.1. The voltage bias begins at zero, and slowly ramps up. At low bias voltages, the amount of current through the

detector is simply given by $I_{det} = V_{bias}/R_{bias}$. This portion of the curve is marked as a); this is the supercurrent branch where DC current flows with no resistance. At the point marked b) we reach J_{ceff} : this is our measure of what critical current density the detector can be biased with. Ideally, this number will be close to the Ginzburg-Landau critical current density[44] for our material/geometry (which in this case is approximately 50 MA/cm² for niobium devices). Above this critical current density, the curve diverts to line c) where a normal zone is nucleated at the point along the detector microstrip where the critical current density is suppressed to its lowest value. In this measurement there is no DC coupled shunt impedance, so all of the current produced by the bias circuit must be carried by the detector (except for at very short times when some power can be transferred to the RF port of the bias tee). This means that the electrothermal feedback discussed in Section 2.3.1.1 drives the evolution of the resistance of the detector. In our IV curves, we find that in most detectors the IV curve past J_{ceff} looks similar to d) in Figure 5.1. The current through the detector drops to a steady state value and does not change even as the bias voltage is varied. Because the resistivity of the detector is so high and the detector is so long (the total residual resistance of an SCDLD at 4.2 Kelvin is several megaohms) compared to the bias resistor, the bias voltage never reaches a value high enough to force the detector to carry more than this steady state current. The steady state current is set purely by the thermal parameters of the device; the steady state resistance is set by solving $\dot{R}_{hs} = 0$ where $\dot{R}_{hs}$ is given by Equation 2.14. The result is fairly intuitive; the steady state resistance of the wire is reached when the power dissipated through joule heating is equal to the power which can be dissipated through the substrate. This means that the current along this branch is the parameter j_p , the minimum required to maintain a stable hot spot, thus allowing us to make an

experimental measurement of α as described in Equation 2.11. Once the bias voltage is lowered such that the current through the detector drops below j_p , the hotspot heals, and the detector returns to the supercurrent state, a). The bias voltage then continues sweeping to negative values. Typically, this curve e) forms a mirror image of the positive curve, making the entire IV curve symmetric around zero.

In the case that there is a significant constriction in the SCDLD wire, the measured critical current density will be more strongly suppressed. However, the value of J_p (or equivalently, the retrapping current) remains unchanged unless the measured J_{ceff} is lower than J_p , in which case the curve will not be hysteretic, and then $J_{retrap} = J_{ceff}$ instead of being equal to J_p . An example of such a measurement is shown in Figure 5.2.

The measured critical current density is a critical measure of the detector's usefulness as an ion detector. As mentioned in Equation 2.3, there is a minimum bias current which needs to be flowing through the detector before it becomes sensitive to particles of a particular energy. If there is a weak point along the length of the detector which causes the measured critical current density to be below this minimum current, it will be impossible to bias the detector in a regime where it is sensitive to the particles of interest. It is challenging to calculate exactly what this minimum current condition is; experimentally, however, it has been found that ion detection can be performed with 750×40 nanometer detectors which have a measured critical current density of at least 10 MA/cm^2 . More in depth measurements of how the rate of ion detection varies with bias current will be described later which paint a slightly more complicated story.

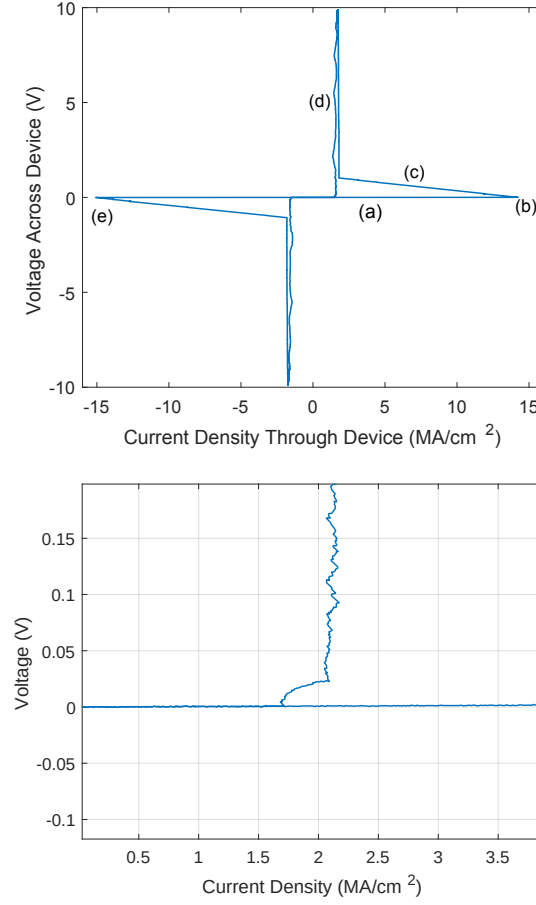


Figure 5.1: Current Voltage Relation of Weakly Constricted SCDLD: Measurement is taken on a device with a cross sectional area of 40 nm by 750 nm at 4.2 K. The bias resistor used is $25\ \Omega$ and the test current is passed through a bias tee with the RF port terminated with a calibrated $50\ \Omega$ termination. Bottom figure is a zoomed in view to show recovery back to superconducting state when bias current falls below J_p (in this case approximately $2\ \text{MA}/\text{cm}^2$).

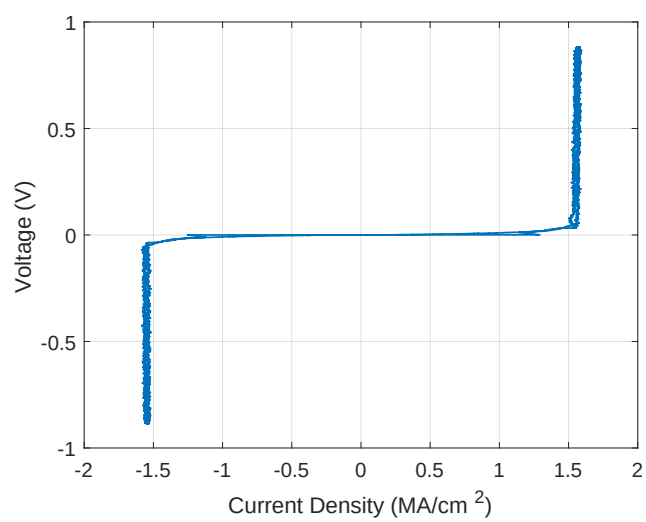


Figure 5.2: Current Voltage Relation of Weakly Constricted SCDLD

5.1.2 Time Domain Reflectometry Results

Measuring a high critical current density in a SCDLD is a necessary but not sufficient condition for the device to be usable as an ion detector. There are ways for a detector to fail while still having a high effective critical current density— for instance, the detector can be shorted to ground. Detectors with these more subtle defects can be detected by performing a Time Domain Reflectometry (TDR), and sometimes Time Domain Transmissivity (TDT) measurement. The setup for taking measurements like this is described in Section 4.3.

An example of TDR data taken from a good detector is shown in Figure 5.3. The different portions of the TDR are labeled as follows: (a) is the TDR response before the voltage edge is generated by the sampling oscilloscope. (b) The sampling scope transmits the initial voltage edge, which then propagates down the coaxial cabling (c) towards the sample, which is immersed in liquid helium. Right before the edge encounters the sample, there are several connections, including bondwires, which show up (d) as small excess series inductances. After (d), the remainder of the trace (e) is a measurement of the impedance of the detector itself. In a good sample, this will be nearly $50\,\Omega$, will be constant, and at the end of the sample (f), the impedance will drop to zero due to the fact that we have shorted the backside of the detector to ground. The length of (e) in time, ideally, is the electrical length/delay of the detector. There can be ambiguity in this final point; if a detector is shorted to ground within the die, it will look the same as the detector being shorted to the sample mount ground via bondwire. The only difference is the measured electrical length/delay (τ_{tot}) will be truncated. This can be ameliorated by having estimates ahead of time of the expected electrical length and

by taking measurements of multiple dies from the same wafer (short circuits faults are almost never correlated from die to die in this fabrication process, so the apparent length will vary randomly if shorts are an issue).

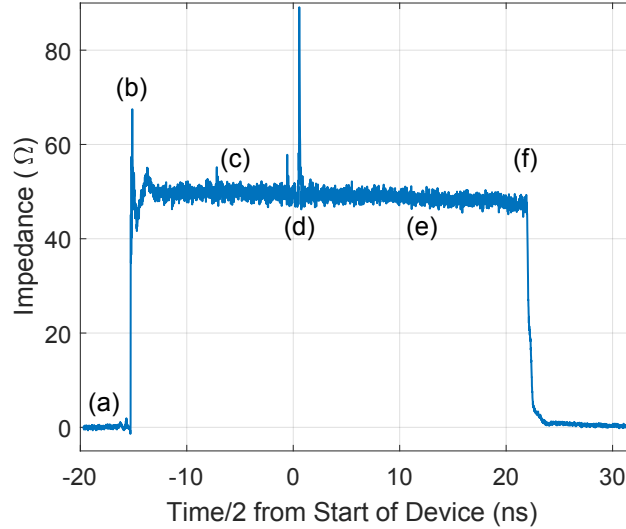


Figure 5.3: Time Domain Reflectometry of Good Detector

It is worth noting that both the characteristic impedance and the propagation velocity which are measured in these dip measurements will change slightly when the detector is operated on the 3K stage of our cryostat due to the change in the London penetration depth (Λ_L) with temperature, given by 2.9. This shift causes the impedance and electrical delay (τ_{tot}) to drop somewhat.

In addition to filtering bad devices out, TDR measurements also provide crucial information for debugging detectors which aren't useful. It is therefore useful to take TDR measurements even if the measured critical current density of a device is low. In Figure 5.5b several different examples of TDR curves taken from detectors with different types of defects are shown. In many cases, it is necessary to take measurement of a large

number of detectors from the same wafer in order to make the conclusions derived from these measurements more certain. This information then drives decisions we make about modifying the fabrication process for the next round of wafers.

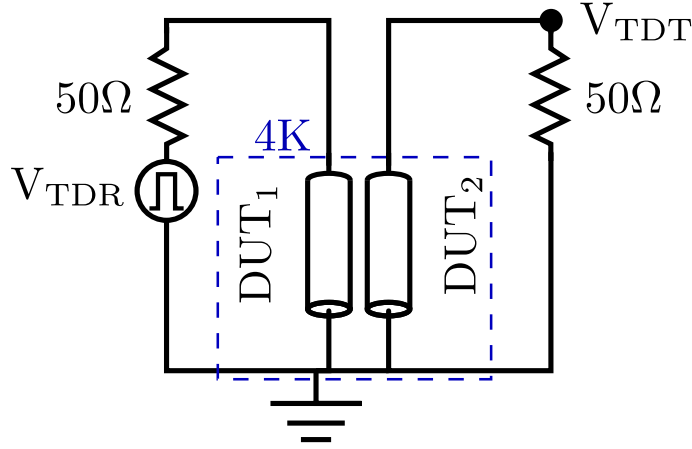
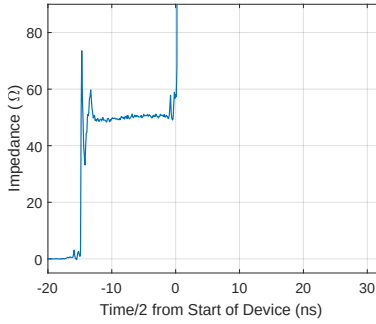


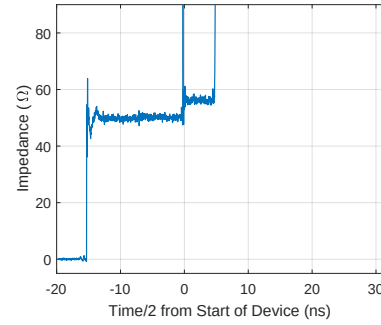
Figure 5.4: Time Domain Transmissometry Equivalent Circuit

The measurement presented in the bottom right of Figure 5.5b shows data taken from both TDR and TDT. In dies which contain two detectors (Pitch Experiment or Pitch Imager designs) it is sometimes observed that the impedance of a detector will drop suddenly, but not to zero before reaching the expected electrical length. This could be caused by a 'weak short' to ground which is not fully superconducting. Alternatively, it could be caused by a short between the two detector wires which are near each other. TDT is useful for distinguishing between these two situations. In this TDT measurement, a voltage edge is sent to DUT₁ and then the reflected voltage is measured as before. Now, though, DUT₂ (the second detector on the chip) is connected to a second 50Ω port on the sampling oscilloscope and the voltage which has been transmitted from DUT₁ into DUT₂ is measured. If there is indeed a superconducting short between the

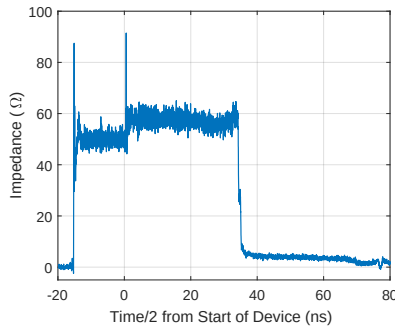
two detectors (and they are the same impedance), a transmitted pulse with a voltage of approximately half the initial TDR edge will be measured coincidentally with the apparent drop in impedance in DUT_1 . If instead the fault is a weak short to ground, there should be little to no transmitted signal coming out of DUT_2 .



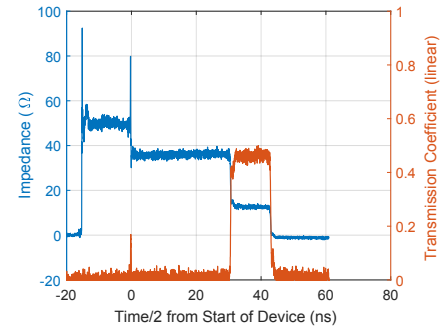
(a) Detector sample with step coverage issue.



(b) Detector with random defect at 4.94 nanoseconds (approximately 11% of the electrical length of the detector).



(c) Detector with excessive loss. This causes the final measured impedance to be greater than $0\ \Omega$.



(d) Pitch Imager sample where two detectors are shorted together at 30.5 nanoseconds.

Figure 5.5: TDR Measurements of various SCDLD Failure Modes: Times are offset to set the beginning of the detector at approximately zero. Curves are calibrated to set the impedance of semi-rigid coaxial cabling to $50\ \Omega$.

5.2 Single Wire Ion Detection Experimental Results

The rare, precious detectors that pass all the previously described screening measurements are then loaded into the cryostat to be tested via ion bombardment. In this section, the data collected from single wire detectors will be described in depth, starting with the actual waveforms coming from the detector. These waveforms will then be processed in order to build up histograms of ion flux and amplitude. Plots of the intermediate steps of this procedure will be presented and explained as well. Finally, information related to the experiment which is not captured by the waveform data alone will be discussed.

5.2.1 Time Domain Waveforms

Under typical conditions, the SCDLDs measured in this thesis generate pulses with amplitudes of order 10s of millivolts. In comparison, the dominant source of noise in our signal chain is the oscilloscope itself which is measured to be of order .5 mV (RMS)¹. Because this signal to noise ratio is quite high, it is not necessary to amplify the SCDLD pulses². Figure 5.6 shows an example of a pulse generated by the impact of an approximately 6 keV helium ion on a single detector SCDLD. This pulse was not amplified or filtered (except by the intrinsic filtering characteristics of the bias tee).

Some important characteristics of these pulses are as follows: typical rise times are

¹For comparison's sake, we expect the Johnson noise to be roughly 50 μ V, taking the room temperature to be 300K, the bandwidth of the oscilloscope to be 4 GHz, and a source resistance of 50 Ω .

²Occasionally solder joints on our coaxial cables will fail after repeated cooldowns yielding a weak connection and excessive loss. In cases like these, it is sometimes beneficial to take measurements anyways and use amplifiers to make up for this excess loss.

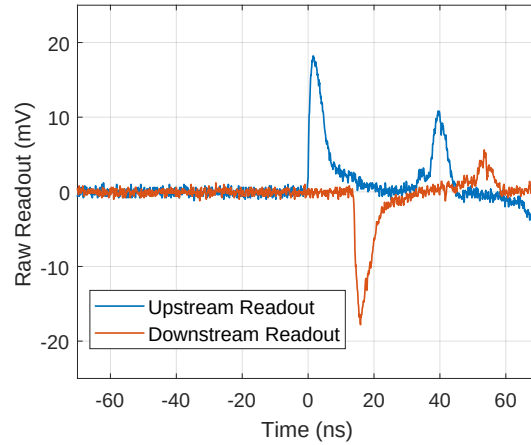


Figure 5.6: Pulse Waveform Example: Measured at 3.1 Kelvin, detector biased at 15.7 mA, $V_{FIM} = 6.6$ kV, event rate 71 Hz.

about 500 picoseconds, amplitude 20-30 mV, pulse width 2 nanoseconds, fall time 3 nanoseconds. The rise time is dominated by the thermal feedback dynamics described in Section 2.3.1.1; in fact, this measured rise time can be used to refine our estimate of v_0 in Equation 2.12. Although we treat this rising edge as an exponential, it is clear from the dynamics outlined in the electrothermal theory that the true functional form is somewhat more complicated, so there is no neat closed form solution for mapping from a rise time back to v_0 ; altering this parameter involves some manual adjustment. The fall time is misleading; it may seem like the detector is recovering, but, in actuality, the detector is remaining latched and the bias tee capacitor (which is in series with our readout/oscilloscope channel) is high pass filtering the signal. This causes the step-edge being emitted by the detector to look like a pulse with an exponential decay.

There are some non-ideal aspects of the signal generated by this detector; for instance, there are severe reflections in both channels and sometimes a low frequency (~ 10 MHz)

ringing artifact. By using the amplitude of the original pulse and its reflections we can estimate the impedance of the detector. Define the following terms: V_{p0} is the initial pulse amplitude inside the detector, V_{p1} is the amplitude of the initial measured pulse, V_{r1} is the amplitude of the first reflected pulse, and Γ is the reflection coefficient between the detector and the readout circuit. We make the assumption that the hotspot impedance ($\mathcal{O}(1000\Omega)$) is large compared to the detector impedance ($\mathcal{O}(50\Omega)$) so that the reflection coefficient off of it is very nearly 1. Then:

$$V_{r1} = \Gamma(1 - \Gamma)V_{p0} \quad (5.1)$$

$$= \frac{Z_0 - Z_{det}}{Z_0 + Z_{det}} \frac{2Z_{det}}{Z_0 + Z_{det}} V_{p0} \quad (5.2)$$

Noting that $V_{p1} = (1 - \Gamma)V_{p0} \dots$

$$V_{r1} = \frac{Z_0 - Z_{det}}{Z_0 + Z_{det}} V_{p1} \quad (5.3)$$

$$Z_{det} = Z_0 \frac{1 - V_{r1}/V_{p1}}{1 + V_{r1}/V_{p1}} \quad (5.4)$$

In the pulses measured in this experiment, this yields an apparent detector impedance of approximately 12Ω . This is much lower than the design impedance and the impedance measured in TDR (both of which were 50Ω). It is unclear where this discrepancy comes from³; however, it is possible to take this measured impedance and insert a resistive matching network between the detector output and the bias tee which suppresses these pulses significantly as shown in Figure 5.7. This comes at the expense of some attenuation which is necessary to achieve impedance matching in this way. The benefit of this

³It is possible that this is due to the fact that our readout port is only AC coupled, so the apparent readout impedance effectively grows in time, worsening the apparent impedance mismatch

method of impedance matching is that it is inherently broadband and is much easier to implement than reactive or transmission line based transformers like the Klopfenstein taper[34][55].

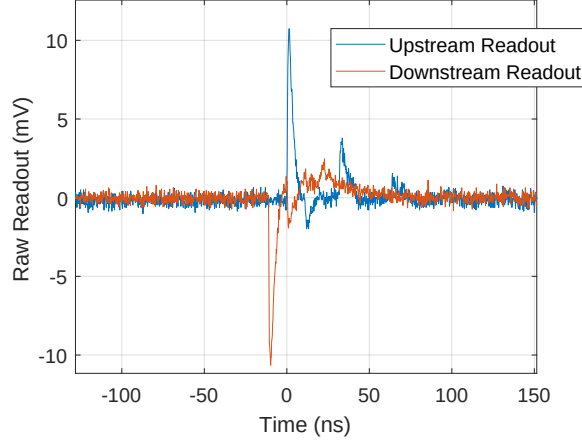


Figure 5.7: Pulse Waveform With Resistive Matching Network Inserted

These reflections somewhat degrade the performance of measuring the relative delay using only the crosscorrelation function between the upstream and downstream readouts by generating extra peaks in the crosscorrelation plot. By constraining our peak finding to only work within the range of expected relative delays ($-\tau_{tot} \leq \tau_{RD} \leq \tau_{tot}$) these issues can be mitigated.

Although our signal to noise ratio is reasonably high, we can improve our τ_{RD} estimation by filtering the waveforms as previously described in Section 4.5.1. Figure 5.8 illustrates the impulse response of the filter used as well as the pulse from Figure 5.6 after applying this filter. The crosscorrelation function is then calculated between the two filtered pulses as shown in Figure 5.9. The coarse peak value is found by finding

the maximum value of the “coarse” crosscorrelation function (which has the same effective sample rate as the original signals). This function is then interpolated between the peak sample and its two neighbors using the Shannon interpolation described in Section 4.5.2.1. The peak value of this new interpolated function is then calculated and is used as our “fine” measure of τ_{RD} .

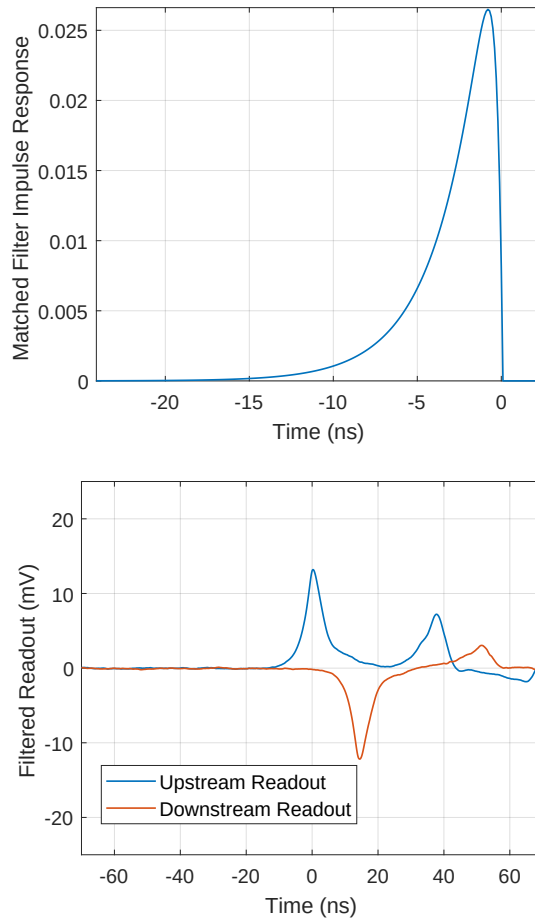


Figure 5.8: Waveform from Figure 5.6 filtered using a matched filter shown in the top plot. The matched filter is most intuitively visualized in the time domain. Its impulse response is illustrated above; the parameters of the matched filter impulse response are as follows: rise time=500 ps, fall time = 3 nanoseconds, pulse width = 2 nanoseconds, normalized to have an integrated area of 1.

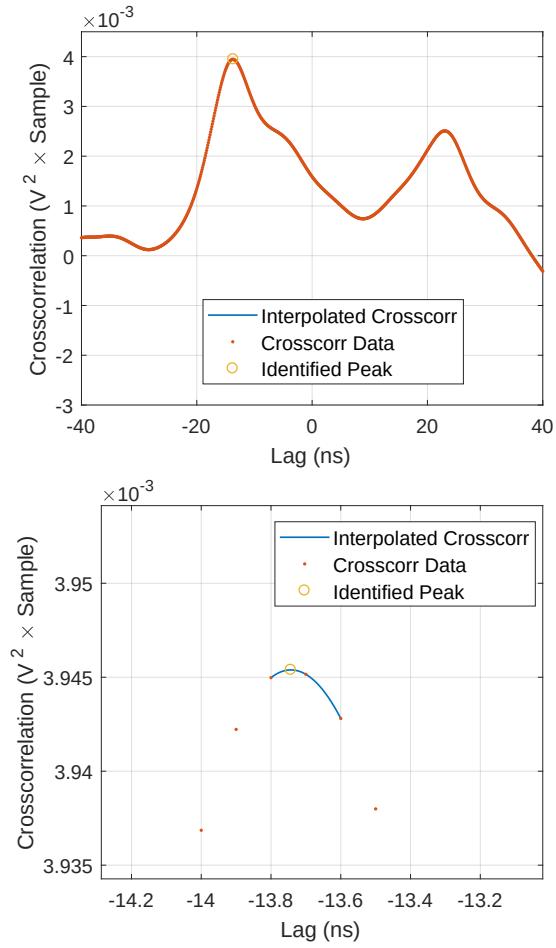


Figure 5.9: Crosscorrelation function for the filtered pulses measured in Figure 5.8. Re-sampled curve is calculated using the Shannon-Whitaker algorithm described in 4.5.2.1.

5.2.2 Relative Delay Histogram

Once the raw waveforms have been processed, histograms of various kinds are generated. First, we bin pulse pairs based on their measured τ_{RD} . This parameter is a measure of the position along the detector where an ion generated an event, so by generating a histogram of this data it is possible to see a folded, 1D image of the ion flux the detector was exposed to. In Figure 5.10 an example of such an image is shown. Zoomed out, the plot clearly reflects the non-uniform ion flux that is expected from the FIM tip. Taking a closer look, it becomes apparent that there are periodic structures in the data. If a Fourier transform is taken of this type of histogram, the dominant peak is always given by $1/(2\tau_{meander})$ (where $\tau_{meander}$ is the time it takes an electrical signal to traverse a single meander) or a harmonic thereof. This is because the image generated by the FIM tip is larger than the meander pitch so that bright places in the image overlap with more than one meander of the detector. This justifies our assumption in the discussion of the Pitch Experiment (Section 2.5.1) that the ion flux does not vary over distances smaller than the pitch between detectors.

5.2.3 Pulse Amplitude Distribution

When processing pulses, information is also stored about the maximum value attained by the pulse in each waveform. This is labeled as the amplitude of the pulse. When using threshold detection for measuring the timing of pulses, it is necessary to have a very narrow distribution of pulse amplitude, as discussed in Section 4.5.2.2. Although this is less important in using the cross correlation method, having a narrow pulse distribution is useful in the case of trying to examine pulses under non-ideal conditions. For instance,

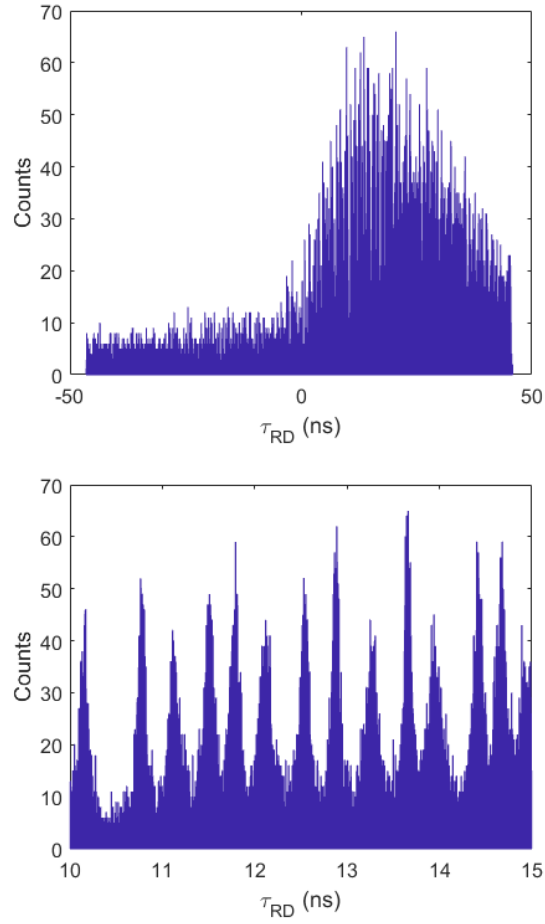


Figure 5.10: SCDLD Relative Delay Histogram: Measured at 3.1 Kelvin, detector biased at 24 mA, $V_{FIM} = 6.6$ kV, event rate 71 Hz.

if two detectors are operating near one another such that there is electrical crosstalk, having narrow amplitude distributions on both can be helpful in filtering out crosstalk signals from real ones.

As shown in Figure 5.11, SCDLDs exhibit a very narrow pulse distribution, having a variance which is smaller than 1/30 of the mean. This is a large improvement over microchannel plates which have a variance which is approximately equal to the mean.

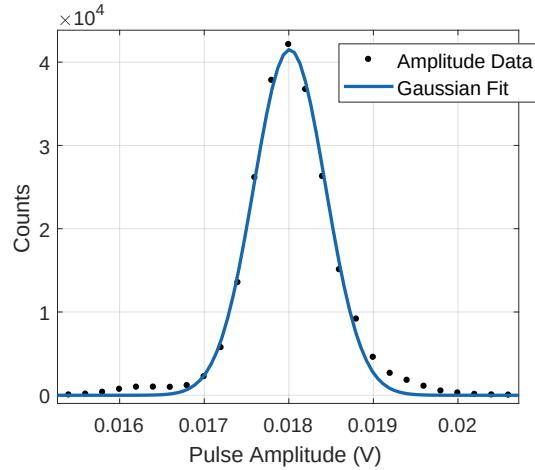


Figure 5.11: Pulse Amplitude Distribution: Fit with gaussian having a mean of 18.01 mV, variance of .42 mV. Measured at 3.1 Kelvin, detector biased at 24 mA, $V_{FIM}=6.6$ kV, event rate 71 Hz.

5.2.4 Event Rates

Each collected waveform only contains a single ion detection event. In order to keep track of the event rate, the measurement code which controls the oscilloscope during an experiment keeps records of how long it takes to fill each FastFrames records. These records typically consist of 100, 1000, or 10000 events depending on the measured event rate and how often the user wants to see feedback on the oscilloscope.

The variation of measured event rate as a function of the bias current through the detector is shown in Figure 5.12. In nanowire based detectors like these, event rates typically follow a pseudo-sigmoidal behavior while varying bias current. At currents below some threshold, the event rate is essentially zero. Around this threshold current, the event rate varies quickly with changing bias current. Well above this threshold current, the event rate, ideally, should saturate to some value. It has been suggested

by multiple authors[35][4][49] that the saturation of this event rate suggests that the internal detection efficiency has also saturated to some “high” value (hopefully near 1). In our SCDLD measurements, we also observe a saturation of the event rate. We are unable to directly measure the detection efficiency of our detectors in the setup currently being described. So, although this is a promising feature in the data, it by no means conclusively shows that the internal detection efficiency of our detector is close to 1.

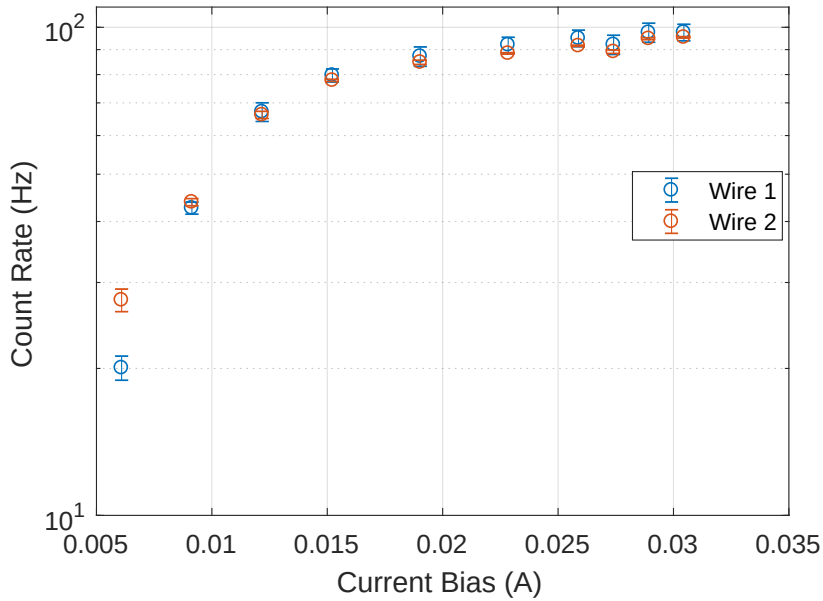


Figure 5.12: Event Rate vs Detector Current Bias: Measured on Pitch Experiment sample with wires biased separately. Error bars are set at 1 standard deviation. The maximum bias current that could be produced in this particular experiment was ≈ 31 mA. Over this bias current range, no dark counts were observed over approximately 1 hour worth of observation time.

The variation of event rate with FIM voltage is somewhat similar to the variation with detector current bias. Again, there is a threshold behavior where below a certain voltage

bias essentially no ions are generated, then near the threshold voltage, the ion current increases very quickly. Well above this threshold, the current continues to increase, but more slowly[40]⁴. This is discussed (and illustrated) in Section 4.4.3

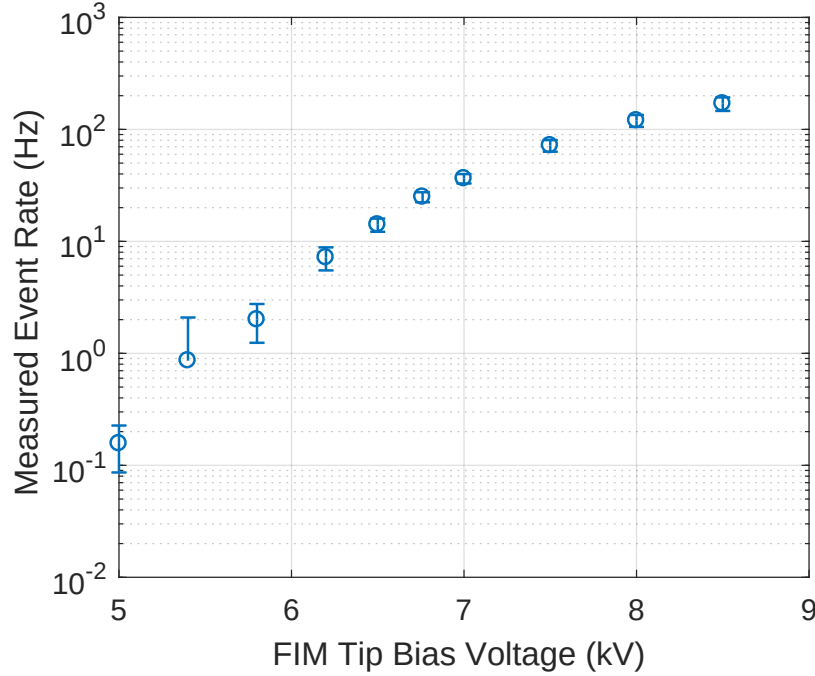


Figure 5.13: Event Rate vs FIM Voltage Bias: Measured on Pitch Experiment sample with only one wire biased. Error bars are set at 1 standard deviation. At 4kV, 3kV, 2kV, and 1kV no counts were measured over the course of a ten minute observation window for each. At 0 kV, no counts were detected over the course of a 60 minute observation window.

The FIM event rate measured with our SCDLD is shown in Figure 5.13. Qualitatively, it has the expected behavior.

⁴Unless the bias voltage reaches the breakdown voltage for the environment in the sample chamber, in which case this relationship breaks down dramatically, as does the FIM tip and possibly your detector

5.2.5 Relative Delay Resolution Results

Using the same detector described in the previous section, it is possible to measure the effective resolution achievable when calculating the relative delay (τ_{RD}) of a pulse pair. The measurement setup is described in Section 4.5.2. The same procedures described in 5.2 to calculate the relative delay for each pair of pulses are carried out. Those relative delay values are then binned into histograms. In Figure 5.14 shows the result of taking this one dimensional relative delay histogram and folding it into a 2D image. This is done by calculating the meander transit time, $\tau_{meander}$, and then mapping τ_{RD} to position (X_{ion}, Y_{ion}) as follows:

$$Y_{ion} = \text{floor}\left(\frac{\tau_{RD}}{2\tau_{meander}}\right)h_{det} \quad (5.5)$$

$$\zeta_{ion} = \frac{\tau_{RD} \bmod 2\tau_{meander}}{2\tau_{meander}} \quad (5.6)$$

$$X_{ion} = \left(\frac{1}{2}(1 + (-1)^{Y_{ion}}) + \zeta_{ion}(-1)^{Y_{ion}+1}\right)w_{det} \quad (5.7)$$

Where h_{det} and w_{det} are the overall height and width of the detector, respectively. This mapping function takes into account the meandered shape of the detector. Previous dip measurements of this detector provide a first estimate for v_{prop} and therefore $\tau_{meander}$. This estimate can then be improved by calculating the Fourier transform of the 1D relative delay histogram. The top plot of Figure 5.14 shows the result of following this procedure. Near the top of this image, there is a clear issue with this mapping. This defect where each row of the image splits farther and farther apart is a symptom that the estimate $\tau_{meander}$ is slightly off. The fact that this defect occurs in only one portion of the detector suggests that the propagation velocity of the detector is not entirely

uniform. This image can be rendered more accurately by allowing the top section to have a slightly $\tau_{meander}$ which is calculated separately from the rest of the detector. This change is reflected in the bottom plot of Figure 5.14.

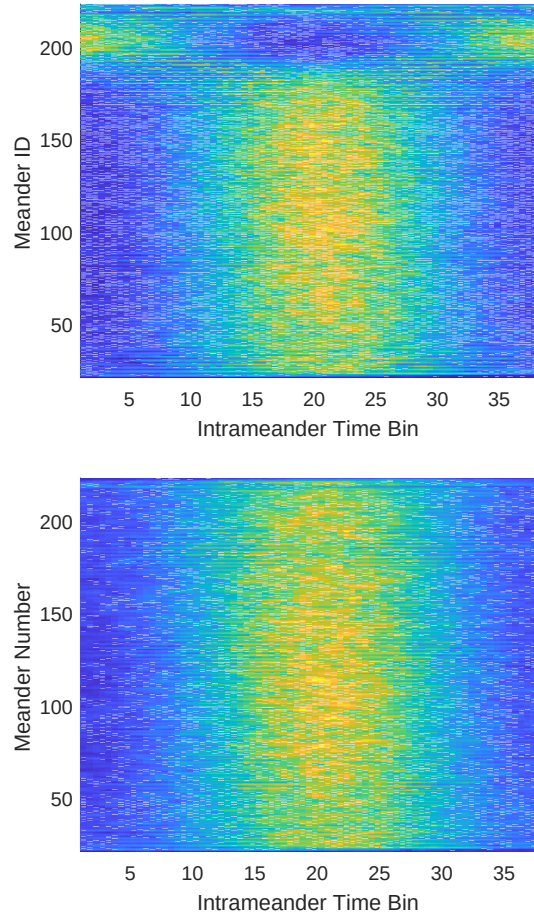


Figure 5.14: SCDLD Relative Delay Resolution Image: Measured on a 22ns long detector using ^{241}Am as the source of ions. Event rate across the detector was approximately 1.5 Hz, data collected over approximately 72 hours. The top image was formed using the average propagation velocity for the entire detector. The bottom image was formed by allowing there to be a different propagation velocity in the top $\sim 20\%$ of the detector.

For the sake of simplicity, calculations of the uncertainty in τ_{RD} will be performed in the one dimensional relative delay histogram. Using the measured value of $\tau_{meander}$ and knowledge of the physical shape and size of the shadow mask (illustrated in Figure 4.9 and described in Section 4.5.2), the ideal relative delay histogram can be calculated.

The model for this histogram is a pair of square waves added together with the duty cycle hard-coded as the width of illuminated region divided by the total width of the detector. The phases of the two square waves are allowed to vary (which allows for the case that the illumination bin is not exactly dead center on the detector), the frequency of the square waves is allowed to vary (but should end up being given by $1/2\tau_{meander}$), and an overall scale factor is a free parameter. This ideal histogram is then convolved with a Gaussian function which is normalized to have an area of 1. The variance of the Gaussian is the only free parameter for this function. This convolved histogram is then fit to the measured data, and the variance calculated for the Gaussian is reported as the uncertainty in relative delay. The results of this procedure are shown in Figure 5.15.

The calculated uncertainty is somewhat higher than expected compared to calculations of how precisely the measured waveforms can be time-tagged. Even in using the less than optimal algorithm of using threshold detection on interpolated data the calculated σ_e (the uncertainty purely due to electrical noise) is approximately 10 picoseconds for this experiment. This type of uncertainty is expected to be completely uncorrelated between measurements of the upstream and downstream ports so if this was the dominant uncertainty, the total relative delay uncertainty would be roughly 14 picoseconds⁵. There must therefore be some other source of uncertainty. In Section 2.6 different archetypal sources of uncertainty are discussed (σ_e is most similar to the

⁵ $10 \times \sqrt{2}$ picoseconds

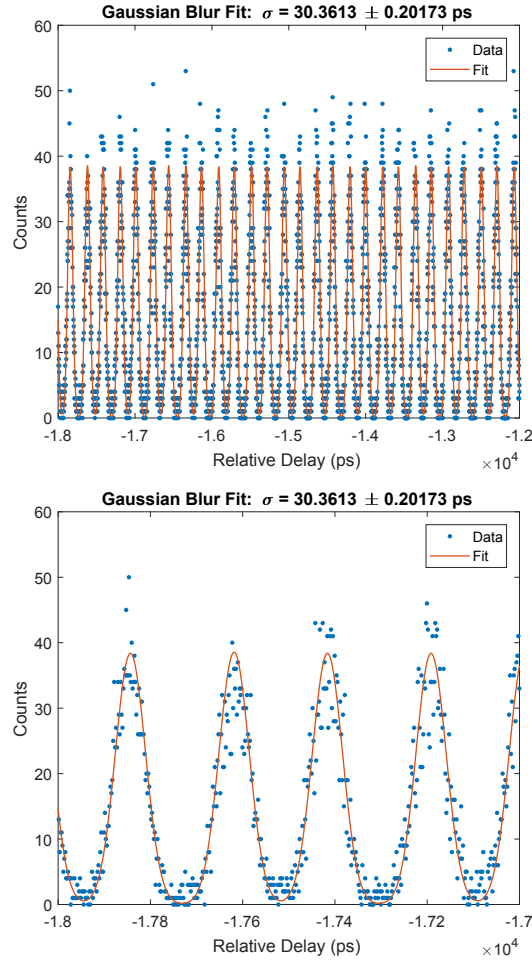


Figure 5.15: SCDLD Relative Delay Resolution Fits: Measured on a 22ns long detector using ^{241}Am as the source of ions. Event rate across the detector was approximately 1.5 Hz, data collected over approximately 72 hours.

velocity noise described in Section 2.6.3). The measurements presented in Figure 5.15 indicate that the uncertainty in τ_{RD} does not vary with position, which suggests that if this uncertainty is due to random fluctuations in the kinetic inductance of the detector, they must be randomly varying in both space and time.

5.3 Pitch Experiment Experimental Results

In order to make some assessments of the imaging capability and the detection efficiency of the SCDLD, the Pitch Experiment was devised. The theoretical underpinnings of the Pitch Experiment are outlined in Section 2.5.1 and the measurement setup is outlined in Section 4.4.5. The results of these efforts are presented in this section.

As an initial check that everything is working properly, the two detectors are biased at separate times (so that they are never on at the same time). In this state, they should operate just as the single detector dies discussed previously, and if they are indeed generating images of the ion emission coming from the FIM tip, they should produce similar relative delay histograms. These histograms are shown in Figure 5.16. They appear to be qualitatively similar.

After confirming that both detectors are operating normally, the experiment is modified so that the pulsed bias currents going into the detectors are synchronized. Pulses are collected as normal, and then processed as previously described in Section 5.2. Our interest lies in the ratio of correlated to uncorrelated events which should allow us to determine the detection efficiency properties of the SCDLDs under test as described in Section 2.5.1. In this experiment, pulses which appeared to be correlated were observed—that is, there were pulses recorded where both wires fired at nearly the same time (within the same ≈ 100 nanosecond window). In Figure 5.17, these correlated events are plotted as τ_{RD1} versus τ_{RD2} . If the events are truly correlated and not just coincidences where two separate events happened in the same time window, they should come from the same location on the two detectors; that is, they should have the same relative delay value. In this plot, we observe that nearly every correlated event pair lies on the line

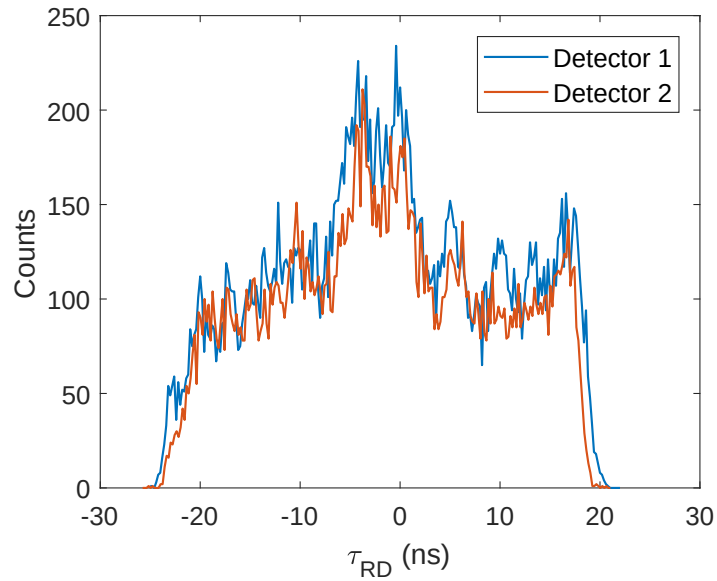


Figure 5.16: Separately Biased Pitch Experiment Relative Delay Histograms: Data was taken with a Pitch Experiment sample with the following geometry: detector thickness = 40 nm, detector width = 750 nm, minimum detector pitch = 1500 nm, maximum detector pitch = 4000 nm. Detectors were biased at separate times with a bias current of 6.1 milliamps and the FIM voltage bias was 10 kV.

of $\tau_{RD1} = \tau_{RD2}$ ⁶, which suggests that these events are indeed correlated. Although this data is plotted as relative delays, it is important to remember that in the Pitch Experiment relative delay values map linearly to the pitch between the two detectors. At $\tau_{RD1} \approx 21$ nanoseconds the pitch between the detectors (measured from center to center) is $1.5 \mu\text{m}$ and $\tau_{RD1} \approx -23$ nanoseconds, the pitch is approximately $4 \mu\text{m}$. The density of correlated events clearly increases in the portion of the Pitch Experiment where the two detectors are closely spaced.

⁶There is a small offset between the two curves due to unequal cable lengths in the readout line

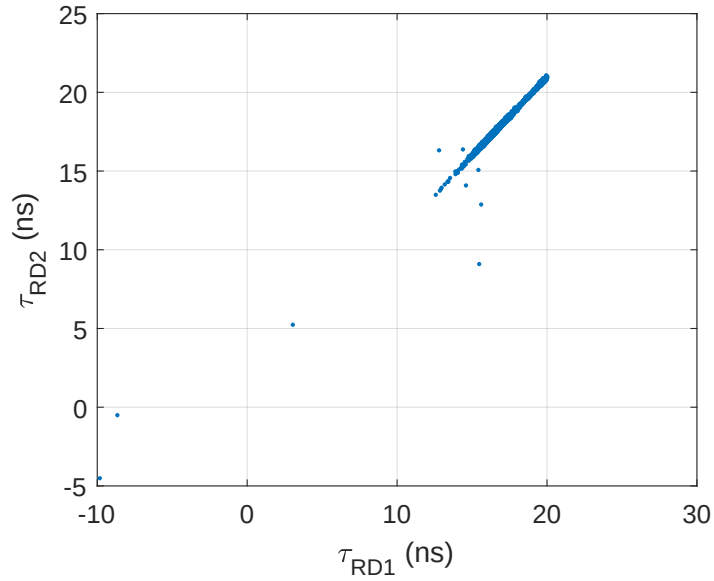


Figure 5.17: Pitch Experiment Relative Delay Correlations: Taken in the same conditions as Figure 5.16 except the bias current is active in both detectors at the same time and is increased to 15.25 mA.

In Figure 5.18 the histograms of relative delays from each detector (denoted as “Wire 1” and “Wire 2”) are plotted. Even in the regions where correlated events are not observed, the image generated by each detector is very closely matched.

In Figure 5.19 the histogram of relative delays from Wire 1 is plotted again. Now, however, it is plotted with respect to the pitch between the detectors, and the second histogram plotted is also Wire 1 events; however, these are events which had correlated Wire 2 events. Dividing the first curve by the second yields Figure 5.20. This plot is the measurement of the Correlation Fraction which was described in Section 2.5.1 which should allow us to estimate the SCDLD’s detection efficiency parameters. However, the functional form of the measured Correlation Fraction does not match our calculated

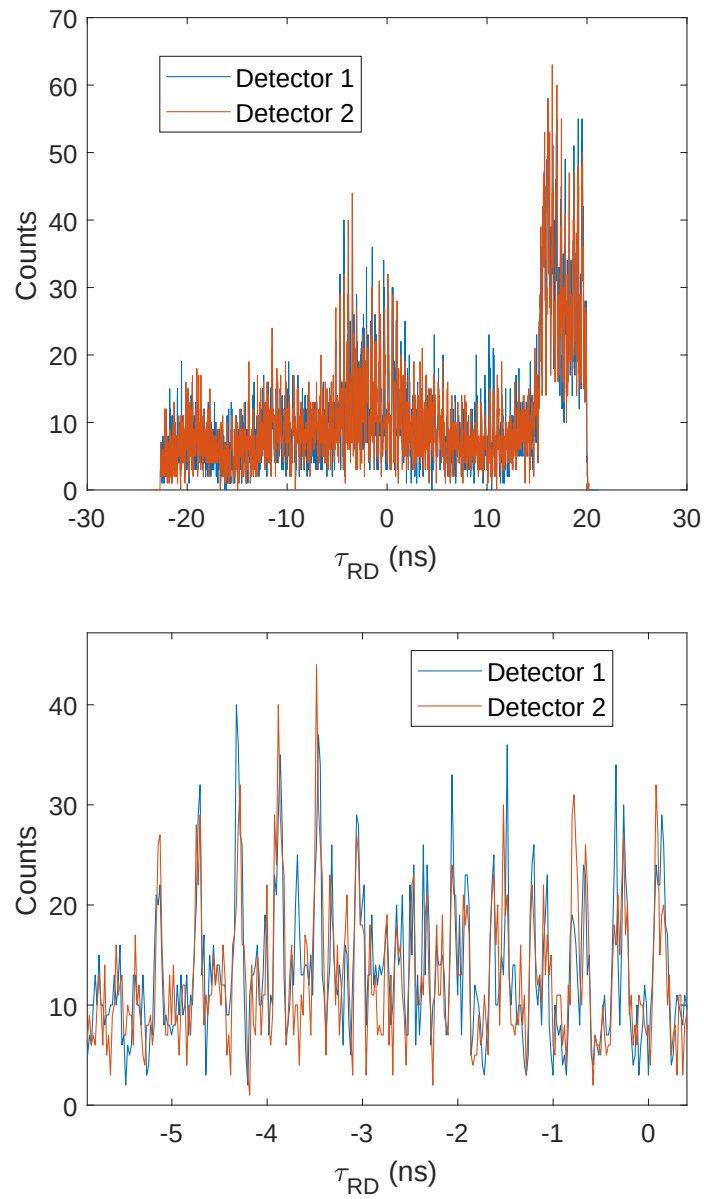


Figure 5.18: Pitch Experiment Relative Delay Histograms: Taken under the same conditions as Figure 5.17

curve from Section 2.5.1. The value should not saturate, and the peak value of the Correlation Fraction should never reach 1. It is therefore apparent that our criteria for

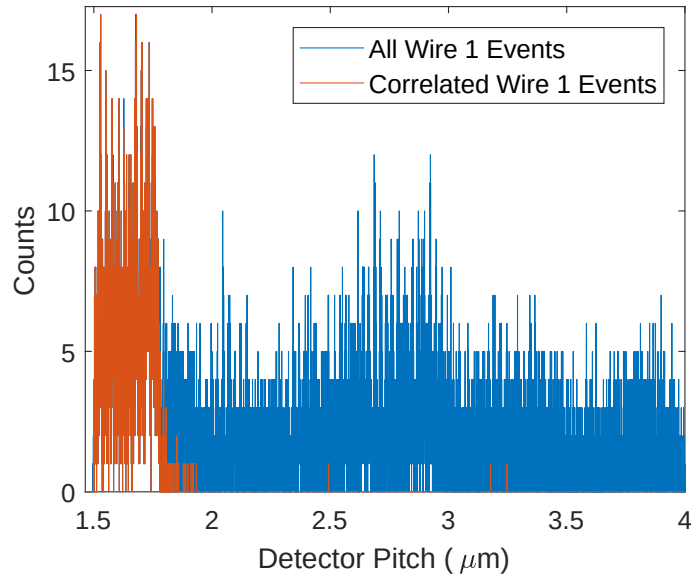


Figure 5.19: Pitch Experiment Correlated vs Uncorrelated Event Histogram: Taken under the same conditions as Figure 5.17

labeling events as ‘correlated’ is flawed.

Let us define a new time scale, τ_l , which will be called the ‘Event Lag’. This is a measure of the time delay between an event occurring on wire one and a subsequent occurring on wire 2. This can be calculated from our waveforms by taking the average time of arrival for the upstream and downstream pulses for each wire, and then taking the difference between those averages. In the case of two wires being triggered by a single ion, this lag parameter should be limited to essentially $\frac{\text{pitch-width}}{v_{\text{sound}}}$ where v_{sound} is the speed of sound in the silicon nitride encapsulation time. This number is of order 100 picoseconds. In Figure 5.21, this lag parameter is plotted as a function of detector pitch. It is immediately apparent that the original criteria for correlation was too loose, and something is causing a large number of psuedo-correlated events to be generated.

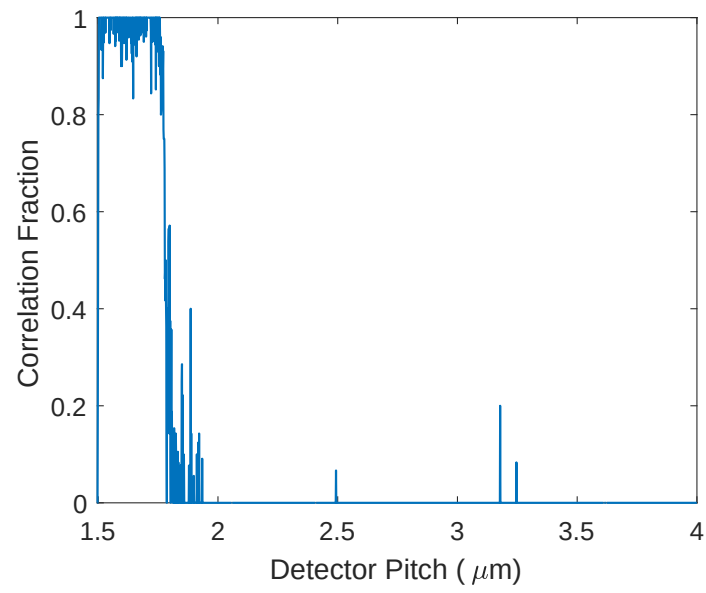


Figure 5.20: Correlation Fraction vs Pitch: Taken under the same conditions as Figure 5.17

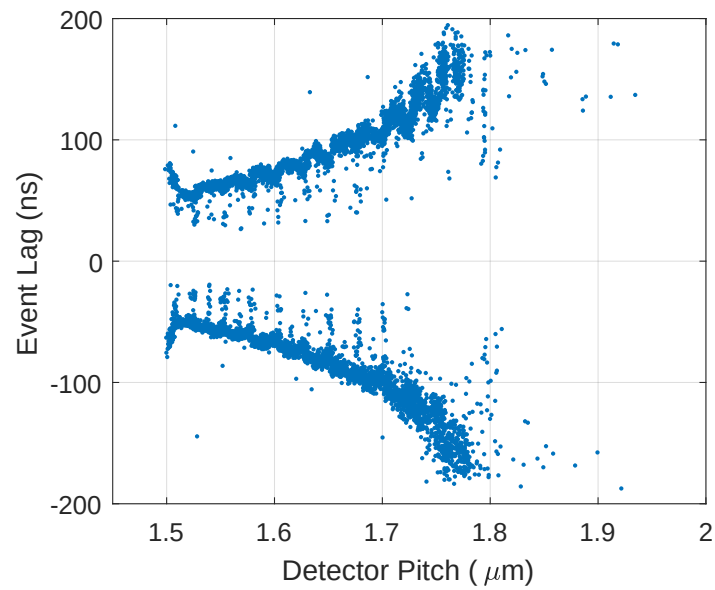


Figure 5.21: Event Lag vs Pitch: Taken under the same conditions as Figure 5.17

Due to the fact that this thesis is not written exactly in chronological order, the source of these false correlated events has already been strongly hinted at in Section 2.5.2. Recall from Section 2.3.1.1 that the SCDLDs measured are expected to latch into the normal state. Therefore, whichever wire is triggered by an ion first will form a stable hotspot, after which any remaining bias current flowing through this detector will dissipate heat. This heat can then increase the temperature of the closely spaced second detector, causing it to also switch into the normal state and generate a psuedo-correlated event. A model was developed in 2.5.2 to predict how the event lag should vary with detector pitch. In Figure 5.22 our measured Event Lag vs Pitch data is fit to this model.

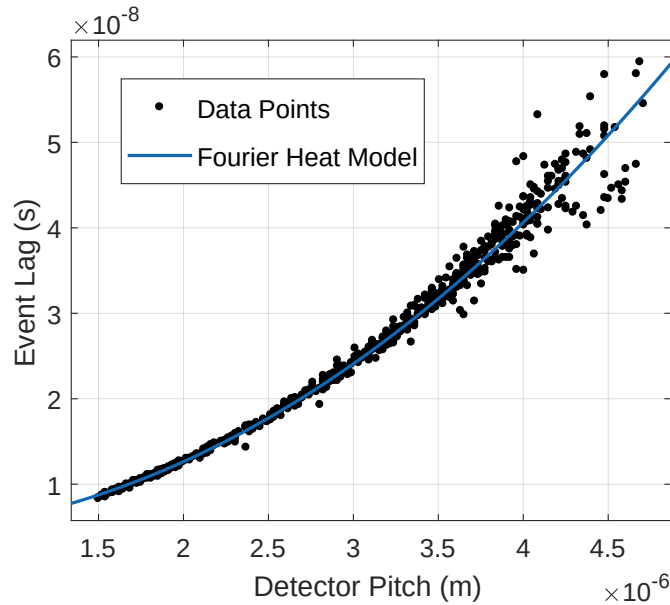


Figure 5.22: Event Lag vs Pitch Model Fit: Taken under the same conditions as Figure 5.17 except bias current increased to 32.9 mA

If the criteria for true correlated events is tightened to eliminate these psuedocorrelated events, it has been found for all Pitch Experiments measured thus far that the

correlation fraction is unmeasurable. The smallest pitch measured thus far has been $1\text{ }\mu\text{m}$. This puts an upper-bound on w_{eff} of $1\text{ }\mu\text{m}$. Unfortunately, this does not allow to make any statements about the internal detection efficiency of the detectors.

5.4 Pitch Imager Experimental Results

In Section 5.2.5 the uncertainty in our measurements of relative delay was found to be approximately 30 picoseconds. Because the relative delay for a pulse emanating from the extreme left side of the meander only differs from a pulse emanating from the extreme right side of the meander by approximately 208 picoseconds in a typical SCDLD, the effective position resolution in the horizontal direction (along the length of a meander) is limited to only approximately 7 pixels. The vertical direction (transverse to the meanders), on the other hand, has as many pixels as it had meanders (which in a typical device is several hundred). It is therefore very difficult to make a useful image of anything other than very simple patterns (such as the resolution measurement experiment image). However, the lessons learned from the Pitch Experiment with respect to thermally driven psuedo-correlated events can be used to enhance our horizontal imaging resolution. The strategy for doing so is outlined in Section 2.5.2.

In Pitch Imager experiments, data is collected in the same way as in the Pitch Experiment runs. The difference arises in the analysis of this data; now the psuedo-correlated events are useful and can be used to measure the horizontal position of an ion event. Because the scale of pitch variations is now smaller than our positional resolution, the transfer function of Event Lag to Pitch cannot be measured in a Pitch Imager device. It is therefore necessary to calculate an approximate transfer function which allows the

measured event lag to be mapped back to horizontal position. The model developed in Section 2.5.2 and the parameters extracted in the fit calculated for Figure 5.22 form the starting point for calculating the transfer function. The minimum and maximum pitch is known from the geometry of the detector, and the minimum and maximum lag can be measured. Small adjustments are made to the starting transfer function to fit these two end points. An example of an approximate transfer function that was calculated for a Pitch Imager is shown in Figure 5.23.

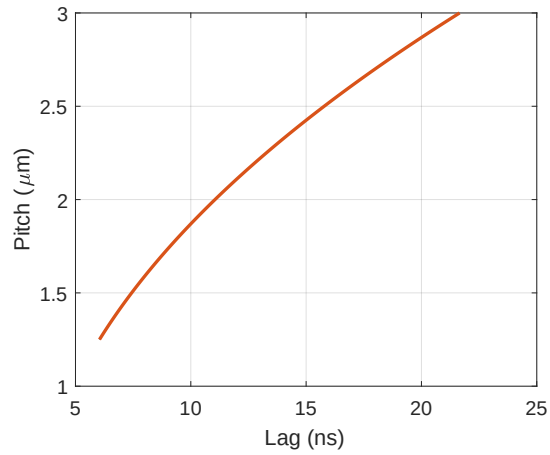
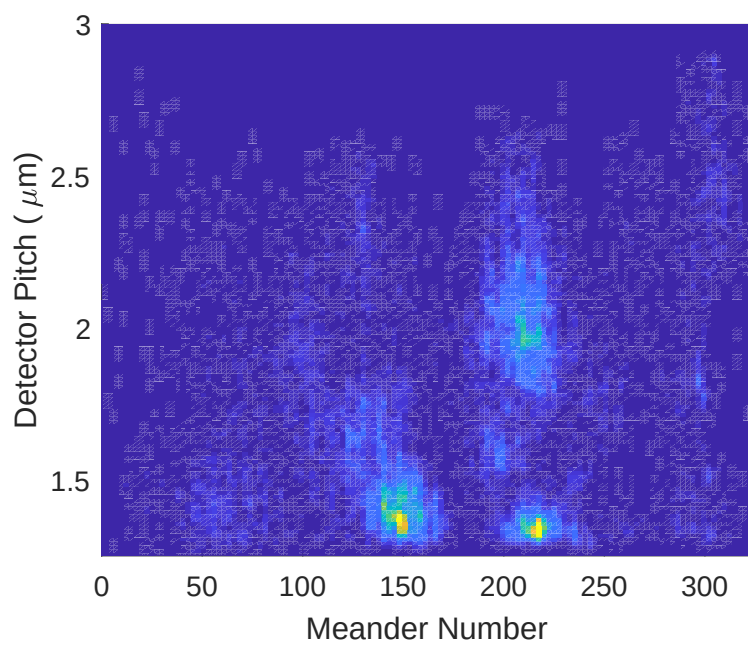
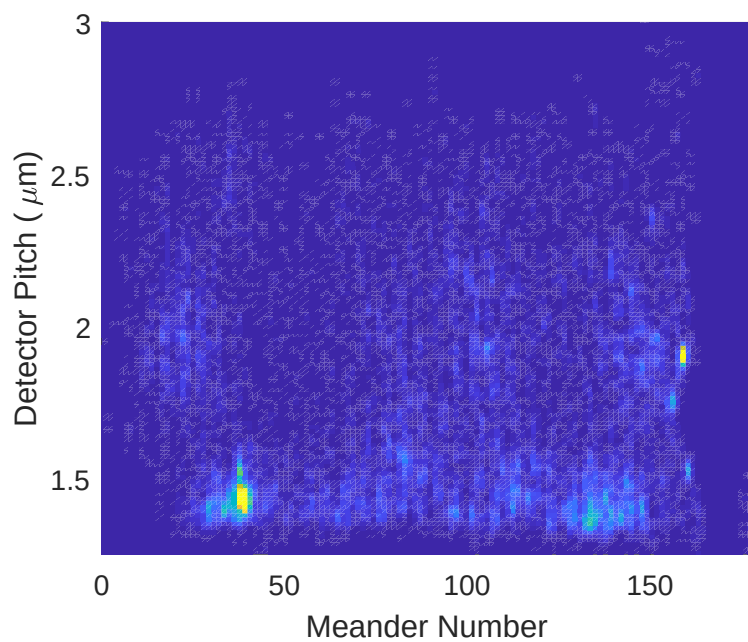


Figure 5.23: Pitch Imager Approximate Transfer Function

Once this transfer function is known, it can be used to map out images of the ion flux detected by the Pitch Imager with greatly improved horizontal resolution. Below is a series of images generated with different tips and at different FIM bias voltages.

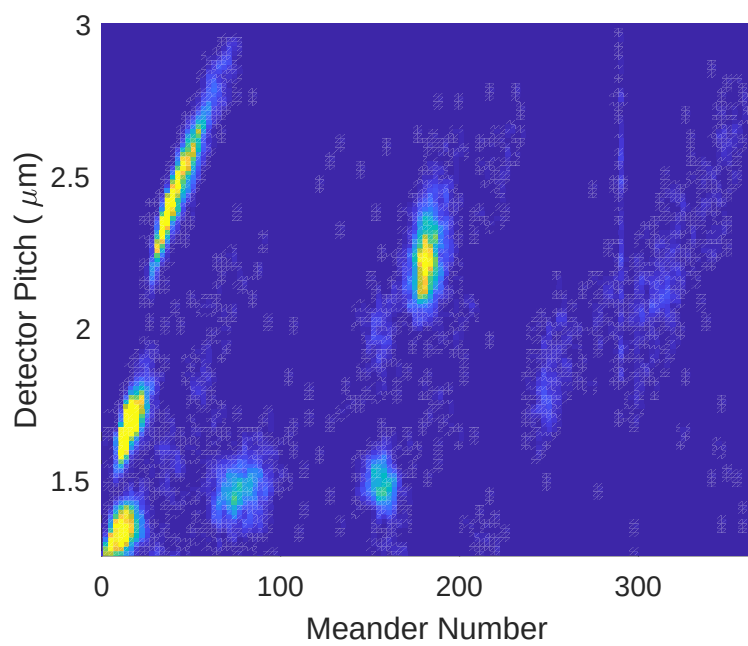


(a) 8 kV FIM Bias

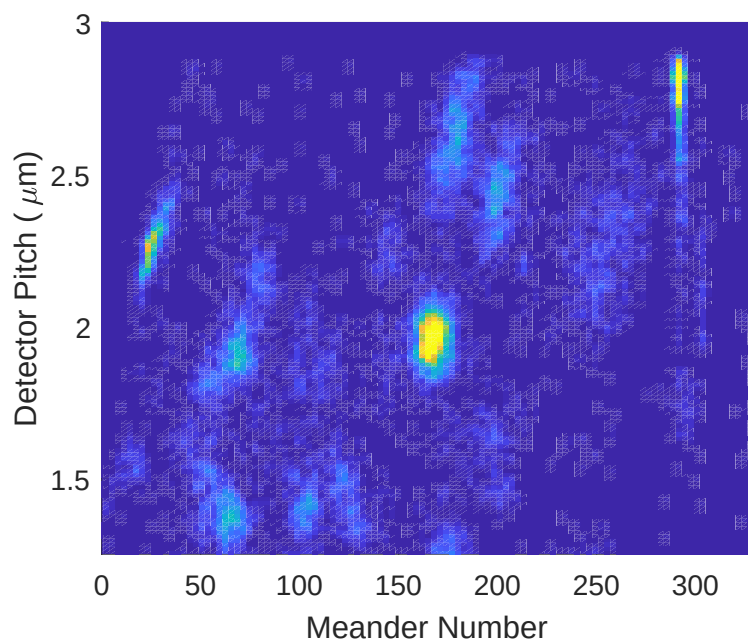


(b) 9.2 kV FIM Bias

Figure 5.24: Pitch Imager FIM Tip Images, Set 1

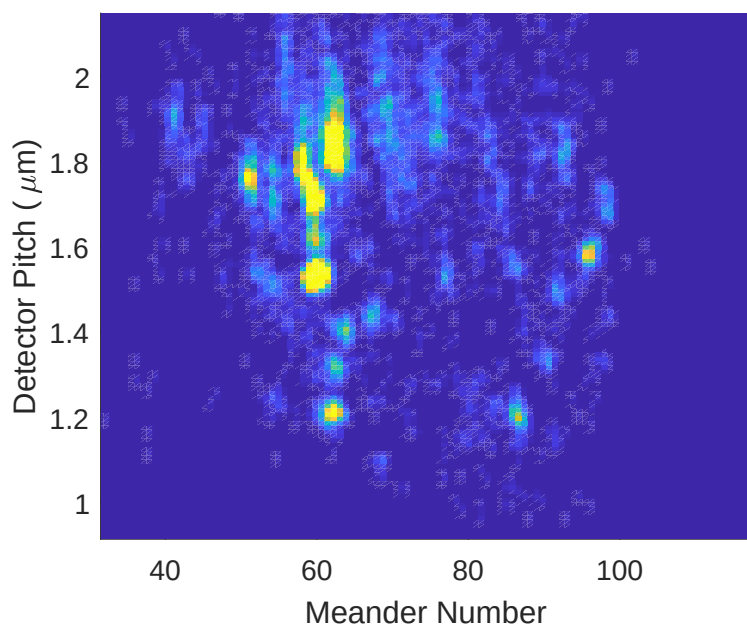


(a) 8 kV FIM Bias

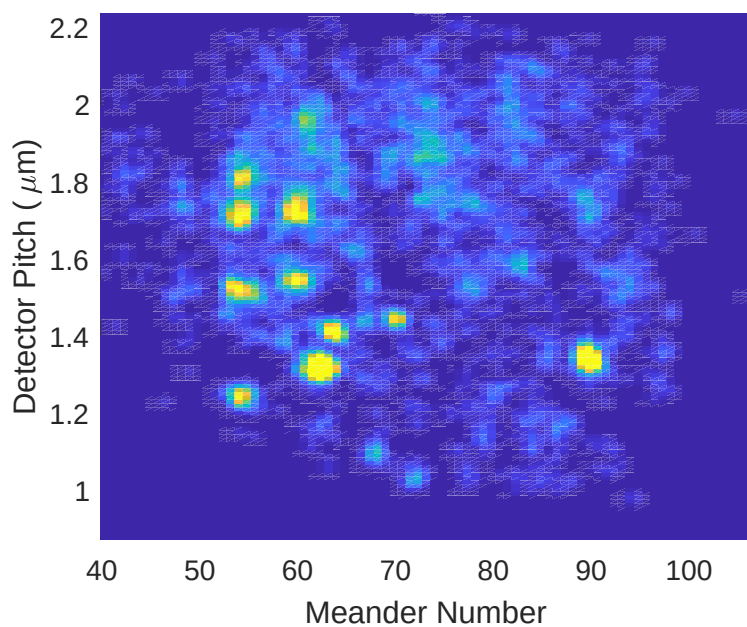


(b) 10 kV FIM Bias

Figure 5.25: Pitch Imager FIM Tip Images, Set 2



(a) 8 kV FIM Bias



(b) 8.4 kV FIM Bias

Figure 5.26: Pitch Imager FIM Tip Images, Set 3. Taken with longer Tungsten emitter to reduce magnification of surface image.

5.5 Conclusions and Outlook

In summary, developed and measured a superconducting delay line detector which is capable of detecting kiloelectron volt ions over a large active area with position sensitivity. The uncertainty of the time of arrival of pulses generated by this detector has been measured to be approximately 30 picoseconds in our samples. That corresponds to a physical length of 1.4 millimeters in our samples. We have shown that these detectors can be designed to have hotspots which self-reset and what design/performance trade-offs are necessary to do so. In the detectors measured in this thesis, however, the hotspots formed by ions latch, and can be used to drive adjacent detectors normal. We have demonstrated a design where this effect is useful to improve the effective timing resolution of our detector and used it to image the surface of tungsten field emitters. We also devised a scheme to probe the detection efficiency of these detectors; however, our fabrication capabilities have not, to this point, been capable of producing detector pairs with a small enough pitch to measure these detection efficiency parameters.

In the future, it would be useful to repeat the Pitch Experiment with these tighter pitches. This could be achieved by either using electron beam lithography to write the gaps which are too small to create using i-line lithography. Another possible approach would be to make the two detectors on separate layers, as shown in Figure 5.27. By doing so, the minimum achievable detector gap (the distance from one edge of the detector to the edge of the adjacent detector) can be reduced from the resolution specification (500 nm) of the i-Line stepper to the overlay specification (100 nanometers). If some inaccuracy in the placement of the detectors can be tolerated (for instance, if the dielectric between the two layers is planarized, as shown in Figure 5.27), this spacing can be

as small as desired. Although the detection efficiency of nanowire detectors has been reported to be unity for kiloelectron volt ions[37], the cross sectional area of our detectors is significantly larger, and so it would be worthwhile to measure the detection efficiency the SCDLD.

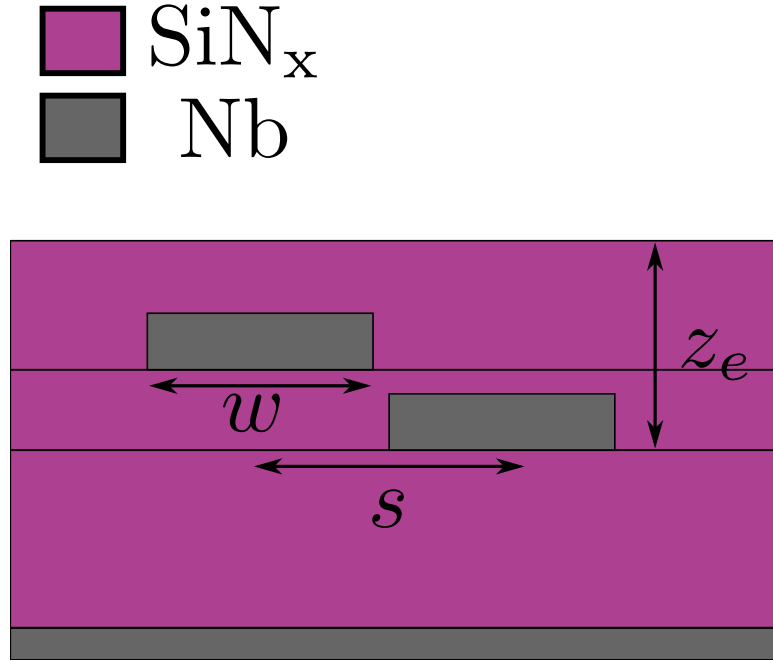


Figure 5.27: Multilayer Pitch Experiment

SCDLDs seem like a promising new technology for ion detection. There are some caveats, however. First of all, they still require a cryogenic environment (although not as sophisticated as the cryogenics necessary for transition edge sensors or microwave kinetic inductance detector). Secondly, they need to be carefully designed to not latch (which also leads to them having long dead times). They do not have proper kinetic energy resolving power (aside from turning down the bias current to alter the threshold energy, which then sacrifices detection efficiency). Although the SCDLD can cover extremely

large areas compared to traditional SNSPDs, in order to cover an area as large as the biggest microchannel plates, it would likely still be necessary to use an array of detectors. We have found that the maximum width an SCDLD can be while maintaining sensitivity to these ions is 1 micron. To cover an entire active area of 9 cm would require an SCDLD which is 90 kilometers long. Obvious fabrication issues aside, the electrical delay across such a detector would be nearly 2 milliseconds. This is unacceptably long and would severely limit the achievable count rate. It is therefore necessary to split this active area into an array of separate detectors. Pre-existing work on the implementation of nanowire detector arrays could be applied to engineering such a solution. In short, there is still a significant amount of work to be done before this technology is ready to be deployed in an industrial setting.

Appendix A

Linear Mode Response

A.1 Voltage Due to time Varying Cooper Pair Density

We begin by writing the probability current for a charged particle in an electromagnetic field.

$$\begin{aligned}\vec{j} &= \frac{1}{2m}[\Psi^*(\hat{p} - q\vec{A})\Psi + \Psi(\hat{p} - q\vec{A})^*\Psi^*] \\ \vec{j} &= \frac{1}{2m}[\Psi^*(\frac{\hbar}{i}\nabla)\Psi - \Psi(\frac{\hbar}{i}\nabla)\Psi^* - 2q\vec{A}|\Psi|^2]\end{aligned}$$

Insert the superconducting macroscopic wavefunction as Ψ :

$$\begin{aligned}\Psi &= \sqrt{n_s(\vec{x}, t)} \exp(i\Theta(\vec{x}, t)) \\ \frac{\partial \Psi}{\partial x_i} &= i \frac{\partial \Theta}{\partial x_i} \sqrt{n_s} e^{i\Theta} + \frac{\frac{\partial n_s}{\partial x_i} e^{i\Theta}}{2\sqrt{n_s}}\end{aligned}$$

Multiplying this macroscopic wavefunction by the charge of each cooper pair (denoted as q here) yields the equation for the supercurrent. Consider a short length of thin, narrow wire (compared to the London penetration depth, λ_L), such that variations in $\vec{j}$ and $\vec{A}$ across the cross section of the wire are essentially zero. This allows the problem to be treated as one dimensional.

$$j_s = \frac{q}{2m} [n_s \hbar \frac{\partial \Theta}{\partial x} + \frac{\hbar}{2i} \frac{\partial n_s}{\partial x} + n_s \hbar \frac{\partial \Theta}{\partial x} - \frac{\hbar}{2i} \frac{\partial n_s}{\partial x} - 2qAn_s]$$

$$j_s = \frac{q}{m} [n_s \hbar \frac{\partial \Theta}{\partial x} - qAn_s]$$

Take the time derivative of this expression...

$$\frac{\partial j_s}{\partial t} = \frac{q}{m} [\hbar (\frac{\partial n_s}{\partial t} \frac{\partial \Theta}{\partial x} + n_s \frac{\partial^2 \Theta}{\partial t \partial x}) - q(n_s \frac{\partial A}{\partial t} + \frac{\partial n_s}{\partial t} A)]$$

Now, noting that: $\vec{E} = -\frac{\partial \vec{A}}{\partial t} - \nabla \phi$

We find:

$$qn_s(E + \frac{\partial \phi}{\partial x}) = \frac{m}{q} \frac{\partial j_s}{\partial t} - \hbar (\frac{\partial n_s}{\partial t} \frac{\partial \Theta}{\partial x} + n_s \frac{\partial^2 \Theta}{\partial t \partial x}) + qA \frac{\partial n_s}{\partial t}$$

$$E = \frac{m}{n_s q^2} \frac{\partial j_s}{\partial t} - \frac{\hbar}{q} (\frac{1}{n_s} \frac{\partial n_s}{\partial t} \frac{\partial \Theta}{\partial x} + \frac{\partial^2 \Theta}{\partial t \partial x}) + \frac{A}{n_s} \frac{\partial n_s}{\partial t} - \nabla \phi$$

This expression looks somewhat hairy, but it can be simplified by assuming the following about our situation:

1. The current remains unchanged throughout the perturbation
2. The phase of the wavefunction does not change significantly compared to the amplitude

$$E = \frac{\hbar}{q} \frac{1}{n_s} \frac{\partial n_s}{\partial t} \frac{\partial \Theta}{\partial x} + \frac{A}{n_s} \frac{\partial n_s}{\partial t} - \nabla \phi$$

Consider the steady state/initial conditions of the problem:

$$\begin{aligned} E(t=0) &\rightarrow 0 \\ \frac{\partial n_s(t=0)}{\partial t} &\rightarrow 0 \end{aligned}$$

Now, solving for $\nabla \phi$ yields:

$$\nabla \phi(t=0) = \frac{A}{n_s} 0 - \frac{\hbar}{q} \frac{1}{n_s} 0 \frac{\partial \Theta}{\partial x} = 0$$

This information, and along with the fact that there is no changes happening in ϕ as a function of time, mean that:

$$\begin{aligned} E &= \frac{\dot{n}_s}{n_s} \left(A - \frac{\hbar}{q} \frac{\partial \Theta}{\partial x} \right) \\ V &= \int_l \vec{E} \cdot d\vec{l} \end{aligned}$$

Now recalling that we are considering a small segment of essentially one dimensional wire, E can be taken to be parallel to the wire. The (arguably unphysical) situation wherein the variation of n_s happens equally across one well defined segment is treated...

$$n_s(t) = \begin{cases} n_0 & : t < t_0 \\ n_0(1 - n_p \mathcal{H}(x-a) \mathcal{H}(a+l-x)) & : t_0 < t < t_{final} \end{cases}$$

Where $\mathcal{H}$ is the Heaviside step function

That is, initially there is some equilibrium population of cooper pairs per unit volume. Then at time t_0 , the perturbation breaks a fraction, n_p , of cooper pairs per unit volume. This happens within the range of a to $a + l$. This makes the line integral very simple to compute:

$$V = \frac{\dot{n}_s}{n_s} \left(A - \frac{\hbar}{q} \frac{\partial \Theta}{\partial x} \right) l \quad (\text{A.1})$$

This suggests that a change in the density of cooper pairs per unit volume does indeed generate a voltage response. It is worthwhile to compare this voltage to what we would expect if we could treat the kinetic inductance of the superconductor as being the same as a geometric inductance. In this case, a simple application of Faraday's law yields:

$$V = \frac{\partial}{\partial t}(IL) = \frac{\partial I}{\partial t}L + I \frac{\partial L}{\partial t}$$

As previously stated, we are considering the case where the current doesn't significantly change.

$$V = I \frac{\partial L}{\partial t}$$

In the previous derivation, only the field generated by the supercurrent itself was considered, so here the inductance is taken as coming entirely from the kinetic inductance.

$$\begin{aligned} L_k &= \frac{m}{n_s q^2} \frac{l}{wt} \\ V &= I \frac{\partial L_k}{\partial t} = wt j_s \frac{\partial L_k}{\partial t} \\ V &= wt j_s \frac{-m \dot{n}_s}{n_s^2 q^2} \frac{l}{wt} \end{aligned}$$

Insert the current derived from the probability current...

$$V = \frac{\dot{n}_s}{n_s} \left(A - \frac{\hbar}{q} \frac{\partial \Theta}{\partial x} \right) l \quad (\text{A.2})$$

Comparing equations A.1 and A.2, we see the voltage derived in each case is exactly the same. This suggests that it is reasonable to consider the total inductance as being the sum of geometric and kinetic inductances for the purposes of calculating the voltage generated by a time varying cooper pair density.

A.2 Voltage Response of Linear Mode Equivalent Circuit

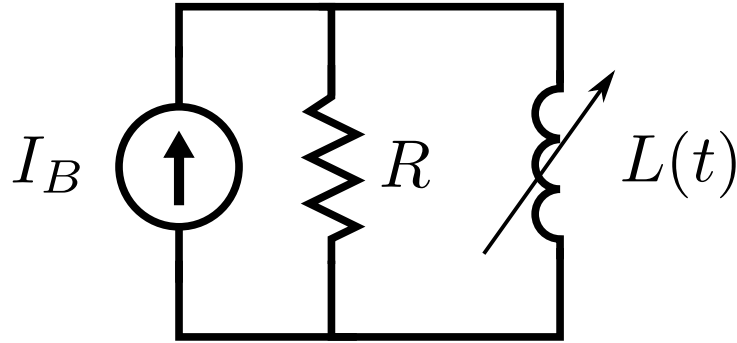


Figure A.1: Simplified Linear Mode Simulation Circuit

Using KCL, we can determine the current going into the load (R) as a function of

time:

$$\begin{aligned}
\frac{\dot{L}I_L + L\dot{I}_L}{R} + I_L &= I_B \\
\frac{\dot{L}I_L}{R} + I_L - I_B &= -\frac{L}{R}\dot{I}_L \\
\frac{R}{L}I_B - \frac{\dot{L}}{L}I_L - \frac{R}{L}I_L &= \dot{I}_L \\
\frac{R}{L}I_B - \frac{\dot{L}}{L}I_L - \frac{R}{L}I_L &= \frac{dI_L}{dt} \\
dt &= \frac{dI_L}{\frac{R}{L}I_B - \frac{\dot{L}}{L}(\dot{I}_L + R)} \\
\int dt &= \int \frac{dI_L}{\frac{R}{L}I_B - \frac{\dot{L}}{L}(\dot{I}_L + R)} \\
t &= -\frac{L}{\dot{L} + R} \log\left(\frac{R}{L}I_B - \frac{\dot{L} + R}{L}\right) \\
I_L(t) &= \frac{I_BR - Ce^{\frac{-(\dot{L}+R)}{L}t}}{\dot{L} + R} \\
I_L(0) &= I_B \rightarrow C = -I_B\dot{L} \\
I_L(t) &= \frac{I_BR + I_B\dot{L}e^{\frac{-(\dot{L}+R)}{L}t}}{R + \dot{L}}
\end{aligned}$$

Consider a perturbation in L that is, to first order, linear in time. We will consider the limit of a fast but small perturbation.

$$L(t) = \begin{cases} L_0 & : t < 0 \\ L_0 + \alpha t & : 0 < t < t_{final} \end{cases}$$

$$\alpha \gg R, \alpha \gg \frac{RL_0}{L_0 + \Delta L}$$

$$L(t_{final}) = \zeta L_0 = L_0 + \Delta L$$

$$\zeta \approx 1, \Delta L \ll 1$$

For the time during which this perturbation is applied ($0 < t < t_{final}$)...

$$I_L(t) = \frac{I_B R + I_B \alpha e^{\frac{-(R+\alpha)}{L_0+\alpha t} t}}{R + \alpha}$$

$$t_{final} = \frac{L_0(\zeta - 1)}{\alpha}$$

$$I_L(t_{final}) = I_B \frac{(R + \alpha e^{\frac{-(R+\alpha)}{\zeta L_0} (\frac{L_0(\zeta-1)}{\alpha})})}{R + \alpha}$$

$$I_L(t_{final}) = I_B \frac{(R + \alpha e^{(\frac{R}{\zeta \alpha} + \frac{1}{\zeta} - 1 - \frac{R}{\alpha})})}{R + \alpha}$$

$$\lim_{\alpha \rightarrow \infty} I_L(t_{final}) = I_B \frac{(0 + \alpha e^{(0 + \frac{1}{\zeta} - 1 - 0)})}{0 + \alpha}$$

$$I_L(t_{final}) = I_B (e^{(\frac{L_0}{L_0+\Delta L} - 1)})$$

$$I_L(t_{final}) = I_B (e^{(\frac{-\Delta L}{L_0+\Delta L})})$$

$$V_{out}(t_{final}) = (I_B - I_L(t_{final}))R$$

$$V_{out}(t_{final}) = I_B R (1 - e^{(\frac{-\Delta L}{L_0+\Delta L})})$$

$$\lim_{x \rightarrow 0} e^x \approx 1 + x + \frac{x^2}{2} + \dots$$

Which finally brings us to the result:

$$V_{out}(t_{final}) \approx I_B R \frac{\Delta L}{\Delta L + L_0} \quad (\text{A.3})$$

Appendix B

Parabolic Heat Equation between Two Closely Spaced Detectors

We seek to solve the forced heat equation in the lower half space of $\mathbb{R}^3$.

$$\frac{\partial \phi}{\partial t} = D_\phi \nabla^2 \phi \quad (\text{B.1})$$

We utilize the Green's function for the forced Fourier Heat transfer problem and use a heat source that is defined as:

$$F(\vec{x}, t) = h_{joule} \Theta(t) \delta(\vec{x})$$

Meaning, the heat source is located at the origin and activates at $t = 0$.

The Green's function in this case is[14]:

$$G(\vec{x}, t; \vec{y}, \tau) = \frac{\Theta(t - \tau)}{(4D_\phi \pi(t - \tau))^{3/2}} \exp\left(-\frac{|\vec{x} - \vec{y}|^2}{4D_\phi(t - \tau)}\right)$$

We assume the system starts at equilibrium so that the entire half space is at the same temperature. Without loss of generality, we can set this temperature to zero. We make the simplifying assumption that all the materials in the region of interest are equivalent in terms of heat capacity and thermal conductivity. This is clearly not true;

however, the majority of the space we are considering is purely silicon nitride, and the heat conduction inside the niobium itself is not of interest in this calculation. Now, we convolve the forcing function with the Green's function:

$$\begin{aligned}
\phi(\vec{x}, t) &= \int_0^\infty \int_{\mathbb{R}^3} G(\vec{x}, t; \vec{y}, \tau) F(\vec{y}, \tau) d^3\vec{y} d\tau \\
\phi(\vec{x}, t) &= \int_0^\infty \int_{\mathbb{R}^3} \frac{\Theta(t - \tau)}{(4D_\phi\pi(t - \tau))^{3/2}} \exp\left(-\frac{|\vec{x} - \vec{y}|^2}{4D_\phi(t - \tau)}\right) h_{joule} \Theta(\tau) \delta(\vec{y}) d^3\vec{y} d\tau \\
\phi(\vec{x}, t) &= \frac{1}{\pi^{3/2}} \int_0^\infty \frac{h_{joule}}{(4D_\phi(t - \tau))^{3/2}} \exp\left(-\frac{|\vec{x}|^2}{4D_\phi(t - \tau)}\right) d\tau \\
\phi(\vec{x}, t) &= \frac{-h_{joule}}{4D_\phi\pi^{3/2}} \int_t^{-\infty} \frac{\exp\left(\frac{-|\vec{x}|^2}{u}\right)}{u^{3/2}} du \\
\phi(\vec{x}, t) &= \frac{h_{joule}}{4D_\phi\pi|\vec{x}|} \operatorname{erf}\left(\frac{|\vec{x}|}{\sqrt{u}}\right) \Big|_{4D_\phi t}^{-\infty} \\
\phi(\vec{x}, t) &= \frac{h_{joule}}{4D_\phi\pi|\vec{x}|} \left(1 - \operatorname{erf}\left(\frac{|\vec{x}|}{\sqrt{4D_\phi t}}\right)\right)
\end{aligned}$$

This solution is true in the case that we are solving for the entirety of $\mathbb{R}^3$. However, in our case, there is a boundary condition that $\frac{\partial\phi}{\partial t}\Big|_{z=0} = 0$ because the silicon nitride ends there. Our solution does not inherently satisfy this boundary condition; however, because this PDE is linear, we are able to use the method of images to add another fictitious heat source that will fix this issue. This source is exactly the same as our original source, but reflected through the plane of $x = y = 0$. Because we set the heat source at $z = 0$, this reflected source is in the exact same place, so we just gain a factor of 2.

$$\phi(\vec{x}, t) = \frac{2h_{joule}}{4D_\phi\pi|\vec{x}|} \left(1 - \operatorname{erf}\left(\frac{|\vec{x}|}{\sqrt{4D_\phi t}}\right)\right)$$

If we consider the situation where we have encapsulated the detectors with some

thickness of nitride (defined as z_e), our unbounded solution shifts just slightly:

$$\text{Define } \vec{w} = (0, 0, -z_e)$$

Then our forcing function is now

$$F(\vec{x}, t) = h_{joule} \Theta(t) \delta(\vec{x} - \vec{w})$$

So the convolution integral is

$$\phi(\vec{x}, t) = \int_0^\infty \int_{\mathbb{R}^3} \frac{\Theta(t - \tau)}{(4D_\phi \pi(t - \tau))^{3/2}} \exp\left(-\frac{|\vec{x} - \vec{y}|^2}{4D_\phi(t - \tau)}\right) h_{joule} \Theta(\tau) \delta(\vec{y} - \vec{w}) d^3\vec{y} d\tau$$

Taking the integral over $\mathbb{R}^3$ gives...

$$\phi(\vec{x}, t) = \frac{1}{\pi^{3/2}} \int_0^\infty \frac{h_{joule}}{(4D_\phi(t - \tau))^{3/2}} \exp\left(-\frac{|\vec{x} - \vec{w}|^2}{4D_\phi(t - \tau)}\right) d\tau$$

And so the solution not considering the boundary imposed by the top of the silicon nitride encapsulation is:

$$\phi(\vec{x}, t) = \frac{h_{joule}}{4D_\phi \pi |\vec{x} - \vec{w}|} (1 - \operatorname{erf}\left(\frac{|\vec{x} - \vec{w}|}{\sqrt{4D_\phi t}}\right))$$

Again, we can implement the boundary conditions by inserting an image heat source at $-\vec{w}$, changing the solution to

$$\phi(\vec{x}, t) = \frac{h_{joule}}{4D_\phi \pi |\vec{x} - \vec{w}|} (1 - \operatorname{erf}\left(\frac{|\vec{x} - \vec{w}|}{\sqrt{4D_\phi t}}\right)) + \frac{h_{joule}}{4D_\phi \pi |\vec{x} + \vec{w}|} (1 - \operatorname{erf}\left(\frac{|\vec{x} + \vec{w}|}{\sqrt{4D_\phi t}}\right))$$

The $\vec{x}$ we are interested in is the closest approach of the adjacent detector. This is located at $(s, 0, -z_e)$, where s is the center to center spacing between the wires. We can now replace the vectors in our expression with geometrical parameters from the detector

$$\phi(\vec{x}, t) = \frac{h_{joule}}{4D_\phi \pi s} (1 - \operatorname{erf}\left(\frac{s}{\sqrt{4D_\phi t}}\right)) + \frac{h_{joule}}{4D_\phi \pi \sqrt{s^2 + 4z_e^2}} (1 - \operatorname{erf}\left(\frac{\sqrt{s^2 + 4z_e^2}}{\sqrt{4D_\phi t}}\right))$$

Taking arbitrary values for the thermal constants h_{joule} and D_ϕ and taking a typical value of z_e to be one fifth of s , this looks like:

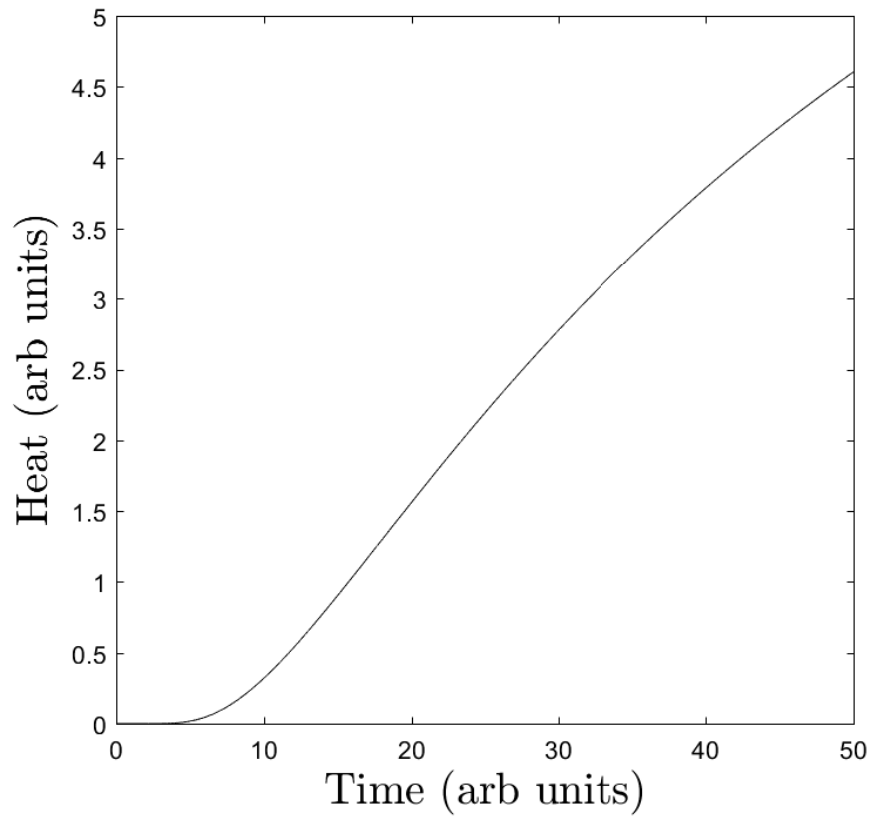


Figure B.1: Typical Heat vs Time Curve: $\phi(s, 0, z_e, t)$ (at Detector 2), $z_e = \frac{s}{5}$

Bibliography

- [1] V. A. ALTOV, *Stabilization of Superconducting Magnetic Systems*, Plenum Press New York, 1977.
- [2] A. J. ANNUNZIATA, *Single-Photon Detection, Kinetic Inductance, and Non-Equilibrium Dynamics in Niobium and Niobium Nitride Superconducting Nanowires*, PhD thesis, 2010. Copyright - Database copyright ProQuest LLC; ProQuest does not claim copyright in the individual underlying works; Last updated - 2016-03-10.
- [3] N. W. ASHCROFT AND D. N. MERMIN, *Solid State Physics*, Harcourt Inc., 1976.
- [4] M. CALOZ, M. PERRENOUD, C. AUTEBERT, B. KORZH, M. WEISS, C. SCHNENBERGER, R. J. WARBURTON, H. ZBINDEN, AND F. BUSSIRES, *High-detection efficiency and low-timing jitter with amorphous superconducting nanowire single-photon detectors*, Applied Physics Letters, 112 (2018), p. 061103.
- [5] A. CASABURI, E. ESPOSITO, M. EJRNAES, K. SUZUKI, M. OHKUBO, S. PAGANO, AND R. CRISTIANO, *A 2 × 2 mm² superconducting strip-line detector for high-performance time-of-flight mass spectrometry*, Superconductor Science and Technology, 25 (2012), p. 115004.
- [6] R. CRISTIANO, M. EJRNAES, A. CASABURI, N. ZEN, AND M. OHKUBO, *Superconducting nano-strip particle detectors*, Superconductor Science and Technology, 28 (2015), p. 124004.

- [7] P. DAHL, *Introduction to electron and ion optics*, Academic Press, 1973. Bibliography: page 142.
- [8] A. ENGEL, J. J. RENEMA, K. ILIN, AND A. SEMENOV, *Detection mechanism of superconducting nanowire single-photon detectors*, Superconductor Science and Technology, 28 (2015), p. 114003.
- [9] B. V. ESTEY, J. A. BEALL, G. C. HILTON, K. D. IRWIN, D. R. SCHMIDT, J. N. ULLOM, AND R. E. SCHWALL, *Time-of-flight mass spectrometry with latching nb meander detectors*, IEEE Transactions on Applied Superconductivity, 19 (2009), pp. 382–385.
- [10] H. FINK, *Mono-atomic tips for scanning tunneling microscopy*, IBM J. Res. Dev., 30 (1986), pp. 460–465.
- [11] G. FRASER AND J. PEARSON, *Soft x-ray energy resolution with microchannel plate detectors of high quantum detection efficiency*, Nuclear Instruments and Methods in Physics Research, 219 (1984), pp. 199 – 212.
- [12] M. S. GOCKENBACH, *Partial Differential Equations: Analytical and Numerical Methods*, Society for Industrial and Applied Mathematics, 2002.
- [13] A. V. GUREVICH AND R. G. MINTS, *Self-heating in normal metals and superconductors*, Rev. Mod. Phys., 59 (1987), pp. 941–999.
- [14] R. HABERMAN, *Elementary Applied Partial Differential Equations With Fourier Series and Boundary Value Problems*, Prentice Hall, 3rd edition ed., 1997.

- [15] K. INOUE, F. YANO, A. NISHIDA, H. TAKAMIZAWA, T. TSUNOMURA, Y. NAGAI, AND M. HASEGAWA, *Dopant distributions in n-mosfet structure observed by atom probe tomography*, Ultramicroscopy, 109 (2009), pp. 1479 – 1484.
- [16] S. B. KAPLAN, C. C. CHI, D. N. LANGENBERG, J. J. CHANG, S. JAFAREY, AND D. J. SCALAPINO, *Quasiparticle and phonon lifetimes in superconductors*, Phys. Rev. B, 14 (1976), pp. 4854–4873.
- [17] R. L. KAUTZ, *Picosecond pulses on superconducting striplines*, Journal of Applied Physics, 49 (1978), pp. 308–314.
- [18] T. F. KELLY AND D. J. LARSON, *The second revolution in atom probe tomography*, MRS Bulletin, 37 (2012), pp. 150–158.
- [19] T. F. KELLY AND M. K. MILLER, *Atom probe tomography*, Review of Scientific Instruments, 78 (2007), p. 031101.
- [20] A. J. KERMAN, E. A. DAULER, W. E. KEICHER, J. K. W. YANG, K. K. BERGGREN, G. GOLTSMAN, AND B. VORONOV, *Kinetic-inductance-limited reset time of superconducting nanowire photon counters*, Applied Physics Letters, 88 (2006), p. 111116.
- [21] A. J. KERMAN, J. K. W. YANG, R. J. MOLNAR, E. A. DAULER, AND K. K. BERGGREN, *Electrothermal feedback in superconducting nanowire single-photon detectors*, Phys. Rev. B, 79 (2009), p. 100509.
- [22] A. KORNEEV, P. KOUMINOV, V. MATVIENKO, G. CHULKOVA, K. SMIRNOV, B. VORONOV, G. N. GOLTSMAN, M. CURRIE, W. LO, K. WILSHER, J. ZHANG,

- W. SYSZ, A. PEARLMAN, A. VEREVKIN, AND R. SOBOLEWSKI, *Sensitivity and gigahertz counting performance of nbn superconducting single-photon detectors*, Applied Physics Letters, 84 (2004), pp. 5338–5340.
- [23] M. KUDUZ, G. SCHMITZ, AND R. KIRCHHEIM, *Investigation of oxide tunnel barriers by atom probe tomography (tap)*, Ultramicroscopy, 101 (2004), pp. 197 – 205.
- [24] F. MARSILI, V. B. VERMA, J. A. STERN, S. HARRINGTON, A. E. LITA, T. GERRITS, I. VAYSHENKER, B. BAEK, M. D. SHAW, R. P. MIRIN, AND S. W. NAM, *Detecting single infrared photons with 93% system efficiency*, Nature Photonics, 7 (2013), p. 210.
- [25] R. J. MATHAR, *Solid Angle of a Rectangular Plate*, tech. rep., Max-Planck Institute of Astronomy, 05 2015.
- [26] B. A. MAZIN, *Microwave kinetic inductance detectors*, PhD thesis, California Institute of Technology, 2005.
- [27] D. MCCAMMON, *Thermal Equilibrium Calorimeters – An Introduction*, Springer Berlin Heidelberg, Berlin, Heidelberg, 2005, pp. 1–34.
- [28] S. MIKI, M. FUJIWARA, M. SASAKI, B. BAEK, A. J. MILLER, R. H. HADFIELD, S. W. NAM, AND Z. WANG, *Large sensitive-area nbn nanowire superconducting single-photon detectors fabricated on single-crystal mgo substrates*, Applied Physics Letters, 92 (2008).

- [29] S. MIKI, T. YAMASHITA, Z. WANG, AND H. TERAJ, *A 64-pixel nbtin superconducting nanowire single-photon detector array for spatially resolved photon detection*, Opt. Express, 22 (2014), pp. 7811–7820.
- [30] M. K. MILLER, T. F. KELLY, K. RAJAN, AND S. P. RINGER, *The future of atom probe tomography*, Materials Today, 15 (2012), pp. 158 – 165.
- [31] I. MILOSTNAYA, A. KORNEEV, I. RUBTSOVA, V. SELEZNEV, O. MINAEVA, G. CHULKOVA, O. OKUNEV, B. VORONOV, K. SMIRNOV, G. GOL'TSMAN, W. SLYSZ, M. WEGRZECKI, M. GUZIEWICZ, J. BAR, M. GORSKA, A. PEARLMAN, J. KITAYGORSKY, A. CROSS, AND R. SOBOLEWSKI, *Superconducting single-photon detectors designed for operation at 1.55-m telecommunication wavelength*, Journal of Physics: Conference Series, 43 (2006), p. 1334.
- [32] E. W. MÜLLER, *Das feldionenmikroskop*, Zeitschrift für Physik, 131 (1951), pp. 136–142.
- [33] C. M. NATARAJAN, M. G. TANNER, AND R. H. HADFIELD, *Superconducting nanowire single-photon detectors: physics and applications*, Superconductor Science and Technology, 25 (2012), p. 063001.
- [34] D. M. POZAR, *Microwave Engineering*, Wiley, third edition ed., 2004.
- [35] D. ROSENBERG, A. J. KERMAN, R. J. MOLNAR, AND E. A. DAULER, *High-speed and high-efficiency superconducting nanowire single photon detector array*, Opt. Express, 21 (2013), pp. 1440–1447.

- [36] D. F. SANTAVICCA AND D. E. PROBER, *Aging of ultra-thin niobium films*, IEEE Transactions on Applied Superconductivity, 25 (2015), pp. 1–4.
- [37] M. SCLAFANI, M. MARKSTEINER, F. M. KEIR, A. DIVOCHIY, A. KORNEEV, A. SEMENOV, G. GOLTSMAN, AND M. ARNDT, *Sensitivity of a superconducting nanowire detector for single ions at low energy*, Nanotechnology, 23 (2012), p. 065501.
- [38] C. E. SHANNON, *Communication in the presence of noise*, Proceedings of the IEEE, 86 (1998), pp. 447–457.
- [39] S. SMITH, *Digital Signal Processing: A Practical Guide for Engineers and Scientists*, Newnes, first edition ed., 2002.
- [40] M. J. SOUTHON, *Field Emission and Field Ionization*, Springer US, Boston, MA, 1968, pp. 6–27.
- [41] K. SUZUKI, S. MIKI, Z. WANG, Y. KOBAYASHI, S. SHIKI, AND M. OHKUBO, *Superconducting nbn thin-film nanowire detectors for time-of-flight mass spectrometry*, Journal of Low Temperature Physics, 151 (2008), pp. 766–770.
- [42] K. SUZUKI, M. OHKUBO, M. UKIBE, K. CHIBA-KAMOSHIDA, S. SHIKI, S. MIKI, AND Z. WANG, *Charge-state-derivation ion detection using a super-conducting nanostructure device for mass spectrometry*, Rapid Communications in Mass Spectrometry, 24 (2010), pp. 3290–3296.

- [43] K. SUZUKI, S. SHIKI, M. UKIBE, M. KOIKE, S. MIKI, Z. WANG, AND M. OHKUBO, *Hot-spot detection model in superconducting nano-stripline detector for kev ions*, Applied Physics Express, 4 (2011), p. 083101.
- [44] M. TINKHAM, *Introduction to Superconductivity: Second Edition (Dover Books on Physics) (Vol i)*, Dover Publications, second edition ed., June 2004.
- [45] G. TURIN, *An introduction to matched filters*, IRE Transactions on Information Theory, 6 (1960), pp. 311–329.
- [46] J. N. ULLOM AND D. A. BENNETT, *Review of superconducting transition-edge sensors for x-ray and gamma-ray spectroscopy*, Superconductor Science and Technology, 28 (2015), p. 084003.
- [47] R. URBAN, R. A. WOLKOW, AND J. L. PITTERS, *Field ion microscope evaluation of tungsten nanotip shape using he and ne imaging gases*, Ultramicroscopy, 122 (2012), pp. 60 – 64.
- [48] J. W. VALLEY, A. J. CAVOSIE, T. USHIKUBO, D. A. REINHARD, D. F. LAWRENCE, D. J. LARSON, P. H. CLIFTON, T. F. KELLY, S. A. WILDE, D. E. MOSER, AND M. J. SPICUZZA, *Hadean age for a post-magma-ocean zircon confirmed by atom-probe tomography*, Nature Geoscience, 7 (2014), p. 219.
- [49] V. B. VERMA, A. E. LITA, M. R. VISSERS, F. MARSILI, D. P. PAPPAS, R. P. MIRIN, AND S. W. NAM, *Superconducting nanowire single photon detectors fabricated from an amorphous $\text{mo}_{0.75}\text{ge}_{0.25}$ thin film*, Applied Physics Letters, 105 (2014), p. 022602.

- [50] C. WU, *Intrinsic stress of magnetron-sputtered niobium films*, Thin Solid Films, 64 (1979), pp. 103 – 110. International Conference on Metallurgical Coatings, San Diego, 1979-Part III.
- [51] T. YAMASHITA, S. MIKI, H. TERAJ, K. MAKISE, AND Z. WANG, *Crosstalk-free operation of multielement superconducting nanowire single-photon detector array integrated with single-flux-quantum circuit in a 0.1 w gifford-mcmahon cryocooler*, Opt. Lett., 37 (2012), pp. 2982–2984.
- [52] J. K. YANG, *Superconducting nanowire single-photon detectors and sub-10-nm lithography*, PhD thesis, Massachusetts Institute of Technology, 2009.
- [53] N. ZEN, A. CASABURI, S. SHIKI, K. SUZUKI, M. EJRNAES, R. CRISTIANO, AND M. OHKUBO, *1 mm ultrafast superconducting stripline molecule detector*, Applied Physics Letters, 95 (2009), p. 172508.
- [54] N. ZEN, S. SHIKI, M. UKIBE, M. KOIKE, AND M. OHKUBO, *Ion-induced dynamical change of supercurrent flow in superconducting strip ion detectors with parallel configuration*, Applied Physics Letters, 104 (2014).
- [55] Q.-Y. ZHAO, D. ZHU, N. CALANDRI, A. E. DANE, A. N. MCCAUGHAN, F. BELLEI, H.-Z. WANG, D. F. SANTAVICCA, AND K. K. BERGGREN, *Single-photon imager based on a superconducting nanowire delay line*, Nature Photonics, 11 (2017), p. 247.
- [56] J. ZIEGLER, J. BIRSACK, AND M. ZIEGLER, *SRIM: The Stopping and Range of Ions in Matter*, SRIM Co, 2015 edition ed., 2015.

- [57] B. ZINK AND F. HELLMAN, *Specific heat and thermal conductivity of low-stress amorphous si-n membranes*, Solid State Communications, 129 (2004), pp. 199 – 204.