

UNIVERSITY OF CALIFORNIA

Los Angeles

**Plasma Density Transition Trapping of
Electrons in Plasma Wake Field Accelerators**

A dissertation submitted in partial satisfaction

of the requirements for the degree

Doctor of Philosophy in Physics

by

Matthew Colin Thompson

2004

© Copyright by
Matthew Colin Thompson
2004

The dissertation of Matthew Colin Thompson is approved.

Chandrashekhar Joshi

Claudio Pellegrini

David Saltzberg

James B. Rosenzweig, Committee Chair

University of California, Los Angeles

2004

To my family, for always believing . . .

TABLE OF CONTENTS

1	Introduction	1
1.1	The Need for Advanced Accelerators	2
1.2	Plasma Based Accelerators	4
1.2.1	The Plasma Wake Field Accelerator	9
1.2.2	The Laser Wake Field Accelerator	11
1.2.3	The Plasma Beat-Wave Accelerator	12
1.2.4	The Self-Modulated Laser Wake Field Accelerator	12
1.3	Conventional Electron Beam Sources	13
1.3.1	Thermionic Direct Current Guns	16
1.3.2	Radio Frequency Photoinjectors	20
1.3.3	Magnetic Compression	25
1.4	Injection and Timing	29
1.5	Plasma Based Electron Beam Sources	33
1.5.1	Random Phase Injection Sources	34
1.5.2	Laser Stimulated Injection Sources	36
1.5.3	Plasma Density Transition Trapping	37
1.6	Chapter Summary	40
2	Theory of Particle Trapping and Acceleration	41
2.1	Trapping in RF Photoinjectors	41
2.2	Plasma Accelerators and Plasma Trapping	45

3 Numerical Simulations of Plasma Density Transition Trapping Regimes	50
3.1 The Strong Blowout Scenario	51
3.2 The Weak Blowout Scenario	55
4 Scaling of the Transition Trapping System	61
4.1 Driver Charge Scaling	61
4.2 Wavelength Scaled Sources	63
5 Design of the Transition Trapping Experiment	68
5.1 Design Goals	69
5.2 The Argon Pulse Discharge Plasma Source	70
5.2.1 General Source Design	71
5.2.2 Plasma Cathode Heater Design	73
5.2.3 Plasma Source Operation and Characterization	77
5.3 Creation of the Plasma Density Transition	79
5.3.1 Transition Production Using Metal Screens	81
5.3.2 Baffled Screens	84
5.3.3 Density Transition Characterization	86
5.4 Simulation of Trapping Performance Under Realistic Experimental Conditions	91
5.4.1 Realistic Density Profiles	91
5.4.2 Impact of Transverse Magnetic Field	94
5.4.3 Imperfect Driving Beams	95

5.5	Comparison of Design Goals and Achieved Results	97
6	Execution of the Transition Trapping Experiment	99
6.1	The FNPL Facility	99
6.1.1	Cathode Drive Laser	102
6.1.2	L-Band RF Photoinjector	103
6.1.3	Nine-Cell Superconducting Accelerating Cavity	104
6.2	The Trapping Experiment Beamline	105
6.2.1	Vacuum Window	105
6.2.2	Diagnostic Positions	108
6.2.3	Beam Transport	109
6.3	Diagnostics	110
6.3.1	Transverse Profile Diagnostics	111
6.3.2	Screen Position Determination	112
6.3.3	Charge Diagnostics	114
6.3.4	Streak Camera	115
6.3.5	The Broad Range Vacuum Spectrometer	117
7	Experimental Results and Analysis	122
7.1	Plasma Focusing	122
7.2	Drive Beam Characterization	124
7.3	Drive Beam Deceleration	132
7.4	The Search for Trapped Electrons	133
7.5	Comparison of Experimental Results with Simulation	135

8	Conclusions and Future Directions	139
8.1	Conclusions	139
8.2	Future Directions	141
8.2.1	Continuation of the Low Density Trapping Experiment . .	141
8.2.2	Underdense Plasma Lens Experiment	143
8.2.3	Prospects for High Density Trapping Experiments	143
8.2.4	Foil Trapping	144
A	Derivation of the Hamiltonian for Particles in Electromagnetic Fields	147
B	Heat Transport in the Radiation Coupled Regime	150
C	Analysis of Electrostatic Probe Signals	154
D	Optical Design of the Broad Range Vacuum Spectrometer . .	159
	References	167

LIST OF FIGURES

1.1	Plot of the plasma wave breaking field and skin depth over the densities of interest for advanced accelerators.	8
1.2	Area occupied by the electron beam in trace space at three different locations along a beamline.	14
1.3	Idealized electron beam sources.	16
1.4	Schematic of an example thermionic DC electron gun	17
1.5	Simplified diagram of an RF photoinjector	20
1.6	The mechanism through which beams are longitudinally focused in photoinjectors by differential phase slippage.	24
1.7	The mechanism of magnetic compression	26
1.8	A typical set of chicane magnets.	27
1.9	Comparison of RF and Plasma accelerator acceptance windows . .	29
1.10	Suppression of laser/RF jitter via beam compression	31
1.11	Configuration space diagram of the plasma density transition trapping mechanism	38
1.12	Diagram of the plasma electron rephasing caused by the plasma density transition	39
3.1	Strong blowout configuration space (r, z)	52
3.2	Trapped particle energy versus temporal position within the bunch for the strong blowout case.	53
3.3	Weak blowout configuration space (r, z)	56

3.4	Plasma density and drive beam current profiles.	58
3.5	Trapped particle energy versus temporal position within the bunch for the weak blowout case.	59
3.6	Trapped particle radial position versus temporal position within the bunch for the weak blowout case.	60
4.1	The effects of driver beam charge scaling on trapped beams. . . .	62
5.1	Schematic of the plasma density transition experiment plasma source.	71
5.2	CAD rendered pictures of the plasma source	72
5.3	Schematic of the plasma source heater.	76
5.4	Photograph of the plasma column.	78
5.5	Measured plasma column density	79
5.6	Simulated dependence of captured charge on transition length . .	80
5.7	Plasma density screen diagram	81
5.8	PIC simulation of density screen performance	83
5.9	Blocking of wake particles by the density screen baffle.	85
5.10	Effect of the beam-baffle distance on trapping	86
5.11	Photograph of the plasma density transition apparatus.	87
5.12	Measured density transition - long scale.	88
5.13	Measured density transition - short scale.	89
5.14	Comparison of gaussian and linear plasma density profiles. . . .	93
5.15	Impact of drive beam length on trapped charge.	96

5.16	The density profile used in the experiment.	98
6.1	The FNPL photoinjector primary beamline	101
6.2	The plasma density transition trapping beamline	106
6.3	Beam charge transmitted past the baffle at beam centroid/baffle edge spacings of interest.	113
6.4	Illustration of streak camera operation.	115
6.5	Example of streak camera data anaylsis.	116
6.6	Mechanical diagram of the spectrometer.	118
6.7	Photograph of a test beam on the spectrometer exit port.	120
6.8	Typical operating configuration of the spectrometer.	121
7.1	Plasma focusing of the uncompressed beam	123
7.2	Variation in charge transport with photoinjector charge - short time scale	125
7.3	Variation in charge transport with photoinjector charge - long time scale	126
7.4	Quadrupole scan emittance measurement	127
7.5	Correlations between the beam length, spot size, and charge . . .	130
7.6	Deceleration of the drive beam due to plasma wake fields	132
7.7	Observation of the baffle edge blocking the drive beam	134
7.8	Weak low energy signal observed on the broad range spectrometer.	135
7.9	Trace 3D simulation of the trapping experiment beamline	138
B.1	Diagram of the radiative heat transfer between two planes. . . .	151

C.1	I-V curve from an electrostatic probe measurement.	155
C.2	I-V curve from an electrostatic probe measurement with curve fits	157
D.1	Variable definitions for the spectrometer optics calculations. . . .	160
D.2	The error expected when observing the captured beam.	165
D.3	The error expected when observing the head of the drive beam. . .	165
D.4	The error expected when observing the middle of the drive beam.	166
D.5	The error expected when observing the tail of the drive beam. . .	166

LIST OF TABLES

1.1	Nominal operating parameters of an example thermionic gun . . .	19
1.2	Nominal operating parameters of an example photoinjector	22
1.3	Theoretical operating parameters of an example plasma electron beam source	34
3.1	Drive and captured beam parameters in the strong blowout case. .	51
3.2	Drive and captured beam parameters in the weak blowout case. .	57
4.1	Simulations of wavelength scaling using MAGIC 2D.	65
5.1	Driving and captured beam parameters for gaussian and linear plasma density profiles	92
6.1	Nominal FNPL operating parameters	100
6.2	Comparison of the scintillators used in this experiment.	112
6.3	Design parameters of the broad range in vacuum spectrometer. . .	117
7.1	Comparison of the design beam parameters with the best drive beam parameters that were achieved	129
7.2	Simulated trapping performance with achieved beam parameters .	136
8.1	Comparison of thin underdense plasma lens experiment parameters	143
8.2	Comparison of foil trapping α parameters.	145
D.1	Simulated parameters of the trapped and drive beam	163

ACKNOWLEDGMENTS

There are many people I wish to thank for their help and support during my graduate career and the work leading to this dissertation. While there is not space enough for me to thank everyone who has helped me over the past several years, I have tried hard to include everyone who has played a major role.

First and foremost I want to thank my advisor James Rosenzweig. Without his guidance and mentorship this dissertation would have been impossible. I have always enjoyed our conversations about physics, politics, history, or (more often) all three at the same time. At this point I should also thank our lab manager Gil Travish who in many ways came in as Jamie's relief pitcher in the later innings of my graduate school career. His enthusiasm and skills as an experimentalist were an immense help.

The experimental portion of this dissertation was performed at Fermilab's A0 photoinjector facility. Many thanks go to Nick Barov, Helen Edwards, Philippe Piot, Rodion Tikhoplav and the rest of the scientific staff there for helping with and hosting my experiment. In the same vein, I would like to thank Hyyong Suk who originated the idea of plasma density transition trapping when he was at UCLA and initiated me into the self-destructing world of plasma sources. Garnick Hairapetian, Chris Clayton, and the other members of Chan Joshi's group at UCLA also have my gratitude for their work on the early development of the plasma source and their help when I inherited it. Thanks also go to Tom Katsouleas of USC, along with Warren Mori and Wei Lu of UCLA for helping with some of the computer simulations.

It is hard for me to include and thank every technician and craftsman who worked on parts of my plasma source and diagnostic spectrometer. First among

these are Harry Lockart, Ted Aliado, Tom Asbury, Ken Luk, John Morrison and the other fine machinists of the UCLA Physical Sciences Machine Shop. Between them, they built most of my apparatus. I also appreciate all the help the A0 vacuum technicians, Wade Muranyi, Mike, and Rocky, gave me. They always went out of their way to help, and even farther out of their way to play a prank. James Santucci, A0's jack-of-all-trades technician, also provide much indispensable help.

I would also like to thank all my colleagues and fellow students at the UCLA Particle Beam Physics Lab, past and present: Claudio Pellegrini, Scott Anderson, Soren Telfer, Pedro Frigola, Ron Agustsson, Sven Reiche, Alex Murokh, Gerard Andonian, and Joel England. They have all helped me, taught me, and offered me their friendship and camaraderie. Special thanks goes to my contemporary Pietro Musumeci with whom I have shared my journey through PBPL from novice to graduate. My deepest gratitude also extends to my two undergraduate assistants Michael Schneider and Chris Muller. They are some of the finest guys I have had the privilege to work with, and I know they will both go far.

Finally, I want to thank my good graduate student friends outside the lab, Gintaras Duda and Derek Walter, for all their support. I never would have made it without you guys, and I highly recommend our tradition of tri-weekly weight lifting support group sessions to all the grad students out there.

VITA

1975	Born, Fullerton, California, USA.
1996-1997	Teaching Assistant, Department of Physics, Stanford University
1997	B.S. with Honors, Physics, Stanford University
1997-1998	Teaching Assistant, Department of Physics, University of California, Los Angeles
1998	M.S., Physics, University of California, Los Angeles
1999	36th Culham Plasma Physics Summer School, Oxford, UK
1998-2004	Graduate Research Assistant, UCLA Particle Beam Physics Laboratory

PUBLICATIONS

S.G. Anderson, R. Agustsson, S. Boucher, A. Burke, C.E. Clayton, J. England, M. Loh, P. Musumeci, J.B. Rosenzweig, H. Suk, M.C. Thompson, Commissioning of the Neptune Photoinjector. *Proceedings of the 2001 Particle Accelerator Conference*, 89 (IEEE, 2001).

S.G. Anderson, M. Loh, P. Musumeci, J.B. Rosenzweig, H. Suk, M.C. Thompson, Commissioning and Measurements of the Neptune Photo-injector, *Advanced*

Accelerator Concepts Ninth Workshop Conference Proceedings, AIP Conf. Proc. **569**, 487 (AIP 2001).

S.G. Anderson, J.B. Rosenzweig, P. Musumeci, and M.C. Thompson, Horizontal Phase-Space Distortions Arising from Magnetic Pulse Compression of an Intense, Relativistic Electron Beam, *Physical Review Letters* **91**, 074803 (2003).

N. Barov, M.C. Thompson, K. Bishofberger, J.B. Rosenzweig, H. Edwards, and J. Santucci, Plasma Wakefield Experiments, *Advanced Accelerator Concepts Tenth Workshop Conference Proceedings*, AIP Conf. Proc. **647**, 71 (AIP 2002).

N. Barov, J.B. Rosenzweig, M.C. Thompson, R. Yoder, Energy Loss of a High Charge Bunched Electron Beam in a Plasma: Analysis, Accepted for publication in *Physical Review Special Topics Accelerators and Beams*(2004).

R.J. England, J.B. Rosenzweig and M.C. Thompson, Longitudinal Beam Shaping and Compression Scheme for the UCLA Neptune Laboratory, *Advanced Accelerator Concepts Tenth Workshop Conference Proceedings*, AIP Conf. Proc. **647**, 884 (AIP 2002).

P. Musumeci, R.J. England, M.C. Thompson, R. Yoder, and J.B. Rosenzweig, Velocity Bunching Experiment at the Neptune Laboratory, *Advanced Accelerator Concepts Tenth Workshop Conference Proceedings*, AIP Conf. Proc. **647**, 858 (AIP 2002).

J.B. Rosenzweig, N. Barov, M.C. Thompson, and R. Yoder, Energy Loss of a

High Charge Bunched Electron Beam in Plasma: Nonlinear Plasma Response and Linear Scaling, *Advanced Accelerator Concepts Tenth Workshop Conference Proceedings*, AIP Conf. Proc. **647**, 577 (AIP 2002).

H. Suk, N. Barov, J. England, J.B. Rosenzweig, M.C. Thompson, Dynamics of a Driver Beam Propagating in an Underdense Plasma with a Downward Density Transition, *Proceedings of the 2001 Particle Accelerator Conference*, 4011 (IEEE, 2001).

M.C. Thompson, C.E. Clayton, J. England, J.B. Rosenzweig, H. Suk, Beam-Plasma Interaction Experiments at the UCLA Neptune Laboratory, *Proceedings of the 2001 Particle Accelerator Conference*, 4014 (IEEE, 2001).

M.C. Thompson and J.B. Rosenzweig, Production and synchronization of electron beams from RF photoinjector/compressor systems for ultra-fast applications, *Advanced Accelerator Concepts Ninth Workshop Conference Proceedings*, AIP Conf. Proc. **569**, 374 (AIP 2001).

M.C. Thompson and J.B. Rosenzweig, Plasma Density Transition Trapping as a Possible High-Brightness Electron Beam Source, *Advanced Accelerator Concepts Tenth Workshop Conference Proceedings*, AIP Conf. Proc. **647**, 600 (AIP 2002).

M.C. Thompson and J.B. Rosenzweig, Plasma Density Transition Trapping as a Possible High-Brightness Electron Beam Source, In *The Physics and Applications of High Brightness Electron Beams*, World Scientific (2003).

M.C. Thompson, N. Barov, W. Lu, W. Mori, J.B. Rosenzweig, G. Travish, The UCLA/NICADD Plasma Density Transition Trapping Experiment, *Proceedings of the 2003 Particle Accelerator Conference*, 1870 (IEEE, 2003).

M.C. Thompson, J.B. Rosenzweig, and H. Suk, Plasma Density Transition Trapping as a Possible High-Brightness Electron Beam Source, *Physical Review Special Topics - Accelerators and Beams* **7**, 011301 (2004).

ABSTRACT OF THE DISSERTATION

Plasma Density Transition Trapping of Electrons in Plasma Wake Field Accelerators

by

Matthew Colin Thompson

Doctor of Philosophy in Physics

University of California, Los Angeles, 2004

Professor James B. Rosenzweig, Chair

Plasma based electron beam sources, which are now under development, will produce beams with much higher particle densities than are currently available. Plasma sources can create beams with brightness (the measure of achievable beam density) orders of magnitude greater than radio frequency photoinjectors, the current state-of-the-art. Plasma density transition trapping is one example of the many plasma electron beam source schemes under development. Plasma density transition trapping is a recently proposed self-injection mechanism for plasma wake field accelerators. The technique uses a sharp downward plasma density transition to trap and accelerate background plasma electrons in a plasma wake field.

This dissertation examines the different regimes in which plasma density transition trapping can operate and the quality of the electron beams captured in terms of emittance, energy spread, and brightness. This is accomplished using two-dimensional Particle-In-Cell (PIC) simulations which show that the captured beam parameters can be optimized by manipulating the overall plasma density, as well as the density profile. A general set of scaling laws is developed that

predicts how the brightness of transition trapping beams, or the beams produced by any plasma system, scales with the plasma density. These scaling laws predict that beam brightness increases linearly with the plasma density of the source.

The design and execution of the first plasma density transition trapping experiment is also documented in this dissertation. Plasma with density on the order of 10^{13} cm^{-3} was used in the experiment. Plasma density transitions steep enough to produce trapping at this density were created and measured. Interaction between the plasma transition and a driving electron beam pulse did not, however, produce trapped electrons. Detailed measurements of the drive beam parameters revealed that it did not meet the trapping experiment design criteria. Simulations using the measured plasma and beam parameters predict zero captured charge in agreement with observations.

CHAPTER 1

Introduction

Accelerating charged particles to very high energies is both an interesting physical problem in its own right and an exceedingly valuable tool in other areas of technology and scientific inquiry. The need to create high energy particles is increasingly coupled with the need to produce concentrated beams of these particles and the correspondingly high densities of energy such intense beams represent. The push towards ever greater beam energy and intensity is straining the physical limits of conventional radio frequency (RF) accelerator technology. Accelerators based on the excitation of large amplitude oscillations in plasmas have been proposed as a means of exceeding the limits of RF technology in the case of electron acceleration. High energy electron accelerators, regardless of the mechanisms they employ, consist of two distinct components working in concert: a relatively low energy source of electrons, and a main accelerator that adds the bulk of the beam's energy. While these two devices can be considered separately to some degree, there is significant coupling between them that must be taken into consideration in order to produce the highest quality electron beam. This dissertation investigates, theoretically, computationally, and experimentally, a new type of plasma based electron beam source mechanism called plasma density transition trapping.

This chapter begins with an examination of the limits of RF technology and the compelling reasons to overcome them. It then examines the many ways that

strong plasma oscillations can be excited with either lasers or intense electron beams. Particular attention is paid to the electron beam driven plasma wake field accelerator (PWFA) because that technique forms the basis of density transition trapping.

The review of RF and plasma based accelerators is followed by a discussion of conventional electron beam sources and their characteristics. The ways in which conventional sources are matched to RF accelerators to produce the highest quality electron beams is also examined. This matching theme is then amplified in a section detailing the challenges of injecting electron beams, especially those from conventional sources, into plasma accelerators. Finally, various types of plasma electron beam sources are discussed and the merits of plasma density transition trapping are presented.

1.1 The Need for Advanced Accelerators

While the technological development of RF accelerating structures has successfully increased accelerating gradients steadily for decades, further advances are becoming more difficult. RF accelerating structures use the conductive walls of an evacuated metallic cavity to create a set of boundary conditions for the propagation of electromagnetic radiation. These boundary conditions are designed so that electromagnetic radiation, typically microwaves, filling the cavity will have a electric field component pointing in the direction of propagation. This allows properly synchronized charged particles to co-propagate with the wave over a significant distance and gain energy from it. The larger the electric field of the radiation, the faster the particles can be accelerated. The rate at which particles are accelerated is generally referred to as the *accelerating gradient*, the energy gained by the particle per unit length.

Increasing accelerating gradients is very desirable since it makes higher beam energies practical. Technical and economic constraints limit the physical size of real accelerators so increasing the accelerating gradient is often the best option for moving to higher energy. High gradients also benefit beam quality, as is discussed later in this chapter. Unfortunately, there is a limit to the magnitude of the electric field, and hence the accelerating gradient, that can be sustained in a RF cavity. As the fields are increased in a RF cavity they become so strong at the metal surfaces that they cause the formation of arcs or sparks in a process referred to as *breakdown*. When a breakdown occurs it effectively short circuits the cavity, ruining the accelerating fields, and depositing a large amount of the cavity's stored energy in a small area of the metal surface. Breakdown must be avoided in operational accelerators since it disturbs accelerating fields and can result in irreversible damage to the structures.

The physics of breakdown are not fully understood. It is clear that breakdown is a complex process that involves field emission of electrons from the metal surface, liberation of absorbed surface gases, and the formation of plasma sheaths [1]. It is also clear that the maximum field which can be sustained in an RF cavity without breakdown increases with the frequency the cavity is operated at. Again, the origin of the frequency dependence is not entirely clear, but researchers have established empirical relations for this dependence starting with Kilpatrick in 1957 [2]. Understandably, the empirical relation has changed with advances in cavity fabrication and for modern cavities is given by

$$E_s = 220(f[\text{GHz}])^{1/3} \text{ MV/m} \quad (1.1)$$

where f is the RF frequency and E_s is the maximum sustainable surface electric field, which is about 2.5 times greater than the maximum accelerating field [1, 3].

The implication of Eq. 1.1 is that the accelerating gradient of RF cavities can

be increased indefinitely if the operating frequency is increased correspondingly. This would be the case if the technical difficulties of cavity fabrication could be ignored. As RF frequency goes up, the size of the cavities needed to utilize it goes down. The state-of-the-art cavities under development for the Next Linear Collider (NLC) operate at 11.4 GHz, and have internal apertures of only 4 mm [4]. The fabrication of thousand of metal cavities, each only a few cm in size, with complex surfaces and very high tolerances is a serious challenge with conventional machining techniques. Going to even higher frequencies and gradients will require new fabrication techniques. Since the generation of high power radiation also becomes extremely difficult as the THz range of frequencies is approached, it is likely that large improvements in accelerating gradient will require a departure from conventional RF accelerator technology.

A vast array of acceleration schemes have been proposed as replacements for RF structures. These include everything from the creation of micro-structures that would use laser radiation in the place of RF, to the inverse of every coherent radiative process: inverse-Cherenkov, inverse-synchrotron, inverse-transition radiation, etc.. The most thoroughly studied alternative to RF structure based accelerators are the schemes based on the use of plasmas.

1.2 Plasma Based Accelerators

While the magnitude of the electric fields that are sustainable in a metallic cavity are ultimately limited by electrical breakdown at the cavity walls, plasmas avoid this problem since they are broken down into free electrons and protons already. The field sustainable in a plasma is ultimately limited only by its density. In order to understand this limitation, let us examine the basic theoretical problem of a one dimensional electric wave in a plasma. Assume that the plasma is infinite,

its ions are fixed, there is no thermal motion, the excited wave is strictly one dimensional, and that there are no magnetic fields. In this limit the plasma can be treated as a fluid and described by the equation of continuity,

$$\frac{\partial n_e}{\partial t} + \vec{\nabla} \cdot (n_e \vec{v}_e) = 0, \quad (1.2)$$

the equation of motion,

$$m_e n_e \left[\frac{\partial \vec{v}_e}{\partial t} + (\vec{v}_e \cdot \vec{\nabla}) \vec{v}_e \right] = -en_e \vec{E}, \quad (1.3)$$

where the convective derivative has been used for the acceleration, and Poisson's equation,

$$\vec{\nabla} \cdot \vec{E} = 4\pi e(n_i - n_e). \quad (1.4)$$

In all of the equations n_e and n_i are the electron and ion density distributions, $\vec{v}_e$ is the electron velocity distribution, $\vec{E}$ is the electric field, m_e is the electron mass, and e is the electron charge. Now linearize the above equations by setting

$$n_e = n_0 + n_1, \quad \vec{v}_e = \vec{v}_0 + \vec{v}_1, \quad \vec{E} = \vec{E}_0 + \vec{E}_1, \quad (1.5)$$

where n_0 is large and constant, $\vec{v}_0 = \vec{E}_0 = 0$ and n_1 , $\vec{v}_1$, and $\vec{E}_1$ are small and variable. Making these substitutions and neglecting any term that is second order or higher in the small quantities n_1 , $\vec{v}_1$, and $\vec{E}_1$, equations 1.2, 1.3, and 1.4 reduce to

$$\frac{\partial n_1}{\partial t} + \vec{\nabla} \cdot (n_0 \vec{v}_1) = 0; \quad (1.6)$$

$$m_e \frac{\partial \vec{v}_1}{\partial t} = -e \vec{E}_1; \quad (1.7)$$

$$\vec{\nabla} \cdot \vec{E}_1 = -4\pi e n_1. \quad (1.8)$$

In order to solve for the wave equation, we differentiate Eq. 1.6 with respect to time to find

$$\frac{\partial^2 n_1}{\partial t^2} + n_0 \vec{\nabla} \cdot \frac{\partial \vec{v}_1}{\partial t} = 0 \quad (1.9)$$

and making substitutions from Eq. 1.7 and then Eq. 1.8 we find

$$\frac{\partial^2 n_1}{\partial t^2} + \left(\frac{4\pi e^2 n_0}{m_e} \right) n_1 = 0. \quad (1.10)$$

From this we deduce that the frequency of the oscillation of the plasma electrons, typically referred to simply as the plasma frequency ω_p , is given by

$$\omega_p^2 = \frac{4\pi e^2}{m_e} n_0. \quad (1.11)$$

The plasma frequency is one of the fundamental parameters that describes a plasma. Numerically the plasma frequency is given approximately by

$$f_p = \frac{\omega_p}{2\pi} = 8.98 \times 10^3 \sqrt{n_0} \text{ [Hz]}. \quad (1.12)$$

Since the plasmas considered for use in advanced accelerator applications typically have a density greater than 10^{12} cm^{-3} , the characteristic frequencies are in the GHz range and higher.

In order to determine the maximum field that the plasma can sustain we will add a driving term to the wave equation above,

$$\frac{\partial^2 n_1}{\partial t^2} + \omega_p^2 n_1 = -\omega_p^2 n_{beam}, \quad (1.13)$$

where the driving term is a delta function sheet beam $n_{beam} = \sigma_{beam} \delta(z - v_{beam} t)$. The solution to this equation is made much easier by switching to a frame moving with the beam by making the substitutions

$$\zeta = z - v_{beam} t \Rightarrow \frac{\partial}{\partial t} = -v_{beam} \frac{\partial}{\partial \zeta} \quad (1.14)$$

so that Eq. 1.13 becomes

$$\frac{\partial^2 n_1}{\partial \zeta^2} + k_p^2 n_1 = -k_p^2 n_{beam} \quad (1.15)$$

where $k_p = \omega_p/v_{beam}$. The homogenous solution to Eq. 1.15 is $n_1 = n_{max} \sin(k_p \zeta)$, which is valid for $\zeta \neq 0$. From Poisson's equation we then have

$$\frac{\partial E_1}{\partial x} = \frac{\partial E_1}{\partial \zeta} = -4\pi e n_1 = -4\pi e n_{max} \sin(k_p \zeta) \Rightarrow E_1 = \frac{4\pi e}{k_p} n_{max} \cos(k_p \zeta). \quad (1.16)$$

Now, in order to satisfy the assumptions made when the original set of equations was linearized, it must be the case that $n_{max} \ll n_0$ and therefore

$$|E_1| \ll \frac{4\pi e n_0}{k_p} = \frac{m_e c \omega_p}{e}. \quad (1.17)$$

However, the case in which $|E_1| = m_e c \omega_p / e$ would theoretically correspond to the maximum possible field of a 100% modulated linear plasma wave. This theoretical maximum is referred to as the *wave breaking limit*. While not a rigorously derived quantity, the wave breaking limit does give the approximate maximum field achievable in a plasma of a given density and an indication of the point where non-linear effects dominate the plasma system. Numerically, this works out to be

$$E_{wave\ breaking} = \frac{m_e c \omega_p}{e} \cong 96 \sqrt{n_0 [cm^{-3}]} \left[\frac{V}{m} \right] \quad (1.18)$$

Fig. 1.1 shows a plot of Eq. 1.18 and the plasma skin depth $k_p^{-1} = c/\omega_p$ over the range of plasma densities typically discussed for advanced accelerator applications. k_p^{-1} is an important quantity because it indicates the characteristic physical dimensions of the plasma accelerating structure. The plasma wavelength $\lambda_p = 2\pi/k_p$ can also be used as the characteristic length, but it can be misleading measure since much less than half the total plasma wave is suitable for accelerating particles. Fig. 1.1 clearly illustrates the fundamental tradeoff between accelerating gradient and accelerated volume that is common to all plasma based accelerators. Plasma accelerators attain clear superiority over conventional accelerators only when operated at gradients above 1 GV/m. Even at 1 GV/m the

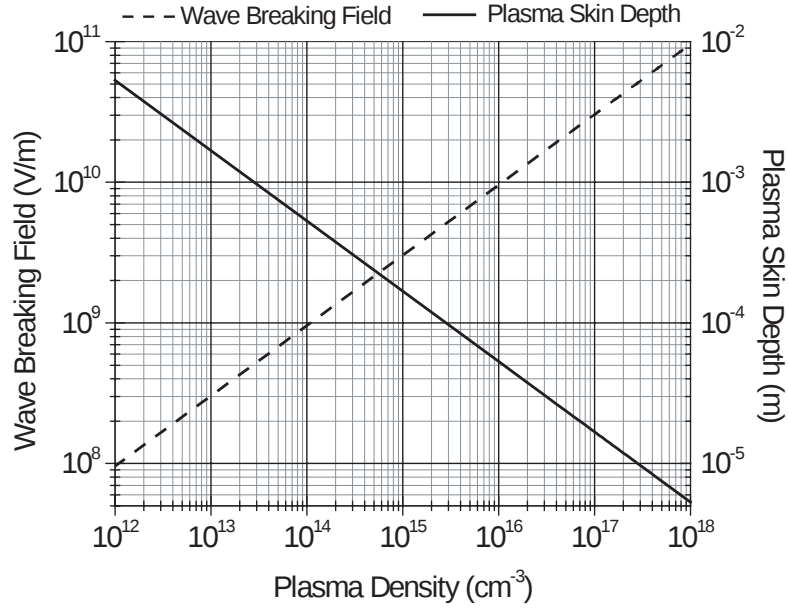


Figure 1.1: Plot of the plasma wave breaking field and skin depth over the densities of interest for advanced accelerators.

size of the plasma wave's accelerating region is well sub-millimeter which leads to serious timing and beam handling problems. These issues will be discussed at length in Section 1.4.

Plasma based accelerators fall into two categories: particle beam driven and laser beam driven. Of these two, the laser driven class has the greatest variety and will be discussed second. Beam driven designs, which are of higher relevance to this dissertation, will be discussed first and in greater depth. A introduction to this material for a general audience can be found in an article by Joshi and Katsouleas [5], and Esarey *et al.* [6] can be looked to for a review of greater technical and historical detail.

1.2.1 The Plasma Wake Field Accelerator

In a plasma wake field accelerator (PWFA) the electric fields of a short, high density electron beam are used to drive large amplitude plasma waves. To lowest order a PWFA is describe by the linear calculations in the preceding section. While real beams cannot be longitudinal delta functions, a beam shorter than a plasma period $1/\omega_p$ will effectively fulfill this condition. The basic concept of using an electron beam driven plasma wave to accelerate particles originated with Fainberg in 1956 [7]. The scenario referred to as the PWFA today was proposed by Chen *et al.* in 1985 [8]. This original paper proposed excitation of waves in the linear regime by multiple electron bunches timed to resonantly excite the wave.

A PWFA operating in the linear regime has some fairly severe limitations. One of these limitations is its *transformer ratio* R_t , which is defined as the ratio between the maximum possible energy gained by the accelerated beam to the energy of the drive beam. Note that the transformer ratio can be greater than 1 without violating energy conservation because the number of particles in the accelerated beam is assumed to be small compared to those present in the driver. Ruth *et al.* showed that $R_t \leq 2$ for a symmetric drive beam, e.g. the delta function like beam discussed above, exciting a PWFA in the linear regime [9]. This limit can be overcome if a asymmetric driving beam is used. Chen *et al.* showed that for a wedged shape driving beam, a linear rise in density ending in a step drop to zero, $R_t = \pi L_{beam}/\lambda_p$ where L_{beam} is the length of the wedge [10]. Electron beams of this shape are difficult to generate, but efforts are underway [11].

Moving from the linear regime of small plasma density oscillations to the non-linear regime of large density oscillations opens up new possibilities for the PWFA. The non-linear regime provides larger accelerating fields and longer

oscillation wavelengths than a linear excitation does at the same plasma density. Rosenzweig also showed that a non-linear one-dimensional PWFA can have $R_t \approx \sqrt{2\pi L_{beam}/\lambda_p}$ even with a symmetric beam [12].

The creation of a non-linear one-dimensional PWFA would require a driving beam whose radial dimension is much larger than k_p^{-1} . The creation of beams that large which still have the charge density to excite a non-linear response is impractical in regimes of interest. For realistic non-linear PWFA parameters the transverse beam dimension is on the order of k_p^{-1} and the system is highly two dimensional. Since no analytic theory is available for the two dimensional non-linear problem, it has been investigated numerically with simulation codes. Rosenzweig *et al.* have shown that it is highly advantageous to operate a two dimensional non-linear PWFA in the underdense, $n_b > n_0$, “blowout” regime [13]. In this regime the fields of the drive beam are so strong that they push *all* the plasma electrons out to the sides leaving only ions behind the driver. The blown out plasma electrons eventually return to the beam axis in a sharp spike of density and the envelope of their motion forms a bubble filled with only uniform density ions behind the driving beam, see Fig. 1.9 for an example. The accelerating and transverse focusing fields inside this bubble are uniform over a significant volume and form an accelerating structure reminiscent of RF based accelerators.

Many experiments have demonstrated plasma wake field acceleration in several different regimes. The first measurements were made by Rosenzweig *et al.* in 1988 [14]. This experiment was performed in the linear regime and measured wake fields of about 1.6 MeV/m in 33 cm long plasma of density $2.3 \times 10^{13} \text{ cm}^{-3}$ using a 21 MeV driving bunch of 2.1 nC with $\sigma_z = 2.4 \text{ mm}$ and $\sigma_r = 2.4 \text{ mm}$. The wake fields in this experiment were probed with a second witness electron beam that was large compared to the wake. This experiment was quickly followed

by one in the non-linear regime employing much of the same apparatus. With plasma parameters nearly identical to those above, 5.3 MeV/m wake fields were observed using a driving beam of 4 nC with $\sigma_z = 2.1$ mm and $\sigma_r = 1.4$ mm [15]. A series of experiments at the KEK facility in Japan during the early 1990s used a train of pulses to excite 30 MeV/m wake fields in plasmas of $2\text{--}8 \times 10^{12} \text{ cm}^{-3}$ [16]. The bunch train contained 6 pulses each with the parameters 10 nC, 500 MeV, $\sigma_z = 3$ mm and $\sigma_r = 1$ mm. In 1999 Barov *et al.* produced wake fields in the blowout regime with gradients of 60 MeV/m [17]. This experiment was conducted in a $1.3 \times 10^{13} \text{ cm}^{-3}$ plasma with a 18 nC beam that was 20-24 psec FWHM in length with $\sigma_r = 250 \mu\text{m}$. An ongoing series of recent experiments at the Stanford Linear Accelerator Center (SLAC) have observed accelerating gradients on the order of 150 MeV/m [18]. The nominal parameters for this experiment are a 1.4 m long $2 \times 10^{14} \text{ cm}^{-3}$ plasma driven by a 30 GeV, 3.2 nC electron beam with $\sigma_z = 0.65$ mm and $\sigma_r = 50 - 100 \mu\text{m}$ [19]. A positron driven PWFA has even been demonstrated [20].

1.2.2 The Laser Wake Field Accelerator

Proposed by Tajima and Dawson in 1979 [21], the laser wake field accelerator (LWFA) is conceptually very similar to the PWFA. In a LWFA, plasma electrons are blown out to the sides by the nonlinear ponderomotive forces [22] created by an intense laser pulse rather than the electric field produced by an intense electron beam. Just as with the electron pulse in a PWFA, the driving laser pulse needs to be about $1/\omega_p$ long to excite the plasma wave optimally. Accelerating gradients of about 1.5 GV/m have been observed in a LWFA operated at a plasma density of $2 \times 10^{16} \text{ cm}^{-3}$ with a 3.5 TW peak power, $1.057 \mu\text{m}$ wavelength laser that produced a maximum intensity of $4 \times 10^{17} \text{ W/cm}^2$ in a pulse 400 fs long at the full width

at half maximum (FWHM) [23].

1.2.3 The Plasma Beat-Wave Accelerator

The plasma beat-wave accelerator (PBWA) was proposed by Tajima and Dawson at the same time as the LWFA [21]. This accelerator operates by using two laser pulses, both long compared to the plasma wavelength λ_p , with frequencies ω_1 and ω_2 chosen so that $\omega_1 - \omega_2 = \omega_p$. The excitation of a plasma wave by beating two laser beams had been suggested earlier in the context of plasma heating for controlled nuclear fusion [24]. Once again the ponderomotive forces produced by the laser push out the plasma electrons and produce an electric plasma wave. In this case, however, the wave is driven resonantly rather than impulsively.

Plasma waves generated by the beat-wave mechanism were first observed in 1985 [25] using a CO₂ operating at two spectral lines, 9.56 and 10.95 μm , in a plasma of about 10^{17} cm^{-3} . The beat-waves were diagnosed through Thomson scattering and found to have peak fields of up to 1 GV/m. This experiment was followed by many others in several research groups. Recently, a photoinjector produced electron beam has been injected into a beat-wave accelerator [26].

1.2.4 The Self-Modulated Laser Wake Field Accelerator

The self-modulated laser wake field accelerator (SMLWFA) is similar to the PBWA in that it uses a laser pulse that is long compared to the plasma wavelength λ_p . The SMLWFA uses a mono-chromatic laser pulse, however, which becomes modulated into a beat-wave like structure through its interaction with the plasma. The mechanism through which the modulation is produced is referred to as the Raman forward-scattering instability. While the SMLWFA has been extensively studied and is one of easiest plasma accelerators to create from a

technical standpoint, its use as an advanced accelerator concept is limited. Since the accelerating fields in the SMLWFA grow from an instability, their phases cannot be controlled. Without any type of phase control, the SMLWFA cannot be used as a normal accelerator. In some cases, however, the fields in SMLWFAs are high enough to capture and accelerate background plasma electrons. In this mode the SMLWFA operates as a single stage electron source and accelerator. Several experiments have produced high energy electrons this way [27].

1.3 Conventional Electron Beam Sources

Before we examine the relative merits of real electron beam sources, it is useful to define the characteristics of an ideal electron beam source and the parameters used to describe an electron beam's quality. The most common parameter used to quantify the transverse quality of an electron beam is the *emittance*, which is typically denoted by ε . The emittance is defined as the area that the beam occupies in trace space, the space of the parameters x and $x' = dx/dz$. Three simplified pictures of a beam's trace space at different points in a beamline are shown in Fig. 1.2. If the area of the ellipse is A , then the emittance is defined as $A \equiv \pi\varepsilon_x$. There are many choices that can be made regarding where to define the boundary of the trace space ellipse. For that reason, the root-mean-squared (rms) emittance is typically used and is defined statistically as,

$$\varepsilon_{x,rms}^2 = \langle x^2 \rangle \langle x'^2 \rangle - \langle xx' \rangle^2. \quad (1.19)$$

In the equilibrium case where $\langle xx' \rangle = 0$ this reduces to $\varepsilon_{x,rms} = \sqrt{\langle x^2 \rangle \langle x'^2 \rangle} = \sigma_x \sigma_{x'}$ which, when multiplied by π , is the area of the rms trace space ellipse.

The significance of the emittance lies in the fact that it remains constant as

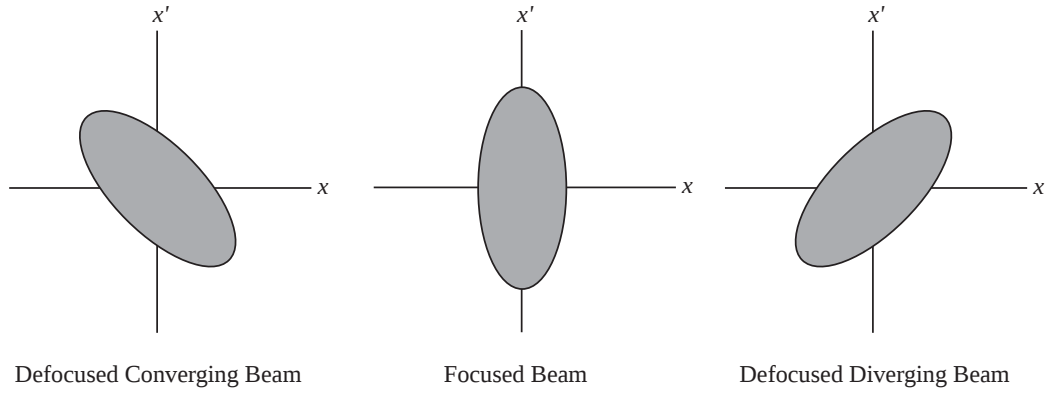


Figure 1.2: Area occupied by the electron beam in trace space at three different locations along a beamline.

an electron beam is transported down a beamline under linear forces, as can be shown from Liouville's theorem. Although the envelope of the electron beam in trace space can be manipulated, e.g. made small in the x dimension and large in the x' when the beam is focused down as in the center plot in Fig 1.2, its overall area is fixed. It follows that beams of low emittance are desirable since they can be focused to smaller spots with smaller angles than larger emittance beams. Smaller beam spots over larger distances lead to higher energy densities, more interactions in particle colliders, and performance enhancement in virtually every demanding particle beam application.

It should be noted that the emittance is not invariant under the action of non-linear forces on the beam. Liouville's theorem states that the area in phase space, the space of the variables x and p_x , corresponding to the motion of a system of particles does not change during the motion. This remains true even under non-linear forces. While the total area in phase space remains constant, however, it can be stretched and filamented by non-linear forces to the point where the beam effectively occupies a larger portion of phase space than its constant area

would suggest. Under such conditions, the emittance as defined in Eq. 1.19 will grow to reflect the effective area of beam in trace space, which is the relevant parameter when one considers beam performance. Emittance growth due to non-linear forces will be discussed at length in the following sections.

The emittance defined in Eq. 1.19 is referred to as the geometric emittance. The normalized emittance is defined as

$$\varepsilon_{norm,x,rms}^2 = (\beta\gamma)^2 [\langle x^2 \rangle \langle x'^2 \rangle - \langle xx' \rangle^2] = (m_0c)^{-2} [\langle x^2 \rangle \langle p_x^2 \rangle - \langle xp_x \rangle^2] \quad (1.20)$$

where $\beta = v/c$ and $\gamma = 1/\sqrt{1-\beta^2}$. Normalized emittance is the preferred quantity typically used because, unlike the geometric emittance, it is invariant under longitudinal acceleration as well as linear transverse focusing forces. This is clear from the fact that $p_x = \gamma m_0 v_x$ is invariant under acceleration in the z direction while $x' = v_x/v_z \cong v_x/c$ decreases since v_x must decrease as the beam γ increases during acceleration. Normalized emittance will be the quantity cited whenever the term emittance is used throughout the rest of the text. A more in depth discussion of emittance can be found in Refs.[28, 29].

A quantity directly related to emittance which is also often used is the beam *brightness* [30]. Brightness combines the emittance and the average current I of the beam into a single figure of merit which is defined as,

$$B = \frac{I}{\varepsilon_{n,x} \varepsilon_{n,y}}. \quad (1.21)$$

Brightness is useful because it indicates the volume density of particles that can be produced at the focus of a particle beam. Particle density is a critical parameter in many high energy density applications including colliders, Compton x-ray sources, and free electron lasers.

An electron beam's emittance is, for the most part, a quantity whose minimum is fixed at the time of its creation. Therefore, emittance is determined by the

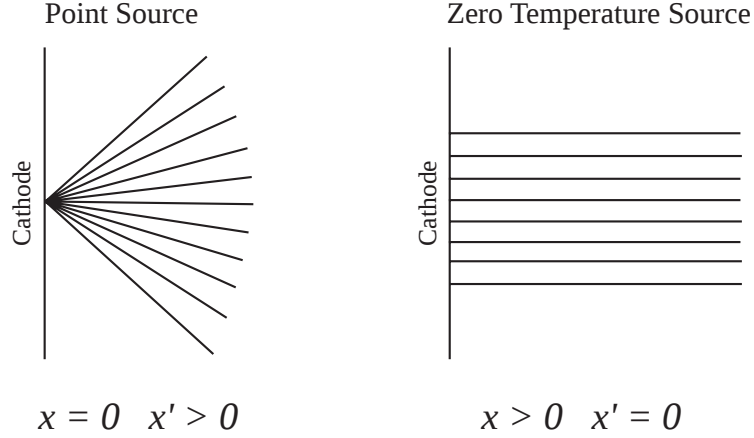


Figure 1.3: Idealized electron beam sources.

nature of the source used to create the electron beam. An ideal electron beam source, at least in terms of transverse dynamics, would create beams with zero emittance. Conceptually, this could be accomplished in either of the two ways depicted in Fig. 1.3. In the first case the beam starts from a true point source so that $x = 0$ for all particles. In the second case the beam originates from a finite source but the electrons have absolutely no transverse momenta so that $x' = 0$. A source of ideal brightness could be made from either case if an arbitrarily large amount of charge could be extracted from the source in an infinitesimal time. While a real systems can never achieve either of these ideal behaviors, electron beam sources try to emulate them as closely as possible while taking into account other realistic effects that can lead to increased emittances.

1.3.1 Thermionic Direct Current Guns

Direct current thermionic guns are one of the oldest and by far the most ubiquitous type of electron beam source. In various forms, they provide the electron beams for particle accelerators, microwave sources such as klystrons, television

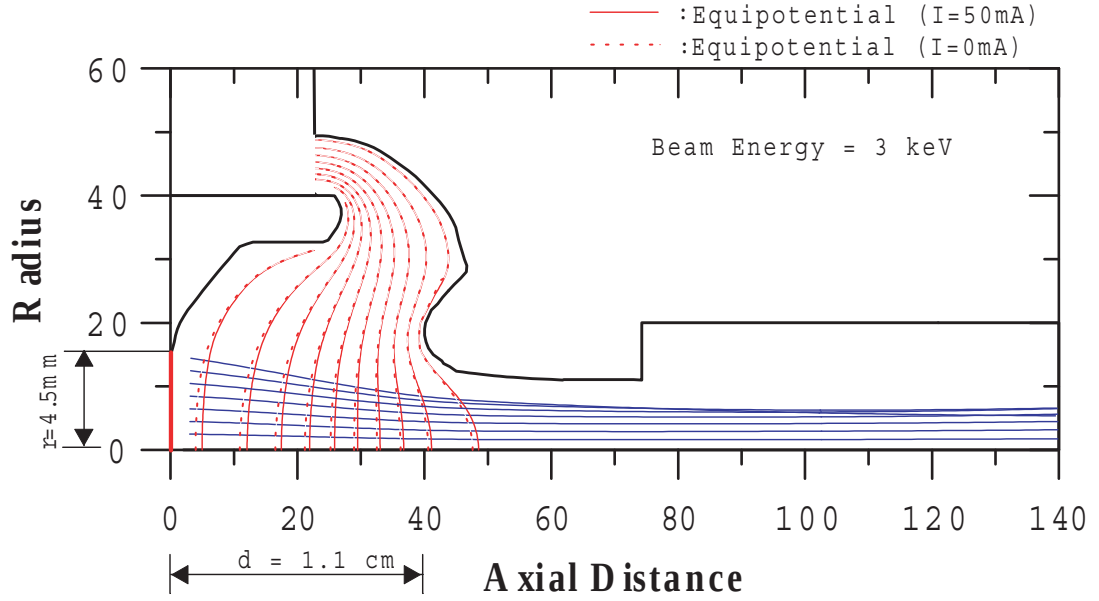


Figure 1.4: Schematic, in r and z , of a cylindrically symmetric Pierce style electron gun. This gun was designed by H. Suk, following the design of similar guns used at the University of Maryland model electron storage ring [32], for an experiment examining electron beam cleaning of photo-cathodes [33]. The dashed lines are equipotentials of the electric field formed by the shaped anode and cathode structures. The fields are designed to provide focusing of the electron beam (solid lines) as it is emitted from the thermionic cathode.

cathode ray tubes (CRTs), and plasma sources.

Thermionic guns are very simple in principle. In general, all that is required is a electron emitter held at potential difference with respect to a grid or hollow conductor toward which the electrons will be accelerated and then pass through. Typically, the cathode and anode are shaped to provide transverse electric field components that will focus the beam and counteract its tendency to expand due to the mutual repulsion of its electrons, which is usually referred to as the *space charge* force acting on the beam. This class of electron guns are referred to as *Pierce-type guns* [31], an example of which is illustrated in Fig. 1.4.

While it is possible to emit electrons directly from a cold metal surface through

the mechanism of Fowler-Nordheim field emission [34], the fields required for the emission of significant charge are very large. The possibility of constructing a Fowler-Nordheim field emission source using the high fields produced in plasma accelerator is discussed in Section 8.2.4. For most DC guns, however, thermionic emitters are a much more practical option. Thermionic emitters work by raising the temperature of a metal and hence its conduction electrons. Since the conduction electrons obey Fermi-Dirac statistics, raising the temperature increases the population of electrons in the Maxwellian tail of the electron energy distribution that can overcome the work function of the metal. The maximum current density $J_{thermal}$ that can be drawn from thermionic emitter is given by the Richardson-Dushman equation [35, 36]:

$$J_{thermal} = AT^2 e^{-\frac{W}{k_b T}}, \quad (1.22)$$

where T is the temperature in Kelvin, W is the cathode work function, k_b is Boltzmann's constant, and $A = 1.2 \times 10^6 \text{ Amp}^{-2} \text{K}^{-2}$ is a derived constant.

Another more universal limit on the current density that can be produced by DC guns comes from the space charge of the beam itself. As the beam emerges from the cathode, its space charge, along with the beam's space charge image in the cathode, creates an electric field that opposes the accelerating electric field. This process ultimately limits the current that can be extracted from a gun at a given voltage. The relationship between the maximum current I_{max} and voltage V is given by the Child-Langmuir law [37, 38]:

$$I_{max} = KV^{3/2}, \quad (1.23)$$

where K is the perveance, a quantity that is measured or derived from the geometry of the gun. In practice, voltage can usually be increased to the point where gun emission is limited by the Richardson-Dushman equation. The effects

Table 1.1: Nominal operating parameters of a high-current-density gun constructed with a LaB₆ cathode for a single-stage microwave FEL experiment at the National Laboratory for High Energy Physics in Japan (KEK) [39]

Cathode Temperature	1900 K	Beam Current	11 A
Cathode Diameter	12 mm	Emittance	72 mm-mrad
Current Density	10 A/cm ²	Brightness	2 x 10 ⁹ A/(m-rad) ²
Anode Voltage	40 kV	Pulse Width	250 ns
Gradient at Injection	< 2 MV/m	Repetition Rate	20 Hz

of space charge on beam longitudinal dynamics are still significant, however, even below the Child-Langmuir limit as will be shown later in this chapter.

A thermionic DC gun is clearly the opposite of the ideal zero temperature cathode presented in Fig 1.3. From the equipartition theorem of thermodynamics the rms velocity of the beam particles in the transverse direction is given by

$$\sigma_{v_x} = \sigma_{v_y} = \sqrt{\frac{k_b T}{m_e}}. \quad (1.24)$$

If the electron emitter is circular of radius r and provides uniform emission then

$$\sigma_x = \sigma_y = \frac{r}{2}, \quad (1.25)$$

and, assuming $v_z \approx c$, the intrinsic geometric emittance of a thermionic gun is

$$\varepsilon_{rms, thermionic} = \frac{r}{2c} \sqrt{\frac{k_b T}{m_e}}. \quad (1.26)$$

While the high temperature of a thermionic source is unavoidable, such sources can be made to resemble the point source in Fig. 1.3 if r is made very small. This is the technique used to produce high quality beams in electron microscopes, but it is only suitable for low current applications.

Table 1.1 gives the parameters associated with an example of a modern thermionic DC gun suitable for use in a high energy accelerator. As we shall

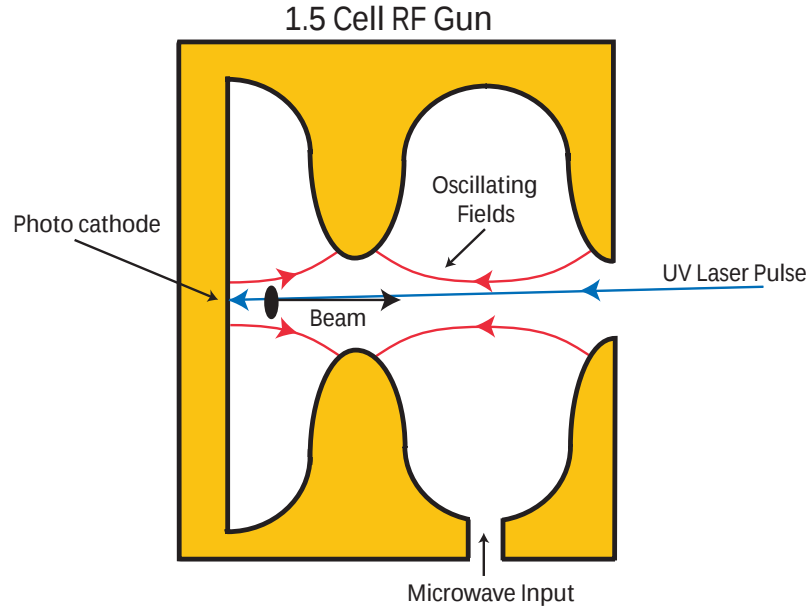


Figure 1.5: Simplified diagram of an RF photoinjector

see in the next section, most of these parameters are inferior to those achievable with photoinjectors. The one exception to this rule is the amount of charge the thermionic gun can produce ($55 \mu\text{C}$ per second in this case), which is orders of magnitude greater than that produced by most photoinjectors. This gives thermionic guns an enduring edge in some high average power applications.

1.3.2 Radio Frequency Photoinjectors

RF photoinjectors are the current state-of-the-art in the production of high quality electron beams. They achieve superiority over thermionic sources by using different methods of both acceleration and electron emission. Accelerating fields are provided by high power microwaves rather than pulsed DC high voltage. Electron emission is achieved through the photo-electric effect rather than thermionic emission.

The schematic of a typical copper photoinjector is shown in Fig. 1.5. The half-cell geometry of the cell in which the photocathode is placed ensures that the accelerating fields will be perpendicular to the cathode. The full cell provides additional acceleration like any other RF structure. Photocathodes are typically made from a durable material with a small work function like copper, magnesium, or cesium telluride [40]. A UV laser is shone on the photocathode from a downstream mirror that is slightly off the beam axis. The laser pulse is short compared the RF period and is timed to arrive at the optimum point in the electric field oscillation for electron beam injection.

RF photoinjectors are subject to the limitations of breakdown as discussed in Sec. 1.1, just like other RF accelerating cavities. The threshold of RF breakdown is, however, much higher than the threshold of DC breakdown. For that reason photoinjectors sustain injection gradients well over an order of magnitude higher than thermionic guns, see Tables 1.2 and 1.1. The injection gradient is important because the effects of beam space charge forces, which inevitably lead to emittance growth, damps like $1/\gamma^2$ [29]. Therefore, the faster the electrons can be accelerated to high energy, the smaller the space charge induced emittance growth.

Somewhat surprisingly, the temperature of the electrons emitted from the photocathode are comparable to that of those emitted from thermionic cathodes. Although photocathodes are typically at room temperature, the photon energy used needs to be larger than the work function to ensure good emission. The temperature of electrons emerging from a photoinjector cathode has not been measured or calculated with high precision but should be in the range of 0.2 to 1.3 eV [42]. For comparison 0.16 eV is the average energy of the electrons in equilibrium with the 1900 K thermionic cathode. The advantage of photoemis-

Table 1.2: Nominal operating parameters of the UCLA Neptune laboratory S-band (2856 MHz) photoinjector [41].

Laser Spot on Cathode σ_r	1 mm
Average Current Density	$\sim 3000 \text{ A/cm}^2$
Gradient at Injection	60 MV/m
Beam Charge	600 pC
Emittance	5 mm-mrad
Brightness	$4 \times 10^{12} \text{ A/(m-rad)}^2$
Pulse Width σ_t	3.6 ps
Repetition Rate	1 Hz

sion does not, therefore, lie in the temperature of the electrons, but rather their quantity. A photocathode is not limited by the Richardson-Dushman equation since as many electrons can be emitted as there are photons available. Practically, this results in photoinjectors producing current densities that are two orders of magnitude larger than thermionic sources, see Tables 1.1 and 1.2. This is a big step closer to the ideal point electron source of Fig. 1.3. The effective size of the photoinjector source is equal to the laser spot ($\sigma_r = 1 \text{ mm}$), which is 6 times smaller than the size of the example thermionic source described in Section 1.3.1. Thanks to the high current density, however, this smaller source produces about 100 Amps, an order of magnitude more than the larger thermionic cathode. The small spot size and strong accelerating fields lower emittance, while the photo-emission produces more current. These simultaneous gains in emittance and current push photoinjector brightness three orders of magnitude above the thermionic source.

As stated above, the UV laser pulse used to photo-emit the electron beam is short, several ps, compared to period of the RF, 350 ps, see Table 1.2. It is necessary to use only a small portion of the oscillating RF wave to ensure that the electron beam is accelerated uniformly. This limitation on bunch length makes it difficult for photoinjectors to compete with thermionic guns on the basis of total delivered charge.

The time during the RF cycle when the laser pulse arrives at the cathode is equally as important as the laser pulse's length. This is referred to as the *injection* time, which is often expressed in terms of the phase ϕ of the RF wave. Naively, it seems like injection should take place at the peak of the wave, equivalent to $\phi = 90^\circ$ if we choose to describe the accelerating wave as a sine function. The reality is, however, that the electrons take time to accelerate from rest to relativistic velocity. Therefore, if we want the electrons to be accelerated out of the photoinjector at the peak field ($\phi = 90^\circ$) they must be injected *before* the peak ($\phi < 90^\circ$) to compensate for the slippage that occurs while they are accelerating. Operationally, the optimum injection phase depends on many variables but is generally in the neighborhood of $\phi = 45^\circ$. Interestingly, the phase slippage that occurs at injection has an impact on the length of the electron beam. Since the RF fields are varying in time, the head of the electron beam will arrive earlier and at a lower gradient than the tail $\phi_{head} < \phi_{tail}$. Since it takes longer for the head to reach relativistic velocity than the tail, the beam will compress slightly in time, see Fig 1.6. The process is referred to a *velocity focusing* and its effect can be describe by

$$\frac{\Delta\phi_f}{\Delta\phi_i} = \sin \phi_{inject} \quad (1.27)$$

where $\Delta\phi_i$ and $\Delta\phi_f$ are the initial and final extent of the beam in phase, respectively, and ϕ_{inject} is the injection phase, see Section 2.1 and Ref. [43]. It is

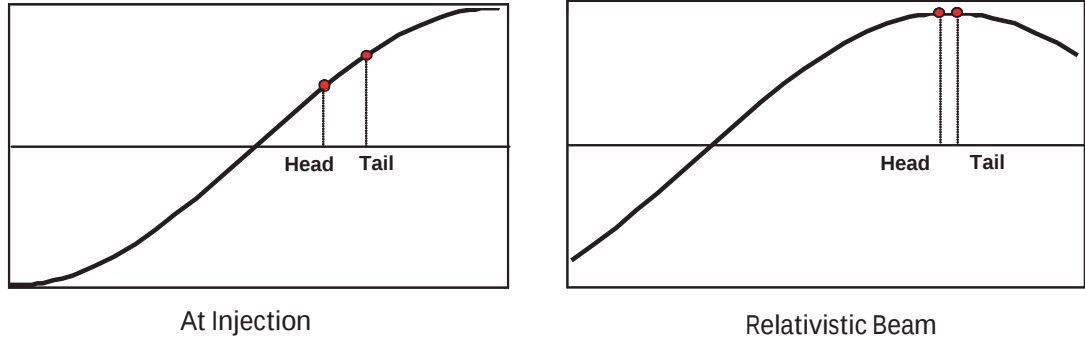


Figure 1.6: The mechanism through which beams are longitudinally focused in photoinjectors by differential phase slippage.

assumed in the derivation of this equation that the final centroid phase of the beam is 90° and that the only force acting on the beam is the accelerating wave. For a typical injection phase of 45° , Eq. 1.27 indicates that the electron beam exiting the photoinjector will be 70% of its length at injection.

Eq. 1.27 has a serious flaw in that it neglects the effects of space charge. Just as in the case of thermionic guns, the space charge of the beam exiting the cathode, along with its image charge in the cathode, creates an electric field that opposes the accelerating fields of the RF. While the currents involved are below the Child-Langmuir limit, Eq. 1.23, the space charge forces at injection are strong enough to cause the beam to expand longitudinally, counteracting the velocity focusing. Detailed analytical and numerical models of these longitudinal photoinjector dynamics have shown that the quantity $\Delta\phi_f/\Delta\phi_i$ can vary from 0.7 to 1.2 or more depending on launch phase, beam charge, laser spot size, gradient, etc. [43]. While the details of longitudinal photoinjector dynamics lead to minor variations in the beam current, and therefore brightness, it is possible to use a much more aggressive form of velocity focusing called *magnetic compression* to significantly boost the beam current after it has left the photoinjector.

1.3.3 Magnetic Compression

Although the current density and emittance produced by photoinjectors is a large improvement over thermionic sources, photoinjectors still fall short of the demands of high performance applications such as short wavelength FELs. Magnetic compression allows the current of a beam to be increased substantially, at the cost of some emittance growth, after it has left the photoinjector and been accelerated to high energy. Typically, magnetic compression can reduce the rms beam length by a factor of 4 [41, 44].

A magnetic compression system can be thought of simply as a longitudinally focusing lens. In a transverse lens the transverse position of the beam particles is correlated to the size of the transverse kick they receive so that the beam will focus to a point downstream of the lens. Similarly, a longitudinal lens must produce a correlation between the position along the length of a beam and momentum so that the beam will shrink longitudinally as it propagates downstream. A longitudinal “lens” is usually an RF cavity through which the beam travels off crest of the sinusoidal accelerating wave so that tail of the beam is accelerated to significantly higher energy than the head, as shown in the center frame of Fig. 1.7. As the beam propagates, the tail will catch up with the head and the beam will focus longitudinally. This type of longitudinal focusing is referred to as velocity bunching [45]. For velocity bunching to be effective however, it must be done at low energy (or on a very long beam line) so that the velocity difference between head and tail is significant. Unfortunately, velocity bunching at low energy is not advisable for high charge beams, due to space charge induced emittance growth, and long beam lines are frequently impractical. Therefore, a series of magnetic bends is used to create a difference in path length for the particles of different momenta. In this way a beam with a momentum slope can be compressed even if

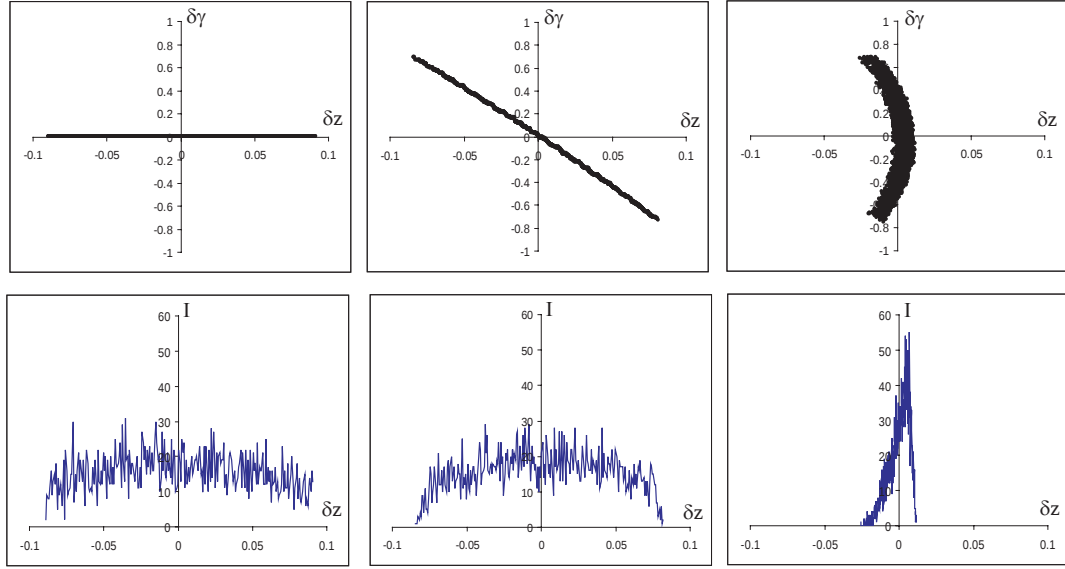


Figure 1.7: The mechanism of magnetic compression. The longitudinal phase space (top) and current profile (bottom) are shown for the beam out of the gun (left), after the linac (center), and after the chicane (right). The beam is propagating to the right in each case. δz is in units of cm and I is in arbitrary units. The deleterious effects of space charge and other collective processes are ignored.

the difference in velocity between the head and tail is minute. This combined use of an RF cavity and magnetic bend system to compress a beam longitudinally is what is referred to as magnetic compression.

Magnetic pulse compression has been studied thoroughly as a means of enhancing beam current and brightness [46]. Mathematically the process of magnetic compression can be described simply as a set of transformations on the longitudinal phase space of the beam. The effect of the linac on the beam is given by

$$\delta p_{linac} = \delta p_0 + \alpha E(\sin(\phi + k\delta z) - \sin \phi); \quad (1.28)$$

$$\delta z_{linac} = \delta z_0. \quad (1.29)$$

where δp_0 and δz_0 are the initial momenta and z offset of each particle relative to

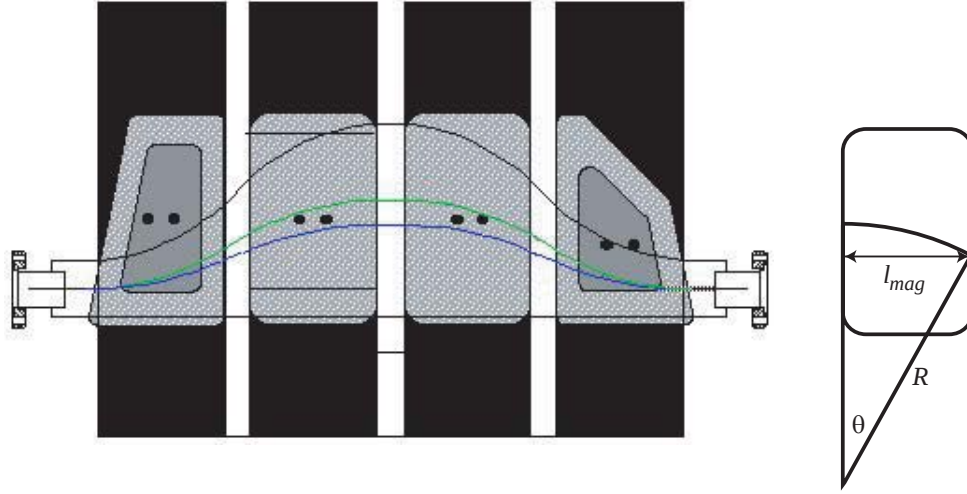


Figure 1.8: A typical set of chicane magnets.

the bunch center, δp_{linac} and δz_{linac} are the final values of these quantities, αE is the peak momentum gain of a particle travelling through the linac on crest, k is the wave number of the accelerating RF, and ϕ is the RF phase of the center of the beam. This is the set of transformations that takes the beam from frame one to frame two in Fig. 1.7.

While there are different types of magnetic bend systems that will produce compression [11], by far the most common configuration is the chicane. A chicane is usually a set of four dipole magnets operated at the same field with the first and last magnets operated with opposite polarity with respect to the middle two, see Fig. 1.8. As the beam passes through the chicane it bends out to the side and then back to its original trajectory. If the small gaps between magnets are ignored, the geometry of the system ensures that the path length of the electron trajectory through the chicane S is given by

$$S = 4R\theta = 4 \left[\frac{p}{eB} \right] \sin^{-1} \left[l_{mag} \frac{eB}{p} \right] \quad (1.30)$$

where R is the radius of curvature of the particle in the dipole field, θ is bend

angle, and l_{mag} is the length of the magnet as shown in Fig. 1.8. Substitutions were made using the equation $p = qBR$ where q is the particle charge and B is the magnetic field. Differentiation of Eq. 1.31 with respect to momentum p leads to the equation for variation of the path length with changes in momentum:

$$\frac{\delta S}{S} = \left[1 - \frac{\tan \theta}{\theta} \right] \frac{\delta p}{p} \equiv \alpha_c \frac{\delta p}{p}, \quad (1.31)$$

where α_c has been defined as the path length parameter. Given this we can define a set of transformations for the chicane that are analogous to Eqs. 1.28 and 1.29 for the linac

$$\delta p = \delta p_0; \quad (1.32)$$

$$\delta z = \delta z_0 + S_0 \alpha_c \frac{\delta p_0}{p_0}. \quad (1.33)$$

When the signs of α_c and δp_0 are chosen correctly, this transformation produces beam compression as shown in the third frame of Fig. 1.7. Note that the compression is limited by the non-linearity of the correlated momentum spread induced on the beam by the linac. The curvature impressed on the beam by the sine function in Eq. 1.28 manifests itself as the comma shaped phase space on the right in Fig. 1.7.

Unfortunately, magnetic compression also has a negative effect on the traverse phase space of the beam. At lower energies space charge forces cause emittance growth during the compression process [47], while at higher energies coherent synchrotron radiation (CSR) produces the same result [48]. Careful design can mitigate these effects and significant brightness gains can be made despite the emittance growth.

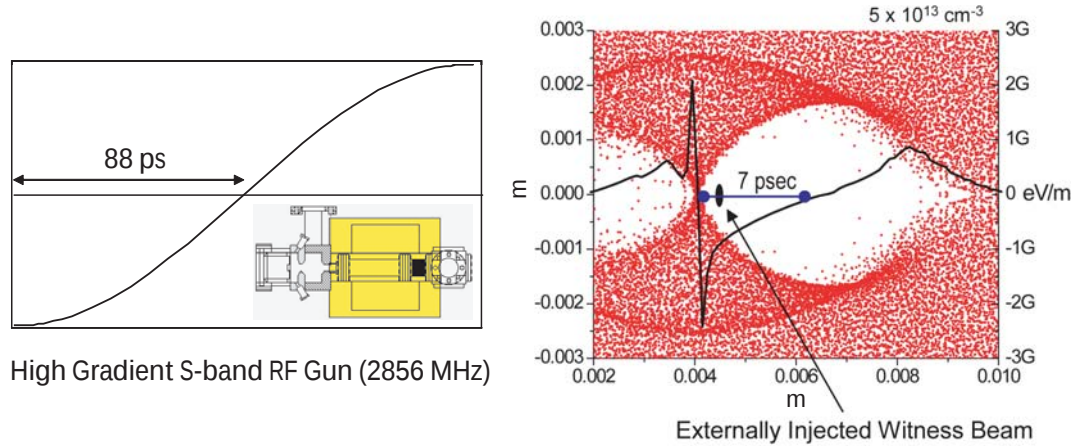


Figure 1.9: Illustration of the difficulties of injecting an external beam into a PWFA as opposed to a S-band RF Gun. At left is the length in picoseconds of the $\lambda/4$ portion of the 2856 MHz RF wave that is suitable for electron injection. At right is an example of a PWFA operated in the highly non-linear blowout regime at $5 \times 10^{13} \text{ cm}^{-3}$ with a frequency of about 75 GHz. The particles shown in the illustration are the plasma electrons. The length of the PWFA electron beam injection window, analogous to that of the RF gun, is indicated.

1.4 Injection and Timing

Now that the conventional mechanisms of beam production and manipulation have been discussed, it is interesting to see how these systems compare to the demands of advanced accelerators. Foreshadowed throughout the preceding sections is the difficulty of coordinating an electron pulse with the plasma wave meant to accelerate it. This is generally referred to as the problem of injection. The nature of this problem is illustrated in Fig. 1.9. The length of time during which an electron pulse can be successfully injected into a typical RF gun is about an order of magnitude longer than the time available for the same purpose in even a modest gradient plasma accelerator. An electron pulse of a few ps length can be accelerated in the RF gun without much energy spread thanks to the fairly slow variation of its the fields on the ps time scale. Likewise, if there

is a timing instability of a few ps between the electron pulse, or equivalently the cathode laser, and the RF, the resulting shot-to-shot energy jitter will be fairly small. By comparison, the accelerating fields of our example plasma accelerator in Fig. 1.9 vary significantly on the ps time scale. The injected electron bunch must be sub-ps to retain a small energy spread and the timing instability between the bunch and the plasma accelerator must be sub-ps as well or the energy jitter will be enormous. As Fig. 1.1 indicates, this problem gets even worse for higher gradient plasma accelerators. The technical difficulty of these beam length and jitter requirements is the reason all the experiments to date that have injected external electrons into accelerating plasma waves have used either continuous electron beams or beam pulses that were long compared to the plasma wave [14, 17, 23, 49]. Progress is ongoing, however, and some recent injection efforts with PBWAs have succeed in injecting pulses that are roughly the same size as the accelerating wave [26].

It is also important to remember that the plasma skin depth values shown in Fig. 1.1 set the scale length for *all* the dimensions of the plasma accelerating structure. This means that not only do injectable electron beams need to be sub-ps in length, they need to be focused to spots on the order of several hundred μm or less. Given the trade off between beam length and emittance inherent in magnetic compression, this is a difficult thing to achieve with photoinjector technology. Another difficulty multiplier in the problem of injection is the fact that for most high energy applications, an electron bunch will have to be accelerated multiple times in successive plasma accelerators, just as conventional RF accelerators use multiple sets of cavities. The use of multiple accelerating structures is referred to as staging.

Whether driven by lasers directly, or by electron beams that originate from

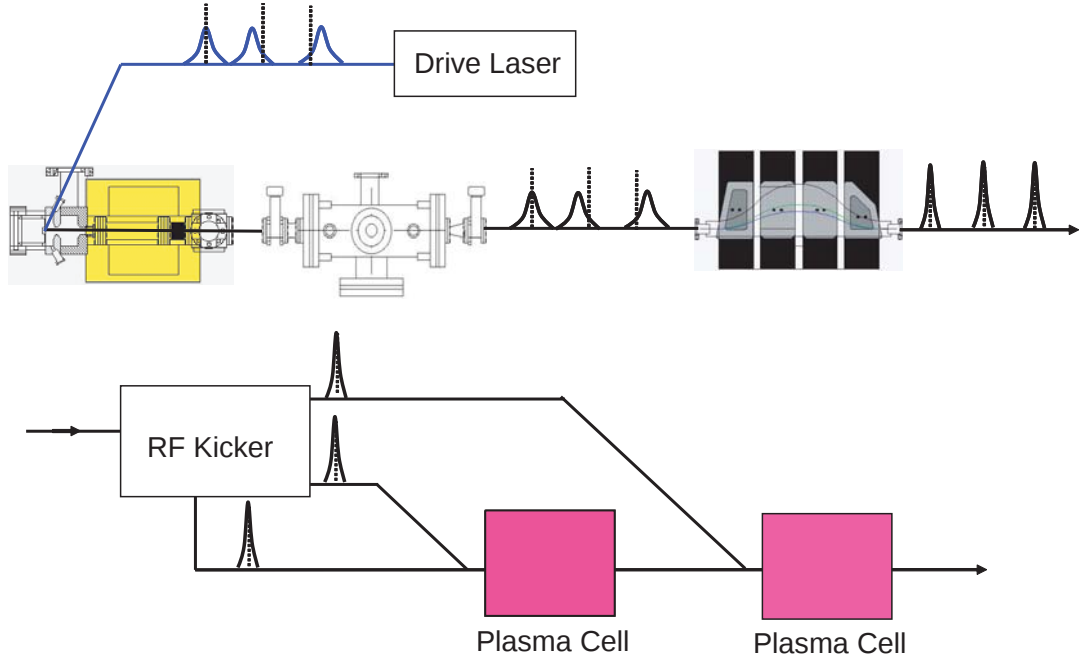


Figure 1.10: Cathode drive laser jitter suppression using magnetic compression of the electron beam.

lasers shining on photocathodes, all the plasma advanced accelerator schemes rely on high power laser pulses. The current technology of high power pulsed lasers cannot reliably provide sub-ps timing accuracy [50]. There are techniques, however, that might allow successful plasma accelerator injection without the need to reduce laser timing jitter to sub-ps levels. In systems where the electron pulse to be injected is provided by a photoinjector, there are a set of schemes based on the time stability of the RF that can be used to circumvent the laser timing jitter [43].

The most obvious solution to the problem of the laser timing jitter is to add some mechanism to the system that will suppress the jitter. Magnetic compression can provide such a jitter suppression mechanism. We can see from Fig. 1.7 that the effect of maximal magnetic compression, ignoring RF curvature and

other higher order effects, is to map any particle which deviates from the centroid phase of the bunch, $\delta\phi_{initial} = k\delta z$, back to the center phase, $\delta\phi_{final} = 0$. In exactly the same way, a bunch whose centroid is displaced in time due to jitter, $\delta\phi_{jitter} = \omega\delta t$, is mapped back to the design phase, $\delta\phi_{final} = 0$. In this mode of compressor operation, initial timing/phase errors do not result in final phase errors,

$$\frac{\delta\phi_f}{\delta\phi_i} = 0. \quad (1.34)$$

This effect can be used in a PWFA to help solve the problems of both injection and staging, as shown in Fig. 1.10. Jitters introduced by the cathode drive laser are suppressed by the chicane, and pulses become reliably locked to the same phase of the RF. Once the pulse to pulse timing is stabilized, the pulses can be separated in an RF kicker and transport path length differences can be used to properly phase the accelerated bunch in multiple PWFA stages.

The exact opposite approach, that of preserving the jitter throughout the system rather than suppressing it, can also be useful. This technique is under consideration for PBWA injection [43]. The idea is use the beat wave laser itself to imprint the beat wave pattern on the cathode drive laser. The beam produced at the cathode will correspondingly be a train of micro-pulses phased perfectly for injection into the plasma beat wave. The longitudinal effects of the gun can be compensated for with appropriate setting of the magnetic compressor so that the whole system behaves according to the equation

$$\frac{\delta\phi_f}{\delta\phi_i} = 1. \quad (1.35)$$

In this way the initial beat wave phase that was impressed on the beam is preserved and electron micro-pulses simply jitter with the beat wave. While this effect insures successful injection into one beat wave structure it is difficult to

stage. There are also a host of difficult technical issues regarding the imprinting of the beat wave on the drive laser using nonlinear optics techniques.

1.5 Plasma Based Electron Beam Sources

The problems surrounding injection of photoinjector produced beams into plasma accelerators, and the desire for ever higher brightness electron beam sources, have lead to the consideration of plasma based electron beam sources. Plasma based electron beam injectors follow the trends established in the move from thermionic source to photoinjectors. With electron temperatures of upwards of 7 eV (equivalent to 81,000 K), see Chapter 5, plasmas are even farther from the zero temperature source of Fig. 1.3 than either of the other source types already described. Plasma sources at high densities can, however, provide effective source sizes orders of magnitude smaller than photoinjectors. As shown in Table 1.3, a high density plasma source can also provide higher gradients, higher current densities, and shorter pulses than a photoinjector. As with the step from thermionic source to photoinjectors, this increased performance comes at the price of a reduction in total beam charge. Perhaps the greatest advantage of plasma electron beam sources, at least from the plasma accelerator point of view, is that they directly provide short, low-emittance, beams suitable for injection into plasma accelerators. Unfortunately, the switch to a plasma based electron beam source does not intrinsically solve the timing problems with injection or staging.

As was the case with plasma based accelerators, there are a host of plasma based beam source schemes. The great promise of plasma electron beam source has inspired the creativity and determination of many researchers. Once again, a brief account will be made of the major types of these source schemes that are being studied. The greatest attention will be given to plasma density transition

Table 1.3: Theoretical operating parameters of a high brightness plasma based electron beam injector. These values are based on simulations of a $2 \times 10^{17} \text{ cm}^{-3}$ plasma density transition trapping source. See Section 4.2 and Table 4.1 for further details.

Region Generating Beam σ_r	10 μm
Average Current Density	$\sim 3 \times 10^7 \text{ A/cm}^2$
Gradient at Injection	$\sim 10 \text{ GV/m}$
Beam Charge	12 pC
Emittance	0.6 mm-mrad
Brightness	$5 \times 10^{14} \text{ A/(m-rad)}^2$
Pulse Width σ_t	28 fs

trapping, the focus of this dissertation.

1.5.1 Random Phase Injection Sources

If one is not concerned with beam quality, there are many ways to produce energetic electrons from a plasma. Even the fundamental plasma process of Landau damping [51], whereby a relatively weak, slowly propagating plasma wave is damped by trapping and accelerating background plasma electrons, will produce a population of energetic electrons with some collective directionality. The problem with this situation is that the electrons are only weakly accelerated by the wave and fill all of its accelerating phases. The accelerated particles produced by this interaction have a huge spectrum of energies and form a beam only in the loosest sense of the word.

In plasma accelerators, the plasma waves are strong, have short wavelengths, and typically have phase and group velocities equal to c . Since the waves are short and moving so quickly, background electrons do not typically have enough

time to be accelerated to significant velocity before the decelerating portion of the plasma wave overtakes them. If the wave's fields are so strong that background electrons can be accelerated to the wave velocity before they move out of the accelerating phase, then they will co-propagate with the accelerating wave and be accelerated to high energy. The energy gained by the accelerated background electrons is lost by the plasma wave, which is thus damped. This self-damping process is another way to conceptualize the wave breaking limit. A wave is said to break, in the plasma accelerator context, when it begins to damp itself through the acceleration of background particles. As with conventional Landau damping, wave breaking is a means of producing energetic particles. These particles will have high energy but will still fill a very large area of phase space since there is no mechanism to restrict what phases of the accelerating wave particles are injected into.

Beams of high energy electrons have been produced using uncontrolled wave breaking. Santala *et al.* have produced 60 nC of electrons at energies above 10 MeV with a SMLWFA operating at 10^{19} cm^{-3} [27]. As is expected from injection through uncontrolled wave breaking, the energy spectrum of the particles was continuous and spanned 10 MeV - 120 MeV. Bulanov *et al.* have suggested a scheme for the LWFA in which a region of gradually declining plasma density is used to induce gentle wave breaking [52]. Once again the predicted energy spread of the captured particles is very large.

While particle trapping and acceleration through wave breaking is convenient in the sense that it occurs automatically once a large plasma wave is produced, the large energy spread of the particles produced is unsuitable for most beam applications. For that reason, research on plasma based electron beam sources turned to finding a way to control the phase space of the particles injected into

the accelerating waves. Many of these schemes use a laser pulse to stimulate the trapping of a small group of particles in a high gradient, but below wave breaking, plasma wave.

1.5.2 Laser Stimulated Injection Sources

In analogy to the operation of an RF photoinjector, an additional laser pulse can be used to stimulate the injection of electrons into a small region of the accelerating wake of a LWFA. This method of stimulated injection eliminates the problems associated with particle capture through wave breaking and allows LWFAs to trap and accelerate high quality beams with low energy spread. Umstadter, *et al.* has proposed a scheme that uses two orthogonal laser pulses that collide in a plasma [53]. The first pulse excites the laser wake field and the transverse ponderomotive forces of the second pulse give a population of the background electrons enough energy to achieve resonance with the wake and be accelerated to high energy. Simulations indicate that this technique can produce a beam of 21.2 MeV electrons with only 6% energy spread in a plasma 250 μm long. Esarey *et al.* have proposed a system that uses three collinear laser pulses. Again, one of these is the LWFA driver pulse. The other two pulses are counterpropagating and collide in the laser wake where they create a slow phase velocity beat wave that injects electrons into the main laser wake. This method creates bunches of similar energy but even lower energy spread, 0.3%, than Umstadter's. Unfortunately for both these methods, the technical challenge of reliably timing sub-ps laser pulse collisions is daunting.

1.5.3 Plasma Density Transition Trapping

Plasma density transition trapping was recently proposed by Suk *et al.* [54] as a new self-trapping system for the use in the blowout regime of PWFAs. In this scheme the beam passes through a sharp drop in plasma density where the length of the transition between the high density in region one (1) and the lower density in region two (2) is smaller than the plasma skin depth k_p^{-1} . As the drive beam's wake passes the sudden transition, there is a period of time in which it spans both regions, see Fig 1.11. The portion of the wake in region 2 has lower fields and a longer wavelength than the portion in region 1. This means that a certain population of the plasma electrons at the boundary will suddenly find themselves rephased into an accelerating portion of the region 2 wake, see Fig. 1.12. When the parameters are correctly set, these rephased electrons are inserted far enough into the accelerating region to be trapped and subsequently accelerated to high energy. This technique replaces the gentle wave breaking induced by a gradual density decline suggested by Bulanov *et al.* [52], with a single wave breaking event stimulated by a sharp density drop. Furthermore, the electrons injected into the region 2 wake come from the recombination point at the end of wake oscillation which has a very short longitudinal extent. Consequently, the injected electrons occupy a fairly small phase area of the accelerating wake and have a reasonable energy spread after acceleration. The energy spread can be reduced even further through manipulation of the plasma density after the transition, as discussed in Chapter 3.

Perhaps the most important feature of plasma density transition trapping is that it achieves injection of a limited phase space of electrons into a plasma wake without external timing requirements. Unlike the optical injection methods discussed in Section 1.5.2, which rely on the precise timing of multiple laser

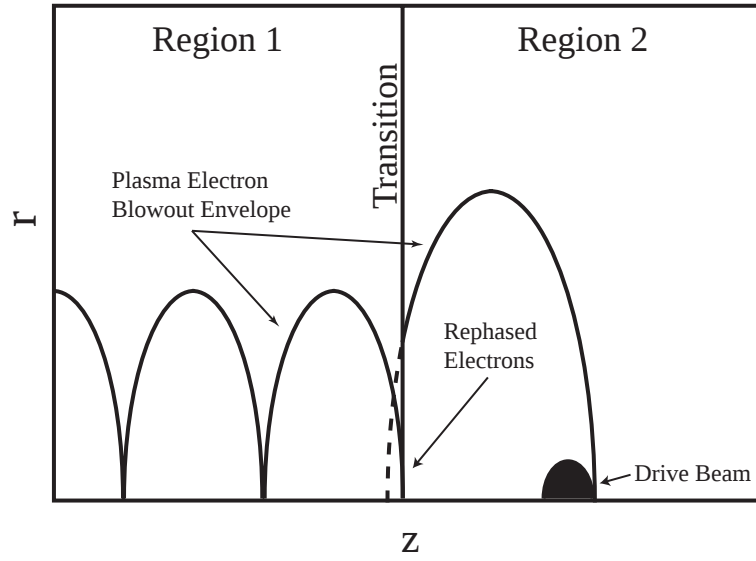


Figure 1.11: Configuration space diagram of the plasma density transition trapping mechanism. When the drive beam passes through the transition from the high density plasma in region 1 to the low density plasma in region 2, the length and breadth of the plasma wake it excites grows instantaneously. This fast change in the plasma oscillation’s wavelength results in the rephasing of plasma electron into the accelerating portion of the region 2 wake.

pulses, transition trapping occurs automatically as a result of the drive beam’s interaction with the static plasma environment. The elimination of the need for sub-ps synchronization gives transition trapping a large advantage over the optical schemes, provided that the production of suitable density transitions does not prove overly difficult (see Chapter 5). It should also be noted that this trapping mechanism can be driven by a LWFA as well.

Since its initial proposal, the merits of density transition trapping as a high quality electron beam source have been studied in detail [55, 56]. The majority of this work is recounted in Chapters 2 and 3. The results of these analyses indicate that plasma density transition trapping has the potential to be a robust and high brightness source of electrons. Its performance is similar to those of other

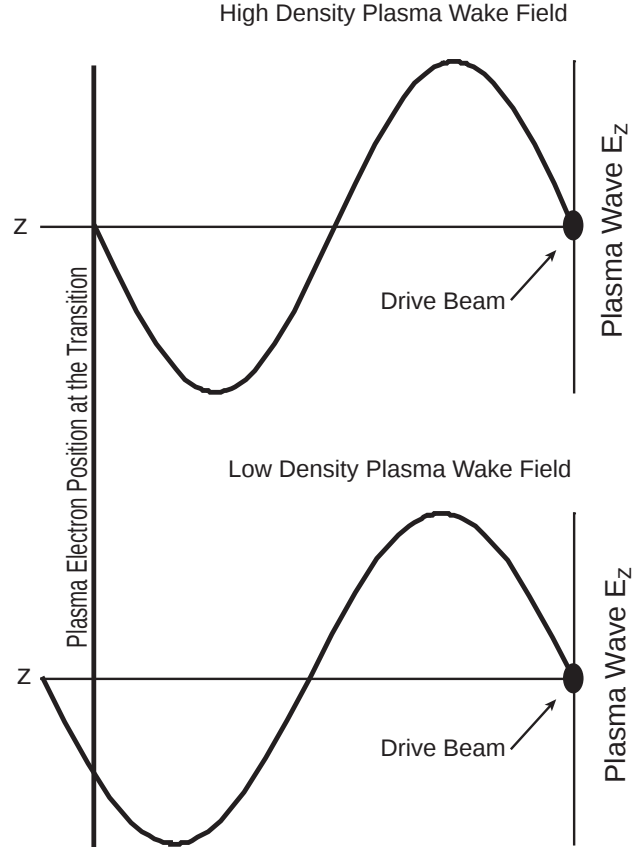


Figure 1.12: Diagram of the plasma electron rephasing caused by the plasma density transition. In a normal blowout regime plasma oscillation, the plasma electrons return to the drive beam propagation axis at the end of one period of the longitudinal wake field oscillation, as in the top diagram of this figure. Since there are no accelerating electric fields at this point, the plasma electrons cannot be accelerated. At the density transition, however, the wavelength of the longitudinal wake field suddenly increases, but the longitudinal distance between the drive beam and the plasma electrons returning to the axis does not. Consequently, the sudden change in wavelength moves the plasma electrons from the non-accelerating phase they had in the high density plasma wave, to an accelerating phase in the low density plasma wave. This rephasing allows the plasma electrons to be trapped and accelerated. Note that the longitudinal electric field oscillations in the blowout regime are not actually sinusoidal, see Section 2.2, but have been represented that way for clarity.

injection schemes at the same plasma density, despite its simpler mechanism. This is to be expected since most of the parameters relevant to high brightness beam production are set by the plasma density.

1.6 Chapter Summary

The current state-of-the-art in electron beam generation and acceleration is based on RF cavity technology. The gains that can be made in the critical parameters of beam energy and brightness by improving RF technology is reaching the point of diminishing returns. Order of magnitude increases in these parameters require new physical systems for acceleration.

Large amplitude plasma oscillations, excited by lasers or electron beams, may be a mechanism through which the limitations of RF technology can be overcome for both accelerators and beam sources. Unfortunately, practical use of plasma accelerators is frustrated by the small physical size of the plasma oscillations and the sub-ps timing requirements those length scales imply. There are ways to circumvent these timing problems by constructing plasma accelerator systems in which electrons are automatically captured by the plasma wave out of the background of plasma electrons.

The mechanism of plasma density transition trapping has been proposed as a high brightness electron beam source based on the PWFA. Unlike other high beam quality plasma injection schemes this technique requires no external timing. The detailed study of plasma density transition trapping, theoretically and experimentally, is the subject of the remainder of this dissertation.

CHAPTER 2

Theory of Particle Trapping and Acceleration

There is a great deal of commonality between the physics of injecting, or equivalently trapping, and accelerating charge in a photoinjector and physics of the same process in a plasma based beam source. The differences between the two cases mainly arise from the size and shape of the accelerating fields, the typical oscillation frequencies of those fields, and the presence in the plasma of a vast quantity of free electrons. This chapter will begin with an examination of trapping in photoinjectors and then extend the results to plasma sources.

2.1 Trapping in RF Photoinjectors

The dynamics of particle trapping and acceleration in a photoinjector can be analyzed using a Hamiltonian formalism. The longitudinal field of the forward traveling RF wave in an accelerating structure can be written simply as

$$E_z = E_0 \sin(\omega t - k_z z) = -\frac{\partial A_z}{\partial t}, \quad (2.1)$$

where,

$$A_z = \frac{E_0}{k_z v_\phi} \cos(k_z(v_\phi t - z)), \quad (2.2)$$

and E_z is the electric field, E_0 is the peak electric field of the wave, A_z is the z element of the electromagnetic vector potential, ω is the accelerating wave frequency, t is time, k_z is the accelerating wave's wave number, z is the particle's

longitudinal position, and v_ϕ is the phase velocity of the wave. The general one dimensional Hamiltonian of a charged particle in an electromagnetic field is

$$H = \sqrt{(\vec{p}_c - q\vec{A})^2 c^2 + (m_0 c^2)^2} + q\phi_e, \quad (2.3)$$

where $\vec{p}_c$ is the particle canonical momentum, q is the charge, and ϕ_e is the electrostatic scalar potential (see Appendix A). Substituting Eq. 2.2 into Eq. 2.3, taking $\phi_e = 0$, and considering only the z dimension gives

$$H = \sqrt{\left(p_{z,c} - q \frac{E_0}{k_z v_\phi} \cos(k_z(v_\phi t - z))\right)^2 c^2 + (m_0 c^2)^2}. \quad (2.4)$$

In order to use this Hamiltonian as a constant of the motion, we need to make a canonical transformation that eliminates time as an explicit variable. Ultimately, we would like to describe the position of the particle by the phase it occupies in the forward wave. The logical choice for the new coordinate is

$$\zeta = v_\phi t - z, \quad (2.5)$$

which is equivalent to a transformation of z into the frame of the forward traveling wave. Taking this new coordinate and equating both the new mechanical momentum to the old $p_\zeta = p_z$, and the new canonical momentum to the old $p_{\zeta,c} = p_{z,c}$, we can write down the transformed Hamiltonian, $\tilde{H}$, so as to preserve the proper equations of motion,

$$\tilde{H}(\zeta, p_{\zeta,c}) = H(\zeta, p_{\zeta,c}) - v_\phi p_{\zeta,c} \quad (2.6)$$

$$\tilde{H}(\zeta, p_{\zeta,c}) = \sqrt{\left(p_{\zeta,c} - q \frac{E_0}{k_z v_\phi} \cos(k_z \zeta)\right)^2 c^2 + (m_0 c^2)^2} - v_\phi p_{\zeta,c}. \quad (2.7)$$

Computing Hamilton's equations from $\tilde{H}$ and making liberal use of substitutions from Eqs. A.2 and A.3 we find

$$\frac{d\zeta}{dt} = \frac{\partial \tilde{H}}{\partial p_{\zeta,c}} = \frac{p_\zeta}{\gamma m_0} - v_\phi = v_z - v_\phi, \quad (2.8)$$

and

$$\frac{dp_{\zeta,c}}{dt} = -\frac{\partial \tilde{H}}{\partial \zeta} = \frac{qE_0 v_z}{v_\phi} \sin(\omega t - k_z z). \quad (2.9)$$

Now we can simplify the transformed Hamiltonian using the following definitions and approximations:

$$\chi = \sqrt{\frac{1 - \beta_z}{1 + \beta_z}}, \quad \beta_\phi \approx 1, \quad \phi = k_z \zeta, \quad (2.10)$$

where we make the assumption that the phase velocity of the accelerating wave is equal to the speed of light. At this point it is also useful to define the important parameter α , which is the ratio of the maximum normalized particle energy gain per unit length to the wave number of the accelerating wave,

$$\alpha = \frac{qE_{\max}}{k_z m_e c^2} = \frac{\left. \frac{d\gamma}{dz} \right|_{\max}}{k_z}, \quad (2.11)$$

where $k_z = \omega/v_\phi$. Again using Eqs. A.2 and A.3 along with the above equations, $\tilde{H}$ reduces to

$$\tilde{H} = m_0 c^2 (\chi - \alpha \cos \phi). \quad (2.12)$$

Here ϕ is the phase of the forward wave that a particle occupies. Since this Hamiltonian is a constant of the motion, it describes the evolution of χ with changes in ϕ for different values of α . We can use this Hamiltonian to derive functions like $\phi_f(\phi_i)$ which maps each particle from its initial position to its final position in the forward wave.

In a typical accelerating structure $v_z = c$ and $\phi_f = \phi_i$ as the beam travels through the structure. In a photoinjector, or a velocity bunching system, $v_z < c$ and $\phi_f \neq \phi_i$. Taking the specific case of a photoinjector, initially $v_z = 0$ at injection, and the final velocity at the exit of the photoinjector is $v_z \approx c$. The electrons must reach relativistic velocity or they will not be able to co-propagate with an accelerating phase of the wave and gain significant energy. In such a

situation the initial and final Hamiltonians are

$$\tilde{H}_{initial} = m_0 c^2 (1 - \alpha \cos \phi_i), \quad (2.13)$$

and

$$\tilde{H}_{final} = -m_0 c^2 \alpha \cos \phi_f. \quad (2.14)$$

Since $\tilde{H}_{initial} = \tilde{H}_{final}$, we can immediately deduce that

$$\phi_f = \cos^{-1} \left(\cos \phi_i - \frac{1}{\alpha} \right). \quad (2.15)$$

This equation gives the change in phase that a particle experiences as it is accelerated in the photoinjector. This equation can also be viewed another way:

$$\alpha = \frac{1}{\cos \phi_i - \cos \phi_f}, \quad (2.16)$$

which gives the value of α needed to accelerate a particle from rest to the speed of light if it starts at ϕ_i and ends up at ϕ_f . Now it can be seen from Eq. 2.1 that the wave is only accelerating between $\phi = 0$ and $\phi = \pm\pi$ (depending on the sign of the charge). The other half of the wave is decelerating. Therefore, the minimal trapping condition is that $\phi_i = 0$ and $\phi_f < \pi$. If the particle is not accelerated to synchronous relativistic velocity before it moves through the first π radians of the accelerating wave it never will be because the remainder of the wave will decelerate it and this acceleration/deceleration cycle will repeat. This minimal trapping condition implies that trapping and acceleration in system with $v_\phi = c$ requires $\alpha > 0.5$. The acceleration of a beam of reasonable quality requires $\phi_f \simeq \pi/2$ so one could say that the realistic trapping condition is more like $\alpha > 1$. Photoinjectors usually meet this criterion. For example, the photoinjector described in Table 1.2 has $\alpha = 1.96$.

As discussed in Section 1.3.2, another interesting quantity is the compression in phase that a group of particles experiences as it is accelerated. We can find this

expression by differentiating Eqs. 2.13 and 2.14 with respect to ϕ_i and solving for $d\phi_f/d\phi_i$, which gives

$$\frac{d\phi_f}{d\phi_i} = \frac{\Delta\phi_f}{\Delta\phi_i} = \frac{\sin\phi_i}{\sin\phi_f}. \quad (2.17)$$

Since it is desirable to have $\phi_f = \pi/2$ under optimal operating conditions this reduces to

$$\frac{\Delta\phi_f}{\Delta\phi_i} = \sin\phi_i. \quad (2.18)$$

This is the basic equation describing velocity compression in a RF photoinjector. In addition, terms can be added to the Hamiltonian to account for the effects of space charge, the backward traveling wave in a standing wave structure, etc. [43].

2.2 Plasma Accelerators and Plasma Trapping

The Hamiltonian formalism developed in the preceding section is sufficiently general that it can be adapted to plasma accelerators and trapping almost immediately. In fact, for the linear, non-relativistic plasma wave case discussed in Section 1.2 the accelerating wave is sinusoidal and the Hamiltonian treatment is identical to that for RF structures. If we use Eq. 1.12, which gives the plasma frequency, and Eq. 1.18, which gives the plasma wave breaking field, to compute α , we find that $\alpha = 1$ for all plasmas oscillations of wave breaking amplitude regardless of density. This is exactly what is expected since wave breaking means that the wave is trapping plasma electrons from the background and strongly accelerating them. This trapping also results in damping of the wave because the wave loses the energy required to accelerate the electrons.

As was discussed in Section 1.2.1, non-linear, relativistic plasma waves are preferred for modern plasma accelerator schemes. Hamiltonian analysis can help

us understand trapping behavior in these wave too, at least in the one dimensional limit. Using Eqs. 2.3 and 2.6 we can write a general form for $\tilde{H}$:

$$\tilde{H} = \sqrt{p^2 c^2 + (m_0 c^2)^2} + q\phi_e - v_\phi p_c = \gamma m_0 c^2 + q\phi_e - v_\phi p_c, \quad (2.19)$$

which reduces to

$$\tilde{H} = m_0 c^2 \chi + q [\phi_e - cA,] \quad (2.20)$$

using the same types of substitutions used in Section 2.1. All that is needed now are the potentials that describe the accelerating wave of interest.

The one-dimensional, relativistic fluid equation of motion for a plasma wave was originally derived by Akhiezer and Polovin and is given by

$$\frac{d^2}{d\tau^2} \left[\frac{1 - \beta_b \beta_p}{(1 - \beta_p^2)^{1/2}} \right] = \beta_b^2 \left[\frac{\beta_p}{\beta_b - \beta_p} + \frac{n_b}{n_0} \right], \quad (2.21)$$

where $\tau = \omega_p(t - z/v_b)$, $\beta_b = v_b/c$, v_b is the driving electron beam velocity which is equal to the wave phase velocity, $\beta_p = v_p/c$, v_p is the plasma electron velocity, n_b is the electron beam density, and n_0 is the unperturbed plasma density [57, 58, 59]. When $\beta_b \approx 1$, which is the case of interest here, and we wish to examine oscillations in the region behind the impulsive driving electron beam, Eq. 2.21 reduces to

$$\frac{d^2 \chi_p}{d\tau^2} = \frac{1}{2} \left[\frac{1}{\chi_p^2} - 1 \right], \quad (2.22)$$

where χ is defined in Eqs. 2.10. This equation describes nonlinear, longitudinal, electrostatic plasma oscillations. It will be shown that Eq. 2.22 is equivalent to Poisson's equation as well as the equation of motion. From Maxwell's equations, and the definition of χ , we can write

$$n = \frac{n_0}{2} \left(1 + \frac{1}{\chi_p^2} \right), \quad (2.23)$$

where n is the plasma fluid density, and n_0 is the unperturbed density. Using this in Poisson's equation we find

$$\nabla^2\phi = 4\pi e(n - n_0) = 4\pi e\left(\frac{n_0}{2}\left(\frac{1}{\chi_p^2} - 1\right)\right). \quad (2.24)$$

Substituting using Eq. 2.22 and converting the Laplacian from cartesian coordinates to τ , following the method of Eq. 1.14, gives

$$\left(\frac{\omega_p^2}{c^2}\right)\frac{d^2\phi}{d\tau^2} = 4\pi en_0\frac{d^2\chi_p}{d\tau^2}. \quad (2.25)$$

Integrating, making the appropriate choice for the constant, and utilizing the definition of the plasma frequency we have

$$q\phi(\chi_p) = -m_e c^2(1 - \chi_p). \quad (2.26)$$

This is the potential for the non-linear relativistic plasma wave. Substituting this potential into the general Hamiltonian, Eq. 2.20, gives the Hamiltonian for an electron interacting with this wave

$$\tilde{H} = m_e c^2[\chi - (1 - \chi_p)]. \quad (2.27)$$

It is very important to remember that χ describes the dynamics of the electron being accelerated, whatever its source, while χ_p describes the dynamics of the plasma electrons producing the wave.

As in Section 2.1, this new $\tilde{H}$ is a constant of the motion. Once again $\tilde{H}_{initial} = \tilde{H}_{final}$, which reduces to

$$\chi_i - \chi_f = \chi_{p,f} - \chi_{p,i}. \quad (2.28)$$

Trapping still requires that the accelerated electron becomes relativistic. This statement is equivalent to the condition $\chi_f = 0$. Now, Eq. 2.22 has been solved to give $\chi_p(\tau)$ in terms of elliptical integrals [58], which can be used to determine the equation for the electric field $E(\tau)$. For our purposes it suffices to note

the behavior of the wave's electric field for the extremes of χ_p : $E = E_{max}$ for $\chi_p = 1$ and $E = 0$ for $\chi_p > 0$, so that the possible values for χ_p are bounded by $0 < \chi_p \leq 1$. If we enforce the trapping condition on Eq. 2.28 and consider the consequences of the boundaries on χ_p we find

$$\chi_i = \chi_{p,f} - \chi_{p,i} < 1, \quad (2.29)$$

which means that electrons initially at rest ($\beta = 0$, $\chi = 1$) cannot be trapped and accelerated.

This limitation on trapping can be circumvented by the presence of a density transition. There is a discontinuity in $\tilde{H}$ at the transition boundary which effectively boosts the electrons that cross it, which were initially at rest, into the trapped regime. Suk *et al.* [54] created a Hamiltonian model of plasma density transition trapping based on this type of analysis. The predictions of this analytic treatment were in reasonable agreement with one dimensional particle-in-cell (PIC) code simulations considering the simplicity of the model.

The one dimensional analysis of Suk *et al.* has a serious shortcoming in that it assumes the plasma fluid is discontinuous at the transition boundary. In this non-physical picture, there is a high density plasma and a lower density plasma that oscillate independently from each other. This fundamental flaw was corrected in the work of England *et al.* [60], which treats the plasma as a continuous fluid which has an initial density which depends on the position $n_0(z)$. This one dimensional analytic analysis agreed fairly well with the results of one dimensional PIC simulations up to the point where the assumptions inherent in the equations broke down.

While the one dimensional analyses of Suk and England shed light on the physical mechanism of plasma density transition trapping, they are both limited models. Plasma density transition trapping operates in the highly two dimen-

sional PWFA blowout regime and, as was discussed in Section 1.2.1, there is no analytic theory available to describe this two dimensional non-linear problem. It is therefore necessary to use numerical computer simulations to further understand the physics of plasma density transition trapping and make accurate predictions about its behavior.

CHAPTER 3

Numerical Simulations of Plasma Density

Transition Trapping Regimes

It is interesting to revisit the prosaic description of the mechanics of plasma density transition trapping presented in Section 1.5.3 in light of the discussion in Chapter 2. In the case presented in the original paper by Suk *et al.* [54] the plasma after the transition has an oscillation frequency of 43 GHz and an accelerating gradient of 470 MV/m which yields $\alpha = 1$. The wave does not break and start accelerating background electrons because operation in the blowout regime insures that no background electrons are present at the accelerating phases of the plasma wave. The function of the density transition is to place background electrons at the accelerating phases of the wave. The electrons rephased by the density transition are trapped and accelerated to high energy as one would expect in a $\alpha = 1$ system.

This chapter reviews the results of detailed numerical simulations of plasma density transition trapping in two regimes: strong blowout and weak blowout. The two regimes are represented by example cases. The first case is the one originally presented in Suk *et al.*, in which the plasma both before and after the transition is strongly blown out. The second is a case originally designed for an experiment at the UCLA Neptune Laboratory in which the underdense condition, $n_b > n_0$, is barely satisfied on the high density side of the transition. The physics

Table 3.1: Drive and Captured Beam Parameters in the Strong Blowout Case. Figures for the captured beam are for the core of the captured beam, which is about 20% of the captured particles, after 12 cm of acceleration.

Drive Beam		Captured Beam Core	
Beam Energy	50 MeV	Beam Energy	56 MeV
Beam Charge	63 nC	Beam Charge	5.9 nC
Beam Duration σ_t	3 ps	Beam Duration σ_t	161 fs
Beam Radius σ_r	500 μm	Beam Radius σ_r	112 μm
Peak Beam Density	$1.2 \times 10^{14} \text{cm}^{-3}$	Normalized Emittance ε_x	155 mm-mrad
		Total Energy Spread	13%

and merits of each case as a beam source are compared.

3.1 The Strong Blowout Scenario

The strong blowout scenario is illustrated in Fig. 3.1 and uses the drive beam parameters presented in Table 3.1. The plasma density profile is a simple step function with a constant density of $n_{\text{region1}} = 5 \times 10^{13} \text{cm}^{-3}$ in the high density region and a constant density of $n_{\text{region2}} = 3.5 \times 10^{13} \text{cm}^{-3}$ in the low density region. The high charge driver produces a very strong blowout.

A series of simulations were performed with the 2D Particle-In-Cell code MAGIC [61] in which the high and low density plasma electron populations are tracked separately. The results of these simulations show that the trapping process begins in the high density region, as can be seen in Fig. 3.1. Electrons from the low density region near the transition are blown out and pushed backward into the high density plasma region. Once they enter the high density region, the oscillation of the region 2 plasma electrons is sped up by the higher ion den-

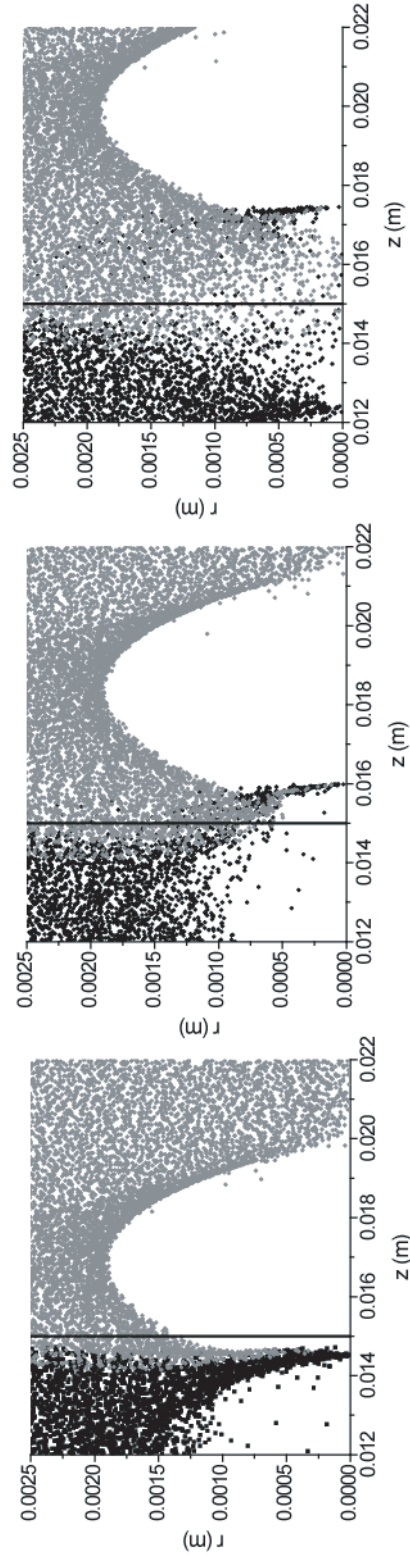


Figure 3.1: Configuration space (r, z) distributions of the plasma electrons illustrating trapping in the strong blowout case. The vertical black line indicates the original position of the density transition. Plasma electron particles originating in the high density region are colored black while particles originating in the low density region are colored grey.

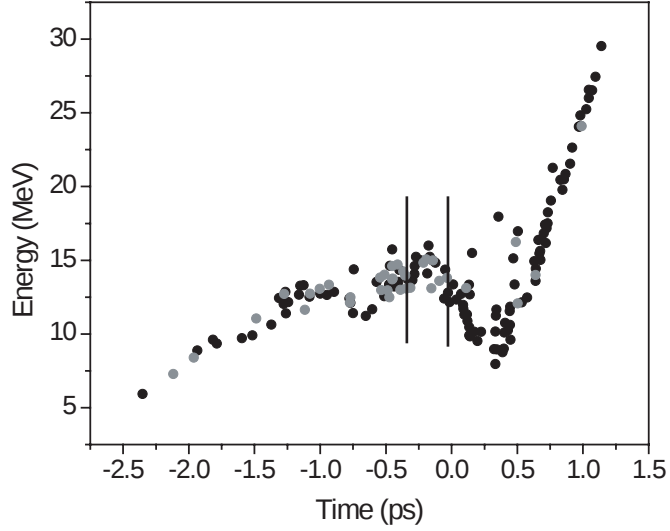


Figure 3.2: Trapped particle energy versus temporal position within the bunch for the strong blowout case. This plot shows the state of the trapped bunch after about 5 cm of acceleration. As in Fig. 3.1, particles origination on the high density side of the transition are black, while those starting on the low density side of the transition are grey. The region between the two vertical lines is the portion of the beam referred to as the core in Table 3.1.

sity and these electrons return to the axis early to mix with electrons from the high density region. As this mixed concentration of plasma electrons crosses the boundary between the high and low density regions many of the electrons are placed in an accelerating phase of the low density plasma wake and are subsequently trapped and accelerated. As is shown in Fig. 3.2, the population that is ultimately accelerated to high energy is a mixture of plasma electrons from both sides of the density transition. This finding is different from that of the initial results of Suk *et al.* which indicated that all the trapped electrons come from region 2, the low density side of the transition [54].

The properties of the beam captured in this scenario are listed in the second column of Table 3.1. The captured beam is very short and has a small radius compared to the beams produced by other sources such as the photoinjector described

in Table 1.2. Both of these features originate from the small accelerating volume of the accelerating plasma wave. The beam also has a high charge that results from the very high concentration of electrons in the oscillation density spike that are injected. Unfortunately, the captured beam has a significant energy spread that results from the fast variation in the plasma wake field accelerating gradient where the particles are captured. This problem of large energy spread can be fixed through manipulations of the plasma density profile as discussed in Section 3.2. The beam also has a poor transverse emittance. The large emittance is the result of several factors. As mentioned in Section 1.5, the plasma electrons have a relatively high temperature. High initial electron temperature is an unavoidable consequence of using plasma electrons to create an electron beam. Since the transition trapping system in this example is operated at the relatively low density of about 10^{13} cm^{-3} , the transverse size of the plasma oscillation producing the beam is relatively large, which contributes to the poor emittance. As will be discussed in Section 4.2, this problem can be solved by using higher plasma densities. The final contribution to the large emittance comes from trapping background plasma particles in the strong blowout regime. The large amplitudes of transverse momenta imparted to the plasma electrons as the drive beam space charge blows them out to the side remains with the particles as they are trapped and accelerated to high energy. This effect can be mitigated by operating in the regime of weaker blowout.

In addition to the undesirable emittance and energy spread properties of the captured beam, this transition trapping scenario is also impractical from an experimental standpoint. The drive beam parameters listed in the first column of Table 3.1 are not currently achievable. Lack of a drive beam for strong blowout case was strong incentive to begin developing trapping experiments which could be done with the more modest driver beams that are available. During these

studies ways were also found to improve both the emittance and energy spread of the captured beams.

3.2 The Weak Blowout Scenario

A great deal can be learned about the mechanism and dynamics of density transition trapping by comparing the strong blowout case, previously described, to a case in which a weak blowout is used, as shown in Fig 3.3. Comparison of Figures 3.1 and 3.3 reveals the dramatic differences between the strong and weak blowout cases. Perhaps most notable is great reduction in transverse plasma motion in the weak blowout case, which leads to reduced emittances.

Our standard example of a weak blowout case is the proof-of-principle experimental case designed for the Neptune Advanced Accelerator Laboratory at UCLA [62]. This case was developed and optimized for parameters achievable at the Neptune Laboratory through extensive simulations with MAGIC. The driving beam parameters of the simulation are shown in Table 3.2. The driving beam has a ramped longitudinal profile as shown in Figure 3.4. Ramped profiles of this type maximize the transformer ratio of the wake field [10] and can be produced using a negative R_{56} magnet compressor system. A negative R_{56} compressor system is under development for the Neptune Laboratory [11]. While the ramped beam profile improves performance, it is not critical to this trapping scenario.

The plasma density profile used in this case is illustrated in Figure 3.4. The plasma density profile is tailored to maximize the amount of charge captured while maintaining an acceptable amount of acceleration. The first centimeter of the profile reflects a realistic finite rise time from zero to the maximum plasma density. After 5 mm of maximum density the transition takes place and the den-

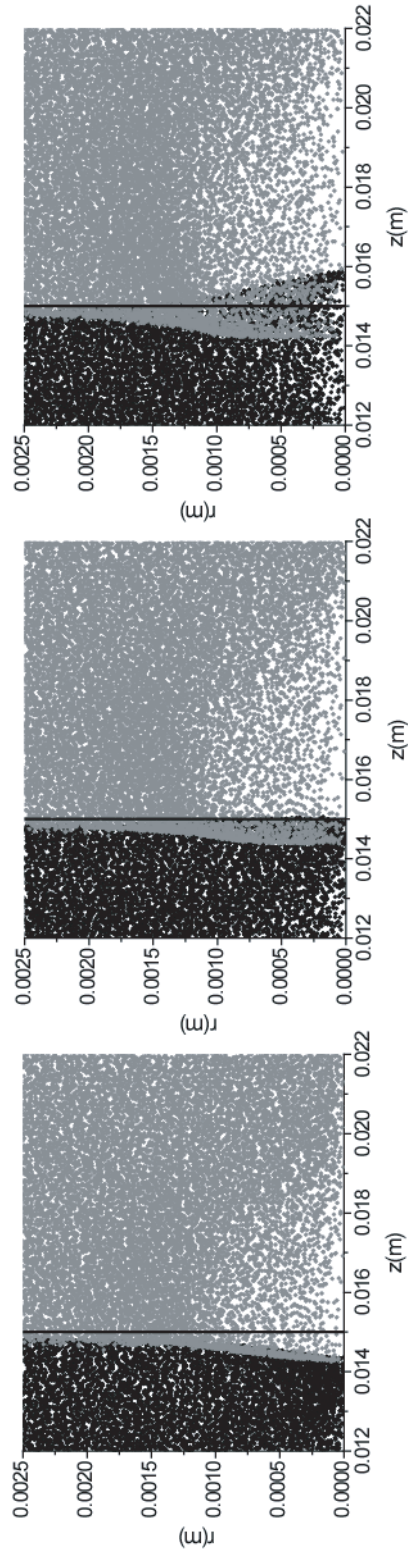


Figure 3.3: Configuration space (r, z) distributions of the plasma electrons illustrating trapping in the weak blowout case. This figure is directly comparable with Figure 3.1. The scale and particle coloring are identical. Note that the trapping mechanism is essentially the same except that it proceeds more slowly due to the low plasma density in the down stream region compared to strong blowout case. The weaker blowout also leads to much less transverse disturbance in the plasma, which in turn yields lower emittance.

Table 3.2: Drive and Captured Beam Parameters in the Weak Blowout Case. Values are for the beam core, which is all particles within a radius of 1 mm.

Drive Beam		Captured Beam Core	
Beam Energy	14 MeV	Beam Energy	1.2 MeV
Beam Charge	5.9 nC	Beam Charge	120 pC
Beam Duration	6 ps	Beam Duration σ_t	1 ps
Beam Radius σ_r	540 μm	Beam Radius σ_r	380 μm
Normalized Emittance ε_x	15 mm-mrad	Normalized Emittance ε_x	15 mm-mrad
Peak Beam Density	$4 \times 10^{13} \text{cm}^{-3}$	Total Energy Spread	11%

sity is reduced to 18% of the maximum. This density drop is near the optimum to maximize charge capture. Decreasing the density of region 2 increases the wavelength of the accelerating plasma wave. Increasing the wavelength has the effect of enlarging the volume of the capture region and enhancing the amount of charge trapped. Lowering the plasma density also reduces the accelerating gradient, however, which reduces the number of initially captured particles that ultimately achieve resonance with the accelerating wave. These two effects compete with the charge capture maximum occurring at about $n_{\text{region2}} = 0.18n_{\text{region1}}$. To quantify the degree of blowout in the weak blowout case, we note that the electron beam density is 2 times larger than the peak plasma density of $2 \times 10^{13} \text{cm}^{-3}$, as can be seen from Table 3.2 and Figure 3.4.

The gradual decline in plasma density after the transition slowly increases the plasma wavelength, and thus the extent of the accelerating phase of the wake field region. The growth in plasma wavelength reduces the peak gradient but rephases the captured charge forward of the peak field of the wake into a

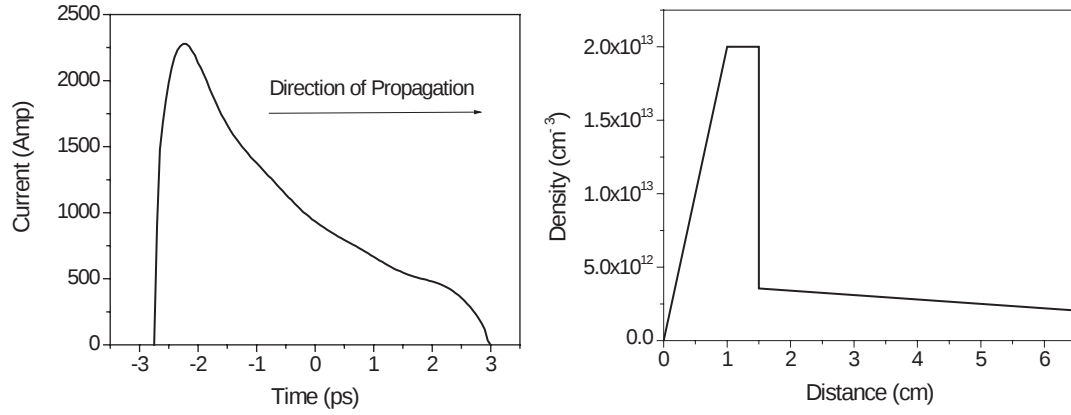


Figure 3.4: Drive Beam Current (left) and Plasma Density (right) Profiles.

region of slightly weaker, but more uniform, acceleration. This rephasing both increases the amount of charge trapped and reduces the energy spread. The rate of density decline can be increased to reduce energy spread even further. Gradually declining post transition plasma densities have been shown to have similar benefits in the strong blowout regime [55]. In cases where high total energy is more important than low energy spread, a post transition density incline can be used. This has the opposite effect of a density decline and rephases the captured charge into a higher gradient portion of the wave. Rephasing using a density incline results in trapped charge loss, however, and the incline cannot be very steep or the trapped charge will be phased out of the wave entirely.

In this weak blowout regime, all the trapped electrons originate in region 1, the high density region, as can be seen in Fig. 3.5. This fact is interesting as it indicates that in this mode of density transition trapping operation the high density region acts much more like a traditional cathode. The perturbation of the high density region is relatively weak so that this region provides a semi-uniform plane of plasma electrons that are accelerated in the wake of the much stronger

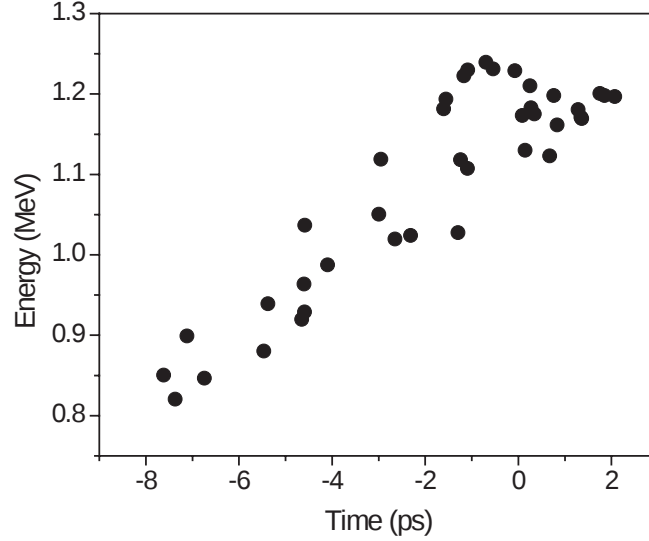


Figure 3.5: Trapped particle energy versus temporal position within the bunch for the weak blowout case. This plot shows the state of the trapped bunch at the end of the acceleration. As in Fig. 3.3, particles origination on the high density side of the transition are black. In this scenario only particles from the high density region are trapped.

blowout in region 2.

The parameters of the core bunch of captured plasma electrons are given in Table 3.2. As can be seen from comparison of Figures 3.5 and 3.6, the low energy tail of the trapped beam corresponds to the particles at large radius. These large radius particles also have significantly higher radial momentum than those at small radius. It is therefore natural to exclude these outlying particles from subsequent analysis since they are unlikely to be transported with the core of the beam. Comparing Tables 3.1 and 3.2 reveals the gains made by using the weak blowout scenario in terms of lowered emittance and relaxation of the driver beam requirements. The comparison also highlights the reduction of trapped charge and trapped beam energy that result from the lower accelerating fields in the weak blowout case. The fairly modest improvement in energy spread results

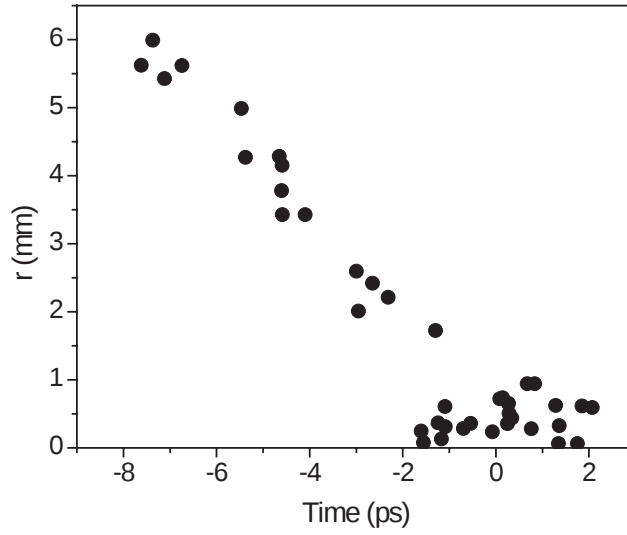


Figure 3.6: Trapped particle radial position versus temporal position within the bunch for the weak blowout case.

from the density decline after the transition. A stepper decline would decrease energy spread even further, but that would mean a reduction in total energy, which is already fairly low in this case.

CHAPTER 4

Scaling of the Transition Trapping System

Chapter 3 presented the performance expectations for plasma density transition trapping in two different regimes with plasma densities in the 10^{13} cm^{-3} range. The parameters of the beams produced under these conditions are not very impressive when compared to conventional sources, compare Tables 3.1 and 3.2 with Tables 1.1 and 1.2, especially in terms of emittance. For that reason a study was undertaken of the changes in plasma density transition trapping performance with scaling of the drive beam charge and plasma density. The behavior found is consistent with the trends established in Sections 1.3 and 1.5 toward better emittance with smaller source sizes. All the simulation results in this chapter are based on the weak blowout case presented in Section 3.2. The trends deduced from, and scaling rules supported by, these simulations also apply to the strong blowout case.

4.1 Driver Charge Scaling

The first scaling examined is variation of the driver beam charge. The strong and weak blowout cases represent two extremes of drive beam charge. In order to find how the captured beam parameters change between these two extremes, the effects of scaling up the drive beam charge without altering the rest of the parameters were simulated. The results of these simulation are shown in Figure

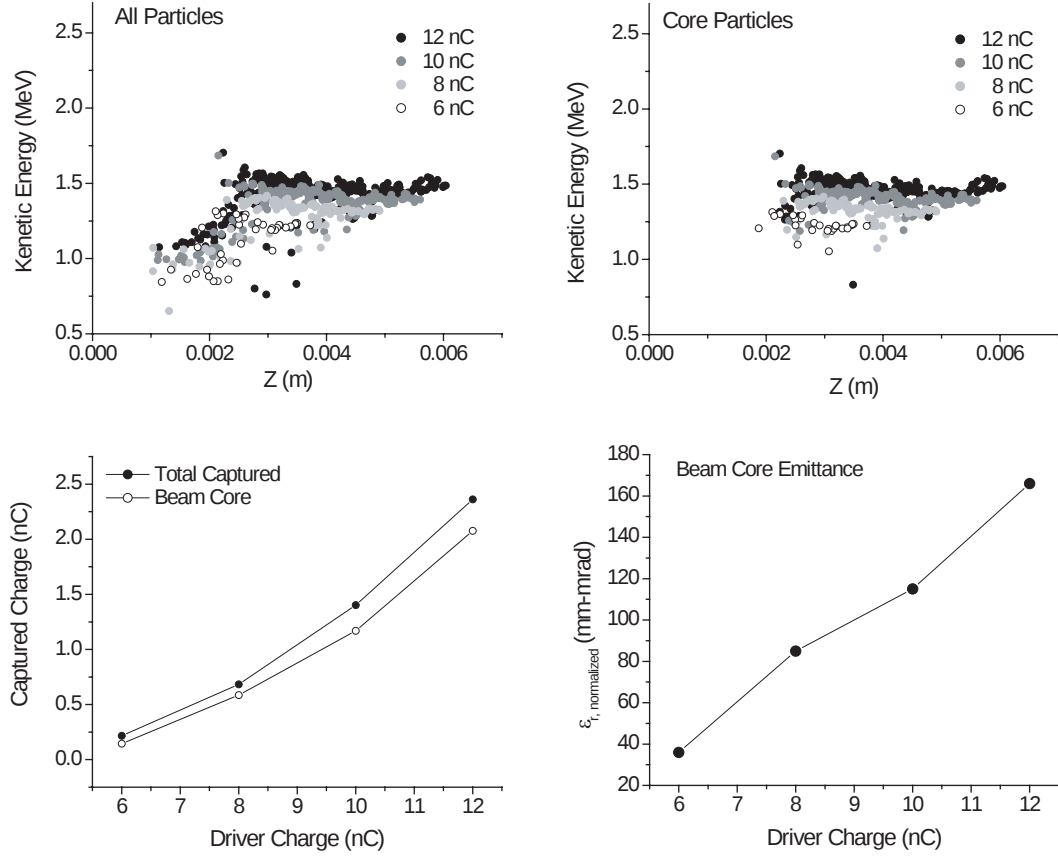


Figure 4.1: The effects of driver beam charge scaling on trapped beams. Note that the lower right plot uses $\epsilon_r = 2 \epsilon_x$.

4.1. Increasing the driver charge increases the strength of the blowout, forming a larger amplitude, more non-linear plasma wave: it follows that all the accelerating fields in the system are increased, as is the size of the accelerating wave. The impact of drive beam scaling on the captured beam is clearly shown in Fig. 4.1, where the beam core is defined in the same way as in Section 3.2. The amount of charge captured, the length of the beam, and the emittance all grow as the driver charge is increased. It is notable that while the other parameters grow linearly with driver charge, the captured charge grows faster. As can be seen from Fig. 4.1, the growth in trapped charge with increasing driver charge is given approximately by $Q_{trap} \propto Q_{driver}^4$. This non-linearity can make scaling to higher driver charge useful in some circumstances, but for the most part simple scaling of the driver charge does not lead to higher quality trapped beams.

4.2 Wavelength Scaled Sources

We have seen the performance of density transition trapping at densities $n_0 \sim 10^{13}\text{cm}^{-3}$ in the preceding sections, as well as how the performance changes with driver charge scaling. From these studies it is clear that transition trapping at $n_0 \sim 10^{13}\text{cm}^{-3}$ produces beams of low brightness when compared to the benchmark of the LCLS photoinjector [63], see Table 4.1. It is therefore interesting to examine how the captured beam performance scales with plasma density or, equivalently, the plasma wavelength. This type of wavelength scaling, and its impact on predicting beam emittance and brightness behavior, has been previously examined in the context of RF acceleration [64] in photoinjector sources, where the beam displays plasma-type behavior.

In order to scale the transition trapping system to a higher plasma density,

n_{high} , all the charge densities in the system must be increased by the ratio,

$$n_{high}/n_0, \quad (4.1)$$

and all the lengths in the system are decreased by the ratio,

$$\frac{\lambda_{p \text{ high}}}{\lambda_{p \text{ 0}}} = \frac{k_p^{-1} \text{ high}}{k_p^{-1} \text{ 0}} = \frac{1/\sqrt{n_{high}}}{1/\sqrt{n_0}} = \sqrt{\frac{n_0}{n_{high}}}, \quad (4.2)$$

where λ_p represents the typical wavelength of oscillations in the plasma and is equal to the plasma skin depth $\lambda_p = k_p^{-1} = c/\omega_p$. In scaling the system we also require that the plasma oscillation be self-similar. This means that both the relative density disturbance $\delta n/n_0$ and the normalized peak field $E_z/E_{wavebreak} = eE_z/m_e\omega_p c$ remain constant (see Section 1.2 for a derivation of $E_{wavebreak}$). It follows that the scaled phase space distributions of the captured electrons will also be self-similar. This can be seen from noting that the above requirement can be written,

$$\frac{eE_z}{m_e\omega_p c} = \frac{e}{m_e c^2} E_z \lambda_p = \text{Constant}, \quad (4.3)$$

so that the captured particle momenta p is given by

$$p \propto E_z \lambda_p = \text{Constant}. \quad (4.4)$$

Consequently, the emittance ε , which is proportional to the product of momenta and the beam size, scales according to

$$\varepsilon \propto \lambda_p p \propto \lambda_p. \quad (4.5)$$

The emittance of the captured beam improves as the system is scaled to higher density as a result of the reduction in the transverse beam size. The result is a formal statement of the trend toward better emittance with smaller source discussed throughout Sections 1.3 and 1.5.

Table 4.1: Simulations of wavelength scaling using MAGIC 2D.

Peak Density	$2 \times 10^{13} \text{cm}^{-3}$	$2 \times 10^{15} \text{cm}^{-3}$	$2 \times 10^{17} \text{cm}^{-3}$	
$\sigma_{t,Driver}$	1.5 psec	150 fsec	15 fsec	LCLS
Q_{Driver}	10 nC	1 nC	100 pC	Photoinjector
$\sigma_{t,Trap}$	2.7 psec	270 fsec	28 fsec	Specifications
Q_{Trap}	1.2 nC	120 pC	12 pC	—————
$I_{Peak,Trap}$	163 Amp	166 Amp	166 Amp	100 Amp
$\varepsilon_{x,norm,Trap}$	57 mm-mrad	5.9 mm-mrad	0.6 mm-mrad	0.6 mm-mrad
$B_{norm,Trap}$	5×10^{10}	5×10^{12}	5×10^{14}	2.8×10^{14}

The amount of charge captured, Q , depends on both the available plasma electron density, n_0 , and the volume of the accelerating portion of the wave, which is proportional to λ_p^3 . This scaling can be written as

$$Q \propto n_0 \lambda_p^3 \propto n_0 \left(\frac{1}{\sqrt{n_0}} \right)^2 \lambda_p \propto \lambda_p. \quad (4.6)$$

While the captured charge goes down as the plasma wavelength is reduced, the current I remains constant since the length of the beam also goes down with the plasma wavelength,

$$I \propto \frac{Q}{\lambda_p/c} = \text{Constant}. \quad (4.7)$$

Finally we can combine the scaling laws for emittance and current to deduce the scaling of the beam brightness, B :

$$B \propto \frac{I}{\varepsilon^2} \propto \frac{1}{\lambda_p^2} \propto n_0. \quad (4.8)$$

Thus the brightness of electron beams produced using density transition trapping increases linearly with the density of the plasma.

These scaling laws were tested using the 2D PIC code MAGIC. The cases examined are scaled versions of the weak blowout case describe in Section 3.2 with a slightly larger driver charge. The results are summarized in Table 4.1. The simulation results follow the scaling laws precisely in the range studied. At $2 \times 10^{17} \text{ cm}^{-3}$ transition trapping can produce an extremely short beam with excellent emittance and a brightness that exceeds state of the art photo-injectors. The drive beams needed at all densities must be of similar length and approximately one order of magnitude greater charge than the beams they capture. The emittance of the driver, however, is irrelevant as long as the driving beam can be focused sufficiently to match into the plasma. This means that plasma density transition trapping might be used as an emittance transformer to produce short, low emittance beam from short beams with high emittances that were produced using extreme magnetic compression or other techniques that produce significant emittance growth. The feasibility of this idea is still under study and may be enhanced by our effort to find new scenarios that produce low emittance trapped beams.

Little can be done to reduce the high temperature of plasma electrons, see Section 1.5, and the contribution of this temperature to the trapped beam emittance. As described previously, however, plasma density transition trapping produces beams with larger emittances than the thermal limit due to the sizeable transverse momenta the plasma particles have at capture. Operating in the weak blowout regime rather than the strong blowout regime reduces the induced transverse momenta but does not eliminate it. Scaling to higher density improves the emittance by reducing the beam size rather than reducing the transverse momentum. There are continuing efforts to explore alternative transition trapping scenarios with the aim of reducing the transverse momentum of the beam further. This may be accomplished by using drive beams that are wide and or

long compared to the plasma skin depth. The development of the technique of foil trapping, which is discussed in the next section, might also lead to lower transverse momenta.

CHAPTER 5

Design of the Transition Trapping Experiment

The body of theoretical work presented in the previous chapters indicates that plasma density transition trapping is an interesting physical process with the potential to become an important advanced accelerator technology. The next logical step in understanding the physics of plasma density transition trapping is creating the effect experimentally. To that end, we had set out to build an apparatus that would allow us to conduct an experiment with parameters similar to those presented in Section 3.2. The goal of this experiment was simply to prove that the principle of transition trapping works. We therefore endeavored to create the most basic experiment that would achieve this aim.

This chapter gives an overview of the process of developing the plasma density transition trapping experiment. This process is recounted in semi-narrative manner and includes elements of the hardware construction, simulation work, and planning for experimental shortcomings and contingencies. The chapter begins with a statement of the experimental design goals. It then moves on to a technical description of the development of the plasma source as a stand alone device in Section 5.2. Once the operational characteristics of the plasma source are established, the successful development of a technique to create plasma density transitions using this source is described in Section 5.3. This section also examines some of the flaws that density transitions produced using this technique have, and how those flaws will impact trapping. Section 5.4 continues along these

lines and analyzes the effect that other real experimental factors will have on the plasma electron trapping. Finally, Section 5.5 describes the experimental plan that emerged from this design process and how it compares to original goals.

It should be noted that this chapter does not deal extensively with the production of the electron drive beam for the experiment. This is because the drive beam was provided by a pre-existing accelerator facility. A description of that facility, the work done to integrate this experiment into its beamline, and the beam diagnostics used to conduct the experiment is given in Chapter 6. Experimental results are given in Chapter 7.

5.1 Design Goals

The fundamental goal of the experimental design is to create a plasma density transition that can interact with an electron beam from an existing accelerator and, through this interaction, produce a significant quantity of captured charge. From this simple statement it is clear that the overall design of the experiment must follow from two choices: the driving beam parameters, and the plasma transition production mechanism.

Since our goal was to produce the simplest possible experiment that would exhibit trapping, we did not want the production of the driving beam to become a research project in its own right. This lead us to plan the experiment around beam parameters that are reasonably achievable at existing facilities (see Section 6.1).

The production of the plasma density transition is a problem that is both novel and complex. Either a plasma must be created with the density transition built in, or the plasma must first be produced and then altered to have a density

transition. The first approach naturally lends itself to the technique of photo-ionization, which is used to produce high density plasmas in many applications. It might be possible to produce density transitions through photo-ionization using a laser with an intensity profile that matches the desired plasma density profile or using a uniform laser to ionize a dual density gas jet. Alternatively, plasmas produced through electric discharges may be more appropriate for producing a density transition by modifying a uniform plasma after it is created. It has been shown that the diameter of a flowing plasma column can be controlled by using solid metallic barriers [65, 66]. If one of these solid barriers is replaced with a metal mesh or screen, the expected result would be a plasma column with high and low density regions and a sharp transition between them. Although electric discharge plasma sources produce lower plasma densities, and consequently lower accelerating gradients, than laser based sources, performing the experiment at low plasma density greatly reduces the technical difficulty of the experiment by relaxing the drive beam requirements. This fact, coupled with our previous experience with PWFAs experiments based on discharge sources [14, 17, 66], lead us to design a low density, discharge source based, plasma density transition trapping experiment.

5.2 The Argon Pulse Discharge Plasma Source

The argon discharge plasma source discussed here was originally designed to produce plasma for a plasma lens experiment. The device was redesigned and rebuilt for the plasma density transition trapping experiment. The mechanism of plasma production, however, remains very similar to earlier versions of the apparatus [65, 66, 67].

The plasma source is, in essence, a very sophisticated version of the common

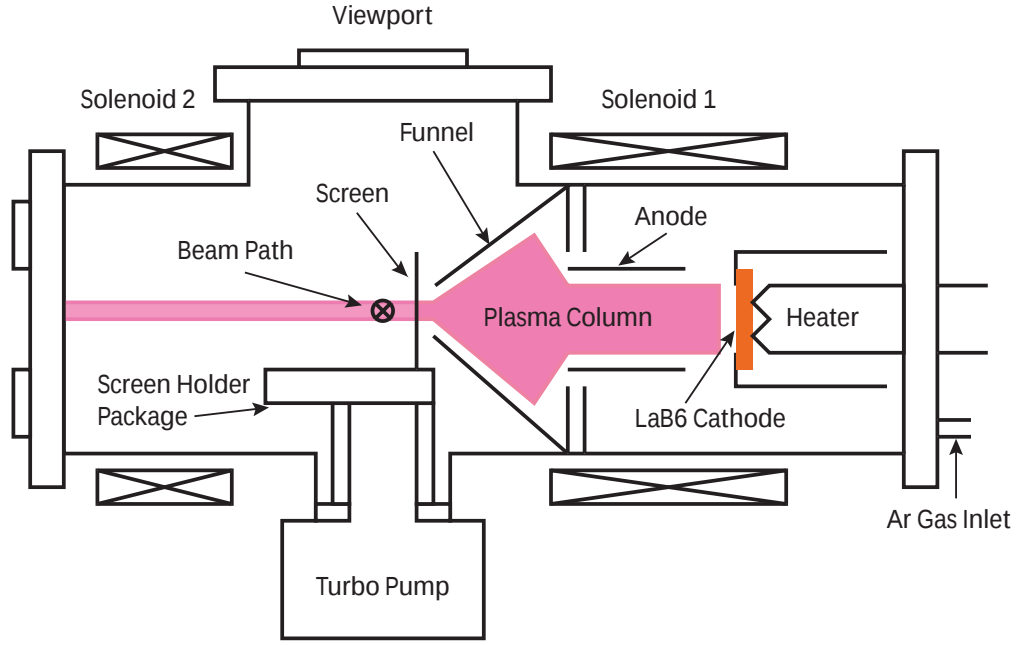


Figure 5.1: Schematic of the plasma density transition experiment plasma source.

neon light. Plasma is produced via a direct current electrical discharge in a diffuse argon gas. The discharge is created using a hot thermionic emission cathode, similar to the ones discussed in Section 1.3.1, and a cold metal anode. The space between the anode and cathode is filled with argon gas at pressure of about 1 mTorr. When a voltage is applied between the anode and cathode the electrons emitted from the cathode collisionally heat the argon into a plasma. An excellent tutorial on this class of plasma sources is given by Braithwaite [68].

5.2.1 General Source Design

The discharge cathode is a 7.5 cm diameter disc of Lanthium Hexaboride (LaB_6) which is heated to about 1300°C in order to give excellent thermionic emission. A nearby hollow anode is then pulsed with 100 volts (2 ms pulse duration, 0.5 - 1

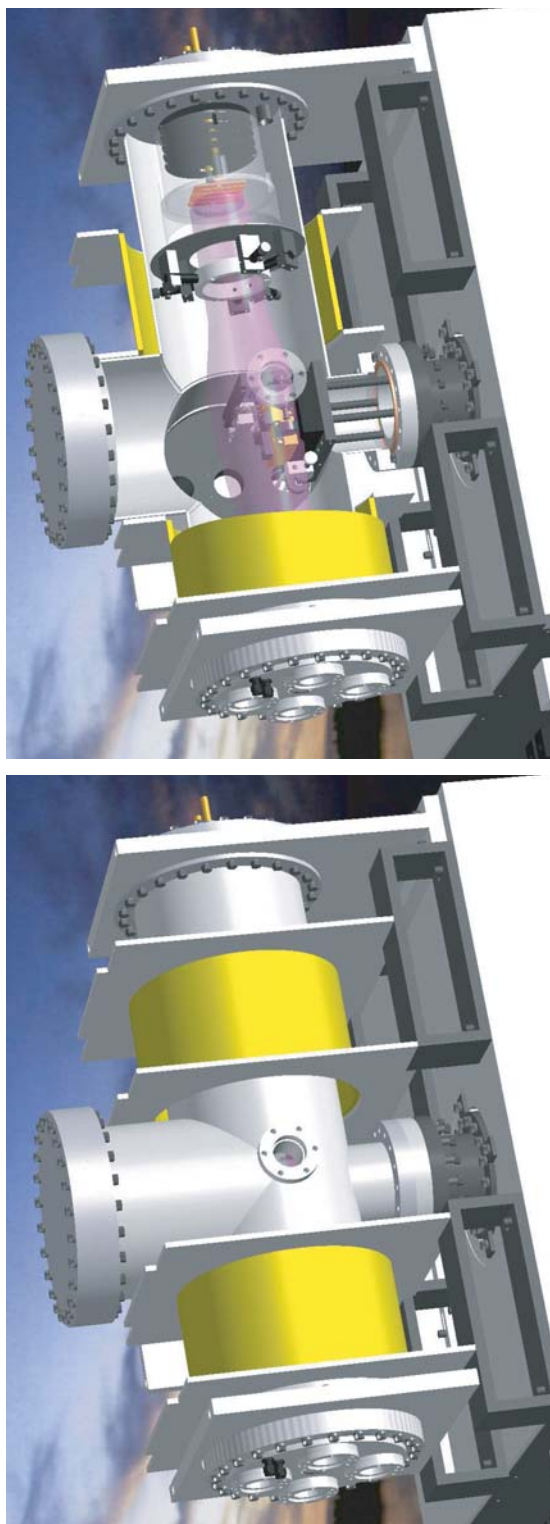


Figure 5.2: CAD rendered pictures of the plasma source. At left is the exterior of source's vacuum vessel and solenoids. At right is a cutaway view showing the internal structure of the source. Some parts, such as the funnel and heater housing, have been removed or made transparent for clarity.

Hz repetition rate) while the space between the anode and cathode is filled with an argon gas at 1 mTorr. The resulting 200 amp discharge produces an argon plasma with a temperature of about 7 eV. See Figures 5.1 and 5.2.

Once produced, the neutral plasma diffuses out through the hollow anode into the area where it will interact with the electron beam. Solenoids provide a weak magnetic field to help confine and guide this flow. The cylindrical plasma column that arrives at the electron beam interaction point has an approximately gaussian density profile with a peak density of about $3 \times 10^{13} \text{ cm}^{-3}$ and a width of about 5 cm FWHM.

5.2.2 Plasma Cathode Heater Design

Creating a system to reliably and repeatedly heat the cathode to 1300° C was a significant technical challenge. The experimental requirements for a wide plasma column set the large disc cathode geometry. This geometry is difficult to heat uniformly through direct resistive heating so black body radiation is used to heat the cathode. The basic theory of heat transfer through black body radiation is treated in Appendix B. Inductive heating was also a suitable option, but was deemed undesirably complex.

The radiation for heating the cathode is provided by a large surface area graphite heater positioned approximately 1 cm behind the back side of the LaB₆ cathode. The heater attains a temperature of over 1500° C when heated resistively with direct current at 235 amp, 16.6 volts. The vast majority of the 3.9 kW dissipated in this system leaves the heater as black body radiation. This radiation is emitted in every direction but a series of molybdenum reflector plates help keep it focused onto the back side of the cathode. The remainder of the heat deposited in the heater is lost through thermal conduction through the heater's

electrical connectors.

Making a survivable electrical connection to the graphite heater turned out to be one of the most technically challenging parts of the experiment. The electrical connection must fulfill three requirements: high electrical conductivity, low thermal conductivity, and tolerance to high temperatures. The design of the connector is dominated by the heat flow equation,

$$\frac{dQ_{input}}{dt} = \frac{dQ_{radiation}}{dt} + \frac{dQ_{conduction}}{dt}, \quad (5.1)$$

where the power dissipated by black body radiation is given by

$$\frac{dQ_{radiation}}{dt} = \epsilon \sigma S T_h^4, \quad (5.2)$$

and the power lost through conduction is given by

$$\frac{dQ_{conduction}}{dt} = \frac{k A_c (T_{tip} - T_{bath})}{l}. \quad (5.3)$$

Here ϵ is the heater surface emissivity, σ is the Stefan-Boltzmann constant, S is the surface area of the heater, T_h is the temperature of the heater in Kelvin, k is the thermal conductivity, A_c is the cross sectional area of the conductor, T_{tip} and T_{bath} are the temperatures of the tip of the conductor attached to the heater and the cooling water, respectively, and l is the length of the conductor. A good radiative heater design will have the characteristics,

$$\frac{dQ_{radiation}}{dt} \gg \frac{dQ_{conduction}}{dt}, \quad (5.4)$$

and

$$T_{tip} \approx T_h < T_{solidus}, \quad (5.5)$$

where $T_{solidus}$ is the temperature where the connector material begins to plasticize and the resulting loss of mechanical rigidity will destroy the electrical connection.

The solidus point is generally slightly below the melting point. For example, the melting point of copper is 1083°C while its solidus point is 1010°C .

Satisfying equations 5.4 and 5.5 simultaneously requires a very finely tuned system. The final heater design, as shown in Fig 5.3, was arrived at after a series of iterations guided by experience and analytic calculations using equations 5.1, 5.2, and 5.3. Computer simulations would have been a very useful tool in the design process, but there was no code available to us that would simulate all the necessary variables. It took several attempts to find a survivable material to make the fastening bolt out of. Early on we naively tried steel and titanium bolts which both melted. Bolts made of molybdenum (melting point 2617°C) seemed to work at first, but it quickly became apparent that the combination of high temperature and low pressure caused the molybdenum to sublime. We finally succeeded in finding a durable fastener by making the bolt out of tungsten (melting point 3410°C). The tungsten, in addition to having a high melting point, had a vapor pressure orders of magnitude lower than molybdenum, which prevents the problem of sublimation.

We also had problems with the tip of electrical connector, which was originally copper, over heating. In a typical failure mode the copper would reach its solidus point and deform, which lead to the electrical contact between the bolt and graphite heater degrading. This contact degradation in turn led to a runaway situation in which the degraded contact produced a higher heat load at the contact point which further degraded the contact, etc. The end result of the runaway was usually vaporization of a significant amount of the graphite heater and meltdown of the end of the copper connector. This problem was fixed by constructing the hybrid molybdenum / copper connector shown in Fig 5.3. The lengths, diameters, and material of each section was determined through care-

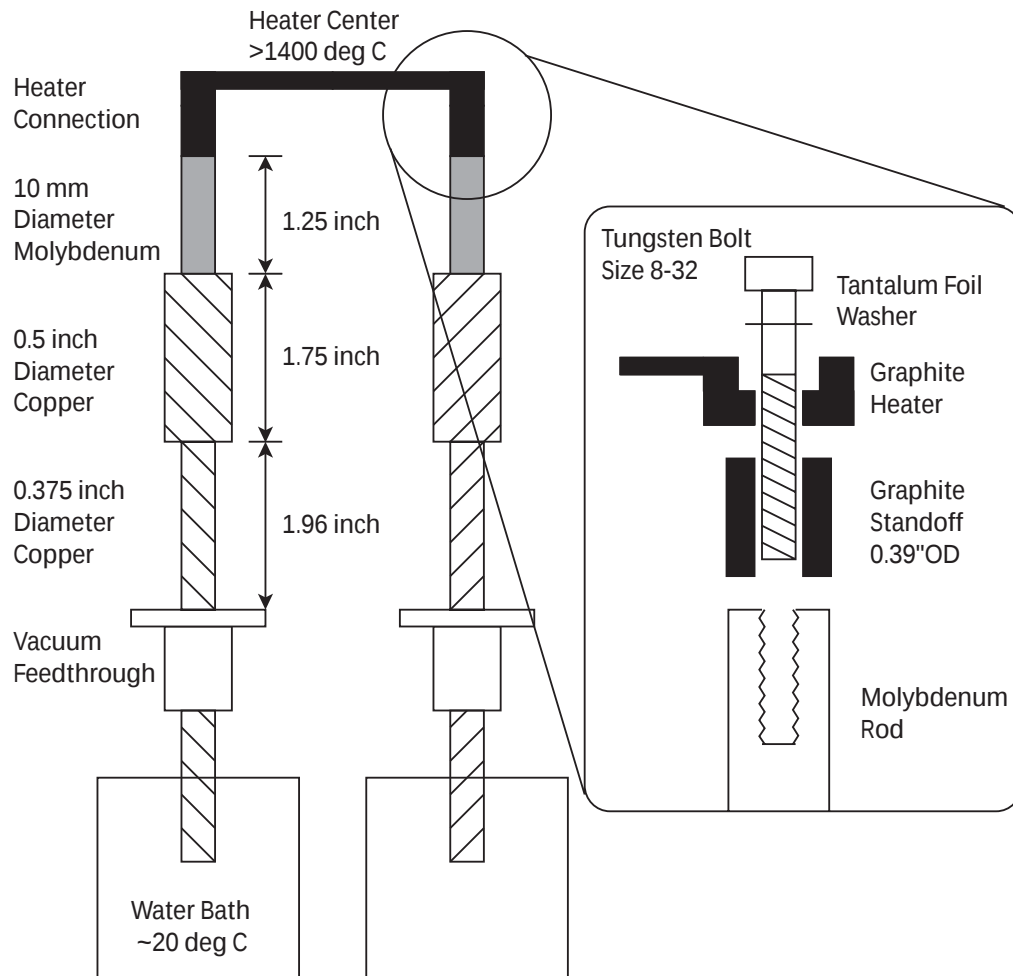


Figure 5.3: Schematic of the plasma source heater.

ful calculation to ensure that all the parts would remain below their respective melting, solidus, and sublimation points.

The final design shown in Fig 5.3 was highly reliable. The design did have one significant failing, however, in its reliance on tungsten bolts. While reliable once installed, these bolts are difficult to manufacture, expensive, fragile, and subject to becoming brittle during prolonged operation. In fact, once installed and operated at high temperature, the bolts cannot be unscrewed without a high

risk of the bolt head shattering.

5.2.3 Plasma Source Operation and Characterization

The plasma source was extensively tested over a large parameter space to determine the conditions for stable production of the maximum plasma density. The parameters relevant to this optimization are: argon gas pressure, magnetic confinement field strength and balance, discharge voltage, ballast resistor size, and cathode temperature. Plasma density was diagnosed using a single electrostatic Langmuir probe [69, 70]. This probe was typically mounted on a movable actuator so that the profile of the plasma density column could be measured.

The plasma discharge is powered by a 48 mF capacitor bank. This bank is charged to 200 V and then connected to the plasma anode using a high current switch (a large Darlington transistor). When the connection is made, the hot LaB₆ cathode starts emitting electrons, which ionize the 1.4 mTorr argon gas, and 200 A flows from the capacitor bank through the 0.22 Ω ballast resistor to the plasma. The ballast resistor stabilizes the discharge formation [71]. The peak power of 40 kW is maintained for 2 ms for a total pulse energy of 80 Joules. Under these conditions the cathode emits about 4.5 Amp/cm², which is a great deal larger than the expected vacuum emission at 1300° C which is about 0.1 Amp/cm² [72]. The difference is explained by the presence of the plasma which enhances the emission since its ions bombard the cathode surface causing additional heating and secondary electron emission. Since these effects dominate, the discharge current is relatively insensitive to cathode temperature. Once the cathode is above a certain threshold temperature it will emit enough electrons to initiate the discharge and relatively little is gained by higher cathode temperature after that point. At 1300° C the cathode is well into this stable saturation region.

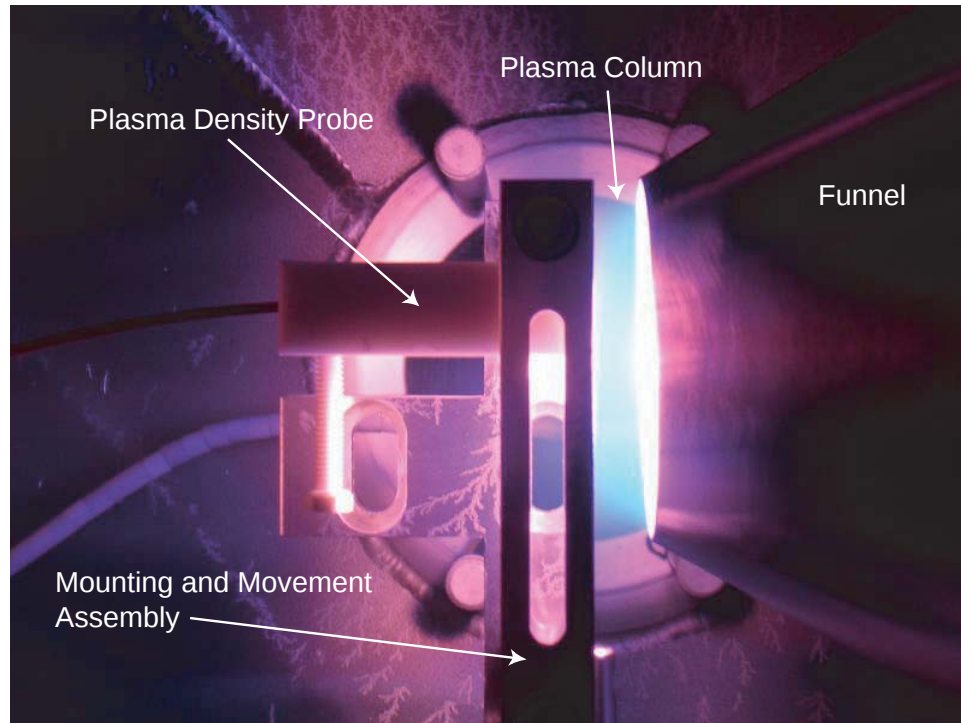


Figure 5.4: Photograph of the plasma column.

The solenoid magnet field, which is 60 gauss under normal operation, is critical to the discharge because it provides confinement and guidance for the plasma flow. The field does not, however, strongly confine the plasma to a tight column. In order to prevent the plasma from diffusing throughout the vacuum chamber a steel funnel was added to the system, see Fig. 5.1. The funnel ensures that the electron beam will encounter a well defined plasma column, which is important for producing the density transition, as will be discussed in Section 5.3.1. A photograph of the plasma column is shown in Fig. 5.4, and the measured plasma density profile of the column is shown in Fig. 5.5. The column is about 5 cm FWHM and has well defined boundaries. Its shape is that of a truncated gaussian. The average plasma temperature, as derived from the I-V curve measured by the Langmuir probe, observed during this measurement was 6.8 eV (see Appendix

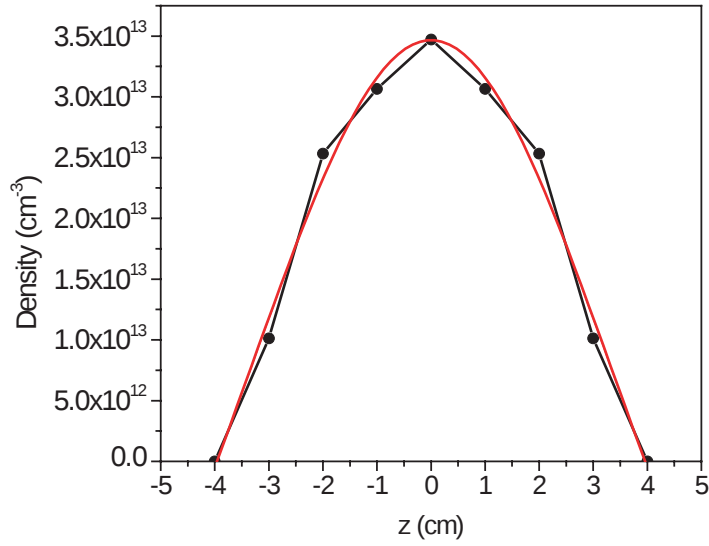


Figure 5.5: Measured plasma column density. The fitted curve is approximately the top half of a gaussian.

C for information on probe signal analysis). Density measurements are derived from the ion saturation current.

While the column displayed in Fig. 5.5 represents the typical operating configuration of the plasma source, it is possible to achieve higher peak densities. The highest peak density achieved with this source was $6 \times 10^{13} \text{ cm}^{-3}$. Producing this density requires the magnetic confinement field to be increased to about 100 gauss, which both increases the peak density and narrows the column to about 3 cm FWHM.

5.3 Creation of the Plasma Density Transition

Once the ability to reliably produce a high density plasma column had been established, the experimental development of density transition could begin in earnest. A large amount of theoretical effort went into the design of transition production

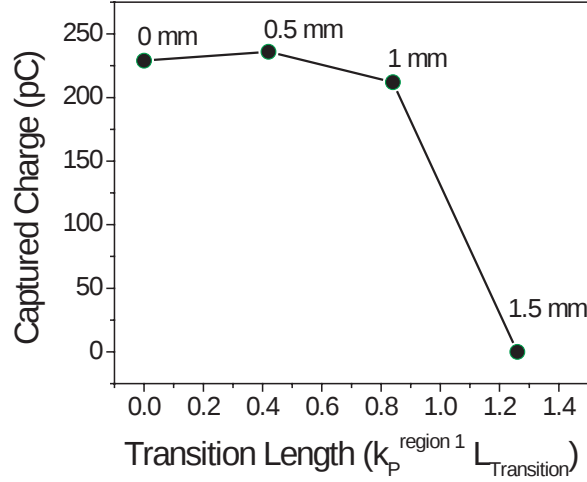


Figure 5.6: Simulated dependence of captured charge on transition length in the weak blowout case presented in Section 3.2. Each point is marked with the length of the transition.

mechanism before it was ever tested. This section begins with those theoretical efforts and ends with the successful demonstration of density transition creation.

Experimental realization of plasma density transition trapping depends on the creation of sharp density transitions. In the preceding chapters dealing with the theory and simulation of density transition trapping, the transition was assumed to be a perfect step function. Such perfection cannot be achieved experimentally. As long as the transition is short compared to the characteristic response length of the plasma, the plasma skin depth k_p^{-1} , the transitions will act like a step function. From this observation, the limit on the sharpness of the transition necessary to produce trapping can be written as the trapping condition;

$$k_p^{region 1} L_{Trans} < 1. \quad (5.6)$$

PIC simulations modelling the transition as a linear ramp of varying length support this condition. As can be seen from Figure 5.6, simulations predict that Eq. 5.6 is a very strict condition. The turn on of the capture in this regime is nearly a

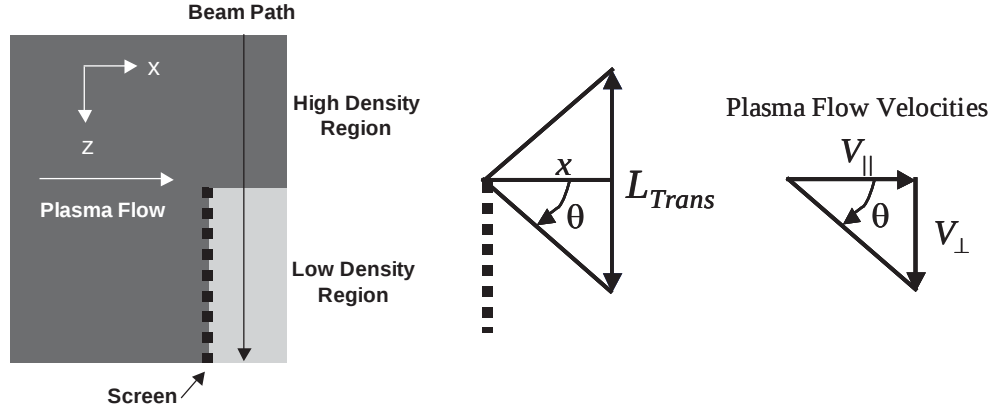


Figure 5.7: On the right is a simplified diagram of a plasma density transition produced by an obstructing screen. On the left is an illustration of the geometry of the plasma flow on the far side of the screen.

step function. Therefore, the primary goal in designing a mechanism to produce a plasma density transition is ensuring that Eq. 5.6 is satisfied.

5.3.1 Transition Production Using Metal Screens

As was discussed in Section 5.1, we choose to pursue metallic barriers and flow obstructions as the mechanism for producing density transitions. This direction of inquiry quickly lead to the idea of using a perforated metal masking screen. The basic concept of the masking screen operation is illustrated on the left in Fig. 5.7. Consider a system in which the plasma discharge is separated from the path of the driver beam. Once the plasma is created in the discharge apparatus it will diffuse and flow towards the beam path. If a perforated metal foil or grid of wires is placed in the path of the plasma flow, it will block a portion of the flow creating a low density region. Unfortunately, the plasma density transition will not remain sharp as the distance from the screen grows as portrayed in the simple picture on the left of Fig. 5.7. In reality, the two plasma regions will

diffuse into one another on the far side of the screen so that the plasma density transition will lengthen and “blur” as the distance from the screen edge increases. This process can be quantified using a simple model based on the velocities with which the plasma diffuses, as shown on the right in Fig. 5.7. On the far side of the screen from the plasma source, the high density plasma will continue to flow past the screen in the direction of the bulk plasma flow with a velocity $V_{\parallel}$ and will begin flowing into the low density region with a velocity $V_{\perp}$. The sum of these two vectors defines the line which marks the end of the transition into the low density plasma region. Symmetry dictates that the start of the transition in the high density region can be defined in the same way so that the total transition length is given by

$$L_{Trans} = 2x \tan \theta = 2x \frac{V_{\perp}}{V_{\parallel}}. \quad (5.7)$$

Since the plasma used for this experiment is weakly magnetized, it is reasonable to assume that the parallel and perpendicular plasma flow velocities are approximately equal. This assumption leads to the conclusion that

$$V_{\perp} \approx V_{\parallel} \rightarrow L_{Trans} = 2x, \quad (5.8)$$

which in turn leads to a new experimental constraint on achieving efficient trapping:

$$x < \frac{k_p^{-1}}{2}. \quad (5.9)$$

This new trapping condition for obstructing screens requires that the drive beam passes within half a plasma skin depth of the boundary. For a $2 \times 10^{13} \text{cm}^{-3}$ plasma the drive beam will have to pass within $600 \mu\text{m}$ of the screen. This is a reasonable level of expected pointing accuracy and stability.

We also explored the validity of this model of the transition through simulations. 2D MAGIC PIC simulations were created in which a plasma of mobile

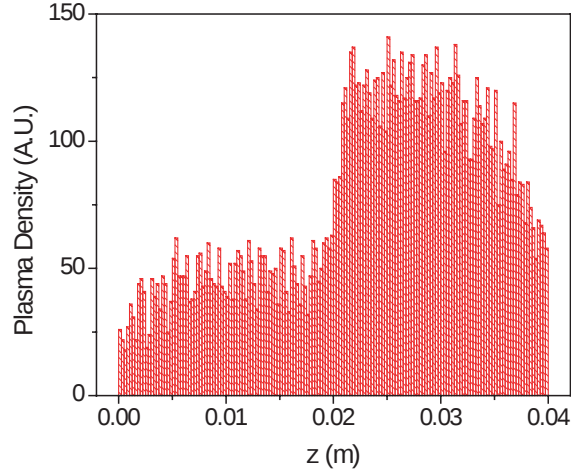


Figure 5.8: PIC simulation of a $2 \times 10^{13} \text{ cm}^{-3}$ plasma modified by a screen with $500 \mu\text{m}$ holes and 50% open area. Density is integrated over the $400 \mu\text{m}$ band between $100 \mu\text{m}$ and $500 \mu\text{m}$ away from the screen. The transition is less than 1 mm in length which equates to $k_p^{\text{region1}} L_{\text{Transition}} < 0.84$ at this plasma density.

elections and protons was initiated in half of a simulation area. A series of objects with metal boundary conditions was placed across the middle of the simulation area to represent the screen. As the simulation progressed, the plasma, which was initiated with a realistic temperature, diffused across the metal boundaries into the unoccupied area of the simulation. The profile of the plasma on the far side of the screen was then analyzed, see Fig. 5.8. The results of MAGIC PIC simulations of the screen-plasma interaction match the predictions of Eq. 5.8 almost exactly.

It should be noted that both the PIC simulations and the simple geometric model ignore collisional effects in the plasma. For a plasma like the one described in Section 5.2.3, with a temperature of 6.8 eV and a density of about $3 \times 10^{13} \text{ cm}^{-3}$, the mean free path between electron/ion collisions is on the order of 50 cm. Since the mean free path in this case is enormous compared to the transition

length scales, is it reasonable to neglect collisional effects when modelling the density transition.

5.3.2 Baffled Screens

Propagating a beam as close to a metallic screen as Eq. 5.9 requires can lead to other difficulties. Interactions with the screen over the entire length of the low density plasma region will completely disrupt the processes of trapping and acceleration. To circumvent this problem many alternative geometries were examined. The best solution was a screen with a solid metal baffle attached to its edge. As shown in at the bottom of Fig. 5.9, this baffle moves the sharp portion of the density transition away from the screen so that the beam and plasma wake will no longer interact with it. During the trapping process at the transition, however, the beam and wake still interacts with the baffle. The primary effect of the baffle is to block a portion of the particles participating in the plasma wake oscillation, as illustrated on the top in Fig. 5.9.

Simulating the effects of the baffle on particle trapping is a complex problem. The baffle breaks the cylindrical symmetry of the problem requiring that any simulations of its effects must be done in three dimensions. The three dimensional version of the PIC code MAGIC was used to simulate this problem. The trapping system was modelled with a metallic baffle at various distances from the beam center. The results of these simulations are summarized in Fig. 5.10. The points in the graph are taken from simulations in which the simulation cells are $0.17k_p^{-1}$ on a side. This level of resolution is considered to be close to the minimum necessary to accurately model the plasma. Higher resolution was not practical because of technical restrictions. MAGIC 3D was not a parallel code and the simulation described above required all the memory available on our

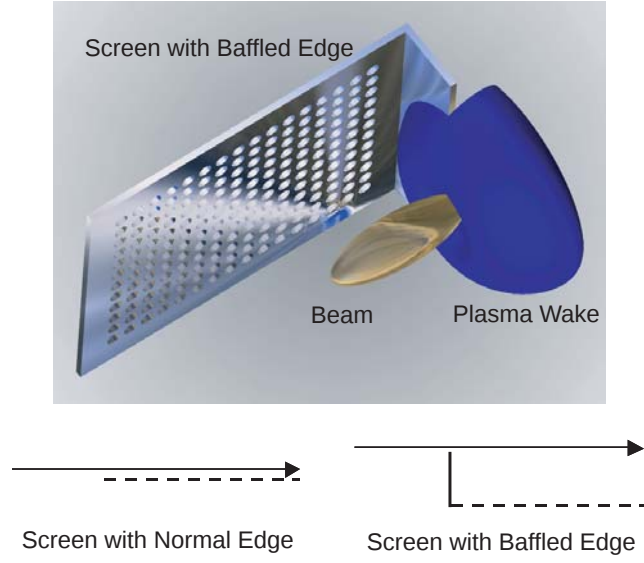


Figure 5.9: Artist's conception of partial blocking of wake particles by the baffle. The prolate spheroid represents the electron beam and the toroid that partially intersects the screen baffle is the plasma wake.

single processor system with 2GB of physical memory. Doubling the simulation resolution would multiply the number of variables by 2^3 . The physical memory needed to store these variable would likewise increase by a factor of 8 to 16 GB. This amount of memory was not technically or economically practical for us. There were no alternative codes capable of simulating the problem available to us at the time.

Since Eq. 5.9 indicates that the beam must pass within $k_p^{-1}/2$ of the baffle edge, the results shown in Fig. 5.10 predict an approximate 50% loss of total captured charge. This may not translate into a 50% loss of particles in the beam core, however, since the large amplitude particles blocked by the baffle are not necessarily the ones that form the beam core. The 3D simulations lacked the resolution to resolve this question.

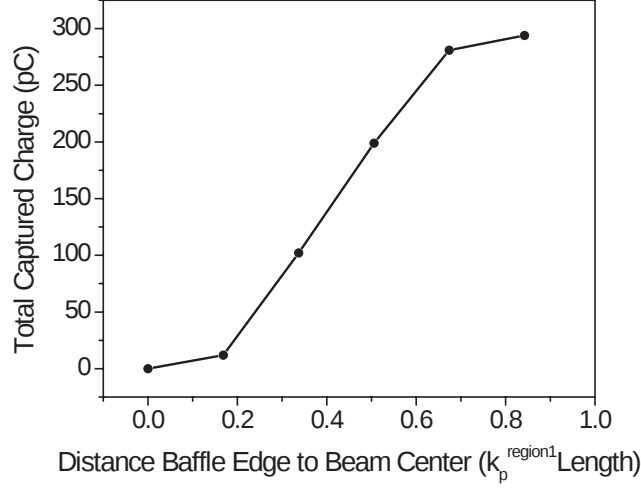


Figure 5.10: Effect of the beam-baffle distance on trapping. These results are from 3D PIC simulations of weak blowout case discussed in Section 3.2.

Another issue with the use of screen produced plasma density transitions is the rapid growth of the transition length with distance from the screen. The growth rate is large enough that there will be a significant transition length gradient over the distance spanned by the plasma wake field. The effect of this transition length gradient is unknown and difficult to simulate. Since the transition will still fulfill the trapping condition, Eq. 5.6, over a substantial portion of the wake, the variation in transition length should be fairly unimportant. We expect this effect to produce another minor degradation of the trapping performance.

5.3.3 Density Transition Characterization

A series of direct measurements of the plasma density transition were made in order to demonstrate the existence of transitions sharp enough to exhibit trapping. The peak density produced by the plasma source is $3 - 6 \times 10^{13} \text{ cm}^{-3}$, which corresponds to a skin depth of 1 - 0.7 mm. As stated in Eq. 5.6, the transition must

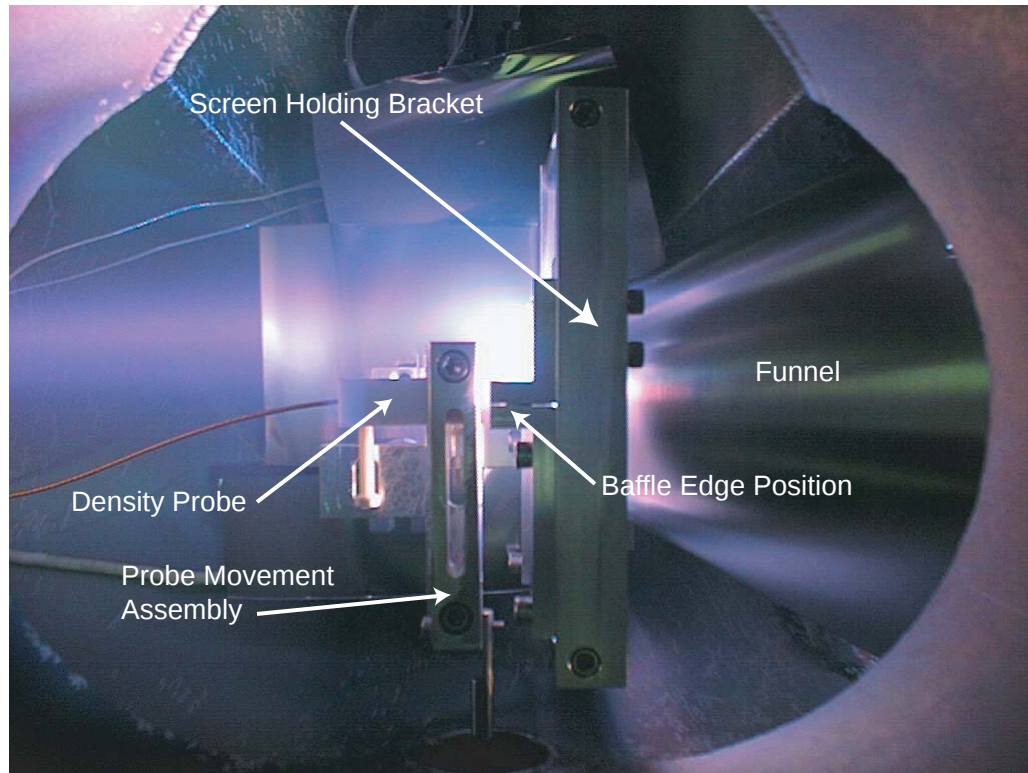


Figure 5.11: Photograph of the plasma density transition apparatus.

be shorter than the skin depth of the peak density. At these densities, transitions shorter than a skin depth can be resolved using a conventional electrostatic probe [69, 70].

The probe used for the transition measurements was constructed from a 0.38 mm diameter tungsten wire that had its end ground flat and mounted flush in the face of a ceramic block. The face of the ceramic block was mounted perpendicular to the plasma column on a precision linear actuator that moved the probe through the plasma on the design path of the electron beam. Since the probe was flat it only sampled one plane of the plasma at once, which is important since the plasma transition has a different length depending on the x plane it is sampled on, see Fig. 5.7. The width of the probe was only half of the shortest expected

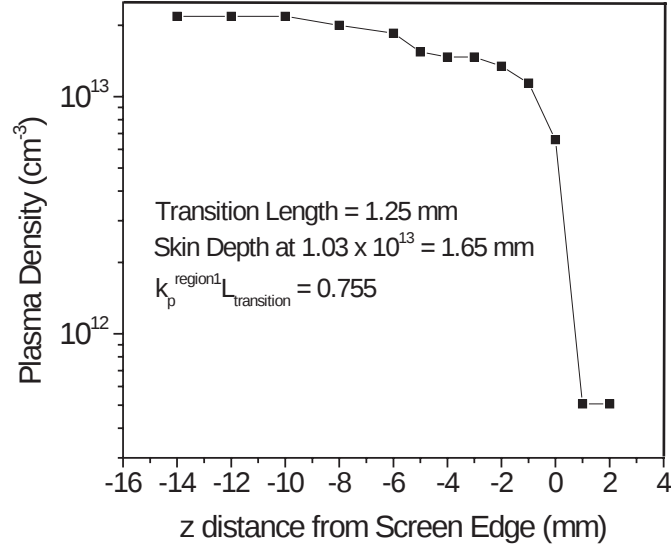


Figure 5.12: Measured density transition - long scale.

skin depth ensuring sufficient resolution to measure sub- k_p^{-1} transitions. The actual screen used to produce the density transition had a configuration exactly like the one shown on Fig. 5.9. It is made out of micro-perforated 42 gauge (78 μm thick) stainless steel sheet with 152 μm holes and 21% open area. The screen and probe apparatus is shown in Fig. 5.11.

Measurements of the transition were made both by moving the baffle edge past a stationary probe mounted in the center of the plasma column, and moving the probe past a stationary screen. These two types of measurements yielded nearly identical results. Fig. 5.12 shows a typical density transition measurement. Note the gradual roll off in density before the sharp transition begins. The peak plasma density of $2.2 \times 10^{13} \text{ cm}^{-3}$ is reduced to $1 \times 10^{13} \text{ cm}^{-3}$ before the step transition begins. This pre-transition density roll off is a feature common to all of the transition measurements. This effect was not anticipated and its origin is unclear. One possible cause for this feature of the transition is the plasma sheaths that form on the metal surfaces of the baffle and screen. While the

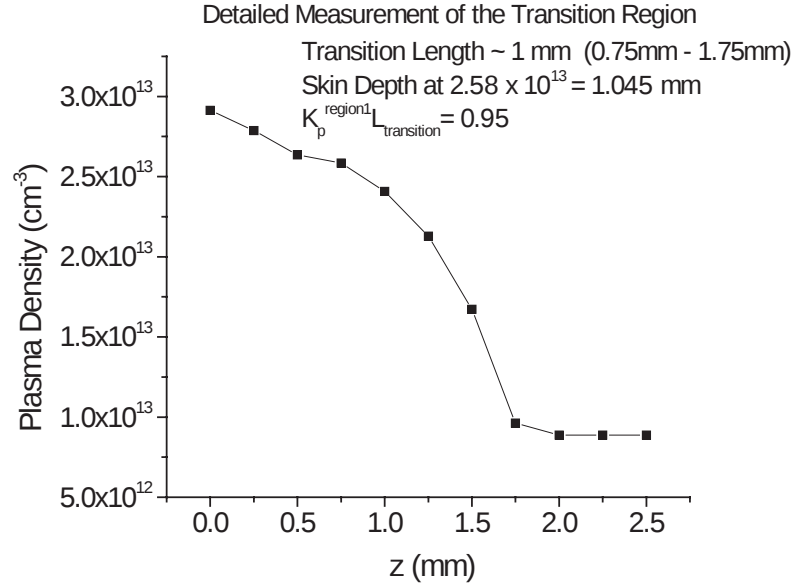


Figure 5.13: Measured density transition - short scale.

actual plasma sheath is quite thin, the typical thickness is less than a skin depth, the pre-sheath reaches much further into the plasma. Effects associated with the pre-sheath might partially blur the baffle edge and produce the gradual density roll off.

The gradual density roll off before the transition reduces the peak density at the top of the steep transition, but it does not impact the sharpness of the density drop. In the example shown in Fig. 5.12, the transition from $1 \times 10^{13} \text{ cm}^{-3}$ to the minimum plasma density occurs in 1.25 mm. Since the skin depth of a $1 \times 10^{13} \text{ cm}^{-3}$ plasma is 1.68 mm this transition fulfills the trapping condition since $k_p^{\text{region1}} L_{\text{Trans}} = 0.74 < 1$.

Finer measurements of the steep transition region were also made. An example of one of these measurements is shown in Fig. 5.13. The measurement in Fig. 5.13 was made with a moving probe under the exact same plasma operating parameters used during the trapping experiment described in Chapters 6 and 7.

The plasma operating parameters were the same as the ones used to produce the raw plasma column shown in Fig. 5.4. The measurement shows that the steepest part of the transition is about 1 mm in width and has a peak density of $2.58 \times 10^{13} \text{ cm}^{-3}$ and a minimum density of $8.9 \times 10^{12} \text{ cm}^{-3}$. Since the skin depth of a $2.58 \times 10^{13} \text{ cm}^{-3}$ plasma is 1.045 mm this transition fulfills the trapping condition since $k_p^{region1} L_{Trans} = 0.95 < 1$. The transition may actually be sharper than this since some blurring due to the probe size can be expected at this level of resolution.

Another factor in the transition measurements is the spacing between the probe and the baffle edge, as discussed in Section 5.3.1. For the measurement in Fig. 5.12, the spacing between probe and baffle was measured to be between 506 μm and 930 μm . The range of measurements originated from a slight transverse instability in the motion of the precision linear actuator. From Eq. 5.8 the expected range of transition lengths given this range of spacings is 1.12 mm - 1.86 mm. The measured value of 1.25 mm for the transition length agrees reasonably well with the prediction of Eq. 5.8. The measurements of spacing and transition length made during the scan shown in Fig. 5.13 exhibited a similar level of agreement. It is therefore reasonable to expect that the transition will be shorter if the probe/baffle spacing is reduced, as predicted by Eq. 5.8. An attempt to make precision measurements of the transition length at different probe/baffle spacings failed because the in-vacuum picomotors designed to actuate the baffle failed in the high temperature environment of the plasma source.

Finally, it should be noted that for a density masking screen to operate correctly there needs to be some sort of barrier placed behind the screen and beam path to prevent back diffusion of the plasma. The purpose of such a barrier is the same as the one that lead to the diffusion control funnel discussed in Section

5.2.3. If no barrier is present, the high density plasma from the portion of the column that is not screened will eventually diffuse back behind the screen, ruining the density modulation. The solution used in this case was to wall off the region behind the density screen, except for a small slit of about 5 mm through which the beam could pass, with thick metal foil. Some of the experiments with the plasma source indicate that a metal backstop much larger than plasma column will prevent diffusion just as effectively if it is placed directly behind the screen. The metal backstop cools the plasma and forces it to recombine into gas before it has a chance to diffuse.

5.4 Simulation of Trapping Performance Under Realistic Experimental Conditions

Section 5.3 documents the imperfections of the density transitions that were produced experimentally and the impact these imperfections are likely to have on trapping performance. There are many other ways in which a real experiment based on the argon pulse discharge plasma source will deviate from the parameters presented for a weak blowout experiment in Section 3.2. These differences stem from the density profile available in the plasma source; the field necessary to guide the plasma in the source; and, the imperfections of real drive beams. An effort was made to understand the consequences of these differences through simulation.

5.4.1 Realistic Density Profiles

It is difficult to produce a plasma density profile with the smooth linear dependencies shown on the right in Fig. 3.4, and again in the lower curve of Fig. 5.14. A more realistic option is to alter the natural profile of the plasma column as

Table 5.1: Driving and captured beam parameters for the two density profiles shown in Fig. 5.14.

Driving Beam Parameters	Profile 1	Profile 2
Beam Energy	14 MeV	14 MeV
Beam Charge	5.9 nC	5.9 nC
Beam Duration σ_t	1.5 ps	1.5 ps
Beam Radius σ_r	362 μm	362 μm
Normalized Emittance	15 mm-mrad	15 mm-mrad
Peak Beam Density	$4 \times 10^{13} \text{ cm}^{-3}$	$4 \times 10^{13} \text{ cm}^{-3}$
Captured Beam Parameters	Profile 1	Profile 2
Beam Energy	1.2 MeV	1.5 MeV
Beam Charge	100 pC	470 pC
Beam Duration σ_t	1.7 ps	0.3 ps
Beam Radius σ_r	250 μm	100 μm
Normalized Emittance	24 mm-mrad	16 mm-mrad
Energy Spread (rms)	4%	4%

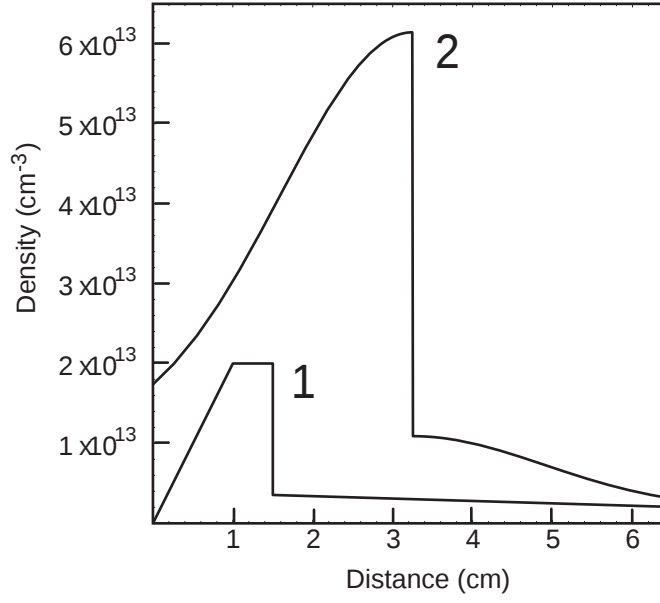


Figure 5.14: Comparison of gaussian and linear plasma density profiles.

little as possible. As can be seen from Fig. 5.5, the raw plasma column produced in our plasma source has a truncated gaussian profile with a higher peak density than we originally anticipated. If we use a simple screen of uniform open area to reduce the amplitude of half the gaussian distribution, we produce the profile labelled number 2 in Fig. 5.14. Note that the profile in Fig. 5.14 is derived from the case of maximum peak density operation described in Section 5.2.3 rather than the more moderate density operation shown in Fig. 5.5.

The gaussian based profile is qualitatively similar to the linear profile and preserves most of its features including the gradual decline in density after the transition. This gradual density decline is critical for enhancement of charge capture and the reduction of energy spread, as discussed in Section 3.2. Simulations using this new cut-off gaussian profile indicate that its performance is superior, especially in terms of captured charge, to that of the original linear density profile. The gain in captured charge is primarily due the higher peak plasma density.

The drive and captured beam parameters from these simulations are presented in Table 5.1 along with results from the linear profile for comparison. While the gaussian density profile is superior from an experimental standpoint, operation at the maximum peak plasma density is not necessarily desirable due to the high magnetic confinement field it requires.

5.4.2 Impact of Transverse Magnetic Field

The magnetic field used for plasma confinement is another problem associated with performing a density transition trapping experiment using the argon pulse discharge plasma source. As shown in Fig. 5.1, the source has two solenoids that generate a substantial field, 60 - 100 gauss on axis, to confine the plasma. This field is perpendicular to the electron beam path through the experiment, and it will cause the beam path to bend during the trapping process. This steering is a problem since the high energy drive beam and low energy trapped electrons will bend with different radii of curvature in this field. This effect could potentially move the trapped electrons out of the accelerating wake and destroy the trapping process.

Simple estimates suggest that the transverse electric field forces that act to keep the trapped electrons in the plasma wake are much stronger than the magnetic steering forces for the values of interest. The three dimensional geometry of this effect requires a three dimensional code to simulate it. The 3D PIC code OSIRIS [73] was used to examine the effect of bending while trapping. The original weak blowout case with a bi-gaussian driver, as described in Section 5.4.1, was simulated with a uniform transverse magnetic field throughout the whole volume of the plasma. At the time of these simulations, OSIRIS did not have the capability to simulate the semi-gaussian profile of the actual magnetic field. In the

case of a 100 gauss transverse field only about 7 pC of charge was captured, as opposed to 100 pC in the zero field case. When the field was reduced to 80 gauss charge capture increased to 15 pC. The captured charge loss indicated in these simulations is overstated since the bending field remains constant throughout the trapping process rather than falling off after the transition. Even though the OSIRIS simulations overstate the charge loss, they clearly indicate that charge loss due to bending while trapping is a significant effect. It is therefore prudent to reduce the plasma magnetic confinement field to as low a level as possible during trapping experiments.

5.4.3 Imperfect Driving Beams

As was stated at the beginning of this chapter in Section 5.1, this experiment was designed around available electron beam parameters. All the drive beam parameters listed in Table 5.1 had been achieved in various laboratories when development of the experiment started, e.g. Table 6.1. It was anticipated that of all the parameter requirements, only the values prescribed in the design for the emittance and bunch length would be challenging to meet. This is because bunch length, emittance, and charge are coupled in complex ways. In general, it is difficult to produce a beam that simultaneously has the charge, bunch length, and emittance listed in Table 5.1 with current photoinjector / magnetic compressor systems. Since higher than design values for emittance and bunch length might be unavoidable, it was important to understand the effect these conditions would have on trapping performance.

The performance of transition trapping does not depend significantly on the driver emittance, as long as the emittance remains within the same order of magnitude as the design. This is because the emittance does not really impact

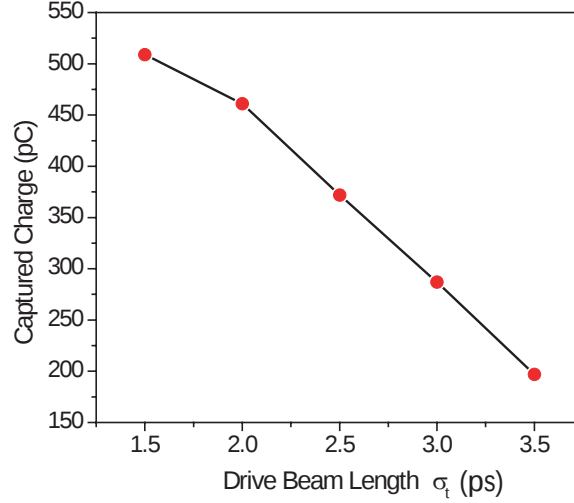


Figure 5.15: Impact of drive beam length on trapped charge.

trapping performance until it become so large that the specified beam spot size cannot be achieved. It was anticipated that the real emittance of drive beam might be two or three times larger than the design value of 15 mm-mrad. Emittances in this range are not high enough to prevent the attainment of the design $362 \mu\text{m } \sigma_r$ spot size.

Small changes in bunch length, however, strongly impact trapping performance because they affect the strength of the plasma wake fields. Since the real bunch lengths might be a few ps longer than planned, the effect of bunch length on trapped charge in the gaussian profile case presented in Section 5.4.1 was simulated, see Fig. 5.15. The amount of charge captured was found to go down linearly with growth in the bunch length. The bunch length could therefore have a serious detrimental effect on trapping if the design value is not achieved.

5.5 Comparison of Design Goals and Achieved Results

The fundamental goal of the plasma density transition trapping experiment design was to create a plasma with a density transition that would produced trapped electrons when excited by a drive beam with previously achieved parameters. This goal was achieved as demonstrated by the measurements in Section 5.3.3. There are, of course, many ways in which the density transition created is imperfect and deviates from an ideal step function. The effect of each of these imperfections on the performance of trapping was simulated to the greatest degree practicable with the codes available.

The development of the plasma density transition trapping experiment started from the nominal experimental weak blowout case described in Section 3.2. When all the factors described in this chapter were considered, a new plan for the experiment emerged. This new plan relied on the same drive beam as the weak blowout case describe in Section 3.2, but used a gaussian plasma profile like the one describe in Section 5.4.1. It was decided to operate the plasma at peak density of $3.5 \times 10^{13} \text{ cm}^{-3}$ with an on axis confinement field of 60 gauss, rather than at high density and field, in order to mitigate the effects of magnetic steering on trapping, see Section 5.4.2. The pre-transition density roll off observed in the system, see Section 5.3.3, reduced the peak plasma at the top of the transition to $2.4 \times 10^{13} \text{ cm}^{-3}$. Taking these facts into account, the plasma profile actually used in the experiment is very close to the one shown in Fig. 5.16.

The simulated trapping performance of the experimental plasma profile is very similar to that of the higher peak density gaussian plasma profile shown in Fig. 5.14, except with regard to charge, when it is driven by the same drive beam, see Table 5.1. With a 5.9 nC drive beam, the profile in Fig. 5.16 produces 126 pC of trapped charge at 1.5 MeV. Note that this value for the trapped beam

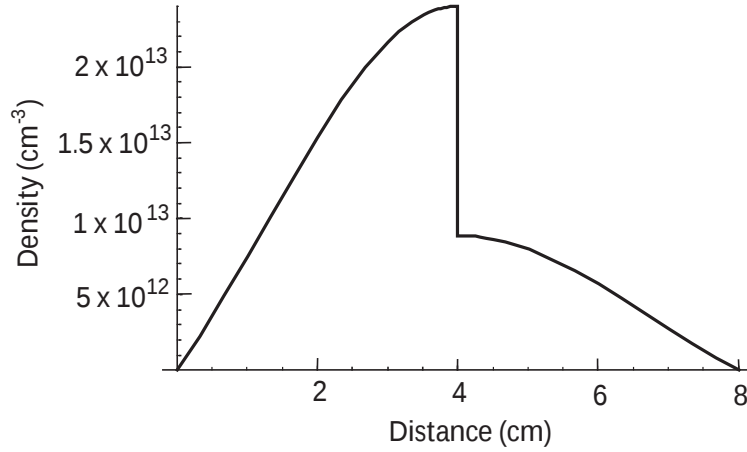


Figure 5.16: The density profile used in the experiment.

charge does not include any losses due to effects like those described in Sections 5.3.2 and 5.4. When losses due to the baffle edge, bending while trapping, and longer than anticipated drive beams are combined it is estimated that about 10 % of the charge listed above, or 12.6 pC, will actually get trapped. As shown in Fig. 4.1, the fast scaling of captured charge with increases in the drive charge can be used to make up for these losses. For example, if the drive beam charge is increased to 8 nC in the simulation of the actual experimental density profile, without altering any other parameter, 548 pC is trapped. Even after 90 % losses, the 8 nC case would still produce about 55 pC of trapped charge, which is not unreasonable from a detection standpoint. The existence of pulses with charges up to 18 nC at the FNPL provided great confidence that the losses caused by flaws in the density transition, no matter how severe, could be overcome by going to higher drive beam charge.

CHAPTER 6

Execution of the Transition Trapping Experiment

The plasma density transition trapping experiment was performed at the Fermilab NICADD Photoinjector Laboratory (FNPL) as part of our larger UCLA/FNAL collaboration on PWFAs. The plasma source (see Chapter 5) and the specialized spectrometer designed for the experiment (see Section 6.3.5) were both constructed at UCLA. The apparatus for the experiment was then shipped to the FNPL and installed as a unit on the end of the existing beamline. The FNPL facility, its capabilities, and the existing and new diagnostic tools used to perform the trapping experiment are described in this chapter.

6.1 The FNPL Facility

The FNPL accelerator is a 18 MeV electron linac [44]. The system consists of a normal conducting L-band RF photoinjector with a cesium telluride photocathode and a 9-cell superconducting accelerating cavity. Bunches with charge in excess of 8 nC can be produced and compressed to $\sigma_t = 1.6$ ps using magnetic compression. The parameters of the system are listed in Table 6.1. A drawing of the beam line is provided in Fig 6.1.

Table 6.1: Nominal FNPL Operating Parameters [44]

Before Compression	
Bunch Charge	8 nC
Laser Pulse Length FWHM	28 ps
Beam Radius at Cathode	3 mm
Peak Field on Cathode (nominal)	35 MV/m
Beam Total Energy	18 MeV
Transverse Emittance Normalized	11 mm-mrad
Longitudinal Emittance (100% rms)	820 deg-keV
Momentum Spread	4.2%
Bunch Length	4.3 mm
Peak Current	276 Amp
After Compression (Theoretical)	
Transverse Emittance Normalized	15 mm-mrad
Bunch Length	1 mm
Peak Current	958 Amp

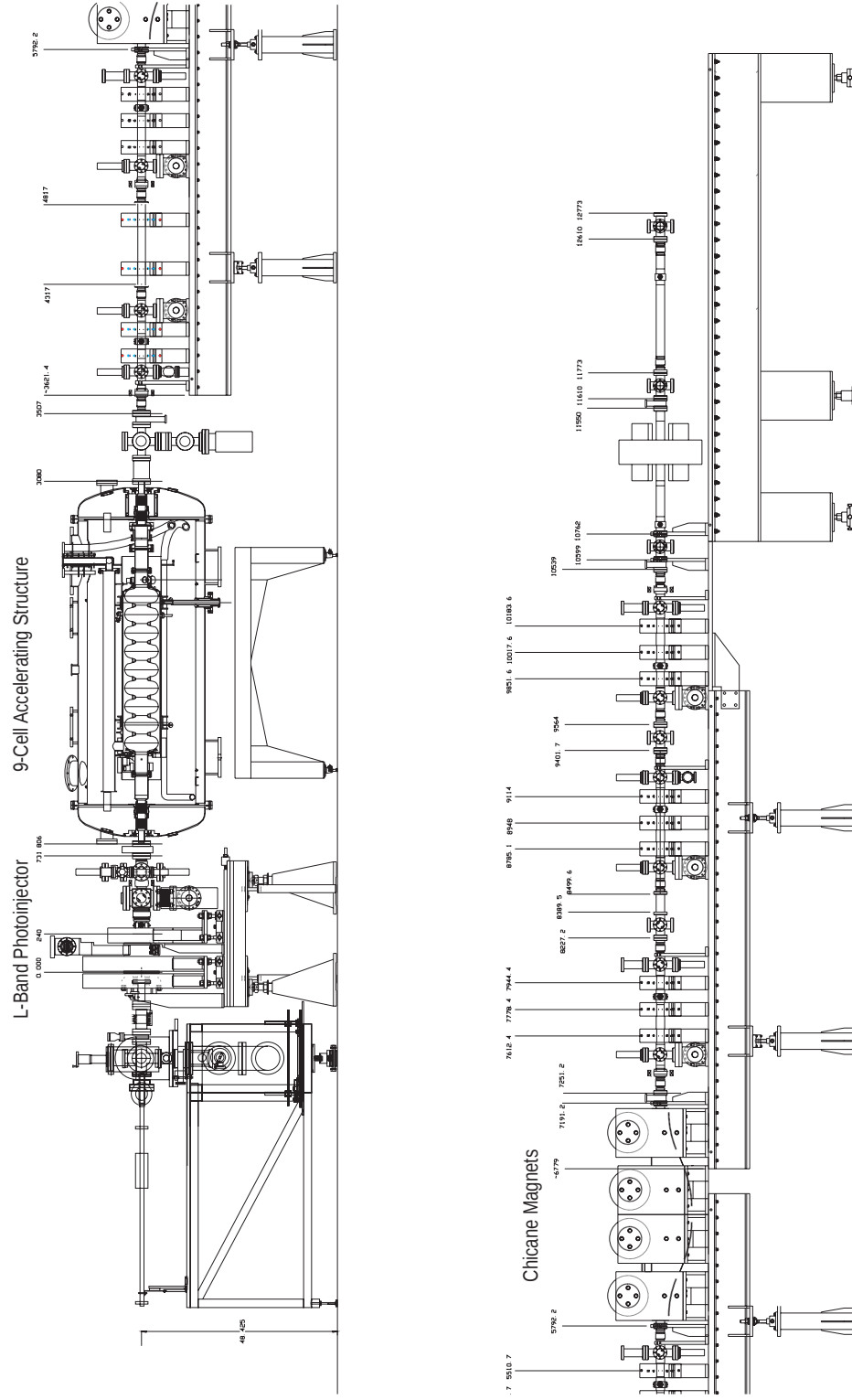


Figure 6.1: The FNPL photoinjector primary beamline

6.1.1 Cathode Drive Laser

Most of the requirements for a photoinjector drive laser are stated in Section 1.3.2. Drive lasers need sufficient photon energy to excite photo-emission, must be synchronized to RF wave, have adjustable timing of arrival at the cathode, and be fairly stable in terms of timing, pointing, and energy. At the FNPL, the driving laser is provided by a Nd:YLF oscillator followed by stages of chirped pulse amplification (CPA) and fourth harmonic generation [74].

The sequence of stages in the production of the drive laser begins with the solid-state Nd:YLF laser oscillator which produces a continuous train of low energy pulses (20 nJ/pulse, 100ps FWHM pulse length, $\lambda = 1054$ nm) at 81.25 MHz. This pulse train is synchronized to the RF by mode-locking of the laser using an acoustooptical modulator that is phase-locked to the FNPL master oscillator. The 1.3 GHz RF, which is the 16th harmonic of 81.25 MHz, for the photoinjector and accelerating cavity is also derived from this master oscillator. The relative timing of the RF and laser at the cathode is achieved by phase shifting the RF. The measured timing jitter of this system is ≤ 2 ps rms.

The pulses produced by the oscillator must be both amplified and shortened before they are ready to be sent to the photoinjector. These effects are achieved through chirped pulse amplification. In this process the laser pulses from the oscillator are sent through 2 km of dispersive optical fiber which stretches and “chirps” the pulse, i.e. sets up a correlation between frequency and position. The fiber also increases the bandwidth of the oscillator pulse through the non-linear process of self-phase modulation so that the final pulse has sufficient bandwidth to eventual be compressed to a shorter length than the input pulse. A subset of the 81.25 MHz pulse train, between 1 and 800 pulses, is selected, using high speed Pockel’s cells, at 1 Hz for amplification. The subset of chirped pulses is then

amplified in a series of flashlamp pumped Nd:glass gain stages. Once amplified the chirped pulses are compressed using a set of parallel diffraction gratings. By adjusting the grating separation, the length of the compressed pulse can be varied from 3 to 30 ps FWHM. Under typical conditions, each pulse emerges from the CPA stage with a length of 20 ps FWHM and an energy of 0.6 mJ.

The wavelength of the photons emerging from the CPA is still 1054 nm which is too long for exciting photo-emission. The pulses are therefore converted from 1054 nm (IR) to 527 nm (green), and then 263 nm (UV) using two non-linear Beta Barium Borate (BBO) crystals. This is called “fourth harmonic generation” which provides photons of increased energy at the expense of pulse energy. Since the conversion efficiency is a non-linear function of the laser intensity, the center of the gaussian pulse is converted more efficiently than the head and tail; which results in pulse shortening. The laser pulses that leave the fourth harmonic stage are 10 ps FWHM long with a design energy of 30 μ J; 10 μ J is more typical in regular operation. The UV laser pulses are transported ~ 20 m to the RF photoinjector through a high efficiency evacuated transport line. The pointing jitter at the cathode was measured to be about 140 μ m rms, which works out to be a 7 μ rad rms angular pointing jitter [75].

6.1.2 L-Band RF Photoinjector

The FNPL photoinjector is a 1.625-cell, iris-coupled, standing wave structure, operated in the π -mode at L-band frequency (1.3 GHz) [76]. Its design incorporates split focusing solenoids that allow the gun to operate in the high quality emittance compensated mode [77, 78] over a wide range of gradients. Perhaps the most important feature of the FNPL photoinjector, at least from the standpoint of the plasma density transition trapping experiment and its high charge

needs, is the cesium telluride (Cs_2Te) photocathode. Cs_2Te photocathodes have a *quantum efficiency*, the percentage of photons striking the cathode that produce electrons, of $\geq 1\%$ and have produced 50 nC pulses in photoinjectors of the same design as FNPL's [79]. Pulses as large as 18 nC were observed during the plasma trapping experiment at the FNPL. Pulses exit the photoinjector with a total energy of about 4 MeV [44].

The L-band photoinjector at FNPL is essentially a scaled up version of the UCLA S-band photoinjector discussed in Section 1.3.2. The S-band source uses 10 cm wavelength microwave radiation while the L-band source uses 23 cm microwaves. The size of these two structures is comparable to the wavelengths they use. It is interesting to compare Tables 1.2 and 6.1 and note that the trends discussed in Sections 1.1 and 1.3 are reflected in even the relatively small difference in scale between the S-band and L-band sources. The smaller source produces higher accelerating gradients and lower emittances than the larger source.

6.1.3 Nine-Cell Superconducting Accelerating Cavity

The 9-cell superconducting cavity, which follows the photoinjector, adds the bulk of the FNPL electron beam's energy. This accelerating structure is one of the cavities built to test the feasibility of TESLA [80], the superconducting accelerator option for the Next Linear Collider (NLC). It is a standing wave structure operated in the TM_{010} , π -mode at 1.3 GHz with an active length of about 1 m and nominal design gradient of 23.4 MV/m. The cavity is constructed of niobium and is cooled to its superconducting operating temperature of 2 K in a superfluid helium bath. The chief advantage of using superconducting accelerator structures is the absence of resistive wall heating, which allows for efficient high average power operation.

The particular cavity in use at the FNPL has an anomalously low quench field believed to be caused by contamination in the cavity welds. This limits the cavity to a gradient of 13 MV/m in CW operation [44]. Slightly higher gradients can be attained in pulsed operation. The phase of the accelerating wave in the structure can be adjusted in order to produce varying levels of magnetic compression using the downstream chicane magnet set, see Fig. 6.1 and Section 1.3.3.

6.2 The Trapping Experiment Beamline

A special beamline was constructed for the plasma density transition trapping experiment. This beamline extended from the end of the primary beamline shown in Fig. 6.1. This type of installation was chosen to facilitate the space needs of the trapping experiment and the need for vacuum isolation between the experiment and the rest of the accelerator. As shown in Fig. 6.2, the beamline consists of a strong focusing solenoid, a vacuum isolation window, two quadrupole magnet triples that first capture the beam and then focus it into the plasma, the plasma chamber, and the spectrometer. Several steering magnets are provided for correcting the beam trajectory. A host of diagnostics are also included, as discussed below.

6.2.1 Vacuum Window

The design of the beamline directly before the trapping experiment is dominated by the need for a robust window to isolate the vacuum of the experiment from the upstream beamline. As discussed in Chapter 5, the plasma source for this experiment requires a constant backfill of 1.4 mTorr of argon gas in its vacuum chamber. Many of the other components of the A0 photoinjector beamline, es-

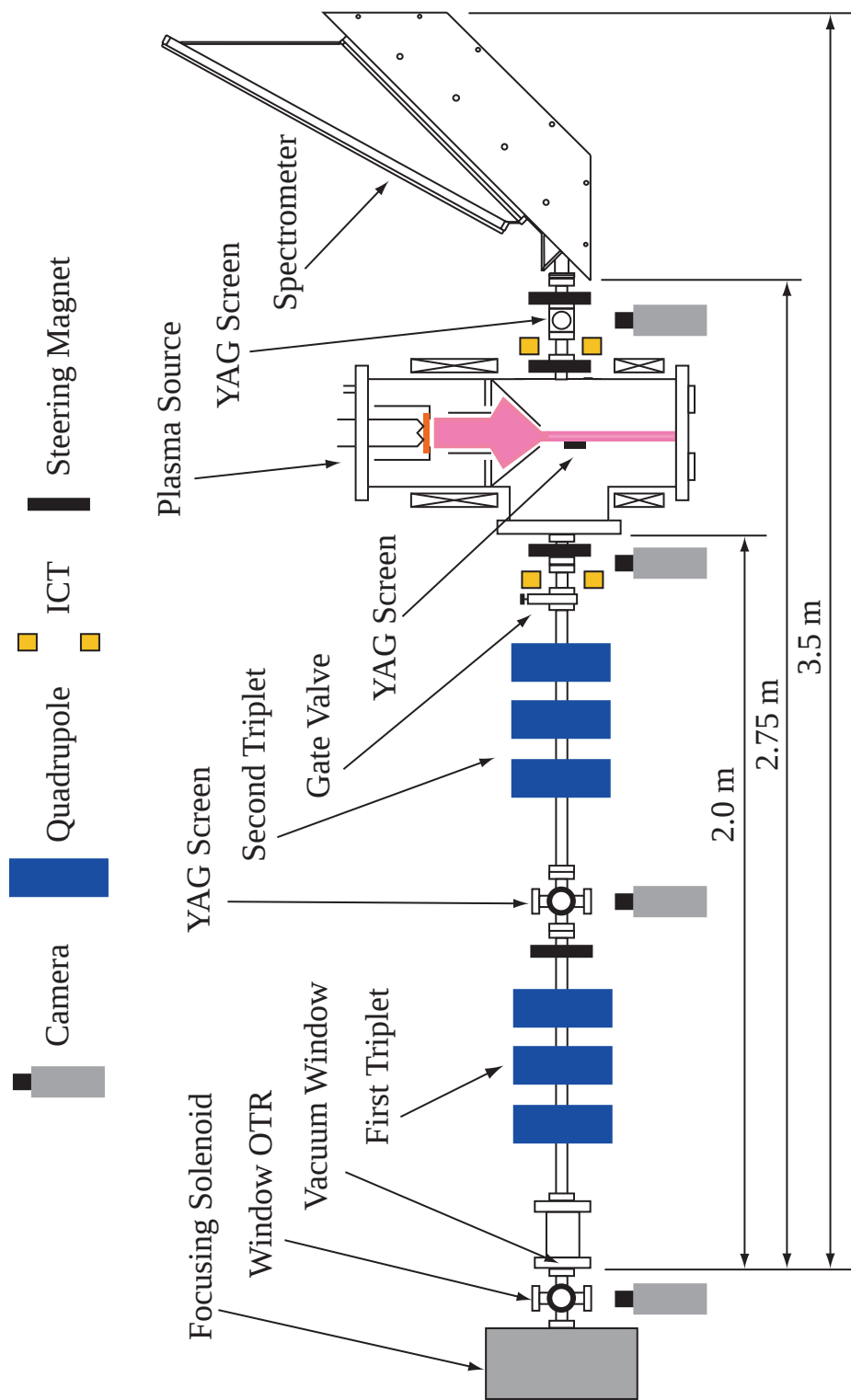


Figure 6.2: The plasma density transition trapping beamline

pecially the superconducting 9-cell accelerating cavity, cannot tolerate such high pressure. In fact, the 9-cell cavity is so sensitive to degradation of its vacuum, which is typically maintained in the 10^{-10} Torr range, that the vacuum isolation technique of differential pumping is not considered a viable option. Therefore, a solid window is used to totally isolate the vacuum of the two parts of the system. The window used for the experiment was $10\mu\text{m}$ thick aluminum. This substantial window was chosen to protect against breakage and contamination of the upstream beamline even in a situation where one side is accidentally vented to atmospheric pressure. The down side of such a thick window is the scattering that the electron beam experiences as it passes through the window, which increases the beam emittance and, therefore, the difficulty of transporting the electron beam to the experiment.

The amount of emittance growth produced by the beam's interaction with the foil can be calculated statistically. The distribution of deflection angles produced by multiple Coulomb scattering is approximately gaussian for small deflection angles. Assuming this small angle limit, and that a zero divergence electron beam with $x' = 0$ enters the foil, the rms angle of the beam electrons emerging from the foil, $\sigma_{x'foil}$, is given by

$$\sigma_{x'foil} = \frac{13.6\text{MeV}}{\beta cp} z \sqrt{\frac{x}{X_0}} \left[1 + 0.038 \ln \left(\frac{x}{X_0} \right) \right], \quad (6.1)$$

where p is the particle momentum, βc is the particle velocity, z is the charge number for the particle ($z = 1$ for the electron), x is the thickness of the foil, and X_0 is the radiation length of the material [81]. The radiation length of a material is the mean distance over which a high energy electron passing through the material loses all but $1/e$ of its original energy. X_0 depends only on the characteristics of the material, such as atomic mass and atomic number, and therefore only needs to be calculated once for a particular material. For aluminum, which has atomic

number 13, atomic mass 26.97 g/mol, and density 2.7 g/cm², the calculated radiation length is 8.57 cm [81]. The nominal electron beam energy, as given by Table 6.1, is 18 MeV. Using Eq. 6.1 to calculate the angular spread that a beam of this energy gains by passing through the 10 μ m aluminum window gives $\sigma_{x'foil} = 5.4$ mrad. In order to translate this angle into an emittance, recall that the normalized beam emittance is given by $\varepsilon_n = \beta\gamma\sigma_x\sigma_{x'}$, see Section 1.3. The nominal beam spot expected at the foil is $\sigma_x = 200$ μ m. Combining these facts gives an estimate for the emittance growth caused by the foil, $\varepsilon_{n,foil} = \beta\gamma\sigma_x\sigma_{x'foil} = 40$ mm-mrad. Since the beam emittance before the foil is expected to be about 15 mm-mrad (again see Table 6.1), foil scattering should dominate the emittance after the foil.

6.2.2 Diagnostic Positions

Diagnostics are positioned along the transition trapping beamline to provide information on the status of the electron beam and plasma, and to detect trapped electrons. Starting from the left, or upstream, in Fig. 6.2 the first diagnostic is OTR light (see Section 6.3.1) from the electron beam hitting a polished metal surface near the vacuum window. This allows measurement of the spot size of the beam hitting the window, which is a critical parameter for emittance preservation. The next beam diagnostic is the 1 inch diameter YAG screen, again see Section 6.3.1, between the two quadrupole triplets which allows observation and tuning of the large, roughly collimated beam between the sets of quads. Right before the entrance to the plasma source is an ICT, see Section 6.3.3, which measures the beam charge entering the experiment. Inside the plasma chamber is a small 1 cm YAG screen position slightly off the beamline center. This screen allows the spot size of the beam striking the plasma to be diagnosed. It is ob-

served, with the aid of mirrors, by the camera shown to the immediate left of the chamber in Fig 6.2. The plasma chamber also contains a capacitance manometer for measuring the pressure of the argon gas and an electrostatic probe for verifying the plasma density. Right after the plasma chamber is another ICT for measuring the charge of the driving electron beam, and any trapped electrons, leaving the plasma. This ICT is followed closely by another 1 inch YAG screen for monitoring the spot sizes and positions after the plasma. Finally, at the end of the line is the magnetic spectrometer, see Section 6.3.5. The spectrometer is the primary diagnostic for observing trapping electrons which are well separated from the drive beam in energy.

6.2.3 Beam Transport

Throughout the plasma density transition trapping experiment it was relatively easy to produce an electron beam that satisfied the design parameters in terms of charge, pulse length, and spot size. It turned out to be difficult, however, to transport this beam all the way to the plasma source without significant degradation. The biggest deleterious effects were the loss of drive beam charge and the larger than expected transverse spot size at the experiment. During most of the experiment it was typical to lose 50% or more of the drive beam charge between the chicane magnets and the plasma source. The minimum achievable transverse beam size at the plasma was several times larger than the design value. Both of these beam problems were clearly related to the unexpectedly high emittance observed at the end of the transport line. See Section 7.2 for the measurements of drive beam properties and further analysis.

Unfortunately, the exact origins of the charge loss and emittance blowup, as well as the best means to remedy these problems, were not clear. The likely

causes of major emittance growth in the system were scattering in the vacuum isolation window, the phase space effects of magnetic compression, and the dynamics of very high bunch charge production in the photoinjector. The relatively large beam energy spread required for magnetic compression probably also played a role in the charge loss. Since the beamline optics have significant chromatic aberration, energy spread increases beam sizes throughout the system. At first, it seemed that scattering in the window was largest problem (see Section 6.2.1). An attempt was made to solve this problem by increasing the diameter of the beamline pipe after the window from 1.5 inch diameter to 2 inch diameter. It was believed that the larger pipe would prevent the edges of the rapidly expanding beam after the foil from striking the vacuum pipe wall and being lost, which is the process assumed to be responsible for the beam losses. When it was found that this upgrade made no obvious improvement, an attempt was made to replace the 10 μm aluminum vacuum window with a thinner polymer window. The thinner, lower density window would significantly reduce emittance growth through scattering. Unfortunately, this window failed its 1 atm pressure hold off test, required to insure the vacuum safety of the upstream beamline (see Section 6.2.1), and no replacement was available in time for use during the experimental run.

6.3 Diagnostics

The success of the plasma density transition trapping experiment hinges on the ability to match the parameters of the real electron beam and plasma to those planned through simulation. Achieving this goal requires many different diagnostics to measure all the parameters of interest. Plasma measurements are described in Sections 5.2 and 5.3. This section discusses the methods by which the electron beam spot size, position, length, charge, and energy are measured.

6.3.1 Transverse Profile Diagnostics

Measurements of the electron beam transverse profile are necessary at many points in the experiment including tuning the beam transport, matching the beam into the plasma, and observing the beam energy on spectrometers. Transverse profile measurements were made in this experiment by using cameras to image light emitted by screens which are inserted into the electron beam path. The screens produce light through either optical transition radiation (OTR) or scintillation. OTR is emitted whenever the electron beam impacts a media boundary, and provides superior resolution for the measurement of sub- $100\mu\text{m}$ beam spots [82]. The signal intensity of OTR is low, however, making it suitable only for diagnosing intense electron beams. Since most of the electron beam spot sizes in the experiment are over $100\mu\text{m}$, scintillating screens, which produce much more light than OTR, were used much of the time. Two types of scintillator were used in the experiment, yttrium aluminium garnet activated by cerium (YAG:Ce), and a plastic referred to as EJ-260.

YAG:Ce is a fast crystalline scintillator that produces a large amount of light [83], see Table 6.2. The material is also mechanically strong and can be fashioned into very thin screens. YAG:Ce also performs well in vacuum. The YAG:Ce screens used in this experiment were typically 1 cm in diameter.

EJ-260 is a green emitting plastic scintillator produced by Eljen Technology [84]. As can be seen from Table 6.2, the light output of EJ-260 is lower than that of YAG:Ce, but still substantial. In addition, very large pieces of EJ-260 can be produced economically, and the solid plastic has better vacuum performance than powder scintillators. For these reasons EJ-260 was used for the scintillator in the large exit port of the broad range in vacuum spectrometer, see Section 6.3.5. The cost of using YAG:Ce in this application would have been prohibitive.

Table 6.2: Comparison of the scintillators used in this experiment.

Material	YAG:Ce	EJ-260
Wavelength of Max. Emission	550 nm	490 nm
Decay Time Constant	70 ns	9.2 ns
Scintillation Efficiency [Photons/(1 MeV e ⁻)]	45,000	9,200

6.3.2 Screen Position Determination

The relative position of the driving electron beam and the edge of the screen baffle is a critical parameter for this experiment, see Section 5.3. It was originally planned to measure this position by making the baffle edge of a scintillating material and observe the edge of the beam's transverse profile on this detector. This method proved technically inconvenient and ultimately unnecessary. The baffle edge that was used in the experiment was made of 42 gage (78 μm thick) stainless steel which effectively blocks the portion of the beam striking it. The spacing between the baffle edge and the electron beam can therefore be deduced from the intensity and shape of the beam spot on a transverse profile screen downstream of the baffle.

The beam intercepting the baffle has a transverse charge density profile that is gaussian in both dimensions:

$$\rho(x, y) = \frac{1}{2\pi\sigma_x\sigma_y} \exp\left(-\frac{x^2}{2\sigma_x^2}\right) \exp\left(-\frac{y^2}{2\sigma_y^2}\right), \quad (6.2)$$

where ρ is the transverse charge density and σ_x and σ_y are the rms sizes of the beam in x and y respectively. Since the baffle is a straight vertical edge in the y axis, the fraction of the beam charge Q_{clear} that clears the baffle is given by

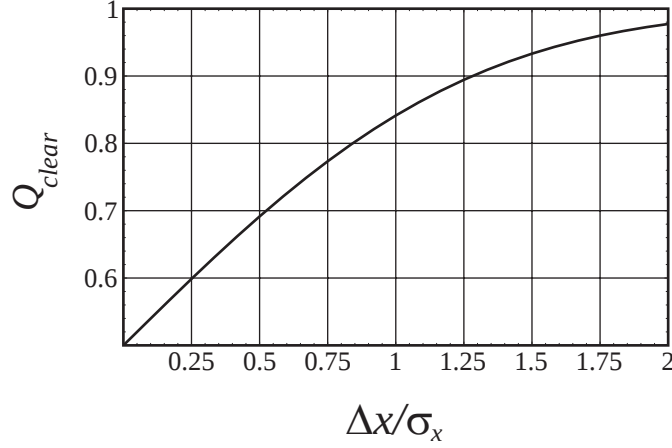


Figure 6.3: Beam charge transmitted past the baffle at beam centroid/baffle edge spacings of interest.

integrating in y over all space but only integrating x up the baffle edge

$$Q_{clear}(\Delta x) = \frac{1}{\sqrt{2\pi}\sigma_x} \int_{-\infty}^{\Delta x} \exp\left(-\frac{x^2}{2\sigma_x^2}\right) dx = \frac{1}{2} \left(\text{Erf}\left[\frac{\Delta x}{\sqrt{2}\sigma_x}\right] + 1 \right), \quad (6.3)$$

where Δx is the difference between the x position of the baffle and the x position of the beam centroid, and Erf is the error function. This function is plotted in Fig. 6.3.

Eq. 5.9 indicates that $\Delta x \leq k_p^{-1}/2$. Under most operating conditions for this experiment $k_p^{-1}/2 \approx \sigma_x$ and it follows that $\Delta x \approx \sigma_x$. The proper operating condition is therefore $Q_{clear} \leq 0.84$. The charge that clears the baffle is directly related to the intensity of the spot on the downstream screen. The drop in intensity is measured by comparing the beam intensity with and without clipping on the baffle.

6.3.3 Charge Diagnostics

Integrating Current Transformers (ICTs) are an excellent way to non-destructively diagnose the charge of short electron bunches. ICTs contain a capacitively shorted transformer that is coupled to a fast read out transformer through a common magnetic circuit. An ICT works by integrating the electron bunch signal (time scale of picoseconds) with a time constant on the order of nanoseconds. This has the effect of slowing down the rise and fall times of the electron beam signal to the point where eddy losses in the transformer are negligible. While this technique makes ICTs very linear integrators of beam charge, it also insures that all information about the beam's time domain structure is lost.

The ICTs used in this experiment are Bergoz Instrumentation company's model ICT-122-070-20:1 [85]. These passive ICTs have a 12.2 cm inner diameter, 70 ns output pulse duration (specified at 6σ), and a 20 to 1 turn ratio. The output of the ICT is given by,

$$\int I_{out} dt = \frac{1}{N} \int I_{beam} dt. \quad (6.4)$$

When the ICT, which has 50Ω output impedance, is connected to a oscilloscope with 50Ω input impedance these resistances act in parallel and $I_{out} = V_{out}/25\Omega$. In our case $N = 20$ so Eq. 6.4 becomes,

$$\int V_{out} dt = \frac{25\Omega}{20} Q_{beam} = 1.25 \left[\frac{V \cdot s}{C} \right] Q_{beam}. \quad (6.5)$$

Volt seconds $[V \cdot s]$ of the ICT pulse is the quantity measured every shot on the oscilloscope. The cables between the ICTs and the oscilloscope are calibrated for losses using a reference pulse of similar magnitude and duration to ensure an accurate measurement.

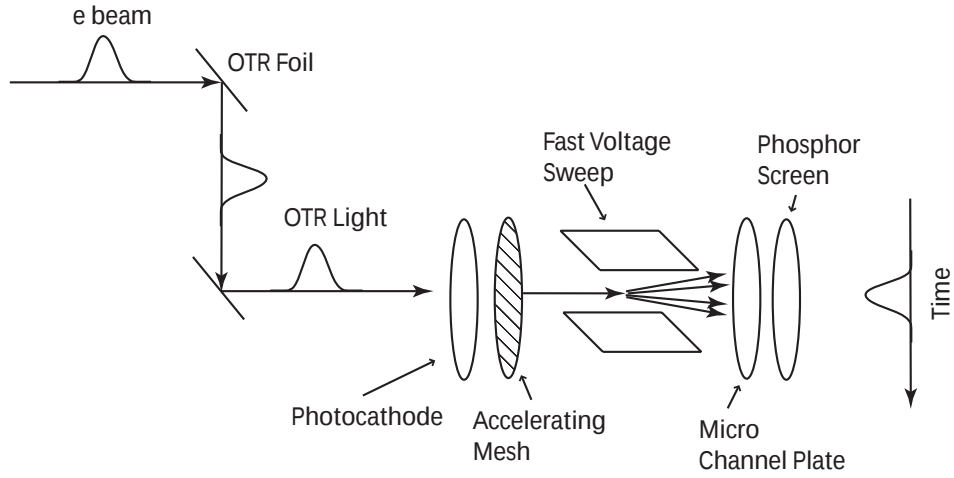


Figure 6.4: Illustration of streak camera operation.

6.3.4 Streak Camera

Diagnoses of the drive electron beam length was accomplished using a Hamamatsu model C5680 streak camera. This streak camera is specified to have less than 2 ps (1.5 ps typical) time resolution. The streak camera operates as illustrated in Fig. 6.4. The electron beam is intercepted by an OTR foil and produces a pulse of light that has the same time domain structure as the electron beam. This light is transported to the entrance slit of the streak camera where a small slice of it illuminates a photocathode. The electron beam produced at the photocathode will have the same time structure as the OTR and is accelerated by a mesh held at potential above the photocathode. This electron beam then passes through an electrode to which a very fast voltage sweep is applied. The fast changing electric fields between the electrodes impart a time dependent steering to the electron beam which results in vertical position being correlated to time when the beam impacts the micro channel plate. The micro channel plate electrically amplifies the electron beam so that a bright image will appear on the

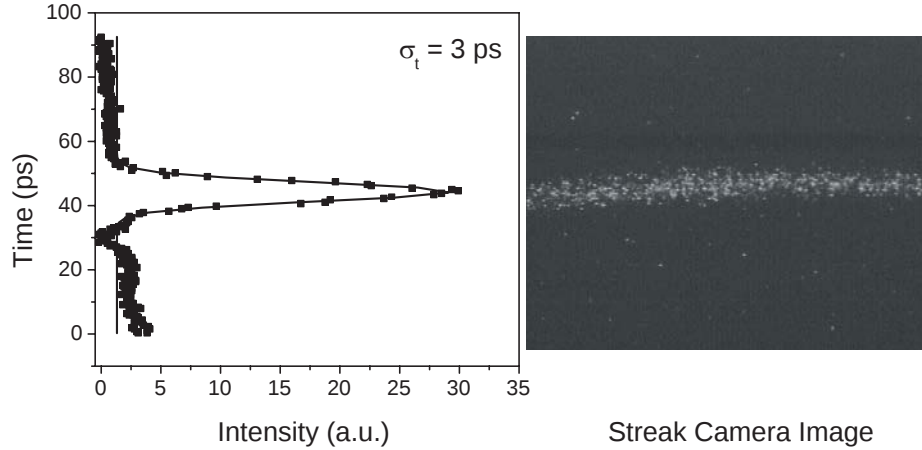


Figure 6.5: Example of streak camera data analysis.

phosphor. The vertical axis of this image gives the time domain structure of the OTR light and hence the original electron beam.

Fig. 6.5 illustrates the analysis of the streak camera signal. On the right hand side is a photograph of the raw image produced by the streak camera. On the left is a graph of the integrated intensity of each horizontal line of the image plotted against the time each horizontal line represents. A gaussian is fitted to the points to give the σ_t of the beam. In this typical case, a beam of approximately 11 nC has been compressed to 3 ps σ_t by magnetic compression.

While streak camera measurements accurately quantify the approximate length of the electron beam, at this level, they do not reveal information about the detailed structure of the electron beam due to the camera's resolution limitations. Resolving the non-gaussian time domain structure of magnetically compressed beams requires more sophisticated methods such as deflection cavities. For the purposes of this experiment, the streak camera measurements, and approximation of the beam profile as gaussian, are sufficient.

Table 6.3: Design parameters of the broad range in vacuum spectrometer.

Gap	2 cm
Maximum Field	1600 gauss
Observable Radii of Curvature	10 cm - 58 cm
Energy Range of Spectrometer at Field =	
180 gauss	740 keV - 3.17 MeV
1550 gauss	4.7 MeV - 27 MeV

6.3.5 The Broad Range Vacuum Spectrometer

Diagnosing the plasma density transition trapping experiment requires measuring the trapped beams, which have low energy (~ 2 MeV) and significant energy spread ($\sim 4\%$ rms), as well as the higher energy drive beams (~ 14 MeV). In order to meet these requirements we designed and constructed a specialized spectrometer magnet. The magnet has a very long output port which allows the observation of a wide range of energies simultaneous. There is a factor of approximately 5 between the minimum and maximum observable energies. In addition, the magnet is equipped with an in-vacuum diagnostic port, which allows the electrons to travel through vacuum all the way up to the scintillator where they are detected. This feature is critical for the accurate detection of low energy electrons, which can not easily penetrate the metal foils that are often used to seal the exits of spectrometer magnets.

The design parameters for the spectrometer are listed in Table 6.3. The construction of the spectrometer is illustrated in Fig 6.6. The center of the magnet consists of two pole pieces which determine the shape of the magnet field. As can be seen from the cross sectional view in Fig. 6.6, the pole pieces also form

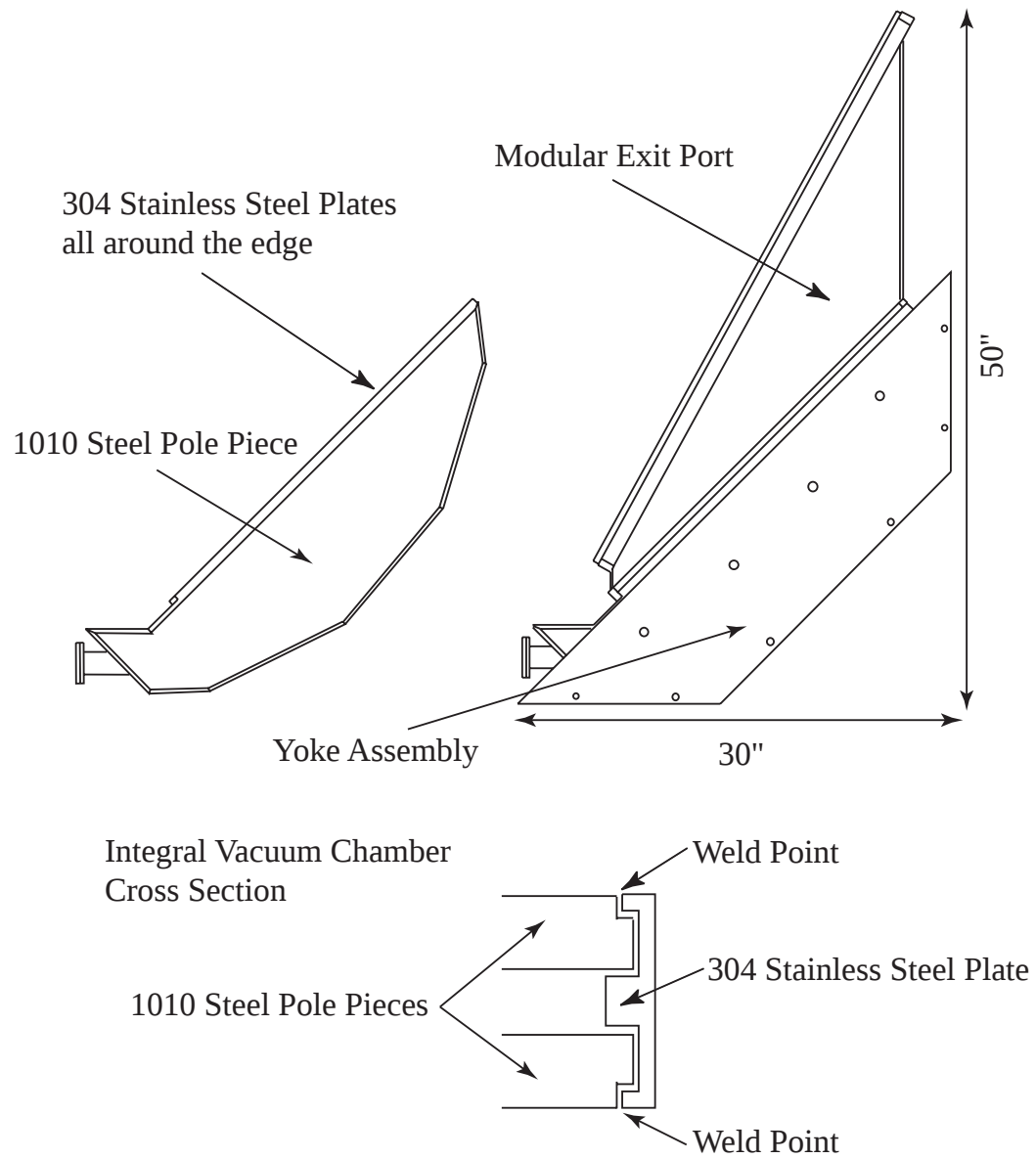


Figure 6.6: Mechanical diagram of the spectrometer.

the spectrometer vacuum vessel which is sealed with a series of non-magnetic stainless steel plates that are welded around the perimeter of the pole pieces. Each stainless steel piece has a raised bar in the center that fits between the two pole pieces. This feature serves as a precision parallel bar and ensures that the magnet gap is uniform. Note that the fit of these pieces in the actual device is very tight, Fig. 6.6 is drawn with greater spacing for clarity. The location of the welds on the outside edges of the stainless pieces was chosen to insure minimum distortion and warping of the magnet pole edges. Although this placement of the welds is not optimum from a vacuum standpoint, the integral vacuum chamber still achieved pressures in the low 10^{-7} Torr range, which was more than adequate for our purposes. The pole assembly goes inside a yoke assembly made of the same 1010 low carbon steel. The magnetic field is generated by two 195 turn coils of 10 gauge solid copper wire. These coils are air cooled and are designed to operate up to at least 10 Amps.

The modular exit port bolts on to the pole assembly and is sealed with a viton gasket. The modular design was chosen so that the spectrometer exit configuration could change in response to the needs of different experiments. The shape of the exit port used in this experiment was dictated by the experiment's beam optics requirements. The beam optics calculations are detailed in Appendix D. The electrons are detected at the end of the exit port using the green emitting plastic scintillator EJ-260 (see Section 6.3.1). The light produced when electrons strike this scintillator escapes the exit port through a glass window that is sealed to the end of the exit port and is imaged using a CCD camera. A fine copper wire mesh is sandwiched between the glass and scintillator to provide charge dissipation.

Magnetic testing of the spectrometer showed that its field is linear with current

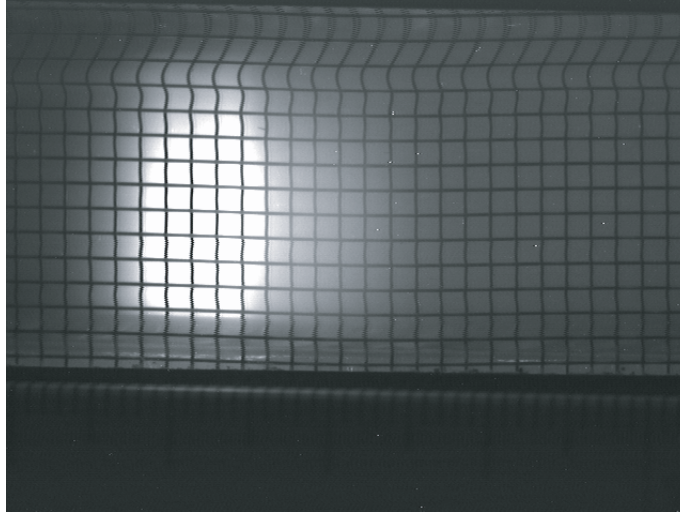


Figure 6.7: 15 MeV nominal energy test beam as seen on the broad range in-vacuum spectrometer. The field of view is centered on the 34 cm radius of curvature, which is in the center of the spectrometer's range, and covers an energy range of 14.3 - 15.2 MeV ($B = 1450$ gauss).

up to at least 9 Amp and is given by $B[\text{gauss}] = 190.1 \times I[\text{Amp}]$. Testing with known electron beams verified the functions and calibrations of the spectrometer. Fig. 6.7 shows one such beam as it appeared on the spectrometer. Note that the grid lines are the shadow of the charge dissipating copper mesh, which also serves as a rudimentary fiducial.

The expected operating configuration for the spectrometer during the trapping experiment is shown in Fig. 6.8. In this scenario the center of the spectrometer is used to search for trapped electrons while the drive beam is dumped into the magnet steel. Depending on the captured beam's energy, it may be possible to view it and portions of the drive beam simultaneously. Measuring the captured and drive beams at the same time would allow correlations between the two to be determined on a shot-to-shot basis.

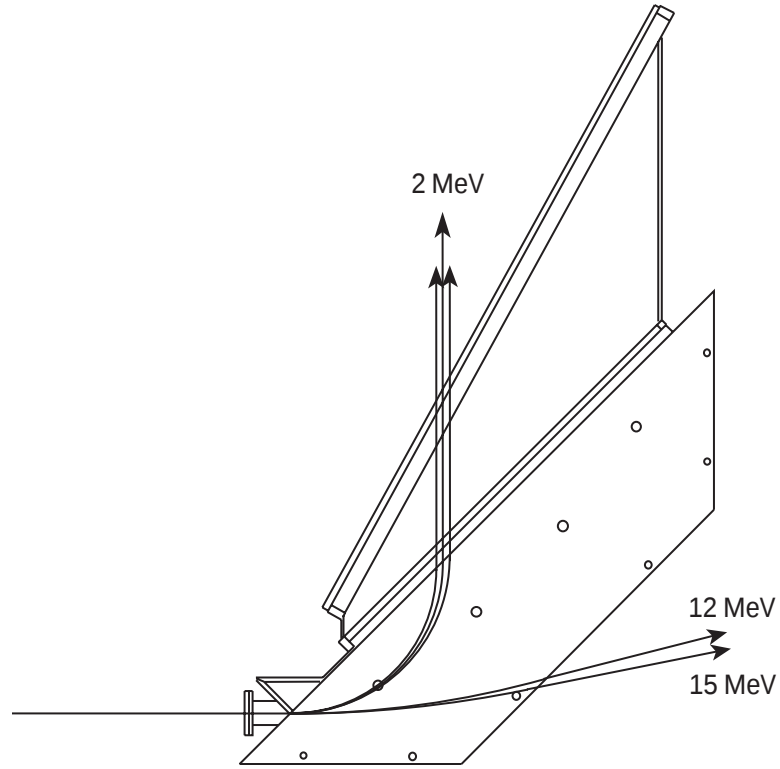


Figure 6.8: Typical operating configuration of the spectrometer. Beam paths are approximately to scale for a magnetic field setting of 352 gauss. The drive beam has a spread of energies typically between 12 and 15 MeV, corresponding to radii of curvature 1.14 m to 1.42 m. The trapped beam has a nominal energy of 2 MeV with 8% energy spread and therefore occupies the radii 17.3 - 18.7 cm.

CHAPTER 7

Experimental Results and Analysis

The first run of the plasma density transition trapping experiment took place at the FNPL between January and May 2004. This was the first time that the plasma apparatus had seen an electron beam. During this five month period, two types of beam-plasma interactions were directly observed, the drive beam parameters were extensively characterized, and a thorough search for trapped plasma electrons was conducted.

7.1 Plasma Focusing

Plasma induced focusing of the electron beam was the first evidence of interaction between the electron beam and plasma. Strong plasma focusing typically occurs when the electron beam radius is comparable to the plasma skin depth but its length is long compared to the skin depth. The mechanism of the focusing depends on the density of the beam relative to the plasma. When the beam is less dense than the plasma, $n_b \ll n_p$, the system is said to be in the overdense regime. In this regime, the electron beam does not strongly perturb the plasma and the main effect of the plasma is to neutralize the beam space charge, which allow the beam's self magnetic field to focus the electron bunch. If the beam is more dense than the plasma, $n_b \gg n_p$, the system is said to be in the underdense regime, just as in the blowout regime of the PWFA. In the underdense plasma

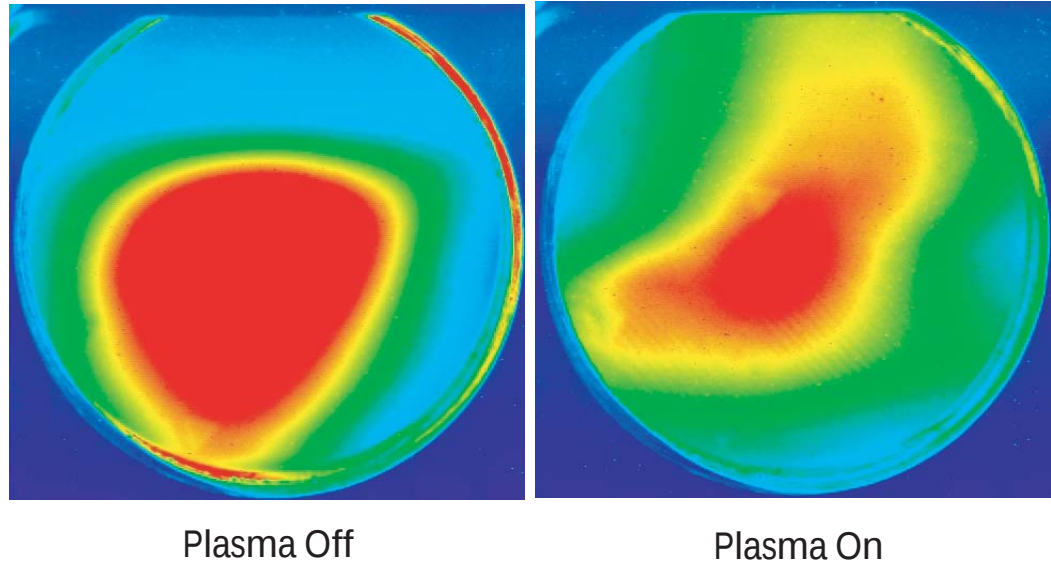


Figure 7.1: Plasma focusing of the uncompressed electron beam as observed on a YAG screen approximately 30 cm downstream of the plasma. Images are displayed in false color.

lens, the strong electric fields of the beam eject all the plasma electrons from the beam volume and the attraction of the beam electrons to the remaining plasma ion column provides the focusing force. Su *et al.* [86] provides a detail review of these two types of plasma focusing.

The development of the plasma density transition trapping source, see Section 5.2, began with a thin underdense plasma lens experiment in mind [65]. The more sophisticated source that evolved from those beginnings still retains the capability to operate as a thin underdense plasma lens. During the commissioning of the trapping experiment at the FNPL, a high charge, but uncompressed, beam was sent through the plasma and strong focusing of the beam was observed, see Fig. 7.1. The electron beam being focused had a charge of 8 nC, a radius σ_r of about 400 μm , and a FWHM length of 15 - 30 ps. These parameters yield a beam density of 1 - 2 $\times 10^{13} \text{ cm}^{-3}$. Given this range of beam density, which is

somewhere in between the densities of high and low density regions of the trapping experiment plasma column, see Fig. 5.13, the beam probably experienced a mixture of overdense and underdense lens focusing. A follow up underdense plasma lens experiment using a unmodulated plasma column is planned.

7.2 Drive Beam Characterization

While extensive work aimed at improving drive beam parameters was something that this experiment explicitly planned to avoid, see Section 5.1, the reality of experiment required it. The first, and most important, beam deficiency noticed was the beam charge. As was discussed in Section 5.5, high beam charges well in excess of the 6 nC design goal were critical to our strategy for the experiment. During the five months of the experimental run, the FNPL photoinjector produced electron pulses with an average charge between 8 and 16 nC. The variability was mainly due to the temperamental nature of the cathode drive laser. Large amounts of the experiment run time were devoted to optimization of the drive laser in order to achieve this level of photoinjector performance. Despite the availability of high charge electron bunches at the photoinjector, it proved exceedingly difficult to transport even 6 nC of this charge to the experiment under any circumstances. Strangely, even after extensive work to optimize the beamline, the amount of charge arriving at the plasma chamber seemed almost independent of the amount of charge coming out of the photoinjector.

This observation lead to systematic studies of the transport efficiency. The results of a study conducted during one run day, Fig. 7.2, and those obtained by examining the transport efficiencies achieved on various run days over many months, Fig. 7.3, are very similar. Both indicate a linear decline in transport efficiency with increasing photoinjector charge. This decrease in transport effi-

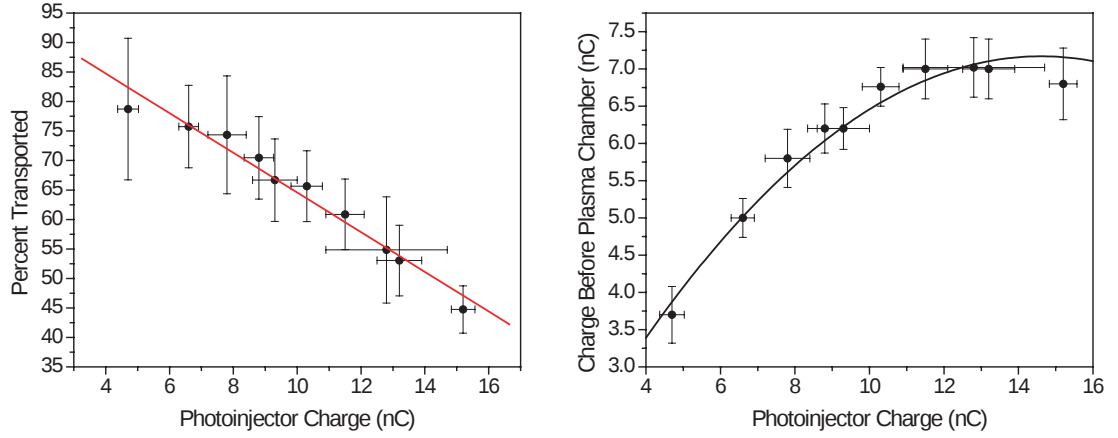


Figure 7.2: Variation in charge transport with photoinjector charge. The measurements were made over a short period of time by using a remote control iris to aperture the the cathode drive laser near the cathode. The plot at left shows charge transport efficiency as a function of charge leaving the photoinjector. A linear fit to the downward trend gives $(\% \text{ transported}) = 98 - 3.36Q_{\text{injector}}$. The plot at right shows the charged transported through the plasma chamber verses the charge leaving the photoinjector. The points are plotted along with the function $Q_{\text{plasma}} = 0.98Q_{\text{injector}} - 0.0336Q_{\text{injector}}^2$, which clearly shows the transported charge reaches a maximum at about 7 nC. The error bars in both plots are the one σ fluctuation in the average value.

ciency effectively limits the charge that can be delivered to the experiment to 7-8 nC. The limit on the charge deliverable to the plasma is not the only thing that can be deduced from this data. These results also hint that something is severely wrong with either the transport line or other aspects of the beam quality. Small improvements to the beamline seemed to have no effect on charge transport, and larger changes were not possible in the time frame of the first run, see Section 6.2.3. Attention therefore turned to accurately measuring the beam parameters.

Initial measurements of the beam emittance after the vacuum window were made during the beamline commissioning using the standard technique of quadrupole scans [87]. The normalized emittance of a 5 nC uncompressed beam was measured to be between 60 mm-mrad and 160 mm-mrad after the vacuum isolation

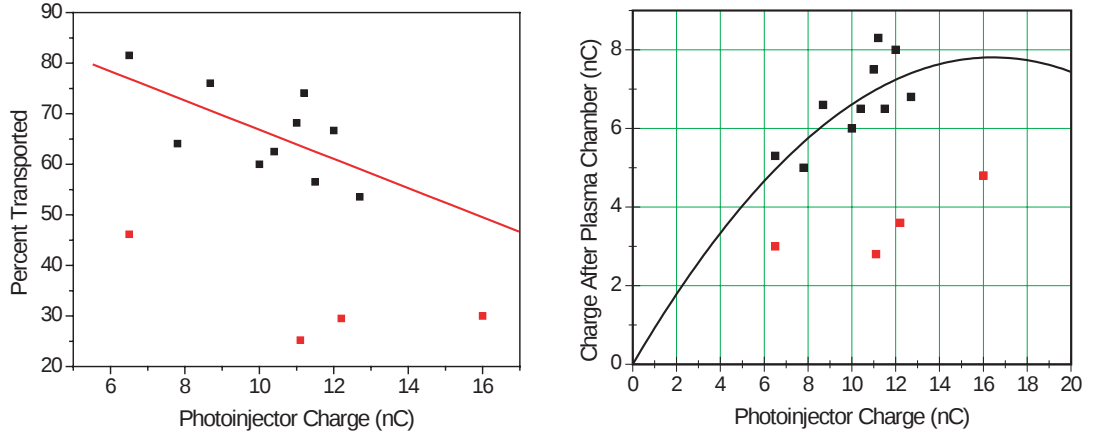


Figure 7.3: The points in the above plot represent the charge transported to the experiment under various conditions, both compressed and uncompressed, during many different run days. The plot at left shows charge transport efficiency as a function of charge leaving the photoinjector. If the four anomalously low points are neglected, a linear fit to the downward trend gives (% transported) $= 95 - 2.89Q_{injector}$. The plot at right shows the charged transported through the plasma chamber versus the charge leaving the photoinjector. The points are plotted along with the function $Q_{plasma} = 0.95Q_{injector} - 0.0289Q_{injector}^2$, which clearly shows the transported charge reaches a maximum at about 8nC.

foil. The emittance measurement depended on how well the beam focus at the foil was optimized. Smaller beam spot sizes at the foil yielded smaller emittances after the foil. These emittance values were substantially larger than the 15 mm-mrad value cited in the original design, but not large enough to seriously impact the trapping experiment (see Section 5.4.3). The measured values were also consistent with the estimate of 40 mm-mrad for the expected emittance growth caused by the foil (see Section 6.2.1). While it was known that compression and running at high charge would increase the emittance somewhat, estimates indicated that the dominate contribution to emittance growth would always come from the foil. When the charge transport problems suggested that the emittance might be much higher than what was believed, more measurements were taken

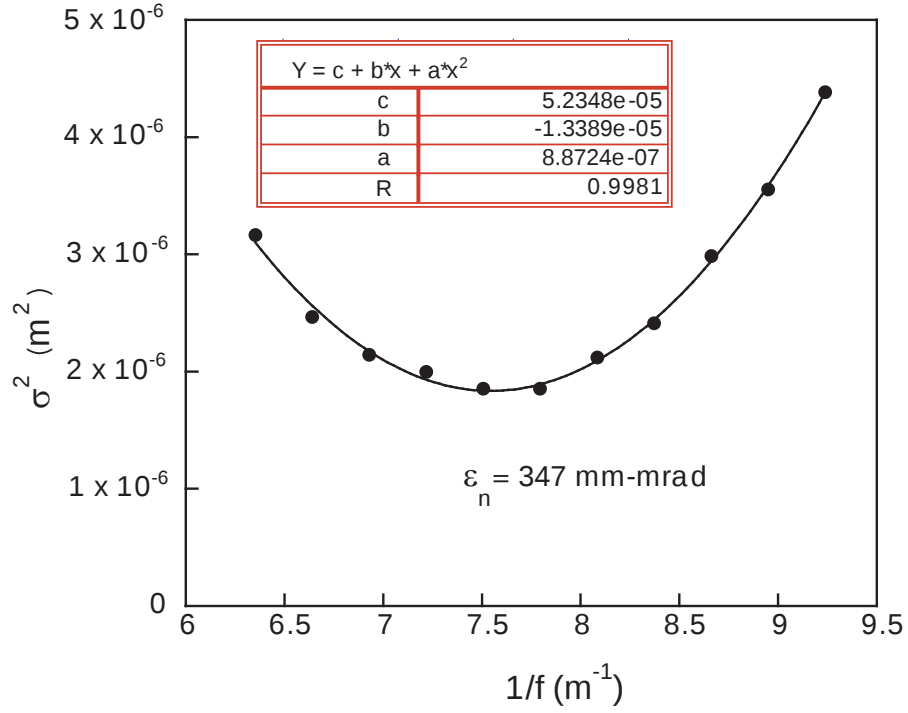


Figure 7.4: Quadrupole scan emittance measurement of the 14 nC electron beam after the vacuum isolation foil. It can be shown that for an electron beam being focused with a thin lens: $\sigma^2 = a(1/f)^2 + b(1/f) + c$ and $\varepsilon^2 = (ac - b^2/4)/L^4$ where L is the drift length between the thin lens and the point at which σ , the rms beam width, is measured [88]. For this measurement $L = 30$ cm and the beam energy, which is necessary for calculating the normalized emittance ε_n , is 12.5 MeV.

after the foil. The normalized emittance of a beam that had been compressed to $\sigma_t = 3$ ps was measured to be 347 mm-mrad, see Fig. 7.4. This beam had exited the photoinjector with 14 nC, of which, only 7 nC arrived at the trapping experiment. This emittance was much larger than expected and was the likely cause of the charge transport losses, as well as the large beam spot sizes at the plasma chamber.

The origins of the very high beam emittance are not yet fully understood. Scattering in the vacuum isolation foil cannot account for emittance values this

large by itself. The average beam spot size at foil was $\sigma_x = 700 \mu\text{m}$ under the compressed beam conditions for which the 347 mm-mrad emittance was measured. Eq. 6.1 predicts that the foil induced emittance growth for a beam of this size passing through the $10 \mu\text{m}$ aluminum foil is $\varepsilon_{n,foil} = 132 \text{ mm-mrad}$. Since the predicted foil emittance growth is only about a third of the observed emittance, there must be significant emittance growth occurring upstream of the foil. This upstream emittance growth is also a factor in the larger than expected beam spot size at the foil.

It is theorized that the beam emittance blowup is the result of a runaway process. Deficiencies in the UV cathode drive laser spot lead to a beam that starts with abnormally high emittance, which is then dramatically increased by first aggressive magnetic compression, and then passage through the thick aluminum vacuum window. In this scenario, the emittance growth at each stage in the beamline is exacerbated by the unexpectedly high emittance growth in the stage preceding it. It was also noted that the overall quality of the cathode drive beam deteriorated between the earlier and later emittance measurements, but this is a purely qualitative assessment. Unfortunately, the FNPL beamline was not instrumented to allow tracking of the beam charge and emittance all along the beamline, which could have verified these theories and allowed better tuning and improvement of the beam. Beam charge measurements were only available immediately after the photoinjector, immediately after the magnetic compressor, and at the very end of the beamline where the trapping experiment was located, see Fig. 6.2. Since all the beam charge losses happened after the compressor and before the experiment, these diagnostics did not aid in determining where the charge was being lost. Beam emittance measurements were only available right after the vacuum isolation foil. By the time it was suspected that emittance upstream of the foil was too high, there was insufficient time left in the first run

Table 7.1: Comparison of the design beam parameters with the best drive beam parameters that were achieved simultaneously. Values given in the achieved column are the mean. The design values are the same as those presented in Table 5.1 except for the emittance, which was updated to include the effects of scattering in the foil.

Driving Beam Parameters	Design	Achieved	$\pm\sigma$	$\pm$ Peak to Peak
Beam Energy [MeV]	14	12.5	0.22	0.49
Beam Charge [nC]	5.9	7	0.5	1.5
Beam Duration σ_t [ps]	1.5	3	1	2.2
Beam Radius σ_r [μm]	362	1220	230	545
Normalized Emittance [mm-mrad]	~ 60	347 (statistics unavailable)		
Peak Beam Density [cm^{-3}]	4×10^{13}	2.1×10^{12}	5.1×10^{12}	3.1×10^{13}
			Peak n_0	Peak n_0

to establish emittance diagnostics before the foil.

The combination of the charge transport limitations and the enormous beam emittances after the foil, when operating with a compressed high charge beam, prevented us from attaining all the parameters originally specified for the drive beam. Table 7.1 lists the design drive beam parameters along with the best, simultaneously optimized, beam parameters that were achieved near the end of the experimental run. The shot to shot fluctuation of the parameters are also listed. While bunches slightly longer than design were anticipated, see Section 5.4.3, the huge emittance, large spot size, and charge transport problems were not. The slightly lower than expected energy contributes to these problems by exacerbating emittance growth in the compressor and chromatic aberration in the focusing optics. The combination of a longer and wider than expected beam, along with the inability to compensate by adding more charge, meant that the

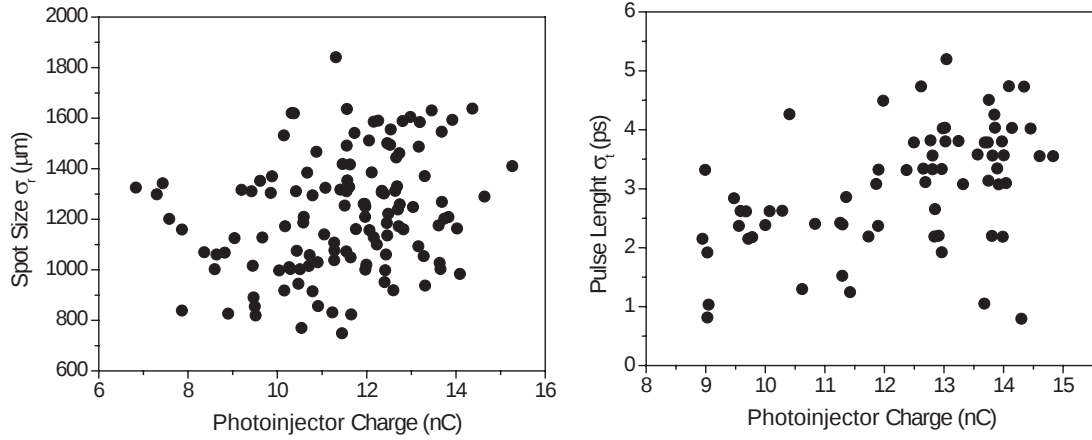


Figure 7.5: Correlations between the beam spot size and the beam charge (left), and between the beam length and the beam charge (right).

average peak beam density achieved is over an order of magnitude lower than design. The low beam density presents a major problem since effective plasma density transition trapping requires the drive beam to satisfy the underdense condition, $n_b > n_0$, at least in the low density region. The average achieved beam density does not fulfill this condition: compare Table 7.1 with Fig. 5.13. When fluctuations are considered, however, there are peak shots that come very close to fulfilling the design value for peak density. These shots are rare and still fall a little short of meeting the design density, mostly due to the large spot size even on the best shots.

It should be noted that the peak values for the peak beam density listed in Table 7.1 are somewhat theoretical since all the beam parameters cannot be measured at once and there is no guarantee that there will be shots in which they all fluctuate together in a favorable way. In fact, since pulse length and emittance usually increase with charge, it is logical to expect that the beams with the smallest spots and shortest length will occur at low charge. Surprisingly, the data does not exhibit a strong correlation between between spot size or beam

length and charge, see Fig. 7.5. There is, perhaps, a weak correlation between high charge and larger spots and longer bunches, but in general the distribution is fairly uniform. The correlations expected between spot size, length, and charge are probably overwhelmed by the large fluctuations in other beam parameters. Table 7.1 shows that there is significant energy jitter, which impacts beam length and emittance since magnetic compression and beam focusing, e.g. at the vacuum isolation foil, have a dependence on beam energy. Also, there was significant fluctuation in the structure of the cathode drive laser spot, which was very far from gaussian. While difficult to quantify, the fluctuations in spot structure certainly lead to fluctuations in emittance, especially at high charge. These sorts of fluctuations probably blur the correlations that would be expected to show up in plots like those shown Fig. 7.5 for a more stable system. From the viewpoint of operating the plasma density transition trapping experiment, the lack of correlation in Fig. 7.5 means that it is likely that occasional shots of high peak beam density exist in this system.

There is one more form of fluctuation which impacts the experiment which is not accounted for in Table 7.1. The pointing jitter of the electron beam at the plasma is important due to critical dependence of the trapping performance on the spacing between the beam and baffle, see Section 5.3.2. The pointing jitter of the beam centroid within the plasma source along the axis of interaction with the baffle was measured to be $\pm 134 \mu\text{m} \sigma$ and $\pm 392 \mu\text{m}$ peak to peak. When expressed in units of plasma skin depth for the high density region of the plasma, these fluctuations are $\pm 0.13 k_p^{-1} \sigma$ and $\pm 0.4 k_p^{-1}$ peak to peak. As can be seen from Fig. 5.10 this much variation in the spacing between baffle and electron beam leads to significant fluctuation in trapped charge.

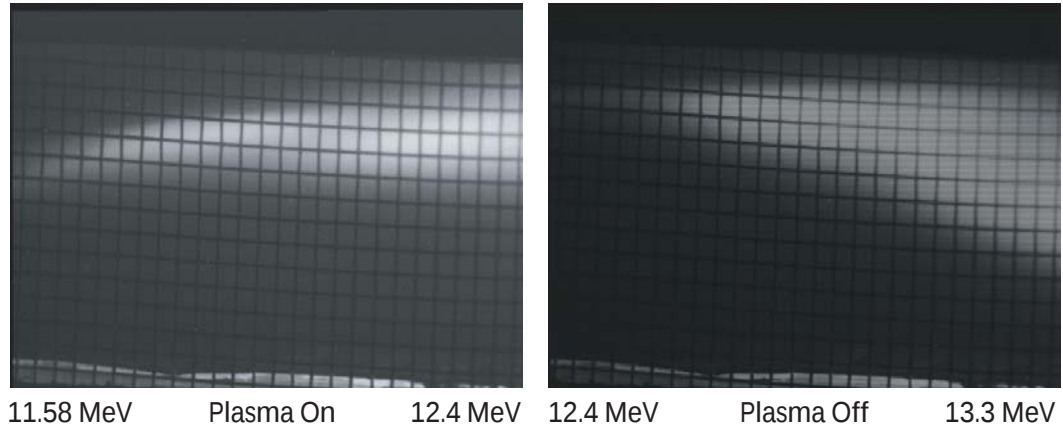


Figure 7.6: Deceleration of the drive beam due to plasma wake fields. On the right is the energy distribution of the driving electron beam as observed on spectrometer. At left is the energy distribution in the presence of the plasma. Note the shift in the energy window being viewed as indicated by the range marked below each window.

7.3 Drive Beam Deceleration

Despite the problems with the drive beam described in Section 7.2, wake fields were observed when the drive beam was interacting with the plasma. The presence of wake fields was indicated by deceleration of the drive beam. The lowest energy drive beam electrons observed were at about 11.6 MeV when the plasma was on, and about 12.6 MeV when the plasma was off, see Fig. 7.6. This means that plasma wake fields within the electron bunch decelerated some of the drive beam electrons by 1 MeV during the passage through the plasma. The average deceleration gradient over the whole plasma profile is therefore about 12 MV/m. These were the strongest wake fields observed on peak shots. 1 MeV is about half to one third of the amount of deceleration that simulations indicate should occur in cases with significant trapping. Therefore, the observation of lower than expected decelerating wake fields corroborates the observations of sub-design beam parameters.

7.4 The Search for Trapped Electrons

The presence of wake fields, even if they were substantially weaker than anticipated, raised the possibility that some small amount of trapped charge might be detectable. With the drive beam parameters optimized as described in Section 7.2, the last variable that needed to be set was the spacing between the beam and the baffle edge. The method for determining the beam/baffle spacing described in Section 6.3.2 ultimately proved only partially effective. The beam pointing jitter and fluctuations in transverse spot size described in Section 7.2 made it essentially impossible to make the precise relative measurement of the percent of the beam being blocked that is describe in Section 6.3.2. On the other hand, the blocking of the beam by the baffle edge could easily be seen, see Fig. 7.7, and the appropriate beam/baffle set approximately. The beam fluctuations which made a more accurate setting of the spacing between beam and baffle difficult also made such precision largely superfluous.

With the beam/baffle spacing approximately set, many searches for trapped electrons were conducted. When the first of these showed nothing, the steering magnets down stream of the plasma were upgraded to ensure that low energy electrons could be successfully steered into the spectrometer. Cylindrical lens were also added to the optical path between the spectrometer exit port and the cameras observing it to enhance light collection. With these improvements more searches were conducted. During these explorations the beam/baffle spacing, post plasma source steering, and spectrometer magnetic field strength were systematically scanned through wide ranges. A weak signal was observed on spectrometer that originally seemed to be correlated to the presence of the plasma. More detailed measurements showed that this weak signal was independent of the presence of the plasma but did correlated to the presence of the electron beam, see Fig. 7.8.

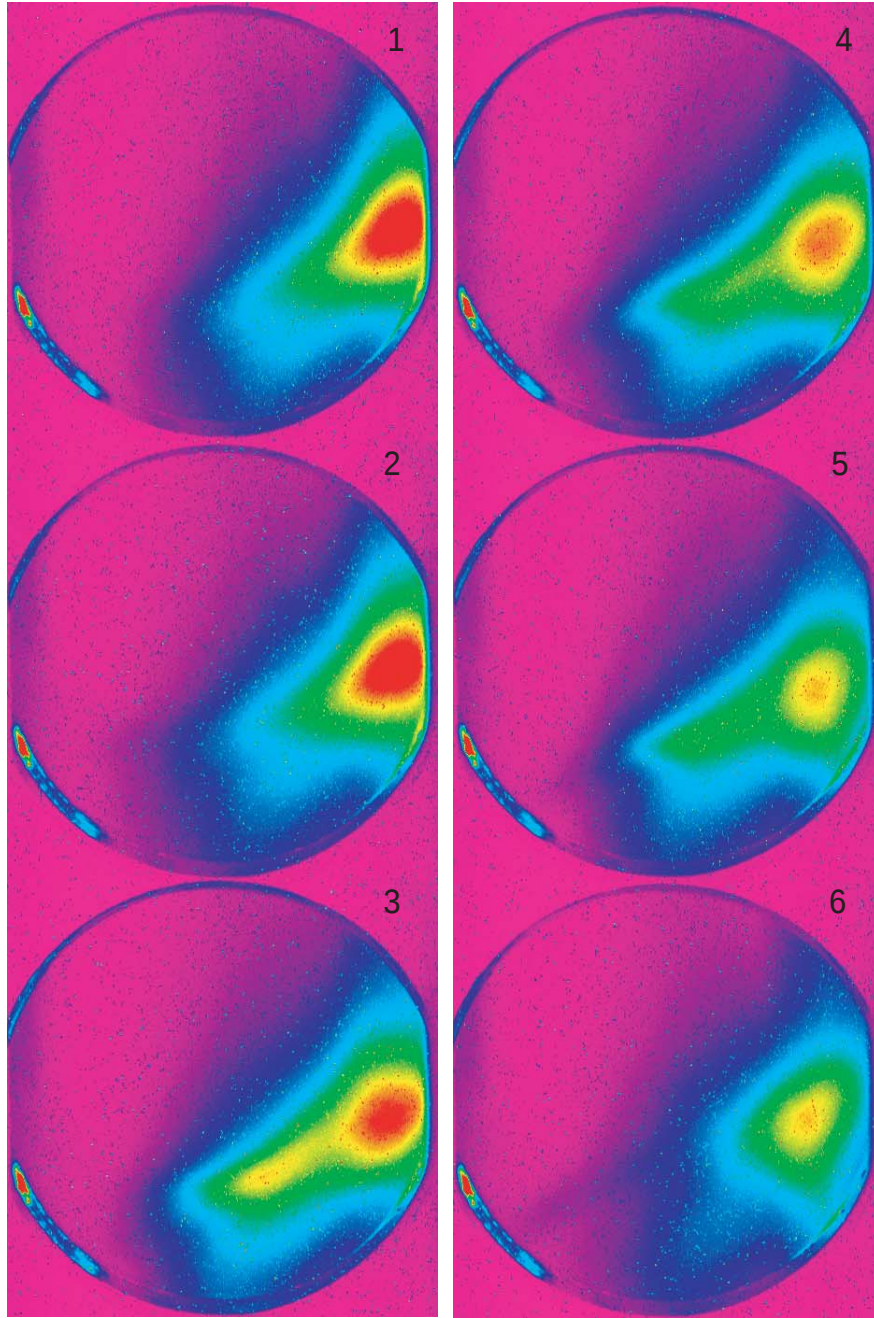


Figure 7.7: Observation of the baffle edge blocking the drive beam. In this sequence of pictures the electron beam, which is observed on a downstream screen, is gradually steered from right to left. The reduction in the brightness of the beam spot is due to the increasing portion of the beam which is intercepted by the baffle edge. Images are displayed in false color.

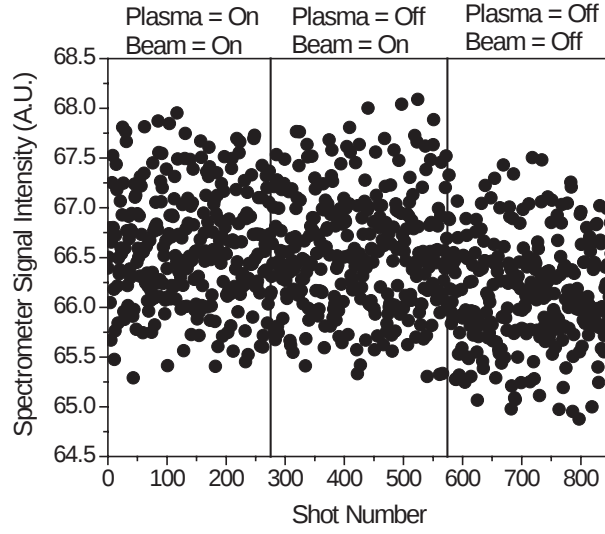


Figure 7.8: Weak low energy signal observed on the broad range spectrometer.

The exact origin of this weak low energy spectrometer signal is unknown but it is believed to be caused by some combination of beam generated x-rays and low energy secondary electrons emitted when the drive beam strikes metal, see Fig. 6.8.

7.5 Comparison of Experimental Results with Simulation

Simulations were run with the achieved beam parameters shown in Table 7.1 and the plasma profile used during the experiment, Fig. 5.16, in order to ascertain whether any plasma electron trapping could be expected. The results of these MAGIC simulations are summarized in Table 7.2. The fact that no trapped plasma electrons were observed during the experiment, see Section 7.4, agrees well with the post-facto simulation results. The mean drive beam parameters simply do not generate wake fields strong enough to produce trapping. The simulation does, however, indicate that the wakes are strong enough to produce about 1 MeV of drive beam deceleration. This amount of deceleration also agrees

Table 7.2: Results of trapping simulations that use the measured plasma column and drive beam parameters. The columns labelled “Best σ ” and “Best Peak to Peak” reflect the best beam possible within those fluctuation ranges. Note that these simulations assume a perfect transition and ignore effects like bending while trapping, etc.

Real Driving Beam Parameters	Mean	Best σ	Best Peak to Peak
Beam Charge [nC]	7	7.5	8.5
Beam Duration σ_t [ps]	3	2	0.8
Beam Radius σ_r [μm]	1220	990	675
Peak Beam Density [cm^{-3}]	2.1×10^{12}	5.1×10^{12}	3.1×10^{13}
Trapping Parameters			
Trapped Charge [pC]	0	0	577
Drive Beam Energy Loss [MeV]	1	2	4.8

well with observation, see Section 7.3.

The best beam that can be expected within the one σ fluctuation of all the parameters is likewise insufficient to excite trapping. This best beam is the one that would exist if the parameters fluctuated to the maximum charge, minimum length, and smallest spot size simultaneously. Simulations of the analogous best beam derivable from the peak to peak fluctuations does show substantial trapped charge. After the estimated 90 % losses from various effects like bending while trapping, see Section 5.5, about 60 pC could be produced by this peak beam case.

Unfortunately, it is very unlikely that shots with these best peak to peak parameters occurred. First of all, the shots meeting any one of the peak parameters are on the tail of the distribution and are fairly rare. Shots that meet all three of the peak parameters simultaneously would be extremely rare, assuming that the fluctuations are uncorrelated. While Section 7.2 indicates that the assumption of uncorrelated fluctuations seems valid for charge and spot size, as well

as charge and beam length, it was not possible to correlate spot size and beam length. Given that emittance typically increases dramatically as the beam length is reduced in a magnetic compression system, it is doubtful that minimum spot size and minimum beam length would occur simultaneously. These arguments against the existence of shots that approach the best peak to peak parameters are substantiated by the fact that the large drive beam deceleration of over 4 MeV that would accompany such shots was never observed.

The beam transport through the plasma trapping experiment beamline was also simulated using the beam dynamics code TRACE 3D [89]. The simulation started at the vacuum isolation foil using the beam conditions measured there, $\varepsilon_n = 347$ mm-mrad and $\sigma_x = 700$ μm , and predicted a optimized beam size of $\sigma_x = 1145$ μm at the plasma. This agrees very well with the measured spot size of $\sigma_x = 1220$ μm . Fig. 7.9 shows an example output plot from the TRACE 3D simulation. Interestingly, Fig. 7.9 shows that the beam FWHM does not exceed about 10 mm. Since the diameter of the beam pipe is 50 mm, this result indicates that scraping of the beam on the pipe wall is not a problem, at least for the well behaved core of the beam. Charge loss must therefore either occur upstream of the foil, or be the result of a very non-gaussian beam trace space distribution. Beam conditions were not known well enough upstream of the foil, due to a lack of diagnostics, to allow accurate beam transport simulations of the beamline between the compressor and the foil.

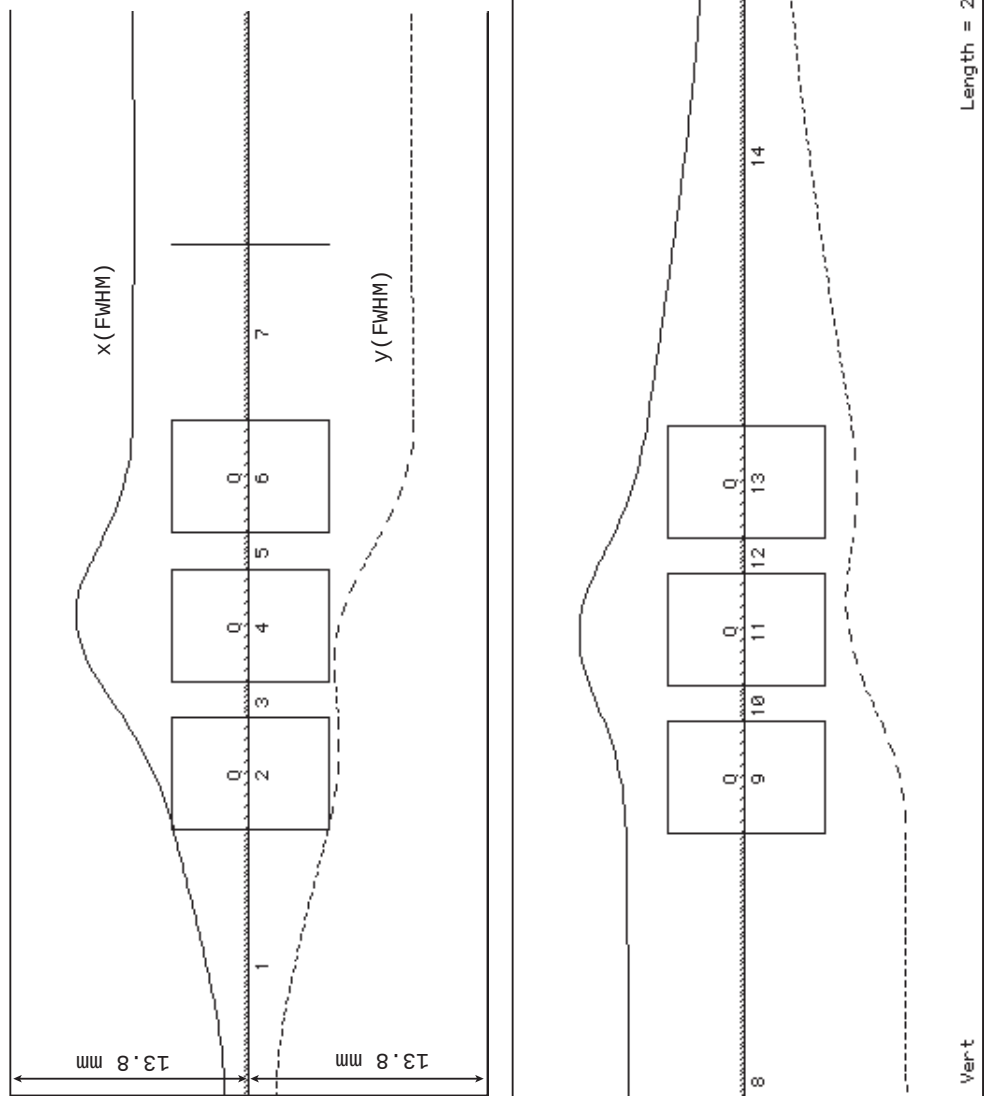


Figure 7.9: Trace 3D simulation of the trapping experiment beamline.

CHAPTER 8

Conclusions and Future Directions

A great deal of progress has been made in developing the concept of plasma density transition trapping in the short time since the idea was introduced by Suk *et al.* [54] in 2001. While much has been learned about plasma density transition trapping, and its characteristics as an electron beam source, many open questions still remain. Several of the things learned during this research also suggest new directions for future experiments.

8.1 Conclusions

The weak blowout regime of plasma density transition trapping, see Chapter 3, was developed as part of the program to create the first transition trapping experiment. In the weak blowout regime, the high density plasma region acts as a plasma cathode and provides all the electrons that are eventually trapped. Operation in the weak blowout mode requires less drive beam charge and produces lower emittance beams than strong blowout operation at similar plasma densities. The amount of plasma charge captured and accelerated is, however, smaller in the weak blowout case than the strong blowout case.

The scaling laws established for plasma density transition trapping in Chapter 4 indicated that the beams produced by this technique at high plasma densities can surpass those produced by state-of-the-art photoinjectors in terms of

brightness. There are practical issues with this type of high brightness operation since the drive beam parameters needed rival those of the captured beam except, perhaps, for the emittance. This raises the possibility of transition trapping operating as an emittance transformer. Since the scaling rules were written without reference to any particular plasma trapping mechanism, they indicate that the general performance of *all* plasma electron beam sources is tied to the plasma density at which they operate. This implies that high beam brightness requires high plasma density for any trapping system, whether it is beam or laser driven.

Plasma density transitions steep enough to fulfill the trapping condition were created and measured, see Chapter 5. This successful demonstration of transition production was achieved using density masking screens to modify a discharge produced plasma with density on the order of 10^{13} cm^{-3} . The plasma density transitions produced using this masking technique form the basis for a fully developed trapping experiment plan that utilizes drive beam parameters previously generated. While it was not possible to model every aspect of the trapping process in a realistic experimental scenario, numerous simulations of the most important effects indicated no insurmountable obstacles to a successful experiment.

The first attempt to carry out the plasma density transition trapping experiment described in Chapter 5 at the FNPL (see Chapter 6) yielded interesting evidence of interaction between the beam and plasma, but it did not produce observable quantities of trapped electrons (see Chapter 7). The primary reason for the lack of trapped electrons was the inability to deliver a drive beam that met the design specifications to the plasma transition. The direct cause of the problems with the drive beam was its disastrously high emittance. The high emittance resulted from a combination of factors. Deficiencies in the UV cathode drive laser spot lead to a beam with abnormally high emittance which was then

dramatically increased by aggressive magnetic compression and passage through the thick aluminum vacuum window. The emittance growth at each stage of this process was exacerbated by the unexpectedly high emittance growth in the stage preceding it. The emittance growth lead to both large charge transport losses and large beam spot sizes at the plasma which made it impossible to produce the beam densities required for the experiment.

8.2 Future Directions

The lessons learned in the first attempt at producing plasma density transition trapping experimentally are directly applicable to future research plans. Both the successes achieved and problems encountered in this experiment suggest the modifications necessary for a revised trapping experiment, as well as new experimental directions that should be pursued.

8.2.1 Continuation of the Low Density Trapping Experiment

The problems with the drive beam are the first set of issues that need to be addressed before a second attempt at the low density plasma trapping experiment is made. Unless the drive beam can be made to better match the design parameters, especially in terms of beam density, there is little hope of achieving plasma electron trapping. Significant improvements to the drive beam will require a coordinated effort to diagnose and reduce the beam emittance.

In the short term, six months to a year, the most significant improvement that can be made is in the vacuum window. The 10 μm aluminum window can be replaced with something thinner that would reduce emittance growth through scattering. The only difficulty with this modification is the requirement of the

FNPL facility that the window be able to hold off 1 atm without leaking. A $2\text{ }\mu\text{m}$ polymer window intended to replace the aluminum window during the first run of the experiment failed this test. It is not clear whether the change to a thinner foil would, by itself, significantly improve the beam parameters delivered to the experiment. It may also be possible to make some improvements to the laser spot in this time frame. Dedicated studies by UCLA and FNAL collaborators aimed at diagnosing and improving the beam are planned.

In the longer term, a year or more, the FNPL is planning several improvements that should significantly improve the driver beam parameters. The first of these is an upgrade to a new drive laser oscillator which should improve both the spot quality and stability of the cathode drive laser. This should lead to better emittance out of the photoinjector. The second improvement planned is a beamline upgrade that will add another stage of acceleration and increase the path length parameter, see Eq. 1.30, of the magnetic compressor. Both these enhancements will allow for the same level of beam compression with less energy spread. Energy spread leads to chromatic aberrations in the focusing optics which produces larger spot sizes at the foil, which in turn leads to larger emittance growth. Reducing energy spread while maintaining the same compression will improve the overall beam parameters. It is also possible that a differential pumping system could be developed on this time scale which would allow the foil to be eliminated.

Work to improve the density transition would also enhance the performance of a future experiment. If the pre-transition density roll off observed in the transition measurements, see Section 5.3.3, can be eliminated, the peak density at the top of the transition should be enhanced. The rise in peak density will increase the magnitude of the density drop and lead to greater charge capture.

Table 8.1: Comparison of the parameters proposed for a thin underdense plasma lens experiment at the Neptune Lab [65] and the parameters achieved at the FNPL.

	Neptune Proposed	FNPL Achieved
Peak Plasma Density	$1.3 \times 10^{12} \text{ cm}^{-3}$	$3 \times 10^{13} \text{ cm}^{-3}$
Plasma Column Width FWHM	2 cm	6 cm
Beam Charge	4 nC	8 nC
Beam Duration FWHM	30 ps	15 - 30 ps
Initial Beam Radius σ_r	400 μm	400 μm
Beam Density	$2.6 \times 10^{12} \text{ cm}^{-3}$	$1 - 2 \times 10^{13} \text{ cm}^{-3}$

8.2.2 Underdense Plasma Lens Experiment

The relative ease with which strong plasma focusing was achieved during the commissioning of the plasma density transition trapping experiment, see Section 7.1, argues strongly for a future, rigorous underdense plasma lens experiment. All of the beam and plasma parameters originally envisioned for a Neptune based underdense plasma lens experiment [65] were matched or exceeded during the first run of the trapping experiment, see Table 8.1. The execution of such an experiment would require little more than the removal of the density screen from the plasma column. An updated version of the Neptune underdense plasma lens experiment could also be developed to take advantage of the higher beam and plasma densities now available at the FNPL.

8.2.3 Prospects for High Density Trapping Experiments

To proceed beyond a proof-of-principle transition trapping experiment will necessarily require scaling to higher plasma densities. This advance will require

improvements to both the driver beam and higher density plasma sources with sharp transitions. The production of very short, high current electron drive beams is a matter discussed at great length elsewhere [45, 46, 88]. High power laser pulses are also being considered as alternative drivers for transition trapping [55]. Ideas for producing plasmas with transition that satisfy Eq. 5.6 at high densities $n \geq 10^{14}\text{cm}^{-3}$ are still in the conceptual phase. Possible techniques for producing these transitions include laser ionization of a dual density gas jet and photo-ionization of Lithium using a laser with a step function intensity profile. The development of diagnostics to measure the plasma transition at these densities will also be necessary.

8.2.4 Foil Trapping

In the extreme limit, one can imagine creating an ultra-sharp transition into a plasma by simply replacing the high density plasma region in a transition trapping scenario with a solid metal foil. Electrons would be provided for trapping from the foil via Fowler-Nordheim field emission [34]. Since this situation is much easier to produce experimentally than sharp plasma density drops, it was examined briefly during the development of the experiment discussed in Chapter 5.

The field values necessary for significant Fowler-Nordeim emission are easy to achieve in current plasma wake field experiments. N. Barov et al. have produced wake fields $\geq 140\text{MeV/m}$ in a 10^{14}cm^{-3} plasma at FNAL [90]. In this experiment, the drive beam enters the plasma through a metal foil, one side of which is immersed in the plasma and experiences the large plasma fields. Taking a reasonable value of $\beta \geq 50$ for the microscope surface field enhancement factor of the foil, Fowler-Nordheim theory predicts a large emission $J \geq 100\text{Amp/mm}^2$ under these conditions. Unfortunately, the emission of charge does not guarantee

Table 8.2: Comparison of foil trapping α parameters.

Accelerating Structure	E_{max}	$frequency$	v_ϕ	α
1.6 Cell Photoinjector	80 MeV/m	2.856 GHz	c	2.6
Barov et al. Wake Field Experiment (7 nC)	300 MeV/m	90 GHz	c	0.3
Experiment with High Charge Driver (70 nC)	1.5 GeV/m	~ 90 GHz	c	1.6

that the emitted charge will be trapped and accelerated, see Chapter 2. The charges emitted from the foil due to the plasma wake fields start essentially at rest and must be accelerated to resonance with the wave within the same period of the plasma wake. This situation is analogous to that in RF photoinjectors and the same dimensionless parameter α , see Section 2.1, can be used to evaluate the plasma wake's potential to capture foil electrons. The capture of electrons starting from rest typically requires $\alpha \geq 1$. If we compare the α parameters of the Barov et al. experiment and a standard 1.6 cell photoinjector, see Table 8.2, we see that a plasma wake is not capable of capturing charge from a foil in this regime since its α is only 0.3. The frequency of the accelerating wave is too high in comparison to the accelerating field and the emitted particles can not achieve resonance with the wave.

The peak accelerating field can be increased by increasing the driver beam charge. If this is done while holding the plasma density constant, the plasma frequency will remain essentially unchanged and α will increase. The driver charge can be increased to the point where $\alpha > 1$ and charge is captured from the foil in the plasma wake. If the driver charge in the Barov et al. experiment is increased by a factor of ten, the α of the system reaches 1.6 and charge is captured. The trapping behavior predicted by the α parameter has been verified

by initial MAGIC 2D simulations. Further work needs to be done to explore the parameter space of foil trapping and characterize the captured beams.

APPENDIX A

Derivation of the Hamiltonian for Particles in Electromagnetic Fields

The electromagnetic forces acting on a charged particle in an accelerating system are described by the Lorentz force equation,

$$F_{Lorentz} = \frac{d\vec{p}}{dt} = q \left(\vec{E} + \vec{v} \times \vec{B} \right) = q \left[-\vec{\nabla}\phi_e - \frac{\partial \vec{A}}{\partial t} - (\vec{v} \cdot \vec{\nabla}) \vec{A} \right] = q \left[-\vec{\nabla}\phi_e - \frac{d\vec{A}}{dt} \right], \quad (\text{A.1})$$

where $\vec{p}$ is the particle momentum, q is the charge, $\vec{E}$ is the electric field, $\vec{v}$ is the particle velocity, $\vec{B}$ is the magnetic field, ϕ_e is the electrostatic scalar potential, and $\vec{A}$ is the electromagnetic vector potential. Since the forces acting on a charged particle are derived from both a scalar potential ϕ_e and vector potential $\vec{A}$, special care must be taken in constructing a Hamiltonian for systems governed by this force equation. Let us define the canonical momenta $p_{i,c}$ for this problem as

$$p_{i,c} = p_i + qA_i. \quad (\text{A.2})$$

The Hamiltonian, which is equal to the sum of the mechanical energy of the particle and the potential energy of the field, is given by

$$H = \gamma m_0 c^2 + q\phi_e, \quad (\text{A.3})$$

where the first term is the relativistic total energy of particle and $\gamma = 1/\sqrt{1 - \beta^2}$, where $\beta^2 = (\vec{v}/c)^2$, is the Lorentz factor. With these definitions, Hamilton's

equation of motion for this canonical momentum gives

$$\frac{dp_{i,c}}{dt} = -\frac{\partial H}{\partial x_i} = -q\frac{\partial \phi_e}{\partial x_i}, \quad (\text{A.4})$$

which when applied to the derivative of Eq. A.2

$$\frac{dp_i}{dt} = \frac{dp_{i,c}}{dt} - q\frac{dA_i}{dt} = q\left[-\frac{\partial \phi_e}{\partial x_i} - \frac{dA_i}{dt}\right], \quad (\text{A.5})$$

reproduces the Lorentz force equation. This demonstrates the validity of the definitions chosen above.

In order to find the form the above Hamiltonian, Eq. A.3, in terms the canonical momentum we begin with definition of these momentum

$$p_{i,c} = p_i + qA_i = \frac{\partial L}{\partial \dot{x}_i} = \gamma m_0 \dot{x}_i + qA_i, \quad (\text{A.6})$$

where use has been made of the Lagrangian L as well as the definition relating it to the canonical momentum. Integration of this equation yields

$$L(\vec{x}, \dot{\vec{x}}) = -\frac{m_0 c^2}{\gamma} + q\vec{A} \cdot \vec{v} - q\phi_e(\vec{x}), \quad (\text{A.7})$$

where the freedom allowed by the integration of the partial differential equation has been used to include the electrostatic potential energy $q\phi_e$, which is a conservative scalar potential that depends only on $\vec{x}$.

Utilizing the definition that relates the Lagrangian and the Hamiltonian gives

$$H = \vec{p}_c \cdot \dot{\vec{x}} - L = \frac{\vec{p}_c \cdot (\vec{p}_c - q\vec{A})}{\gamma m_0} - L = \frac{(\vec{p}_c - q\vec{A})^2}{\gamma m_0} + \frac{m_0 c^2}{\gamma} + q\phi_e, \quad (\text{A.8})$$

where substitutions have been made using Eqs. A.2 and A.7. Multiplying both sides of this equation by $\gamma m_0 c^2$ and using Equ A.3 gives

$$(H - q\phi_e)^2 = (\vec{p}_c - q\vec{A})^2 c^2 + (m_0 c^2)^2, \quad (\text{A.9})$$

and finally

$$H = \sqrt{\left(\vec{p}_c - q\vec{A}\right)^2 c^2 + (m_0 c^2)^2} + q\phi_e. \quad (\text{A.10})$$

This is the one dimensional Hamiltonian of a charged particle in an electromagnetic field described by the Lorentz force equation.

APPENDIX B

Heat Transport in the Radiation Coupled Regime

The problem of transporting heat between two objects, e.g. the graphite heater and LaB₆ cathode described in Section 5.2, using only thermal radiation can be understood analytically. It is essentially the classic problem of the heat transported between two infinite plains, one at temperature T_1 with emissivity ϵ_1 , and the other at temperature T_2 with emissivity ϵ_2 , see Fig B.1. The two plates are assumed to be isolated by a perfect vacuum so that the only exchange of heat between them is through black body radiation. In this analysis we must consider both the radiation emitted by each surface and the many reflections of radiation between the surfaces. Let us begin with the radiation emitted from surface 1 per unit area, which is given by Eq. 5.2 if the surface area is omitted,

$$\left. \frac{dQ}{dt} \right|_{emitted1} = \epsilon_1 \sigma T_1^4. \quad (\text{B.1})$$

This radiation is then partially absorbed at the second surface. The rest of the radiation is reflected so that,

$$\left. \frac{dQ}{dt} \right|_{absorbed\ at\ 2,\ 1st\ order} = \epsilon_2 \epsilon_1 \sigma T_1^4, \quad (\text{B.2})$$

and,

$$\left. \frac{dQ}{dt} \right|_{reflected\ at\ 2,\ 1st\ order} = (1 - \epsilon_2) \epsilon_1 \sigma T_1^4. \quad (\text{B.3})$$

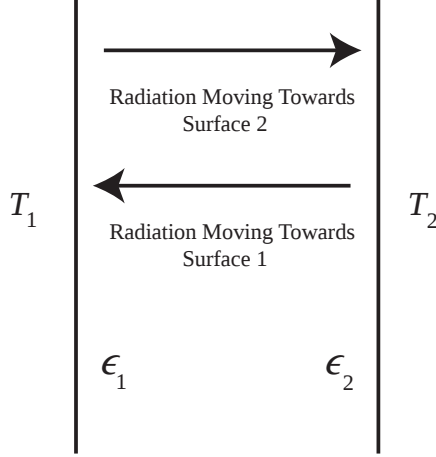


Figure B.1: Diagram of the radiative heat transfer between two planes.

When the radiation reflected from the second surface hits the first surface, a portion of this incident reflected radiation is reflected back towards the second surface again,

$$\left. \frac{dQ}{dt} \right|_{\text{reflected back at 2, from 1}} = (1 - \epsilon_1)(1 - \epsilon_2)\epsilon_1\sigma T_1^4. \quad (\text{B.4})$$

Again only a portion of this reflected radiation arriving from surface 1 is absorbed by surface 2,

$$\left. \frac{dQ}{dt} \right|_{\text{absorbed at 2, 2nd order}} = (1 - \epsilon_1)(1 - \epsilon_2)\epsilon_2\epsilon_1\sigma T_1^4, \quad (\text{B.5})$$

and the remainder is reflected back at surface 1 again. This cycle of reflection and partial absorption will continue indefinitely producing ever higher terms. We can write the total radiation absorbed at surface 2 as an infinite series,

$$\left. \frac{dQ}{dt} \right|_{\text{total absorbed at 2, from 1}} = \epsilon_1\epsilon_2\sigma T_1^4 [1 + (1 - \epsilon_1)(1 - \epsilon_2) + (1 - \epsilon_1)^2(1 - \epsilon_2)^2 + \dots]. \quad (\text{B.6})$$

At this point we recall that,

$$\frac{1}{(1 - x)} = 1 + x + x^2 + \dots, \quad (\text{B.7})$$

so that Equation B.6 can be written,

$$\left. \frac{dQ}{dt} \right|_{total\ absorbed\ at\ 2,\ from\ 1} = \frac{\epsilon_1 \epsilon_2 \sigma T_1^4}{1 - (1 - \epsilon_1)(1 - \epsilon_2)}. \quad (B.8)$$

Precisely the same equation can be written for the heat absorbed by surface 1 from surface 2. The only change needed is to replace T_1 with T_2 . We can then take the difference between these two absorbed heats to find the net heat flow in the system per unit area,

$$\left. \frac{dQ}{dt} \right|_{net} = \frac{\epsilon_1 \epsilon_2 \sigma (T_2^4 - T_1^4)}{1 - (1 - \epsilon_1)(1 - \epsilon_2)}. \quad (B.9)$$

Clearly, the best transport is achieved when both surfaces are perfect absorbers $\epsilon_1 = \epsilon_2 = 1$. As one would expect, no heat is transported in the case when either surface is a perfect reflector, $\epsilon_1 = 0$ or $\epsilon_2 = 0$.

Now, in the case of the cathode heater discussed earlier, we have for LaB_6 $\epsilon_{\text{LaB}_6} = 0.82$ and for carbon $\epsilon_C = 0.81$. From Section 5.2, the LaB_6 cathode has a temperature of about 1300°C and a surface area of 44 cm^2 . The graphite heater has roughly the same surface area and a temperature of about 1500°C . Applying these values directly to Eq. B.9 gives 650 W for the heat transferred between the heater and the cathode. Since both sides of the heater emit radiation, however, and most of the radiation of the back side is reflect towards the cathode, we can effectively double this figure and estimate the heat transported as 1300 W . The heat lost from the cathode is simply the heat radiated by its front surface, which must balance the heat transported to the cathode when the system is in equilibrium. If this number is calculated using Eq. 5.2 and the above values, we find 1250 W for the power lost by the cathode, which agrees very well with the estimate of heat transported to the cathode. The total power emitted by the heater, once again as given by Eq. 5.2 using a temperature of 1500°C and an area of 88 cm^2 , is 3993 W which agrees well with the 3.9 kW input power,

see Section 5.2. The 2.7 kW of this total power which is not transported to the cathode is deposited throughout the system as waste heat.

APPENDIX C

Analysis of Electrostatic Probe Signals

Electrostatic probes operate by drawing electron and ion currents out of the plasma. The probe itself is usually the tip of a wire, or some other metal object with known surface area, which is placed into the plasma. In order to make a measurement, a voltage sweep, typically about -100 to +100 Volts, is applied to the probe using an external power supply and the current drawn out of the plasma is monitored. Analysis of the current versus voltage, or I-V, curve gives both the plasma temperature and density [69, 70].

Figure C.1 shows a typical I-V curve from electrostatic probe measurements of the plasma produced in the trapping experiment plasma source (see Section 5.2). The I-V curve is composed of three distinct regions, which are labelled A, B, and C in Fig. C.1. When the probe is positively biased, it attracts plasma electrons. The plasma electron current drawn into the probe does not, however, simply increase linearly with the voltage applied to the probe. This deviation from linearity occurs because the plasma electrons redistribute so as to shield the bulk of the plasma from the electrostatic fields of the probe. The thin region of modified plasma density that provides this shielding lies between the probe surface and the undisturbed plasma and is referred to as the *plasma sheath*. Because of the shielding effect of the plasma sheath, the plasma electron current drawn out of the plasma at high positive probe biases saturates. Due to this saturation effect, region A of Fig. C.1 is referred to as the electron saturation

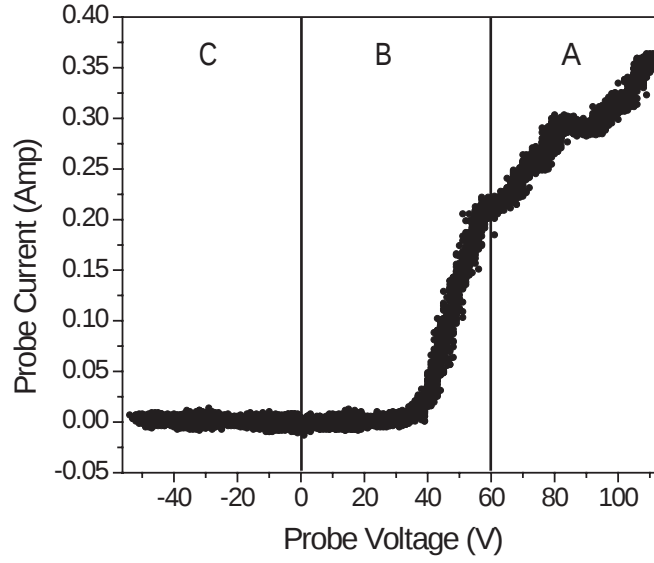


Figure C.1: I-V curve from an electrostatic probe measurement. Region A is the plasma electron saturation current. Region B is the plasma electron current fall off. Region C is the ion saturation current.

current. The electron saturation current drawn in by an ideal probe is given by calculating the current of plasma electrons that hits the probe due to random thermal motion, which can be written as,

$$I_{eSat} = j_{random} A_{probe}, \quad (C.1)$$

where I_{eSat} is the plasma electron saturation current, j_{random} is the current density of plasma electrons hitting the probe through random thermal motion, and A_{probe} is the area of the probe. The quantity j_{random} can be written in terms of the average velocity with which plasma electrons cross a plane in space. Substituting the result of a calculation of this average velocity for a Maxwellian plasma into Eq. C.1 gives,

$$I_{eSat} = \frac{1}{2} q_e n_0 \sqrt{\frac{2k_b T_e}{\pi m_e}} A_{probe} \quad (C.2)$$

where q_e is the electron charge, n_0 is the unperturbed plasma density, m_e is the electron mass, T_e is the electron temperature, and k_b is Boltzmann's constant.

Note the Eq. C.2 indicates that the plasma electron saturation current is a constant independent of the probe voltage, which is clearly not the behavior shown in Fig. C.1. The growth of the current in region A of Fig. C.1 results from the fact that the area of plasma sheath A_{sheath} , not the area of the probe A_{probe} , is what actually determines the plasma electron current collected. For small probe voltages $A_{sheath} \approx A_{probe}$, but as the probe voltage increases so does the sheath area.

The voltage at the boundary between regions A and B in Fig. C.1 is referred to as the plasma *space potential*. This the point at which the probe is at the same potential as the plasma so there are no electric fields between the probe and the plasma. There are no sheaths under these conditions. When the probe potential drops below the space potential, the probe has an effective negative bias with respect to the plasma, even through the probe voltage is still positive with respect to ground. As the probe becomes negative with respect to the plasma it begins to repel the plasma electrons. A significant number of the plasma electrons can, however, overcome this potential barrier, until it becomes quite large, due to their high temperature. The fall off of plasma electron current with decreasing probe potential is shown in region B of Fig. C.1. Assuming that the thermal distributions of the plasma electrons is Maxwellian, the fall off of the plasma density, and hence the plasma electron current density, with the probe potential is given by Boltzmann's relation,

$$n = n_0 \exp \left[\frac{q_e \Delta \phi}{k_b T_e} \right], \quad (C.3)$$

where n is the plasma density at a particular point, and $\Delta \phi = V_s - V_{probe}$ the difference between the space potential and the probe potential. If this modification to the plasma density is included in Eq. C.2, we arrive at an equation for

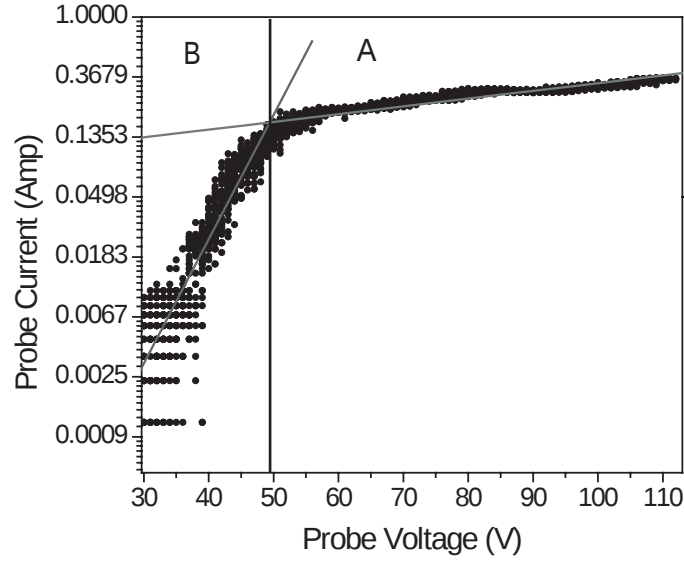


Figure C.2: I-V curve from an electrostatic probe measurement with curve fits. As in Fig. C.1, region A is the plasma electron saturation current, and region B is the plasma electron current fall off.

the plasma electron current in region B,

$$I_e = \frac{1}{2} q_e n_0 \sqrt{\frac{2k_b T_e}{\pi m_e}} A_{probe} \exp \left[\frac{q_e \Delta \phi}{k_b T_e} \right]. \quad (C.4)$$

Eq. C.4 can be used, in conjunction with the I-V curve, to find the plasma temperature and density. As can be seen from Eq. C.4, the rate of the plasma electron current exponential growth is determined by the temperature. If we take the natural logarithm of both sides of Eq. C.4 we find that,

$$\ln I_e = C + \left(\frac{e}{k_b T_e} \right) V_{probe}, \quad (C.5)$$

where C is a constant composed of all the terms that do not depend on V_{probe} . It is clear from this equation that if the the natural logarithm of I_e is plotted verses V_{probe} , the slope of the exponential growth region in this plot is the inverse of the plasma electron temperature in eV. Fig. C.2 shows this type of log-linear plot for the I-V trace in Fig. C.1. The fit to the exponential growth in region B gives a

slope of 0.207 [1/eV]. Inverting we find that the plasma electron temperature is 5 eV. The log-linear plot also clearly shows the point of transition between exponential plasma electron current growth and plasma electron saturation current. If linear fits are made to these two regions, the intersection of these two fits can be taken as a accurate value for I_e at the point where $\Delta\phi = 0$. Substituting these values into Eq. C.4 and solving for n_0 in practical units gives,

$$n_0 = \frac{3.73 \times 10^{13} I_e [\text{Amp}]}{A_{sheath} [\text{mm}^2] \sqrt{k_b T_e [\text{eV}]} } \quad (\text{C.6})$$

The intersection of the linear fits in Fig C.2 gives $I_e = 0.1675$ Amp at $\Delta\phi = 0$. Substituting this value for I_e , along with the measured value of 5 eV for the plasma electron temperature, and the known value of 0.11 mm² for the probe area into Eq. C.6 gives a plasma density of 2.5×10^{13} cm⁻³.

Region C in Fig. C.1 is called the ion saturation current. It is the analog to electron saturation current for negative probe biases. The ion saturation current is weaker than the electron saturation current due to the lower thermal velocity of the ions, and requires more amplification to accurately measure. The I-V curve in Fig. C.1 is not optimized for observing ion saturation current and the current therefore erroneously appears to be zero. Measuring the ion saturation current is generally considered to be a more accurate means to determine the plasma density because it is less susceptible to effects like secondary electron emission, which can skew electron saturation current measurements. Once the ion saturation current I_{ionSat} is found, the density of an argon plasma can be determined by using the equation,

$$n_0 = \frac{8.03 \times 10^{15} I_{ionSat} [\text{Amp}]}{A_{sheath} [\text{mm}^2] \sqrt{k_b T_e [\text{eV}]} } \quad (\text{C.7})$$

APPENDIX D

Optical Design of the Broad Range Vacuum Spectrometer

The beam optics associated with the broad range in-vacuum spectrometer are complicated by the large range of beam energies produced by the density transition trapping experiment. The optics of the system can be analyzed using matrix equations of the form,

$$\begin{pmatrix} x_{final} \\ x'_{final} \end{pmatrix} = M_{Transport} \begin{pmatrix} x \\ x' \end{pmatrix}. \quad (D.1)$$

A detailed treatment of the matrix description of electron beam transport can be found in Ref [28]. Let us begin with the matrices that describe the beam optics associated the spectrometer:

$$M_x = \begin{pmatrix} 1 & B \\ 0 & 1 \end{pmatrix} \begin{pmatrix} 1 & 0 \\ -\frac{\tan \theta_E}{R} & 1 \end{pmatrix} \begin{pmatrix} \cos \theta_b & R \sin \theta_b \\ -\frac{\sin \theta_b}{R} & \cos \theta_b \end{pmatrix} \begin{pmatrix} 1 & 0 \\ \frac{\tan \theta_E}{R} & 1 \end{pmatrix} \begin{pmatrix} 1 & A \\ 0 & 1 \end{pmatrix}, \quad (D.2)$$

$$M_y = \begin{pmatrix} 1 & B \\ 0 & 1 \end{pmatrix} \begin{pmatrix} 1 & 0 \\ \frac{\tan \theta_E}{R} & 1 \end{pmatrix} \begin{pmatrix} 1 & R \theta_b \\ 0 & 1 \end{pmatrix} \begin{pmatrix} 1 & 0 \\ -\frac{\tan \theta_E}{R} & 1 \end{pmatrix} \begin{pmatrix} 1 & A \\ 0 & 1 \end{pmatrix}, \quad (D.3)$$

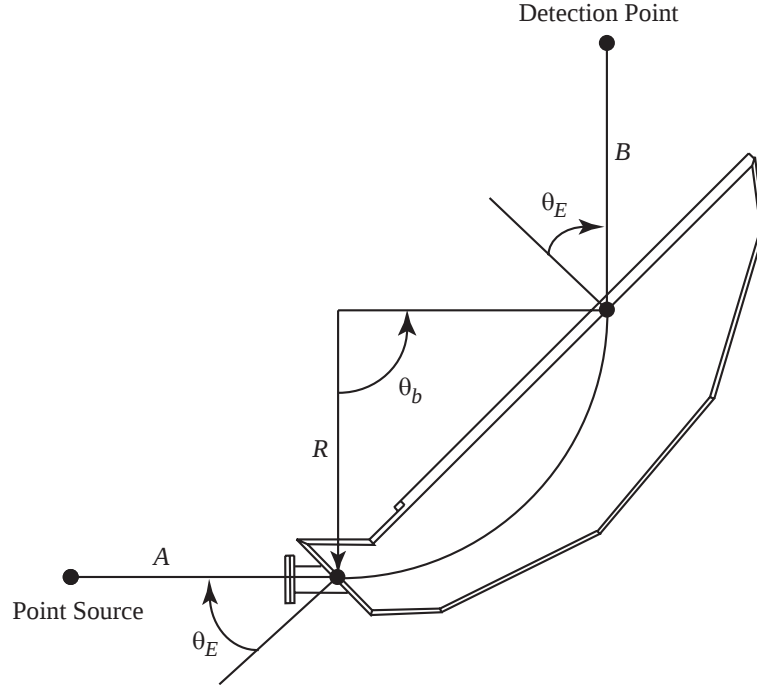


Figure D.1: Variable definitions for the spectrometer optics calculations.

where A is the distance from the point source to the entrance of the spectrometer, θ_E is the angle from the outward line normal to the magnet edge to the beam trajectory which is taken as positive when it points toward the center of magnetic curvature, θ_b is the bend angle, and B is the distance from the spectrometer exit to the point where the beam is detected, see Fig D.1. Now in our magnet,

$$\theta_b = \frac{\pi}{2}, \quad \theta_E = \frac{\pi}{4}, \quad (\text{D.4})$$

so that Eq. (D.2) and Eq. (D.3) reduce to,

$$M_x = \begin{pmatrix} 1 - \frac{2B}{R} & A - B - \frac{2AB}{R} + R \\ -\frac{2}{R} & -1 - \frac{2A}{R} \end{pmatrix}, \quad (\text{D.5})$$

$$M_y = \begin{pmatrix} 1 - \frac{\pi(B+R)}{2R} & \frac{R(B(2+\pi)+\pi R)-A(B\pi+(\pi-2)R)}{2R} \\ -\frac{\pi}{2R} & \frac{1}{2}(2 + \pi - \frac{A\pi}{R}) \end{pmatrix}. \quad (\text{D.6})$$

In order to accurately measure energy with the spectrometer, we want to satisfy the point-to-point condition in x , $M_{x12} = 0$, so that the initial angle of the beam will not produce the appearance of additional energy spread. Satisfying the point-to-point condition requires A , B , and R to fulfill the relationship,

$$M_{x12} = A - B - \frac{2AB}{R} + R = 0 \implies B = R \frac{A + R}{2A + R}. \quad (\text{D.7})$$

Typically, we would want to send a parallel beam into the spectrometer, so that $A \rightarrow \infty$ and Eq. (D.7) reduces to $B = R/2$. Unfortunately, this is a difficult thing to do in the plasma transition trapping experiment since the drive beam and captured beam have very different energies and both have significant energy spread. This means that chromatic aberrations will be significant in any focusing systems used to collimate the beams, especially the low energy captured beam. Therefore, we would like to avoid the use of focusing optics between the experiment and spectrometer. From the practical considerations of available hardware, and the need to observe the diverging beams while they are small, the value of A has been set at 30cm.

With $A = 30\text{cm}$, the point-to-point condition is described by the equation,

$$B = R \frac{30\text{cm} + R}{60\text{cm} + R} \quad \text{with} \quad 10\text{cm} \leq R \leq 58\text{cm}. \quad (\text{D.8})$$

Since one of the main design goals of this new spectrometer is the ability to keep the beam in vacuum until the point of detection, strict fulfillment of the point-to-point condition would require the vacuum vessel to have the complex figure described by Eq. (D.8). This is an extremely difficult technical problem, hence we would like to examine the errors introduced by approximating Eq. (D.8) with a linear equation. The simplest way to do this is to Taylor expand Eq. (D.7)

about the center of the spectrometer $R = 30\text{cm}$,

$$B \approx \frac{30(30 + A)}{30 + 2A} + \left[\frac{60 + A}{30 + 2A} - \frac{30(30 + A)}{(30 + 2A)^2} \right] (R - 30) + \dots \quad (\text{D.9})$$

If we take $A = 30\text{cm}$ and linearize we find,

$$B = \frac{7}{9}R - \frac{10}{3}. \quad (\text{D.10})$$

We will also want to examine errors arising from uncertainty in A since the exit of the plasma is not precisely defined.

To begin this analysis we will make the substitutions

$$A = A_0 + a \quad \text{and} \quad B = B_0 + b, \quad (\text{D.11})$$

in Eq. (D.5) and Eq. (D.6), where A_0 and B_0 are the design distances and a and b are the error in the distances. For the moment we will only examine M_x , which becomes

$$M_x = \begin{pmatrix} (1 - \frac{2B_0}{R}) - \frac{2b}{R} & (A_0 - B_0 - \frac{2A_0B_0}{R} + R) + a - b - \frac{2(A_0b + aB_0 + ab)}{R} \\ -\frac{2}{R} & (-1 - \frac{2A_0}{R}) - \frac{2a}{R} \end{pmatrix}. \quad (\text{D.12})$$

Now, we still want the design distances to fulfill the point-to-point condition so

$$M_{x12} = A_0 - B_0 - \frac{2A_0B_0}{R} + R = 0 \implies B_0 = R \frac{A_0 + R}{2A_0 + R}, \quad (\text{D.13})$$

and Eq. (D.12) becomes

$$M_x = \begin{pmatrix} (1 - \frac{2B_0}{R}) - \frac{2b}{R} & a - b - \frac{2(A_0b + aB_0 + ab)}{R} \\ -\frac{2}{R} & (-1 - \frac{2A_0}{R}) - \frac{2a}{R} \end{pmatrix}. \quad (\text{D.14})$$

What we are really interested in is the final size of the beam, in the x dimension, do solely to the optics of the system,

$$x_{\text{optical}} = \left[M_x \begin{pmatrix} x \\ x' \end{pmatrix} \right]_1 = \left[\left(1 - \frac{2B_0}{R} \right) - \frac{2b}{R} \right] x + \left[a - b - \frac{2(A_0b + aB_0 + ab)}{R} \right] x', \quad (\text{D.15})$$

Table D.1: Simulated parameters of the trapped and drive beam

	σ_x	$\sigma_{x'}$	E_{ave}	ΔE_{rms}
Drive Head	$759\mu\text{m}$	-0.002	14 MeV	0.75%
Drive Middle	$702\mu\text{m}$	0.012	13 MeV	3.4%
Drive Tail	2 mm	0.025	12.5 MeV	2.3%
Trapped Beam Core	$596\mu\text{m}$	0.028	1.22 MeV	4.6%

where $A_0 = 30\text{cm}$, B_0 is a function of R given by Eq. (D.13),

$$-2.5\text{cm} \leq a \leq 2.5\text{cm}, \quad \text{and} \quad b = \left[\frac{7}{9}R - \frac{10}{3} \right] - \left[R \frac{30 + R}{60 + R} \right], \quad (\text{D.16})$$

the difference between the linearized and exact point-to-point conditions.

In order to determine the error introduced by linearizing the point-to-point condition we need to compare $x_{optical}$ to the size that the beam will have at the exit of the spectrometer due to energy spread. To good approximation this quantity is given by

$$\Delta x_{energy} = \frac{\sqrt{2} \left(\frac{v}{c} \right) \Delta E [\text{MeV}]}{299.8 B_{mag} [\text{T}]}, \quad (\text{D.17})$$

where the familiar formula for the bend radius of an electron in a magnetic has been used (with suitable generalization for low energy electrons) and the $\sqrt{2}$ has been added to account for the 45° angel of the spectrometer edge. Since the exit port will probably not be parallel to the spectrometer edge, this formula is not strictly correct, but the correction should be slight.

The relevant beam parameters for both the drive beam and the captured beam can be derived from simulations of the trapping process and are listed in Table D.1. The parameters used for these simulation were essentially the ones presented in Section 3.2.

Taking these values we can use Eqs. D.15, D.16, and D.17 to plot the er-

ror, expressed as $x_{optical}/\Delta x_{energy}$, over the parameter space of interest which is $-2.5\text{cm} \leq a \leq 2.5\text{cm}$ and $10\text{cm} \leq R \leq 58\text{cm}$, see Figs. D.2-D.5.

From Figures D.2-D.5, it is clear that for both the captured and drive beams the maximum value of $x_{optical}/\Delta x_{energy}$ over the entire parameter space of interest is 0.009, or only 0.9%. Over most of the parameter space the size of the optical spread, as compared to the energy spread, is much less. Since even the maximum error produced by the linearization of the exit port appears to be negligible, a linear exit port based on Eq. D.10 was chosen as the best design choice.

The errors induced by linearizing the exit port have no significant impact on the beam dynamics in the y direction. Calculations using Eq. D.3 and data from the simulations indicates that all elements of the beams should be visible at the spectrometer exit without clipping, although it may be necessary in some cases to use the bottom third of the R range in order to provide sufficient vertical focusing.

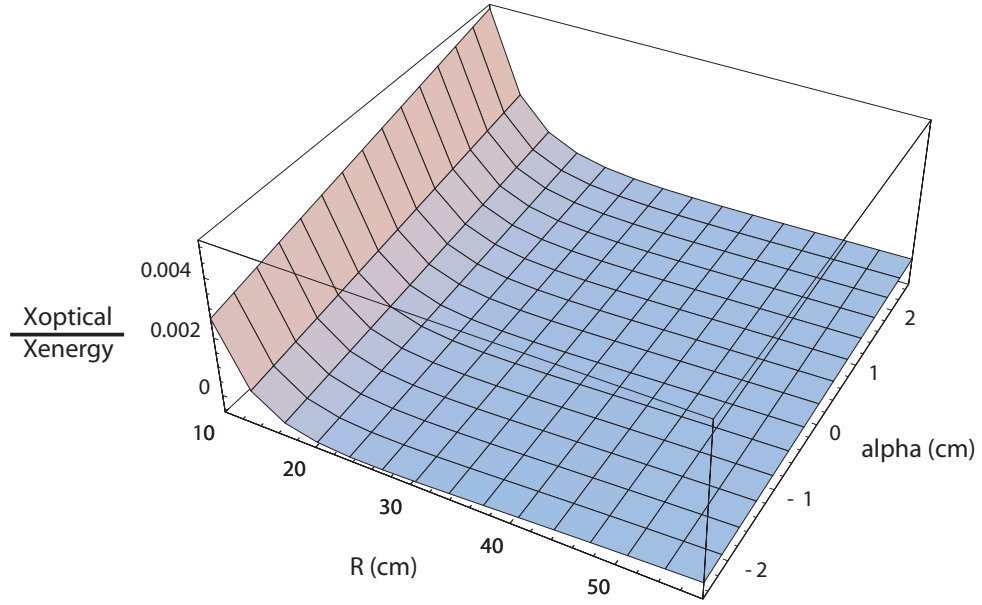


Figure D.2: The error expected when observing the captured beam.

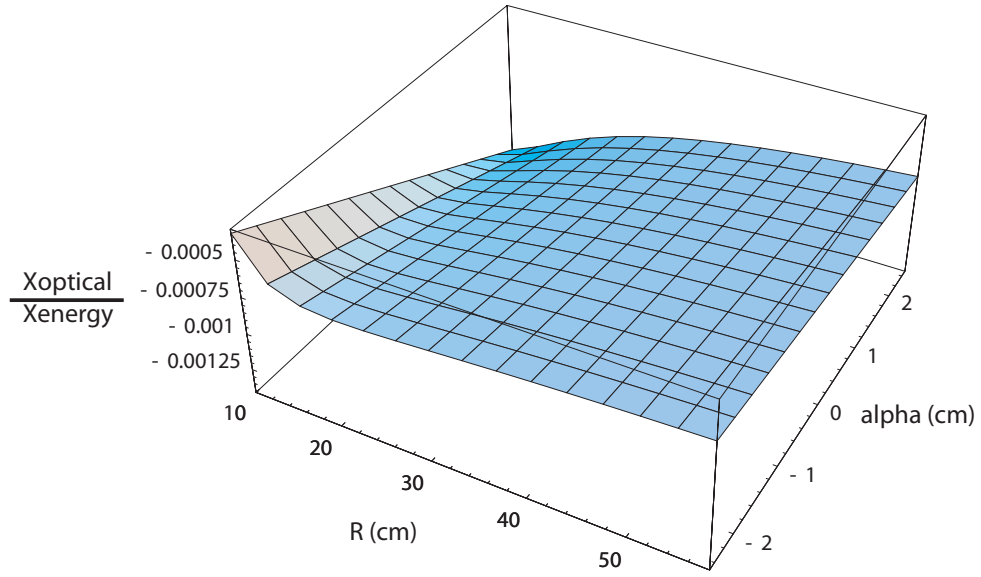


Figure D.3: The error expected when observing the head of the drive beam.

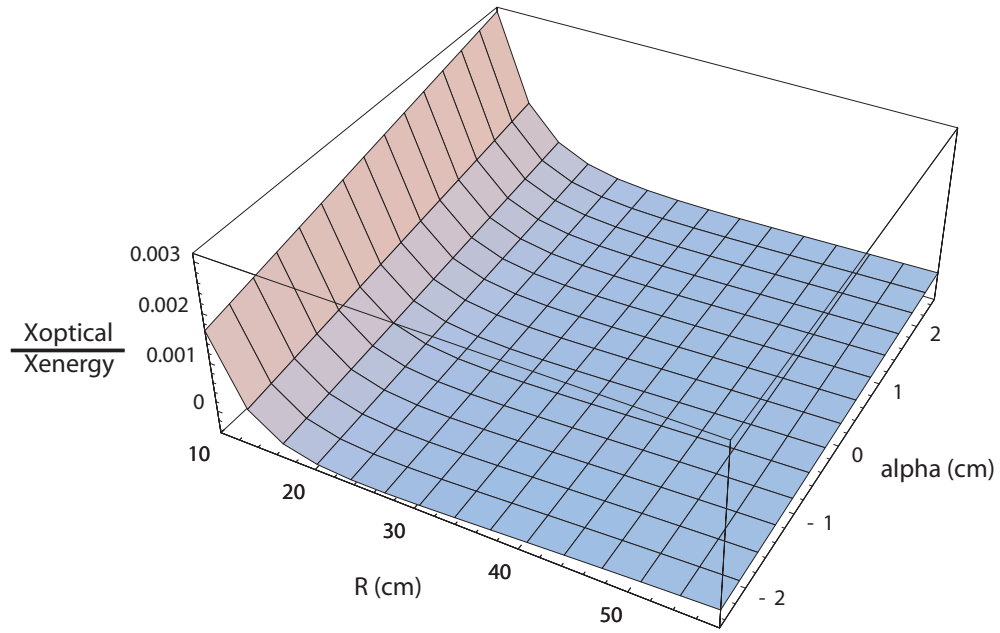


Figure D.4: The error expected when observing the middle of the drive beam.

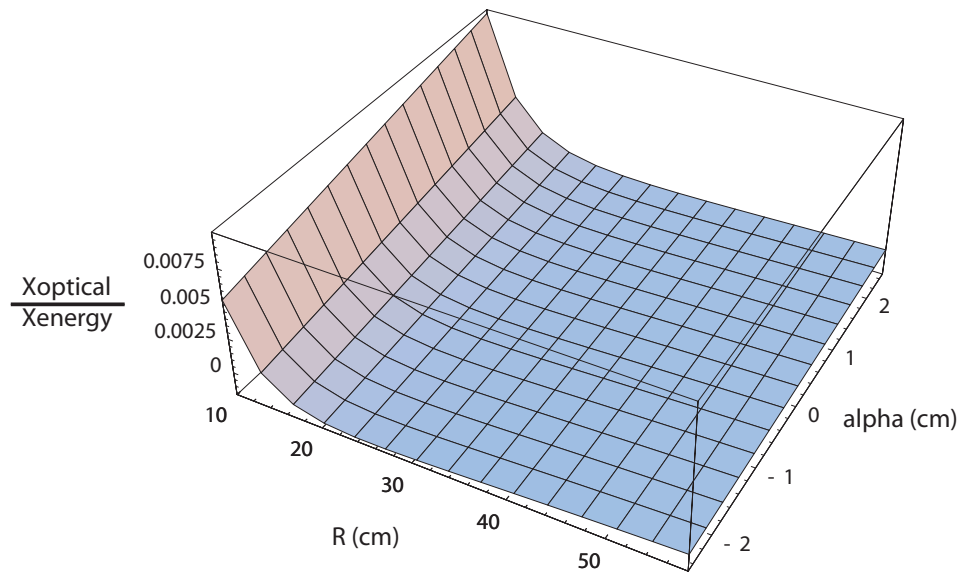


Figure D.5: The error expected when observing the tail of the drive beam.

REFERENCES

- [1] J. W. Wang and G. A. Loew. Field Emission and RF Breakdown in High-Gradient Room-Temperature Linac Structures. October 1997. SLAC-PUB-7684.
- [2] W. D. Kilpatrick. Criterion for Vacuum Sparking Designed to Include Both rf and dc. *Review of Scientific Instruments*, 28(10):824, October 1957.
- [3] G. A. Loew and J. W. Wang. RF Breakdown and Field Emission. January 1989. SLAC-PUB-4845.
- [4] David L. Burke. The NLC Technical Program. In Robert H. Siemann, editor, *Symposium on Electron Linear Accelerators In Honor of Richard B. Neal's 80th Birthday*, page 35, 1998. SLAC-Report-526.
- [5] Chandrashekhar Joshi and Thomas Katsouleas. Plasma Accelerators at the Energy Frontier and on Tabletops. *Physics Today*, pages 47–53, June 2003.
- [6] E. Esarey, P. Sprangle, J. Krall, and A. Ting. Overview of Plasma-Based Accelerator Concepts. *IEEE Transactions on Plasma Science*, 24(2):252, April 1996.
- [7] Y. B. Fainberg, V. A. Balakirev, I. N. Onishchenko, G. L. Sidel'nikov, and G. V. Sotnikov. Wake field excitation in plasma by a train of relativistic electron bunches. *Plasma Physics Reports*, 20(7):606, July 1994. Along with references therein.
- [8] Pisin Chen, J. M. Dawson, Robert W. Huff, and T. Katsouleas. Acceleration of Electrons by the Interaction of a Bunched Electron Beam with a Plasma. *Physical Review Letters*, 54(7):693, February 1985.
- [9] R. D. Ruth, A. W. Chao, P. L. Morton, and P. B. Wilson. A Plasma Wake Field Accelerator. *Particle Accelerators*, 17:171, 1985.
- [10] Pisin Chen, J. J. Su, J. M. Dawson, K. L. Bane, and P. B. Wilson. On energy transfer in the plasma wake field accelerator. *Phys. Rev. Lett.*, 56:1252–1255, 1986.
- [11] R. J. England, J. B. Rosenzweig, and M. C. Thompson. Longitudinal Beam Shaping and Compression Scheme for the UCLA Neptuen Laboratory. In *Advanced Accelerator Concepts: Tenth Workshop*, volume 647 of *AIP Conf. Proc.*, page 884, 2002.

- [12] J. B. Rosenzweig. Nonlinear Plasma Dynamics in the Plasma Wake-Field Accelerator. *Physical Review Letters*, 58(6):555, February 1987.
- [13] J.B. Rosenzweig, B. Breizman, T. Katsouleas, and J.J. Su. Acceleration and focusing of electrons in two-dimensional nonlinear plasma wake fields. *Phys. Rev. A*, 44:R6189, 1991.
- [14] James Benjamin Rosenzweig, D. B. Cline, B. Cole, H. Figueroa, W. Gai, R. Konecny, J. Norem, P. Schoessow, and J. Simpson. Experimental Observation of Plasma Wake-Field Acceleration. *Phys. Rev. Lett.*, 61(1):98, 1988.
- [15] J. B. Rosenzweig, P. Schoessow, B. Cole, W. Gai, R. Konecny, J. Norem, and J. Simpson. Experimental measurement of nonlinear plasma wake fields. *Physical Review A*, 39(3):1586, February 1989.
- [16] Atsushi Ogata. Plasma Lens and Wake Experiments in Japan. In *Advanced Accelerator Concepts*, volume 279 of *AIP Conf. Proc.*, page 420, 1993.
- [17] N. Barov, J.B. Rosenzweig, M.E. Conde, W. Gai, and J.G.Power. Observation of plasma wakefield acceleration in the underdense regime. *Phys. Rev. Special Topics - Accel. Beams*, 3:011301, 2000.
- [18] C. D. Barnes, C. O'Connell, F.-J. Decker, P. Emma, M. J. Hogan, P. Iverson, P. Krejcik, R. H. Siemann, D Walz, B. Blue, C. E. Clayton, C. Huang, D Johnson, C. Joshi, K. A. Marsh, W. B. Mori, S. Deng, T. Katsouleas, P. Muggli, and E. Oz. Improvements for the Third Generation Plasma Wakefield Experiment E-164 at SLAC. In *Proc. PAC 2003*, page 1530. IEEE, 2003.
- [19] M.J. Hogan, R. Assmann, F.-J. Decker, R. Iverson, P. Raimondi, S. Rokni, R.H. Siemann, D. Waltz, D. Whittum, B. Blue, C.E. Clayton, E. Dodd, R. Hemker, C. Joshi, K.A. Marsh, W.B. Mori, S. Wang, T. Katsouleas, S. Lee, P. Muggli, P. Catravas, S. Chattopadhyay, E. Esarey, and W.P. Leemans. E-157: A 1.4-m-long plasma wake field acceleration experiment using a 30 GeV electron beam from the Stanford Linear Accelerator Center Linac. *Physics of Plasmas*, 7:2241, 2000.
- [20] B. E. Blue, C. E. Clayton, C. L. O'Connell, F.-J. Decker, M. J. Hogan, C. Huang, R. Iverson, C. Joshi, T. C. Katsouleas, W. Lu, K. A. Marsh, W. B. Mori, P. Muggli, R. Siemann, and D. Walz. Plasma-Wakefield Acceleration of an Intense Positron Beam. *Physical Review Letters*, 90(21):214801, 2003.
- [21] T. Tajima and J. M. Dawson. Laser Electron Accelerator. *Physical Review Letters*, 43(4):267, July 1979.

- [22] R. D. Hazeltine and F. L. Waelbroeck. *The Framework of Plasma Physics*. Perseus Books, 1998.
- [23] F. Amiranoff, S. Baton, D. Bernard, B. Cros, D. Descamps, F. Dorchies, F. Jacquet, V. Malka, J. R. Marques, G. Matthieussent, P. Mine, A. Modena, P. Mora, J. Morillo, and Z. Najmudin. Observation of Laser Wakefield Acceleration of Electrons. *Phys. Rev. Lett.*, 81(5):995, 1998.
- [24] M. N. Rosenbluth and C. S. Liu. Excitation of Plasma Waves by Two Laser Beams. *Physical Review Letters*, 29(11):701, September 1972.
- [25] C. E. Clayton, C. Joshi, C. Darrow, and D. Umstadter. Relativistic Plasma-Wave Excitation by Collinear Optical Mixing. *Physical Review Letters*, 54(21):2343, 1985.
- [26] S. Ya. Tochitsky, R. Narang, C. V. Filip, P. Musumeci, C. E. Clayton, R. B. Yoder, K. A. Marsh, J.B. Rosenzweig, C. Pellegrini, and C. Joshi. Enhanced Acceleration of Injected Electrons in a Laser-Beat-Wave-Induced Plasma Channel. *Physical Review Letters*, 92(9):095004, 2004.
- [27] M.I.K. Santala, Z. Najmudin, E. L. Clark, M. Tatarakis, K. Krushelnick, A. E. Dangor, V. Malka, J. Faure, R. Allott, and R. J. Clarke. Observation of a Hot High-Current Electron Beam from a Self-Modulated Wakefield Accelerator. *Physical Review Letters*, 86(7):1227, 2001.
- [28] J. B. Rosenzweig. *Fundamentals of Beam Physics*. Oxford University Press, 2003.
- [29] D. A. Edwards and M. J. Syphers. *An Introduction to the Physics of High Energy Accelerators*. John Wiley and Sons, Inc., 1993.
- [30] C. A. Brau. What Brightness Means. In *The Physics and Applications of High Brightness Electron Beams*, World Scientific, page 20, 2003.
- [31] J. R. Pierce. *Theory and Design of Electron Beams*. D. Van Nostrand Company, Inc., Princeton, New Jersey, second edition, 1954.
- [32] S. Bernal, A. Dragt, M. Reiser, M. Venturini, J. G. Wang, and T. F. Godlove. Matching section and inflector design for a model electron ring. *Fusion Engineering and Design*, 32-33:277–281, 1996.
- [33] Charles A. Brau and Warren B. Mori. Working group 3 summary: Exotic source physics: photon-assisted field emission, spin polarization, cathodeless (plasma) sources, etc. In James Rosenzweig and Luca Serafini, editors, *The Physics of High Brightness Beams*, page 393. World Scientific, 2000.

- [34] R. H. Fowler and L. Nordheim. Electron Emission in Intense Electric Fields. *Proceedings of the Royal Society of London*, A119(781):173–181, May 1928.
- [35] O. W. Richardson. *Phil. Mag.*, 28(5):633, 1914.
- [36] Saul Dushman. Electron Emission from Metals as a Function of Temperature. *Physical Review Series II*, 21:623, 1923.
- [37] C. D. Child. *Phys. Rev. Ser. I*, 32:492, 1911.
- [38] I. Langmuir. *Phys. Rev. Ser. II*, 2:450, 1913.
- [39] K. Ebihara and S. Hiramatsu. High-current-density gun with a LaB₆ cathode. *Review of Scientific Instruments*, 67(8):2765, August 1996.
- [40] S. H. Kong, J. Kinross-Wright, D. C. Nguyen, and R. L. Sheffield. Photocathodes for free electron lasers. *Nuclear Instruments and Methods in Physical Research A*, 358:272, 1995.
- [41] S. G. Anderson, M. Loh, P. Musumeci, J. B. Rosenzweig, H. Suk, and M. C. Thompson. Commissioning and measurements of the Neptune photoinjector. In *Advanced Accelerator Concepts: Ninth Workshop*, volume 569 of *AIP Conf. Proc.*, page 487, 2001.
- [42] Patrick G. O’Shea. RF Photoinjectors. In James Rosenzweig and Luca Serafini, editors, *The Physics of High Brightness Beams*, page 17. World Scientific, 2000.
- [43] M. C. Thompson and J. B. Rosenzweig. Production and synchronization of electron beams from RF photoinjector/compressor systems for ultra-fast applications. In *Advanced Accelerator Concepts: Ninth Workshop*, volume 569 of *AIP Conf. Proc.*, page 347, 2001.
- [44] J.-P. Carneiro, R. A. Carrigan, M. S. Champion, P. L. Colestock, H. T. Edwards, J. D. Fuerst, W. H. Hartung, K. P. Koepke, M. Kuchnir, J. K. Santucci, L. K. Spentzouris, M. J. Fitch, A. C. Melissinos, P. Michelato, C. Pagani, D. Sertore, N. Barov, and J. B. Rosenzweig. First results of the fermilab high-brightness rf photoinjector. In *Proc. PAC 1999*, page 2027. IEEE, 1999.
- [45] P. Musumeci, R. J. England, M. C. Thompson, R. Yoder, and J. B. Rosenzweig. Velocity Bunching Experiment at the Neptune Laboratory. In *Advanced Accelerator Concepts: Tenth Workshop*, volume 647 of *AIP Conf. Proc.*, page 858, 2002.

- [46] J. B. Rosenzweig, N. Barov, and E. Colby. Pulse Compression in RF Photoinjectors: Applications to Advanced Accelerators. *IEEE Transactions on Plasma Science*, 24:409, 1996.
- [47] S. G. Anderson, J. B. Rosenzweig, P. Musumeci, and M. C. Thompson. Horizontal Phase-Space Distortions Arising from Magnetic Pulse Compression of an Intense, Relativistic Electron Beam. *Physical Review Letters*, 91(7):074803, August 2003.
- [48] S. Heifets, G. Stupakov, and S. Krinsky. Coherent synchrotron radiation instability in a bunch compressor. *Physical Review Special Topics - Accelerators and Beams*, 5:064401, 2002.
- [49] C. E. Clayton, K. A. Marsh, A. Dyson, M. Everett, A. Lal, W. P. Leemans, R. Williams, and C. Joshi. Ultrahigh-gradient acceleration of injected electrons by laser-excited relativistic electron plasma waves. *Phys. Rev. Lett.*, 70(1):37, 1993.
- [50] H. Dewa, A. Mizuno, T. Taniuchi, H. Tomizawa, T. Asaka, T. Kobayashi, S. Suzuki, K. Yanagida, H. Hanaki, and M. Uesaka. Beam Quality and Stability Improvements for a Single-Cell Photocathode RF Gun. In *The Physics and Applications of High Brightness Electron Beams*, World Scientific, page 28, 2003.
- [51] E. M. Lifshitz and L. P. Pitaevskii. *Physical Kinetics*, volume 10 of *Landau and Lifshitz Course of Theoretical Physics*. Butterworth-Heinemann, 1981.
- [52] S. Bulanov, N. Naumova, F. Pegoraro, and J. Sakai. Particle injection into the wave acceleration phase due to nonlinear wake wave breaking. *Phys. Rev. E*, 58(5):R5257, 1998.
- [53] D. Umstadter, J.K. Kim, and E. Dodd. Laser Injection of Ultrashort Electron Pulses into Wakefield Plasma Waves. *Phys. Rev. Lett.*, 76:2073, 1996.
- [54] Hyyong Suk, Nick Barov, James Benjamine Rosenzweig, and E Esarey. Plasma Electron Trapping and Acceleration in a Plasma Wake Field Using a Density Transition. *Phys. Rev. Lett.*, 86(6):1011, 2001.
- [55] Hyyong Suk. Electron acceleration based on self-trapping by plasma wake fields. *Journal of Applied Physics*, 91(1):487, 2002.
- [56] M. C. Thompson, J. B. Rosenzweig, and H. Suk. Plasma density transition trapping as a possible high-brightness electron beam source. *Physical Review Special Topics - Accelerators and Beams*, 7:011301, 2004.

- [57] A. I. Akhiezer and R. V. Polovin. Theory of Wave Motion of an Electron Plasma. *Soviet Physics, JETP*, 3(5):696, 1956.
- [58] Robert J. Noble. Plasma Beat-Wave Accelerator. In *Proceedings of the Twelfth International Conference on High Energy Accelerators, Batavia, IL, 1983*, page 467. Fermilab, 1984.
- [59] J. B. Rosenzweig. Trapping, thermal effects, and wave breaking in the non-linear plasma wake-field accelerator. *Physical Review A*, 38(7):3634, October 1988.
- [60] R. J. England, J. B. Rosenzweig, and N. Barov. Plasma electron fluid motion and wave breaking near a density transition. *Physical Review E*, 66:016501, 2002.
- [61] Bruce Goplen, Larry Ludeking, David Smithe, and Gary Warren. User-configurable MAGIC for electromagnetic PIC calculations. *Computer Physics Communications*, 87:54–86, 1995.
- [62] M. C. Thompson, C. E. Clayton, J. England, J. B. Rosenzweig, and H. Suk. Beam-Plasma Interaction Experiments at the UCLA Neptune Laboratory. In *Proc. PAC 2001*, page 4014. IEEE, 2001.
- [63] M. Ferrario, P. R. Bolton, J. E. Clendenin, D. H. Dowell, S. M. Gierman, M. E. Hernandez, D. Nguyen, D. T. Palmer, J. B. Rosenzweig, J. F. Schmerge, and L. Serafini. New Design Study and Related Experimental Program for the LCLS RF Photoinjector. In *Proc. EPAC 2000*, page 1642. Austrain Acad. Sci. Press, 2000.
- [64] J.B. Rosenzweig and E. Colby. Charge and Wavelength Scaling of RF Photoinjector Designs. In *Advanced Accelerator Concepts*, volume 335 of *AIP Conf. Proc.*, page 724, 1995.
- [65] H. Suk, C. E. Clayton, G. Hairapetian, C. Joshi, M. Loh, P. Muggli, R. Narang, C. Pellegrini, J. B. Rosenzweig, and T. C. Katsouleas. Underdense Plasma Lens Experiment at the UCLA Neptune Laboratory. In *Proc. PAC 1999*, page 3708. IEEE, 1999.
- [66] H. Suk, C. E. Clayton, C. Joshi, T. C. Katsouleas, P. Muggli, R. Narang, C. Pellegrini, and J. B. Rosenzweig. Plasm Source Test and Simulation Results for the Underdense Plasma Lens Experiment at the UCLA Neptune Laboratory. *IEEE Transactions on Plasma Science*, 28(1):271, February 2000.

- [67] D. M. Goebel, Y. Hirooka, and T. A. Sketchley. Large-area lanthanum hexaboride electron emitter. *Review of Scientific Instruments*, 56(9):1717, September 1985.
- [68] N St J Braitwaite. Introduction to gas discharges. *Plasma Sources Science and Technology*, 9:517–527, 2000.
- [69] H. M. Mott-Smith and Irving Langmuir. The Theory of Collectors in Gaseous Discharges. *Physical Review*, 28:727–763, October 1926.
- [70] Francis F. Chen. *Plasma Diagnostic Techniques*, chapter 4, pages 113–200. Academic Press, 1965. Huddleston and Leonard editors.
- [71] A. Bauer. The Arc Cathode. In S. C. Haydon, editor, *An Introduction to Discharge and Plasma Physics*, page 319. University of New England, Armidale, N.S.W., Australia, 1964.
- [72] J. M. Lafferty. Boride Cathodes. *Journal of Applied Physics*, 22(3):299–309, March 1951.
- [73] R. A. Fonseca, L. O. Silva, F. S. Tsung, V. K. Decyk, W. Lu, C. Ren, W. B. Mori, S. Deng, S. Lee, T. Katsouleas, and J. C. Adam. OSIRIS: A Three-Dimensional, Fully Relativistic Particle in Cell Code for Modeling Plasma Based Accelerators. In P.M.A. Sloot, C.J. Kenneth Tan, J.J. Dongarra, and A.G. Hoekstra, editors, *Computational Science - ICCS 2002: International Conference, Amsterdam, The Netherlands, April 21-24, 2002. Proceedings, Part III*, page 342. Springer-Verlag Heidelberg, 2002. LNCS 2331.
- [74] A. R. Fry, M. J. Fitch, A. C. Melissinos, N. P. Bigelow, B. D. Taylor, and F. A. Nezrick. Laser System for the TTF Photoinjector at Fermilab. In *Proc. PAC 1997*, page 2867. IEEE, 1998.
- [75] Michael J. Fitch. *Electro-Optic Sampling of Transient Electric Fields from Charged Particle Beams*. PhD thesis, University of Rochester, Rochester, New York, 2000.
- [76] Eric R. Colby. *Design, Construction, and Testing of a Radiofrequency Electron Photoinjector for the Next Generation Linear Collider*. PhD thesis, University of California, Los Angeles, Los Angeles, California, 1997.
- [77] B. E. Carlsten. New Photoelectric Injector Design for the Los Alamos National Laboratory XUV FEL Accelerator. *Nuclear Instruments and Methods in Physical Research A*, 285:313, 1989.

- [78] Luca Serafini and James B. Rosenzweig. Envelope analysis of intense relativistic quasilaminar beams in rf photoinjectors: A theory of emittance compensation. *Physical Review E*, 55(6):7565, 1997.
- [79] D. Sertore, S. Schreiber, K. Floettmann, F. Stephan, K. Zapfe, and P. Michelato. First operation of cesium telluride photocathodes in the TTF injector RF gun. *Nuclear Instruments and Methods in Physical Research A*, 445:422, 2000.
- [80] TESLA, The Superconducting Electron-Positron Linear Collider with an Integrated X-Ray Laser Laboratory, Technical Design Report Part II: The Accelerator. Technical Report 2001-011, DESY, 2001.
- [81] K. Hagiwara *et al.* Review of Particle Physics. *Physical Review D*, 66:010001, 2002. Section 26: Passage of Particles Through Matter.
- [82] A. Murokh, J.B. Rosenzweig, I Ben-Zvi, X. Wang, and V. Yakimenko. Limitation on Measuring a Transverse Profile of Ultra-Dense Electron Beams with Scintillators. In *Proc. PAC 2001*, page 1333. IEEE, 2001.
- [83] Crytur Ltd. Palackeho 175, 51101 Turnov, Czech Republic. See <http://www.crytur.cz>.
- [84] Eljen Technology. PO Box 870, 300 Crane Street, Sweetwater, TX 79556, USA. See <http://www.eljentechnology.com>.
- [85] Bergoz Instrumentation. Espace Allondon Ouest, 01630 Saint Genis Pouilly, France. See <http://www.bergoz.com>.
- [86] J.J. Su, T. Katsouleas, J. M. Dawson, and R Fedele. Plasma lenses for focusing particle beams. *Physical Review A*, 41(6):3321, 1990.
- [87] J. T. Seeman. Transverse and longitudinal emittance measurements. In *Accelerator Physics and Engineering*. World Scientific, 1999.
- [88] Scott G. Anderson. *Creation, Manipulation, and Diagnosis of Intense, Relativistic Picosecond Photo-electron Beams*. PhD thesis, University of California, Los Angeles, Los Angeles, California, 2002.
- [89] K. R. Crandall and D. P. Rusthoi. TRACE 3-D Documentation, Third Edition. Technical Report LA-UR-97-886, Los Alamos National Laboratory, 1997.

- [90] N. Barov, K. Bishofberger, James Benjamine Rosenzweig, J. P. Carneiro, P. Colestock, H. Edwards, M. J. Fitch, W. Hartung, and J. Santucci. Ultra High-Gradient Energy Loss by a Pulsed Electron Beam in a Plasma. In *Proc. PAC 2001*, page 126. IEEE, 2001.