



UNIVERSITA' DEGLI STUDI DI TRENTO

Department of Physics



Doctoral School in Physics - XXVI Cycle

---

FINAL THESIS

**Positronium in the AEgIS  
experiment: study on its  
emission from nanochanneled  
samples and design of a new  
apparatus for Rydberg  
excitations**

Candidate:

**Lea Di Noto**

Supervisor:

**Prof. Roberto S. Brusa**

Defended on 24th January 2014



*a mio papà*



# Contents

|                                                              |            |
|--------------------------------------------------------------|------------|
| <b>Abstract</b>                                              | <b>vii</b> |
| <b>1 Introduction</b>                                        | <b>1</b>   |
| 1.1 Antimatter history . . . . .                             | 1          |
| 1.2 The importance of antimatter studies . . . . .           | 3          |
| 1.2.1 CPT tests . . . . .                                    | 3          |
| 1.2.2 Gravity and antimatter . . . . .                       | 4          |
| 1.3 AEgIS Experiment . . . . .                               | 6          |
| 1.4 Positronium in the AEgIS experiment . . . . .            | 11         |
| 1.5 Content description . . . . .                            | 13         |
| <b>2 Positronium physics</b>                                 | <b>15</b>  |
| 2.1 Ps formation in solids from slow positron beam . . . . . | 15         |
| 2.1.1 Ps formation in metals . . . . .                       | 17         |
| 2.1.2 Ps formation in semiconductors . . . . .               | 18         |
| 2.1.3 Ps formation in insulators . . . . .                   | 18         |
| 2.2 Nanochanneled silicon samples . . . . .                  | 22         |
| 2.2.1 Electrochemical etching . . . . .                      | 22         |
| 2.2.2 Ps yield . . . . .                                     | 24         |
| 2.2.3 Ps cooling . . . . .                                   | 27         |
| <b>3 Ps Time of Flight measurements</b>                      | <b>35</b>  |
| 3.1 The high intensity positron beam NEPOMUC . . . . .       | 35         |
| 3.2 The experimental apparatus . . . . .                     | 37         |
| 3.2.1 The transport line . . . . .                           | 37         |
| 3.2.2 Beam intensity selection by turnable disc . . . . .    | 41         |
| 3.2.3 The electrostatic lenses . . . . .                     | 41         |
| 3.2.4 The sample holder . . . . .                            | 43         |
| 3.3 The acquisition chain . . . . .                          | 44         |
| 3.4 The channeltrons signal . . . . .                        | 46         |
| 3.5 The NaI(Tl) scintillators signal . . . . .               | 48         |

|          |                                                             |            |
|----------|-------------------------------------------------------------|------------|
| 3.6      | ToF and Energy spectra . . . . .                            | 50         |
| 3.7      | Preliminary results . . . . .                               | 55         |
| 3.7.1    | Measurements by using the Trento positron beam . . .        | 56         |
| 3.7.2    | Measurements at the NEPOMUC facility . . . . .              | 67         |
| 3.8      | Parameters optimization . . . . .                           | 70         |
| 3.8.1    | The signal to background ratio . . . . .                    | 70         |
| 3.8.2    | Error analysis and parameters optimizations . . . . .       | 76         |
| 3.9      | Conclusions . . . . .                                       | 82         |
| <b>4</b> | <b>Apparatus for Ps excitations in Rydberg states</b>       | <b>85</b>  |
| 4.1      | AEgIS accumulator . . . . .                                 | 85         |
| 4.2      | Overview of the apparatus . . . . .                         | 87         |
| 4.2.1    | The magnetic transport line . . . . .                       | 88         |
| 4.2.2    | The electron optical line . . . . .                         | 92         |
| 4.2.3    | The final focusing on the target . . . . .                  | 99         |
| 4.3      | The electronic design . . . . .                             | 101        |
| 4.3.1    | Technical construction . . . . .                            | 106        |
| 4.3.2    | Tests on the switching circuit . . . . .                    | 108        |
| 4.4      | The sample chamber . . . . .                                | 110        |
| 4.4.1    | Magnetic field shield . . . . .                             | 111        |
| 4.4.2    | Switching circuit of the last electrode VL . . . . .        | 111        |
| 4.4.3    | Switching electric field near the sample . . . . .          | 115        |
| 4.5      | Conclusions . . . . .                                       | 116        |
| <b>5</b> | <b>Ps excitations in Rydberg states and Ps spectroscopy</b> | <b>119</b> |
| 5.1      | Ps levels in free field . . . . .                           | 119        |
| 5.2      | Ps levels in magnetic field . . . . .                       | 122        |
| 5.2.1    | Zeeman effect . . . . .                                     | 124        |
| 5.2.2    | The Motional Stark effect . . . . .                         | 124        |
| 5.3      | Ps excitations in Rydberg states . . . . .                  | 127        |
| 5.3.1    | Excitation from $n=1$ to $n=3$ . . . . .                    | 130        |
| 5.3.2    | Excitation from $n=3$ to $n=20-30$ . . . . .                | 131        |
| 5.4      | Ps Rydberg experiments . . . . .                            | 132        |
| 5.5      | Planned measurements on Rydberg Ps . . . . .                | 136        |
| 5.6      | Laser cooling technique . . . . .                           | 138        |
| <b>A</b> | <b>Trasmission line theory</b>                              | <b>143</b> |
|          | <b>Bibliography</b>                                         | <b>149</b> |
|          | <b>Ringraziamenti</b>                                       | <b>157</b> |

# Abstract

This experimental thesis has been done in the framework of AEgIS (Anti-matter Experiment: Gravity, Interferometry, Spectroscopy), an experiment installed at CERN, whose primary goal is the measurement of the Earth's gravitational acceleration on anti-hydrogen. The antiatoms will be produced by the charge exchange reaction, where a cloud of Ps in Rydberg states interacts with cooled trapped antiprotons. Since the charge exchange cross section depends on Ps velocity and quantum number, the velocity distribution of Ps emitted by a positron-positronium converter as well as its excitation in Rydberg states have to be studied and optimized.

In this thesis Ps cooling and emission into vacuum from nanochannelled silicon targets was studied by performing Time of Flight measurements with a dedicated apparatus conceived to receive the slow positron beam as produced at the Trento laboratory or at the NEPOMUC facility at Munich. Measurements were done by varying the positron implantation energy, the sample temperature and the nanochannel dimensions, with the aim of finding the best parameters to increase Ps fraction having velocity lower than  $5 \cdot 10^4$  m/s. Preliminary data were analyzed in order to extract the Ps velocity distribution and its average temperature. More, an original method for evaluating the permanence time of Ps inside the nanochannels before being emitted into vacuum, was described. A first rough evaluation based on the performed measurements is reported and this result will be useful to investigate the Ps cooling process and to synchronize the laser pulse for Ps excitation in the AEgIS experiment.

In order to perform measurements of Ps excitation in Rydberg states, a new apparatus for bunching positron pulses, coming from the AEgIS positron line, was designed and built. COMSOL and SIMION softwares were used for designing a magnetic transport line and an electron optical line, which extracts positrons from the magnetic field and focus them on the nanochannelled Si sample. More, a buncher device, which spatio-temporally compresses the positron bunches, was built and a fast circuit for supplying the 25 buncher electrodes with a parabolic shaped potential was designed and

tested. According to the simulations, at the target position the device will deliver positrons with an energy ranging from 6 to 9 keV, in bunches of 5 ns duration and a spot of 2.5 mm in diameter.

By using this apparatus, first measurements for the optimization of Ps excitation in Rydberg states and studies on the Ps levels with and without magnetic field will be performed. At a later stage investigations of Ps spectroscopy or of Ps laser cooling with the same apparatus could be achievable for the first time.

# Chapter 1

## Introduction

### 1.1 Antimatter history

In December 1927 Paul Dirac developed a relativistic equation for the electron, now known as the Dirac equation. It describes fields corresponding to elementary spin-1/2 particles (such as the electron) as a vector of four complex numbers (a bispinor), differently from the Schrödinger equation, which described a field of only one complex value. It was the first theory to account fully for relativity in the context of quantum mechanics and it predicted the existence of antimatter for the first time. In fact the equation was found to have negative-energy solutions in addition to the normal positive ones, with the nonsensical consequence that the energy has not a lowest level anymore. As a way of getting around this, Dirac introduced the hypothesis, known as *hole-theory*, that the vacuum is the many-body quantum state, called *Dirac sea*, in which all negative-energy electron eigenstates are occupied. Since the Pauli exclusion principle forbids electrons from occupying the same state, any additional electron would be forced to occupy a positive-energy eigenstate, and positive-energy electrons would be forbidden from decaying into negative-energy eigenstates. Thus if the negative-energy eigenstates are incompletely filled, each unoccupied eigenstate (called a hole) would behave like a positively charged particle, with the same mass of the electron. This hole, identified with the positron, has a positive energy, since some energy is required to create a particle-hole pair from the vacuum. Few years later, in 1932, Carl Anderson, confirmed the Dirac's prediction by observing the positron for the first time, by studying cosmic rays passing through a gas chamber. In the following years many experiments discovered antiprotons (Segrè, 1954) and antineutrons (Cork, 1956) and other heavier antiparticles. Today's Standard Model shows that every particle has an antiparticle, for

which each additive quantum number has the negative of the value it has for the normal matter particle. More for particles, whose additive quantum numbers are all zero, the particle may be its own antiparticle, like photons and the neutral pions. It was in 1965 that A. Zichichi using the Proton Synchrotron at CERN, and L. Lederman using the Alternating Gradient Synchrotron (AGS) accelerator at the Brookhaven National Laboratory, at New York, discovered for the first time the possibility of anti-matter formation, by observing the antideuteron, an anti-nucleus made out of an antiproton plus an antineutron. This fact opened the possibility to try to produce not only antiparticles (that is a quite usual operation in particle accelerators), but more complicated systems of anti-particles to study how subatomic antiparticles behave when they come together. The first goal was the production of the anti-hydrogen that is the simplest atom of all, made of a single antiproton orbited by a positron. A first production of some antihydrogen atoms was performed in 1996 [1] at LEAR (Low Energy Antiproton Ring) at CERN and few years later at Fermilab [2]. The few atoms were produced in collisions between an antiproton beam on Xenon atoms, where the positron of the electron-positron pair generated was captured by the antiproton and it formed anti-hydrogen. Despite the relevant result, the production rate was too small and the atom energy too much high to allow measurements on antimatter properties to be carried out at an interesting level of accuracy. For this reason in 1999 a new machine AD (Antiproton Decelerator), totally dedicated to antiprotons deceleration, was built at CERN. AD, delivering dense beams of  $10^7$  antiprotons per minutes with low energy (100 MeV/c) in bunches of  $\sim 80$  nanoseconds, allowed two CERN experiments in 2002, ATHENA and ATRAP, to trap antiprotons and positrons very close each other and to create thousands of atoms of anti-hydrogen at very low temperature (few tens Kelvin) [3, 4]. The successful results, obtained in producing cold antihydrogen inside electro-magnetic traps, opened the possibility to have antimatter atoms in considerable quantities suitable for gravity or spectroscopic tests. Ever since five different experiments, still present, use the antiprotons delivered by AD at CERN with the purpose of studying antimatter properties: AEGIS (Antimatter Experiment: Gravity, Interferometry, Spectroscopy), ALPHA, ATRAP (Antihydrogen Trap Collaboration) and ASACUSA (Atomic Spectroscopy And Collisions Using Slow Antiprotons). In 2010 ALPHA experiment trapped for the first time antimatter atoms for over 16 minutes, [5] showing that the future measurements of anti-hydrogen spectroscopy, which is the ALPHA and ATRAP primary goal, will be possible. ASACUSA collaboration has a proposal to create hybrid atoms, such as *antiprotonic helium* (that is a helium atom with an antiproton instead of one of the two electrons), for hyperfine measurements and in 2011

the same collaboration measured the antiproton's mass with an accuracy of one part on a billion [6]. On the contrary AEgIS, the experiment in which this thesis is collocated, and GBAR, a new experiment not yet installed at CERN, have as primary goal the study of the gravitational interaction by measuring the Earth's gravitational acceleration on anti-hydrogen. Now the AEgIS experiment is under construction and it is performing the first test with antiprotons.

In all the antimatter experiments the precision on the measurements depends on the anti-hydrogen energy and for this reason ELENA, a new antiproton ring, will be connected in the next years to AD and it will deliver antiprotons at very low energy ( $\sim 5$  keV). In this way the number of trapped antiprotons could be increased and thus the gravity measurement on anti-hydrogen with higher precision will be possible.

## 1.2 The importance of antimatter studies

The importance of antimatter studies is due to the possibility of relevant comparisons between matter and antimatter to confirm or deny fundamental theories, like the CPT theorem through high precision spectroscopy and the Weak Equivalence Principle (WEP) through gravity experiment.

On one hand, since the hydrogen atom is one of the best understood systems in physics, comparison with antihydrogen can be done at high level of precision, and thus a verification of CPT invariance at the same level could be achieved. On the other hand the anti-hydrogen can be used for testing gravity theories on antimatter-matter interactions.

### 1.2.1 CPT tests

The CPT theorem asserts that any Lorentz invariant local quantum field theory with a Hermitian Hamiltonian must have CPT symmetry. This means that all physical phenomena have to be invariant under the combined operations of charge conjugation (C) (replacing of particle with anti-particle), parity (P) (reflection of coordinate respect the axis origin  $\vec{x} \rightarrow -\vec{x}$ ), and time reversal (T). The theorem, though it is embedded in relativistic and quantum theories, concerns all fundamental interaction (except the gravitational one) and thus predicts antimatter properties. For instance, it implies that particle and antiparticle have to be the same inertial mass, same lifetime, gyromagnetic ratio and energy levels and that, in general, a mirror image of our universe, in which all momenta and positions are reflected, all matter is replaced by antimatter and time is reversed, will evolve like our universe.

Despite CPT invariance seems to be held in all experiments performed up to now, being it so fundamental, it is important to test it with the highest level of accuracy using all type of particles (baryons, mesons and leptons). For leptons the direct measurement of the magnetic moment of positron and electron has achieved a precision of  $2 \times 10^{-12}$  [7], while for baryons the charge-to-mass ratio of proton and antiproton has been found equal within  $10^{-10}$  [8]. Finally for mesons, the maximal sensitivity of  $10^{-8}$  was achieved [9] by measuring the difference between kaon and antikaon mass, but in this last case it was not a direct measurement and its sensitivity was model dependent [10]. In this context a precise confirmation of CPT could come from measurements on gross structure, fine structure, Lamb shifts and hyperfine structures in hydrogen and anti-hydrogen. The CPT theory predicts that all these constants will be identical for matter and antimatter systems. Spectroscopic measurements of the energy levels of anti-hydrogen could be, in principle, performed with very high precision and a very good test of CPT could be achieved by a comparison with the analogous transitions of hydrogen. In fact the precision of spectroscopic studies of hydrogen, achieved in the last decades about the ground state hyperfine structure ( $10^{-13}$ ) [11] or the 1S-2S transition frequency ( $10^{-15}$ ) [12], could lead to the most accuracy CPT test on baryons ever made, if the same accuracy would be obtained on the anti-hydrogen atom.

### 1.2.2 Gravity and antimatter

Gravity takes a special place amongst the four fundamental interactions. While the electromagnetic force and the strong and weak nuclear forces are described by Quantum Field Theories, gravity is described by the General Relativity from a geometric point of view: the bodies travel along geodesics in four-dimensional spacetime, which is distorted by the presence of massive objects. The Equivalent Principle, which is the cornerstone of the General Relativity, can be formulated in different forms and it summarizes three aspects: firstly the independence of the trajectories of falling particles from their internal composition (the Weak Equivalence Principle), secondly the independence of the outcome of any local non-gravitational experiment from the velocity of the laboratory (Local Lorentz Invariance) and finally from its location in the space-time (Einstein Equivalence Principle). All these aspects and in particular the WEP, have been tested with ordinary matter with precision down to  $10^{-13}$  by Eötvös-type experiments [13], but the WEP principle has never been tested with antimatter particles. Being the General Relativity a classical theory, the WEP is expected to be violated at some level, when a quantum theory of gravity comes into play.

In the attempt to unify gravity and other fundamental forces, some theories predict that, in principle, antimatter can fall differently from normal matter in the Earth's gravitational field. According to CPT, an antimatter system, let's say an anti-apple, fall into a field generated by a massive antimatter-system (for example a hypothetic anti-earth), exactly as an apple falls to the earth, but CPT does not give any indication on how an anti-apple falls to the Earth leaving room for the possibility that  $g(H) \neq g(\bar{H})$ .

An ultimately provided Theory of Everything was formulated in a quantum context where the attraction (or repulsion) between two bodies is due to the exchange of virtual bosons, which determine, with their spins and their masses, the properties of the force itself. The gravitational force would be mediated by three gravitons [14]: the ordinary Newtonian gravity corresponds to the exchange of massless tensor (spin-2) gravitons while additionally, vector (spin-1) and scalar (spin-0) gravitons may exist. These two added forces would lead both to attractive/repulsive forces if even/odd-bosons are exchanged between interacting massive bodies. The potential of interaction between two massive particles  $m_1$  and  $m_2$  in that case would be:

$$V = -\frac{G}{r} m_1 m_2 (1 \mp v e^{-r/\lambda_v} + s e^{-r/\lambda_s}) \quad (1.1)$$

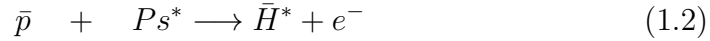
where  $v, s$  represents the coupling strengths and  $\lambda_v$  and  $\lambda_s$  are the distance ranges of vector and scalar potentials. In equation 1.1 tensor and scalar forces are always attractive, while vector forces are attractive between opposite charges and repulsive between same charges, where charge means matter or antimatter. Thus the effect of gravitational forces could have, in principle, macroscopic effects on antimatter falling in the gravity field produced by Earth and, at the same time, could be extremely small in any interacting matter-matter system. Using data from Eötvös-like experiments, that measured the correlation between inertial mass and gravitational mass, limits on the range of the scalar and vector fields have been obtained [15], but anyway the possibility of differential gravitational interactions between matter and antimatter is still open. Such an anomalous anti-gravity could lead to a repulsion between matter and antimatter with the consequential segregation of matter and antimatter in different regions of the cosmos and hence it could also explain the apparent absence of antimatter in the observable universe [16].

The investigation of gravity on antimatter has never been performed until now and thus the measurement on anti-hydrogen could allow the first direct measurement of gravity on antimatter. Two unsuccessful attempts were done with antiprotons [17] and positrons [18], but much high was the electromagnetic interaction of a charged particle to be considered negligible with respect

to the weak gravitational interaction: an electric field lower than  $10^{-11}$  V/m would be necessary to give a positron an acceleration equal to that of gravity. The anti-hydrogen is thus the simplest neutral antimatter system useful for gravity measurements, that are the primary goal of the AEGIS collaborations.

### 1.3 AEGIS Experiment

The AEGIS Experiment (Antimatter Experiment: Gravity Interferometry Spettroscopy) at the Antiproton Deceleration (AD) room at CERN, has as primary goal the direct measurement of the Earth's gravitational acceleration on anti-hydrogen with an initial precision of 1%. The set up is designed for the formation of an anti-hydrogen beam, on which the gravity measurement can be performed by observing the vertical deflection with a classical two grating Moiré deflectometer, coupled with a position sensitive detector. The anti-hydrogen production occurs by the charge-exchange reaction [19, 20]:



where trapped antiprotons interacts with a cloud of excited Positronium atoms, formed by positron implantation into a nanoporous Si sample.

In the AEGIS apparatus (see Fig.1.1) two magnetic lines, one for antiprotons and the other one for positrons, terminate into the two main magnets of 5 T and 1 T magnetic field and in different times positrons and antiprotons bunches enter the same Malmberg-Penning trap, built in the 5 T magnet and kept at a temperature of 4 K.

Firstly, the antiprotons, coming from AD in bunches of  $3 \times 10^7$  particles every 100 s, at a kinetic energy of  $\sim 5$  MeV, lose their energy down to few keV, by traversing a set of thin degrader foils (170  $\mu\text{m}$  of aluminum and 55  $\mu\text{m}$  of silicon) and then are caught in the 5 T trap. Here they are cooled down to few K with sympathetic electron cooling [21, 22] and then are transferred in the 1 T magnet, where, after cooling down to 100 mK by innovative techniques as resistive cooling [23] or evaporative cooling [24], wait for the Ps atoms in the recombination trap for anti-hydrogen production.

Secondly, the positrons, emitted from a radioactive source ( $^{22}\text{Na}$  50 mC) in a beta decay, are slowed down to few eV by a solid Neon moderator [25] and then stacked in a two-stage Surko trap accumulator. Every 100 s the second stage of the Surko-trap (see section 4.1) can deliver positron bunches of  $10^8$  positrons in 5 ns at an energy in the 40-400 eV range. When antiprotons are in the recombination trap in the 1 T magnet, positrons are transported by a magnetic line into the 5 T trap, where they are accumulated up to obtain a very high density plasma, formed by around  $10^8$  positrons in a cylindrical

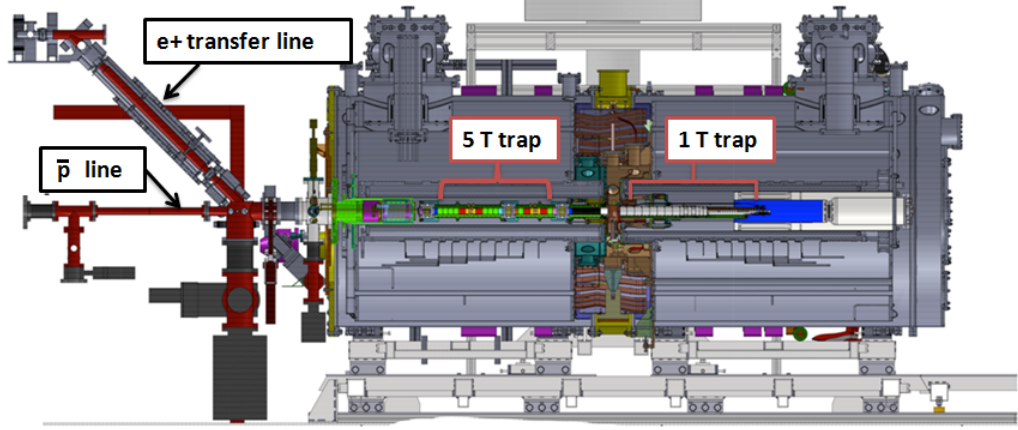


Figure 1.1: The two beam lines for positrons and for antiprotons terminate in the main magnets of 5 T at 4 K (at left) and 1 T with a region at 100 mK (at right). In the 5 T a unique trap is used for antiprotons and for positrons, at different time; in the 1 T region two traps, one of them out of axis, are present. In the on axis recombination trap, the antiprotons are cooled down to 100 mK and anti-hydrogen can be formed and accelerated toward the Moiré deflectometer.

cloud of  $\sim 1$  mm radius. Then the positrons plasma is transferred in the 1 T region, where it is excited to a diocotron mode [26], off-axis displaced and trapped again inside a smaller trap, whose axis is shifted with respect to the symmetry axis of the system (see Fig.1.2). Finally positrons are accelerated up to 7-10 keV and focused on a Silicon nanochanneled sample, where Ps formation occurs. The ortho-positronium (o-Ps) (spin 1, 145 ns lifetime into vacuum) and the para-positronium p-Ps (spin 0, 142 ps lifetime into vacuum) atoms, formed in the material, are emitted into nanochannels with high kinetic energy (1-2 eV). Through collisions with the inner walls of the nanochannels, o-Ps, that has a longer lifetime and it is formed at sufficiently high depth, can thermalize at the sample temperature and diffuse in vacuum at lower energy. In order to increase the o-Ps lifetime together with the cross section for antihydrogen production (eq.(1.2)), which depends on the Ps quantum number and on its velocity as [27]:

$$\sigma \sim \frac{n^4}{v} \quad (1.3)$$

the Ps has to be produced with low velocity and it has to be excited to a Rydberg state with  $n_{Ps} > 20$ . The excitation occurs by means a two steps transition, where two laser pulses excite positronium atoms firstly from

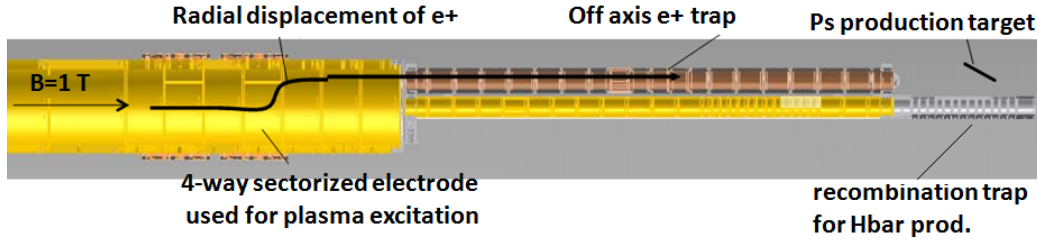


Figure 1.2: Traps in the 1 T region: the big trap on axis and the two smaller traps for positrons (out-of-axis) and for antiprotons (on-axis) are shown. The path of positron bunches is indicated and the target, where Ps atoms are formed is shown. At right the recombination trap, where the antiproton cloud cooled at 100 mK waits for Ps, for anti-hydrogen formation, is visible.

ground state to  $n_{Ps} = 3$ , then from  $n_{Ps} = 3$  to Rydberg state with  $n_{Ps} > 20$ . Thus, when the flying excited Ps atoms interact with the cloud of trapped antiprotons at 100 mK temperature, the charge exchange reaction can occur. The anti-hydrogen atoms are produced with a distribution of levels, related to the energy levels of the excited Ps, like  $n_{\bar{H}} \cong \sqrt{2}n_{Ps}$ , while the  $\bar{H}$  velocity distribution is dominated by the antiproton temperature and it should be in the range of 20-80 m/s if an antiprotons temperature of  $\sim 100$  mK is achieved.

In order to form an anti-hydrogen beam, the anti-atoms must be accelerated along the symmetry axis in last section of the 1 T magnet. Since the  $\bar{H}$  is formed in an excited state, this acceleration can be achieved with an electric field gradient acting on the anti-atom dipole. Since the dipole moment of an atom scales approximately as  $\sim n^2$ , Rydberg atoms are especially suited to be accelerated in this way [28]. In fact because an atom in an electric field  $F$  changes its energy levels as:

$$E = -\frac{1}{2n^2} + \frac{3}{2}nkF \quad (1.4)$$

where the energy  $E$  is expressed in atomic unit and  $k$  is the quantum number running from  $-(n-1-|m|)$  to  $(n-1-|m|)$  in steps of two (with  $m$  azimuthal quantum number), a change on the amplitude of the electric field implies an opposite variation on the kinetic energy to conserve the total energy. If an electric field of some hundreds of V/cm is used, whose value is limited by the field ionization of the Rydberg atoms, an acceleration of some  $10^8 m/s^2$  is generated on the anti-hydrogen, so that it reaches around 500 m/s velocity in 1 cm.

The summary scheme of the anti-hydrogen production and the picture of the

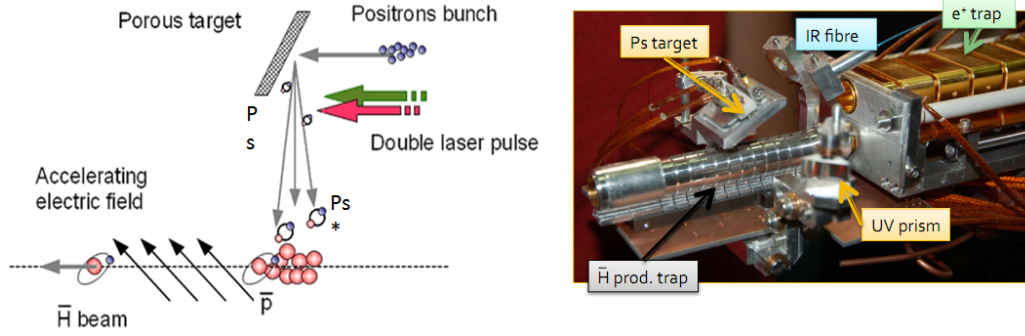


Figure 1.3: Left: scheme of antihydrogen beam production: antiprotons are caught in the trap and cooled down to 100 mK; positrons are focused on a porous Si nanochanneled sample and positronium is produced. After laser-exciting of cooled Ps, antihydrogen is produced via charge exchange reaction (eq.1.2). With an appropriate electric field,  $\bar{H}$  beam is formed and accelerated toward the Moiré deflectometer.

Right: a picture of the Si sample holder with the prism for UV laser and the fiber for IR light are shown. The out-of-axis positron trap and recombination antiprotons trap, where  $\bar{H}$  is formed, are visible.

sample holder with the recombination trap are shown in Fig.1.3.

Once the anti-hydrogen beam is formed, in principle, by knowing the production point of the atom and its horizontal velocity, in order to determine the value of  $g$ , it would be sufficient to measure the vertical displacement of the atom after a certain length of flight. On the contrary this is non feasible in practice, due to the radial velocity spread of about 50 m/s of the anti-hydrogen beam and to the beam spot size, in comparison with the small vertical displacement of around  $35\mu\text{m}$  expected with  $v = 500\text{m/s}$  and a flight path of about 80 cm.

To overcome this problem a Moiré deflectometer is foreseen. The device [29] consists of two equally spaced and parallel gratings, together with a position sensitive microstrip detector to register the impact point of anti-hydrogen atoms (see Fig.1.4). In order to work in the classical regime, the slit widths  $a$  are chosen sufficiently large that diffraction can be neglected in such a way that the condition:

$$a \gg \sqrt{\lambda_B L} \quad (1.5)$$

between  $\lambda_B$ , the de Broglie wavelength of the anti-hydrogen and the geometrical gratings parameters  $L$  and  $a$  (see Fig.1.4), is satisfied.

Thus an interference pattern with the same period of the gratings is formed on the detector and a vertical shift  $\delta$  of the fringe pattern is obtained, de-

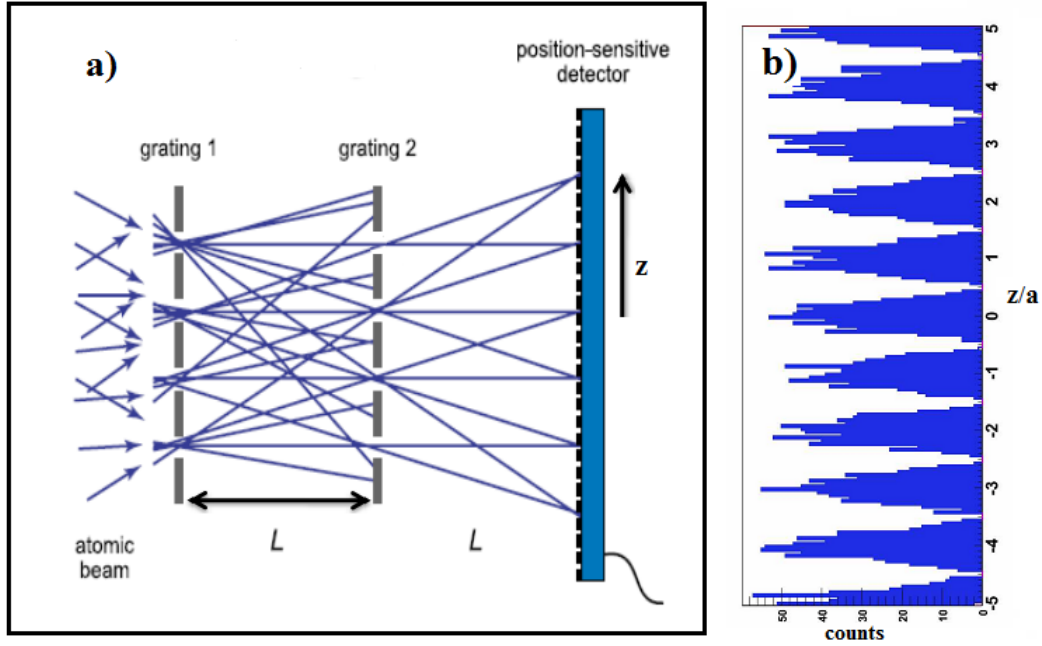


Figure 1.4: a) Scheme of the Moiré deflectometer.  $L$  is the distance between the gratings and  $a$  is the grating period. The anti-hydrogen beam arrives from the left and an interference pattern is shaped on the detector on the right. b) Simulated interference pattern as shaped on the position-sensitive detector. Grating with  $L = 40$  cm and  $a = 80\mu\text{m}$  have been considered in this simulation. On the vertical axis is reported the vertical coordinate  $z/a$ , while counted atoms are present on the horizontal axis. The gravity influences the displacement of fringe pattern on the vertical axis  $z$ .

pending on the anti-hydrogen velocity, on the gravity acceleration and on the period of the gratings  $a$  as:

$$\delta = gt^2/a \quad (1.6)$$

In order to obtain  $g$  from eq.1.6, the time of flight between the time of the electric field switching on and the arrival time of  $\bar{H}$  on the detectors surface has also to be measured and the impact point of antiatoms has to be determined with a spatial accuracy of  $1-2\mu\text{m}$ . By extracting an average time of flight and an average shift for different anti-hydrogen velocities and then by fitting these data by using the equation 1.6, the gravity measurement with an accuracy of 1 % can be obtained if  $10^5$  anti-hydrogen atoms are collected.

## 1.4 Positronium in the AEgIS experiment

Due to the cross section of the charge exchange reaction which depends on Ps velocity and its quantum number (see eq.1.3), the Positronium physics plays an important role for the anti-hydrogen production in the AEgIS experiment. Thus the Ps yield from the sample after positron implantation, its cooling and finally its excitation in the Rydberg state have to be studied and optimized in order to maximize the  $\bar{H}$  formation.

In this context the Trento group, belonging to the AEgIS collaboration, has directed its work towards two main directions: on the one side the study of the production and of the collisional cooling of Ps emitted from nanochanneled Si sample and on the other side its excitation in the Rydberg states. In this Ph.D thesis, developed with the Trento group, both the research lines are dealt with. In particular the Ps energy distribution after its emission in vacuum was investigated by preliminary measurements performed with a Time of Flight apparatus, while a new apparatus embedded in the AEgIS set up was designed and built for future measurements on Ps excitation in Rydberg states.

As regards the study on the Ps emission, it is fundamental for AEgIS experiment for the following reasons:

- the maximum possible Ps yield from the material has to be obtained for increasing the number of anti-hydrogen atoms, produced for each positrons bunch,
- the Ps cooling inside nanochannels has to be optimized in order to achieve the maximum yield in correspondence of Ps velocities for which the charge-exchange cross-section is maximum. In the AEgIS proposal [27] it was demonstrated that the cross-section depends on the relative

velocity  $v_{rel}$  between the  $\bar{p}$  and Ps and it was found that the maximum of the cross section is obtained if  $v_{rel} < v_{orb}$ , where  $v_{orb}$  is the orbital velocity of the positron (or electron) while moving on the Ps orbit considered at rest in the lab-frame. Thus if  $n \sim 30$ ,  $E_{CM} \cong 16$  meV, and the Ps cooling has to be optimized in order to achieve as much as possible Ps with a velocity lower than  $\sim 5 \cdot 10^4$  m/s,

- the dimension of Ps cloud has to be estimated, because it should fit the power of the laser used for exciting the atoms,
- the permanence time of Ps inside nanochannels before its emission and thus the Ps energy distribution as a function of the time elapsed since the positrons bunch implantation, has to be known for the laser synchronization in the excitations,
- the spatial distribution has to be investigated, in order to maximize the excited Ps cloud, which crosses the antiprotons trap,
- the Ps temperature has to be evaluated and it should be as low as possible, in order to minimize the transfer of energy to anti-hydrogen atoms.

In my thesis, preliminary measurements done by using the Trento continuous slow positron beam and then performed at the NEPOMUC facility at Munich will be presented. In particular the influence of the size of the nanochannels in the Si target on the final Ps energy distribution was investigated as a function of the positron implantation energy and the sample temperature, (also at cryogenic temperature). A fitting function for extracting an average Ps temperature will be shown and finally a possible method of data analysis for the estimation of the permanence time inside the nanochannels will be presented. In the future by using the same apparatus but with a multi-ring detector an estimation on the spatial distribution of Ps emission might be done.

On the other side studies on Ps excitations will be possible thanks to the new apparatus described in this thesis. The main requirements on the positron bunches for reproducing the AEgIS conditions and for the planned measurements were:

- a tunable positrons energy between 7-9 keV for the Ps formation in depth,
- a spot dimension lower than 3 mm diameter on the target surface, so that the emitted Ps cloud into vacuum is as small as possible,

- a bunch time duration of 5 ns at the target, to increase the density of the Ps cloud.

More it was necessary to built:

- a system of two parallel grids connected to a switching circuit, for producing an electric field of few kV/cm in the target region for Ps ionization after its excitation,
- a proper shield for the magnetic field due to the AEgIS transport line and to the 5 T AEgIS magnet, in order to reduce to a few gauss the magnetic field in the target region,
- two switching coils for generating a tunable magnetic field perpendicular to the target surface, for observing the magnetic quenching effect on Ps excited in  $n = 3$  level.

This apparatus will be useful to reproduce the AEgIS conditions during the Ps formation and excitation process in order to study and optimize the sample properties and the laser synchronization. At the same time the new setup is conceived to be exploited for Positronium spectroscopy and for new physics on Positronium. This simple atom system in fact is also a good test for many theories, mainly because:

- the Ps atom is very light compared to the lightest stable atom, hydrogen, and effects related to the atom velocity, such as Doppler and motional Stark effects, are unusually high compared with it,
- the Ps atom is composed by two elemental leptonic particles (an electron and a positron), bonded together by pure electromagnetic forces and thus the fine structure splitting of ground state can be investigated,
- the Ps atom spontaneously annihilates in vacuum (in 125 ps for para-positronium, 142 ns for ortho-positronium in absence of external fields), offering an useful observable, which can be measured, in comparison with the hydrogen atom that is stable.

## 1.5 Content description

In the following chapters these items will be discussed in details:

- **Positronium physics**

In this chapter an overview on the physics of the Positronium formation

and emission from different material will be given. The Ps emission from nanochanneled Si sample will be discussed in more detail, also giving a classical and a quantum description of the cooling process occurring inside the nanochannels.

- **Time of Flight measurements**

The Time of Flight (ToF) apparatus realized for measurements at NEPOMUC facility will be described. The preliminary results obtained using Trento slow positron beam will be shown together with the first measurements performed at Munich with the sample kept at 50 K. A possible method of data analysis will be explained in order to investigate the energy spectrum of emitted Ps and secondly to extract an estimation of the Ps permanence time inside the nanochannels. I personally took part to all measurements and I performed the data analysis.

- **Apparatus for Ps excitations in Rydberg states** The new apparatus connected to the AEgIS transfer line will be presented. The electro-optical and the electronic design necessary to spatio-temporally compress the positron bunches, coming from the Surko-trap accumulator through a magnetic transport line of about 60 cm, will be presented. I personally did the design and the simulations for the beam transport in magnetic and electric field and for the final focusing of positrons on the target. The characteristic of time-dependent potential applied to the electrodes of the buncher system for the time compression will be explained and discussed, taking into account the final requirement of obtaining bunch of 5 ns at the target with energy of 7-9 keV. The final electronic circuit will be described and the results of the first tests will be shown.

- **Ps excitations in Rydberg states and Ps spectroscopy**

The Positronium physics which could be investigated by the new apparatus will be discussed. A description of the Ps levels in presence of magnetic field will be dealt with in order to analyze the possibility of spectroscopy measurements. The experimental results achieved in the last year in the Positronium spectroscopy will be described and taken as reference for the proposed measurements. The set up necessary for observing the excited Ps will be presented and in the final part a discussion on the possibility of implementing the laser cooling technique on Ps will be shown.

## Chapter 2

# Positronium physics

### 2.1 Ps formation in solids from slow positron beam

Positronium (Ps), the bound state between a positron with an electron, was discovered for the first time by Martin Deutsch in 1951 [30]. It is a quasi-stable hydrogen-like atom [31] that can be formed with a probability of  $\frac{1}{4}$  as para-positronium p-Ps, (mean lifetime in vacuum of 125 ps) and with a probability of  $\frac{3}{4}$  as ortho-positronium orto-Ps, (mean lifetime in vacuum of 142 ns). Due to the different spins, the o-Ps with total spin 1, decays in three  $\gamma$ -rays, while only two  $\gamma$ -rays are emitted by the p-Ps with total spin 0. Positronium can not be found in nature, but abundant emission of Ps in the vacuum has been observed by bombarding surfaces of different solid materials with slow positron beam [32].

When positrons reach the surface of the sample can backscatter or enter into the material. The fraction of backscattered positrons from the surface depends on the atomic number  $Z$  of the material and on the positron implantation energy; for example in Silicon, it has been found to be about 5% constant up to 20-30 keV positron implantation energy [33]. On the contrary positrons which enter in a solid, lose their energy very quickly with respect to their mean lifetime in that material and thus they can often thermalize [34].

The implantation profile of positrons in a sample is described by the Mahkavian function, which depends on positrons initial kinetic energy  $E_+$  and on the material density  $\rho$  [35]. The mean stopping depth, in nanometer, is

$$\bar{z} = \frac{40}{\rho} E_+^{1.6} \quad (2.1)$$

where  $\rho$  is expressed in  $g/cm^3$  and  $E_+$  in keV and the constants 40 and 1.6 are material dependent. The Mahkopian profile can be written as:

$$P(z, E) = \frac{mz^{m-1}}{z_0^m} e^{-\left(\frac{z}{z_0}\right)^m} \quad (2.2)$$

where  $m$  is a dimensionless parameter, depending on the atomic number  $Z$  of the material and  $z_0 = \frac{\bar{z}}{\Gamma[1/m+1]}$  results to be  $z_0 \cong 1.13\bar{z}$  for  $m = 2$ . The implantation profiles for  $\rho = 1.9g/cm^3$  (the density of Silicon as in the samples described in section 2.2) and for different positron implantation energies of 3 keV (blue line), 5 keV (red line), 7 keV (green line) are reported in Fig.2.1. The figure shows that the positron implantation profile is broadened and it

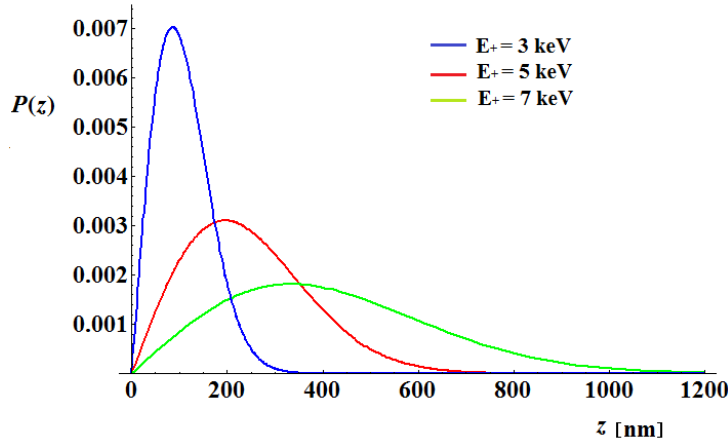


Figure 2.1: Implantation profiles as eq.2.2 calculated for positron implantation energy of 3,5,7 keV.

is distributed from the surface to a few hundreds of nanometers in depth. In particular the broadening increases with higher positron implantation energies. Therefore the positrons implanted nearer to the surface, might not thermalize (epithermal positrons), but they could be directly emitted in vacuum as  $e^+$  (if its work function is negative  $\phi_+ < 0$ ) or as Ps, while positrons implanted deeper thermalize up to the sample temperature. After the thermalization process, the positron essentially remains in a “free” or delocalized state and it starts to diffuse around. In this phase it can:

- be emitted directly in vacuum as  $e^+$  (if  $\phi_+ < 0$ ) or as Ps.
- be trapped in a new surface state inside the material, from which it can annihilate or be emitted in vacuum as  $e^+$  (if  $\phi_+ < 0$ ) or as Ps,

As regards the positronium emission from the materials, it involves the positron work function  $\phi_+$  and the electron work functions  $\phi_-$ , together with the Ps binding energy in vacuum  $E_B = 6.8$  eV, but the production mechanisms is different for metals or semiconductors or insulator materials.

### 2.1.1 Ps formation in metals

In the bulk of metals, incoming positrons cause a strong enhancement of the local free electron density, which prevents the Ps formation [36]. On the contrary near the surface, the region of decreasing electron density, the Ps formation can take place if the condition on the Ps formation potential  $\epsilon_{Ps}$

$$\epsilon_{Ps} = \phi_+ + \phi_- - E_B \leq 0 \quad (2.3)$$

is satisfied [34] or if the positron energy is above the  $\epsilon_{Ps} > 0$  value. In many cases the 6.8 eV of the Ps binding energy are enough to ensure that  $\epsilon_{Ps}$  is negative and adiabatic emission occurs; anyway Ps emission can occur in small quantity by epithermal or backscattered positrons.

Therefore Ps at the surface of a metal can be formed by different mechanisms: [37]:

- by backscattering [38] or epithermal positrons with energy of the order of a few tens of eV, which capture an electron from the conduction band,
- by direct formation by thermalized positrons which capture an electron,
- by thermalized positrons, which, after getting trapped in a surface state, are desorbed as Ps by thermal activation.

In the first case the positron energy is always much higher than the threshold of electron capture or the Ps formation potential  $\epsilon_{Ps}$  and the emission is favored [39].

In the second case, if the formation potential is negative (eq.2.3), the adiabatic Ps emission occurs at any temperature of the sample, often being in competition with the bare positron emission, which can occurs if also the positron work function is negative. In this case Positronium, formed by direct electron capture, has a continuum energy spectrum, where the upper limit correspond to  $|\epsilon_{Ps}|$  value [40].

Finally in the third case, thermal positrons have insufficient energy to escape from a solid, either because  $\phi_+ > 0$  or because their energy  $E$  is too low to satisfy the condition  $E - \phi_+ > 0$ . Thus positrons become trapped in a surface state with a binding energy of the order of  $E_b^+ \sim 1$  eV and an energy

$E_a = E_b^+ + \phi_- - E_B$ , however lower than  $E_b^+$ , is required for the Ps emission. By thermal activation [41] this energy can be provided and Ps can be emitted with an energy spectrum along the normal to the surface given by:

$$\frac{dN}{dE_\perp} \propto e^{-\frac{E_\perp}{k_B T}} \quad (2.4)$$

that is known as a beam maxwellian distribution [42, 43]. It was experimentally observed by using an Al sample at a temperature ranging from 100 K to 200 K [44].

### 2.1.2 Ps formation in semiconductors

Ps formation process in semiconductors appears to be similar to those taking place in metals. The main difference is that in such materials the work function, that here corresponds to the ionization potential, is higher than in the metals and the relation of eq.2.3 is often unsatisfied. Anyway the Ps formation induced by epithermal positrons can take place. In semiconductors as Si or Ge, Ps formation in the bulk has never been observed, but Ps formation at the surface occurs in both the materials [45]. In particular in Ge, where the Ps formation potential is negative, the Ps emission at high temperature in the low implantation energy limit was found to be near the 100 % of the implanted positrons. On the contrary in Si the yield is lower due to the positive Ps formation potential and only thermally activated emission or electron capture from epithermal emission can be held [34].

### 2.1.3 Ps formation in insulators

In insulator materials as molecular solids or ionic crystals, where the electron work function (i.e. the ionization potential) is very high, the Ps formation by thermal positrons returning to the surface is always energetically forbidden. Thus Ps can be formed in the bulk as electron-positron pair and it has to reach the surface to be emitted as Ps in vacuum. Two processes are possible: the *ore process*, induced by not thermalized positrons and the *spur process* induced by thermalized positrons.

In the *ore process* Ps is formed by the capture of a valence electron, induced by a positron, whose kinetic energy  $E$  is not enough to create an electron-hole pair, but it is higher than the thermal one [46]. Thus the minimum positron kinetic energy for Ps formation is:  $E_{min} = E_{gap} - E_b$  where  $E_b$  is the Ps binding energy in the solid (which is always lower than the corresponding value in vacuum  $E_B$ ) and  $E_{gap}$  is the gap energy between the valence and the conductive band in the insulator. The excess in energy

given by  $E - E_{min}$  is available as kinetic energy of the Ps center of mass, but, if  $E - E_{min} > E_b$ , the pair will dissociate in the collision. Thus a stable formation of Ps occurs only if the condition of eq.2.5 on positron kinetic energy is satisfied.

$$E_{gap} - E_b < E < E_b + E_{min} = E_{gap} \quad (2.5)$$

On the contrary in the *spur process* the electrons, created during the positrons slowing down, thermalize within a characteristic distance which is material-dependent. Thus the Ps formation occurs between the thermalized electrons attracted by the thermalized positrons. This process was studied in crystalline ice [46] and was found to become more and more relevant at higher positron implantation energy  $E_+$ , where the Ps yield increases. The Ps yield probability can be expressed as:

$$p(E_+) = p_{max} + (p_0 - p_{max})e^{-\frac{1}{2}\left(\frac{E_+}{E_1}\right)^\beta} \quad (2.6)$$

where  $p_0$ ,  $E_1$  and  $\beta$  are parameters, while  $p_{max}$  is the maximum Ps yield in the bulk [47].

In both processes after formation in bulk, Ps quickly thermalizes and diffuses inside the material until it eventually reaches the free surface and if it finds favorable conditions it is emitted into vacuum.

In particular among insulators, quartz and amorphous silica, whose ionization potential is of the order of 10 eV, are known as efficient Ps emitters. In such materials, in fact, bulk formation (Ore and spur processes) proper of insulator materials was observed together with a thermally desorption of Ps formed at the surface from capture of excited electrons [48].

Time of Flight measurements were performed on  $SiO_2$  at different positron implantation energy [49], where the relative Ps energy distribution was found to be characterized by two components: one is nearly monoenergetic at 3.27 eV, consistent with a bulk exciton-like Ps emitted from the surface, while a broader component of  $\sim 1.5$  eV full width at half maximum, is also observed due to the thermally activated process.

Finally a special situation occurs in porous insulators materials (as aerogels, porous silica films,..), where Ps can be also emitted in subnanosize or nanosize open volume defects [50]. Indeed, after emission, Ps can be hosted in the pores and, if a network of connected open volume defects is present, it can be emitted into the vacuum with a different energy distribution after a hopping diffusion. In general for a powder or a porous solid, many positrons are implanted near the internal surfaces of the pores or of the grain and here they do not reach the complete thermalization. Therefore the contribution of

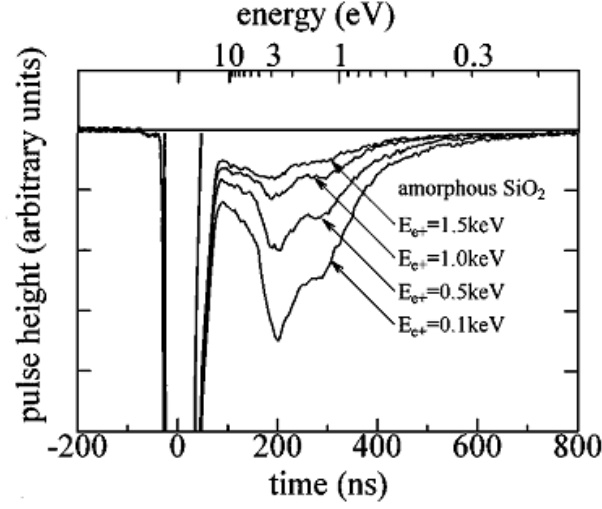


Figure 2.2: Time of Flight spectrum of Ps emitted from amorphous  $SiO_2$  for different positron implantation energies  $E_+$  (after [49]).

Ps emission given by epithermal formation is quite high and Ps is formed with high energy. Anyway Ps atoms remain trapped in an empty space where, by colliding with the walls of the trap, lose their energy until they annihilate or escape in the vacuum [37].

Despite p-Ps has too short lifetime, o-Ps lives enough to cool down or to thermalize and to reach the vacuum, even if a fraction of this undergoes to a *pick-off* annihilation process. In the hopping diffusion, in fact, the o-Ps, hitting the surface of the walls can annihilates via 2  $\gamma$ -rays with an electron of the material [51]. This fact causes a reduction of the o-Ps lifetime from 142 ns to 30-50 ns, depending on the material and on the pores dimensions, as observed by Positron Annihilation Spectroscopy measurements [52]. In the last years several porous silica-based materials with disordered porosities were investigated [53, 50] and a fraction of Ps was found to escape from the samples with a mean energy near the thermal one.

In this contest Ps emitted in vacuum with velocity of the order of  $10^4$  m/s, as required from AEgIS experiment (see section 1.4), could be achieved, in principle, by different strategies:

1. by increasing the Ps yield at low temperature by using a thermal desorption process, where Ps temperature seems to be closely related to the temperature of the target [54], as observed with Al sample (see section 2.1.1).

2. by applying innovative techniques (as laser cooling [55] or others) to cool down o-Ps after its emission in vacuum,
3. by producing samples with nanostructure connected to the vacuum, in such a way that Ps, formed in the pores, can cool down by colliding with the walls of the pores and thus reducing its energy.

Since the first possibility requires a sample heating for providing the thermal activation energy, this was avoided in the AEgIS experiment where the 100 mK antiprotons trap is placed near the sample. More due to the limitation on the final temperature achievable by the cooling techniques (see section 5.6) and also because of the difficulties in applying a further technique on Ps cloud before the Rydberg excitation, the first two choices were excluded in the AEgIS context. On the contrary the third possibility, which could allow Ps to be cooled just inside the sample, was taken into account.

In the past, in fact, many disordered porous materials were investigated [56, 57] and in particular several experiments, performed on different silica-based porous materials kept at room temperature, showed a high Ps yield, where a fraction of emitted o-Ps had energy close to the thermal one.

Therefore the possibility of realizing pores with different dimensions and shapes in silica-based films was investigated and finally the formation of more regular structure as nanochannels, directly connected to the surface, was realized by the Trento group, through an electrochemical process [58]. In the last years the electrochemical process was studied and optimized in order to obtain as high as possible Ps yield in vacuum, while the Ps cooling process inside the channels is nowadays under investigation.

This is the framework of this thesis, where nanochanneled silicon samples were considered and used for Time of Flight measurements on the emitted Ps.

In the next section an overview of the properties of the nanochanneled silicon samples is given and in the next chapter the preliminary Time of Flight measurements performed also by keeping the sample at low temperature are reported.

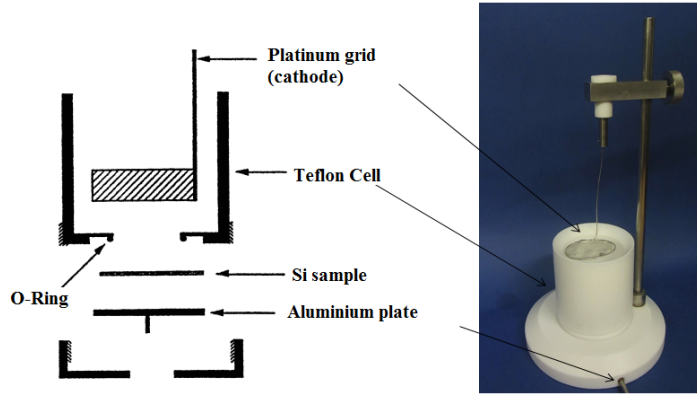


Figure 2.3: Left: cross-sectional scheme of a “single-tank” anodization cell (after [60]). Right: picture of the cell built and used to make nanochannels into the Si sample. In the picture the platinum grid, the teflon body and the connection to the aluminum plate are visible, while the inner part, where the sample is placed on, is only shown in the left scheme.

## 2.2 Nanochanneled silicon samples

### 2.2.1 Electrochemical etching

In Si samples, nanochannels perpendicular to the surface can be obtained by electrochemical etching in hydrofluoric acid HF solution [59]. The Si surface dissolution is obtained by a constant anodic current, acting on the sample immersed in an aqueous HF solution, ethanol added to increase the surface wetting. The anodization cell, where the etching process takes place, is shown in Fig.2.3. In this cell, whose body is made of teflon, a highly acid resistant polymer, the Si sample is placed on a metal disk at the bottom and sealed through an O-ring, so that only the front side of the sample is exposed to the electrolyte. In this way the Si sample acts as the anode and the cathode, made of platinum, a HF-resistant and conductive material, is also immersed in the aqueous HF solution. When a potential is applied, a measurable external current flows through the system, inducing a precise chemical reaction, the nature of which is fundamental to the porous formation. The process in fact is not homogeneous on the sample surface, but pore initiation occurs at surface defects or irregularities, with the consequence of pore formation, during the electrochemical process.

All the properties of pores, such as density, diameter and length, depend on anodization conditions as HF concentration, current density, Si type and resistivity and anodization duration [59]. The current value in particular de-

termines the pores size and density, while the anodization duration influences mainly their length. By optimizing all the above production parameters a precise control both on the nanochannels diameter and length was achieved [58], in order to increase as much as possible the surface area for high Ps yield. Si p-type (100) wafers with resistivity between 0.15 and 0.21  $\Omega\cdot\text{cm}$  were used and the solution was realized by adding absolute ethanol to a commercial aqueous solution at 48% of HF with volume ratio of 1:3=HF:ethanol. In particular nanochannels with smaller diameters in the 5-20 nm range and extending for  $\sim 2\ \mu\text{m}$  in depth were obtained with etching current in the 10-30  $\text{mA}/\text{cm}^2$  range and anodization duration in the 10-20 min range. The achieved pore density was quite high, because the distance between two close channels was comparable with their diameter.

However after the nanochannels formation, the surface of the walls is highly chemically reactive and an air contamination of the inner channel surface could strongly reduce the Ps yield. Thus after the electrochemical etching, the samples were cleaned in absolute ethanol 99.8% and then oxidized in air at temperature of 100  $^\circ\text{C}$ , as surface passivation process. In this phase a thin (2 nm) layer of silica  $\text{SiO}_2$  is formed on the nanochannels walls, which is essential for Ps formation and emission. This oxidation procedure makes the sample and the o-Ps formation and emission stable for many weeks, if the samples are maintained in vacuum. A reduction in pore dimensions after annealing at temperature higher than 100  $^\circ\text{C}$ , or their complete filling can occur with the consequent reduction in o-Ps fraction out diffusing from the channels.

A SEM picture of the surface of a nanochanneled Si sample after the manufacturing process is shown in Fig.2.4, where the regular pore distribution on the sample is clearly visible.

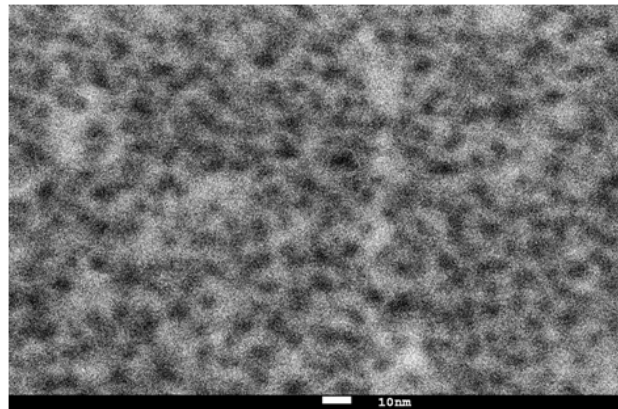


Figure 2.4: A SEM picture of the surface of a nanochanneled Si sample

## 2.2.2 Ps yield

In the nanochanneled Si samples the Ps yield in vacuum principally depends on the pores size and on the ratio between the volumes of the inner silicon pillars and of the surface silica layer [58].

The high Ps amount is in fact due to a cascade process: thermalized positrons in Si bulk, having diffusion length of 200 nm [34], higher than the distance between each nanochannels (few nanometers), diffuse to the Si/ $SiO_2$  interface and then pass into the silica  $SiO_2$ , because this is energetically favored [61]. Here positrons can efficiently form Ps (see section 2.1.3) which, after a very short diffusion path, reaches the inner surface of the nanochannels, where it is emitted.

An evaluation of the Ps yield was done in [58], where the annihilation gamma rays were detected by two high-purity germanium (HpGe) detectors. The data were collected as function of the gamma energies and were divided into two regions:

- the 511 keV peak area (P), in the region defined as  $|511 \text{ keV} - E_\gamma| \leq 4.25 \text{ keV}$ , attributable prevalently to  $2\gamma$  annihilations,
- the valley area (V)  $410 \text{ keV} \leq E \leq 500 \text{ keV}$ , due to o-Ps  $3\gamma$  annihilations.

Then the  $2\gamma/3\gamma$  ratio  $R(E_+) = P(E_+)/V(E_+)$  was calculated as the ratio between the valley area and the peak area at each positron implantation energy  $E_+$ .

By calibrating the  $R_1$  parameter as the value of  $R(E_+)$  for 100% Ps formation as measured in Ge held at 1000 K (see section 2.1.2) and  $R_0$  as the value of  $R(E_+)$  for 0% Ps formation at the highest positron implantation energy (implantation in Ge bulk), the  $F_{3\gamma}$  function of the calibrated fraction of implanted positron annihilating as o-Ps, can be obtained.

$$F_{3\gamma}(E_+) = \frac{3}{4} \left( 1 + \frac{P_1[R_1 - R(E_+)]}{P_0[R(E_+) - R_0]} \right)^{-1} \quad (2.7)$$

where  $P_0$  and  $P_1$  are the values of the 511 keV peak area measured at 0% and 100% o-Ps formation, respectively [45, 62]. In the Fig.2.5  $F_{3\gamma}(E_+)$  as a function of the positron implantation energy for samples with different pores sizes is shown [58]. The sizes of nanochannels were estimated to be around 4-7 nm in sample #0 and 8-12 nm, 8-14 nm, 10-16 nm, 14-20 nm, and 80-120 nm in samples #1, #2, #3, #4 and #5, respectively. Due to the relation between the positron kinetic energy and the average implantation depth (see eq.2.1 and 2.2 in section 2.1), the Ps emission process in the samples can be

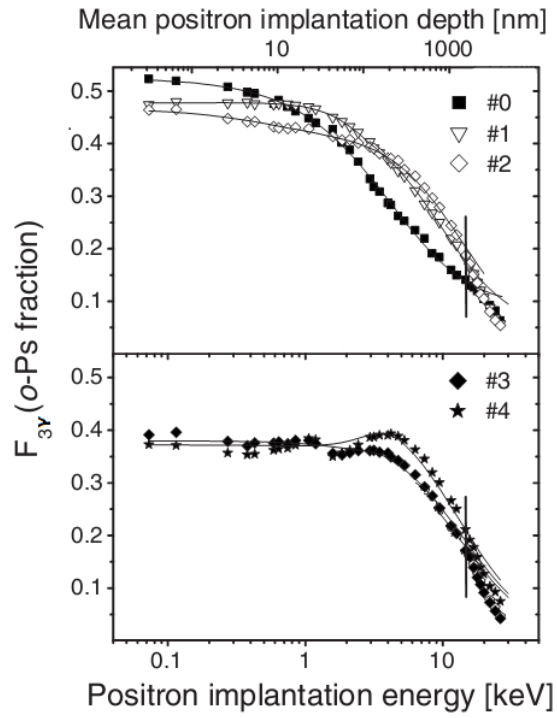


Figure 2.5: o-Ps fraction  $F_{3\gamma}(E_+)$  as a function of the positron implantation energy. Data were corrected for the fraction of fast o-Ps atoms outdiffusing from nanochannels. The continuous lines through the points are best fits obtained by eq.2.8 derived from a diffusion model. The vertical lines mark the border between the region with nanochannels and the silicon bulk (after [58]).

studied as a function of different depths. The data were corrected by taking into account that at low positron implantation energy, a larger fraction of Ps escapes into vacuum with high kinetic energy, and since it annihilates on the chamber wall it is detected with low probability, due to the smaller solid angle. In Fig.2.5 the  $3\gamma$  annihilations were observed to be constant for lower positron implantation energies and to decrease at higher energies for the  $2\gamma$  pick-off annihilations enhancement. The o-Ps yield in vacuum is, in fact, strictly related to the pick-off process, which depends on the number of collisions inside the pores and so on the pores dimensions. As shown in Fig.2.5 the pick-off influence is more relevant for smaller channel size, where Ps yield decreases just starting from 1 keV energy.

In order to extract the real fraction of Ps which escape into vacuum, a deep analysis based on a diffusion model is required. At each positron implantation energy, the o-Ps fraction, observed in the peak area by  $3\gamma$  annihilations, is in fact composed of two contributions: the first one given by the o-Ps atoms that annihilate inside the channels and the second one given by o-Ps atoms that escape into the vacuum. Thus the relative  $F_{3\gamma}(E_+)$  curve can be fitted by the equation:

$$F_{3\gamma}(E_+) = A(E_+) \cdot B(E_+) + C(E_+) \quad (2.8)$$

which takes into account the fraction  $A(E_+) \cdot B(E_+)$  of formed o-Ps that annihilates in vacuum and  $C(E_+)$  the fraction that annihilates into pores. In particular:

$$A(E_+) = p_{max} + (p_0 - p_{max})e^{-\frac{1}{2}\left(\frac{E_+}{E_1}\right)^\beta} \quad (2.9)$$

is the phenomenological profile for o-Ps formation in a spur process, (see section 2.1.3) while:

$$B(E_+) = \frac{1}{1 + \left(\frac{E_+}{E_0}\right)^{1.6}} \quad (2.10)$$

is the probability of o-Ps out-diffusion from the nanochannels with  $E_0 = \sqrt[1.6]{L_{Ps}\rho/40}$  [34, 46].

Finally the  $C(E_+)$  term is expected to increase for increasing positron implantation energy up to reach a constant value at higher positron implantation energies. For this reason it was assumed proportional to the pick-off annihilations observed by lifetime measurements:

$$C(E_+) = K \left[ \frac{C_1 - C_2}{1 + \left(\frac{E_+}{E_2}\right)} + C_2 \right] \quad (2.11)$$

where  $K$  is the constant of proportionality and  $C_1$ ,  $C_2$  and  $E_2$  are the characteristic parameters for the pick-off process [58].

In Fig.2.5 the best fits of the equation 2.8 are shown for all measured samples. It can be observed that the agreement between the model and the data is good. From the model it was found that the fraction of Ps reaching the vacuum is very high for lower energy (at 1 keV is up to 40%), and it decreases down to 25% (for samples #2, #3, #4, and #5 with bigger nanochannels) or 17% (for sample #1 with smaller nanochannels diameter) at 10 keV positron implantation energy (corresponding to about 800 nm depth) due to pick-off annihilations.

As regards the Ps energy distribution, being formed from silica, Ps is expected to have an initial energy inside nanochannel characterized by two distributions centered on 1 and 3 eV (see Fig.2.2 in section 2.1.3). Anyway the Ps velocity distribution in vacuum after the collisional cooling in nanochannels was found to be considerably different [63] and the dependence of the final distribution on the channel properties is nowadays under investigations.

In the next section a theoretical description of the cooling process inside nanochannels is done.

### 2.2.3 Ps cooling

The cooling process due to o-Ps collisions on the nanochannels walls was described by mainly two models in the classical or quantum regime, respectively, depending on the o-Ps energy  $E$ . If the de Broglie wavelength of Ps atom is small in comparison to the diameter  $d$  of the confinement cavity

$$d \gg \lambda_{dB} = \frac{h}{\sqrt{2m_{Ps}E_{Ps}}} \sim 0.9 \text{ nm}(1 \text{ eV}/E_{Ps})^{1/2} \quad (2.12)$$

Ps can be considered as a free object with velocity  $v_{Ps}$ , bouncing against the wall of the cavity. On the contrary if  $d \sim \lambda_{dB}$  a classical description will not be suitable anymore, because the Ps diffusion will be dominated by quantum effects.

In the next section both the models will be presented, in order to achieve the o-Ps energy evolution as a function of the time spent inside nanochannels.

#### Classical description

In a classical model, elastic collisions between o-Ps with mass  $m_{Ps} = 2m_e$  and free atoms with mass  $M$ , placed at the surface of the wall are considered. In this case the mean kinetic energy lost in each collision, calculated by

averaging on the scattering angles and on the nuclei velocities [57] can be written as:

$$\langle \Delta E \rangle = -\frac{2m_{Ps}M}{(m_{Ps} + M)^2} \left( \frac{m_{Ps}v_{Ps}^2}{2} - \frac{M \langle V^2 \rangle}{2} \right) \quad (2.13)$$

where  $\langle V^2 \rangle$  is the mean square velocity of the sample atoms.

Thus the Ps energy as a function of the time  $t$  elapsed from Ps emission can be derived by:

$$\frac{dE}{dt} = \langle \Delta E \rangle f \quad (2.14)$$

where  $f = \frac{v_{Ps}}{\lambda}$  is the collision frequency against the wall, if  $\lambda$  is the mean free path of Ps inside the pores.

By assuming that the nuclei have thermal energy, with mean value

$\frac{M \langle V^2 \rangle}{2} = \frac{3}{2}k_B T$  where  $T$  is the sample temperature and by replacing  $v_{Ps} = \sqrt{\frac{2E}{m_{Ps}}}$ , the eq.2.14 becomes:

$$\frac{dE}{dt} = \frac{2m_{Ps}M}{(m_{Ps} + M)^2} \sqrt{\frac{2E}{m_{Ps}}} \frac{1}{\lambda} \left( E(t) - \frac{3}{2}k_B T \right) \quad (2.15)$$

whose solution is

$$E(t) = \left( \frac{1 + Ae^{-bt}}{1 - Ae^{-bt}} \right)^2 \frac{3}{2}k_B T \quad (2.16)$$

where  $A = \frac{\sqrt{E_0} - \sqrt{\frac{3}{2}k_B T}}{\sqrt{E_0} + \sqrt{\frac{3}{2}k_B T}}$  with  $E_0 = E(t = 0)$ , the Ps initial energy and  $b = \frac{2}{M\lambda} \sqrt{3m_{Ps}k_B T}$ .

This equation predicts three stages of Ps thermalization: at the beginning Ps gradually loses its energy in few picoseconds, then it starts to cool down more quickly and reaches a quasi-thermalized state with kinetic energy a few times longer than the thermal energy. Finally Ps reaches the thermalization very slowly [57]. The plot of the  $E(t)$  function, calculated for  $M = 25$  a.m.u. and for  $T=1$  K, is shown in Fig.2.6.

Anyhow if a thermalization up to very cold temperature is investigated, a quantum description is required because the condition expressed in eq.2.12 becomes early unsatisfied.

## Quantum description

A quantum model, which considers the Ps interaction with the acoustic phonons of the material has been developed [64]. The longitudinal acoustic phonons interaction was treated in terms of the acoustic deformation

potential which was written in a perturbative way as:

$$W = E_d \sum_q \left[ \sqrt{\frac{\hbar}{2NM\omega_q}} i q \left( a_q e^{i\vec{q} \cdot \vec{r}} - a_q^\dagger e^{-i\vec{q} \cdot \vec{r}} \right) \right] \quad (2.17)$$

where  $N$  is the number of atoms in the sample with mass  $M$ ,  $\vec{q}$  the phonon momentum and  $\omega_q$  its angular frequency, while  $\vec{r}$  is the Ps vector position and  $a$ ,  $a^\dagger$  the destruction and creation operators, respectively.  $E_d$ , the strength of the deformation potential, was estimated to be 3.6 eV [65].

In this context the state of Ps, placed in  $\vec{r}$  position, with momentum  $\vec{k}$  and energy  $E_{Ps}$  can be written in the quantum formalism by the ket:  $\left| Ps \left( \vec{r} \right)_k \right\rangle$  and the initial state of Ps and the phonons can be expressed by:

$$|k\rangle = \left| Ps \left( \vec{r} \right)_k \right\rangle \prod_q |n_q\rangle \quad (2.18)$$

where  $|n_q\rangle$  is the wave function of the  $n_q$  phonon with momentum  $\vec{q}$ .

Thus the transition probability of Ps and phonons from the initial state to the final one, where Ps has momentum with magnitude  $k^f$  and energy  $E_{Ps}^f$  and phonons have momentum  $q^f$ , is given by:

$$P_{kk^f} = \left| -\frac{i}{\hbar} \int_0^t \sum_q \langle k^f | W | k \rangle e^{\frac{i(E_{Ps}^f - E_{Ps}^i + \sum_q \hbar \omega_q (n_q^f - n_q))t}{\hbar}} dt \right|^2 \quad (2.19)$$

Taking into account the energy and the momentum conservation in the interaction Ps-phonons, expressed by the condition  $\frac{\hbar^2(k^f)^2 - \hbar^2 k^2}{2m_{Ps}} = \hbar \omega_q$  a maximum on the variation of the momentum magnitude is found:

$$\Delta k_{max} = \pm \frac{2m_{Ps}v}{\hbar} \quad (2.20)$$

where the dispersion relation between  $\omega_q$  and  $q$  was assumed to be linear  $\omega_q = v_s q$ . In particular the eq.2.20 carries to important consequences in the case of Ps confined in a cubic box. In such case some transitions between two close energy levels could be not allowed if the maximum change in momentum magnitude corresponds to a change in energy smaller than the levels separation. In fact if  $d$  is the side of the cubic box, and the wave function of Ps confined in the box can be written as:

$$\left| Ps(\vec{r})_k \right\rangle = \sqrt{\frac{2}{d}} \sin(k_x x) \sqrt{\frac{2}{d}} \sin(k_y y) \sqrt{\frac{2}{d}} \sin(k_z z) \quad (2.21)$$

with the energy levels given by  $E = \frac{\hbar^2(n_x^2+n_y^2+n_z^2)}{4m_{Ps}d^2}$ , ( $n_x, n_y$  and  $n_z$  are the natural quantum numbers) the minimum distance between two energy levels is  $\Delta E = \frac{\hbar^2}{4m_{Ps}d^2}$  which increases by decreasing the side  $d$ . Thus the condition (2.20) on the momentum magnitude implies a condition on a maximum energy variation  $\Delta E_k$ , which could be smaller than  $\frac{\hbar^2}{4m_{Ps}d^2}$  if the side  $d$  is too small. That is the case in which the ground state cannot be reached by phonons-Ps interaction [64] and the Ps remains in an higher energy state until its annihilation.

On the contrary in case of collisions in channels this fact does not happen, because the non quantized dimension parallel to the channel axis allows that no maximum is present on the transition. Thus the condition  $\Delta E_k \geq \frac{\hbar^2}{4m_{Ps}d^2}$  comes to be always satisfied and Ps can reach the ground state by phonons interactions. In this case the energy on the parallel direction can be totally canceled and the Ps emission from nanochannels occurs with wider angles with respect to the normal to the sample.

In order to derive an expression of the Ps energy as a function of the time spent in the pores, the averaged energy variation in each collision was obtained in the one-dimensional case:

$$\langle \Delta E(k) \rangle = \frac{d}{\pi} \int_0^\infty \frac{\hbar^2}{2m_{Ps}} (k^2 - k^{f2}) P_{kkf} dk^f \quad (2.22)$$

where the transition probability  $P_{kkf}$  was calculated from eq.2.19.

From the equation:

$$\frac{dE}{dt} = \left\langle \frac{\Delta E(k)}{\tau} \right\rangle \quad (2.23)$$

analog to the eq.2.14 of the classical description, where  $\tau = d/v_{Ps}$  is the time spent to cross the channel, an expression of  $\langle \Delta E(k) \rangle$  as a function of the time spent by Ps in the interaction region  $R$  was found. This equation is analog to eq.2.15, but in the quantum context.

In Fig.2.6 the eq.2.16 as obtained in the classical context and the solution of eq.2.23 in the quantum context for Ps in a one-dimensional well are plotted, in order to compare them, without taking into account their range of applicability, which depends on the pore size. In particular the figure shows that Ps thermalization takes more time in the quantum context in comparison with the classical process.

Finally it is worth noting that due to the quantization of the energy levels, the ground state in the quantum description does not coincide with the thermalized state of the classical description. It means that the ground state in the quantum context can be related to a Ps minimum temperature different from the sample one. In the case of Ps in nanochannels, the quantization

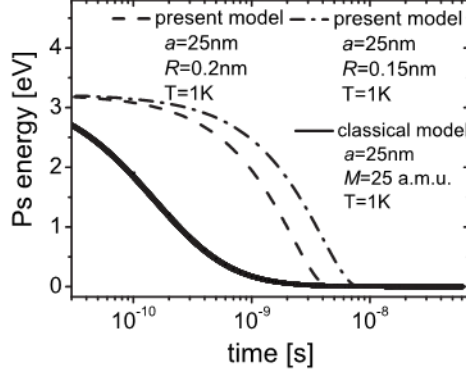


Figure 2.6: Ps energy as a function of the time in an infinite one-dimensional potential well with size  $d = 25$  nm in a matrix of  $\text{SiO}_2$  at the temperature of 1 K. Dash-dotted and dashed curves correspond to quantum model calculations with interaction region  $R=0.15$  nm and  $R=0.20$  nm, respectively. The solid curve is calculated by the classical model [57], assuming an effective mass  $M = 25$  a.m.u (after [63]).

on the energy levels associated with two of the three space-dimensions forces the minimum energy to be  $E = \frac{h^2}{4m_{Ps}d^2}$ , which sometime can be bigger than the thermal one  $\frac{3}{2}k_B T$ . Thus the Ps minimum temperature can be found as a function of the channels diameter  $d$  by:

$$T_{min} = \frac{h^2}{6m_{Ps}d^2k_B} \quad (2.24)$$

For instance in nanochannels with side of 5 nm, o-Ps cannot be cooled at temperature lower than 116 K, but increasing the size of the nanochannels the minimum accessible temperature decreases and, in nanostructures of 20 nm, it becomes around 7 K [64] (see Fig.2.7, in which  $T_{min}$  is plotted as a function of the nanochannel dimension  $d$ ).

Many experimental evidences have confirmed the Ps minimum temperature is determined by the pores size. For instance in mesoporous thin films [66] Ps emitted from a sample with small pores size was found at higher minimal energy (73 meV) in comparison with Ps going out from a sample with bigger pores (43 meV). In porous materials with 2.7 nm pores size, a minimum o-Ps kinetic energy of 42 meV was measured via Doppler broadening of the  $1^3S - 2^3P$  transition [67].

In conclusion during the thermalization process Ps starts in a classical regime and loses its energy by elastic collisions with the material molecules till it reaches an energy (called  $E_{trans}$ ) for which the de Broglie wavelength

is comparable with the nanochannels diameter (see eq.2.12). From that point Ps loses its energy by interacting with phonons and the minimum temperature given by the sample temperature or by the quantum confinement temperature  $T_{min}$  can be reached. In the plot of Fig.2.7 the  $E_{trans}$

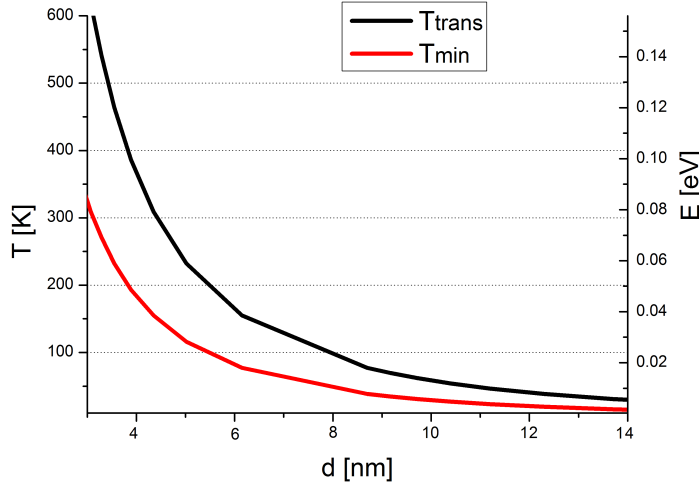


Figure 2.7: Red curve: minimum cooling temperature  $T_{min}$  of a Ps atom as a function of the nanochannel dimension  $d$  (eq.2.24).

Black curve: energy  $E_{trans}$  (or temperature  $T_{trans}$ ), at which the Ps de Broglie wavelength is comparable with  $d$ , as a function of the nanochannel dimension  $d$  (from eq.2.12). For both the plots the energy and the correspondent temperature scales, calculated by  $T = \frac{2E}{3k_B}$ , are shown.

values as a function of the nanochannel dimensions  $d$  are reported and in the same plot both the energy and the correspondent temperature scales, calculated by  $T = \frac{2E}{3k_B}$ , are shown. Thus at a fixed  $d$ , depending on the sample temperature  $T_{sample}$ , the cooling process could be only a classical process, if  $T_{trans} < T_{sample}$ , or initially classical and then quantum one, if  $T_{sample} < T_{trans}$ . More by comparing the sample temperature and the quantum minimum temperature  $T_{min}$ , it could be known which of them (the higher one) will be reached by Ps atoms. This plot was taken into account in the choice of the nanochannel diameter for observing Ps cooling up to cryogenic temperature (see section 3.7.2).

In conclusion silicon samples with ordered array of oxidized nanochannels can be suitable for the AEgIS experiment.

Firstly in fact it was shown that this material can be used as efficient positron-o-Ps converter, where implanted positrons give rise to high yield of Ps into

the vacuum, thanks to the  $\text{SiO}_2$  layer present on the nanochannels walls. Secondly, differently from not porous materials, a cooling process takes place inside nanochannels, with a consequent strong modification of the final Ps energy distribution. More, the possibility of tuning the nanochannels size allows o-Ps to be emitted in vacuum at cryogenic temperature.

Anyway, because a minimum size is imposed by the desired temperature, a too large diameter could increase too much the thermalization time and thus an accurate optimization has to be done for the final application in the AEGIS experiment.

In the last years first Time of Flight measurements on this kind of samples were done and Ps cooling and thermalization at cryogenic temperature was observed [63]. Further measurements with the same apparatus are reported in the chapter 3 of this thesis. The scope of these measurements was the study of the Ps velocity distribution in samples kept at lower temperature and with larger pores size.



## Chapter 3

# Ps Time of Flight measurements

The Time of Flight (ToF) apparatus was designed to measure the elapsed time between the Ps formation inside the nanochannels and its annihilation in flight at a given distance from the sample. The measurement allows one to extract the energy distribution of Ps emitted into vacuum. The apparatus was conceived to work with a continuous positron beam: it was preliminary tested with the Trento laboratory positron beam of  $7 \cdot 10^3$  positrons/s and then it was assembled at the neutron-induced positron source NEPOMUC at FRM-II reactor in Munich, where the positron beam intensity is  $\sim 2 \cdot 10^6$  positrons/s.

In the first part of this chapter an overview of the NEPOMUC beam and of the experimental apparatus is presented. Then the preliminary data acquired at Trento and at Munich together with a possible method of data analysis are presented and finally some considerations on the signal to noise ratio and on its influence on the final spectra are reported.

### 3.1 The high intensity positron beam NEPOMUC

The NEutron-induced POsitrone source at MÜniCh (NEPOMUC), located at the research reactor Heinz Maier Leibnitz FRM II, is the world's strongest source of mono-energetic positrons [68], from which a high intensity positron beam is extracted and then magnetically guided to five beam ports for physics experiments.

The positron source is mounted as an in-pile component of the beamtube SR 11, close to reactor core inside the moderator tank of the reactor. The cross-

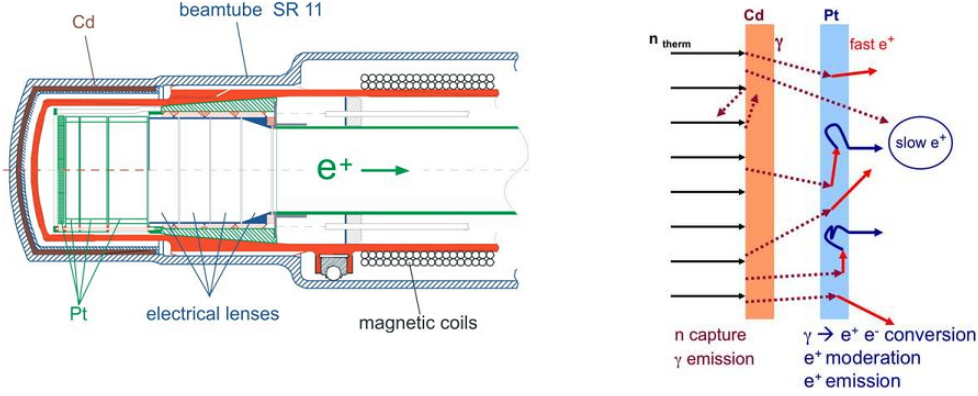


Figure 3.1: Left panel: cross-sectional view of the positron source NEPOMUC. After [70]. Right panel: Principle of the NEPOMUC source.

sectional view of the tip of the beam tube is shown in Fig.3.1 [69]. In the reactor core, neutrons, after thermalization in  $D_2O$  of the moderator tank, are absorbed in  $^{113}Cd$ , by the nuclear reaction  $^{113}Cd(n, \gamma)^{114}Cd$ , from which high energy gamma rays are emitted. The  $^{113}Cd$ , whose natural abundance is of 12.2%, has been chosen, due to its high cross section of 26000 barn for thermal neutron capture. The cadmium cap with 110 mm diameter and 95 mm length is installed near the maximum of the thermal neutron flux, inside the beamtube SR11 (see Fig.3.1). The de-excitation energy of 9.05 MeV is released by  $^{114}Cd$  as a  $\gamma$ -cascade, where on average 2.3  $\gamma$ , with energies higher than 1.5 MeV per captured thermal neutron, are emitted [71].

A structure of platinum foils is used for the conversion of the high-energy  $\gamma$ -radiation into electron-positron pairs, whose energy spectrum has a maximum positron energy of about 800 keV. At the same time the platinum, having a negative work function, acts as moderator, from which thermalized positrons are emitted with an energy spread of 2 eV. The principle of the NEPOMUC source is shown in the right panel of Fig.3.1.

The moderated positrons are extracted and accelerated by electrostatical lenses, which reduce the diameter of the positron beam and release it into a region with a magnetic guiding field of  $\sim 60$  G [70]. The maximum beam intensity for a low-energy positron beam of 15 eV amounts to  $\sim 4 \cdot 10^6$  moderated positrons per second distributed inside a diameter of 20 mm [69]. By the magnetic transport line the beam is guided through three bendings in the concrete wall of the reactor pool in order to screen the accessible part of NEPOMUC from fast neutrons,  $\gamma$ -radiation, fast electrons and fast positrons. A further remoderation stage with a tungsten single crystal with a total effi-

ciency of 5% [70] in reflection geometry has been installed. It decreases the beam diameter to  $\sim 2$  mm (FWHM), and the energy spread is reduced to the thermal energy spread of the remoderating tungsten crystal. The energy of the remoderated beam can be adjusted between 20 and 200 eV and the beam diameter is less than 2 mm (FWHM) in the 60 G guiding magnetic field [72].

Through the magnetic line positrons are transported for many meters towards one of the five experimental setups: three of these are in routine operation for positron annihilation induced Auger spectroscopy [73], for Doppler broadening studies [74] and for positron lifetime measurements [75]. In addition an interface for the connection of the scanning positron microscope has been installed [76] and finally an open multi-purpose beam port is available, where various experimental setups can be connected to the positron beam line for short-term experiments [77]. Our Time of Flight apparatus was connected to this free beam port for Ps velocity measurements.

## 3.2 The experimental apparatus

### 3.2.1 The transport line

At the open beam port of the Munich facility the positron beam extracted from the reactor arrives with a diameter of about 8 mm after the transport in a guiding magnetic field of 60 G. In such an axial magnetic field, positrons undergo helical motion with a pitch and diameter determined by the field strength and by the positrons momentum. In particular if  $v_T$  and  $v_L$  are the transversal and the longitudinal component, respectively of the positron velocity with respect to the magnetic field axis, which corresponds to the symmetry axis, the gyration radius  $r_g$  is given by:

$$r_g = \frac{mv_T}{qB} \quad (3.1)$$

while the pitch or gyration length  $h = v_L T$ , where  $T$  is the revolution period, can be written as:

$$h = 2\pi \frac{mv_L}{qB} \quad (3.2)$$

However the transport line of our apparatus was designed in order to deviate the magnetic field and to focus positrons on the target only by electric field generated by electrostatic lenses. Despite an electrostatic lens system appears more complex than a magnetic line for beam transport, it offers the means for beam manipulation, focusing as well as the possibility to control the final energy and the angular divergence.

In this context a restriction in the beam parameters during the transport is

imposed by the Liouville's theorem, which states that the volume occupied in phase space of the beam will be a constant under influence of the conservative forces. Thus in an electrostatic line a parameter, which takes into account the beam properties as the beam intensity  $I$  [particle/s], its diameter  $d$ , the particle kinetic energy  $E$  and the angle  $\Theta$  between the positron trajectory and the longitudinal axis, was introduced. It was called *brightness*  $B$ , and it was defined as:

$$B = \frac{I}{\Theta^2 d^2 E} \quad (3.3)$$

Due to the Liouville's theorem, the brilliance is constant inside an electrostatic line and for this reason the final beam properties at the target are strictly limited by the initial beam brilliance, occurring after moderator or after magnetic field extraction, as in our case. More it is worth taking into account that along the transport, given that  $I$  is constant, each modification on the positrons energy or on the beam diameter, and so on, is reflected in an opposite modification on the other beam parameters [78].

Since essentially 100 % transmission through an electrostatic system is possible if optimum design is adopted, the most critical point in our transport line becomes the positrons extraction from the magnetic field into electrostatic lenses. Because positrons in magnetic field follow the magnetic flux lines as long as their gyration length is small compared to the length in which the magnetic field heavily changes, two possible solutions can be applied in order to make positrons released from a magnetic field without changing the transversal momentum significantly [79].

Firstly, for a fast beam with some keV energy, the magnetic field can be changed abruptly, such that the gyration length  $h$  is much larger than the region where the field changes.

Secondly, if the positrons are initially slow and no losses are acceptable, it is possible to change smoothly the magnetic field, and to compensate this by accelerating positrons along the longitudinal axis. The latter principle has been used for ToF apparatus, since the re-moderated beam of the facility has a limited energy range from 20 to 200 eV and intensity losses are not acceptable.

With this goal at the entrance of the ToF apparatus a  $\mu$ -metal electrode with a proper shape was placed in order to terminate the magnetic line, and to ensure that no magnetic guiding field penetrates into the sample chamber. The sketch of the  $\mu$ -metal terminator is shown in Fig. 3.2. It consists of an aperture, which is shaped as a cylindrical electrode with diameter slowly changing from 10 mm diameter on the top and 30 mm on the bottom with a thickness of 14 mm. The cone shape is generally used to avoid saturation effects due to the increasing magnetic flux for larger radial distances. It is ly-

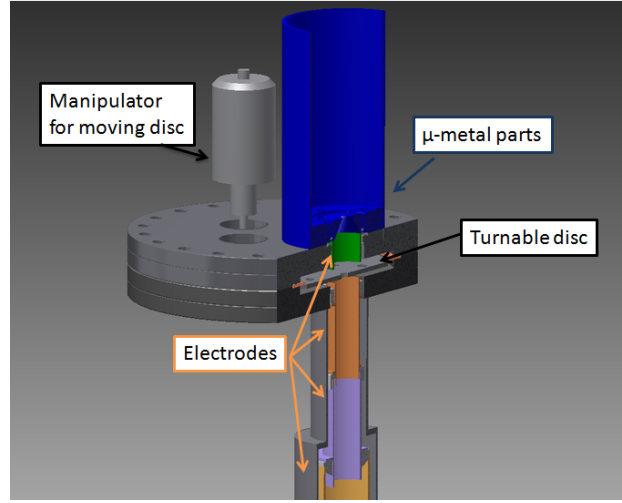


Figure 3.2: Drawing of the top of the ToF apparatus. The  $\mu$ -metal electrode and the  $\mu$ -metal shield are colored in blue, the turnable disc with the manipulator and the first electrodes of the lens system are indicated. The coils before the  $\mu$ -metal electrode are not shown here.

ing on a ceramic ring for insulation purposes and is separated from the outer part of the  $\mu$ -metal shield with as small as possible gap of 1 mm. This gap allows voltages of 1 kV maximum to be applied to the electrode and at the same time it allows magnetic field lines to be carried through the electrode towards the outer shield. This last is in fact formed by a  $\mu$ -metal flange and a  $\mu$ -metal tube (see Fig.3.2, in touch with each other, in order to guide the magnetic flux lines on backward with respect to the ToF chamber.

The multipurpose FEM-program COMSOL Multiphysics [80] has been used to simulate the magnetic flux density distribution near the  $\mu$ -metal electrode, if a guiding field of roughly 6mT (60 G) is applied. The results are shown in Fig. 3.3. Due to the cylindrical symmetry of the geometry, only the half space with positive distance  $r$  from the center has been taken into account in order to save computing time. The last coil structure together with the field terminator are sketched in the same figure with black lines. According to this simulation, the flux density drops from 6 mT to 0.8 mT within a distance of 10 mm from the top of the field terminator.

In the same area an electric field is generated by applying an electric potential to the  $\mu$ -metal electrode (1 kV maximum) in contact with the adjacent cylindrical steel electrode (in green in Fig.3.2). This allows positrons to be accelerated along the axial direction so that the positron gyration length of around 10 cm in that zone results an order of magnitude larger than

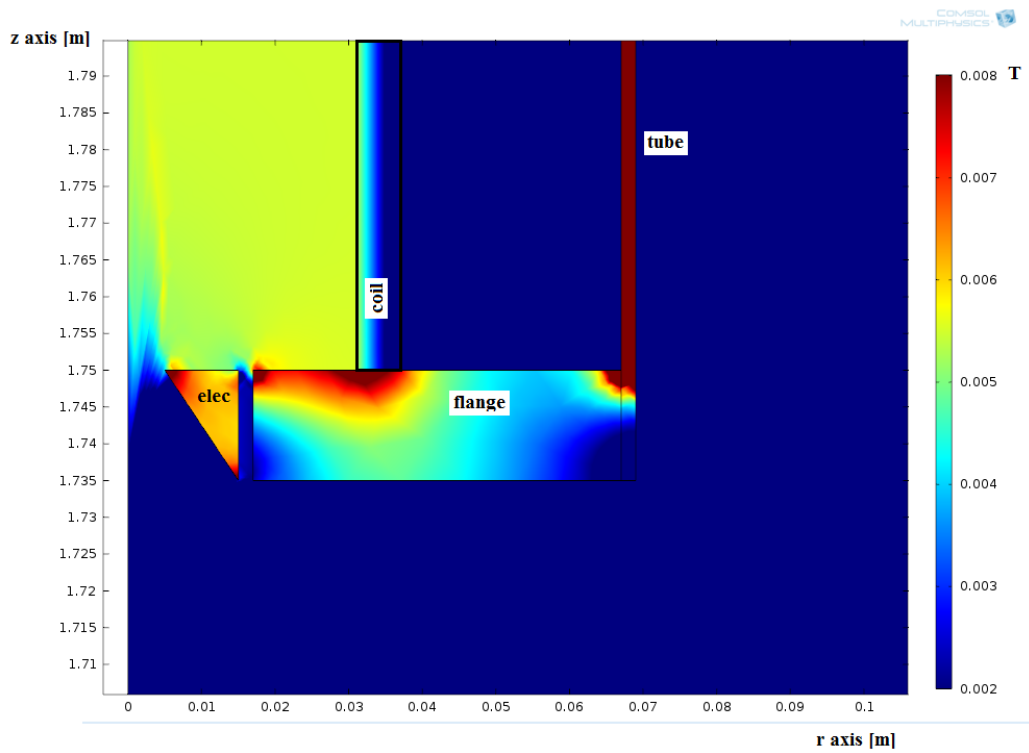


Figure 3.3: COMSOL-simulation of the magnetic flux density distribution on the top of the ToF apparatus. The magnetic field is generated by the last coil of the facility magnetic transport line and it is mostly deviated by the  $\mu$ -metal parts (electrode, flange, tube). Near the  $\mu$ -metal electrode in fact the flux density drops from 6 mT to 0.8 mT within 10 mm from the  $\mu$ -metal aperture, which allows a non-adiabatic release of the positron beam coming from above along the z-axis.

the length scale of the breakdown of the magnetic flux density, satisfying the condition for a non-adiabatic release of the positrons from the magnetic field. The overall trajectories are shown in Fig.3.5 in section 3.2.3, where the magnetic release is visible near the first electrode.

### 3.2.2 Beam intensity selection by turnable disc

Between the first electrode and the other lenses a beam intensity selector was placed, in order to choose the desired beam intensity as required by the signal to background optimization (see section 3.8.1). The selector is a turnable disc ( $\varnothing = 185$  mm) with the different holes of 0.5, 1, 2, 3, 4, 6, 8, 10 mm diameter (see Fig. 3.2) mounted about 5 mm below the steel electrode in contact with the magnetic field terminator. The variable hole is manipulated via a rotary feed through from outside. By turning the disc, a desired part of the positron beam can be shielded and the beam diameter at the entrance of the focusing system is adjusted. Since it is placed inside an electric field region, given by the magnetic field terminator and the other lenses system, the turnable disc is in touch with the neighboring second electrode and thus kept at the same voltage, in order to ensure that the geometry of the aperture does not distort the passing trajectories of the positrons. Because the reactor beam diameter was estimated to be around 8 mm in diameter, the beam intensity can be adjusted from 2 to 100 % of the original intensity.

Anyway in the event that a hole with a small diameter is used to reduce the beam size, strong background radiation can be generated at the entrance of the lens system and it is therefore favorable to put the sample far enough in order to keep the background in the detectors as small as possible. For this reason the sample was placed at a distance of  $\sim 140$  mm from the disc and the beam is transported over the target by a system of electrostatic lenses, as explained in the next section.

### 3.2.3 The electrostatic lenses

Five insulated steel tubes as well as the  $\mu$ -metal electrode shown in Fig.3.4 were necessary to guide positrons into the chamber and to focus them on the sample. In general in an electrostatic line a sets of cylindrical tubes in series along the symmetry axis focus positrons in intermediate points along the path, allowing them to be transported up to the target [81, 82]. In our case the first three accelerating potentials together with the last two slowing down ones work as a unique big lens, which focuses positrons on the target. Since their defocusing/focusing effect, the beam diameter increases in the first part (where the electrodes diameter increase too) while it is reduced in

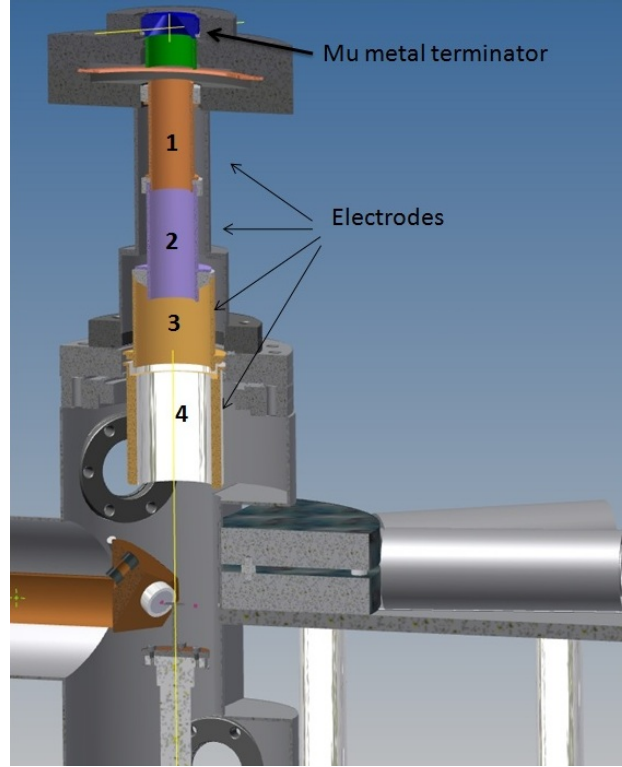


Figure 3.4: Sketch of the sample chamber of the ToF apparatus. The magnetic field terminator, the variable aperture, the lens system, the sample holder and the scintillators and channeltrons detector (one of these) are shown.

the final region up to the target surface. The final spot dimension is also determined by the additional focusing effect due to the sample set on negative potentials down to -10 kV.

In order to find the best electrodes potentials, the software package SIMION 8.1 [83] has been used to simulate the positrons trajectories. In each simulation 100 positrons trajectories can be calculated and the absolute beam diameter along the whole transport line has been determined for different values of the target potential. In these simulations the magnetic guiding field together with the electric field were taken into account. The result for -10 kV sample potential is shown in Fig.3.5, where a 5 mm positron beam, starting with 200 eV axial kinetic energy and with 0.1 eV transversal energy and entering the magnetic field terminator from the left side was simulated. In this simulation the variable aperture of the turnable disc was set to 10 mm, so that the full incoming positron beam was transported up to the target. Inside the transport line the positrons gyrate through the magnetic guiding

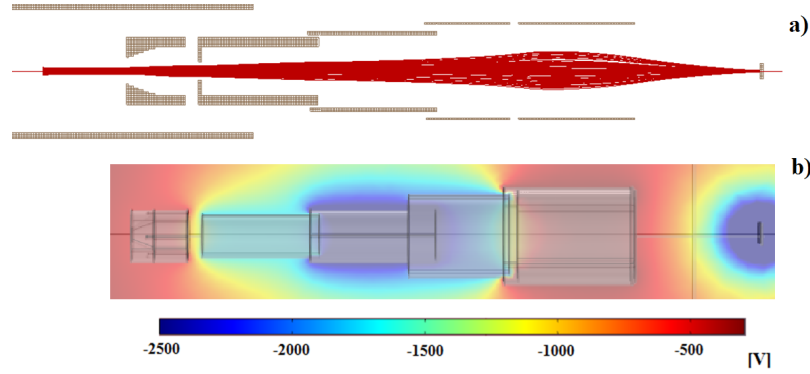


Figure 3.5: a) Resulting ray-tracing simulation obtained by SIMION 8.1 for -10 kV sample potential and for 100 positron starting from a circle of 5 mm diameter with 200 eV axial kinetic energy and with 0.1 eV transversal energy. The electrode potential were set to -500 V, -1700 V, -2700 V, -2000 V, -300 V, -5000 V, respectively. b) Electric potential plot for electrode set at -500 V, -1700 V, -2700 V, -2000 V, -300 V, -5000 V, respectively obtained by COMSOL simulation.

field at the beginning, increase their gyration radius due to the acceleration toward the field terminator, get non-adiabatically released from the magnetic field near the  $\mu$ -metal electrode, then fly electrostatically guided through the lens system and are finally focused on the sample, held at the voltage, where a focal diameter of 3 mm is achieved. As results of these simulations, the electrodes 1, 2 and 3 have to be set to the constant voltages of -500 V, -1.7 kV and -2.7 kV respectively, for all the sample potentials, while lens 4 has to be adjusted dependently on the sample potential, since its value is critical for the final focusing. The optimized values for the fourth electrode ( $V_4$ ), which allow to achieve a constant spot of 2 mm diameter on the target, are shown in Tab.3.1. The resulting electric potential inside the line, as simulated by COMSOL by setting the electrode  $V_4$  at a voltage of -300 V, is shown in Fig.3.5.

### 3.2.4 The sample holder

The sample holder has to meet two main requirements. It has to insulate the target, set to a maximum voltage of -10 kV for adjusting the positrons energy and it has to be connected to a cryostat to cool sample down to cryogenic temperatures, in order to investigate Ps cooling process as a function of the sample temperature.

A ceramic ( $Al_2O_3$ ) disc, on which the sample is placed, was chosen, because

| Target (kV) | V4 (V) |
|-------------|--------|
| -1          | -200   |
| -2          | -250   |
| -3          | -250   |
| -4          | -300   |
| -6          | -400   |
| -8          | -400   |
| -10         | -500   |

Table 3.1: Table of values of the V4 electrode potential for different target potentials. With these configurations a final spot on the target of 2 mm diameter was achieved.

it allows the electric insulation between the target and the cryostat, but it guarantees the heat conductivity with the cryostat. The ceramic disc was covered on the top surface by a silver paste, in order to embed the target on the disc, achieving an as homogeneous as possible electric field. Like the sample chamber, the sample holder is fully rotationally symmetric in order to obtain a symmetric focusing: the disc, placed on the focal point of the lens system connects the cryostat by a cylindrical support. This last is about 77 mm high, has a diameter of 10 mm and it is manufactured from copper to ensure a good thermal coupling to the cryostat Coolpower 2/10 (Leybold). The cooling of the sample down to desired temperature is performed by the cryostat working at its minimum temperature (10 K) and by a heater connected to the sample holder. In this way the sample temperature can be tuned to the desired value as monitored by a diode attached to the sample holder itself. The accuracy on the measured temperature is  $\pm 1$  K for  $T < 100$  K and  $\pm 1$  % for  $T > 100$  K.

### 3.3 The acquisition chain

The Ps ToF technique needs a start signal, when the Ps is formed in the sample, and a stop signal, when it decays in flight at a fixed distance  $z_0$  from the sample. Ps formation can be well approximated by the arrival time of positrons into the sample, which is acquired by two channeltrons, detecting secondary electrons emitted from the sample surface after positron implantation. The two channeltrons utilized have 10 mm in diameter and are put around 15 mm far from the sample, so that a solid angle of 0.34 rad each one and a total angular acceptance of  $\alpha_c = 5.5\%$  are achieved. Ps annihilation rays are detected with 5 NaI(Tl) scintillators,  $d_s = 30$  mm in diameter, cou-

pled to photo-multiplier tubes (PMT), positioned at the desired distance  $z_0$  from the sample surface with a precision of 1 mm. They are arranged in an arc at a radial distance of  $L + D = 130$  mm from the sample axis and each of them is shielded by lead tubes and placed behind a  $D = 100$  mm long lead shield with a  $s = 3$  mm slit (see Fig.3.6). The slit dimension is properly chosen in order to obtain a sufficient resolution on the annihilation position while keeping a high total angular acceptance  $\alpha_s = 0.21\% = \frac{5s d_s}{4\pi(L+D)^2}$  to have a high scintillators counting rate.

From the slit dimension, the error on the annihilation position can be geometrically calculated as:

$$\Delta z = \frac{s}{D} \left( L + \frac{D}{2} \right) \quad (3.4)$$

and it results to be  $\Delta z = 2.4$  mm for a slit of 3 mm.

In order to obtain the start-stop delay, the channeltrons signal is amplified by an octal fast linear amplifier FL80000 (NES) with gain 12 and both the signals from channeltrons and NaI(Tl) detectors are transformed in squared pulses by an octal constant fraction discriminator CF8000 (EG&G-ESN). As the amplitude and the slope of the detectors signal are strictly correlated, the lower threshold value of the CF8000 affects the time resolution of the measurements: with a low threshold the time resolution worsens, while with a high threshold too many events are lost. For the NaI(Tl) detectors, the lower threshold is set by observing the 356 keV annihilation peak of  $^{133}\text{Ba}$  and increasing the threshold value till the signal disappears, so that false counts coming from Compton scattering of 511 keV-rays are excluded. On the contrary for the channeltron a lower threshold was chosen in order to have as many as possible counts.

Two logic fan-in/fan-out LF4000 modules (EG&G-ESN) are fed by the 5 scintillators and the 2 channeltrons signals, respectively and as outputs, two or-signals are obtained. The time delay between these two signals is measured by using an ultra fast Multiscaler P7887 (FastComTEC) with 250 ps resolution. When multiscaler receives a start signal, it waits for the arrival of a stop signal during a window of  $T = 1 \mu\text{s}$  and the start-stop time delay is acquired.

Because in our system (see section 3.4, 3.5) the counting rate of channeltrons signal is usually bigger than the counting rate of scintillators, in order to maximize the acquired counts per unit of time, the scintillators signal is used as start and the channeltrons signal is delayed by about 700 ns and used as stop signal. The delay is obtained by an Octal Gate Generator GG8000 (EG&G-ESN).

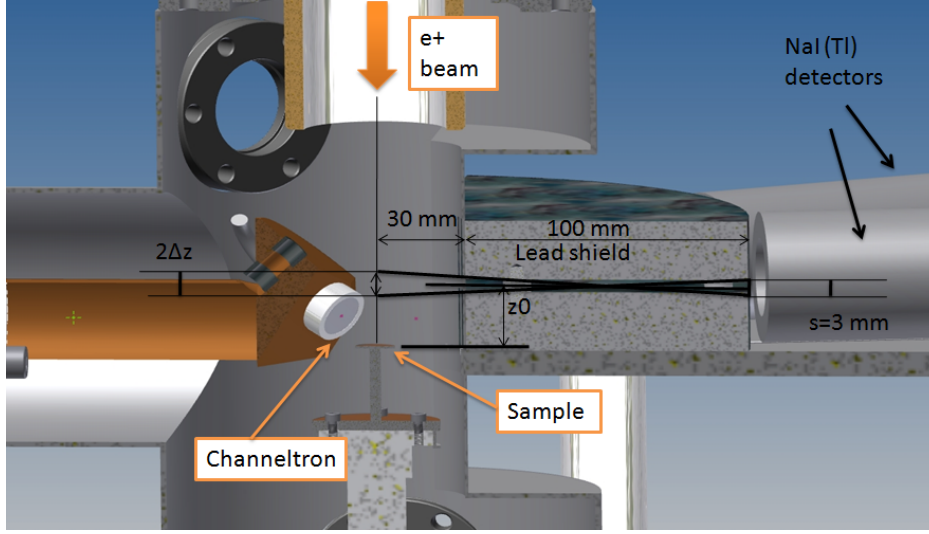


Figure 3.6: Cross-sectional view of the ToF chamber near the sample: one of the two channeltrons and two of the five scintillators are visible. The lead shield dimension, the slit and the error on the annihilation positions are shown.

### 3.4 The channeltrons signal

Channeltrons are durable and efficient detectors of positive and negative ions as well as electrons and photons. The devices are formed by three main parts (see Fig. 3.7): a cone, manufactured from a lead silicate glass, where the incoming particles, striking the inner surface, produce 2-3 secondary electrons; a multiplier channel, set at high positive potential with respect to the cone, where each secondary electron, hitting the inner surface produces other electrons in a cascade process and the final collector, where the output, given by a pulse of  $10^7 - 10^8$  electrons, emerges. Usually to detect negative particles the cone is set to ground or a low positive potential, while the tail of the multiplier channel is powered by a very high positive potential (2200-2600 V), in order to force the secondary electrons to travel down the channel until they reach the collector. After a capacitor, which cuts the DC voltage of the tail polarization, the final signal can be acquired. The circuit scheme for the signal acquisition is shown in Fig. 3.7, where the capacitor together with a resistor, which connects the collector and the tail to the same high positive potential, are visible. In order to increase the collection efficiency a grid can be placed in front of the cone and it can be set to the same potential, to avoid that outer electric field could cause the escape of the secondary electrons from the channeltron cone. The typical signal of our channeltron

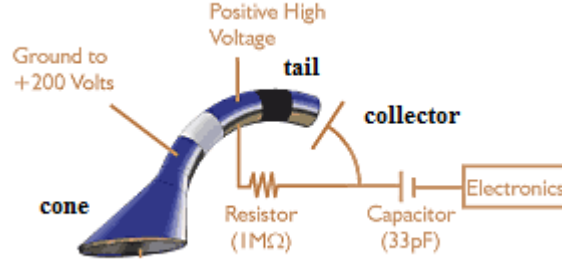


Figure 3.7: Channeltron structure given by the cone, the tail and the collector. The circuit scheme for detecting negative particles is also shown.

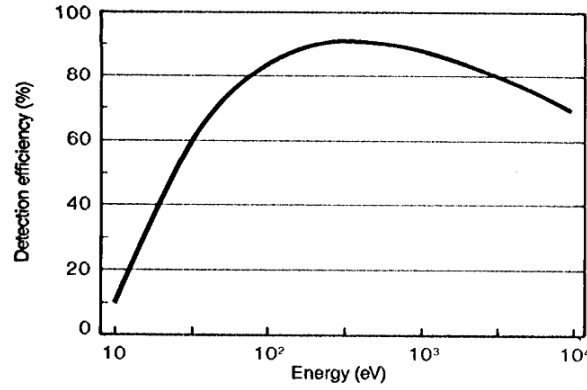


Figure 3.8: Detection efficiency for electron as a function of the electron incident energy.

(DETEC 102N-D) has a duration of 150 ns and an amplitude of 50-100 mV and it does not strongly depend on the energy of the incident electron.

The detection efficiency, defined as the probability that a charged particle (or photon) incident on the channeltron cone will produce an output pulse, is strongly dependent upon the energy, mass and velocity of the incident particle as well as the charges of the particle and the angle of incidence. As shown in Fig.3.8 it is near 90 % for electrons with energy ranging from 100 to 1000 eV.

Moreover the dark counts, defined as signal outputs, when no real event occurred, are very low in this device and they can be often neglected.

In our apparatus the two channeltrons overlook the sample, in order to detect the secondary electrons emitted after the positrons hit on the sample surface. Due to the energy transfer from positrons to the material molecules, these last can be excited and if the energy is sufficient to overcome surface

energy barriers (e.g. work function or electron affinity) electrons can escape from the material. In general if electrons hit the surface, the emitted electrons can be due to: i) Secondary electrons with lower energy ( $<50$  eV by convention) that originate within the material, produced by numerous inelastic scattering events of the incident particles; ii) or to backscattered electrons with higher energy ( $>50$  eV by convention) that originate from the incident electron source, but scatter either elastically or inelastically before leaving the target material.

In our material, since positrons are implanted, only secondary electrons can be observed by channeltrons, while backscattered positrons annihilate on the chamber wall.

Thus the  $p$  parameter, which includes only the secondary electrons, is defined as the ratio of the number of electrons emitted from the sample to the total incoming positrons and it is found to be around  $p=0.1-0.2$  for Si sample with incident electrons of 7 keV energy [84].

In conclusion the rate of the channeltrons signal, strictly dependent on the secondary electron emission, can be written as:

$$f_c = fp\epsilon_c + b_c \quad (3.5)$$

where  $f$  is the positron beam intensity,  $\epsilon_c$  takes into account the solid angle of all channeltrons and their intrinsic efficiency and  $b_c$  is the dark counts rate, when it is not negligible.

In our measurements the channeltron signal rate was found to vary from 30 to 400 cps depending on the channeltrons efficiency and on the beam intensity, as it will be discussed in the section 3.8.1.

### 3.5 The NaI(Tl) scintillators signal

Five NaI(Tl) scintillators are used to observe the  $\gamma$ -rays coming from the Ps annihilation. The photons interacting with the scintillator crystal produce a pulse of light that is converted to an electric pulse by a photomultiplier tube. The PMT consists of a photocathode, a focusing electrode, and 10 or more dynodes that multiply the number of electrons striking at each dynode. The NaI(Tl) crystal has a very high light output, among scintillators and it exhibits no significant self-absorption of the scintillation light, producing one of the highest signals in a PMT per amount of radiation absorbed (an average of  $1 \cdot 10^4$  photoelectrons per MeV  $\gamma$ -rays under optimum condition). It has the highest luminescence efficiency in comparison with the other scintillators, because it has a light emission range in coincidence with the maximum efficiency region of PMT with alkali photocathodes. The 5 NaI(Tl) scintillators

placed in the ToF apparatus are single crystal (by Scionix) of  $d_s = 30$  mm diameter coupled to Hamamatsu H105080 photo-multiplier tubes PMT. The typical signal has a decay time of 230 ns and the counting rate is mainly due to:

- the fraction  $y$  of positrons, which forms Ps and annihilates at the distance  $z_0$ ,
- the fraction  $pt$  of positrons, which annihilates on the chamber wall or on the sample surface. These events, called prompt peak events, (see section 3.6) are due to  $\gamma$ -rays that have crossed the lead shield and have been detected by the scintillators.
- the background counting rate due to dark counts or background radiation coming from the reactor  $b_s$ .

Thus, taking into account the three contributions, the total scintillator counting rate  $f_s$  can be written as:

$$f_s = f\epsilon_s(y + pt) + b_s \quad (3.6)$$

where  $\epsilon_s$  takes into account the solid angle of all scintillators and their intrinsic efficiency.

In particular the background counting rate was measured by switching off the incident beam, while the prompt peak events are acquired (as described in section 3.6) by using a Si single-crystal as sample, where Ps is not formed, when positrons are implanted in the bulk.

The scintillator counting rate is strictly dependent on the distance  $z_0$ : in fact  $pt$  decreases with the increase of  $z_0$  due to the different solid angle with respect to the target, but also  $y$ , that is the product of Ps total yield from the sample and the probability of Ps annihilation at distance  $z_0 \pm \Delta z$  decreases at higher  $z_0$ , due to the Ps finite lifetime.

In our set up configuration with scintillator placed at  $z_0 = 1$  cm and with  $s = 3$  mm, a counting rate of  $0.4 - 70$  cps was observed with a beam intensity ranging from  $7 \cdot 10^3 - 2 \cdot 10^6$   $e^+$ /s and the parameters  $\epsilon_s pt = 0.0069$  and  $\epsilon_s y = 0.0039$  were obtained.

Since the channeltrons counting rate is bigger than the scintillators counting rate, the scintillator signal was used as start (as explained in section 3.3), while the channeltrons signal was delayed by a fixed value  $T_D$  and it was used as stop. In this way the real Ps time of flight  $t_f$  is related to the acquired time  $t_{measured}$  by the relation  $t_{measured} = T_D - t_f$ .

### 3.6 ToF and Energy spectra

Once the times  $t_{measured}$  between scintillators and delayed channeltrons signals, had been acquired and reported in a histogram, some operations were done in order to obtain the *ToF Spectrum* describing the time annihilations distribution at a given distance from the sample and the correspondent *Energy Spectrum*, from which the Ps velocity distribution can be extracted.

Firstly the background gamma rays not coming from Ps has to be subtracted from the rough data. To do this a background measurement has to be acquired for each positron implantation energy  $E_+$  and for each distance  $z_0$  by using a Si single-crystal as sample, in which Ps is not formed when positrons are implanted in the bulk. In these measurements in fact background gamma rays are mainly due to prompt positron annihilations into the sample, or to positrons annihilating on the chamber wall after backscattering, or to background rays coming from the reactor or from the selector disc. In particular the time of flight distribution of these annihilations depends on the beam properties and detector placement and shielding and it is characterized by a prompt peak, placed at very low time (hundreds of ps after positron implantation). The position of this peak is related to  $T_D$ , while its width shows the resolution of the measurement.

A background measurement with  $E_+ = 7$  keV and at  $z_0 = 1$  cm is shown in Fig. 3.9. The prompt peak was fitted by a Gaussian function (continuous line in Fig. 3.9), from which it results to be centered on  $t_{prompt} = 714$  ns and with a FWHM of 8 ns. Thus  $t_{prompt}$  corresponds to the delay time set on the channeltron signal  $t_{prompt} \cong T_D$ , while the FWHM value depends on the back-scattered positrons and on the acquisition chain and it represents the final resolution of the Time of Flight measurement. For this reason rebin or mobile average procedures will be done on the rough data by considering bins inside the temporal window correspondent to the prompt peak FWHM, while the knowledge of  $t_{prompt}$  allows the effective time of flight  $t_f$  to be calculated for each bin by the relation:

$$t_f = t_{prompt} - t_{measured} \quad (3.7)$$

In the Fig. 3.10 the rough data acquired with nanochanneled sample in comparison with the data acquired with Si-single crystal for background measurement, reported in the  $t_f$  Time of Flight scale are shown. For the comparison the data were renormalized to the same measurement duration. In this plot the emission of Ps from nanochannelled sample is well evidenced by the increase of counts above 50 ns. From 50 ns the two distributions departs from each other, while at times of flight larger than 180 ns they overlap again near zero.

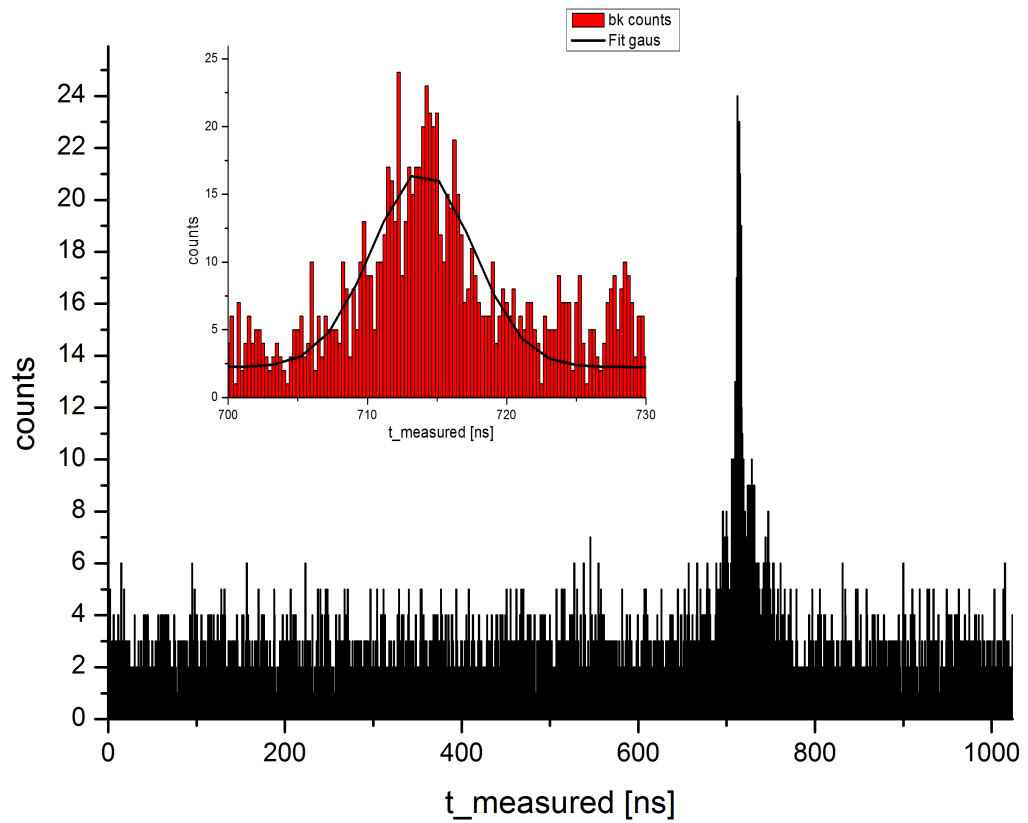


Figure 3.9: The background data acquired by using a Si single-crystal as sample. In the inset a zoom of the spectrum near the peak is shown. The prompt peak was fitted by a gaussian and the FWHM of 8 ns was found.

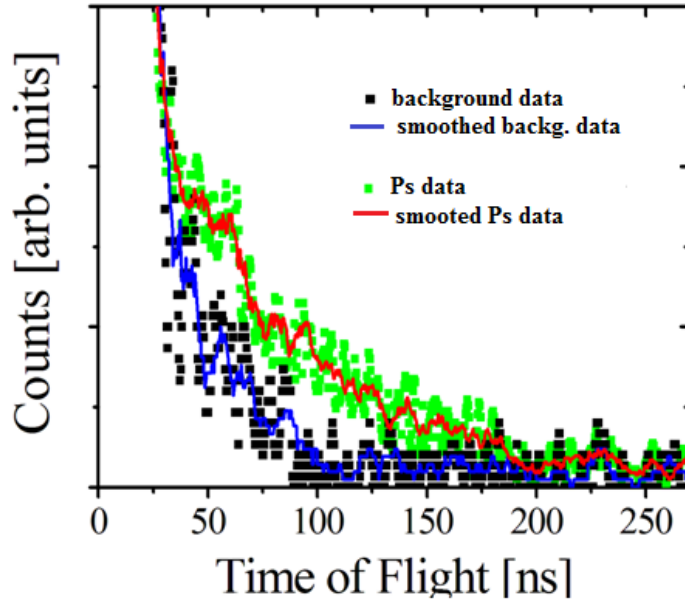


Figure 3.10: The rough data acquired with nanochanneled sample (Ps data in green), in comparison with the data acquired with Si-single crystal for background measurement (background data in black). The data are renormalized to the same measurement duration and reported in the time of flight  $t_f$  scale. Due to the high statistical fluctuation, a moving average on 21 channels (0.405 ns/channel) was done. The smoothed data are shown in the plot as continuous lines (red for Ps data and blue for background data).

Therefore in the interval where the Ps signal is well separated from the background signal (for example from 50 ns to 180 ns in the plot of Fig. 3.10), the background data can be subtracted and the *Rough Spectrum* can be obtained by reporting the subtraction results in the time of flight  $t_f$  scale. In fact too few data (compatibly with the noise) are present for higher times of flight, while at lower times of flight the over imposition of the tail of the prompt peak on the fast emitted Ps, prevents the possibility to distinguish the two components.

From the *Rough Spectrum*, in order to take into account the o-Ps finite lifetime  $\tau_{Ps}$ , which heavily affects the acquired time of flight distribution and by considering that slow Ps spends more time in front of the slit of the detectors with respect to fast Ps, each point centered on the time  $t_f$  was multiplied by the factor:

$$\frac{e^{\frac{t_f}{\tau_{Ps}}}}{t_f} \quad (3.8)$$

In this way the *Time of Flight Spectrum* (see for example Fig. 3.14 or 3.16b), in which the Ps time of flight distribution is reported, was obtained.

Finally the *Energy Spectrum*, not depending on the distance  $z_0$  and more useful for the Ps energy distribution analysis, can be easily achieved. The Ps energy along the axis perpendicular to the sample surface can be obtained as

$$E_{\perp} = \frac{1}{2} m_{Ps} \frac{z_0^2}{t_f^2} \quad (3.9)$$

and if  $N(t)$  is the time of flight distribution, by the relation:

$$N(t_f) dt_f = N(E_{\perp}) dE_{\perp} \quad (3.10)$$

and calculating:

$$\frac{dE_{\perp}}{dt_f} = -\frac{m_{Ps} z_0^2}{t_f^3} \quad (3.11)$$

the energy distribution  $N(E_{\perp})$  can be obtained:

$$N(E_{\perp}) = \frac{N(t_f) t_f^3}{m_{Ps} z_0^2} \quad (3.12)$$

As examples in Fig. 3.12, or 3.17 some *Energy Spectra* are shown.

Usually in these spectra the data are presented in a semi-logarithmical scale, where the  $\ln[N(E_{\perp})]$  is reported as a function of the energy. In this way the exponential distribution:

$$\frac{dN}{dE_{\perp}} = A e^{-\frac{E_{\perp}}{k_B T}} \quad (3.13)$$

becomes a straight line and it can be easily enlightened and fitted. This distribution, called *Beam Maxwellian distribution* (BM) [42], is related to the desorbed process of thermal Ps (see 2.1) as it was evidenced in some previous experiments. For example the presence of a single exponential energy distribution was found in the case of Ps emitted from Al surface [43] while the sum of two exponential distributions was used for fitting previous *Energy Spectra* as obtained by ToF measurements [63] performed on Ps emitted from nanochanneled sample.

In this context it is worth investigating how the Maxwell-Boltzmann velocity distribution of a possible Ps thermalized component appears in this semi-logarithmical scale. Because the Time of Flight apparatus detects only one component of the Ps velocity, the one-dimensional Maxwell-Boltzmann energy distribution has to be considered. If the energy probability distribution given by the Maxwell-Boltzmann theory is defined as:

$$f_E dE = 2\sqrt{\frac{E}{\pi}} \left(\frac{1}{k_B T}\right)^{3/2} e^{-\frac{E}{k_B T}} dE \quad (3.14)$$

where  $E$  is the sum of the squares of the three normally distributed momentum components, the one dimensional distribution can be obtained by the equipartition theorem. Assuming in fact that the energy  $E$  is equally distributed among all three degrees of freedom, the energy per degree of freedom is distributed as [85]:

$$f_{E_\perp} dE_\perp = \sqrt{\frac{1}{E_\perp \pi k_B T}} e^{-\frac{E_\perp}{k_B T}} dE_\perp \quad (3.15)$$

which represents the *one-dimensional Maxwell-Boltzmann Energy distribution* (M1D). The plot of eq.3.15 in comparison with the Beam Maxwellian distribution of eq.3.13 with the same temperature at  $T=300$  K is shown in Fig. 3.11 in the semi-logarithmical scale. The difference between the two functions appears in the different slopes and in the behavior at lower energy and it increases at higher temperature  $T$ .

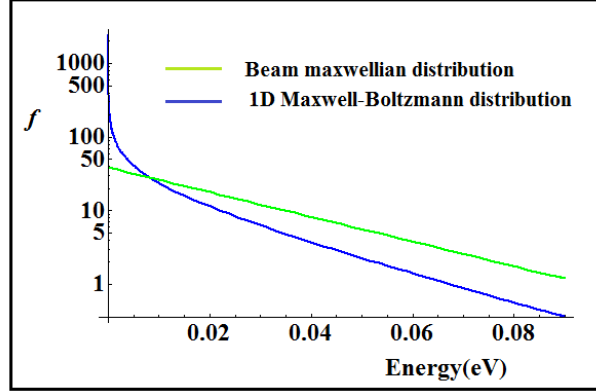


Figure 3.11: Plot of one dimensional Maxwell-Boltzmann energy distribution of eq.3.15 (in blue) and of beam maxwellian distribution of eq.3.13 (in green). Both the functions were calculated for  $T=300$  K.

### 3.7 Preliminary results

As described in section 2.2.3, the Ps cooling process inside nanochannels depends on:

- the number of collision  $N_c$  related to the positron implantation energy  $E_+$  and to the nanochannel dimension  $d$ ,
- the energy lost for each collision, which is related to the sample temperature  $T_{sample}$ .

The ToF data could show how the final Ps energy distribution depends on these three parameters ( $E_+$ ,  $d$  and  $T_{sample}$ ), evidencing the best condition to obtain as much as possible cooled Ps at a given temperature, necessary for anti-hydrogen production in the AEgIS experiment. To reach this final goal some preliminary measurements were firstly performed with the Trento positron beam and later at the NEPOMUC facility.

In the two next sections the preliminary *ToF* and *Energy spectrum* will be presented and described and a possible method of data analysis will be shown in order to extract the average Ps temperature and to estimate the permanence time of Ps inside the nanochannels. The scope of this work is primarily to explain which kind of data analysis will be necessary to study the Ps emission and cooling using the ToF apparatus and not to examine in detail the measured *Energy* and *ToF spectra*. In addition improvements and optimizations necessary on the set up for more relevant measurements will be evidenced and discussed.

### 3.7.1 Measurements by using the Trento positron beam

The Trento slow positron beam is located at the Department of Physics of the University of Trento and, at present, it delivers around  $7 \cdot 10^3$  moderated positrons/s. Fast positrons emitted by  $\beta^+$  decay from a  $^{22}\text{Na}$  source (2.5 mCi on 1/4/2011) with a broadened spectrum (0-500 keV), are moderated by a tungsten foil to an energy of few eV and transported by a completely electrostatic system up to the target surface [86]. The electric scheme was designed so that the target remains at ground potential, while the source is set at higher potential [87]. The final positrons energy, determined by the different voltages present between the source and the target, can be varied from 20 eV to 25 keV, with an energy spread of  $\pm 1$  eV.

Being the Trento beam focused on the target with variable kinetic energy in a free magnetic field region, the electron optical system of the ToF apparatus explained in section 3.2.3 was not necessary here. Only the main chamber with the target holder and the cryostat were connected to the beam line and mounted in such a way that the target surface coincides with the focal point of the beam design. In this configuration the target was kept at ground for all measurements and the positron kinetic energy was determined by the source potential.

In the Trento measurements the nanochannels diameters were chosen in the 5-8 nm range. They were obtained by etching a Si sample with a current of 10 mA and an anodization time of 17 min.

Measurements at implantation energy of  $E_+ = 6$  or 7 keV, with the sample kept at a temperature of 150 K and 300 K and with scintillators placed at different distances  $z_0$  of 1, 2, 3 cm were done.

The counting rate of the start-stop coincidence events was around 0.0125 counts/s and to record at least  $10^4$  coincidence events for each spectrum about 9 days were needed.

In the following the fitting function needed to extract the Ps energy distribution and its average temperature will be described; then the fitted parameters, as obtained by analyzing the preliminary spectra, will be studied as a function of the sample temperature and of the positron implantation energy. At the end a further analysis on the evaluation of the permanence time of Ps inside nanochannels will be presented, in order to evidence the possible influence on the extracted Ps temperature.

#### The fitting function

The *Energy Spectra* in the semilogarithmical scale were fitted by eq.3.16, written as a superimposition of two exponential distributions, in which the

exponents  $n$  and  $m$  were varied from  $-1/2$  to  $1$  in steps of  $1/2$ :

$$\ln [N(E_{\perp})] = \ln \left[ \alpha E_{\perp}^m A(m, T_1) e^{-\frac{E_{\perp}}{k_B T_1}} + \beta E_{\perp}^n B(n, T_2) e^{-\frac{E_{\perp}}{k_B T_2}} \right] \quad (3.16)$$

where  $A(m, T_1)$  and  $B(n, T_2)$  are normalization factors such that:

$$\begin{aligned} \int_0^{\infty} E_{\perp}^m A(m, T_1) e^{-\frac{E_{\perp}}{k_B T_1}} dE_{\perp} &= 1 \\ \int_0^{\infty} E_{\perp}^n B(n, T_2) e^{-\frac{E_{\perp}}{k_B T_2}} dE_{\perp} &= 1 \end{aligned} \quad (3.17)$$

and  $\alpha$  and  $\beta$  the weight parameters of the two distributions, satisfying the condition:  $\alpha + \beta = 1$ . The temperatures  $T_1$ ,  $T_2$  of the two distributions and the factors  $\alpha$  and  $\beta$ , were free parameters in the fitting procedure.

In the Fig. 3.12 an example of fit is shown for the *Energy Spectrum* measured at 7 keV positron implantation energy, 150 K sample temperature and at a vertical distance  $z_0 = 1$  cm from the sample surface. The data were obtained after a moving average of 21 channels. Due to the big error on the data, the confidence level of the fit is low and more data should be acquired to obtain a fit with higher significance (see section 3.8); anyway some preliminary considerations can be asserted and some important results are achieved with a view to the future measurements.

As the fitting procedure applied to all the *Energy Spectra* here presented does not support  $\alpha = 0$  or  $\beta = 0$ , a single distribution seems not to be sufficient to describe the Ps energy distribution. On the contrary the best fits were obtained with  $m = -1/2$  and with  $n = 0$ . In this case the two distributions correspond respectively to an one-dimensional Maxwellian distribution (with  $m = -1/2$ ), describing the coolest part of the spectra:

$$F(E_{c\perp}) = \frac{1}{\sqrt{\pi k_B T_1 E_{\perp}}} e^{-\frac{E_{\perp}}{k_B T_1}} \quad (3.18)$$

whose average kinetic energy is given by:  $E_{c\perp} = \frac{k_B T_1}{2}$  and to a Beam Maxwellian distribution (with  $n = 0$ ) describing the warm part of the spectra:

$$G(E_{w\perp}) = \frac{1}{k_B T_2} e^{-\frac{E_{\perp}}{k_B T_2}} \quad (3.19)$$

whose average kinetic energy is given by:  $E_{w\perp} = T_2 k_B$ . Thus the resulting fitting equation is:

$$\ln [N(E_{\perp})] = \ln \left[ \alpha \frac{1}{\sqrt{\pi k_B T_1 E_{\perp}}} e^{-\frac{E_{\perp}}{k_B T_1}} + \beta \frac{1}{k_B T_2} e^{-\frac{E_{\perp}}{k_B T_2}} \right] \quad (3.20)$$

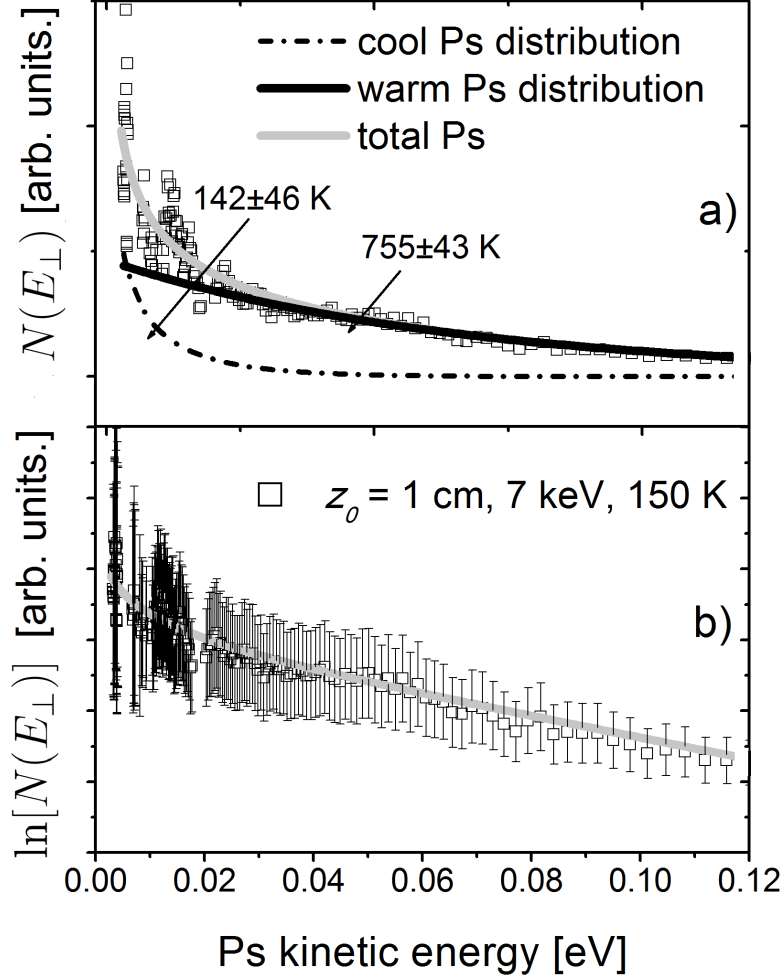


Figure 3.12: Ps *Energy Spectrum* measured at 7 keV positron implantation energy, 150 K sample temperature and at a vertical distance  $z_0 = 1$  cm from the sample surface. The shown data were obtained after a moving average of 21 channels. The spectrum is represented in: a) linear scale, b) semi-logarithmical scale. The light gray curves are the best fit as obtained by using a function sum of a one-dimensional Maxwellian and a Beam Maxwellian distribution (see eq.3.16). The one-dimensional Maxwellian distribution corresponding to the cool Ps fraction and the Beam Maxwellian distribution corresponding to the warm Ps fraction are reported in a) as dash-dot and continuous black lines, respectively. Their characteristic temperatures are also indicated.

From now on, these distributions will be called cool distribution and warm distribution, respectively.

In Fig.3.12, where the two distributions are also shown singly, it is possible to notice how the Beam Maxwellian dominates at higher energy, while the one-dimensional Maxwellian is crucial to describe the steep slope of the data at lower energy. This fit procedure is an evolution of the procedure applied in the first ToF measurements [63] and it could be used for better investigate Ps energy distributions. In that case two slopes were more evident and two linear fits in the semi-logarithmical scale were in agreement with the data. In fact if the temperatures of the two distributions are very different, as in that case, the approximation that the sum of two exponentials can be fitted separately can be done; but, as in this case, when the temperatures are not so different a single fit function including both the distributions is necessary. This effect is evidenced in Fig. 3.13 where the sum of two Beam Maxwellian distributions with different temperatures is shown together with the single distributions separately. In the case of plot a), where the two temperatures are nearer, the sum function yet shows the two same slopes of the single distributions, while in plot b), where the two temperatures are farther apart, the sum results very different from the single distributions.

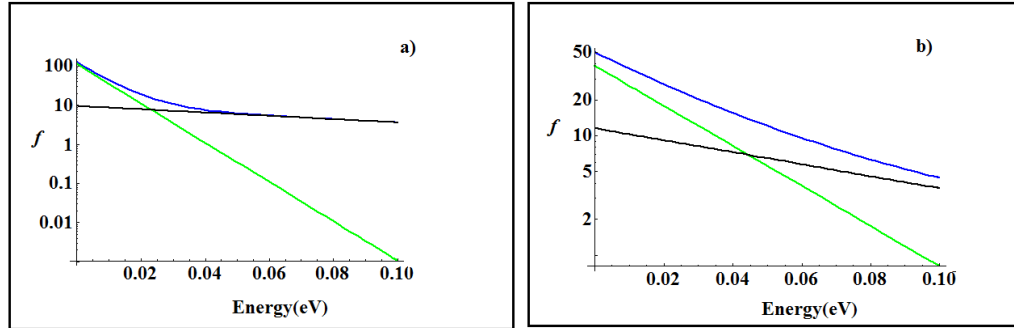


Figure 3.13: Plot of two Beam Maxwellian distributions (in green and in black) with two different temperatures, and plot of their sum (in blue). In plot a) the temperatures are 100 K (green) and 1200 K (black), while in plot b) the temperature are 300 K (green) and 1000 K (black). Only in plot a) the two distributions are dominant each one in a different energy range and the sum function shows the two slopes of the two distributions. In plot b) the contribution of both functions is relevant in the whole energy range and the final slope of the sum function is very different from the slopes of the single distributions.

Anyway, if the current fitted distribution of eq.3.20 would be confirmed by

more significative fits, its origin could be explained by the following argument. In a classical description (see section 2.2.3) the energy loss  $\Delta E$  of a Ps hitting with energy  $E$  the atoms of mass  $M$  on the surface of the nanochannels is  $\frac{\Delta E}{E} \cong \frac{2m_{Ps}}{M}$  with the consequence that the slowing down is very fast during the first collisions, while it relaxes within a few nanoseconds with increasing the number of collisions. Because the number of collisions that Ps can undergo before escaping into the vacuum is strongly related to the positron implantation energy that determines the depth of Ps formation into the sample, a fraction of Ps, emitted from the nanochannels region in the proximity of the sample surface, escapes hot into the vacuum with energies higher than 200 meV [37]. The warm distribution, well approximated by a beam maxwellian, could come from Ps that has reached the conditions such that the energy loss  $\Delta E$  for each collision is in the release region. This distribution takes into account Ps atoms which, undergoing a different number of collisions, are emitted randomly with similar energy. At last, the presence of the one-dimensional Maxwellian distribution could be explained by considering Ps atoms residing inside the nanochannels long enough for approaching thermal equilibrium with the medium through collisions with the walls of the nanochannels.

### Fitted parameters from preliminary measurements

The temperature  $T_1$  and  $T_2$  of the two distributions and the weights  $\alpha$  and  $\beta$  present in the final distribution of eq.3.20 were studied for different values of the positron implantation energy  $E_+$  and sample temperature  $T_{sample}$  in order to investigate the cooling process. In Tab.3.2 the fitted temperatures and the parameters are shown as resulted for the spectra measured at 6,7 keV positron implantation energies, 300 K, 150 K sample temperatures and at a vertical distance  $z_0 = 1$  cm from the sample surface.

Firstly the positron implantation energy was chosen as variable for different spectra. The goal of these measurements was to know the right positron

| $E_+$ [keV] | $T_{sample}$ | $\alpha$ | $T_1$          | $\beta$ | $T_2$            |
|-------------|--------------|----------|----------------|---------|------------------|
| 6           | 300 K        | 0.37     | $684 \pm 50$   | 0.63    | $1800 \pm 50$ K  |
| 7           | 300 K        | 0.45     | $290 \pm 17$   | 0.54    | $1688 \pm 150$ K |
| 7           | 150 K        | 0.25     | $142 \pm 46$ K | 0.75    | $755 \pm 43$     |

Table 3.2: Table of fitted parameters for spectra measured at 6,7 keV positron implantation energies, 300 K, 150 K sample temperatures and at a vertical distance  $z_0 = 1$  cm from the sample surface.

implantation energy at which a fraction of Ps emitted has energy near the thermal one. Because the average positron implantation depth is related to the positron implantation energy by eq.2.1 at lower  $E_+$  Ps is formed near the surface and it does not make enough collisions with the nanochannel walls to reach thermal energy. Increasing the positron implantation energy from 6 to 7 keV, the temperature  $T_1$  of the cool distribution strongly decreases from  $684 \pm 50$  K to  $290 \pm 17$  K, while the temperature of the warm distribution varies only from  $1800 \pm 50$  K to  $1688 \pm 150$  K. Since the temperature of the cool distribution is compatible with the target temperature, a fraction of Ps can be considered thermalized and the positron implantation energy of 7 keV can be fixed as the minimum energy of the positron beam energy for cooling Ps up to thermal energy, in samples with nanochannels diameter of 5-8 nm.

The same analysis was done by changing the sample temperature. By keeping the positron implantation energy at  $E_+ = 7$  keV and by setting the sample temperature at 150 K the parameters  $\alpha = 0.25$ ,  $T_1 = 142 \pm 46$  K and  $\beta = 0.75$ ,  $T_2 = 755 \pm 43$  K were obtained. Here the temperatures of both the two distributions decreased with respect to the values obtained with the same positron implantation energy and with  $T=300$  K. In particular the  $T_1$  value was also compatible with the target temperature, showing that a fraction of Ps is emitted thermalized into vacuum also at lower sample temperature. Since the  $\alpha$  parameter was smaller with respect to the  $\alpha$  obtained at 300 K, the emitted fraction was less abundant, but anyway this measurement confirms again that, if  $E_+ = 7$  keV, the Ps atom spends enough time inside the nanochannels to be thermalized near the sample temperature.

More it points out the possibility to tune one of the two average temperature of Ps energy distribution by changing the sample temperature and by choosing the nanochannel dimension large enough to avoid the quantum confinement conditions. For investigating this effect, more measurements at different and lower temperatures were planned. In fact 150 K was the lowest possible value by using Trento positron beam, because the long acquisition time hinders measurements at lower temperature, where residual gases, in particular water vapor, can contaminate the sample surface in a few hours also in a  $10^{-9}$  Torr vacuum. For this reason the apparatus was set up at the Munich facility, where with higher beam intensity a lower temperature of 50 K was investigated and future measurements at few kelvin were planned.

### Estimation of the Ps permanence time inside nanochannels

The knowledge of the two distributions, which describe our data, allows a deeper analysis by calculating and using the average Ps time of flight and the Ps average velocity.

Firstly the average time of flight for each distribution  $f(t)$  can be analytically calculated by the relation:

$$\int_0^{+\infty} t f(t) dt \quad (3.21)$$

after expressing eq. 3.19 and 3.18 as a function of the time, by using eq.3.12. The average time of flight  $\bar{t}_{BM}$  and  $\bar{t}_{M1D}$  for the Beam Maxwellian distribution and for the one-dimensional Maxwell-Boltzmann distribution respectively are:

$$\begin{aligned} \bar{t}_{BM} &= \sqrt{\frac{\pi m_{Ps}}{2k_B T}} z_0 \\ \bar{t}_{M1D} &= 2.68 \sqrt{\frac{m_{Ps}}{2\pi k_B T}} z_0 \end{aligned} \quad (3.22)$$

Similarly the velocity distribution can be calculated by:

$$N(v)dv = N(E)dE \quad (3.23)$$

with

$$\frac{dE}{dv} = m_{Ps}v \quad (3.24)$$

and the average velocities  $\bar{v}_{BM}$  and  $\bar{v}_{M1D}$  can be obtained.

$$\begin{aligned} \bar{v}_{BM} &= \sqrt{\frac{\pi k_B T}{2m_{Ps}}} \\ \bar{v}_{M1D} &= \sqrt{\frac{2k_B T}{\pi m_{Ps}}} \end{aligned} \quad (3.25)$$

With these values a new analysis can be done in order to extract an average cooling time (also called *permanence time*) of Ps inside the nanochannels. In fact the average time  $t$  is the sum of the time spent by Ps inside nanochannels  $t_p$  and of the real time of flight  $t_f$  of Ps between the target and the annihilation at a distance  $z_0$ :

$$t = t_p + t_f \quad (3.26)$$

In the analysis of the ToF measurements carried out up to now the time spent by Ps inside the nanochannels was considered negligible with respect to the time of flight, but this approximation can be acceptable for large distances  $z_0$  or for very fast Ps, which have done only few collisions inside the nanochannels, before being emitted into vacuum.

In any case a more precise evaluation of the Ps velocity requires to take into account the permanence time of Ps inside the sample and we have pointed out

a possible method to estimate the Ps permanence time inside the nanochannels. This method consists in measuring *ToF spectra* at different distances  $z_0$  from the surface of the sample and in observing the change of the average time of the *ToF spectra* as a function of the distance  $z_0$ .

The *ToF Spectra* as obtained by placing scintillators at the distances  $z_0 = 1$  cm,  $z_0 = 2$  cm and  $z_0 = 3$  cm, are shown in Fig. 3.14. From each spectrum the average times of flight  $\bar{t}$  were calculated for the cool and warm components (by using eq.3.22).

The  $\bar{t}$  values are reported as a function of  $z_0$  in Fig. 3.15, where the error on  $\bar{t}$  arises from the statistical error on the temperature of the two distributions, while the error on  $z_0$  is due to the spatial resolution of  $\Delta z$  at the target axis, as determined by the geometry of the set-up (see section 3.3). In particular it was evaluated as the standard deviation,  $\Delta z/\sqrt{3}$  mm, associated to a uniform distribution. By fitting the data with the equation

$$\bar{t} = t_p + \frac{z_0}{w} \quad (3.27)$$

the values of  $t_p$  and  $w$  can be obtained, where the intercepts of the two sets of values  $t_p = t_{c1}$  and  $t_p = t_{c2}$  represent the average Ps cooling time of the cool and warm distribution, respectively, while the slopes of the two straight lines  $\frac{1}{w}$  are linked to the average velocities  $\bar{v}$  of the two distributions. In fact by considering the two distributions and the eq.3.22 and 3.25, a relation between the average time of flight and the average velocity can be achieved for each distributions:

$$\begin{aligned} \bar{t}_{BM} &= \frac{z_0}{\bar{v}_{BM}} \frac{\pi}{2} \\ \bar{t}_{M1D} &= \frac{z_0}{\bar{v}_{M1D}} \frac{2.68}{\pi} \end{aligned} \quad (3.28)$$

and thus the  $w$  parameter is linked to the mean velocities by:

$$\begin{aligned} \bar{w}_{BM} &= \bar{v}_{BM} \frac{2}{\pi} \\ \bar{w}_{M1D} &= \bar{v}_{M1D} \frac{\pi}{2.68} \end{aligned} \quad (3.29)$$

The fit gives an average cooling time  $t_{c1} = 18 \pm 6$  ns for the Ps emitted with the one-dimensional Maxwellian distribution and an average time of  $t_{c2} < 7$  ns for the warm distribution.

In particular the time  $t_{c1}$  could be compared with the theoretical value obtained by a diffusion model proposed in [88], where a cooling time of around 17 ns was found for positron implantation energy of 7 keV. The diffusion

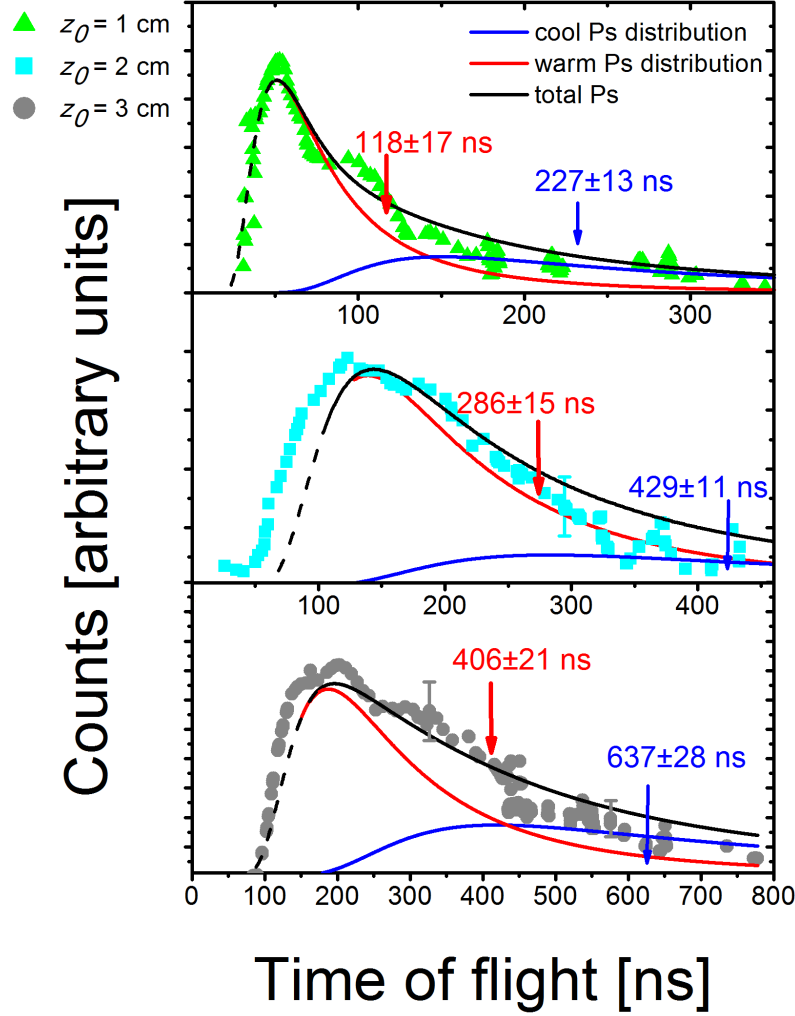


Figure 3.14: *ToF Spectra* measured at 7 keV positron implantation energy and at three different distances  $z_0$  from the sample surface. The black curves are the time distributions corresponding to the energy distributions obtained by best fitting the *Energy Spectra* by using eq.3.16. The blue curves are the time distributions of the cool Ps corresponding to the one-dimensional Maxwellian distributions. Similarly the red curves are the time distributions of the warm Ps corresponding to the Beam Maxwellian distributions. The average time of flight  $\bar{t}$  of each distribution is reported and its position marked by an arrow.

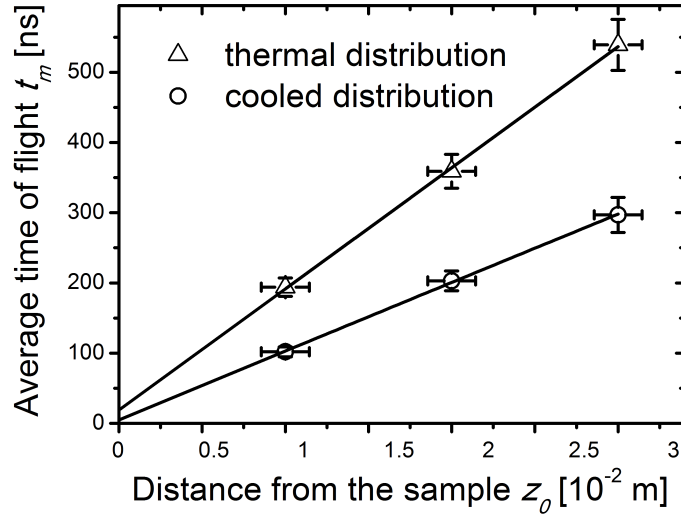


Figure 3.15: Average time of flight calculated from the time distributions of Fig. 3.14, as a function of the distance  $z_0$  from the sample surface. Triangles were used for cool distributions and circles for warm distributions. The intercepts of the two best fit lines give the permanence time  $t_p$  of Ps before being emitted into vacuum. The slopes of the lines are related by eq.3.29 to the average perpendicular flying velocities  $\bar{v}_{M1D}$  and  $\bar{v}_{BM}$  of the two distributions.

model describes Ps with a constant velocity, which thermalizes in a random walk process inside the nanochannels. Since in the real cooling process Ps decreases its kinetics energy from some eV to some meV, the knowledge of the permanence time with more accuracy could allow a deeper understanding of the Ps cooling by comparing the experimental result with more accurate models.

It can be evaluated that more precise results can be obtained by acquiring around five spectra at five different distances, in order to validate the linear fitting procedure. More this result could be important in the AEgIS context, because, by knowing the permanence time for each Ps energy distribution, the synchronization of the laser on the Ps cloud could be possible.

As further improvement for the study of the Ps velocity distribution, the extracted velocity  $v_{BM}$  and  $v_{M1D}$  can be achieved and compared with the previous one obtained by the mean time of flight and calculated as eq.3.29 and shown in Tab. 3.3. From the fitting procedure the values  $v_{M1D} = 4.1 \pm 0.2 \cdot 10^4$  m/s for the cool distribution, and  $v_{BM} = 11 \pm 0.4 \cdot 10^4$  m/s for the warm distribution are obtained. It can be noted, especially in the cold component

|            | Beam Maxwellian     |                               | 1-D Maxwell-Boltzmann |                                |
|------------|---------------------|-------------------------------|-----------------------|--------------------------------|
| $z_0$ [cm] | $\bar{t}_{BM}$ [ns] | $\bar{v}_{BM}$ [ $10^4 m/s$ ] | $\bar{t}_{M1D}$ [ns]  | $\bar{v}_{M1D}$ [ $10^4 m/s$ ] |
| 1          | 126                 | 12                            | 220                   | 3.7                            |
| 2          | 286                 | 10                            | 429                   | 3.9                            |
| 3          | 406                 | 11                            | 637                   | 4.0                            |

Table 3.3: Table of the data shown in Fig.3.15 as acquired with  $E_+ = 7$  keV and  $T_{sample} = 300$  K at different distances  $z_0$ . The average times were calculated by using eq.3.22, while the velocity values were calculates by the using eq.3.25, where the  $T$  parameters were obtained by the fitting procedure explained in section 3.7.1.

(M1D), that the mean velocity is underestimated in the measurement at  $z_0 = 1$  cm due to the permanence time. The difference between the fitted value and the value obtained at  $z_0 = 3$  cm becomes negligible and in particular the relevance of the permanence time is bigger for the cold component. With the described procedure a real estimation of the effect of the permanence time on the ToF measurement can be achieved. The data also confirm that the permanence time can be neglected if  $z_0 \geq 3$  cm or for fast Ps which undergoes few collisions inside the nanochannel.

The presented methodology used to extract the permanence time  $t_p$  of Ps inside the tunable nanochannels will allow, in the future, to study in detail the cooling process of Ps by measuring  $t_p$  for each distribution as a function of the positron implantation energy.

In order to shorten the acquisition time, the described ToF apparatus was assembled at the NEPOMUC facility, where a higher beam intensity is delivered. In the next section other measurements acquired in about 2 days of beam time will be shown. They were performed with lower sample temperatures in order to obtain the fraction of Ps emitted in vacuum with lower energy.

### 3.7.2 Measurements at the NEPOMUC facility

With the final goal of studying Ps cooling up to cryogenic temperature, a sample with nanochannels of 10-16 nm in diameter was chosen in such a way that Ps thermalization is expected to be not hindered by the quantum confinement (see section 2.2.3). The sample was obtained with an etching current of 10 mA for 17 min and with 2 cycles of etching and successive re-oxidation.

The re-moderated positron beam was tuned through the magnetic transport line and the effectiveness of the magnetic field terminator was checked. A tune-up of the electrostatic accelerator's lenses, needed to focus the positrons at the target position, was as well performed with target voltage in 5-10 kV range and an optimized value of  $\sim 7$  mm in diameter was achieved on the target for any sample potential. This value, bigger than that obtained from the SIMION simulations, was due to the bigger diameter of the positron beam coming from the NEPOMUC facility.

Three spectra were acquired at the positron implantation energy of 7 keV, with the scintillators placed at 1 cm above the sample surface: one spectrum was obtained with the sample kept at room temperature (RT), while the second one with the sample kept at 50 K as well as the background spectrum obtained using Si-single crystal sample. To check the absence of condensation on sample surface at low temperature, a further fast measurement was also done after warming the sample at RT and re-cooling it again at 50 K. Measurements at the lowest temperature (10 K as programmed) were prevented by a problem with the cryostat circuit that did not work properly.

The *Energy* and the *ToF Spectra*, as resulted after two days of measurements, are shown in Fig. 3.16 and in Fig. 3.17. A rebin of 24 channels was done on the rough data and a 5-bins moving average was performed. The number of bins involved in the rebinning and smoothing processes were

chosen in order to consider a correspondent time (1 bin=0.25 ns) smaller than the time resolution of 8 ns, as extracted by the background analysis (see section 3.6).

Both the spectra measured at 50 K and at 300 K (see Fig.3.17) were fitted with eq.3.16, and the extracted parameters are shown in Tab.3.4. As in the Trento measurements the *Energy Spectrum* of the emitted Ps can be fitted by two exponential functions: one is the one-dimensional Maxwell-Boltzmann distribution (with  $m = -1/2$ ), where the temperature is the same of the sample and the other one is the Beam Maxwellian component (with  $n = 0$ ) at temperature higher than the sample one. Despite the low reliability of the fits, due to the big errors on the low energy data, the obtained parameters show an increment on the component at lower energy in the 50 K spectrum (see parameter  $\alpha$  in Tab.3.4) with respect to the same component at 300 K. Moreover an evaluation of Ps emitted with velocity lower than  $5 \cdot 10^4$  m/s (i.e.  $t_f < 200$  ns) can be done, by considering the *ToF Spectrum* in Fig. 3.16b, where the vertical line marks the time corresponding to a Ps velocity of  $5 \cdot 10^4$  m/s. In this plot a 30 % of the formed Ps was found with velocity

| $T_{sample}$ | $\alpha$ | $T_1$          | $\beta$ | $T_2$          |
|--------------|----------|----------------|---------|----------------|
| 300 K        | 0.41     | $300 \pm 23$ K | 0.59    | $988 \pm 80$ K |
| 50 K         | 0.85     | $50 \pm 5$ K   | 0.14    | $955 \pm 50$ K |

Table 3.4: Values of the fitting parameters obtained by fitting the *Energy Spectra* measured with sample held at 50 K and 300 K and with detectors placed at the vertical distance  $z_0 = 1$  cm from the sample surface. The positron implantation energy was  $E_+ = 7$  keV.

lower than  $5 \cdot 10^4$  m/s (i.e.  $t_f < 200$  ns) and by taking into account that from nanochannelled silicon samples with 10-16 nm channel size about 30 % of implanted  $e^+$  give Ps in vacuum [58], we can estimate that about 9% of implanted  $e^+$  are emitted with a velocity  $5 \cdot 10^4$  m/s in a sample held at 50 K. This result is particularly important for AEgIS experiment, where a high fraction of cold Ps with velocity lower than  $5 \cdot 10^4$  m/s is demanded for high production rate of anti-hydrogen. Compared to the previous measurements performed at higher temperature (see section 3.7.1), the velocity distribution of Ps emitted from sample with 10-16 nm in diameter, kept at 50 K seems to be shifted to lower velocity showing that the cooling process inside the nanochannel with this size seems to be more efficient. Despite the low statistic of these data, the measurement at 50 K is the first ToF measurement at cryogenic temperature on a *Si-SiO<sub>2</sub>* system and it gives important informations about the Ps cooling process inside the nanochannels. Anyway

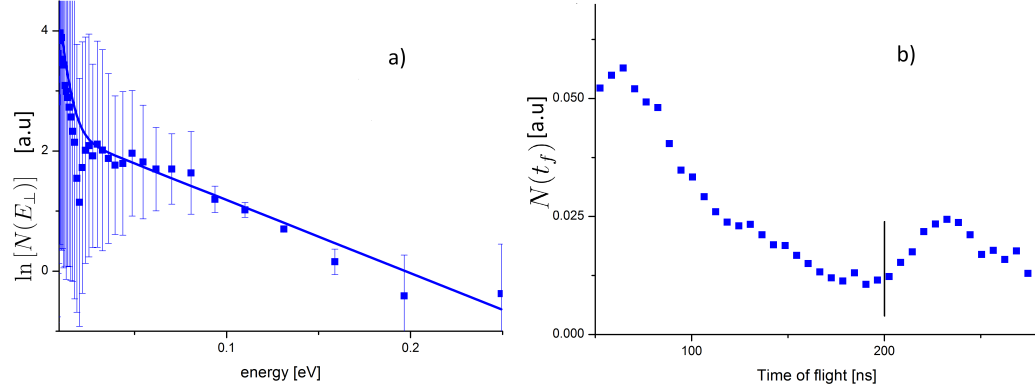


Figure 3.16: The *Energy Spectrum* (in a) and the *ToF Spectrum* (in b), as obtained with sample kept at  $T=50$  K and with  $E_+ = 7$  keV. The eq.3.20 was used for the fit. The values of the fitting parameters are reported in Tab.3.4.

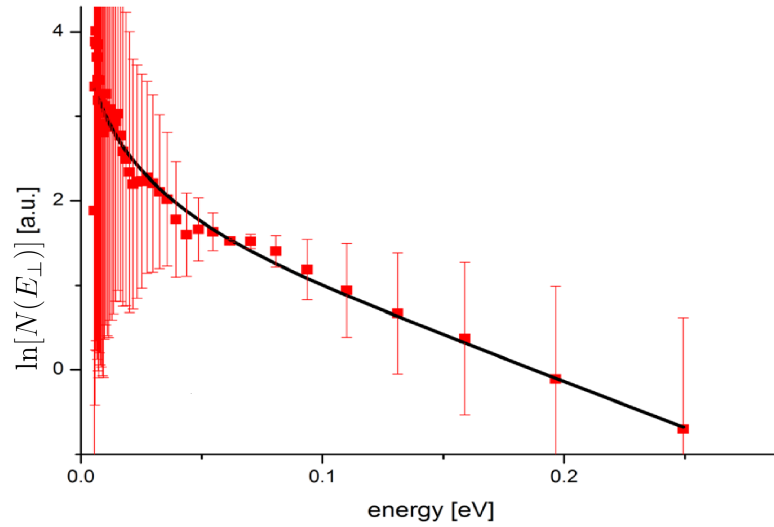


Figure 3.17: The *Energy Spectrum* as obtained with sample kept at  $T=300$  K and with  $E_+ = 7$  keV. The eq.3.20 was used for the fit, whose fitted parameters are reported in Tab.3.4.

more measurements at lower temperature with the same sample have to be performed in order to confirm the result. For each measurement more events have to be accumulated in order to decrease the final error on the *Energy Spectrum*, allowing the fit to be more reliable and the final parameters to be extracted with more accuracy.

In conclusion the data as acquired by Trento and NEPOMUC positron beam give some preliminary results about the Ps thermalization process, but they are yet insufficient for a more quantitative analysis. Despite the measurement conditions and in particular the beam intensities were very different, a large error on the final spectrum was however obtained. Since the error on each data is related to the total counts of Ps annihilation and to the background counts, that depend on the beam intensity and on the measurement duration, an analysis in terms of signal to background ratio has to be done and an optimization on these parameter has to be achieved. In the next section a correlation between the beam intensity, the measurement duration and the final data errors is discovered and a best condition of measurement is indicated.

## 3.8 Parameters optimization

### 3.8.1 The signal to background ratio

The ToF measurement is based on the presence of two digital signals (start and stop) inside a temporal window  $T$ , which begins with the start signal. Only if both signals appear in the right start-stop order, during the interval  $T$ , the elapsed time between the two events is registered, otherwise the next start is waited.

However because of the prompt annihilations or the  $\gamma$ -rays coming from the reactor or the high beam intensity (see below), false events can be acquired. In Fig.3.18 the rough data, as obtained with nanochanneled Si sample (in red) and with Si-single crystal for background measurement (in blue), renormalized to the same measurement time, are shown. By comparing the two spectra, three zones, filled by events coming from different physical processes, can be identified:

- the prompt peak zone (indicated with B in Fig.3.18), which regards counts inside the peak. It is present in both spectra and it is centered on  $T_D$ , the delayed time set by the acquisition chain. It is due to the prompt peak annihilations with time of flight  $t_f$  lower than 50 ns, as explained in section 3.6,

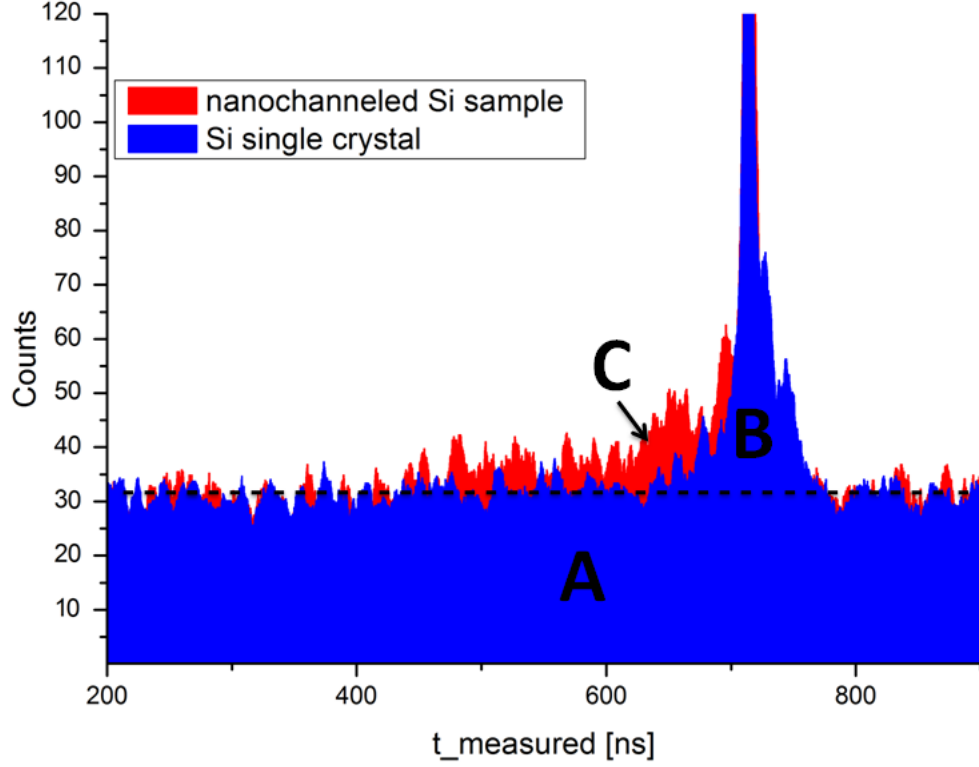


Figure 3.18: The rough data as obtained at NEPOMUC facility by measuring with nanochanneled Si sample (in red) and with Si-single crystal (in blue) for background measurement. With A, B, C letters are indicated the false counts with flat distribution, the prompt peak zone of the background events, and the Ps annihilation events, respectively.

- the flat region, which includes counts uniformly distributed in the time window. It is indicated with A in Fig.3.18, in which all counts below the dashed line are included. These events are due to random counts occurring in the coincidence measurements, as will be explained in this section and are present in both the spectra.
- the counts due to Ps annihilations (indicated with C in Fig.3.18), present only in the spectrum obtained with nanochanneled Si sample.

In order to study how the measurement conditions influence the three regions, the signal to background ratio  $SB$ , defined as the ratio between the Ps annihilations and the false counts uniformly distributed (i.e.  $\frac{C}{A}$ ) was studied as a function of the beam intensity  $f$  and of other parameters characterizing our

acquisition system.

In particular if  $f_c$  and  $f_s$  are the channeltrons and the scintillators counting rate, respectively, the false counting rate (involving the A zone in Fig.3.18), if the time window  $T$  is used, can be evaluated as [89]:

$$r_F = f_c f_s T \quad (3.30)$$

while the true counting rate is:

$$r_T = f \epsilon_s y \epsilon_c p \quad (3.31)$$

Because the equations 3.6 and 3.5 the false counting rate can be written as:

$$r_F = (f \epsilon_c p + b_c)(f \epsilon_s(y + pt) + b_s)T \quad (3.32)$$

and the function  $SB(f)$  defined as  $SB(f) = \frac{r_T}{r_F}$  is obtained:

$$SB(f) = \frac{f \epsilon_s y \epsilon_c p}{(f \epsilon_c p + b_c)(f \epsilon_s(y + pt) + b_s)T} \quad (3.33)$$

From eq.3.33 the  $SB$  function results strictly dependent on the beam intensity  $f$  as well as on the characteristic parameters (efficiency  $\epsilon$  and background counts  $b$ ) of our detectors.

By using the Trento parameters, shown in Tab.3.5 where the values of  $\epsilon_c p$

|        | $f [e^+/s]$      | $f_c [cps]$ | $b_c [cps]$ | $\epsilon_c p$ | $f_s [cps]$ | $b_s [cps]$ | $\epsilon_s(y + pt)$ |
|--------|------------------|-------------|-------------|----------------|-------------|-------------|----------------------|
| Trento | $7 \cdot 10^3$   | 30          | 0           | 0.004          | 0.44        | 0.3         | $2 \cdot 10^{-5}$    |
| Munich | $1.8 \cdot 10^6$ | 400         | 80          | 0.0003         | 70          | 20          | $2 \cdot 10^{-5}$    |

Table 3.5: Table of measurement parameters of Trento and Munich set up. The values of  $\epsilon_c p$  and  $\epsilon_s(y + pt)$  were deduced from  $f_c$  and  $f_s$  by using eq.3.5 and 3.6, respectively. The Trento parameters were used for the plot of SB function shown in Fig.3.19, while the Munich parameters were used for Fig.3.21.

and  $\epsilon_s(y + pt)$  were calculated by measuring the counting rate  $f_c$  and  $f_s$  and by using eq.3.5 and 3.6 respectively, the  $SB(f)$  as a function of the beam intensity  $f$  was calculated, whose plot is shown in Fig.3.19. It is interesting to note that when the beam intensity increases, both true and accidental coincidences rates increase (because  $r_T \propto f$  and  $r_F \propto f^2$ ), while the ratio  $SB$  decreases as  $\sim 1/f$ . The  $SB$  value for the Trento measurement was around 20, a high value which allowed us to neglect the false counts (called A counts

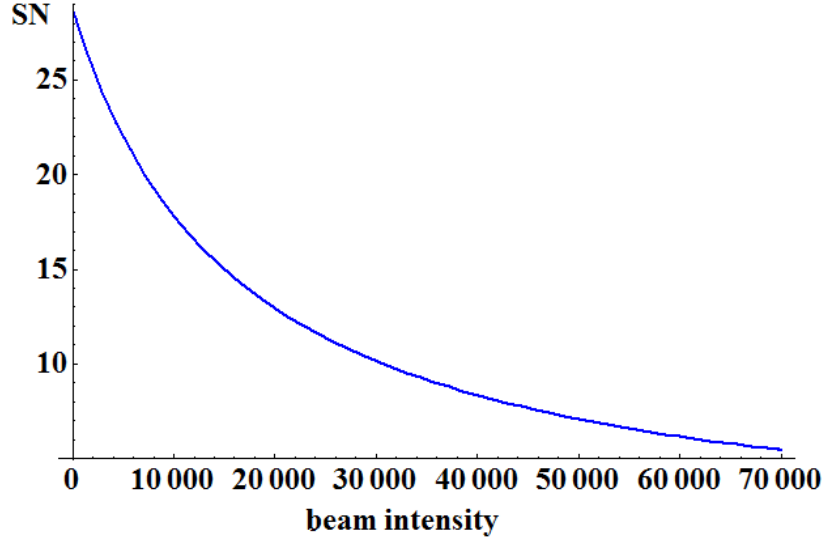


Figure 3.19: The SB function of eq.3.33 as a function of the beam intensity calculated using Trento parameters shown in Tab.3.5. The SB value for the Trento measurements, where the beam intensity was estimated to be around  $7 \cdot 10^3 e^+/s$ , was around 20.

in Fig.3.18) and thus not to subtract them during the analysis.

Anyway so far the beam has been considered as a flux of particles with constant frequency  $f$ , where every positron arrival on the target is separated from the next one from the time interval  $1/f$ . In the reality the beam is generated by a nuclear process, described by the Poissonian statistic. Thus the probability  $P(k; \mu)$  of having  $k$  positrons arriving inside the time window  $T$ , if the events average value in the same window is  $\mu = fT$  can be expressed as in eq.3.34:

$$P(k; \mu) = \frac{\mu^k}{k!} e^{-\mu} \quad (3.34)$$

whose plot calculated for  $\mu = 1$  and  $f = 10^6 e^+/s$  is shown in Fig.3.20.

From the plot it results that  $P(k; \mu)$  is not negligible for  $k \leq 3$  and as a consequence the Poissonian effect has to be considered. In order to take into account it, a corrected signal to background ratio  $SB_P(f)$  was introduced. In this case the real beam can be considered as a weighted superimposition of different beams of intensities  $f_i = nf$  with  $n \in N$ , whose weights are given by the Poissonian distribution  $P(k, fT)$ . For each beam of intensity  $f_i$ , the  $SB(f_i)$  describes the respective signal to background ratio.

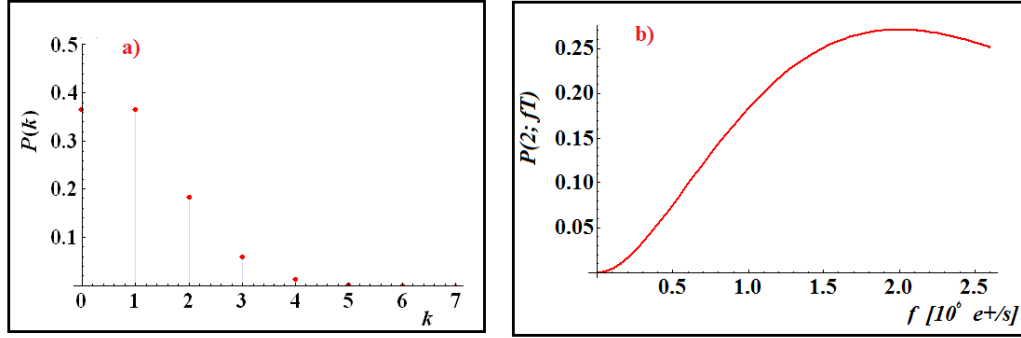


Figure 3.20: The Poissonian probability  $P(k; \mu)$  of having  $k$  positrons arriving inside the time window  $T = 1 \mu s$ .

- a) The values of  $P(k; \mu)$  for  $k \leq 7$  if the mean event value in the window  $T$  is  $\mu = fT = 1$ , calculated with a beam of  $f = 10^6 e^+/s$ .  
b)  $P(k = 2; \mu = fT)$ , the poissonian probability of having 2 positrons in the window  $T$  as a function of the beam intensity  $f$ .

Thus the final  $SB_P(f)$  function can be written as:

$$SB_P(f) = \frac{P(1; fT)SB(f) + P(2; fT)SB(2f) + P(3; fT)SB(3f)}{P(1; fT) + P(2; fT) + P(3; fT)} \quad (3.35)$$

where the terms for  $k \geq 4$  were considered negligible. It is worth noticing that if the beam intensity decreases, the probability of having more than one positron decreases and in particular because  $\lim_{f \rightarrow 0} SB_P(f) = SB(f)$ , the  $SB$  and  $SB_P$  functions tend to overlap each other for  $f < 2 \cdot 10^5 e^+/s$  as shown in Fig.3.21.

Using the Munich parameters reported in Tab.3.5 the  $SB_P$  function was calculated and it was compared with  $SB$  as obtained in the Trento condition shown in the Fig.3.21. The parameters are very different from that used for the Trento conditions: firstly in Munich due to radiations coming from the reactor hall, the scintillators counting rate measured without positron beam hitting on the sample was found  $\sim 20$  cps and for this reason  $b_s$  parameter was higher than in Trento; secondly the channeltrons efficiency was found lower, because of the deterioration of the device kept in air for a long time and thirdly their dark counts  $b_c$ , measured without positron beam hitting on the sample was found higher, whose cause is not yet clear.

Therefore in the plot of Fig.3.21, due to background counts of scintillators and channeltrons ( $b_s, b_c \neq 0$ ) the  $SB$  behavior at low intensity beam is different from that shown in Fig.3.19. In fact in this case there is a region, for  $f$  lower than  $\sim 1.5 \cdot 10^5 e^+/s$ , where  $SB$  increases, until it reaches a maximum, where

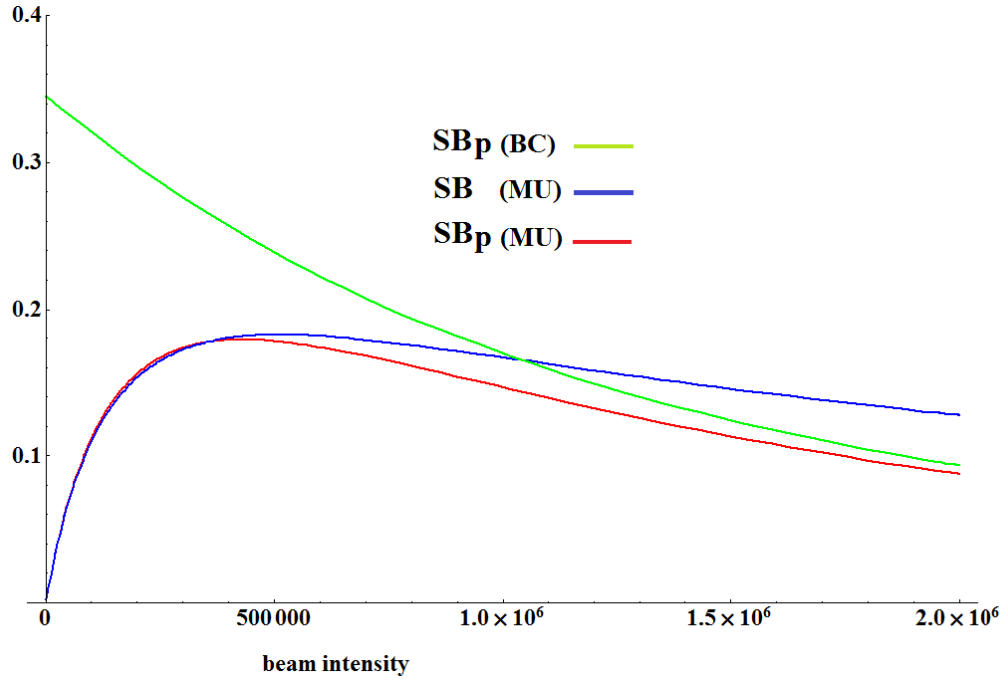


Figure 3.21: signal to background ratio  $SB$  of eq.3.33 (in blue) and the corrected  $SB_P$  of eq.3.35 (in red), calculated using Munich (MU) parameter reported in Tab.3.5 as a function of the beam intensity. The green plot of  $SB_p$  is the signal to background function in the best possible conditions (BC) that could be obtained by using the Trento values for  $\epsilon_c p$  and  $b_c$ , but by using  $b_s$  as in the Munich conditions, because of the background radiations that cannot be reduced.

the noise due to the accidental rate begins to dominate. More the comparison between  $SB_P(f)$  and  $SB(f)$  shows that if  $f$  is low, the Poissonian effect is not relevant, while it increases at higher beam intensity where the signal to background ratio is heavily reduced. The plot is in agreement with the experimental results obtained by measuring with a beam intensity of  $1.8 \cdot 10^6$   $e^+$ /s, where a SB of 0.1 was experimentally found on the rough data and with beam intensity of  $2.5 \cdot 10^6$   $e^+$ /s, where SB resulted to be 0.06.

In conclusion the signal to background condition in the Munich measurement was heavily worst than in Trento, because of the higher beam intensity and the lower efficiency of channeltrons but, above all it is due to the background on channeltrons and on scintillators (see Tab.3.5).

A further estimation of the  $SB_P$  value can be done by considering the best condition that could be reached by considering channeltrons parameters as in the Trento measurements and by taking into account only the background counts present in the reactor, which affect the  $b_s$  value. From now on these optimized parameters will identify the *best condition* (BC) of the set up. In this case the  $SB_P$  function would result as in Fig.3.21: it would be better than in the previous case, but however it would not reach the range of values like in the Trento conditions.

Therefore in order to discover which signal to background ratio could be acceptable for our measurements, it must be taken into mind that the final goal is the *ToF* and the *Energy Spectra*. For this reason in the following section, a relation between the error on the data and the system set up parameters will be done, in order to know in advance which beam intensity and detectors efficiency are request to have data with acceptable error.

### 3.8.2 Error analysis and parameters optimizations

In this section the relative error on each point of the *ToF* and *Energy Spectra* is calculated as a function of the beam intensity, the acquisition time and the other system parameters, introduced in the previous section.

The relevance of the  $SB$  function and of the false counts (A region in Fig.3.18) on the final spectrum can be evaluated by considering that the relative error on each point is mostly related to the false counts. In fact due to the poissonian statistic on the counting measurement, the error on each point is  $\sqrt{N_i}$  where  $N_i$  is the counts for each bin. Because in each bin both the counts due to Ps annihilation  $N_i^{Ps}$  (C region in Fig.3.18) and the false counts  $N_i^F$  are present, the final error results to be:

$$\Delta N_i = \sqrt{N_i^{Ps} + N_i^F} \quad (3.36)$$

where the prompt annihilations can be neglected, if data with time of flight  $t_f$  bigger than 40-50 ns are considered.

In particular since the false counts are uniformly distributed on the overall window, the value of  $N_i^F$  can be estimated to be the same for all points and it can be extracted from the rough data by considering the region in the window where neither the prompt annihilations nor the Ps annihilations are present (i.e. for  $t_f$  very high or for  $t_{measured} > t_{prompt}$  (see section 3.6)). Thus the relative error of each bin can be calculated as:

$$\frac{\Delta N_i}{N_i} = \frac{\sqrt{N_i^{Ps} + N_i^F}}{N_i^{Ps}} \quad (3.37)$$

and because of eq.3.8 and eq.3.12, it is the same for the points of the *ToF* and *Energy Spectra* as well as for the *Rough Spectrum*.

At this point, in order to study  $\frac{\Delta N_i}{N_i}$  as a function of the system parameters and the beam intensity, the  $N_i^F$  and  $N_i^{Ps}$  terms can be differently expressed. On one hand by remembering that the false counts are described by flat distribution, the estimation of the false counts in each bin  $i$  which covers a time interval  $\Delta t$ , can be expressed by the relation:

$$N_i^F = r_F \Delta t = (fp\epsilon_c + b_c)[f(y + pt)\epsilon_s + b_s]T\Delta t \quad (3.38)$$

On the other hand an estimation of the Ps annihilations in each bin  $N_i^{Ps}$  can be achieved by considering the total acquired Ps annihilations during the measurement time  $T_m$  and by estimating which fraction of the total counts is found in each bin. Since the total Ps annihilations can be written as  $N^{tot} = r_T T_m$ , the  $N_i^{Ps}$  can be expressed as:

$$N_i^{Ps} = r_T T_m \beta_i \quad (3.39)$$

where  $\beta_i = \frac{n_i}{N_{tot}}$  is defined as the ratio between the counts  $n_i$  in each bin and  $N_{tot}$  the total Ps counts in the overall spectrum.

Thus one of the *Rough Spectra* acquired in Trento or in Munich can be used for calculating the  $\beta_i$  value in order to calculate  $\frac{\Delta N_i}{N_i}$  for each bin. With the view to optimizing the measurements with sample kept at low temperature, the  $\beta_i$  values shown in Fig.3.22 were calculated using the *Rough Spectrum* acquired with sample at 50 K (relative to *ToF Spectrum* of Fig.3.16) as obtained after the flat region and prompt peak subtraction. Moreover it was assumed that no Ps can be observed for time of flight higher than 300 ns, because the detection probability is too low, and thus the condition  $\sum_{t=100ns}^{t=300ns} \beta_i = 1$  is satisfied for the points reported in Fig.3.22.

Thus by substituting  $r_T$ , the Ps counts for each bin  $i$  can be written as:

$$N_i^{Ps} = f y \epsilon_c p \epsilon_s T_m \beta_i \quad (3.40)$$

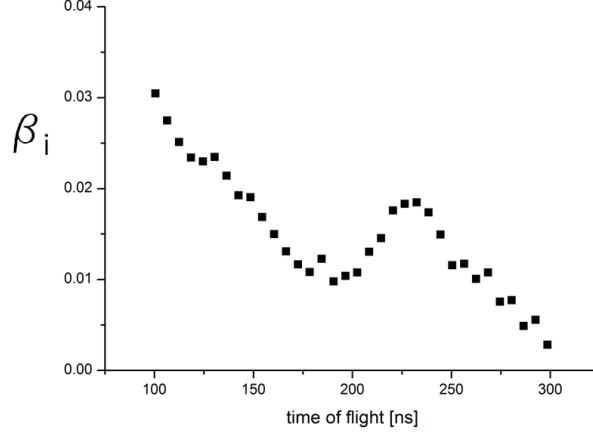


Figure 3.22: The fraction  $\beta_i = \frac{n_i}{N_{tot}}$  of Ps observed in each bin as a function of the time of flight is shown. It was calculated by using the acquired data of Munich measurement obtained with sample kept at 50 K.

Now it is possible to estimate the total counts in each bin where the prompt annihilations are negligible and then the relative errors on each one can be expressed as a function of the acquisition time  $T_m$  and of the other system parameters:

$$\begin{aligned} \frac{\Delta N_i}{N_i} &= \frac{\sqrt{r_F \Delta t + r_T T_m \beta_i}}{r_T T_m \beta_i} \\ &= \frac{\sqrt{(f p \epsilon_c + b_c)(f(y + p t) \epsilon_s + b_s) T \Delta t + f y \epsilon_c p \epsilon_s T_m \beta_i}}{f y \epsilon_c p \epsilon_s T_m \beta_i} \end{aligned} \quad (3.41)$$

These values are different for each point and are bigger in the zone where  $\beta_i$  are low. In Fig.3.23 the  $\frac{\Delta N_i}{N_i}$  values were calculated for each point of the spectrum by using the Trento (in green) and the Munich (in red) parameters and with the acquisition time of 9 and 2 days, respectively. For all plots, due to eq.3.41, the  $\frac{\Delta N_i}{N_i}$  appears to be inversely proportional to the  $\beta_i$  trend, shown in Fig.3.22: in fact if  $\beta_i$  is bigger  $\frac{\Delta N_i}{N_i}$  is lower and vice versa. Anyway the range of values of  $\frac{\Delta N_i}{N_i}$  depends on the measurement condition and the estimated values are in agreement with those experimentally obtained and shown in section 3.7.1 and 3.7.2. As in the previous section where the  $SB$  values were investigated, the statistical relevance of the Trento data is better than the Munich data, but the distance between each other from the relative error point of view is less dramatic with respect to the signal to background ratio point of view. Despite the difference in the beam intensity and

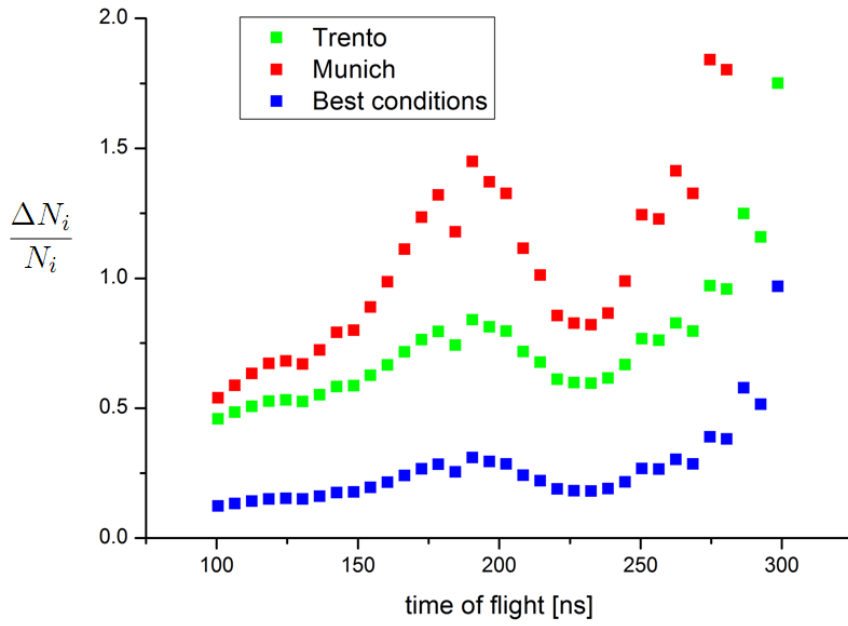


Figure 3.23: The relative error on each point as a function of the time of flight is shown for Trento (green), Munich (red) conditions. Due to eq.3.41 the  $\frac{\Delta N_i}{N_i}$  behavior is inversely proportional to the  $\beta_i$  trend shown in Fig.3.22, because where  $\beta_i$  has a maximum,  $\frac{\Delta N_i}{N_i}$  has a minimum and vice versa. In blue the plot as obtained for the best conditions of measurement (see previous section), with a beam intensity of  $8 \cdot 10^5 e^+/s$  and an acquisition time of 2 days is shown.

in the acquisition time, the range of values of the relative error in the final spectrum in Trento and in Munich spectra is almost the same. In Trento measurement in fact too small was the acquisition time with respect to the low beam intensity and too low was the total acquired counts, while in Munich measurements, despite the number of collected events was very high, a bigger fraction of collected data were due to false counts that enlarged the error on the final data. For this reason in both the cases, as it have just been evidenced in the section before, the high relative error, above all at the highest times of flight (lowest Ps energies), limited the reliability of the fitted functions and of the extracted parameters.

In order to discover if the best measurement condition, discussed in the previous section, could be sufficient to obtain a spectrum with a statistical relevance, and in particular in order to optimize the beam intensity and the acquisition time at the NEPOMUC facility, the relative error of eq.3.41 was studied as a function of the main system parameters.

Because for AEGIS experiment Ps with velocity lower than  $5 \cdot 10^4$  m/s is required, a value of  $\beta_i = 0.01$  relative to a time of flight of  $t_f = 200$  ns measured with detector placed at a distance  $z_0 = 1$  cm from the sample was considered. In the Fig.3.24 the relative error calculated by eq.3.41 was evaluated as a function of the beam intensity, by considering 2 days of acquisition time. It can be observed that a decrement of the relative error occurs with the increasing of the beam intensity up to  $4 \cdot 10^6$   $e^+$ /s even if the slope of the plotted function appears to decrease at higher beam intensity. Thus a bigger advantage is obtained if the beam is increased up to  $8 - 9 \cdot 10^5$ , while a smaller advantage is gained by increasing the beam intensity above to this value. However it is worth noticing that in this analysis the poissonian effect, which increases the false counts if the beam intensity is higher, was not considered and for this reason too high beam intensity has to be avoided.

In conclusion the high beam intensity causes a reduction on the final relative error, because it increases the total counts in the spectrum for the same acquisition time, but however a beam intensity higher than  $8 - 9 \cdot 10^5$   $e^+$ /s has to be avoided for the Poissonian effect and for the worst signal to background ratio.

A similar behavior was found by studying  $\frac{\Delta N_i}{N_i}$  as a function of the acquisition time  $T_m$  calculated with  $f = 9 \cdot 10^5$   $e^+$ /s. As in the previous plot a decrement of  $\frac{\Delta N_i}{N_i}$  occurs by increasing the measurement duration and a bigger advantage is gained during the first hours of data acquisition, while a saturation effect is evidenced after 20 hours of measurement, where the relative errors do not change significantly.

From these results, by considering the best condition of measurement, an

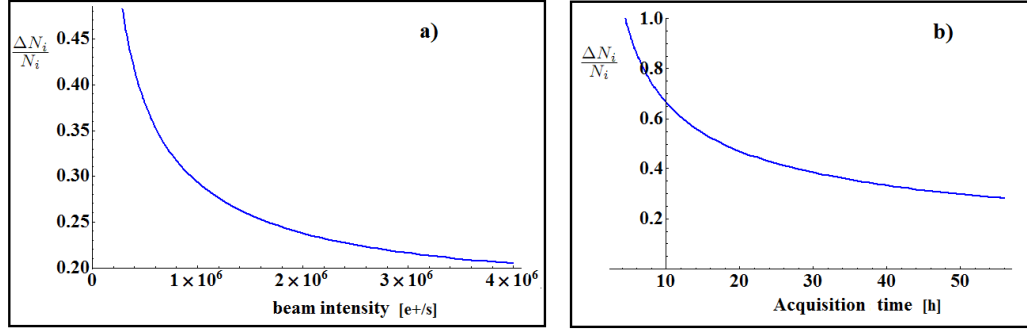


Figure 3.24: The plot of  $\frac{\Delta N_i}{N_i}$  as a function of the beam intensity  $f$  and on the acquisition time is reported.

estimation of the optimum beam intensity can be found if Ps with velocity of around  $5 \cdot 10^4$  m/s ( $\beta_i \sim 0.01$ ) is investigated. In particular by considering two days of acquisition time, which would allow the acquisition of around two spectra for each beam time of 6 days and by fixing the condition  $\frac{\Delta N_i}{N_i} \leq 0.4$  for  $\beta \sim 0.01$ , a necessary beam intensity of around  $7 - 8 \cdot 10^5$  was found. As confirmation the plot of  $\frac{\Delta N_i}{N_i}$  in the final spectrum was calculated for the optimized parameters and it is shown in Fig.3.23, where it results to be lower than 0.4 for  $\beta_i > 0.01$ . Therefore these parameters would allow reaching enough statistic in the region where Ps is emitted with a high yield, while more time would be necessary to reach an acceptable error on that data at higher time of flight where  $\beta_i$  is lower than 0.01.

Moreover when slow Ps is detected the finite resolution  $\Delta z$  on the annihilation position (see section 3.3) must also be considered. In fact since the time spent in front of the slit depends on the Ps velocity  $v$  and it is bigger for slower Ps, the times of flight  $t_f$  of Ps annihilating at  $z - \Delta z$  and at  $z + \Delta z$ , could be very different each other, if the velocity is low. In particular the difference  $\Delta t$  between the two times of flight  $t_f(z - \Delta z)$  and  $t_f(z + \Delta z)$  can be expressed as a function of  $v$ , of the slit dimension  $s$  and on the shield dimension  $L$  and  $D$  (see Fig.3.6) as:

$$\Delta t = \frac{2\Delta z}{v} = \frac{s}{D} \frac{2L + D}{v} \quad (3.42)$$

Thus with  $s$ ,  $L$  and  $D$  fixed as in the presented measurements, for  $v \sim 2 \cdot 10^4$  m/s,  $\Delta t$  results  $\sim 180$  ns and it cannot be neglected. However, in order to better investigate slow Ps, the dimensions  $s$ ,  $L$  and  $D$  must be optimized, by minimizing  $\Delta t$ , but anyway the angular acceptance of the scintillator would decrease and thus the acquisition time would be heavily increased.

### 3.9 Conclusions

In conclusion some preliminary measurements were done using both low and high beam intensity at Trento and at Munich. Some preliminary data as acquired by using different positron implantation energy  $E_+ = 6, 7$  keV, or sample temperatures 300, 150, 50 K and nanochannel dimensions  $\varnothing = 5 - 8$  nm or  $10 - 16$  nm were analyzed.

A useful method of data analysis was presented in order to obtain the *ToF* and the *Energy Spectra* necessary to extract the Ps energy distribution, and its average temperature and velocity. More an original method for the estimation of the permanence time of Ps inside the nanochannels was shown. All the described methods of analysis are not dependent on the fitting function and on the behavior of the data, but they can be applied for all Time of Flight measurements, when the energy distribution, the average temperature and velocity or the permanence time are investigated.

Thanks to the preliminary data, even if they were affected by a large error, it was possible to discover that at 7 keV of positron implantation energy, a fraction of Ps, formed deeper, undergoes a suitable number of collisions to be emitted thermalized in vacuum. More it was found that the temperature of the cool distribution  $T_1$  is strictly related to the sample temperature and in the Munich measurements, an increment on the fraction of Ps with velocity lower than  $5 \cdot 10^4$  m/s was observed emitted by the sample with nanochannel  $\varnothing = 10-16$  nm kept at 50 K. Despite the results are based on fit procedure with no so great statistical significance, these measurements were the first results as obtained with sample at very low temperature and they were however important for testing the apparatus and for discovering the optimal working conditions.

The measurement conditions in terms of beam intensity, background radiations and detectors efficiency were studied in order to find how a spectrum could be obtained with an acceptable error on the data. For this reason an analysis of the dependence of the signal to background ratio on the beam intensity and on the system parameters was performed and the experimental errors as resulted in Trento and in Munich spectra were justified by an analytical description.

More the preliminary measurements were used to estimate the expected Ps time of flight annihilations distribution, in order to discover the measurement conditions which optimize the error on the data in each part of the spectrum. The best condition for the ToF set up, where channeltrons with high efficiency and with low dark counts were considered, was analyzed for different beam intensities and acquisition times, taking into account the background radiation of the reactor hall as well.

As conclusion, with the best conditions for the ToF set up, the optimized beam intensity of  $\sim 7 - 8 \cdot 10^5 e^+/s$  and the acquisition time of 2 days were found to assure a relative error lower than 0.4 in the region of Ps emitted with velocity higher than  $4 - 5 \cdot 10^4$  m/s and this result could be acceptable for the analysis presented before.

On the contrary for discovering slower Ps with the same relative error of around 0.4, a longer acquisition time should be provided. In this context in fact the limitation on this set up is given by the low probability of observing Ps with very low energy, due to Ps finite lifetime, in comparison with the high probability of detecting other events (as prompt peak events or fast Ps or false events). As said, for very slow Ps, the resolution on the annihilation positions must be increased by changing the shield and slit dimensions with a consequent increase in the acquisition times.

Therefore with the scope of investigate very slow Ps, a different apparatus could be designed, where bunches of positrons instead of continuum beam could be employed. In this case the false counts would not be present and the annihilations distributions as a function of the elapsed time from the positrons arrival on the target, could be acquired by using very fast scintillators. In that spectrum the prompt peak events present at low times of flight could disturb the detection of fast Ps atoms, while the investigation of slow Ps could be considerably improved.

In principle the apparatus described in the next chapter, designed and built for delivering bunches of  $10^7$  positrons with the desired kinetic energy within 5 ns, could be useful also for this kind of measurements. Anyway it was designed for Ps spectroscopy measurement and in the next chapter a description of the apparatus and of the future measurements of Ps excitations in Rydberg states and Ps spectroscopy will be presented.



## Chapter 4

# Apparatus for Ps excitations in Rydberg states

A new apparatus was designed and built in order to receive positron bunches sent from the two-stage Surko accumulator of AEgIS and to re-bunch and focus them on a porous target to form Ps in vacuum and to study its excitation in Rydberg states. The apparatus connected to the Surko trap, through the first section of the AEgIS positrons transfer line, was designed to accelerate positrons to high energy (from 6-9 keV) and to obtain positrons pulses at the target with a time duration of 5 ns, inside a spot of 3 mm in diameter. These goals were realized compatibly with the constraint of keeping the target in a free field region, allowing one to produce Ps in absence of electric and magnetic field for spectroscopy measurements.

In the first part of this chapter an overview of the accumulator will be presented, in order to define the initial positron bunch conditions which were considered for the design and the simulations of the following magnetic and electron optical lines. Then a detailed description of the different parts of the bunching system will be explained, with taking particular attention to the electrical requirements for the bunches compression. At the end the electrical tests performed on the electrodes will be described and the first results will be reported.

### 4.1 AEgIS accumulator

The Surko-type accumulator is a commercial systems which is able to produce high density bunches by trapping positrons spontaneously emitted by radioactive source. The  $^{22}\text{Na}$  AEgIS positron source has a diameter of  $\sim 4$  mm and an activity of 50 mCi and it is installed on a metallic support cooled

down to 5.5 K via an helium-pump refrigerator. In front of the source, a parabolic cone is installed, on whose surface a thin layer of frozen Neon is grown at a temperature of 7 K with ultra-pure Neon admitted at a pressure of  $10^{-4}$  mbar for a few minutes. After positrons are moderated in the thin Ne layer, a slow positron beam of  $\sim 9 \cdot 10^6$   $e^+$ /s with an energy of 2-6 eV is obtained, due to the moderation efficiency of 0.1% with respect to the total amount of positrons emitted by the radioactive source. The positron beam is then conducted by a magnetic transport line into the positron trap to be cooled down and accumulated in order to form a non-neutral plasma. The principle used for cooling the positrons is the energy loss by random collisions with inert gas molecules (ordinary molecular nitrogen  $N_2$ ). As it can be seen in Fig.4.1, the trap is divided in three regions with decreasing pressure of cooling gas and decreasing electric potential. Once trapped the positrons

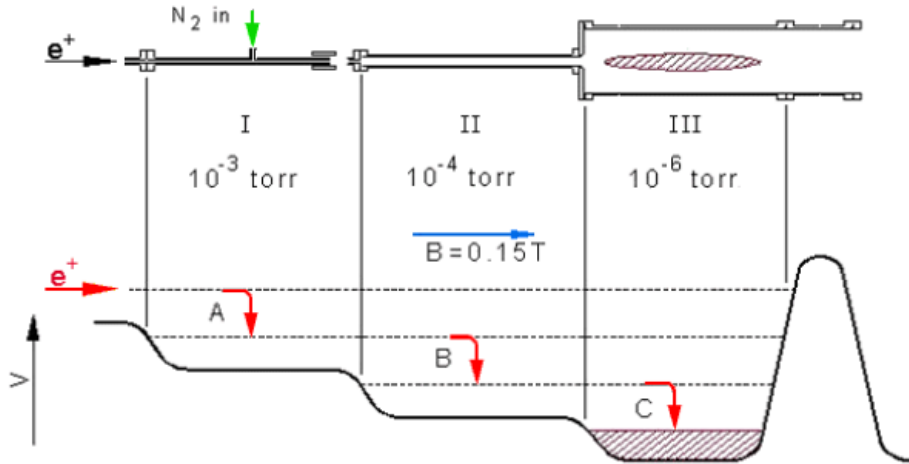


Figure 4.1: Scheme of the positrons Surko-trap. The three regions with different pressure and diameter are visible. The different electric potentials allow cold positrons to be confined in the biggest trap on the right.

continue to lose their energy in collisions with the gas, finally residing in the potential well, formed by the voltages applied to the large diameter trap electrodes. A segmented electrode, split into six segments, is used to apply a time varying electric field. This technique, called *rotating wall*, is used to compress positrons. Performing the procedure described in [90], up to  $2 \cdot 10^8$  positrons can be trapped in the accumulator in a few hundreds seconds and there accumulated with a storage time of hundreds of seconds. After the period of accumulation, the accumulator can be rapidly opened and the positrons are extracted by applying a linear-increasing potential difference

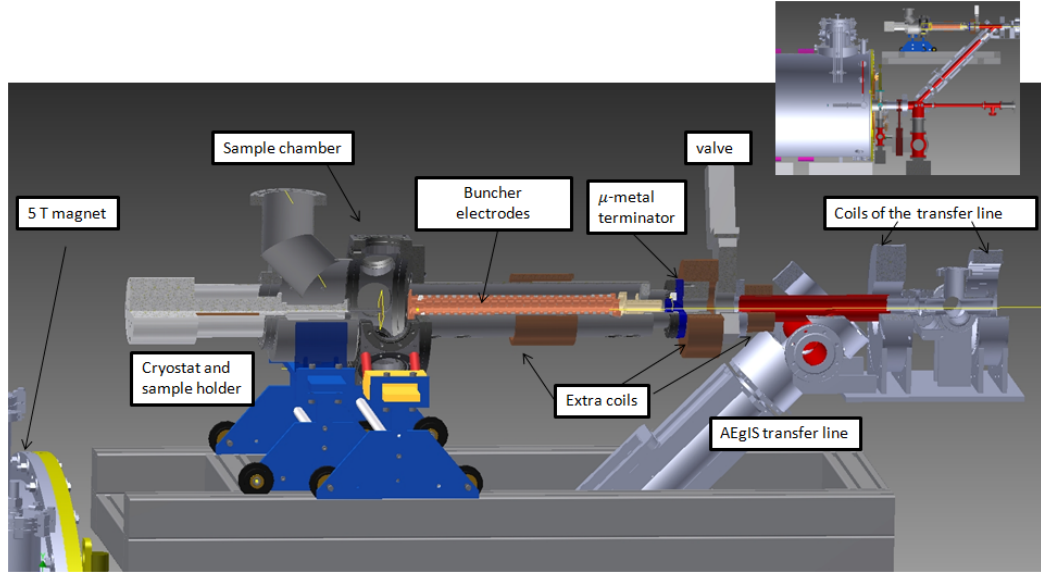


Figure 4.2: Overview of the apparatus. The valve connects the bunching system with the AEgIS transfer line. The magnetic transport line given by the two coils of the AEgIS transfer line and by the extra coils is shown. The electrodes forming the electron optical line are visible inside the tube, which connects the valve to the sample chamber. The cryostat, the sample holder are indicated together with the turbo-pump connection on the top of the apparatus. The AEgIS 5 Tesla magnet, placed not far from the sample chamber, is also visible at the left corner. In the inset at the right corner the new apparatus embedded in the AEgIS set up is shown.

on the last electrodes. In this way every 50-100 s,  $10^7$  positrons are delivered in a bunch of 5 ns duration, with a diameter of 1 mm and with a kinetic energy in the 40-400 eV range. At the exit of the accumulator a magnetic line of 1000 G (called *AEgIS transfer line*), matched with magnetic field of the accumulator is provided for transporting positrons from the accumulator to the AEgIS main magnets. Before the first bend of the transfer line, a valve was placed to connect the positron line to the new apparatus (see Fig.4.2).

## 4.2 Overview of the apparatus

The apparatus is mainly formed by a magnetic line, which transports positrons beyond the valve and by an electron optical line, given by many electrodes, which transports, compresses, focuses and accelerates positrons up to the target. In the following a description of the two lines, which are overlapping

in the region immediately after the valve will be reported. The simulations, as obtained by using COMSOL Multiphysics software for the magnetic field and by using SIMION software for the ray-tracing of positrons trajectories, will be shown.

### 4.2.1 The magnetic transport line

In order to transport positrons beyond the valve a guiding magnetic field was designed by using two coils of the AEgIS transfer line and by adding two other extra coils built on the two sides of the valve (see Fig.4.3). Since no magnetic field is request in the target zone, the transport line was designed in such a way that the magnetic flux density is slowly reduced along the path and then deviated after the valve, when an electrostatic transport is also provided. Thus the currents of the four coils are set in order to match the 0.1 T magnetic field of the accumulator and to decrease it in adiabatic way down to 600 G at the entrance of the Ps apparatus. In the narrow space before the valve, only a small coil can be built and only a magnetic field of 100 G can be generated, while a higher magnetic field (600 G) is produced by a bigger coil after the valve. The magnetic field produced by the four coils was simulated by using the COMSOL Multiphysics software [80]. In Fig.4.3 the magnetic flux density distribution is shown, where the adiabatic reduction from 1000 G to 600 G is visible.

Due to the invariance of the positron magnetic moment  $\mu$  given by [91]:

$$\mu \equiv \frac{mv_{\perp}^2}{2B} = \frac{E_{\perp}}{B} \quad (4.1)$$

where  $v_{\perp}$  and  $E_{\perp}$  are the velocity and the kinetic energy perpendicular to the magnetic field direction, respectively, an increment in the parallel kinetic energy is caused by a decrease in the magnetic flux density and vice versa. From COMSOL simulations a set of values, corresponding to the simulated magnetic field on a mesh, were extracted and were used for the ray-tracing simulations performed by the SIMION software [83]. A positron bunch starting inside the first coil of AEgIS transfer line with a diameter of 1 mm was simulated with the initial condition expressed in Tab.4.1 in agreement with the output conditions from the Surko-trap. The initial positrons kinetic energy was optimized by considering the release from the magnetic field, the time compression and the final focusing; as it will be explained in the following, an optimum value centered around 80 eV was finally chosen for the positrons energy.

A plot of the beam diameter as resulted along the magnetic line is shown in Fig.4.3, where it appears to be  $\sim 1.5$  mm inside the last coil of 600 G,

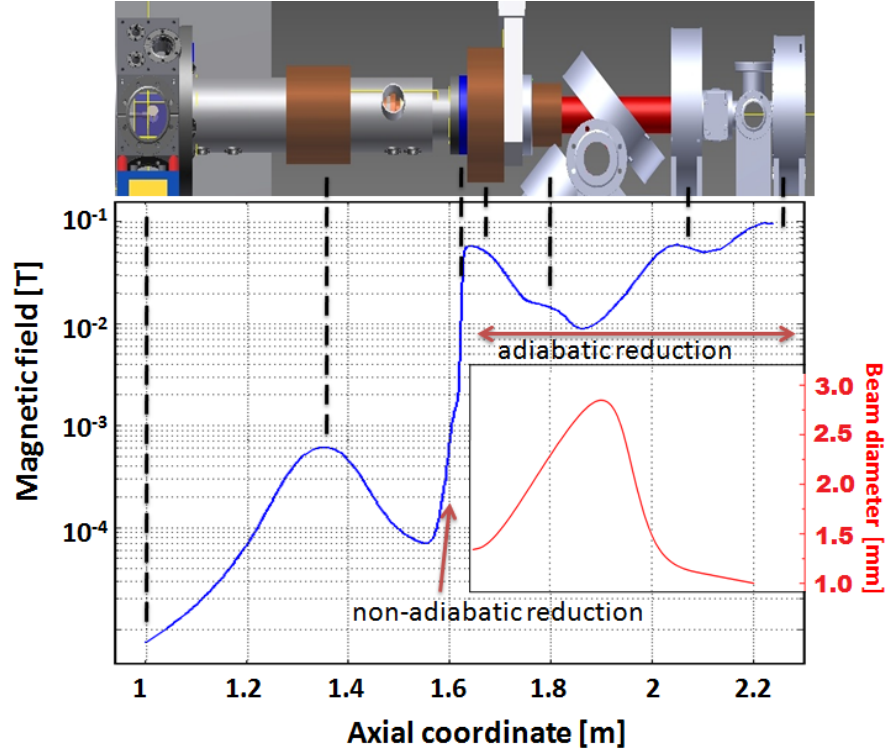


Figure 4.3: Plot of magnetic flux density as a function of the axial coordinate. The target position corresponds to  $z = 1$  m, while the  $\mu$ -metal electrode position corresponds to  $z = 1.62$  m. The adiabatic reduction from 0.1 T to 0.06 T together with the non-adiabatic change near the  $\mu$ -metal electrode are evident. In the red plot the beam dimension along the magnetic line, as obtained by SIMION simulation, where the magnetic field as extracted by COMSOL simulations was considered, is shown. The increment on the beam dimensions in correspondence of the lower magnetic flux density is visible. As result, at the end of the magnetic line, just before the  $\mu$ -metal electrode, a beam diameter of 1.5 mm is achieved.

|                    | Initial condition |
|--------------------|-------------------|
| diameter           | 0.9 mm            |
| angle              | 0-10 °            |
| pulse duration     | 5 ns              |
| Min kinetic energy | 73 eV             |
| Max kinetic energy | 88 eV             |

Table 4.1: Table of initial conditions of positron bunches used for SIMION simulation. The positron bunches are considered to start inside the first coil of the AEgIS transfer line (see Fig.4.3).

before the end of the magnetic line near the entrance in the first electrode. Moreover in the same position the bunch time duration was found to be  $\sim 15$  ns, if initially it is 5 ns, due to the positrons energy spread of 1 % at the exit of the accumulator, and to the distance of  $\sim 60$  cm between the exit of the Surko-trap and the valve connection.

Reached the entrance of the new apparatus, positrons are released by the magnetic field, thanks to the same technique adopted in the ToF apparatus (see section 3.2). In fact a  $\mu$ -metal electrode, with the same shape described in section 3.2, was connected to a  $\mu$ -metal flange and placed close to the last coil of the magnetic line in such a way that a non-adiabatic magnetic field reduction from 600 G to less than 20 G in 5 mm was achieved. In the Fig. 4.3, where the magnetic flux density as a function of the axial coordinate is reported, the non-adiabatic change in correspondence with the  $\mu$ -metal electrode position at  $z = 1.62$  m, is shown. The electrode, being very close (1 mm) to the a  $\mu$ -metal flange connected through a collar to a  $\mu$ -metal shield, transports the magnetic streamline out from the chamber (see Fig.4.4) and thus deflects the magnetic field, acting as a magnetic terminator.

As in the ToF apparatus, in the same region where the magnetic field drops, the  $\mu$ -metal electrode together with the nearer electrode, produce an electric field which accelerates positrons along the direction pointing to the sample. The potential values of the two electrodes are -500 V and -1500 V, respectively and they were optimized in order to extract positrons trajectories from the magnetic streamline and to make them as parallel as possible to the axis.

However due to the higher intensity of the magnetic field with respect to the ToF apparatus, a further coil placed between the magnetic terminator and the sample chamber was necessary for producing a low magnetic field of around 10 G. In this way the magnetic field streamlines near the  $\mu$ -metal zone can be divided into two groups: the farer ones from the axis, which are deflected and carried into the  $\mu$ -metal electrode and the closer ones to the

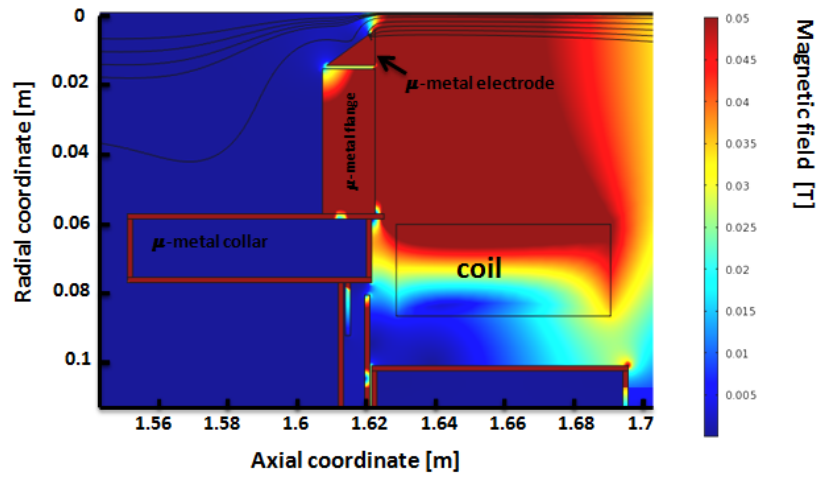


Figure 4.4: Magnetic field streamline and magnetic flux density in the region near the  $\mu$ -metal electrode as obtained by COMSOL simulation. In the figure the positrons fly from the right to the left. The cross sectional view of the coil, of the  $\mu$ -metal electrode and of a part of the external  $\mu$ -metal shield connected to the flange through a collar (see section 4.4.1) is visible. The magnetic field streamlines can be divided into two groups: the farer ones from the axis, which are deviated and carried into  $\mu$ -metal electrode and the closer ones to the axis, which are kept parallel to it by the further coil placed after the  $\mu$ -metal electrode (not visible in figure).

axis, which are kept parallel by the new small coil (see Fig.4.4). This further magnetic field only helps the low energy positrons to reach the buncher system with a smaller angle, along more parallel trajectories and in this way an optimization on the final beam spot found on the sample was performed. In any case the further coil put after the magnetic termination, being far from the sample and generating magnetic field of low intensity, does not affect the sample zone.

### 4.2.2 The electron optical line

After being released from magnetic field the positrons trajectories are mainly related to the electric field generated by the optical line. This line is shown

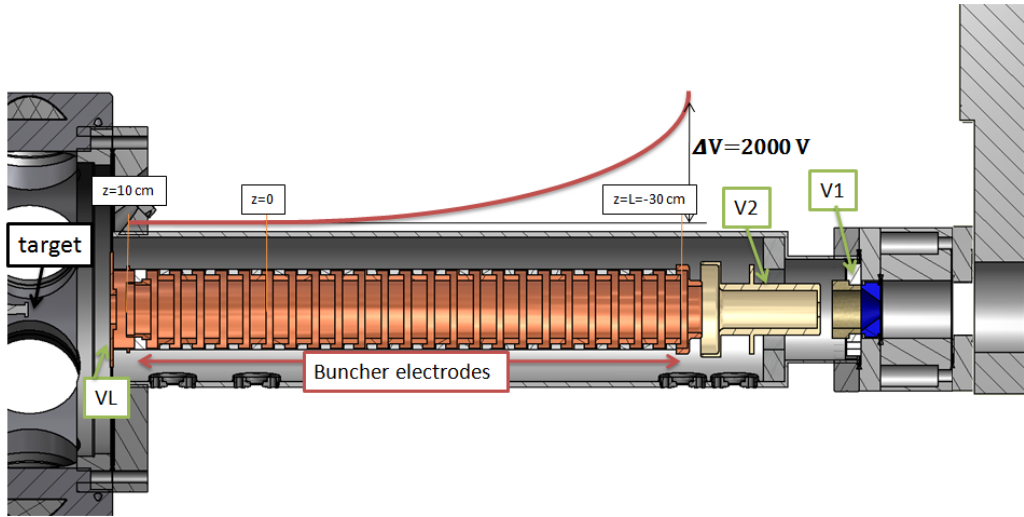


Figure 4.5: The cross sectional view of the electron optical line which extracts positrons from the magnetic field, and which transports, compresses and focuses positron bunches on the sample. The  $\mu$ -metal electrode (V1), the second electrode (V2) and the series of 25 buncher electrodes are visible. The last electrode (VL) necessary for the final focusing is also indicated. The switching parabolic and constant potential applied to the buncher electrodes for the bunch time compression is shown. The  $z$  axis with the reference point at  $z = 0$  is also indicated.

in Fig.4.5 and it is composed by:

- the  $\mu$ -metal electrode, above discussed (called V1 from now on),
- the second electrode set at potential of -1500 V, just presented (called V2 from now on),

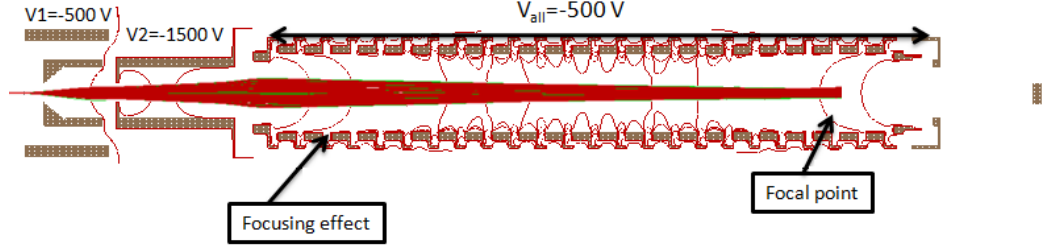


Figure 4.6: The equipotential lines of the electron optical line as obtained with all buncher electrodes set at the potential  $V_{all}$  are shown. The two lenses necessary for the magnetic field extraction (given by  $V1$  and  $V2$  values) and for the transport inside the buncher (related to  $V_{all}$  value) are indicated. The ray-tracing simulation performed with these statical potentials is also visible. The focal point in the last part of the buncher is evidenced.

- 25 electrodes, which constitute the buncher device,
- the final electrode (called VL from now on), which, together with the electric field generated by the sample, focuses positrons on the target.

While the first electrodes ( $V1, V2$ ) are powered by a statical potential, the buncher electrodes are subjected to time dependent potentials, which spatio-temporally compress positron bunches, after all positrons are inside the buncher. The positrons transport in the electron optical line is mainly assured by two lenses: the first one is given by electrode  $V2$  and the buncher electrodes set at the same potential  $V_{all}$ , the second one given by the last electrode (VL) and by the sample kept at ground potential.

In order to transport positron inside the buncher device, before the time compression, the focal point of the first lens ( $V2-V_{all}$ ), determined mainly by  $V_{all}$  potential value, was chosen in the second half of the buncher device, so that the positrons trajectories are made as parallel as possible. The final value of  $V_{all}$  potential was optimized by SIMION simulations and it resulted to be -500 V. The SIMION simulation performed with all buncher electrodes set at the potential  $V_{all}$  is shown in Fig.4.6. The focal point chosen in the final part of the buncher system is visible from both the equipotential lines, evidencing the lens effect and from the ray-tracing simulation. From the simulation, the positrons properties at the entrance of the first buncher electrode were extracted. In particular the bunch has 16 ns duration and it is distributed along the axial coordinate for 18 cm, with a diameter of 6 mm and a kinetic energy ranging from 540 eV to 700 eV.

Once all positrons are inside the buncher device, a time-dependent potential can be applied to the 25 electrodes, in order to reduce the bunch duration and to give positrons the desired kinetic energy. In the two next sections the methods for achieving these two goals will be explained.

### Time compression of positron bunches

Because of the Liouville's theorem, the only way to change the bunch time duration is distorting the phase space volume occupied by the ensemble of positrons [92]. Since the positron bunch is distributed inside many electrodes, a different electric field can be generated along the bunch length, in order to accelerate more the last positrons with respect to the first ones. The idea can be realized by setting different potentials to each electrode in such a way that a higher electric field is generated on the region occupied by the last positrons. If  $z$  is the axial coordinate, the potential  $V_i$  applied to the electrode  $i$ , placed at the  $z_i$  coordinate, has to be set as a function of the axial coordinate  $V_i(z_i)$ .

In particular if a harmonic oscillator potential is considered:

$$V(z) = \frac{kz^2}{2} \quad (4.2)$$

it can be demonstrated that all positrons, starting at  $z < 0$  with initially low velocity, arrive almost simultaneously at  $z = 0$  independently from their initial positions. In fact by observing the differential equation for particles in a harmonic oscillator potential:

$$m\ddot{z} = -kqz \quad (4.3)$$

where the  $k$  parameter can be written as  $k = \frac{2\Delta V}{L^2}$  and it is related to the amplitude of the parabolic function  $\Delta V = V(z = -L) - V(z = 0)$  defined as the difference between the maximum potential value  $V(z = L)$  set on the electrode placed at  $z = L$  and the minimum potential value  $V(z = 0)$  correspondent to the electrode placed at  $z = 0$  (see Fig.4.5), the general solution can be calculated:

$$z(t) = z(0)\cos(\omega t) + \frac{\dot{z}(0)}{\omega}\sin(\omega t) \quad (4.4)$$

with  $\omega = \sqrt{\frac{qk}{m}}$  and  $z(0)$  and  $\dot{z}(0)$  the initial position and velocity, respectively. Thus the time  $t_f$ , spent by a positron to reach the coordinate  $z = 0$ , can be obtained as a function of the initial conditions  $z(0)$  and  $\dot{z}(0)$  and on the  $k$

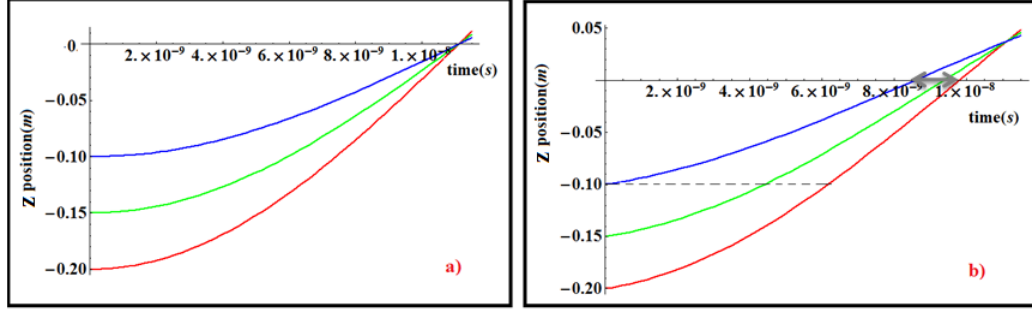


Figure 4.7: a) Solution  $z(t)$  of the equation of motion (eq.4.4) for particles starting with  $\dot{z}(0) = 0$  in three different positions  $z(0) = -0.10$  m,  $z(0) = -0.15$  m,  $z(0) = -0.20$  m. They take the same time  $t_f = 1.2 \cdot 10^{-8}$  s to reach the end of the buncher at  $z = 0$ .

b) Solution  $z(t)$  of the equation of motion (eq.4.4) for particles starting with  $\dot{z}(0) = 1.3 \cdot 10^7$  m/s from the same position of the plot a). In this case they arrive at different times at the coordinate  $z = 0$ , but they are temporally compress with respect to the initial bunch duration. The dashed line indicates a rough approximation of the initial bunch duration ( $\sim 6$  ns), while the grey arrow, shorter than the dashed line, indicates the bunch duration at the  $z = 0$  position ( $\sim 1$  ns). In both the case a  $\Delta V = 2000$  V, correspondent to  $k = 44000$  if  $L = -0.3$  m, was considered.

parameter. In particular as shown in Fig.4.7a, if the initial particles velocity is zero, the results is  $t_f = \frac{\pi}{2\omega}$ , not depending on the initial position  $z(0)$  and all positrons arrive at the same time at  $z = 0$  coordinate. On the contrary if particles have an initial velocity  $\dot{z}(0) \neq 0$  the solution  $t_f$  is more complicate and it depends on both the initial velocity and the initial position. As shown in Fig.4.7b, where the solution  $z(t)$  was obtained for particles starting at different positions and with the same velocity, the positrons arrive at  $z = 0$  more temporally compressed with respect to the starting condition (see grey arrow in comparison with the dashed line), but not in the same instant as in Fig.4.7a.

Therefore an ensemble of particles, uniformly distributed in 18 cm length with different velocities, as resulted after the electrostatical transport inside the buncher, can be considered and the compression effect by using a buncher device 30 cm long ( $L = -0.3$  m) and by applying a parabolic potential of  $\Delta V = 2000$  V ( $k = 44000$ ) can be investigated. Since  $z = 0$  is the coordinate relative to the minimum of the parabolic function, the positrons inside the bunch are uniformly distributed between  $z = -0.3$  m and  $z = -0.12$  m and in order to study the time compression, taking into account the energy spread

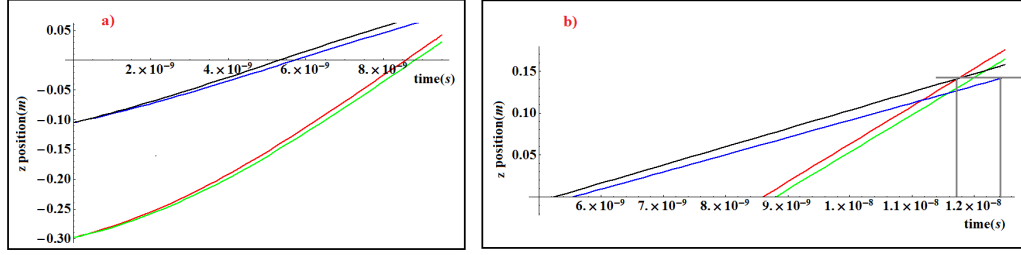


Figure 4.8: a) Solution  $z(t)$  of the equation of motion (eq.4.4) of four particles starting at  $z = -0.3$  m with the highest and the lowest velocity ( $v = 1.6 \cdot 10^7$  m/s (red) and  $v = 1.4 \cdot 10^7$  m/s (green)) or starting at the shortest coordinate  $z = -0.1$  m, with the same highest (blue) and lowest (black) velocity. As before  $k = 44000$  was considered. From this plot the times  $t_f$  spent by each positron to reach the  $z = 0$  point, were extracted. b)  $z(t)$  solution as obtained by considering an uniform motion in a free field region for particles starting at the time  $t_f$  and with velocities  $\dot{z}(t_f)$  extracted from Fig.4.9.

of the particles, the motion of four positrons was calculated. The solution  $z(t)$  of the equation of motion (eq.4.4) for two particles starting at  $z = -0.3$  m and for other two particles starting at  $z = -0.12$  m with the highest and the lowest velocity ( $v = 1.6 \cdot 10^7$  m/s and  $v = 1.4 \cdot 10^7$  m/s correspondent to 540 eV and 700 eV, as extracted by SIMION simulations) are shown in Fig.4.8 a. From the figure it is possible to extract the bunch time duration of 3.5 ns at the coordinate  $z = 0$  m. Moreover the velocities  $\dot{z}(t)$  of the four positrons can be extracted: their values as a function of time are plotted in Fig.4.9. From this plot it appears that due to the parabolic potential the final velocity depends more on the starting position than the initial velocity. By looking at Fig.4.9, the following conclusion can be obtained: at the  $z = 0$  coordinate, the positrons starting from  $z = -0.3$  m (red and green in Fig.4.8 or in 4.9), that are the last in the bunch, have higher velocity than the first positrons started at  $z = -0.12$  m, which arrive before at  $z = 0$ , but with a lower velocity.

Therefore an easy improvement on the final bunch duration can be achieved by producing a free electric field region after the parabolic function region, where positrons can fly in an uniform motion. Considering this configuration, an estimation of the final bunch duration was obtained by calculating the solution  $z(t)$  for particles, starting with velocities  $\dot{z}(t_f)$  at the time  $t_f$  and undergoing an uniform motion in a free field region. By extracting from the plot of Fig.4.8 the time  $t_f$  spent by each positron to reach  $z = 0$  and

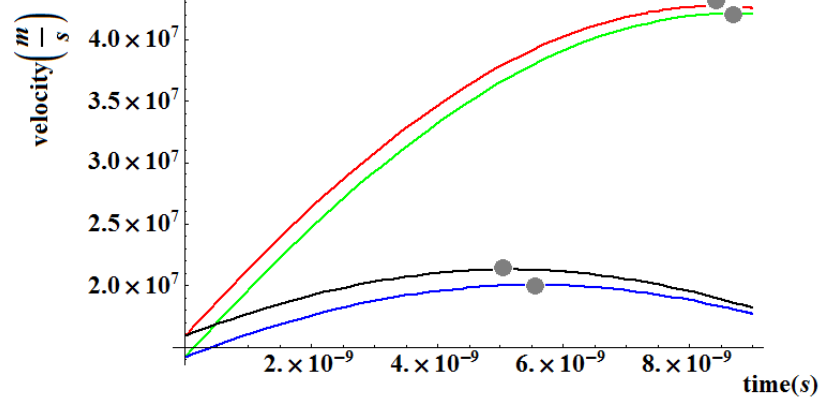


Figure 4.9: Velocity  $\dot{z}(t)$  as a function of the time as extracted by the solution  $z(t)$  of the four positrons considered in the plot of Fig.4.8a. The red and the green plots correspond to particles started with different velocities at  $z = -0.12$  m, while the blue and the black ones correspond to particles started at  $z = -0.3$  m. Due to the parabolic potential the final velocity depends more on the starting position than the initial velocity.

the relative velocity  $\dot{z}(t_f)$ , from Fig.4.9, the plot b) of Fig.4.8 was obtained, where a minimum pulse duration of less than 1 ns occurs at the coordinate  $z = 0.14$  m. From calculations done with different values of the  $k$  parameter and for different initial velocities, the final time compression was found to be smaller for higher  $k$  parameter and thus for higher amplitude of the parabolic function. It was also found that the final bunch duration is slightly influenced by an increase of the initial velocity, provided that the region with constant potential is made longer. Balancing these parameters with the final goal of obtaining a bunch duration of 5 ns at the sample position, an optimization on the  $k$  parameter and on the choice of the initial positrons velocity in the range 40-400 eV (see section 4.1) was done; the best result was found with  $k = 44000$ , correspondent to  $\Delta V = 2000$  V and with an initial positron kinetic energy  $E = 80$  eV.

In conclusion in order to minimize the bunch duration, the potential applied to the buncher electrodes must be divided in two parts: in the first one 0.3 m long the parabolic function of amplitude 2000 V has to be applied after the positrons entrance, while in the following length of 0.14 m zero electric field has to be applied. The buncher system was designed with 25 electrodes, each 11 mm long separated by 5 mm, such as the parabolic zone (30 cm length) is bigger than the initial pulse length and the final zone (10

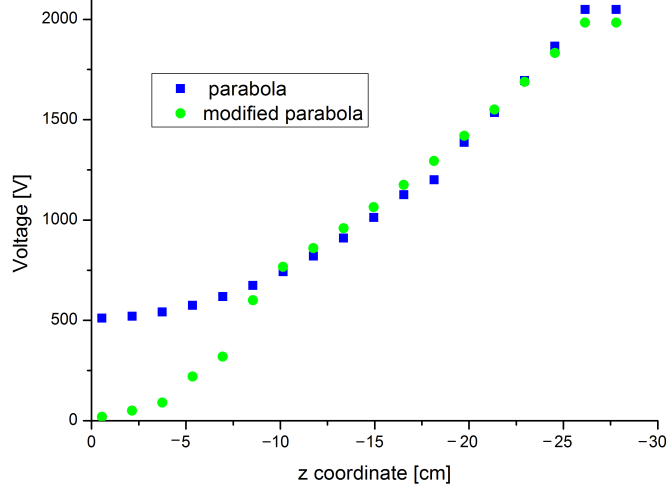


Figure 4.10: Theoretical parabola (in blue) and modified parabola (in green). The voltage values as a function of the axial coordinate, correspondent to the electrodes position, are shown. The theoretical parabola was modified during the optimization of the positrons transport performed by the SIMION simulations.

cm length) is in agreement with the results shown in Fig.4.8 b.

This final configuration was simulated with SIMION software in order to check the positrons losses on the electrodes walls. In fact, differently from [93], where a guiding magnetic field securing the positrons transport was present, the electric field has both to compress and to transport positrons. By SIMION simulations the ray tracing of the particles inside the buncher was studied and the parabolic potential applied to the electrodes was modified and optimized, in order to transport all particles to the target. The parabolic function was empirically modified, by inserting a steeper decrement on the potential values in the middle of the parabolic section, for producing a sort of focusing lens. The optimized potentials and the original parabolic potentials are shown in Fig.4.10, where the maximum value of  $V_{max} = 1983$  V is applied to the first two electrodes and  $V_{min} = 0$  V (at  $z = 0$  coordinate) is applied on the 19th electrode. Moreover the SIMION simulations confirm the compression effect calculated before, and they show that, with the modified parabola, transport efficiency of almost 100% is achieved. The time duration resulted from the final simulations is a little worse than that calculated before because it results around 4 ns inside the last electrode of

the buncher; however this value is compatible with our request.

Finally a consideration on the rise time of the switching potential has to be done, because it has to be small with respect to the time elapsed by positrons inside the system. In particular if the average positrons velocity is around  $1.5 \cdot 10^7$  m/s a rise time of maximum 5 ns can be tolerated, because in this time positrons travel for 4.5 cm (three electrodes length), that is a small length with respect to the total buncher length. This value was confirmed by simulations, in which since the rise time was fixed to 1 ns, the switch on time of the potentials was delayed up to 6 ns, without having any loss of positrons. More the crossing time of the positrons inside the 25 electrodes resulted of the order of 20 ns and this value was taken as reference for the duration of the pulse, which must be applied to form the almost-parabolic function.

### Offset voltage for accelerating positrons to 6-9 keV

The second goal, realized with the buncher electrodes, is to tune the positrons kinetic energy, allowing the target to be kept at ground potential in order to form Ps in free electric field. Since also the positrons source is set at ground and because positrons come from the accumulator with low energy ( $< 400$  eV), it is necessary to increment the potential on the buncher electrodes, before positrons exit the buncher. In fact, as long as they are inside the buncher they are only subjected to the internal electric field, which is given by the different potentials of the adjacent electrodes. Thus if all electrodes are simultaneously elevated to a high voltage, the positrons will be accelerated only between the last electrode and the target. In this way the final energy (6-9 keV) can be set by changing the amplitude of an offset voltage (4-6.5 kV) applied to all the electrodes of the buncher for producing a high electric field between last electrodes of the buncher and the sample.

As regards the rise time of the offset, it has to be lower than the time spent by positrons into the buncher electrodes, that is around 20 ns. In conclusion the electrodes potentials have to be switched from  $V_{all} = -500$  V to the value  $V(z_i)$  different for each electrode, added to an offset of some kilovolts. The electronic circuit necessary to realize this potential configuration will be described in section 4.3.

#### 4.2.3 The final focusing on the target

Between the target and the last buncher electrode a further electrode is placed to focus positrons on the sample surface in a spot smaller than 3 mm diameter. This electrode (VL) is 56 mm far from the sample with an

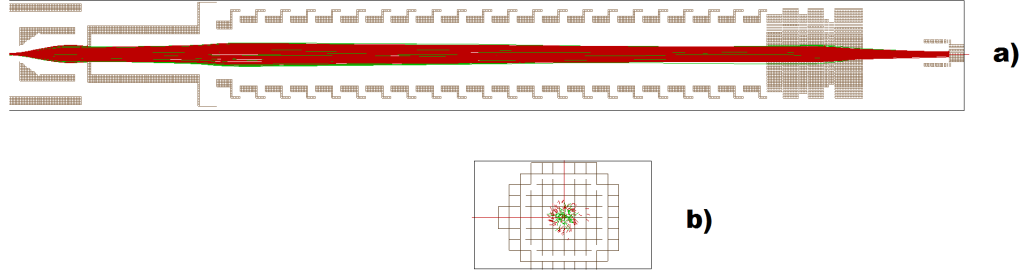


Figure 4.11: a) Ray-tracing simulation of positrons inside the buncher system and through the last lens. Two groups of positrons starting inside the first coil of AEgIS transfer line with different energies (73-88 eV) and with the initial condition, as expressed in Tab.4.2 were simulated.

b) Spot of  $\sim 2.5$  mm diameter achieved at the target position.

The final results extracted from this simulation are reported in the Tab.4.2. The simulation took into account the presence of the two grids kept at ground potential placed near the sample position. For this reason the simulation was performed by using two different geometrical sketches, matched together near the last buncher electrodes: one was drawn in a 2D plane with cylindrically symmetry and the other one in the 3D space.

inner diameter smaller than the buncher electrodes and its voltage potential was optimized for the final focusing. Because of the initial positrons energy spread (see Tab.4.1) and because the focal point depends on the positrons energy, the voltage value of the last electrode was chosen by simulating two groups of particles with 73 eV and 88 eV energy and minimizing the spot dimension at the sample position. In the Fig.4.11a the overall optimized ray-tracing simulation in the electron optical line is reported. In Fig.4.11b the sample surface as resulted after the ray-tracing simulation is shown, where a spot of  $\sim 2.5$  mm diameter was achieved thanks to a potential of 500 V, applied to the last electrode.

Since two grids will be necessary to ionize the excited Ps in some planned measurements (see chapter 5), the final simulation was performed by taking into account the potential generated by them, set at ground potential like in the measurement conditions. By introducing the two grids, the cylindrical symmetry of the simulated system was lost in the region near the sample; for this reason the simulation was performed by using two different geometrical sketches, matched together near the last buncher electrodes. One sketch was drawn in a 2D plane because of the cylindrical symmetry, while the other one was drawn in the 3D space (see Fig.4.11).

|                    | Initial condition | Final condition |
|--------------------|-------------------|-----------------|
| diameter           | 0.9 mm            | 2.5 mm          |
| angle              | 0-10 °            |                 |
| pulse duration     | 5 ns              | 5 ns            |
| Min kinetic energy | 73 eV             | 7.2 keV         |
| Max kinetic energy | 88 eV             | 8.7 keV         |

Table 4.2: Table of initial and final conditions of positron bunches. The reported final kinetic energy was obtained by considering an offset potential of 6300 V.

As regards the final bunch duration at the target, despite the presence of a further acceleration necessary for the final focusing, a positrons pulse of 5 ns was found at the target position. The final characteristics of positron bunches at the target position are shown in the Tab.4.2, from which it can be concluded that the described system allows positrons to arrive to the target, kept at ground potential, with an energy, a spot and a time spread, compatible with the planned Ps spectroscopy measurements, that will be described in chapter 5.

In the next section the design of the electronic circuit to drive the potential of the electrodes of the buncher will be illustrated.

### 4.3 The electronic design

The electronic circuit was designed in order to change the buncher electrodes potential from the initial value of -500 V to the parabolic voltage over-imposed on the offset voltage. The idea was to use a single switch and to realize the offset and the parabolic function with proper resistors, used in a voltage divider configuration. The switch connects a HV capacitor bank to the buncher electrodes (see Fig.4.12) and it was designed to supply a signal with an amplitude of  $V_S = 8.5$  kV, which is the sum of the amplitude of the parabola ( $\sim 2kV$ ) and the maximum offset (6.5 kV). Because the voltage drop on the internal switch resistor  $V_{sw} \sim 2kV$ , the HV capacitor bank was chosen in order to have around 11 kV, when charged. Thus, when the switch is closed on the buncher electrodes, a discharge of the HV capacitor on the buncher resistors occurs and for this reason the capacitance value  $C$  was chosen to provide a small reduction on the potential value  $V(t)$  during the pulse duration. In fact if the pulse duration is 30 ns long and the discharging time  $\tau$  is chosen to be  $\tau = R_{tot}C = 3000$  ns, where  $R_{tot}$  is the sum of the resistors in the line, only a voltage drop of 0.1% (around 90 V) is expected during

the pulse duration. As it will be explained in the following,  $R_{tot} = 75 \Omega$  was determined by the impedance matching constrain and thus a value of 40 nF was chosen for the capacity  $C$ .

The parabolic function is realized by applying 2 kV between the first and the 19th electrode and by connecting the buncher electrodes in pair through proper resistors, in such a way that the series of the voltages drops generate the desired function. In Tab.4.3 the potential of each electrode and the values of the resistors put between adjacent electrodes, are reported. The equivalent resistor of the buncher system is the sum of the 18 resistors,

| N. electrode | Potential [V] | Potential drop [V] | Resistance value [ $\Omega$ ] |
|--------------|---------------|--------------------|-------------------------------|
| 1-2          | 1983          | 150                | 1.2037                        |
| 3            | 1833          | 144                | 1.153                         |
| 4            | 1689          | 138                | 1.107                         |
| 5            | 1551          | 131                | 1.050                         |
| 6            | 1420          | 125                | 1.000                         |
| 7            | 1295          | 118                | 0.947                         |
| 8            | 1175          | 112                | 0.897                         |
| 9            | 1065          | 106                | 0.850                         |
| 10           | 959           | 99                 | 0.793                         |
| 11           | 860           | 93                 | 0.747                         |
| 12           | 767           | 166                | 1.330                         |
| 13           | 600           | 280                | 2.245                         |
| 14           | 320           | 100                | 0.800                         |
| 15           | 220           | 130                | 1.043                         |
| 16           | 90            | 40                 | 0.320                         |
| 17           | 50            | 30                 | 0.240                         |
| 18           | 20            | 20                 | 0.160                         |

Table 4.3: Table of the potentials applied to each buncher electrode and of the values of the corresponding resistors connecting adjacent electrodes. The correspondent potential drops between adjacent electrodes, for generating the almost-parabolic function, are also reported. The electrodes are numbered from right to left in Fig.4.5.

which is  $R_b \approx 16 \Omega$ . To produce the offset voltage a further resistor,  $R_f$  in Fig.4.12, is placed between the last buncher electrode and ground, in order to create a voltage divider formed by the total buncher resistor  $R_b$  and the resistor  $R_f$ . In this way, by changing the value of  $R_f$ , the offset potential can be varied as desired and thus the positrons energy can be selected.

The complete circuit is schematized as in Fig.4.12, where the buncher

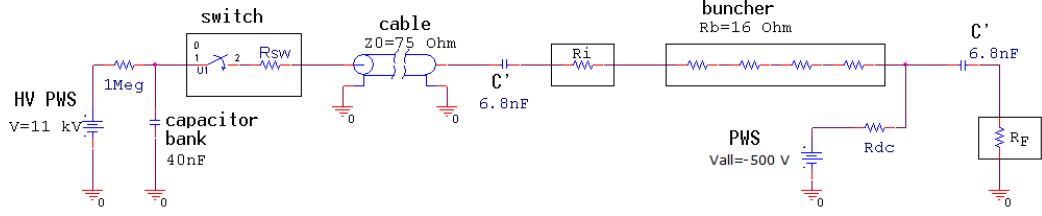


Figure 4.12: Circuit scheme to supply the buncher electrodes. The power supplies of  $-500$  V for the static potential  $V_{all}$ , decoupled from the switching network through two capacitors  $C'$ , is shown. The switching network with the capacitor bank and with the power supply of  $11$  kV for supplying the parabolic function and the offset voltage to the buncher electrodes is also shown.

device is connected through the switch to a capacitor bank and to a high voltage power supply (HV PWS) for the parabolic pulse. More the buncher electrodes are powered by a statical potential of  $V_{all} = -500$  V given by the second power supplies (PWS) that is decoupled from the switching circuit by two capacitors  $C'$ .

At the beginning the switch is open and no voltage (neither offset nor parabolic potentials) is applied between the first and the last buncher electrodes, which are all kept at the static potential; when the switch is closed, the pulse signal propagates inside the buncher, producing at the same time the parabolic function superimposed to the offset. To have a pulse of  $V_S = 8.5$  kV and a fast rise time of  $\sim 4$  ns a switch BEHLKE HTS-12-150-UF, whose inner resistance is  $\sim 15 \Omega$ , was chosen.

Due to the fast rise time with respect to the buncher dimension, the switching circuit was designed by taking into account that the system works as a transmission line, where reflection effects of the propagating signal have to be considered and minimized. At high frequencies in fact, the wavelength of the propagating pulse can be comparable with the circuit size, and for this reason the signal results in different phases at different locations in the circuit. In this condition the quasi-static circuit theory cannot be applied, but the transmission line theory is necessary, where the voltages and the currents depend on time and on the longitudinal coordinate  $z$ . The state variables of such a system are  $v(z, t)$  and  $i(z, t)$  and the main equations of the transmission line theory are described in appendix A. In the apparatus the voltage signal has to propagate from the switch through the cable and the buncher electrodes up to the last resistor at the end of the line. Therefore the chamber and the buncher electrodes behave like a coaxial cable with a

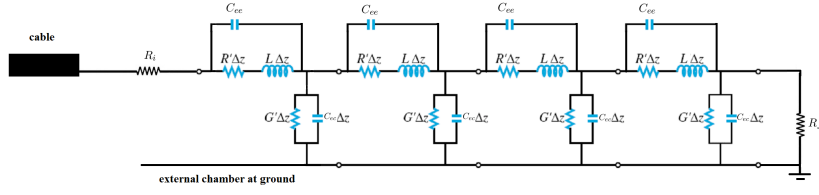


Figure 4.13: Schematic representation of the transmission line of the switching network. The characteristic parameters of the buncher electrodes are shown. The  $R_f$  resistor is necessary for applying the offset voltage to all the electrodes, while  $R_i$  was placed to satisfy the impedance matching condition  $R_{tot} = R_i + R_f + Rb = 75 \Omega$  for whatever  $R_f$  value.

peculiar inner conductor, where the mechanical layout (resistors, chamber and electrodes dimensions and materials) establishes the characteristic parameters of the transmission line. In particular in the buncher system, the electrodes have a capacitance  $C_{ee}$  between each other and a capacitance  $C_{ec}$  between themselves and the outer vacuum chamber. More each resistor for the parabolic function (whose value are reported in Tab.4.3) can be schematized as a resistor in series with an inductance  $L$  while the conductance  $G$ , usually expressed in a transmission line, can be neglected in our case. A schematic representation of the transmission line of the switching network is shown in Fig.4.13, where the capacitance  $C_{ee}$ ,  $C_{ec}$  and the resistor and the inductance are shown.

By solving the wave propagation equation in this system (see Appendix A), the voltage and the current inside the line can be obtained:

$$\begin{cases} v(z, t) = v^+(z - v_{ph}t) + v^-(z + v_{ph}t) \\ i(z, t) = \frac{1}{Z_\infty}v^+(z - v_{ph}t) + \frac{1}{Z_\infty}v^-(z + v_{ph}t) \end{cases} \quad (4.5)$$

where the wave propagation properties inside the line are fixed by the phase velocity  $v_{ph} = \frac{1}{\sqrt{LC}}$ , together with the characteristic impedance  $Z_\infty = \sqrt{\frac{R+i\omega C}{G+i\omega L}}$  ( $\mathcal{R}$ ,  $\mathcal{L}$ ,  $\mathcal{C}$ ,  $\mathcal{G}$  are all values per unit of length).

As explained in Appendix A, the impedance of the buncher system related to the impedance of the cable determines the reflections effects. The values of  $C_{ee}$  and  $C_{ec}$  and  $L$ , that are determined by geometrical properties were optimized in order to improve the system performances, compatibly with the geometrical constrains fixed by the ray tracing simulations discussed before. In particular:

- the capacitances  $C_{ee}$  between electrodes were reduced as much as possible, by reducing the thickness of each electrode. The final value is

$C_{ee} \sim 4 \cdot 10^{-12}$  F and it allows the current to flow mainly in the resistors;

- the capacitances  $C_{ec}$  were reduced by increasing the chamber diameter, in order to increase the propagation velocity of the pulse. The final value was  $C_{ec} \sim 1.5 \cdot 10^{-12}$  F;
- the characteristic impedance of the buncher system  $Z_\infty = 75 \Omega$ , related to the ratio between the diameter of electrodes and chamber (that is the outer conductor of the coaxial line) was matched with the impedance of the rest of the system (see discussion below);
- the inductance  $L$  was reduced as much as possible with a proper choice of resistors design and fabrication ( $L < 2 - 3$  nH) to avoid electric oscillations.

In the Fig.4.14 the behavior of the characteristic parameters of the coaxial cable versus the geometrical dimensions are shown. The plot shows that an increment on the  $D/d$  ratio favours a decrease on the  $C_{ec}$  value, if a higher impedance can be adopted in the circuit.

In our line, as first approximation, the characteristic impedance of the buncher system can be calculated as the sum of the resistors values  $Z_\infty = R_{tot}$  and thus  $R_{tot}$  becomes the parameter that has to be matched with the cable impedance in order to reduce the reflections between the cable and the first buncher electrode. Thus, because the choice on  $Z_{cable} = R_{tot}$  fixes the current value on the line, a cable with  $Z_{cable} = 75 \Omega$  was chosen so that the current flowing into the circuit becomes  $\sim 120$  A if 11.5 kV is applied. Finally in order to fix the impedance independently from each value of the offset voltage, related to the  $R_f$  value, a further resistor  $R_i$  was placed before the buncher itself so that  $Z = R_b + R_f + R_i = 75 \Omega$  for whatever  $R_f$  values. On the other hand this condition imposes that, independently on the value of the offset voltage, since  $R_{tot} = 75 \Omega$ , the pulse amplitude has to be always at the maximum value  $V_s = 8.5$  kV at the output of the switch, in order to produce the desired parabolic function with  $\Delta V = 2000$  V, for whatever  $R_f$ .

In this way the system results to be globally matched with exceptions of some possible reflections that could rise between the switch, whose impedance is not known and the cable or between the cable and the buncher, due to the characteristic impedance of the buncher itself  $Z_b$  that can be related to  $L$ ,  $C_{ee}$ ,  $C_{ec}$  values. Therefore in order to avoid the problem of multiple reflections on the electrodes potential, while the positrons are flying inside the buncher, a time delay was introduced between the switch and the first electrode of the buncher, so that the signal takes more time to cross the line. A

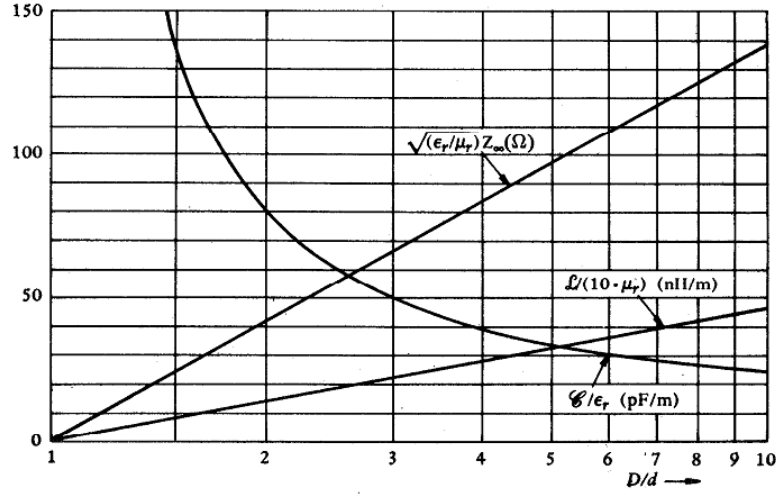


Figure 4.14: Parameters of the coaxial cable versus the ratio between the outer conductor diameter ( $D$  that is the chamber diameter in our line) and the inner conductor diameter  $d$  (that is the electrodes diameter in our line).

time delay  $T_D$  bigger than the pulse duration  $\Delta t = 20$  ns was chosen so that, since the transit time of the reflected signal inside the line is  $2T_D$ , the reflections are disjoint (echo conditions) from the pulse itself and the reflected signal appears on the electrode only when the pulse is off and the positrons have already arrived at the target.

Because the switch has a maximum operating voltage of 12 kV, in order to protect it from high voltage negative reflections, the best solution would be to put at the output of the switch a resistance equal to the characteristic impedance of our network. In this case in fact the switch would not absorb an eventually reflected wave, but in order to obtain however 7-9 kV at the first buncher electrodes a switch with double maximum operating voltage (24 kV) and with rise time of 4 ns would be necessary. This last solution, although it was the best for the safety of the switch, was discarded because a 24 kV switch with 4 ns rise time is too expensive.

### 4.3.1 Technical construction

Due to the high current value flowing in the circuit (123 A), each resistor of Tab.4.3 was realized by using 12 resistors in parallel such as in each resistor a current of 10 A flows. A platelet with four resistors in parallel was realized with a gold electro-deposition on a high resistivity silicon support, where the width of each gold line (see Fig.4.15b) is related to the desired value for

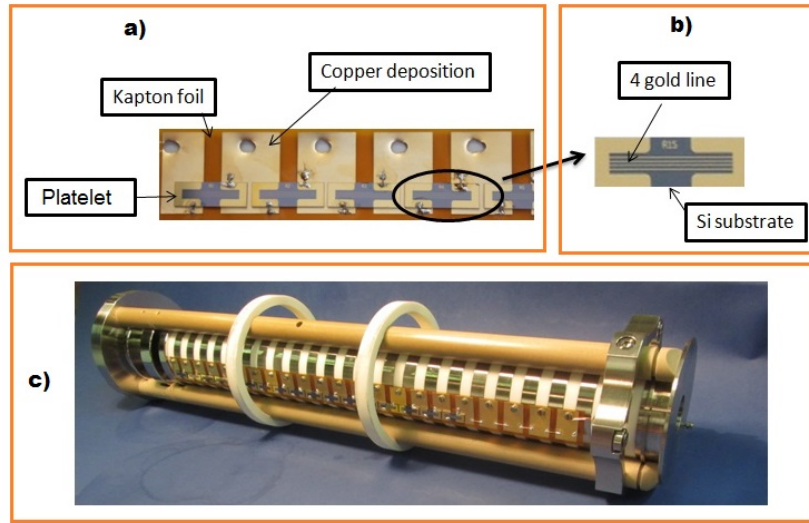


Figure 4.15: a) Picture of a kapton-gold strip. The kapton foil and the platelets soldered on the copper deposition are shown. b) Picture of a platelet where 4 resistors are fabricated by 4 gold deposited line on a Si-substrate. c) Picture of the electron optical line after assembling. The electrodes V1, VL and the 25 electrodes and macor rings of the buncher system were kept together by three peek bars, tangential to the external electrode surface. A small screw connects each electrode to the platelet through the copper deposition. One of the three strips with the resistors and the screws is visible.

reproducing the quasi-parabolic potential. Each platelet was soldered on a kapton-gold strip with indium (see Fig.4.15a) and three of these strips were wired to the external electrodes along the buncher and were placed at  $120^\circ$  on the external surface of the buncher electrodes, to reduce the magnetic field on the symmetry axis (i.e. the positrons path), generated by the current that goes through the resistors. A small screw connects each electrode to the platelet through the copper deposition.

The electrodes, whose geometry was mainly determined by the SIMION simulation, were built in arcap material, in order to avoid a possible magnetization in the region, where the magnetic field is present and were separated by each other by macor rings. A particular shape in the external part of the electrodes was chosen, in order to stack the electrodes with the macor rings, but at the same time in such a way that the insulator macor material does not affect the electric field in the inner zone. The many stacked macor and arcap rings constituting the whole electron optical line, were kept together by three peek bars tangential to the external electrode surface (see Fig.4.15 c). At the

end the cylindrical stuff, shown in Fig.4.15, was inserted into the external chamber shown in Fig.4.5, where some connectors entering from the lateral flanges were in touch with the V1, V2, VL electrodes, separately. More the first and the last buncher electrode were connected to the switch and to the  $R_i$  and  $R_f$  resistors, like the circuit configuration shown in Fig.4.12. The resistors  $R_f$  and  $R_i$  were placed near the connectors inside a shielded box.

### 4.3.2 Tests on the switching circuit

Preliminary tests with voltages ranging from 0 kV to 5 kV, with a preliminary switch [BELKE-HTS-50-12] with 5 ns rise time and 150 ns FWHM pulse duration were done. The configuration of maximum offset in which  $R_i = 0$  and  $R_f = 59 \Omega$  was chosen and a delay time of  $T_D = 500$  ns bigger than the pulse duration  $\Delta t$  (that in this case is 150 ns), was obtained with a 100 m long coaxial cable with a velocity factor of 0.66. In the following figures pulses acquired by a probe connected to the output of the switch or to the first or the last electrode are shown. In these tests the capacitor before the switch was charged up to 100 V and a pulse amplitude of 98 V was found at the first electrode due to the voltage drop on the inner switch resistor. In Fig.4.16 the signal as acquired at the output of the switch is shown, where a reflection is evident at  $1\mu s$  after the signal starting. It is due to the expected reflection at the first electrode of the buncher because a mismatch between  $Z_{cable}$  and  $Z_b$  is present. The amplitude of acquired reflected signal is around 33 V to be compared with the signal amplitude of 98 V. It is worth noticing that the reflected signal is acquired on the switch output when the switch is open again. Thus in that point a total reflection of the signal reflected at the beginning of the buncher occurs and as a consequence the real reflected signal amplitude is 16 V. By using the reflection coefficient that is the ratio between the reflected signal over the incident signal, the impedance of the buncher system can be evaluated. By the relation  $\Gamma = \frac{Z_b - Z_{cable}}{Z_b + Z_{cable}}$  demonstrated in the appendix A, the buncher impedance was found around  $106 \Omega$  and this is due to the mismatch between the box, where the resistor  $R_f$  and the other circuit element are placed, and the buncher system. Moreover the  $1\mu s$  delay of the reflected signal is due to the cable length of 100 m, appropriately chosen to have a delay time of 500 ns (that here is doubled because the cable is crossed twice from the incident and from the reflected signal) not to disturb the electrodes voltage, when positrons are flowing inside the buncher system. On the contrary the reflections between the switch and the cable can be seen in the small oscillations on the top of the pulse.

The pulse as acquired at the first and at the last electrodes is shown in Fig.4.17. The reflected signal above discussed is also present, assuming here

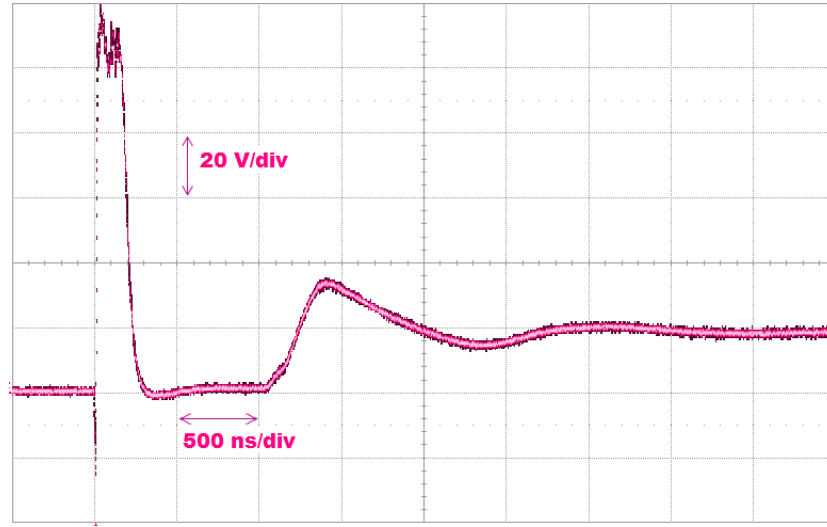


Figure 4.16: Signal acquired on the switch output, when a pulse amplitude of  $V_S = 100$  V was set. The reflected signal after  $1 \mu s$  is visible.

its real value of around 16 V. As in the previous case the reflected signal crossed twice the cable length before being acquired and thus a delay time of  $1 \mu s$  was found.

The pulse amplitude acquired at the last electrode (see Fig.4.17b) results to be around 80 V as expected due to the voltage divider between  $R_b = 16 \Omega$  and  $R_f = 56 \Omega$  because  $V_f/V_{in} = \frac{R_f}{R_{tot}} = 0.78$ .

More due to the time constant of the switching network, the voltage decrement of 0.1 % inside the pulse duration is negligible for this values and not observable. As regards the rise time and the pulse duration, they were found

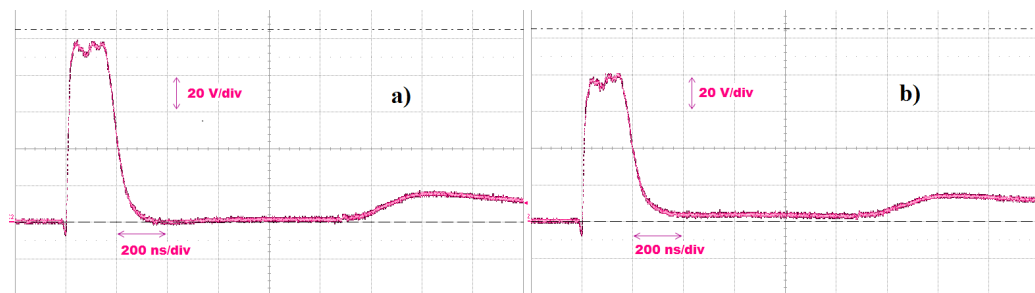


Figure 4.17: Voltage pulse signal as acquired at the first electrode a) and at the last electrode b).

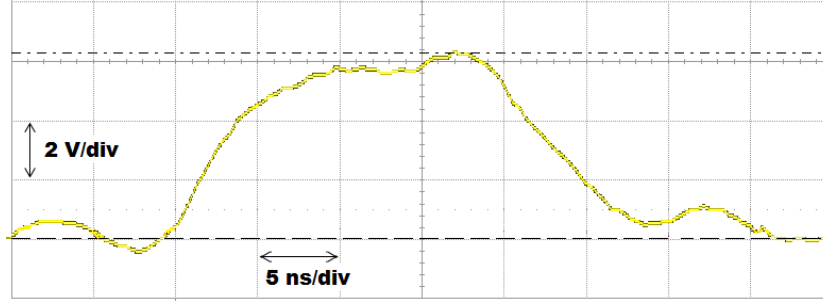


Figure 4.18: Acquired signal at the first electrode with  $V_S = 8kV$ . A value of 6 kV was found at the first electrode as acquired with a probe 1000:1.

compatible with those reported in the datasheet.

In a second phase tests with a switch having 3 ns rise time and 20 ns full width half maximum (BEHLKE HTS-12-150-UF) were performed by using a pulse amplitude of  $V_S = 8kV$ . The signal as acquired at the first electrode is shown in Fig.4.18, where the rise time was found slightly larger than 3 ns due to the different slope in the signal rise. The final value of 6 kV, as expected due to the inner switch resistor of  $15\ \Omega$ , which produce a voltage drop of around 2000 V, was found together with a full width half maximum of 20 ns. At the same time the current value was acquired by a High Frequency Current Transformers and a value of around 100 A was found in agreement with the expected results of  $\frac{8kV}{75} = 106\ A$ .

## 4.4 The sample chamber

In order to obtain a larger number of apertures for detectors or other measurement devices, the chamber was designed with an octagonal shape in such a way that 8 flanges overlook the sample, from the lateral side. More three smaller flanges tilted of  $45^\circ$ , with respect to the sample surface were designed. These could be useful if the target surface had to be directly observed or hit by laser beam, as for example for laser cooling measurements or Doppler measurements.

Anyway in order to perform the planned measurements on Ps in Rydberg states (see chapter 5), some conditions had to be satisfied in the sample chamber. Firstly due to the adjacency of the AEgIS magnet of 5 T to the sample chamber (see Fig.4.2) a residual magnetic field of 300 G had been found in that region that had to be shielded not to perturb the energy levels of the formed Ps. Secondly, despite the sample was designed to be at ground potential, the last electrode placed 56 mm away from the target surface and

set at 500 V, produced an electric field of 90 V/cm near the sample that had to be reduced for the same reason discussed before. Thirdly for ionizing the Ps after Rydberg excitation an electric field of some kV/mm had to be produced in front of the target. For this reason two grids were designed one on the top and one on the bottom of the sample and two switches were provided in such a way that a potential of around 1 kV will be applied only after the Ps formation and excitation.

In the following sections the ploys used to fulfill these requests will be shown together with the simulations and the performed tests.

#### 4.4.1 Magnetic field shield

In order to reduce to 2 G at maximum the magnetic flux density in the target chamber, a proper  $\mu$ -metal shield to be put all around the new apparatus, was designed. The shield is formed by two boxes, built one inside the other and separated by a gap of 10 mm, in such a way that a reduction from 300 to  $\sim 20$  is achieved thanks to the external shield and another reduction from 20 to few gauss thanks to the internal one. Each shield has a thickness of 1.5 mm and the two shields must not touch each other for allowing the maximum screening effect. Moreover some overlays of the adjacent panels forming the box were provided along each edge, in order to ensure the continuity of the magnetic lines inside the  $\mu$ -metal shield. A particular geometry was chosen near the  $\mu$ -metal flange, in order to connect it only to the inner box. Finally the holes necessary for the electric or vacuum connections were provided by proper collars, which ensure a reduction of the magnetic field penetrating in the box. The final drawn of the  $\mu$ -metal box is reported in Fig.4.19, where the tunnel necessary for the AEgIS transfer line on the bottom panel is visible.

The design of two boxes with a simplified shape, placed near the 5 T magnetic field like in the AEgIS geometry, was simulated using COMSOL Multiphysic software. The magnetic flux density on a plane crossing the sample is shown in Fig.4.20, where the request of having at maximum 2 G near the sample results to be fully satisfied.

#### 4.4.2 Switching circuit of the last electrode VL

To produce Ps in free electric field the potential of the last electrode (VL) has to be switched from its focusing value (500 V) towards 0 V (i.e. the same potential of the sample) after positrons have passed through it. To do this a switching circuit was designed, in order to lower the VL voltage in around 5 ns that is less than the time spent by Ps inside the nanochannels, before

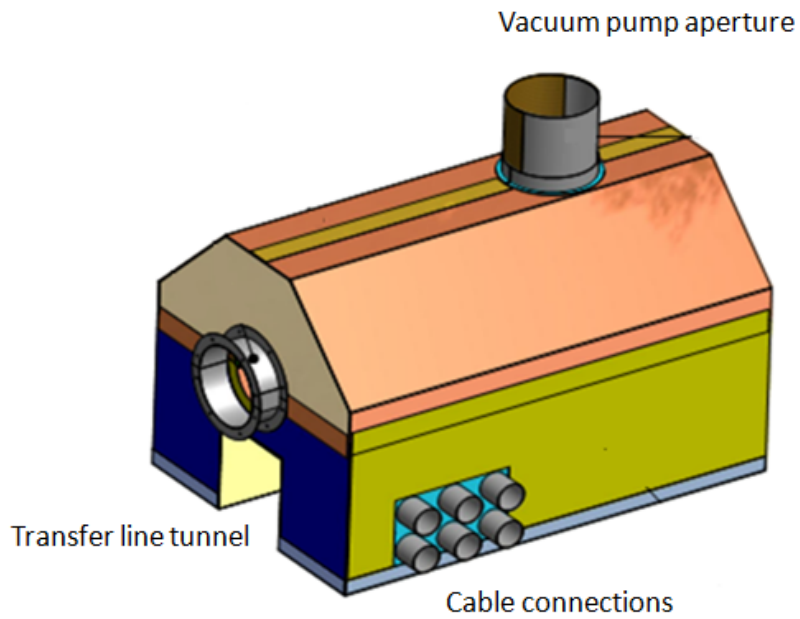


Figure 4.19: Sketch of the  $\mu$ -metal shield necessary to screen the residual magnetic field generated in the sample area by the AEgIS 5 T magnet. The tunnel for the AEgIS transfer line, the apertures for electronic or detectors cables and for vacuum pump are visible. Every hole of the box has its collar to ensure the magnetic screening. Two holes (not shown in the figure) will be added for laser entrance in correspondence of the sample chamber position.

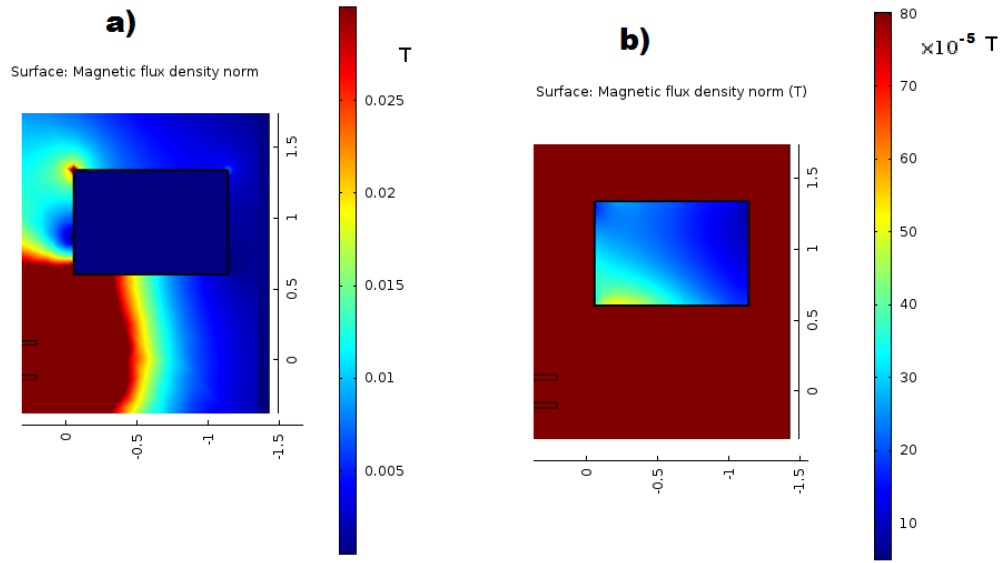


Figure 4.20: Magnetic flux density on a plane crossing the sample. The magnetic field is due to the 5 T magnet (visible at the left corner) in presence of two  $\mu$ -metal boxes. Two different ranges for the magnetic flux density were chosen in plot a) and b), in order to evidence the outer or the inner zone, respectively. A maximum field of 7 Gauss inside the box can be estimated and a lower value of  $\sim 2$  Gauss was found on the sample position, as demanded for Ps experiments.

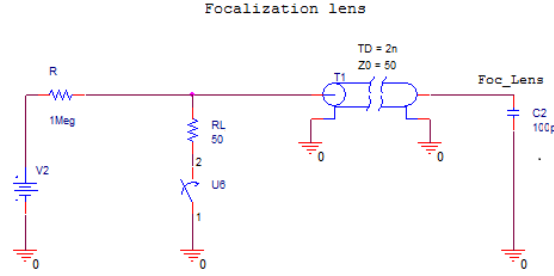


Figure 4.21: Circuit scheme to switch-off the VL electrode potential. The voltage drops from 500 V to 0 V in 10 ns after positrons have passed through it.

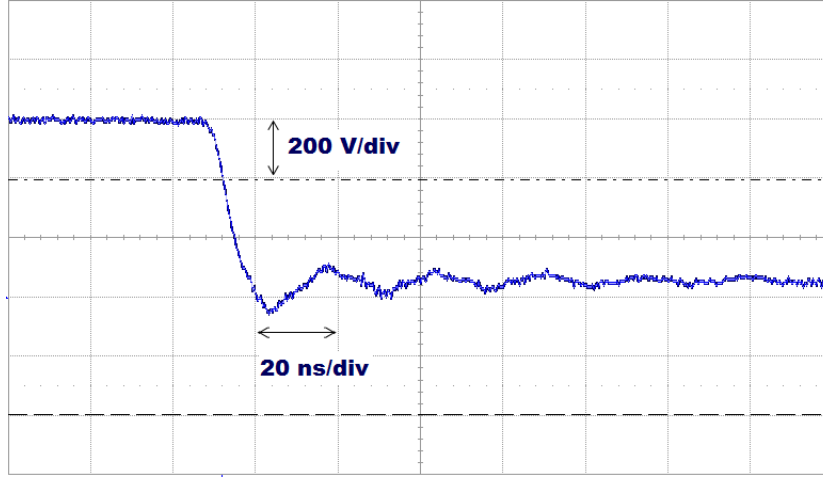


Figure 4.22: Signal acquired by a probe connected to the last focusing electrode surface, when the switch was closed.

being emitted into vacuum (see chapter 3). In Fig.4.21 the circuit scheme is shown, where the electrode is represented as a capacitor of  $C \sim 60$  pF. When the switch is open, after a transient of around  $3\tau = 3R_L C = 300 \mu\text{s}$ , no current flows and the electrode potential is set to the nominal value given by the power supply voltage; on the contrary when the switch is closed, the capacitor discharges itself on the resistor  $R = 50 \Omega$  with a time constant given by  $3\tau = 3RC = 9$  ns. In order to minimize the reflections, the cable was chosen with characteristic impedance of  $50 \Omega$  like the resistor value. The signal as acquired by a probe connected to the electrodes is shown in Fig.4.22. In this figure the fall time results to be around 10 ns as designed, while the potential value, starting from 600 V, is reduced near zero potential (80 V).

### 4.4.3 Switching electric field near the sample

In order to switch an electric field for the Ps ionization, two parallel grids were designed in front of the target position (see Fig.4.23). Being the grids

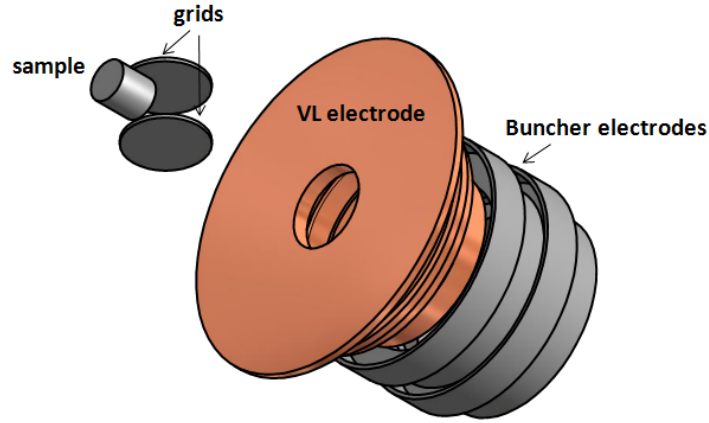


Figure 4.23: Sketch of the two grids for producing an electric field of some kV/cm for Ps ionization. The last buncher electrodes and the VL electrode are also visible.

1.4 cm far from each other, a potential of +700 V and -700 V was chosen for each grid in order to produce an electric field of 1000 V/cm just in front of the target. Two High-Voltage Q-Switch Driver devices (BEHLKE FQD 40-03), whose circuit is shown in Fig.4.24, were bought in order to be connected separately to two different power supplies, one with positive polarity and the other one with negative polarity. Each grid, that has not yet been built, will be attached to the correspondent switch output and the same TTL trigger signal will be sent to the input of both switches to cause the potential change. The rise time of the signal related to the  $R_S$  value (see Fig.4.24), was chosen  $\sim 3$  ns, because the electric field has to be produced only after the Ps emission and the laser excitation in the Rydberg states, but before that the Ps cloud has fled too far from the target. The pulse duration depending on  $R_L$  value, was chosen around  $1\mu s$ , time interval comparable with the lifetime of the Ps in Rydberg state. Due to the high time constant given by the product between  $R_S$  and  $C_{grid}$ , the voltage drop due to capacitor discharge is negligible during the pulse duration and in this interval the grid potential can be considered stationary.

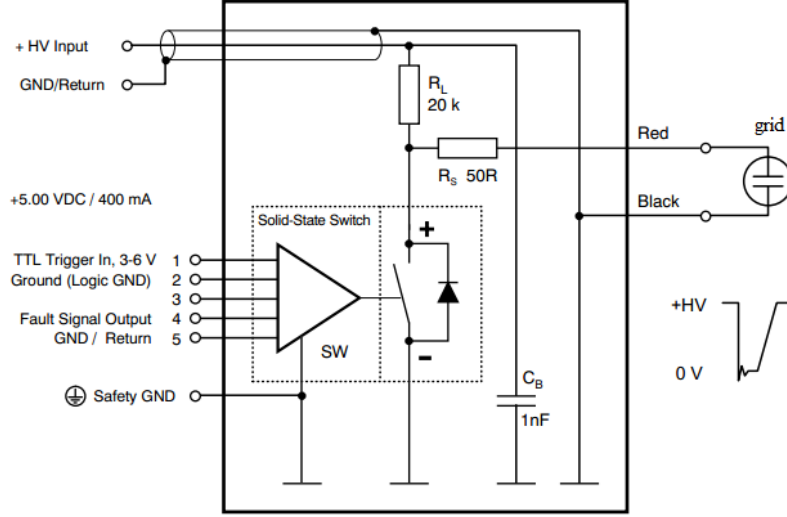


Figure 4.24: Circuit scheme embedded in the High-Voltage Q-Switch Driver device (BEHLKE FQD 40-03)

## 4.5 Conclusions

The apparatus described in this chapter was designed for producing positron bunches at the porous sample with properties compatible for producing Ps in vacuum and measuring its excitation in Rydberg states. The magnetic line of the AEgIS transfer line was adapted, extended and finally connected through a  $\mu$ -metal electrode to an electron optical lines appositely designed and built. Electric fields generated by electric potentials applied to many electrodes allowed positron bunches to be temporally compressed, accelerated to an energy ranging from 6 to 9 keV and focused on the sample with the demanded conditions. The sample chamber was shielded from residual magnetic field of AEgIS set up and a further electric circuit acting on VL electrode, was designed and tested in order to obtain a free electric field region in the sample zone. Moreover the sample chamber was designed with many open ports which overlook the sample, allowing many detectors or lasers or other devices to be set simultaneously.

Finally in order to produce a tunable and controlled magnetic field in the sample zone for Ps measurements spectroscopy, two coils were built and located just out from the chamber, in order to produce a magnetic field of 200 G, perpendicularly directed to the target surface. On the contrary inside the sample chamber nothing has been yet mounted, even if two switching drivers were bought for changing the potential of the two grids provided near

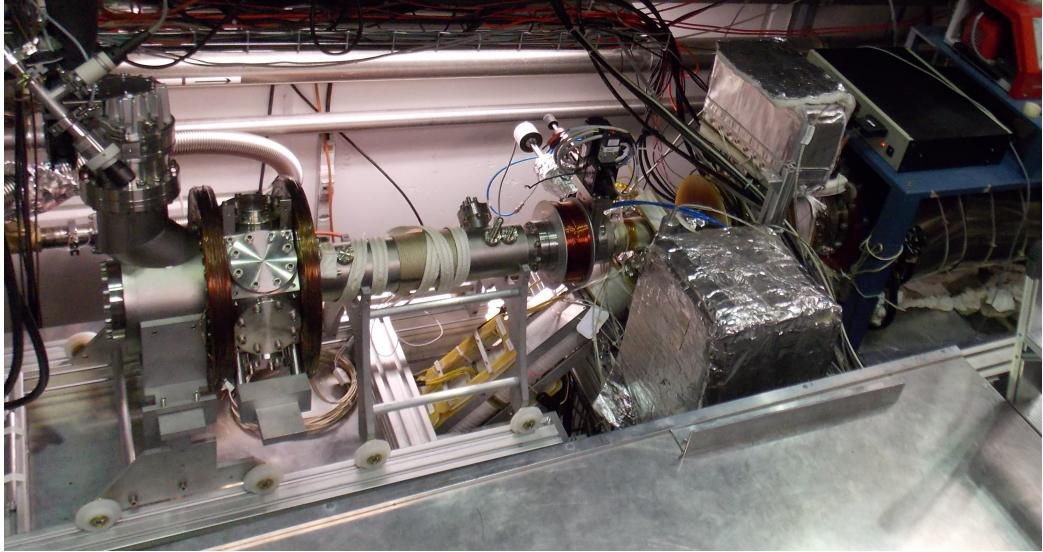


Figure 4.25: Picture of the apparatus, embedded in the AEgIS set up. The AEgIS transfer line and the final part of the Surko-trap are visible. The  $\mu$ -metal shield and the cryostat with the sample holder are not present in the picture. The two coils for tunable magnetic field are visible around the sample chamber.

the sample for Ps Rydberg ionization.

On the last October the apparatus was connected to the valve of the AEgIS transfer line in the AD zone at CERN and now it is ready for the first tests with positron bunches. In the Fig.4.25 the apparatus embedded in the AEgIS set up is visible.

In the following chapter, measurements that could be performed with this apparatus in the next years will be presented.



## Chapter 5

# Ps excitations in Rydberg states and Ps spectroscopy

In this chapter the Ps levels in free field or in the presence of a magnetic field are explained and in particular the two excitations from  $n = 1$  to  $n = 3$  and from  $n = 3$  to the Rydberg states are investigated, in order to point out the laser parameters (as wavelength, fluence or energy) necessary for a Ps efficient excitation in the AEgIS experiment (see section 1.4).

The first experimental observations of the Ps excitation in Rydberg states, as realized by Ziock's [94] and Mill's [93] groups, are described. At the end the planned measurements that will be performed, by using the apparatus described in the previous chapter, are presented and moreover, the Ps laser cooling technique is explained and taken under consideration for next measurements.

### 5.1 Ps levels in free field

The Hamiltonian describing Ps atom in zero external fields can be written as:

$$H_0 = H_f + H_C + H_F \quad (5.1)$$

where  $H_f$  is the Hamiltonian of free positron and electron,  $H_C = -\frac{e^2}{4\pi\epsilon_0 r}$  is the Hamiltonian of a two particles system with electric charge  $+e$  and  $-e$  at distance  $\vec{r}$ , interacting via the Coulomb potential, and finally  $H_F$  is the correction due to the spin-orbit and spin-spin hyperfine relativistic interactions. The energy levels of Ps are primarily determined by the spherical symmetric Coulombian potential, hence are similar to those of a hydrogen atom but

with a reduced mass  $\mu = m_e/2$ :

$$E_C(n) = -\frac{\mu c^2 \alpha^2}{2n^2} \cong -\frac{6.8 \text{ eV}}{n^2} \quad (5.2)$$

where  $\alpha = \frac{e^2}{4\pi\epsilon_0\hbar c}$  is the fine-structure constant and  $n$  is the principal quantum number.

Taking care of spin-orbit coupling and hyperfine interaction  $H_F$ , the energy levels, exact to the order  $\alpha^4$ , can be expressed as usual in terms of  $n$  and the angular quantum numbers  $s$  (of the total spin  $\vec{S}$ ),  $l$  (of the total angular momentum  $\vec{L}$ ), and  $j$  (the quantum number of the sum of the total angular momentum and total spin  $\vec{J} = \vec{S} + \vec{L}$ ) [95]. If  $m = -j, \dots, j$ , the corrections to the levels expressed by eq.5.2 can be written as:

$$E_F(n, s, l, m_l) = \frac{\mu c^2 \alpha^4}{4n^3} + \left[ \frac{11}{8n} - \frac{4}{2l+1} + \delta_{s,1} \left( \delta_{l,0} \frac{14}{3} + \frac{2(1 - \delta_{l,0})}{l(2l+1)(l+1)} + \left\{ \begin{array}{ll} -\frac{(l+1)(3l-1)}{2l-1} & \text{if } j = l-1 \\ -1 & \text{if } j = l \\ +\frac{l(3l+4)}{2l+3} & \text{if } j = l+1 \end{array} \right\} \right) \right] \quad (5.3)$$

where the possible total spin quantum numbers are  $s = 0$  (singlet states) and  $s = 1$  (triplet states). It is worth noting that  $E_F$  decreases with the higher  $n$  and in particular for  $n = 3$ , the subset of triplet states  $^3D_2$  and  $^3P_2$  have the same energy and an *accidental degeneration* occurs. On the contrary, for  $n = 2$  any different value of  $l, s, j$  corresponds to different energies [96]. In Fig.5.1 the first 3- $n$  levels of Ps in zero external fields described by eq.5.2 and 5.3 are shown, where the accidental degeneration at  $n = 3$  is visible.

The Ps eigenstates corresponding to the energy levels of eq.5.3 and 5.2 can be written as a ket  $|nsljm\rangle = |R_{nl}\rangle |sljm\rangle$ , where  $|R_{nl}\rangle$  represents the radial eigenfunction (in the case of a spherical symmetric potential, as in the present case) and  $|sljm\rangle$  is the spin and angular momentum part of Ps states. In the following the ket  $|sljm\rangle$  will be expanded in terms of the ket product  $|nlm_l\rangle |sm_s\rangle$ , which is the natural description in absence of hyperfine interactions.

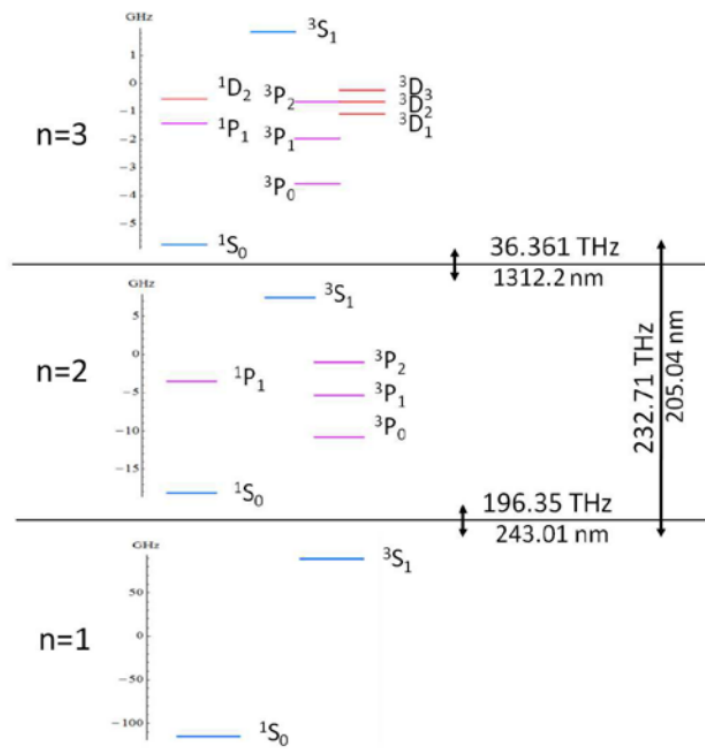


Figure 5.1: First 3- $n$  levels of Ps in zero external fields. The spectroscopic notation  $^{2S+1}L_J$  is indicated for each one.

## 5.2 Ps levels in magnetic field

In presence of external magnetic field, the Hamiltonian can be written as:

$$H = H_0 + H_B \quad (5.4)$$

where

$$H_B = H_Z + H_{dia} + H_{MS} \quad (5.5)$$

In particular  $H_Z$  and  $H_{dia}$  represent the first order and the second order (diamagnetic or quadratic) Zeeman interactions, respectively. These terms are due to the interaction between magnetic field and Ps at rest, while  $H_{MS}$  takes into account the contribution of the Motional Stark effect that is present only if charges are in motion within the magnetic field.

In order to calculate the magnitude of linear Zeeman splitting  $\Delta E_Z(s, m_s)$ , the magnetic dipole momenta associate to orbital angular momentum  $\vec{L} = \vec{r} \times \vec{p}$  and spin  $\vec{S} = \frac{\hbar}{2} \vec{\sigma}$  (where  $\vec{\sigma}$  are the Pauli spin operators) for both electron and positron can be written as:

$$\vec{\mu}_{L(e^+, e^-)} = \pm \frac{e}{m_e} \vec{L}_{e^+, e^-} \quad (5.6)$$

and

$$\vec{\mu}_{S(e^+, e^-)} = \pm \frac{e}{m_e} \vec{S}_{e^+, e^-} \quad (5.7)$$

where the gyromagnetic factors are assumed to be 2. If  $\mu_B = e\hbar/2m_e$  is the Bohr magneton, the linear Zeeman Hamiltonian comes out to be:

$$H_Z = -(\vec{\mu}_{L, e^+} + \vec{\mu}_{L, e^-} + \vec{\mu}_{S, e^+} + \vec{\mu}_{S, e^-}) \cdot \vec{B} = \mu_B \vec{B} (\vec{\sigma}_{e^-} - \vec{\sigma}_{e^+}) \quad (5.8)$$

where the equality of eq.5.8 was obtained considering that in the center of mass  $\vec{\mu}_{L, e^+} = -\vec{\mu}_{L, e^-}$  and thus there is no energy contribution from the orbital motion.

On the contrary the Hamiltonian for quadratic Zeeman effect can be written as [97]:

$$H_{dia} = \frac{e^2}{8\mu} (\vec{r} \times \vec{B})^2 \quad (5.9)$$

where  $\vec{r}$  is the relative position of the particles. Its contribution results to be relevant only for  $n$  levels greater than 40 and in any case it turns out to be very small with respect to other interactions [98] in the Rydberg excitation where  $n = 20$ ; for this reason it will be neglected in the following.

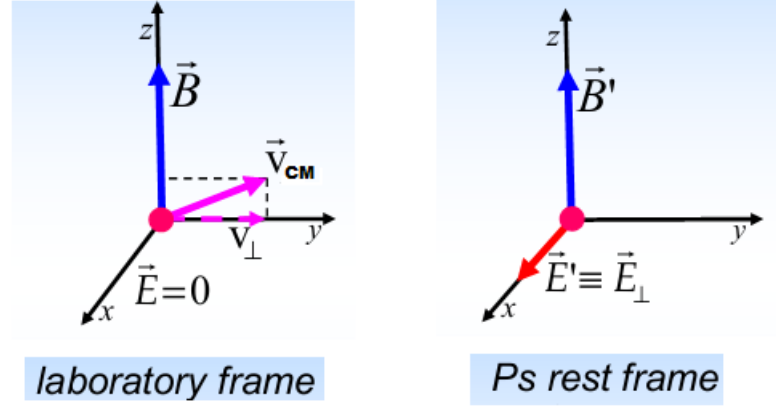


Figure 5.2: Scheme of magnetic field transformation from the laboratory reference frame, where  $\vec{v}_{CM} \neq 0$  and no electric field is present, to the Ps rest frame ( $\vec{v}_{CM} = 0$ ), where an electric field, orthogonally directed to the magnetic field, arises.

Finally the Motional Stark effect contribution can be calculated for Ps moving with center-of-mass velocity  $\vec{v}_{CM}$  in the laboratory reference frame. In the (instantaneous) Ps atom rest frame, the magnetic field can be transformed, in the non-relativistic approximation, as:

$$\begin{cases} \vec{B}' = \vec{B} \\ \vec{E}' = \vec{v}_{CM} \times \vec{B} = \vec{E}_{MS} \end{cases} \quad (5.10)$$

where a new electric field  $\vec{E}_{MS}$  arises. This electric field can be interpreted as a field, self-induced by the Ps motion in the magnetic environment, acting on Ps exactly as an external electric field. Therefore in the center of mass reference frame, where Ps is at rest, both magnetic and electric field are present. It is worth noticing that the direction of this electric field results to be orthogonal to the magnetic field direction (see Fig.5.2), thus destroying the axial symmetry. Because the interaction between an atom and an external electric field is described as Stark Effect, the Hamiltonian describing the so called *Motional Stark Effect*, due to the  $\vec{E}_{MS}$  field, can be written as [99]:

$$H_{MS} = -e \vec{r} \cdot \vec{E}_{\perp} = -e \vec{r} \cdot (\vec{v}_{CM} \times \vec{B}) \quad (5.11)$$

In conclusion by neglecting the diamagnetic contribution and by choosing the reference system in which the  $z$  axis is along the direction of magnetic

field ( $\vec{B} = B \vec{z}$ ) (see Fig.5.2), the global hamiltonian of Ps in magnetic field can be written as:

$$H = H_0 + \mu_B B (\sigma_{e^-z} - \sigma_{e^+z}) - e \vec{r} \cdot \vec{E}_{MS} \quad (5.12)$$

In the following a description of the eigenvalues of the Zeeman and the Motional Stark Effect hamiltonians will be done for obtaining Ps levels in the presence of a magnetic field.

### 5.2.1 Zeeman effect

As regarding the Zeeman contribution, it is possible to observe that  $H_Z$  is not diagonal in the singlet and triplet basis of Ps state  $|s, m_s\rangle$ . In fact because the total spin  $s$  can be equal to 0 or 1, the matrix elements

$$\langle s', m'_s | \sigma_{e^-z} - \sigma_{e^+z} | s, m_s \rangle \quad (5.13)$$

over the total spin states  $|00\rangle$  and  $|1m_s\rangle$  ( $m_s = 0, \pm 1$ ), is non zero only for  $m'_s = m_s = 0$  and  $s \neq s'$ . Therefore the magnetic perturbation mixes ortho and para-Ps states with  $m_s = 0$ , while it leaves unchanged the triplet states with  $m_s = \pm 1$  of the o-Ps. As a consequence the lifetime of o-Ps with  $m_s = \pm 1$  does not change, while, due to the different lifetime of o-Ps with respect to p-Ps, the levels mixing leads to the well known enhancement of the average annihilation rate of the Ps in ground state  $n = 1$ , whose effect is called *magnetic quenching* [100]. The lifetime of o-Ps with  $m_s = 0$  as a function of the external magnetic field, calculated as in [101], is shown in Fig.5.3. Finally the constraint  $m_s = m'_s$  implies that the Zeeman interaction does not change the total angular momentum projection on the  $z$  axis.

The maximum value of the matrix element can be evaluated as

$$\Delta E_Z^{max} = 4\mu_B B \quad (5.14)$$

that, for a magnetic field of 1 T, results to be  $2.4 \cdot 10^{-4}$  eV (58 GHz) independently from  $n$ .

### 5.2.2 The Motional Stark effect

As regarding the Stark effect, by calculating the matrix elements:

$$\langle nsljm | -e \vec{r} \cdot \vec{E}_\perp | n's'l'j'm' \rangle \quad (5.15)$$

it is possible to demonstrate [96] that the interaction acts on states with different  $l$  (such as  $\Delta l = l - l' = \pm 1$ ), and with  $j$  and  $m$  such as

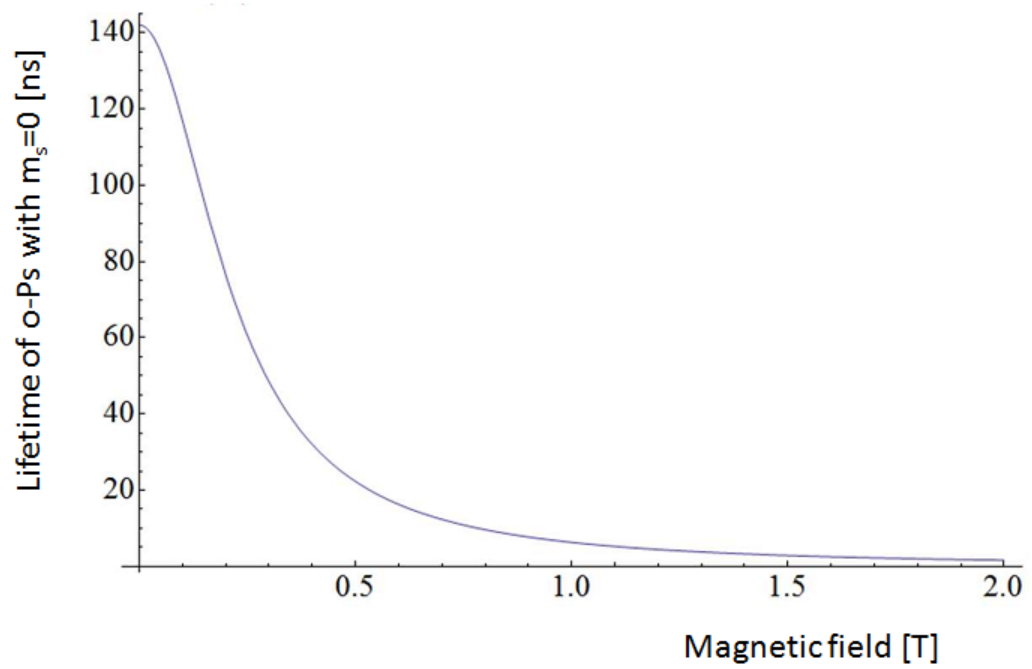


Figure 5.3: Lifetime of o-Ps with  $m_s = 0$  as a function of the magnetic flux density. The reduction of the lifetime is due to the magnetic quenching effect. After [102].

$\Delta j = j - j' = 0, \pm 1$  and  $\Delta m = 0, \pm 1$ . For this reason a complete breaking of the original axial symmetry around the  $\vec{B}$  axis occurs, with the consequence of a complete mixing of the  $m$  and  $l$  state belonging to the same  $n$ , which remains the unique good quantum number.

The maximum energy splitting  $E_{MS}$  between the  $n^2$  sub-levels can be estimated as [96]:

$$\Delta E_{MS}^{max} = 3ea_{Ps}n(n-1) \left| \vec{E}_{\perp} \right| = 3ea_{Ps}n(n-1)Bv_{\perp} \quad (5.16)$$

where  $v_{\perp}$  is the transverse component of the Ps center of mass velocity  $\vec{v}_{CM}$  with respect to the  $B$  direction and  $a_{Ps} = 2a_0$  is the Bohr radius of the Positronium, that is double with respect to the hydrogen radius  $a_0$ . Because  $\Delta E_{MS}$  increases both with  $B$  and  $n$  and because the mass of Ps is around 1000 time lighter than the hydrogen mass, this splitting effect becomes relevant on Ps atoms even for small  $n$  values and it is largely predominant on Rydberg levels in the working conditions of the AEgIS experiment where  $B=1$  T. Values of the Motional Stark splitting are 12.6 GHz for  $n = 3$ , 800 GHz for  $n = 20$ , 1800 GHz for  $n = 30$  with 1 T magnetic field and the most probable Ps thermal velocity  $v_{\perp} = \sqrt{k_B T / m_{Ps}}$  calculated with  $T = 100$  K.

In this context, because an electric field of  $E_{\perp} = 275$  V/cm is obtained if  $B = 1$  T and  $T = 100$  K, the possible Ps ionization induced by the electric field  $E_{\perp}$  itself has to be considered as well. In fact the transition between the bottom sub-level (called the *red state*) of each  $n$  level to the unbound state occurs with the minimum electric field of [103]:

$$E_{min} = \frac{e}{4\pi\epsilon_0 a_{Ps}^2} \frac{1}{9n^4} \quad (5.17)$$

Thus, for the working conditions of the AEgIS experiment ( $B = 1$  T,  $T = 100$  K) the ionization starts affecting some sub-levels from  $n > 27$ , and as a consequence,  $n = 27$  will be the highest  $n$  value, usable for the Ps excitations in Rydberg states.

Finally, once the magnetic flux density and the Ps velocity have been chosen, the different contributions on the splitting energy, (given by Zeeman effect or Motional Stark Effect or fine interaction) can be plotted as a function of the quantum number  $n$ , as shown in Fig.5.4 where the AEgIS conditions ( $B = 1$  T and a  $T = 100$  K) were adopted. The dominant effect in each  $n$  region can be observed, and more it is possible to conclude that the Zeeman and diamagnetic energy splitting can be neglected at higher  $n$ , where the energy splitting due to the Motional Stark effect becomes more and

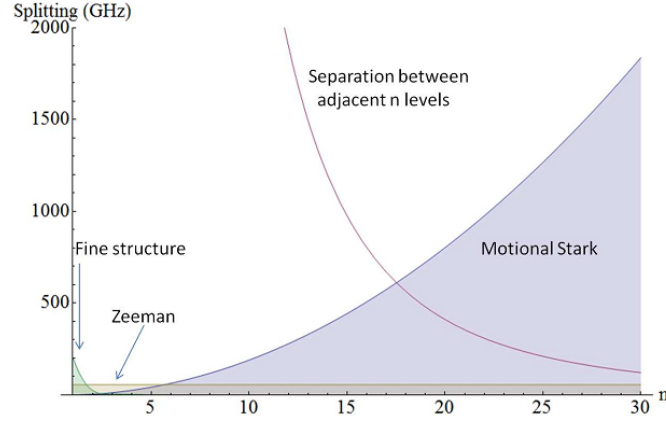


Figure 5.4: Maximum energy splitting (in GHz) due to fine structure, Zeeman effect and Motional Stark effect for 1 T magnetic field and a Ps thermal velocity calculated with  $T = 100$  K. After [102].

more relevant. In particular at  $n = 17$  the energy spread between adjacent unperturbed  $n$  levels (eq.5.2) becomes smaller than the maximum Motional Stark splitting and an interleaving between the  $n$  levels occurs. In these conditions, considering a cloud of positronium atoms with thermal distribution of velocity, the laser pulses see Rydberg states as a band rather than separate discrete levels and this approximation is used to calculate the transition probability (see the next section).

### 5.3 Ps excitations in Rydberg states

The direct photo-excitation of Ps, from ground state to the Rydberg band, requires photon energies close to 6.8 eV, which corresponds to a wavelength of 180 nm. Since such a type of laser is not commercially available, a two-step excitation with two simultaneous laser pulses is required. In the AEgIS the excitation from  $n = 1$  to  $n = 3$ , involving radiation at 205 nm, and then from  $n = 3$  to  $n = 20 - 30$ , where radiation in the range from 1600 nm to 1700 nm is necessary, was proposed. This choice seems more adequate than the other possible two steps sequence  $1 \rightarrow 2$  and  $2 \rightarrow \text{high-}n$ , because the level  $n = 2$  has three times shorter lifetime than the  $n = 3$  level (3 ns against 10.5 ns). In addition, while the power required for the transitions  $1 \rightarrow 2$  and  $1 \rightarrow 3$  up to saturation are close to each other, the transition  $2 \rightarrow n$  (up to saturation) requires one order of magnitude higher power when compared with the transition  $3 \rightarrow n$ , mainly due to the broadening of the Rydberg

levelband of the Ps excitation lines, as explained below.

Anyway both the optical excitations are strictly dependent on the laser properties, as bandwidth and fluence, that have to be optimized for an efficient production of Ps in Rydberg states.

As regards the bandwidth, since the excitation occurs on Ps in motion, the Doppler broadening has to be considered and its value has to be compared with the energy splitting due to the Zeeman and Stark effects above discussed. In particular the resonance frequency of a transition depends in the first approximation on the component of Ps velocity parallel to the propagation direction of the laser field as  $f = f_0(1 + v_{\parallel}/c)$  where  $f_0$  is the resonance frequency for Ps at rest and  $c$  is the light velocity. Thus for a cloud of Ps atoms with a Maxwellian velocity distribution at temperature  $T$ , the resonant wavelengths become distributed according a Gaussian function, whose full width half maximum can be written as:

$$\Delta\lambda_D = \lambda \sqrt{\frac{8K_B T \ln 2}{m_{Ps} c^2}} \quad (5.18)$$

depending on the wavelength of the transition  $\lambda$  and on the Ps temperature  $T$ . As a consequence in the first transition the Doppler broadening enlargement on the wavelength results to be around  $\Delta\lambda_D = 4.4 \cdot 10^{-2}$  nm (0.7 eV) for  $T = 100$  K and thus it dominates over the Zeeman effect, which affects only states with  $m_s = 0$  (due to electric dipole selection rules for optical transitions) and which amounts to  $1.2 \cdot 10^{-4}$  eV for  $B = 1$  T, much lower than the Doppler broadening. On the contrary by comparing the Doppler and the Motional Stark effect for  $n = 20 - 30$ , levels, it is possible to discover that the second transition is dominated by the Motional Stark effect, if a magnetic field is present as in the AEgIS experiment [104].

Therefore by describing the spectral profile of the lasers intensity with the following gaussian function:

$$g(\omega - \omega_0) = \frac{1}{2\pi\sigma_\omega} e^{-\frac{(\omega - \omega_0)^2}{2\sigma_\omega^2}} \quad (5.19)$$

centered on the proper frequency of the transition  $\omega_0$ , the full width half maximum (FWHM),  $\Delta\omega_L = \sigma_\omega 2\sqrt{2\ln 2}$ , has to overlap the correspondent line shape of the two resonances, given by the Doppler or by the Motional Stark effect, respectively.

Finally once the broad laser wavelength  $\Delta\omega_L$  has been fixed, the coherence time defined as  $\Delta T_c = \frac{2\pi}{\Delta\omega_L}$  can be calculated. Since the average laser pulse duration is 5 ns, the coherence time  $\Delta T_c$  results three orders of magnitude shorter than that for both the transitions and for this reason they can be

considered completely incoherent.

In this context for calculating the laser fluence that optimize the Rydberg excitation, the rate equation model [105] can be used. Thus, by considering the dipole allowed transition from a state  $a$  to a higher level state  $b$ , the rate equation for the population of the  $b$  state, is:

$$\frac{dP_b}{dt} = -P_b W_{SE} - P_b W_{ba}(t) + P_a W_{ab}(t) \quad (5.20)$$

where  $W_{SE}$  is the total spontaneous emission rate,  $W_{ab}(t)$  the absorption probability rate and  $W_{ba}(t)$  the stimulated emission probability rate. By neglecting  $W_{SE}$  if the  $b$  state lifetime is longer than the pulse laser duration, and by taking into account the condition  $P_a + P_b = 1$ , the equation eq.5.20 becomes:

$$\frac{dP_b}{dt} \cong (1 - 2P_b)W_{ab}(t) \quad (5.21)$$

If the laser pulse has the total intensity  $I_L(t) = \int_{\Delta E_L} d\omega I(\omega, t)$  where  $I(\omega, t)$  is taken as a time dependent Gaussian spectral intensity and its energy bandwidth is  $\Delta E_L = \hbar\Delta\omega_L$ , the transition probability for unit of time is given by:

$$W_{ab}(t) \int_{\Delta E_L} d\omega \frac{I(\omega, t)}{\hbar\omega} \sigma_{ab}(\omega) \quad (5.22)$$

with

$$\sigma_{ab}(\omega) = \frac{\hbar\omega}{c} g(\omega - \omega_{ab}) B_{ab}(\omega) \quad (5.23)$$

The factor  $B_{ab}(\omega)$  is the absorption Einstein coefficient, appropriate to the dipole allowed transition, calculated in the first approximation as:

$$B_{ab}(\omega) = \frac{\left| \left\langle \psi_b \left| e \vec{r} \cdot \vec{\epsilon} \right| \psi_a \right\rangle \right|^2 \pi}{\epsilon_0 \hbar^2} \quad (5.24)$$

where  $\vec{\epsilon}$  is the radiation polarization vector.

Thus the solution of the eq. 5.21 with the initial condition  $P_b(0) = 0$  is:

$$P_b(t) = \frac{1}{2} (1 - e^{-2F(t)/F_{sat}}) \quad (5.25)$$

being  $F(t) = \int_{-\infty}^t dt' I_L(t')$  the laser pulse fluence and

$$F_{sat}(a \rightarrow b) = \frac{c\sqrt{2}}{B_{ab}g(0)} \quad (5.26)$$

the saturation fluence. From the eq.5.25 the saturation fluence is defined as the fluence such as the 43% of the atoms are in the excited state and its value is taken as reference for the laser fluence in the optical incoherent transitions.

In the following the laser system providing radiation for the two excitations will be described and the correspondent saturation fluences and energies will be calculated for both transitions.

### 5.3.1 Excitation from $n=1$ to $n=3$

In order to perform the Ps excitation from  $n = 1$  to  $n = 3$  a laser system formed by a Nd:YAG laser and an Optical Parametric Generation crystal (OPG) and an Optical Parametric Amplifier (OPA) devices was proposed. Since the Nd:YAG laser provides three pulses at three different wavelengths (1064 nm, 532 nm and 266 nm, corresponding to the first, second and fourth harmonic of the Nd:YAG laser radiation), an OPG-OPA system was introduced in order to produce a 894 nm impulse, which, thanks to the relation  $\frac{1}{205nm} = \frac{1}{266nm} + \frac{1}{894nm}$  involving the fourth harmonic of the Nd:YAG radiation, allows the 205 nm light to be produced. The laser bandwidth has to cover the Doppler broadening of  $\Delta\lambda_D = 4.4 \cdot 10^{-2}$  nm calculated before, while the saturation fluence for this excitation can be obtained by using eq.5.26 with  $g(0) = \frac{1}{2\sigma_\omega\pi}$  as obtained by eq.5.19:

$$F_{sat}(1 \rightarrow 3) = \frac{c^2}{B_{13}} \sqrt{\frac{2\pi^3}{\ln 2}} \frac{\Delta\lambda_D}{\lambda_{13}^2} \quad (5.27)$$

By substituting the  $B_{13}$  value proper of the  $n = 1 \rightarrow 3$  transition, the result  $F_{sat}(1 \rightarrow 3) = 93.3 \mu J/cm^2$  can be achieved.

Finally for calculating the total energy of the laser pulse, an estimation of the laser spot size, which has to be larger than the Ps cloud, has to be done. Considering Ps going out from a circle of 2.5 mm diameter (due to the positrons focusing on the target as reported in chapter 4) with velocity of  $4 \cdot 10^5$  m/s, a Ps cloud with an area of about  $6 mm^2$  after 15 ns of expansion time can be considered [27]. Because the laser pulse is sent parallel to the sample surface, assuming a Ps cloud of FWHM dimension  $\Delta r \sim 2.8$  mm and fixing the fluence value as  $F = 2F_{sat}$ , to saturate the transition and to retain a security factor for maximum excitation, the laser pulse energy becomes  $E = \frac{\pi F}{2} \left(\frac{\Delta r}{1.177}\right)^2 \sim 14.4 \mu J$ , that is a value experimentally achieved by the laser facility, mounted in the AEgIS apparatus.

### 5.3.2 Excitation from $n=3$ to $n=20-30$

By considering the excitations from  $n = 3$  to  $n = 20 - 30$ , infrared radiation of wavelength between 1650 and 1700 nm has to be produced. In the AEgIS set up the demanded light is obtained by the conversion of the first harmonic (1064 nm) of the Nd:YAG laser radiation. In particular the correct wavelength is generated within an OPG and then, amplified by two OPAs; at the end a dichroic mirror filters the non-useful 1064 nm radiation and leaves only the amplified impulse with the desired wavelength.

In the  $n = 3 \rightarrow 20 - 30$  transition, where the Motional Stark effect is dominant, the laser energy bandwidth  $\Delta\lambda_L$  has to be greater than Doppler broadening, but smaller than the Motional Stark effect  $\Delta\lambda_{MS} = \Delta E_{MS}\lambda^2/hc \sim 10$  nm (correspondent to  $\sim 1$  THz), so that all mixed sub-levels with transition energy under the laser bandwidth can be populated.

In order to derive an expression for the second step saturation fluence, it is possible to observe that since many sub-levels are distributed in a narrow energy spread, a quasi continuum band of levels can be considered. Thus the density of levels for unit of angular frequency ( $\omega$ ) can be calculated as:

$$\rho(\omega) = \frac{n^2 \Delta E_{MS} / \Delta E_n}{\Delta E_{MS} / \hbar} = n^5 \frac{\hbar}{13.6 \text{ eV}} \quad (5.28)$$

where the number of the total levels inside  $\Delta E_{MS}$  range is calculated as the products between the number of  $n$  levels inside the  $\Delta E_{MS}$  and their multiplicity  $n^2$ . It is worth noticing that  $\rho(\omega)$ , which substitutes the  $g_D(0)$  in eq.5.26, is independent from the induced Stark field and consequently from the Ps velocity and from  $B$ .

Thus by generalizing the theory of the incoherent excitation for continuum levels, the eq.5.23 becomes:

$$\sigma_{3n}(\omega) = \hbar \omega \rho(\omega) B_{MS}(\omega) / c \quad (5.29)$$

with  $B_{MS}(\omega)$  the absorption coefficient, appropriate for the transition to the quasi-continuum level band modified by the Motional Stark effect. Because in first approximation  $B_{MS}(\omega) \cong \frac{1}{n^2} B_{3n}(\omega) \propto \frac{1}{n^5}$ , the absorption probability, calculated by eq.5.22, is found not depending on the quantum number  $n$  and on the transverse Ps velocity [104].

Finally by using the eq.5.26 modified in the Rydberg continuum levels context, the saturation fluence can be calculated as:

$$F_{sat}(3 \rightarrow n) \cong \frac{c \rho(\omega)}{B_{MS}} = \frac{c 13.6}{B_{3n} \hbar n^3} \quad (5.30)$$

which is approximately constant in the range  $n = 20, 30$ . For example for  $n = 25$  the saturation fluence is  $\cong 0.98 mJ/cm^2$  and for the same spot condition of the excitation from  $n = 1$  to  $n = 3$  the total laser energy results  $E = 174 \mu J$ .

In conclusion the Rydberg excitation can be performed by two near simultaneous laser pulses, whose bandwidths overlap the Doppler and the Stark broadening and whose energies satisfy the transition saturation conditions. In this context, the excitation efficiency of each transition is related to the laser spot dimension in comparison with the expanding Ps cloud and to the incoherent transition dynamics, which set a maximum value of 43% of excited atoms. On the contrary the global efficiency of the  $1 \rightarrow n$  transition is mainly related to the losses on the intermediate level population due to the not-negligible spontaneous emission as well as the efficiency of both the transitions. For this reason a maximum of 33% of excited Rydberg Ps was evaluated [106] in the limit of no losses, while a lower value of 30% was estimated by considering the spontaneous emission from the  $n = 3$  state. This value is however greater with respect to that obtained by considering the  $n = 2$  level as the intermediate level; in that case a 25% efficiency for the Rydberg excitation was found due to the higher spontaneous decay rate.

In the next section the first experimental results on Ps excitations in Rydberg states, performed by Ziock's [94] and Mill's [93] groups, will be described.

## 5.4 Ps Rydberg experiments

From an experimental point of view the first observation of resonant excitation of high- $n$  state in Ps was performed in 1990 by Ziock et al. [94], through a two-step excitation, where the  $n = 2$  state was used as an intermediate state.

The diagnostic of the first transition from  $n = 1$  to  $n = 2$  was based on the deviation of the Ps annihilation rate with laser on with respect to the Ps annihilation rate obtained with laser off. In fact, due to the presence of 200 G magnetic field, the sub-levels of  $n = 2$  resulted mixed by the Zeeman effect and the excitation from the ground state to  $n = 2$ , obtained using ultraviolet light, can populate both the triplet  $2^1P_1$  and the singlet state  $2^3P_{0,1,2}$ . Due to the lower lifetime of the singlet state (0.12 ns) an enhancing of the annihilation rate, proportional to the excited population in  $n = 2$ , was observed by plastic scintillators.

Moreover the excitations to the high- $n$  states were investigated by using at the same time two lasers (ultraviolet for  $1 \rightarrow 2$  excitation and red light for  $2 \rightarrow n$ ) and by discovering a reduction on the annihilation rate in comparison with the data obtained with only the first laser on. In fact when

the red light was on, a fraction of the  $n = 2$  population was excited to the long-living high- $n$  state, where it remained after the laser was turned off. In Fig.5.5 the three plots of the annihilation rate as a function of the time, as acquired with laser off, or with one/two lasers switched on, are reported. The subtraction between the data obtained with laser on and with laser off is shown in Fig.5.5b and 5.5c, where in particular in b) only one laser was on and in c) both lasers were on. The data were compared with theoretical results obtained from the solution of the time-dependent rate equation and the evidence of a Stark mixing effect, due to an electric field of  $\sim 40$  V/cm found in the chamber, was guessed. This effect could allow Ps to populate all  $l$  sub-levels, by breaking the dipole selection rule which would limit the number of populated sub-levels in the high- $n$  state to a value which is too small to explain the observed decrease in the annihilation rate.

In order to make a spectrum of Ps excited in different energy levels, some measurements at several wavelengths were performed. The plot in Fig.5.6 shows the ratio of counts in the laser on-laser off difference spectrum as a function of the second laser wavelengths. Resonances between  $n = 13$  and  $n = 19$  are recognizable and a large increase in the yield at resonance wavelengths for  $n = 13, 14, 15$  can be observed. The peaks of these resonances were found centered on the wavelength expected from theory. On the contrary at high  $n$ , the data seem to confirm the prediction of a  $n^{-3}$  scaling of the transition rate, even if the resonances peaks are not well evidenced, due to the low level of precision related to the laser linewidth of 0.4 nm.

After these preliminary results, the first efficient production of Rydberg Positronium was performed by Cassidy et al. in 2012. The same two-steps process of Ziock was carried out, in order to excite Ps to levels with  $n > 10$  [93]. The yield of formed Ps and its laser induced changes, were observed by using the single shot positron lifetime technique, in which the output from a fast gamma-ray detector is recorded as a function of the time. For each lifetime spectrum the delayed Ps annihilation fraction  $f_d$ , defined as the ratio between the annihilations in the time interval 50-300 ns [107] after the positron implantation and the total annihilations, was extracted and it was evaluated for different measurement conditions. As in Ziock the first excitation was detected using the magnetic quenching present in weak magnetic field, while the excitation to high- $n$  was performed for different values of the magnetic field.

In this way it was discovered that at higher magnetic flux density, a more efficient production of Rydberg Ps atoms occurs, because the magnetic quenching effect on the  $n = 2$  level decreases. On the contrary for magnetic field higher than 2 T, the amount of long-living Ps decreases, because the ground state singlet and triplets becomes mixed. Thus the optimum value of the

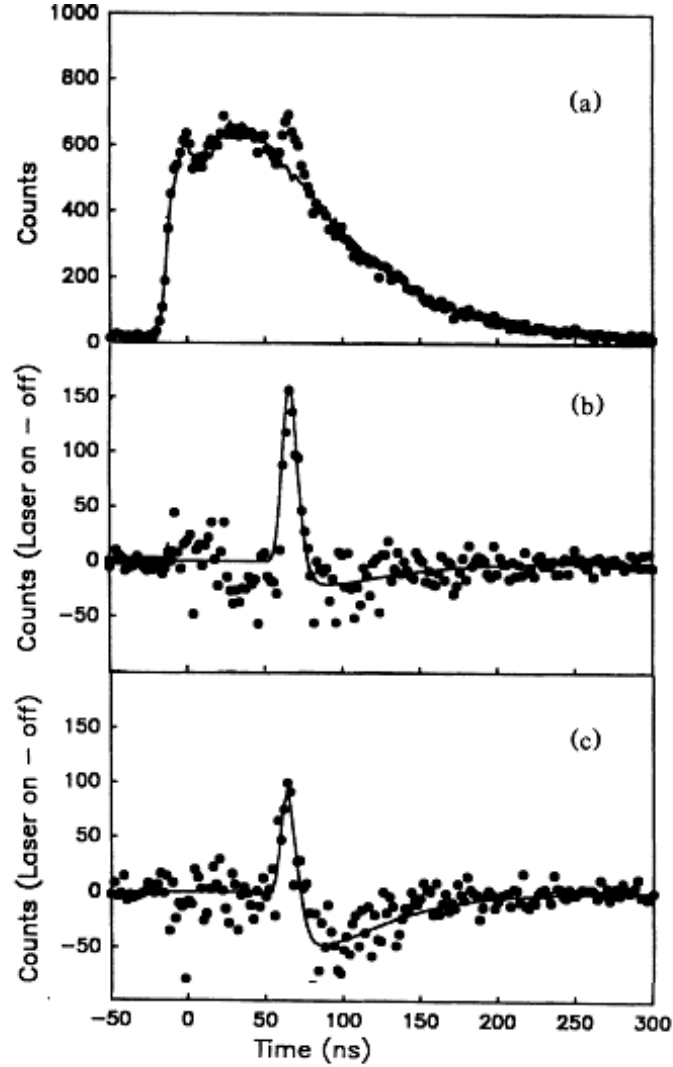


Figure 5.5: a) Ps annihilation rate as a function of the time (in ns) as recorded with ultraviolet laser on (solid circle), or with laser off (solid line). b) The subtraction between the data of plot a) obtained with laser on and with laser off. These counts are due to the Ps excited in  $n = 2$ . c) The subtraction between the data as obtained with two lasers on and with laser off is reported: a reduction on the annihilation rate with respect to plot b) is visible due to the Ps excitation at high- $n$  levels. After [94].

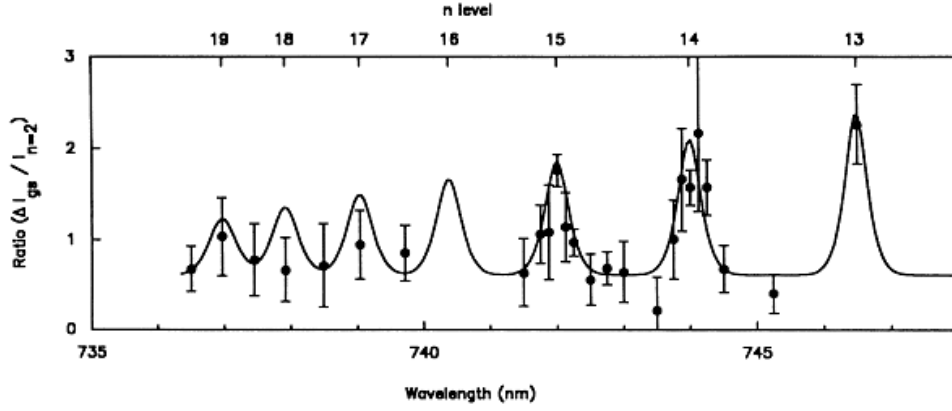


Figure 5.6: Ratio of counts in the laser on-laser off difference spectrum as a function of the wavelengths of the second laser used for the Rydberg excitation (where the first laser was the red light necessary for the  $1 \rightarrow 2$  transition). The excitations of Ps at high- $n$  levels from from  $n = 13$  to  $n = 19$  are recognizable and the resonance wavelength for  $n = 13, 14, 15$  can be extracted. After [94]

magnetic field for Ps excitation in Rydberg states was found to be 1 T where both the 2P and 1S quenching losses are reduced.

In Fig.5.7 the parameter  $f_{Ryd}$  as a function of the IR laser wavelength is shown. The  $f_{Ryd}$  is defined as  $f_{Ryd} = \frac{f_d[UV+IR] - f_d[UV]}{f_d[off] - f_d[UV]}$  where  $f_d[UV + IR]$ ,  $f_d[off]$ ,  $f_d[UV]$  refers to the delayed fraction measured with the UV and IR laser beam, or with laser off or with only UV light. Thus  $f_{Ryd}$  is equal to 1, if the change in the long-living Ps yield caused by the UV is completely counteracted by the application of the IR beam, that it occurs only when all the 2P states, that would otherwise have been quenched, are excited to long-living Rydberg states. In Fig.5.7 the transitions levels are resolved up to  $n = 18$ , while for higher  $n$  the excitations cannot be resolved. Anyway in these measurements the production of Rydberg Ps from 2P states was surprisingly found about 90 % efficient, in contrast with theoretical descriptions which set the limit to 40%. This result was explained by the fact that the high- $n$  Rydberg states can be split into Stark sub-levels by the motional and static electric fields, which inhibit the stimulated emission taken into account in the theoretical description. In this case in fact the IR light cannot drive the excited states back to the 2P level as it occurs for the 1S-2P excitation, which is almost 100% due to the removal of the excited atoms through the quenching effect.

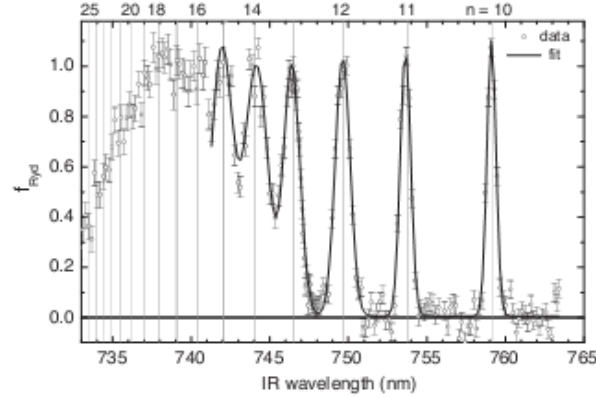


Figure 5.7: Line shapes scanned from 733 to 764 nm at  $B=0.16$  T. The solid line is a multi-Gauss fit. The vertical lines indicate the expected peak position for various Ps states based on the energy difference between each state and  $2^3P_1$  level. After [108].

## 5.5 Planned measurements on Rydberg Ps

The apparatus described in Chapter 4 will be used firstly for performing for the first time the Ps excitation at  $n = 3$  level and from  $n = 3$  to Rydberg states, then for other Ps spectroscopy studies, as will be described in the following.

The positron bunches will allow to use the single shot positron annihilation lifetime technique as in the Cassidy measurements, where fast detectors with few ns of time resolution has to be used.

Firstly the excitation to the  $n = 3$  level in presence of moderate magnetic field (200 Gauss), which induces the quenching effect, will be observed, by using fast  $PbF_2$  Cherenkov radiators, coupled to photomultipliers, whose time resolution is of the order of few nanoseconds.

On the contrary the excitations to high- $n$  levels will be observed in absence of magnetic field, by producing a proper electric field after the excitation, which ionizes only Ps at high  $n$  level. To produce this electric field two grids were designed near the sample (see section 4.4.3) to create an ionizing electric field greater than 1000 V/cm for observing the excitation to  $n = 10$ -20 levels and around 900 V/cm for the excitations to higher  $n$ . The positron and electron charges of ionized Ps will be collected by channeltrons or MicroChannelPlates, whose signal will be analyzed as a function of the wavelength and of the linewidth of the second laser pulse. In this way, for the first time, the detailed investigation of the complicated structure of Rydberg levels in absence of magnetic field could be achieved, compatibly with the

fact that the laser linewidth and the natural linewidth of each transition have to be small with respect to the distance between adjacent levels.

Moreover the study of Ps excitation to Rydberg levels in presence of the motional Stark effect could be a research of great interest in Atomic Physics also because it can be studied only in the case of an atom as light as positronium. In our system this effect can be studied without the magnetic field, by adding an electric field of suitable amplitude between the grids. In this way it will be possible to reproduce the electric field seen by the Ps in its reference frame, when in the laboratory frame a strong magnetic field ( $\sim 1$  T) is present. The required fields are of the order of hundreds of V/cm, thereby easily applicable with the same grid system used for the diagnostic in the Rydberg experiment.

An example of the set up of the sample chamber for this kind of measurements is shown in Fig.5.8. In particular the two horizontal windows facing each other, were provided for the laser beams, coming from the laser table just mounted in the AEgIS set up. In this way the light enters along the direction perpendicular to the sample surface for overlapping the Ps cloud. Then two grids will be placed at the top and at the bottom of the sample as in Fig.5.8 and their potentials will be switched on for Ps ionization after the laser excitation. The charge coming from the Ps ionization will be collected by channeltron detectors mounted behind the grids. At the same time fast  $PbF_2$  detectors could be used for observing the gamma ray annihilation from the Ps decay for a single shot annihilation lifetime spectroscopy.

Finally some measurements aimed to Ps spectroscopy and to atomic physics could be performed. Firstly the excitation between Rydberg levels could be studied by exciting  $n = 25$  level with laser system (UV and IR) and then driving higher states with a source of 200 GHz and less.

Secondly, while the excitation to  $n = 2$  is Doppler broadened and it prevents any study on the sub-levels and on the hyperfine structure, the investigation on the hyperfine structure with a 10 GHz-200 GHz microwave source on the level  $n = 3$  could be performed by observing singlet-triplet mixing in a magnetic field or by using laser with 20 GHz waves.

Finally the Ps laser cooling technique could be performed for the first time. In the next section a discussion about this interesting topic will be proposed.

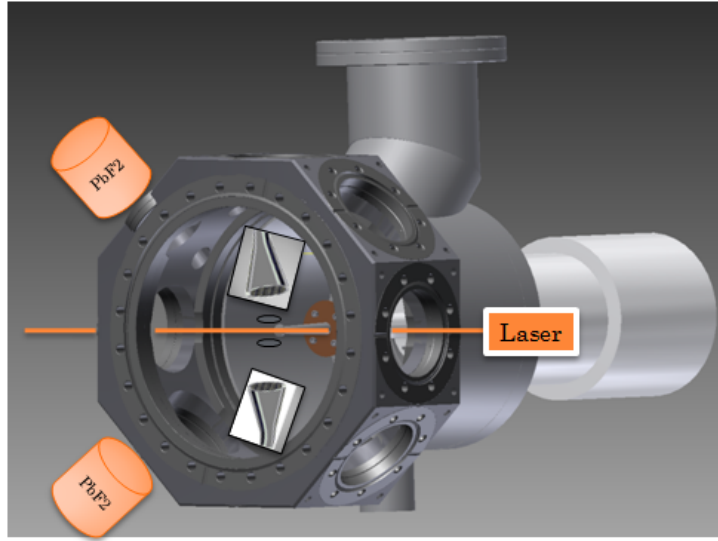


Figure 5.8: Set up of the sample chamber for Ps measurements: the laser beam and the two grids and channeltrons for Ps excitation and detection are shown. The other chamber flanges are suitable for fast detectors as  $PbF_2$  for acquiring positron annihilations signal.

## 5.6 Laser cooling technique

The laser cooling technique of ordinary atoms was theoretically proposed by Wineland [109] and Hänsch [110] in 1975 as a method of cooling atoms down to sub-K temperature. The mechanism is based on the absorption from the atom of photons coming from laser and on their reemission in random directions. Since the emitted photon has a higher average energy than the one it has absorbed from the laser, and because the velocity of the atom is greater than the recoil velocity of photon emission, the overall atom velocity is reduced. Due to the Doppler effect on moving atoms, the laser frequency for the absorption has to be tuned slightly below the electronic transition in an atom and thus the technique selectively targets atoms moving in certain directions and at certain speeds. In fact when the apparent frequency of the light with respect to certain atoms is correctly tuned, the atom absorbs the incoming photons, losing a momentum equal to the momentum of the photon and after temporarily reaching an excited state, it emits a photon. For each absorbed photon an atom receives a net change of momentum  $\Delta p = \frac{h}{\lambda}$  where  $\lambda$  is the laser wavelength, but, since the reemission of the photon has no preferred direction, an average of many scattering events gives a net scattering force along the direction of the light [111]. In conclusion the atom

performs a random walk in momentum space with steps equal to the photon momentum. This fact constitutes a heating effect, which counteracts the cooling process and imposes a limit on the amount by which the atom can be cooled. By considering this effect, the minimum temperature  $T_{recoil}$  can be calculated as the temperature associated with the momentum gain from the spontaneous emission of a photon and it can be written by considering that each translational degree of freedom has an energy  $\frac{1}{2}k_B T_{rec}$ :

$$\frac{1}{2}k_B T_{rec} = \frac{p^2}{2m} \quad (5.31)$$

where  $p$  is the momentum of the emitted photon.

Moreover, a further limitation on the minimum temperature is imposed by the finite frequency linewidth of the optical transition used for cooling. In fact the frequency linewidth given by the Doppler effect or by the natural linewidth of the transition limits the velocity discrimination, and therefore the final temperature.

In particular considering Ps with a thermal velocity distribution, whose standard deviation in the x direction is  $\sigma_x = \sqrt{\frac{k_B T}{m_{Ps}}}$ , the first-order Doppler full width at half-maximum of the Gaussian distribution of the resonant frequencies is:

$$\Delta\nu(T) = 2\sqrt{2\ln 2}\nu_0 \sqrt{\frac{k_B T}{m_{Ps}c^2}} \quad (5.32)$$

Being the natural linewidth of the cooling transition  $\Gamma/h$  where  $\Gamma = \frac{1}{\tau}$  and  $\tau$  is the lifetime of the excited state, a temperature  $T_{linewidth}$  such as the Doppler width is equal to the natural linewidth can be defined [112]. It can be calculated by the eq.  $\Delta\nu(T_{linewidth}) = \Gamma$  and it results:

$$T_{linewidth} = \frac{\Gamma^2 m_{Ps} c^2}{8\ln 2 \nu_0^2 h^2 k_B} \quad (5.33)$$

and if  $T > T_{linewidth}$  the Doppler width is larger with respect to the natural linewidth.

Finally a further limit on a random walk cooling process is due to the Doppler limit and it is expressed by the temperature  $T_{Dopp}$  [113] such as:

$$k_B T_{Dopp} = \frac{1}{2}\Gamma \quad (5.34)$$

Depending on atoms,  $T_{rec}$  can be lower or higher than  $T_{Dopp}$  and the higher of them can be achieved by laser cooling, while further cooling techniques have to be applied for reaching a lower temperature.

Furthermore only certain atoms and ions, which have optical transitions with wavelength larger than 300 nm, can be compatible with the laser cooling process, since it is extremely difficult to generate the amounts of laser power needed at wavelengths much shorter than 300 nm. At the same time, for many atoms the technique could have low efficiency, because, depending on the energy levels of each atoms, it could be possible for an atom emit a photon from the upper state and not return to its original state, putting it in a dark state and removing it from the cooling process.

The laser cooling technique was performed for the first time on Mg ions trapped in a Penning trap [114] and cooled down to less than 40 K, and then it was unambiguously demonstrated by Chu et al. in 1985 on a cloud of Na atoms that was confined in a set of three mutually perpendicular counter-propagating pairs of laser beams, used to cool the three motional degrees of freedom of the atom. Nowadays laser cooling of positive atomic ions [115, 116] and neutral atoms [117] is commonplace and has opened the avenue for exciting new fields such as particles condensation [118] and ion crystals [119].

On the contrary laser cooling on Ps has not yet been demonstrated experimentally, although it has been theoretically predicted [55]. In general for ordinary atoms the three temperatures  $T_{linewidth}$ ,  $T_{Dopp}$ ,  $T_{rec}$  are in a descendant order, while for Ps the situation is different. Ps in fact is the opposite of ordinary atoms because the Doppler cooling limit is  $T_{Dopp} = 2.4\text{mK}$ , while the recoil limit is a much higher temperature  $T_{rec} = 295\text{ mK}$  [112] and for this reason this last temperature could be reached by the laser cooling technique. Moreover Ps would have to be less than  $T_{linewidth} = 14\text{ mK}$  temperature before the first-order Doppler width would be less than the natural linewidth. The Ps laser cooling was proposed by utilizing the 1S-2P transition, whose energy interval is 5.1 eV corresponding to the wavelength of 243 nm [120]. As a lifetime of the spontaneous transition from the 2P to 1S is 3.2 ns, it is difficult to cool p-Ps by the Doppler cooling due to the extremely short lifetime (0.125 ns). On the contrary by considering o-Ps, an estimation of the number of cooling cycles (excitation-de-excitation) occurring on each atoms can be done. In fact by assuming that the intensity of the cooling laser is high enough to have half of o-Ps atoms occupying the 2P state, spontaneous emission occurs every 6.4 ns, which is two times longer than the lifetime of the spontaneous transition from the 2P state. At the same time the o-Ps lifetime is doubled i.e. 284 ns, because the direct decay time (100  $\mu\text{s}$ ) from the 2P state is extremely longer than that from the 1S state and thus the excitation-de-excitation cycle occurs no faster than two spontaneous emission lifetimes (one lifetime to excite at the saturation intensity and one lifetime to de-excite) and in the case of the 1S-2P transition it takes 6.4 ns. In con-

clusion since each cycle takes 6 ns, the cooling of Ps with lifetime of 284 ns will give 50 transitions on the average. If the cooling process would be performed in the three dimensional coordinate each cartesian momentum could be reduced of around  $17h\nu_0/c$  and thus the recoil limit would be achieved starting from Ps with a temperature of  $300 T_{recoil}$  (about 90 K).

In the new apparatus the laser cooling process could be investigated by using three open ports especially built, which are facing toward the sample surface with an angle of  $45^\circ$ , along three orthogonal axes. In absence of electric and magnetic field the pre-cooled Ps, going out from the porous sample, could be further cooled down and the velocity distribution could be monitored by observing changes on the spatial distribution of the Ps cloud, monitored by MultiChannelPlate coupled to a phosphor screen.

This cooling technique is difficult to be implemented in the AEgIS experiment, because the presence of high magnetic field could prevent the 1S-2P transition, due to the change of the Ps levels and to the magnetic quenching. In any case its experimental realization would be an extremely important result in the Ps and atoms physics, because it could open the port to Ps spectroscopy measurements with high level of accuracy and to new perspectives in the Bose Einstein Condensation field.



## Appendix A

### Transmission line theory

Depending on the frequency of the electromagnetic signal with respect to the circuit dimensions in which the signal propagates, a system can be described by the lumped parameter circuit theory or by the transmission line theory [121]. Even if an electromagnetic signal can have an arbitrary time dependance, it is well known that it can be represented as a summation of sinusoidal fields with frequencies contained in a certain band (Fourier theorem). Thus the size  $L$  of the structures with which the electromagnetic field

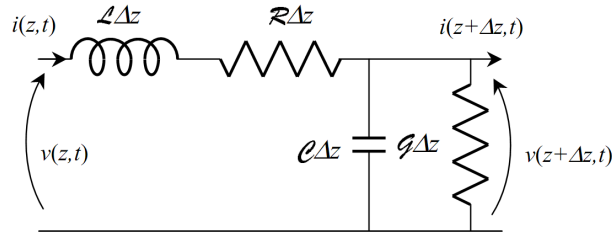


Figure A.1: Equivalent circuit of a transmission line.

interacts must always be compared with its wavelength  $\lambda$  (as obtained by the Fourier analysis in the frequency domain). Differently from a lumped circuit, which is made of elements of negligible electrical size, a transmission line, as schematized in Fig.A.1, has larger dimensions. Here the propagation time of the signal has to be considered and the propagation delay is not negligible with respect to the period of the oscillations of the signal. Typically, in the transmission lines, the transverse size is small with respect to wavelength but the length can be very large. Thus, while a lumped parameter circuit is modeled as point like, a transmission line is a one dimensional system, in which voltage and current depend on time and on a longitudinal coordinate that will be indicated with  $z$ . The state variables of such a system are then

the voltage  $v(z; t)$  and the current  $i(z; t)$ . A circuit, containing transmission lines is often called *distributed parameter circuit* to underline the fact that electromagnetic energy is not only stored in specific components, e.g. inductors, capacitors, but also in the space surrounding the conductors of a line. As shown in Fig.A.1 in each transmission line four parameters per unit of length have to be considered.

Firstly in order to take into account the linked magnetic flux, produced by the current flowing in the conductor, the inductance of the element has to be considered. It can be expressed as  $\mathcal{L}\Delta z$  where  $\mathcal{L}$  (measured in H/m) is the inductance per unit length of the line.

Secondly, through the resistance value  $\mathcal{R}\Delta z$ , the power dissipated in the metal conductor is taken into account, where  $\mathcal{R}$  (measured in  $\Omega/\text{m}$ ) is the resistance per unit length.

Thirdly, as a consequence of the potential difference maintained between the inner and outer conductors, a charge is induced on them and thus the capacitance of the element,  $\mathcal{C}\Delta z$ , where  $\mathcal{C}$  (measured in F/m) is the capacitance per unit length of the line, has to be calculated.

Finally, since the dielectric between the conductors has a non zero conductivity, a current flowing from one conductor to the other through the insulator can be present and this phenomenon is accounted by the conductance  $\mathcal{G}\Delta z$ , (measured in S/m), where  $\mathcal{G}$  is the conductance per unit length of the line. However, the line can be subdivided in a large number of sufficiently short elements, and a lumped equivalent circuit for each of them can be derived and then analyzed by the usual methods of the circuit theory.

Therefore by applying the Kirchhoff laws to each part  $\Delta z \ll \lambda$ , the equations:

$$\begin{cases} v(z, t) - v(z + \Delta z, t) = \mathcal{R}\Delta z i(z, t) + \mathcal{L}\Delta z \frac{\partial i(z, t)}{\partial t} \\ i(z, t) - i(z + \Delta z, t) = \mathcal{G}\Delta z v(z + \Delta z, t) + \mathcal{C}\Delta z \frac{\partial v(z + \Delta z, t)}{\partial t} \end{cases} \quad (\text{A.1})$$

can be obtained.

By letting  $\Delta z \rightarrow 0$ , these lead to coupled equations:

$$\begin{cases} -\frac{\partial v(z, t)}{\partial z} = \mathcal{R}i(z, t) + \mathcal{L}\frac{\partial i(z, t)}{\partial t} \\ -\frac{\partial i(z, t)}{\partial z} = \mathcal{G}v(z, t) + \mathcal{C}\frac{\partial v(z, t)}{\partial t} \end{cases} \quad (\text{A.2})$$

called General Transmission Line Equations or Telegrapher's equations.

Fourier transforming both sides, the real transmission line equations are

obtained in the spectral domain,

$$\begin{cases} -\frac{dV(z, \omega)}{dz} = (\mathcal{R} + j\omega\mathcal{L})I(z, \omega) \\ -\frac{dI(z, \omega)}{dz} = (\mathcal{G} + j\omega\mathcal{C})V(z, \omega) \end{cases} \quad (\text{A.3})$$

where the relation between the voltage and current in the time domain is:

$$\begin{cases} v(z, t) = \text{Re}\{V(z)e^{j\omega t}\} \\ i(z, t) = \text{Re}\{I(z)e^{j\omega t}\} \end{cases} \quad (\text{A.4})$$

After decoupling:

$$\begin{cases} -\frac{d^2V(z)}{dz^2} = \gamma^2 V(z) \\ -\frac{d^2I(z)}{dz^2} = \gamma^2 I(z) \end{cases} \quad (\text{A.5})$$

with  $\gamma = \sqrt{(\mathcal{R} + j\omega\mathcal{L})(\mathcal{G} + j\omega\mathcal{C})}$  the solutions can be obtained:

$$\begin{cases} V(z, \omega) = V_0^+(\omega)e^{-jkz} + V_0^-(\omega)e^{jkz} \\ I(z, \omega) = Y_\infty V_0^+(\omega)e^{-jkz} - Y_\infty V_0^-(\omega)e^{jkz} \end{cases} \quad (\text{A.6})$$

The solutions, as expressed in eq.A.6, are the sum of two forward and backward travelling waves, whose amplitudes at  $z = 0$  are  $V_0^+$ ,  $V_0^-$  and  $Y_\infty V_0^+$ ,  $Y_\infty V_0^-$  and  $k$  is the complex *propagation factor*:

$$k = \omega \sqrt{\left(\mathcal{L} - j\frac{\mathcal{R}}{\omega}\right) \left(\mathcal{C} - j\frac{\mathcal{G}}{\omega}\right)} \quad (\text{A.7})$$

while  $Y_\infty$  is the complex *characteristic admittance*:

$$Y_\infty = \frac{1}{Z_\infty} = \sqrt{\frac{\mathcal{G} + j\omega\mathcal{C}}{\mathcal{R} + j\omega\mathcal{L}}} \quad (\text{A.8})$$

where  $Z_\infty$  is defined as the *characteristic impedance* of the line.

By anti-trasforming the solutions of eq.A.6 and by considering a monochromatic signal with frequency  $\omega_0$ , the  $v(z, t)$  and  $i(z, t)$  functions can be obtained:

$$\begin{cases} v(z, t) = |V_0^+| \cos(\omega_0 t - k_0 z + \arg(V_0^+)) + \\ \quad |V_0^-| \cos(\omega_0 t + k_0 z + \arg(V_0^-)) \\ i(z, t) = Y_\infty |V_0^+| \cos(\omega_0 t - k_0 z + \arg(V_0^+)) \\ \quad - Y_\infty |V_0^-| \cos(\omega_0 t + k_0 z + \arg(V_0^-)) \end{cases} \quad (\text{A.9})$$

where  $k_0 = \omega_0 \sqrt{\mathcal{L}\mathcal{C}}$ . Thus from eq.A.9 the propagation velocity of a wave (phase velocity), that is the velocity that an observer must have in order to see the wave phase unchanging, can be calculated as  $v_{ph} = \frac{1}{\sqrt{\mathcal{L}\mathcal{C}}}$ .

In conclusion the general solution of the transmission line equations is expressed as linear combination of two waves, a forward one propagating in the direction of increasing  $z$  and a backward one, moving in the opposite direction. Each wave is made of voltage and current that, in a certain sense, are the two sides of a same coin. Moreover it is worth observing that the two waves are absolutely identical since the transmission line is uniform and hence is reflection symmetric. In fact the proportional factor between voltage and current of the same wave (called impedance relationship)

$$I_0^+(\omega) = Y_\infty V_0^+(\omega) \quad \text{and} \quad I_0^-(\omega) = -Y_\infty V_0^-(\omega) \quad (\text{A.10})$$

is only apparently different in the two cases, because the minus sign in the impedance relation for the backward wave arises because the positive current convention of the forward wave is used also for the backward one.

Therefore forward and backward waves on the line are the two normal modes of the system and they are independent (uncoupled) if the line is of infinite length, whereas they are in general coupled by the boundary conditions (generator and load) if the line has finite length.

Thus along the line the *local voltage reflection coefficient* in a point  $z$  can be defined as the ratio of the backward and forward voltages in that point:

$$\Gamma(z) = \frac{V^-(z)}{V^+(z)} = \frac{V_0^- e^{jkz}}{V_0^+ e^{-jkz}} = \Gamma_0 e^{2jkz} \quad (\text{A.11})$$

where  $\Gamma_0 = \frac{V_0^-}{V_0^+}$ .

By considering a transmission line with characteristic impedance  $Z_\infty$  and phase constant  $k$ , loaded by the impedance  $Z_L$ , excited by a generator that produces a forward wave incident on the load as:

$$\begin{cases} V^{inc}(z) = V_0^+ e^{-jkz} \\ I^{inc}(z) = Y_\infty V_0^+ e^{-jkz} \end{cases} \quad (\text{A.12})$$

the voltage reflection coefficient between the line and the load can be calculated.

If the load coordinate is  $z = 0$  the relation  $V(0) = Z_L I(0)$  is forced and thus it is satisfied only if  $Z_L = Z_\infty$ , because in this case the forward wave alone is capable of satisfying the boundary condition. On the contrary if the load impedance is arbitrary, necessarily a backward (reflected) wave must be

generated on the load.

$$\begin{cases} V^{ref}(z) = V_0^- e^{+jkz} \\ I^{ref}(z) = -Y_\infty V_0^- e^{+jkz} \end{cases} \quad (\text{A.13})$$

with a suitable amplitude  $V_0^-$  in such a way that the total voltage and current, sum of the forward and backward components, satisfy the boundary condition:

$$V^{inc}(0) + V^{ref}(0) = Z_L(I^{inc}(0) + I^{ref}(0)) \quad (\text{A.14})$$

From the eq.A.14, by substituting eq.A.12 and eq.A.13,  $V_0^-$  is immediately deduced as:

$$V_0^- = \frac{Z_L Y_\infty - 1}{Z_L Y_\infty + 1} V_0^+ \quad (\text{A.15})$$

and the reflection coefficient that relates the backward voltage to the forward one called voltage reflection coefficient between the line and the load is calculated as:

$$\Gamma_L = \frac{Z_L - Z_\infty}{Z_L + Z_\infty} \quad (\text{A.16})$$



# Bibliography

- [1] G.Baur et al., *Phys. Rev. B* **368**, 251 (1996)
- [2] G.Blanford et al., *Phys. Rev. Lett.* **80**, 14 (1998)
- [3] Athena coll., *Nature* **419**, 456 (2002)
- [4] Atrap coll. *Phys. Rev. B* **578**, 2332 (2004)
- [5] Alpha coll., *Nature* **468**, 673 (2010)
- [6] M.Hori et al., *Nature* **475** 484 (2011)
- [7] R.S.Van Dyck,Jr, P.B.Schwinberg and H.G.Dehmelt, *Phys. Rev. Lett.* **59**, 26 (1987)
- [8] G.Gabrielse et al., *Phys. Rev. Lett.* **82**, 3198 (1999)
- [9] Particle Data Group, *Phys. Rev. B* **667**, 1 (2008)
- [10] M.Hayakawa, *Proceedings of the First KEK meeting on CP violation and its origin*, Vol. **193**, (1997)
- [11] T.Pask et al., *Phys. Rev. B* **678**, 55 (2009)
- [12] C.G. Parthey et al., *Phys. Rev. Lett.* **107**, 203001 (2011)
- [13] B.Heckel et al., *Adv. Space Res.* **25**, 1225 (2000)
- [14] M.Nieto et al., *Phys. Rep.* **205**, 5 (1991)
- [15] S.Bellucci, V.Faraoni *Phys. Rev. D* **49**, 2922 (1994)
- [16] G.Steigman, *Annu. Rev. Astron. Astrophys.* **14**, 339 (1976)
- [17] M.H.Holzschneider et al., *Nucl. Phys. A* **558**, 709c (1993)

- [18] W.M.Fairbank et al., in B.Bertotti, *International School of Physics Enrico Fermi*, Academic Press, New York (1974)
- [19] C.H.Storry et al., *Phys. Rev. Lett.* **93**, 263401 (2004)
- [20] M.Charlton, *Phys. Rev. A* **143**, 143 (1990)
- [21] Athena coll. *Nucl. Instr. and Meth. A* **518** 679 (2004)
- [22] L.Spitzer *Physics of fully ionized gases* Interscience Pub. New York (1962)
- [23] W.M.Itano et al., *Phys. Scripta* **1995**, 106 (1995)
- [24] Alpha coll., *Phys. Rev. Lett.* **105**, 013003 (2010)
- [25] A.P.Mills Jr. and E.M.Gullikson, *Appl. Phys. Lett.* **49**, 1121 (1986)
- [26] C.Canali et al., *Eur. Phys. Jour. D* **65**, 499 (2011)
- [27] AEgIS coll. *Proposal for the AEgIS experiment at the CERN antiproton decelerator* CERN-SPSC (2008)
- [28] E.Vliegen and F. Merkt, *J. Phys. B: At. Mol. Opt. Phys.* **39**, L241 (2006)
- [29] M.K.Oberthaler et al., *Phys. Rev. A* **54**, 3165 (1996)
- [30] M.Deutsch, *Phys. Rev. Lett.* **83**, 866 (1951)
- [31] S.Berko and H.N.Pendleton, *Annu. Rev. Nucl. Part. Sci.* **30**, 543 (1980)
- [32] K.F.Canter, A.P.Mills,Jr and S.Berko, *Phys. Rev. Lett.* **33**, 7 (1974)
- [33] J.A.Baker and P.G.Coleman, *J. Phys. C* **21**, L875 (1988)
- [34] P.J.Shultz and K.G.Lynn, *Rev. Mod. Phys.* **60**, 701 (1988)
- [35] S.Valkealahti and R.M.Nieminen, *Appl. Phys. A* **35**, 51 (1984)
- [36] K.G.Lynn, P.J.Schultz, *Phys. Rev. B* **22**, 1 (1960)
- [37] R.S.Brusa and A.Dupasquier *Positronium emission and cooling in Physics with Many Positrons, Proceeding of the Interantional School of Physics "Enrico Fermi"*, Course CLXXIV, edited by A.Dupasquier and A.P.Mills jr and R.S.Brusa IOS,Amsterdam, Oxford,Tokio,Washington DC p.245 (2010)

- [38] D.W.Gidley et al. *Phys. Rev. Lett.* **58**, 595 (1987)
- [39] T.D.Steinger and R.S.Conti, *Phys. Rev. B* **34**, 3069 (1986)
- [40] A.P.Mills Jr, L.Pfeiffer and P.M.Platzmann *Phys. Rev. Lett* **51**, 1085 (1983)
- [41] K.G.Lynn and D.O.Welch, *Phys. Rev. B* **22**, 99 (1980)
- [42] A.P.Mills, Jr and L. Pfeiffer, *Phys. Rev. B* **32**, 53 (1985)
- [43] A.P.Mills, Jr., M.Leventhal, P.M.Platzman, R.J.Chichester et al., *Phys. Rev. Lett.* **66**, 6 (1991)
- [44] A.P.Mills, Jr., E.D.Shaw, M.Leventhal, P.M.Platzman, R.J.Chichester et al., *Phys. Rev. Lett.*, **66** 735 (1991)
- [45] A.P.Mills Jr, *Phys. Rev. Lett.* **41**, 1828 (1978)
- [46] M.Eldrup, A.Vehanen, P.J.Schultz and K.G.Lynn *Phys. Rev. B* **32**, 7048 (1985)
- [47] M.Eldrup, A.Vehanen, P.J.Schultz and K.G.Lynn *Phys. Rev. Lett.* **51**, 2007 (1983)
- [48] P.Sferlazzo, S.Berko, K.F.Canter *Phys. Rev. B*, **35**, 5315 (1987)
- [49] Y. Nagashima et al., *Phys. Rev. B* **58**, 19 (1998)
- [50] A.P.Mills Jr., E.D.Shaw, R.J.Chichester and D.M.Zuckerman, *Phys. Rev. B* **40**, 2045 (1989)
- [51] I.A.Ivanov and J.Mitroy, *Phys. Rev. A* **75**, 034709 (2002)
- [52] A.Dupasquier et al, *Positrons in solids* Springer-Verlag, Berlin (1979)
- [53] R.S.Vallery, P.W.Zitzewitz and D.W.Gidley, *Phys. Rev. Lett.* **90**, 203402 (2003)
- [54] N.N.Mondal et al., *Appl. Surf. Scien.* **149**, 269 (1999)
- [55] E.P. Liang, C.D. Dermer, *Opt Common.* **65**, 419 (1988)
- [56] D.W.Gidley, W.E.Frieze et a., *Phys. Rev. B* **60**, R5157 (1999)
- [57] C.He, T.Ohdaira, N.Oshima, et al., *Phys. Rev. B* **75**, 195404 (2007)

- 
- [58] S.Mariazzi et al., *Phys. Rev. B* **81**, 235418 (2010)
- [59] O.Bisi, S.Ossicini and L.Pavesi, *Surf. Scien. Rep.* **38** (2000)
- [60] A.Halimaoui, in: L.T. Canham (Ed.), *Properties of Porous Silicon*, IEE INSPEC, The Institution of Electrical Engineers, London, p. 12 (1997)
- [61] G.Brauer et al., *Phys. Rev. B* **66**, 195331 (2002)
- [62] E.Soininen, A.Schwab and K.G.Lynn, *Phys. Rev. B* **43**, 10051 (1991)
- [63] S.Mariazzi, P.Bettotti and R.S.Brusa *Phys. Rev. Lett.* **104**, 243401 (2010)
- [64] S.Mariazzi, A.Salemi and R.S.Brusa, *Phys. Rev. B* **78**, 085428 (2008)
- [65] Y.Nagai et al., *Phys. Rev. B* **62**, 5531 (2000)
- [66] P.Crivelli, U.Gendotti A.Rubbia, L.Liszkay, P.Perez and C.Corbel, *Phys. Rev. A* **81**, 052703 (2010)
- [67] D.B.Cassidy, P.Crivelli, T.H.Hisakado, L.Liszkay, V.E.Meligne, P.Perez, H.W.K.Tom, and A.P.Mills,Jr., *Phys. Rev. A* **81**, 012715 (2010)
- [68] C.Hugenschmidt, B.Löwe, J.Mayer, C.Piochacz, P.Pikart, R.Repper, M.Stadlbauer, K.Schreckenbach, *Nucl. Instr. Meth A* **593**, 616 (2008)
- [69] C.Hugenschmidt, T.Brunner, S.Legl, J.Mayer, C.Piochaz, M.Stadlbauer and K.Schreckenbach, *Phys.Status Solidi C* **4**, 3947 (2007)
- [70] C.Hugenschmidt, G.Kögel, R.Repper, K.Schreckenbach, P.Sperr, B.Straßer and W.Triftshäuser, *Nucl. Instr. Meth B*, **192**, 97 (2002)
- [71] C.Hugenschmidt, G.Kögel, R.Repper, K.Schreckenbach, P.Sperr, B.Straßer and W.Triftshäuser, *Nucl. Instr. and Meth. B* **221**, 160 (2004)
- [72] C.Piochaz, G.Kögel, W.Egger, C. Hugenschmidt, J.Mayer, K.Schreckenbach, P.Speer and W.Triftshäuser, *Appl. Surf. Sci.*, **255**, 98 (2008)
- [73] C.Hugenschmidt, J.Mayer and K.Schreckenbach, *Surf. Sci.* **601**, 2459 (2007)

- [74] C.Hugenschmidt, N.Qi, M.Stadlbauer and K.Schreckenbach, *Phys. Rev. B* **80**, 224203 (2009)
- [75] P.Sperr, W.Egger, G.Kögel, G.Dollinger, C.Hugenschmidt, R.Repper and C.Piochaz, *Appl. Surf. Sci.* **255**, 35 (2008)
- [76] A.David, G.Kögel, P.Speer and W.Triftshäuser, *Phys. Rev. Lett.* **87**, 067402 (2001)
- [77] C.Hugenschmidt, *Positron sources and positron beam* in *Physics with Many Positrons, Proceeding of the Interantional School of Physics “Enrico Fermi”*, Course CLXXIV, edited by A.Dupasquier and A.P.Mills jr and R.S.Brusa IOS,Amsterdam, Oxford,Tokio,Washington DC p.399 (2010)
- [78] P.G.Coleman, *In Positron Beams and their applications*, edited by P.G.Coleman, World Scientific, Singapore-New Jersey-London Hong Kong 2000, p. 1-40
- [79] L. D. Landau and E. M. Lifschitz. *The classical theory of fields*, Pergamon Press, Vol.2, page 55 (1971).
- [80] [www.comsol.com](http://www.comsol.com)
- [81] P.Grivet, M.Y.Bernard, *Electron Optics*, Pergamon Press (1972)
- [82] P.Dahl, *Introduction to electron optics*, Academic Press, New York, NY (1973)
- [83] [www.simion.com](http://www.simion.com)
- [84] G.D.Ruano, J.Ferron, R.R.Koropecski *Jour.Phys.: Conf. Ser.* **167**, 012006 (2009)
- [85] N.M.Laurendeau, *Statistical thermodynamics: fundamentals and applications*, Cambridge University Press., p.434 (2005)
- [86] A.Zecca, M.Bettonte, G.P.Karwasz and R.S.Brusa, *Meas. Sci. Technol.* **9**, 409 (1998)
- [87] R.S.Brusa, G.P.Karwasz, M.Bettonte and A.Zecca, *Appl. Surf. Sci.* **116**, 59 (1997)
- [88] D.B.Cassidy, T.H.Hisakado, V.E.Meligne, H.W.K.Tom, and A.P.Mills,Jr., *Phys. Rev. A* **82**, 052511 (2010)

- 
- [89] G.F.Knoll, *Radiation detection and measurement*, John Wiley & Sons Inc. (2010)
  - [90] L.V.Jorgensen et al., (Athena collaboration), *Phys. Rev. Lett.* **95**, 025002 (2005)
  - [91] J.D.Jackson, *Elettrodinamica Classica*, Zanichelli, (2001)
  - [92] A.P.Mills,Jr, *Appl.Phys.***22**, 273 (1980)
  - [93] D.B.Cassidy, S.M. Deng, R.G.Greaves, A.P.Mills, Jr, *Rev.Sci.Instrum.* **77**, 073106 (2006)
  - [94] K.P.Ziock, R.Howell, F.Magnotta, R.A.Failor and K.M.Jones, *Phys. Rev. Lett* **64**, 2366 (1990)
  - [95] H.A.Bethe, E.E.Salpeter, *Quantum mechanics of one- and two-electrons atoms*, Dover Publications (1977)
  - [96] C.D.Dermer and J.C.Weisheit, *Phys.Rev.A***40**, 5526 (1989)
  - [97] R.H. Garstang, *Rep. Prog. Phys.* **40**, 105 (1977)
  - [98] F.Castelli, *Eur. Phys. J. Special Topics* **203**, 137 (2012).
  - [99] S.M. Curry, *Phys. Rev. A* **7**, 447 (1973)
  - [100] A. Rich, *Rev. Mod. Phys.* **53**, 127 (1981)
  - [101] O. Halpern, *Phys. Rev. Lett* **94**, 904 (1954)
  - [102] F.Villa, *Laser system for Positronium excitation to Rydberg levels for Aegis experiment*, Ph.D thesis (2011)
  - [103] T.F.Gallager, *Rydberg Atoms*, Cambridge University Press (2005)
  - [104] F.Castelli, I.Boscolo, S.Cialdi, M.G.Giammarchi and D.Comparat, *Phys.Rev. A* **78**, 052512 (2008)
  - [105] K. Blushs and M. Auzinsh, *Phys. Rev. A* **69**, 063806 (2004)
  - [106] B.W.Shore, *The Theory of Coherent Atomic Excitation*, Wiley, New York (1990)
  - [107] D.B.Cassidy, T.H.Hisakado, H.W.K.Tom and A.P.Mills Jr. *Phys. Rev. B* **84**, 195312 (2011)

- [108] D.B.Cassidy, T.H.Hisakado, H.W.K.Tom and A.P.Mills,Jr. *Phys. Rev. Lett.* **108**, 043401 (2012)
- [109] Wineland, D.J.; H. Dehmelt, *Bull. Am. Phys. Soc.* **20**, 637 (1975)
- [110] Hänsch, T.W.; A.L. Shawlow, *Optics Communications* **13**, 68 (1975)
- [111] A.Ashkin, *Phys. Rev. Lett.* **24**, 156 (1970)
- [112] A.P.Mills,jr *Physics with many positrons* in *Physics with Many Positrons, Proceeeding of the Interantional School of Physics “Enrico Fermi”*, Course CLXXIV, edited by A.Dupasquier and A.P.Mills jr and R.S.Brusa IOS,Amsterdam, Oxford,Tokio, Washington DC p.77 (2010)
- [113] J.P.Gordon, A.Ashkin, *Phys.Rev.A* **21**, 1606 (1980)
- [114] D.J.Wineland, R.E. Drullinger and F.L. Walls, *Phys. Rev. Lett.* **40**, 1639(1978)
- [115] D. J. Wineland, R. E. Drullinger, and F. L. Walls, *Phys. Rev. Lett.* **40**, 1639 (1978)
- [116] W. Neuhauser, M. Hohenstatt, P. Toschek, and H. Dehmelt, *Phys. Rev. Lett.* **41**, 233 (1978)
- [117] S. Chu et al., *Phys. Rev. Lett.* **55**, 48 (1985)
- [118] M.H. Anderson et al., *Science* **269**, 198 (1995)
- [119] A. Fischer, C. Canali, U. Warring, and A. Kellerbauer, *Phys. Rev. Lett.* **104**, 073004 (2010)
- [120] H.Iijima, T.Asonuma, T.Hirose, M.Irako, T.Kumita, M.Kajita, K.Matsuzawa, K.Wada *Nucl. Instr. Meth. A* **455**, 104 (2000)
- [121] P.C. Magnusson, G.C. Alexander, V.K. Tripathi, *Transmission lines and wave propagation*, Boca Raton, CRC (1992)



## Ringraziamenti

Era l'8 dicembre 2010 quando ho deciso di partire per fare il dottorato a Trento. Ci tenevo a fare il dottorato, ma non avevo proprio mai considerato di partire da casa!

Grazie a Roberto per avermi accolto nel suo laboratorio e per le numerose letture e correzioni di questa tesi.

Grazie agli amici del laboratorio, con cui ho condiviso ogni giorno della vita trentina: a Marco e a Seba, persone speciali ed esempi da cui imparare, a Luca per la pazienza nel lavorare con me, all'Ale per la sua originalità e simpatia, a Max per i numerosi consigli di inglese, a Mara per tutti gli spritz e brulé che mi ha fatto bere e al Pippa per il supporto informatico.

Grazie al Tenga, compagno di vita da dottorandi.

Grazie a chi é venuto a trovarmi: alla nostra Steffy, sempre presente, agli zii (c'è mancato poco!) e a Marco e Ida.

Grazie a chi mi ha lasciato partire...al mio Carlo, che mi é stato vicino (anche da lontano!!), a Paoletta, che fa sempre il tifo per me e alla mia mamma, che mi vuole un bene infinito.

*“...senza qualcuno, nessuno può diventare uomo...”*