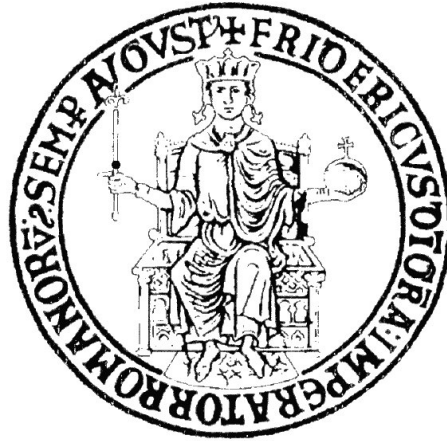




UNIVERSITÀ DEGLI STUDI DI NAPOLI
“FEDERICO II”



DOTTORATO DI RICERCA IN FISICA
FONDAMENTALE ED APPLICATA

XXV CICLO

Dipartimento di Fisica

Ph. D. Thesis

The CHarged ANTI Counter for the
NA62 experiment at CERN SPS

Defended by

DOMENICO DI FILIPPO

Coordinator: Prof. Raffaele Velotta

Advisor: Prof. Fabio Ambrosino

Academic year 2012/2013

Contents

Introduction	1
1 The $K \rightarrow \pi \nu \bar{\nu}$ decay and R_K	3
1.1 The CKM framework	3
1.2 The unitary triangle	4
1.3 $K \rightarrow \pi \nu \bar{\nu}$ decays	6
1.4 $K^+ \rightarrow \pi^+ \nu \bar{\nu}$ experimental status	8
1.5 R_K in Standard Model and beyond	10
1.6 NA62 Phase I	11
1.6.1 Beam, Apparatus and Trigger Logic	11
1.6.2 Data samples and measurement strategy	13
1.6.3 Signal selection and main background	14
1.6.4 Result	16
1.6.5 Future perspective	16
2 The NA62 experiment	18
2.1 NA62 overview	18
2.2 Kinematic rejection and backgrounds	19
2.3 Beam line	21
2.4 CEDAR	23
2.5 GTK	24
2.6 CHANTI	25
2.7 LAV	25
2.8 STRAW	26
2.9 RICH	28
2.10 CHOD	29
2.11 LKr	29
2.12 MUV	30
2.13 IRC and SAC	31
2.14 Trigger and data acquisition	31
3 CHANTI design	35
3.1 Detector requirements	35
3.2 General layout	36
3.3 Scintillators	37

3.4	Optical fibers	41
3.5	SiPM	41
3.6	Readout electronics	44
3.7	Construction	45
3.8	SiPM Characterization	47
3.9	Bars characterization	47
4	Study of beam induced background	50
4.1	Beam induced background	50
4.2	Background rejection cuts	51
4.3	Background estimation	69
5	Technical run	75
5.1	Hardware setup	75
5.2	Threshold calibration	77
5.3	Stand alone acquisition	82
5.4	Time resolution	83
5.4.1	Time walk correction	84
5.4.2	Position correction	85
	Conclusions	94
A	Silicon Photo Multiplier	96
A.1	Silicon detectors	96
A.2	Geiger mode Avalanche Photo Diode	97
A.3	Photon detection efficiency	98
A.4	GAPD electric model	99
A.5	Bias voltage and gain	102
A.6	Time properties	103
A.7	Other properties	104
B	TEL62 network latency	107
B.1	Ethernet monitoring	107
B.2	Working latencies	107

Introduction

The work described in this Thesis was performed in the context of the NA62. NA62 is an international collaboration involving 29 institutes from 11 countries which is carrying on a fixed target experiment on a positive hadron beam extracted from the Super Proton Synchrotron (SPS) at CERN. The goal of the experiment is the measurement of the Branching Ratio (BR) of the ultra-rare kaon decay $K^+ \rightarrow \pi^+ \nu \bar{\nu}$ with a 10% error.

The Standard Model (SM) prediction is $BR(K^+ \rightarrow \pi^+ \nu \bar{\nu}) = (7.81_{-0.71}^{+0.80} \pm 0.29) \times 10^{-11}$ while the current best measurement performed by the E787/E949 collaboration are $BR(K^+ \rightarrow \pi^+ \nu \bar{\nu}) = 1.73_{-1.05}^{+1.15} \times 10^{-10}$. More details on the theoretical estimation and on the state of the art for the measurement is reported in chapter 1

The BR of the decay studied by NA62 is one of the best possible observables in order to test the flavor structure of SM and, given its extremely low value, is very sensitive to possible extensions of the SM. In fact most of the possible SM extensions which have been proposed (including Supersymmetry) are expected to have a non-trivial flavor structure which should be observed as a deviation from SM predictions on flavor observables. However, due to hadronic uncertainties only few, selected processes are predicted with enough accuracy so to allow sizable deviations to be observed experimentally: and among these the Flavor Changing Neutral Current decays $K_L \rightarrow \pi^0 \nu \bar{\nu}$ and $K^+ \rightarrow \pi^+ \nu \bar{\nu}$ play a crucial role, being sensitive to the V_{td} element of the Cabibbo-Kobayashi-Maskawa (CKM) flavor mixing matrix, and giving thus information which is complementary to what may be obtained in B meson physics. This motivates the strong effort which is being carried on in the world to measure these experimentally challenging modes: several experiments (NA62, KOTO, ORKA) are actually on going or have been proposed at CERN, in Japan and in the U.S.

NA62 will be performed at CERN North Area and will use a 75 GeV/c unseparated beam with a 6% kaon component. The apparatus is designed to collect the kaon decays in flight within a 65 m long evacuated fiducial volume.

In two years of data taking (starting at end of 2014) NA62 will collect about 9×10^{12} kaon decays and assuming a 10% signal acceptance, we expect about ~ 100 $K^+ \rightarrow \pi^+ \nu \bar{\nu}$ events. The experimental strategy and the detectors which form the NA62 apparatus will be described in chapter 2.

To suppress backgrounds which are up to 10^{10} times higher than the signal we will use a kinematic rejection that needs an accurate measurement of the

momentum and direction of the only two detectable particle in the process of interest: the primary K^+ and the product π^+ . To this purpose we placed a differential Cherenkov and a Silicon tracker at beginning of the decay volume, to tag and measure the kaon momentum and direction; in the same way at end of the volume there are a spectrometer and a Ring Imaging Cherenkov to identify the π^+ and to measure its momentum and direction. The kinematic rejection alone cannot reject all the dangerous backgrounds so a complex system of veto detectors was designed. A Liquid Krypton calorimeter and 12 stations (Large Angle Veto) of leadglass rings form a γ veto to reject, for example $K^+ \rightarrow \pi^+ \pi^0$; a Muon Veto system will help the identification of muons in final states. Last but not least, we need a detector able to identify the inelastic interaction between the beam and the kaon tracker.

The Naples group has in fact proposed in 2009 a detector to be used as veto for charged particles coming from inelastic interactions of the beam with the beam Si-tracker, called CHarged ANTICounter (CHANTI). The role of this detector in NA62 and its design, construction and performance evaluation are the main subjects of the present work.

Though it is evident that the beam inelastic interactions can potentially rise a relevant obstacle towards reaching the goal of the experiment, an accurate Monte Carlo evaluation of this background had never been done before this work. This was mainly due to the technical problems in simulating the huge amount of Monte Carlo statistics needed to this aim, and to the lack of a complete simulation of the detector. The first attempt to do so is described in chapter 4. We worked in the official NA62 simulation framework, but with some optimization to be able to produce and analyze more than 10^9 events through which we characterized the inelastic interactions and estimated the NA62 suppression factor for this background. This work has stressed the key role played by the CHANTI (and the Large Angle Veto system) in this context, in order to reach the Signal/Background goal of NA62.

Guided by the physics needs and the requested performances for the detector, during this Thesis work we have fixed the CHANTI design and defined all of the construction steps and semi-automatic testing procedures needed to realize the detector. Moreover we have developed¹ a dedicated front-end electronic board and defined the full readout chain of the CHANTI. The construction is steadily on going and the first of the six stations composing the CHANTI has been completely assembled in July 2012.

The description of the design choices and of the realization of the CHANTI is done in chapter 3.

In November 2012 there was a testbeam involving most of the NA62 sub-detectors. We participated to this “Technical Run” with the first station of the CHANTI, as described in chapter 5. This represented the first possibility to test the station in realistic conditions and allowed us to validate the CHANTI operation and, as well, to measure one of the most important features of the detector, namely its time resolution.

¹In collaboration with “Servizio Elettronica, Laboratori Nazionali di Frascati”

Chapter 1

The $K \rightarrow \pi \nu \bar{\nu}$ decay and R_K

The flavor physics program allows to explore the possible extensions of the Standard Model (SM) with an approach complementary to the one adopted in direct searches e.g. at the LHC. It is based on the idea that through virtual contributions physics at high energy scales can manifest itself also in low energy phenomena provided that the observables are carefully chosen, precisely measured and compared to accurate predictions. In this context it is quite natural that highly suppressed processes, where the possible New Physics (NP) corrections can alter SM predictions by a sizable relative amount play a very important role. It is in this framework, which we will describe with some more detail in the following, that the NA62 experiment has been proposed and is being pursued.

1.1 The CKM framework

The Cabibbo 2x2 matrix expresses how the weak charged current couples $u \rightarrow d$ or $u \rightarrow s$. The d and s quarks represent the mass eigenstates (physical particles) and two new state, d' and s' , are introduced as the ones effectively involved in the weak interactions. They are two orthogonal combination of d and s

$$\begin{pmatrix} d' \\ s' \end{pmatrix} = \begin{pmatrix} \cos\theta_C & \sin\theta_C \\ -\sin\theta_C & \cos\theta_C \end{pmatrix} \begin{pmatrix} d \\ s \end{pmatrix}$$

described by a single real parameter, θ_C , which is called the Cabibbo angle. It can be experimentally determined, using the information on the corresponding quark flavor transition e.g. studying the semileptonic decays of the kaons.

The 3×3 quark-mixing CKM matrix [1] generalizes the Cabibbo one by including the third generation of quark states:

$$\begin{pmatrix} d' \\ s' \\ b' \end{pmatrix} = \begin{pmatrix} V_{ud} & V_{us} & V_{ub} \\ V_{cd} & V_{cs} & V_{cb} \\ V_{td} & V_{ts} & V_{tb} \end{pmatrix} \begin{pmatrix} d \\ s \\ b \end{pmatrix}$$

The matrix can be expressed in the Wolfenstein parametrization [2]:

$$v = \begin{pmatrix} 1 - \lambda^2/2 & \lambda & A\lambda^3(\rho - i\eta) \\ -\lambda & 1 - \lambda^2/2 & A\lambda^2 \\ A\lambda^3(1 - \rho - i\eta) & -A\lambda^2 & 1 \end{pmatrix} + O(\lambda^4) \quad (1.1)$$

where λ is used as an expansion parameter; A and λ are defined to be positive and satisfy the relations: $\lambda = \sin \theta_{12}$, $A\lambda^2 = \sin \theta_{23}$ and $A\lambda^3(\rho - i\eta) = \sin \theta_{13} \exp(-i\phi)$, with θ_{ij} denoting three real parameters (Cabibbo-like angles) and $\exp(-i\phi)$ a phase factor. Thanks to the presence of the non-multiplicative complex phase the CKM matrix is able to explain CP violation observed in kaon and B meson decays.

The current status of the experimental situation is summarized in the following [3]:

$$v = \begin{pmatrix} 0.97425 \pm 0.00022 & (2.252 \pm 0.009) \times 10^{-1} & (4.15 \pm 0.49) \times 10^{-3} \\ (2.30 \pm 0.11) \times 10^{-1} & 1.006 \pm 0.023 & (40.9 \pm 1.1) \times 10^{-3} \\ (8.4 \pm 0.6) \times 10^{-3} & (42.9 \pm 2.6) \times 10^{-3} & 0.89 \pm 0.07 \end{pmatrix} \quad (1.2)$$

The values reported above are obtained by averaging various measurements (but without imposing unitarity in the fits). It is worth noticing that the diagonal elements are the dominant ones, reflecting the transitions $u \rightarrow d$, $c \rightarrow s$ and $t \rightarrow b$ which are the most allowed.

The off-diagonal elements represent transitions suppressed at a certain level, depending on their amplitude. The value reported for $|V_{us}|$ comes from the measurement $BR(K^+ \rightarrow \mu^+\nu(\gamma))$ [4] and a combined result of $K_L \rightarrow \pi e\nu$, $K_L \rightarrow \pi\mu\nu$, $K^\pm \rightarrow \pi^0 e^\pm\nu$, $K^\pm \rightarrow \pi^0 \mu^\pm\nu$ and $K_S \rightarrow \pi e\nu$ decays [5].

The determination of $|V_{td}|$ and $|V_{ts}|$ is based on the measurements of the mass difference of two neutral B_0 meson mass eigenstates performed by the CDF [6] and LHCb [7] experiments while the ratio $|V_{ts}/V_{cb}|$ can be extracted from the ratio $BR(B \rightarrow X_s\gamma)/BR(B \rightarrow X_c e\bar{\nu})$ [8] or the $B_s \rightarrow \mu^+\mu^-$ decay rate [9].

A theoretically clean and independent measurement of $|V_{td}V_{ts}^*|$ is possible from the $K^+ \rightarrow \pi^+\nu\bar{\nu}$ decay [10] as we will see in the following.

1.2 The unitary triangle

Imposing the unitarity condition [11] on the CKM matrix leads to the nine conditions $\sum_i V_{ij}V_{ik}^* = \delta_{jk}$ and $\sum_j V_{ij}V_{kj}^* = \delta_{ik}$, where $i = u, c, t$ and $j, k = d, s, b$, and δ_{ij} is the Kronecker delta.

The three equations on the diagonal elements constrain the CKM elements magnitude and express the universality of the weak interaction, while the six vanishing combinations can be thought as triangles in a complex plane. The unitarity property of CKM matrix can be used to test the SM flavor sector: a way to do this is to measure the CKM matrix elements and monitor any

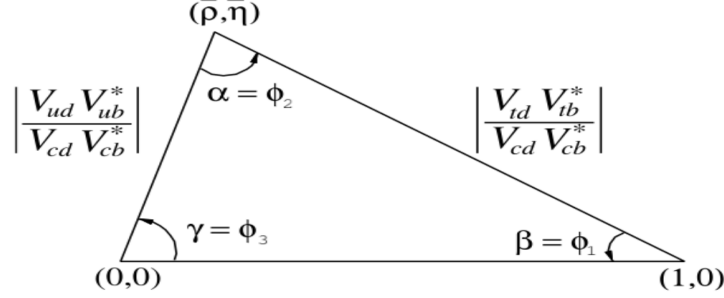


Figure 1.1: Unitarity triangle corresponding to the condition 1.3.

significant deviation from unitarity. There is only one non-degenerate triangle corresponding to the combination

$$\begin{aligned} 0 &= V_{ud}V_{ub}^* + V_{cd}V_{cb}^* + V_{td}V_{tb}^* \\ &= 1 + \frac{V_{td}V_{tb}^*}{V_{cd}V_{cb}^*} + \frac{V_{ud}V_{ub}^*}{V_{cd}V_{cb}^*} \end{aligned}$$

where in the last line we divided the expression by the best known term $V_{cd}V_{cb}^*$. In the last form, the expression can be seen as the sum of three vectors in a complex plane, represented with the triangle in figure 1.1; the angles of the unitarity triangle in figure 1.1 are:

$$\begin{aligned} \alpha = \phi_1 &= \arg \left(-\frac{V_{td}V_{tb}^*}{V_{ud}V_{ub}^*} \right) \\ \beta = \phi_1 &= \arg \left(-\frac{V_{cd}V_{cb}^*}{V_{td}V_{tb}^*} \right) \\ \gamma = \phi_1 &= \arg \left(-\frac{V_{ud}V_{ub}^*}{V_{cd}V_{cb}^*} \right) \end{aligned}$$

In the Wolfenstein parametrization 1.1 the previous CKM unitarity condition is written as:

$$1 + \frac{V_{td}V_{tb}^*}{V_{cd}V_{cb}^*} = -\frac{V_{ud}V_{ub}^*}{V_{cd}V_{cb}^*} = \rho \left(1 - \frac{\lambda^2}{2} \right) + i\eta \left(1 - \frac{\lambda^2}{2} \right) \quad (1.3)$$

As a consequence, the vertex coordinates of the unitarity triangle in figure 1.1 are exactly $(0,0)$, $(1,0)$, and $(1 - \lambda^2/2)(\rho, \eta)$.

1.3 $K \rightarrow \pi\nu\bar{\nu}$ decays

The transitions $K^+ \rightarrow \pi^+\nu\bar{\nu}$ and $K_L \rightarrow \pi^0\nu\bar{\nu}$ are very interesting in SM because the measurement of their decay rates provides important information about some of the less well-known fundamental physics parameters of the model. In fact, for these transitions the branching ratios are theoretically predicted in the SM and the purely theoretical (i.e. not related to experimentally measured quantities) relative uncertainties are well known, $\approx 4\%$ both for $BR(K^+ \rightarrow \pi^+\nu\bar{\nu})$ and for $BR(K_L \rightarrow \pi^0\nu\bar{\nu})$ [12]. The calculations show the sensitivity of these decay rates to the magnitude of the V_{td} element of CKM matrix, which can be determined with few percent accuracy without relying on unitarity constraints. Moreover since, as well known, Flavor Changing Neutral Currents (FCNC) processes are strongly suppressed in the SM, they can be used to test the occurrence of new physics (NP). Finally, simultaneous BR measurements of $K^+ \rightarrow \pi^+\nu\bar{\nu}$ and $K_L \rightarrow \pi^0\nu\bar{\nu}$ decays provide determinations of CKM parameters and the unitarity triangle in a complementary and independent way with respect to the study of B decays.

In the SM the quark level process that contribute to the $K \rightarrow \pi\nu\bar{\nu}$ decay is the flavour changing quark transition $s \rightarrow d\nu\bar{\nu}$ described, at first non-null order, by the one-loop diagram shown in figure 1.2: penguin diagrams with Z exchange and box diagrams with W exchange. In fact flavour-changing transitions such as $s \rightarrow d$ are forbidden at tree-level. Separating the contribution of the u , c , and t quarks that appear as internal lines, one has:

$$A(s \rightarrow d\nu\bar{\nu}) = \sum_q V_{qs}^* V_{qd} A_q$$

with leading order quantum loops contributing to the amplitude with terms which are positive power of $x_q \equiv m_q^2/M_W^2$. Because of its mass, the top-quark contribution becomes the dominant term and the transition $s \rightarrow d$ is described by short-distance quark dynamics. As a consequence, the QCD corrections are small and calculable in perturbation theory [13].

In this scenario we can describe the $s \rightarrow d\nu\bar{\nu}$ process by means of a Fermi-like coupling between a quark and a lepton neutral weak currents [14]:

$$H_{eff} = \frac{G_F \alpha}{\sqrt{2} 2\pi \sin^2 \theta_W} \sum_{l=e,\mu,\tau} [V_{cs}^* V_{cd} X^l + V_{ts}^* V_{td} Y(x_t)] \cdot \bar{s} \gamma^\mu (1 - \gamma^5) d \bar{\nu}_l \gamma^\mu (1 - \gamma^5) \nu_l$$

where G_F is the Fermi coupling constant, α is the fine-structure constant, θ_W is the Weinberg angle (weak mixing angle), s, d, ν are the Dirac spinors for the respective particles and γ^μ, γ^5 the Dirac matrices. The two function Y and X^l are:

- $Y(x_t)$ is a function encoding dominant (exclusive, in case of the neutral decay) top-quark loop contribution to the $K \rightarrow \pi\nu\bar{\nu}$ decays modes . It

is known in the QCD at next-to-leading order (NLO) [15], while for the electro-weak interactions the two loop contributions are computed. The theoretical value and uncertainty is

$$Y(x_t) = 1.469 \pm 0.017(QCD) \pm 0.002(WEAK)$$

- X^l with $l = e, \mu, \tau$ are known at NLO [16] and they are functions describing the charm-quark contribution in the loop, conveniently described by the relation [14, 17]:

$$P_0(X) = \frac{1}{\lambda^4} \left[\frac{2}{3}X^e + \frac{1}{3}X^\tau \right] = 0.42 \pm 0.06$$

with $X^\mu = X^e$ and using the value $\lambda = \sin\theta_{12} = 0.2240 \pm 0.0036$ for the Wolfenstein parameter. Its contribution is about 30% of $A(s \rightarrow d\nu\bar{\nu})$ and the NLO uncertainty translates into an error of about 10% in the SM estimate of $K^+ \rightarrow \pi^+\nu\bar{\nu}$. The Next-to-Next-to-Leading Order (NNLO) calculation reduces the pure theoretical uncertainty below 4% [10, 17].

Therefore the $BR(K^+ \rightarrow \pi^+\nu\bar{\nu})$ can be written as [16, 18]:

$$BR(K^+ \rightarrow \pi^+\nu\bar{\nu}) = \frac{\bar{k}_+}{\lambda^2} (Im\lambda_t)^2 Y^2(x_t) + \frac{\bar{k}_+}{\lambda^2} (\lambda^4 ReV_{cs}^* V_{cd} P_0(X) + ReV_{ts}^* V_{td} Y(x_t))^2$$

where the coefficient k_+ depend on $V_{us} = \lambda$, the kaon mass, charge and lifetime as well as the kaon and pion decay constant and phase space factor. Since the same factors appear in the BR of the Ke3 process ($K^+ \rightarrow \pi^0 e^+ \nu$) we have [19]:

$$\bar{k}_+ = r_{K^+} BR(K^+ \rightarrow \pi^0 e^+ \nu) = r_{K^+} \times (5.08 \pm 0.05) \times 10^{-2}$$

where $r_{K^+} = 0.901$ is a parameter necessary to relate the final states $\pi^+\nu\bar{\nu}$ and $\pi^0 e^+ \nu$ that includes the isospin breaking corrections. Since $|V_{cd}|$ is well measured from charmed particles decay [10, 17], and $|V_{cs}|$ is determined from leptonic and semileptonic D decays, the $BR(K^+ \rightarrow \pi^+\nu\bar{\nu})$ is directly linked to $|V_{td}V_{ts}^*|$ without using unitarity conditions. Using the theoretical value for the BR we get an error of $\approx 5\%$ – 7% on the determination of $|V_{td}|$ [12].

The theoretical prediction for $K_L \rightarrow \pi^0\nu\bar{\nu}$ is even more precise. Since the process is CP-violating, and for the structure of the CP symmetry, only the imaginary part of the effective hamiltonian contributes to the amplitude. The charm-quark contribution in the loop is negligible and the principal contribution comes from the top-quark, for which the QCD corrections are suppressed and rapidly convergent [13]. The result is [16, 18]:

$$BR(K_L \rightarrow \pi^0\nu\bar{\nu}) = \frac{\bar{k}_L}{\lambda^2} (Im\lambda_t)^2 Y^2(x_t)$$

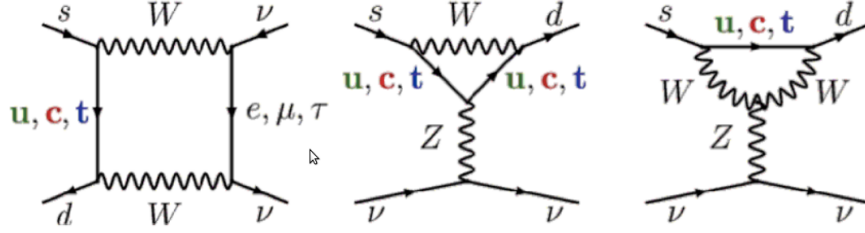


Figure 1.2: One-loop diagram (Leading Order) contributing to $s \rightarrow d\nu\bar{\nu}$ process. The one on the left is a box diagram with W exchange, in the middle there is the penguin diagram with Z exchange and the last is the boson vertex diagram.

where k_L have the same nature of k_+ and can be analogously written as [20]:

$$\bar{k}_L = r_{K_L} BR(K^+ \rightarrow \pi^0 e^+ \nu)$$

with $r_{K_L} = 0.944$.

Therefore in the SM, the theoretical expectations for the BR of $K \rightarrow \pi\nu\bar{\nu}$ processes are [12]:

$$\begin{aligned} BR(K^+ \rightarrow \pi^+ \nu\bar{\nu}) &= (7.81^{+0.80}_{-0.71} \pm 0.29) \times 10^{-11} \\ BR(K_L \rightarrow \pi^0 \nu\bar{\nu}) &= (2.43^{+0.40}_{-0.37} \pm 0.06) \times 10^{-11} \end{aligned}$$

where the first error is related to the uncertainties in the input parameters and is dominated by the CKM parameter V_{cb} while the second error arise from intrinsic theoretical uncertainties.

As it is well known the measurements of CP asymmetries in $B_d \rightarrow J/\psi K_S$ and $B_d \rightarrow \pi\pi$ allow to extract $\sin 2\beta$ and $\sin 2\alpha$, thus constraining ρ and η .

The structure of the CKM makes possible an important consistency check of SM, by comparing the determination of the unitarity triangle obtained in the B sector with the one obtained from kaon rare decays, as it is shown in fig 1.3

1.4 $K^+ \rightarrow \pi^+ \nu\bar{\nu}$ experimental status

In figure 1.5 is shown an historical view of the $BR(K^+ \rightarrow \pi^+ \nu\bar{\nu})$ measurements. The first upper limit was found in year 1969 by heavy-liquid bubble chamber experiment at the Argonne Zero Gradient Synchrotron [21]: $BR(K^+ \rightarrow \pi^+ \nu\bar{\nu}) < 10^{-4}$ at the 90% C.L. . In 1973 a spark chamber experiment at the Berkeley Bevatron improved the limit down to $BR(K^+ \rightarrow \pi^+ \nu\bar{\nu}) < 5.6 \times 10^{-7}$ [22]. About ten years later, with an experiment at KEK Proton Synchrotron [23], the limit was improved to $BR(K^+ \rightarrow \pi^+ \nu\bar{\nu}) < 1.4 \times 10^{-7}$.

The first event candidates were observed by the E787/E949 collaboration at Brookhaven National Laboratory (BNL). E787 started in the '80s to study

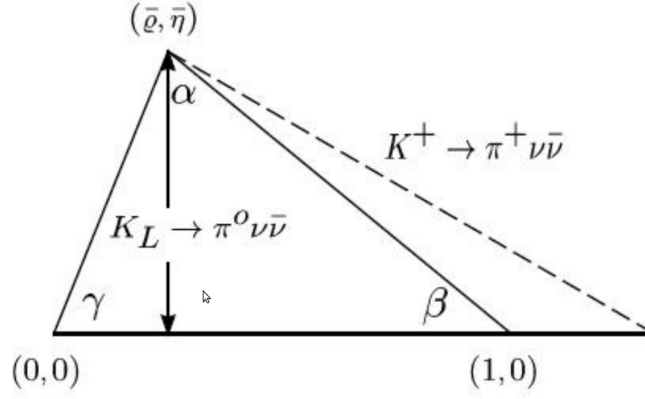


Figure 1.3: Constraints on unitarity triangle from Kaon rare decays. The $K^+ \rightarrow \pi^+ \nu \bar{\nu}$ decay rate defines the dashed side and the displacement of the right down vertex is due to the charm-quark contribution; the $K_L \rightarrow \pi^0 \nu \bar{\nu}$ decay rate gives the triangle height.

the $K^+ \rightarrow \pi^+ \nu \bar{\nu}$ decay from kaons at rest: a proton beam of ~ 20 GeV/c was extracted from the Alternating Gradient Synchrotron (AGS), and by its interaction with a target a secondary beam containing K^+ was generated. Then ~ 700 MeV/c kaons were selected by means of electromagnetic separators, Cherenkov counters, tracking devices, and energy loss counters. Ultimately the beam was directed on a scintillating fiber target which stops $\sim 27\%$ of kaons. The signal region was defined in terms of the range (R), momentum (M) and kinetic energy (E) of charged decay products. The target was designed as a large acceptance solenoid spectrometer and an hermetic photon veto.

The $K^+ \rightarrow \mu^+ \nu(\gamma)$ background was rejected by time-constraining the decay chain $\pi^+ \rightarrow \mu^+ \rightarrow e^+$, while the $K^+ \rightarrow \pi^+ \pi^0$ was suppressed by detecting the two photons from the $\pi^0 \rightarrow \gamma\gamma$. The usage of the simulation was limited to the estimation of the geometrical acceptance of the signal.

E787 with a series of runs between 1995 and 1998 extracted the first signal candidate from an initial sample of 10^{12} stopped K^+ , as shown in 1.4: the analysis was divided in two pion momentum regions, so the result is $BR(K^+ \rightarrow \pi^+ \nu \bar{\nu}) = 1.57^{+1.75}_{-0.82} \times 10^{-10}$ [24] based on the red-circle markers in the region 211 MeV/c - 229 MeV/c, and $BR(K^+ \rightarrow \pi^+ \nu \bar{\nu}) < 42 \times 10^{-10}$ at 90% C.L [25] in the region 140 MeV/c - 195 MeV/c.

The E949 experiment [26] was an upgrade of E787 with a sensitivity enhanced by a factor of 5, due to improvements on the efficiencies of photon veto detection, tracking and trigger and data acquisition. The results are reported in the same figure 1.4. Combining the result of E787 and E949, one obtains the current estimation of the $BR(K^+ \rightarrow \pi^+ \nu \bar{\nu})$ based on 7 observed events [26]:

$$BR(K^+ \rightarrow \pi^+ \nu \bar{\nu}) = 1.73^{+1.15}_{-1.05} \times 10^{-10}$$

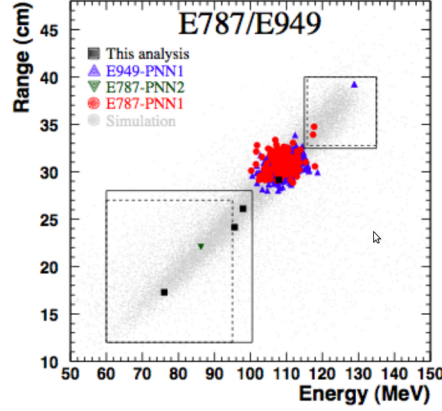


Figure 1.4: Distribution of pion range versus pion kinetic energy for events passing the analysis selection in E949 [27]. The solid (dashed) squares are analysis region of E949 (E787) and they contain the observed $\pi\nu\bar{\nu}$ events. The markers outside the boxes correspond to $K^+ \rightarrow \pi^+\pi^0$ events passing the analysis selection while the light gray points are the simulated $K^+ \rightarrow \pi^+\nu\bar{\nu}$ events.

This result have a large statistical uncertainty and is compatible with the SM expectation.

1.5 R_K in Standard Model and beyond

The $R_K = BR(K_{e2})/BR(K_{\mu 2})$ (were K_{l2} means $K \rightarrow l\nu_l(\gamma)_{IB}$) ratio in SM framework is a very well determined quantity [28]:

$$R_K^{SM} = \frac{m_e^2}{m_\mu^2} \cdot \frac{m_k^2 - m_e^2}{m_k^2 - m_\mu^2} \cdot (1 + \delta R_K^{Rad.Corr.}) = (2.477 \pm 0.001) \times 10^{-5} \quad (1.4)$$

where $\delta R_K^{Rad.Corr.} = (3.79 \pm 0.04)\%$ is an electromagnetic correction due to the IB and structure dependent effects. Any significant deviation from this value could signal new physics. In Minimal Supersymmetric Standard Model (MSSM) scenario R_K value is modified due to Lepton Flavor Violating (LFV) terms in charged Higgs exchange diagrams (Figure 1.6). Using reasonable SUSY parameters values (the mixing parameter between the superpartners of the right-handed leptons, $\Delta_{13} = 5 \times 10^{-4}$, the ratio of the two Higgs vacuum expectation values, $\tan(\beta)$, and the Higgs mass, $m_H = 500$ GeV) sizable deviations from SM value have been predicted [29]:

$$R_K^{LFV} = 2 \frac{\Gamma_{SM}(K \rightarrow e\nu_e) + \Gamma_{LFV}(K \rightarrow e\nu_\tau)}{\Gamma(K \rightarrow \mu\nu_\mu)} = R_K^{SM}(1 + 0.013) \quad (1.5)$$

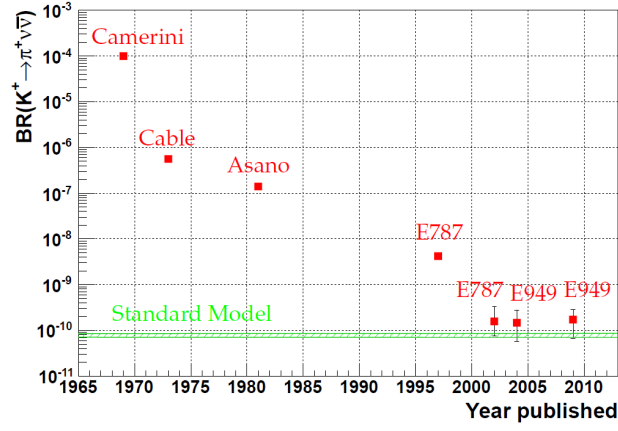


Figure 1.5: Summary of the $\text{BR}(K^+ \rightarrow \pi^+ \nu \bar{\nu})$ experimental status over the years.

R_K PDG 2012 value is computed using five measurements, three dating back to the 70s:

$$R_K^{70s} = (2.45 \pm 0.11) \times 10^{-5} \quad (1.6)$$

In 2008 a result from KLOE [30] experiment improved the measurement:

$$R_K^{KLOE} = (2.493 \pm 0.031) \times 10^{-5}. \quad (1.7)$$

In 2011 a result from NA62 experiment, using a partial data set, gave:

$$R_K^{NA62} = (2.487 \pm 0.013) \times 10^{-5}. \quad (1.8)$$

The world average was:

$$R_K^{NA62} = (2.487 \pm 0.012) \times 10^{-5}. \quad (1.9)$$

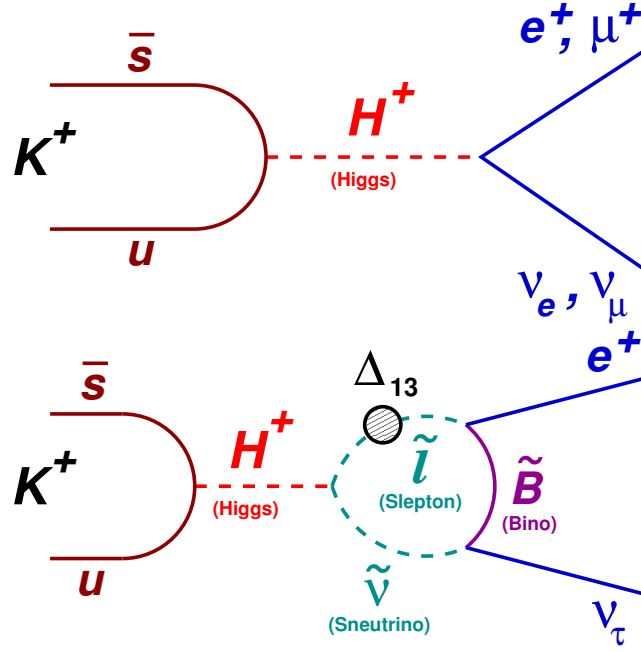
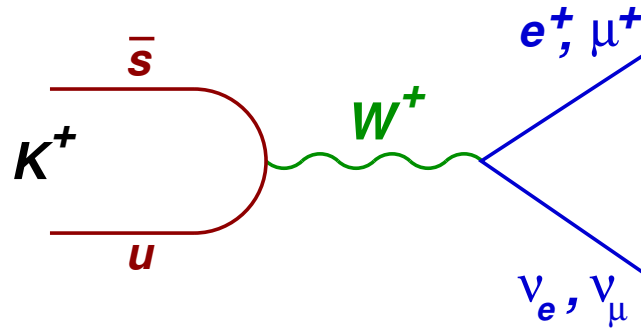
During year 2013, a new and more accurate determination of R_K has been published by NA62 “phase I”. This result is described in the following section.

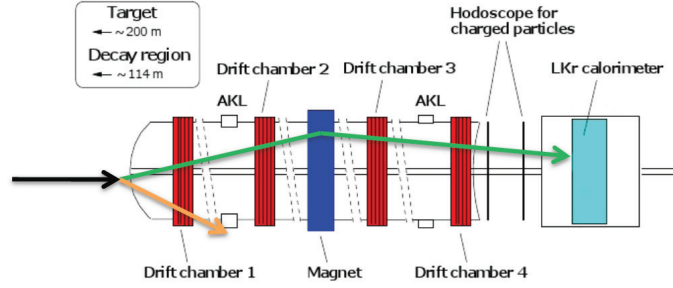
1.6 NA62 Phase I

NA62 phase I took place in 2007 when we collected data in order to measure the ratio R_K (see section 1.5) at few per mill level. A brief description of the experimental layout, the analysis strategy and of the results obtained is reported in the following [31].

1.6.1 Beam, Apparatus and Trigger Logic

Data have been taken in the June-October 2007 period. The beam was given by either simultaneous or single positive and negative secondary hadrons, with

Figure 1.6: SUSY contributions to $BR(K \rightarrow l \nu_l)$.Figure 1.7: SM contributions to $BR(K \rightarrow l \nu_l)$.

Figure 1.8: Scheme of apparatus for R_K measurement.

central momentum of 74 GeV/c and momentum spread of ± 1.4 GeV/c (rms). The apparatus used is reported in Figure 1.8.

It is composed by a charged Hodoscope (called HODO), used as fast trigger, a Drift CHamber (DCH) spectrometer and the NA48 Liquid Krypton (LKr) calorimeter. The DCH was composed by four drift chambers, each consisting of 8 planes of sense wires, and a dipole magnet located between the second and the third DCH which gave a horizontal transverse momentum kick of 265 MeV/c to charged particles. The measured spectrometer momentum resolution was $\sigma_p/p = 0.48\% \oplus 0.009\% \times p$, where the momentum p is expressed in GeV/c. The LKr is a 127 cm ($27X_0$) thick liquid krypton electromagnetic calorimeter, used for lepton identification and as a photon veto detector. It is located further downstream. The energy resolution was $\sigma_E/E = 3.2\%/\sqrt{E} \oplus 9\%/E \oplus 0.42\%$ (E in GeV).

We used a minimum bias hardware trigger in order to select simultaneously K_{e2} and $K_{\mu2}$ events to minimize the systematics. The two samples only differ for energy release in LKr. Common logical conditions used are: activities in DCHs and energy release into both the hodoscope planes. K_{e2} events have to satisfy a further condition: energy released in LKr higher than 10 GeV. The $K_{\mu2}$ trigger is downscaled by a factor $D=150$. In order to achieve an accuracy better than 0.5% about 150K events of K_{e2} have been collected.

1.6.2 Data samples and measurement strategy

The data-taking strategy was optimized to measure the two main backgrounds in the K_{e2} sample, which are due to the beam halo muons and to $K_{\mu2}$ decays with a muon misidentified as an electron. To measure the muon halo background directly from data, the K^+ and K^- data samples were collected alternately by blocking the negative or the positive beam, respectively. Therefore, 65% (8%) of the total 2007 beam flux corresponded to K^+ (K^-) decays collected in single-beam mode. In addition to being the signal samples (i.e. providing the K_{l2}^+ data), these data sets are used as control samples to measure the muon halo background to the decays of opposite sign kaons. To estimate the $K_{\mu2}$

background, the probability to misidentify a muon as an electron due to large energy deposition in the LKr calorimeter has been measured. This required the collection of a muon sample free from the typical $\sim 10^{-4}$ electron contamination due to $\mu^\pm \rightarrow e^\pm \nu_e \nu_\mu$ decays in flight. To this end, 55% of the kaon flux in 2007 was collected with a transverse horizontal lead (Pb) bar installed below the beam pipe between the two HOD planes, approximately 1.2 m in front of the LKr calorimeter. The bar was $9.2 X_0$ thick in the beam direction (including an iron holder) and shadowed 11 rows of LKr cells (about 10% of the total number of cells). For a 50 GeV electron traversing the Pb bar, the probability of depositing over 95% of its initial energy in the LKr is $\sim 5 \times 10^{-5}$, as estimated with a simulation.

The experimental quantity to be measured is:

$$R = \frac{1}{D} \cdot \frac{N_{Ke2} - N_{Ke2}(BG)}{N_{K\mu2} - N_{K\mu2}(BG)} \cdot \frac{A_{K\mu2} \times \epsilon_{K\mu2} \times PID_{K\mu2}}{A_{Ke2} \times \epsilon_{Ke2} \times PID_{Ke2}} \quad (1.10)$$

where N_{Kl2} ($l=e,\mu$) is the number of selected events, $N_{Kl2}(BG)$ is the number of background (BG) events, A_{Kl2} the geometrical acceptance, ϵ_{Kl2} and PID_{Kl2} the trigger and selection efficiencies respectively. The ratio R has been evaluated in 10 momentum bins, ranging from 13 to 65 GeV.

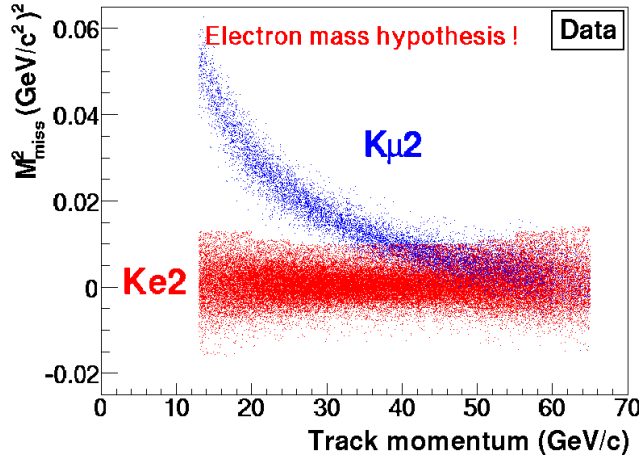


Figure 1.9: M_{miss}^2 vs track momentum in electron mass hypothesis.

1.6.3 Signal selection and main background

In order to separate the two samples we exploited the kinematic separation (using $M_{miss}^2 = (p_K - p_l)^2$ see Figure 1.9), which is optimal for tracks with energy up to 25 GeV, and particle identification using E/p ratio (energy released in LKr over measured track momentum, see Figure 1.10). The selection criteria are:

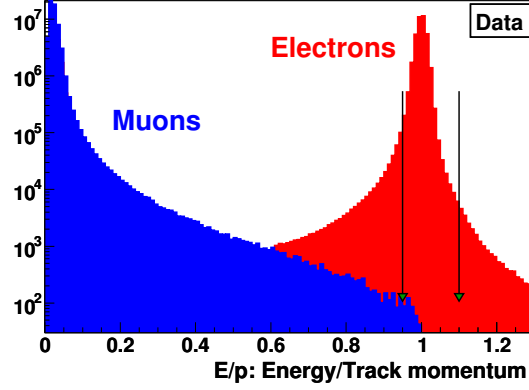


Figure 1.10: E/P distribution: in red the Ke_2 signal, in blue the $K\mu_2$ background. The arrows define the Ke_2 signal cuts.

- electron: $(E/P)_{min} \leq E/P \leq 1.1$;
- where $(E/P)_{min} = 0.95$ for $p \geq 25$ GeV/c and $(E/P)_{min} = 0.90$ otherwise
- muon: $E/P \leq 0.85$.

The number of Ke_2 candidates in the signal region is $N(Ke_2) = 145958$. The main background source for Ke_2 sample are $K\mu_2$ events in which the muon loose all of its energy into LKr (catastrophic bremsstrahlung) thus emulating an electron (therefore Ke_2 event). These events were expected to contribute at O(10%) level to the final sample and we thus decided to directly measure their fraction in order to validate Monte Carlo estimates. This measurement was done using a pure (electron contamination was evaluated to be $\sim 5 \times 10^{-5}$) muon sample obtained interposing a $10X_0$ thick Pb bar between the two hodoscope planes. A MC simulation was made with and without the Pb bar. The first, corrected for the ionization energy loss (dominant at low momentum) and bremsstrahlung (dominant at high momentum) in the Pb bar, was compared with data finding a very good agreement. The second, corrected for the effect above indicated, was used to evaluate the real background contamination which is of the order of 6%.

The number of $K\mu_2$ candidates collected with a trigger chain involving down-scaling by a factor of 150 is $N(K\mu_2) = 4282 \times 10^7$. The main background source for $K\mu_2$ sample is due to the beam halo muons. This effect has been measured directly by reconstructing the $K_{\mu_2}^+$ from a K^- data sample collected with the K^+ beam (but not its halo) blocked, and a special data sample collected with both beams blocked. The real background contamination was: $(0.50 \pm 0.01)\%$

1.6.4 Result

The result of R_K measurement is:

$$R_K = (2.488 \pm 0.007_{stat} \pm 0.007_{syst}) \times 10^{-5} = (2.488 \pm 0.010) \times 10^{-5} \quad (1.11)$$

the precision reached is 0.5%, see Figure 1.13.

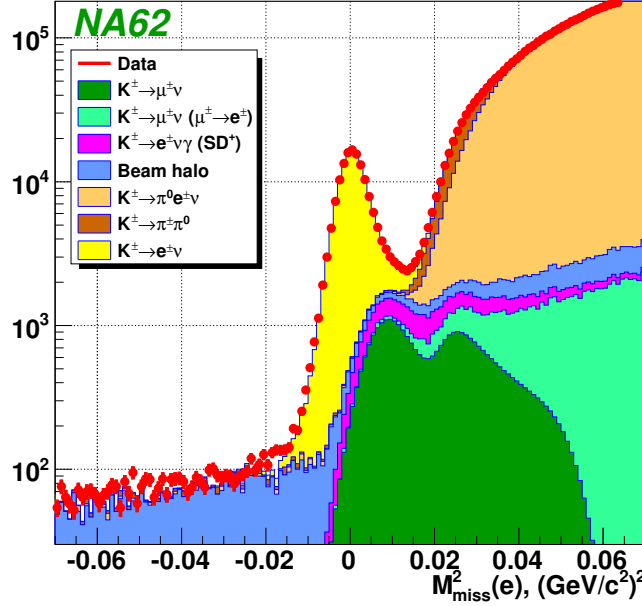
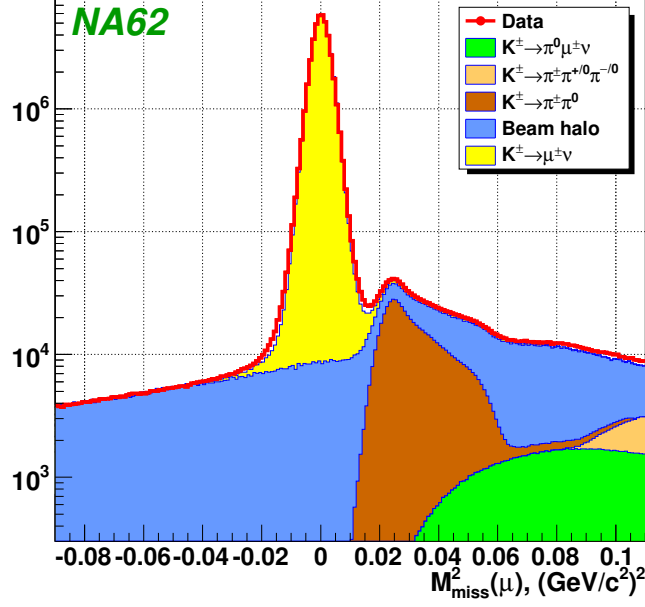
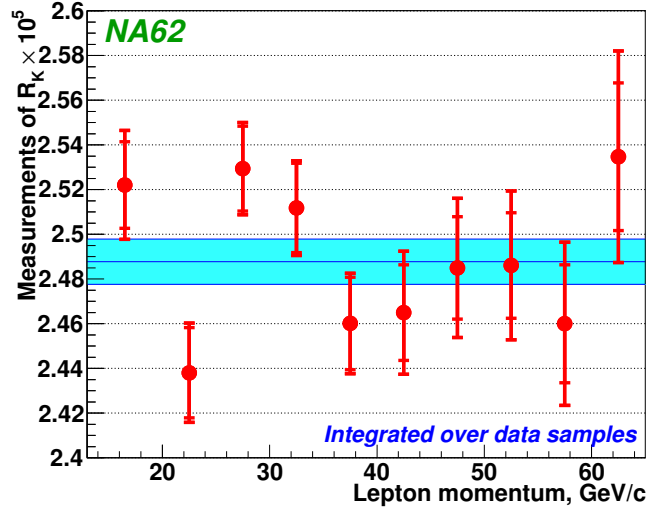


Figure 1.11: M_{miss}^2 distribution in electron mass hypothesis.

The main source of systematic uncertainty is due to the evaluation of the $K_{\mu 2}$ background in the $K_{e 2}$ sample $\delta R_K \times 10^5 = 0.004$.

1.6.5 Future perspective

In the framework of NA62, phase II, the uncertainties on the measurement of R_K can be reduced, both the statistical one and the systematic one. During the first year of data taking more than 1200k $K_{e 2}$ candidates will be collected while the use of RICH for the electron-muon discrimination will reduce contamination to negligible level. The expected total uncertainty is below 0.2%.

Figure 1.12: M^2_{miss} distribution in muon mass hypothesis.Figure 1.13: R_k evaluation for different track momentum bins.

Chapter 2

The NA62 experiment

2.1 NA62 overview

NA62 represents the current kaon physics program at CERN and offers a complementary approach, with respect to the Large Hadron Collider high energy frontier, to probe new physics at short distances, corresponding to energy scales up to ~ 100 TeV.

The purpose of the experiment is to measure the Branching Ratio (BR) of the K decay $K^+ \rightarrow \pi^+ \nu \bar{\nu}$. The Standard Model prediction of this BR is very accurate ($BR = (7.81_{-0.71}^{+0.80} \pm 0.29) \times 10^{-11}$, [12]) and it is very interesting to be measured because this process is quite sensitive to the flavour structure of possible new physics beyond the Standard Model [32, 33]. The current measurement of this BR is $1.73_{-1.05}^{+1.15} \times 10^{-10}$, performed by E787/E949 collaboration at Brookhaven National Laboratory [27].

NA62 aims to collect ~ 100 events in 2 years (starting at end of 2014) and to get a 10% background/signal ratio [35, 34] measuring the kaon decays in flight. The strategy is to generate an high intensity K beam and to detect both the decaying kaon and the final π^+ coming from a big fiducial volume (≈ 60 m), and to exclude all the other channels. In figure 2.1 there is a scheme of the apparatus.

The K beam arises from a 400 GeV proton beam hitting on a beryllium target (CERN T10 line). The beam optics selects particles with an average momentum of 75 GeV/c producing a 750 MHz unseparated positive hadron beam with a 6% kaon component. This system will be able to produce $\approx 10^{14}$ K decays in two years but only $\approx 10\%$ of them will decay in the fiducial region and only $\approx 10\%$ of the ones in the fiducial region will pass the analysis cuts.

The K is identified by a differential Cherenkov detector (CEDAR) while the final π^+ is identified in a Ring Imaging Cherenkov (RICH) and its momentum is measured in a Straw Chamber Spectrometer (STRAW).

In order to apply kinematic rejection [35] we need to measure the K momentum, to this aim we use a silicon tracker based spectrometer (the so-called

GigaTracKer - GTK) to precisely measure the beam momentum and timing. However, the interaction of the beam with the GTK is itself a possible source of background. In particular, inelastic scattering events can mimic the signal if a produced pion falls into the RICH acceptance, if it is badly reconstructed inside the fiducial volume, and if no other tracks are detected.

The charged anti counter (CHANTI) is a detector placed in vacuum just after the GTK (27 mm) to help the rejection of this background, by covering hermetically the region between 49 mrad and 1.2 rad wrt the third and last GTK station. This detector is being designed, constructed and tested in Naples.

To exclude the other decays, there are an electromagnetic calorimeter (LKr) and several veto counters: large angle photon vetoes (LAV), low angles vetoes (IRC and SAC), muon veto detectors (MUV1-2-3).

2.2 Kinematic rejection and backgrounds

In figure 2.2 the kinematics of the decay $K^+ \rightarrow \pi^+ \nu \bar{\nu}$ is sketched. Only the momentum of the incoming kaon P_K and of the daughter pion P_π and the angle between them $\theta_{K\pi}$ can be measured. Calling m_K and m_π the mass of the two particles, and p_K^α and p_π^α their four-momentum, we can define the kinematic variable m_{miss}^2 as

$$\begin{aligned} m_{miss}^2 &= (p_K - p_\pi)^\alpha (p_K - p_\pi)_\alpha = \\ &= m_K^2 + m_\pi^2 + 2P_K P_\pi \left[\cos\theta_{K\pi} - \sqrt{\left(1 + \frac{m_K^2}{P_K^2}\right) \left(1 + \frac{m_\pi^2}{P_\pi^2}\right)} \right] \end{aligned}$$

As said before, P_K will be measured with a beam spectrometer, the Gigatracker (GTK, see section 2.5), placed upstream the kaon decay region, while P_π will be measured by another spectrometer (STRAW, see section 2.8) at end of the decay region. The two spectrometers are able to measure also the direction of the the two particles so with the combined information we can get also a measurement of the angle $\theta_{K\pi}$. With the whole system we can thus determine the m_{miss}^2 of the decay.

The m_{miss}^2 distributions for the principal backgrounds (92% of the total, the $K^+ \rightarrow \mu^+ \nu$ alone is $\sim 63\%$) are shown in figure 2.3 together with the distribution for the signal $K^+ \rightarrow \pi^+ \nu \bar{\nu}$. By selecting the events in Region I ($0 < m_{miss}^2 < 10^4$) and Region II ($2.6 \times 10^4 < m_{miss}^2 < 6.8 \times 10^4$), these background are completely eliminated (neglecting resolution effects) since for them the m_{miss}^2 are limited. Instead the remaining background (8% of the total) shown in figure 2.4 (assuming a 75 GeV/c kaon) are not strongly rejected by this selection. In order to guarantee the background rejection at the required level of 10^{-12} , we need anyway to rely also on veto system and accurate particle identification. [36].

In fact in the discussion above we completely neglected the effects of the resolution on m_{miss} , but obviously we have to take it into account since, for

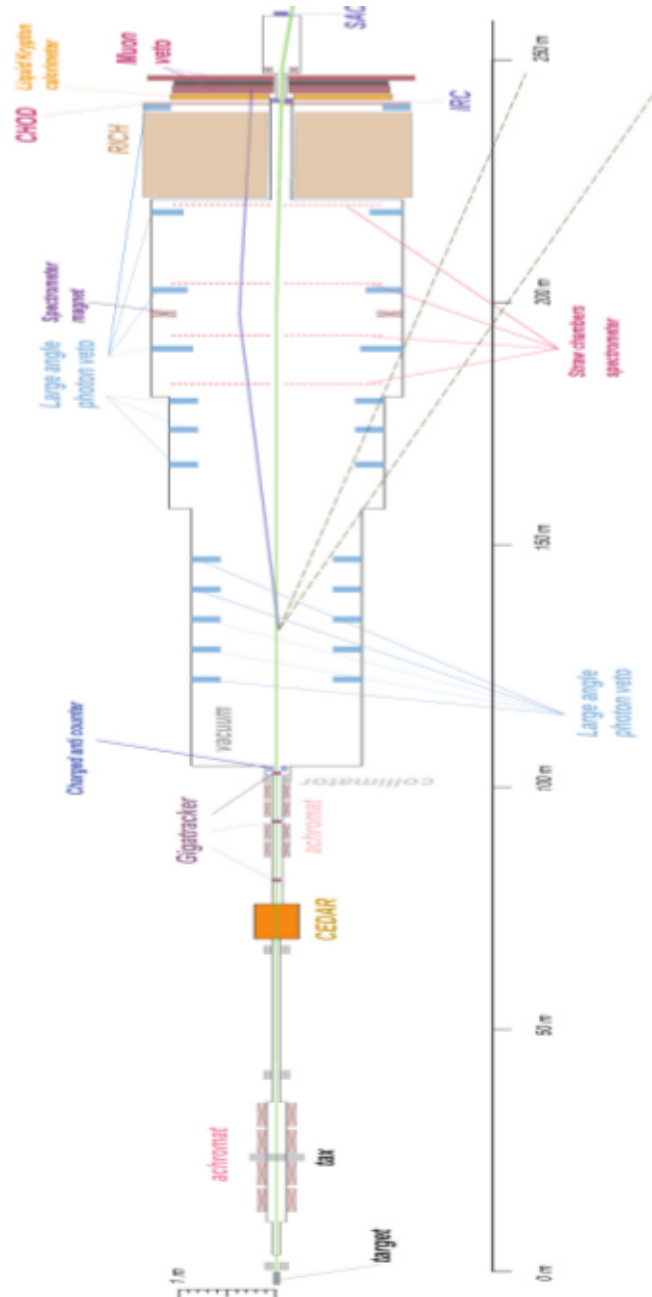
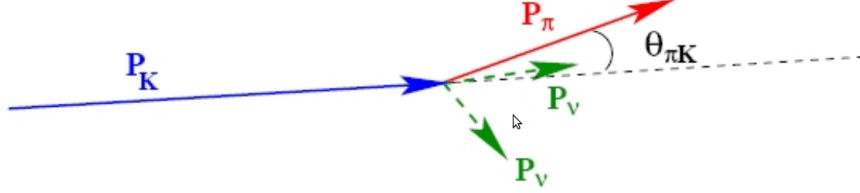


Figure 2.1: NA62 detector layout

Figure 2.2: Kinematics of the $K^+ \rightarrow \pi^+ \nu \bar{\nu}$ process.

example, it allows a fraction of $K^+ \rightarrow \pi^+ \pi^0$ decays to be reconstructed inside Region I or Region II. The plot in figure 2.5 shows the m_{miss}^2 resolution as a function of the daughter particle momentum, as well as the various contributions arising from the resolution on the determination P_k , P_π , θ_K and θ_π .

2.3 Beam line

A primary proton beam at 400 GeV is extracted from the SPS and interacts with a beryllium target, generating a non separated secondary beam. The high energy protons are used in order to exploit the trend of the kaon production cross-section which increases with the proton energy; this helps also to reduce the non-kaon related activity in the detector reducing the number of protons needed.

The measured particle production data [37] established that the number of K^+ (K^-) decays in a given fiducial volume is maximum for $p_K \approx 0.23 \times p_0$ ($\approx 0.15 \times p_0$), where p_K and p_0 are the kaon and proton beam momenta respectively, however the highest K momentum achievable by the various stages of the beam optic and K tagging to be fitted in the available space is 75 GeV/c. Ultimately, we chose to use a positive rather a negative beam because the cross section is higher ($K^+/K^- \approx 2.1$) while the ratio of the kaon component is approximately the same for both charges.

After the production target quadrupole magnets are used to focus hadrons towards the center of a beam dump stage (labeled with TAX). The latter is an achromatic corrector composed by four dipole magnets and a momentum-defining slit in the middle, which allows to select a narrow momentum band (1% $\Delta p/p$). More quadrupoles focus the beam towards two collimators acting in both vertical and horizontal planes and align the beam to the optical axis of the CEDAR counter. The next stage on the beam line is a particle tracking system composed by a second achromatic corrector, made of four dipole magnets and three stations of the GTK.

Using an high momentum beam means that the kaons cannot be efficiently separated from other charged hadrons, mainly pions and protons, to obtain a pure kaon beam. The total rates of the secondary beam will be 750 MHz while

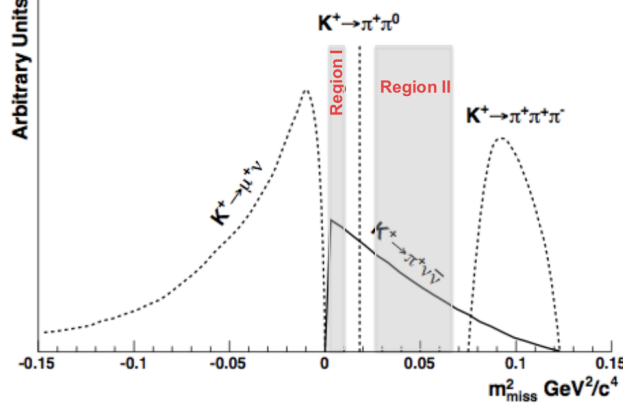


Figure 2.3: m^2_{miss} distribution comparison of the $K^+ \rightarrow \pi^+ \nu \bar{\nu}$ and of the main backgrounds (92% of the total) for a 75 GeV/c kaon. Limiting the analysis to the events in Region I or Region II, these backgrounds are (ideally) fully rejected. The m^2_{miss} of $K^+ \rightarrow \mu^+ \nu$ is not a delta, like $K^+ \rightarrow \pi^+ \pi^0$, because it is evaluated deliberately under the wrong assumption that the daughter track is a pion; the distribution have however an upper bound at $m^2_{miss} = 0$, so it can be fully rejected.

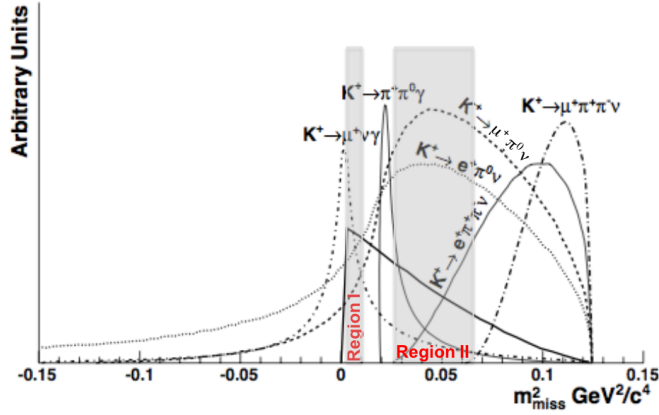


Figure 2.4: m_{miss} distribution for background not kinematically constrained.

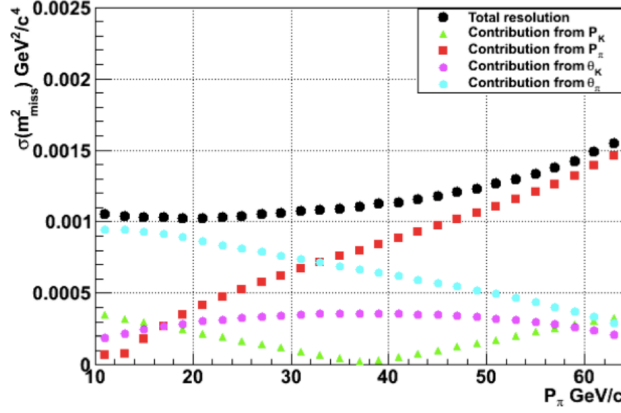


Figure 2.5: Resolution on m_{miss}^2 as a function of the π^+ momentum. The black circle is the full resolution while the other markers are the contribution coming from the uncertainty on various terms.

its components are: protons at 173 MHz, K + at 45 MHz, pi + at 525 MHz, mu+ at 6 MHz and e+ 0.3 MHz. The beam will produce about 4.5×10^{12} decays per year [35] in the fiducial region, which is 62 m long and starts at 105 m from the beryllium target.

2.4 CEDAR

The interactions of pions, kaons and protons with the residual gas in the vacuum decay tank were studied by a FLUKA simulation and the probability that these interactions can cause fake triggers was computed. The conclusion was that, in absence of kaon tagging, the vacuum should be better than 6×10^{-8} mbar: in this way the background is kept to less than one fake event per year. This very challenging requirement can be relaxed by at least an order of magnitude by tagging the kaons. To this aim, an upgraded form of the CEDAR built for the SPS secondary beams (CERN Report CERN-82-13) will be used by NA62. The CERN CEDAR (see figure 2.6) is designed to work as a particle mass selector: for a given momentum the Cherenkov angle of the light emitted by a particle traversing a gas of a given pressure is a unique function of the particle mass and emitted light wavelength. The Cherenkov light emitted by particles of different mass is then blocked by the CEDAR optics through the diaphragm slit onto the light detectors. The necessity of using CEDAR derives from the requirements to identify the particles of interest, i.e. Kaons, and distinguish them from pions and proton at a beam level. In order to obtain this, the NA62 CEDAR will be filled with hydrogen gas to reduce the beam multiple scattering. The gas will be at a pressure just below 4 bar. It is important to stress that given the extremely high rate of particles in input to the CEDAR the Kaon identification

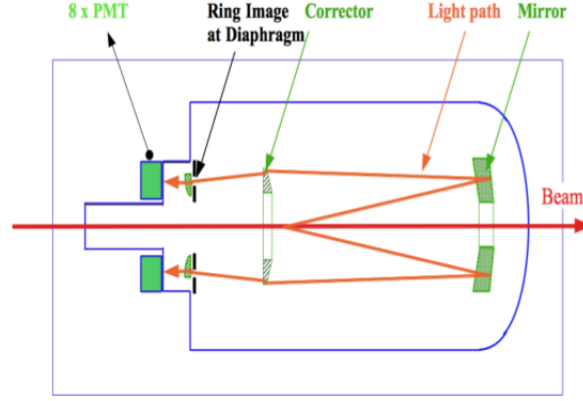


Figure 2.6: CEDAR schematic layout.

process must also be done with a very good timing resolution in order to reduce accidental background overlap. The kaon component rate in the high intensity beam for NA62 is 50 MHz. The CEDAR detector is required to achieve a kaon efficiency of at least 95% with a time resolution of 100 ps. CEDAR optics produce 8 light spots of eight $30 \times 8 \text{ mm}^2$ rectangular areas. Photodetectors mounted on CEDAR are Hamamatsu R7400U-03. The photon rate must be limited at 3 kHz/mm^2 ($\sim 50 \text{ MHz}$ per PM) in order to avoid system paralysis.

2.5 GTK

The GigaTracKer (GTK) has to measure the momentum and direction of the K^+ track. It will be subject to an high and non uniform beam rate of about 750 MHz over 1620 mm^2 with peak of 1.3 MHz/mm^2 . To preserve the beam divergence and limit the hadronic interactions, it should use a minimum amount of material ($0.01 X_0$). It also must have an excellent time resolution ($\sim 150 \text{ ps}$) in order to match the information from the downstream detectors. The requirement on the reconstructed track is of $\sigma(P_K)/P_K \sim 0.2\%$ on the momentum and $\sigma_\theta \sim 16 \mu\text{rad}$ on the angle with the beam axis.

The GTK is composed of three stations (GTK1, GTK2 and GTK3) placed along the beam line and mounted in between four achromat magnets (figure 2.7). Each GTK station is a hybrid silicon pixel detector with a total size of $63.2 \times 29.3 \text{ mm}^2$ containing 90×200 pixels of size $300 \times 300 \mu\text{m}^2$. The pixel thickness is $200 \mu\text{m}$ and it corresponds to $0.22\% X_0$. Including the material budget for the pixel readout and cooling, the total amount per station is below $0.5\% X_0$.

The silicon sensors are connected to the readout chip ($100 \mu\text{m}$ thick) by Sn-Pb solder bumps. Any connection to read-out chips is kept outside the beam in order to minimize the materials seen by it.

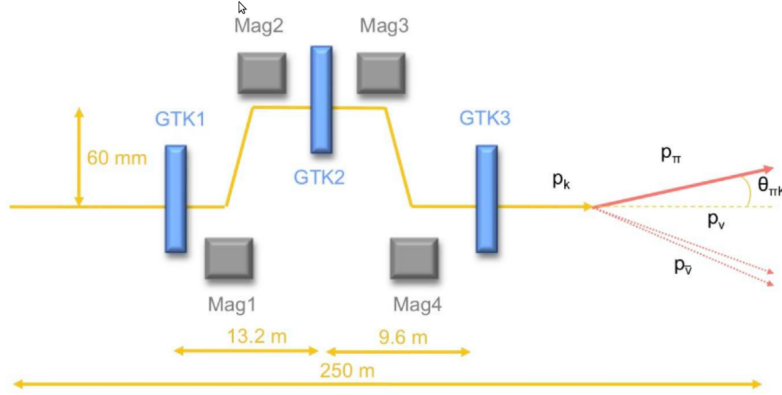


Figure 2.7: Schematic layout of the Gigatracker stations.

The GTK implements a compact micro channel cooling system that will keep the operative temperature at 5°C.

2.6 CHANTI

The CHarged ANTI counter (CHANTI) aims to detect and reject inelastic interactions between the beam and the last GTK station: these are the most harmful because can mimic the signal signature if a secondary pion is emitted in the STRAW acceptance and, for effect of STRAW resolution tails, is badly reconstructed as coming from the decay region. It has been proposed, designed and constructed by the Naples group of NA62 and is the main subject of this work.

See chapters 3, 4.1 and 5 for more details.

2.7 LAV

The Large Angle Veto (LAV) together with the LKr (2.11), (2.7), the IRC and the SAC (2.13), form a γ veto system for the decay products of the secondary π^0 ; the LAV has to detect γ coming from the decay volume with angles from 8.5 mrad to 50 mrad wrt the beam axis.

The required time resolution should be better than 1 ns in order to keep to an acceptable level the fake veto rate, while the inefficiency should be of the order of 10^{-4} for energies above 50 MeV in order to reach an inefficiency of 10^{-8} or less on π^0 detection (since the π^0 decays into two γ and at least one has to be detected). Energy resolution has a less stringent requirement of $10\%/\sqrt{E(\text{GeV})}$.

The LAV is composed of 12 stations, like the one shown in figure 2.8, located all along the vacuum tank from 120 m to 240 m (starting from the Beryllium

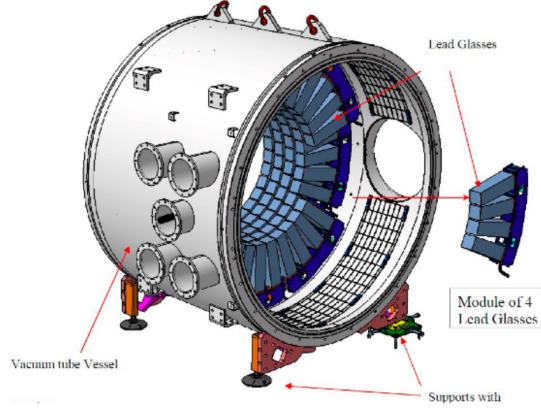


Figure 2.8: Schematic layout of a LAV station.

target position). Only the station 12 will be placed outside the vacuum tank. The LAV covers a photon energy distribution from 50 MeV to 30 GeV. The LAV building blocks are lead glass crystals from the former OPAL electromagnetic calorimeter [38]. Four crystal detectors, made of lead glass crystals and PMTs, are mounted on a common support structure forming an azimuth segment. The segments are assembled inside the vacuum tube in order to form a lead glass blocks full ring. Each LAV station is made up of 4 or 5 rings, which are staggered in the azimuthal direction providing complete hermeticity.

2.8 STRAW

A downstream magnetic spectrometer (STRAW) is used to measure the momentum and direction of secondary charged particles originating from the decay region. The kinematic rejection of the backgrounds require a $\sigma(P_K)/P_K < 1\%$ and a resolution on $\theta_{K\pi} < 60\mu\text{rad}$.

The STRAW, as shown in figure 2.9, is composed of a high aperture dipole magnet, providing a vertical B-field of 0.36 T, and four tracking chambers, working in vacuum in order to minimize the multiple scattering effects. Each chamber is equipped with 1792 straw tubes, a technology needed to cover large areas and operate in vacuum. Each chamber has 4 views, as shown in figure 2.10, providing measurements of two couple of orthogonal coordinates (x,y and u,v) rotated by 45° : this is necessary to eliminate the ambiguity in the track assignment in case of multiple hits.

The tubes are manufactured from $36\mu\text{m}$ thin PET (PolyEthylene Terephthalate) foils, coated (on the inside of the tube) with two thin metal layers ($0.05\mu\text{m}$ of Cu and $0.02\mu\text{m}$ of Au) to provide electrical conductance on the cathode. The anode wire (diameter $30\mu\text{m}$) is gold-plated tungsten.

In the central part of the chamber there is room for the beam pipe.

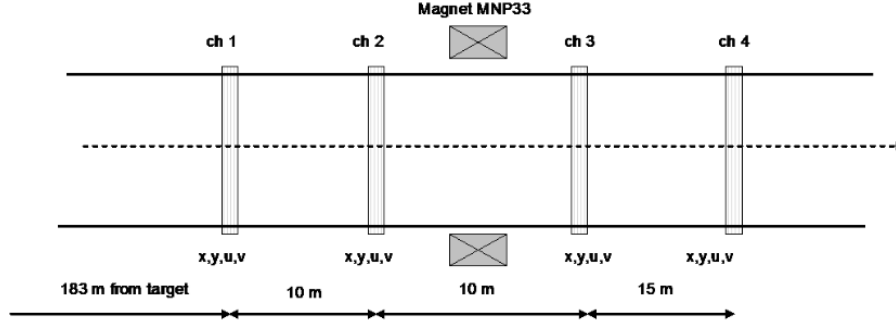


Figure 2.9: Layout of the STRAW chambers and magnet.

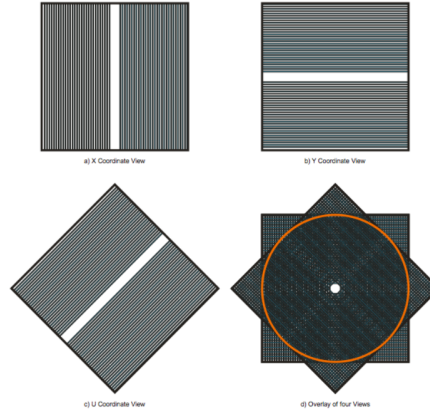


Figure 2.10: The 4 view forming a STRAW chamber (the V view is not shown but it has the straw tubes orthogonal to the ones of the U view).

Experimental requirements for the STRAW are:

- use of minimum amount of material ($\leq 0.5\%X_0$ for each chamber) along the particle trajectory to minimize multiple Coulomb scattering;
- a spatial resolution $\leq 130 \mu m$ per coordinate and $\leq 80 \mu m$ on the final reconstructed point;
- operation with an average particle rate of 40 kHz and up to 500 kHz per straw
- capability to veto events with multiple charged particles.

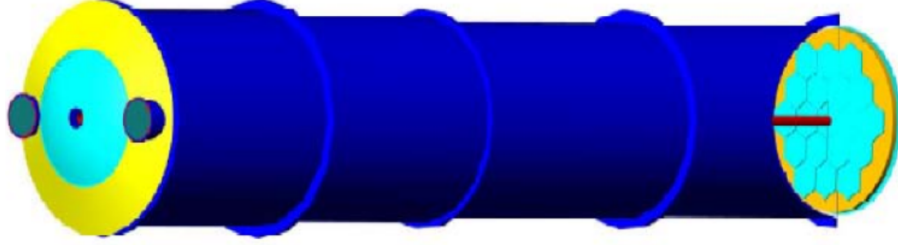


Figure 2.11: Schematic view of the RICH. On the downstream end there is the mirror array. The light will be focused on the opposite end where an array of PMTs will reconstruct the ring image.

2.9 RICH

The Ring Imaging CHerenkov (RICH) is part of the L0 trigger and it has to identify the secondary pions and muons in a energy range of 15-35 GeV. It must provide a muon suppression factor of at least 10^{-2} and a time resolution of 100 ps.

It is a Cerenkov detector that focuses the light emitted in the filling gas on a readout plane by means of a mirror system placed at the end of the vessel (see figure 2.11). The readout plane is made of a photomultiplier array that can reconstruct the image of collected light: it is seen as a circle whose radius is related to the Cerenkov angle.

Since the RICH will be about 17 m long, the mirror focal is constrained, so, in order to see the light cone produced by ~ 15 GeV pions, the refractive index of the medium must differ from 1 by 6.2×10^{-6} . For this reason, the neon at atmospheric pressure was chosen as medium.

To satisfy the time resolution request, we will use fast single anode photomultipliers. We have also to take into account the compactness of the photodetectors since we need to arrange about 2000 of them on the readout plane which will be of some (~ 5) m^2 . All these requirements lead to the choice of Hamamatsu R7400U-03 photomultiplier tubes (PMTs).

The mirror surface was segmented in 20 exagonal sub-surfaces (with a side of 35 cm) arranged to form a spherical mirror with curvature radius of 34 m. This solution was used in order to cover a big area of 6 m^2 .

In 2009 a test beam was performed with a RICH prototype equipped with 414 Hamamatsu R7400U-03 PMTs. A time resolution of 65 ps was measured and the π/μ separation was scanned from 15 to 35 GeV/c. The μ suppression factor was estimated to be 10^{-2} integrated in the momentum range of interest [39]. In order to suppress $K^+ \rightarrow \mu^+ \nu$ events, in addition to the STRAW and RICH contribution, we need a muon veto system which should provide a further muon reduction of the order of 10^{-5} with respect to pions.

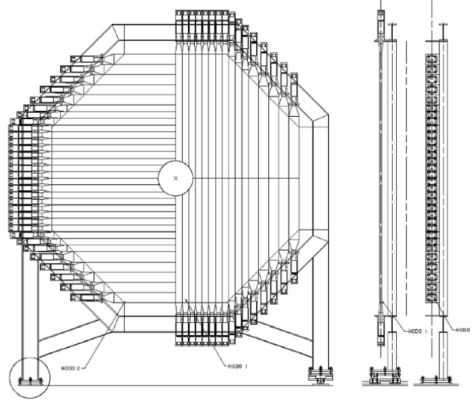


Figure 2.12: CHOD scintillator arrangement.

2.10 CHOD

The Charged HODoscope is part of the L0 trigger and it provides the timing to charged decay product with a resolution of 200 ps. It will be used also to detect the products of photo-nuclear interactions in the RICH mirror plane.

It is located, in front of the electromagnetic calorimeter and it is composed by 128 plastic scintillation counters (BC498), 2 cm thick ($\sim 0.005 X_0$). The light will be collected by plexiglas light guide and it will be read by Photonis XP2262B photo-multiplier.

As shown in figure 2.12, there are two planes, spaced by 75 cm: in the first plane the counters are arranged horizontally while in the second they are placed vertically. The length of the scintillators vary from 60 cm up to 121 cm, and the width from 6.5 cm in the region close to the beam pipe, to 9.9 cm.

The CHOD is the same used in NA48 but the front-end and readout electronics have to be entirely redone in order to cope with particle flux estimated to be around 11 MHz.

2.11 LKr

The Liquid Krypton calorimeter (LKr), together with the LAV (2.7), the IRC and the SAC (2.13), form a γ veto system for the decay products of the secondary π^0 ; the LKr has to detect γ coming from the decay volume with angles from 1 mrad to 8.5 mrad wrt the beam axis. It is also part of the L0 trigger.

The photon veto inefficiency is mainly due to π^0 decays with one photon emitted at low energy which is either outside the acceptance of the LAV or inside it, but with high inefficiency. The other photon is mainly impinging the LKr calorimeter, with an energy larger than 35 GeV, or the IRC with an energy larger than 60 GeV [35]. The requested inefficiency is lower than 10^{-5} at those

energies, to reduce the contribution from these events to the average π^0 rejection inefficiency.

The LKr is the same calorimeter of the NA48 [40] experiment performed at CERN to study the direct CP-violation in the neutral kaon system [41]. In the NA48/2 experiment it was possible to measure the LKr inefficiency that results in accordance with the NA62 requirements: $(2.8 \pm 1.1_{stat} \pm 2.31_{sys}) \times 10^{-5}$ [42]. It is positioned outside the vacuum tank, and will see photons with minimum energy of 1 GeV. The LKr is a quasi-homogeneous ionization chamber that, due to the short Kr radiation length ($X_0 = 4.7$ cm) and Molière's radius ($R = 6.1$ cm), it is able to contain high energy electromagnetic showers (50 GeV showers are $>99\%$ contained) but maintaining a relative compactness.

It contains about $10m^3$ of liquid krypton, it has an octagonal section, as shown in figure 2.13, and it is 127 cm long ($27 X_0$) with an hole of 80 mm radius to house the beam pipe. The readout consists of copper/beryllium ribbons extending from the front to the back of the detector, dividing the volume into 13248 readout cells with a transverse size of approximately 2×2 cm² each, as shown in figure 2.13.

After calibration and corrections, the energy resolution of the LKr calorimeter can be parametrized as $\sigma(E) = 3.2\% \sqrt{E} + 9\% E + 0.42\%$, in units of GeV. The first term depends on stochastic sampling fluctuations, the second one on electronic noise and Kr natural radioactivity and the last one on inhomogeneities, material in front of the calorimeter and the non perfect intercalibration of the cells. The spatial and time resolutions are $\sigma_x = \sigma_y = 42\% \sqrt{E} + 5\%$ and $\sigma_t = 2.5 \sqrt{E}$ in units of GeV, cm and ns.

Inside the the LKr, at $\sim 9.5 X_0$ (the dept where the development of 25 GeV shower is maximum), are also positioned 256 vertical bundles of scintillating fibers that can be used as a "Neutral" horscope (NHOD) with a resolution of 150 ps. It will provide a control signal to measure the efficiency of the main data acquisition trigger.

2.12 MUV

The fast MUon Veto (MUV) is part of the L0 trigger and it has to partially suppress muon event. It should have a time resolution below 1 ns to reject events with coincident signals in the GTK and the CEDAR. The detector consists of 3 distinct parts (MUV1, MUV2 and MUV3) placed one just after the other along the beam axis. The MUV1 and MUV2, shown in figure 2.14, follow directly the LKr calorimeter and work as an hadronic calorimeter for the measurement of deposited energies and shower shapes of incident particles. These modules are a sequence of iron and scintillating strips (24 layer for the MUV1 and 22 for the MUV2) placed alternatively horizontally and vertically, to form X-Y views. In the MUV1 the light is read by wavelength shifting fibers and then by photomultipliers, while for the MUV2 the photomultiplier are placed directly on the scintillators (coupled by a light guide).

The MUV2 is followed by a 80 cm thick iron wall and then the MUV3

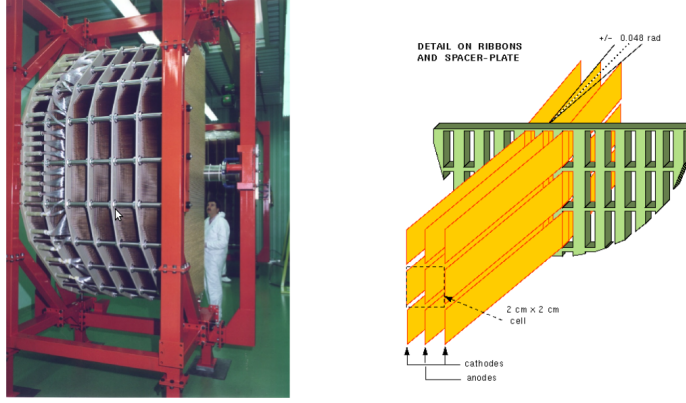


Figure 2.13: LKr calorimeter. Left: structure while under construction. Right: details of the cell structure.

module, which is a Fast Muon Veto and detects non-showering muons, reaching the desired time resolution. To achieve the required time resolution, the MUV3 is made of scintillator tiles arranged to minimize the differences in the light path. Also for the MUV3 the light is read by a photomultiplier placed on the scintillators.

2.13 IRC and SAC

The Inner Radius Calorimeter (IRC) and the Small Angle Calorimeter (SAC) are parts of the γ veto system for the decay products of the secondary π^0 . The SAC has to detect γ coming from the decay volume with angles from 0 mrad to 1 mrad wrt the beam axis, while the IRC will cover a dead region at radii between 7 and 14 cm in front of the LKr.

The IRC will be placed in the front of the LKr as close as the to the beam pipe, while the SAC at end of the experimental setup, so the only geometric requirement is not being hit by the deflected undecayed beam. These two small detectors will have a single-photon detection efficiency better than 10^{-5} for photon energies higher than 5 GeV.

Both the detectors will use “Shashlyk” technology, shown in figure 2.15 which is based on alternate layers of lead and plastic scintillator read by wavelength shifting fibers [43]. The fiber are read in bundle by classic photomultiplier tubes, as shown in figure 2.16.

2.14 Trigger and data acquisition

Rare decay experiments need intense flux so the performance and the dead time of the Trigger and Data Acquisition system (TDAQ) is a crucial factor. The

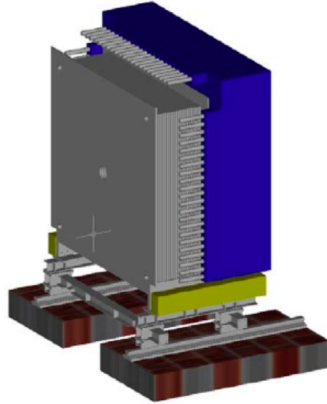


Figure 2.14: Schematic view of MUV1 (gray) and MUV2 (blue). The beam is coming from the left.

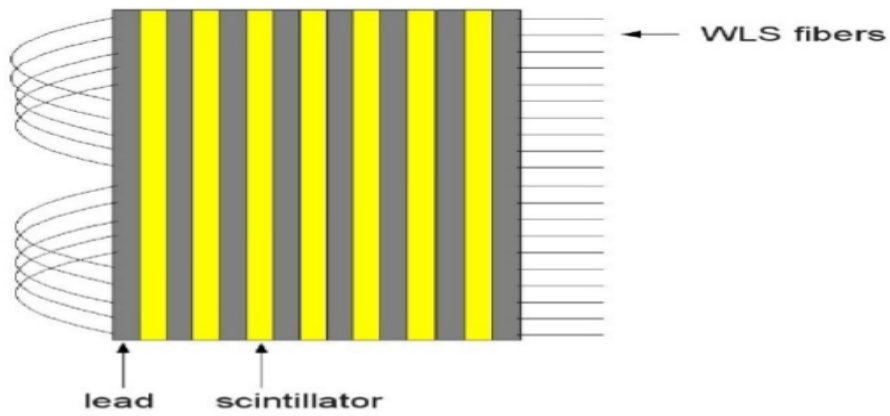


Figure 2.15: Shashlyk layout.

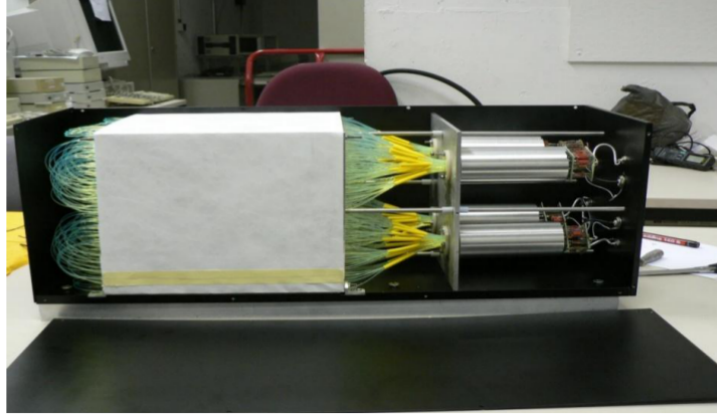


Figure 2.16: SAC prototype.

NA62 solution [44] uses a common clock, with a frequency of ~ 40 MHz, optically distributed to all systems by means of a Timing, Trigger and Control (TTC) electronic module designed for LHC experiments [45]. The only on-line trigger is the Level-0 (L0) one which uses the input from the CHOD, the MUV, the LKr and, optionally, from the RICH, the LAV, and the STRAW.

The data acquisition and the elaboration of the trigger primitives relies for all the detector on TDC and a custom board, the TEL62. The only exception is the LKr that uses a dedicated board in place of the TDC. The whole system, sketched in figure 2.17, is designed to work with time data, indeed the first stage of acquisition is made of fast Time to Digital Converter (TDC) that accept signals in the LVDS standard, so all the detector are supposed to provide this kind of signals. The TDC are arranged in boards containing 128 channels, and each group of 4 TDC board are connected to one TEL62 powered by 5 Altera FPGA.

We chose a technology based upon FPGA so that each detector, with small firmware changes, can make some online data elaboration in order to modify the data before sending it to the farm (ex. calibration) or to produce some online data (ex. CHOD data for L0 Trigger). However the main task of four of the FPGA on the TEL62 (PP FPGA) is to manage the communication with the TDC boards (one PP for each TDC board) and to store them temporary in a 1 Gb of external RAM (for each PP). The fifth FPGA (SL FPGA) has to wait and decode trigger signals from an optical line and propagate the trigger to the PP that then have to search in the RAM the corresponding data and to send it back to the SL. The final task of the SL is to pack this data and to send them to the PC farm through four gigabit Ethernet boards.

On the TEL62 is also present a credit card PC (CCPC), a fully featured PC running Linux, that is able to flash the FPGA firmware and to communicate with them by an ECS bus. This permits us to set and read FPGA registers and to automate various tasks by means of simple scriptings.

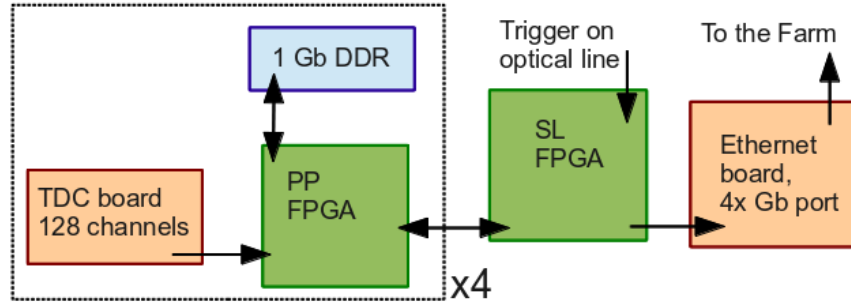


Figure 2.17: Schematic view of the TEL62 and TDC board.

The farm will merge the data from all the detector TEL62 board; it will perform some basic L1 Trigger and reconstruction task and will store the result on disk. This architecture needs, in order to properly work, a fast switch and a dispatch mechanism that will ensure that all the pieces of data of each event will be sent to the same PC.

Chapter 3

CHANTI design

3.1 Detector requirements

A GEANT 4 simulation (see chapter 4) has shown that the interactions between the K^+ beam and the GTK material can generate a potentially harmful background for the $K^+ \rightarrow \pi^+ \nu \bar{\nu}$ measurement. In fact in these inelastic interactions a π^+ can be emitted in the STRAW acceptance and simulate the $K^+ \rightarrow \pi^+ \nu \bar{\nu}$ signature. This can happen when all the other particles are emitted at large angle with respect to the beam axis, so that no one can be seen by any NA62 detector.

The interactions on the last station of the GTK are particularly dangerous since they could be completely invisible to the tracker itself.

Simply moving the beginning of the allowed fiducial region far from the GTK is not a solution since the spatial resolution on the reconstructed vertex, which is performed by means of the momentum measured by the GTK and STRAW, in case of inelastic interactions, has long non Gaussian tails (see chapter 4).

In order to suppress (and study) this background we designed the CHarged ANTI counter (CHANTI) which has to detect the charged products of the inelastic interactions that are emitted at large angle. To this purpose the CHANTI has to be as close as possible to the last station of the GTK (see figure 3.1).

We need the CHANTI inefficiency in detecting charged particles below percent level (see chapter 4) and it has to provide hermeticity for the angles not covered by the other detectors of NA62 (mainly the LAV system), which can detect particles produced on the GTK-3 at a maximum angle of 49 mrad. For this reason, CHANTI has to stay close to the beam so that the muon halo will be a major issue; considering the contribution of muon halo and the beam inelastic interactions, the estimated total rate on the detector will be about 2-3 MHz.

Although if the CHANTI is not part of the level 0 trigger, we need a good time resolution (< 2 ns) in order to reduce the probability of rejecting a signal event due to accidental coincidences. With a 2MHz rate, and supposing to reject

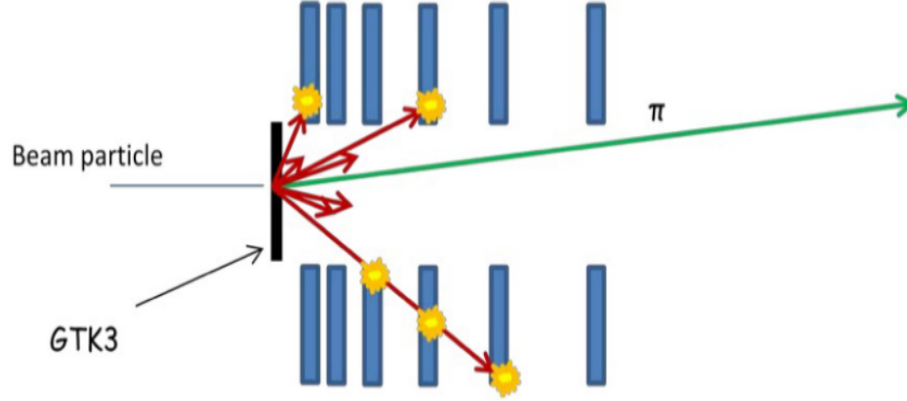


Figure 3.1: The beam, interacting with the last station of the GTK, generates inelastic products. The CHANTI will detect particles emitted at large angles.

all the event in a time window of $\approx 5 \times \sigma_t \approx 10ns$, around the event time we introduce a 2% inefficiency on the signal.

A minor issue is the capability to perform some particle tracking in order to validate our inelastic event simulation and to study the beam halo. This will also improve the time resolution allowing to correct the raw times for xy position time offsets.

Other issues are a low outgassing rate and a low power consumption since the CHANTI will work in vacuum, and a radiation hardness up to some Gy/year.

3.2 General layout

The detector is made of six square stations, the first five placed in the same vacuum chamber that we are designing and constructing in Naples and that will house also the GTK-3, and the last one in a dedicated vacuum chamber. The distances of the stations from the GTK-3 surface are: 27-85-200-430-890-1810 mm (see figure 3.2).

Each station is a $300 \times 300 \text{ mm}^2$ square with an hole of 65 mm along Y and 95 mm along X; it is needed to let the CHANTI not being hit by the beam and its size reflects the beam asymmetry. The CHANTI covers the angles between 26.2 mrad to 1.38 rad with respect for particles coming from the center of the GTK-3, and angles from 49 mrad and 1.16 rad for tracks arising from a corner of the GTK-3. This geometry assures the hermeticity down to 49 mrad since the LAV and the other detectors will cover the angles below 49 mrad.

The vessel housing the first five CHANTI station and the GTK-3 is made by AISI 304L stainless steel plates welded together to form a rectangular box, with both top and bottom removable plates coupled via a suitable O-ring system to the rest of the box in order to keep the vessel vacuum tight and to provide at

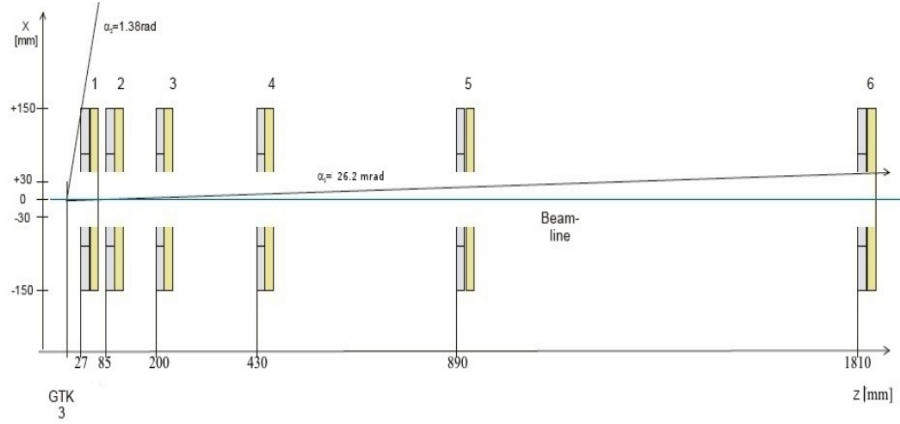


Figure 3.2: Position of the six CHANTI stations along the beam line.

the same time easy access to the detector for maintenance if needed. The vessel has two ISO-F 200 compliant holes closed by two flanges which host the signal feedthroughs. These are SUB-D37 male-to-male T.I.G. welded connectors. Out of the 37 pins in one connector, one is unused, four are used to read Pt100 temperature probes and 32 are used to connect 16 channels. Each couple of pins is connected to coaxial cables on both sides of the flange and these lines are used both to polarize the SiPMs (see 3.5) and to read the output signals. A 3D drawing of the vessel is shown in fig. 3.3.

As shown in figure 3.4, each station is made of 46 scintillating bars in the shape of prisms with triangular base arranged in two view, X and Y, each containing two layers. In order to be able to correctly shape the hole, the bars have to be of variable length:

- 22 “Long” bars, 300 mm long, 10 in the X view and 12 in the Y view;
- 14 “Medium” bars, 102.5 mm long, used only in the X view;
- 10 “Short” bars, 117.5 mm long, used only in the Y view;

Inside each bar there is a wavelength shifting fiber (WLS) collecting the scintillation light, and mirrored from one side. On the other side a silicon photomultiplier (SiPM) readout is coupled to the fiber by means of a dedicated mechanical housing and its signals are readout through a dedicated coaxial cable connector.

3.3 Scintillators

We use scintillators produced by extrusion at the NICADD facility at Fermilab [46, 47] (the same used in the MINERVA and D0 collaboration [49, 48]). They

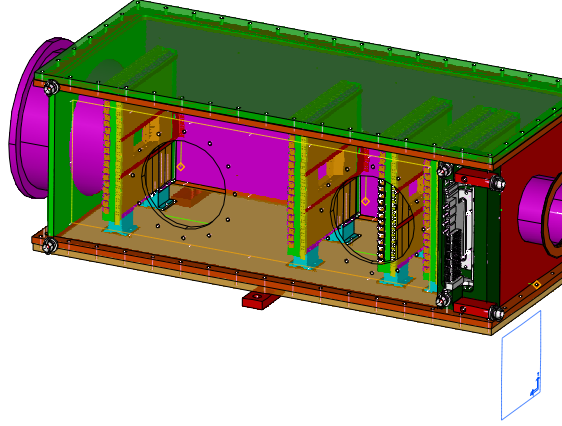


Figure 3.3: 3D Drawing of the vessel hosting GTK-3 and the first five stations of the CHANTI

have a polystyrene core doped with fluorescent compounds which emit in the blue (1% PPO and 0.03% POPOP), and a co-extruded 0.25 mm coating made of TiO₂ which, diffusing the light, helps to increase the portion of collected light. Their main characteristics are:

- good light yield (100% of Kuraray SCSN-81);
- radiation hardness (5% light yield loss for 1Mrad irradiation);
- fast response (1-2 ns);
- low cost.

In figure 3.5 the transmittance and fluorescent spectra of the scintillator are shown[46]: there is a cut-off in the absorption at $\sim 400\text{nm}$ and an emission peak at 420 nm (blue light).

The transverse section of a scintillator bar is an isosceles-rectangle triangle with an hypotenuse of 33 mm. In the center, at 8.5 mm from the hypotenuse, there is a hole that houses the WLS fiber. In our setup we fill this hole with an optical glue (SCIONIX Silicon Rubber Compound glue) with a refractive index carefully chosen to optimize the coupling between scintillator and fiber cladding (1.406).

Thanks to the triangular shape we can construct an almost self-sustaining plane (pratically) without dead regions, as shown in figure 3.6. Furthermore a track passes always in two bars emitting light proportionally to the length traveled in each bar so weighting the fibers position with appropriate coefficients we can estimate the track hit position with a precision higher than the granularity of the detector. In fact, despite the fiber being 17 mm distant, we can reach a space resolution of 3 mm.

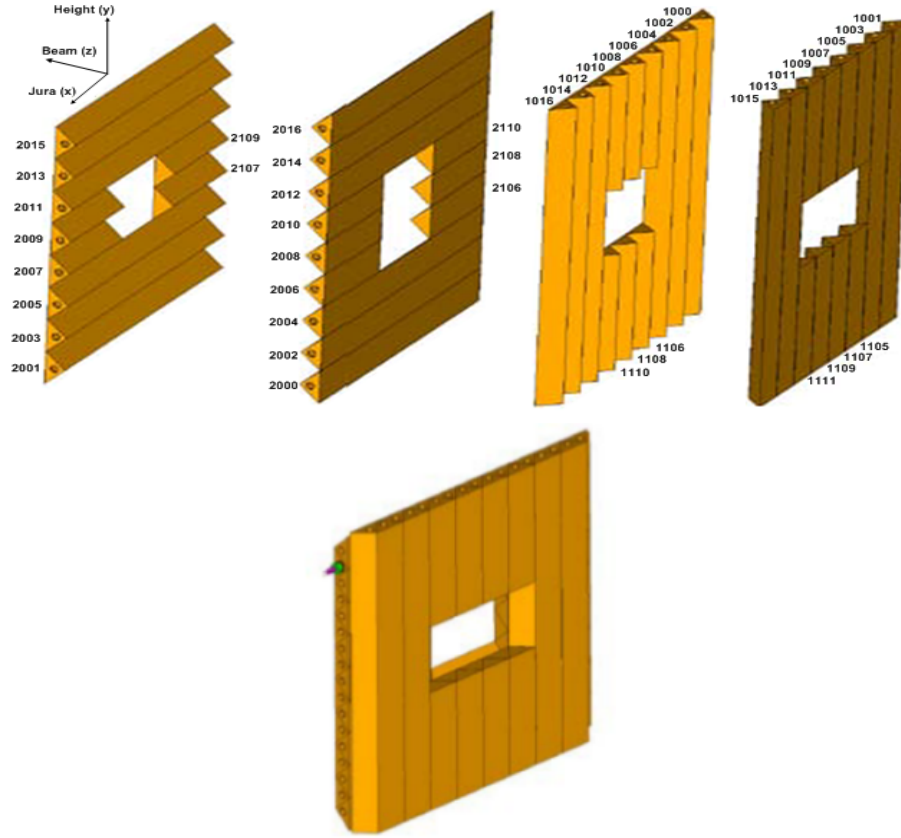


Figure 3.4: Layout of a single CHANTI station. Up: different bars forming the views with channel numbering scheme (including a 2 digit station id prefix). Down: the assembled station.

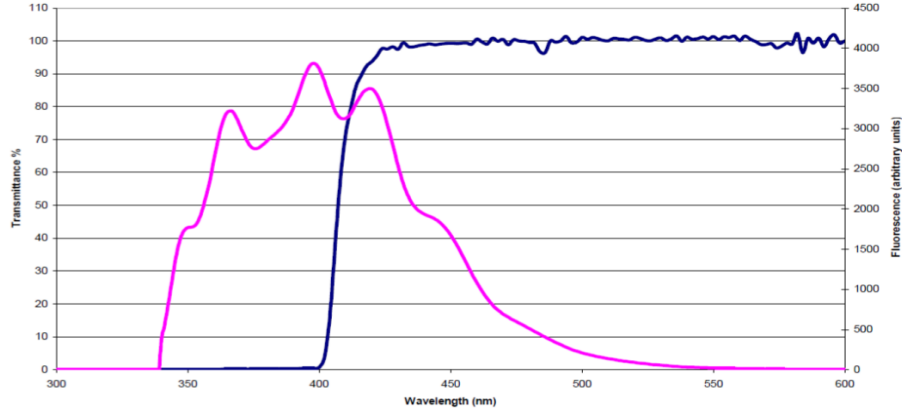


Figure 3.5: Transmittance (dark blue line) and fluorescent (light blue line) spectra for the CHANTI scintillator.

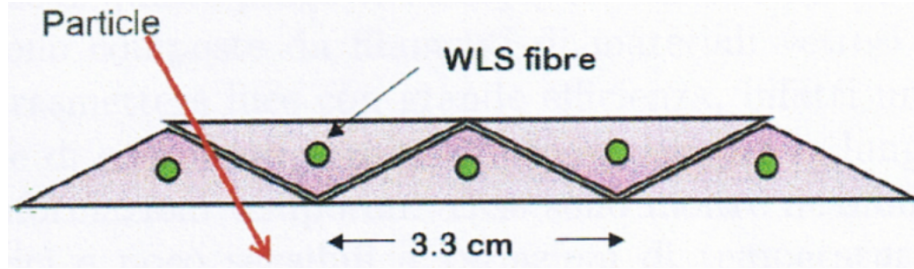


Figure 3.6: Section of a CHANTI view. The structure does not present dead regions and allows to interpolate the position of the crossing particle by means of different amount of light collected by neighbour fibers.

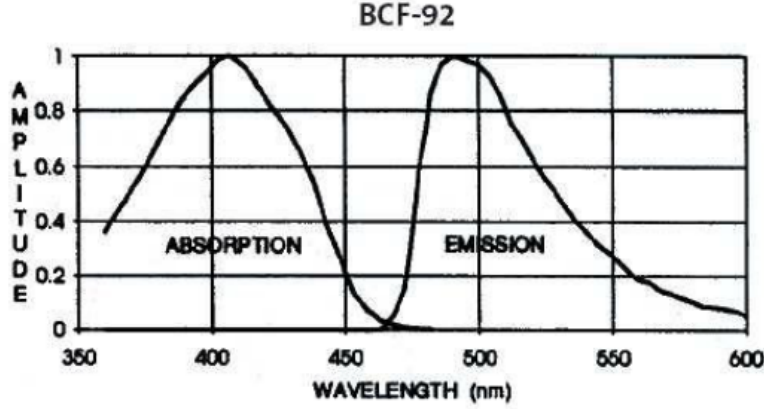


Figure 3.7: Absorption and emission spectra for the Saint-Gobain BCF-92 wave-length shifter fiber.

3.4 Optical fibers

We use the BCF-92 WLS fiber produced by Saint-Gobain which has a total radius of 0.5 mm and it is constituted by a polystyrene core and an acrylic coating ($C_5H_8O_2$, 3% of the radius). The core has a refractive index of 1.60 and a density of 1.05 g/cm^3 , while the coating has a refractive index of 1.49 and a density of 1.19 g/cm^3 .

The fiber core is doped with a particular combination of compounds that let them absorb the light at one wavelength and re-emit it at a larger wavelength. In figure 3.7 [50] it is shown that the absorption has a maximum at 420 nm (blue light) while the emission have the maximum at 492 nm (green light). In this way the coupling between scintillator and fiber is optimized. The green light is captured in the fiber, due to the total reflection, when it has an angle with the normal at the interface core-coating of 68.6° .

The BCF-92 have also a long attenuation length, about 3.7 m, and a fast re-emission time $\tau \approx 2.7 \text{ ns}$.

To increase the portion of light collected, one side of the fiber is mirrored by means of a aluminum deposition technique which was developed for the ALICE calorimeter (Al sputtering in vacuum [51]). On the other side there is a mechanical housing for the SiPM, shown in figure 3.8, to which the fiber is glued.

3.5 SiPM

The silicon photo-multipliers (SiPM) are becoming a widespread solution in particle detectors where high number of channels or high level of integration is needed [52, 53, 54]. A SiPM is an array of p-n junctions, electrically connected



Figure 3.8: Fiber glued to the SiPM housing connector.

in parallel, that work in Geiger regime, each of them quenched by a dedicated circuit. More details on how this device works are given in the Appendix A, here we will just recall that their junctions must be inversely polarized to $\sim 70V$, and that their dark current is typically about hundred nA.

We use Hamamatsu SiPM with passive quenching system (resistors), an active area of $1.3 \times 1.3 mm^2$ and a fill factor of about 62%. In figure 3.9 there is a photo of the device with a zoom on the junction arrays. Typical photon detection efficiency (PDE) for this product are in figure 3.10

The advantages of using this kind of photodetector are:

- Extremely compact - they could be placed directly on the bars avoiding the transportation of the light far away through longer fiber (that introduce additional light loss);
- Low power/heat dissipation - they are alimented by 70V and absorbing current of a few nA;
- Good rates - they work well at rates above 1 MHz;
- Good radiation hardness - they have a good radiation hardness [55], but there is a known issue with proton/neutron flux; however we checked trough a dedicated FLUKA simulation that it will not be a major issue in our environment [56]

For the following discussion we recall the concept of “Dark rates”: a SiPM generates signals also when no light reaches the active area. The principal effect contributing to the dark rate is the thermally produced electron-hole pair, but a relevant secondary effect is the tunnel effect that could move an electron from the valence band to the conduction one. For different devices and applied bias voltage the dark rate (i.e. the rate of counts at 0.5 p.e. threshold) may range from 100 kHz to several MHz per mm^2 at 25° C. For small value of $(V_{bias} - V_{bd})^2$,

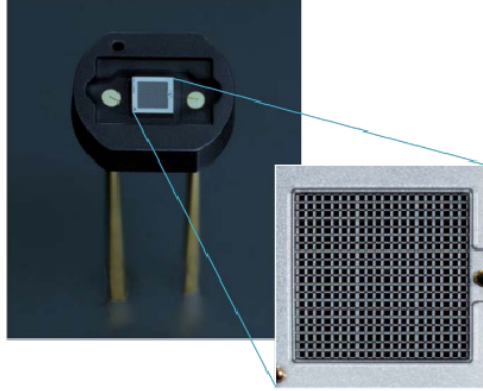


Figure 3.9: Photo of a Hamamatsu SiPM with a zoom on the active area showing the junction array .

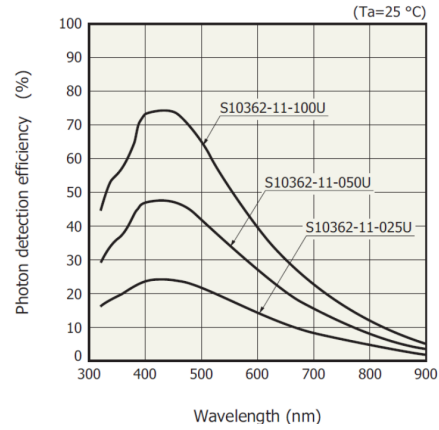


Figure 3.10: Typical PDE for different Hamamatsu SiPM and their wavelength dependence.

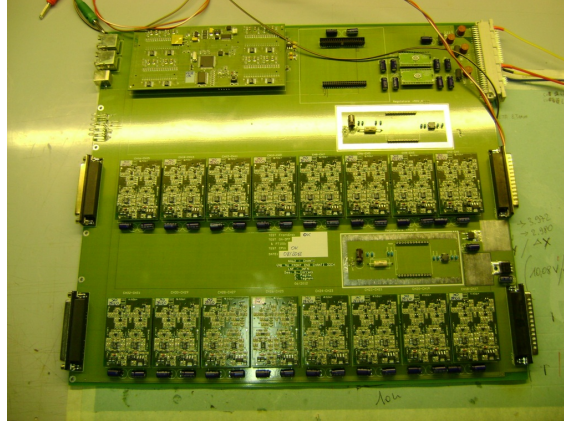


Figure 3.11: View of the CHANTI readout board.

where V_{bd} is the breakdown voltage of the SiPM and V_{bias} is the applied voltage, the dark rate is proportional to $(V_{bias} - V_{bd})^2$. Hamamatsu Photonics suggests an “operative” value for the bias voltage at 25°C: adopting this value all devices have a typical dark rate of about 800 kHz and a typical gain of 7×10^5 .

3.6 Readout electronics

For the powering and reading of the SiPM, we developed in collaboration with “Servizio Elettronica-Laboratori Nazionali di Frascati” of INFN an all-in-one solution. Each board, shown in figure 3.11, is a standard VME 9U with CAN interface and it can control up to 32 SiPM. Its main features are:

- it can independently power up each SiPM by setting voltage bias with 10 mV precision;
- it can read absorbed current with few nA precision;
- it has a very good temperature stability (25 ppm/K);
- it has a fast signal amplification (x25);
- it can read temperature through pt100 probes.

The analog signal generated by this board will be processed by a Time Over Threshold board already designed for the LAV detector of NA62 [57]. For the digitization we will use the standard NA62 data acquisition chain (see section 2.14): a fast TDC (100 ps resolution) and a dedicated board powered by 5 Altera FPGA that fetch the data and send it to the pc farm through fast Ethernet.

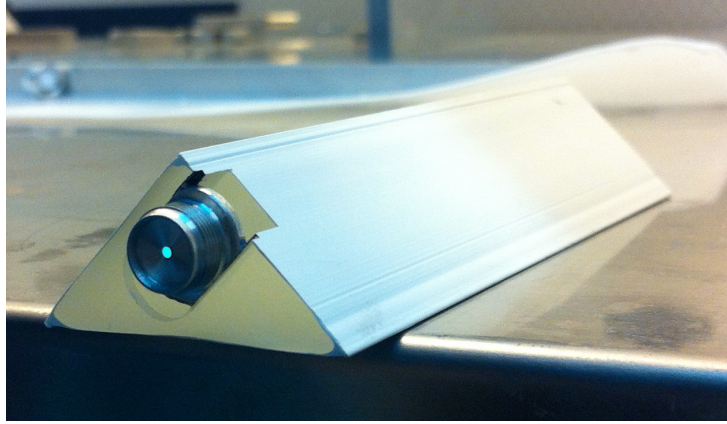


Figure 3.12: Fully assembled bar.

3.7 Construction

The construction phase starts by cutting the scintillator bars at 3 different lengths: 30 cm, 11.75 cm and 10.25 cm. One side of each bar was machined to create the seat of the connector coupling the fiber to the SiPM (fig 3.12).

Fibers are mirrored at one side (opposite to SiPM), cut at the right length and glued to the connectors. Then, each fiber is inserted in the hole of the scintillator bar, kept vertical by a dedicated support. The free space between the fiber and the scintillator is filled by an optical glue (GE Silicones - Silicone Elastomer), injected slowly from the bottom to prevent air mixing. The glue was previously placed in a vacuum pump to extract the air in it that otherwise could create bubbles during the fixing.

A structural glue (Epoxy 3M-DP490) is used to fix the connector to the scintillator. The final aspect of the bar is shown in figure 3.12.

When all the bars are ready, we perform a quality check (see section 3.9) and a characterization of each of them. We also characterize every SiPM (see section 3.8).

After the tests, we used an alignment jig and the structural glue to assemble the station; the result is shown in figure 3.13. At the end the station is housed in an aluminum frame and the SiPM are placed in their connectors. Since the SiPMs could be damaged by relatively small force applied to their pins we designed a cabling solution based on a custom connector, a frame support and nylon bands. The first station, fully assembled, is shown in figure 3.14. Thanks to the adoption of appropriate design features and material choices, the overall outgassing of one station has been kept very low: it has been measured to be less than $3 \cdot 10^{-5}$ mbar l/s. All the construction procedure takes about 20 days per station.

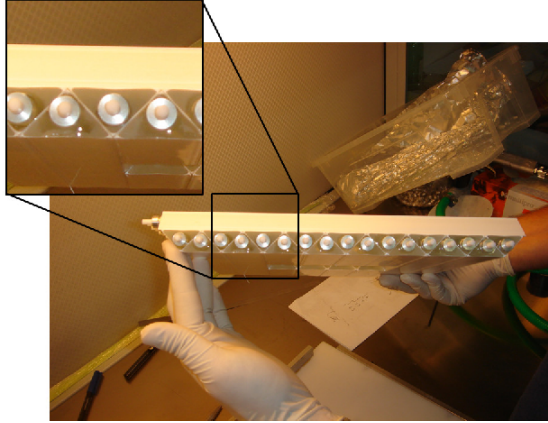


Figure 3.13: Station prototype just assembled.

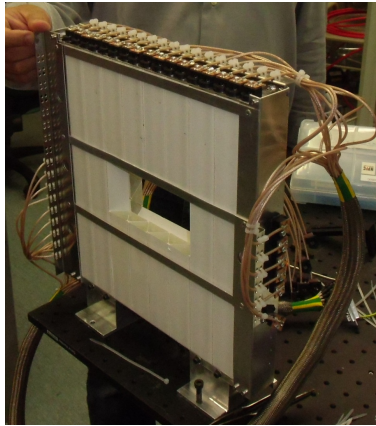


Figure 3.14: First station with mechanical frame and full cabling

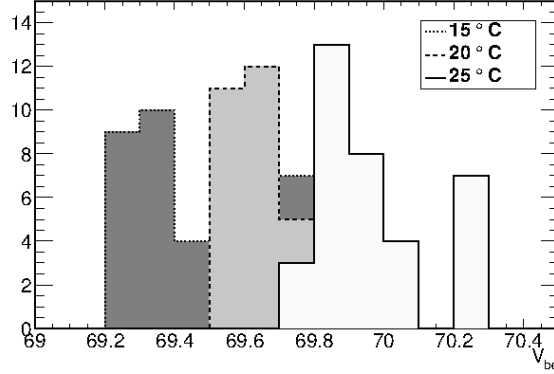


Figure 3.15: Breakdown voltage of the SiPM of the first station measured at 3 temperatures.

3.8 SiPM Characterization

The SiPM gain (and so the output signal height) depends linearly on $V - V_{bd}$, where V is the voltage applied to the SiPM (70V is a typical value) and V_{bd} is the breakdown voltage. So we need to know and monitor V_{bd} which is a function of temperature (69.5 V is a typical value at 20 °C). We can do this using the Voltage-Current response of the SiPM. Near the breakdown region, this curve is approximated by $I = \alpha(V - V_{bd})^2$ where I is the absorbed current. So we can measure V_{bd} by a fit to the V-I curve. We designed an automatic system to test 32 SiPM at the same time; it uses a power supply with a nano-amperometer and an analog multiplexer, both controlled by a computer and LabView software. We performed the test in thermal chamber at 3 different temperatures: 20, 25 and 15 °C while the Hamamatsu specifies the SiPM properties only at 25 °C. The result for the first 36 SiPM used in the NA62 technical run in November 2012 are shown in figure 3.15. In the final NA62 setup we will be able to routinely check V-I curves online through the front-end electronics.

3.9 Bars characterization

Once bars are glued inside the station they cannot be substituted if faulty. In order to get the required level of efficiency we have to check the response of each individual bar before assembling the detector. We did thus checks using cosmic rays. As the low rate of events did not allow for a study of the response of the bars as a function of the longitudinal coordinate, we performed two kind of tests.

In the first kind of test, the signal, read by a simple circuit, was sent to a large bandwidth, ≈ 11 ns, amplifier and the whole waveform was acquired using

a fast digital oscilloscope connected to a PC. The test was performed in auto-trigger mode: we acquired each event with a signal higher than 50 mV. This threshold was chosen in order to have an acquisition rate dominated by cosmic ray signals, in fact, with this threshold, the contribution from electronic noise and SiPM thermal noise is negligible (at level of 1% or less). The bars and electronics were put in a thermal chamber that allowed us to perform the test at a constant and fixed temperature of 25 °C. With this test we studied the global response of the bar. Smaller (larger) signals were mostly due to cosmic muons crossing the bar perpendicularly (not perpendicularly) to its axis. As a variable to be used for the quality check we computed the ratio R of the number of signals exceeding 150 mV over the total.

In the second kind of test, the apparatus was the same but we had also two small scintillator counters and phototubes above the bars to form a telescope, in order to select almost straight cosmic rays passing next to the edge of the bar opposite to the SiPM. We acquired the signal only when there was a coincidence between the signals of the two counters. In this way we can have an idea of the response of the bar to a roughly vertical MIP muon, in the worst case scenario with respect to the fiber attenuation. For this kind of test we expressed the result as the mean number of generated photoelectrons N_{pe} which was obtained comparing the signal to the single photoelectron signal, that was previously determined for each SiPM in a dedicated run measuring the dark noise.

We show R vs N_{pe} for the first 96 bars in figure 3.16. In this figure the two triangles refer to two “Reference” bars intentionally produced following a “Wrong” procedure: one without the optical glue and the other with a bad gluing process (we let the glue to flow out during the fixing obtaining a non uniform gluing). The ellipses shown in figure 3.16 correspond to the 1σ , 2σ and 3σ contours obtained after a rotation to find out the two uncorrelated variables. We decided to reject all the bars under-performing out of 2σ contour.

In figure 3.17 it is shown the time resolution obtained with the same acquisition system described above, in an external triggering test with a long bar. While of course the complete signal shape is acquired in this case via the waveform digitizer, we emulate via software the expected performance of the final readout electronics by selecting the time the signal crosses a given threshold as the “Leading time” of the signal, and the time it remains above the same threshold (ToT) to correct for the time-walk effect. This correction is based on the functional (average) dependence of the time over threshold versus the delay between the signals of the bars and the triggers.

The final single channel time resolution is estimated to be ≈ 0.9 ns : of course this must be checked in a realistic environment and with the final front-end electronics: this has been the main result obtained at the Technical Run in November 2012 and will be discussed in Chapter 5.

A perpendicular muon track always crosses two neighboring opposed bars. We report in figure 3.17 the sum (over all the channels) of the collected charge in terms of number of photoelectrons. This was obtained by a 5 bar prototype using an external cosmic ray trigger and by reading two opposite bars by a prototype electronics and a digital oscilloscope.

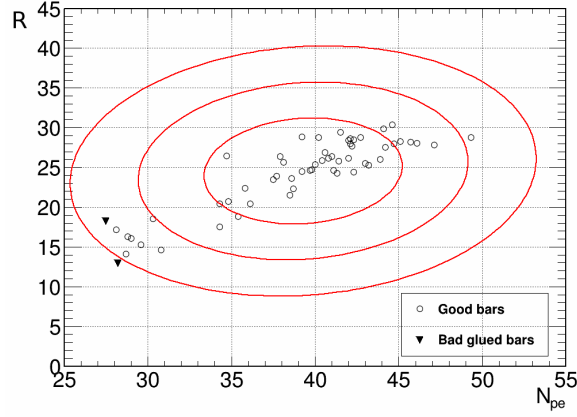


Figure 3.16: Autotrigger R versus External-trigger N_{pe} for the first 96 bars. The ellipses represent the 1σ , 2σ and 3σ bound. The triangles are the bad glued bars.

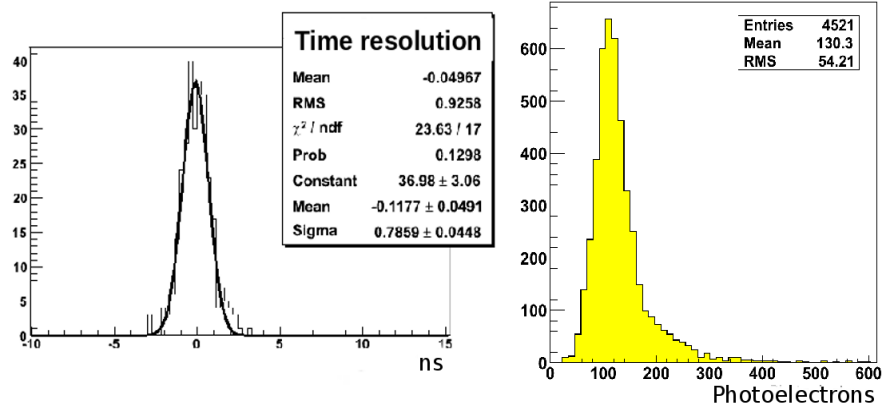


Figure 3.17: Left: Time distribution with time-slewing correction estimated by software emulation of the front-end electronics on acquired waveforms with external trigger. Right: Collected photo-electrons in two neighbors bars. It corresponds to a cosmic ray crossing about 1.7 cm of scintillator.

Chapter 4

Study of beam induced background

4.1 Beam induced background

The CHANTI aims to reject inelastic interactions between the beam and the GTK that could be misidentified as a $K^+ \rightarrow \pi^+ \nu \bar{\nu}$. The most harmful beam interactions are the ones with the GTK station 3, indeed the last GTK station lies close to the decay volume and the GTK could be potentially blind to them, especially when they occurred in the very downstream part of the station. To stress the relevance of this problem it is probably worth to note here that the upper limit set as final result [58] of the experiment E391A at KEK-PS, aiming at the search for the decay $K_L \rightarrow \pi^0 \nu \bar{\nu}$, has been limited by a background induced by inelastic interactions of the beam halo neutrons.

In order to fix the order of magnitude of the problem we have to compare the number of kaons suffering an inelastic interaction on GTK-3 to the number of expected signal events in two years data taking. If one considers roughly the GTK-3 station as a $400 \mu\text{m}$ thick Si target, that is a $0.06\% \lambda_{int}$ (pion interaction length) thick target, one expects that a fraction of order 6×10^{-4} of the kaon flux impinging the GTK-3 will suffer an inelastic scattering. Since about 8×10^{13} kaons will cross the GTK in two years data taking at the nominal NA62 beam intensity, one expects roughly 5×10^{10} inelastic interactions to be produced in the same amount of time, to be compared to 100 signal events detected. The combined effect of the analysis cuts based on the kinematical reconstruction, and of the detector vetoing must reach at least a rejection factor of order (few) 10^{-10} to be able to keep the inelastic background at a reasonable fraction of the total number of observed events. It is thus straightforward to see that any tentative estimate of this background requires a very large Monte Carlo generation. The results presented here refer to a simulation performed using the NA62 official Monte Carlo framework, based on GEANT4 package [59], to estimate this background and the rejection factors obtained by the various NA62

detectors used in veto.

This simulation includes all the detector but the CEDAR since only K mesons are simulated, thus no K tagging is needed. The signal simulation is based on 0.992×10^6 events where the beam kaons are forced to decay in $\pi^+ \nu \bar{\nu}$: the beam kaons have the same lifetime as the real kaons but every decay is forced to be a signal one. The background simulation sample is composed by 1.22082×10^9 events where the beam kaons are not allowed to decay, so only interactions with matter are possible. Due to the need to produce such a huge Monte Carlo sample of the background some simplifications have been adopted to parametrize the response of the detectors, and some of these simplifications (namely the LAV response matrix) have been inserted in the official Monte Carlo of the experiment.

It is clear that also pions in the NA62 beam will in general suffer comparable interactions with the GTK-3 : however, they have not been considered at this stage, since their contribution to the final inelastic background will be highly suppressed by the Kaon tagging performed by the CEDAR. They will however be analyzed in a subsequent work.

4.2 Background rejection cuts

The official NA62 Monte Carlo is still in a development phase, thus not all of the NA62 detectors have a fully working simulation and reconstruction. Furthermore for some of them the reconstruction is not yet optimized to reduce the CPU time and this is of course a concern given the huge amount of events we need to produce. It is for this reason that for some detectors we have made mainly a geometrical study and/or educated assumptions on their response to the particles. We also use some simplifications, where feasible, in order to reduce the simulation time.

In this section we describe the implementation of each detector and the detector veto cuts, some criterion we chose to tag an event as rejected by the detectors. In the same way we will also introduce some analysis veto cut which we consider to reject an event according to some global event property, for example according to the reconstructed missing mass m_{miss}^2 .

STRAW simulation

We fully activate the STRAW tracker in the simulation and reconstruction. The reconstruction provides two algorithms to fit the track, a “Tight” one and a “Loose” one¹. For each fit the χ^2 are calculated and the number of hits associated

¹Both the fits are performed iteratively: they start fitting the hits in the first chambers, then they extrapolate the track position on the next chamber containing an hit. They calculate the distance between the extrapolated hit and the new measured one and if it is small they start a new iteration using also the new hit. The fits differ in what to do in case of bad agreement of the two hits: the “Loose” fit will discard the measured and it will proceed with the hit on the next chamber, while the “Tight” one will use anyway the measured hit but assigning to it a fake variance equal to the distance with the extrapolated point.

Component	Thickness (μm)
Si sensor	200
Si-Pb bump bonds	25
Si readout chips	100
Si cooling plate	60
C_6F_{14} cooling fluid & Si channel wall	70
Glue ($C_{15}H_{44}O_7$) Epoxy	30
G10/FR-4 (PCB) (25% epoxy, 75% SiO_2)	2000

Table 4.1: Material budget of one GTK station. The 2 mm thick PCB frame is outside the beam acceptance.

to the track per each chamber is reported. The track fit reconstructs the charge of the particle, its momentum and direction. This is used in all subsequent reconstruction, e.g. to evaluate the kaon decay vertex position and the missing mass for the event.

RICH simulation

The RICH is fully described and enabled in the simulation but for processing time reasons we do not run its reconstruction, the only information we use being obtained by Monte Carlo truth.

GTK simulation

We fully activate the GTK in the simulation and reconstruction, but, since the reconstruction is in a early stage of development, we only use the information about the reconstructed K momentum (no reconstructed energy deposit). However the geometrical structure of the tracker in the simulation is extremely accurate and includes all the passive material the beam interacts with, including the sensor itself, its bonding, the micro channel cooling, the epoxy glue and the PCB frame. The material budget is reported in table 4.1

Preselection

The “preselection cut” aims to select a sample of events that are good candidates for the $K^+ \rightarrow \pi^+ \nu \bar{\nu}$ decay channel. The ratio of events surviving cuts, reported in the following plots and tables are normalized to the events sample that fulfill these requirements:

- one, and only one, positive track was reconstructed in the STRAW tracker (the “Main Pion” in the following);
- momentum reconstructed for the Main Pion in the STRAW tracker was in the range $15 \text{ GeV}/c < P < 35 \text{ GeV}/c$;
- one, and only one, kaon candidate was reconstructed in the GTK;

- one, and only one, pion track was found in the RICH acceptance.

In figure 4.2 we report the ratio of events on the full sample that survive these selections, grouped by detector. From the plot we can get an estimation of the ratio of inelastic on the full sample, which is about 6×10^{-4} , as expected.

For the background we restricted the analysis to the inelastic interaction happened on GTK station 3, thus, just for these events we add the request

- inelastic interaction occurred on the GTK3 (Monte Carlo truth);

while for the signal only we request

- primary particle must have decayed in the fiducial volume ($105 \text{ m} < z < 165 \text{ m}$ from Be target) .

This request is also useful in order to simplify the comparison of all the selection procedure with the “standard” signal analysis, where this request is always present.

We refer to the selection described above as “Presel” in the plots.

We point out that we excluded the so called “quasi elastic” kaon interactions from the background sample. That is, we do not consider an interaction to be tagged as “Inelastic” if in the final state there is another K, different from the primary one (according to G4), but with a momentum closer than $3\sigma \approx 3\%$ to the original K momentum . In fact, in these cases, the final K is still a good potential candidate for a $K \rightarrow \pi\nu\nu$ decay. In figure 4.1 there is a plot of the difference between the primary K momentum and the secondary one; it show that $\approx 6\%$ of the inelastic interaction are actually quasi-elastic events.

In figure 4.3 we show the resolution on the Z component (i.e. the beam axis component) of the position of the reconstructed interaction vertex for the events in the preselection sample for the signal and the background (with superimposed a Gaussian fit of the core). The plots are the distribution of the quantity $Z_{reco} - Z_{true}$, where Z_{reco} is the Z position of the reconstructed vertex, while Z_{true} is the position of the K decay/interaction obtained by Monte Carlo truth. To quantify the long non-Gaussian tail we report also the fraction of events with $Z_{reco} - Z_{true} > 3m$, which for the background corresponds to the fraction of events that are erroneously reconstructed inside the fiducial volume, starting 3 m downstream GTK-3. The difference between the two plots is purely due to the different geometrical distribution of the true vertexes. In fact, while the signal vertexes are almost uniformly distributed all along the fiducial volume, the background vertexes are concentrated on the GTK-3 surface.

Quality

This is a check on the quality of the reconstruction of the track and vertex determination. The event is kept only if all of the following conditions are met:

- the reconstruction algorithm satisfy some minimal checks²

²Convergence of the STRAW track fits.

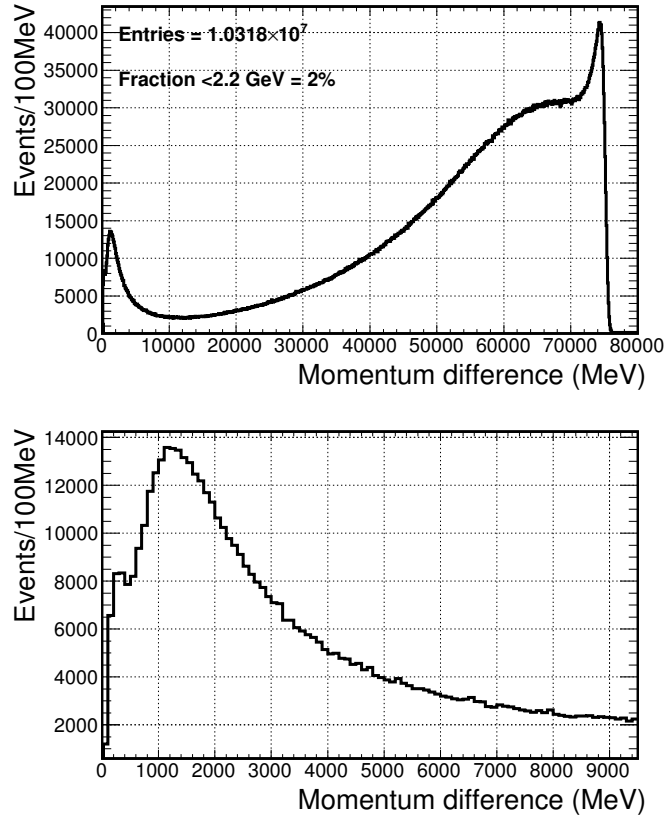


Figure 4.1: Difference between the momentum of the “beam kaon” and the “secondary kaon” produced in the inelastic interactions. The near-zero peak corresponds to the “quasi-elastic” events described in the text.

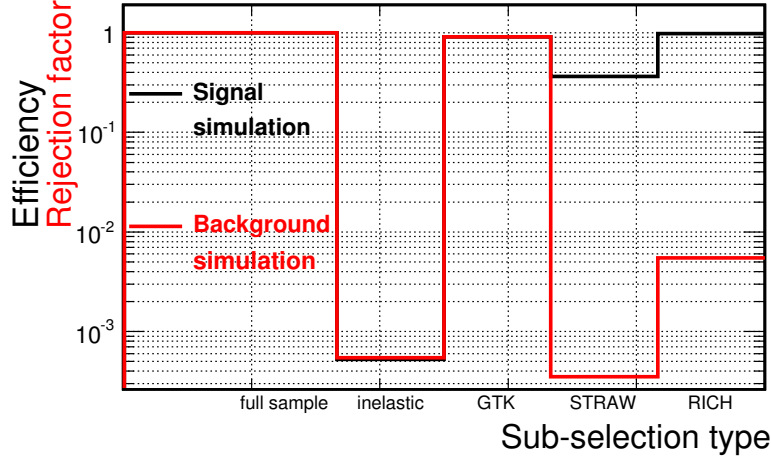


Figure 4.2: Effect of the different detector requests in the definition of the preselection sample. The normalization is to the full sample of 1.22082×10^9 events for the background and 0.992×10^6 events for the signal. It is shown that the fraction of inelastic interactions on GTK-3 is about 6×10^{-4} as expected.

- the χ^2 of the track candidates in the GTK and STRAW is lower than a threshold³;
- the minimum distance between the GTK track and the STRAW tracks does not exceed a threshold⁴;
- the reconstructed vertex is compatible with the beam position⁵.

We need to search for the minimum distance among the tracks because, quite obviously, due to resolution effects the two reconstructed tracks actually do not intersect in a point. If the minimum distance is too large, this could be an hint of a bad reconstruction or of the matching of particles coming from different vertexes.

We refer to this selection as “Qual” in the plots.

In figure 4.4 is shown the resolution on the Z component of the interaction vertex for the events in the preselection sample that pass the Qual cut for the signal and the background. The selection seems to act in the same way on the two distributions reducing both tails by one order of magnitude.

³40 for the GTK and 1.25 for the STRAW

⁴The minimum approaching distance between the two tracks must be less than 4 mm.

⁵Independent cuts are imposed on the X and Y component of the distance in order to follow the beam direction: the former is $0.0012 \times Z_{reco} - 153 < X_{reco} < 0.0012 \times Z_{reco} - 88$, while the latter requests $|Y_{reco}| < 25$ (all units are mm).

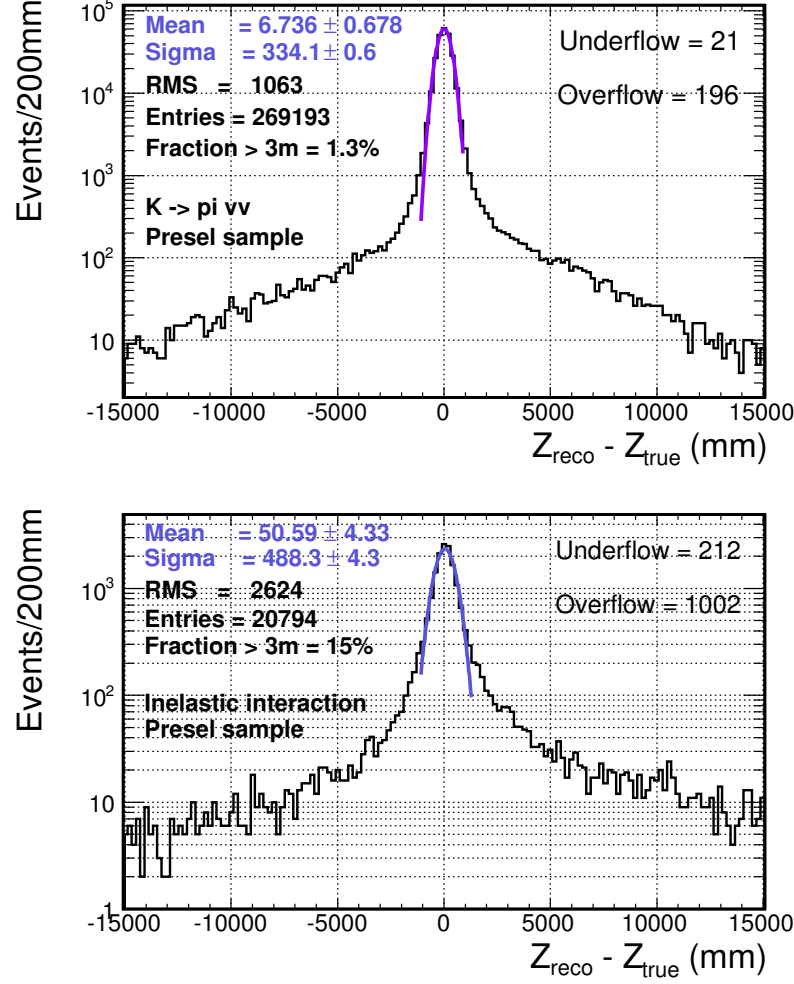


Figure 4.3: Resolution on the Z of reconstructed vertex obtained for the preselection sample for signal (up) and background (down). The distributions have a Gaussian core with $\sigma \approx 40$ cm, but long non Gaussian tails. Since they are computed as $Z_{\text{reco}} - Z_{\text{true}}$ the ratio above 3 m (distance between GTK and the start of decay region) in the background plot represents the fraction of dangerous events. The difference between the two distribution is due to the different geometrical distribution of the true vertexes.

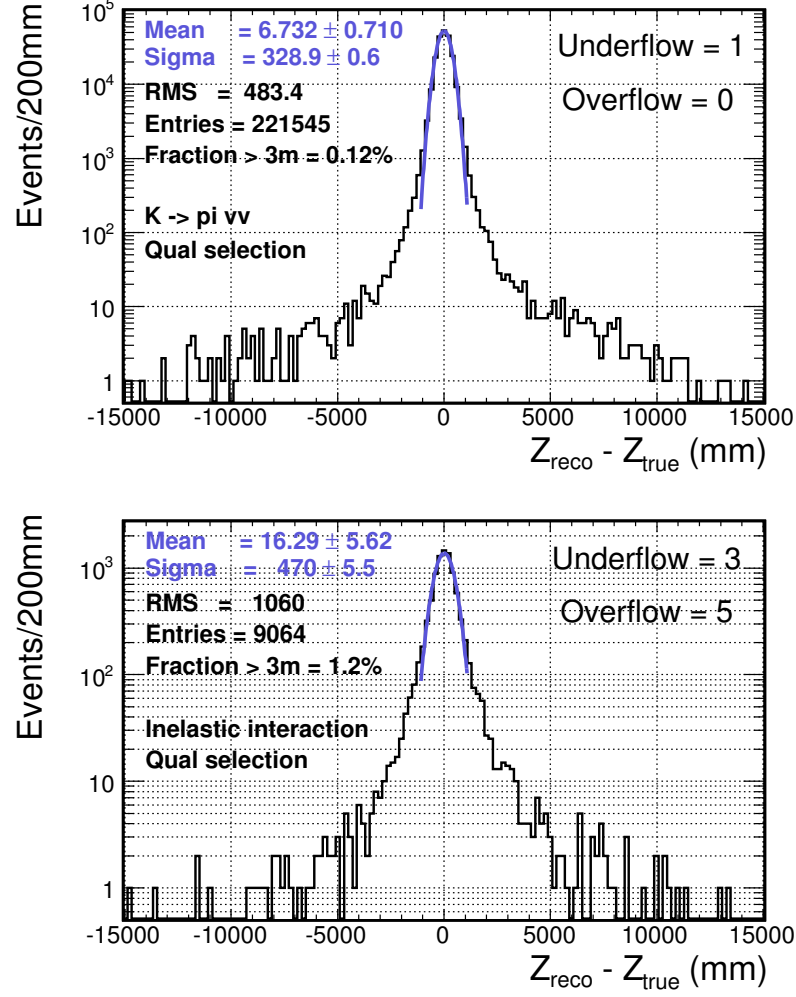


Figure 4.4: Resolution on the Z of reconstructed vertex obtained after Qual selection for signal (up) and background (down).

Kine

As already mentioned in chapter 2 the signal selection is mainly based on the so called missing mass of the event, whose definition we recall here:

$$\begin{aligned} m_{miss}^2 &= (p_K - p_\pi)^\alpha (p_K - p_\pi)_\alpha = \\ &= m_K^2 - m_\pi^2 + 2P_K P_\pi \left[\cos\theta_{K\pi} - \sqrt{\left(1 + \frac{m_K^2}{P_K^2}\right) \left(1 + \frac{m_\pi^2}{P_\pi^2}\right)} \right] \end{aligned}$$

The event is kept only if the missing mass is inside one the regions $0 < m_{miss}^2 < 10^4$ or $2.6 \times 10^4 < m_{miss}^2 < 6.8 \times 10^4$, which are defined to select signal against $K^+ \rightarrow \mu^+ \nu$, $K^+ \rightarrow \pi^+ \pi^0$ and $K^+ \rightarrow \pi^+ \pi^+ \pi^-$ backgrounds, as described in section 2.2.

In addition to this kinematical selection, in order to ensure that all kaon decays other than the signal are properly vetoed by the vetoing system a decay fiducial region is defined for the signal ranging from $105m$ to $165m$ (wrt Be target). Therefore a cut on the position of the reconstructed vertex is done, requiring it to be within the decay fiducial region.

We refer to the selection combining missing mass and fiducial volume cuts as “Kine” in the plots.

In figure 4.5 is shown the resolution on the Z component of the interaction vertex for the events in the preselection sample that pass the Kine cut for the signal and the background. The effect on the signal is to reduce by a factor ≈ 2 the tails, while on the background it simply selects the dangerous events badly reconstructed in the fiducial volume (and with a m_{miss} such as to be compatible with $K^+ \rightarrow \pi^+ \nu \bar{\nu}$).

CHANTI

We developed a complete G4 simulation of the CHANTI materials and geometry, containing scintillating bars, fibers and aluminum frame, and this is now included in the official NA62 Monte Carlo which we used for this study.

The digitization is still under development (see Chapter 5). It will include the full signal shape simulation and the determination of the Time over Threshold for each fired bar, and profiting of the data collected during the November 2012 Technical Run, it will be fine tuned to match with real data. For the aim of this study, however, we simply parametrized the bar response by means of a cut on the energy deposit: a CHANTI bar is considered “Fired” if there is an energy deposit larger than $1/3$ of the energy released, on average, by a MIP crossing orthogonally the bar. This is indeed a conservative estimate, since we have been able to safely operate the detector with a threshold considerably lower than that corresponding to $1/3$ of a MIP during the Technical Run. Moreover, given the anticorrelation of the response of two adjacent bars in one view, a more efficient way to select events could be addressed by summing up the responses (as evaluated by means of their ToT) of couples of nearby bars.

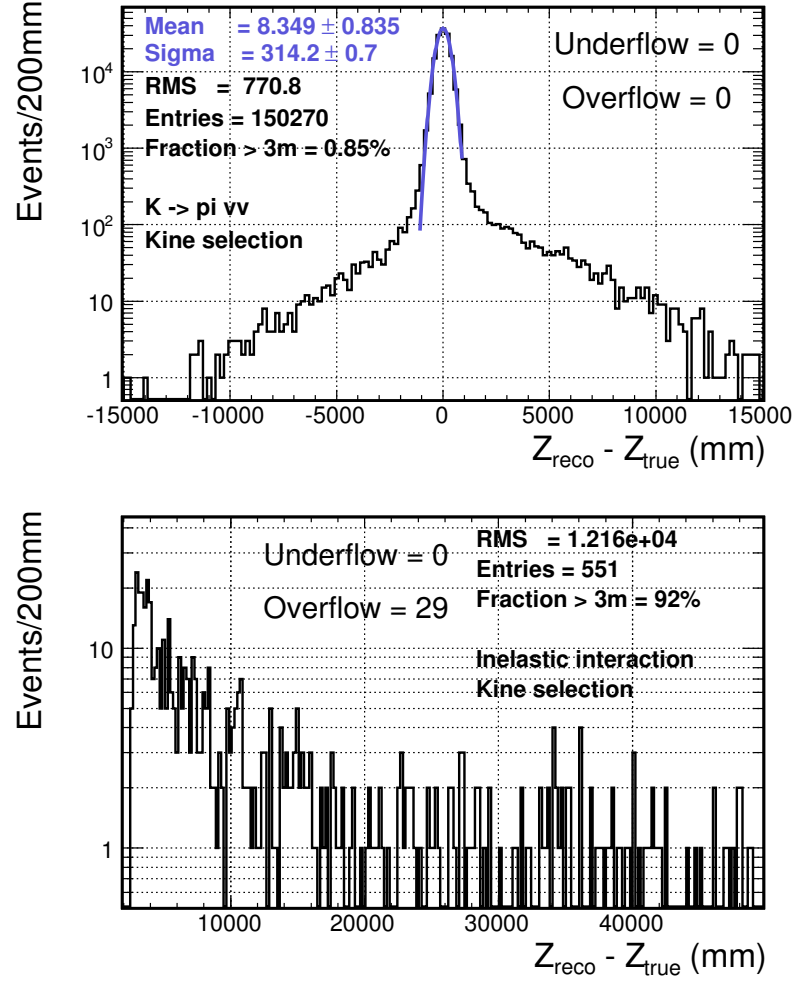


Figure 4.5: Resolution on the Z of reconstructed vertex obtained after Kine selection for signal (up) and background (down).

The number of bars fired for the events surviving preselection is shown in figure 4.6. If we calculate the cumulative of this distribution (normalized to 1) we get the plot in figure 4.7 which allows us to evaluate the “rejection factor” of a cut on the number of fired bars, i.e. the fraction of background surviving the CHANTI veto. This kind of plot allows us to easily show and compare the effect of the cut on both background and signal. In fact it is possible to report on the same plot both cumulative distributions: these two curves actually correspond to $1 - \alpha$ (where α represents type I error) and β (type II error) of a signal vs background hypothesis test using the number of fired bars as test statistics. Of course, requiring more and more bars to be fired to veto the event will increase the vetoing inefficiency, and, as seen in 4.7 the rejection factor will eventually approach 1 (no rejection at all) when increasing this number to very high (>70) values.

In figure 4.8 we report a zoom of the curve in the region relevant to the cut, namely the one of few fired bars: we chose a threshold of at least two CHANTI bars fired to reject the event and this request will be called “CHANTI cut” in the following plots. With this cut the CHANTI rejection factor on the background preselection sample is about 1.4%. In figure 4.9 we show that a big part of this vetoing inefficiency is due to an acceptance effect: 54% of the events surviving the cut did not release any energy in any of the six stations (i.e. have 0 Monte Carlo true track hits in the CHANTI). If we were less conservative, and asked for a single bar⁶ in the CHANTI to reject an event, 86% of residual vetoing inefficiency would have been geometrical.

As we will show in the following sections, inelastic events surviving the CHANTI veto may be recovered by suitable cuts on the downstream detectors, with a key role played by the LAV system, which covers the angular region below 49 mrad with respect to GTK-3. It is worth to be noted here that while apparently the cut on the number of fired bars has a negligible impact on the signal, the main mechanism through which this cut will reduce the signal efficiency is the accidental coincidences of the CHANTI activity (which will be dominated by the muon halo) in the time window around a given trigger. It is for this reason that the CHANTI time resolution plays a very important role in the detector design. The time resolution of the CHANTI on real data collected during the November 2012 Technical Run will be discussed in the next Chapter.

In the following we call “CHANTI selection” the request of having less than two fired bars in the CHANTI.

GTK

The GTK plays a double role in the inelastic background reduction. As already shown, it is used in the kinematical reconstruction of the event; however we can also use its response to reject the background events. In fact, since the sensors of the GTK are placed in the downstream part of the detector, one expects that an

⁶As already mentioned, to improve the CHANTI veto performance, we could exploit the anti-correlation in adjacent bars response due to their triangular shape. A lower CHANTI bars threshold give us an idea of the performance we can reach in that condition.

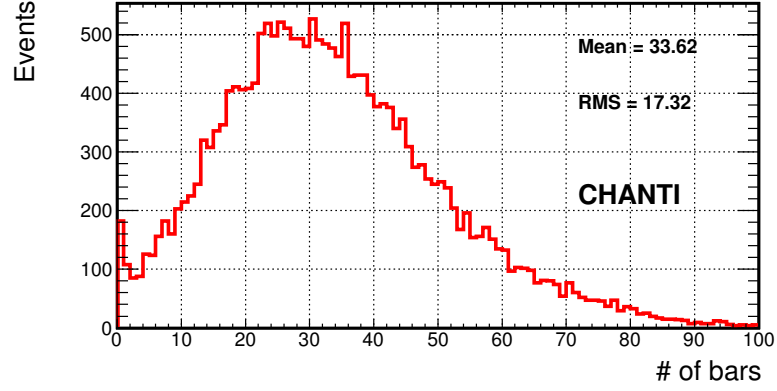


Figure 4.6: Number of CHANTI bars fired for the inelastic background events in the preselection sample. We consider a bar “fired” if there is an energy deposit into it larger than $1/3$ of the MIP one.

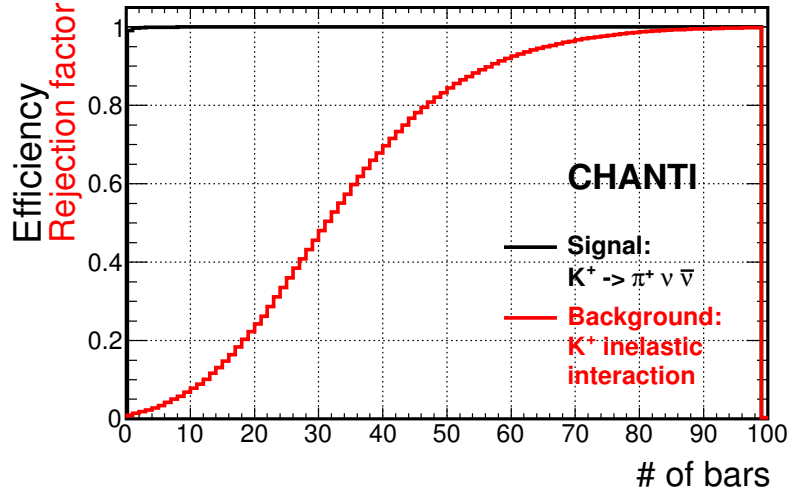


Figure 4.7: Cumulative distribution of the number of fired bars. The plot can be read as follows: setting a threshold of 10 CHANTI bars or more to tag an event as vetoed by the CHANTI, we get a background reduced to 5% of the original sample and we select almost 100% of the signal.

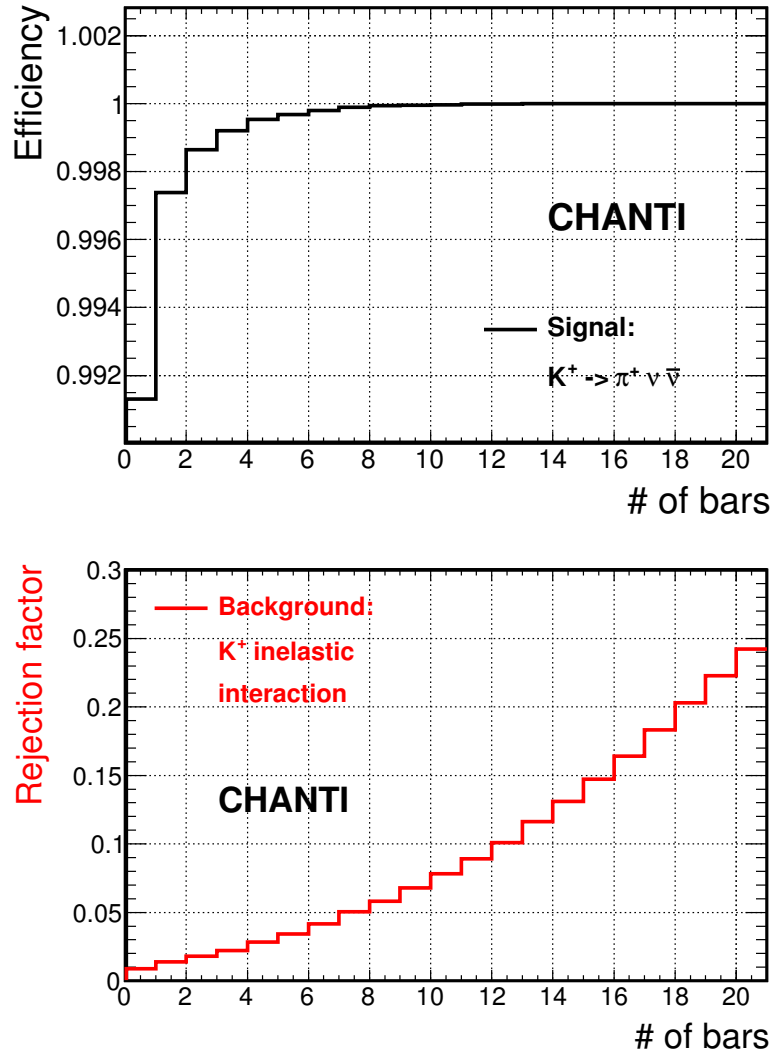


Figure 4.8: Zoom of the rejection factor and of the signal efficiency as function of a cut on the minimum number of CHANTI bars fired.

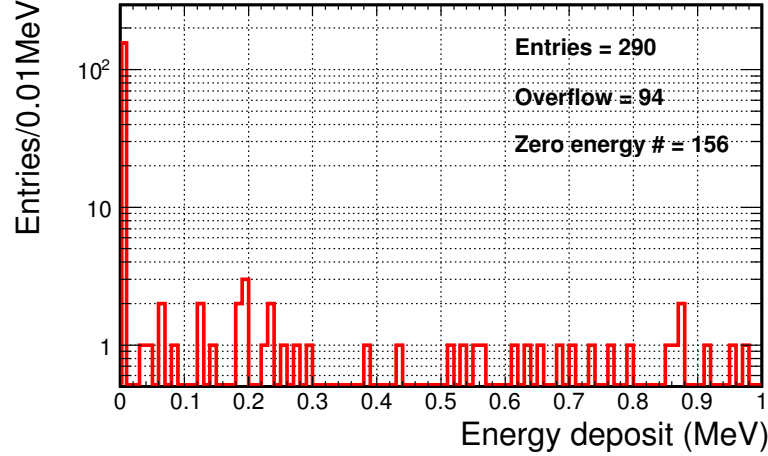


Figure 4.9: Energy deposit in the CHANTI by Monte Carlo truth of the preselection sample not vetoed by the CHANTI cut. The zero energy bin represent events with no charged particle hitting the detector.

inelastic interaction which happened in its upstream part may be observed by the GTK itself, and should be characterized on average by a higher multiplicity of fired hits and a larger energy release given the higher multiplicity of charged particles crossing the detector with respect to the non-inelastic case.

The usual cumulative distributions for the reconstructed number of pixels fired (figure 4.10) and for the (Monte Carlo true) energy release (figure 4.11) show that signal and background can be quite easily separated using these two variables.

We call “GTK selection” the request of less than 2 pixels being fired in the GTK and an energy deposit in it lower than 0.4 MeV .

LAV

The LAV system has been described in some detail in Chapter 2. To the aim of the present study the most important fact to keep in mind is that the basic building block of the LAV is a lead glass detector where the light production mechanism is due to the Cherenkov effect. For this reason the response of a single block to high energy photons and to MIPs is quite different even for the same amount of energy released inside the block. It turns out that for a realistic estimate of the LAV (in)efficiency one has to properly take care of the different light production mechanisms inside each block. For this reason we developed a detailed simulation of all the optical properties of a LAV block, including the lead glass absorption length, the diffusive Tyvek wrapping, the light guide, the optical glue and the photocathode window, and we have also parametrized the average quantum efficiency of the photocathode as a function of the optical

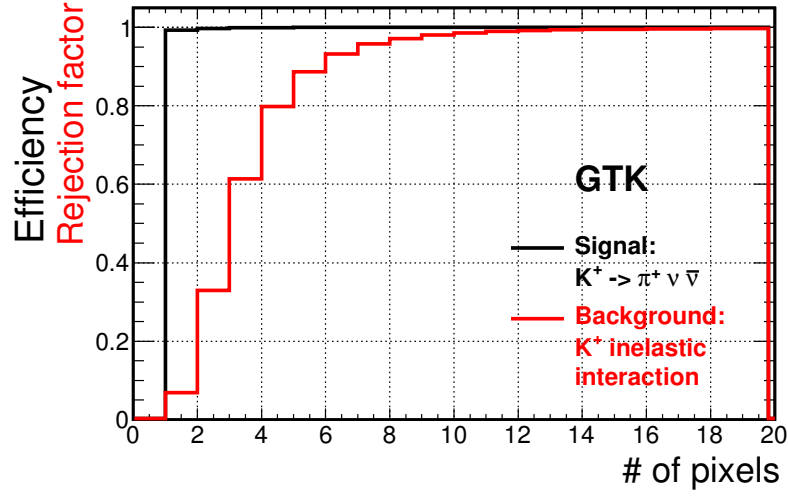


Figure 4.10: Rejection factor (and signal efficiency) as a function of a cut on the maximum number of GTK-3 pixels fired. The plot refers to the events surviving the preselection cut.

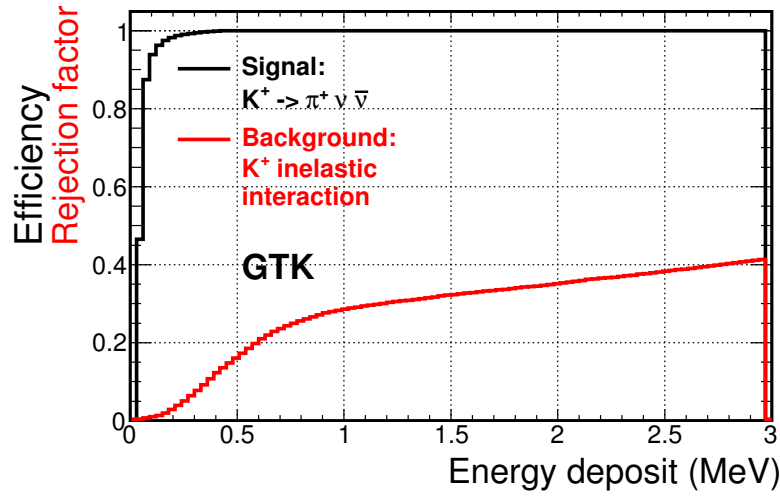


Figure 4.11: Rejection factor (and signal efficiency) as a function of a cut on the maximum energy deposit in the GTK-3. The plot refers to the events surviving the preselection cut.

photon wavelength, as given by Hamamatsu data sheet.

This simulation has been able to reproduce all of the relevant features of the LAV block response to both MIPs and electrons, and has been compared to real data taken in a cosmic ray test stand in Naples and at the Frascati Beam Test Facility [60]. Despite the fact that this accurate simulation has been inserted in the official NA62 Monte Carlo, it is clear that such a detailed description, including the tracking of every optical photon inside each LAV block fired would be too much time consuming and is therefore not an option for the present study.

However, in parallel with the accurate optical simulation we have also implemented a fast parametrization which improves the CPU time per event by a factor of 50 while having comparable accuracy. This is obtained by sampling the probability for an optical photon to reach the photocathode as a function of all possible variables (namely photon initial position, initial direction, and energy) in a six-dimensional space by means of an appropriate six-dimensional matrix with adaptive bin size.

During the simulation all the optical photons are generated according to the G4 Cherenkov simulation, but they are not tracked: instead their initial position, direction and energy is used to assign them a probability to reach the photocathode and (using a separate matrix) the mean arrival time. A simple random extraction gives in the end the number of optical photons collected and their arrival time at the photocathode, which can in turn be used to reproduce the signal shape and the Time over Threshold. In this analysis we did not use the full digitization of the signal but only the number of optical photons collected, and we compared the number of optical photons collected in a block with the number of (optical) photons collected when a MIP crosses orthogonally the same block.

In figure 4.12 there is the rejection factor for a LAV block response (i.e. number of optical photons) normalized to the MIP response.

The plots refer to the LAV block which have the third highest response, so imposing a minimum threshold on this variable (to reject an event) corresponds to ask at least three blocks having a response as high (or higher). It is interesting to note that this selection introduces about 15% inefficiency on the signal, which may seem not intuitive. This inefficiency is mainly due to the production of secondary δ rays generated inside the RICH material which reach the last LAV station (LAV12). This can be easily demonstrated by removing the RICH from the simulation: in this way the amount of signal events which are lost becomes negligible. This fact could suggest to change the selection strategy treating LAV12 on a different foot with respect to the other LAV stations, but since a cut on LAV12 multiplicity is anyway necessary in order to reduce the $K \rightarrow \pi^+ \pi^+ \pi^-$ background, this has not been done.

The “LAV selection” excludes events with three or more LAV blocks above the threshold equal to 1/3 of the MIP response .

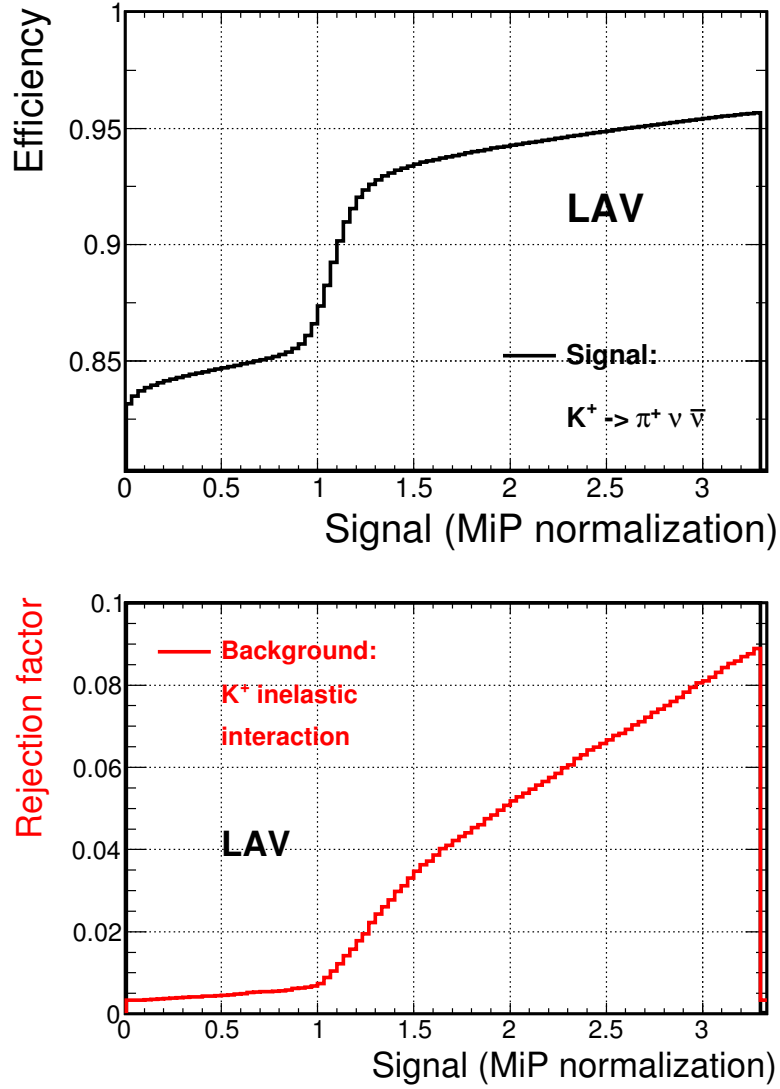


Figure 4.12: Rejection factor (and signal efficiency) as a function of a cut on the minimum LAV response. The response is normalized to the one obtained in the case of a MiP and refers to the LAV block with the third highest response: asking this block to have a response above a certain threshold implies two more blocks exceeding the same threshold. This means we ask to have at least 3 blocks fired in the LAV system to veto the event.

Downstream detectors: LKr , MUV, CHOD, IRC and SAC

The detectors downstream LAV12 cover a very small fraction of the solid angle with respect to the GTK-3, and, thus, their contribution to the overall rejection factor is expected to be relatively small. Since the simulation of the interaction with the LKr is extremely time consuming in the version of the Monte Carlo used for this work, it has been excluded from the simulation. It would have been of no real physical meaning thus to include in the simulation the detectors which are further downstream like the MUV or the SAC. For all of these “downstream detectors” the strategy we have used is then to try to simulate their response by a geometrical extrapolation (including of course the kick of the magnetic field) of the particles present in the event and an educated guess of the response of the detectors to these particles. Although this method is clearly a bit crude, it is believed to be sufficient at this stage of the analysis to assess the contribution of these detectors to the overall rejection, which remains anyhow quite small.

The procedure used for the LKr is as follows: first of all we track the “Main Pion” downstream to the LKr surface: this is done to be able to define an “isolation region” i.e. a circle of radius $R = 2 R_M$ (where $R_M = 6.1$ cm is the Molière radius for the LKr) around the extrapolated pion track where we look for further activity of the LKr which may be used to veto the event. Then we compute the maximum visible energy in the LKr for the event, which is either the energy of the most energetic photon, or the one of the most energetic pion (corrected for a scale factor typical of pion showers) in the event hitting the LKr outside the isolation region.

In figure 4.13 the rejection factor of a cut on the maximum visible energy described above is shown. The rejection power of this cut is not particularly high, since the solid angle under which the LKr is seen from the GTK-3 is obviously quite small: as already mentioned this feature is common to all the downstream detectors.

In the following, an event is considered to be rejected by the “LKr selection” when the variable defined above is larger than 500 MeV.

The effect of punch-through of pions from the LKr is completely neglected in this analysis, since the LKr simulation is missing. Thus the MUV is used only to veto events where muons are produced. In figure 4.14 we report the rejection factor for a threshold on the energy of the most energetic muon hitting the MUV. It is evident that, since the muons hitting the MUV are essentially produced in the decay of the “Main Pion” both for the signal and for the background, the rejection factor for the background turns out to be almost the same as the signal efficiency, so that there is not much to be gained on the ratio of the signal to the inelastic background from this cut. This selection is anyhow considered since it is used to suppress the $K^+ \rightarrow \mu^+ \nu$ background.

The “MUV selection” rejects an event when the most energetic muon in the detector acceptance has an energy larger than 200 MeV.

The CHOD multiplicity could be of some help in suppressing the inelastic background. Indeed one can easily see that any charged particle hitting the CHOD, other than the “Main Pion” has to be produced either in a kaon decay

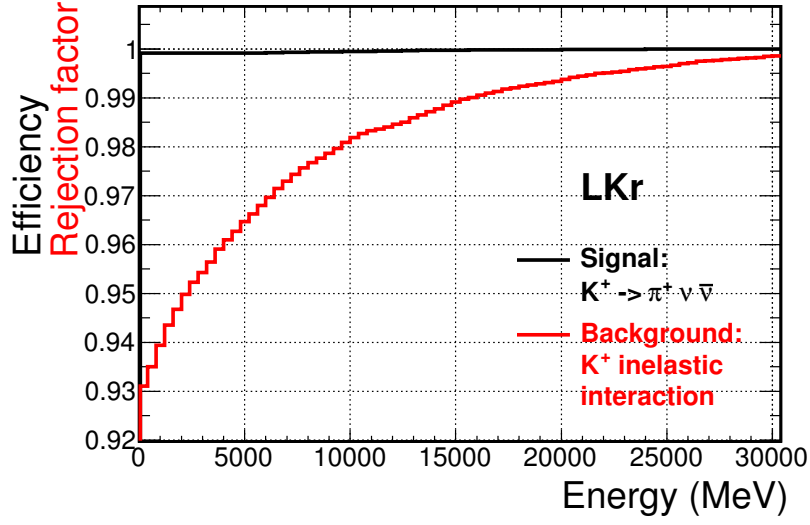


Figure 4.13: Rejection factor (and signal efficiency) as a function of a cut on the maximum visible energy hitting the LKr in a region separated from the Main Pion.

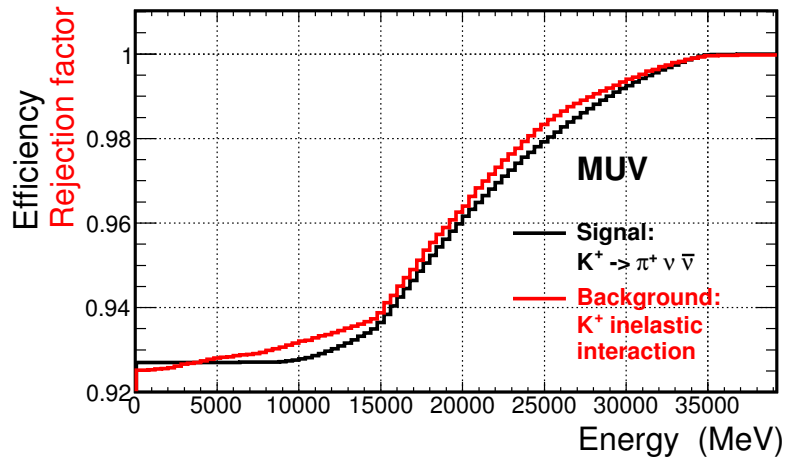


Figure 4.14: Rejection factor (and signal efficiency) as a function of a cut on the maximum muon energy hitting the MUV.

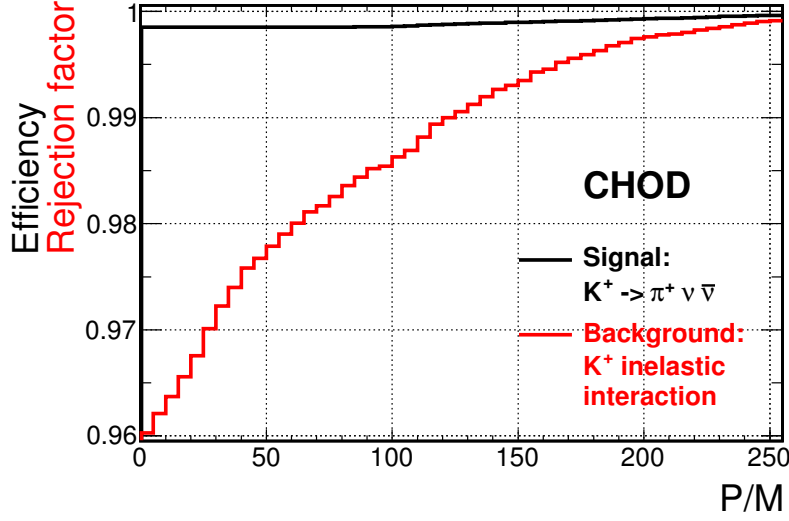


Figure 4.15: Rejection factor (and signal efficiency) as a function of a cut on the maximum momentum-to-mass ratio of a charged particle hitting the CHOD in a region separated from the Main Pion.

different from the signal one, or in an inelastic interaction.

In figure 4.15 we report the rejection factor for a cut on the ratio momentum to mass ratio (P/M) of the charged particle that hits the CHOD having the highest value of P/M in the event. Of course, since the CHOD will be used to trigger on the “Main Pion” an isolation region around the “Main Pion” is needed when counting additional particles to reject the event. We do not consider the charged particles with an entry point on the CHOD being far from the entry point of the “Main Pion” less than the pad diagonal (8.5 cm).

The “CHOD selection” rejects events where the value of the variable P/M defined above is above 2, since we assume the CHOD to be nearly 100% efficient for MIPs.

The IRC and SAC cover a very small acceptance region with respect to GTK-3. In figure 4.16 there is the rejection factor for a cut on the energy of the most energetic γ hitting one of these detectors. The “IRC&SAC selection” will reject events where a γ of at least 100 MeV hits one of these detectors.

4.3 Background estimation

A summary of the effect of each of the individual selection cuts described above is shown in figure 4.17: the normalization is to the “Preselection cut”, and the number of event that pass this cut is 2.69193×10^5 (corresponding to a selection efficiency of about 27%) for the signal and 2.0794×10^4 (corresponding

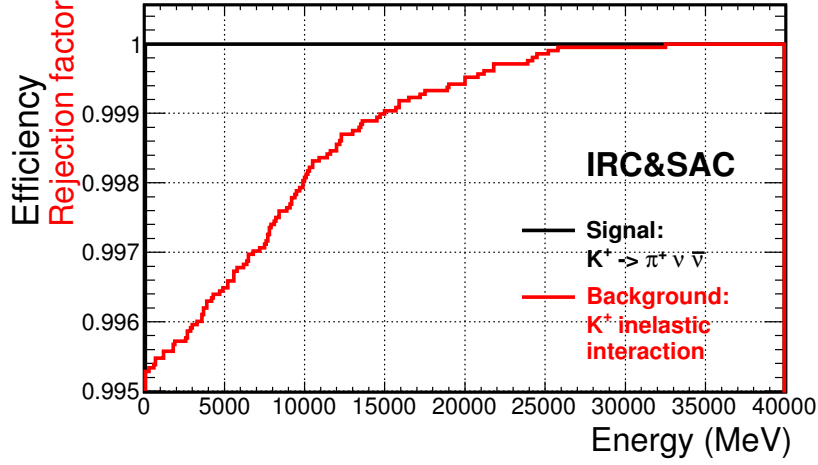


Figure 4.16: Rejection factor (and signal efficiency) as a function of a cut on the maximum energy of a γ hitting the IRC or the SAC.

to a rejection factor of about 1.8×10^{-5}) for the background. The “OtherDet” selection refers to the AND of all the selection cuts for the downstream detectors, namely LKr, MUV, CHOD, IRC and SAC.

For the signal, we get 783576 decays in the fiducial volume by Monte Carlo truth; out of these 103530 pass all the selection cuts, giving $\approx 13\%$ acceptance on the signal, well in line with the results obtained in other independent analysis. Since in NA62 we expect $K_{fv} \approx 9 \times 10^{12}$ decays in fiducial volume [35] in two years of data taking, and considering a (Standard Model inspired) branching ratio for $K^+ \rightarrow \pi^+ \nu \bar{\nu}$ of 8×10^{-11} , we expect to collect ≈ 94 signal events.

For the background, we have not enough statistics to estimate how many events pass all the cuts, so the only choice we have is to make some factorization assumptions. In order to assess whether a factorization assumption for two selection cuts holds or not we consider all possible combinations of two cuts C_1 and C_2 . We perform the selection putting them together in AND C_{12} ; we can then calculate the rejection factors ϵ_1 , ϵ_2 and ϵ_{12} . If $\epsilon_{12} \cong \epsilon_1 \times \epsilon_2$ we have indication that the factorization assumption for these cuts is reasonable.

We expect that, the more two selections rely on independent physical quantities the more likely they will factorize, and thus the more ϵ_{12} will be close to $\epsilon_1 \times \epsilon_2$. In figure 4.20 there is the value of $\epsilon_{12}/(\epsilon_1 \times \epsilon_2)$ for all the possible combinations of cuts; values much larger than one are hints of highly correlated selections, which tend to reject the same events, values significantly lower than one are hints of selections which acts almost orthogonally thus increasing significantly the overall rejection factor.

It is evident that GTK and CHANTI selections are quite correlated, as they are expected to be, since both selections rely on the multiplicity of the secondary

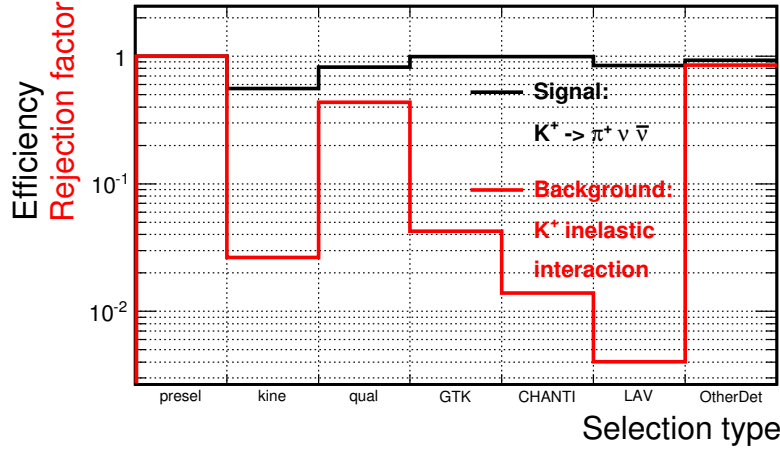


Figure 4.17: Summary of the effects of the cuts described in the text. The “rejection factor” is the fraction of events that pass the cut. For example the CHANTI selection alone let $\approx 1.4\%$ of the background survive and does not affect the signal. The values are normalized to the “Presel” cut that, as described in the text, represent an initial selection of $K^+ \rightarrow \pi^+ \nu \bar{\nu}$ candidates.

particles generated in the inelastic interaction of the K.

On the contrary the Kine and Quality selections are almost orthogonal, and this is also well understood since the fiducial volume request (in Kine) is at odd with the “good reconstruction” required by the Quality selection, as we have already noticed that for inelastic events the vertex may be found inside the fiducial volume only if it is reconstructed far away from its true position. This can be clearly seen also by comparing figs 4.4 and 4.5.

The correlation between the LAV cut and the CHANTI cut could not be assessed with the current statistics since no event passes both the selections, however we get an hint in figure 4.18 where there is a plot of the number of CHANTI bars fired for events that survive the LAV veto. When compared with the one in figure 4.6, it shows that the distribution of the LAV-passing events is slightly shifted toward the high CHANTI multiplicity region (the mean is about 3σ above). This suggests that the factorization hypothesis could be indeed conservative. This behavior is probably due to the fact that the detectors cover complementary angular regions with respect to the GTK-3, as happens also for the CHANTI (or the LAV) with respect to the downstream detectors, where the factorization assumption appears verified.

Now since both CHANTI and GTK selections seem to reasonably factorize with respect to both Quality and Kine; and assuming the CHANTI and LAV factorization, we decide to estimate the overall rejection by evaluating the product of three rejection inefficiencies, namely:

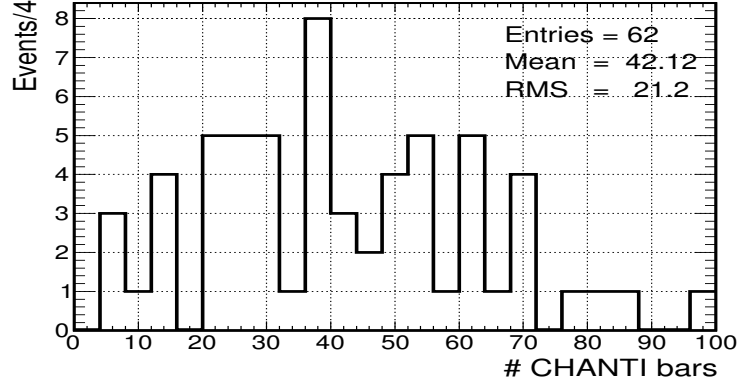


Figure 4.18: Number of CHANTI bars fired (energy deposit $> 1/3$ MIP) for the events that pass the LAV selection.

- GTK+CHANTI rejection factor: $\sim 0.47\%$
- LAV+OtherDet rejection factor: $\sim 0.29\%$
- QUAL+KINE rejection factor: $\sim 0.16\%$

So, multiplying all these factors together, with the factorization assumption we made, the final rejection factor estimate is $\epsilon = 2 \times 10^{-8}$. Putting this together with the Preselection, which adds a suppression factor of $1.8 \times 10^{-5}/6 \times 10^{-4} = 0.03$ for the inelastic events, we reach an overall suppression of the inelastic events of about 6×10^{-10} , which is of the order of magnitude of what is needed.

We plot in figure 4.19 the reconstructed Z for the background event that pass both the QUAL and the KINE selection (and obviously the PRESEL selection).

Although the statistics is still quite low, it is clear that some gain in S/B ratio can be obtained by means of an optimization of the choice of the start of the fiducial decay region in order to reduce the inelastic background. For instance if we move the beginning of the fiducial region 4 m downstream we reduce the background to 40% of its value (but of course in this way we also loose about 7% of the signal that is uniformly distributed in the decay region): this kind of optimization will be performed when more Monte Carlo statistics will be available and is therefore not yet assumed here.

To estimate the background using the previous shown results, we note that in 2 years of data taking, the number of kaons hitting GTK 3 is

$$F_{gtk} = \frac{\exp(-Z_{gtk}/L)}{\exp(-Z_{start}/L) - \exp(-Z_{end}/L)} \times K_{fv} \equiv G \times K_{fv}$$

where $G \approx 9.4$ since $L=533$ m is the attenuation length of a 75 GeV kaon beam, $Z_{gtk} = 102m$ is the position of the GTK-3 and $Z_{start} = 105m$ and

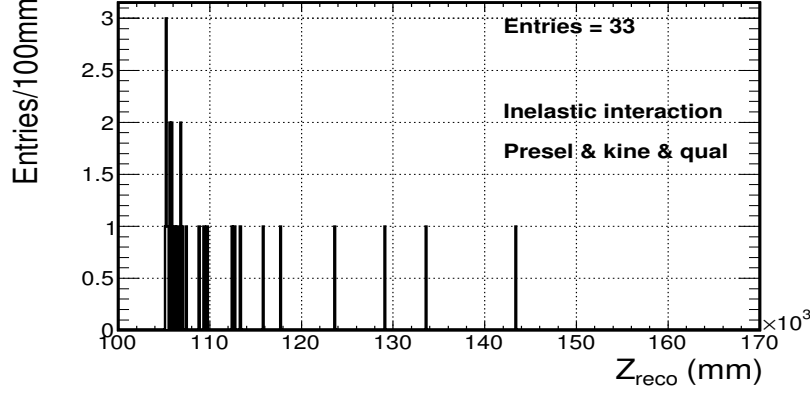


Figure 4.19: Reconstructed Z position of the decay vertex for the background events that pass the Presel, Kine and Qual selections. About 60% of the entries has $Z_{reco} < 109m$.

$Z_{end} = 165m$ are the start and end position of the nominal fiducial volume. The simulation refers to these since, as we said, we turned off all the beam K decays, so every Kaon generated reaches the GTK. From the simulation we know that only $R_i = 1.8 \times 10^{-5}$ of this do inelastic interaction and pass the preselection cut. So our estimation of the number of bad reconstructed inelastic interaction is $B = G \times K_{fv} \times R_i \times \epsilon \approx 30$ which correspond to $\approx 32\%$ of the signal.

We note here that this is a relatively simple analysis scheme based on a rough cut and count mechanism, where none of the cuts has been statistically optimized. Apart from the aforementioned optimization of the decay volume, which could easily improve by a factor of 2 the S/B ratio, one could of course improve the selection by means of multivariate analysis, introducing a Likelihood ratio estimator or a neural network classifier: This approach is however left to a later and more mature stage of the simulations and of the understanding of the detector performances, and was not the main aim of this study, which was intended to state the adequateness (or not) of the NA62 apparatus (including the appositely designed CHANTI detector), to suppress the inelastic background.

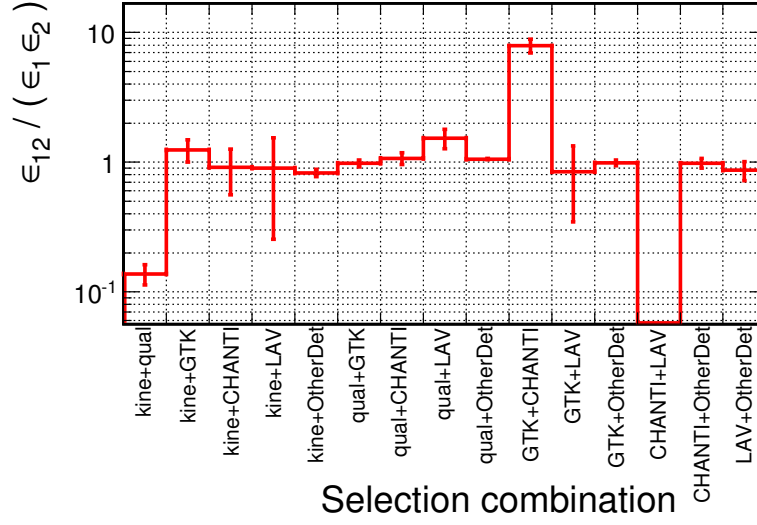


Figure 4.20: Correlation of the different selections. For each combination we report the ratio between the rejection factor imposing both cuts (ϵ_{12}), and the product of the rejection factors considering them separately (ϵ_1 and ϵ_2). When the selections factorize the ratio should be close to 1.

Chapter 5

Technical run

5.1 Hardware setup

In November 2012 there was a testbeam on the K12 beam line at CERN (in the NA62 experimental area) with the aim to test the hardware setup of several detectors and the common acquisition chain. Most of the NA62 detectors took part to this test even if some of them with a limited functionality and/or with prototype setup: CEDAR, CHANTI, LAV, STRAW, CHOD, LKr, MUV, and SAC. During the testbeam two kinds of beam have been provided, one with muon halo only, and one generated as the final NA62 beam, with a 6% kaons component. A very short summary of the installation and operation of the various detector in the Technical Run is shown in Table 5.1.

The CHANTI was present in the Technical Run (TR) with a single station, fully equipped (SiPMs, frame, cables, etc), outside the vacuum tube but in a dedicated (vacuum capable) chamber, placed as close as possible to the beam axis (see fig 5.1) in order to detect beam halo muons. Each SiPM was connected by a coaxial cable to one of three SUB-D 37 feedthroughs on a flange. The external part of the feed through was connected to one readout prototype board (ChantiFE, see section 3.6) which has to power up and read the SiPMs.

This setup reproduces the cabling the CHANTI will have in the final setup. The main difference with respect to the final experimental setup is that the ChantiFE board was missing the voltage regulators which allow it to be powered in a standard crate, and was thus powered by means of linear test-bench power supplies and placed on a deck instead of being placed inside its crate. This of course may have changed somehow the level of electronic noise with respect to what will be in the final configuration, however, given the quite large amount of light collected by the bars of the CHANTI, which allows to operate the detector at relatively high levels of signal threshold, this effect should not significantly affect the results shown here.

Since only one prototype board was available for the TR, it was possible to power-up and readout simultaneously only 32 out of the 46 channels of the

Sub - System	Installation	Operation in Technical Run
CEDAR /KTAG	Installed. 4/8 sectors equipped to 50% with PMTs	All sectors operational and readout. Time synchronization, pressure scan and detector alignment successful
GTK	Not installed	Si wafer placed in the beam to emulate the GTK effecton rates
CHANTI	First CHANTI station positioned outside the beam used with and without vacuum	Readout of all 46 channels successful. (Up to 32 channels simultaneously)
STRAW	Chamber 1 installed with 2 modules (u/v) fully equipped	Readout limited to 8 frontend boards
LAV	Eight LAV stations installed. 3 fully connected to HV, LV and readout	LAV A1, A2 A3 included in common readout
RICH	Not installed	
CHOD	Reconditioned and equipped with LAV ToT FE boards. Small prototypes for new CHOD tested	CHOD readout and used for the trigger
IRC	Not installed	
LKr	Fully installed	Readout with CPD/SLM system (limited to 40% of channels due to availability of FASTBUS power supplies)
MUV	MUV2 and MUV3 fully installed	Full readout and used in the trigger
SAC	Final detector installed; provisional read out	TDC readout based on LAV ToT boards. Standalone test with a commercial CAEN 1GHz FADC
Vacuum system	85m (out of 115m in total) vacuum tank and corresponding detector elements installed. Primary pumping system and 3/7 cryo pumps installed	Leak tightness at desired level. Vacuum reached 4×10^{-6} mbar with two cryopumps

Table 5.1: Summary of detector installation and operation in November 2012 Technical Run

CHANTI station. However all of the channels have been read, since we changed the configuration during the testbeam in order to inspect all the bars of the station. The data presented in the following refers to the configuration with the maximum number of channels read, as shown in figure 5.3.

The output of the prototype board was sent to the a Time Over Threshold (ToT) board which generates one LVDS signal in correspondence to the time when the input exceeds a threshold (Leading signal) and switches the logical level at the time when it returns under that threshold (Trailing signal). For each channel in input it can set two different threshold levels and generates two output signal corresponding to these thresholds, so the number of physical channels is doubled by the ToT board. The information which can be reconstructed from this data is the starting time of the signal and its width, which is related logarithmically to the charge. Referring to figure 5.4, a first approximation of these quantities are the Leading and the difference between Trailing and Leading called “Time over Threshold” (ToT). The signals generated by the ToT board were processed by a standard NA62 acquisition chain (see section 2.14) using two HPTDC on one TDC daughter board placed inside a TEL62 board.

The ChantiFE board and the ToT board are designed to be configured trough CAN bus since the CERN has a central DCS system that permits us to control crates and electronics with this standard from the control room, however for the ChantiFE board prototype the firmware to use this infrastructure was not ready yet. We have thus designed and operated a temporary standalone slow control system. The ToT board was connected to a notebook that configured the board by a serial communication over USB with the help of simple scripts. The same PC was linked (serial communication over USB too) to the ChantiFE and controlled the board by a dedicated LabView software (figure 5.5) by which it was possible easily to setup the SiPM bias voltage and to monitor the absorbed currents channel by channel. The ChantiFE board uses a system of DACs and ADCs to convert respectively a digital input to a bias voltage and, vice-versa, an analog readout voltage or current into a digital output. These DACs and ADCs were calibrated before the TR by means of a voltage meter and a picoamperometer, for each of the 32 channels of the board. A satisfactory linear behavior was found for all channels; the (linear) calibration constants were inserted in the LabView slow control program so that all of the settings and readings were directly expressed in mV and nA respectively. During the last days of the TR the CHANTI chamber has been evacuated and the detector has been also successfully operated in vacuum.

5.2 Threshold calibration

Before the beam was on, we performed inside the experimental hall a threshold study using a LVDS-NIM converter and a NIM scaler, measuring the rate of signals passing a certain threshold. We put the SiPM at two voltage: their operational bias voltage, (as suggested by Hamamatsu for each device, and corrected for the difference between the 25 °C nominal value and the actual

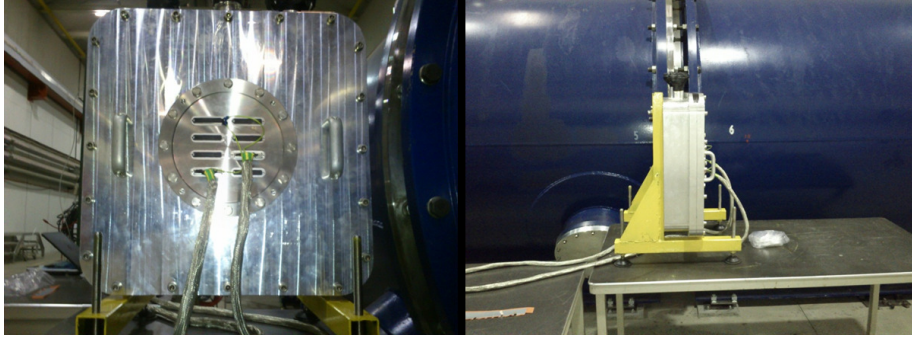


Figure 5.1: The CHANTI setup for the TR. Left: front view of the box hosting the station with the outer cabling. Right: side view of the setup. The detector is placed few meters upstream LAV station ANTI-A1.

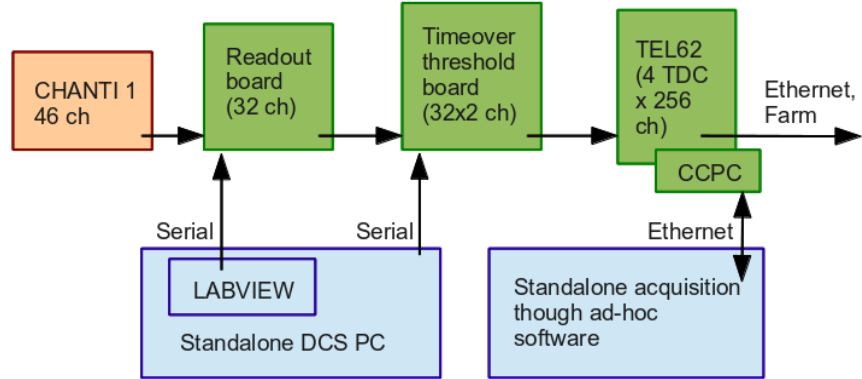


Figure 5.2: Schematic view of the CHANTI data acquisition and slow control system setup for the TR.

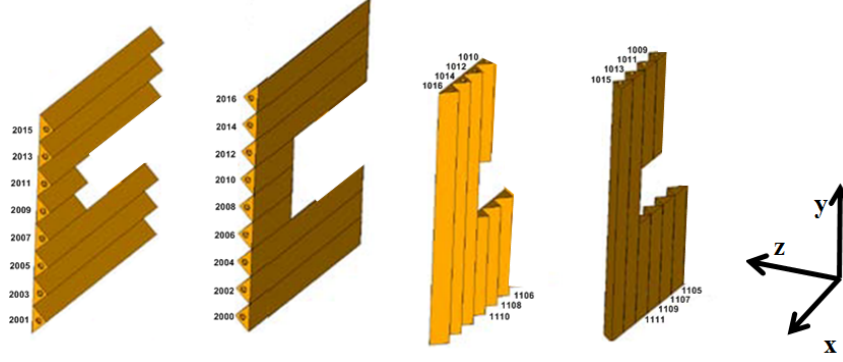


Figure 5.3: Geometrical distribution of the channels readout in one of the configurations used for the testbeam. All the following data refers to this configuration.

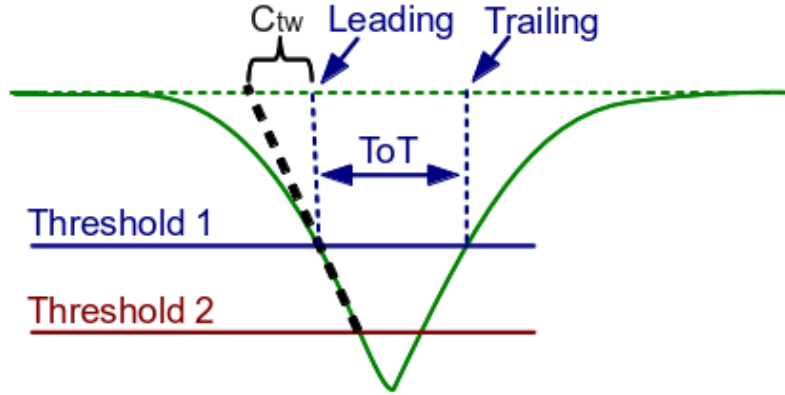


Figure 5.4: Definition of the Leading time, Trailing time and ToT of a signal, for a threshold “Threshold 1”; we could associate to “Threshold 2” the same kind of information. C_{tw} is the “Time walk correction”, a linear correction on the starting point of the signal (wrt the Leading at Threshold 1) that we can easily calculate when the signal passes both thresholds.

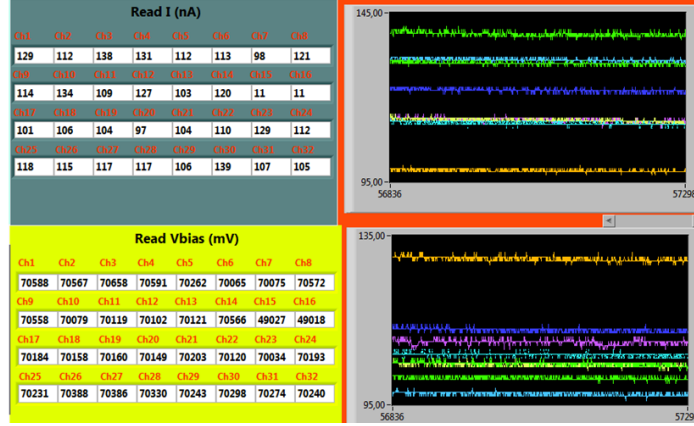


Figure 5.5: Screenshot of the LabView Software controlling the prototype Chan-tiFE board.

temperature in the cavern); and 49 V at which the SiPM are off since are well below the breakdown voltage. The result for two SiPMs and for different threshold value are shown in figure 5.6. The counts with the SiPM off gives us an idea of the electronic noise, which is completely negligible with respect to the intrinsic SiPM dark rate already at 30-40 mV threshold, while the ones with SiPM on show the typical exponential behavior of the dark rate of a SiPM with respect to the threshold set.

We decide to set the low threshold to 80 mV for most of the acquisition time (though dedicated runs with both lower and higher threshold were acquired, too) since this value should correspond to about 3-5 photoelectrons. In fact it is known from Hamamatsu (and lab tests) that at operational bias the dark rate of the device used is about 700 kHz at 0.5 pe threshold. Since the rate scales down roughly by one order of magnitude per photoelectron (e.g. is 70 kHz at 1.5 pe threshold) and since at 80 mV we observe (depending on the channel) a rate between 50 Hz and 500 Hz we deduce the threshold to be 3 and 5 pe, which is well below the expected number of photoelectrons released on average by a straight MIP in one bar, as shown in Chapter 3.

The high threshold was typically set to 250 mV (the maximum possible value) for two reasons: first of all we had a limitation on the length of the buffer during the acquisition (see below for explanation) and we were thus interested in reducing the data rate in input to the TDCB. As a second reason it was realized that a larger “lever arm” for the time walk correction was going to provide a more precise determination of this effect.

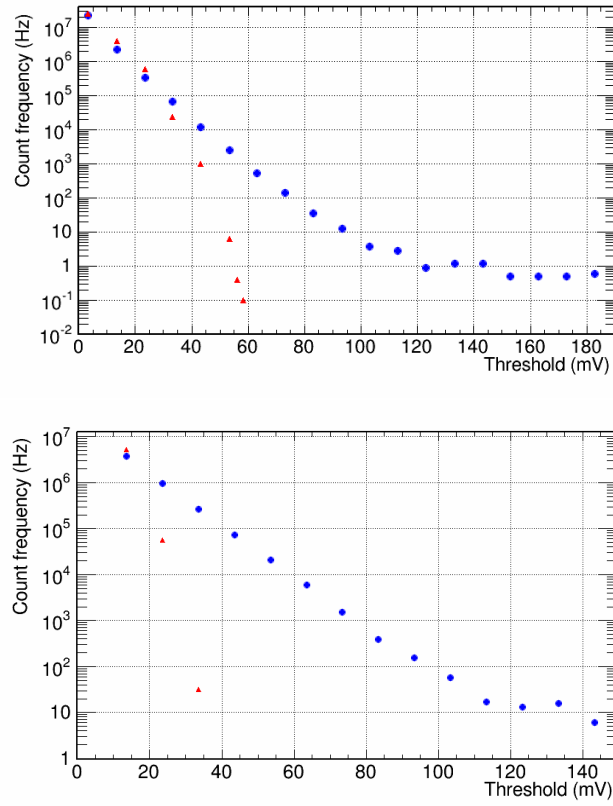


Figure 5.6: Rates of signals above a certain threshold for two channels. Red triangles: SiPM OFF. Blue circle: SiPM ON.

▣ 7-bit overlap between **Timestamp** and **Fine time**

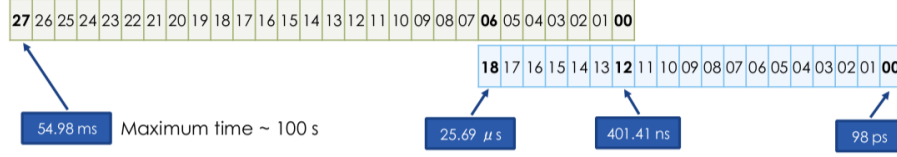


Figure 5.7: The data present in the monitor FIFO in the FPGA is of two kind: data stamps and fine times. The fine time is the information from the TDC which have 19 bits and a 98 ps tick so every $\sim 51 \mu s$ they rollback. To follow event in a longer time window, the time stamp information is regularly pushed into the fifo. These have 28 bit with a 400 ns tick so that 7 bit are common between timestamps and fine time. The lower bit number in the fine time is due to the fact that part of the word has to encode other information like the channel number or the TDC event type (leading or trailing).

5.3 Stand alone acquisition

The 6 CHANTI stations will require 2 TEL62 (see section 2.14) equipped with 3 TDC boards each (if we will decide to read each channel with 2 threshold) while in the testbeam only 1 TDC board was necessary (and installed). However, the SL and PP firmware was not ready until the last days of the testbeam, so for most of the time we needed an alternative way to acquire data. To this purpose we used the possibility of connecting through Ethernet to the CCPC of the TEL62 to monitor the PP FPGA: we take advance of a monitor FIFO between the TDC and the rest of PP firmware to fetch the data and send them to a PC which has to decode the data format and to save them in a ROOT tree format. In this operation we take care to carefully treat the Timestamps and Fine time information as shown in figure 5.7.

All the elaborations are then performed off line grouping as “Single event” all the hits collected within a 100 ns time window. A limitation of this system was that it was “Blind” to the trigger: when the acquisition starts it simple begins to save “Absolute” times and when the FIFO is full, it stops, so we were limited initially to 2048 words per burst summed over all the channels. During the TR a firmware upgrade permitted to increase this limit to 8192 words to allow for a more efficient data taking.

In the last days of the runs we also successfully tested the CHANTI operation in the official common acquisition chain. However most of the data was taken with our stand alone acquisition, so the following results are based on that data. As an example we report in figure 5.8 the distribution of the Leading time in the cases of two acquisitions, one with the beam on and one with the beam off. While the starting time of the burst is clearly visible, due to the buffer limitation the acquisition was possible only for a very limited amount of time (typically 100 ms /burst).

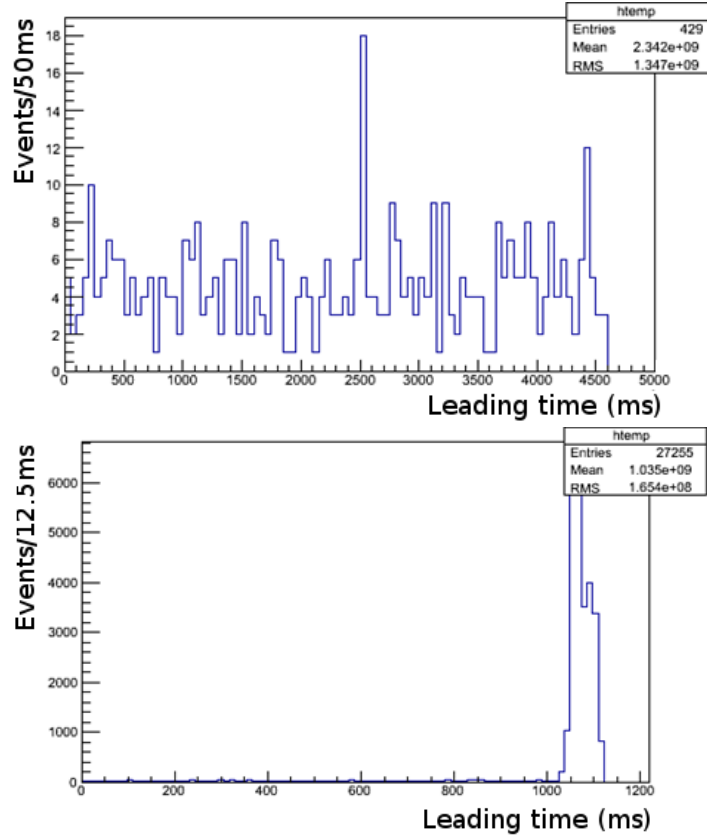


Figure 5.8: Distribution of the leading time of the hits read by the standalone acquisition. Up: Without beam. Down: With beam turned on. When the FIFO is full the acquisition stops so the acquisition follows the burst only for a short time.

In figure 5.9 there is a plot of the hit multiplicity versus the nominal x-y position of the channel. In the latter plot, as in the rest of plots in this chapter, we do not consider the full data sample but only the events in which there are exactly 4 bars producing signals at least above the lower threshold: two adjacent bars in x view and two adjacent bars in y view.

5.4 Time resolution

As it has already been discussed in Chapters 3 and 4, the time resolution is the figure of merit which will determine, in the end, the amount of signal which will be lost due to the CHANTI veto. As already stated, in order to keep the random veto coincidences at few percent level, and given a total rate of 2-3

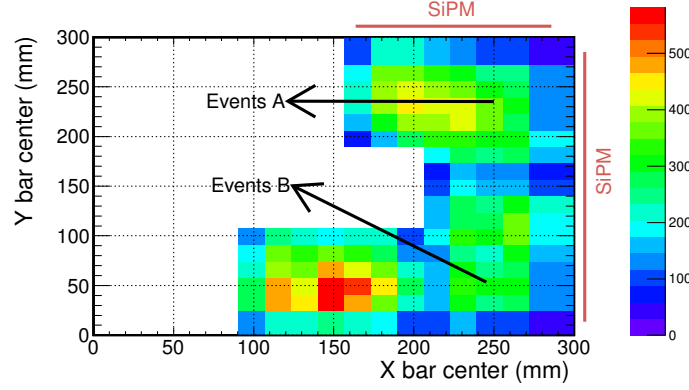


Figure 5.9: Hit multiplicity versus the x-y nominal position of the channel. It shows the position of the silicon photomultiplier (SiPM) and two kind of events. “Events A” are events hitting long bars in a point close to the SiPM in both the view. “Events B” are events hitting long bars in a point close to the SiPM of the Y view and far by the SiPM of the X view.

MHz on the CHANTI due to muon halo events and inelastic scattering, a ± 5 ns window around the trigger fine time will be the one in which a CHANTI signal will be looked for to veto the event. So, to keep the efficiency of the veto at good level, the time resolution of the CHANTI must clearly not exceed 2 ns. While a time resolution of order 1 ns or better has been already demonstrated on a few channels prototype in laboratory, before the TR it had still to be measured with the final readout electronics in the full NA62 framework.

During the TR the CHANTI was placed outside the acceptance of all other detectors: we can not perform a direct estimation of the CHANTI time resolution looking at a single channel, since we have no external reference trigger. However we can get this information by studying the difference in Leading times among two fired channels in the same event (see figure 5.10).

This distribution is made with “raw” times, before any correction. Of course a correction for the time walk of the signal and for the position of the hit with respect to the SiPM can improve the timing performance of the detector. These corrections will be discussed in the following.

5.4.1 Time walk correction

When a signal passes both thresholds for the channel we can make a correction using a linear approximation of the rising edge of the signal to obtain an estimation of the start time of the signal (C_{tw} , “Time walk correction”, see figure 5.4); the resolution on the extrapolated time is shown in figure 5.11.

This kind of correction is possible only when we have two thresholds per channel and when the signal actually passes both thresholds but we can use a

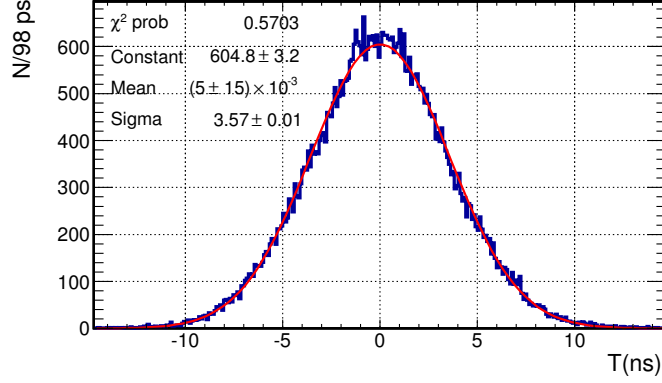


Figure 5.10: Leading time difference distribution for pair of bars without any correction.

statistically averaged correction to handle the other cases: in figure 5.12 we show the ToT corresponding to the first threshold versus the time walk correction for the events that pass both thresholds. The distribution has been fitted to a logarithmic dependence, as in figure, and we use the result to correct the time of hits passing the low threshold only. Although the very simple function used for the fit may be clearly optimized, as we will see, even this simple approximation gives good results. In figure 5.14 we show the distribution of the Leading time difference with this “Mean walk correction”.

Of course all of this “time difference” distribution include the fluctuation due to the time resolution twice, once per each measured time. If we make the reasonable assumption that the time resolution is the same for all channels, a good estimate of the time resolution is then given by the σ of the time difference distribution divided by $\sqrt{2}$.

Thus for the “raw” times we get an estimate for the single channel of $\sigma = 2.5\text{ns}$ which is marginally compliant to the specifications for the CHANTI. However the situation clearly improves for the resolution after the time walk correction: in this case we get, in fact, $\sigma = 1.34\text{ns}$, which is already in line with what requested for the CHANTI. However this value includes a source of fluctuation which is not intrinsic to the detector itself, but is related to the distribution of the impact points of the muon halo on its surface. This effect is addressed in the following section.

5.4.2 Position correction

The time resolution measured above is influenced by the different path lengths that light (collected and re-emitted by the WLS fiber) has to travel inside each bar before reaching the SiPM. To understand the magnitude of this effect we can consider two classes of events separately:

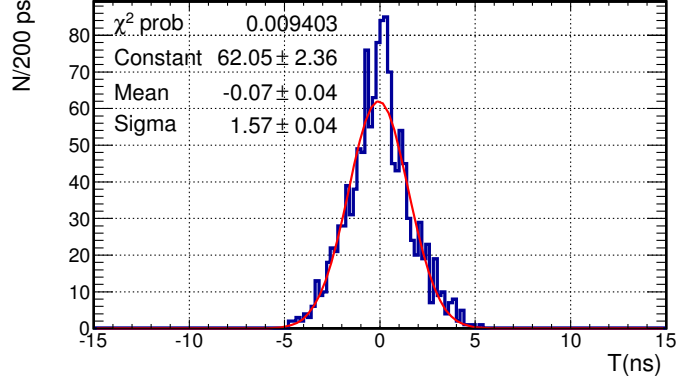


Figure 5.11: Difference of signal start time distribution for two bars. The start time is estimated by mean of the Leading time and the C_{tw} correction.

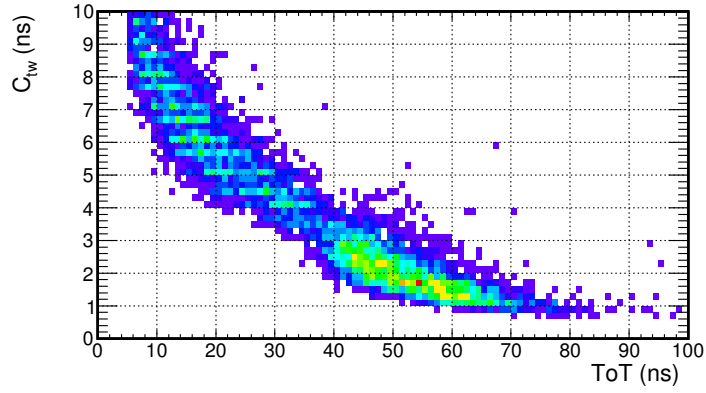


Figure 5.12: Time walk correction vs ToT for the events passing both thresholds.

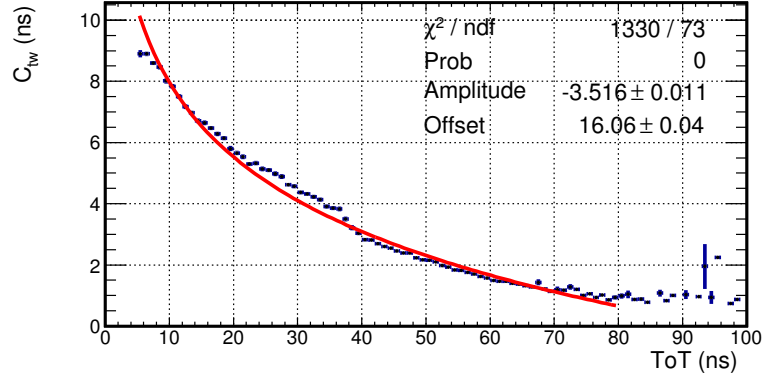


Figure 5.13: Logarithmic fit to the Time walk correction vs ToT dependence.

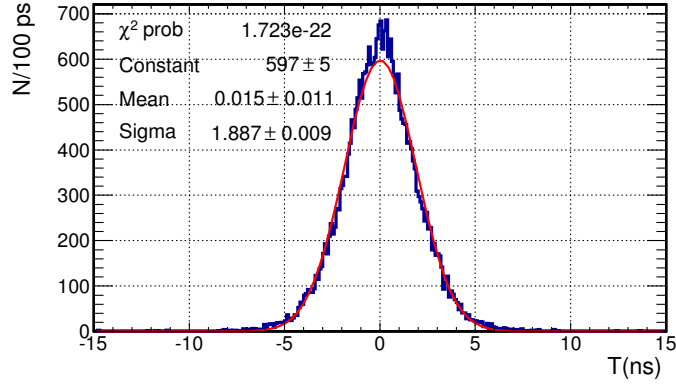


Figure 5.14: Difference distribution of the Leading time of two channel, with the mean Time walking correction. The correction was applied to all the signals, also the ones that pass only the lower threshold.

- Parallel bars: we study the distribution of the Leading time difference in adjacent bars lying both in the same view.
- Crossing bars: we study the distribution of the Leading time difference in two bars on different views and that intersect themselves.

The time difference among parallel bars is not affected by position effects, since in this case the distance covered by the light should be the same for the two bars, and in the difference the position contribution cancels out. The Leading time difference (after the mean time-walk correction) for parallel bars is shown in figure 5.16 and has a σ significantly smaller than the one shown in fig 5.14.

When we register a hit in two crossing bars, we have the possibility to know the two dimensional hit position by means of the xy view information. In these cases we expect to be able to correct, on average, the times for the position effect and obtain an intrinsic resolution as the one in 5.16. At a first sight the implementation would imply just to consider a linear correction related directly to the effective light speed inside the WLS fiber. Despite its apparent simplicity, however, the position correction deserves in our case some more thought. As it has already mentioned in Chapter 3, the WLS fibers used inside the CHANTI are mirrored at the end opposed to the SiPM. This allows to recover the fraction of light re-emitted inside the WLS fiber in the direction opposed to the photodetector, but has also some important effect on the shape of the signal. As a matter of fact, whenever a threshold larger than zero is set on the signal amplitude this implies that, statistically, a part of the reflected light may reach the photodetector and contribute to the signal formation before it reaches the threshold itself. This effect becomes of course more and more relevant as far as the threshold is set to higher values since, in that case, a larger fraction of reflected light will be collected before the signal reaches the threshold level. When we come to evaluate the leading time of the signal, (i.e. the time it reaches the threshold) it may have in general a non linear dependence on the position, since the “direct light” and the “reflected light” which contribute to the signal formation have opposite time dependencies as function of the position along the fiber where the light is re-emitted.

In order to evaluate this effect, and also in view of a more accurate description of the CHANTI detector inside the NA62 Monte Carlo we developed [61] a detailed Monte Carlo simulation of the CHANTI bar. This simulation includes all optical surfaces including the TiO_2 coating, the optical glue, the fiber clad and core, and its mirrored end. By tracking the optical photons inside the bar and by applying a simple circuit analog to describe the SiPM, we have been able to simulate individual signals inside a bar which reproduce all of the main features of the real signals (see fig. 5.15), and, among other things, may be used to evaluate the effect of the time correction.

With the help of this simulation, we have studied the dependence of the Leading time as a function of the distance from the SiPM of a MIP crossing the bar (with a 80 mV threshold), and we see (figure 5.17) that for small distances the Leading time does not depend on the position, while for larger distances shows an approximately linear trend with a coefficient that depends on how

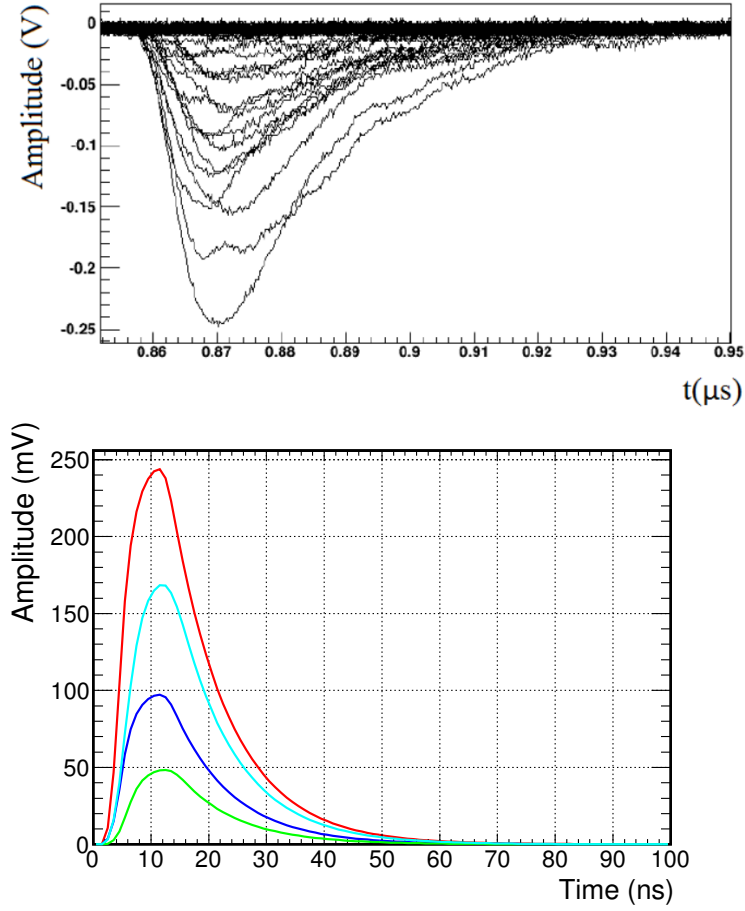


Figure 5.15: Accurate Monte Carlo description of the bar. Top: real signals; Bottom: simulated signals.

the light was generated and on the chosen threshold, and is not equal to the mean light speed in the fiber. Since we are estimating the resolution through a time difference, and a constant delay for all the bars clearly does not affect this quantity, we can perform the correction only when the hit point is far from the SiPM, namely more than halfway to the end of the bar. We want also to find a way to estimate on data the angular coefficient of the linear correction.

We can write a generic time measurement which we want to correct as

$$t_m = t_t + T(x, q) + C(x)$$

where t_t is the true time that we want to estimate, T is a stochastic delay due to physical effect like light generation and propagation, and C is the correction we applied. T depends on the hit position x along the bars and from other random variables q , while C depends only on x . Supposing q and x to be uncorrelated, the variance of this quantity is

$$\sigma_{tm}^2 = (\partial_q T)^2 \sigma_q^2 + (\partial_x (T + C))^2 \sigma_x^2 \equiv \sigma_t^2 + (\partial_x T(x, q) + \partial_x C(x))^2 \sigma_x^2$$

where we called σ_t^2 the intrinsic variance of the quantity i.e. the one due to effects that we are not trying to correct. This shows that the resolution (σ_{tm}) is minimal when the derivative of the correction C is the opposite of the derivative of the effect T which we are correcting. Now, supposing to correct a linear effect $T(x, q) = \frac{x}{v_t}$ with $C(x) = -\frac{x}{v_m}$, we get

$$\sigma_{tm}^2 = \sigma_t^2 + \left(\frac{v_m - v_t}{v_m v_t}\right)^2 \sigma_x^2$$

and so at its minimum the v we are using in the correction is the right one:

$$\partial_x \frac{x}{v_t} = \partial_x \frac{x}{v_m} \quad \leftrightarrow \quad v_t = v_m$$

To apply this result to our case we notice that a constant correction does not affect the resolution and that we have to apply the correction only to the bars fired far away from the SiPM. We consider only long bars so the exact point of switching from the constant behavior to the linear one is not a issue since we have two well separated classes of events, as shown in figure 5.9:

- events A: for both the view the hitting point is near the SiPM, so we do not apply the correction.
- events B: in a view the hitting point is near the SiPM, in the other is far, so we apply the linear correction with a parameter v .

Estimating by the data the value of the σ_{tm}^2 and performing a scan for various values of the parameter v we get the plot in figure 5.18; then we perform a fit with the shape $(v - v_t)^2 / v^2 v_t^2$ to extract the parameter v_t . A plot of the Leading time distribution corrected with $v = v_t$ is in figure 5.19, and it shows that the σ obtained is the same of the event with two parallel bars.

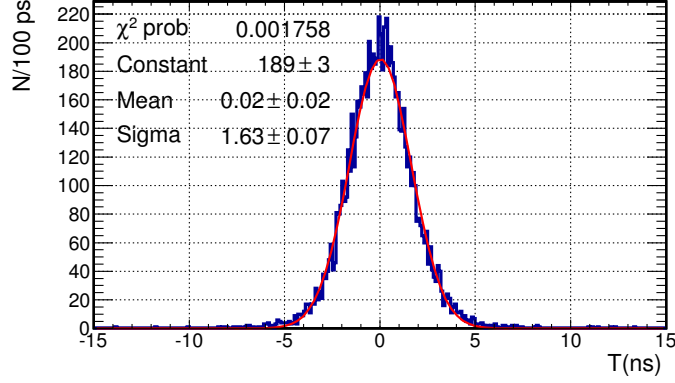


Figure 5.16: Leading time difference distribution for two adjacent parallel bars after mean time walk correction.

This indeed is exactly what one expects if the position correction is well done. By applying the usual $\sqrt{2}$ factor to the measured sigma, we get for a single hit time resolution, after both time walk and position correction an “intrinsic” time resolution $\sigma = 1.15\text{ns}$. This is quite satisfactory, also considering the fact that, given the geometry of the detector, on average one expects that two to four bars will be fired by each charged particle crossing it, and, thus the determination of the time of the interaction will be averaged, allowing to reach in most cases a resolution below 1 ns.

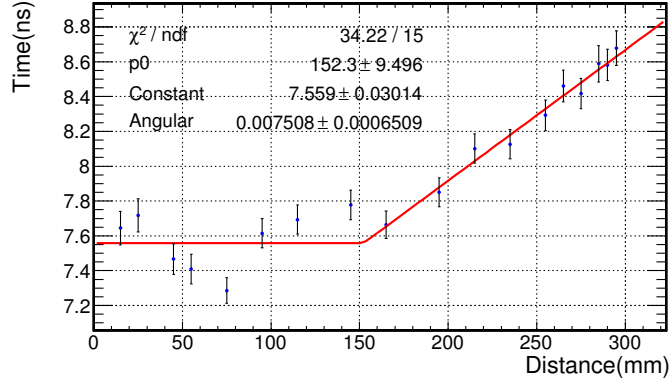


Figure 5.17: Position correction to Leading time by Monte Carlo simulation supposing a threshold of 80 mV. “Constant” is the correction (in ns) in the constant part while “Angular” is the time to distance linear coefficient (in ns/mm)

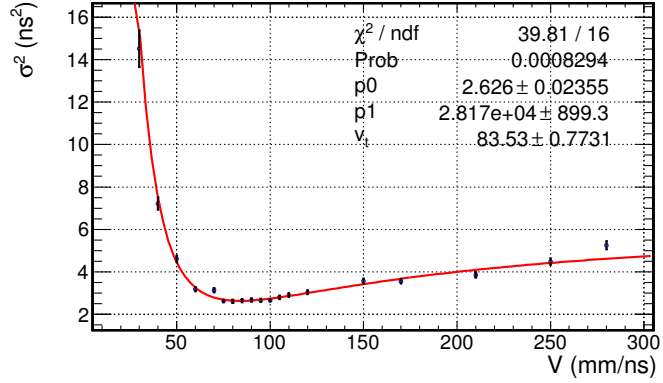


Figure 5.18: Sigma of the time difference after correcting the position effect as a function of the slope correction v . The minimum corresponds to the best correction v_t . A fit using the model function described in the text is used to estimate v_t .

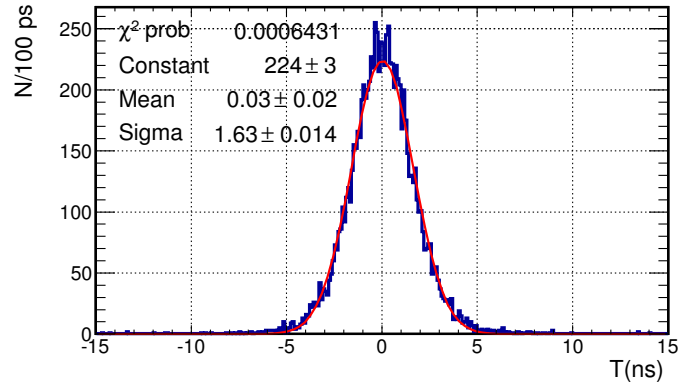


Figure 5.19: Distribution of the time differences corrected for both time walk and the position effect with the fitted v_t .

Conclusions

The NA62 group of Naples has proposed, designed and tested a detector (CHANTI) based on scintillator bars, wavelength shifting fibers and silicon photomultipliers to reduce the background induced by beam inelastic interactions in the NA62 experiment. In this thesis we described the design philosophy, the construction procedure and the quality tests we adopted during the assembly of the detector. It is shown that the CHANTI can safely operate in vacuum, has a high response to charged particles, reaches a time resolution of less than one nanosecond and is capable to suppress the inelastic background by two orders of magnitude, as shown by a detailed analysis of simulated data. Moreover the tracking capabilities of the detector will help in validating the Monte Carlo simulations of the inelastic interactions by measuring charged particles multiplicities and angular distributions, and will allow to track the beam halo muons in the region close to the beam.

It is worth to note that the study of inelastic background performed in this work suggests that the inelastic interactions on GTK-3 will be one of the major contributions to the overall background in the $K \rightarrow \pi^+ \nu \bar{\nu}$ analysis. A pure Monte Carlo based estimation of this background is thus not an option for at least two reasons. First of all the events passing the cuts are related to the far tails of the resolution in vertex position which are of course simulated with limited accuracy, and could cause large systematic uncertainties on the estimated background. Moreover, even from a purely statistical point of view, the amount of inelastic interactions on GTK-3 to be produced in order to get a 1:1 ratio with the data is so large (roughly a factor of 10^5 times the ones studied in this work) to be impractical (not to speak of reaching MC/data ratio of order 10 or more to reach a satisfactory Monte Carlo statistical error on the estimated background counts). For this reason a dedicated study on how to evaluate the inelastic background directly on data will have to be addressed. While at this stage this kind of study cannot be detailed yet, it is clear that in this context the CHANTI will play a crucial role in order to select a clean and almost unbiased sample of inelastic events to evaluate the effect of the other selection cuts.

The results obtained with the first prototype board of the CHANTI front-end electronics during the NA62 Technical Run in November 2012 confirm the earlier laboratory tests, and allow us to proceed with the finalization of the construction of the detector and of its electronics. When writing these lines 90% of the bars of the CHANTI had been assembled and more than half of them

have been individually tested, while the CHANTI vessel is being finished in the Mechanical Workshop of INFN, Naples. The completion of the construction, and the commissioning of the front-end board is foreseen by end 2013- beginning 2014. The full detector will be installed at CERN during year 2014, in time for the start of the NA62 data acquisition.

Appendix A

Silicon Photo Multiplier

A.1 Silicon detectors

Silicon detectors work according to the functioning principle of the photodiode, whose general scheme is reported in figure A.1. The photodiode is constituted by a simple P-N junction, where P and N indicate the kind of extrinsic doping of the semiconductors. The energy gap between the valence and conduction band allows the absorption of photons in the visible spectrum (from 1.5 eV to 3.5 eV). When this happens, a valence electron will move to the conduction band, i.e. creating a free hole-electron pair (h-e pair). Both electron and hole start to diffuse inside the silicon until they will recombine with a typical lifetime that, depending on the crystal quality, can range from 10^{-9} s to 10^{-3} s.

When an electric field is present, we can collect the opposite charges before they recombine. This happens when the photo-electron is generated in the depletion region: in such a case, the electric field associated to space-charge regions can force the electron to drift from the N to P (and the holes from P to N). As the electron (and/or the hole) reach the diode cathode (or anode) a current is detected. To enhance this effect the junction is inversely polarized in order to increase the depletion area and electrical field. This reduces the signal noise since it reduces the parasite capacity of the junction.

Another way to improve this kind of detector is to use a PIN scheme, i.e. putting a nominally intrinsic layer (I) between the two heavily doped P+ and N+ layers. In this configuration a depletion region will be created all along the intrinsic semiconductor thickness, also with no external polarization. This, as before, reduces the parasite capacity of the junction and increases the “Active” volume of the photo-detector.

The previously discussed detectors have not any internal amplification process so they need additional electronics to actually read the signal. However increasing the bias voltage above $1.75 * 10^5 V/cm$ [63] the electrons generated in the P layer can gain enough energy to make secondary excitation, and so on, generating an avalanche. However if the electron is generated in the N layer,

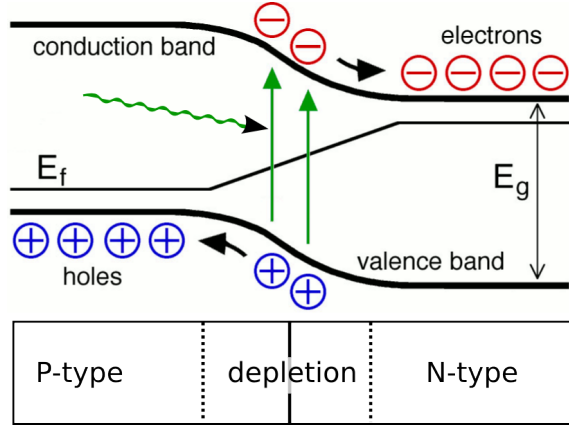


Figure A.1: Generic P-N junction diode layouts.

where the field is lower, they can not trigger any avalanche. For the same reason, when an avalanche reach the N layer, the secondary electrons can not continue the ionization, so the avalanche stops. Such device are called Avalanche Photo Diode (APD).

With this technique we can reach an amplification up to 10^4 ; also with this system we still need some relative complex amplification electronics to actually read the signal.

If we increase the field above $2.5 \times 10^5 V/cm$ [64], also the holes can gain enough energy to produce other h-e pairs. In these conditions the avalanche is self-sustaining. For example consider a electron triggered avalanche: when it reaches the N layer the electrons stop to produce other e-h pairs, but the holes starts to move backward and to produce h-e pair. So the avalanche continues indefinitely. This working regime is called “Geiger mode”, and give an upper limit to the gains that we can obtain with an APD, because in this situation the device has a namely infinite dead time.

A.2 Geiger mode Avalanche Photo Diode

At beginning of the millennium a new kind of silicon device started to be developed. This is designed to work in the Geiger regime i.e. the one in which the kinetic energy of both holes and electron suffices to trigger an avalanche. In this way we can reach very high amplification (up to 10^7) and the signal output can be read without external amplification electronics (or with very simple one).

As discussed previously (section A.1), in these condition the avalanche is self-sustaining so, to make this devices suitable, it must be introduced an external quenching mechanism that suppresses the avalanche, such as a quenching resistor. This configuration permits us to use higher field and to get higher gain like in the first dynode of a standard photo-multiplier. For this reason some-

body prefer to call this kind of device Silicon Photo-Multiplier (SiPM) instead of Geiger Avalanche Photo Diode (GAPD).

There are two kind of problem with these device:

- Since they work in a Geiger regime, the signal they generate is not proportional to the number of photons actually crossing the active area
- They can work only at relative low rates: also using active quenching circuits the maximum working rates may not exceed 1 MHz

An attempt to solve both problems was done redesigning this devices as a parallel of many GAPS (figure A.2). Current technology permits us to develop devices with up to 40000 cells/ mm^2 . In this way the output signal is the sum of signals from each cell, so it is proportional to the number of cells fired (see section A.4). Despite the fact that the maximum working rate does not actually increase, with the new layout we can know if more photons reach the device at same time by looking at the pulse height.

A issue with the multicell design is the optical crosstalk: during the avalanche development inside a cell, a photon could be emitted having enough energy to start another avalanche (for each 10^5 carriers in the cell there are about 3 photon emitted in the visible spectrum [65]). If this photon reach another cell it could trigger another avalanche that is not related to the incident light. Narrow grooves filled by an optical absorber are placed between the cells to reduce this effect.

Another issue to be solved can be represented by avalanche photons that go outside the SiPM and then reenter due, for example, to the reflection inside the external crystal. Some studies show that this process increases the count of the multiple photo-electron signals of about $\sim 18\%$ [62].

We note that the avalanche may be triggered by carriers generated by any process. They can be generated also by thermal excitation, so there are some signal also without any crystal and isolating (optically) the device. A typical rate of this kind of signal is about 100 kHz per mm^2 at 25°C and it get a factor 2 of reduction every 8°C of temperature drop [62]. This is called “Dark rate”. By optical cross talk this process may also produce multiple firing cells, however the count decreases of about one order of magnitude per fired cell. So increasing the detection threshold in order to reject events with less than 4 cell fired, the dark rate at 25°C typically drop to 1 kHz [62].

A.3 Photon detection efficiency

The photon detection efficiency ϵ_{pd} for a GAPD cell depends n several variables, namely:

- The amount of active area (fill factor), i.e. the ratio between the active area and full SiPM area (F_f)

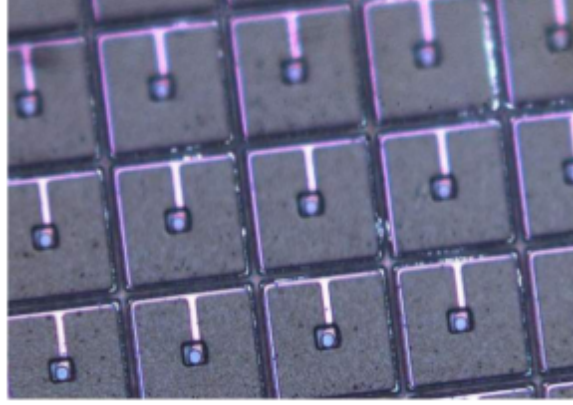


Figure A.2: Magnified photo of a GAPD cells array.

- The efficiency of the h-e photon production (quantum efficiency), i.e. the probability that an h-e pair is generated by a photon entering in the silicon (ϵ_q)
- The avalanche trigger probability, i.e. the probability that a released carrier produces an avalanche (P_t)

So we can write $\epsilon_{pd} = F_f \times \epsilon_q \times P_t$. To increase F_f we can not simply put the cells closer since we need the space to place individual quenching resistors (see section A.4) on the cells and to place the photon absorber (in order to reduce the optical crosstalk). The only other way to increase F_f is to use few big cells but this reduces the signal dynamic range. So we have to find a compromise on the basis of our needs: common fill factors range from 40% up to 60%.

ϵ_q depends by the used technologies and typical value go from 80% to 90% depending on the photon wavelength. An example is in figure A.3.

P_t depends on the ionization coefficient for the electrons and the holes. These coefficient depends on the electric field, so P_t and ϵ_{pd} depends on the over-voltage applied to the junction. A typical value for the photon detection efficiency in a case of 1V over-voltage is about 20% (in figure A.4 it is shown the wavelength dependence) however already with 2 V it could reach the 40% (typical voltage dependence is shown in A.5).

We highlight that also the dark rate depends over P_t .

A.4 GAPD electric model

In figure A.6a there is a model of a single SiPM cell, with its power-up circuit and quenching circuit. We considered only a passive quenching or rather a resistor R_q ($\sim 300k\Omega$); more complex design substitutes this resistor with an active load.

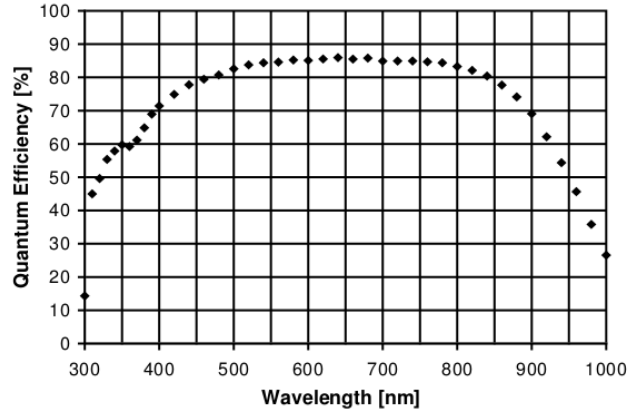


Figure A.3: Quantum efficiency ϵ_q of an APD produced by Hamamatsu (type S8148) [66].

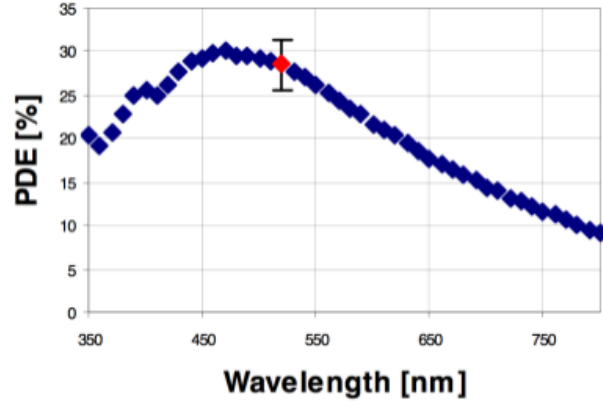


Figure A.4: Wavelength dependence of photo detection efficiency. Bars on the red point indicate the systematic error. The SiPM was operating at 1V over the breakdown voltage [62]

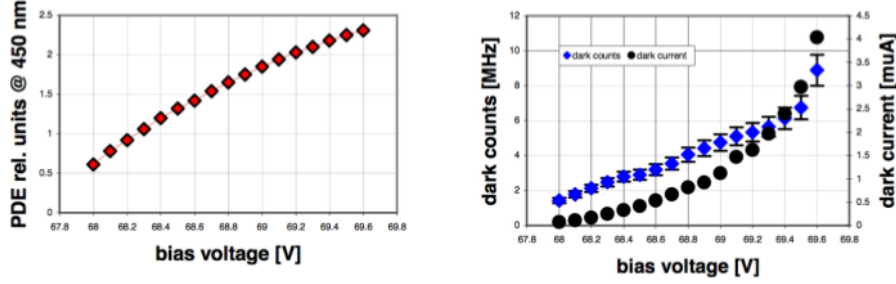


Figure A.5: ϵ_{pd} of Hamamatsu MPPC-33-050C as function of the applied bias voltage in relative units (left) and the dark currents and dark counts measured at 25°C in the same voltage range (right). [62]

V_{bd} represents the diode breakdown voltage and V_{bias} is the power-up voltage, that we assume to be $> V_{bd}$. R_d is the internal resistance ($\sim 1k\Omega$) of the cell.

We want our model to describe correctly the signal shape, so its prediction is explained in the following and it is shown in figure A.7. We consider the two stationary states: S_0 with the switch S open and S_1 with S closed. In the former there is no absorbed external current I_d , so the voltage V_d at the capacitor C_d is equal to V_{bias} . Instead in the latter we have $I_d = I_{max} = (V_{bias} - V_{bd})/R_q$ and $V_d = V_{bd}$.

Suppose to start in S_0 , when a photon reach the active area there is a probability P_t (trigger probability) that S closes. In this case it starts to move toward S_1 with a time constant $R_d \times C_d$. Thus, calling $t_M = 0$ the moment when I_d is maximum, the current goes as $I_d = I_{max}(1 - \exp(-t/R_d C_d))$ and the voltage to C_d accordingly drops until the field inside the silicon can not sustain the avalanche anymore. In our model this mean that S will open again (actually this is a stochastic event that may happen with probability P_{off} that depends by V_d). So somewhere near $V_d = V_{bd}$ (and maximum I_d) the circuits open and C_d starts to charge with a time constant $R_q \times C_d$. Now the current trend changes in to $I_d = I_{max} \exp(-t/R_q C_d)$. When C_d is fully recharged the system returns in S_1 again. A typical signal height for a single cell is several mV with an 50Ω load with a rise time of about $1ns$ and a tail of several $10ns$ (see section A.6).

As we see, the quenching resistor has an important role in the signal shaping, so it should be as stable as possible. Indeed the device dependence on the temperature are due also to this resistor.

In a SiPM we want that the signal from every cell add together in order to get a response linear with the number of fired cells. Since the output signal is the absorbed current, and since we also want all the cells to be at same voltage to get an homogeneous response, we have to design the cells in parallel wrt the power-up and reading point.

The quenching resistor (or circuit) could not be placed upstream otherwise the discharge (quenching) time constant depends by the number of fired cells.

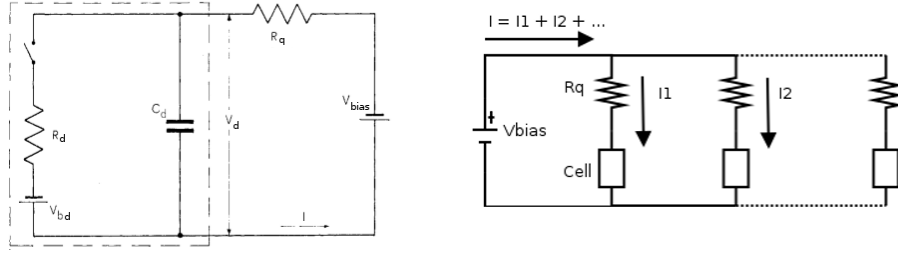


Figure A.6: SiPM electrical model. Left (A.4a): inside the in dashed box there is the single cell model while outside there is an power supply model. Right (A.4b): multiple cells connection.

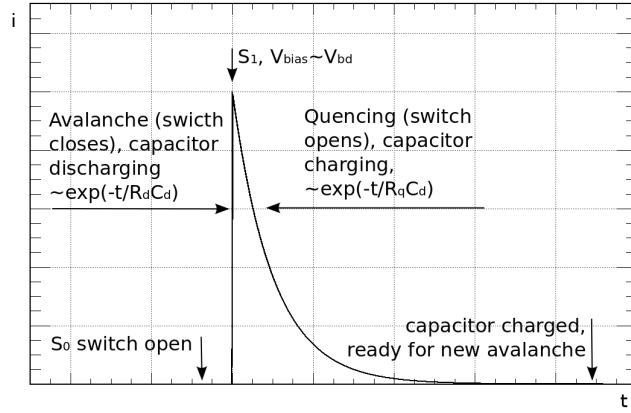


Figure A.7: Signal shape from the simple electric model.

So every cell must have a dedicated quenching resistor (or circuit). The final layout is in figure A.6b.

We highlight that, in this configuration, the output signal is proportional to the number of cells fired so the dynamic range of the output is equal to the number of cells.

A.5 Bias voltage and gain

We can define the gain as $G = \int I_a dt / q$, where I_a is the absorbed current and q is the electron charge. For the simple electrical model we discussed before, neglecting the fast rising phase, we have $G = I_{max} R_q C_q / q = (V_{bias} - V_{bd}) C_d / q$. So, notwithstanding the GAPDs work at relative small voltage (with respect to classical phototubes), they have a good gain: normally the power supply is below 100 V, and the gain goes from 10^5 to 10^7 .

Regardless of the number of photon that generate carriers in the same cell

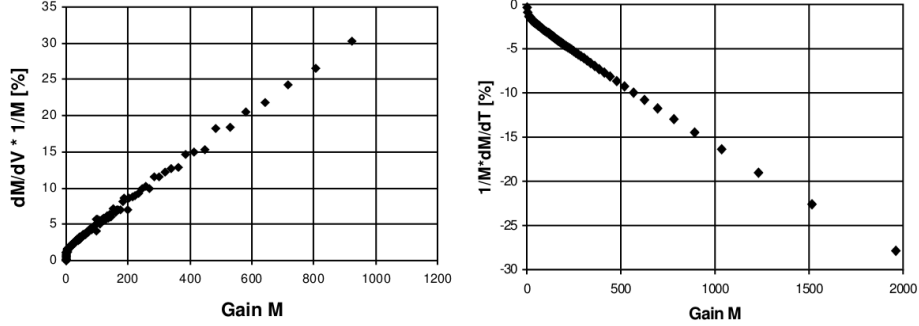


Figure A.8: The total gain M (over all SiPM cells) is proportional to A . Left: Voltage coefficients k_V (V) of a GAPD from Hamamatsu as gain function. Right: Temperature coefficients of a GAPD from Hamamatsu as function of the gain. [66]

during the avalanche, the single cell signal will be always the same, only the trigger probability may increase. As said the signal amplitude (I_{max}) generated by the cell will be proportional to the gain and so $A_i \propto C_d(V_{bias} - V_{bd})$ (i is the cell index). When summing the signal from all the cells, if the cells fire with a delay compatible with the rise time (few ns), the output amplitude is simply the sum $A = \sum^{fired} A_i \propto N_{fired} C_d(V_{bias} - V_{bd})$. So the signal is proportional to the number of cells fired N_{fired} .

Considering a small light pulse, N_{fired} depend over the number of photon that reach the SiPM N_{ph} and the number of cells N_{tot} . A good approximation of this dependence is $N_{fired} = N_{total}(1 - \exp(-\epsilon_{pd} N_{ph}/N_{tot}))$, for example, when $\epsilon_{pd} N_{ph}$ exceeds 50% the deviation from linearity is more than 20% [62]. Thus the SiPM response dependence over the photon detection efficiency is much stronger wrt traditional photo tubes device.

Since N_{fired} depends on ϵ_{pd} , it change with V_{bias} . Furthermore A depend over V_{bd} that change with the temperature T . In the figure A.8 is shown the sensitivity to these parameters for a Hamamatsu. The plots refer to the dependent coefficient $k_X = d \log(M)/dX \times 100\%$ where X could be V_{bias} or T .

We note that this dependence to the temperature introduces a negative feedback to the absorbed power: if the SiPM heats up the V_{bd} increases and if the V_{bias} is the same, the gain decreases and they start to absorb less current. So the power absorption cannot go out of control.

A.6 Time properties

As explained from the simple electrical model of the GAPD cell, the rising slope of single photon signal depends over the the internal capacity of the junction and internal resistance (the topic is slightly more complex because for the multi-cell design). The capacitance strongly depends on the SiPM design; as reference,

a $50 \times 50 \mu m^2$ Hamamatsu has a capacitance of about 10 pF. Instead A typical value for the resistance is about $1 k\Omega$. So we can expect a rise time of few *ns*.

However the time resolution depends not only over the rise time but also over the signal jitter. We expect this to be a minor effect since the GAPD has a very thin active layer ($2\text{--}4 \mu m$) and the avalanche process is fast. Indeed, as shown in A.9, for a single photons we can reach time resolution below the 50 ps.

The signal rise time of more photons could be slower for the effect of the convolution where the signal tails could became important. They, as said before, depends on the internal capacity and the quenching resistor that have to be big ($\sim 300 k\Omega$) in order to reduce the time needed for the quenching and so the signal width; so the tails can be greater than 100ns. A properly design of the readout electronics can help to reduce this effect; test with common setups shown that the rise time in response to short ($\sim ns$) light pulse could be of few *ns*.

The signal tail also affects the dead time of the device since we need to wait that the junction is fully recharged before detect efficiently the next photon. As shown in figure A.9 the signal width is ~ 100 ns.

We highlight that short time (few ns) after an avalanche stops and the recharge starts, the field strength inside the cell is enough to trigger another avalanche. However in this case the conditions are not optimal for the signal development, so the generated currents are lower. Only when C_d is fully charged (~ 100 ns) we can get signals having the full height.

This effect is evident on the afterpulses. The avalanche generates a high temperature (few $1000^\circ C$) plasma that left the silicon in an excited state. Each trap can return in the base state emitting a carrier that can trigger a second avalanche. These signals are called “Afterpulses”. Typically this phenomena could have more component (e.g. the Hamamatsu S10362-33-050C has one component with a mean time around $50 ns$ and one around $150 ns$ [68]). Plotting the afterpulses signal height versus the time distance from the first discharge, we can see the recovery feature previously discussed as shown in figure A.10.

A.7 Other properties

The principal construction property of the SiPM is its compactness. Also if the small area is a disadvantage when we have to monitor big surface, it helps in the design of detectors with higher granularity. Furthermore the kind of material used and construction technologies involved make them ideally perfect for low-cost mass production.

Studies were performed on the SiPM radiation hardness, in fact after exposition to γ [70], neutrons [71], protons [72] and electrons [73] it can produce defects in the silicon structure. These defects could be electrically active by changing the doping concentration or by becoming a charge trap. The γ produce defect on the interface with the isolation layer on the surface; the sensitivity to γ radiation could be eliminated by a proper design of this layer. The hadron instead

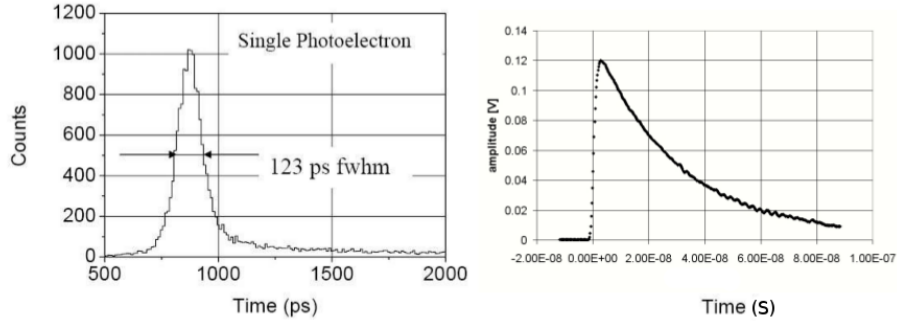


Figure A.9: Left: Time resolution for single photon. Right: Signal for a short ($\sim$ ns) light pulse. Reprinted from [67].

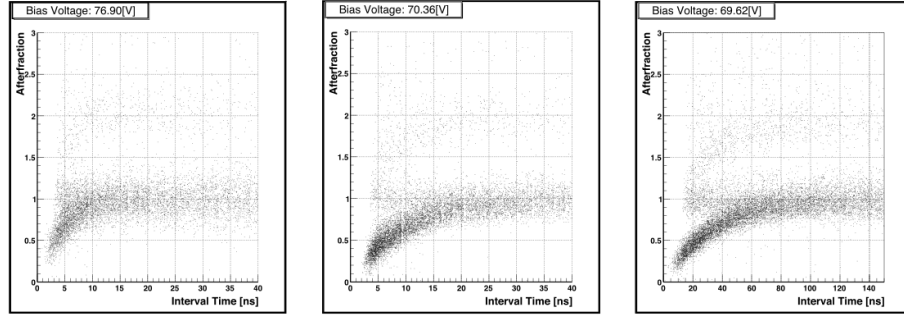


Figure A.10: Recovery curve for a GAPD at different bias voltage. Reprinted from [69].

produce defect in the bulk silicon and they can increase the dark rate: after a dose of 2×10^{10} neutrons/cm² it could change of one order of magnitude.

There are tests of the SiPM long term stability that use process of “Accelerated aging” [74]: the devices are placed for 30 days in environment at 80°C. After that, their response to a LED was monitored for more than 7 month. With few exceptions all the detector show no changing in the response wrt tests before heating.

Another important feature of the SiPM is that the response to a ionizing particle passing across the whole junction is the same of the single photon. This is explained by the fact that the only active volume for such particle (where a carrier could actually trigger an avalanche) is few μm of thick. Sometimes it refers to this properties as “Small nuclear counter effect”.

Others good properties of the SiPM are:

- Low power consumption ($< 50 \mu W/mm^2$) and heat generation
- Magnetic field insensitivity (up to 7 TeV tested)
- Tolerance to accidental illumination

Many different design are possible for the GAPD and there are many parameters that could be change to adapt these device to our needs:

- Semiconductor material - It influences the Photon Detection Efficiency and wavelength sensitivity dependence
- Thickness of the depleted layer (and PIN geometry) - It has effects on accepted wavelength, gain, optical crosstalk and dark counts
- Doping concentration - It could change the operating voltage and range
- Active cell area and cell number - It influences the gain, the dynamic range and recovery time
- Quenching resistor - recovery time, count rate and temperature stability depend on it

We highlight also that we have to keep the crystal impurity as low as possible in order to reduce dark counts and afterpulses.

Appendix B

TEL62 network latency

B.1 Ethernet monitoring

We describe here a small project we worked on to acquire the know-how to manage the CHANTI TEL62 and contribute to the general data acquisition framework. The main issue here is the LKr calorimeter. Since the amount of information handled by the calorimeter is huge the collaboration has decided to keep data buffer locally and send them to the TEL62 only when a trigger is detected. In order to work, this solution needs a custom board that will replace the standard TDC boards on the TEL62, the “Talk” board. However in this scheme the latency time between LKr and Talk is critical since the LKr has to empty its buffer when full. The communication between LKr electronics and Talk board is performed through Ethernet and the project we realized was the measurement of the latency time of this communication.

We use the Gbit Ethernet card on a Tell1 board, which is a early prototype of the TEL62, to send data to the Talk on the same board: this is not actually the final layout but helped us to easily define sending and receiving time using probes on the Tell1. We also had to make some changes to the SL FPGA firmware, to be able to send fake data (timestamps) through the Gbit Ethernet and monitoring the start of the communication. The firmware is controlled by the CCPC to setup data size, fake event frequency, and so on.

We define the latency as the time between the start of the first packet of data and the arrival time of the last packet (see figure B.1), so it depends upon the data size: we make this choice to have a full view of the communication when we start to stress the Ethernet.

B.2 Working latencies

In figure B.2 we show a plot of latency time versus the number of byte sent for various time distance between two consecutive triggers. Instead in figure B.3 there is the same plot for $S=25\mu s$ with a description of the various operative

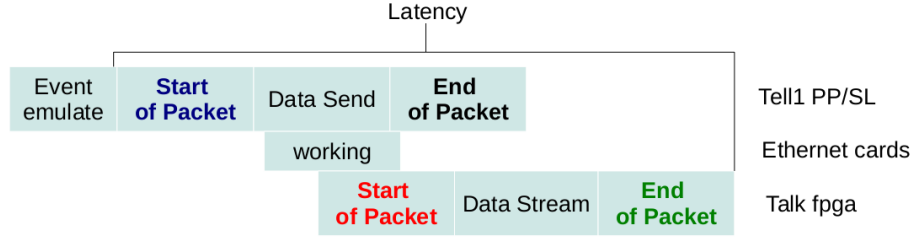


Figure B.1: Operation timing between Tell1, Ethernet cards and FPGA. We define the latency as the time needed to receive the full data stream.

regions. For low frequencies the dependencies is linear however when we overload the Ethernet boards they start caching causing wiggling non linear dependencies. We call X the point where the dependence stops to be linear.

In figure B.4 there is a plot of the latency and byte sent for the X point at different trigger frequencies. The line joining all the X point surround the region where the couple trigger-frequency / byte-sent is in in the linear working point of the Ethernet cards. The blue line is the nominal rate defined by $T = B/(V \times G)$, where T is the time between packets, B is the byte in a packet, V is the primitive rates (10 MHz), and G is the size of a primitive (8 byte). This confirms that the Ethernet latency is adequate to handle the LKr rates.

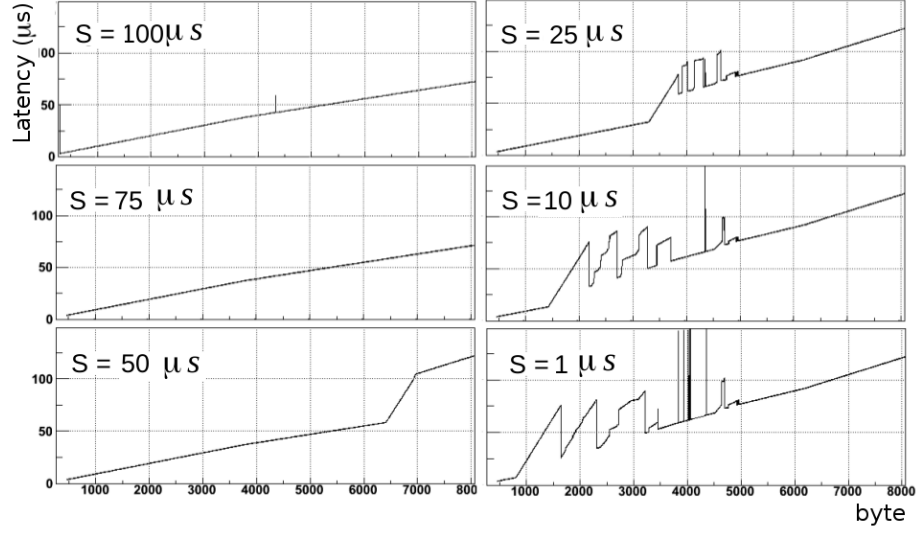


Figure B.2: Latencies obtained for several fake trigger period as a function of the number of bytes sent. When the Ethernet card is stressed the dependence stops to be linear.

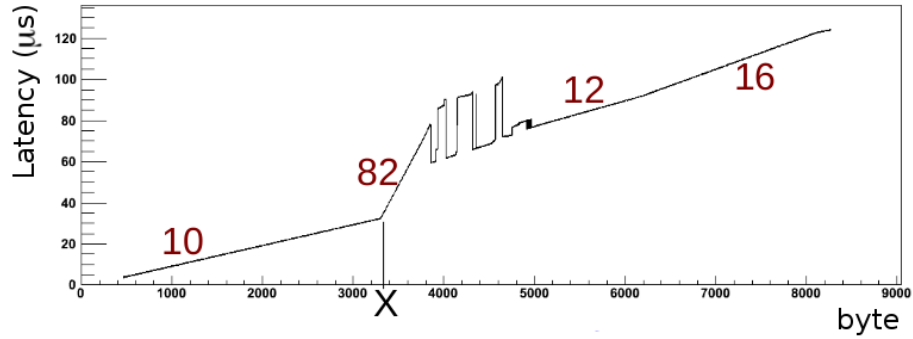


Figure B.3: Latency versus number of bytes sent. The red number are the slopes in $ns/byte$. In the first linear part, the $10ns/byte$ correspond to the Ethernet speed ($1Gb/s$). We call X the point where the dependence of the Latency on the number of bytes sent stops to be linear.

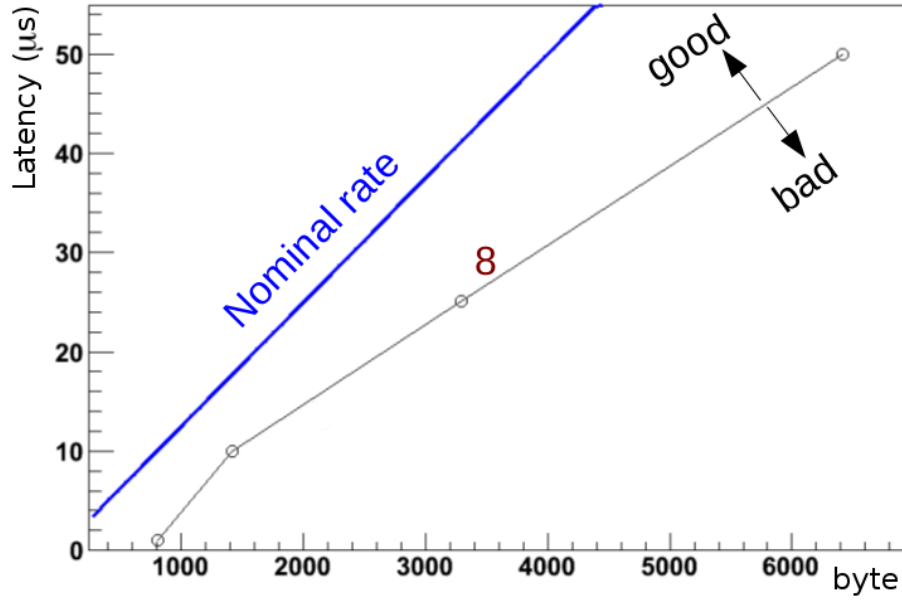


Figure B.4: The black circles are the points (for different trigger frequencies) where the dependence between latency and number of bytes sent stops to be linear. The blue line is the nominal rate expected in NA62, that is in the “Linear” region of the plane. The slope of the black line is about 8 ns per byte.

Bibliography

- [1] M. Kobayashi, T. Maskawa, Prog. Theor. Phys. 49, 652 (1973).
- [2] L. Wolfenstein, Phys. Rev. Lett. 51, 1945 (1983).
- [3] A. Ceccucci, Z. Ligeti and Y. Sakai, The CKM quark-mixing matrix, PDG(2012).
- [4] F. Ambrosino et al., Phys. Lett. B632, 76 (2006) [hep-ex/0509045].
- [5] E. Blucher and W. J. Marciano, $V_u d$, $V_u s$, the Cabibbo Angle and CKM Unitarity, PDG(2012).
- [6] A. Abulencia et al., Phys. Rev. Lett. 97, 242003 (2006) [hep-ex/0609040].
- [7] Aaij et al., Phys. Lett. B709, 177 (2012) [arXiv:1112.4311].
- [8] M. Misiak et al., Phys. Rev. Lett. 98, 022002 (2007) [hep-ph/0609232].
- [9] R. Aaij et al. [LHCb Collab.], arXiv:1203.4493; S. Chatrchyan et al. [CMS Collab.], arXiv:1203.3976; G. Aad et al. [ATLAS Collab.], arXiv:1204.0735.
- [10] A. J. Buras et al., Phys. Rev. Lett. 95, 261805 (2005) [hep-ph/0508165].
- [11] M. Battaglia et al., The CKM matrix and the unitarity triangle [hep-ph/0304132].
- [12] J. Brod, M. Gorbahn, and E. Stamou, Phys. Rev. D83, 034030 (2011).
- [13] M. Misiak and J. Urban, Phys. Lett. B451 161(1999) [hep-ph/9901278].
- [14] G. Buchalla and A. J. Buras, Nucl. Phys. B548 309 (1999) [hep-ph/9901288].
- [15] A. J. Buras, M. Gorbahn, U. Haisch and U. Nierste, Phys. Rev. Lett. 95, 261805 (2005) [hep-ph/0508165].
- [16] G. Buchalla and A. J. Buras, Nucl. Phys. B412 106 (1994).
- [17] G. Isidori et al., Nucl. Phys. B718 319 (2005).
- [18] G. Buchalla and A. J. Buras, Nucl. Phys. B398 285 (1993); 400 225 (1993).

- [19] J. Beringer et al. [Particle Data Group], PR D86, 010001 (2012) (<http://pdg.lbl.gov>)
- [20] F. Mescia and C. Smith, Phys. Rev. D76, 034017 (2007).
- [21] U. Camerini, et al., Phys. Rev. Lett. 23, 326 (1969).
- [22] G. D. Cable et al., Phys. Rev. D 8, 3807 (1973).
- [23] Y. Asano et al., Phys. Lett. B 107, 159 (1981).
- [24] S. Adler et al., Phys. Rev. Lett. 88, 041803 (2002).
- [25] S. Adler et al., Phys. Rev. Lett. 79, 2204 (1997).
- [26] V. V. Anisimovsky et al., Phys. Rev. Lett. 93, 031801 (2004).
- [27] A. V. Artamonov et al., Phys. Rev. D 79, 092004 (2009).
- [28] Cirigliano, Vincenzo and Rosell, Ignasi, Two-Loop Effective Theory Analysis of $\pi(K) \rightarrow e\bar{\nu}_e\gamma$ Branching Ratios, Phys. Rev. Lett. 99 (2007), 10.1103/PhysRevLett.99.231801
- [29] Masiero, A. and Paradisi, P. and Petronzio, R., Probing new physics through μ e universality in $K \rightarrow l \nu$, Phys. Rev. D74, 2006, 10.1103/PhysRevD.74.011701, hep-ph/0511289
- [30] F. Ambrosino et al., Precise measurement of $\Gamma(K \rightarrow e \nu (\gamma))/\Gamma(K \rightarrow \mu \nu (\gamma))$ and study of $K \rightarrow e \nu \gamma$, The European Physics Journal C, 2009, 10.1140/epjc/s10052-009-1177-x
- [31] C. Lazzeroni et al., Precision measurement of the ratio of the charged kaon leptonic decay rates, Physics Letters B 719 (2013) 326–336, 2013, 10.1016/j.physletb.2013.01.037
- [32] Battaglia M., Buras A.J., Gambino P. and A. Stocchi (eds.), CERN report CERN 2003-002-corr [hep-ph/0304132]
- [33] Buras A.J., Schwab F. and Uhling S., arXiv:hep-ph/0405132
- [34] Ceccucci A. et al., CERN-SPSC-2005-013,SPSC-P-326
- [35] E. Cortina et al., NA62: Technical design document, NA62-10-07, CERN, (2010), <http://na62.web.cern.ch/na62/Documents/TechnicalDesign.html>.
- [36] G. Ruggiero, "NA62 Physics Working Group", NA62 Collaboration Meeting, (07-09-2011)
- [37] H. W. Atherton et al., Precise Measurements of Particle Production by 400 GeV/c Protons on Beryllium Targets, CERN Yellow Report: CERN 80-07 (1980).

- [38] K. Ahmet et al., Nucl. Instrum. Meth. A **305**, 275 (1991).
- [39] B. Angelucci et al., Pion-Muon separation with a RICH prototype for the NA62 experiment, Nucl. Instrum. Meth. A **621**, pp. 205-211 (2010).
- [40] G. Barr et al., Nucl. Instrum. Meth. A **370**, 413 (1996).
- [41] V. Fanti et al., Proposal for a Precision Measurement of ϵ'/ϵ in CP violating $K^0 \rightarrow 2\pi$ decays CERN-SPSC-90-22 (1990).
- [42] Ruggiero G. Measurement of the LKr calorimeter efficiency using NA48/2 data, Nucl. Instrum. Meth., **A621**:205-211,2010
- [43] G. S. Atoian et al., Nucl. Instrum. Meth. A **531**, 467 (2004).
- [44] Sozzi, M. A concept for the NA62 Trigger and Data Acquisition. NA62 Note 07-03 (2007).
- [45] [<http://ttc.web.cern.ch/TTC/intro.html>], TTC: Timing, Trigger and Control Systems for LHC Detectors.
- [46] A. Pla-Dalmau, A. Bross and V. Rykalin, Extruding plastic scintillator at Fermilab. Prepared for IEEE Nuclear Science Symposium Conference Record (2003)
- [47] Beznosko D. et al., *Nuclear Science Symposium conference record IEEE vol.2* (2004) 790-793
- [48] MINERVA Collaboration, *MINERVA document 218-v4* (2006)
- [49] D0 Collaboration, *Nucl. Instrum. Methods* **A565** (2006) 413
- [50] <http://www.detectors.saint-gobain.com/fibers.aspx>, Saint-Gobain Crystals
- [51] Ronchetti (on behalf of the ALICE collaboration), *Journal of Physics Conferences Series* 160012012 (2009)
- [52] Buzhan P. et al., *Nucl. Instrum. Methods* **A504** (2003) 48
- [53] Golovin V. and Saveliev V., *Nucl. Instrum. Methods* **A518** (2004) 560
- [54] Sadygov Z. et al., *Nucl. Instrum. Methods* **A504** (2003) 301
- [55] Angelone M. et al., arXiv:1002.3480
- [56] Vito Palladino, Simulation, realization and test of veto system for the NA62 experiment, Univerista degli Studi di Napoli "Federico II" Ph. D. Thesis, 2010
- [57] A. Antonelli et. al, arXiv:1111.5768
- [58] J.K. Ahn et al., Search for the Decay $K_L^0 \rightarrow \pi^0 \nu \bar{\nu}$, Phys. Rev. Lett. **100**, 201802 (2008)

- [59] CERN toolkit, GEANT4, <http://geant4.cern.ch>
- [60] Domenico Di Filippo, Studio e simulazione del sistema di rivelazione di fotoni a grande angolo per l'esperimento NA62 al CERN, Univerista degli Studi di Napoli "Federico II" Master Thesis, 2009
- [61] Marco Mirra, Simulation and test of the CHANTI detector for the measurement of the ultra-rare branching ratio $K^+ \rightarrow \pi^+ \nu \bar{\nu}$ with the NA62 experiment at CERN, Univerista degli Studi di Napoli "Federico II" Master Thesis, 2013
- [62] D. Renker and E.Lorenz, Advances in solid state photon detectors, IOP Science 2009 JINST 4 P04004, <http://iopscience.iop.org/1748-0221/4/04/P04004>
- [63] R.J. McIntyre, A new look at impact ionization-Part I: A theory of gain, noise, breakdown probability, and frequency response, IEEE Trans. Electron. Dev. 46 (1999) 1623.
- [64] C.A. Lee et al., Ionization Rates of Holes and Electrons in Silicon, Phys. Rev. 134 (1964) A761.
- [65] A. Lacaita et al., On the bremsstrahlung origin of hot-carrier-induced photons in silicon devices, IEEE Trans. Electron. Dev. 40 (1993) 577.
- [66] Dieter Renker, Properties of avalanche photodiodes for applications in high energy physics, astrophysics and medical imaging, Nuclear Instruments and Methods in Physics Research A 486 (2002) 164–169.
- [67] P. Buzhan et. al., Silicon photomultiplier and its possible applications, Nucl. Instrum. Meth. A 504 (2003) 48.
- [68] Th. Kraehenbuehl, G-APD arrays and their use in axial PET modules, Diploma thesis, ETH ZueRICH (2008).
- [69] H. Oide, H. Otonoa et al., Study of afterpulsing of MPPC with waveform analysis, International workshop on new photon-detectors PD07 June 27-29 2007.
- [70] T. Matsubara et al., Radiation damage of MPPC by gamma-ray irradiation with Co 60, PoS(PD07)032.
- [71] I. Nakamura, Radiation Damage of Pixelated Photon Detector by Neutron Irradiation, NDIP08, to be published in Nucl. Instrum. Meth. A.
- [72] T. Matsumura et al., Radiation damage to MPPCs by irradiation with protons, PoS(PD07)033.
- [73] Y. Musienko et. al., Radiation damage studies of multipixel Geiger-mode avalanche photodiodes, Nucl. Instrum. Meth. A 581 (2007) 433.

- [74] O. Mineev et. al., Scintillator counters with multi-pixel avalanche photodiode readout for the ND280 detector of the T2K experiment, Nucl. Instrum. Meth. A 577 (2007) 540.