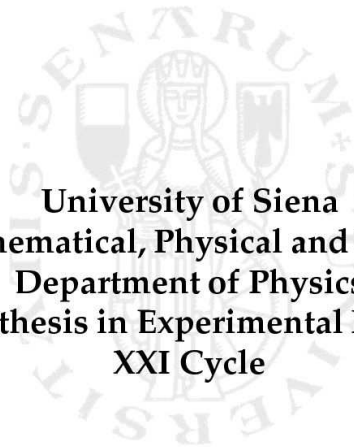




University of Siena
Faculty of Mathematical, Physical and Natural Sciences
Department of Physics
Ph.D thesis in Experimental Physics
XXI Cycle

The forward inelastic telescope T2 for the TOTEM experiment at the LHC.

PhD Thesis in Experimental Physics

CERN-THESIS-2010-178
2010

Thesis Advisor
Dr. Lami Stefano

Tutor
Dr. Turini Nicola

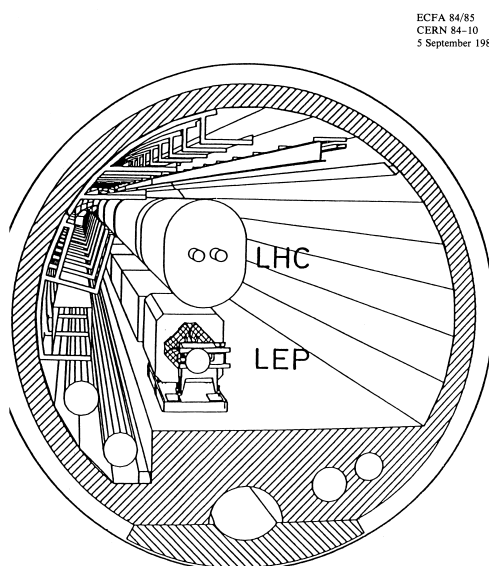
Candidate
Dr. Oliveri Eraldo

Abstract: The TOTEM Experiment will measure the total pp cross-section with the luminosity-independent method and study elastic and diffractive scattering at the LHC. To achieve optimum forward coverage for charged particles emitted by the pp collisions in the interaction point IP5, two tracking telescopes, T1 and T2, will be installed on each side in the pseudorapidity region $3.1 \leq |\eta| \leq 6.5$, and Roman Pot stations will be placed at distances of 147 m and 220 m from the Interaction Point 5 (IP5). The triple Gas Electron Multiplier (GEM) technology has been chosen by the collaboration for its T2 telescope which will provide charged track reconstruction in the rapidity range $5.3 < |\eta| < 6.5$ and a fully inclusive trigger for diffractive events. Results from the preliminary tests on the prototypes of the TOTEM Triple GEM detectors up to the data taking during the first pp collisions at the LHC will be described in this thesis.

Ad Eva

Preface

En mars 1984, un atelier tenu à Lausanne est l'occasion pour la communauté de la physique des particules de réfléchir à l'étape faisant suite à la construction et à l'exploitation du LEP. C'est là qu'on parle pour la première fois de construire un collisionneur de hadrons dans le tunnel du LEP. [CERN Courier Sep 19, 2008]



ECFA 84/85
CERN 84-10
5 September 1984

LARGE HADRON COLLIDER IN THE LEP TUNNEL

Vol. I

PROCEEDINGS OF THE ECFA-CERN WORKSHOP

held at Lausanne and Geneva,
21-27 March 1984

Figure 1: Proceeding of the March 1984 ECFA CERN Workshop. In fact, according to the LHC project leader, Lyn Evans: *"It is generally considered that the starting point for the Large Hadron Collider (LHC) was an ECFA meeting in Lausanne in March 1984, although many of us had begun work on the design of the machine in 1981"*. CERN Courier Jan 27, 2004

In spring 1992, a meeting in Evian, France, marked the beginning for the LHC experiments where the "expressions of interest" were presented to the community. Among the experiments, a group of physicists expressed the interest to perform measurements of the total cross section, the elastic scattering, the diffraction dissociation and the very forward particle production at LHC (fig.2). The aim was to obtain accurate information on the basic properties of proton-proton collision at the maximum accelerator energy. The experiment took the name TOTEM.

Expression of Interest Total Cross Section, Elastic Scattering and Diffraction Dissociation at LHC

2 March 1992

M.Bozzi, Dipartimento di Fisica dell'Università and INFN, Genova, ITALY

J.Bourret and M.Haguenauer, Ecole Polytechnique, Palaiseau, FRANCE

G.Sanguinetti, Sezione INFN Pisa, Pisa, ITALY

G.Matthiae, Dipartimento di Fisica dell'Università Roma II and INFN, Roma, ITALY

J.Velasco, IFIC, Centro Mixto Universidad de Valencia-CSIC, Valencia, SPAIN

We wish to express our interest to perform the following measurements at LHC.

Elastic Scattering. Elastic scattering events will be observed with high-resolution detectors placed inside "Roman pots" in an insertion having high-value of β , the betatron function at the crossing point. This technique was pioneered at the CERN ISR many years ago and then successfully employed at the SPS Collider by experiment UA4 and at Fermilab by experiment E710. The extrapolation to the LHC from present Collider energies was already discussed in some detail in the LHC Workshops [1,2].

We plan to make use of a "mini" version of the standard "Roman pots" with scintillating fibers, silk on pixels or gaseous microstrips. These detectors have already sufficient spatial resolution. The choice will be imposed by the requirement of radiation hardness and operational reliability. One has to keep in mind that the detectors for this experiment must be operated very close (distance of the order of 1 mm) to an intense proton beam.

The basic requirements on the high- β insertion were defined at the Arden Workshop [2] and recently discussed at the Experimental Requirement Committee [3]. We expect that an insertion with $\beta = 1 \pm 2$ km should be suitable to detect elastic scattering down to a momentum transfer squared $-t = 5 \cdot 10^{-4} \text{ GeV}^2$. A first study of a high- β insertion with $\beta = 50$ m was presented at the Arden Workshop [4]. In order to measure elastic scattering in the Coulomb interference region, an insertion with higher β is needed. The β of the insertion should be tunable, the high- β mode allowing detection of b- $\bar{b}$ events with low luminosity and the medium or low- β mode of large-t events

with high luminosity.

Total Cross Section. The total cross section will be obtained with the so called luminosity independent method, i.e. by a simultaneous measurement of elastic scattering at low $-t$ and of the rate of the inelastic interactions. With minimum momentum transfer of $5 \cdot 10^{-4} \text{ GeV}^2$, assuming a forward slope $b \approx 20 \text{ GeV}^{-2}$, the extrapolation to the optical point will be of only 10%. As discussed in ref. 2, the measurement of the inelastic rate should be done with a vertex detector which observes charged tracks and reconstructs the interaction point. Extrapolating from UA4 [5] and E710 [6] to LHC, one finds that the coverage of the vertex detector should be at least three pseudorapidity units (from 6 to 9), corresponding to production angles in the range from 0.25 mrad up to about 6 mrad.

In practice the design of the detectors to measure the inelastic rate will follow the final design of the insertion itself.

The error on the measurement of σ_{tot} could be of 2-3%. Because the LHC beams cross each other at an angle, the feasibility of implementing the Van der Meer method to measure the luminosity as at the ISR, is worth studying. The measurement of σ_{tot} would be considerably simplified.

Diffraction Dissociation. We plan to study the process $pp \rightarrow pX$ by detecting the proton scattered quasi-elastically with the "Roman pots" and measuring the decay products of the system X in a vertex detector. The mass distribution of X will be measured together with the pseudorapidity distribution of its decay products.

Diffraction dissociation was extensively studied at the ISR and at the SPS Collider by UA4 and UA8. It has a substantial cross section, $\sigma_{\text{DD}} \approx 0.2 \sigma_{\text{tot}}$, but the basic mechanism and its connection with elastic scattering still remain rather obscure today.

Very Forward Particle Production. The centre-of-mass energy of the LHC is equivalent to $1.2 \times 10^{11} \text{ eV}$ in incident $p-p$ interaction, which corresponds to the high energy tail of the cosmic ray spectra. The cosmic ray experiments are mostly sensitive to particles produced very forward. As shown by UA7 [7], the measurement in an accelerator experiment of the particles produced at $\eta \approx \eta_{\text{max}}$ is useful to fix the energy scale for the cosmic ray experiments allowing a much better understanding of the production mechanism. In addition the amount of scaling violation in the rapidly distribution in the forward region can be measured. We then propose to install in the "Roman pots" at the LHC small calorimeters to measure particle production, in particular photons and electrons at very small angles.

References

1. M. Haguenauer and G. Matthiae, Proceedings of the ECFA-CERN Workshop on the LHC, Hadron Collider in the LEP Tunnel, Lausanne 1984, Vol. 1, pag.303.
2. G. Matthiae, Proceedings of the Large Hadron Collider Workshop, Arden 1990, Vol.2, pag.100.
3. Experimental Requirement Committee, September 16th, 1991.
4. W. Scandale, private communication.
5. M. Bozzi et al. Phys. Lett. 147B (1984) 392.
6. N. A. Amos et al. Phys. Lett. 243B (1990) 158.
7. E. Pare et al., Phys. Lett. 242B (1990) 531.

Figure 2: TOTEM Expression of Interest. 2 March 1992, Proceedings of the Evian Meeting (1992)

The following year LHC moved his first formal steps inside the CERN Council.

CERN Press Release 17.12.1993: Large Hadron Collider is presented to CERN Council

The CERN Council, where the representatives of the 19 Member States of the Organization decide on scientific programmes and financial resources, held its 98th session on 17 December under the chairmanship of Sir William Mitchell (UK).

Large Hadron Collider In December 1991 CERN's Council delegates agreed unanimously that the Large Hadron Collider (LHC) was the right machine for further significant advance in the field of high energy physics research and for the future of CERN and asked the CERN management to prepare a full technical, scientific and financial proposal for the accelerator for December 1993. Accordingly Prof. Christopher Llewellyn Smith, Director General designate, presented to Council a complete outline of the LHC project.

The Large Hadron Collider (LHC) is an accelerator which will bring protons into head-on collision at higher energies (14 TeV) than ever achieved before to allow scientists to penetrate still further into the structure of matter and recreate the conditions prevailing in the Universe just 10^{-12} seconds after the "Big Bang" when the temperature was 10^{16} degrees. The accelerator will produce not only higher energy but also a higher luminosity - the probability of collision between particles - than has been achieved before and it will reveal the behaviour of fundamental particles of matter which has never been studied.

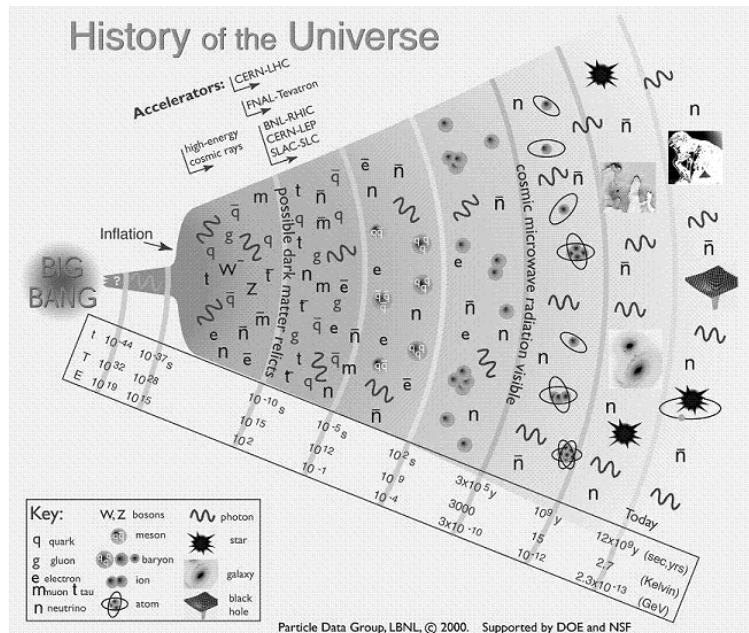


Figure 3: Evolution of the universe.

Let me conclude this preface with September 2008, a very important month for this story and what will follow.

CERN Press Release 10.09.2008: First beam in the LHC - accelerating science

Geneva, 10 September 2008. The first beam in the Large Hadron Collider at CERN was successfully steered around the full 27 kilometers of the worlds most powerful particle accelerator at 10h28 this morning.

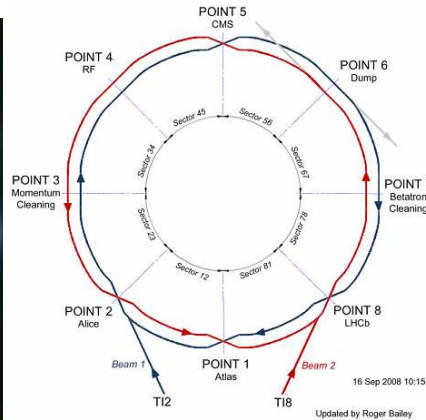


Figure 8: LHC First Circulating Beam. Two Beam spots appeared on the screen of the CERN Control center (CCC), indicating that the beam 1 had completed the full circle, from injection at Point 2 round to the same point.

Few days before, the Technical Coordinator of the TOTEM experiment sent to the collaboration the latest news on the installation of the TOTEM detectors.

—Original Message—
From: Ernst Radermacher
Sent: gioved 14 agosto 2008 8.37
To: TOTEM All

Subject: Installation of second Roman Pot
Dear colleagues,
Yesterday the second Roman Pot has been successfully installed in sector 5-6. Therefore, we met the minimum goal to install 2 Roman Pots before the closure of LHC. Let's hope that we can install 2 more pots during the running-in phase of the LHC.
We congratulate all people which where involved in this critical task.
Regards
Radi

—Original Message—
From: Ernst Radermacher
Sent: mercoled 3 settembre 2008 11.45
To: TOTEM All

Subject: T2 installed and CMS closed
Dear colleagues,
Yesterday evening CMS was closed on the - side where with a lot of effort a quarter of T2 was installed and even a water leak finally repaired in the last minute. Thanks go to all people which participated in the effort to get the T2 detector assembled, tested and installed.
Regards
Radi

Figure 9: Mails from the TOTEM's technical coordinator to the collaboration, few days before the first circulating beam in LHC, about the detectors installation.

The preparation of the LHC machine and of the experiments is almost ready to bring protons to the stage. Unfortunately, after the "first circulating beams", some machine problems have postponed the collisions in 2009. TOTEM needs anyway more time to complete its telescopes and this additional time will be used for this purpose. Finally proton collisions will come and they will hopefully provide us with good data for interesting and maybe unforeseen discoveries.

Contents

Introduction	1
1 Total Cross Section, Elastic Scattering and Diffraction Dissociation Measurements at the LHC.	5
1.1 Physics Objectives [3] [10]	5
1.1.1 The Total Cross Section	5
1.1.2 Elastic Scattering	7
1.1.3 Diffraction Dissociation and very forward particles production [24].	11
1.2 Technical Aspects	14
1.2.1 The inelastic detectors	15
1.2.2 The "leading" protons detectors	15
1.2.3 Elastic Measurement	16
1.2.4 Inelastic Measurement.	22
1.2.5 Diffractive Measurement.	27
2 The forward inelastic telescope T2	29
2.1 Telescope requirement	29
2.2 Triple GEM Detector	33
2.2.1 The TOTEM Triple GEM	35
2.3 On detector electronics	39
2.4 TOTEM electronics overview	45
2.4.1 T2 Readout and Control System	47
3 TOTEM Triple GEM and Simulation	51
3.1 Detector Simulation	51
3.1.1 The simulation Tools.	51
3.1.2 The TOTEM Triple GEM Simulation	52
3.2 Digitization	60
3.2.1 The Strip Case	61
3.2.2 The Pad Case	68
3.2.3 Strips and pads Digitization algorithms.	74
3.2.4 Strips and pads integration.	76
3.2.5 Digitization and Measurements	77
4 No Beam Detector Test and Front End Electronics Integration	79
4.1 Detector Test	79
4.1.1 Preliminary Tests	79
4.1.2 Absolute Gain Calibration	81
4.1.3 Energy Resolution Studies	85
4.1.4 Charge Sharing	87
4.1.5 Uniformity, Stability and Charging-Up	87
4.2 Front End Integration	90
4.2.1 Noise Measurement Tools	90
4.2.2 Noise Measurement Results	94

5	Test with particles	105
5.1	Test Beam Data	105
5.1.1	Timing Measurements	105
5.1.2	Latency Scan	111
5.1.3	Preliminary Efficiency Measurements	113
5.1.4	Cluster Size	115
5.1.5	Spatial Resolution	115
5.1.6	LV1 induced noise	116
5.2	Cosmic Ray Data	117
5.2.1	Data	117
5.2.2	Simulation	121
5.3	LHC pp Collisions Data	124
	Conclusions	135
	Acknowledgments	139
	Appendix Chapter 1	141
.1	Total Cross Section	141
.1.1	From the Optical Theorem to the Luminosity-Independent Total Cross Section Measurement [14]	141
.1.2	Dispersion Relations [10]	142
.1.3	Accelerator physics of colliders [27]	143
.1.4	Elastic Scattering: Optics Condition [5]	143
	Bibliography	i

List of Figures

1	Proceeding of the March 1984 ECFA CERN Workshop.	I
2	TOTEM Expression of Interest (1992).	II
3	Evolution of the universe.	II
4	TOTEM Letter of Intent(1997).	III
5	TOTEM Technical Proposal(1999).	III
6	TOTEM Technical Design Report(2004).	III
7	LHC Article(2007)	III
8	LHC First Circulating Beam	IV
9	TOTEM Status Before LHC starts.	IV
10	Structure within the atom and their dimension.	1
11	Energy time line from the Big Bang to the present.	3
1.1	The total cross section for $\bar{p}p$ and pp scattering.	6
1.2	The ratio of the elastic to the total cross section for $p\bar{p}$ interactions as a function of energy.	7
1.3	Graphical illustration of the "expanding protons" in the <i>Impact Picture</i> of Cheng and Wu [16].	7
1.4	The ratio ρ of the real to the imaginary part of the forward amplitude for $\bar{p}p$ and pp scattering.	8
1.5	Behaviour of the ρ parameter for different assumptions on the energy dependence of the total cross section [11]	8
1.6	Proton-proton elastic scattering as measured at the CERN ISR.	9
1.7	The forward slope parameter B for the $\bar{p}p$ and pp elastic scattering.	9
1.8	The local slope parameter $B(t)$ for the $\bar{p}p$ and pp scattering at different energies.	9
1.9	Dip-shoulder region in the differential elastic cross section of $\bar{p}p$ and pp	10
1.10	Donnachie and Landshoff three gluons exchange process.	10
1.11	pp elastic scattering at large t. Three Gluons Exchange.	11
1.12	pp elastic scattering at large t. <i>Impact Picture</i> and Regge model.	11
1.13	Rapidity Structure of single diffraction dissociation events.	12
1.14	Average charge multiplicity of the diffractive cluster.	12
1.15	Diffractive process classes [7].	14
1.16	CMS-TOTEM coverage in the η - ϕ plane, charged particle multiplicity and energy flow for inelastic events at $\sqrt{s} = 14TeV$	15
1.17	View of the inelastic forward trackers T1 and T2 inside the CMS detector.	16
1.18	The forward Inelastic Telescope T1	16
1.19	The forward Inelastic Telescope T2	17
1.20	The LHC beam line and the Roman pots at 147m and 220m.	17
1.21	The forward Inelastic Telescope T1	18
1.22	Scheme for a complete study of elastic scattering at the LHC [3].	19
1.23	Optics for elastic scattering: the magnification v and the effective length $L^{eff}(m)$	20
1.24	Acceptance for elastically scattered protons at the Roman Pot station at 220 m.	21
1.25	Differential cross-section of elastic scattering at $14TeV$ as predicted by various models.	21
1.26	Elastic scattering cross section measurement with TOTEM up to large $ t $	22
1.27	Differential cross-section of elastic scattering at $\sqrt{s} = 14TeV$ as predicted by various models [25].	22
1.28	The t-distribution of $\bar{p}p$ elastic scattering at $\sqrt{s} = 540GeV$ in the region of the Coulomb interference [26].	23
1.29	$-t$ corresponding to an acceptance of 50% versus the distance to the beam center.	23

1.30	The p_T distributions for NSD (left) and SD (right) events, as generated by Pythia.	24
1.31	The η distributions for NSD (left) and SD (right) events, as generated by Pythia.	24
1.32	The T1(left) and T2(right) track multiplicities for 1000 NSD events.	25
1.33	Results of the simulation on the efficiency of the inclusive trigger.	25
1.34	Ratio of detected Single Diffractive events as a function of the diffractive mass $M_{\perp}$ and Single Diffractive distribution as function of $1/M^2$	26
1.35	Single Diffraction and Double Pomeron Exchange process with respect to the TOTEM detectors coverage.	27
2.1	Two quarters of the forward inelastic telescope $T2$ installed on the Plus End of the Interaction Point $IP5$ of the LHC	29
2.2	Telescope assembly of the forward inelastic telescope $T2$	30
2.3	Readout pattern of the forward inelastic telescope $T2$	31
2.4	Trigger pattern of the forward inelastic telescope $T2$	31
2.5	Average particle fluxes and radiation dose in the region of the forward inelastic telescope $T2$	32
2.6	Flux of charged particles through $T2$ at $L = 10^{28} cm^{-2} s^{-1}$	32
2.7	Measurement of the discharge probability as a function of gain.	33
2.8	The GEM foil.	33
2.9	The GEM multiplication process.	34
2.10	Real gain of a GEM foil: external and internal field effects.	34
2.11	Triple GEM.	35
2.12	TOTEM triple GEM for the $T2$ telescope.	36
2.13	Drift and GEM foils of the TOTEM Triple GEM.	36
2.14	Readout plane of the TOTEM Triple GEM.	37
2.15	H.V. distribution of the TOTEM triple GEM.	38
2.16	Block diagram of the VFAT2 chip [7].	40
2.17	The VFAT2 readout chip.	40
2.18	Triple GEM strip and pad capacitance.	41
2.19	Triple GEM strip and pad capacitance.	42
2.20	VFAT2 Input Impedance [34].	42
2.21	VFAT2 input circuit.	43
2.22	VFAT2-GEM.	43
2.23	Time Walk and Jitter effects.	44
2.24	VFAT2-monostable.	44
2.25	VFAT2 preamplifier output simulated for ideal and GEM signal.	46
2.26	TOTEM electronics functional overview.	46
2.27	11th Card.	47
2.28	HORSESHOE card.	48
2.29	T2-electronics overview.	49
3.1	The signal formation in the internal structure of the TOTEM Triple GEM.	52
3.2	Simulation Scheme.	53
3.3	First Ionization: Gas properties Table.	54
3.4	Ar/CO_2 drift velocity and diffusion coefficients as a function of the electric field.	54
3.5	Ar/CO_2 Townsend and attachment coefficients.	54
3.6	First Ionization: Cluster formation and released energy.	55
3.7	GEM foil: Meshing and fields plot.	56
3.8	GEM foil: 3D hole's field intensity map.	56
3.9	GEM foil: Multiplication Process.	57
3.10	Electrons transport properties in Ar/CO_2 70/30 at stp.	57
3.11	Transport Properties: spatial and timing distribution.	57
3.12	First Ionization: spatial and timing properties of the drifted first ionization cluster.	58
3.13	First Ionization: First GEM foil hole occupancy.	58
3.14	Transport Properties: Hole occupancy of the third GEM foil.	58
3.15	Signal Induction: pad weighting field.	59
3.16	Signal Induction: Direct and Cross signal on pad.	60

3.17	Signal Induction: Pads and Strips	60
3.18	Geometrical superimposition between the electron cloud and the strip readout pattern.	62
3.19	Charge collected by one strip in function of the distance from the center of the electron cloud.	63
3.20	Strips and pads absolute gain and charge sharing.	63
3.21	Electrical Field Lines and equipotential curves in proximity of a strip.	65
3.22	Cross Induced Signal Area.	65
3.23	210 μm strips effective width.	66
3.24	Polynomial 9th order n functions obtained for various σ_{ch} of the Gaussian electron cloud obtained for an effective strip width of 210 μm	67
3.25	$SC(\sigma_{ch})/a_0$ and $ST(\sigma_{ch})$ needed to transform $n_0(d, \sigma_0)$ in $n(d, \sigma_{ch})$ with eq. 3.9 for an effective strip width of 210 μm	67
3.26	Scaled and stretched <i>electron collection curve</i> results.	68
3.27	Track of an ionizing particle and relative charge collection curve.	69
3.28	Geometrical computation of the electrons collected by one pad in differen position with respect to the clouds.	70
3.29	Collected charge versus the distance between the center of the electron gaussian cloud and the center of a pad.	70
3.30	Left: Fit of the charge collection curve with the eq. 3.17. Right: Residuals	71
3.31	Behavior of the four parameter $m1, m2, m3, m4$ of eq.3.17 obtained fitting various electron cloud with different sigma.	72
3.32	Power fit of the parameter $m3$ in function of σ_{ch}	72
3.33	Test Distribution of d_x, d_x and σ for 1K events	73
3.34	Difference between the counted charge collected and the calculated one with the eq.3.22 for 1K events.	73
3.35	MultiTest Distribution of d and σ	74
3.36	Difference[%] between the counted charge collected and the calculated one with the Polynomial approach.	74
3.37	Difference[%] between the counted charge collected and the calculated one with the Erf approach.	75
3.38	Collected charge obtained with the use of the erf function.	76
3.39	Strips and Pads integration scheme	77
4.1	Resistive H.V. Distribution Boards.	80
4.2	Cu X-rays interaction in the drift zone.	82
4.3	Current readout scheme.	83
4.4	Rate and spectrums readout scheme.	83
4.5	X-ray Tube Calibration curve.	84
4.6	Gain Calibration.	84
4.7	Voltages and Fields in the Triple GEM.	84
4.8	Bremsstrahlung.	85
4.9	Resolution as a function of the high voltage applied.	85
4.10	Cu tube X-Ray Strips Spectrum as a function of the H.V.	86
4.11	Strips-pads correlation spectrums.	86
4.12	Strips and pads charge sharing.	87
4.13	Gain measurements.	87
4.14	Cu-tube X-Ray Spectrums.	88
4.15	Stability test.	88
4.16	Charging Up.	89
4.17	The calibration pulse amplitude.	91
4.18	The pulse scan.	91
4.19	Pads Calibration Pulse Scan	91
4.20	Calibration Pulse as a system improvement tool.	92
4.21	Threshold Scan.	92
4.22	Comparison between Calibration Pulse Scan and Threshold Scan results.	93
4.23	Trigger bits as a noise tool.	94
4.24	Digital and Analog Ground connections on the hybrid.	94

4.25	Test of the Digital and Analog Ground connections on the hybrid (stand alone).	95
4.26	Test of the Digital and Analog Ground connections on the hybrid (on Triple GEM).	95
4.27	Metal Cover on the VFAT2 Chip.	95
4.28	Test of the Metal Cover on the VFAT2 Chip.	96
4.29	Low Voltage Power Supply Configuration.	96
4.30	Analog and Digital power supply measurement.	97
4.31	Power Supply Decoupling Board.	97
4.32	Detector Shielding.	98
4.33	Detector Shielding: Calibration Pulse Scan.	98
4.34	Detector Shielding: Top foil effects	99
4.35	Detector Shielding: Direct measurement.	99
4.36	Detector Shielding: Shields Quality	99
4.37	Detector Shielding: Final Configuration.	100
4.38	Cross talk between pads and strips: one extreme case.	101
4.39	Cross talk between Strips	101
4.40	Cross talk between pads and strips.	102
4.41	Cross talk: External and Internal Strips.	102
4.42	Fully Equipped Detector: Threshold Scan	103
4.43	Fully Equipped Detector: Calibration Pulse Scan	103
4.44	Threshold Scan at different stretching of the monostable pulse length on a set of detector ready to be installed.	104
5.1	TOTEM Test Beam Area	106
5.2	VFAT2 monostable block diagram.	106
5.3	VFAT2 Block Diagram.	107
5.4	TDC acquisition system.	107
5.5	<i>Negative and Positive Monostable Input Polarity</i>	108
5.6	Triple GEM rising time distribution.	108
5.7	Rising time.	109
5.8	TDC measurement at -4.2kV.	109
5.9	<i>Monostable Pulse Length</i> effects on the triple GEM efficiency.	110
5.10	Triple GEM falling time distribution.	110
5.11	Falling time.	111
5.12	Time Length of the signal.	111
5.13	Comparison between TDC and Scintillators Coincidence Counts versus the high voltage applied to the Triple GEM.	111
5.14	Latency scan with beam for MSPL=1clk.	112
5.15	TDC and latency scan comparison.	112
5.16	Latency Scan with the Monostable Pulse Length of 1clk.	113
5.17	Latency Scan with the Monostable Pulse Length of 2clk.	113
5.18	Telescope set up for efficiency studies in H8.	114
5.19	Strip clusters radial hit separation between chamber HG08 and HG02.	114
5.20	HG02 Efficiency measurement versus the high voltage applied.	115
5.21	Cluster Size.	115
5.22	Cluster Size versus high voltage and threshold.	116
5.23	Radial distance for strip clusters belonging to aligned hits in two different T2 planes.	116
5.24	LV1 Induced noise.	117
5.25	Cosmic Rays Set Up and Tracks Polar Distribution.	118
5.26	Cosmic Ray Efficiency Measurement.	119
5.27	Cosmic Ray Efficiency Measurement: strips and pads.	119
5.28	Efficiency Measurement of a Triple GEM with the VFAT2 during the 2009 October RD51 Test Beam [41].	120
5.29	Cosmic Rays Efficiency Measurement at different thresholds.	120
5.30	Cosmic Rays Efficiency Measurement at different thresholds and high voltages.	120
5.31	Cosmic Rays Measurement: Pad and strip cluster size.	121
5.32	Cosmic Rays Efficiency: Data and Simulation.	123
5.33	Cosmic Rays Cluster Size: Data and Simulation.	123

5.34	Cosmic Rays efficiency for different hardware setting: Data and Simulation.	124
5.35	IP5: T2 before 2008 LHC run.	125
5.36	IP5: First Forward Region Closure.	125
5.37	IP5: T2 before 2009 LHC run.	126
5.38	LHC Collisions: First T2 Particles signature	127
5.39	LHC Collisions: track multiplicity and η distribution.	127
5.40	LHC Collisions: Primary Vertex XY Reconstruction.	127
5.41	LHC Collisions: Primary Vertex Z Reconstruction.	128
5.42	Position Sensors.	128
5.43	The overlapping region.	129
5.44	TT2 Triggering Pattern	129
5.45	Trigger Sector Acceptance	130
5.46	η Distributions without alignment correction.	130
5.47	η Distributions with alignment correction.	131
5.48	Background	131
5.49	Signal and Background	131
5.50	Signal and detector alignment	132
5.51	T2 Event Display: Goods Events	132
5.52	T2 Event Display: Noise pattern	133
5.53	T2 Event Display: multi hits events	133
5.54	T2 Event Display: Histograms	134

List of Tables

1.1	Events cross section at the LHC with $\sqrt{s} = 14TeV$	24
2.1	Material budget in one TOTEM triple GEM detector.	37
2.2	Overview of electronics requirements from the different detectors.	39
2.3	VFAT2 specification.	41
2.4	VFAT2 Dynamical Range [33].	45
3.1	Charge Sharing between strips and pads.	64

Introduction

Particle Physics studies the constituents of matter and their fundamental interactions, and the progress of knowledge relies on the interplay between theory and experiment. In order to reach smaller and smaller scales of distance associated with the most elementary constituents, experimental research in the last decades has been carried out with high energy particle accelerators and the parallel development of highly sophisticated technologies for the associated detection equipment.

Why High Energy Accelerators and Colliders?

The structure within the atom is shown in fig.10. The size of quarks is smaller than $10^{-18} fm$. In order to have experimental information on these structures, the scientists have to use a microscope with a resolution compatible with these incredibly small spatial dimensions.

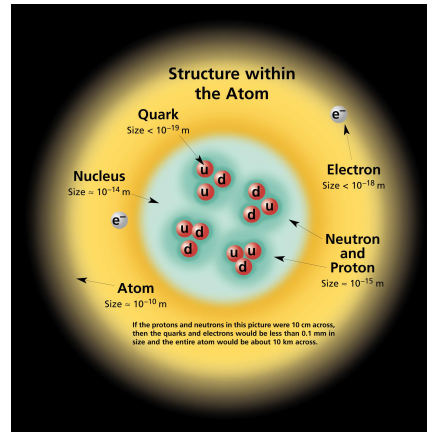


Figure 10: Structure within the atom [<http://cpepweb.org/>].

The minimum separation d that can be resolved by any kind of a microscope is given by eq.1, where n is the refractive index and λ is the wavelength. .

$$d = \frac{\lambda}{2n \sin \theta} \quad (1)$$

In 1924 De Broglie claimed that all matter has a wave like nature. This means that particles can be used as a probing microscope and this is the basis for accelerator physics. Quantitatively, the relation between the momentum and the related wavelength of a particle can be expressed by eq.2.

$$\lambda = \frac{h}{p} = \frac{1.2}{p[GeV/c]} fm \quad (2)$$

In order to probe the system on the quark scale, we then need a particle with a momentum $p \approx 10 TeV/c$ and this is more or less what the LHC will foresee.

This technique, in reality, is more than a standard microscope. It doesn't leave the probed system intact. If the energy is enough, according to the Einstein eq.3 that links mass and energy, we will observe the creation of particles. The types of particles, that can be produced during the collision, depend on the available energy.

$$E = mc^2 \quad (3)$$

This assumption motivates the development of colliders as the LHC, where the interactions happen between two colliding beams (instead of one beam that collide on a fixed target). The available cms energy for a fixed-target machine rises as the square root of the incident energy, while in a collider, with equal energy in the two beams, rises linearly with E. *(The great advantage of particle colliders, where two counter-rotating beams of particles collide in several intersection regions around the ring, is in term of the large center of mass energy available for the creation of new particles. (Perkins)).*

Concerning the detection of these particles, Einstein and relativity help us again. Some of the unstable particles produced will have a very short mean life. If the velocity of a particle is close to c and its mean life is about $10^{-12}s$, as it often happens, the travelling length would only be about 1mm neglecting relativistic effects. This is too small for detection, and it would limit the possibility to study the process with an adequate accuracy. Most of the decay processes will be inside the colliding beams and only the final steps of possible fragmentation will reach the detectors. Moreover, more than one interaction can happen during the collision and then, in the detectors, there will be tracks from all these interactions. This will obviously limit the capability of distinguish between interactions and possible decay processes. Fortunately, according again to Einstein theory on relativity effects on the space-time, the mean life of the particle in the laboratory system, where the detectors are at rest, will be $\tau = \gamma \cdot \tau_0 = (1/\sqrt{1 - \frac{v^2}{c^2}}) \cdot \tau_0$, and then bigger. With the energy available, the time dilatation will be very effective. For example, if a B^0 Meson with a rest mass of $\sim 5GeV$ and a mean life of $\sim 1.5ps$ will be produced with an energy of $500GeV$, it will be able to travel for $\sim 5cm$, enough to be outside the beam-beam interaction zone and to be properly reconstructed by detectors.

Collider physics is moreover very important from the point of view of knowledge on the universe and its history. Increasing the collision energy is equivalent to recreate conditions closer to the Big Bang. According to the theory, a relation exists between the expansion time of the universe after the Big Bang and the instantaneous temperature or density of energy. In the interaction point of a collider experiment, a very high density of energy can be reached. Higher the energy, closer in time to the Big Bang will be the condition that is reproduced. At LHC, with 14TeV in the center of mass system, it will be recreated the condition $\approx 10^{-12}s$ after the Big Bang, with a temperature of $\approx 10^{16}K$. In fig. 11 the time scale of the universe, with the relative temperature and energy scale, is shown together with specific particles and interactions.

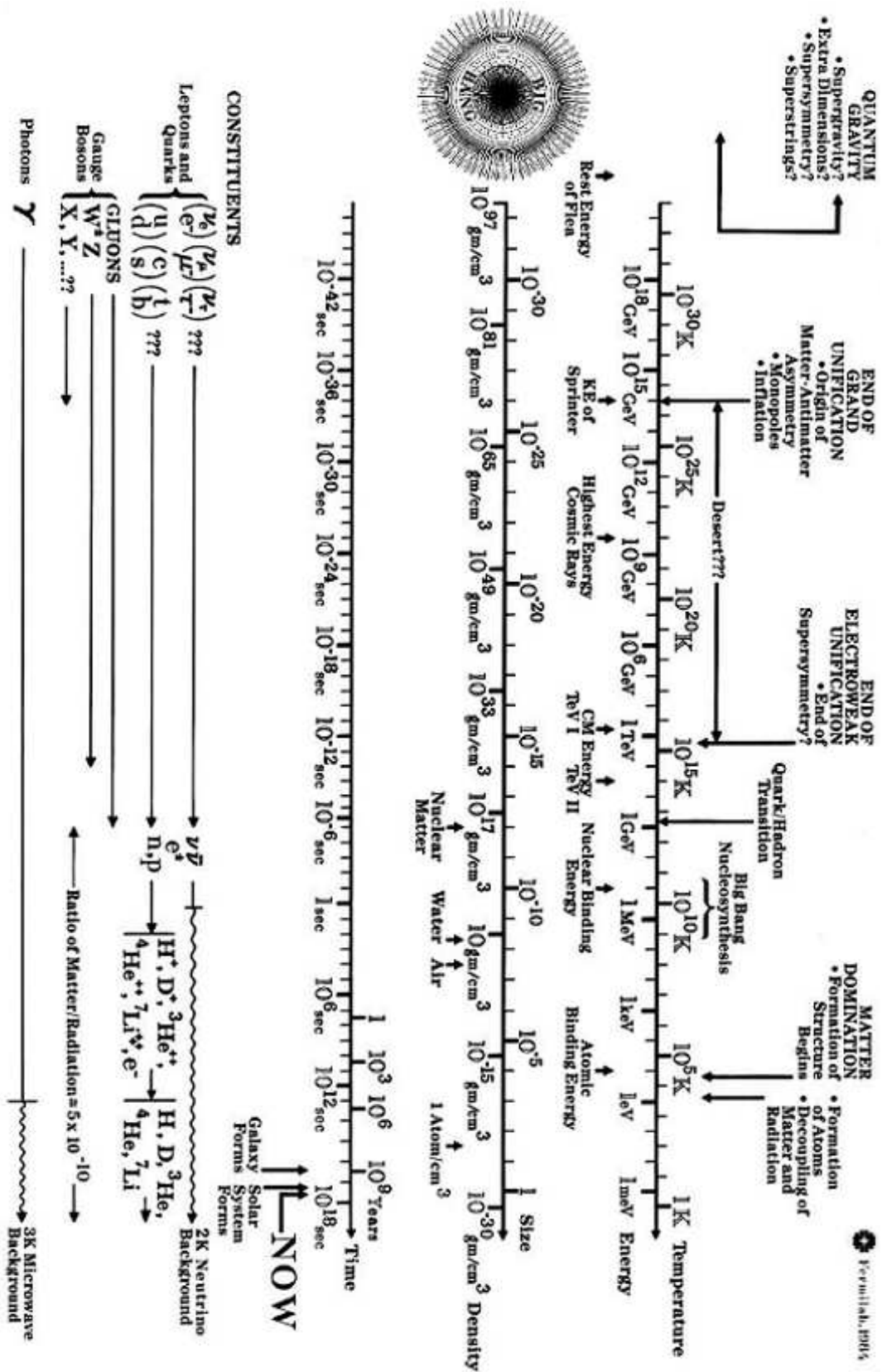


Figure 11: Energy time line from the Big Bang to the present.

Chapter 1

Total Cross Section, Elastic Scattering and Diffraction Dissociation Measurements at the LHC.

1.1 Physics Objectives [3] [10]

The TOTEM experiment was designed to perform TOTAl Cross Section, Elastic Scattering and Diffraction Dissociation Measurements at the LHC. The reasons that motivate the interest on this kind of measurements will be briefly reported in this section.

1.1.1 The Total Cross Section

The total cross section can be thought as the effective area seen by the impinging particles in a scattering process. It is consequently expressed in unit of area. In the field of particles physics, where the typical dimensions are on the scale of the fm ($10^{-15}m$), the *barn* ($1b = 10^{-24}cm^2$) or *millibarn* is normally used. Fig.1.1 shows the total cross section for $\bar{p}p$ and pp scattering, measured by colliders and cosmic ray experiments at various cms energy $\sqrt{s}$ ¹. In the LHC regime we expect a $\sigma_{TOT} \sim 110mb$ that can be roughly associated to a sphere with a radius of $\sim 2fm$. The aim of TOTEM is to measure it.

Two reasons, that motivate the interest on this measurement, are:

The importance of a direct measurement: the total cross section reflects the various interactions between the colliding particles and their constituents. It is useful to have a precise measurement to better understand the mechanism of hadron collision at high energy. It is not possible to describe with calculations ab-initio the behavior of the σ_{TOT} in the pp scattering at the LHC energy but there are various phenomenological models that try to give account of it. The current large theoretical uncertainty on the extrapolation of the proton-proton total cross section at the LHC energy is due to the current lack of a fully satisfactory theoretical explanation of the cross section in low momentum transfer collisions. Therefore their description needs to rely on phenomenological models. TOTEM aims to measure σ_{TOT} with a precision of 1% or $1mb$, finally discriminating among the different models.

The normalization of the single processes: the total cross section is the sum of all the possible interactions that can occur during a collision. Each single process will have its own cross section. TOTEM measures the sum of all of them to get the normalization for all the physics process at LHC, including

¹In a two body scattering $a + b \rightarrow a + b$, defining the 4-momentums of ingoing ($p1, p2$) and outgoing ($p3, p4$) particles, the kinematics can be described using the Lorentz invariant Mandelstam Variables (s, t, u), that are defined as:

$$s = (p1 + p2)^2 = (p3 + p4)^2$$

$$t = (p1 - p3)^2 = (p2 - p4)^2$$

$$u = (p1 - p4)^2 = (p2 - p3)^2$$

In the process $a + b \rightarrow a + b$ then, s represent the square of the cms energy, while t is the 4-momentum transfer squared. In the text s will be then used as a reference for the cms energy and t for the momentum transfer.

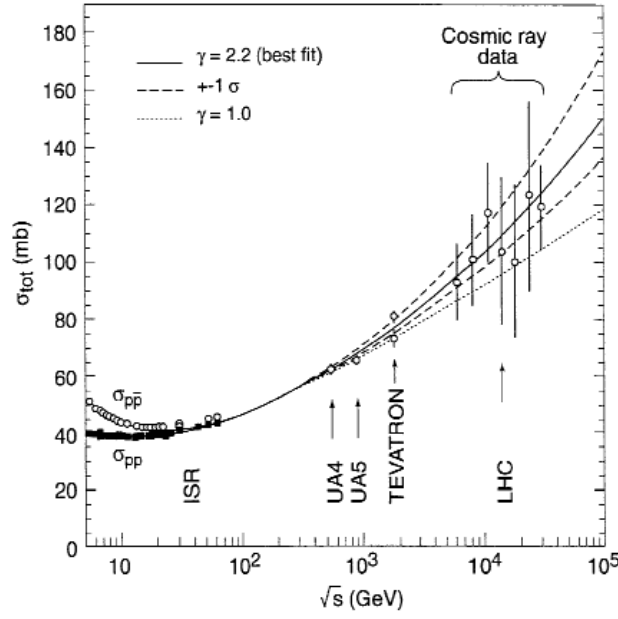


Figure 1.1: The total cross section for $\bar{p}p$ and pp scattering is shown together with the prediction of the dispersions relations fit of ref. [8]. The best fit (solid line) corresponds to $\gamma = 2.2$. The region of uncertainty is delimited by the dashed lines. The result obtained with $\gamma = 1$ is shown as a dotted line. The expected uncertainty of the TOTEM measurement at LHC will appear in this plot as large as the data point itself.

for example the Higgs production. This is mandatory for having precise measurements and not only observation.

The TOTEM collaboration will perform this measurement by using a method that is independent from the luminosity² of the LHC. The method, based on the Optical Theorem (see appendix .1.1 for more details), allows us to define the total cross section as:

$$\sigma_{TOT} = \frac{4\pi}{p} \Im F(t=0) = \frac{16}{(1+\rho^2)} \frac{(dN_{el}/dt)_{t=0}}{N_{el} + N_{inel}} \quad (1.1)$$

where:

- i. t is the 4-momentum transfer squared³.
- ii. p is the momentum in the center of mass and $\Im F(t=0)$ is the imaginary part of the forward amplitude.
- iii. $\rho = \Re F(t=0)/\Im F(t=0)$ is defined as the ratio of the real to imaginary part of the forward amplitude.
- iv. N_{inel} is the inelastic rate, consisting of non-diffractive minimum bias events and diffractive events, measured with an adequate acceptance in the forward region.
- v. N_{el} is the total nuclear elastic rate.

²The luminosity L relates the reaction rate R of a process to its cross section σ and is defined by $R = L \cdot \sigma$ with L normally expressed in units of $cm^{-2}s^{-1}$. Knowing the rate and the luminosity is then possible to obtain the cross section. Equivalently, knowing the rates and the cross section is possible to obtain the luminosity. For the LHC $L = (fk_b N^2)(4\pi\epsilon_n \frac{\beta^*}{\gamma})^{-1} F$ [9]. It depends exclusively on beam parameters: N is the number of protons per bunch, k_b the number of bunches, f the revolution frequency, γ the relativistic factor, ϵ_n the normalized transverse emittance, β^* the value of the betatron function at the IP and F (0.9) the reduction factor caused by the crossing angle. LHC will foresee $10^{34} cm^{-2}s^{-1}$.

³Refers to footnote in sec. 1.1.1

- vi. $(dN_{el}/dt)_{t=0}$ is the nuclear part of the elastic cross-section measured down to the four-momentum transfer of $-t = 10^{-3} GeV^2$ and extrapolated to $t = 0$, i.e. to the *optical point*.

The previous quantities, necessary for the σ_{TOT} measurement, have their own physical interest. Moreover, the experimental configuration needed for the cross section measurements, give us the possibility to exploit other additional measurements.

1.1.2 Elastic Scattering

The hadron-hadron interaction strongly depends on the transfer momentum and on the energy of the collision. This is reflected on the goodness of the predictions, that can be limited only to particular energy-momentum ranges. To understand the various mechanisms involved is useful to exploit the measurements of the total and differential cross section of elastic processes at different transfer momentum and different energies. Particular behaviors and structures of the forward differential cross section has been theorized and found experimentally. Their evolution as a function of energy has also been predicted. The LHC interaction rate will give us enough statistics to describe accurately these behaviors and to be sensitive to small structures.

The ratio of the elastic to the total cross section

The dependence on energy of the elastic to total cross section ratio (fig.1.2) has allowed to discriminate between different phenomenological models. The "geometrical scaling" model for example, used at the time of ISR, predicts a constant value and for this reason has been abandoned when higher energy measurements were performed. The observed increase has been therefore explained with the *Impact Picture* model of Cheng and Wu [16] (fig.1.3). They associated the energy dependence to the increase of the effective "opacity" of the two colliding particles. The measurement at the LHC gives the possibility to confirm the predictions of this model or to show a different behavior.

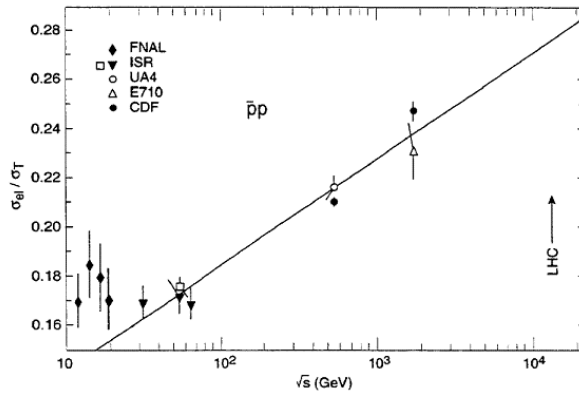


Figure 1.2: The ratio of the elastic to the total cross section for $p\bar{p}$ interactions as a function of energy. The line is a linear extrapolation to guide the eye.

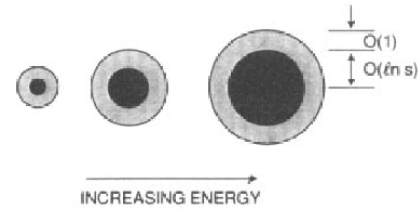


Figure 1.3: Graphical illustration of the "expanding protons" in the *Impact Picture* of Cheng and Wu [16]. They predicted that, at very high energy, the effective radius of interaction of the two colliding hadrons would increase (as $\log s$) and that also the "opacity" would increase. This is in agreement with the measurements of fig.1.2 where the ratio σ_{el}/σ_{TOT} , and then the "opacity", is increasing. The limit predicted by the *Impact Picture* model is $1/2$.

The real part of the amplitude at $t \simeq 0$.

It is interesting to measure the real part of the forward amplitude because it provides information on the energy dependence of the total cross section. It is important to note that this relation is derived using only fundamental principles. According to relativistic QFT (Quantum Field Theory), scattering amplitude must satisfy general properties like analyticity, unitarity, crossing symmetry and polynomial boundedness.

The dispersion relations, that establish the relation between real and imaginary part of the scattering amplitude, derive directly from these properties. A simplified form of dispersion relations, in proton-proton scattering at high energy, is the *Derivative Dispersion Relations*. It links in a straightforward way the parameter ρ , ratio of the real to imaginary part of the forward amplitude, with the dependence of the total cross section (see appendix .1.2 for more details):

$$\rho \approx \frac{\pi}{\sigma_{TOT}} \frac{d\sigma_{TOT}}{d \log s} \quad (1.2)$$

In fig.1.4 is shown the ratio ρ , measured at various cms energy $\sqrt{s}$ and in fig.1.5 the expected [11] behavior for different dependencies of the total cross section on energy.

An independent measurement of ρ and σ_{TOT} gives the possibility to test the validity of eq. 1.2. A failure

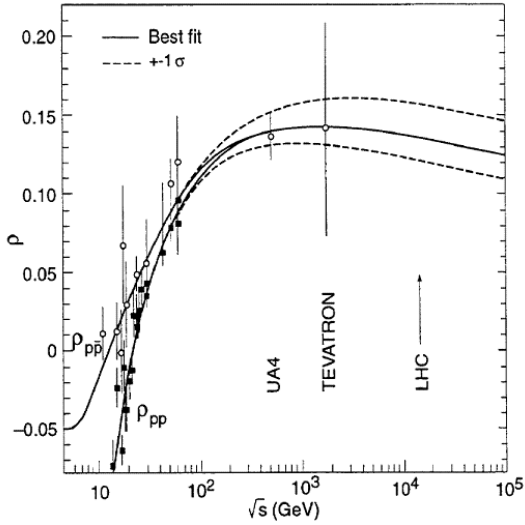


Figure 1.4: The ratio ρ of the real to the imaginary part of the forward amplitude for $\bar{p}p$ and pp scattering. The prediction of the dispersion relations fit of ref. [8] (solid line) is shown together with the region of uncertainty delimited by the dashed lines.

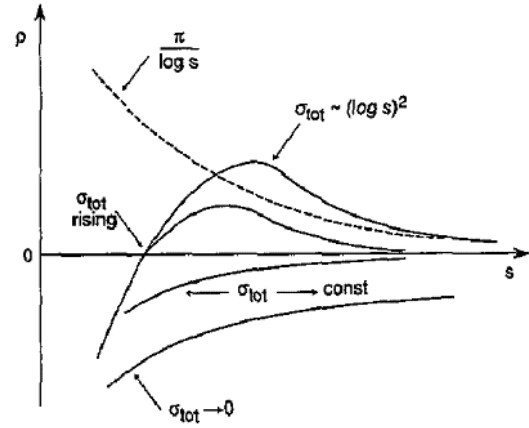


Figure 1.5: Behaviour of the ρ parameter for different assumptions on the energy dependence of the total cross section [11]

on the prediction has to be assumed as a failure of one of the fundamental principles used to obtain the dispersion relations. This is more stringent than a failure of one of the phenomenological models that wants to describe the interactions between the protons during the collision. N.N. Khuri [12] and J.Wess [13] theorized for example the appearance of a fundamental length in the theory that could violate the actual dispersion relations. Khuri proposed the existence of extra compact dimensions for the space-time. Wess proposed a discretisation of space due to quantum groups. From our point of view, it is interesting that the value of ρ , predicted at LHC by these theories, differs substantially from what is expected from the standard dispersion relations. If these fundamental length exist, its effect should be detected.

The forward peak and its slope B .

A measurement of the differential elastic cross section for pp scattering is shown in fig.1.6, where at low t is present the forward or *diffraction peak*. Near the forward direction ($|t| < 0.1 \text{ GeV}^2$) the differential cross section can be described by the simple exponential $e^{-B|t|}$. Measurements of the slope B as a function of energy and transfer momentum can be used to test the prediction of the phenomenological models. The $\log s$ rise of $B(s)$ (fig.1.7) has been predicted for example in the classical Regge model. This simple exponential form is not an accurate description, as can be seen in fig.1.8 where B has been plotted at different cms energies as a function of the squared momentum transfer $-t$. A change of shape from the simple exponential of the forward peak is manifested by a non constant value of the slope. At the ISR

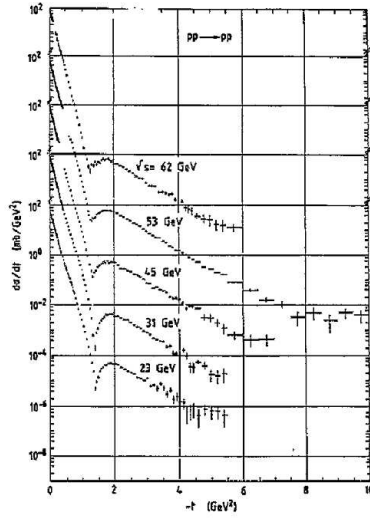


Figure 1.6: Proton-proton elastic scattering as measured at the CERN ISR. The differential cross section at different CMS energies from 23 GeV up to 62 GeV is shown as a function of the momentum transfer.

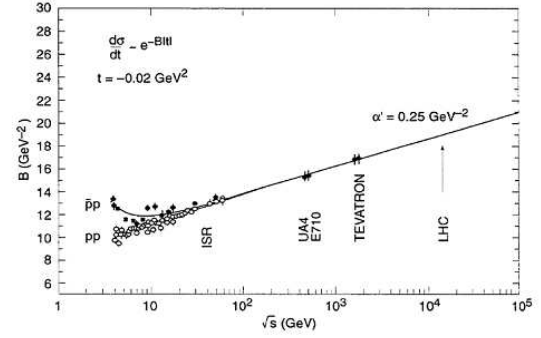


Figure 1.7: The forward slope parameter B for the $\bar{p}p$ and pp elastic scattering. The solid line refers to the Pomeron trajectory of the Regge Model.

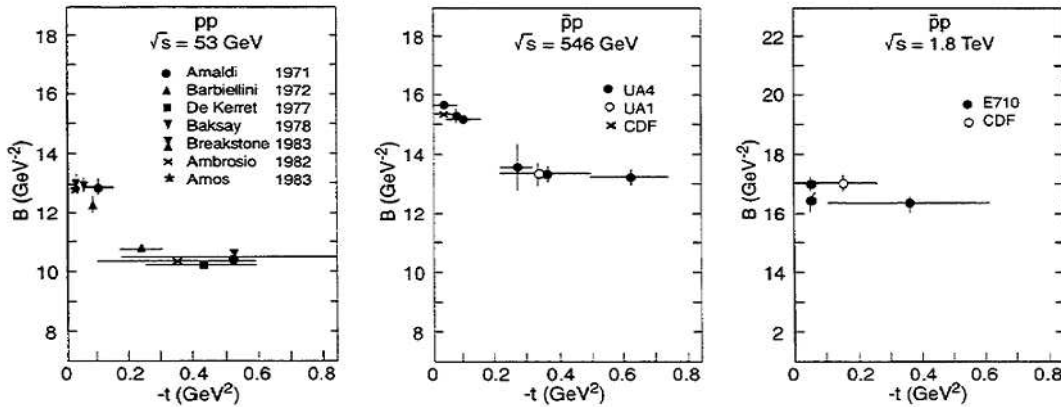


Figure 1.8: The local slope parameter $B(t)$ is plotted as a function of t for the $\bar{p}p$ and pp scattering at different energies. The horizontal bar indicate the interval in t where the exponential fit was actually performed. A simple exponential form should show a constant value of B as a function of t . The presence of steps indicate a presence of positive or negative curvature superimposed at the exponential form. These data show a positive curvature in the energy range of ISR (first graph) and SPS colliders (second graph) and an absence of curvature at the Tevatron's energy (third graph).

and SPS energies a concave structure is observed in the forward peak (higher B at lower $-t$). At the Tevatron energy a simple exponential is a well description (B is nearly constant).

This behavior can be reproduced with the *Impact Picture* of Wu [16] [17]. This model predicts that the curvature will change from positive to negative at some finite energy. At low energy, the shape of the t -distribution, determined by the proton electromagnetic form factor, has a positive curvature, as it has been observed. Asymptotically it should have a form similar to the Fraunhofer formula for diffraction by a completely absorbing disk in optics (i.e. a negative curvature). This feature can be reproduced also in the Regge approach, due to the interplay of two amplitudes for single and double Pomeron exchange, which have different energy and momentum-transfer dependence. A convex shape is what is expected at the LHC energy.

Another interesting observation would be the presence of oscillations at small momentum transfer

and high energy. The possibility of oscillation in the hadron-hadron differential cross section has been indicated using asymptotic theorems derived from general principles. According to Nicolescu et al. [18] these oscillations are compatible with existing measurements and the existence of these oscillations would suggest a general mechanism of saturation of axiomatic bounds.

The dip-shoulder region.

A dip followed by a broad maximum, typical of a diffraction pattern, has been found in the differential elastic cross section of pp scattering. As can be seen from fig.1.6 and fig.1.9 it is located in the region between $t = 1\text{GeV}^2$ and $t = 2\text{GeV}^2$. The depth of the minimum reflects the energy dependence of the real part of the scattering amplitude because the minimum of the diffraction pattern corresponds to a zero of the imaginary part. These results are well reproduced by the various phenomenological models. In fig.1.9 the same measurement for $\bar{p}p$ collision is shown. There is no dip, but is present instead only a break followed by a shoulder. The different results for $\bar{p}p$ and pp are well explained by the three-gluon exchange mechanism of Regge model of Donnachie and Landshoff [19]. The amplitude of this process, with different signs for $\bar{p}p$ and pp , interfere destructively with the complex amplitude of the diffraction peak in pp and constructively in $\bar{p}p$ scattering. At the LHC a dip is expected. The measurement of its energy position and depth could be used to test the prediction of the phenomenological models.

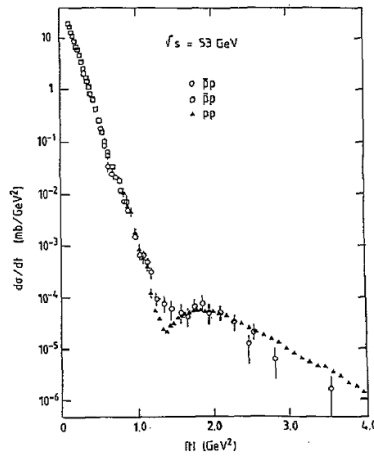


Figure 1.9: Differential cross section of $\bar{p}p$ and pp elastic scattering at $\sqrt{s} = 53\text{GeV}$.

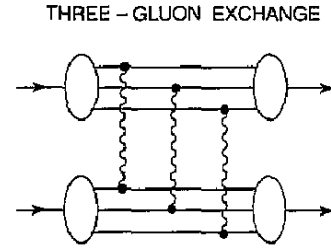


Figure 1.10: Three gluons exchange diagram for large- t proton-proton elastic scattering in the model of Donnachie and Landshoff [19]. The amplitude is proportional to t^{-4} . At smaller value of t other diagrams have to be included. Being these diagrams not calculable, an adjustable parameter has been introduced.

The large momentum transfer zone.

At large transfer momentum, the differential cross section shows a rather gentle fall off. This behavior is compatible with various models. The dependence on energy and transfer momentum measured and reported in fig.1.11 has been well explained by the three gluon exchange mechanism [19]. Also the *Impact Picture* [16] [17] and the Regge model of Desgrolard [20] fit well the large t regime of the differential cross section as it shown in fig.1.12.

The large- t measurement is interesting because the phenomenological models have different predictions at the LHC cms energy. If the three gluon exchange is really the dominant mechanism at high energy, the t -distribution will be smooth without structures. If the nature follows the *Impact Picture* or the Regge model we will have to find a diffraction pattern with several dips. This difference on the prediction reflects a basic question on the nature of the elastic scattering of hadrons. Is the scattering at large momentum transfer a single elementary process or is it dominated by a complicated interplay of various multiple elementary processes at low momentum transfer?

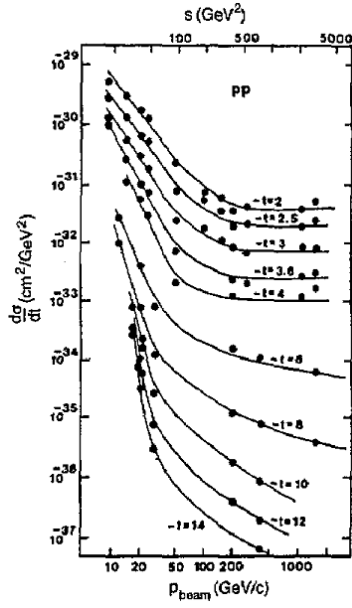
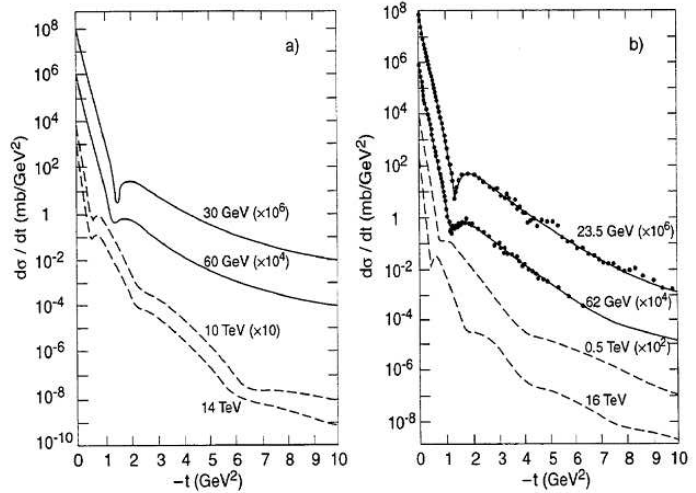


Figure 1.11: pp elastic scattering at large t . At fixed t the cross section is a fast decreasing function of the collision energy up to $\sqrt{s} = 10\text{GeV}$ and then seems to flatten off and remain constant. When energy and momentum transfer are sufficiently large, the differential cross section becomes only a function of t and no longer of s . This feature of the data suggests the onset of a specific dynamic mechanism which, according to Donnachie and Landshoff [19], is provided by the exchange of three gluons between the valence quarks of the two colliding protons (fig.1.10). In this model the proton is regarded as a three-quark state. In the elastic collision each quark in one proton scatters on one of the quarks of the other proton so that after the three elementary scattering processes of the constituent quarks, each triplet of quarks is again moving almost in the same direction and may thus recombine to form a proton. In the limit of large s and t , with $-t \ll s$, the cross section of the three-gluon exchange diagram can be approximately calculated and it is well described by $d\sigma/dt = C \frac{1}{t^8}$. For $|t| > 3\text{GeV}^2$ the data follows this power law.

Figure 1.12: Predictions on pp elastic scattering at high energy for (a) the *Impact Picture* and (b) the Regge model of Desgrolard. Both models predict a diffraction pattern with several dips at high energy (dashed lines).



1.1.3 Diffraction Dissociation and very forward particles production [24].

In general a diffraction dissociation is a process where (excluding the elastic scattering) one or both the ingoing protons are excited into a system X which fragments in a certain number of particles. It is distinguished from processes with a similar output, because the diffractive has distinctive kinematical features. In the simplest case of single diffraction $p + p \rightarrow X + p$, an event characterized by one proton in one hemisphere around the interaction point and a certain number of particles in the other, is a diffraction event if:

- i. The system X has the same quantum number (same intrinsic numbers) as the ingoing proton. Spin and parity could be different because some orbital angular momentum can be transferred to X in the collision.
- ii. The coherence condition [23] is maintained between the ingoing and outgoing proton's waves.

The coherence condition limits the change of momentum Δp that should be smaller than the inverse of the proton size. It implies that:

- ii.1 the mass M of the system X has an upper limit⁴.
- ii.2 The kinematical structure of the diffractive events has to be characterized by large rapidity⁵ gaps between the proton quasi-elastic scattered in the forward direction and the decay products of the system X .

The rapidity topology of a single diffractive dissociation is described in the caption of fig.1.13.

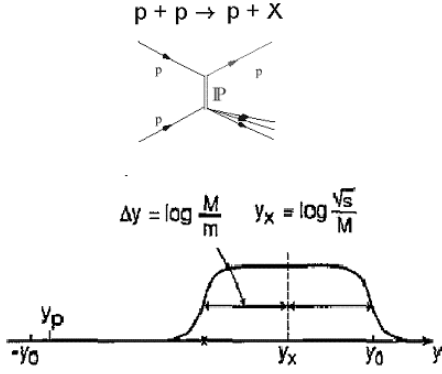


Figure 1.13: Rapidity Structure of single diffraction dissociation events [10]. The scattered p , having lost in the collision only a small fraction of its momentum, emerges with rapidity very close to the beam rapidity $y_0 = \log(\sqrt{s}/m)$. The rapidity distribution of the decay products of the system X will be centered in $y_x = \log(\sqrt{s}/M)$ and their spread will be $\Delta y = \pm \log(M/m)$. According to these relations, lower transfer momentum means lower mass of the X system with lower emerging angle and higher angle spread.

Another observed characteristic of diffractive events is that the average multiplicity of the diffractive cluster of invariant mass M was found to be the same as the one measured in hadronic collision with cms energy equal to M . The comparison is shown in fig. 1.14. Following the previous consideration, at the

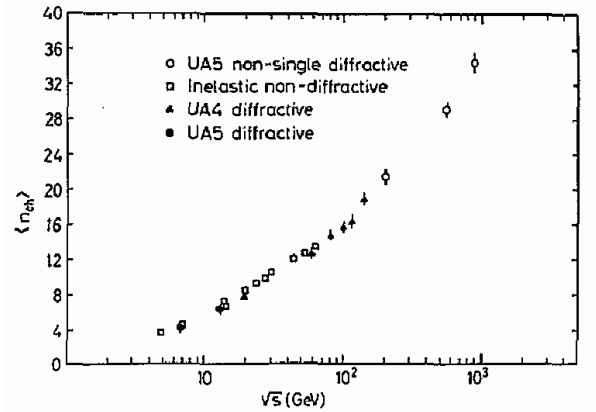


Figure 1.14: The average charge multiplicity of the diffractive cluster is compared with the overall charge multiplicity of non diffractive events [10].

LHC, a single diffractive event will have a maximum available mass for the system X of about $3TeV$ and the decay products will emerge at a rapidity of about 0.7, with a spread of ± 3 (see caption of fig.1.13). The average multiplicity will be more than 40 (see fig.1.14).

Diffractive classes and their cross-section measured at Tevatron and estimated for the LHC are shown in fig. 1.15. Their weight in the total cross section is significant and TOTEM has to perform an accurate measurement of these processes.

Nevertheless, these processes bring with them information on the mechanism of hadronic collisions at very high energy and thus, their importance is not exclusively related to the measurement of the

⁴At the high energy colliders, the transferred momentum is directly related to the mass M of the system X by $M^2 = \frac{\Delta p}{p}s$. The coherent condition implies that $M^2/s \leq 0.15$. High energy data provide clear evidence for diffractive production up to $M^2/s \sim 0.05$. At the LHC, this mean that X should have a mass M as large as $3TeV$.

⁵The rapidity of one particle with energy E and longitudinal momentum p_l is defined as $y = \frac{1}{2} \log(E + p_l/E - p_l)$. For $p \gg m$ the rapidity y can be approximated with pseudo-rapidity $\eta = -\ln \tan(\theta/2)$, where $\theta = p_l/p$

total cross section. The measurements and analysis, performed by some collider experiments [21] [22], have explained such diffractive events in terms of the existence of a colorless gluon-dominant cluster (Pomeron) in the partonic sea of the proton, with a most likely momentum fraction of its host proton near zero. More recently, a TOTEM/CMS joint on diffractive and forward physics has pointed out the interest on these measurements starting from the current understanding on diffractive and forward physics [24]:

"Diffractive events are characterized by the fact that the incoming proton(s) emerge from the interaction intact, or are excited into a low mass state, with only a small energy loss. Diffractive processes, for proton energy losses up to a few per cent, are mediated by an exchange with quantum numbers of the vacuum, the so-called Pomeron, IP, now understood in terms of partons from the proton. For larger energy losses, mesonic exchanges Reggeons and pions become important. The topology of diffractive events is characterized by a gap in the rapidity distribution of final-state hadrons caused by the lack of color and the effective spin of the exchanged object. Events with a fast proton in the final state can also originate from the exchange of a photon. In particular, tagging one leading proton allows the selection of photon-proton events with known photon energy; likewise, tagging two leading protons gives access to photon-photon interactions of well known center of mass energy. The average proton energy loss is larger and the proton scattering angle smaller in photon exchanges than for the diffractive case. This can be used to establish relative contributions of these two processes."

For this kind of physics it is mandatory to have a good coverage for charged particles at high rapidity η . This explains the collaboration between the two experiments TOTEM and CMS. Their coverage in η and ϕ , as shown in fig.1.16, is very promising for these studies.

In an exhaustive report [24], the working group outlined all the interests in diffraction and forward physics accessible at the LHC. A list of intents can be summarized as follow:

- i. At instantaneous luminosities, $L \lesssim 10^{30} \text{ cm}^{-2} \text{ s}^{-1}$: access to fundamental aspects of the strong interaction. The cross section for inclusive single diffraction $p + p \rightarrow p + X$, the cross section for inclusive double-Pomeron exchange $p + p \rightarrow p + X + p$ and their t and mass dependencies at the LHC center-of-mass energy.
- ii. At low luminosities: study of the low- x structure of the proton. This opens up the possibility of investigating the behavior of QCD in the high-density, saturation regime already probed for the first time in heavy-ion collisions and at HERA.
- iii. At higher luminosities, $L \sim 10^{32} \text{ cm}^{-2} \text{ s}^{-1}$: investigation of the proton structure, accessing the diffractive parton distribution functions, as well as the so-called rapidity-gap survival probability. The latter is closely linked to soft re-scattering and the features of the underlying event at the LHC.
- iv. At the highest available luminosities, $L \gtrsim 10^{33} \text{ cm}^{-2} \text{ s}^{-1}$: access to the generalized (or skewed) parton distribution functions. This is given by central exclusive production $p + p \rightarrow p + \phi + p$, that may become a discovery channel for particles with appropriate quantum numbers that couple to gluons (central exclusive production of a (SM or) MSSM Higgs boson).
- v. At all luminosities: a rich program of photon-photon and photon-proton physics can be pursued.
- vi. Test of the models used to simulate the development of cosmic-ray showers in air. The LHC center-of-mass energy indeed is approximately equal to the center-of-mass energy of a 100 PeV fixed-target collision in air.

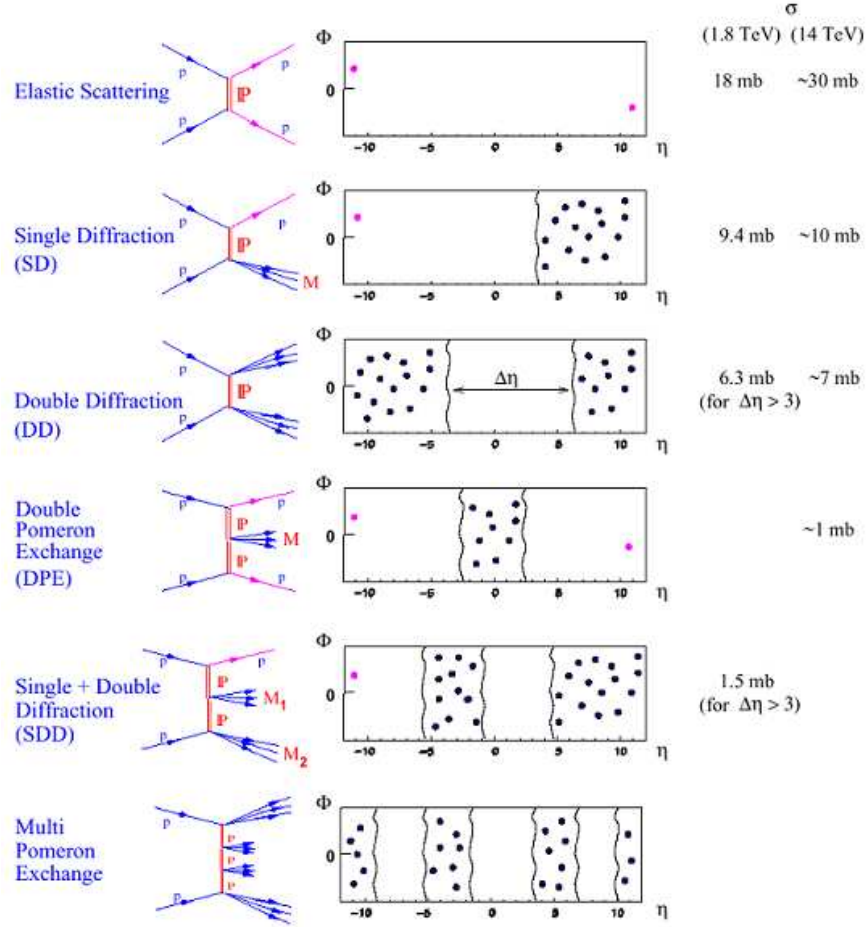


Figure 1.15: Diffractive process classes and their cross-section measured at Tevatron and estimated for the LHC [7].

The gaps in the rapidity distribution of final-state, caused by the lack of color and the effective spin of the exchanged object, is clearly visible in the topology of the diffractive classes that are shown in the plots, that are characterized by forbidden range in the pseudo-rapidity (output angle). Moreover, in diffractive dissociations, non-exponentially suppressed rapidity gaps have to be found.

1.2 Technical Aspects

The prerogatives of the TOTEM collaboration are the measurement of the total cross section and the study of elastic and diffractive dissociation processes at the LHC. According to eq. 1.1⁶, the total cross section is obtained measuring the inelastic and elastic rate, extrapolating the elastic differential cross section at the optical point (i.e. $t = 0$) and fixing the value of the parameter ρ . The design of the detectors has been done according to these intents and scaled on the quantity that have to be measured. In particular detectors for the measurement of intact "leading" protons (elastically scattered or with low momentum loss $\xi = \frac{|\Delta p|}{p}$) and for the inelastic processes have been foreseen. As it will be shown more in detail in this section, the specific observed quantities fix the requirements for this two type of processes, that can be briefly summarized in:

Elastic Scattering:

- A wide range in the momentum transfer, from $-t \approx 10^{-3} \text{ GeV}^2$ up to $-t \approx 10 \text{ GeV}^2$.
- The capability to reconstruct the kinematics of the scattered proton, i.e. the scattering angles $\Theta_{x,y}$ and the fractional momentum eventually lost (as for diffracted protons) $\xi = |\Delta p|/p$.

⁶ $\sigma_{TOT} = \frac{16}{(1+\rho^2)} \frac{(dN_{el}/dt)_{t=0}}{N_{el}+N_{inel}}$

Inelastic Scattering:

- A fully inclusive trigger (loss of events below $\sim 1\%$), obtained with a wide range covered in the pseudo-rapidity η^7 .
- The capability to reconstruct the tracks of the detected particles.

The total coverage that TOTEM will foresee in the forward direction is shown in fig. 1.16 together with the central coverage of CMS. In these plots $T1$ and $T2$ are the inelastic and RP the elastic TOTEM detectors. The effectiveness of this coverage can be understood looking at the charged particle multiplicity and at the energy flow for inelastic events at $\sqrt{s} = 14TeV$ shown on the right side. The pseudo-rapidity gap between RP and $T2$ is $\sim 3mrad$, while the one between $T1$ and $T2$ is $\sim 8mrad$.

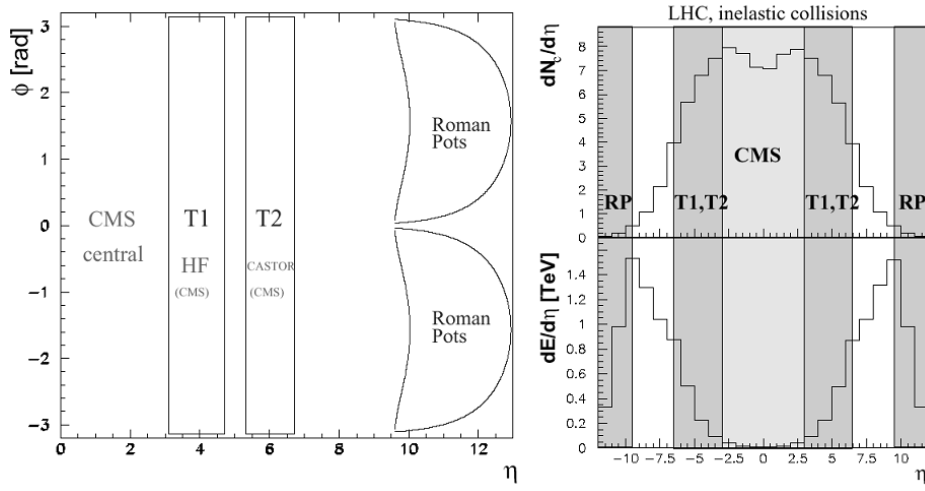


Figure 1.16: Left: coverage of different detector in the pseudo-rapidity (η) - azimuthal angle (ϕ) plane. Right: charged particle multiplicity and energy flow as a function of pseudo-rapidity for inelastic events at $\sqrt{s} = 14TeV$. As it can be seen from the plots on the right, the energy flow and the particle multiplicity of inelastic events peak in the forward region. About 99.5% of all non diffractive minimum bias events and 84% of all diffractive events have charged particles within the acceptance of the TOTEM detectors, T1 and T2.

1.2.1 The inelastic detectors

Gaseous detector have been chosen for the inelastic events detection. This type of detector give us the possibility to have large sensible area to cover a wide angle range, with a good spatial resolution for tracking purposes. They are moreover sufficiently robust to sustain radiation hard environment of the forward region. The regions inside CMS, where are located the TOTEM forward inelastic telescopes $T1$ and $T2$, are shown in fig. 1.17.

Two different type of gaseous detector have been used. Cathode strips chamber (*CSC*) for the telescope $T1$ and gas electrons multiplier (*GEM*) chamber for $T2$. The telescope $T1$ is centered at $9m$ on both side of the interaction point and it covers θ angles between $\approx 18mrad$ and $\approx 90mrad$ ($3.1 \leq |\eta| \leq 4.7$). In fig. 1.18 a drawing of one arm of this telescope and the readout layout are shown.

The telescope $T2$ is centered at $13.5m$ on both side of the interaction point and it covers θ angles between $\approx 3mrad$ and $\approx 10mrad$ ($5.3 \leq |\eta| \leq 6.5$). In fig. 1.19 a drawing of one arm of this telescope and the readout layout are shown.

1.2.2 The "leading" protons detectors

Silicon detectors inserted in movable pots (Roman Pots) have been chosen for the elastic scattering measurement. This configuration give us the possibility to place the sensitive area very close to the outgoing

⁷ $\eta = -\ln \tan(\theta/2)$

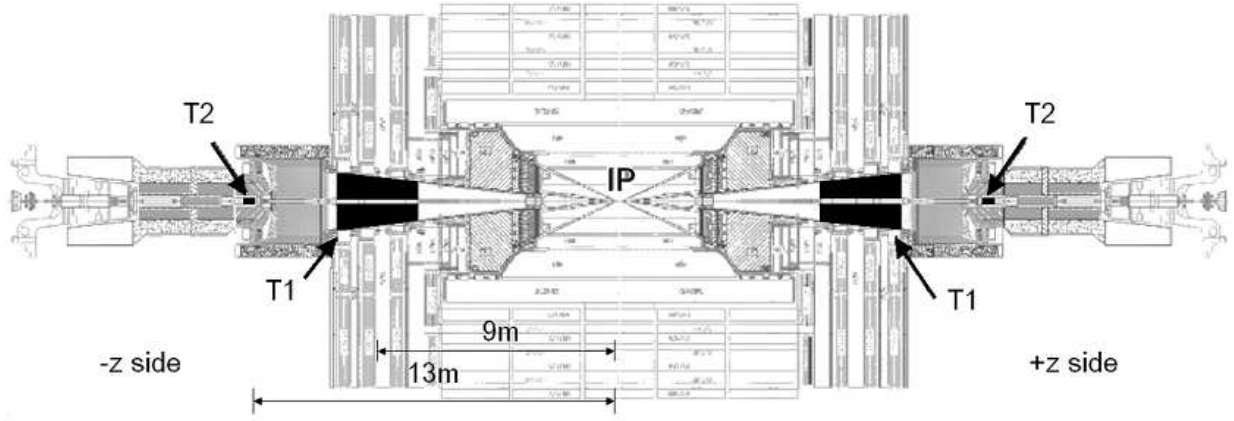


Figure 1.17: View of the inelastic forward trackers T1 and T2 inside the CMS detector.

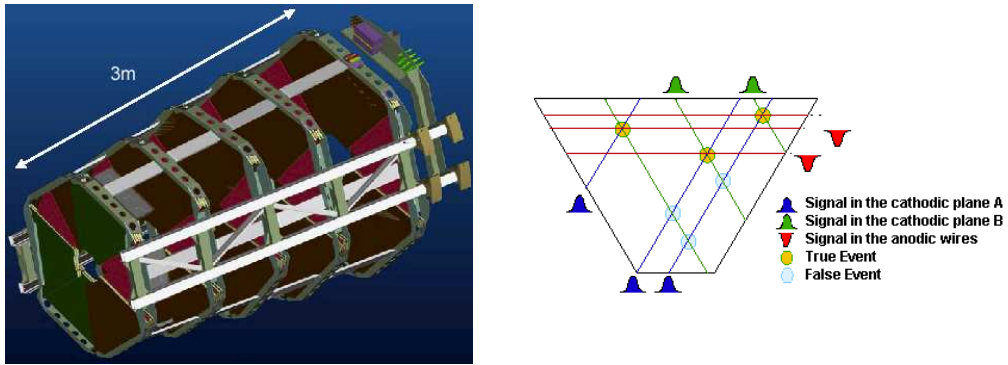


Figure 1.18: The forward Inelastic Telescope T1. Left: CAD drawing of one arm. Each arm is made by five planes. On each plane there are six trapezoidal CSC with full coverage. Right: readout scheme of one CSC of the TOTEM T1 telescope. Three different coordinates are acquired, two from the cathode strips and one from the anode wire. This allows unique identification of real tracks in case of multiple synchronous particles. In the picture it is shown the example of three synchronized particles that pass through the detector. With the anode readout the points of real passage can be isolated.

beam, where the main part of the quasi or elastically scattered protons will emerge. The use of silicon detector is a good choice when high spatial resolution is required and when the size is limited by insertion constraints. Two *RP* stations, on both side of the interaction point, will be installed at 220m and 147m, as indicated in fig. 1.20 where the LHC beam line with the *RPs* positions are shown.

The θ coverage is between $\approx 4\mu\text{rad}$ and $\approx 90\mu\text{rad}$ from the beam center. This coverage is not uniform on the azimuthal angle ϕ (see fig. 1.16) as it is for the inelastic detector. This is due to the design of the stations, whose drawing is shown in fig. 1.21. Each station is made by two units. In each units three pots are installed: two with vertical and one with horizontal movement. The elastic protons will pass mainly through the vertical ones, while the diffractive ones, with a small fraction of momentum lost, will be in the region covered by the horizontal pot. The sensitive planes and their overlapping are shown in the right part of the figure.

1.2.3 Elastic Measurement

In hadrons colliders, where the cms coincide with the laboratory system, elastic scattering is characterized by two emerging particles with back to back angular correlations and without other particles in the final state. At the LHC the scattering angles are quite small (fractions of $m\text{rad}$). This requires detectors on both side of the interaction point placed very close to the circulating beam.

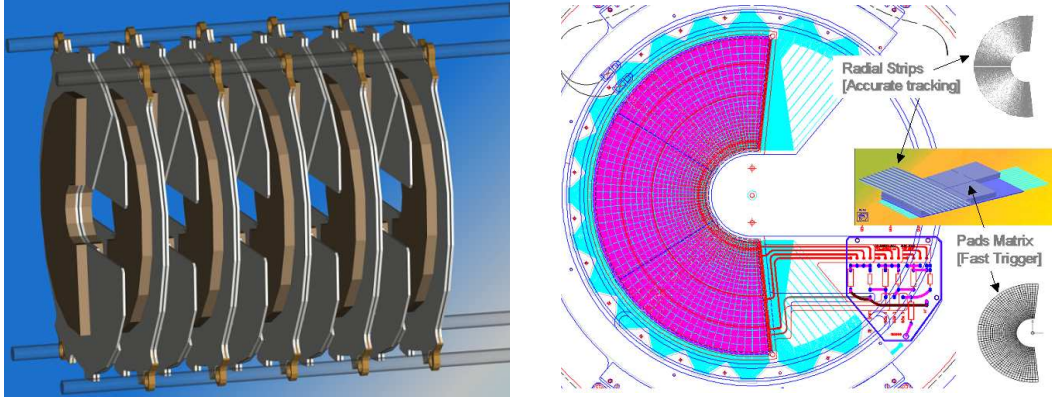


Figure 1.19: The forward Inelastic Telescope T2. Left: CAD drawing of one arm. Each arm is made by two two semi-telescopes, consisting in five pairs of Triple GEM chambers, mounted back-to-back. In each chamber, the sensitive area has a coverage of 192° in order to avoid any dead zone. Right: readout scheme of the TOTEM T2 telescope. The readout plane is a multilayered pcb, with a pattern of pads and strips for a two-dimensional readout. The angular coordinate of the particle's trajectory in the detector will be extracted from the pads readout, while the strips will foresee the accurate radial coordinate. The signals coming from groups of pads are used for fast triggering.

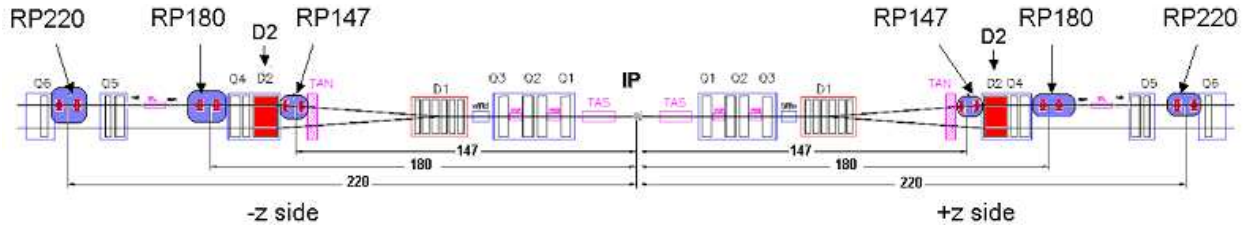


Figure 1.20: The LHC beam line and the Roman pots at 147m(RP147) and 220m(RP220). RP180 at 180m is another possible location but presently not equipped. It is important to note that the 147m RP are located before the D2 magnet, while the 220m tracking station is well behind it. This geometry naturally implements a magnetic spectrometer in the standard insertion, permitting TOTEM to measure particle momenta, with an accuracy of a few part per thousand. This will allow the accurate determination of the momentum loss of quasi-elastically scattered protons in diffractive processes.

The measurement of the elastic scattering requires particular beam conditions. Hadron colliders are used normally to look for rare events and then high interaction rate has to be foreseen. This is normally obtained reducing the size of the beams as much as possible in the interaction point, with the focusing effects of the quadrupoles. As a consequence the beam angular divergence increases. It follows that a large fraction of the scattered particles will emerge inside the acceptance of the beam, without the possibility to be detected. To measure an elastic process the opposite scheme is needed: a relatively large beam size in the interaction point with an angular divergence as small as possible.

In the elastic measurement, where the momentum is conserved, the most physically relevant quantity is the emerging angle of the outgoing particles. In general, the emerging angle that will be detected depends on the transversal position of the collision point in the interaction area. The uncertainty on this position could worsen the measurement. Once the beam characteristics have been defined, it is possible to achieve a particular condition in which particles scattered with the same angle are brought together to the same point on the detector, independently of the position of the collision point. In this way the measurement of the angle becomes decoupled from the measurement of the position of the collision point. The beam size at the crossing point becomes irrelevant and the scattering angle is uniquely determined by the displacement y at the detector. This condition is called *parallel-to-point focusing* optics and it is obtained placing the detector at the right distance from the interaction point (precisely at one

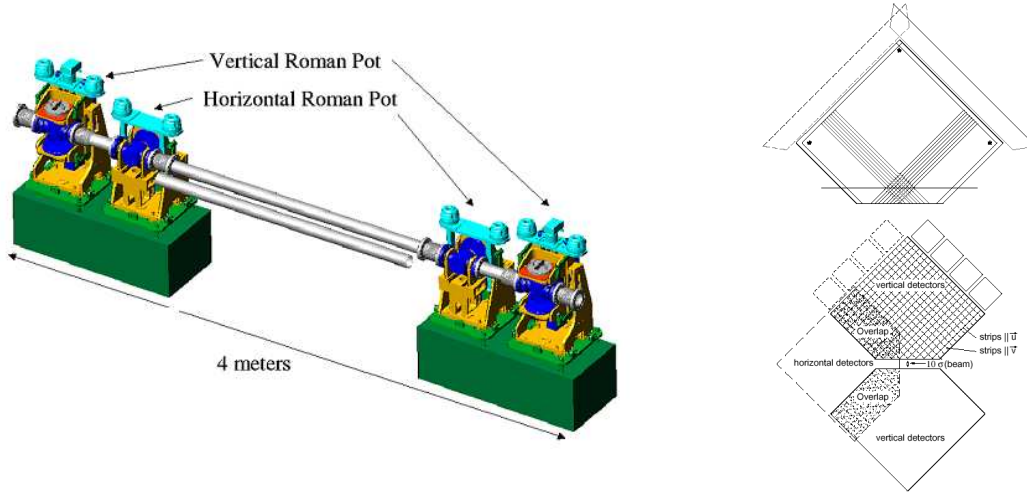


Figure 1.21: The forward elastic detectors RP. Left: CAD drawing of one Roman Pot station composed of two units each consisting of two pots that move vertically and one that moves horizontally. Right-Top: Two detectors planes mounted back-to-back defining an x - y coordinate plane by means of the strips at a 45° angle. Each pot is equipped with ten detectors plane. Right-Bottom: Overlap of the vertical and horizontal detectors for alignment. The 10σ of the beam is also shown.

quarter of the period of the betatron oscillation⁸).

In order to treat this subject in a quantitative way, it is necessary to spend few word about typical parameters and quantities of a collider (refer to appendix .1.3, .1.4 and to [27] for a more detailed description). In particular it is useful to define:

- the nominal transverse beam emittance ϵ . It is a beam quality concept reflecting the process of bunch preparation, extending all the way back to the source for hadrons.
- the amplitude function β . It is a beam optics quantity and it is determined by the accelerator magnet configuration. It is locally defined and the amplitude function in the interaction point is indicated as β^*

With these two quantities is possible to characterize the beam size $\sigma_{x/y}(s)$ at a distance s from the interaction point and the luminosity of the accelerator:

$$\sigma_u(s) = \sqrt{\epsilon_u \beta_u(s) / \pi} \quad (1.3)$$

$$L \propto \frac{1}{\sqrt{\epsilon_x \beta_x^* \epsilon_y \beta_y^*}} \quad (1.4)$$

Other useful definitions are the magnification $v_{x/y}$, the effective length $L_{x/y}^{eff}$ and the phase advance $\delta\mu_{x/y}(s)$:

$$v_u = \sqrt{\frac{\beta_u(s)}{\beta^*}} \cos \delta\mu_u(s) \quad (1.5)$$

$$L_u^{eff} = \sqrt{\beta_u(s) \beta^*} \sin \delta\mu_u(s) \quad (1.6)$$

$$\delta\mu_u(s) = \int ds \frac{1}{\beta_u(s)} \quad (1.7)$$

The *parallel-to-point focusing* optics previously introduced, ensuring that the scattering angle is uniquely determined by the displacement y at the detector, is satisfied if the phase magnification is zero (see

⁸The betatron oscillation is the motion of the particle that undergoes oscillation with respect to the designed trajectory of the collider, for the action of the alternating gradient focussing of the quadrupole magnetic fields. See appendix .1.3.

app. .1.4), i.e. the phase advance is $\pi/2$ (one quarter of the period of the betatron oscillation). If this condition is satisfied it follows that the angular divergence $\sigma_{x/y}^{I*}(s)$ in the interaction point can be expressed as:

$$\sigma_u^{I*} = \sqrt{\epsilon_u^*/(\beta_u^*\pi)} \quad (1.8)$$

Using these relations, the previous consideration about beam size and beam divergence can be summarized as follows: the elastic scattering measurement requires small beam divergence in the interaction point, i.e. large β^* and small ϵ . At the LHC the nominal β^* is $0.5m$. This is done to have high luminosity for searching rare events. For the forward elastic measurement higher values of β^* will be foreseen (on the scale of km). The nominal value of ϵ is instead around $0.5 \cdot 10^{-9} \pi \text{ rad m}$ or, in terms of normalized emittance $\epsilon_N = \gamma \frac{v}{c} \epsilon$, $\sim 3.75 \cdot 10^{-6} \pi \text{ rad m}$. Under particular conditions however it will be possible to work even with smaller values ($\sim 1 \cdot 10^{-6} \pi \text{ rad m}_r$).

Actually for the elastic measurement it is not useful to have only very large β^* . TOTEM wants to investigate a wide range of momentum transfer, from $-t \approx 10^{-3} \text{ GeV}^2$ up to $-t \approx 10 \text{ GeV}^2$. If the momentum transfer is large enough to scatter the particles at angles bigger than the angular divergence of the beam even at lower β^* , reducing the amplitude function it is useful in terms of statistics (lower β^* means higher luminosity L and then higher rate). For this reason, as it is shown in fig. 1.22, the use of different optics for various t range, improves the results that can be obtained.

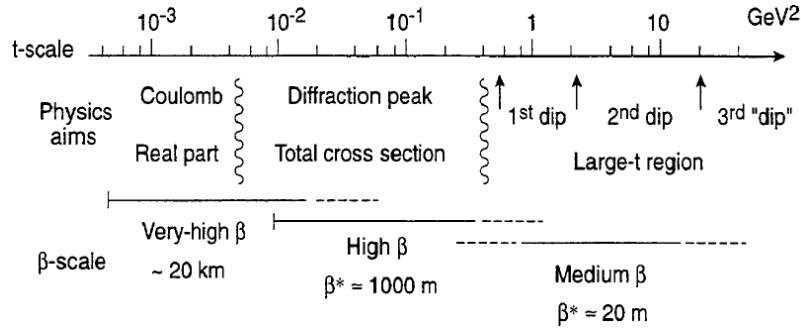


Figure 1.22: Scheme for a complete study of elastic scattering at the LHC [3].

At this point we are able to determine the minimum scattering angle that it is possible to measure with our RPs . First of all we have to define the detectable minimum distance u_{min} from the center of the beam where the RPs can be placed. Its value depends on the beam halo extension because we have to avoid any damage of the detector. We can express it terms of the σ_u of the beam, i.e. $u_{min} = K\sigma_u$. The multiplicative factor K will be chosen in the range $10 - 15$. In *parallel-to-point focusing* optics there is an unique correspondence between the emerging angle in the interaction point and the displacement u at the detector. This means that $u_{min}(s) = K\sigma_u(s)$ is directly translated into $u_{min}^{I*} = K\sigma_u^{I*}$ where the angular divergence has been used. According to eq. 1.8 we will have:

$$u_{min}^{I*} = K\sigma_u^{I*} = K\sqrt{\epsilon_u^*/\beta_u^*} \quad (1.9)$$

For example, if we want to detect protons with emerging angles of about $10\mu\text{rad}$, β_u^* has to satisfy the following requirement:

$$\beta_u^* \geq \left(\frac{K}{u_{min}^{I*}}\right)^2 \times \epsilon \approx \left(\frac{15}{10^{-5}}\right)^2 \times 5 \cdot 10^{-10} \approx 1000m \quad (1.10)$$

An emerging angles of about $10\mu\text{rad}$ is equivalent to a $t \approx 5 \cdot 10^{-3} \text{ GeV}^2$, that is a value that TOTEM has to reach. This shows explicitly why TOTEM needs special beam conditions with respect to the nominal ($\beta^* = 0.5m$). In fig.1.23 the value of the magnification and effective length have been calculated for $\beta^* = 1540m$ and $\beta^* = 90m$ around the interaction point IP5 in the LHC.

Extrapolation of the elastic scattering distribution to $t = 0$.

The extrapolation is performed according to a distribution e^{-Bt} , that is valid up to $|t| \approx 0.25 \text{ GeV}^2$. The slope B expected at the LHC is $\approx 20 \text{ GeV}^{-2}$ as reported in fig. 1.7. An extrapolation of about 20% means

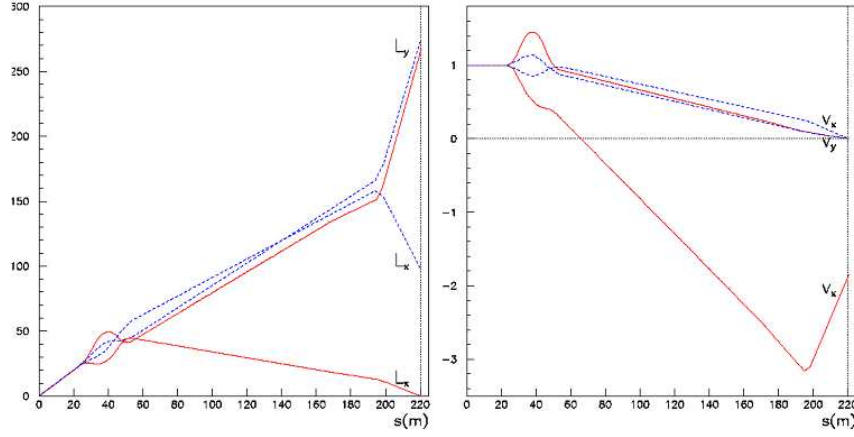


Figure 1.23: The effective length $L^{eff}(m)$ and the magnification v versus the distance to the IP along the beam (solid lines $\beta^* = 90m$, dashed lines $\beta^* = 1540m$). In order to be in a *parallel-to-point focusing* optics, the magnification v must tend to zero. From the right plot, if the detectors is placed at $220m$ with a $\beta^* = 1540m$, v is practically zero and we have $L_y^{eff} \sim 270m$. If v is sufficiently small, the displacement from the beam center of the scattered particle can be roughly calculate by the products of the effective length times the emerging angle (ref. app. .1.4). According to eq. 1.8, with $\beta^* = 1540$ and with a reduced emittance $\epsilon_N \sim 1 \cdot 10^{-6} \pi rad m$, the angular divergence of the beam is $\approx 0.3 \mu rad$. This means that the divergence of the beam σ_{beam} in the detector region will be of about $270m \times 0.3 \cdot 10^{-6} \approx 0.08mm$. If we want to get a safe separation between detector and beam ($\sim 15 \times \sigma_{beam}$), we have to stay at $1.2mm$ from the center of the beam. Particles scattered at $5 \mu rad$ ($-t \sim 10^{-3} GeV^2$) will pass in the detectors at about $1.35mm$ from the center of the beam, i.e. in a region where can be detected.

that, for the exponential behavior and for the value of the slope, the minimum $|t|$ should be $10^{-2} GeV^2$. On this extrapolation two aspects have to be considered. The acceptance of the detector and the non perfect exponential form for the elastic scattering distribution.

The acceptance of the detector limits the minimum measurable $|t|$ and it has to be used to convert the reconstructed elastic events $(dN/dt)_{rec}$ into the real (dN/dt) . In fig. 1.24 the acceptances for the RP station at $220m$ from the interaction point are shown from different optics settings. The uncertainties related to the acceptance function can be reduced using a special selection of the reconstructed events (for which a purely geometrical acceptance can be used).

The deviation from a perfect exponential form of the elastic scattering distribution can be described by an analytical function. According to various phenomenological models, fig. 1.25 shows that a parabolic form could be a reasonable solution. Moreover, other source of uncertainty are related to beam and optics properties (like beam divergence, energy, crossing angle and effective length L_{eff}) and to Detector-Beam position alignment. All these parameters influence the extrapolation and have to be taken under control.

Large $|t|$ Elastic Scattering.

For large momentum transfer (above $|t| \approx 1 GeV^2$) the high- β optics (i.e. low luminosity) is not suitable because of the reduced proton acceptance, and therefore an optics with a relatively low value of β^* is needed. Fig. 1.26 shows the differential cross section up to $|t| = 10 GeV^2$. The measurable ranges of $|t|$ with different β^* values are also shown, together with the main physics process in each part of the $\frac{d\sigma}{dt}$ distribution. Fig. 1.27 shows the predictions of various models [25]. It is interesting to observe as the various models give different predictions at large t . The proper optics, i.e. the proper β^* , will permit a measurement with good statistics for solving those different expectations.

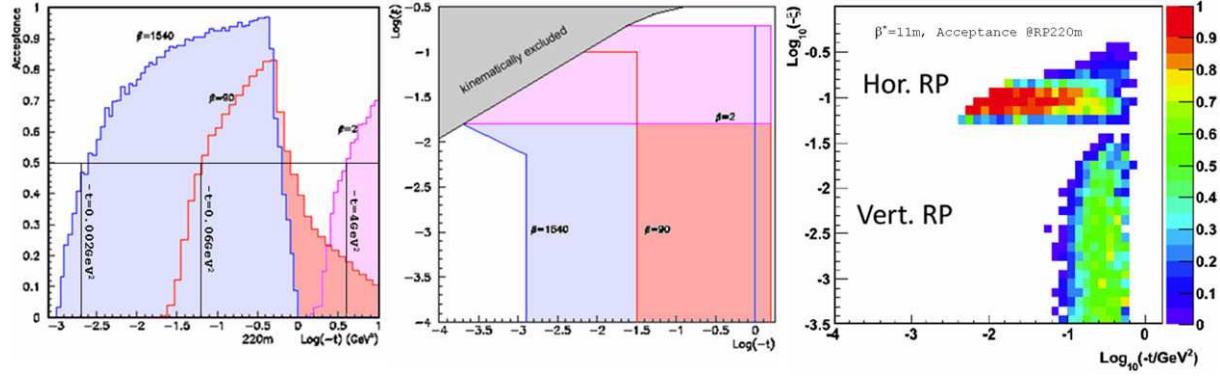


Figure 1.24: Left: Acceptance in $\log_{10}|t|$ for elastically scattered protons at the Roman Pot station at 220 m for different optics settings. Middle: acceptance in $\log_{10}|t|$ and $\log_{10}\xi$ for diffractively scattered protons at the same RP station for different optics settings. The contour lines represent the 10% level. Right: Simulated acceptance for $\beta^* = 11m$.

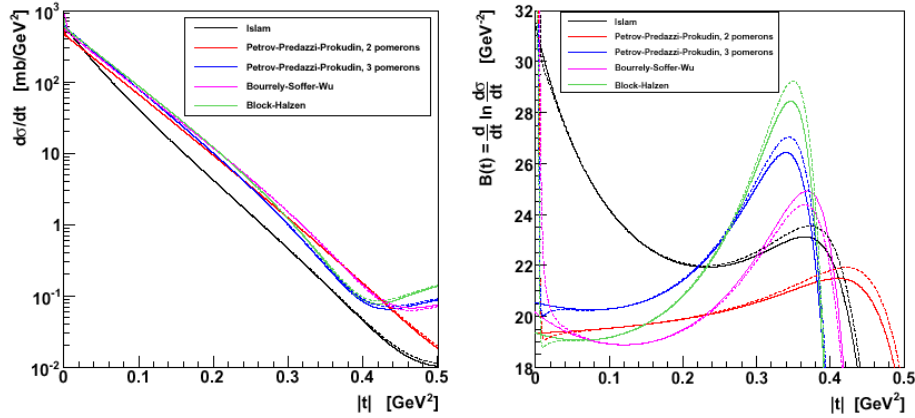


Figure 1.25: Left: differential cross-section of elastic scattering at $14TeV$ as predicted by various models [25], focussing on the quasi-exponential domain at low $|t|$. Right: exponential slope of the differential cross-section. The deviations from a constant slope show how the cross-sections differ from a pure exponential shape. Continuous (dashed) lines: with (without) Coulomb interaction. Most theoretical models predict an almost exponential behavior up to $|t| \approx 0.25 GeV^2$. For all the models considered, except for the one by Islam et al., the deviations are small. In the t -range mentioned, the slope $B(t)$ can be well described by a parabola, which is therefore used for the fitting function and the extrapolation. Since this quadratic behavior of the slope characterizes all the models, the extrapolation method is valid in an existing-model independent way.

The $\rho = \Re F(t=0)/\Im F(t=0)$ [10].

The real and imaginary parts of the scattering amplitude in the forward direction are needed to measure ρ . According to the optical theorem in eq. 1.1, the imaginary part is the total cross section. The real part of the hadronic amplitude is obtained observing the interference with the Coulomb amplitude which is known. This means that the measurements have to be done in a region where the Coulomb and hadronic amplitude are comparable in magnitude. The two processes have the same probability when the momentum transfer is $t_0 \simeq (8\pi\alpha)/\sigma_{TOT}$ where α is the fine-structure constant. At the LHC cms energy, assuming a total cross section of $\approx 110mb$, this means $t_0 \approx 0.7 \cdot 10^{-3} GeV^2$.

The differential cross section, with the Coulomb amplitude and with a low- t parameterization nor-

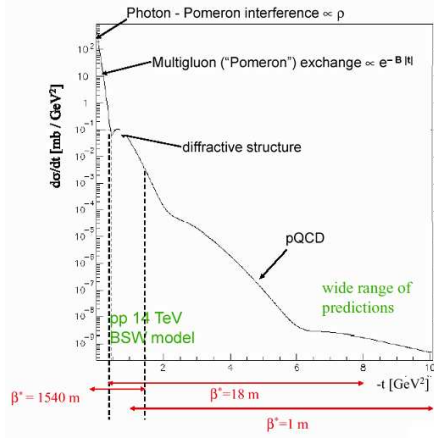


Figure 1.26: Elastic scattering cross section measurement with TOTEM up to large $|t|$. The β^* with their relative t -range investigation are also shown.

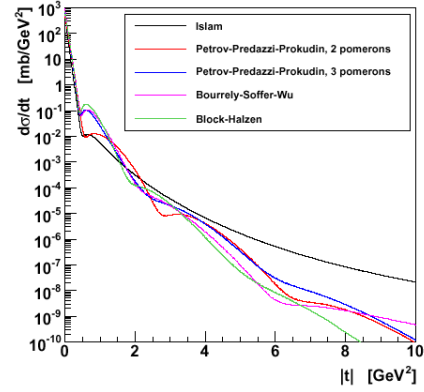


Figure 1.27: Differential cross-section of elastic scattering at $\sqrt{s} = 14\text{TeV}$ as predicted by various models [25]. On this scale, the cross-sections with and without the Coulomb component (continuous and dashed lines respectively) cannot be distinguished.

mally used for the hadronic one, can be written as:

$$\begin{aligned} \frac{d\sigma}{dt} &= \frac{16\pi}{s^2} |F_C e^{(\mp i\alpha\phi)} + F_h|^2 \\ F_C &= \pm \frac{1}{2} \alpha s \frac{G^2(t)}{|t|} && \text{Coulomb Amplitude} \\ F_h &= \frac{s}{16\pi} \sigma_{TOT}(\rho + i) e^{Bt/2} && \text{hadronic Amplitude} \\ \frac{d\sigma}{dt} &= \frac{4\pi\alpha^2(\hbar c)^2 G^4(t)}{|t|^2} + \frac{\alpha(\rho - \alpha\phi)\sigma_{TOT} G^2(t)}{|t|} e^{-B|t|/2} + \frac{\sigma_{TOT}^2(1+\rho^2)}{16\pi(\hbar c)^2} e^{-B|t|} \end{aligned} \quad (1.11)$$

where $G(t)$ is the proton electromagnetic form factor and the upper-lower sign refer to $\bar{p}p/pp$ scattering respectively. The interference term in the differential cross section is proportional to the quantity $(\rho \pm \alpha\phi)$. The relative Coulomb-hadronic phase $\alpha\phi$ can be reasonably approximated for $B \sim 15\text{GeV}^{-2}$ with $\phi = \log(0.07/|t|) - 0.577$. At $t = t_0$, $\alpha\phi \simeq 0.027$. In fig. 1.28 data in the Coulomb region from the SPS are shown. The best fit, at the SPS energy, is $\rho = 0.135$.

In fig. 1.34 the possibility for the TOTEM detectors to perform the measurement of the differential cross section at low $|t|$ is shown. In these plots, the $-t$ corresponding to an acceptance of 50% are shown as a function of the distance from the beam, the emittance ϵ and the cms energy. It is not possible for TOTEM to perform this measurement at $\sqrt{s} = 14\text{TeV}$ with the nominal emittance $\epsilon \sim 0.5 \cdot 10^{-9} \pi \text{ rad m}$. With a reduced emittance ($\approx 1/4$ of the nominal) it is not excluded, but it will be necessary to go much closer to the beam center with respect to the designed $10\sigma + 0.5\text{mm}$. This possibility will be investigated more deeply, when information on the real characteristics of the beam will be available. Lowering the cms energy of the collision, the ρ measurement will be accessible as can be seen from the right graph. Here, the $10\sigma + 0.5\text{mm}$ (which fix the minimum distance that the TOTEM detector can safely reach) and the $\sigma_{hadronic} = \sigma_{Coulomb}$ (which fix the minimum distance to measure the needed $|t|$) lines are plotted.

1.2.4 Inelastic Measurement.

The inelastic events are the main contribution to the total cross section of pp interactions at the LHC, with $\sqrt{s} = 14\text{TeV}$. They are indeed more than 70% of the σ_{TOT} . They are characterized by different inelastic physics processes that can be classified in two classes, with very different topologies:

- i. the non-single-diffractive events (NSD) where the secondaries have a rapidity range distribution

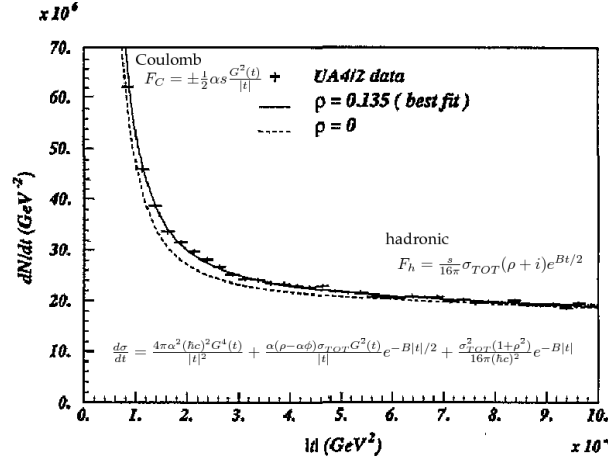


Figure 1.28: The t -distribution of $\bar{p}p$ elastic scattering at $\sqrt{s} = 540 \text{ GeV}$ in the region of the Coulomb interference [26]. The broken curve indicates the result which would have been obtained for $\rho = 0$.

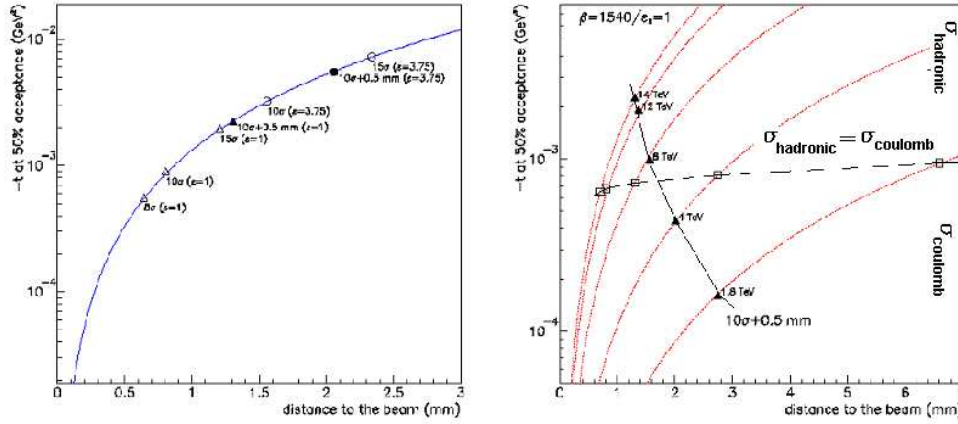


Figure 1.29: $-t$ corresponding to an acceptance of 50% versus the distance to the beam center. Left: for nominal LHC center-of-mass energy. The corresponding σ^{beam} value for two emittances is also marked. Right: for different center-of-mass energies, the $-t$ value where $\sigma_{hadronic} = \sigma_{coulomb}$ is also marked (square).

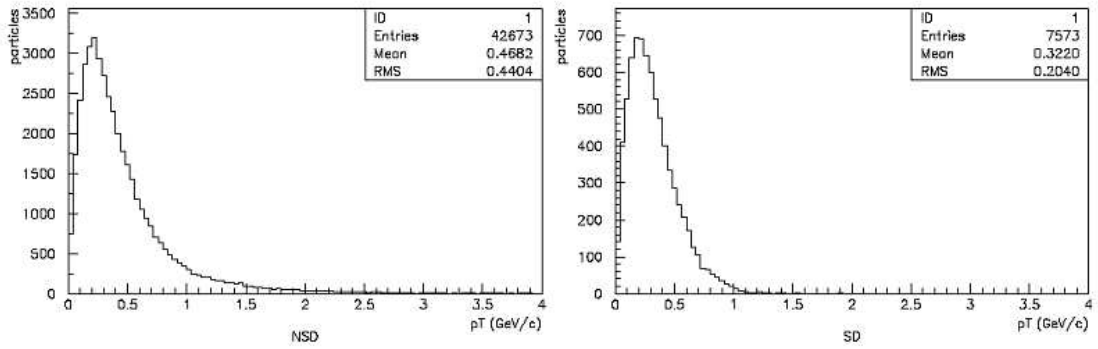
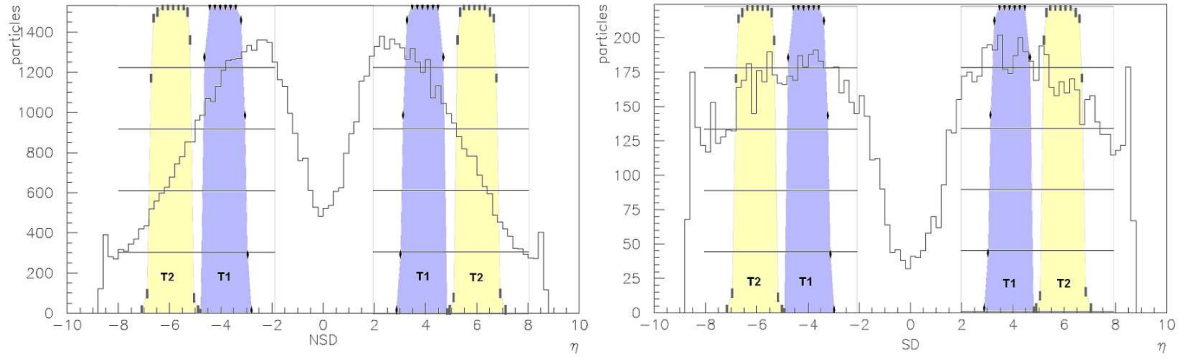
extending over the full phase space; they include the contributions of the non-diffractive minimum bias and of the double diffractive processes.

- ii. the single diffractive events (SD): $p + p \rightarrow p + X$, where one proton is scattered quasi-elastically at very small polar angle with energy close to the beam energy. The system X , which recoils against the quasi-elastically scattered proton, fragments into secondaries having a typical rapidity distribution concentrated in the opposite hemisphere.

The NSD and SD events are respectively the $\sim 60\%$ and the $\sim 10\%$ of the total cross section, as reported in table 1.1.

Simulations for these events have been performed. The p_T and η distributions of the particles at the generation level for NSD and SD events are shown in fig. 1.30 and 1.31. The most probable number of charged particles generated per hemisphere has been also calculated. The multiplicities are quite small and in the pseudo-rapidity coverage of T1, we have $n_{T1}^{NSD} \approx 12$ and $n_{T1}^{SD} \approx 2$ and for T2, $n_{T2}^{NSD} \approx 9$ and $n_{T2}^{SD} \approx 2$. In fig. 1.32 the results for NSD events are shown. These small multiplicity for SD could cause uncertainty on the identification of SD events, for possible losses or contamination with background.

	$\sigma[mb]$		$\sigma[mb]$
Elastic	≈ 30		
Inelastic	≈ 80	Non-diffractive	≈ 58
		Single diffractive	≈ 14
		Double diffractive	≈ 7
		Double Pomeron	≈ 1
		Multi Pomeron Exchange	$\ll 1$
σ_{TOT}	≈ 110		

Table 1.1: Events cross section at the LHC with $\sqrt{s} = 14TeV$.Figure 1.30: The p_T distributions for NSD (left) and SD (right) events, as generated by Pythia for 1000 events.Figure 1.31: The η distributions for NSD (left) and SD (right) events, as generated by Pythia for 1000 events (a $500MeV$ cut on $|p|$ has been applied). The geometric acceptance of the two inelastic telescope $T1$ and $T2$ is also shown.

Particular attention has to be paid on the reduction of the uncertainty on the rate of these inelastic events N_{inel} , that is used in eq. 1.1⁹ for the measurement of the total cross section. The main sources of error are the loss of good event and the contamination from background. For this reason, the inelastic detectors have been designed in order to guarantee:

- A *Fully Inclusive Trigger* or *Minimum Bias Trigger* with expected loss on N_{inel} at the 1% level.
- The capability of identification of beam-beam events against background.

⁹ $\sigma_{TOT} \propto 1/(N_{el} + N_{inel})$

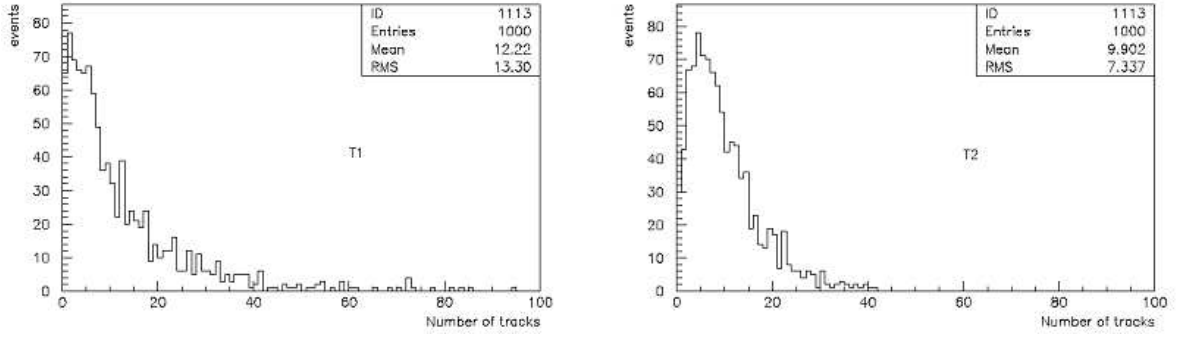


Figure 1.32: The T1 (left) and T2 (right) track multiplicities for 1000 *NSD* events.

The *Fully Inclusive* requirement is obtained choosing the proper angular coverage. The fraction which is not seen (not included in the angular coverage) has to be evaluated properly and accounted for. Lines of constant fractional loss on N_{inel} predicted by Monte Carlo simulation are shown in fig. 1.33 [3]. They have been calculated at $\sqrt{s} = 14\text{TeV}$ and have been plotted as a function of the lower and upper extremes of pseudo-rapidity interval covered. In order to keep the loss of the inelastic events to the level of 1%, the coverage has to be at least 3 pseudo-rapidity unit, close to the beam rapidity. The coverage of the inelastic detectors of TOTEM ($3.1 \leq |\eta| \leq 4.7$ and $5.3 \leq |\eta| \leq 6.5$) have been reported in the plots.

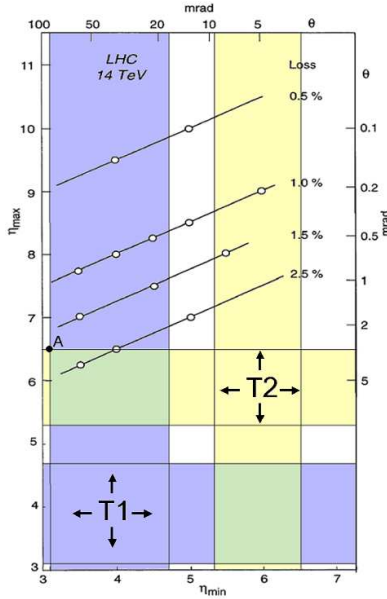


Figure 1.33: Results of the simulation on the efficiency of the inclusive trigger. Lines of constant fractional loss of the inelastic events are shown as function of η_{min} and η_{max} , the extremes of the interval covered by the inelastic detector. The corresponding polar angles are also indicated. The geometric acceptance of the two inelastic telescopes T1 and T2 it is shown in both axes. The point A, corresponds to having choose the the minimum η of T1 and the maximum η of T2. This point doesn't take into account the gap between the two telescope and this means that the effective loss will be a little bit worse than the one indicated. This plot is extracted from the TOTEM "letter of intent" where it has been shown just as a qualitative guide line. More accurate studies have been done and with a proper choice of the optics, this level can be reduced under the 1%. In particular $\sim 0.8\%$ is expected with $\beta^* = 1540m$ and $\sim 1\%$ with $\beta^* = 90m$.

The major loss in the inelastic rate is related to single and double diffractive events and the lost events are mainly those with a very low mass (below $\sim 10\text{GeV}/c^2$), since all their particles are produced at pseudo-rapidity beyond the T2 acceptance. To obtain the total inelastic rate, the fraction of events lost, because of the incomplete angular coverage, has to be estimated by extrapolation. A good extrapolation is given assuming $d\sigma/dM^2 \sim 1/M^2$ as shown in fig. 1.34 for single diffractive events. The ratio of detected single diffractive events is also shown as a function of $1/M^2$.

For the identification of real beam-beam events against background, the TOTEM collaboration has decided to follow two strategies:

- i. the *vertex reconstruction* of the detected tracks, that allows the rejection of events not originated in the interaction point;
- ii. the use of *particular triggering condition*, that helps in the events identification, according to their topology.

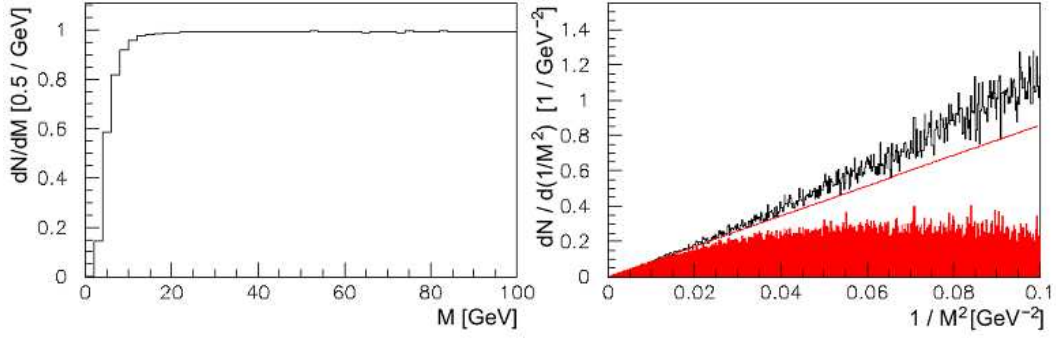


Figure 1.34: Left: ratio of detected Single Diffractive events as a function of the diffractive mass M . Right: simulation (unshaded) and acceptance corrected (shaded) Single Diffractive distribution as function of $1/M^2$. The line shows a linear fit based on the acceptance corrected events in the mass region above $10 \text{ GeV}/c^2$. The extrapolated number of events differs from the simulation expectations by 4%, corresponding to a 0.6 mb uncertainty on the total cross-section. The same estimate for the Double Diffraction and Double Pomeron Exchange gives a 0.1 mb and 0.2 mb uncertainty, respectively.

In the inelastic detectors, mainly three types of background are present.

1. The showering of the very forward particles produced in beam-beam interaction.
2. The beam-gas interactions that takes place when a beam proton collides with a residual molecule in the beam-pipe vacuum. It will be one of the main contamination sources, because the topology of these events is very similar to Single Diffractive events.
3. The interaction of scattered particles with the beam-pipe. The number of charged particles generated from these secondary interactions in the beam-pipe material can increase the number of hits per events up to one order of magnitude.

The *vertex reconstruction* is based on the tracking capability of each TOTEM sub-detector. It is realized in three steps: the *Pattern Recognition* where space regions (roads) with a high probability of having a track are identified, the *Track Reconstruction* where the tracks in the selected roads are reconstructed and finally, the *Vertex Reconstruction* where the reconstructed tracks are fitted to a primary vertex.

Particular triggering conditions are foreseen by the capability of each TOTEM sub-detector to send fast trigger-signals. The possible combinations for inelastic events, allow at least three main types of trigger:

- *Inelastic Single-Arm Trigger*: at least one particle in either hemisphere has to be detected by the inelastic detector. With this condition there is the best efficiency, but the measurement could suffers from background contamination. The vertex reconstruction will be particularly useful to discriminate real events from background in this trigger condition.
- *Inelastic Double-Arm Trigger*: at least one particle per hemisphere. With this condition the measurement is protected against background contamination but there is inefficiency at small M .
- *Inelastic Triggers and Proton*: at least one proton in one of the elastic detectors and at least one particle in one of the inelastic detectors in the opposite hemisphere. In this case the trigger is very clean, but the proton inefficiency has to be extrapolated.

Different acquisition runs can be realized with different trigger conditions. In this way it will be possible to focus the measurements on specific processes. For a better understanding of the background events, it will be possible to trigger on non colliding bunches. In this way the rate of the background not generated during beam-beam interaction could be directly measured and the errors on the inelastic rate could be better defined.

1.2.5 Diffractive Measurement.

A study of diffractive dissociations¹⁰ requires a set of detectors that are able to observe and distinguish the typical kinematics and topology of this kind of processes.

Two different final states may emerge from the interactions: groups of particles created by the fragmentation of excited protons or intact (LEADING) protons that have lost part of their momentum. In order to measure these dissociations, detectors with a very large rapidity coverage in the first case, and detectors with kinematics reconstruction capabilities for the second case are needed.

The large coverage requirement is due to the complete pseudo-rapidity range covered by the fragmen-

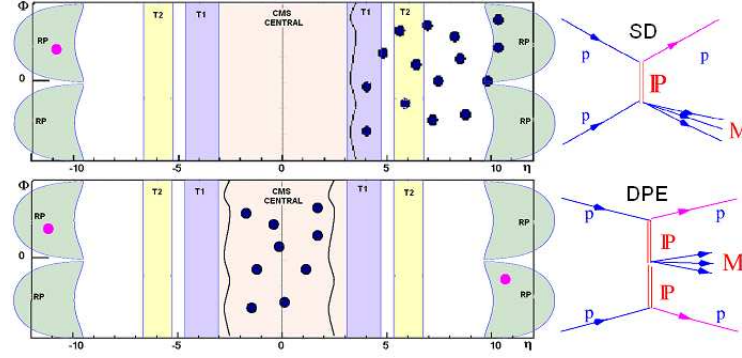


Figure 1.35: Single Diffraction (SD) and Double Pomeron Exchange (DPE) processes. The TOTEM ($T1$, $T2$ and RP s) and central CMS coverage is shown.

tation of the various topologies of dissociations (fig 1.15), and by the fact that, for these kind of processes, it is necessary to be able to locate the rapidity gaps that characterize these interactions. As an example, a possible Single Diffraction SD and a Double Pomeron Exchange DPE are shown in fig. 1.35. In these plots the very large coverage foreseen by the detectors of TOTEM and CMS together has been superimposed to the topology of the processes. The case of SD is an example of an event where all the TOTEM detectors are involved. For the DPE case instead, it is evident why a collaboration between the two experiment is required for the study of this kind of physics.

The kinematics of the diffracted proton will be reconstructed with the measurements performed by the RP stations. In particular the scattering angles $\Theta_{x,y}$ and the fractional momentum loss $\xi = |\Delta p|/p$ will be foreseen. This information will be useful to determine the mass of the system X , produced during the interaction.

The measurement of the momentum loss by the diffractively scattered protons will be obtained using two different approaches:

- i. the backtracking of the observed positions and angles of the protons through the accelerator optics elements. The full set of kinematic variables, $(\Theta_x, \Theta_y, x^*, y^*, \xi)$, can be reconstructed inverting the proton transport equations (34) by χ^2 minimization procedure;
- ii. the measurement of the bending of the protons in the beam separation dipole D2. This dipole is inserted between two TOTEM's RP stations (see fig. 1.20) and its dispersion difference can be used for the proton momentum reconstruction.

¹⁰See sec. 1.1.3.

Chapter 2

The forward inelastic telescope T2

My work for the TOTEM experiment has been focused on the forward inelastic telescope $T2$. The aim of this section is to give an overview of this telescope. As previously anticipated (sec. 1.2.2), the collaboration has decided to use a gaseous triple GEM detector and the VFAT2 chip, a front end ASIC designed for this experiment, to process the signals.

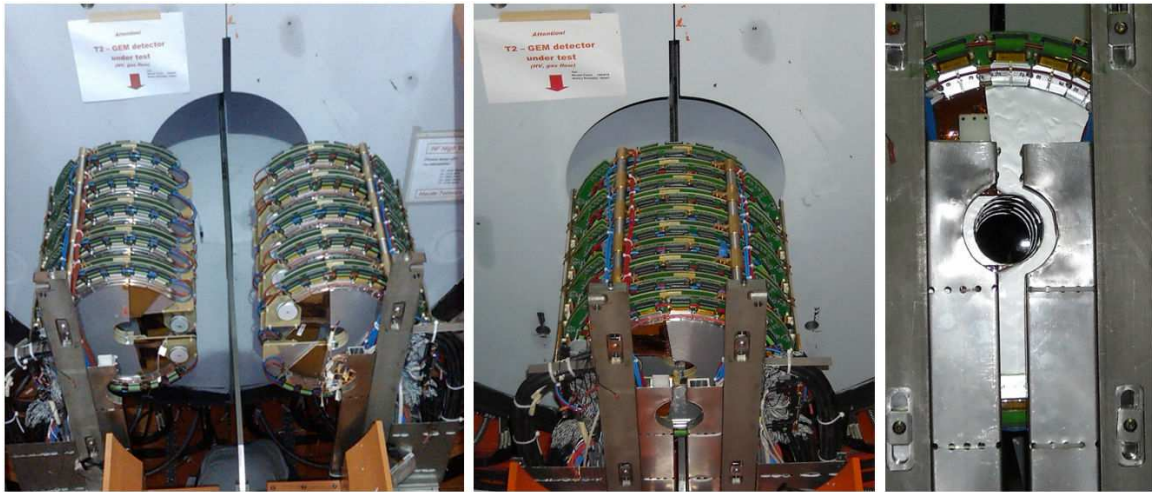


Figure 2.1: One arm of the forward inelastic telescope $T2$ installed on the plus side of the interaction point $IP5$ of the LHC . Each arm is characterized by two quarter (left picture) of ten triple GEM each, that will be closed around the beam pipe as shown in the middle-right picture.

2.1 Telescope requirement

The configuration and the technology chosen to realize the $T2$ telescope is motivated by the physics that has to be performed and the operating condition that the telescope has to sustain.

Physics Requirements

The forward telescopes for the physics aimed by the TOTEM experiment has to provide a *fully inclusive trigger* in order to minimize loss of events in the computation of the inelastic rate, a *vertex reconstruction* capability to reduce the background contamination (ref. sec. 1.2.4 and 1.2.5) and a measurement of the forward *process topology*. A *left-right* installation of the detectors around the interaction point it is moreover used to control systematic uncertainties (the pp collision in the LHC machine is indeed left-right symmetric).

The *fully inclusive trigger* puts constraints on the pseudo-rapidity coverage and on the efficiency of the detector used. The angular coverage is fixed by the mechanical insertion inside the experiment CMS. To have the maximum possible coverage the TOTEM inelastic detection it has been split in two telescopes T1 and T2. T2 is placed at $13.5m$ on both side of the interaction point (IP) and covers the range $5.3 \leq |\eta| \leq 6.5$. In fig. 2.2 a drawing of the assembly of half telescope placed on one side of the IP is shown. The total efficiency of T2 has to be practically on the scale of 100% because, as indicated in sec. 1.2.4, it has to detect events that could have a very small number of particles produced. Due to the aligned detector redundancy (ten detector per quarter), the minimum efficiency demanded to each single detector, to guarantee the full efficiency of the full telescope, is of about the 80%. Despite this a final single detector efficiency higher than 95% will be conservatively required.

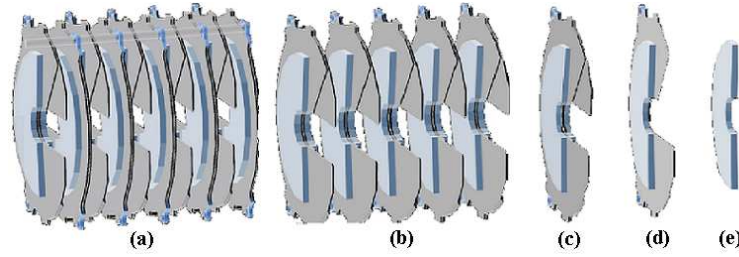


Figure 2.2: Telescope assembly of the forward inelastic telescope T2. Each T2 telescope is divided in two part, one per hemisphere, with twenty Triple GEM (a) each, equally spaced in $\sim 40cm$. The coverage of each half telescope is $\Delta\phi = 360^\circ$ and $5.3 \leq |\eta| \leq 6.5$. It is divided in two semi-cylindrical quarter (b) that are close on the beam pipe. In each quarter five couple of Triple GEM chambers mounted in a *back to back* configuration (c) are installed. An emerging particle from the IP will travel through 10 detectors. The sensitive volume (e) of a single Triple GEM (d) has a coverage of $\Delta\phi = 192^\circ$, with a radial estension between $\sim 42.5mm$ and $144.5mm$ from the center. Drawing elaborated from a CERN presentation of E. David.

The *vertex reconstruction* requires an adequate granularity of the readout pattern and a sufficient number of measured points per track. All particles that emerge from the interaction with a pseudo-rapidity covered by T2 will travel trough a quarter of telescope (ten triple GEM detectors). It will be possible therefore to collect a maximum of ten point per particle uniformly distributed in $40cm$. The spatial resolution foreseen for the *vertex reconstruction* will provide also important information on the study of forward inelastic processes, measuring their multiplicity and the rapidity gaps in diffractive dissociation. A description of the readout, a bi-dimensional pattern of strips and pads, is given in fig. 2.3.

The readout granularity of pads will be used for triggering, giving a first level of background suppression. It is possible indeed to combine (in a roughly pointing arrangement) the signals coming from aligned *super-pads*¹ of ten collinear chambers of on quarter of telescope. Fig. 2.4 shows the pattern of these *super-pads* with one example of a low level trigger road. A background event, not generated close to the interaction point, will be outside of any of the possible roads and it will be excluded.

Operating Condition [6]

The basic measurements of the TOTEM experiment doesn't requires high luminosity. It is however useful to have the possibility to perform measurement up to $L \approx 10^{33} cm^{-2} s^{-1}$ where studies on hard diffraction and forward physics can be done with the joint CMS/TOTEM experiment [24]. In these conditions, the environment in the forward direction where T2 is placed will be strongly radiation hard. This means that high rate capability and radiation hardness of detectors and electronics are required.

A simulation at $L = 10^{34} cm^{-2} s^{-1}$ of the hadronic and neutron rates in the T2 region is shown in fig. 2.5 together with the evaluation of the dose received by the T2 electronics. Form these plots a charged hadrons rate of $\approx 1 - 2 MHz/cm^2$ and a dose received by electronics of $\approx 5 Mrad/year$ is expected at $L = 10^{33} cm^{-2} s^{-1}$. Actually, the large number of high-energy neutral particles that emerge in the direc-

¹The 1560 pads of each detector are grouped in 104 *super-pads* of 15 pads each. The trigger signal of the single *super-pads* will be

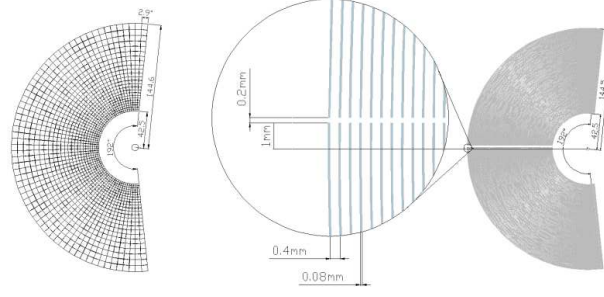


Figure 2.3: Readout pattern of the forward inelastic telescope $T2$. The readout is realized on a multilayered pcb, with a pattern of pads and strips for a two-dimensional readout. It covers a $\Delta\phi = 192^\circ$ and a $\Delta\eta = 1.2$ ($5.3 \leq |\eta| \leq 6.5$). The pads pattern (left) is realized with 1560 pads. Each pad is characterized by a $\delta\phi \sim 2.9^\circ$ and a $\delta\eta \approx 0.05$. They are subdivided in 65 sectors of 24 pads with area that varies between $\sim 2\text{mm} \times 2\text{mm}$ and $\sim 7\text{mm} \times 7\text{mm}$.

The strip pattern (middle-right), is realized with two sectors of 256 concentric rings, $80\mu\text{m}$ wide with $400\mu\text{m}$ pitch, providing the radial coordinates of traversing tracks with good precision ($\approx 100\mu\text{m}$). The strips have been divided into two parts to reduce the occupancy. The different $\Delta\phi$ of the two strips sectors avoids any dead zone when the detectors are mounted in *back to back* configuration.

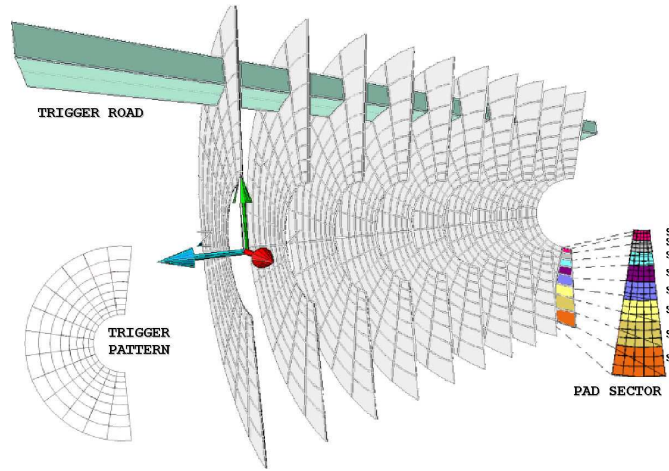


Figure 2.4: Trigger pattern of the forward inelastic telescope $T2$. The pads of the readout plane are grouped in 104 *super-pads* with the pattern shown in the left side. The *super-pads* of the ten cambers that constitute a quarter of telescope are aligned. A logical combination of the 10 collinear *super-pads* can be chosen in order to select good tracks, defined with a minimum number of hits inside a road. In the picture one road is shown. The grouping in 8 *super-pads* ($S1\dots S8$) of a sector of 120 pads is shown on the right side. Each readout plane has is divided in 13 sectors like this one.

tion of the $T2$ telescope, in particular neutrons and photons, interacts with the material and generates charged secondary particles (mainly electrons). It follows that a particular attention has to be payed on the quantity of material used to realize the detector. In fig. 2.6 a simulation at $L = 10^{28}\text{cm}^{-2}\text{s}^{-1}$ has been made using the final configuration for the $T2$ telescope. The secondary particles represent a large fraction of the charged particles hitting $T2$ ($\approx$ one order of magnitude larger than the charged hadrons) and should be considered as the main background source for the inelastic detectors.

An estimate of the rate requirement can be done using the results of the simulation shown and the specification of the $T2$ readout foil. At $L = 10^{33}\text{cm}^{-2}\text{s}^{-1}$, the average rate of charged hadrons is $\approx 2\text{MHz}/\text{cm}^2$ (fig. 2.5). Considering also the secondary charged particles produced by the neutral ones,

generated by a fast-OR combination of the relative pads.

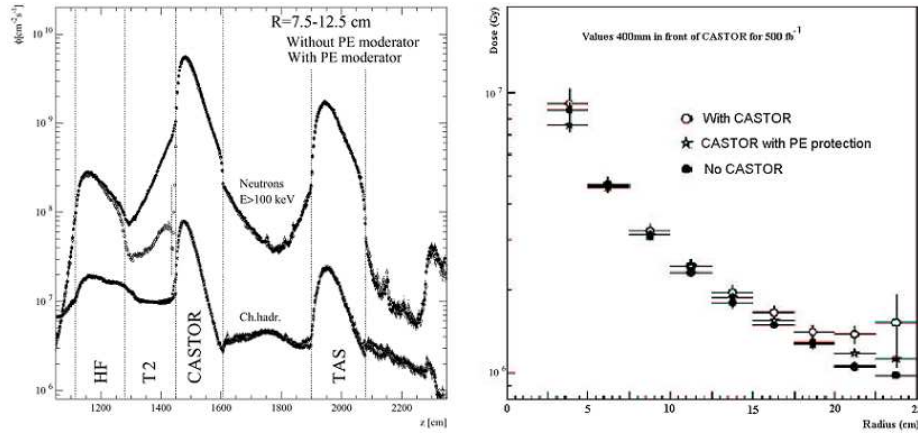


Figure 2.5: Average particle fluxes and radiation dose in the region of the forward inelastic telescope T2. Left: Average particle fluxes (between $R = 7.5$ and $R = 12.5$) per second at $L = 10^{34} cm^{-2} s^{-1}$ with and without the polyethylene disk between *CASTOR* and *T2* (*T2* at ~ 1350 to $1400 cm$). Extrapolation to the inner radius of the *GEM* ($\approx 42 mm$) increases the flux by a factor of 2. Right: Dose calculated for silicon in the *T2* region for an integrated luminosity of $500 fb^{-1}$. This correspond to about 5 years of operation at the maximal *LHC* luminosity $L = 10^{34} cm^{-2} s^{-1}$. The expected dose for *T2* electronics of $\approx 5 Mrad/year$ is obtained scaling the doses reported on the plot to one year at $L = 10^{33} cm^{-2} s^{-1}$.

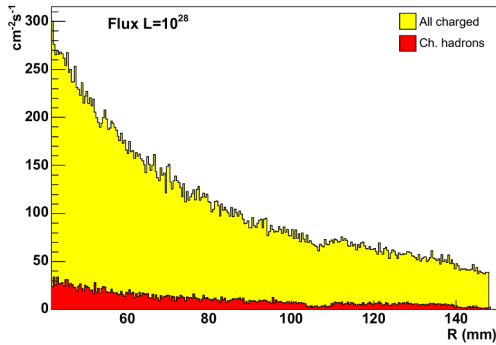


Figure 2.6: Flux of charged particles through *T2* at $L = 10^{28} cm^{-2} s^{-1}$: hadrons (dark gray) and all (light Gray) charged particles. For the larger pads (with an area of $\sim 0.5 cm^2$) there is a rate of charged particles of about $\approx 40 cm^{-2} s^{-1} \times 0.5 cm^2 = 25 Hz$. This means that it will be probably hit every $\approx 40 ms$. For the smaller the result is similar. The area is one order of magnitude smaller but the rate is one order of magnitude bigger. For a luminosity $L = 10^{33} cm^{-2} s^{-1}$, this time have to be divided by a factor 10^5 , i.e. two hit on the same pad will be every hundreds of *ns*.

we have to multiply by a factor ~ 10 (fig. 2.6). The sensitive area covered by one readout foil is nearly $320 cm^2$. This area has been divided in 1560 pads. If we consider that for each particle only one pad is hit, there will be a rate of $20 MHz/cm^2 \times 320 cm^2 / 1560 pad \approx 4 MHz/pad$. This means that a pad will be hit roughly every $250 ns$, i.e. every $10 pp$ collisions at the LHC. A similar result is obtained for the 512 strips, considering two strips hit per particle. The time response of the detector and of the readout electronics should be as short as needed to sustain these rates. The VFAT2 chip, used as front-end by all the TOTEM sub-detectors, has its time constants below $25 ns$ and the GEM gaseous detectors, chosen for the *T2* telescope, have shown no variation in performances at rate of $\sim 20 MHz/cm^2$.

High particles rate is not only a question of readout capabilities. Detector and electronics have indeed to sustain these fluxes without degrading. In the triple GEM detectors no aging effects have been observed for time period equal to the one needed to perform all the relevant measurements for TOTEM and the front end electronics chosen has been tested up to $10 Mrad$ preserving their normal operation.

With the use of a gaseous detector it is necessary to take into account the occurrence of discharges and their possible damages. With the particular triple GEM configuration, where three GEM foils are used as amplification stages, the rate of discharge is strongly reduced. Moreover, thanks to the fact that the gain (GEM foil) is decoupled from the readout (2-D readout foil), the deterioration after discharge is minimized. In fig. 2.7, tests performed on the *COMPASS* [30] Triple GEM² shows that with gain of about

²The TOTEM triple GEM detector has been developed on the basis of the COMPASS one. The detector structure is practically

10^4 , the discharge rate is less than 10^{-12} per particle. In our environment, 10^{-12} is a huge value and the detector has to be well below this level. This will consequently limit the possibility to increase the gain unconditionally. For a charged particles flux of $20\text{MHz}/\text{cm}^2$ for example, $\sim 6.4 \times 10^9$ part./s will pass in the $\approx 320\text{cm}^2$ area of T2. In one year ($\sim 3.15 \times 10^7\text{s}$) this means 2×10^{17} particles. If a discharge rate of $10^{-12}\text{disch.}/\text{part.}$ is considered, ≈ 500 disch./day will happen.

Even if the rate is maintained well below $10^{-12}\text{disch.}/\text{part.}$, there will be discharges inevitably and a protection circuit is required on the input channels of the readout chips.

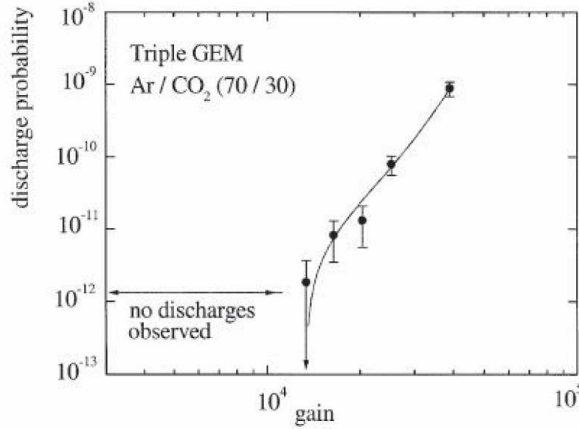


Figure 2.7: Measurement of the discharge probability as a function of gain on a *COMPASS* triple GEM. The gain needed in TOTEM is around 10^4 , where the discharge rate is less than 10^{-12} .

All the previous results has been computed for a luminosity $L = 10^{33}\text{cm}^{-2}\text{s}^{-1}$. For the TOTEM basic measurements, everything have to be re-scaled. This means that in general, the basic requirements are less stringent than the ones listed.

2.2 Triple GEM Detector

A gaseous detectors use the ionization produced in the gas by charged particles or photons to detect their passage. The signal induced by the number of electrons released naturally is too small to be read directly. An amplification is needed. The simplest way is to give to the electron produced during the first ionization enough energy to ionize again the gas molecules. This energy can be transmitted to these electrons with high electric field. In a GEM-based detector this is done with the GEM foil (fig. 2.8).

The high electric field is produced between the copper cladding on both side of the dielectric foil and the entry and exit paths to these high electric field regions for the ingoing electrons are created etching true holes on it. A simulation that maps the field intensity inside the hole is shown in fig. 2.9 where one example of an electrons multiplication is given.

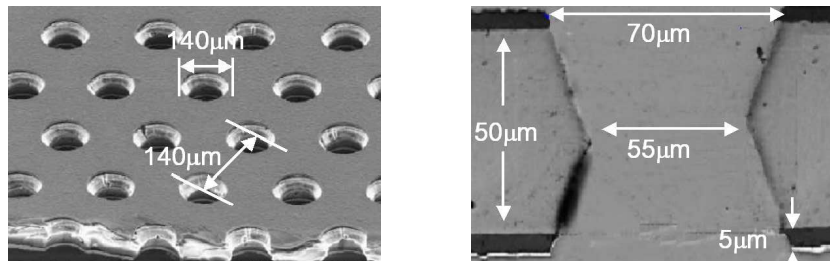


Figure 2.8: The GEM foil: the foil shown is made with a $50\mu\text{m}$ polyimide foil with $5\mu\text{m}$ copper cladding on both sides. The diameter of the holes is less than $100\mu\text{m}$ and the pitch is $140\mu\text{m}$. Different thickness, hole dimensions and patterns can be realized according to the specific application. In this foil the bidirectional wet etching process caused the double conical shape of the holes.

the same except for the semi-cylindrical shape and for the readout pattern

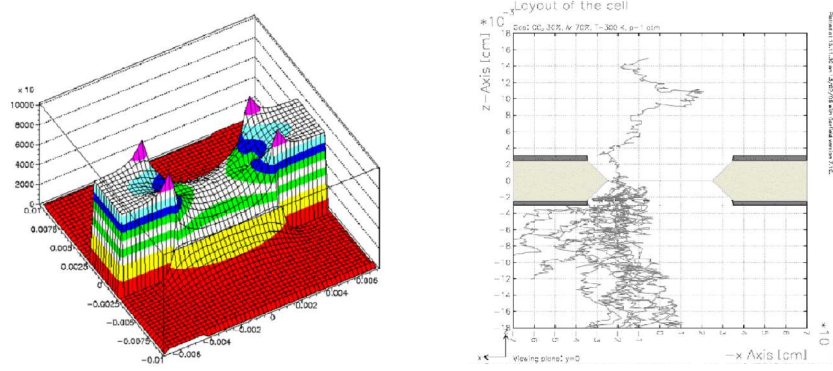


Figure 2.9: Left: Simulation of the electric field intensity inside one hole (left). A ΔV of few hundreds of Volts across the $50\mu\text{m}$ GEM foil produce very intense field in the hole that can reach a lot of tenths of kV/cm . Right: one simulation of an electrons multiplication process. The corresponding ions produced are not shown.

The multiplication factor of one electrons that travel in the hole is fixed by the electric field in the holes. The real gain of one GEM foil instead has to take into account also of the transparency³ of the foil, that can be described by its efficiency of electrons collection and extraction. These parameters depend on the field lines generated by the electric field outside and inside the hole as it is shown in fig. 2.10.

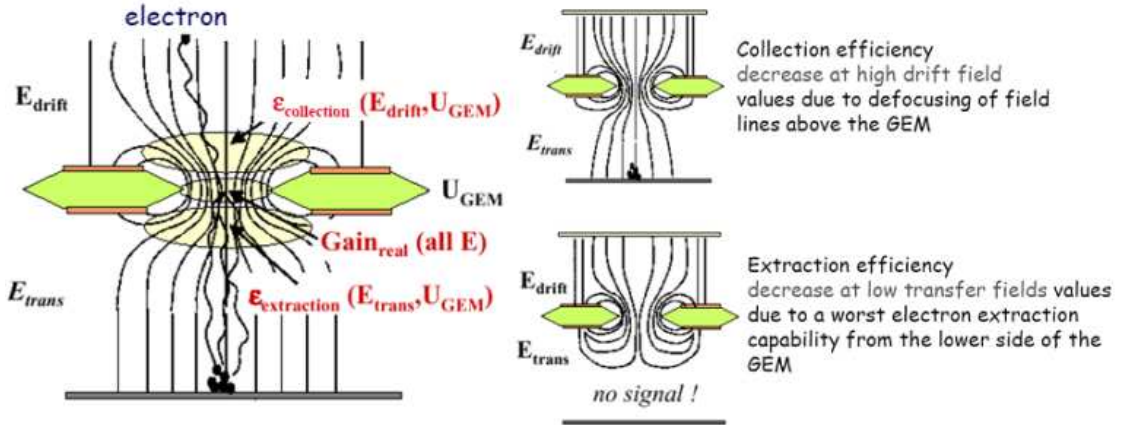


Figure 2.10: The real gain of a GEM foil [29] depends on the intensity of the electric field inside the hole and on the surrounding field lines. A reduction of the gain indeed is due to the capture of incoming electrons from the top metal layer of the GEM foil if the external electric field is so high that the electrons are hardly focused in the hole. Analogous scenario is for the outgoing electrons that can be captured by the GEM bottom layer if the external field is not enough high to move them away.

The properties of a GEM foil can be used in all the gaseous detectors where an electron multiplication is needed. The simplest GEM-based detector configuration is a chamber with a sensitive volume, where the first ionization takes place, a GEM foil to amplifying the signal and a readout plane. The use of several foils however is preferred because improves the performance of the detector. The gain sharing between the foils enhances indeed the reliability of operation at high gains and reduces the effects of the ions feedback. The internal structure of a triple GEM detector is schematized in fig. 2.11. It is characterized by a *drift zone* (sensitive volume), followed by the alternation of three GEM foils and two *transfer zones* through which the electrons drift before reaching the *induction zone*, where a signal will be induced on the readout foil. As previously reported, one of the advantages of the use of more than one GEM foil is the reduction of the discharge probability. If only one foil is used, the multiplication factor requested to each hole could be quite high. The high fields and the high local electron density inside the hole could

³ratio between ingoing and outgoing electrons in unitary gain condition.

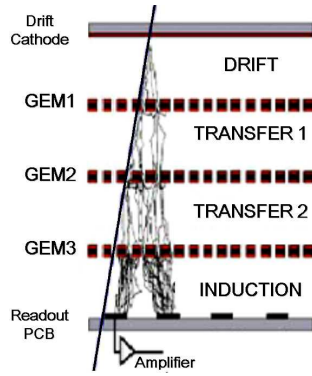


Figure 2.11: A triple GEM configuration where the total gain is shared among three GEM foils. The sensitive volume corresponds to the *drift zone* because the induced signal on the readout will depend mainly on the number of primary electrons released by the charged particle in this zone. The primary electrons produced in the other zones indeed induce a smaller signal because they are not submitted to all the three amplification stages.

produce discharges too frequently. This rate can be strongly reduced dividing the gain between more foils. In this way there will be a sharing of the total gain between different holes, reducing the discharge probability. The other reported improvement of multiple foils configurations is the suppression of spurious signals generated by ion-feedback currents. With a single GEM foil, the neutralization of the large amount of ions produced during the multiplication in the holes will occur mainly on the drift cathode. For energy balance, this process could be followed by secondary emission of another electron that will drift toward the GEM where it will be multiplied producing a delayed spurious signal on the readout. The probability of this phenomena is low and therefore it is evident only when a large quantity of ions is neutralized on the drift foil (as in the single GEM foil case). In a triple GEM, practically only the ions produced in the first GEM will reach the drift foil, while the others will be neutralized mainly in the bottom of the first and of the second GEM. In this way the greatest part of the secondary electrons produced by the ions will be multiplied by only two or one GEM foil and their signal will be smaller than the one produced by primary electrons.

2.2.1 The TOTEM Triple GEM

A gaseous detector based on GEM (gas electrons multiplier) technology is typically characterized by:

- high rate capabilities ($> 10^5 \text{ Hz/mm}^2$);
- radiation hardness;
- good time ($< 10 \text{ ns}$) and spatial ($< 60 \mu\text{m}$) resolution;
- active area up to 10^3 cm^2 ;
- reduced constraints on the detector shape;
- readout fully decoupled from amplification (very interesting aspect because it allows an independent optimization of the functional parts of the chamber).

All these properties are compatible with the TOTEM requirements for the telescope *T2*. The collaboration has decided therefore to realize its inelastic telescope with a GEM detector, based on the successfully triple GEM of COMPASS [30].

The TOTEM triple GEM (fig. 2.12) is realized using the structure shown in fig. 2.11. Three GEM foils (fig. 2.13) are used to achieve a nominal total gain of about 8000. Higher values can be safely used if needed. The voltages to the drift and GEM foils are provided by a resistive high voltage divider (see fig. 2.15).

The gas mixture used is Ar/CO_2 70/30 even if it has been foreseen the possibility to add CF_4 (to have faster signals).

The sensitive area has a semi-cylindrical shape with an azimuthal coverage of $\Delta\Phi = 192^\circ$ and a radial extension from $R_{\min} \sim 42.5 \text{ mm}$ and $R_{\max} \sim 144.5 \text{ mm}$. The sensitive volume of the *drift zone* has a thickness of 3 mm to have the maximum efficiency of detection for a MIP (minimum ionizing particle). The two *transfer* and the *induction zone* are 2 mm thick.

A two dimensional readout foil (see fig. 2.14) is foreseen with a pattern of strips for accurate tracking and pads for azimuthal measurement and triggering. The front and back plates are honeycomb structures that ensures a robust mechanical structure with low material budget. In table 2.1 the list of the low-Z

material used for a single triple GEM with the relative radiation length is given.

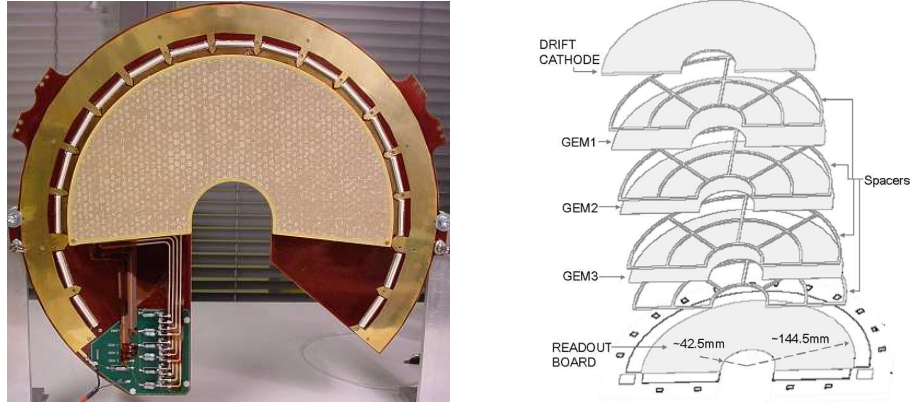


Figure 2.12: Left: a TOTEM triple GEM used in the *T2* telescope. The high voltage divider (see fig. 2.15) is placed on the bottom-left part of the detector. The gas connectors are installed on the readout board through channels engraved in the fiberglass frame beneath the polyimide foil. The frames of the GEM foils contain holes in two corners for a uniform distribution of the gas inside the chamber. The readout front end will be connected on the seventeen high density connectors all around the chamber. The front and back plates are honeycomb structures, in which a honeycomb sheet of thickness 3mm (Nomex) is sandwiched between two thin *FR4* sheets of thickness $125\mu\text{m}$ and enclosed inside a supporting frame made of the same material as the frames of the foils. Right: A drawing of the internal components of the TOTEM semi-cylindrical triple GEM.



Figure 2.13: Left: The drift foil. It is realized with a polyimide foil with a $5\mu\text{m}$ copper cladding on a single side with thickness of $50\mu\text{m}$ glued to the front plate. Right: The GEM foil. It is realized with a $5\mu\text{m}$ copper cladding on both side with thickness of $50\mu\text{m}$. The shape of the holes is double conical like in fig. 2.8. The diameters of the holes in the middle of the foil and on the surface are 65 and $80\mu\text{m}$ respectively. The distance between them is $140\mu\text{m}$. The foil is characterized by a single plane in one side and by four electrically separated ring segments (visible in the picture) in the other. This reduce the risk of damage in case of discharge because not all the energy stored in the GEM foil will be released, but only the part of the sector involved. The foils are stretched and glued over supporting frames, which are manufactured by Computer Numerical Control (CNC) machining from fiberglass reinforced epoxy plates (Permaglas) with thicknesses of 2mm . Two additional supporting spacers of thickness 0.5mm are designed in the middle of the frames. Their position is slightly asymmetric to minimize dead areas. In the picture it is shown one foil with its spacer.

In the next section a description of the front end electronics will be given. It is therefore helpful to close this section with a brief discussion on the expected signals from a detector like the TOTEM triple GEM.

Firstly a very rude estimate of the charge expected on the input of the front end is given. The numbers of electrons produced in the *drift zone* depends on the energy released by the charged particles and on

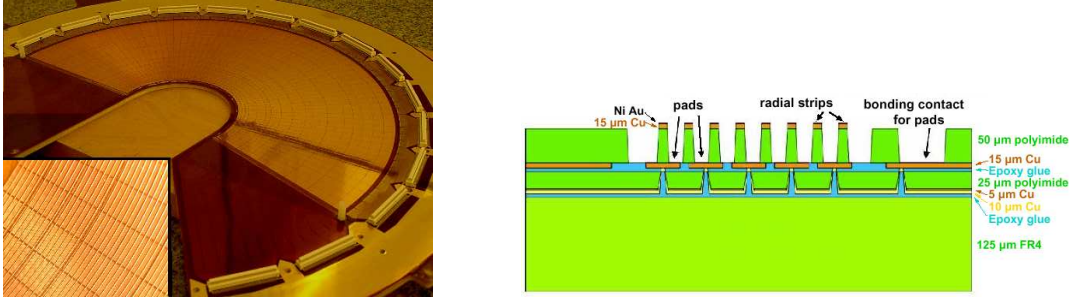


Figure 2.14: Left: The readout plane is realized on a printed circuit board covered by polyimide foil with the pattern of strips and pads shown in fig. 2.3. The readout board contains 2×256 concentric strips for the radial coordinates and a matrix of 1560 pads for azimuthal coordinates and for the T2 local trigger. The width and spacing of the strips are 80 and $400 \mu\text{m}$ respectively. To reduce the occupancy, the strips are divided into two parts. The pads are divided into 65 radial sectors each containing 24 pads with sizes ranging from $2 \times 2 \text{ mm}^2$ in the internal side (the one that is close to the vacuum chamber of the beam pipe) to $7 \times 7 \text{ mm}^2$ on the external side. The charge collected is read at the outer edge of the readout board where 17 high density 130pins connectors are foreseen. In the corner a view of the strips and pads is shown and a complete description of the readout pattern is given in fig. 2.3. A measurement of strips and pads capacitance is shown in fig. 2.18. Right: Detailed view of the strips and pads structure with dimensions and materials used. The strips lie on top of the pads and are isolated from the pads by a thin layer of polyimide, which is removed between the strips by wet etching.

Item	Details	$^{\circ}/_{oo}$ of Xo
3 GEMs	$6 \times 5 \mu\text{m}$ copper [0.7]	1.68
	$3 \times 50 \mu\text{m}$ kapton [0.7]	0.42
	TOTAL	2.10
1 Drift	$5 \mu\text{m}$ copper	0.35
	$50 \mu\text{m}$ kapton	0.17
	TOTAL	0.52
3 Grid spacers	$3 \times 2 \text{ mm}$ fiberglass [0.008]	0.25
1 Readout board	$80 \mu\text{m}$ strips: $5 \mu\text{m}$ copper [0.2]	0.07
	1560 pads: $5 \mu\text{m}$ copper [0.85]	0.26
	$50 \mu\text{m}$ kapton [0.2]	0.03
	$120 \mu\text{m}$ fiberglass	0.62
	$60 \mu\text{m}$ epoxy	0.30
	TOTAL	1.28
1 Shielding	$10 \mu\text{m}$ aluminium	0.11
2 Honeycombs	$2 \times 3 \text{ mm}$ Nomex	0.46
4 Fiberglass foils	$4 \times 120 \mu\text{m}$ fiberglass	2.47
TOTAL		7.19

Table 2.1: Material budget in one TOTEM triple GEM detector.

the energy required to ionize. The gas used in our detector is Ar/CO_2 (70/30). Weighting properly the properties of the two components, the mean energy released by a *MIP* in this mixture is of about $2.6 \text{ keV}/\text{cm}$. The sensitive volume (*drift zone*) has a thickness of 3 mm and therefore the energy lost by charged particles is $\sim 800 \text{ eV}$. The Ar/CO_2 averaged energy required to ionize is roughly 28 eV and a mean of ~ 30 electrons will be produced. Actually, the energy loss follows a Landau distribution and the most probable value of 10 – 15 electrons will be considered in the next instead of the mean.

The charge collected by each electrode depends on the triple GEM gain and on the layout of the readout board. With a gain of 8000 the numbers of electrons that will induce signal on the readout board will be $\approx 10^5$. The transversal diffusion coefficient of one electron in Ar/CO_2 with a longitudinal electric field of $3 \text{ kV}/\text{cm}$ (compatible with the ones used) is $\sim 260 \mu\text{m cm}^{-1/2}$. The electrons produced in the *drift zone*

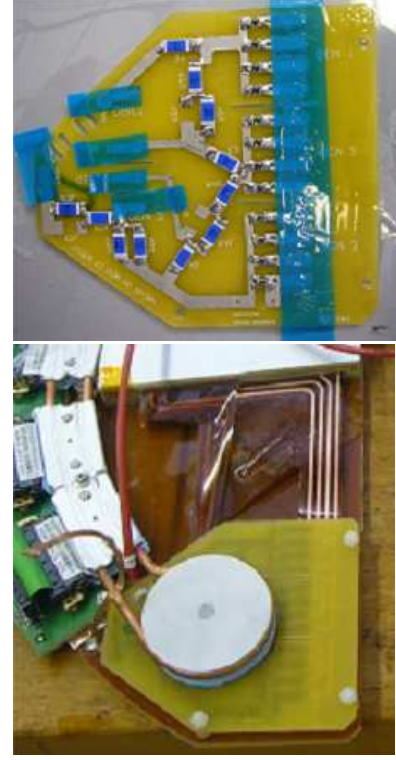
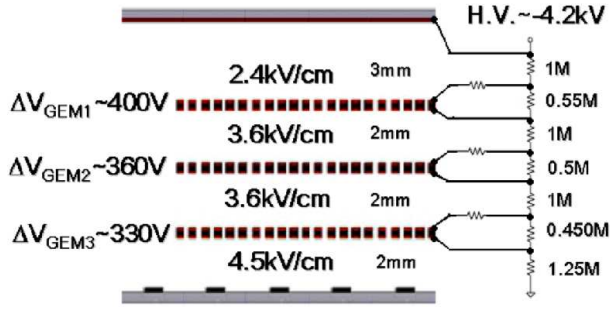


Figure 2.15: Left: Electric fields inside the chamber. The ΔV across the GEM foils is decreased from the first to the last foil to reduce the discharge probability. The drift field is taken smaller than the others to have a good charge collection in the first foil. An higher value is instead supplied across the induction zone to increase the extraction from the last foil and the rising time of the induced signal.

All the electric field inside the chamber are provided by the resistive divider shown on the right. It is possible to use multi-channels floating high voltage power supply to apply/control the voltage to each part independently (a passive divider fixes univocally the voltages ratios once the resistor are placed). The solution adopted however is preferred because it maintains the right voltage differences between all the electrodes in the on/off procedure. Each sector of the GEM foil is connected with a series $10M\Omega$ resistors to the divider, that will limit the current in case of discharge. Moreover if a discharge causes a short circuit in one sector, these resistors preserve the working operation of the other maintaining the right ΔV across the foil. Left: H.V. divider realized with *SMD* resistors. A cooling of the board is foreseen to remove the heat ($5/6Watts$ dissipated in normal operations) produced by the resistors.

have a distance from the readout board between 6 and 9mm. Using the maximum value, the transversal distribution of the electron cloud will be $\sigma_x \sim 260\mu m cm^{-1/2} \times \sqrt{0.9cm^{-1/2}} \sim 250\mu m$. The pitch between strips is $400\mu m$ and the area of pads is between 2×2 and $7 \times 7mm^2$. From the founded σ_x it follows that roughly 2 – 3 strips and 1 – 2 pads will be reached by the electron cloud. The readout pcb has been designed in order to have a charge sharing of $\approx 50\%$ between pads and strips. The total charge collected by one electrode (strip or pad) will be few tenths of thousands of electrons: $\approx 1 - 5 \times 10^4$ electrons/electrode $\approx 1 - 5fC$ /electrode for the most probable value of 10 – 15 electrons/MIP.

The temporal length of the signal depends on the travelling times of the electrons produced in the *drift zone*. The induction of the signal starts when the electrons closest to the readout (i.e. the electrons close to the top of the first GEM foil) will exit from the last GEM foil. This distance is 4mm. The induction of the signal ends when the electrons more far from the readout board (i.e. the electrons close to the bottom of the drift cathode) are collected on the electrodes. This distance is 9mm. The drift velocity of one electron in Ar/CO_2 with a longitudinal electric field of $3kV/cm$ is $\sim 7cm\mu s^{-1}$. The duration⁴ of the signal will be therefore $(0.9 - 0.4)cm/7cm\mu s^{-1} \approx 70ns$.

A signal of $1 - 5 \times 10^4$ electrons, with a duration of $\approx 70ns$ is what is expected for the most probable

⁴This has not to be confused with the rising time of the signal, that is obviously smaller

energy loss in the *drift zone* (i.e. 10 – 15 electrons). Efficiency requires that even the signal produced by few electrons produced in the sensitive volume have to be detected and this means that the signal of few thousand of electrons collected on the electrode must be outside the noise. Moreover according to the Landau energy loss distribution, even 100 electrons or more can be produced in the *drift zone*. This imposes a strong condition on the dynamics of the readout electronics that has to cover a large range, between few thousand to hundreds of thousand of electrons.

2.3 On detector electronics

The measurements of the TOTEM experiment are performed by three different detectors: RPs (silicon), GEM and CSC (gaseous). As reported in table 2.2 they have different characteristics and requirements.

	RP	T1	T2
No. and type of detectors	240 Si strip	60 CSC	40 GEM
No. of channels	122880	11124 anodes, 15936 cathodes	62400 pads, 20480 strips
Typical input charge ^a	$\sim 4fC$	$\sim 50fC$	$\sim 50fC$
Signal Polarity	+	+/-	-
Occupancy	< 1%	anodes: < 10%, cathodes: < 20%	pads: < 5%, strips: < 30%
Radiation Dose	< 10Mrad	< 50krad	< 50Mrad

Table 2.2: Overview of electronics requirements from the different detectors. (a) This charge is a MIP mean value (i.e. ≈ 30 electrons/MIP) and it has to be divided by the number of the readout channels involved to have the total charge for a single front end input.

Despite this, the collaboration has decided to use a common front end with all the three types of detectors. This choice is motivated by the fact that a common chip means identical control, trigger and data readout chains. The front-end ASIC designed for this purpose is the VFAT2.

The two basic functions of the chip are:

- Triggering: provide fast regional hit information to aid the creation of a first level trigger LV1.
- Tracking: provide an high density of channels to obtain a precise spatial hit information for a given triggered event.

Fig. 2.16 shows the block diagram of the VFAT2. *The VFAT2 chip (as reported in the Operating Manual [32]) has 128 identical channels. It is a synchronous chip designed for sampling sensors at the LHC clock frequency of 40MHz. Each channel consists of a preamplifier and shaper followed by a comparator⁵. If a particular channel receives a signal greater than the programmable threshold of the comparator a logic 1 is produced for one clock cycle only by a monostable⁶. This logic 1 is written into the first of two SRAM memories (SRAM1). All other channels that do not go over threshold record a logic 0 in SRAM1. This occurs in parallel for all 128 channels at 40MHz. At the same time a fast OR function can be used to set a flag which can immediately be used for creating a trigger. It is foreseen to have up to eight programmable sectors which can be flagged with the fast OR in this way. ... On receiving a LV1A signal, data corresponding to the triggered time slot is transferred to a second SRAM memory (SRAM2). ... As soon as SRAM2 contains data the Read cycle begins. ... The chip operates with a continuous write/read operation without dead time.*

Fig. 2.17 shows a photograph of the chip with its layout and in table 2.3 the technical specifications of this front end ASIC are listed. Detailed documentations on the VFAT2 can be found in the electronics pages of the TOTEM home-page. In this section few aspects related to the use of this front end with the TOTEM triple GEM are discussed.

⁵The comparator is an asynchronous comparator without hysteresis. On passing a programmable threshold the comparator output goes high and returns low again when descending back through the threshold. For very large signals the comparator output may remain high for more than one clock cycle. Also if the signal barely passes threshold the comparator output may go high for less than one clock period.

⁶One clk cycle is the default setting. It is possible to stretch the output of the monostable up to eight clock cycles.

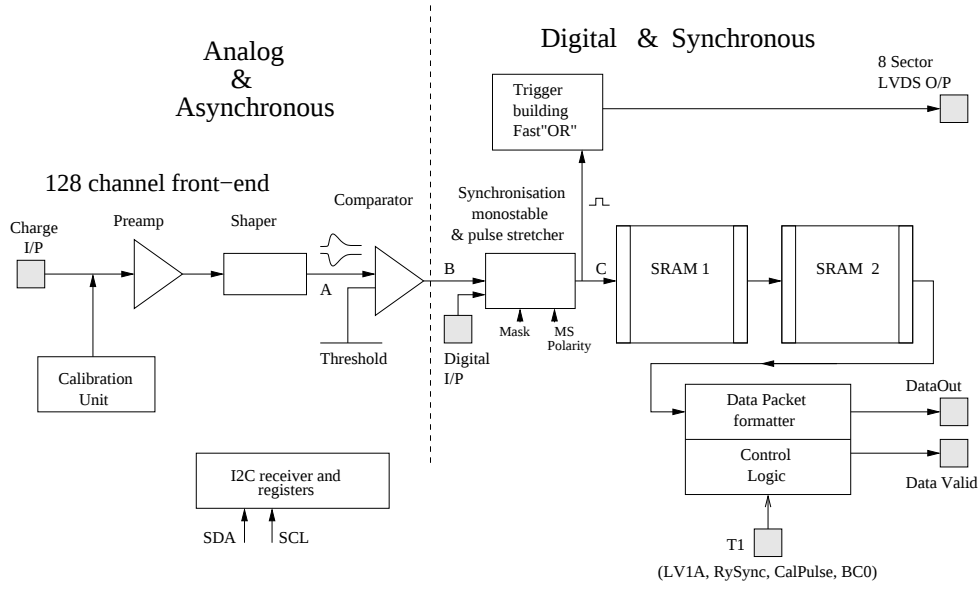
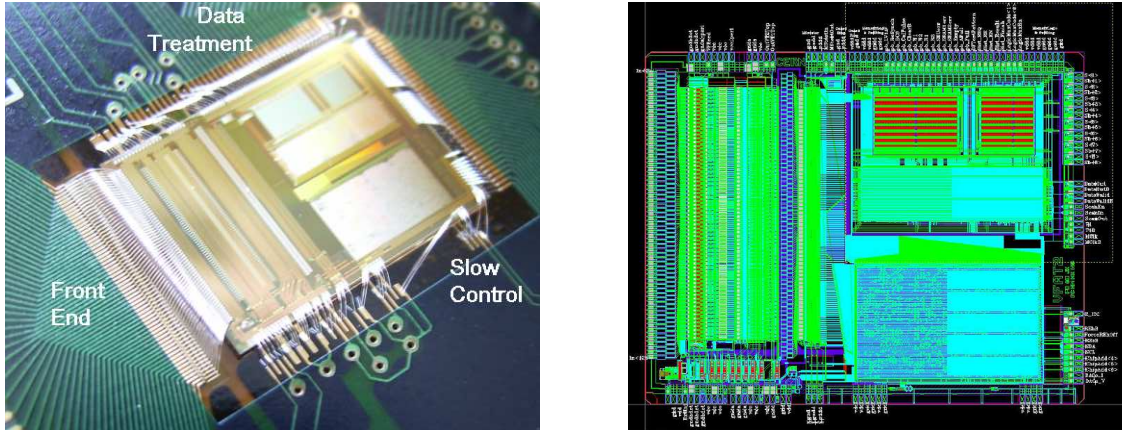


Figure 2.16: Block diagram of the VFAT2 chip [7].

Figure 2.17: A photograph of the VFAT2 chip designed in $0.25\mu\text{m}$ CMOS with its layout.

The first point concern the theoretical noise level expected for a detector with a readout layout like the one used for the TOTEM triple GEM. The input equivalent noise charge of the chip depends on the capacitance of the detector's strips and pads following the relation indicated in table 2.3 ($ENC \sim 400e^- + 40 - 60e^-/pF$). In fig. 2.18 the measured capacitance for the two type of electrodes of the TOTEM triple GEM readout board are reported. The ENC should be between 800 and $2200e^-$ if the detector capacitance varies from 10 to $30pF$.

The results obtained experimentally have shown higher value (even if acceptable) and this means that a more realistic model of the readout circuit could be useful to better understand the noise performance of VFAT2 when used with the TOTEM triple GEM. The noise picked-up electromagnetically is not included in the ENC relation and this is surely an important components for a readout with the dimension of the T2 telescope. A big improvement on the noise level have been obtained indeed with a proper shielding of the detector. Nevertheless, explanations of the measured noise could be probably found analyzing more accurately the input stadium of the preamplifiers with the detector readout electrodes connected. The capacitive coupling between the strips and their underlying pads (see the readout layout in fig. 2.3 and 2.14 and fig. 2.19) is the main component on the measured capacitance (fig. 2.18). This will influence the cross-talk between the channels. The input stage of the VFAT2 has been designed to have a low input impedance in order to reduce the cross talk and the attenuation of the signal. Fig. 2.20 shows a simulation

Number of channels	128
Gain	$60mV/fC$
Front- end shaping time	$22ns$
Time walk (for 1.2 to $10fC$ with $1fC$ threshold)	$12ns$
ENC	$\sim 400e^- + 40 - 60e^-/pF$
Linearity	$\pm 12fC$
Internal Calibration test pulse	$-2fC$ to $18.5fC$, $LSB = 0.08fC$ with $\sigma(LSB) = 0.3fC$
Sampling frequency	$40MHz$
LV1A Latency	Up to $6.4\mu s$
Storage capacity	128 triggered events
Slow Control interface	I^2C
Testability features	Scan Chain, BIST, Probe pads, Auto test patterns, Auto Data Packet
Power Consumption	$168mW$ (Sleep Mode), $572mW$ (Run Mode)
Radiation Resistance	$< 10MRad$

Table 2.3: VFAT2 specification.

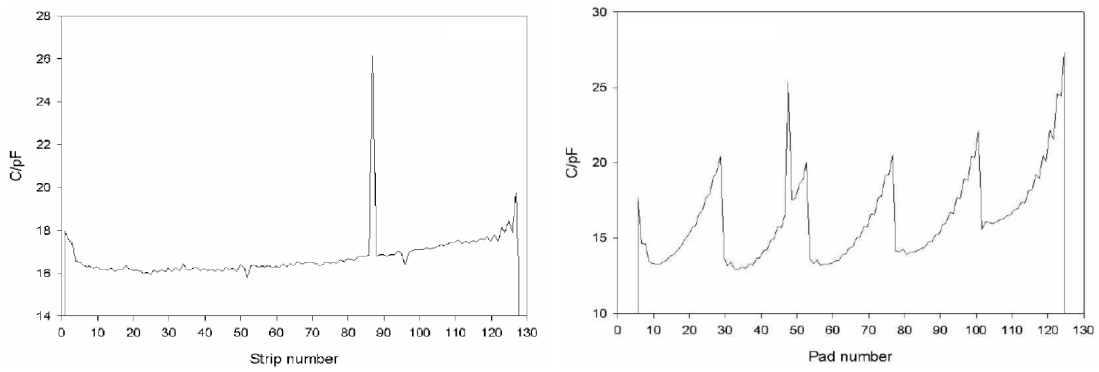


Figure 2.18: Capacitance measurement for strips (top-left) and pads (top-right). The anomalous spike is caused by a short between strips and pads. The capacitance has been measured between each strip or pad of one sector and gnd, with all the others electrodes floating. One sector is made by 128 strips and 120 pads. The length of the strips under the sensitive area increase as the number of the channel while the length of the strips fan-out (lines that connect the strips to the readout front end) decrease increasing the channel number. This explain why there aren't big differences on the capacitance of the various strips. Completely different is the behavior for pads. Each sector with 120 pads has five sub-sector with common ϕ and increasing θ . The numbering of the channels is made sector by sector increasing θ , i.e. from smaller up to larger area pad. Each sub-sector contains 24 pads and this explain the shape of the capacitance measurement with five repeated structures that reflects the area of the pad. In this case the contribute of the pads fan out is lower.

of the input impedance for the ABCD ASCIC that has the same analog front end of the VFAT2.

The noise modelling is moreover affected by the additional DC coupled protection circuit⁷ added between the detector and the input channels of the analog front end [35]. In a gaseous detectors discharges could occur. Even if the most probable discharge in a GEM-based detector is between the two electrodes of the GEM foil, it is possible to have discharges that reach the readout pcb. The protection circuit (fig. 2.21) has been therefore inserted to avoid possible damages. It has been realized using two diodes and a resistors. High currents produced by discharges will flow through the diodes without entering the pre-amplification stage of the VFAT2. The value of the impedance added in series (a resistor

⁷The COMPASS detector has also a protection circuit in the input. They used two diodes in back to back configuration and AC coupled the input with a series capacitors of about $220pF$.

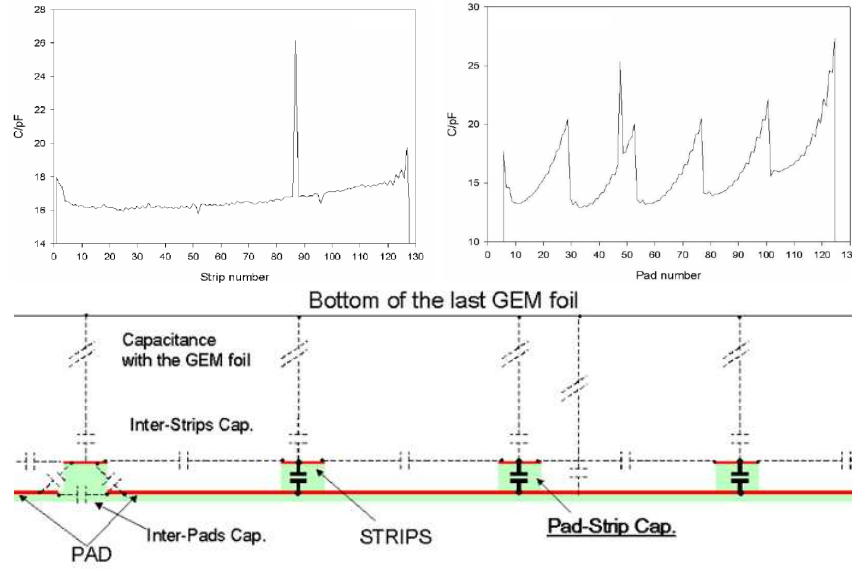


Figure 2.19: For the geometrical layout, the major contribute to the capacitance (see fig. 2.18) of a pad or a strip is given by the capacitance between them as schematized in the drawing.

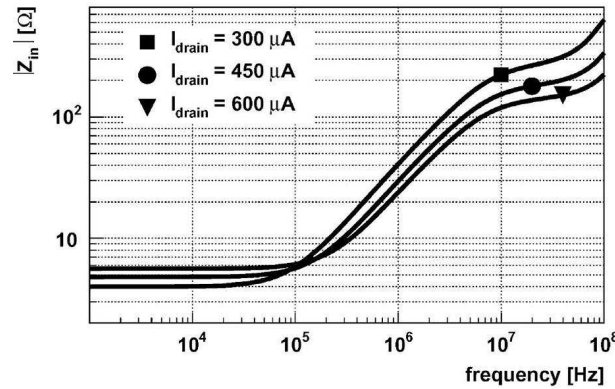


Figure 2.20: Input resistance of the preamplifier simulated for $800nA$ bias current of the feedback transistor and various input transistor bias currents [34]. The impedance of the preamplifier is of the order of 120Ω to 200Ω at the central frequency of the shaper bandwidth (Shaping time $\sim 22ns$). The plot is referred to the ABCD ASIC, from which the VFAT2 has been developed.

smaller than 10Ω) doesn't affect practically the input impedance previously shown.

Finally, the level of the noise is affected by the connection between the input channels of the chip and the detector readout electrodes that is not realized directly with bonding, as it was for example in the COMPASS triple GEM, but with an high density connector inserted in between.

Fig. 2.22 shows the hybrid where the VFAT2 is bonded to the female 130-pin white connector that will fit on the male mounted on the bottom plane of the triple GEM.

Strip-pad capacitance, protection circuit and input connection through connectors could affects the noise. If properly inserted in the modelling to define the theoretical noise level, possible improvement could be found.

As reported in the description of the VFAT2 (see footnote), the internal monostable output can be stretched up to 8 clock cycles. This possibility is very important for the time properties of our signals. At the end of the previous section it has been shown that the typical signals of the TOTEM triple GEM

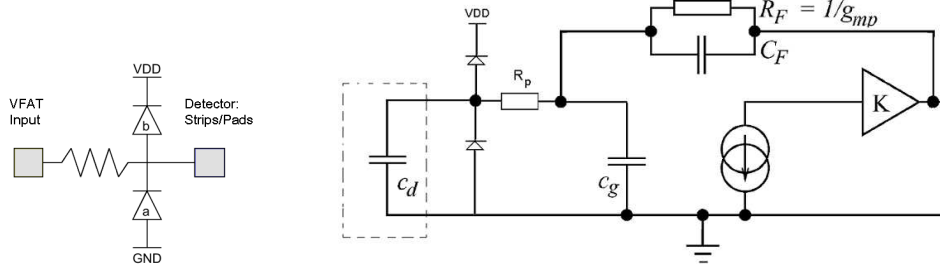


Figure 2.21: Left [35]: VFAT2 input protection circuit realized with two diodes between VDD and GND and with a series resistors. Diode a is made by a well of n-doped silicon into the p-type substrate; diode b is made by a p-type diffusion into an n-type well. The small resistor of about 7.5Ω in series increase the injectable charge without consequent damages on the chip. Right [34]: equivalent schematic of the MOS transimpedance preamplifier with the input protection and the detector equivalent capacitance C_d . The total gate capacitance C_g of the input transistor for the nominal bias is in the range of $4pF$ (the nominal bias current of the input transistor is in the range 300 to $600\mu A$, which allows for operating that device close to weak inversion.). The feedback is active and is realized with a transistor instead of a conventional resistor. It works in saturation and is biased in moderate inversion and at nominal condition its transconductance is equivalent to a $120k\Omega$ feedback resistor.

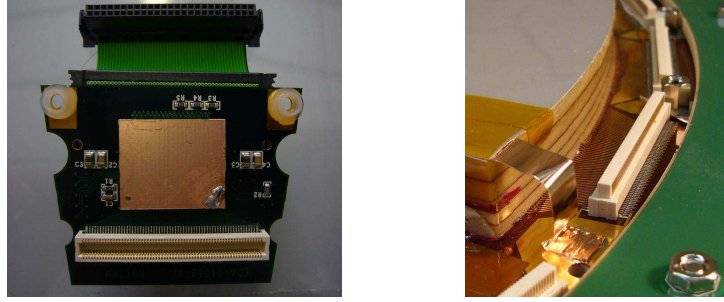


Figure 2.22: VFAT2 GEM hybrid (right). Two different hybrid versions have been realized for pads and strips because in the strip the trigger signals are not extracted from the VFAT2 and because each pad sector has 120 channels while the strip 128. Over the VFAT2 a grounded metal cover has been added to protect the chip and reduce the noise. The green flat cable brings all the digital data/trigger/control lines of the VFAT2. The white female 130-pin connector is connected to the input channels of the front-end chip and all the analog signals from the detector will flow through it. It will fit on the male (right) soldered to the pads and the strips on the readout plane of the triple GEM. Few lines on this connector are dedicated to the gnd connection between the detector and the front end. The digital and analog ground of the VFAT2 are merged together on the hybrids with zero ohm resistor.

have duration on the scale of $70ns$. The duration actually is not a problems because this affects only the occupancy of one channel and in this case it would mean a sustainable rate of $\approx 10MHz/ch..$ If a mean number of channels per square centimeter of about $3 - 5ch./cm^2$ is considered⁸, this would mean a sustainable rate of $30 - 50MHz/cm^2$ that is higher than the expected at luminosity up to $L = 10^{33}cm^{-2}s^{-1}$. The rising time distribution of the signal is however extremely important for us. On receiving a $LV1A$ signal, the VFAT2 transfer the interesting data contained in the triggered time slot of the first memory $SRAM1$ to the second $SRAM2$ memory from which the read cycle begin. The VFAT2 has written the information in $SRAM1$ at the clock cycle after the crossing of the threshold by the output signal of the shaper. If the *threshold crossing time* distribution is larger than one clock cycle, this means that events that occur at the same time will be stored in different time slots. Moreover, the $LV1A$ signal is generated with the same input signal and with the same analog front-end of the data and therefore it will suffer of the

⁸The readout plane covers $\sim 320cm^2$. Considering one pad hit per particle this means $1560/320ch./cm^2 \sim 5ch./cm^2$. With two strips hit per particle we will have $(2 \times 512)/320ch./cm^2 \sim 3ch./cm^2$

same time uncertainty. The stretching is fundamental for taking into account all these aspects. In a gaseous detector like the triple GEM these time spreads are present and they depend mainly in time walk and jitter effects, as shown in fig. 2.23.

The time walk is caused by the energy loss distribution of a passing charged particle. The number of

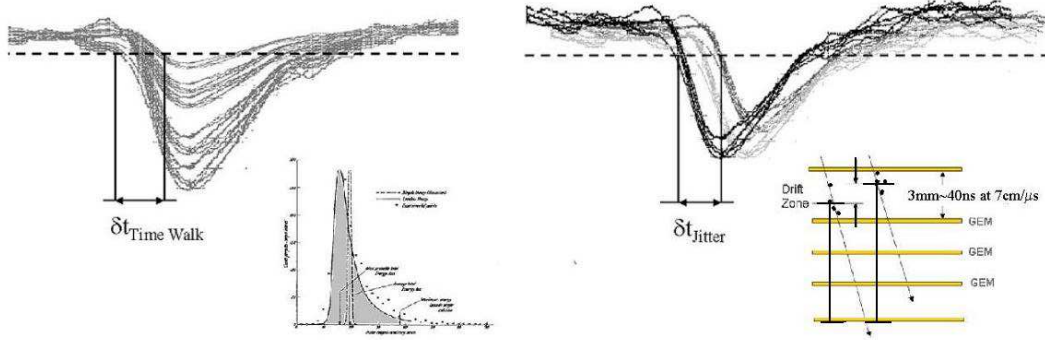


Figure 2.23: Right: Time spread induced by Time Walk and energy loss distribution. From few up to hundred electrons could be released. Left: Jitter effects and primary ionization.

electrons released in the sensitive volume has a very large range (at least one order of magnitude) and this is proportionally reflected on the intensity of the read signal. An increase of the signal to noise ratio is helpful for minimize these effects because the relative position of the threshold can be lowered.

The jitter of the signal produced in the detector is instead related to the distance between the readout plane and the groups of electrons that will start the induction of the signal. These electrons have been generated from the primary clusters closest to the first GEM foil. If their distance from the GEM has large variations from event to event, this will be manifested as a Jitter effects. With an electron drift velocity of about $7\text{cm}/\mu\text{s}$ ($E \sim 3\text{kV}/\text{cm}$, gas $\text{Ar}/\text{CO}_2 70/30$) the 3mm thickness of the drift zone is equivalent to $\sim 40\text{ns}$ and therefore a spread in the threshold crossing time of few ns is possible. A reduction of the noise doesn't help on the jitter. Neither the thickness of the sensitive volume has effects, because this time spread depends only the electrons that begin the induction. An improvement can be obtained instead increasing the number of primary clusters produced by the passing particle. If necessary, this can be done modifying the gas mixture. In general, higher will be the number of clusters produced by the passing particle, higher will be the probability that the position of the primary electrons that are closest to the readout plane doesn't vary to much. This means that the jitter will be smaller.

Even if the distribution of the rising time is pushed to σ on the scale of $5 - 10\text{ns}$, the synchronization

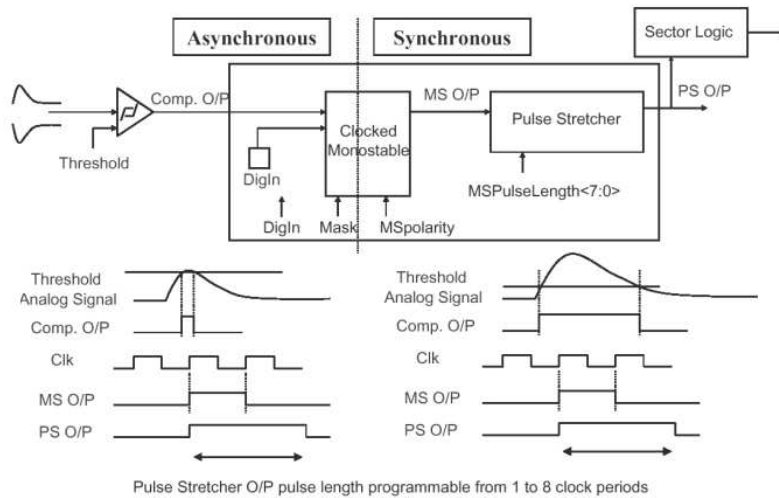


Figure 2.24: Pulse stretching of the monostable output [31].

with the *LHC* clock will probably require the use of at least two clock cycles as a temporal coverage of each memory slot. With the VFAT2 it is possible to take account of this possibility. The output of the monostable (that is driven by the comparison between the shaped signal and the threshold) can be stretched from one to eight clock cycles (fig. 2.24). The output of the programmable pulse stretcher is the information that is stored in the memory slot. Having defined a latency (the position in the memory where the triggered time slot have to be found) if a stretch of three clock cycles is used this mean that in the slot is stored the information on what happens in that clock cycle and in the two earlier. The limit on the stretching is imposed by the occupancy because a mixing of different events has to be avoided. An interesting particle that travel through *T2*, will pass ten detectors. A compromise will be found between the efficiency of each single detector (increased using large stretching to collect all the events) and the temporal "integration" of each memory slot.

Before closing this section, few comments about the dynamical range of the VFAT2 are given. The triple GEM detector requires a large range for the big variation on the number of primary electrons that can be released by the particle that has to be detected. In table 2.3 the linearity of the signal is guaranteed up to $\pm 12fC$. This is equivalent to 7.5×10^4 electrons. With a total gain of 8000 the GEM signal could reach in one channel 4^5 ($60fC$) electrons or more. Actually we are interested on presence of hits and not on the quantity of charge collected. The leak of linearity therefore is not a problem. Dead time related to saturation have instead to be carefully considered. A simulation on dynamical range and recovery time after saturation has been performed by the VFAT2 designers [33]. In table 2.4 the results for *RPs* and *GEM* are summarized.

The time duration of the signal for the GEM is helpful from this point of view. The total charge produced

VFAT2 Dynamic Range for RP	$\sim 18fC$ (5 <i>MIPs</i>)
VFAT2 Dynamic Range for GEM	$35fC$ to $45fC$ ≈ 2 <i>MIPs</i> for pad, ≈ 6 <i>MIPs</i> for strip ^a
Threshold range	RP $\sim 18fC$ Gem $\sim 30fC$
Trim-DAC range	$\sim 40\%$ of dynamic range
Recovery time from very large signals	$\sim 1\mu s$ with a $100fC$ ideal input signal. $\sim 10\mu s$ following $10pC$ of ideal input charge. $\sim 0.8\mu s$ for $1pC$ of GEM input charge

Table 2.4: VFAT2 Dynamical Range [33]. (a) Rude estimate obtained considering 30 electrons released by a *MIP* in the *drift zone*, a pad and strip gain of 4000 and 1 pad and 3 strips hit.

for each event is indeed distributed in a relatively long time period (long referred to the time constant of the analog front-end). For this reason the linearity range shown in table 2.3 is increased and the saturation effects are reduced. Fig. 2.25 shows the differences in the preamplifier output for an ideal signal and for the GEM one for different input charge. Using a reasonable shape for a GEM signal, it has been found that the distortion on the shaper output occurs at about $40fC$ (against the $18fC$ found for an ideal fast signal). Moreover, a simulation of a GEM signal with large charge released ($1pC$) has shown that the recovery time for the input of the comparator is about $\sim 0.8\mu s$. This number even if large is acceptable for our purposes.

2.4 TOTEM electronics overview

The VFAT2 front end is used by all the TOTEM sub-detectors. A common system architecture is therefore used for readout and control. Fig. 2.26 shows its functional block diagram. Three different region are considered: *On Detector*, *Local Detector* and *Counting Room Region*. The system can be divided in two blocks: the readout (trigger and tracking data) and the control.

The readout path flows from the detector to the counting room region. The trigger path starts with the fast-OR signals of sector⁹ outputs (S-bits) generated by VFAT2s. These signals are logically combined

⁹For the trigger pattern of the *T2* readout plane refers to fig. 2.4

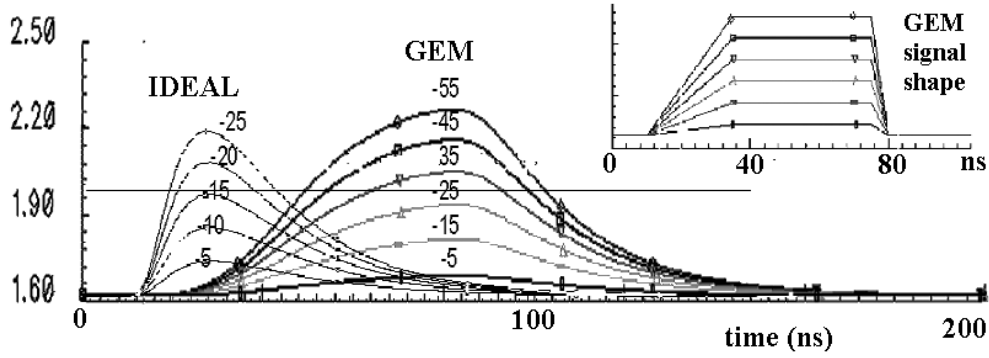


Figure 2.25: Comparison between the simulated [31] outputs of the preamplifier for a fast ideal negative signal (groups of signals from $-5fC$ to $-25fC$ on the left) and for a reasonable GEM signal (groups of signal from $-5fC$ to $-55fC$ on the right). To reach the same preamplifier output as for the ideal case, the GEM signal need more charge (roughly the double) and more time. The shape of the signal used for the GEM is shown in the right corner.

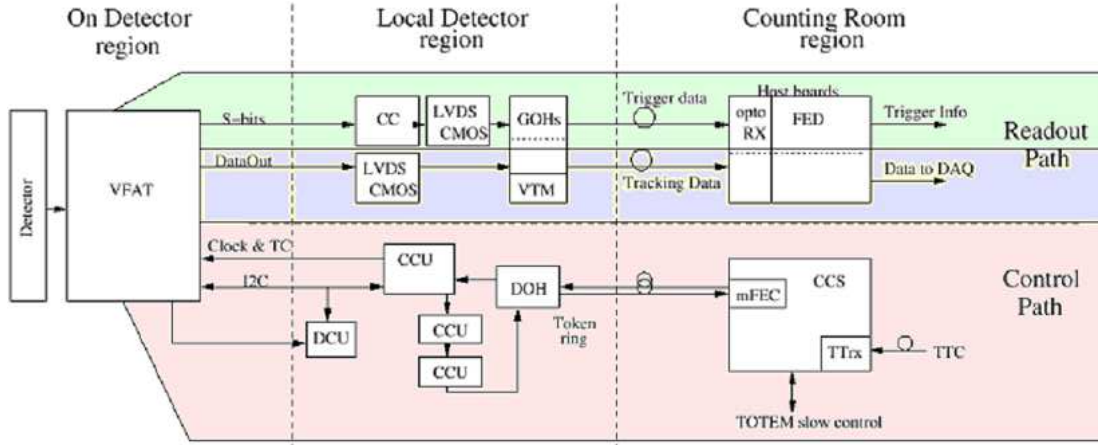


Figure 2.26: TOTEM electronics functional overview.

in the programmable coincidence chip (CC) and the output are transmitted to the the Front-End Drivers (FEDs). In the TOTEM FED Host board coincidence logic functions and more complex algorithms (via FPGA) between all trigger signals coming from all detectors can be performed in order to prepare trigger primitives for the global L1 TOTEM trigger. The global trigger synchronization is based on $BC0$ signal (Beam Crossing Zero) broadcasted by the Timing, Trigger and Control system (TTC). This signal is related to the first bunch of a LHC beam revolution cycle and it is issued every 3564 bunch crossings. On the local region a trigger mezzanine card (TMC) or VFAT2 trigger mezzanine (VTM) transmits the trigger data stream ("a photograph of the outputs of the coincidence chips in the local region"). In order to have the global synchronization with the LHC, the VTM superimpose to this stream the decoded $BC0$ signal.

Data used for tracking (i.e. the status of all the strips and pads) are continuously buffered within the VFAT2. Once a $LV1A$ is received the corresponding Data Packet is transmitted with a continuous write/read operation without dead time. This allow high readout rate capabilities. These data will be sent to the TOTEM FED Host board, where they will be elaborated in an *Event Builder* operation. The transmission of tracking and trigger data to the counting room region is performed with the Gigabit Optical Hybrid (GOH) optical links, that serializes and transmit data at 800 Mb/s via optical links.

The control path is used for two purposes: the transmission of the LHC machine clock and trigger commands (TTC) and of the TOTEM slow control instructions. The Clock and Control System (CCS)

transmits this information onto a token ring by the front end control (FEC) modules FEC. The optical signal at 40 Mb/s from the counting room region is converted in the local region by the Digital Opto Hybrid Module (DOHM). Fast clock/TC and slow control information are then distributed respectively via Low Voltage Differential Signals (LVDS) and I^2C protocol to and from the VFAT2s by the Communication and Control Units (CCU) sit in the token ring. Detector Control Units (DCU) to perform current, voltage and temperature measurement are foreseen in the ring.

2.4.1 T_2 Readout and Control System

The *On Detector* and *Local Region* of the T_2 telescope have some differences with respect to the general structures for some constraints due to the environmental conditions. In particular the *local detector region* has been spatially divided in two zones because the GOH modules for triggering and tracking data transmission doesn't sustain high radiation hard environment. They have been therefore placed far from the detector and a dedicated control token ring has been provided for them. All the other functionality associated in fig 2.26 to the local region are hosted on four identical dedicated cards, directly mounted on each quarter of the T_2 telescope. In particular, this card has been designed to:

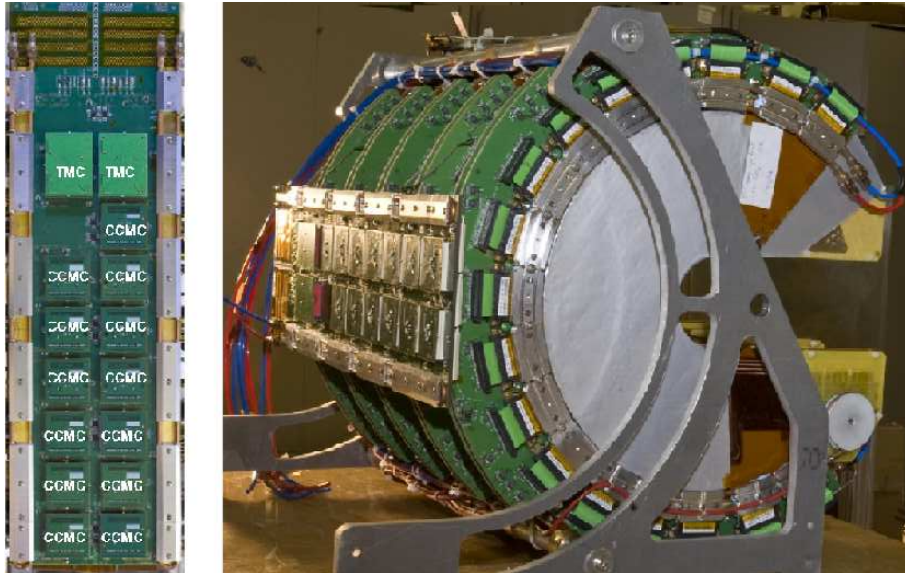


Figure 2.27: Left: 11th card with the 13 CMCCs and the two TMCs plugged. Right: Half telescope with the 11th card connected.

- Send the data lines from all the 170 VFAT2s of one quarter of telescope to the data opto transmitter board (DOB) where the GOHs for tracking data are placed.
- Logically combine the triggers lines coming from the $10 \times 13 \times 8$ *super-pads*¹⁰ with the 13 coincidence chips mezzanine card (CCMCs) and send their outputs to the Trigger opto transmitter board (TOB) where the GOHs for triggering data are hosted.
- Generate the trigger information (i.e. the 104 trigger signals from CCMCs) that will be inserted in the event data packet with the two Trigger Mezzanine Cards (TMC) and send it to the Data Opto-Board.
- Provide the second token ring for the on-detector slow control and for the LHC Clock and Trigger Commands distribution.

¹⁰10 is the number of aligned detectors, 13 is the number of pad sector per chamber and 8 is the number of *super-pad* per sector. Each sector is read by one VFAT2.

- Send the signals from temperature, humidity and radiation monitors mounted on detectors and on the 11th card to the Detector Control System (DCS).

One of this card and its assembly on a quarter of telescope is shown in fig. 2.27. The 11th card, this is its name, has been realized with the multi-wire technology. For the huge number of hosted signal lines indeed the size of this card with standard pcb technology would be too large with respect to the space constraints requested by the insertion of the telescope inside the Hadron Forward detector of CMS.

The data, trigger and control lines from the 17 VFAT2s (4 for Strips and 13 for pads) mounted on each triple GEM are transferred to the 11th card with a multilayered pcb mounted on each detector, the HORSESHOE board. On this board one Detector Control Unit (DCU) and few clock repeaters are present. Fig. 2.28 shows this card with a description of its architecture.

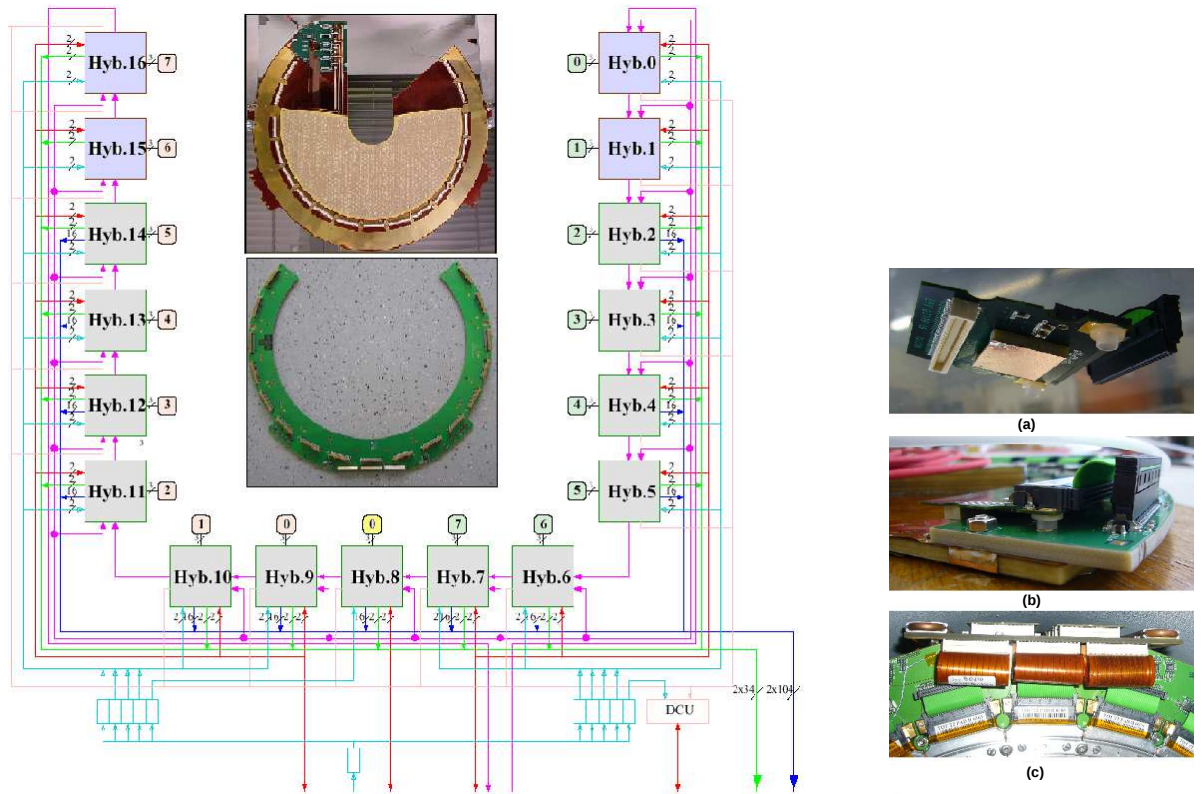


Figure 2.28: Left: The HORSESHOE card with its particular shape to fit with the detector. Borders and lines colors of the block diagram: red for I²C lines, cyan for clock and trigger lines and LVDS buffers, pink for DCU and monitoring lines, blue for trigger hybrid outputs, green for data and data valid hybrid outputs, magenta for scan test lines, dark green for pad VFAT2 hybrids, dark brown for strip VFAT2 hybrids. Right: The VFAT2 hybrids (a) connected on the readout plane and on the HS (b). HS connection with the 11th card realized with three kapton cables(c).

The readout and control system in the Counting Room Region of the T2 telescope has no differences from the others TOTEM detectors. Front End Driver (FED) host boards for tracking and triggering data and Clock and Control System (CCS) with Front End Control (FEC) mezzanines for trigger, time and control are provided. The functionality of this system has been previously described in section 2.4 where a general introduction to the TOTEM electronics has been given. A summary of the T2-specific architecture from the detector up to the counting room for the readout and control electronics of half telescope is shown in fig. 2.29.

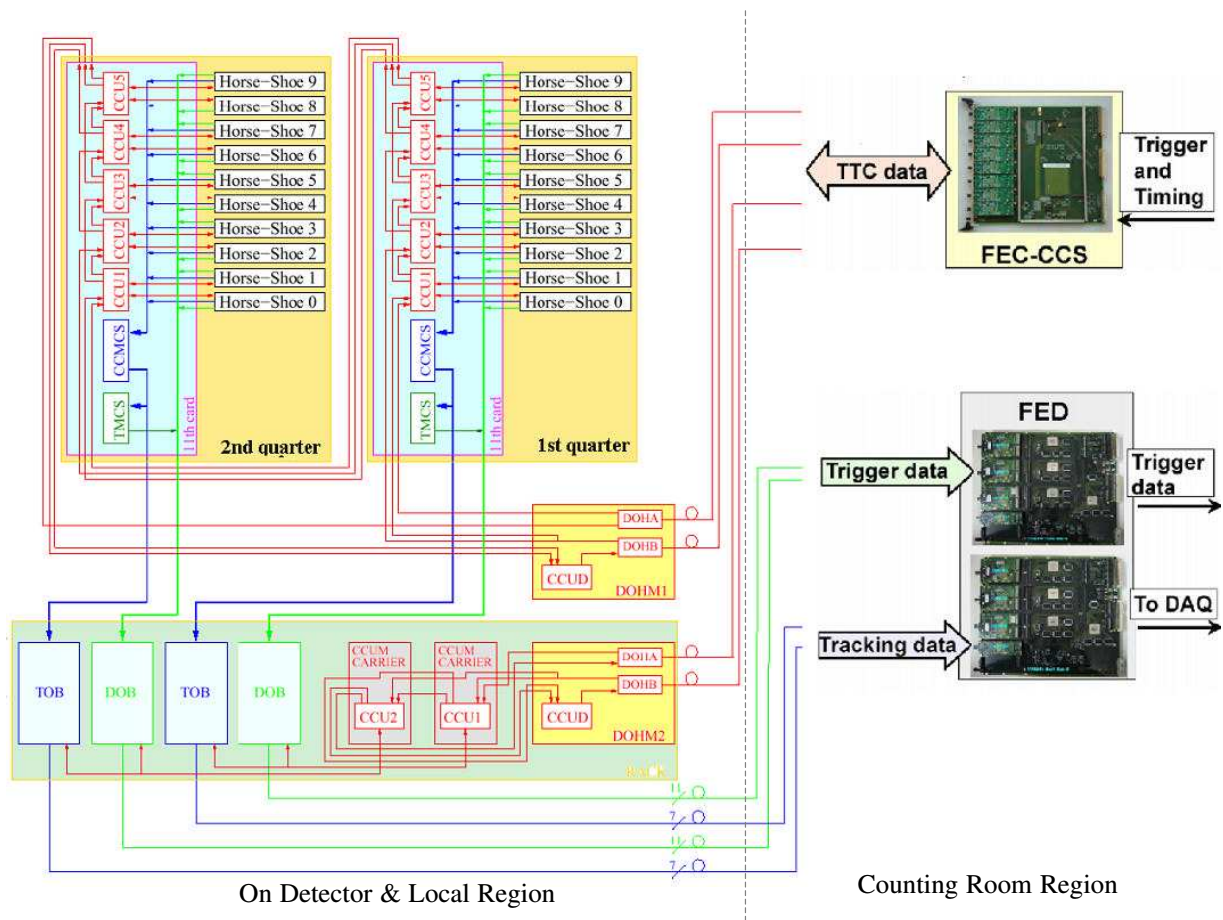


Figure 2.29: Architecture of the readout electronics of half of the detector in the *On Detector*, *Local* and *Counting Room Region*. Borders and lines: black for Horse-Shoe cards, magenta for 11th cards, red for CCUMs, DOHMs and control, clock and trigger lines, blue for trigger hybrid outputs, CCMCs and Trigger Opto-Boards, green for data outputs and Data Opto-Boards, dark green for TMCs and TMC outputs to Data Opto-Boards.

Chapter 3

TOTEM Triple GEM and Simulation

Since the invention of the GEM foil (F. Sauli, CERN, 1997) and the development of detectors based on this technology, a lot of efforts have been spent to develop tools able to simulate their behavior. In this chapter, the simulations done on the TOTEM triple GEM chamber are presented. The aims of this work were the understanding of the detector and the development of the T2 digitization (i.e. the complete T2 telescope simulation, that can reproduce the measured data from simulated pp collision). In sec. 3.1, a description of the simulation and of the used tools is presented together with some general considerations about GEM-based detectors. The algorithm, on which the T2 digitization is based, is however described in sec. 3.2. It has been developed according to detector measurements done without the final electronics and Garfield simulation results. The final tuning and validation will be done during the T2 commissioning in the experimental area at the Interaction Point 5 of the LHC. Preliminary results, obtained from cosmic ray measurements, will be shown in sec. 5.2.

3.1 Detector Simulation

In sec. 3.1.1 the used tools are summarized while brief descriptions or examples of the simulations done and of the detector functioning are given in sec. 3.1.2. The characteristics and principles of operation of a gaseous detectors, exposed (used) in the following sections, are extracted from [?], where a good and complete description is given. As previously noted, the *Garfield* simulation has been used to better understand the way of functioning of a GEM based detector. A more scrupulous analysis is needed if precise and accurate results are aimed. We started following the precious suggestions of D.Pinci¹ and using his codes, that we have modified according to our aims and our detector.

3.1.1 The simulation Tools.

The program used to simulate the TOTEM Triple GEM was *Garfield*², that has been developed by Rob Veenhof³.

It has been used together with finite elements software when accurate descriptions of the electrostatic fields were needed (i.e. inside the GEM holes and in proximity of the readout plane). In particular we used *Maxwell 2D SV*⁴ to create field maps specific to our detectors. This version of *Maxwell* is free and available in the *Ansoft* web site. Other finite elements programs such as *Ansys*, *Maxwell 3D*, *Tosca*, *QuickField* and *FEMLAB* are supported by *Garfield*. The main advantage of these commercial versions is the capability of analyzing the problem in a real three-dimensional way (while symmetry operations

¹D.Pinci, Universita' "La Sapienza" and INFN, Roma, Italy.

²<http://garfield.web.cern.ch/garfield/>

³R.Veenhof, CERN PH/DT and University of Wisconsin.

⁴<http://www.ansoft.com/maxwellsv/>

on bi-dimensional maps are used in the *Maxwell 2D SV* version) giving a more precise description of the system.

Inside *Garfield*, the simulation of the transport and the ionization properties in gas mixtures is done using the interface with *Magboltz*, *Imonte* and *HEED*. *Magboltz* solves the Boltzmann transport equations for electrons under the influence of electric and magnetic fields in various gas mixtures. *Imonte* is used to calculate the Townsend and attachment coefficients for a given gas mixture. *HEED* computes the energy loss of fast charged particles or the absorption of photons in gases.

3.1.2 The TOTEM Triple GEM Simulation

The scheme of a standard Triple GEM detector is shown in fig. 3.1. The use of a triple GEM structure is particularly useful because the discharge rate is reduced (thanks to the sharing of the gain between a higher number of GEM's holes) and the ions feedback effects are suppressed (i.e. fake signals arising from the ions neutralization and consequent release of electrons in proximity of the drift foil). The detector is

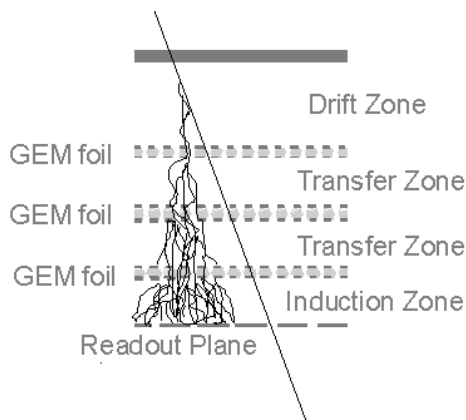


Figure 3.1: TOTEM Triple GEM internal structure. One charged particle (or a photon) that travel trough the detector release energy in the sensitive volume. This energy is converted into electrons that drift inside the detector toward the readout plane. They will be multiplied passing through the three GEM foils in order to have a measurable signal in the front end electronics. The *Garfield* simulation is done fragmenting the problem. Each zone shown in the right picture is treated individually and its output becomes the input of the next one.

characterized by the following parts:

- The *Drift Zone*: the relevant first ionizations happen in this volume. Primaries and secondaries ionizations are simulated and the clusters are diffused and drifted toward the first GEM foil according to the electrostatic field and to the gas mixture. The thickness of this volume affects the detection efficiency and the time length of the signal.
- The *GEM foils*: they perform the multiplication of the incoming electrons. The electrons capture of the top and of the bottom foil is also considered in the simulation. According to the gas used, special effects as the *Penning transfer corrections* should be added to obtain meaningful quantitative results.
- The *Transit Zones*: the diffusion and drift of the electrons is simulated in these volumes.
- The *Induction Zone*: when the electrons enter this volume, they start to induce signal on the readout electrodes according to the *Ramo Theorem*. The thickness of this volume affects the rising time of the signal.

The *Garfield* simulation fragments the full process, solving independently the various parts. The needed steps to obtain the induced signals are shown in fig. 3.2.

The gas mixture

The gas choice follows general criteria. Any gas can be used in a proportional counter because the avalanche multiplication happens in all. The choice is based on the specific requirement. Energy resolution, high rate, radiation hard environment etc. etc.

Nobles gas are normally used as the main component because the multiplication happens at lower field. Polyatomic molecules have indeed many mode of energy dissipation that don't involve the ionization as the rotational and vibrational ones.

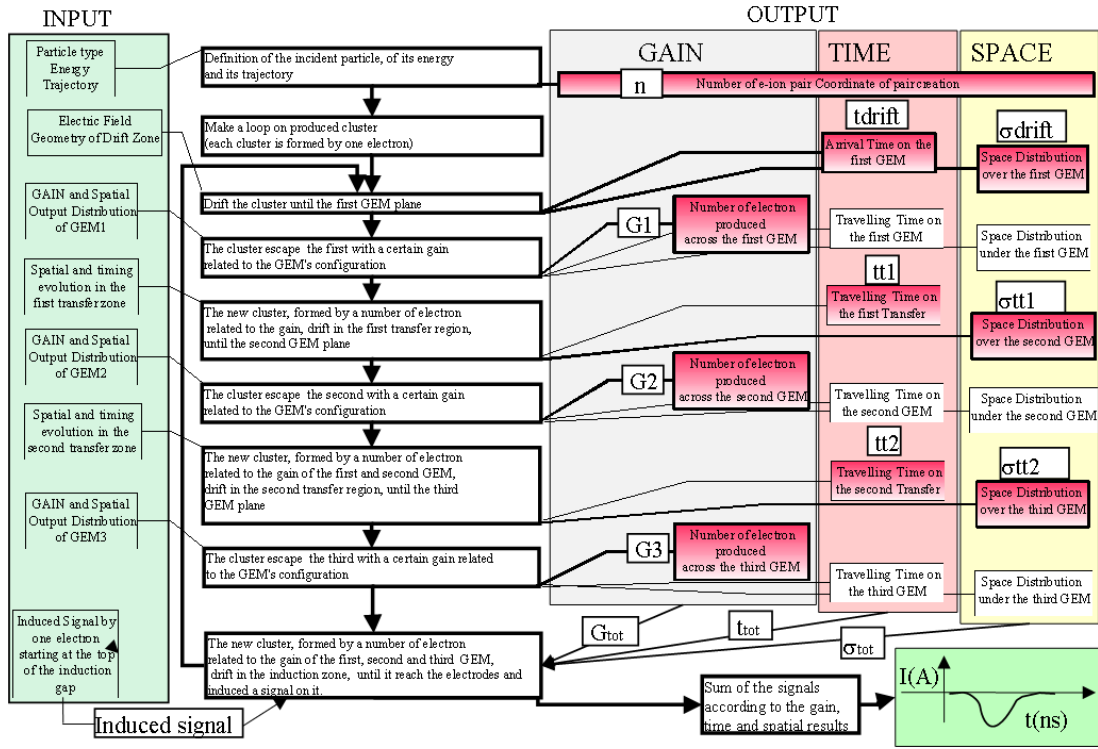


Figure 3.2: Simulation Scheme.

During the avalanche process atoms will be excited and ionized. Both, excited and ionized atoms, can induce emission of secondary electrons. The excited atoms return to the ground state through photon emission. The photons can have enough energy to extract photoelectrons from the metal foils inside the detector. The ionized atoms, instead, can induce secondary emission of electrons during the neutralization on the metal foils. If the electron is emitted in the drift volume, a fake signal will be induced in the readout plane with an intensity comparable with the signals produced by the primary electrons released by passage of the particle.

For this reason polyatomic molecules are normally added to noble gases. The radiationless modes of energy dissipation of these molecules are now useful, because they are able to quench the previous processes.

The gas mixture properties affect directly what happens inside the detector. Ionization, gain, electron's drift and diffusion depend on the used gas. Argon and Carbon Di-Oxide are used in the TOTEM triple GEM, in a 70/30 mixture, whose properties have been computed with *Magboltz* and updated or replaced with more accurate data coming from the literature or computed with other softwares. Ionization properties of different gases for minimum ionizing particles (MIP) are shown in fig. 3.3. They affect the first ionization processes in the drift zone of the detector. Drift velocity and diffusion coefficients are shown in fig. 3.4 for different mixture concentration of *Ar/CO2* as a function of the electric field. These quantities affects the timing properties and the spatial distribution of the charge produced inside the detector. Finally, in fig. 3.5 the Townsend and attachment coefficients are shown. Dependence on field and mixture concentration have been simulated. Those parameters are involved in the gain properties of the GEM foil.

The sensitive volume

The sensitive volume of the detector is the *Drift zone* of fig. 3.1. The electrons released in this region will be multiplied by the three GEM foils. All the other primary electrons released in the other zones will have a smaller gain factor and the produced signal will be therefore smaller. The choice of the thickness of this zone is related to the intrinsic efficiency of the detector. For minimum ionizing particles the primary

Gas	Density [g/cm ³]	Ex [eV]	Ei [eV]	wi [eV]	(dE/dX) _{mip} [keV/cm]	Np (# of primary electrons) [1/cm]	Nt (# of total electrons) [1/cm]	Radiation Length [m]
Ar	1.66E-003	11.6	15.7	26	2.44	23	94	110
CO ₂	1.86E-003	5.2	13.7	33	3.01	35.5	91	183
CF ₄	3.93E-003	12.5	15.9	54	7	51	100	92.4
Ne	8.39E-004	16.67	21.56	36.3	1.56	12	43	345

Figure 3.3: E_x and E_i are the excitation and ionization energies. w_i is the average energy required to produce one electro-ion pair. $(dE/dx)_{mip}$ is the most probable energy loss by a MIP (minimum ionizing particle). n_p is the number per centimeter of primary ionizations and n_T is the total number per centimeter of electrons released.

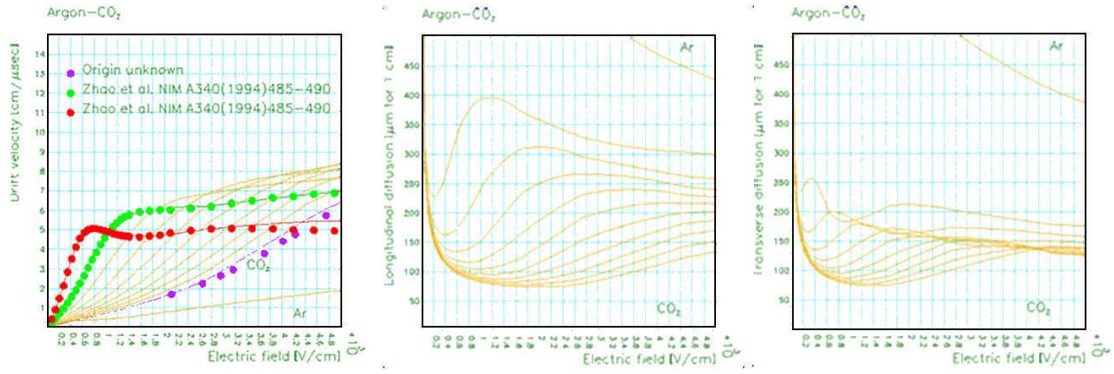


Figure 3.4: The plots show the drift velocity and the diffusion coefficients for Argon-CO₂ mixtures as a function of the electric field. The curves are obtained varying from pure Argon to pure CO₂ in steps of 10%. They have been computed using Magboltz and assuming a temperature of 20 C and a pressure of 760 Torr.

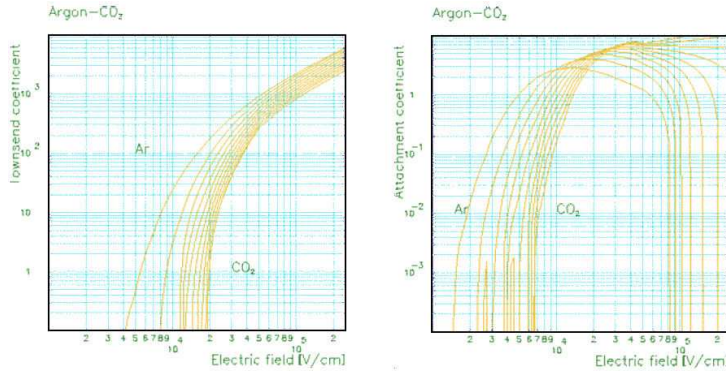


Figure 3.5: The plots show the Townsend and attachment coefficients for Argon-CO₂ mixtures as a function of the electric field. The curves are obtained varying from pure Argon to pure CO₂ in steps of 10%. They have been computed assuming a temperature of 20 C and a pressure of 760 Torr.

ionization follows a Poisson statistics. The probability of k primary ionizations, if the average number of primary interactions is n_p , is:

$$P_k^n = \frac{n_p^k}{k!} e^{-n_p} \quad (3.1)$$

The intrinsic efficiency of the detector is therefore:

$$\epsilon = 1 - e^{-n_p} \quad (3.2)$$

The averaged number per centimeter of primary ionization n_p for Ar and CO₂ (fig. 3.3) is of roughly 30. Our drift volume has a thickness of 3mm and therefore ~ 9 primary ionizations are expected. The intrinsic efficiency is therefore $\epsilon = 1 - e^{-9} = 99.99\%$. If a thickness of 1mm would have been used, the

intrinsic efficiency ϵ would have been of about the 95%.

With the previous statistics it is possible to compute also probability of finding the closest pair at a distance x from the first GEM foil:

$$P_{c.p.}^n(x) = n \cdot e^{-nx} \quad (3.3)$$

where n is the average number of primary ionizations per unit length. If the time performance of the detector is important, this quantity has to be taken into account. The time of detection is indeed the time required from the cluster closest to the GEM foil to reach the induction zone. If the number of primary ionization cluster is small, the distance x has large fluctuation that are reflected on a jitter of the signal between different events or different detectors (aligned on the particle trajectory). This represent an intrinsic limit of the time resolution, that is expressed with the $\sigma(t) = 1/n \cdot v_{drift}$ of the time probability distribution of the first GEM closest cluster. The biggest improvement is done acting on the gas and on the electric field.

In fig. 3.6 it is shown one example of first ionization process for a MIP with the clusters formation along the track of a particle passing through the drift zone. The distribution of total number of electrons (primaries and secondaries) created has a mean value of 27. This is in agreement with the expectation for the configuration that we are using (a drift zone of 3mm filled with Ar/CO_2 70/30 at stp) as can be checked using the data reported in fig. 3.3. The Landau distribution is due to the presence of delta electrons. From the plot of the energy released to the electrons it has been found a peak at 26eV, that corresponds to the average energy required to produce one electron-ion pair in the gas (see table in fig. 3.3). The same simulation can be done with different particles and sources. In ch.4 Cu X-Ray photons are simulated and the results are compared with the measurement done during the TOTEM triple GEM test.

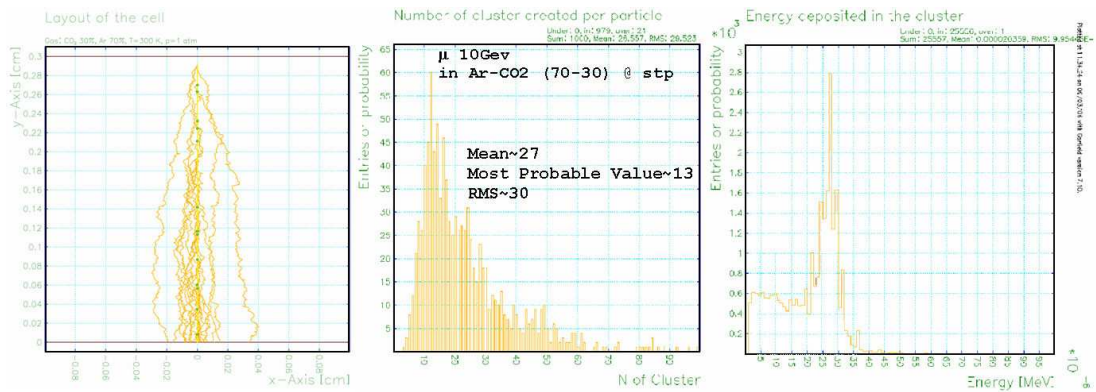


Figure 3.6: Left: MIP (Minimum Ionizing Particle) ionization with the formation of clusters along the track. Middle: Distribution of the number of electrons released in the drift zone. Right: Energy released to the electrons. The simulations have been done for a drift zone of 3mm filled with Ar/CO_2 70/30 at stp.

The GEM Foil

The charge released in the sensitive volume is not enough to be easily detected and it needs to be amplified. The electrons multiplication is done by the GEM foil. With few hundreds of volts across the foil, the electric field inside the holes is sufficiently intense to provide to the incoming electrons the necessary energy to release other electrons through other ionization processes. In terms of total effective gain, not only the internal field has to be considered. The external fields play also their role and have to be properly fixed. The field over the foils has to be not too much high in order to increase the electrons focusing in the hole region and to reduce the number of electrons trapped on the metal plane on the top of the foil. The field under the GEM has to be enough intense to reduce the trapping of electrons on the metal plane on the bottom of the foil.

The simulation of the multiplication process requires high accuracy in the fields definition. External map fields, created with finite element software, have been imported inside Garfield. In fig. 3.7 the

meshing used and the field maps obtained with *Maxwell 2D SV* are shown. The meshing granularity has to be properly done in the more critical region as in the border of the hole.

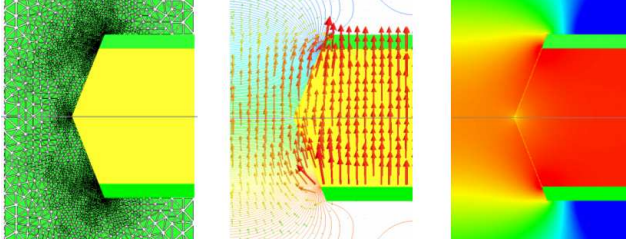


Figure 3.7: *Maxwell 2d SV* elementary cell used to describe the GEM hole. Left: Meshing. Electric field (middle) and its intensity maps (right).

These field maps have been done in 2D with *Maxwell SV*, specifying the type of symmetry that is needed to reproduce the three dimensional object under study. The elementary cell is the one that has been shown in fig. 3.7. In fig. ?? the used axial symmetry is shown with a 3D intensity map of the field obtained for the GEM hole. We were not too much interested on the accuracy of the gain simulation and therefore the solution used was satisfying. A three dimensional elementary cell is more indicated if a more precise simulation is needed. In fig. ?? one example is shown. The map is obtained using translational symmetry. One of the advantages of this cell is that the hole can be simulated taking into account the effects of the nearest ones, while with the 2D elementary cell with the axial symmetry, only one hole can be considered. In fig. 3.9 a simulated multiplication process is shown. Electrons and ions

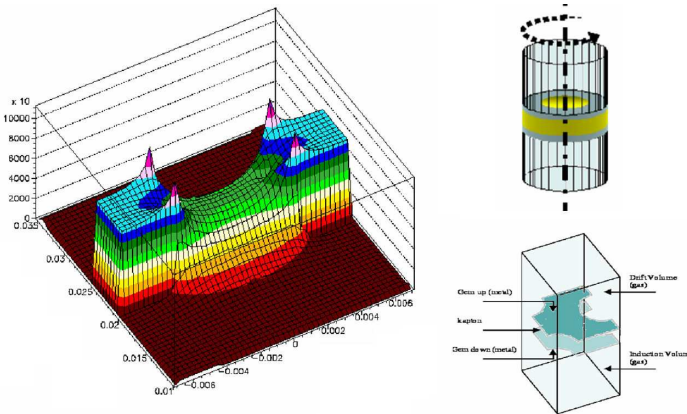


Figure 3.8: Field intensity map (left) in the hole region, obtained with the 2D elementary cell of fig. 3.7 and an axial symmetry operation (top-right). Bottom-right: elementary cell that can be used if a 3D finite element software is used. The symmetry operation will be translational in this case.

produced during the avalanche are drowned. Not all the electrons reach the bottom of the simulated volume, due to a capture in the bottom foil. In the same figure it is shown the distribution of the number of electrons produced per incoming electron. The count at zero corresponds to electrons collected by the top metal foil. The input electrons simulated were one thousand and roughly the 38% with the field configuration used has been trapped before entering the hole (i.e. the transparency of the foil was about the 62%). The effective gain for the tested configuration was about 13. It is important to stress that this result is not very accurate. If an improving of the reliability of these results is required, more detailed and precise field maps (3D finite element sw) and proper corrections on the Townsend and attachment coefficients are needed. Moreover additional effects (as the Penning effect for the Argon) have to be considered.

The electrons transfer

The electrons transfer is related to all the volumes of the detector and it is simulated according to the gas used and the electromagnetic field mapping inside the detector. In fig. 3.10 the Drift velocity and the Diffusion Coefficients are shown as a function of the electric field applied. With these information spatial and temporal distribution of the electron clouds can be simulated. One example is shown in fig. 3.16 where the spatial and arrival time distribution of the electrons drifted for 2mm of $Ar/CO_2(70/30)$ in an electric field of $5kV/cm$ are shown. The time distribution has a rms less than 1ns. The mean is of about

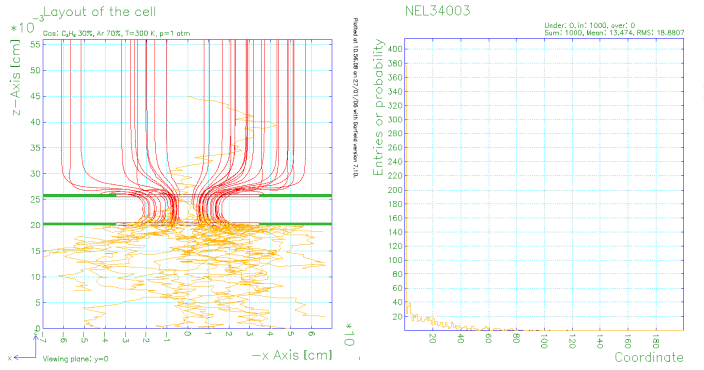


Figure 3.9: Left: Multiplication of an incoming electron in one hole of the GEM foil. Ions are also shown (red lines). Right: Distribution of the number of electrons produced per incoming electron. The field external fields were of $3kV/cm$ on both side. The voltage across the GEM foil was $400V$.

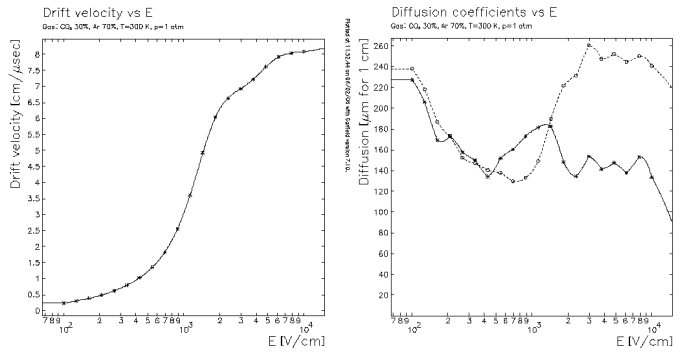


Figure 3.10: Drift velocity and Diffusion Coefficients as a function of the electric field for Ar/CO_2 70/30 at stp.

$26ns$ as expected. with the field used in the simulation (i.e. $5kV/cm$) the drift velocity is nearly $8cm/\mu s$ in $Ar/CO_2(70/30)$ and this means that the expected time to drift for $2mm$ is $25ns$. The drift velocity change between $\sim 6cm/\mu s$ and $\sim 8cm/\mu s$ for fields between $\sim 2kV/cm$ and $\sim 10kV/cm$. The field used inside a detector is normally inside this range and therefore these are typical values for this kind of detector. The spatial distribution can be described with a Gaussian and the rms found ($115\mu m$) can be easily checked using the gas properties. The diffusion coefficients in $Ar/CO_2(70/30)$ with a field of $5kV/cm$ is of about $260\mu m$ for $1cm$. The σ of the spatial distribution is therefore $\sigma = \sqrt{L} \cdot \text{Diffusion Coefficients} = \sqrt{0.2cm} \cdot 260\mu m/cm^{\frac{1}{2}} = 116\mu m$.

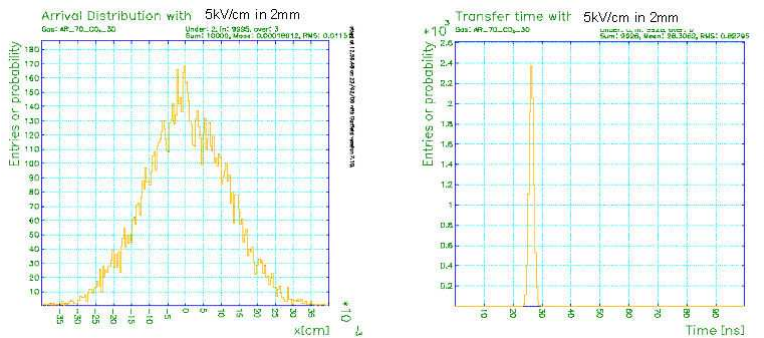


Figure 3.11: Spatial and timing properties of the drifted electrons. Left: Space(left) and time(right) distribution of the electrons transferred in $2mm$ of $Ar/CO_2(70/30)$ with an electric field of $5kV/cm$.

Another example is shown in fig. 3.12 where the spatial and arrival time distribution, in proximity of the first GEM foil, are simulated for electrons released in the drift volume. One thousand MIPs (i.e. roughly $26k$ electrons) have been simulated. The nearly flat structure of the time distribution reflects the fact that the electrons have been created uniformly in the drift volume. The spatial distribution in this case cannot be described with a single Gaussian as before. The found distribution is actually a sum of gaussian with different σ . Each σ corresponding to the distance between the first GEM foil and the position where the electron has been created. The electrons drift simulation can be used moreover to collect approximative information about the hole occupancy. In fig. 3.13 it is shown how it has been done for the first foil. As a first approximation, the fields around the holes have been neglected. The foil

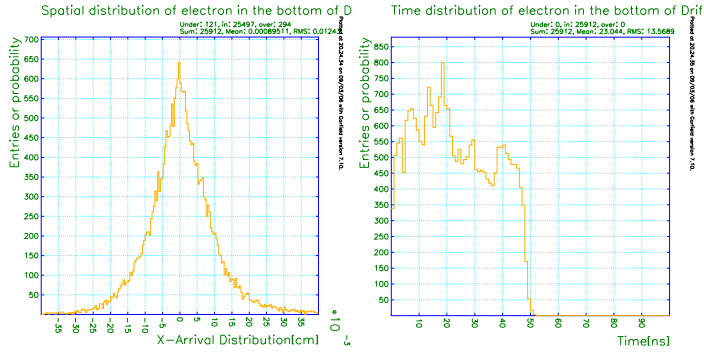


Figure 3.12: [Spatial and timing properties of the drifted first ionization clusters. Left: Space(left) and time(right) distribution of the electrons released in the drift zone when they arrive in proximity of the first GEM foil.

has been divided according to the holes pattern (in our case, with hexagonal shape). The electrons have been therefore drifted and the arrival position has been recorded. In the figure the found occupancy is shown. More interesting is the same analysis on the last GEM foil. In that case indeed, information about

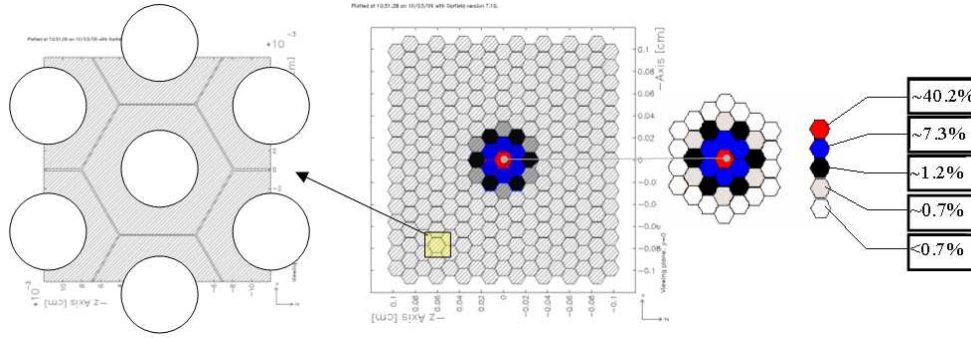


Figure 3.13: First Ionization: First GEM foil hole occupancy for electrons produced in the drift volume by a charged MIP for a gas mixture of Ar/CO_2 70/30 at stp and an electric field of about $2kV/cm$.

the density of charge that could be reached in a single hole can be evaluated. This information is useful for discharging studies. In fig. 3.16 one example is shown. The other foils have been neglected and the electrons have been directly drifted from the drift region until the last GEM foil. The simulation has been done with 1k electrons uniformly created in the $3mm$ drift zone and drifted until the third GEM foil. A total distance of $7mm$ (drift region included) has been considered. The electric field was $3kV/cm$ and the gas mixture $Ar/CO_2(70/30)$.

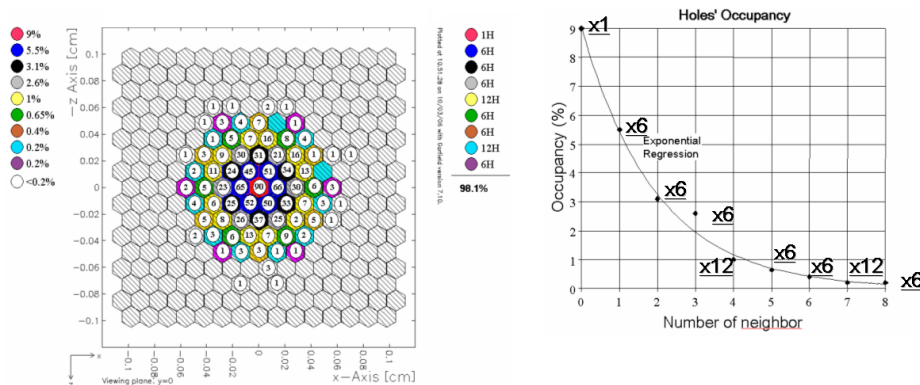


Figure 3.14: Holes occupancy of the last GEM foil. The electric field was $3kV/cm$ and the gas mixture $Ar/CO_2(70/30)$

The signal induction

The current induced by a charge q moving on a trajectory $\vec{x}(t)$ on a generic electrode k of a multi-electrodes system is described by the Ramo's Theorem:

$$i_k(t) = -q \cdot \vec{E}_k^W(\vec{x}) \times \vec{v}(\vec{x}, t) \quad (3.4)$$

where E_k^W is the weighting field in the point $\vec{x}$. Eq. 3.4 is valid when the potential of the other electrodes doesn't change for induction effects. This can be considered true when they are connected to ground with low impedance. The weighting field is defined as the field obtained when the k electrode is at a potential of 1V and all the other electrodes are grounded. Accurate fields maps are needed and they can be obtained, as in the case of the field inside the hole of the GEM foil, with finite element software.

In fig. 3.15 the weighting field computed for a pad is shown. A 2D version of *Maxwell* has been used and the readout plane has been obtained using a translational symmetry. In this way it is not possible to take into account the bi-dimensional structure of pads and the results found have to be considered obviously as an approximation.

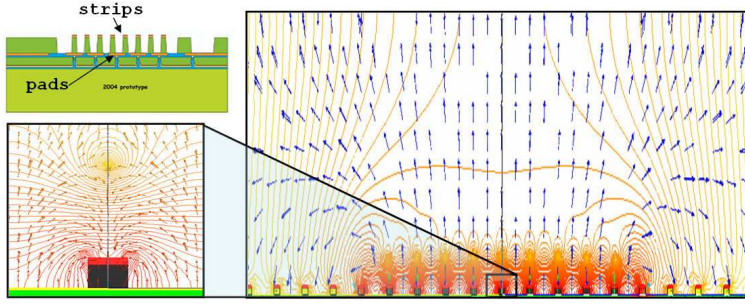


Figure 3.15: TOTEM Triple GEM readout foil structure (top-right) and weighting field map for one pad. The zoom in the right shows the fields lines in proximity of the readout foil. The field has been computed with *Maxwell* 2D SV.

According to the Ramo's theorem, the signal induced by one electron, that is drifting toward the readout plane, start as soon as it exits from the last GEM foil and ends when the electrons is collected. Actually a signal will be induced also in the other electrodes according to eq. 3.4. The integral of the signal induced is the charge q for the electrode that will collect the electron, and zero for all the other. These two type of induced signals are normally called *Direct* when the electron is collected and *Cross* in the other case. Direct signal has the same polarity for all the induction of the signal. Cross signal starts as the direct one and change polarity when the electrons is on the region of inversion of the weighting field (when the product $\vec{E}_k^W(\vec{x}) \times \vec{v}(\vec{x}, t)$ change sign).

Once the weighting field has been computed, the induced signals can be simulated. We have followed a simple approach to obtain these signals. We have uniformly created one thousand of electrons in the drift zone. They have been drifted and diffused down to the last GEM foil and the spatial and timing distribution of the electrons in proximity of the foil have been used as the initial condition of the electron cloud that will induce the signals. This approach cannot describe accurately the signal produced by a MIP because a MIP will never release 1k electrons uniformly distributed in the drift volume. The results are nevertheless a sufficiently good description of the averaged characteristics of real signals. In fig.3.16 one example of signals induced on pad is shown. The Total induced current is the sum of the Direct and the Cross induced currents. In the pad case, the Cross induce component is mainly related to electrons that will be collected by strips. The detector scheme is given as a reference for understanding the signal development. It starts when the electrons originated from the primary ones closest to the first GEM exits from the last foil, i.e. after t_{dt} . The Direct and Cross signals behaves roughly in the same way until the first electrons start to be collected, i.e. after $t_{dt} + t_i$. The signals will be non zero until the electrons originated from the primary ones closest to the drift foil are collected, i.e. after $t_d + t + dt + t_i$. The total length of the signal will be $t_i + t_d$. The previous simulation has been done with the could centered on the pad. In fig. 3.17 the cloud is moved over the readout. Strip and pads are shown and the Direct and Cross Induced signals are summed. This kind of simulation can be used for readout cluster size studies

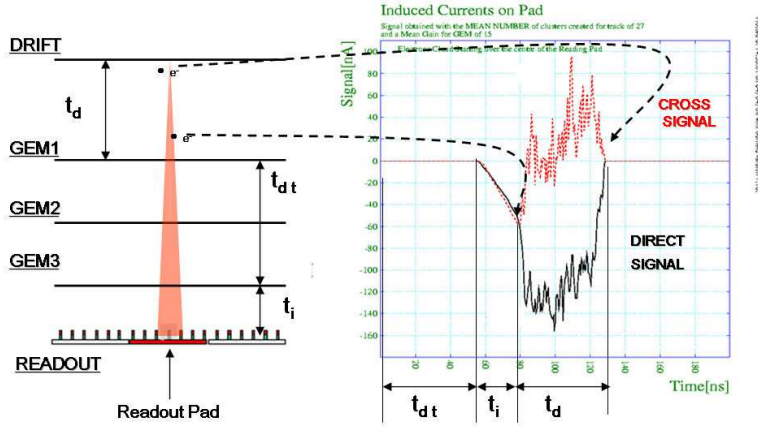


Figure 3.16: Simulated Direct and Cross Induced signals on pad.

or for the signal intensity dependence on the relative position between the first ionization cluster and the readout pattern.

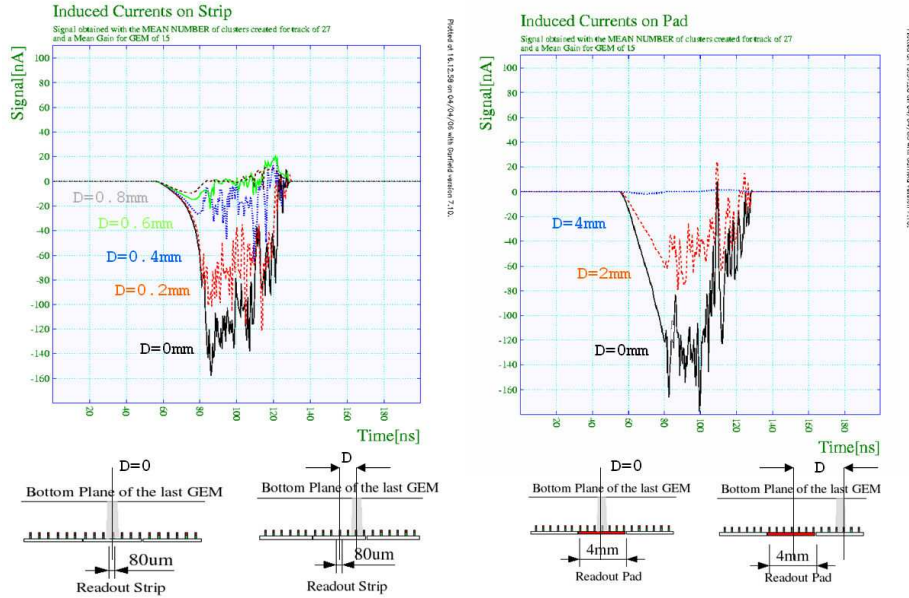


Figure 3.17: Simulated Direct and Cross Induced signals on different pads and strips obtained shifting the initial electron cloud of a quantity D with respect to the center of the readout electrode.

Once the induced signal is simulated, it can be used as input of the transfer function of the readout electronics⁵. A complete simulation is therefore achievable. In the next section the T2 digitization will be presented. The tools described in this section have been used for testing and checking some the assumptions that have been done on developing it.

3.2 Digitization

In this section the algorithms developed to reproduce and predict the response of the TOTEM T2 telescope to an ionizing event will be described.

⁵The readout input impedance has to be properly taken into account. The Ramo's theorems has been used in the approximation of low impedance.

In order to start with a simple approach, the time development of the signal has been neglected and only the amount of charge collected by each readout electrode has been considered.

The output of the digital *V FAT2* front end chip describes the status of its channels. One channel is in the high state when the amplified input signal overcomes a programmable threshold. The digitization will do the same, comparing the simulated amount of charge that the electrode will pick up with the threshold that corresponds to the one programmed on the chip.

Two different algorithms have been developed, based on numerical and analytical approach respectively. The first version of the digitization has been done using the former, that has been coded inside the CMS framework by Erik Brucken⁶. The final release instead is based on the latter.

The basic idea was to find a function that provides the amount of charge collected by one electrode through a *geometrical superimposition* between the readout pattern and the electron cloud that will reach the readout plane. The results obtained with this method will be referred in the next to the *Geometrical Approach*. This function will be described by an *electron collection curve* n_i that represents the percentage of the electrons collected by the electrode i for an ionizing event. It will depend on the shape of the electron cloud and on the place where the ionizations will happen with respect to the electrode i . To formally make explicit these dependencies, a form $n_i = n_i(d, \sigma_{ch})$ will be used, where d is the relative distance in the readout plane between the electrode i and the cloud spatially characterized by σ_{ch} . The total number of electrons N_i collected by the electrode i will be then calculated using:

$$N_i = N_{e-} \cdot G \cdot n_i(d, \sigma_{ch}) \quad (3.5)$$

where N_{e-} is the total number of electrons released by the charged particle in the sensitive volume (the *drift zone* of the TOTEM Triple GEM) and G is the total gain of the detector. Once the charge collected by the electrode is computed, the noise can be added and the found charge can be compared with the threshold level. This will give us the simulated status of strips and pads for an ionizing event, solving the aim of the digitization.

The arguments presented in the next sections will follow the chronology of development of the algorithms used to obtain n_i .

3.2.1 The Strip Case

The size of the electron cloud (few hundreds of micron) with respect to the strips pattern allowed us to face the problem in one dimension, treating the strips as parallel straight lines. The procedure to make explicit the dependencies of the n function on its parameters was:

- i. Creating an electron cloud that describes the electrons that will reach the readout plane.
- ii. Assigning to one electrode the electrons that are spatially over it (i.e. the *Geometrical Approach*).
- iii. Plotting the amount of charge picked up as a function of the parameter under analysis (the relative distance d between the electrode and the projection in the readout plane of the ionization coordinates or the σ of the cloud).
- iv. Reproduce these curves with polynomial fits or analytical functions (that will be the sought *electron collection curve* n).

For simplicity, the initial cloud used in point (i) was produced by one ionizing particle that travels perpendicularly to the detector and that release a quite large amount of primary electrons⁷. The cloud has been realized adding up Gaussian distributions with different sigma, associated to the different positions where the first ionization clusters has been generated in the sensitive volume. The values assumed by σ , neglecting any effect due to the GEM foils, have been determined using:

$$\sigma = \sqrt{L} \times \text{Diffusion Coefficient} \quad (3.6)$$

⁶E. Brucken, Helsinki Institute of Physics HIP and Department of Physical Sciences, University of Helsinki, Helsinki, Finland

⁷The n function, that will be obtained from this particular cloud, cannot be used for non perpendicular trajectory or events with a small amount of electrons released in the *drift zone*. The extension to these generic events will be done after the study of the dependence of the n function on the electron cloud shape.

where L is the distance between the origin of the cluster and the readout plane. The Diffusion Coefficient is obtained with the Magboltz software supplied with the Garfield Package.

As indicated in point (ii), the meaning of the *Geometrical Approach* is that each electron is collected from the underlying electrode. Fig. 3.18 shows two examples of collected charge on strips for two different position of the electron cloud with respect to the strips pattern.

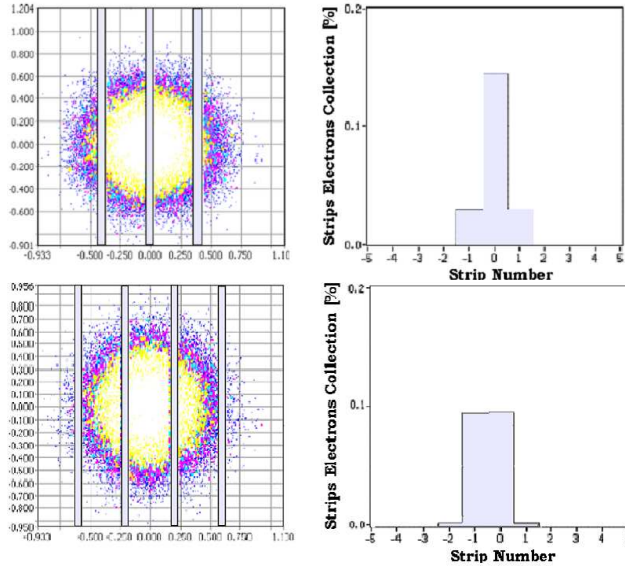


Figure 3.18: Geometrical superimposition between the electron cloud and the strip readout pattern. The electron cloud used in these examples is a sum of $100k$ concentric gaussian distribution of $1k$ electrons each. They describe $100k$ clusters uniformly created in the drift zone. Each σ has been defined by eq.3.6 with L that ranges between $9mm$ and $6mm$. The diffusion coefficient used is $260\mu m$ for $1cm$ that is reasonable for fields between $3kV/cm$ and $5kV/cm$ in $Ar/CO_2/70/30$. Every strips is large $80\mu m$ and the pitch between them is $400\mu m$. In the top plots the cloud is exactly over the middle of one strip that collects nearly the 15% of the total charge. In the bottom plots the center of the cloud is between the two strip that collect less than $\sim 10\%$ each. The abscissa in the histogram plots is the identification number of the strips and the normalization is done with respect to the total charge of the electron cloud.

The Strips Electron Collection Curve: the dependence on d .

The dependence of n_i on d has been made explicit using the curve obtained shifting the electron cloud over the readout plane and calculating, for each position, the collected charge with the *Geometrical Superimposition* (as in fig. 3.18).

The data are plotted in fig.3.19 where also the 9th order polynomial fit that represent the $n_i(d)$ function is shown. In this case the parameter d is the distance between the middle of the strip and the center of the cloud⁸.

The fit is good⁹ and the dependence of the *electron collection curve* n on the relative position between readout and the used cloud can be expressed using:

$$n_i(d) = a_0 + a_1 \cdot d + a_2 \cdot d^2 + a_3 \cdot d^3 + \dots + a_9 \cdot d^9 \quad (3.7)$$

The initial cloud used to obtain the n function of fig. 3.19 is a particular case but it should describe quite well a mean behavior of the detector. The prediction of n has been checked therefore before facing the dependence of n on the electron cloud shape. A preliminary test on this *Electron Collection Curve* $n_x(dx, \sigma_{ch})$ for strips can be done with experimental tests and simulations performed with Garfield. We focused the attention on the sharing of electrons between strips and pads. Even if the algorithm of eq. 3.7 is referred only to strips, it can be used to analyze the charge sharing because the charge on pads can be obtained by difference.

Performing an absolute gain measurement (as it will be described in sec. 4.1.4) we have found that the total signal on the strips is nearly 10% larger than for pads. This is in agreement with the TOTEM

⁸This is true for the used cloud that is produced by one ionizing particle that is travelling perpendicularly to the detector

⁹The polynomial function will show obviously oscillations at large d . This is not a problem because they start to appear far from the end of the electron cloud and can be therefore cut out.

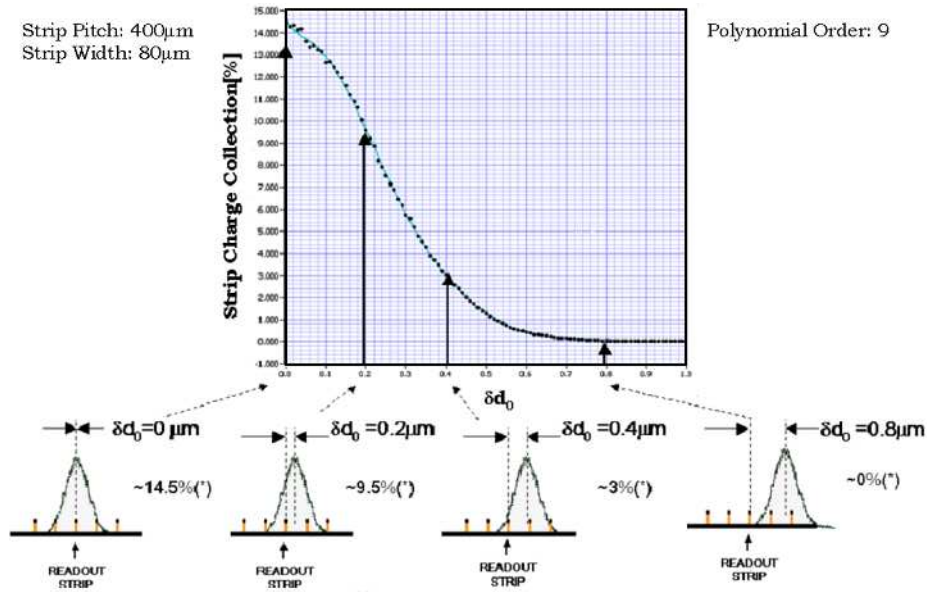
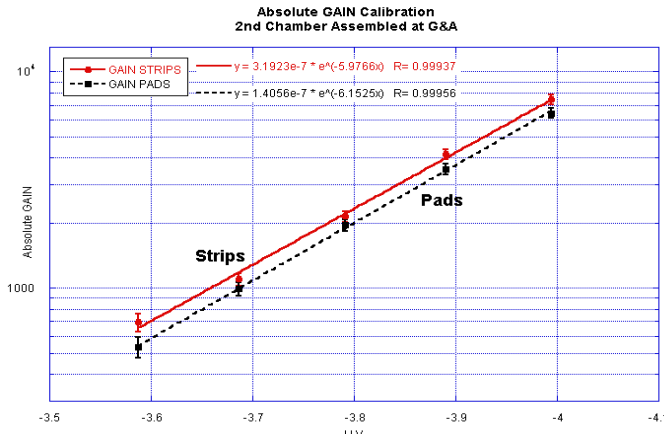


Figure 3.19: Charge collected by one strip in function of the distance from the center of the electron cloud. The dots are obtained with a geometrical superimposition between strip and cloud. The curve is a 9th order polynomial fit.

Triple GEM readout board that has been designed to have the total signal collected by strips per event larger than for pads. In fig.3.20 it is shown the comparison between this measurement and the prediction of the n function (eq. 3.7) that has been obtained with the *Geometrical Approach*. Although the sharing



	Geometrical Approach [%]
All Pads	~80
All Strips	~20
Central Strip	~14
Right neighbor strip	~3
Left neighbor strip	~3

Figure 3.20: Left: Absolute gain of strips and pads in function of the high voltage provided to the triple GEM. The Gain ratio between strips and pads is related to the charge sharing. The measurement has been done exposing one of the chamber assembled by the G&A company to the radiation emitted by a Cu X-Ray tube. Right: One example of the Charge sharing between strips and pads obtained from the n function derived from the data of fig. 3.19. The results are referred to the particular case of electron clouds centered over the middle of the strip.

predicted by n is referred to a particular condition (i.e. the clouds centered on the middle of one strip), the difference is too high. The charge collected on the strips is just $\frac{1}{5}$ of the total. This is actually what we have to find using the *Geometrical Approach* that is based only on the geometric superimposition of the electron cloud and the readout pattern. The width of the strips ($80\mu\text{m}$) is indeed $\frac{1}{5}$ of the pitch ($400\mu\text{m}$). The discrepancy with the measurements has to be linked to the fact that in the *Geometrical Approach* the information on the electrostatic configuration inside the chamber is not included. The simplest way to add it is to use an effective width for the strips that is larger than the physical one. The width became an

additional degree of freedom that can be properly fixed according to the measurements.

The reasons that motivate the use of this effective width have been analyzed with the Garfield software, that takes into account the field maps. Two different methods have been used. The first one gives the collected charge counting the electrons have been picked up by one strip or a pad, after the drifting and diffusion of the electrons in the electrostatic configuration of the chamber. It will be called the *Garfield Test 1* (GT1). This is very close to the *Geometrical Approach*, except for the presence of the electrostatic field configuration inside the detector. The second one gives the charge collected by the time integration of the simulated induced signal on a strip or a pad. It will be called the *Garfield Test 2* (GT2). The induced signal of GT2 involves two components: the *Direct Signal* produced by the electrons that will be picked up by the electrode and the *Cross Signal* due to all the others. The integration of the *Direct Signal* represents the collected charge. The integration of the *Cross signal* should be zero because is generated by electrons that are not picked up by the electrode. The results obtained are summarized in table 3.1.

There are some interesting observation to do:

	Geometrical Approach [%]	GT1: Garfield electrons Collection [%]	GT2: Total Signal Integration	GT2: Direct Signal Integration
All Pads	~80	~57	~47	~59
All Strips	~20	~40	~53	~41
Dielectric between pad and strips		~3		
Central Strip	~14	~28	~38	~30
Right neighbor strip	~3	~7	~7.5	~5.5
Left neighbor strip	~3	~5	~7.5	~5.5

Table 3.1: Charge Sharing between strips and pads obtained with the simple geometrical approach and with the garfield tests GT1 and GT2. In the last two columns the integration of the *Total* and *Direct Signals* are shown.

- i. The GT1 results are closer but not yet compatible with the the measurements (fig. 3.20).
- ii. The integration of the *Total signal* in GT2 is different from the value obtained integrating the *Direct signal*. This should not happen, because it means that the integral of the *Cross signal* is not zero.
- iii. The integration of the *Direct Signal* is in agreement with the results of the GT1 method, which perform a simple count of the electrons collected.
- iv. The values obtained from the integration of the *Total Signal* are very close to measurements where the signal from all the strips is $\sim 10\%$ bigger than the one from pads.

The results obtained with the GT1 test underline that the field lines (see fig. 3.21) in proximity of the strips are enough intense to move the electrons toward them, increasing the number of electrons collected. The use in the geometrical approach of an effective width for the strips, as previously suggested, takes into account of this focusing effects and it could render the *Geometrical Approach* reliable.

The GT1 result is however different from what we have found experimentally and with the GT2 method. Theoretically GT1 and GT2 should give the same result. This doesn't happen for a problem of dynamics and sensitivity of the space discretization in the simulation. An *Induction Zone* with a thickness of $2mm$ is used when the induction of the signal has been simulated in Garfield. The induced signal, according to the Ramo's Theorem¹⁰, depends on the electric field lines. In proximity of the electrodes, where the

¹⁰Ramo's Theorem: $I_i = e\vec{v} \cdot \vec{E}_{wi}(\vec{x})$ where I_i is the instantaneous current on a given electrode i induced by an electron at the position $\vec{x}$, e is the electronic charge, $\vec{v}$ is the instantaneous velocity of the electron, $\vec{E}_{wi}(\vec{x})$ is called the weighting field of that particular electrode. It is defined as the electric field that would prevail at the position $\vec{x}$ with all electrodes grounded except electrode i , which is set to unit potential.

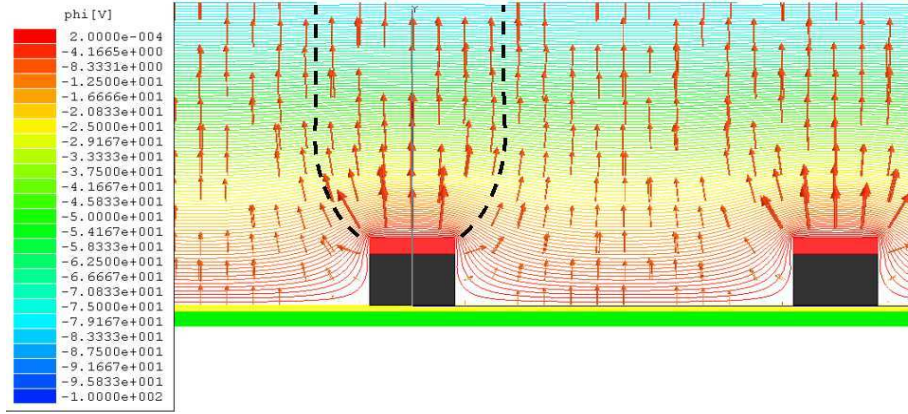


Figure 3.21: Electrical Field Lines and equipotential curves in proximity of a strip.

intensity of the induced signal is higher, the variations of the field lines (see fig. 3.21) follow the typical dimensions of the readout board (tenths of microns). An accuracy on the micron scale cannot be used in the simulation because we have to cover the $2mm$ of the *Induction Zone*. In fig. 3.22 it is shown an example of a *Cross Signal* that has been induced in the middle strips by one electron collected on the underlying pad. The integral of this signal, as it is shown in the right side of the figure, is not zero. The reason is that the last part of the signal is not accurately sampled and part of the signal is lost.

This could explain why the integration of *Direct* and *Total* are different in table 3.1. The fact that the GT1

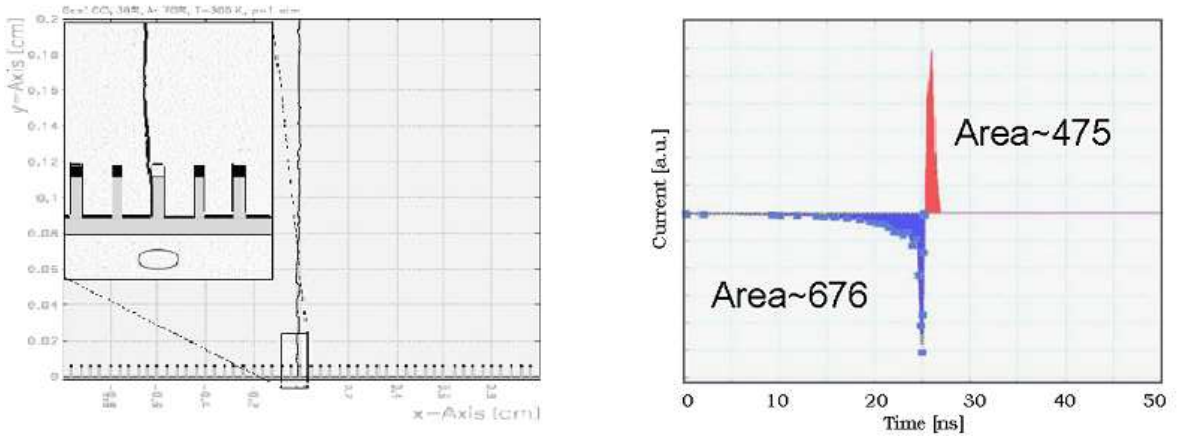


Figure 3.22: Cross Induced Signal Area: The two areas are different and the one corresponding to negative signal is larger. The total integral is therefore not zero.

result is similar to the integration of the *Direct Signal* suggests that this inaccuracy has more effects on the *Cross Signal*. This is reasonable because these signals have more sharp and rapid variation with respect to the *Direct* one.

The charge sharing obtained with the integration of the *Total Signal* is due to the limit in the simulation that is manifested as a sort of filtering effect on the fast parts of the signal. Something similar happens also for real measurements. If the transfer function of the front end electronics used to measure the collected charge is constant in the frequency range of the signal, the time integration of the *Direct Signal* is actually the collected charge while the one of the *Cross Signal* is zero. If this is not the case, the measured integral of the *Direct Induced Signal* is different from the charge collected and the *Cross Induced* one is not zero. Simulation and measurement have the same behavior because in both there is a sort of filtering effect (real for the measurements). It is important to stress however that the found agreement between the integration of the *Total Signal* and the measurements has to be considered only qualitatively. The filtering effects of the simulation is related to its limits and not to the front end transfer function.

Also in this case, the use of an effective width for strip in the geometrical approach gives us the possibility to take into account of these effects. For both cases (the focussing due to the field lines and the non zero area of cross induced signals) the geometrical approach requires a larger width to match the results of the simulations. This request is in agreement with the results obtained with the gain measurements.

In fig. 3.23 the results obtained using an effective width of $210\mu m$ are shown.

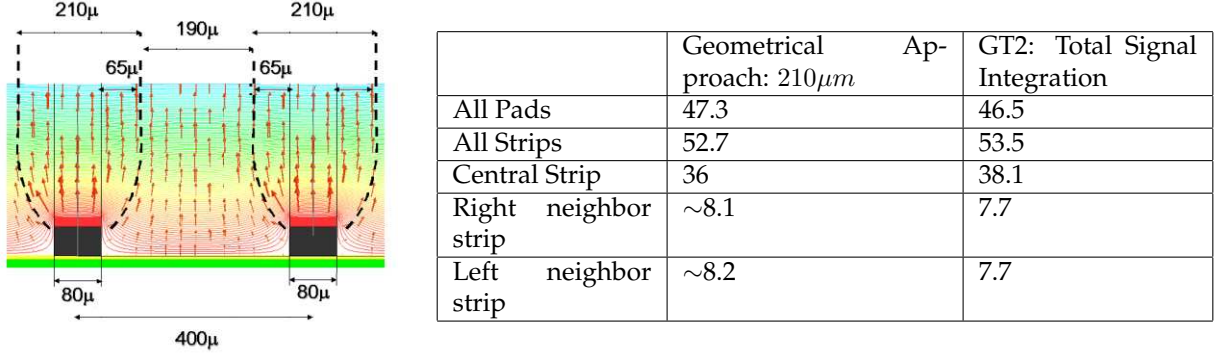


Figure 3.23: Left: $210\mu m$ strips effective width. Right: Comparison between the results obtained with the geometrical approach with an effective strip width of $210\mu m$ and the GT2 Garfield test. The electron cloud was centered over the middle of the central strip.

The proper choice of the effective width has to be fixed according the measurement preformed with the final electronics. The VFAT2 front end used in the TOTEM experiment, with shaper time constant less than $25ns$, is fast with respect to the signal produced by the triple GEM. The instantaneous output of the shaper for a *Cross Signal* could overcome the threshold and this will be reflected into an increased number of strips and pads in the ON state per event.

The Strips electron collection curve: the dependence on σ

The use of a polynomial function n with the additional degree of freedom of the effective width seems to be the simplest solution to predict a value compatible with the expected one. We have started therefore the analysis of the dependencies of n on the electron cloud shape. This will allow the extension to generic processes (i.e. ingoing particles not perpendicular to the triple GEM and with any number of first ionization clusters).

The n function will be associated from now to each first ionization cluster and not to the total number of electrons that will reach the readout plane. The prediction for one event will be provided by the normalized sum of all the n functions of each cluster. The initial electron cloud used in the *Geometrical Approach* to determine n will be now a simple Gaussian cloud (and not a sum of Gaussian as before). The geometrical approach with an effective strip width of $210\mu m$ has been used to compute the $n(d, \sigma_{ch})$ for Gaussian clouds with a σ_{ch} in proximity of the readout plane that varies between 0.14 and $0.28mm$ (fig. 3.24). The effective strip width of $210\mu m$ has been chosen looking at the compatibility between the measured and the simulated charge sharing. These n functions can be described by the general equation:

$$n(d, \sigma_{ch}) = a_0(\sigma_{ch}) + a_1(\sigma_{ch}) \cdot d + a_2(\sigma_{ch}) \cdot d^2 + a_3(\sigma_{ch}) \cdot d^3 + \dots \quad (3.8)$$

A generic n function is therefore defined fixing the a_i coefficients that depend on the particular σ_{ch} of the cloud. This can be done using tables that link each a_i to the σ_{ch} or using analytic functions that link them. we decided to avoid these solutions and we have looked instead for operations as scaling and stretching that move the various $n(d, \sigma_{ch})$ of fig. 3.24 into one $n_0(d, \sigma_0)$, computed for $\sigma_{ch} = \sigma_0$ and used as a reference. The target is to express $n(d, \sigma_{ch})$ using the following form:

$$n(d, \sigma_{ch}) = SC(\sigma_{ch}) \cdot \{a_0(\sigma_0) + a_1(\sigma_0) \cdot [ST(\sigma_{ch}) \cdot d] + a_2(\sigma_0) \cdot [ST(\sigma_{ch}) \cdot d]^2 + a_3(\sigma_0) \cdot [ST(\sigma_{ch}) \cdot d]^3 + \dots\} \quad (3.9)$$

where the a_i are the polynomial coefficients of n_0 . SC and ST are the scaling and stretching factors needed to transform $n_0(d, \sigma_0)$ in $n(d, \sigma_{ch})$.

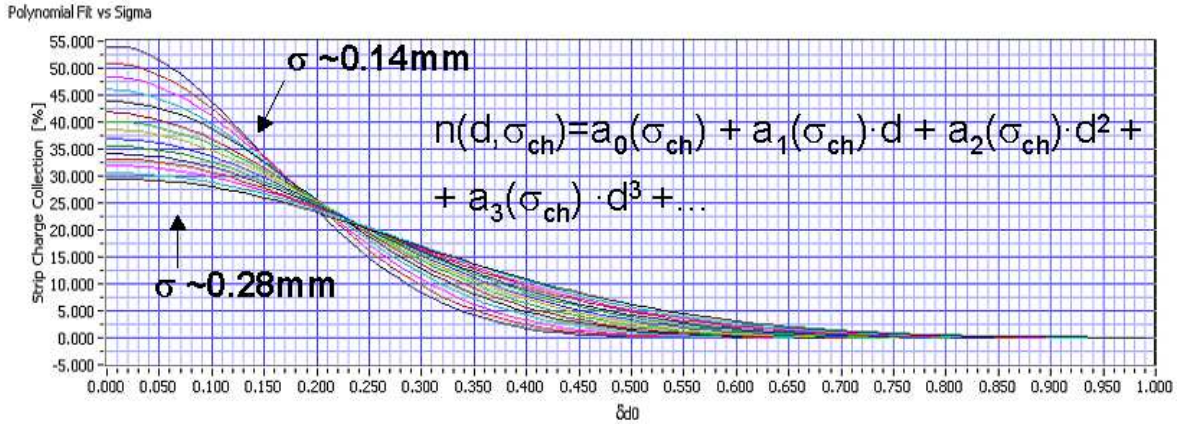


Figure 3.24: Polynomial 9th order n functions obtained for various σ_{ch} of the Gaussian electron cloud obtained for an effective strip width of $210\mu m$.

Using $\sigma_0 = 0.275mm$ the reference curve n_0 is described by:

$$n_0(d, \sigma_0) = a_0(\sigma_0) + a_1(\sigma_0) \cdot d + a_2(\sigma_0) \cdot d^2 + a_3(\sigma_0) \cdot d^3 + \dots \quad (3.10)$$

with:

$$\begin{aligned} a_0 &= 29.574 & a_5 &= -12586.673 \\ a_1 &= -3.021 & a_6 &= 15217.238 \\ a_2 &= -51.516 & a_7 &= -10984.9 \\ a_3 &= -1261.465 & a_8 &= 4408.285 \\ a_4 &= 5989.083 & a_9 &= -756.550 \end{aligned} \quad (3.11)$$

The stretching and scaling operations, that satisfy the request of moving any $n(d, \sigma_{ch})$ to the reference $n_0(d, \sigma_0)$ or viceversa, exist. In fig. 3.25 there are the $SC(\sigma_{ch})/a_0$ and $ST(\sigma_{ch})$ factors needed to obtain all the curves of fig. 3.24 from the reference n_0 according to eq. 3.9. The dependence of the two factors on

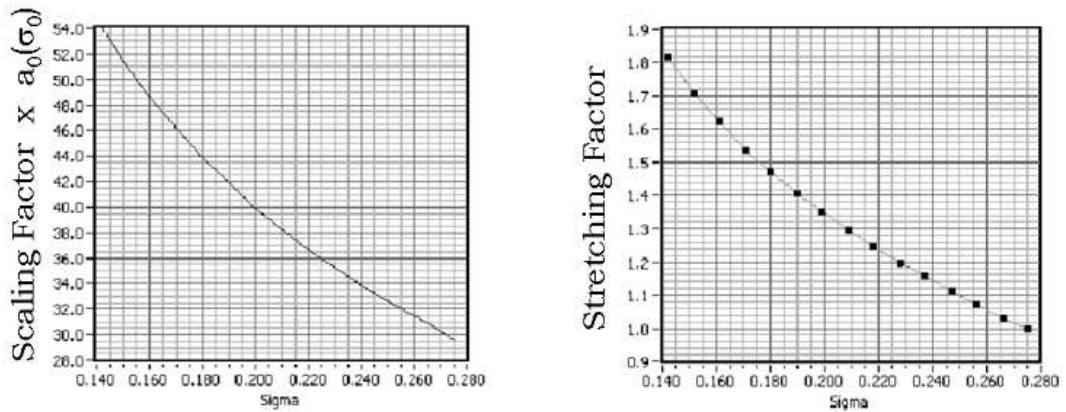


Figure 3.25: $SC(\sigma_{ch})/a_0$ and $ST(\sigma_{ch})$ needed to transform $n_0(d, \sigma_0)$ in $n(d, \sigma_{ch})$ with eq. 3.9 for an effective strip width of $210\mu m$.

σ_{ch} can be expressed with two 4th order polynomial functions that fits the data and that are:

$$SC(\sigma_{ch}) = SC_0 + SC_1 \cdot \sigma_{ch} + SC_2 \cdot \sigma_{ch}^2 + SC_3 \cdot \sigma_{ch}^3 + SC_4 \cdot \sigma_{ch}^4 \quad (3.12)$$

$$ST(\sigma_{ch}) = ST_0 + ST_1 \cdot \sigma_{ch} + ST_2 \cdot \sigma_{ch}^2 + ST_3 \cdot \sigma_{ch}^3 + ST_4 \cdot \sigma_{ch}^4 \quad (3.13)$$

with:

$$\begin{aligned} SC_0 &= 4.3955 \\ SC_1 &= -25.5952 \\ SC_2 &= 47.5814 \\ SC_3 &= 72.4109 \\ SC_4 &= -255.3941 \end{aligned} \quad (3.14)$$

$$\begin{aligned} ST_0 &= 8.150 \\ ST_1 &= -100.455 \\ ST_2 &= 594.662 \\ ST_3 &= -1671.203 \\ ST_4 &= 1793.838 \end{aligned} \quad (3.15)$$

To check the reliability of what it has been found, the $n(d, \sigma_{ch})$ predicted with eq. 3.9 has been compared with the same curve obtained directly from the geometrical superimposition of the cloud and the readout pattern. The errors are below $5^\circ/_{oo}$ and are shown in fig. 3.26.

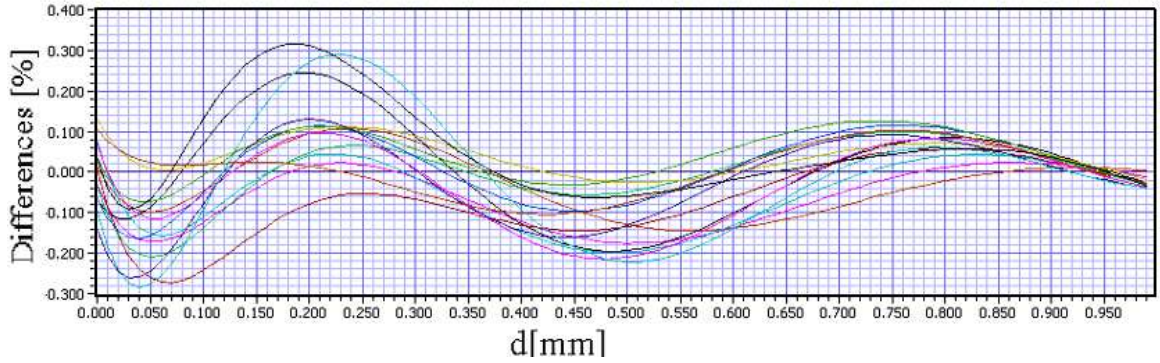


Figure 3.26: Differences between the *electron collection curves* calculated collecting electrons from initial electron clouds with different σ and the collection curves extrapolated with the sigma parameterization (i.e. using an *electron collection curve* relative to a fixed σ , properly scaled and stretched to describe the other cases. The error is less than 1% for all the analyzed cases.

The Strips electron collection curve: a generic ionization process

The parameterization in σ allows to describe a generic event summing over the j first ionization clusters created by an ionizing particle. The *electron collection curve* n can be defined as:

$$n = \sum_j \{ SC(\sigma_j) \cdot \{ a_0(\sigma_0) + a_1(\sigma_0) \cdot [ST(\sigma_j) \cdot d_j] + a_2(\sigma_0) \cdot [ST(\sigma_j) \cdot d_j]^2 + a_3(\sigma_0) \cdot [ST(\sigma_j) \cdot d_j]^3 + \dots \} \} \quad (3.16)$$

where:

- n is the percentage of charge collected by one strip that is placed at distance d_j from the cluster j .
- The σ_j are defined by eq. 3.6¹¹ for each first ionization cluster j .
- The a_i are the polynomial coefficients of the reference $n_0(d, \sigma_0 = 0.275mm)$ of eq. 3.10.
- The SC and ST are the scaling and stretching factors defined in eq. 3.12 and eq. 3.13.

The j clusters can be a result of a simulation or can be defined arbitrarily. A uniform distribution in the drift zone using the entry and exit point of the ingoing particle is shown for example in fig. 3.27.

3.2.2 The Pad Case

The electrons collected by strips are defined by eq. 3.16. Now the *electron collection curve* n for pads has to be found. The main differences are the bi-dimensionality of the problem and the fact that the pads are partially covered by the strips. The latter problem has been solved calculating the electrons collected by the pads without considering the strips and viceversa. The charge of strips has been then subtracted from the charge picked up by the underlying pad. The bi-dimensionality has been faced checking if $n(d, \sigma_{ch})$

¹¹ $\sigma_{ch} = \sqrt{L} \times \text{Diffusion Coefficient}$ where L is the distance between the cluster origin and the readout plane.

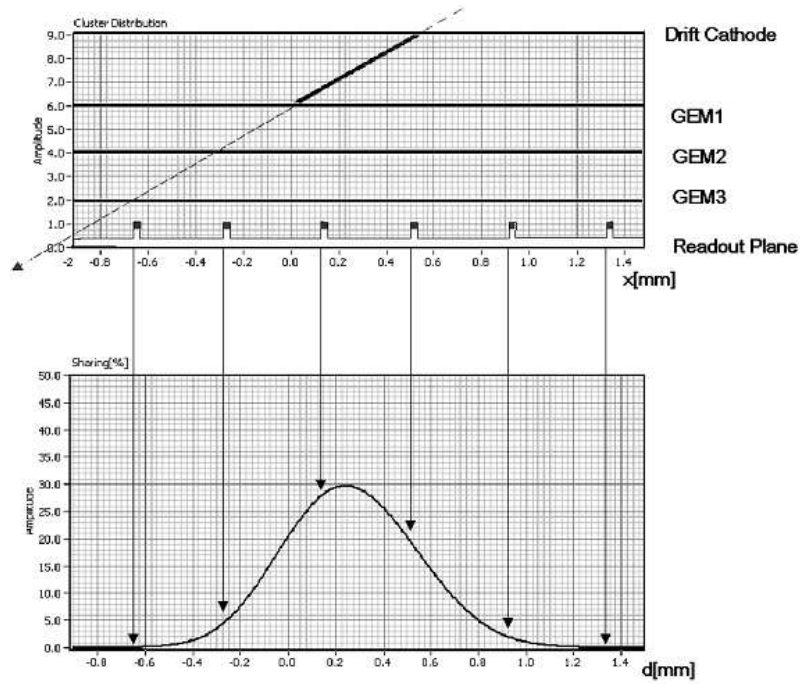


Figure 3.27: Track of an ionizing particle and relative charge collection curve. The abscissa in the plot of the *electron collection curve* n (bottom) has been properly transformed to have a direct correspondence with the spatial coordinate of the plot on the top. The percentage collected by one strip can be graphically extracted projecting the strip position from the top to the bottom plot.

could be factorized as $n_x(dx, \sigma_{ch}) \cdot n_y(dy, \sigma_{ch})$. A square shape for the pads will be used in the following instead of the real trapezoidal one. With this sufficiently good approximation, the n_x and n_y should be an unique $n = n_x = n_y$.

The Pads electron collection curve

The followed approach is the same as the one used for strips, with the difference that we started directly with simple gaussian clouds as in sec.3.2.1. Using the linearity of the problem, the result for a complete ionization process will be obtained adding the various *electrons collection curves* obtained from single gaussian cloud with different σ_{ch} (as expressed in eq. 3.16).

An electron cloud characterized by a $\sigma_{ch} = 0.275$ has been generated and divided, according to the pads geometry. Fig.3.28 shows one pad in different position with respect to the clouds together with the extraction of the electrons that are geometrically over the pad. The amount of charge collected by the pad, as a function of the two spatial shift dx and dy , has been obtained changing the relative position between the cloud and the readout plane.¹² Fig.3.29 shows the percentage of charge collected by one $2.2\text{ mm} \times 2.2\text{ mm}$ pad as a function of the relative position between the cloud and the pad itself. The $n_x(dx, \sigma_{ch})$ and $n_y(dy, \sigma_{ch})$ has been obtained fitting polynomially the charge computed for displacements along x and y respectively. From these data (fig. 3.29), obtained by the geometric superimposition between cloud and readout pattern, the polynomial coefficients of the n_x and n_y functions have been extracted. The prediction of these polynomials have been superimposed to the data in the same figure. Tests have been performed also in different conditions (different σ_{ch} and pad area). They have confirmed that the polynomial fits are reliable to predict the amount of charge collected and that $n(d, \sigma_{ch}) = n_x(dx, \sigma_{ch}) \cdot n_y(dy, \sigma_{ch})$ can be actually expressed by the product of two identical mono-dimensional polynomial function. Until now, the distance between the electron cloud and the center of the pads has been used to compute the collected charge. A cloud with a $\sigma = 275\mu\text{m}$ is practically totally included in the $2.2\text{ mm} \times 2.2\text{ mm}$ pad (i.e. the one with the smaller size in the TOTEM triple GEM). This relationship between electron cloud and pad size suggests the use of the distance from the border of the pad instead of the

¹²In this context dx and dy represent the distance between center of the cloud from the center of the pad

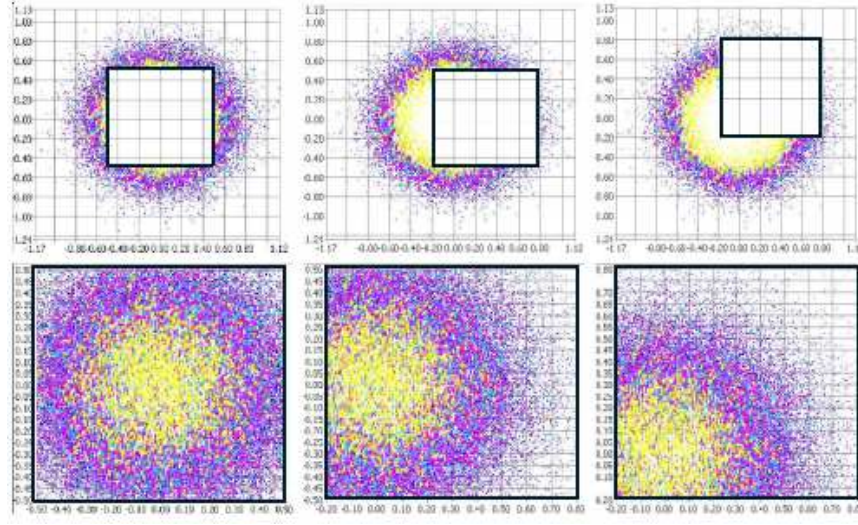


Figure 3.28: Geometrical computation of the electrons collected by one pad in different position with respect to the clouds. The pad area is $1 \times 1 \text{ mm}^2$ just to better visualize the various position. The smaller pads of the triple GEM have an area of $2 \times 2 \text{ mm}^2$ that include completely the cloud.

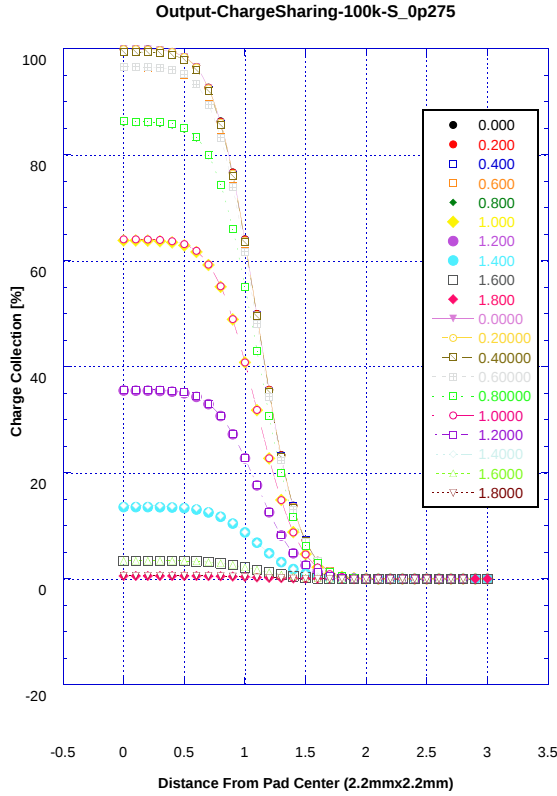


Figure 3.29: Collected charge for different relative position between the gaussian cloud and a pad of $2.2 \text{ mm} \times 2.2 \text{ mm}$. The σ used was 0.275 cm . Relative shifts in the x and y direction are shown. Each point style corresponds to a fixed value for one coordinate and the corresponding curve represents the scan that has been done in the other. There is an obvious superimposition between curves obtained moving along x or y due to the symmetry of the problem. This justify the assumption that two *electron collection curves* n_x and n_y will be actually the same function.

center. In this way pads with different area are treated exactly in the same way and no other geometrical information, except the distance from the border, have to be supplied. The dependencies on the σ of the clouds has been analyzed and solved as for strips (i.e. finding scaling and stretching factors and using a solution of the type expressed in eq. 3.9).

The plots in fig. 3.29 suggested however another possibility to face the problem of predicting the charge collected. The n function could be indeed described by the *erf* function instead of the 9th order polynomial. The computation of the charge collected in the *Geometrical Approach* is actually the integra-

tion of a gaussian cloud, that is exactly what is done by the erf function. The next section will describe this approach.

The Erf function

The Erf Function, used to fit the data in fig. 3.29, is:

$$n(x) = m1(1 + m2 \cdot \text{erf}(m3 \cdot (x - m4))) \quad (3.17)$$

where x represent the distance between the center of the gaussian cloud and the border of the pad. The *electron collection curve* n computed with eq. 3.17 and compared with the one geometrically calculated is shown in fig.3.30 with the residuals. The result obtained is good and therefore it is possible to study the

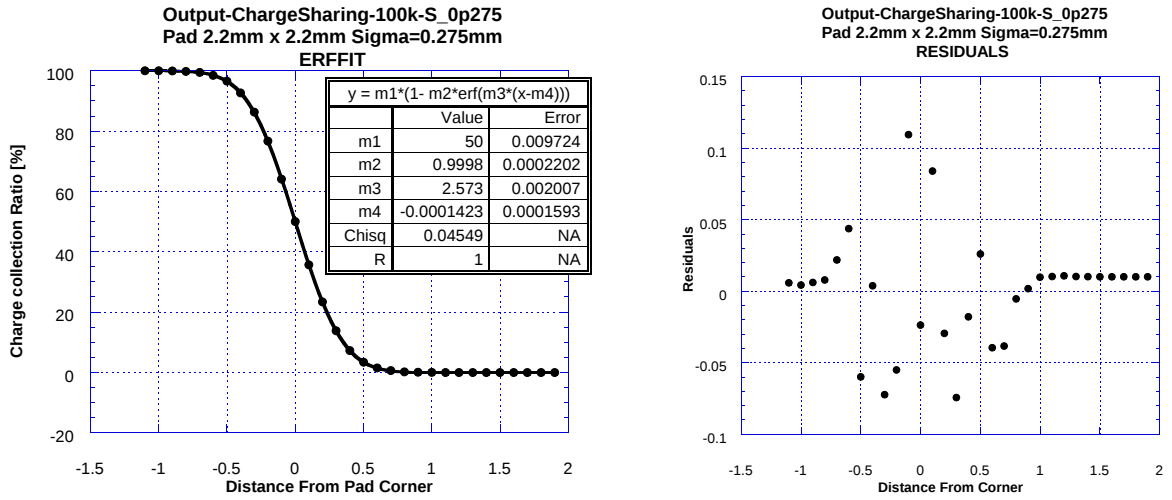


Figure 3.30: Left: Fit of the charge collection curve with the eq. 3.17. Right: Residuals

parameterization in term of σ_{ch} . According to the definition of the *erf* function, it has to be expected that it is related to the parameter $m3$. Many initial electron clouds have been created with different σ and the number of electrons, that have fallen inside the pad, have been computed changing the relative position between pad and cloud. Each curve has been therefore fitted with the eq. 3.17 and the variation of each parameter of the fit have been extracted. They are shown in fig. 3.31. The initial hypotheses that the most affected parameter is $m3$ was obviously right. In fig.?? we have fitted the behavior of $m3$ versus sigma with a power fit (that actually should be expressed by $m3 = (\sqrt{2}\sigma_{ch})^{-1}$).

Being the σ_{ch} dependence expressed by the parameter $m3$, we have decided to fix the others and to use the following equation instead of eq. 3.17:

$$n(x) = 50(1 - \text{erf}(m3 \cdot x)) \quad (3.18)$$

In fig.3.32 the behavior of $m3$ versus σ_{ch} , with $m1, m2, m4$ fixed, has been fitted with the following power law:

$$m3(\sigma_{ch}) = 0.7043 \cdot \sigma_{ch}^{-1.004} \quad (3.19)$$

This confirm the expected relation $m3 = (\sqrt{2}\sigma_{ch})^{-1}$, that will be used in the next to test the algorithm.

Pad electron collection curve Test

In this section the previous results will be tested. The mono-dimensional equation that describe the charge collection is:

$$n(d) = 50(1 - \text{erf}(m3 \cdot (d))) \quad (3.20)$$

where d , expressed in *mm*, is the smallest distance (projected on x or y) between the border of the pad and the center of the initial electrons gaussian cloud. The parameter $m3$ will be obtained for each σ_{ch} , expressed in *mm*, following this equation:

$$m3(\sigma_{ch}) = (\sqrt{2} \cdot \sigma_{ch})^{-1} \quad (3.21)$$

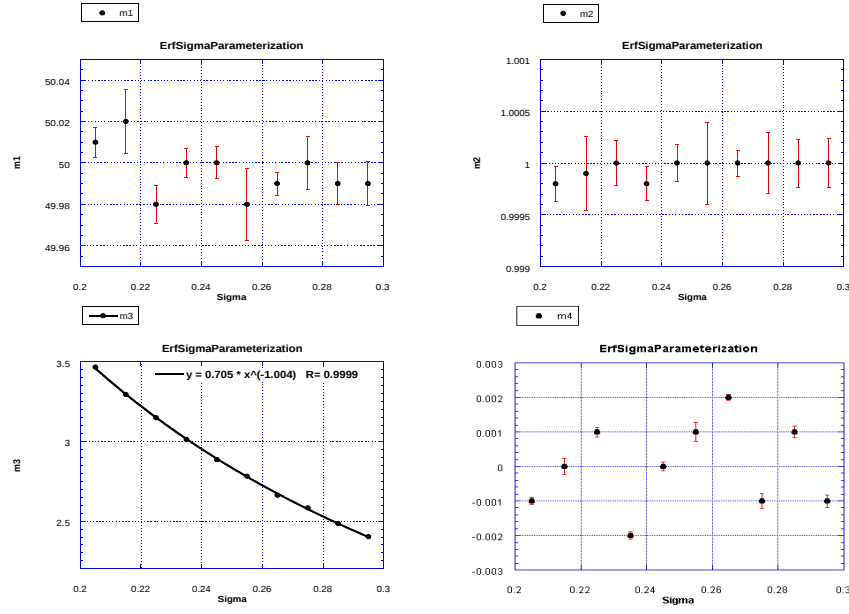


Figure 3.31: Behavior of the four parameter $m1$, $m2$, $m3$, $m4$ of eq.3.17 obtained fitting various electron cloud with different sigma.

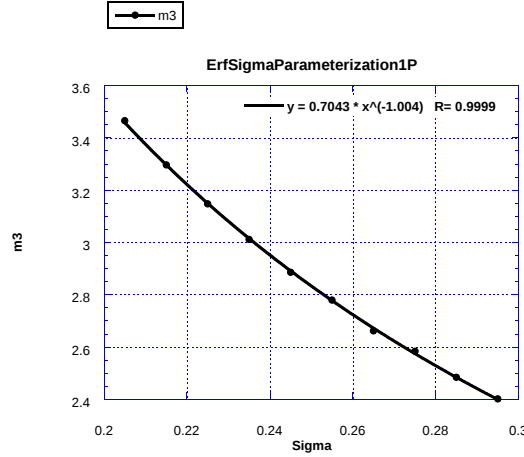


Figure 3.32: Power fit of the parameter $m3$ in function of σ_{ch} . These values have been obtained fitting with eq. 3.18 the geometrically computed *electron collection curve*.

The Charge Collected will be obtained using the following bi-dimensional function:

$$n(d_x, d_y, \sigma_{ch}) = \left(\frac{1}{100}\right) \cdot n(d_x, \sigma_{ch}) \cdot n(d_y, \sigma_{ch}) \quad (3.22)$$

One thousand of events have been generated for the test, with randomly distributed d_x , d_y and σ as shown in the plots of fig.3.33. The difference [%] between the "geometrically counted" and the analytically calculated (with eq.3.22) Charge Collection is shown in fig.3.34. All the points are under the 0.6%. Just to have an idea of what it means: if we have a total gain (strips+pads) of 16000 and if we have 30 primary electrons (mean value of the Landau for the energy loss in the sensitive volume), the 0.6% of this total charge is nearly 3000 electrons, that is something like 5-6 DAQ channels for the VFAT2. If we consider instead the sigma of the residuals distribution, its value is $\sim 0.09\%$ and this means <1 bin. It's important to observe that this error is a relative one and that we are doing a "digital" readout. For these reasons we are sure that this error will not influence the information if there is an hit or not. Eventually it

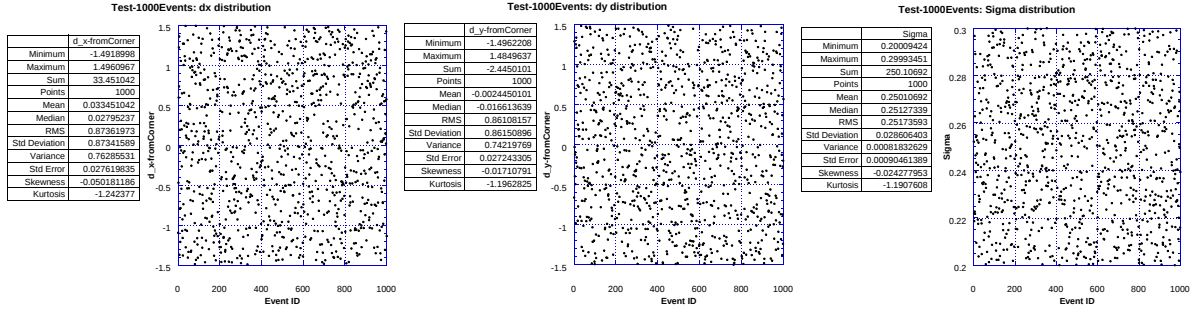
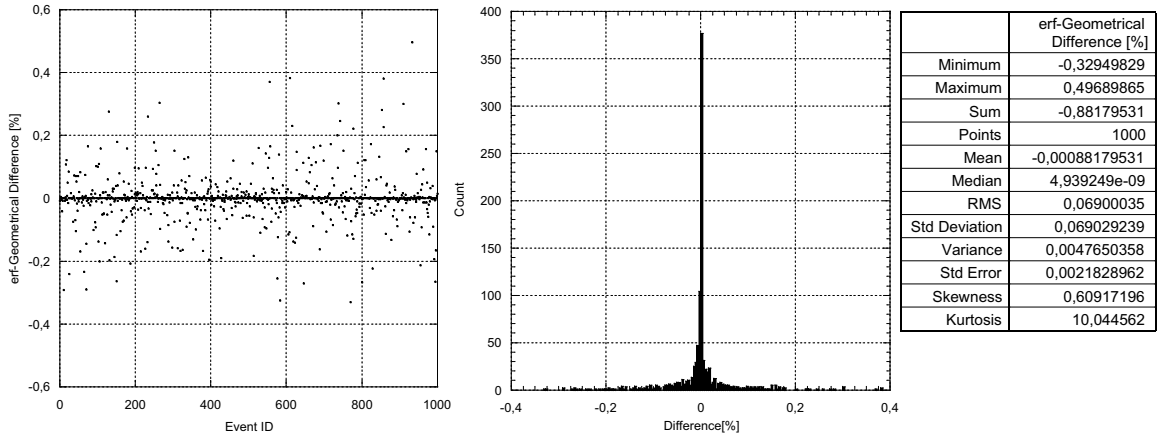
Figure 3.33: Test Distribution of d_x , d_y and σ for 1K events

Figure 3.34: Difference between the counted charge collected and the calculated one with the eq.3.22 for 1K events.

can influence the strips/pad cluster size when the number of primary electrons is high and when some strips or pads have a small overlap with the electron cloud (i.e. when they are in the external part of the clouds and when the number of collected electrons is small but comparable with the threshold). Actually we have to remember the presence of the noise that will be higher than this error.

The erf function applied to the strips case

The results that we have obtained for pads seem to be good, for its simplicity and for the reason that it is enough intuitive (surely more than a polynomial fit of order n) and realistic. The fact that the bi-dimensional description is practically an iteration of the mono-dimensional one opens the possibility to describe also the strips in the same way.

In particular the idea is that the charge collected by a strip can be described as:

$$n(d) = 50(1 - \operatorname{erf}(m3 \cdot (d - \frac{\text{StripWidth}}{2}))) - 50(1 - \operatorname{erf}(m3 \cdot (d + \frac{\text{StripWidth}}{2}))) \quad (3.23)$$

Where d is the radial distance between the center of the gaussian electron cloud and the middle of the strips. In the *StripsWidth* we have to insert the "effective width" and not the geometrical one, for the reason discussed during the strips study.

Practically in this way we calculate:

- The total charge in the region after the most close (with respect to the electron cloud) border of the strip.
- The total charge in the region after the most distant (with respect to the electron cloud) border of the strip.

- The difference between the values obtained in the previous steps (that is the charge collected from the strips).

This procedure could be used also for pads if we have a cloud that is bigger than the dimensions of the pads. The two methods have been tested comparing their prediction with the data obtained directly counting the electrons that fall into a strip. In fig.3.35 the distributions of d and σ used in the test are shown. The goodness of the prediction obtained with the Polynomial and the ERF function is shown in

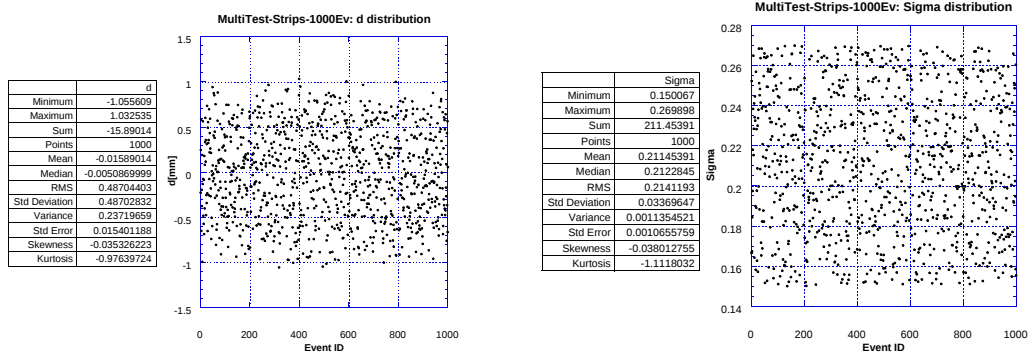


Figure 3.35: MultiTest Distribution of d and σ

fig.3.36 and fig.3.37 respectively. The erf results are more symmetric and closer to the data obtained from the geometrical superimposition between cloud and readout pattern. This could be related to the parameterization in σ that is surely better in the erf than in the polynomial. Even if the errors are acceptable for both, the erf seems to be the preferred solution because simpler, elegant, versatile and reliable.

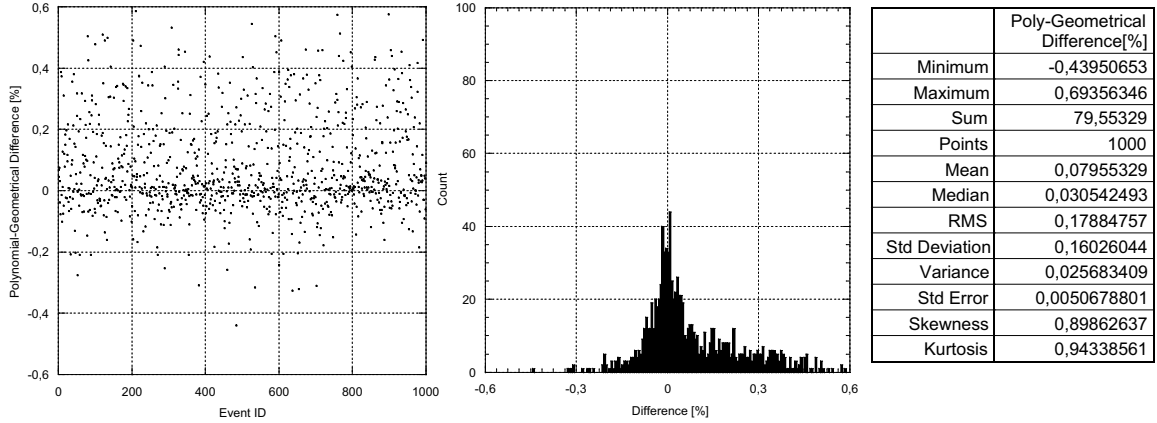


Figure 3.36: Difference[%] between the counted charge collected and the calculated one with the Polynomial approach.

3.2.3 Strips and pads Digitization algorithms.

In this section the final algorithms used for strips and pads are summarized. The method used is the one relative to the use of the erf function. The following formulas are referred to strips and pads separately. An additional algorithm is needed to combine them (i.e. remove from the charge collected by pads the part that it has been picked up by strips). In sec.3.2.4 a possible approach is given.

The algorithm, that provide the charge collected when a particle release N_C first ionization clusters of $N_C^{e^-}$ electrons each in the sensitive volume of a detector with total gain G , is for a pad α :

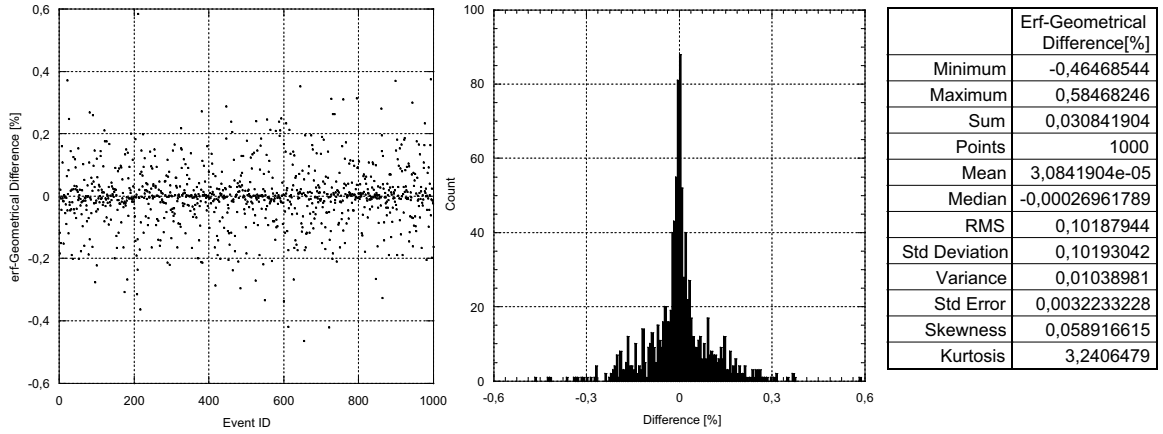


Figure 3.37: Difference[%] between the counted charge collected and the calculated one with the Erf approach.

$$N^\alpha = \sum_C N_C^{e^-} \cdot G \cdot n^\alpha(X_C, Y_C, \sigma_{ch}(Z_C)) \quad (3.24)$$

$$\sigma_{ch}(Z_C) = \sqrt{Z_C} \times \text{Diffusion Coefficient} \quad (3.25)$$

$$n^\alpha(X_C, Y_C, \sigma_{ch}) = \frac{1}{100} \times n(d_{x1}^\alpha, d_{x2}^\alpha, \sigma_{ch}) \cdot n(d_{y1}^\alpha, d_{y2}^\alpha, \sigma_{ch}) \quad (3.26)$$

$d_{i1,i=x,y}^\alpha$ = Transversal distance between the first ionization cluster and the first border of the pad α .

$d_{i2,i=x,y}^\alpha$ = Transversal distance between the first ionization cluster and the second border of the pad α .

$$n(d_1, d_2, \sigma_{ch}) = 50 \times |1 \mp 1 \pm [(1 - \text{erf}(\frac{d_1}{\sqrt{2}\sigma_{ch}})) \mp (1 - \text{erf}(\frac{d_2}{\sqrt{2}\sigma_{ch}}))]| \quad (3.27)$$

The upper (bottom) sign is used when the first ionization cluster is inside (outside) the pad α (see fig.3.38).

For a strip β only eq. 3.26 has to be modified because now the problem is solved in only one dimension:

$$N^\beta = \sum_C N_C^{e^-} \cdot G \cdot n^\beta(X_C, Y_C, \sigma_{ch}(Z_C)) \quad (3.28)$$

$$n^\beta(X_C, Y_C, \sigma_{ch}) = n(d_1^\beta, d_2^\beta, \sigma_{ch}) \quad (3.29)$$

d_1^β = Radial distance between the first ionization cluster and the first border of the strip β .

d_2^β = Radial distance between the first ionization cluster and the second border of the strip β .

Expressing d_1^β and d_2^β as a function of the effective strip width ESW^{13} and of the distance d between the cluster and the middle of the strip:

¹³The effective strips width is larger than the geometrical width of the strips. This is needed to take into account the electrostatic fields in the *induction zone*.

$$\begin{aligned} d_1^\beta &= d - ESW/2 \\ d_2^\beta &= d + ESW/2 \end{aligned}$$

Eq. 3.27 will be therefore expressed for strips as:

$$n(d, ESW, \sigma_{ch}) = 50 \times |1 \mp 1 \pm [(1 - \operatorname{erf}(\frac{d - ESW/2}{\sqrt{2}\sigma_{ch}})) \mp (1 - \operatorname{erf}(\frac{d + ESW/2}{\sqrt{2}\sigma_{ch}}))]| \quad (3.30)$$

In fig. 3.38 the previous equations are explained graphically. When the transversal coordinate of the first ionization cluster is inside (left picture) the investigated pad or strip, the charge collected is obtained subtracting from the integral of the whole gaussian the two areas over the two edges (at distance d_1 and d_2 from the cluster). When the transversal coordinate of the ionization cluster is outside the pad or strip (right picture), the collected charge is obtained subtracting the area over the most distant edge from the area over the nearest one.

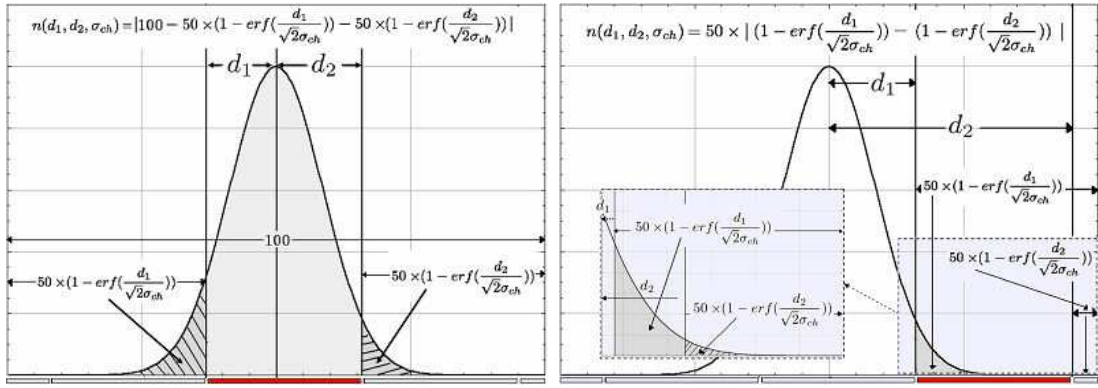


Figure 3.38: Collected charge obtained with the use of the erf function.

3.2.4 Strips and pads integration.

The previous algorithms have been developed individually for strips and pads. The strips and pads have to be treated together now, taking into account the readout pattern of the TOTEM triple GEM. One possible approach is:

- i calculate the charge collected by strips;
- ii calculate the charge collected by pads without considering the overlying strips;
- iii subtract from pads the charge collected by the overlying strips.

In point [iii.] it's important to take into account that each strip is over more than one pads and therefore it has to be removed from the pad only the amount of charge corresponding to the piece of strip that is exactly over it. One possibility to do it is schematized in fig.3.39. This procedure use the fact that for the cloud dimensions the readout can be described as orthogonal (square pads and parallel straight strips) and that the charge collected by pads is calculated applying the mono-dimensional eq.3.20 twice in two orthogonal direction (that practically cuts the cloud selecting only the electrons overlying the pad). In particular applying the first time the eq.3.20, with a cut parallel to the orientation of the strips, the charge collected by all the pads that are under the same group of strips is computed. The charge collected by the overlying strips is then subtracted and the remaining part is properly divided according to the position and dimension of the pads with the application for the second time of the eq.3.20, this time with cuts in a direction perpendicular to the strips.

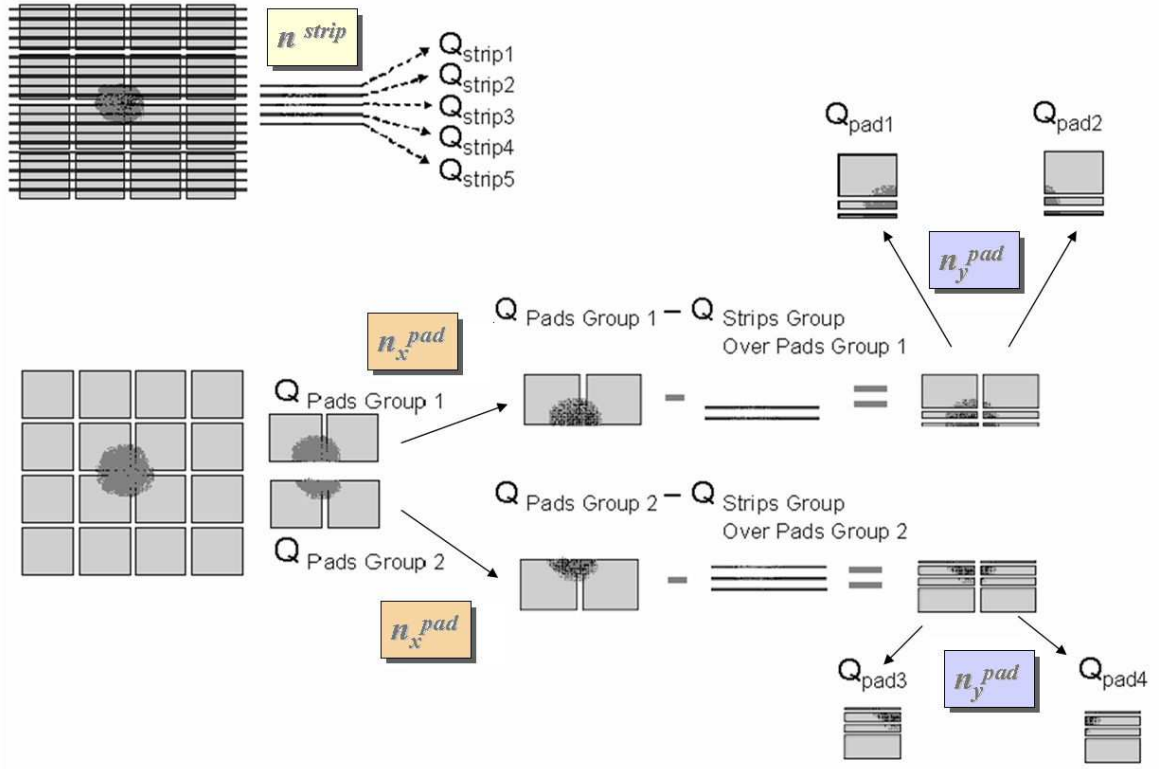


Figure 3.39: A possible scheme for the Strips-Pads Integration. i) The *electron collection curve* for strips n^{strip} will provide the charge collected on strips Q_{strip_i} . ii) The *electron collection curve* for pads n_x^{pad} will give the charge collected in all the pads with the same x (r) coordinate. iii) From this charge, that is linked to pads with the same radial coordinate, the charge of the overlying strips computed in i) is subtracted. iv) The n_y^{pad} will provide finally the charge collected by each pads Q_{pad_j} .

3.2.5 Digitization and Measurements

A preliminary comparison between the digitization and measurements will be given in sec. 5.2. Here the needed input (I) and the degree of freedom (DOF) available to tuning the digitization will be briefly summarized:

- | | | |
|--------|-----|---|
| (i) | I | Particle trajectory in the sensitive volume of the detector. |
| (ii) | I | Energy released and eventually the location of the clusters of first ionization. |
| (iii) | DOF | Detector Gain. |
| (iv) | DOF | Gas Mixture and the operating conditions (pressure, temperature, humidity, ...). They will be used to characterize spatially (σ) the electron cloud that will be picked up by the readout plane. |
| (v) | DOF | Readout Pattern and the effective strip width ¹⁴ . |
| (vi) | I | Noise Level in terms of electrons. |
| (vii) | DOF | The VFAT2 DAC step in terms of electrons. |
| (viii) | I | The threshold level, in terms DAC steps, programmed inside the VFAT2 chip. |

The found algorithms are very versatile and they can be easily adjusted to detectors and readout different from the TOTEM ones, setting properly these parameters.

Chapter 4

No Beam Detector Test and Front End Electronics Integration

In this chapter the measurements done for testing the TOTEM triple GEM detector will be presented. In sec. 4.1, tests with standard electronics will be shown. These tests have been done for checking the performances and the quality of the detectors. The followed procedure was the one used for the Triple GEM of the COMPASS experiment. The results obtained have been well within the expectation. Once the VFAT2 has been available, measurements with the final chip have started and will be presented here. The main issue was the noise level that was absolutely unacceptable at the beginning. Actually, the more critical point were the evidence of a noise dependence on the chip digital activity and the propagation of this noise in the readout plane through the pad and strip cross talk. This problem has been solved after having spent a lot of efforts and time. An acceptable level of noise has been reached working on the chip grounding, detector shielding and reducing, where it has been possible, the effects of the cross talk between channels (grounding capacitive coupling of the strip fan out). In sec. 4.2 a brief summary of what it has been done is given. Measurements performed with the triple GEM and VFAT2 on beam, cosmic rays and pp collision will be presented after, in ch.5.

4.1 Detector Test

Two triple GEM detectors for the TOTEM experiment have been assembled by the Italian private company G&A [37]. I've done some tests on these two chambers at the CERN laboratory of the Gas Detector Development (GDD¹) group at the beginning of my PHD. The aim of these tests was to give me the possibility to learn how to work with this kind of detector and to understand the laboratory equipment needed to start a research line in this field in our laboratory. It was moreover a good chance to verify the quality of the detectors assembled by the Italian company. In this section I will describe the various tests performed that have followed the procedures described in ref. [38].

4.1.1 Preliminary Tests

Preliminary tests have to be performed before and after the assembly of the detector and each component has to be singularly tested to guarantee its quality and functionality. In this section I will report some of the preliminary tests that have been done on the two detectors that have been assembled by the Italian company.

¹The research leader of this group is Leszek Ropelewski. Before him this role was covered by Fabio Sauli, the inventor of the GEM foils.

Before giving a description of these test, in which I have been involved, I will briefly mention two preliminary test that I've have not done but that I consider useful. The first is related to the test of the quality of holes of the GEM foil. The method that our colleagues of Helsinki (that have realized the 50 triple GEM for the TOTEM experiment) have developed for this purpose is based on a commercial scanner and on the elaboration of the acquired images. This system is able to check the quality of the holes revealing for instance the presence of holes that are blind or that have a wrong size or shape. The second test is on the readout pcb. It is possible that there would be shorts between strips and pads, broken electrodes or missing contacts that has to be found before the assembly of the detector. The check can be done with a noise measurement on the readout electrodes. The anomalous capacitance associated to the previous defects influence strongly the noise level of the involved channels that will be easily recognized in the noise pattern of the readout plane.

The single foil discharge

The single GEM foil discharge test was performed at the G&A company during and at the end of the detector assembly. The test was made with the Keithley Model 237 High Voltage Source-Measurement Unit. The result was a current $I < 100pA$ for 5 minutes at 550V for every foil.

A typical procedure for this test is :

- Flow the chamber for 2/3 hours with nitrogen (2-16 l/h).
- Set the power supply current limit to 50nA.
- Link to ground the single bottom plane and the not tested sectors on the top of the foil .
- Increase the H.V. up to 550V with a rate of nearly 10V/s.
- Record the leakage current and the number of sparks, if any.

Note: if a spark occurs, wait some minutes until repeating the ramp-up.

The validation of the single foil is obtained if : *the leakage current per sector at 550V and for 5 minutes is lower than 1nA in dry air, lower than 0.5nA in N_2 and lower than 5nA in $Ar - CO_2$ [38].*

The H.V. Distribution Board

The high voltage is provided to each GEM foil and to the drift cathode with a resistive divider. The H.V. boards used for the two detectors assembled by the G&A have been mounted in our lab. No coating or heating processes have been done on them. The resistors chains used are shown in fig. 4.1.

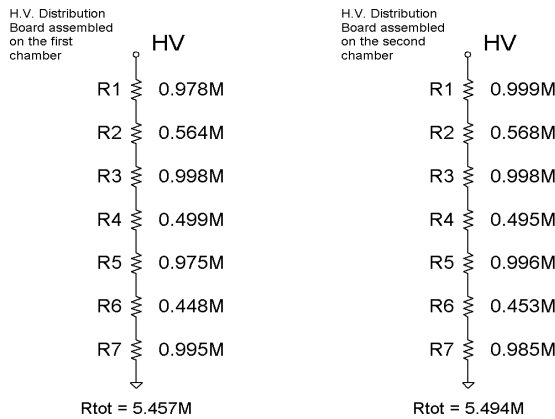


Figure 4.1: Resistive divider of the H.V. Distribution Boards mounted respectively on the first and on the second triple GEM assembled at G&A. The gain of the detector has an exponential dependence on the voltages applied to the GEM foils. The knowledge therefore of the exact values of the resistors used in the H.V. distribution board is necessary for comparing different chambers.

Before soldering the foils, each H.V. divider has been tested up to 5kV and characterized with a voltage-current curve. These data have been used as a reference when the detector has been fully assembled. In particular, once the detector was finished, the current limit of the H.V. power supply has been normally fixed 1 – 2 μA above the current previously measured. This limitation is done to avoid damages caused by internal discharges.

The external and internal discharge

Discharges test are performed once the detector is assembled and the gas tightness verified. The first test is done on the external ones that are not caused by the electrons multiplication inside the GEM foil. For this purpose the chamber is flushed with only CO_2 . The second test is done instead with an $Ar - CO_2 70/30$ mixture to test the internal discharges that are instead related to the electrons multiplications. Before starting the test, the chamber has been left under a flux of 5l/h of CO_2 for nearly 12 hours to clean it and then the voltage has been applied. It was carefully raised up to 4kV, reducing each step from 500V to 50V increasing the high voltage. Every time, the voltage has not been increased until the discharges disappeared. I've monitored the presence of discharges looking with an oscilloscope at the output of a preamplifier'amplifier readout chain connected to a group of strips and pads. I've seen only external discharges in the first of the two chamber that have been probably caused by dust on the H.V. distribution board.

The validation is obtained if: *after an hour of monitoring discharge does not occurs, the high voltage distribution network is validated. It is quite usual that some discharges occur in the first minutes with high voltage on, due to dust and metal splinters, but their frequency should decrease quickly, since these impurities are burned away [38].*

4.1.2 Absolute Gain Calibration

Once the detector has passed the basilarly test, a characterization of the performances has been done. In particular I've realized measurement of the absolute gain, charge sharing and energy resolution. The radiation source used to perform these tests has been a Cu x-ray tube. This source is particularly indicated because the interactions rate is regulable, the numbers of electrons released per photon in the sensitive volume is large and the x-ray beam can be collimated in a sufficiently small spot.

Measurement Description

The absolute² gain has been obtained using the relation:

$$GAIN = \frac{I_{tot}}{n \cdot e \cdot f} \quad (4.1)$$

where:

- I_{tot} = Total current collected by strips and/or pads.
- n = Mean number of electrons produced by the incident particle (or photon as in our case) in the sensitive volume.
- e = Electron Charge.
- f = Rate of Interaction.

Strips and pads can be considered individually or together inserting in eq. 4.1 the current and the rate on strips and pads or the sum.

The mean number n of electrons produced per the incident photon

The mean number n of electrons produced in the sensitive volume can be computed considering the energy released by the incident photons in the sensitive volume. The Cu X-ray tube emits two lines $K_\alpha \sim 8KeV$, $K_\beta \sim 8.9KeV$. In Ar, considering a fluorescence ($\sim 2.9KeV$) yield³ of $\approx 15\%$ and an average energy required to produce one el-ion pair of $26eV$, there will be for the two Cu K-lines:

²The adjective absolute is related to the fact that the measurement expressed in eq. 4.1 is not based on calibrations coefficients of the readout system as it is for example for a gain measurement based on the output level of a charge preamplifier.

³The photoelectric absorptions can be followed by an emission of a photon. If an Argon atom will follow this de-excitation mechanism, an X-ray photon at $2.9KeV$ will be emitted. Its mean free paths is very long and it can escape the detection volume without converting its energy in primary electrons.

- $n(K_\alpha : 8keV) \sim 0.85 \cdot \frac{8keV}{26eV} + 0.15 \cdot \frac{8keV-2.9keV}{26eV} \sim 290 K_\alpha$ primary electrons;
- $n(K_\beta : 8.9keV) \sim 0.85 \cdot \frac{8.9keV}{26eV} + 0.15 \cdot \frac{8.9keV-2.9keV}{26eV} \sim 325 K_\beta$ primary electrons;

With an intensity ratio between the two lines of about $\frac{I(K_\beta)}{I(K_\alpha)} \sim 0.135$, the total number of electrons released will be:

$$- n = (1 - 0.135) \cdot n(K_\alpha) + (0.135) \cdot n(K_\beta) = 0.865 \cdot 290 + 0.135 \cdot 325 \sim 293 \text{ Primary Electrons.}$$

In fig. 4.2 it is shown a simulation that I've done with HEED and GARFIELD to obtain this mean number n where I've used the right gas mixture (Ar/CO_2 70/30) and the intensity ratio $\frac{I(K_\beta)}{I(K_\alpha)} = 0.135$ for the two lines K_α and K_β of Copper. The mean number of electrons produced per incident photon is 260 instead

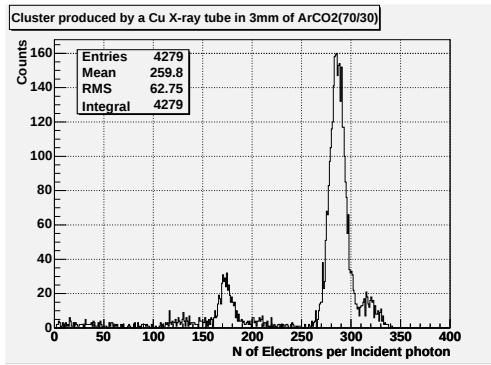


Figure 4.2: Total number of electrons released in 3mm of Ar/CO_2 70/30 by photons produced by a Cu X-ray tube obtained with an HEED/Garfield simulation. In this simulation I've have obtained 4279 interaction over 100000 initial photons ($\sim 5\%$). The mean number of electrons produced by the incident particle is ~ 260 .

of 293 as previously calculated. The difference is due to the fact that before the presence of CO_2 has not been considered. Even if the simulation result is closer to the real value, I will use in the next the value of 293 as it has been done normally in other papers.

The Current measurement

The measurement of the total current I_{tot} from strips and pads has been done grouping together 128 Strips and 120 Pads respectively and reading the current induced. The readout configuration is shown in fig. 4.3. The large number of strips and pads used for each group avoids problems related to edge effects and ensures that all the electrons produced per incident photon are collected. The current measurement is performed with a relatively high x-ray flux to have a voltage across the 1M resistors measurable with a standard multimeter.

The rate measurement

The system used to measure the interaction rate in the sensitive volume is shown in fig. 4.4. The signals coming from a group of strips and a group of pads has been read with an ORTEC Charge Sensitive Preamplifier (Model 142IH) and an ORTEC Research Amplifier (Model 450). The output signals have been sent to a discriminator unit, whose output have been used to measure the rate of interaction with a scaler. The level of the discriminator has to be fixed sufficiently low to measure all the signals produced by a photon interaction in the *drift zone* and sufficiently high to be outside the noise. This can be checked looking at the output signal of the amplifier on the oscilloscope and with the energy spectrums that are provided by a LeCroy ADC (Model 2249A) module. The relatively high interaction rate needed for current measurement may lead to an inaccurate measurement of the rate for the possible pile up of events. The rate measurement was made therefore inserting an absorber (a thin Cu foil) between the X-Ray tube and the detector. The absorber efficiency has been measured at low rate (fig. 4.5) and the ratio of the interaction rates without and with the absorber has been found. The measurement of the high interaction rate used during the current measurement and needed in eq. 4.1 has been calculated scaling with this ratio the rate measured with the absorber inserted.

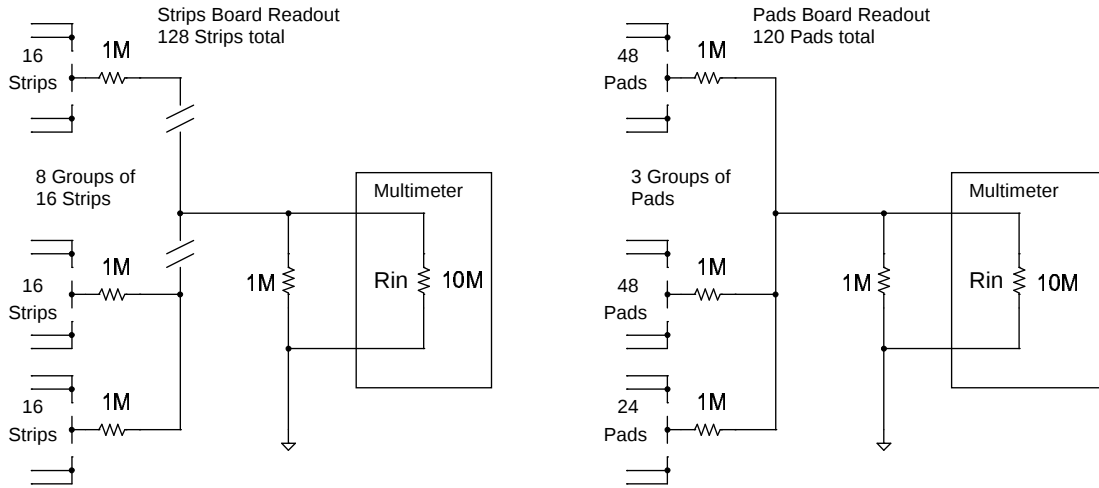


Figure 4.3: Readout scheme for the measurement of the current collected from Strips or Pads Group for the test on the second chamber assembled at G&A. The $10M\Omega$ input impedance of the multimeter has been taken into account computing the current.

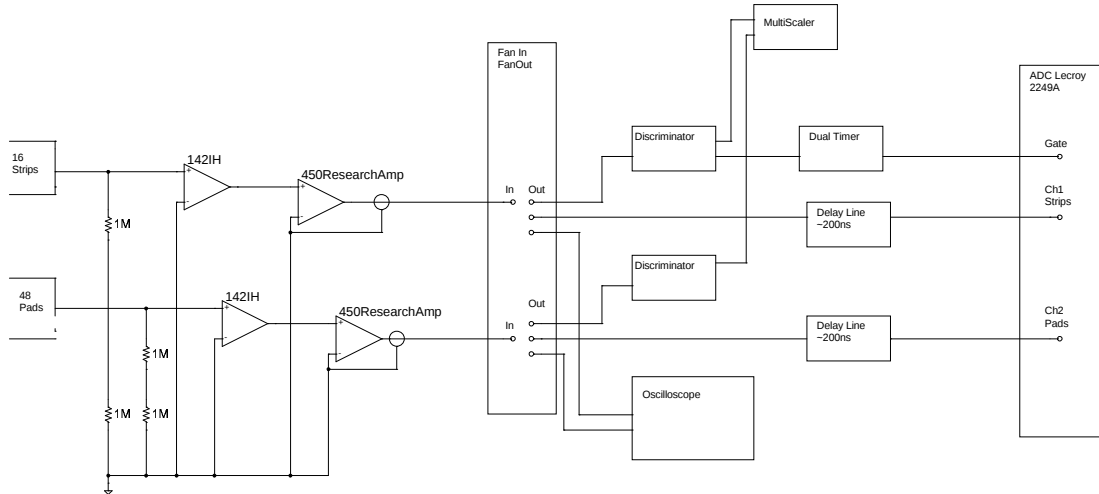


Figure 4.4: Readout Scheme for rate and spectrum measurements.

The absolute gain

Using the previous measurements and eq. 4.1 it is possible to characterize the gain performance of the detector. Fig.4.6 shows the total gain (strips and pads together) calibration curve of the two triple GEM assembled by the G&A. The value obtained are compatible with the expectation and with the a total gain of 8000 requested by the TOTEM experiment. An increase of 100V of the high voltage ($\approx 10V$ on each GEM foil) at $-4kV$ is roughly equivalent to doubling the gain. Even if a limitation due to discharge rate doesn't allow to increase the gain too much, there is still a margin to increase the induced signal if the ratio with the noise is not sufficiently high. In the same figure it is shown the relative gain of strips and pads for the second chamber it is also shown. The ratio between the two types of readout electrodes is coherent with the readout pcb design as will be discussed in the section on the charge sharing measurement.

According to the H.V. divider shown in fig. 4.1, the voltages across the GEM foils and the internal fields for the tested high voltages applied to the detectors are shown in fig. 4.7.

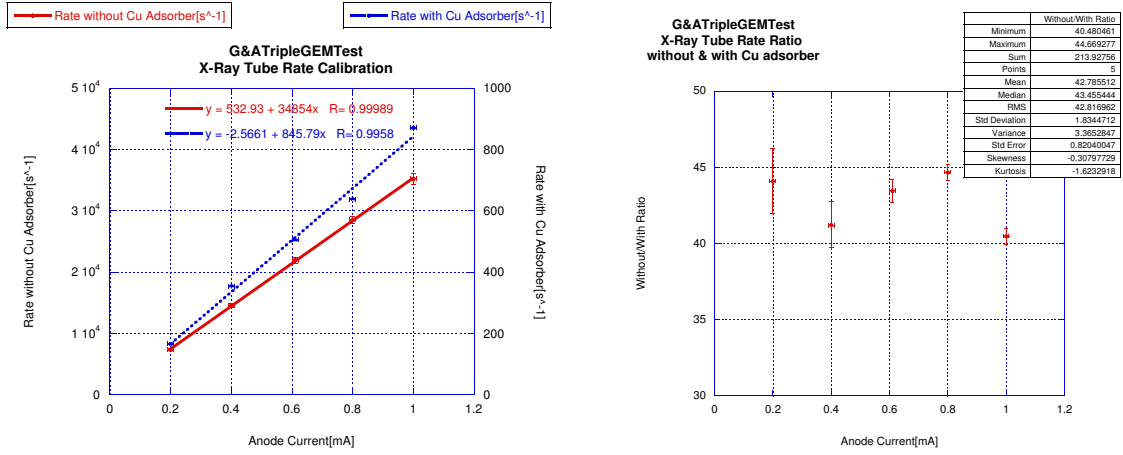


Figure 4.5: X-Ray absorption rate versus the anode current of the anode Cu Tube (right) and Ratio of the X-Ray absorption rate without and with the Cu Absorber Foil on the X-Ray Collimator, measured at various anode X-Ray tube current(left).

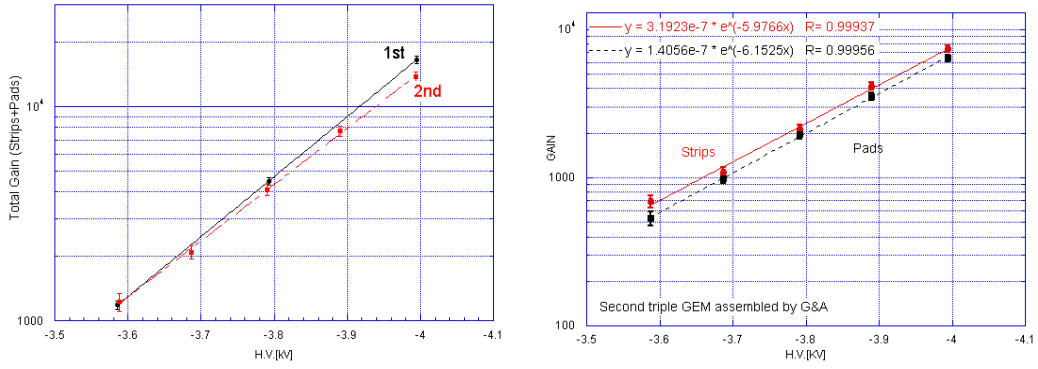


Figure 4.6: Left: Total (Strips+Pads)gain comparison between the first and the second triple GEM assembled by the G&A company. The GAIN is obtained using Eq.4.1. The current has been collected from a group of 16 strips and 48 pads. The interaction rate evaluation has been done on the strips signal. The number of electrons produced from the radiation is assumed equal to 293. Right:

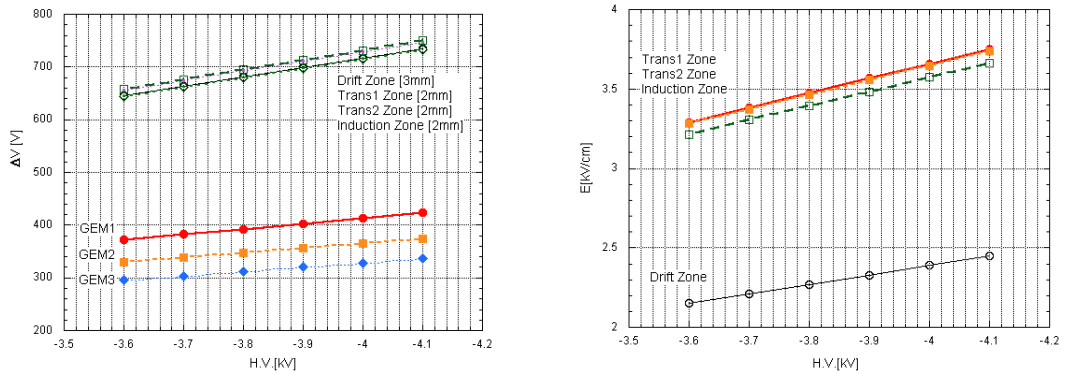


Figure 4.7: Left: Internal ΔV versus the high voltage applied to the first resistive divider shown in fig. 4.1. Right: Electric fields inside the drift, the two transit and the induction zone versus the high voltage applied.

4.1.3 Energy Resolution Studies

The energy resolution is obtained from the energy spectrums acquired with the readout system used is shown in fig. 4.4. This measurement it is particularly useful because it is strictly related to the quality and uniformity of the GEM foils. If the gain is not uniform for the low quality of the GEM foil in the irradiated zone there will be an anomalous broadening of the spectrum's peaks with a worsening of the resolution. All the spectrums have been acquired with a Cu foil on the collimator output of the X-Ray gun to have a better sensitivity. This absorber indeed reduces the bremsstrahlung components and increments the K peaks of Cu (fig. 4.8).

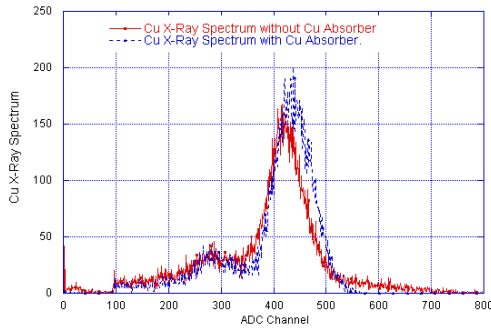


Figure 4.8: Cu tube X-Ray Spectrum acquired without and with a Cu Absorber Foil on the X-Ray collimator. The Cu absorber reduces the bremsstrahlung and increases the $K\alpha$ (and $K\beta$) emission of Cu. This is due to the conversion of the high energy part of the bremsstrahlung absorbed by the Cu foil.

The energy resolution of the two tested chamber in function of the H.V. applied is shown in fig. 4.9. For both the resolution is $\sim 21\%$ for $H.V. \geq -3.7kV$. Fig. 4.10 shows examples of energy spectrum at various

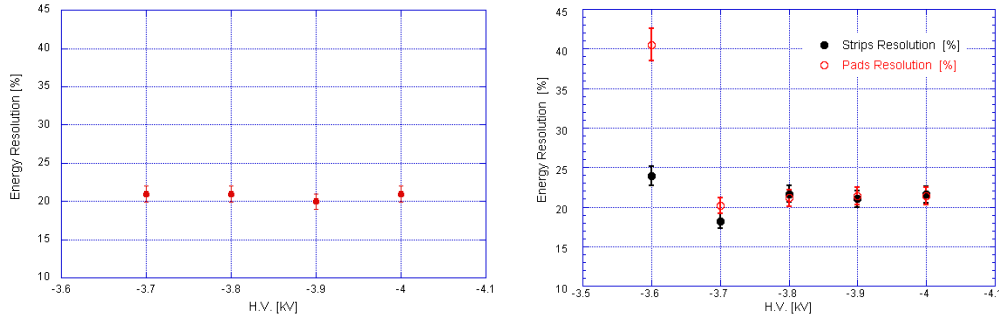


Figure 4.9: Left: Energy Resolution of the $K\alpha$ (and $K\beta$) emission peaks of Cu versus the H.V. applied to the triple GEM for the first chamber (left: strips readout) and the second (right: strips and pads readout).

H.V. applied. In particular they have been measured from a group of strips in the first tested chamber. In each spectrum two peaks are visible. The one at higher energy is generated by the not distinguishable $K\alpha$ and $K\beta$ photons at respectively $8keV$ and $8.9keV$. The energy resolution is measured on this peak. The one at lower energy is instead related to the argon escape and it has to be found at $\approx 2/3$ of the other (the x-ray emission from argon remove $2.9keV$ from the energy of $8/8.9keV$ that will be converted in electrons).

In fig. 4.11 strips and pads correlation spectrums measured with an applied H.V. of $-3.6kV$ (first row) and $-4kV$ (second row) are shown, where the increase in resolution with higher voltage (gain) is evident. From the bi-dimensional plots (left column) information about the charge sharing can be directly obtained from the slope of the linear fit of the data. A slope ≤ 1 means an larger charge collection on strips that are on the abscissa of the plots. Obviously this is true if the two amplification chains of strips and pads are equivalent. The same information can be extracted measuring the $K\alpha$ (and $K\beta$) position in the strips and pads spectrums like the ones shown in fig. 4.10.

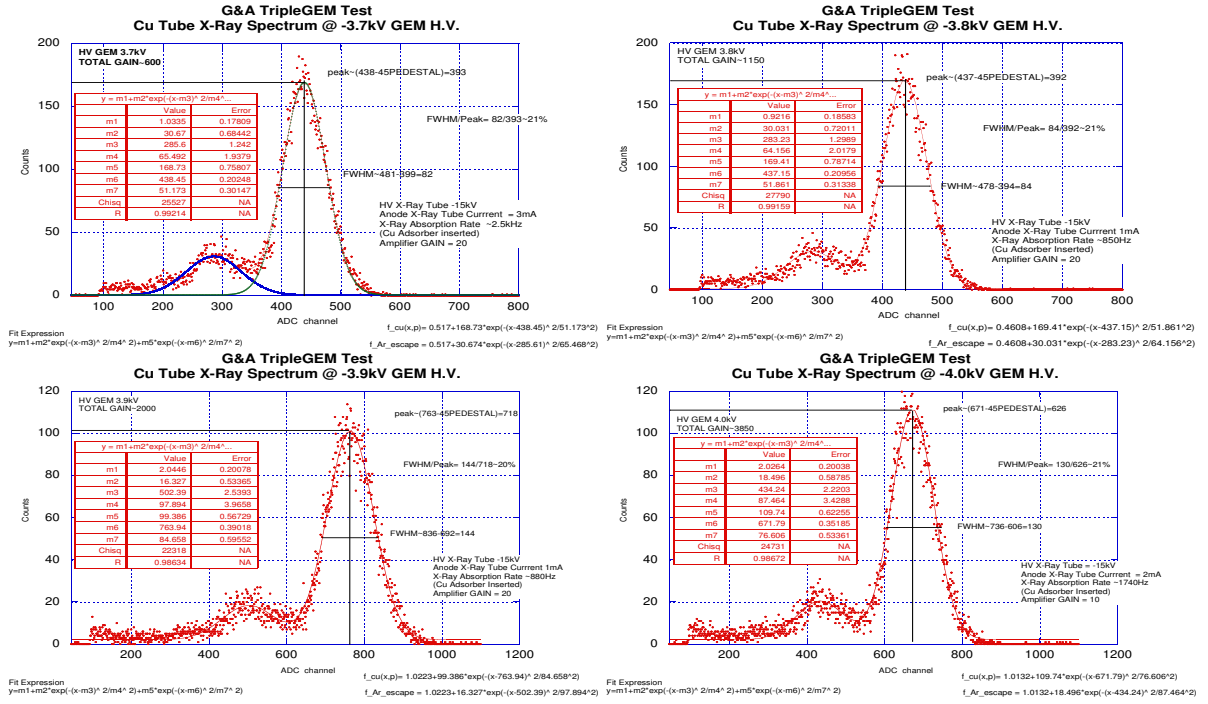


Figure 4.10: Cu tube X-Ray Strips Spectrum as a function of the H.V. applied to the first triple GEM tested. The readout has been done on a group of 8 strips.

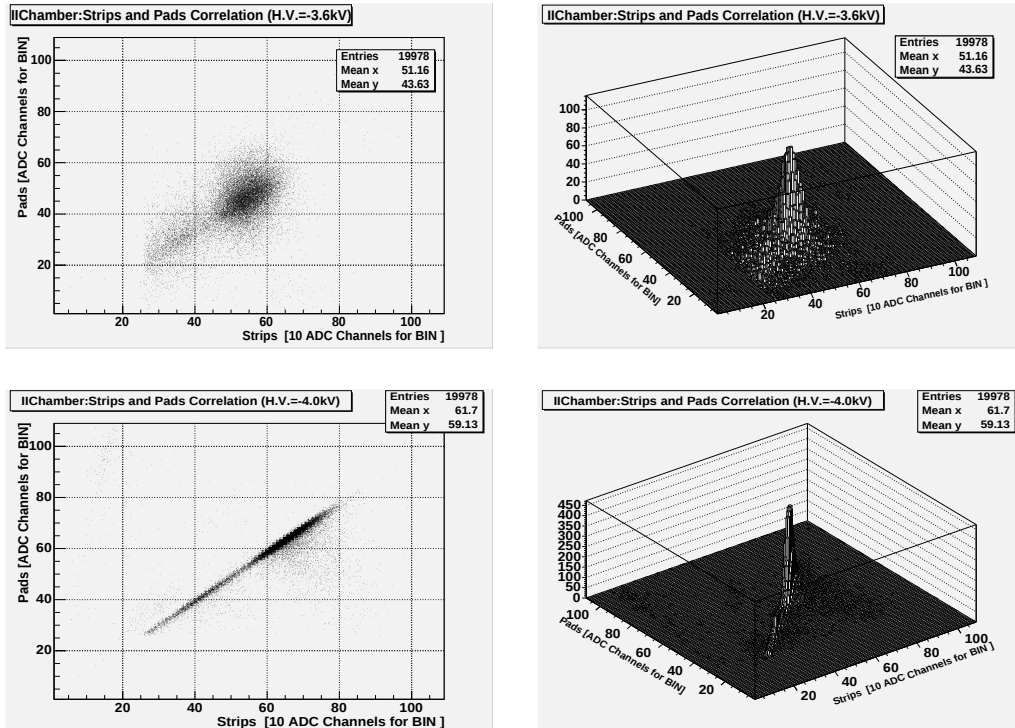


Figure 4.11: Strips-pads correlation spectra for the triple GEM exposed to Cu x-ray tube. Measures for the second chamber tested with an applied H.V. of -3.6kV and -4kV. In the left column plots the abscissa is related to strips and the ordinate to pads.

4.1.4 Charge Sharing

In Fig.4.12 it is shown the charge sharing between strips and pads for the first and the second chamber. The sharing has been obtained from the measurement of the current collected by a group of strips and a group of pads. The behavior is nearly the same for the two chamber with a pad signal about 10 – 15% lower than the strip signal.

This result confirm the expectation. Pads and strips are read with the VFAT2 front end. When a particle ionizes the gas, there will be normally involved one or two pads and two or three strips. For this reason the readout plane was designed to have a pad signal about 10% lower than the strips one. This charge sharing should provide a good signal to noise ratio for both.

The sharing can be obtained also from spectrums measurements as discussed in the previous section. It can be evaluated indeed from the measurement of the peak position in energy spectrums (fig. ??) of strips and pads or directly from the correlation plots (fig. 4.11). It is important to take into account that the ratio between the peak position for strips and pads reflects their charge sharing only if their two amplifications chains are equivalent. In Fig.4.13 we have compared the difference[%] of the current collected from strips(128) and pads(120) with the difference[%] of the position of the K_α peak in the strips(16) and pads(48) spectrums. In the same plot it is also shown the difference[%] of the gain of the two readout electrodes. According to eq. 4.1 the difference with respect to the ratio of the measured currents is related to the fact that in the gain measurement also the rates are involved. These measurements could be affected by various errors as the presence of shorts between strips and pads, a wrong rate measurement or a not completely equivalent amplification chains for strips and pads. Nevertheless the results obtained confirm the expected charge sharing with more or less a charge collected by strips 10% larger than for pads.

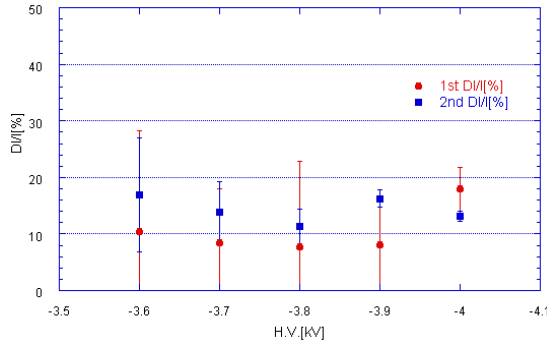


Figure 4.12: Charge sharing between strips and pads for the first and the second chamber assembled by the G&A company. In the plot the quantity $(I_{strips} - I_{pads})/I_{strips}$ is shown versus the high voltage applied to the detector.

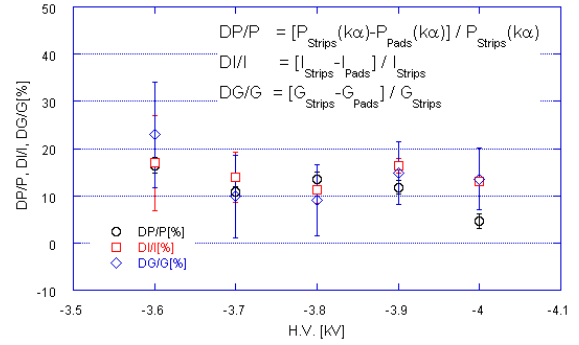


Figure 4.13: Comparison between the relative position of the K_α peak, the relative current collection and relative gain of strips and pads for the second chamber assembled by the G&A company.

4.1.5 Uniformity, Stability and Charging-Up

There are other tests that have to be done before the validation of the detector. In particular they are related to the study of uniformity of the performances in the sensitive area, stability and charging up of the GEM foils. During the period that I've spent at the CERN I had no time to complete all of these measurements. I've tried however to understand the procedures followed to perform them.

Uniformity Test

The aim of this test is to verify the homogeneity of the response of the chamber in the active area. I had no time to perform a complete test of uniformity and I've just looked to three different position on the detector. This was done moving the x-ray gun with respect to the chamber and looking for the maximum of the signal for strips as well as for pads. The overlapping of three spectrums obtained from

three different groups of eight strips is shown in Fig.4.14. Even if a complete measurement of uniformity

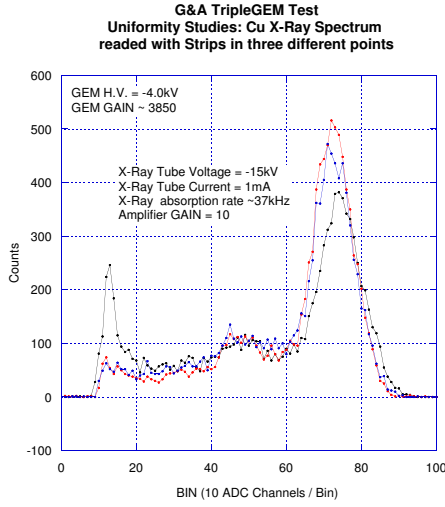


Figure 4.14: Cu-tube X-Ray Spectrums acquired on three different point over the chamber from a group of eight strips for an applied H.V. of $-4.0kV$ on the triple GEM.

has not been done, there are some useful considerations that follow from these measures. In this kind of test it is important to take into account the presence of sector boundary or spacer that can influence the charging up and the response of the detector. Moreover it useful the comparison between different chambers to understand if there are common behavior. A good solution could be the use of a mask placed in front of the detector that fixes univocally the positions that have to be investigated. In this way the measurement is faster and a correlation of the data with the detector structure and with others chambers can be easily done. Moreover a large group of strips and pads is suggested (here we used only eight pads or strips) to be sure to really collect all the electrons produced after the interaction.

Stability Test

We have done this test making spectrum acquisitions every 240 seconds. The software that we used saves the position of the Cu K_{α} peak. For these measurements we recorded 10 ADC Channels for each bin. In Fig.4.15 we have the peak position versus time and the overlap of the eighteen spectrums acquired during this test. The test that we have performed doesn't show any evident variation on the gain, but it

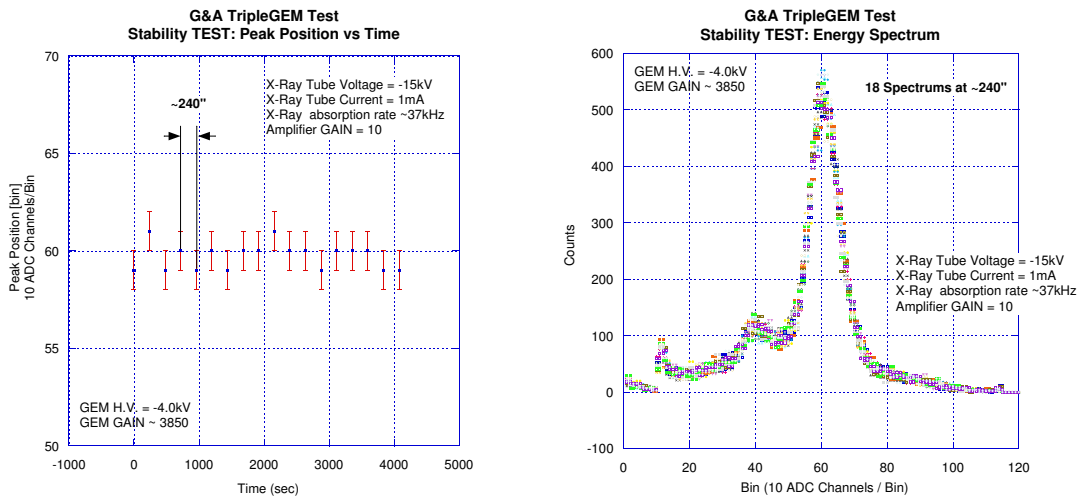


Figure 4.15: Left: Acquisitions of the K_{α} (and K_{β}) emission peaks of Cu every ~ 240 sec from a group of eight strips at $-4.0kV$ of TripleGEM H.V. Right: Eighteen Cu-tube X-Ray Spectrums acquired every ~ 240 sec from a group of eight strips at $-4.0kV$ of TripleGEM H.V.

was done just for ~ 1 hour. Actually longer acquisition in time are needed. In that case it could be useful

to monitor the humidity and temperature of the ambient to correlate variations of the gain with these parameter that will obviously change during long time period and that could affects the behavior of the detector.

Charging-Up Studies

A variation of the gain should be seen when the chamber is irradiated after a long period of no incident radiation. This is due to the charging up of the dielectric zones of the foils (hole, sector boundary, spacer). The time development of this gain variation should be faster for the charging-up of the inner part of the holes than for example for the sector boundary one and it could range from some tenth of seconds to some tenth of minutes.

In Fig.4.16 the variation of the K_{α} peak is plotted for spectrums acquired every $\sim 90\text{sec}$. No evidences of charging effects is expected because I've irradiated this zone to maximize the signal before starting the test and the time between each acquisition is too long (some minutes). In the same figure also the spectrums acquired during the measurement are shown.

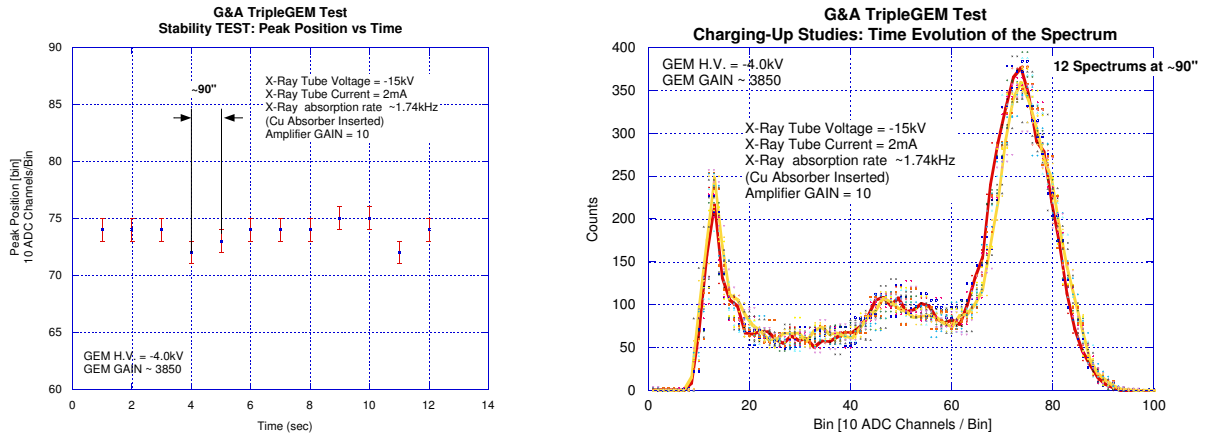


Figure 4.16: Left: Acquisitions of the K_{α} (and K_{β}) emission peaks of Cu every $\sim 90\text{sec}$ from a group of eight strips with -4.0kV applied on the triple GEM. Right: Twelve Cu-tube X-Ray Spectrums acquired every $\sim 90\text{sec}$ from a group of eight strips with -4.0kV applied on the triple GEM. The red line is the first and the yellow the last spectrum acquired.

To make this measurement in the right way we need to be very fast in the acquisition of the spectrum (few seconds) and to be able to center the X-Ray gun over the readout group of pads and strips without irradiating them before. This is a problem if a maximization of the signal is needed before starting the test. A mask as the one suggested for the uniformity test, that fix the relative position between source and detector, would solve this point. Otherwise the signal has to be maximized and the charging up measurement has to be done after a period of no incident radiation without having moved the X-ray Gun and the detector.

4.2 Front End Integration

In this section the first measurements made with the VFAT2 chip on the TOTEM Triple GEM Detector will be described. As previously reported, the VFAT2 is a chip used in all the sub-detectors of TOTEM. The chip has shown since the beginning good performance with the silicon detectors of the Roman Pots. The gaseous detectors (i.e. the CSC of the T1 telescope and the T2's Triple GEMs) have instead encountered few problems related to the noise level. Even if those detectors have enough gain to improve the signal to noise ratio, there is a limit to the quantity of charge that can be produced per incident particle due to the dynamic range of the chip. The reduction of noise with respect to the initial condition was therefore mandatory in order to provide the triggering and tracking efficiency required. For both T1 and T2 the final results have been in agreement with the requirements. This section will describe the noise studies done on T2. After a brief description of the available tools used to analyze the noise of our system, what it has been done to improve the noise for the T2 triple GEM will be presented.

4.2.1 Noise Measurement Tools

This introduction on the used tools should be sufficient to understand the plots and the measurements that will be presented in this section. A more detailed description is given in the VFAT2 manuals [32], [39] if needed.

The measurement of the noise level has been performed using three methods:

- The Calibration Pulse Scan.
- The Threshold Scan.
- The Trigger S-bit.

The Calibration Pulse Scan

The Calibration Pulse Scan uses the internal pulse generator of the VFAT2 (fig. 4.17) to measure the noise S-curve (fig. 4.18) of one channel. This curve is obtained fixing the threshold in the internal comparator of the VFAT2 and rising, for each step of the scan, the amplitude of the calibration pulse sent to the channel under test. A measurement of the threshold level and of the noise is provided respectively by the mean and by the σ of the S-curve erf fit. These quantity can be converted in equivalent electrons for signals with the same characteristics of the calibration pulse (a nearly ideal fast pulse). The VFAT2 datasheet gives an estimate value of 400-600 electrons per VFAT2 DAC step (the unit reported in the following plots). A direct measurement can be done to have accurate values. If we want to compare this equivalent number of electrons with the expected value for an incident particles in our gaseous detector, it has to be taken into account that the time properties of a real signal are different from the calibration pulse ones. Our signal is indeed spread over many clock cycles and more charge is needed therefore to reach the same level of voltage (in the comparator stage inside the VFAT2) as the one obtained with a fast pulse. The injection of charge can be done in parallel in many channels at the same time. The capacitance that is seen by the internal pulse generator is the parallel of the $100fF$ series capacitors of all the selected channels. This will reduce the charge collected by each channel and it will modify the time development of the calibration pulse. For all the reported measurements we have chosen to test only one channel at a time (using in this way a calibration pulse closer to an ideal pulse).

The Calibration Pulse Scan offers the most reliable quantitative measurement of the noise level. One example is given by the measurement performed on pads and by the information that can be extracted from it. In fig. 4.19 it is shown the S-Curve erf fit σ measured on pads of different area. The correlation between the noise and the readout pattern (i.e. area or capacitance of pads) of our detector is good. From this correlation, it is possible to evaluate the noise added by the pads fan-out (i.e. the connection between the pad and the input of the chip). Using the extrapolation to zero-area of the fit in fig. 4.19 and comparing it with the noise level obtained from the same chip alone it has been found that the noise added by the pads fan-out was negligible with respect to the pad itself. The σ of the S-Curve is a good quantity that can be followed during the improvement of the system. In fig. 4.20 one example of system monitoring is shown during one of the first tests on the detector with the VFAT2.

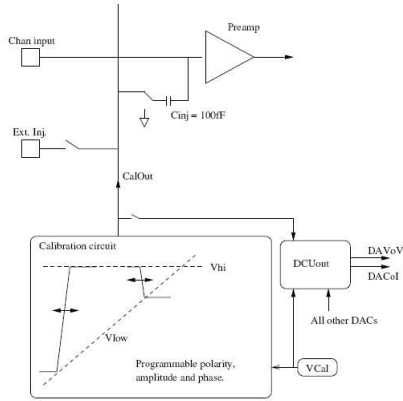


Figure 4.17: The calibration pulse amplitude. The calibration circuit generates a step between V_{Hi} and V_{Low} with a very fast rising edge and this pulse is delivered to the CalOut. Each channel has a 100fF series capacitors connected to CalOut via a switch. Hence a selected channel will receive an injected charge to the pre-amp of approximately 0.1fCmV [39].

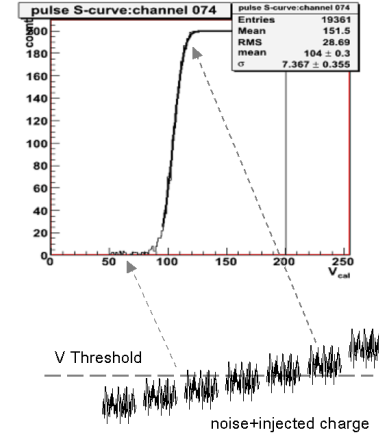


Figure 4.18: The pulse scan. The S curve is an histogram of hits whilst varying the injected charge (V_{Cal}) for a given threshold ($V_{t2V_{t1}}$)...The noise on the Scurve is the σ . The threshold is the mean. Since V_{Cal} adjusts the amplitude of the calibration pulse (the difference between V_{hi} and V_{low}) and the test pulse injection capacitor is 100fF we know the value of charge injected Q_{inj} and the corresponding number of electrons. Hence the noise and the threshold can be expressed in terms of electrons. [39].

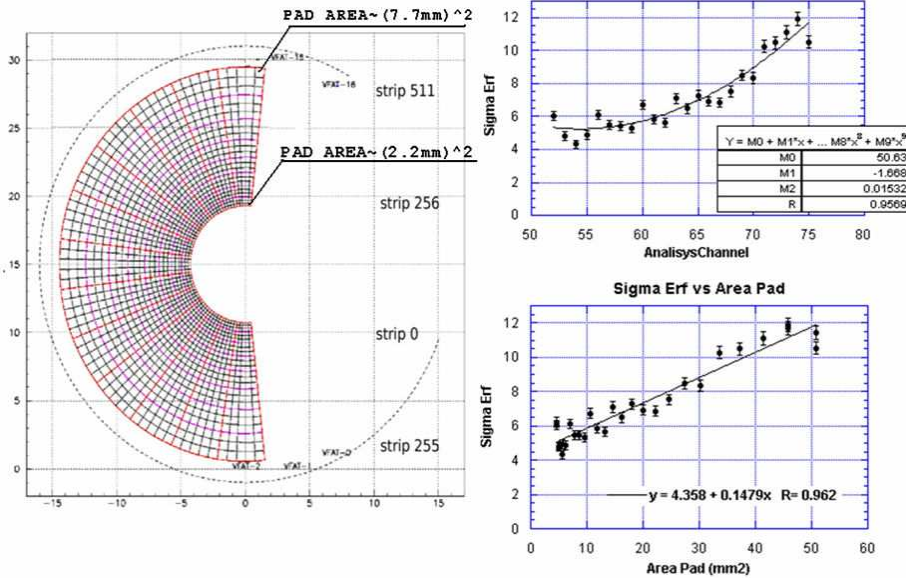


Figure 4.19: Pads Calibration Pulse Scan: (Left) Pads Readout Pattern of the TOTEM Triple GEM. The 1560 pads are organized in 65 radial columns. In each column the smaller pad area is of about $2.2 \times 2.2\text{mm}^2$, while the larger is $\sim 7.7 \times 7.7\text{mm}^2$. (Right-Top) Erf Fit σ versus the VFAT2 channel number (the selected pads belong to one column of the readout pattern). A quadratic dependence is fitted. (Right-Bottom) The same data are plotted versus the pad's area and linearly fitted.

The Threshold Scan

In preliminary studies, the time required to perform a calibration scan on a sufficient number of channels is too much. We used therefore the threshold scan measurements that can be faster and qualitatively reliable, even if less accurate. During a Threshold Scan, the value of the threshold is changed and the counts per threshold of all the VFAT2 channels are recorded. The test can be performed for all the channels in

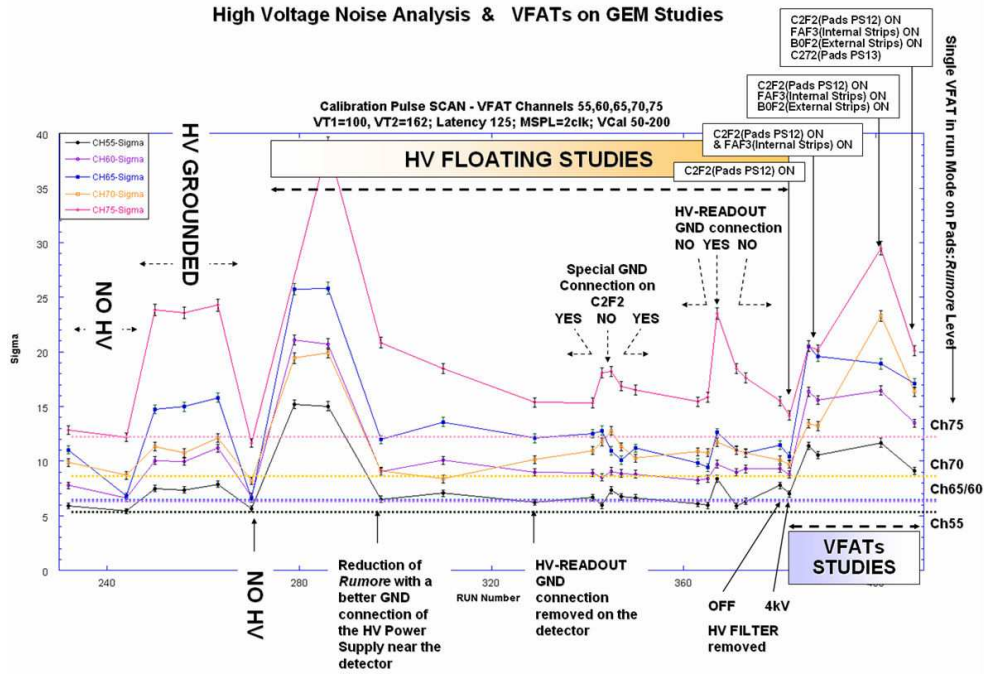


Figure 4.20: Calibration Pulse as a system improvement tool: in this example the σ obtained with the CalPulse Scan has been used as a reference for the initial studies on the effects of the High Voltage configuration on the noise level. Few channels of the same chip has been investigated.

parallel without degrading the result. The standard summary plots that we used as a reference for the Threshold Scan were the "Cumulative" (total number of counts versus the channel number), the "Noise S-Camel"⁴ (counts versus threshold for the selected channel) and the "Noise Start (5%)" (threshold level at which the count are equal to the 5% of the sent triggers). These plots are directly provided by the TOTEM software. In fig. 4.21 one example of these plots is shown.

The number of triggers per threshold that we used in these scans was fixed normally to one hundred.

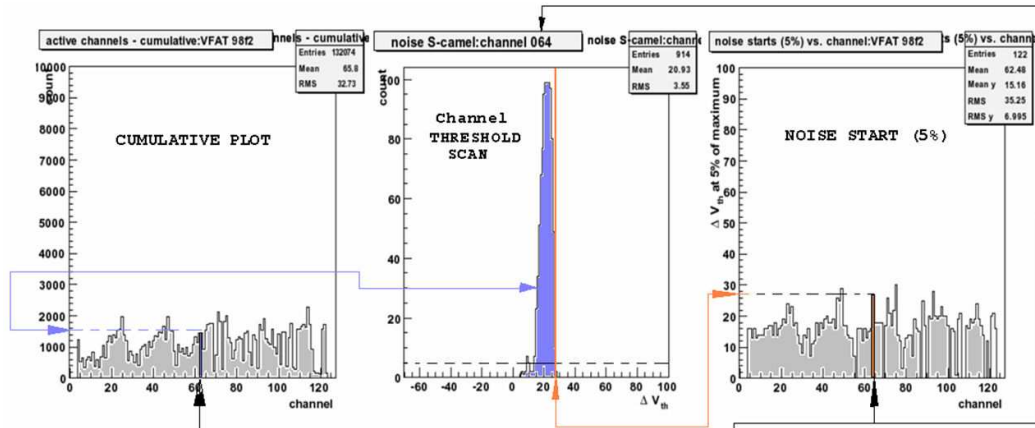


Figure 4.21: Threshold Scan: for each step of the scan the threshold of the comparator of the VFAT2 is changed and a fixed number of trivial triggers are sent to retrieve the status of the channels. For each channel, the number of counts per thresholds is plotted in the noise S-curve (middle plot). The integral of these curves is available, for all the channels, in the cumulative plot (left). The plot on the right however shows the thresholds where the counts have been the 5% of the sent triggers as a function of the channel number.

We have weighted more the testing time than the statistical quality of data. This affects the quantitative analysis of the results, giving mainly a qualitative indication of the system status. In order to evaluate

⁴The term "Camel", if I'm not wrong, has been introduced when strange structures have been observed in these curves. The proper name should be S-Curve as for the Calibration Pulse case.

the reliability information given by a threshold scan with a low statistics (i.e. a low number of triggers per threshold), we have done a comparison with the results of the Calibration Pulse Scan. One of these tests is shown in fig. 4.22 and it confirm the validity of the test. The measured trends of the compared quantity (i.e. the S-Curve σ of the calibration Pulse with the mean and the RMS of the Threshold Scan 5% Noise Start) are compatible.

One of the utility of the Threshold Scan can be appreciated on the plot of fig. 4.22, where the information

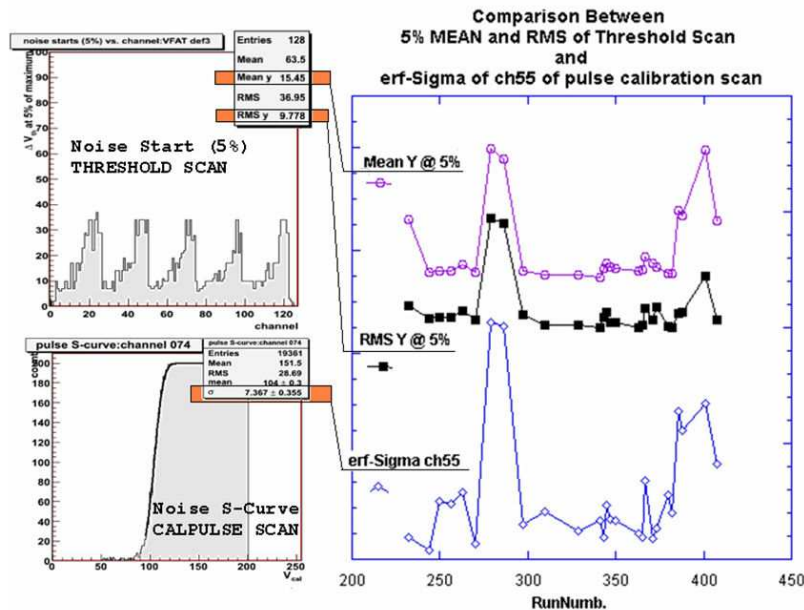


Figure 4.22: Comparison between Calibration Pulse Scan and Threshold Scan results: The information obtained by the left plots are compared in the right one. In particular the erf fit σ of "Calibration Pulse Scan S-Curve" of one channel (left bottom) is plotted with the mean and rms of the "Threshold Scan Noise Start (5%)". Actually the TOTEM software has an analogous "Threshold Scan Noise Start" computed at 1% instead of 5% that could follow also better the behavior of the Calibration Pulse Output. It's however less reliable than the 5% because with the used triggers per threshold (i.e. 100) only one hit (i.e. 1%) is totally useless.

about all the channels are measured at the same time. The structure, repeated five times, reflects the five columns of pads (each structure with 24 pads with increasing area) with their capacitance. If something different from this happens, it is a clear indication that other sources of noise are degrading the standard noise pattern of our detector. From the beginning of our tests, we discovered that the noise level of one chip is affected by the status of the other chips installed on the detector (i.e. their threshold or equivalently their digital activity). When we have encountered this kind of problems, the previous structures have shown anomalies that can be immediately observed with this kind of scan.

The Trigger S-bit

Even if the threshold scan is acceptably fast, there exists an even faster way to check roughly the noise status of the system: the status of the S-Bit lines of the VFAT2. These lines are driven by an OR combination of programmable groups of the VFAT2 channels. The status of the S-bit is high when one of the logically combined channels is over threshold. They are used for triggering purposes. A direct measurement performed with an oscilloscope of these signals provide a real-time measurement of the improvement obtained by an action done on the system (shielding, grounding, ...). Fig. 4.23 describes the internal scheme of the VFAT2 where those signals are formed. During the initial stage this possibility reduce enormously the efforts required to understand what it is happening on the system. One of the main advantages of this method is the fact that it represents the only way to understand or to see frequency characteristics of the noise pattern. This signal will be used (and it will be shown in the next sections) also for timing studies of the detector signal.

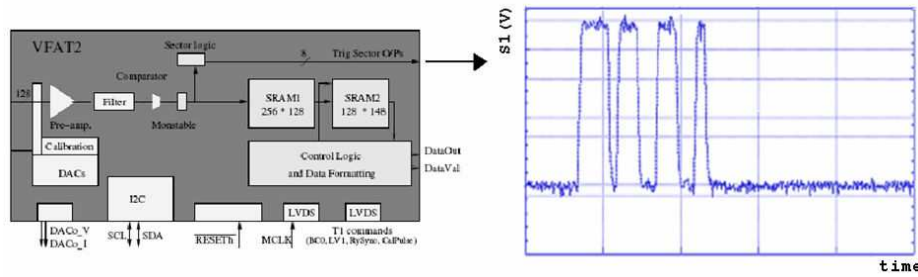


Figure 4.23: Trigger bits as a noise tool: The trigger SBit generated internally on the VFAT2 by the Sector Logic (right) has been used to have a real time and qualitative measurement of the noise level. A rough quantitative estimate can be done using counter, but we didn't it because the previous tools are more indicated to that purpose.

4.2.2 Noise Measurement Results

Improvements done on the VFAT2 and on the detectors to improve their performance are described here. The first results were completely unacceptable but at the end we were able to satisfy the requirements for the measurements that have to be done with the T2 telescope. The main problems were surely the noise dependence on the digital activity of the VFAT2 and the cross talk between chips mounted on the same detector. These problems have not been fully solved, but their effects have been reduced to a level where the signal to noise ratio was acceptable. In the following it will be summarized what it has been done.

The VFAT2 Hybrid

The Analog and Digital powers are supplied to the VFAT independently. Two places are provided on the hybrid where it is possible to join together the two grounds (see fig. 4.24). Noise test has been done for all the possible configurations (i.e. no connection, connection on one place, connection on both places). The results (see fig. 4.25) have shown that the noise level is degraded in the case of two connections. This could be explained taking into account the ground loop that is very close to the chip (and to its digital high frequency activity). The other two configurations are equivalent for the chip stand alone (i.e. not connected to the detector) as can be seen from the results shown.

The single ground connection is instead the preferred solution when the chip is mounted on the Triple

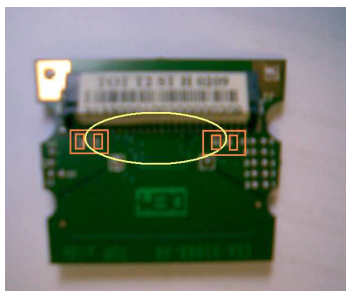


Figure 4.24: Digital and Analog Ground connection on the hybrid: Four places (two on the left and two on the right) are available on the hybrid to join the digital and analog ground. The connections of all the four point produce a very small, but also very close to the chip, ground loop that has been considered as the responsible of the worsening of the noise level for the chip stand alone.

GEM (see fig. ??). This solution has been therefore chosen for the final set up. The connection between digital and analog ground, directly on the hybrid, should be on of the best solution because in this way the two parts (analog and digital) are joined together in the closest place to the chip and to the ground plane on the detector. This is in agreement with the leading idea of using the detector ground plane as the main common reference for the detector subsystem (the same reference will be given to the high voltage line).

The second point, related to the VFAT2 hybrid, is the reduction of the noise level that we have found covering the chip with a grounded metal foil (see fig. 4.27). The effects is about a factor of two. The outputs of a threshold scan performed in two chips, without and with the cover, are shown in fig. 4.28.

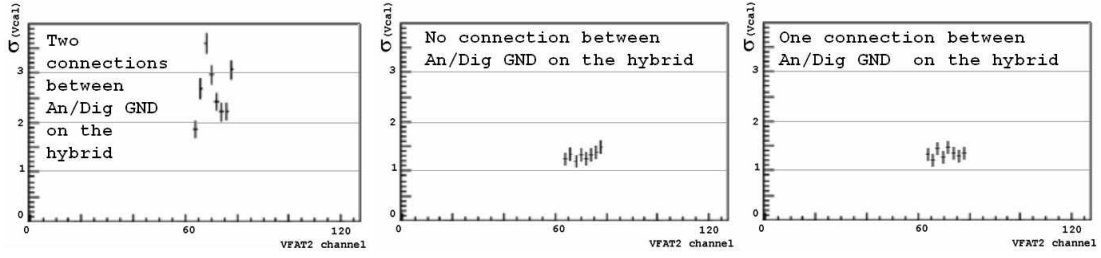


Figure 4.25: Test of the Digital and Analog Ground connections on the hybrid: Erf fit σ of the Calibration Pulse S-Curve for connections on both side (left), none (middle) and only on one side (right) for a stand alone VFAT2 chip.

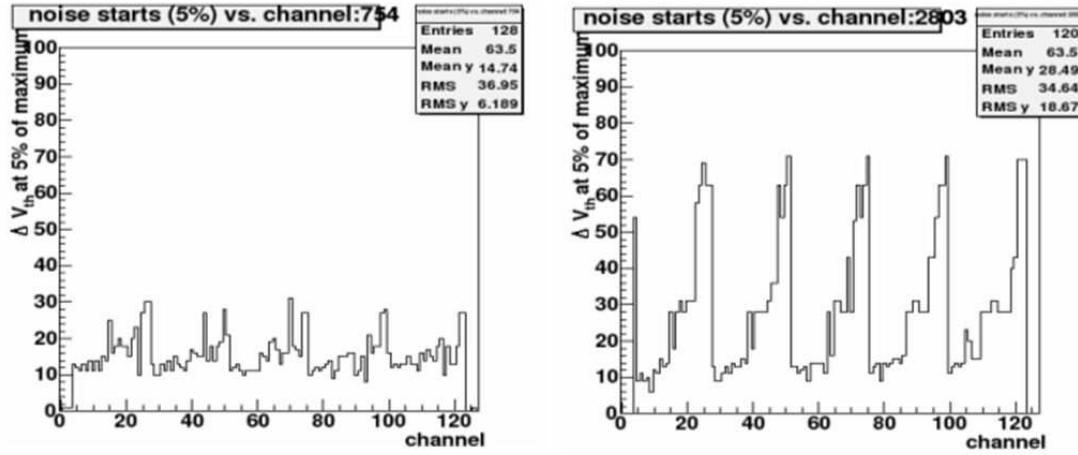


Figure 4.26: Test of the Digital and Analog Ground connections on the hybrid: "Threshold Scan Noise Start (5%) measurement for a VFAT2 chip mounted on the Triple GEM with (left) and without (right) the connection on the hybrid between the digital and the analog ground.

A test has been performed also inserting completely the hybrid into a grounded metal box. It didn't improve the noise level and therefore the previous solution has been chosen for the final installation.

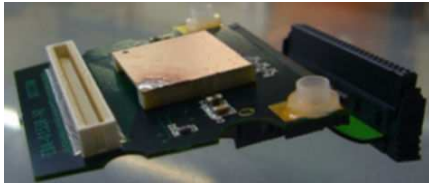


Figure 4.27: Metal Cover on the VFAT2 Chip: the metal cover is glued on the hybrid and referred to the analog and digital ground of the chip.

Low Voltage Power Supply

Measurements done with different power supply configurations will be shown in this section. We had the possibility to supply the analog and digital power directly to the VFAT2s or through voltage regulators mounted on the transition board⁵. In both cases the two power lines were in common to all the chips. The main power has been moreover provided by dedicated channels (analog, digital, beam board) or through a single one. The best solution, even if the difference was not so big (less than 10 – 20%), was the configuration with independent floating channels, with the ground reference defined on the detector. In fig. 4.29 the sigma obtained with a calibration pulse scan on few channels is shown. This configuration

⁵G. Antchev:EDA-01440, EDMS-CERN. The transition Board was used at the beginning for controlling and reading the VFAT2 chips.

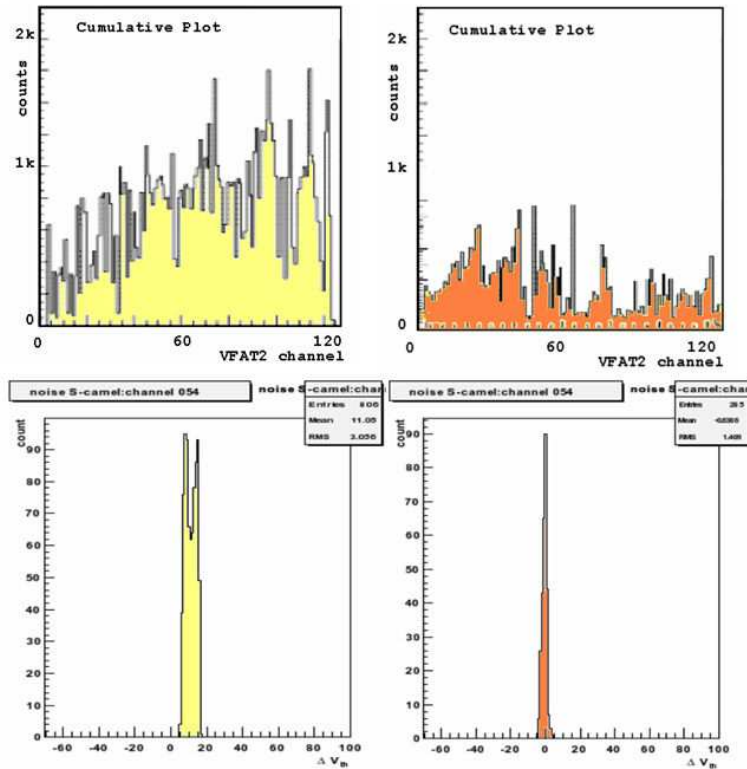


Figure 4.28: Test of the Metal Cover on the VFAT2 Chip: this cover, useful to protect the chip against unwanted damages, has reduced the noise level by roughly a factor two. In the plots are shown the results of a Threshold Scan (Cumulative and single channel counts) performed on the same hybrid without (left) and with (right) the cover. The hybrid was plugged on a Pads sector as can be seen from the repeated structure in the noise plots.

is the one used in the final setup, where the analog, the digital powers of the VFATs are provided by different floating channels of a Maraton (WIENER) power supply. A direct measurement of the analog

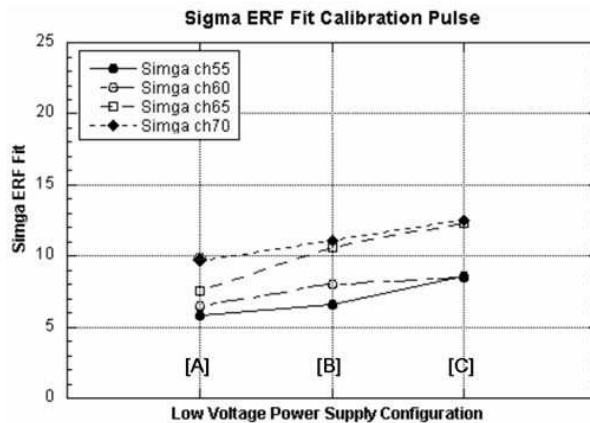


Figure 4.29: Low Voltage Power Supply Configuration: [A] CAEN FLOATING POWER SUPPLY (+5V, +2.5V, +2.5V), [B] CAEN FLOATING POWER SUPPLY (only 5V, 2.5V supplied from Regulators on TB), [C] LOCAL LOW VOLTAGE PS (only 5V , 2.5V supplied from Regulators on TB) GND connected to Table GND

and digital power lines has been performed to understand if they could be involved in the dependence of the noise level of one chip on the status of the other (actually on their digital activity). It was found that in both, analog and digital, the clock and the digital activity are reflected on the power lines. The effects was about the 5% for the trigger bit (largest effect) as it is shown in fig. 4.30. On the digital side this result should not affect the behavior of the chip, while on the analog one it was not clear and we decided to investigate more deeply this point. Before the power lines were common to all the chips. We decided to use independent low dropout voltage regulators for each VFAT2. In this way, if the noise coupling between chips is related to the common power lines, it should be eliminated or highly reduced. Decoupling boards (see fig. 4.31 have been inserted between each VFAT2 and the main power so that each chip has its own regulator (actually two, one for analog and one for digital). All the lines entering this board are left untouched except the power lines that are truncated and replaced by the on-board

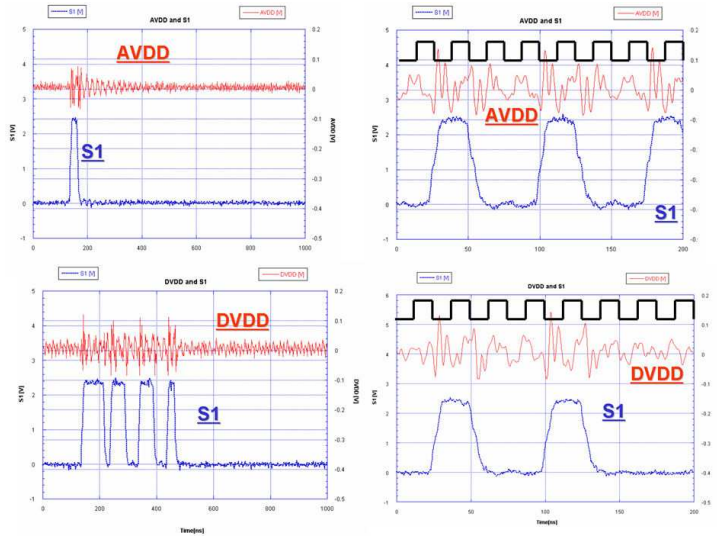


Figure 4.30: Analog and Digital power supply measurement: the analog (top) and digital (bottom) power lines in during high states of the trigger bits (that reflects the digital activity of the chip). A signal with the clock periodicity is superimposed. Clock's effects are visible even if less effective than the trigger bit (S1) ones.

regulators output. No improvement has been observed. This underlines that the cross-talk between the chips propagates only through the coupling of the channels in the readout plane of the detector.

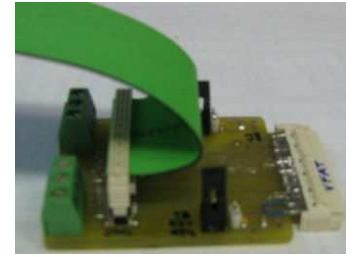
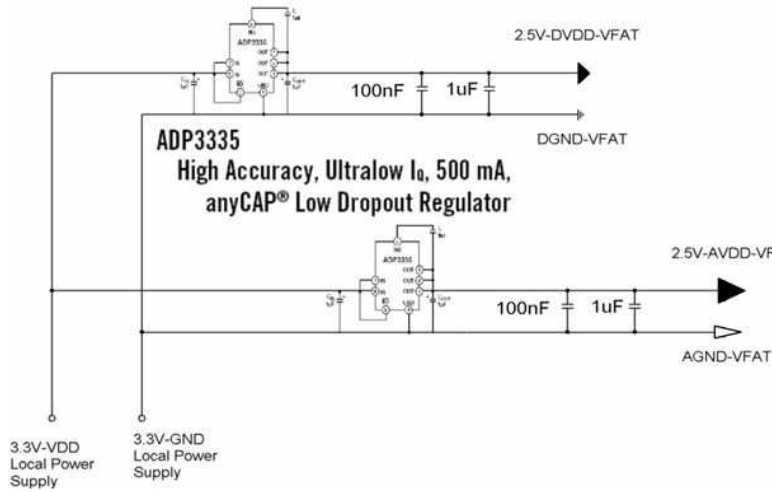


Figure 4.31: Power Supply Decoupling Board: The board (top), with two low voltage regulators on it, is inserted between each VFAT2 and the test beam board. The schematics related to the regulators are shown in the left.

Shielding

The effects of the detector shielding on the noise level have been investigated. In fig. 4.32 the tested shields are shown with one example of their effects on the Threshold Scan "Noise S-Camel Curves".

A very noisy starting condition has been chosen to magnify the effects of each shield. All the shields have a clear effect on the noise level. Among them, the shield on the strips fan out and on the bottom are the most useful. This is quite reasonable. The bottom foil is expected to be more useful than the top one for example, because on the top side the three GEMs foils and the drift one of the detector are shields themselves. The foil placed on the strips fan out instead is more than an e.m. shield. Its role is related also to the spurious capacitance added between the strips fan out and ground. Without this foil, the dominant capacity in this area is the inter-strip one. This could have effects on the cross talk between them. The coupling with ground reduces these effects. It will affect obviously also the signal, but the improvement on the noise was more important (the intensity of the signal that can be provided by the Triple GEM is enough for the VFAT2 readout chip).

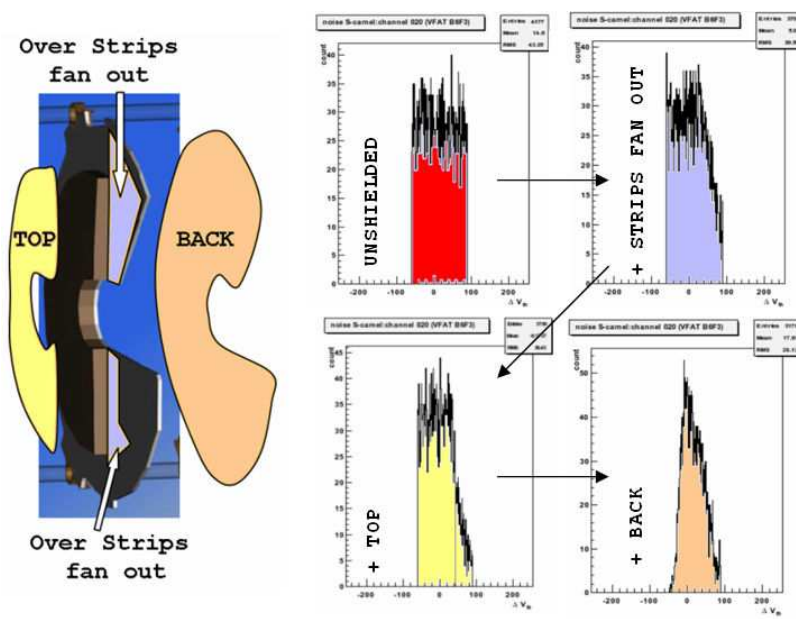


Figure 4.32: Detector Shielding: The tested shields has been placed on the strips fan out (connection between the strips and the VFAT2 connector), on the top and on the back of the detector (left drawing). In the right plots threshold scan measurements for single channel are shown. The scan has been done, adding in the shown order the various shields. Only the noise curve of the last plot is fully inside the threshold range scanned. In the unshielded one for example we found a plateau without any falling edge.

A measurement made with the calibration pulse (see fig. 4.33) shows that the level of noise (σ of the S-Curve) with a shielded detector can be reduced by more than a factor of two.

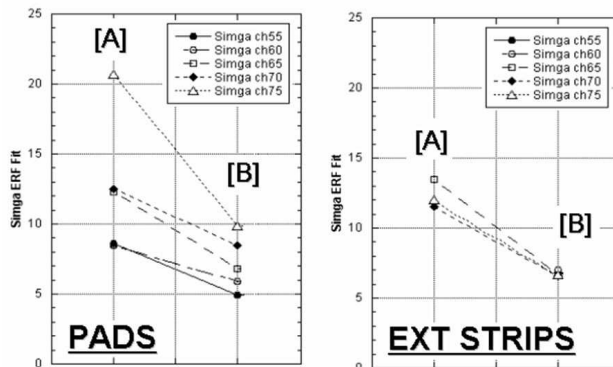


Figure 4.33: Detector Shielding tested by the Calibration Pulse Scan: One example of the noise level reduction obtained adding to the unshielded detector [A] the shields shown in fig. ?? [B].

In fig. 4.34 a more accurate threshold scan measurement to evaluate better the effects of an aluminum foil on the top of the detector is shown. Even if the improvement achieved with this shield is smaller than for the other foils, it clearly reduces the noise level and it has been added therefore also in the final configuration.

The same behavior has been found from our Helsinki colleagues, who studied the noise improvement due to the various shields looking directly with the oscilloscope at the noise on 128 strips grouped together and on a group of 120 pads (see fig. 4.35).

It has to be stressed that particular careful has to be used doing the shielding of the detector. In fig. 4.36 it is shown one example of different improvements obtained with two shields that cover exactly the same parts of the Triple GEM. The second one has been done taking much care of the connections to ground of the shields (multiple and with a good contact) and making uniform and smaller (as much as possible) the separation between the metal foil and the strips fan out.

In order to guarantee a good shielding with uniform improvement, our colleagues of Helsinki provided aluminized foil with proper shape and multiple points for connecting them to ground. In fig. 4.37 details are shown. The thickness of the metallization was about $100\mu m$ in order to have a small amount of material in the sensitive area of the detectors.

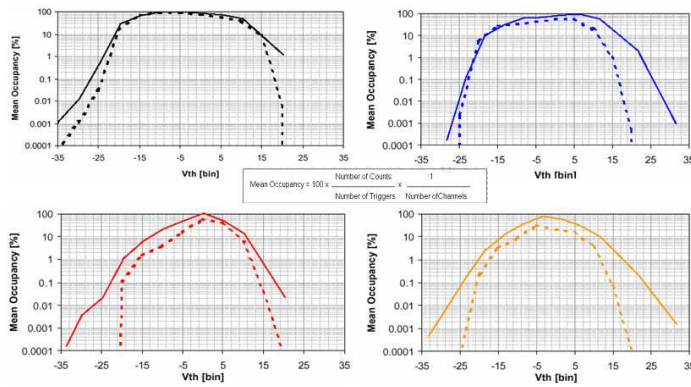


Figure 4.34: Threshold Scan measurements (40k triggers per threshold to have more accuracy) on four different chips. In each plot is shown the comparison between the same detector without (continuous lines) and with the metal foil on the top. All the other shields were installed. The mean occupancy in the plot is defined as the total number of counts for the chip per threshold, divided by the number of triggers and channels.

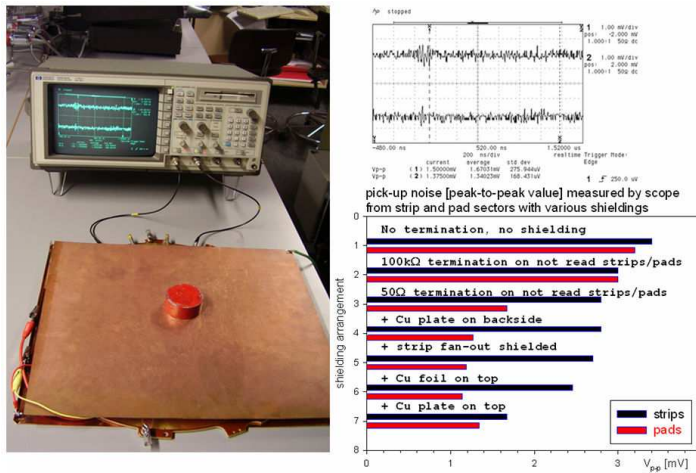


Figure 4.35: Detector Shielding checked by a direct measurement. Noise level on a group of 128 strips and one of 120 pads, read directly with the oscilloscope. The test of the shields has been done plugging a 50Ω termination on the connectors of the readout plane for the VFAT2. Helsinki Group Measurement.

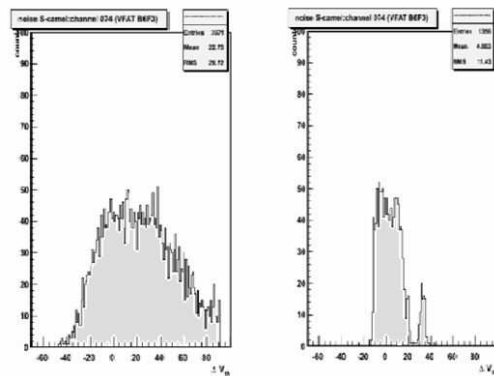


Figure 4.36: Noise reduction effectiveness for the same detector and the same shielded zones having done in the second case (right plot) a more accurate grounding (better and multiple contacts) of the foils and fixing as much adherent as possible the foils on the strips fan out.

Cross Talk and Noise Injection

As previously mentioned, cross-talk between channels have been found when the VFAT2 chips are mounted on the Triple GEM. It has been observed a dependence of the level of noise of one VFAT2 on the status (digital activity) of the others chips mounted on the same detector. The capacitance between strips and pads in our readout plane is the way how this noise propagate inside the detector. The effects of this cross-talk are more drastic than the normal coupling between channels that could affects the cluster size in case of the passage of particles or that could produces randomly localized clusters of strips and pads

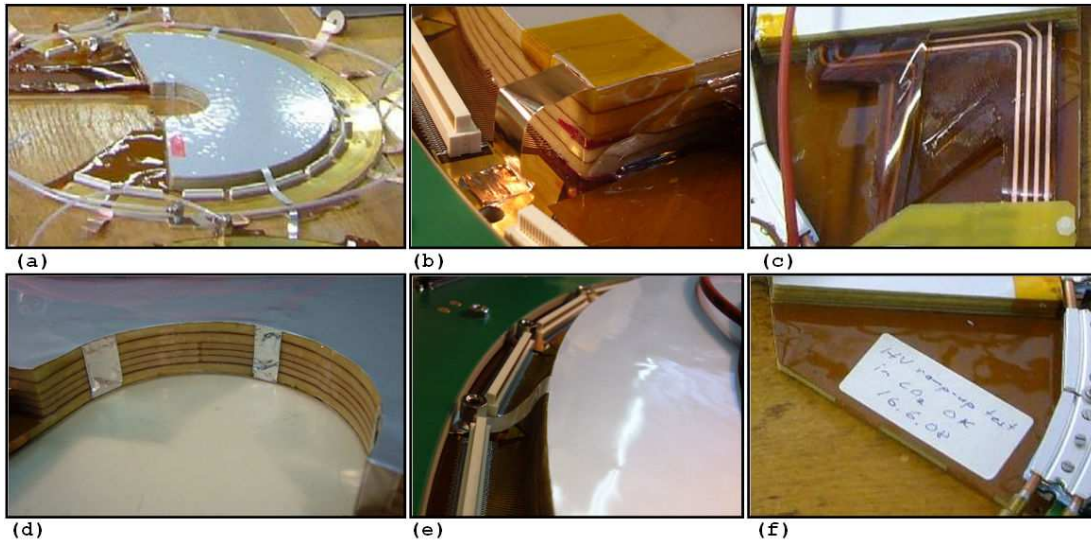


Figure 4.37: Detector Shielding: In the picture (a) a Triple GEM detector is shown with its shields. In (b), (d) and (e) the multiple connections to grounding the foils are shown. In picture (c) and (f) the metal foil placed under the strips fan out is shown. This part has been done using kapton foil with copper on one side.

induced by the noise.

A description of the behavior of the system can be done using the concept of positive feedback. When one chip starts its digital activity for the noise that it has on its channels, it starts to inject additional noise in the readout plane. Other chips, coupled to the first one via the inter-channel (strips-pads) capacitance, start to feel this injected noise, increasing it with additional noise produced by their digital activity. This link with the digital activity is only one possible hypothesis.

There were no possibility to eliminate the noise source or to reduce the pad-strips capacitive coupling. It has been possible however to reduce of the level of noise at which this positive feedback starts (i.e. the threshold level that can be used during normal runs has been reduced, that is actually what we need). Being the noise level dependent on the status of the chips installed, all the measurements in the next pages are related to the particular configuration and testing procedure used. During normal data taking, all the situation, that will be presented, will never happen because thresholds that inhibits this kind of positive feedback will be chosen for each VFAT2. These effects will nevertheless affect the system because, for particularly noisy channels, higher threshold are needed to maintain the system in a quite state (i.e. loose in efficiency of the area covered by those channels).

An extreme cross-talk case will be shown first. The cumulative plots recorded by one chip plugged on the external strips (the VFAT2 under scan) and by one chip mounted on a sector of pads partially underlying the strips are shown in fig. 4.38. This chip on pads has not been scanned and its threshold has been left fixed. In this case, the standard structure of the noise pattern for pads (that follows their area) is absent, the pads under the internal strips (not scanned in this test) are silent and all the pads under the external strips (the scanned ones) had exactly the same counts of the external strips, with a nearly flat shape. No noise has been found on the same chip at the same threshold if the threshold in the external strips chip is such that its digital activity is suppressed (i.e. no noise hits for the strips).

In the previous example it was very easy to identify the cross talk between pads and strips. The extremely flat structures in the cumulative plots are normally related to cross talk. Another example is reported in fig. 4.39 where the strips are analyzed during a threshold scan. The counts of two different channels of the chips are practically the same, showing identical structures. Even if the strips are more similar and no evident structures in term of noise has to be found as for pads, the found curves are too much close one to the other, underlining that the channels are practically moving together. The previous examples have been obtained doing a threshold scan on the strips. It is useful to see what happens if we perform a threshold scan on pads with the strips at a fixed threshold. In this way it is possible to

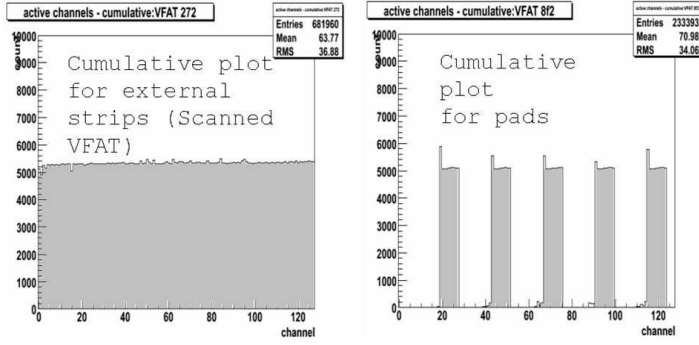


Figure 4.38: Cross talk between pads and strips: Cumulative plots for two chips mounted respectively on the external strips (left) and on pads (right). These data have been collected during a threshold scan on the strips and with the chip on pads left at a fixed threshold, where there aren't counts if the strips are not scanned and left at a threshold over their noise level.

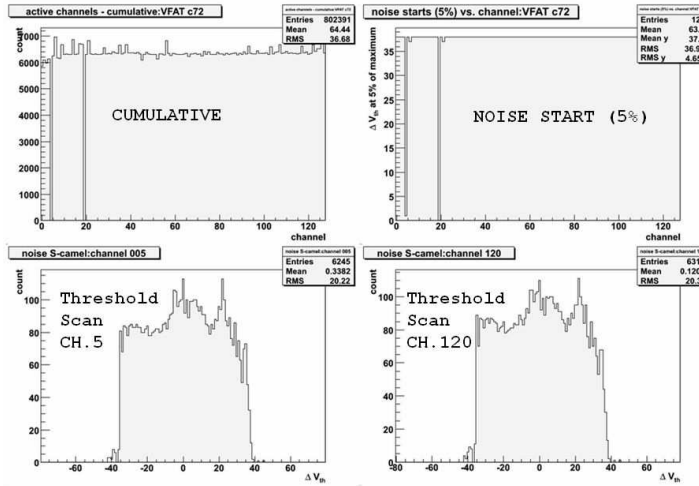


Figure 4.39: Cross talk between strips: Threshold Scan plots for the external strips. Top: Cumulative and Noise Start (5%) plots with a flat structure on both. Bottom: Counts relative to two different channels (strips with different length, underlying pads) measured during the scan.

monitor the noise level of pads from their threshold scan. The results are shown in fig. 4.40. If the strips are moved from a threshold that is under the noise level to a threshold over it (i.e. no counts on the cumulative plots), the results obtained on the pads chip are improved quite a lot and the noise level is very close to the theoretical one obtained from the readout electrode capacitance.

In all the tests that we have done, we found that the more problematic chips involved in this cross talk are the strips. Normally the worsening of the system noise that is due to pads is negligible. This could be related to their different inter-channel coupling. Among the two type of strips that we have in our readout plane⁶, the external one are the more critical. Nevertheless also the internal ones show the same behavior, but with smaller effects. In fig. 4.41 is shown how the digital activity on external strips turn on only the channels associated to pads that underlie these strips (blue area on the plot) and not the ones under the internal (yellow area). The opposite happens if we forced the digital activity only on the internal strips. For the latter, a zoom in the plot has been done to make the counts visible.

Full Equipped Detector

The previous studies of noise has driven us to the final configuration of the detectors. In this section some results achieved on the detectors, before mounting them in the final telescope, will be shown. The following measurement have been done with all the VFAT2 mounted on the detector (i.e. seventeen) in

⁶Each readout plane, as described in the T2 telescope section, has 512 strips, with 256 per side. On each side the 128 strips with smaller radius are defined as internal and the other as external. The main differences between external and internal is the strips length and the area (i.e. the coupling) of the underlying pads. Considering also the strips fan out, the different length is partially compensated because it is bigger for the internal strips that are shorter in the sensitive area of the detector. The two regions however (the fan-out and the sensitive area) have different width and thickness of the strips.

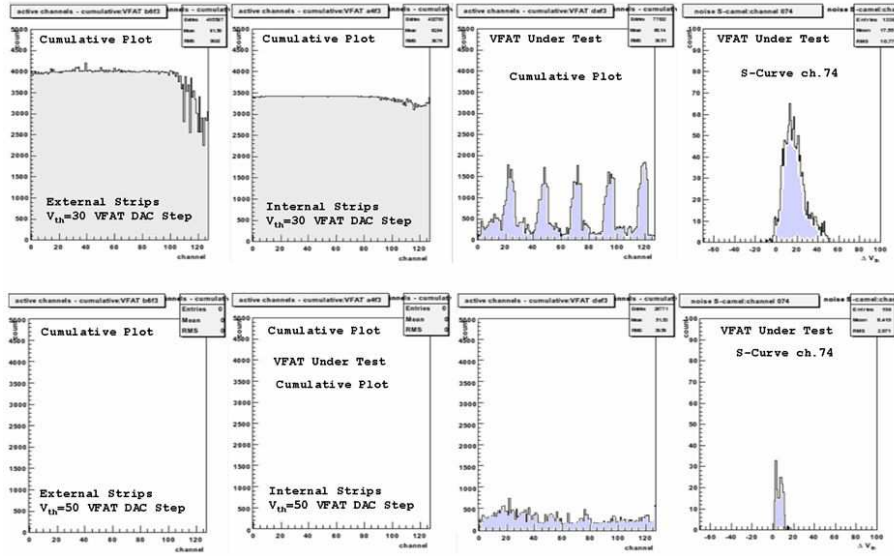


Figure 4.40: Threshold scan measurement performed on pads, using for the strips two fixed threshold: (top) Threshold under the noise level, (bottom) Threshold over the noise level. The results obtained for pads (plots on the right) is largely affected by strips and where the strips are silent (no counts in their cumulative plot) the noise on pads is extremely reduced.

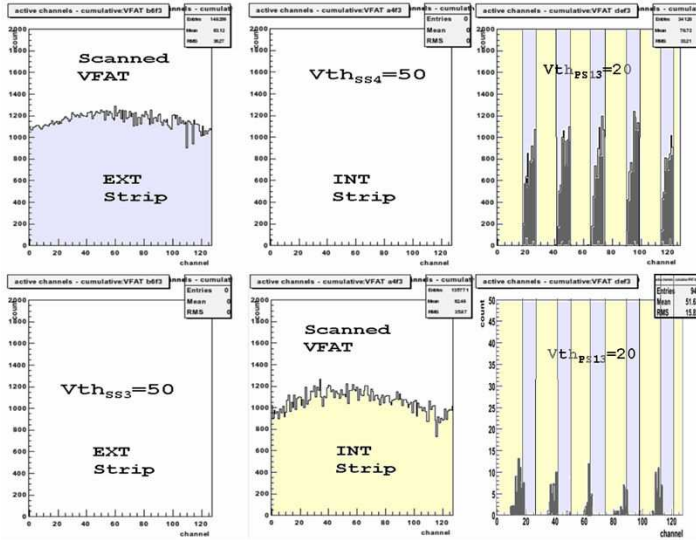


Figure 4.41: Cumulative plots of a threshold scan on pads (right-top, right-bottom) with different digital activity on the strips of the same detector (top: digital activity forced on external strips, bottom: digital activity forced on the internal ones). The status of the strips is shown through their cumulative plots (left for external and middle for internal).

run mode and making a threshold scan on fourteen of them. In fig. 4.42 is shown a summary plot for ten detectors belonging to one T2 quarter. It shows also a picture of one detector (ready to be mounted, with shields and cooling on it) and the noise curve obtained from the threshold scan, from which the quantity of interest (i.e. the negative and positive starting point of the noise curve) have been extracted. The measured quantity are related to the strips (external and internal). The pads show a smaller level of noise.

In fig. 4.43 the erf fit σ of the S-Curve obtained with a Calibration Pulse in few strips and pads is shown. In the pads case, the dependence on the capacitance is evident. The values obtained have been acceptable for all the detectors tested (i.e. less than 8 VFAT DAC step for all the tested channels).

The data have been obtained using a four clock stretching of the VFAT2 monostable (MSPL [32])

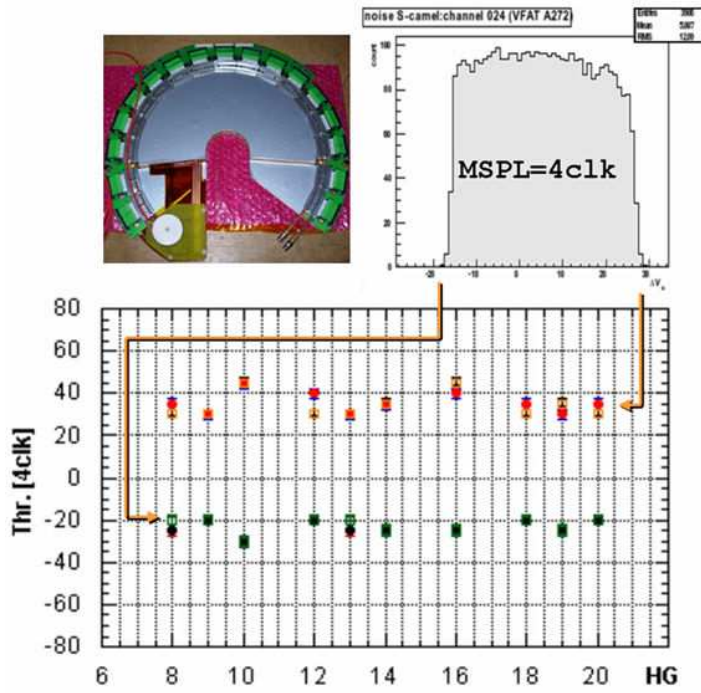


Figure 4.42: Noise measurement of a fully equipped detector. Top-Left: a Triple GEM detector before the measurement and ready to be mounted on the T2 Quarter. Top-Right: Threshold Scan measurement for a single channel from which the negative and positive starting points of the noise have been extracted. Bottom: Summary plot of the ten detectors tested. For each detector (labelled with the HG id) eight points are shown. Four negative values and four positive value for the two external and the two internal strips respectively. All these measurements have been obtained using a four clock stretching of the VFAT2 monostable [32].

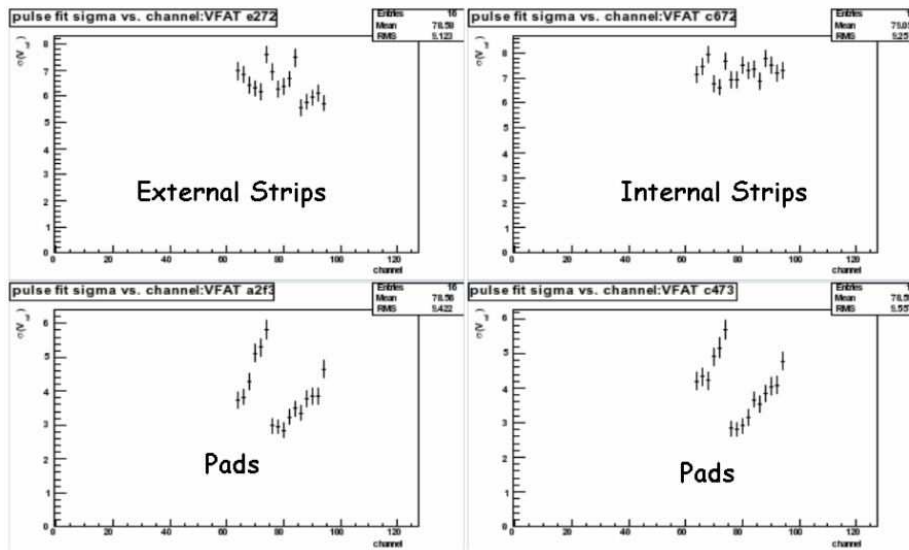


Figure 4.43: Noise measurement of a fully equipped detector. Calibration Pulse Scan on strips and pads for one detector ready to be installed on one T2 Quarter.

monostable pulse length⁷) that is a reasonable working point in normal operation. We investigated also other stretching lengths to analyze their effects on the noise level. In fig. 4.44 the plots of fig. 4.42 is shown together with a MSPL of one, two and eight clock cycles. The GEM signal is negative and we are therefore more interested on that side. Except the case of one clock cycle (that will never be used with our detector because it would degrade too much the efficiency), the negative starting point is below 40 VFAT2 DAC Step. This is enough to have the necessary efficiency.

⁷The stretching of the monostable affects the information that is stored on the memory cells of the vfats. If the MSPL is equal to two clock cycles it means that the channel status that we retrieve from the VFAT2 memory is relative to what happened in a time interval of two clocks. The possibility to stretch the monostable pulse length is necessary for signals produced by a Triple GEM detector. With a MSPL of one clock cycle, hits induced by one particle will be stored at different latency and it will be therefore impossible to retrieve all of them at the same time (i.e. loose of efficiency).

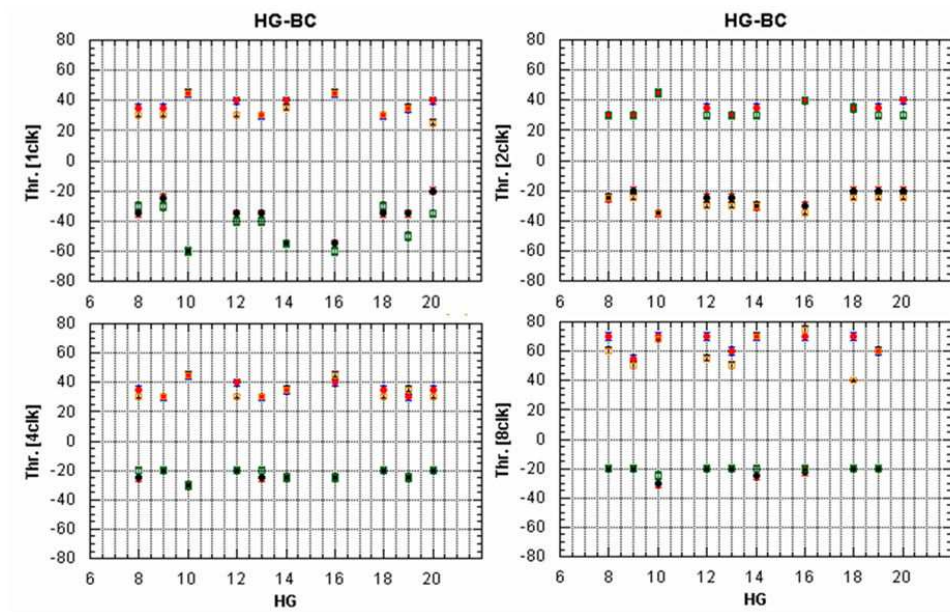


Figure 4.44: Summary plot (as the one in fig. 4.42) of the ten detectors using a stretching of the monostable pulse length of one, two, four and eight clock cycles.

In the next chapter, measurements with particles will be shown. Starting from the results obtained during a test beam, cosmic ray data taking and pp collision at the LHC will be presented.

Chapter 5

Test with particles

In the following section a collection of measurements performed with particles will be shown. The first data, presented in sec. 5.1, are referred to a test beam on the extraction lines of the SPS. They have been used to check, for the first time, the performance of the triple GEM detector with the VFAT2 chip. These tests have been done on partially equipped detector, because parts of the final electronics were still missing. The results, even if preliminary, and the monitoring of the status of the detectors in normal operation for long time periods, have driven some updates that has been done before the final installation.

No other beams have been available at the SPS before the installation of the quarters in the experimental area at the Interaction Point 5. Each quarter has been therefore tested under cosmic rays in the TOTEM area in Preveessin. The aim of the test was the minimum check of functionality of the quarter, because the time schedule (T2 completely installed by May 2009) has not allowed us to perform long runs and optimizations.

Once the quarters have been installed in the IP5 area, test with cosmic ray were not anymore possible, because of the longitudinal layout of the TOTEM experimental apparatus. We had to wait for the pp collision of the 2009 LHC runs. Preliminary results will be shown in sec. 5.3

5.1 Test Beam Data

The most important measurements of this test beam in the H8 zone of the North area of the SPS extraction line (fig. 5.1) were related to the timing properties of the Triple GEM - VFAT2 integrated system and to the detector efficiency. In addition we studied the cluster size and the spatial resolution.

5.1.1 Timing Measurements

During the first test beam we measured the timing performance of the TOTEM Triple GEM. The gas mixture used was Ar/CO_2 (70/30). We were interested in the measurement of the spread in the rising time of the signal due to jitter and time walk effects. This information is necessary to understand the degree of coincidence between all the channels involved when a particle travels through the T2 telescope. Channels of different detectors will have different time responses to the passage of the particle according to this spread. The test has been done measuring with a TDC (time to digital converter) module the delay between a coincidence of two scintillators (used to define the passage of a particle) and the trigger S-bit generated by the VFAT2. The time response of the scintillator is much faster than the Triple GEM one and it doesn't affect therefore the measurement.

When the charge released by the passage of a particle (and multiplied by the GEM foils) induces a signal that overcomes the threshold in one channel of the VFAT2 chip, the internal monostable (fig. 5.2) will change state. The outputs of all the monostables integrated in each VFAT2 chip are stored in an

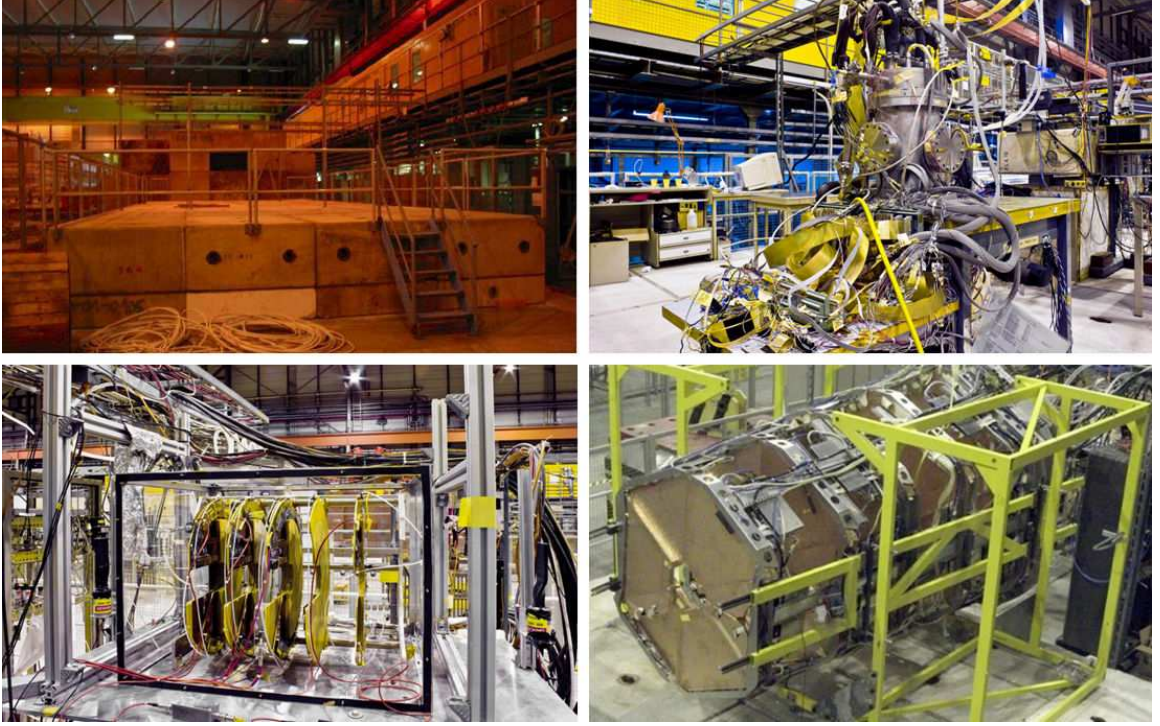


Figure 5.1: H8 Totem test beam zone in the North area of the SPS extraction lines. Before the detectors installation (top-left). RP (top-right), T2 (bottom-left) and T1 (bottom-right).

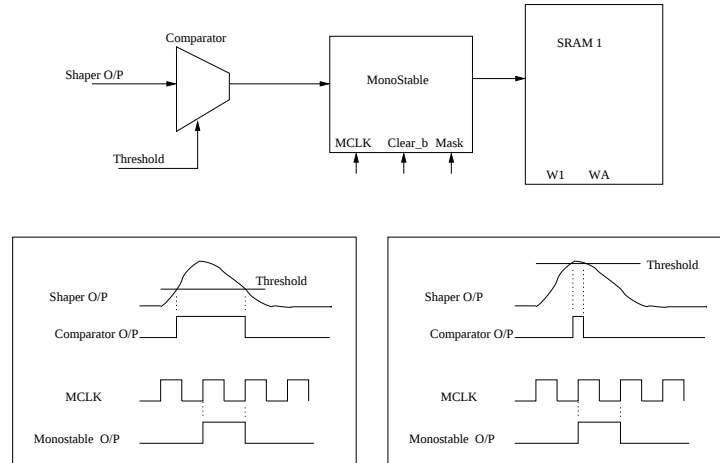


Figure 5.2: VFAT2 monostable block diagram: *The comparator is an asynchronous comparator without hysteresis. On passing a programmable threshold the comparator output goes high and returns low again when descending back through the threshold. For very large signals the comparator output may remain high for more than one clock cycle. Also if the signal barely passes threshold the comparator output may go high for less than one clock period. Both cases should trigger the monostable which should provide a pulse of one clock period only (in the default mode). It is this pulse which is sampled by SRAM1 [32].*

internal memory and are transmitted to the outside world if a *LV1* trigger is sent to the chips. The information on the monostables of pads are moreover continuously processed and sent to the external to provide signals for the generation of the *LV1* trigger (see sec. 2.1 and 2.3). In particular 8 trigger-bits are extracted from each VFAT2 on the Trigger Sector Lines as schematized on fig.5.3. For the triple GEM, each bit is generated with a fast-OR combination of the pads that belong to one of the 8 *super-pad* in

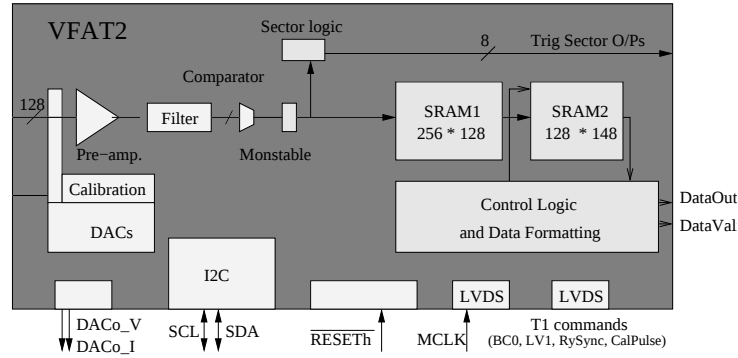


Figure 5.3: VFAT2 Block Diagram: the outputs of the pads monostable are stored in the SRAM1 memory and contemporarily processed in the *Sector logic* block that will generate 8 trigger sector outputs.

which the pads sectors read by one VFAT2 has been divided¹. Actually these combinations of pads can be partially modified. For the measurement performed with the TDC we have used one of the *Trigger Sector Lines* on which the fast-OR of all the 120 pads connected to the VFAT2 has been sent. This signal has been sent to the TDC unit and it has been measured the time delay with respect to the coincidence of two scintillators, placed before and after the detectors as shown in fig.5.4.

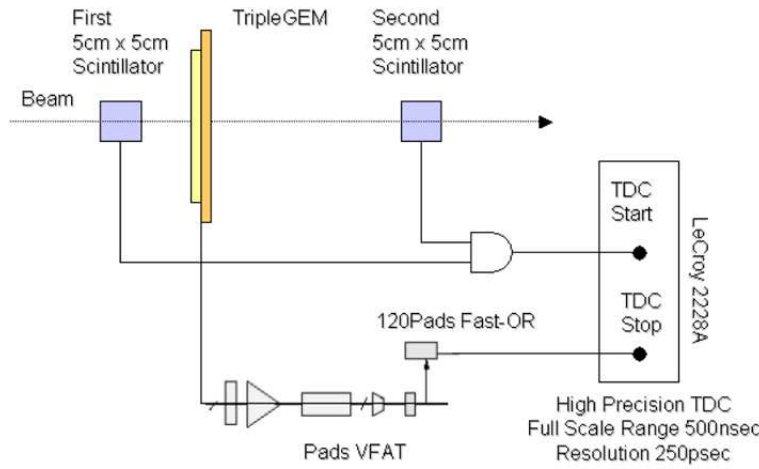


Figure 5.4: TDC acquisition system. For this measurement it has been used a LeCroy 2228A TDC unit. The full scale range has been set to 500ns with a resolution of 250ps. The coincidence of the discriminated scintillator signals was sent to the TDC start input while the S-bit (trigger output from the VFAT) to the TDC stop. In this measurement it has to be taken into account that inside the VFAT there is a synchronization with the external clock between the comparator and the monostable output.

The measurement of the time distribution of the rising and falling time of the signal can be done choosing the right setting of the VFAT2. The triple GEM signal is negative and therefore a negative threshold has to be chosen. For the rising time measurement it is necessary that the monostable output goes high when a signal with negative slope crosses the threshold. The falling time instead requires that the output of the monostable goes high when a signal with positive slope cross the threshold. This is done selecting respectively a *Negative* and *Positive Monostable Input Polarity* as explained in fig.5.5.

The Negative Monostable Input Polarity measurements. The delays measured between the rising front of the triple GEM signal and the scintillators coincidence are shown in fig. 5.6. The spread reflects the timing properties of the signal induced on the triple GEM readout plane and the internal synchronization of the VFAT2. Each plot is relative to a different high voltage applied to the triple GEM with in series an H.V. filter of 95.9kΩ. A larger number of counts, a reduction of the spread and a displacement toward

¹Each VFAT2 read one sector of 120 pads that is divided in 8 *super-pad* with 15 pads each as shown in fig. 2.4.

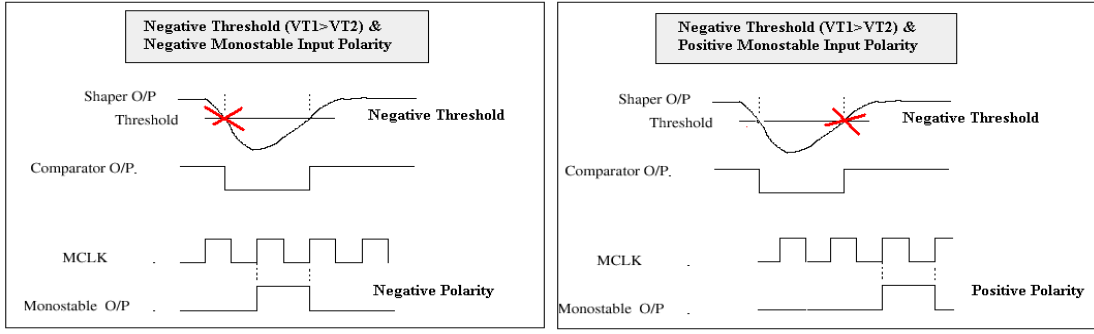


Figure 5.5: *Negative and Positive Monostable Input Polarity*. Left: with a *Negative Polarity* a signal that cross the threshold with a negative slope will be triggered. For the negative signal of the triple GEM this means that the monostable becomes high in the rising front of the signal. Right: with a *Positive Polarity* a signal that cross the threshold with a positive slope will be triggered and therefore the monostable becomes high in the falling front of the signal.

smaller delays are observed increasing the voltage. These effects are due to the higher amplification of the signal and to the increased drift velocity of the electrons under the larger internal electric field. The

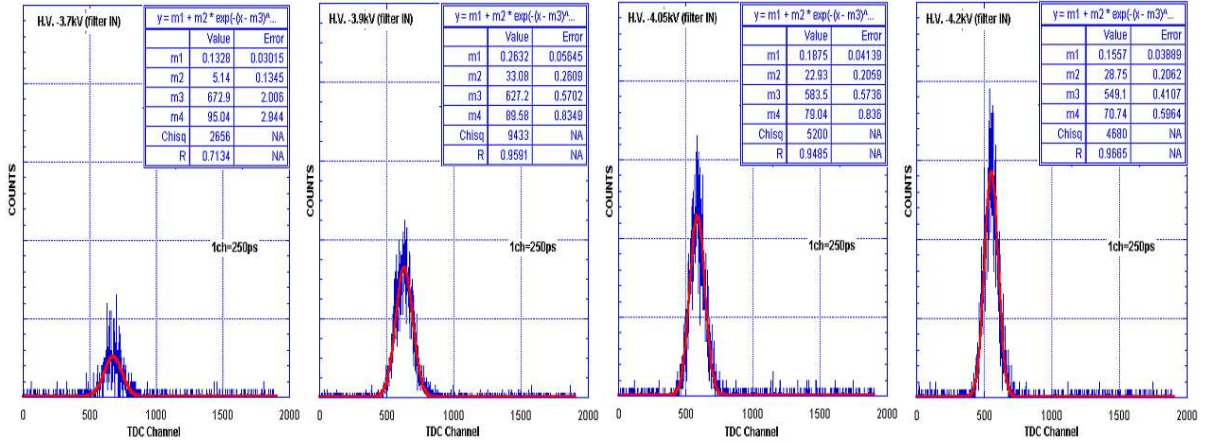


Figure 5.6: Triple GEM rising time distribution: delay distribution between the rising front of the triple GEM signal (defined by the fast-OR of 120 pads extracted on the *Trigger Sector Line* bit) and the scintillators coincidence. Each plot represent the output at different high voltage applied to the triple GEM with in series an H.V. filter of 95.9k Ω : in order from $-3.7kV$, $-3.9kV$, $-4.05kV$ and $-4.2kV$.

dependence on the applied high voltage of the mean and the standard deviation of these distributions is shown in fig. 5.7. The value assumed by the mean is not of particular interest. The standard deviation instead is very important because it reflects the different times of response of the VFAT to different signals produced by the triple GEM. This is the quantity that is of our interest. As it can be seen in the plot, the standard deviation has not reached the plateau at the maximum high voltages tested (plateau that is due to the VFAT2 internal synchronization). This means that the time of response can be improved working on the detector. The best condition is achieved when the standard deviation is short enough to have all the signals triggered in one clock cycle (25ns). This is not possible with a Triple GEM like the TOTEM one. It is possible however to reduce it if necessary, changing the used gas mixture and the fields. The time distribution measured at nearly $-4.1kV$ has a standard deviation of 0.7 clock cycles. This means that at least three clock cycles are needed to include all the signals (fig. 5.8). If the VFAT2 stores the status of its input channel during one clock cycle in one memory cell, the same event will be stored in at least three different memory cells with this spread. The latency, used to retrieve from the memory the status of the channels when an external trigger comes, cannot be univocally fixed. Moreover the VFAT2 itself has

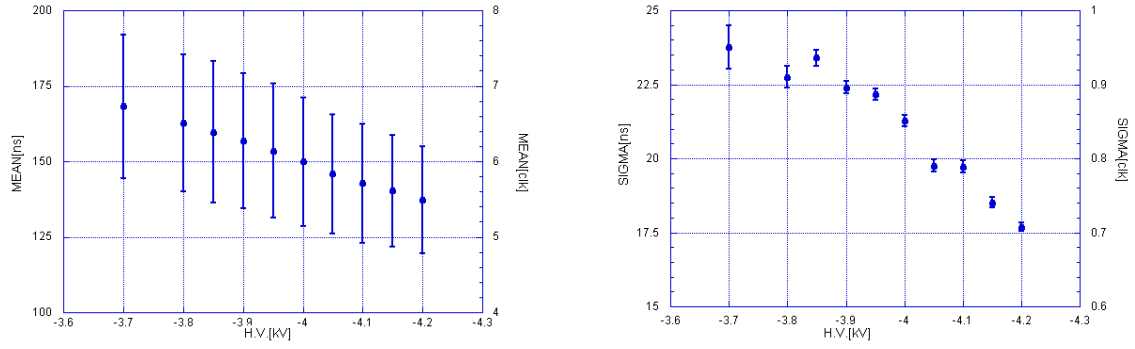


Figure 5.7: Rising time. Mean (left) and standard deviation (right) of the delay distributions between the rising front of the triple GEM signal and the scintillators coincidence versus the high voltage applied to the triple GEM with in series an H.V. filter of $95.9k\Omega$.

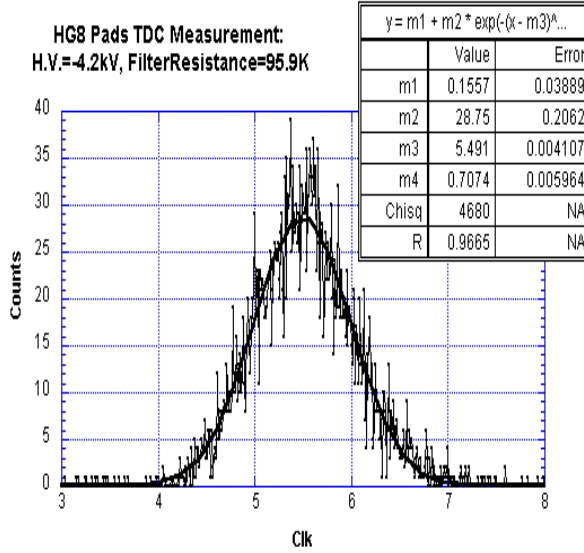


Figure 5.8: TDC measurement at $-4.2kV$. Considering the $95.9k\Omega$ high voltage filter this means that $\approx -4.1kV$ are applied on the triple GEM. The standard deviation is represented by the parameter $m4$ in the table.

to sent signals that will be used for the *LV1* trigger generation that would be obviously affected by the same temporal spread. To avoid this problems the VFAT2 can be programmed to store in each memory cell the information associated to more than one clock cycle. This can be done changing the length of the monostable output that can be stretched² up to 8 clock cycles. (see fig. 2.24 in sec. 2.3). With a stretch of three clock cycles for instance, each memory cell and each pad signal used for triggering purposes will be referred to the previous three clock cycles³. In fig. 5.9 the normalized number of events that fall into a shifted interval of one, two three and four clock cycles has been computed for the distribution of fig. 5.8. This integration has been normalized to the total number of signals collected from the TDC module. This is a rude estimate of the percentage of events seen by the VFAT2 for a fixed latency. With a *Monostable Pulse Length MSPL* of one clock cycle the maximum number of events that can be found in one memory cell of the VFAT2 is nearly the 70% of the collected. Only with a stretching of three clock the 100% in a small latency interval is achieved. The VFAT2 latency can be varied in integer clock cycles. This means that in this case it is necessary to fix properly external delays to maximize the number of events that will be collected. With four clock cycles instead the 100% is found for more than one clock cycles and this means that the maximum is obtained just changing the latency inside the VFAT2.

²The pulse from the monostable can also be stretched over many clock periods. The length is programmable ... If the monostable pulse is stretched (to say X clock periods) it is possible that the output of the comparator returns below threshold and a second signal makes the comparator go back over threshold. In this case the monostable output pulse remains high until X clock periods after the last over threshold result from the comparator. [32]

³The *LV1* trigger is decoded in the clock. This fast command requires three clock cycles to be transmitted to the front end electronics. This means that the time between two consecutive *LV1* triggers has to be equal to or higher than $3clk$ cycles.

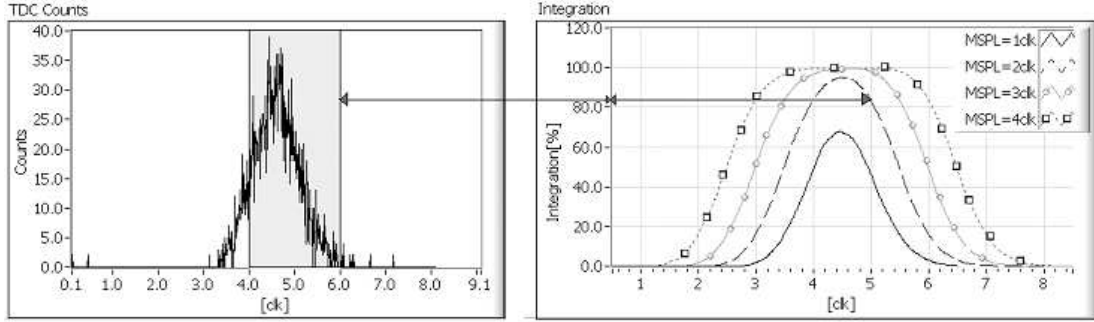


Figure 5.9: Percentage of the collected events stored in one memory cell using different *Monostable Pulse Length*: The right plot is obtained summing all the events of the left plot that fall inside a moving interval centered in the latency shown in the abscissa and with a length of one, two, three or four clock cycles. The time scales on the two plots are both shifted. In the second plot it has been chosen a shift that gives in abscissa the position of the center of the integration interval in the first plot. This plot shows how the efficiency measurement will be influenced by the *Monostable Pulse Length* for signals generated by the triple GEM. In the plots it is also shown as an example the maximum achievable efficiency (considering 100% for the detectors) that is nearly 85% with an integration interval (or a *Monostable Pulse Length*) of two clock cycles and a latency of five clock cycles⁵.

The Positive Monostable Input Polarity measurements. In fig. 5.10 the delays measured between the falling front of the triple GEM signal and the scintillators coincidence are shown. A double peak structure (fig. 5.6) has been found, whose meaning has to be investigated. The means and the standard deviations

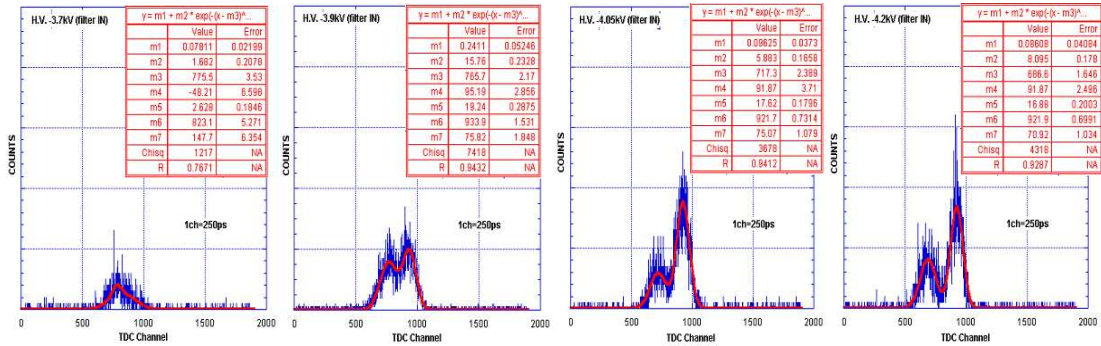


Figure 5.10: Triple GEM falling time distribution: delay distribution between the falling front of the triple GEM signal (defined by the fast-OR of 120 pads extracted on the *Trigger Sector Line* bit) and the scintillators coincidence. Each plot represents the output at different high voltage applied to the triple GEM with in series an H.V. filter of 95.9kΩ: in order from $-3.7kV$, $-3.9kV$, $-4.05kV$ and $-4.2kV$.

of the two gaussian fits of these curves are plotted in fig. 5.11.

The data reported in fig. 5.11 have been compared with fig. 5.7 and the differences on the means (related to the length of the signal) are shown in fig. 5.12. Considering the second peak (at higher delay) of the *Positive Monostable Input Polarity*, a length between three and four clock cycles is found, that, taking into account the spread for the VFAT2 internal synchronization, is compatible with the simulation of the triple GEM signal length. The first peak of the *Positive Monostable Input Polarity* is instead between one and two clock cycles after the rising front. The origin of the double structure on the plots of fig. 5.10 cannot be associated to the internal synchronization of the VFAT2 because otherwise a maximum difference of one clock cycle should be found.

The total number of counts of the VFAT2 trigger output line for *Negative* and *Positive Monostable Input Polarity* over the total number of counts of the scintillators coincidence is shown in fig. 5.13. This plot gives

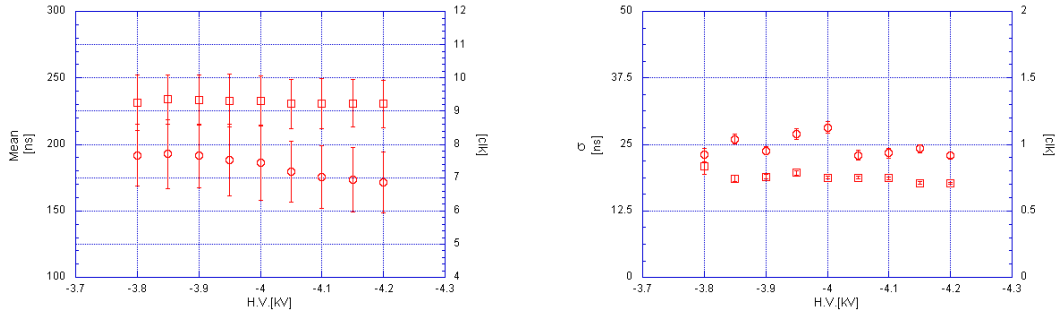


Figure 5.11: Falling time. Mean (left) and standard deviation (right) of the delay distributions between the falling front of the triple GEM signal and the scintillators coincidence versus the high voltage applied to the triple GEM with in series an H.V. filter of 95.9k Ω .

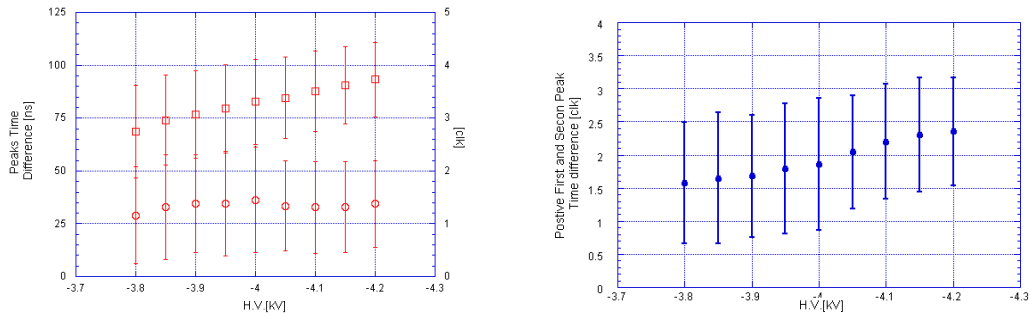


Figure 5.12: Left: Difference between the rising time shown in fig.5.6 and the falling time shown in fig.5.10. For the fact that the plots in fig.5.10 are fitted with a sum of two gaussian distribution, in this plot we have plotted the temporal difference with each peak. Right: Difference between the two line of the left plot.

us information about the trend of the detector efficiency. The saturation value cannot be 100% for the larger area of the scintillators with respect to the active GEM area connected to the VFAT2. Nevertheless at the higher voltage there is a beginning of a plateau, that means that we are close to the maximum efficiency achievable.

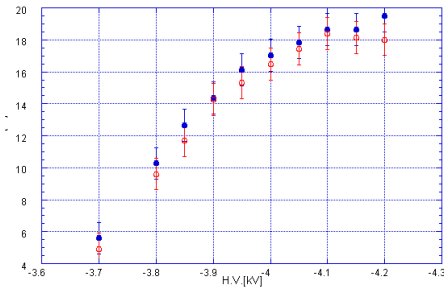


Figure 5.13: Comparison between TDC and Scintillators Coincidence Counts versus the high voltage applied to the Triple GEM. The counts of the VFAT2 trigger output line have been divided by the total number of counts of the scintillators coincidence and plotted in percentage.

After these measurements, we decided to increase the field in the induction zone of the Triple GEM detector. We passed from an electric field of nearly 2.5kV/cm to 3kV/cm in this zone. The gas mixture will be Ar/CO_2 (70/30), even if it will be foreseen the possibility of adding CF_4 .

5.1.2 Latency Scan

Additional information about timing can be obtained from latency scans. The Read cycle from the VFAT2 begins with a continuous write/read operation as soon as the memory SRAM2 contains data (fig. 5.3). Data are transferred in SRAM2 from the first memory SRAM1 on receiving a LV1A signal. The Write

State machine provides the *Write address* W_a for the data stored in the SRAM1 with a *write pointer counter* that increments every clock period. On receiving an LV1A the channel data corresponding to the trigger is identified by the *Read address* R_a that is obtained by subtracting the programmed latency number from the W_a . In the test performed on beam in the H8 area, the LV1A signal has been provided by the coincidence of two scintillators. The right latency has been located scanning different latencies and searching for the maximum number of coincidences between the LV1A signal and the VFAT2 readout from the triple GEM.

In fig. 5.14 latency scans performed on two detectors exposed to a beam, with and without the high voltage applied, are shown. The measurement on the chamber with the high voltage turned on defines the proper latency. The measurement on the detector with the high voltage turned off is equivalent to a noise measurement. For the tested detector it has been found that the probability that one channel is over-threshold even if it has not been hit is less than $\sim 1^\circ/\text{oo}$. Both the measurements have been done with a *Monostable Pulse Length* (MSPL) of 1clk and the level of the threshold has been fixed to 35 DAC step (one step $\approx 600e^-$).

The analysis of the first plot has to take into account the area of strips and pads covered by scintillators

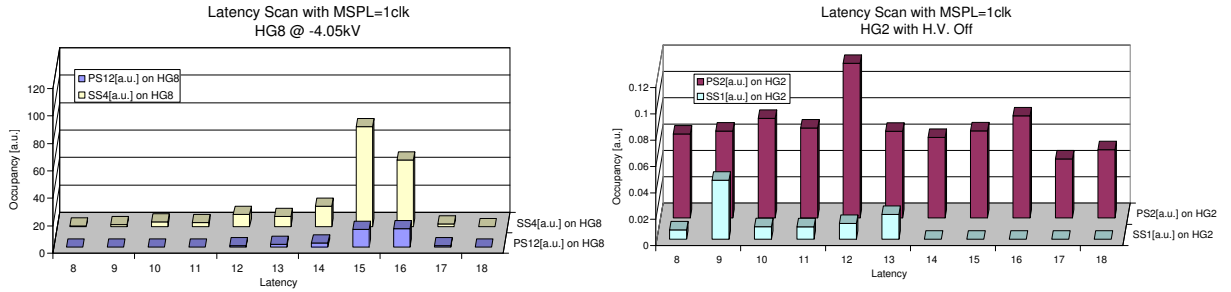


Figure 5.14: Latency Scan with beam for a *Monostable Pulse Length* (MSPL) of 1clk. Left: $H.V. = -4.05kV$ applied to the triple GEM HG0 with an H.V. filter with a series resistor of $95.9k\Omega$. Right: $H.V.$ off for the triple GEM HG2. The ratio between the ordinate scales of the two plots is 10^3 . The labels PS12 and PS2 mean Pad Sector 12 and 2 while SS4 and SS1 mean Strips Sector 4 and 1.

and the usual definition of occupancy ($\text{Occupancy[a.u.]} = 100 \cdot \frac{\text{Number of channels ON}}{\text{Number of Triggers}}$). For each event, two or three strips and one or two pads are involved. This means that it has to be expected an higher value of occupancy for the formers. Moreover if the induced signal is shared between more than one electrode, the ones with a lower collected charge will overcome the threshold level after the others (i.e. at smaller latency) for time walk effects. It will happen therefore that one event increases the occupancy level in more than one latency slot. Fig. 5.15 shows the comparison between the TDC measurement and a latency scan performed on the triple GEM HG8 at the same high voltage applied. The time distribution obtained

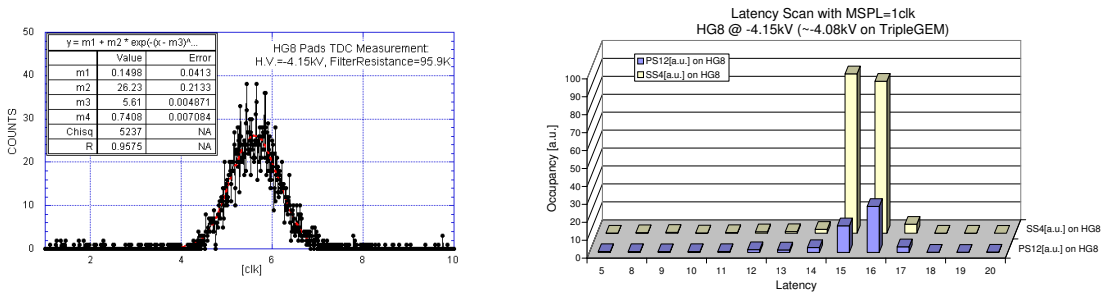


Figure 5.15: TDC and latency scan measurements performed on beam with an high voltage of $-4.15kV$ applied to the triple GEM with in series an H.V. filter of $95.9k\Omega$. Left: TDC measurement. Right: Latency scan for a *Monostable Pulse Length* (MSPL) of 1clk.

with the TDC measurement is symmetric and the standard deviation of the distribution is $\sigma \sim 0.7$ clock cycles. In the latency scans there is instead a longer tail for smaller latency. Smaller latencies correspond

to signals that happen later in time. We can explain this long tail with signal produced by pads and strips of a cluster that have less electrons collected. For smaller signals therefore the threshold will be crossed later than the strip and pads of the clusters that collect the main part of the electron produced inside the GEM. In the TDC measurement we don't have this effect because the TDC unit measures the first hit over threshold of the clusters.

The comparison between two different detectors with roughly the same high voltage is shown in fig. 5.16 while in fig. 5.17 the results have been obtained from two triple GEMs with different high voltages, $-4kV$ and $-3.6kV$ respectively.

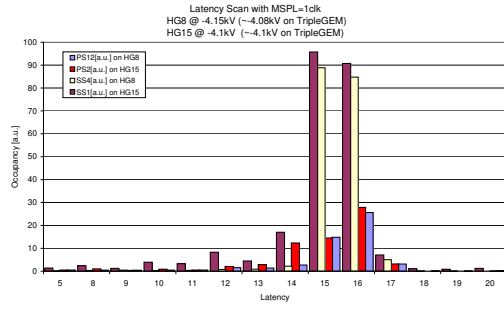


Figure 5.16: Latency Scan with the Monostable Pulse Length of 1clk. Same plot as in fig.5.14 with the two chamber HG08 and HG15 nearly at the same high voltage. The label PS12 means Pad Sector 12 and SS4 means Strips Sector 4. PS2 means Pad Sector 2 and SS1 means Strips Sector 4. The chamber in which they were plugged is shown on the plot.

Even if the gain at $-3.6kV$ is $\approx 10 - 15$ times lower than at $-4kV$, it is still possible to see some events and it is interesting to observe that the position of the maximum is shifted by 1clk cycle. This is coherent with the TDC measurement shown in fig. 5.7 where a relative delay of about one clock cycle is found between the two used voltages.

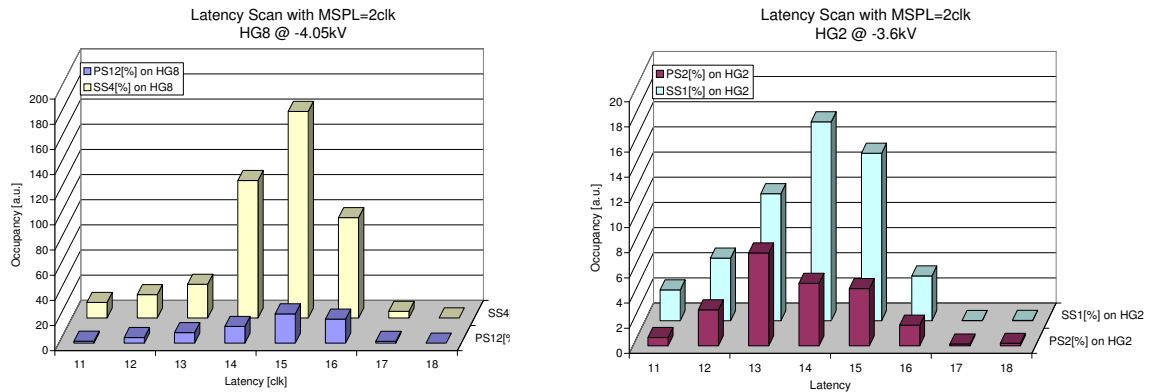


Figure 5.17: Latency Scan with the Monostable Pulse Length of 2clk. Left: HG8 with $-4.05kV$ applied to the chamber with an H.V. filter of $95.9k\Omega$ (i.e. $\sim -4kV$ on the triple GEM). HG2 with $-3.6kV$ applied to the chamber with an H.V. filter of 150Ω (i.e. $\sim -3.6kV$ on the triple GEM). The z axis of the second plot is expanded with respect to the one in the left plot by a factor 10.

5.1.3 Preliminary Efficiency Measurements

A preliminary measurement of the efficiency with the final electronics has been done during the test beam. The available electronics for the readout gave us the possibility to equip at the same time only two detectors, each with one VFAT on 128 strips and one on 120 pads. The experimental set up is shown in fig.5.18. The coincidence of the two scintillators was used as external trigger for the readout. One of the two detectors tested has been used as a reference to select good tracks among the triggers of the scintillators, removing any dependence on their coverage.

A good track has been defined in the reference chamber as an event with one cluster of pads collinear with a cluster of strips. No more than one cluster of pads and one of strips has been required. Moreover

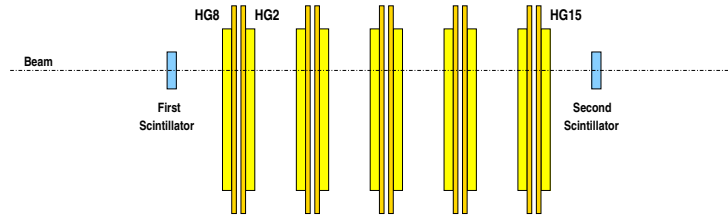


Figure 5.18: Telescope set up for efficiency studies in H8.

the maximum size of the clusters has been imposed to be less than 4 pads and 5 strips. The aim of this definition was to accept only real tracks.

In the following data the detector under test was HG02, while HG08 has been used for tracks selection (see fig.5.18). The two detectors were installed in a back to back configuration. HG08 had a filter of $95.5k\Omega$ in series with the high voltage divider board. The thresholds used were 35 VFAT DAC step ($\approx 21000 e^-$) for HG08 and 25 ($\approx 15000 e^-$) for HG02. The monostable pulse length of the VFAT was fixed to $2clk$. Good events have been selected in the data of HG08 according to the selection previously described. The efficiency of HG02 has been deduced analyzing the data that correspond to these events. In particular the collinearity with the clusters found in the reference chamber has been imposed. A cluster measured on pads in HG02 is considered real if its center is inside the pads that correspond to the center of the reference cluster or in one that is adjacent to it. The resolution of strips require more attention to properly define the acceptable radial distance between clusters of the two detector. A preliminary plot has been therefore extracted, where no cuts have been imposed to fix offsets and spreads. This plot is shown on the left side of fig.5.19 while on the right side are shown the central events selected to choose the offset and the σ (fixed in this case to 0.3 and 0.6 strips respectively). The event was accepted if the cluster distance (with the offset added) was less than 3.5σ . In fig.5.20 the calculated efficiency for pads, strips

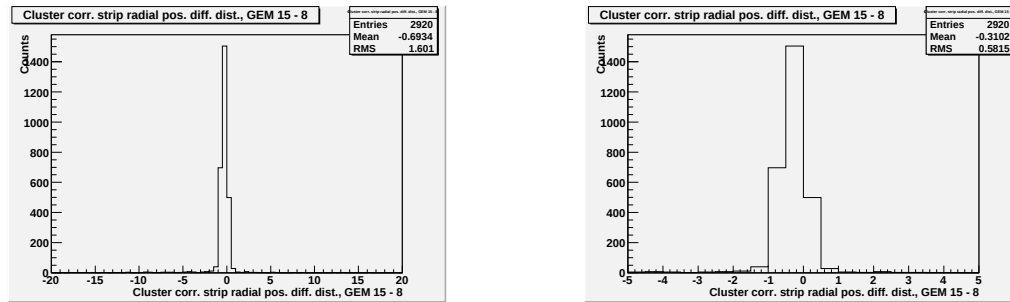


Figure 5.19: Strip clusters radial hit separation between chamber HG08 and HG02. Left: full set of data corresponding to the events defined by the reference chamber HG08. Right: Subset used to fix the offset and the σ of the acceptable distances.

and both is shown versus the high voltage applied to HG02. The plateau of the efficiency has not been reached at the high voltages tested (i.e. up to the nominal high voltage).

The measurements performed with a TDC unit have shown that the σ of the Gaussian distribution of the time in which the signal overcomes the threshold was nearly $0.7/0.8clk$ cycles ($25ns/clk$) at the nominal triple GEM high voltage and with the thresholds used for the VFATs of $\sim 30-40$ VFAT DAC step, i.e. $\sim 18000-24000 e^-$. This implies a total spread of about three/four clock cycles. This will heavily affects the efficiency measurement that will depends on the monostable pulse length (MSPL) used if it is not enough large. The 2 clock cycles stretching used in the measurement was therefore too low. Higher values are accepted and will be used in the final configuration. There is moreover large margin on the gain, that can be safely increased with respect to the nominal one. We decided however to modify the internal field configuration of the detector as noted at the end of sec. 5.1.1 in order to improve the time performance and consequently the detection efficiency for smaller MSPL stretching.

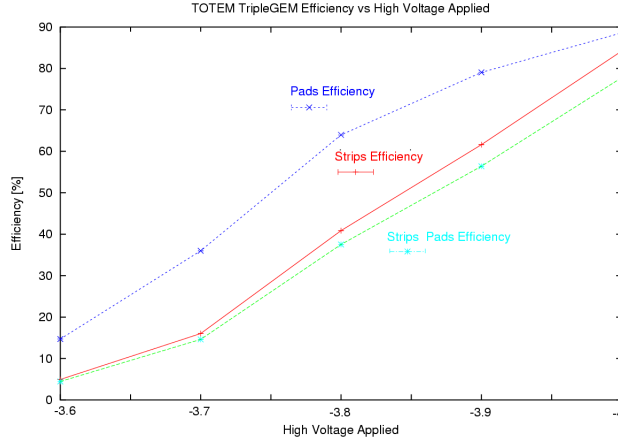


Figure 5.20: HG02 Efficiency measurement versus the high voltage applied. The measurement has been performed using a threshold of 25 VFAT DAC step ($\approx 15000 e^-$) for HG02. The monostable pulse length of the VFAT was fixed to $2clk$.

5.1.4 Cluster Size

Fig. 5.21 shows the cluster size distribution for strips and pads at the typical HV of -4 kV. In fig. 5.22 the mean cluster size is shown as a function of the high voltage and the VFAT threshold. The cluster sizes found are compatible with the results obtained by a simulation.

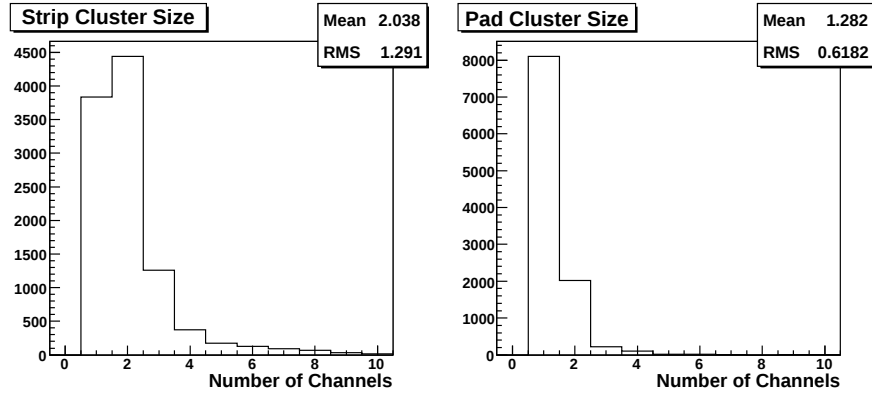


Figure 5.21: Strip and pad cluster size at $HV \approx -4kV$ with a threshold of 40 DAC steps (or 24000 electrons) and a signal sampling time window of 2 clock cycles in the VFAT.

5.1.5 Spatial Resolution

The cluster size distribution has an impact on the spatial resolution of the detector. For the operating parameters chosen for the tests underlying fig. 5.21, pad clusters consist predominantly of only one pad, which leads to a resolution of $p/\sqrt{12}$ (where p is the pad width), ranging from $2\text{ mm}/\sqrt{12} \approx 580\text{ }\mu\text{m}$ to $7\text{ mm}/\sqrt{12} \approx 2\text{ mm}$. The strip clusters on the other hand contained mainly one or two strips, with approximately equal probability. Since tracks passing near the center of a strip produce preferably 1-strip clusters whereas tracks passing between two strips will rather give 2-strip clusters, the resolution will be better than for a pure 1-strip cluster population. While the precise value of the resolution depends on details of the charge sharing mechanism, one can expect a range between $0.5 d/\sqrt{12} \approx 58\text{ }\mu\text{m}$ and $d/\sqrt{12} \approx 115\text{ }\mu\text{m}$, where $d = 400\text{ }\mu\text{m}$ is the strip pitch. In the test beam setup at hand, no external reference detector was available, which excludes a direct measurement of the resolution. However, with the hit measurements in two GEM planes and an approximate knowledge of the beam parallelism, a very rough consistency check is possible.

Fig. 5.23 shows the distribution of the radial distance between the strip cluster centers belonging to projective track hits in the two GEM planes with a distance $\Delta Z_{GEM} = 395\text{ mm}$ along the beam. The

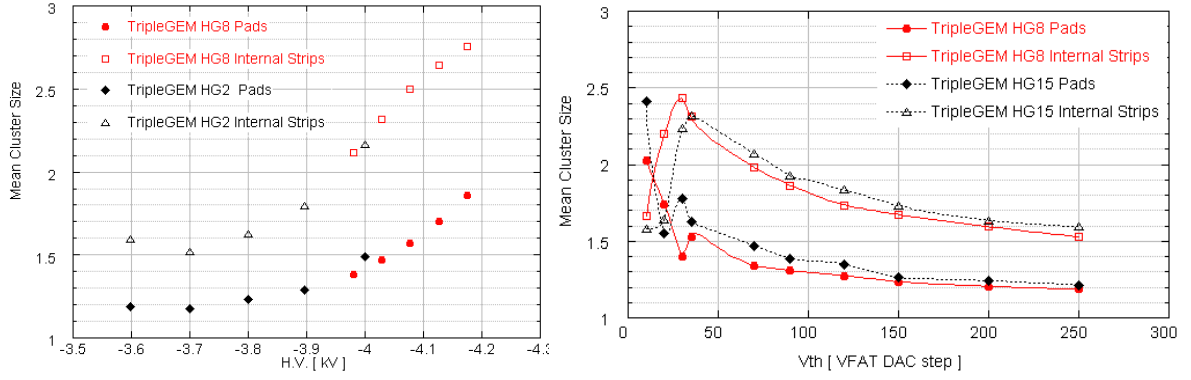


Figure 5.22: Cluster size for strips (open marker) and pads (filled marker) as a function of the high voltage for a threshold of 35 DAC steps (top), and as a function of the VFAT threshold for $HV = -4.1\text{ kV}$ (bottom) for different GEM planes (HG8, HG15). One DAC step corresponds to about 600 electrons. In all cases, the signal sampling time window in the VFAT was 2 clock cycles wide.

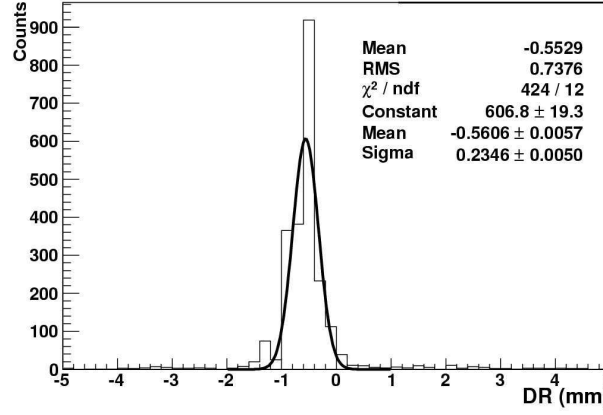


Figure 5.23: Radial distance for strip clusters belonging to aligned hits in two different T2 planes.

observed standard deviation, $\sigma_{obs} \approx 235\text{ }\mu\text{m}$, of this distribution can be decomposed according to the relationship

$$\sigma_{obs}^2 = 2\sigma_{GEM}^2 + \Delta Z_{GEM}^2 \sigma^2(\theta_{beam}) . \quad (5.1)$$

Solving eq.(5.1) with $\sigma_{GEM} \sim 115\text{ }\mu\text{m}$, the angular spread of the beam $\sigma(\theta_{beam})$ has a value between 0.43 and 0.56 mrad which is consistent with the beam characteristics.

5.1.6 LV1 induced noise

During the first test beam we observed an induction of noise when the LV1 signal reach to the VFAT2. A direct measurement of this induced noise has been done monitoring the trigger S-bit on the oscilloscope. In fig. 5.24 a picture of the scope display with this signal is shown. The timing of the various signals are compatible with a noise injection in correspondence with the arrival time of the LV1 on the VFAT2.

Increasing the VFAT2 threshold to a level of 60-65 bins, this noise disappears. This is not acceptable (the increase of the threshold) and therefore a VETO will be used in the final setup to overcome this problem.

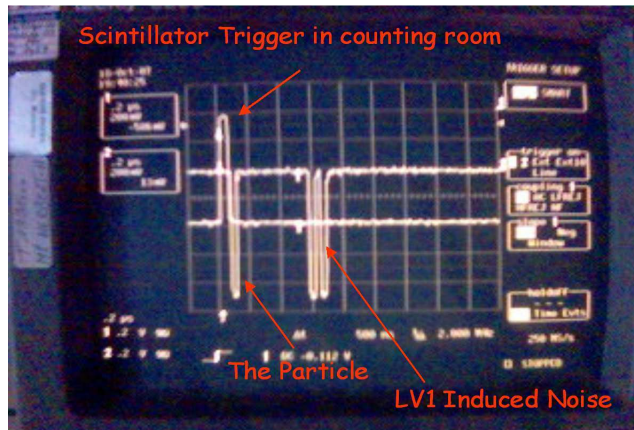


Figure 5.24: The first two signals in the display are the scintillators coincidence and the trigger S-bit. The time difference between them is related to the different time response of the scintillators and of the Triple GEM. After nearly 400ns other trigger S-bit arrive. The LV1 signal (produced in counting room after the arrival of the scintillator coincidence) takes approximately 200ns to reach the test beam zone. The same time is spent from the trigger S-bit S1 to come up in the counting room.

5.2 Cosmic Ray Data

The 2008 test beam was followed by runs of cosmic rays data taking in the H8 area in Preveessin (CERN) due to the lack of beam for the first half of 2009, which was unfortunately the period when we needed to test the T2 before the complete installation foreseen for May 2009. The results will be presented in this section. The data analysis has been followed by M. Berretti and a detailed description of the results can be found on his talks [42] [43]. I will report here only the parts where I've been directly involved. Other aspects (as the studies of detector alignment, spatial resolution, etc. etc.) can be found in his works.

The Triple GEM detectors have been tested in the final mechanical assembly (the T2 quarter) with their final electronics. The system has been equipped with the same low voltage and high voltage power supplies, cables, DAQ modules and boards that will be used in the CMS experimental area (IP5) of the LHC. All the detectors have been upgraded according to the previous noise and test beam data (i.e. grounding, shielding, new high voltage board and new internal field configuration). The installation schedule has not given us the possibility to continue on the optimization of the single detector and we spent our efforts on the test of the performance of the full quarter. The Triple GEMs have been therefore used at the nominal working point. Improvement, if needed, will be done directly in IP5 (i.e. increasing gains and fields inside the detector, VFAT2 thresholds optimization and changes of the gas mixture). Each quarter has ten collinear detectors and, thanks to this redundancy, the performances of the single Triple GEM are not so critical in terms of tracking and triggering capabilities of the telescope.

The data collected have been used to validate the status of the quarter. We started moreover to test and tune the digitization (see sec. 3.2) developed for the T2 telescope. Other data are needed on this issue, but the first results have indicated that the simple approach used could be satisfactory for our aims.

5.2.1 Data

The particle trajectories in cosmic ray runs are totally different from the ones expected for particles in the LHC collisions. In the very forward region covered by the T2 telescope (~ 40 cm of length at nearly 14 m from the interaction point, with an absolute pseudo-rapidity range between 5.3 and 6.5) the particles of interest are practically perpendicular to the detectors planes. The tracking methods, optimized for this types of tracks, was not the best solution to analyze cosmic rays data. It has been therefore modified to accept, as much as possible, all the real tracks collected in order to have enough statistics for the data analysis. In particular an $r - \Phi$ linear fit has been replaced with an $X - Y$ one and the cuts on the measured angles have been released according to the geometric acceptance of the T2 telescope for cosmic rays. In fig. 5.25 (left) a picture of one quarter of T2, vertically oriented to increase the cosmic rays acceptance, is shown together with a polar angle distribution of the tracks measured during one run. The polar range

found is in agreement with the detector geometry and the number of hits per tracks required. The radial coverage of each detector is about 10 cm and the height of the full telescope is about 40 cm. The accepted polar angles should be therefore less than 0.25 radians if high number of hits per tracks is required.

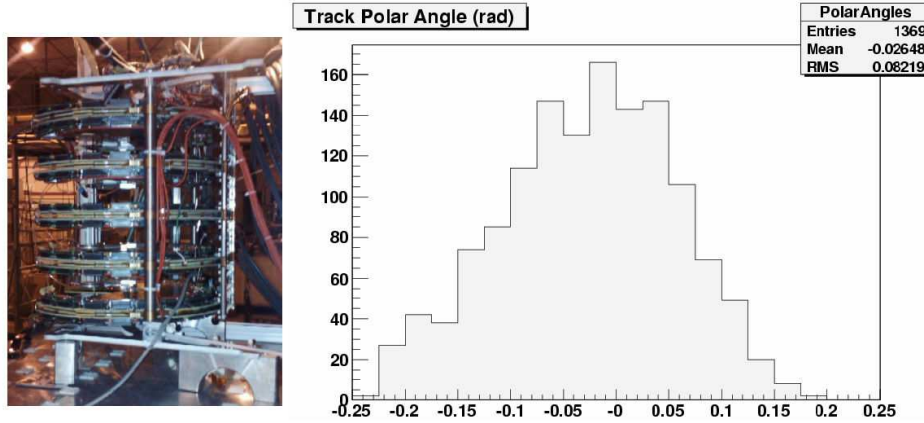


Figure 5.25: Left: One T2 Quarter during cosmic ray data taking in the H8 test beam zone (CERN-Preveessin site). Right: Polar distribution of measured tracks.

Before starting the data taking, noise scans have been done on the fully assembled T2 quarter in order to define the threshold configuration of all the VFAT2 chips installed. The found level of noise was worse than the one obtained from the single detectors but we were anyway satisfied because the thresholds required were acceptable (in terms of expected efficiency) and because the system was, for the first time, exactly in the final configuration (except the external environment obviously). It was moreover the first time that we had found the system in a working condition since the beginning, without the necessity of efforts to reduce the noise level. The upgrade that followed the previous tests have been demonstrated to be absolutely helpful.

In the following, preliminary efficiency measurement, used to validate the status of the T2 quarters, will be shown. For time constrain we decided to validate each quarter if the overall performance was within the requirements. This makes less critical the requests on the single detector. Three regions have been therefore defined in the efficiency measurement. If the efficiency of each single detector was in between the 80% (value at which simulations have been successfully done) and the 95%, the quarter has been considered ready to be installed in the Interaction Point at the LHC. An optimization of the parameters (threshold, gain and latency) will be done "in situ" during the commissioning to improve the performances of the detectors found in this region. If the efficiency of the detector was higher than the 95%, the detector has been considered well within the requirement and no further optimization will be required. If it was less than the 80%, a preliminary optimization (enough to overcome the 80% limit) has been done immediately. These regions will be indicated in the plots.

If not explicitly reported, the high voltage used in the following tests was the nominal working point of $-4.2kV$. It differs from the initial nominal value of $-4kV$ because the divider has been modified to increase the induction field. The procedure used to evaluate the efficiency follows the basic idea that has been described in sec. 5.1.3. This time however we have ten detectors fully equipped and therefore the results are more accurate. The data have been analyzed roughly in the following way. One detector is selected as detector under test (DUT). The other nine detectors are used as reference to define the passage of the particle and its track. The cuts on the reference detectors have been very tight, in order to have a clean and reliable sample of tracks. A high number of hits (defined as strip cluster overlying a pad cluster in the same detector) per tracks in the reference detectors (at least 7, depending on the found number of tracks) was required as well as a high collinearity between hits of different detectors to have a good definition of the track parameters. In addition a limited number of hits per detector was also required to avoid any noise contamination in the tracks tagging. No cuts have been however required for the DUT. Once the track has been found, it has been checked if in the DUT there was a hit (pad cluster over strip cluster) in the expected position with a reasonable spread. If the hits were found the numerator in the efficiency formula is increased. The denominator is fixed by the number of tracks found. In fig. 5.26 it is shown the efficiency measurement obtained under cosmic rays for ten detectors mounted on one T2

quarter. The used threshold are reported (all between 40 and 15 VFAT DAC Step) and are expressed in the format `detID(threshold for strips in VFAT DAC steps,threshold for pads in VFAT DAC Steps)`. The high voltage was the nominal one for all the Triple GEM. The stretching of the monostable (MSPL [32]) was fixed to three clock cycles. The difference between the detectors is explainable taking into account the intrinsic differences between detectors and the different threshold used. The three different regions of validation are superimposed in the plot. From the point of view of the quarter, it has been considered "ready to be installed" if all the detectors were over the 80% of efficiency.

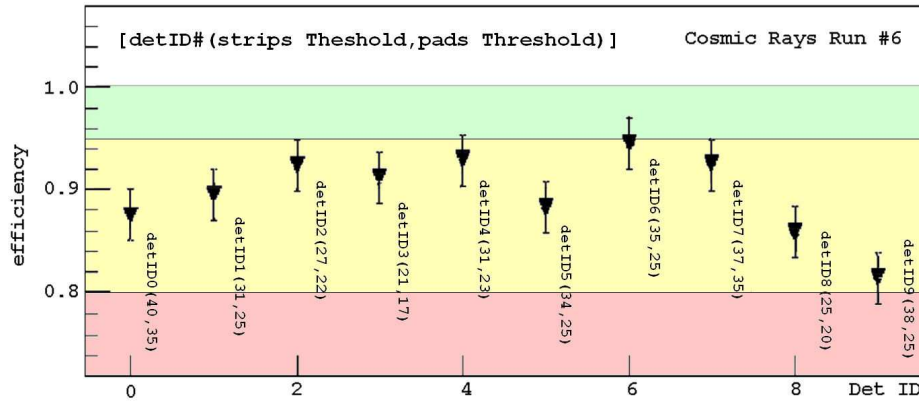


Figure 5.26: Cosmic Rays Efficiency Measurement: HV=-4.2kV, MSPL=3clk, Threshold on the plot. The three regions used to validate the status of the quarter are also shown.

The shown efficiencies in fig.5.30 of each individual detector has been computed requiring a coincidence between pad cluster and strip cluster in that detector along the track defined by all the other. The efficiency for the overall track reconstruction asking at least 7 out of 10 planes, was always above 98% for all the T2 quarters. More than satisfactory. The pad-strip coincidence was required in order to reject the readout chain noise. Without this requirement, but still excluding big size clusters to remove the chip intrinsic noise, each individual plane efficiency for pads and strips respectively is shown in fig. 5.26.

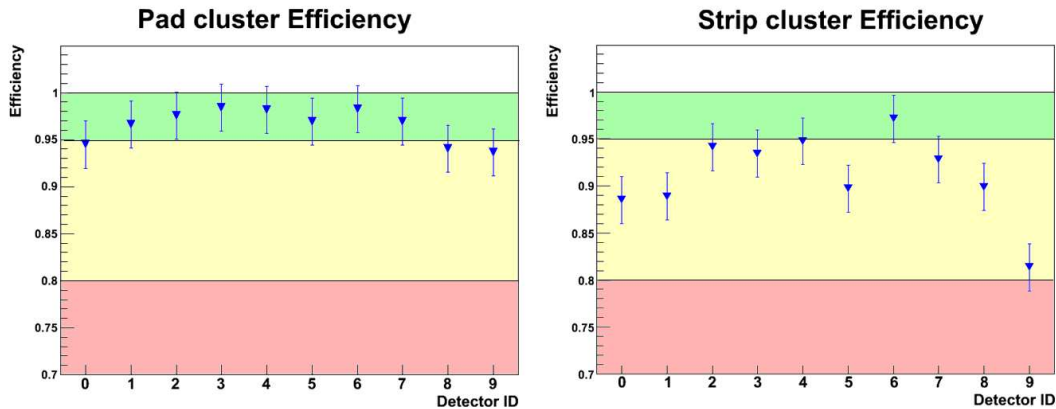


Figure 5.27: Cosmic Rays Efficiency Measurement for pads and strips for the measurement of fig. 5.26 .

While for pads the 95% is reached, the strips have to be still optimized during the commissioning. This is consistent with the higher threshold normally used for the more noisy strips. The lack of efficiency will be optimized during the commissioning at the Interaction Point 5. This will be done fixing properly the VFAT2 latency and thresholds and the detector gain. One example of the dependency of the efficiency on the high voltage applied to a Triple GEM detector, read with the VFAT2 chip, is shown in fig. 5.28. We participated to this measurement that has been done during the RD51 test beam in October 2009 together with the CERN GDD and CMS groups. In order to have a direct comparison with our triple GEM the different size of the electrodes and the particular high voltage filter used have to be taken into account. Nevertheless, relative quantities can be compared. In particular, the increase of voltage needed to pass from 80% of efficiency to the plateau is of about 150 – 200V. This is well within our possibility in terms of discharges and power supply rates.

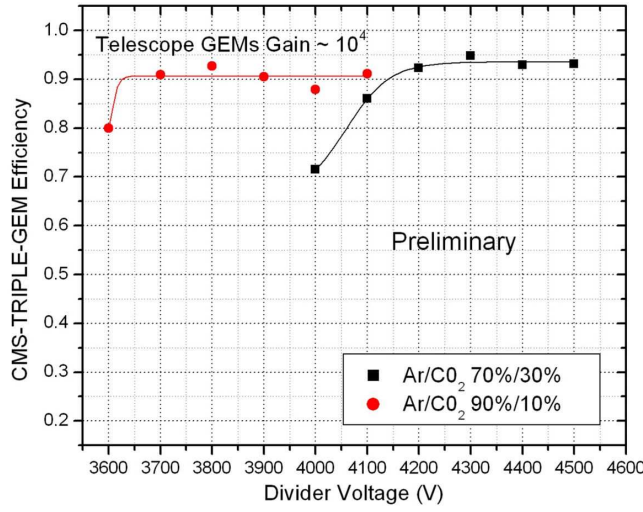


Figure 5.28: Efficiency Measurement on a Triple GEM detector with the VFAT2 during the October 2009 RD51 test beam [41]. Two different mixture are shown: Ar/CO₂ 90/10 (red line with circle) and 70/30 (black line with square). The latter mixture is the one used in TOTEM.)

Looking at the found efficiencies from the point of view of the single detector, it is clear that the performances are not optimized. A Triple GEM can reach higher value, remaining in a safe region of operation. A reduction of the threshold (i.e. of the noise level) or an easier higher gain or eventually a fastening of the signal could improve the results that have been shown. Nevertheless, if the efficiencies are analyzed from the T2 telescope point of view (i.e. ten detector aligned), the results are well inside the requirement and the telescope capabilities are not degraded.

Few tests, changing the working point of the detectors shown in fig. 5.26 have been done. Only few different points have been chosen. In fig. 5.29 the variation of efficiency for three detectors with two different sets of thresholds are shown, while in fig. 5.30 also the high voltage has been changed. The values are reported on the plot. The high voltage range is between $-4.15kV$ and $-4.25kV$, while the thresholds between 20 and 60 VFAT2 DAC steps.

No strong variation are observed with the threshold and high voltage values tested, taking into account

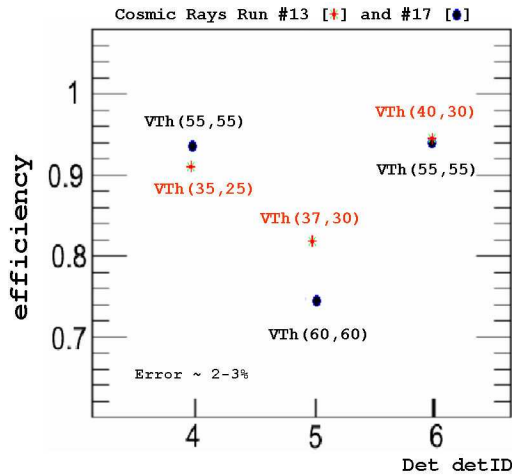


Figure 5.29: Efficiency measurement for different detectors (identified by the *detID* in abscissa) with the same high voltage applied and different threshold configuration. The error on the efficiency, not plotted, is of about 2 – 3%.

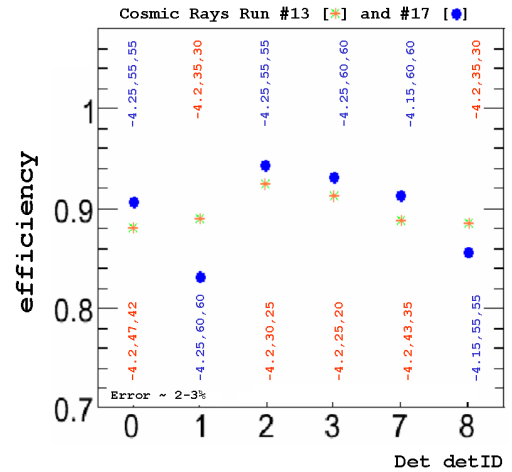


Figure 5.30: Efficiency Measurement for different detectors (identified by the *detID* in abscissa) with different High Voltage and Threshold configuration. The error on the efficiency, not plotted, is of about 2 – 3%.

the error of about 2 – 3% on the measurement. Detectors with higher efficiencies are less sensitive to variations of gain and threshold. In fig. 5.29 for instance, detector 5, with an initial efficiency less than

85%, is affected by the increase of the threshold much more than detectors 4 and 6.

The statistics of the collected data doesn't allow an accurate evaluation of these dependencies in the scanned range. Anyway, it is clear that up to a threshold of 60, the degrading of the efficiency is acceptable and compensation with the gain are still possible.

In fig. 5.31 typical pad and strip cluster size distributions, relative to the previous measurements, are shown. The hardware configuration (noise, threshold, gain) of the three detectors was different. These distributions will be used in the tuning of the simulation together with the efficiency.

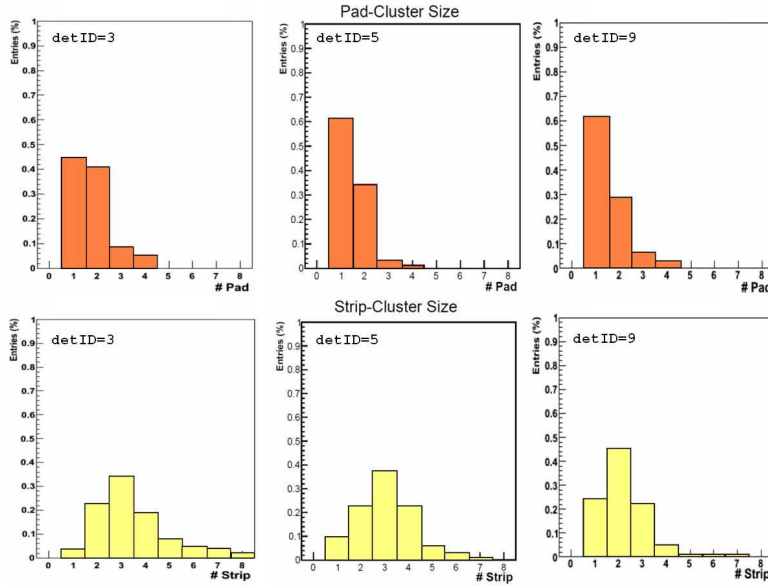


Figure 5.31: Cluster size measurement for three detectors in different configurations (noise level, VFAT2 thresholds, high voltage applied).

5.2.2 Simulation

Even if strong and evident correlations have not been found between the tested parameters (VFAT2 thresholds and high voltage applied) and the quantity observed (efficiency and cluster size), the data have been used to begin a preliminary validation of the T2 digitization. Complete scans versus high voltage and threshold have to be done for a complete validation of our simulation. With the information collected is anyway possible to understand if the developed tools are able to reproduce the measured data. The validation of the capability on predicting them will be left for the next data. Few runs at the LHC, during the commissioning, will surely provide enough data to face this point.

Let me here recall the core of our digitization, that is the algorithm that is used to compute the charge collected by one electrode when a particle releases energy in the sensitive volume of the Triple GEM. In eq. 5.2 the strips case is shown. The parameters used for pads are the same. Refer to sec. 3.2.3 for details. The algorithm, that provides the charge collected N by one strip when a particle releases C first ionization clusters of $N_C^{e^-}$ electrons each in the sensitive volume of a detector with a total gain G , is:

$$N = \sum_C N_C^{e^-} \cdot G \cdot n(d, ESW, \sigma_{ch}) \quad (5.2)$$

$$\sigma_{ch}(Z_C) = \sqrt{Z_C} \times \sigma_{DC} \text{Diffusion Coefficient} \quad (5.3)$$

$$n(d, ESW, \sigma_{ch}) = 50 \times |1 \mp 1 \pm [(1 - \text{erf}(\frac{d - ESW/2}{\sqrt{2}\sigma_{ch}})) \mp (1 - \text{erf}(\frac{d + ESW/2}{\sqrt{2}\sigma_{ch}}))]| \quad (5.4)$$

where d is the distance between the first ionization cluster and the middle of the strips, while ESW ⁶ is the effective strip width. The charge collected is therefore converted in VFAT DAC steps (or viceversa)

⁶The effective strips width is larger than the geometrical width of the strips. This is needed to take into account the electrostatic fields in the *induction zone*.

using the equivalent number of electrons per bin ($n_{el/bin}$) and the value obtained is compared with the threshold (Vth) used on the relative VFAT2 chip. At this point the simulation gives the status of one channel after the passage of a particle. The noise can be equivalently inserted, but at the moment it is still pending and it will be developed as soon as possible.

The status of one channel is therefore expressed (in the case of signal induced by particle) by:

$$\text{VFAT2 Channel ON if: } N \geq Vth \times n_{el/bin} \quad (5.5)$$

From eq. 5.2 and eq. 5.5, the available parameters in the digitization are:

- C and $N_C^{e^-}$: The number of first ionization clusters (with their spatial coordinates) and the energy (charge) released.
- G : The gain of the detector.
- σ_{DC} : The diffusion coefficient.
- ESW : The effective strip width.
- $n_{el/bin}$: The equivalent number of electrons per VFAT DAC step (bin).
- Vth : The VFAT2 threshold used in the internal comparator stage.

The approach that we used can be summarized as follows. A fixed number of first ionization clusters has been chosen as a first approximation. The clusters have been uniformly distributed along the trajectory of the particles in the drift volume. The energy released by each particle is obtained by the GEANT4 simulation. These parameters therefore are fixed. The same is for the threshold of the VFAT2s that is defined by hardware, as it is the value written in the chip register. It cannot be changed except if the hardware configuration has been also modified. All the other parameters are degrees of freedom that we can use to tune the digitization.

The gain G is the only one that can change between different detectors even if the high voltage applied is the same. The triple GEM detectors with the same high voltage applied can indeed show different gain⁷. It is not useful therefore to hardware define the G (unless a complete lookup table for all the detectors that links high voltage and gain is available). All the other parameters are common and have to be fixed once and used for all the detectors.

The σ_{DC} , ESW are useful to act on the charge sharing between electrodes. They modify the electron cloud shape and how it is divided between strips and pads. Their is important also for taking into account cross talk effects between channels. If for instance, the cross talk increases the cluster size, the σ_{DC} can be used to follow it. It is clear that the significance assumed by the parameter is not anymore the electron diffusion coefficient inside the gas. Its meaning is more complex.

The gain G and the $n_{el/bin}$ behave both as an amplification factor. Being both a multiplicative factor, it seems that there is a redundancy on the parameters. This is not totally true if the noise is considered. The G acts only on signal, while $n_{el/bin}$ acts on both, noise and signal. In the following test we neglect the noise and we fixed $n_{el/bin}$ to a value of 400 el per VFAT DAC step.

The three parameters, that we have left free in order to compare the data with the simulation results, were therefore G , σ_{DC} and ESW . The quantity observed were the efficiencies and the cluster size.

In fig. 5.32 the measured and simulated efficiency for one T2 quarter (ten Triple GEM) is shown. In this case the ESW has been fixed to $0.15mm$ and the σ_{DC} to $0.38mm/(cm^{\frac{1}{2}})$.

These plots are meaningless without the cluster size comparison obtained with the same setting. With the available parameters indeed, if only the efficiency is checked, any kind of data can be reproduced. Actually the tuning has been done looking to both, efficiency and cluster size, and finding the best value for the parameters that gives the best agreement between simulation and data. In fig. 5.33 two examples of cluster size comparison (one of the best and one of the worse agreements) for two different detector are shown. The used settings are the ones used to obtain the efficiencies shown in fig. 5.32.

It is important to stress that at the moment there is no noise integration inside the simulation and an

⁷The gain is an exponential function of the voltages applied across the foils. Small differences on the resistive divider that provide these voltages can induce for instance not negligible differences on the gain. The foil itself can be slightly different according to the production bunch

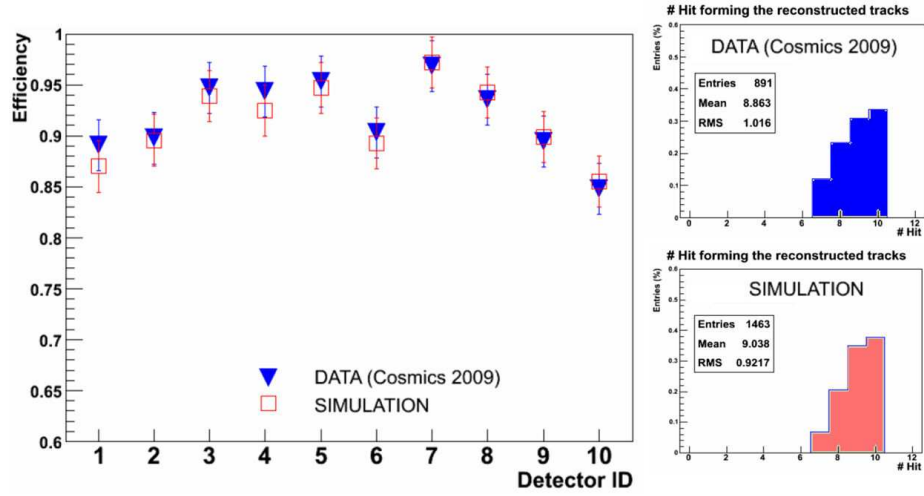


Figure 5.32: Efficiency comparison between cosmic ray data and simulation. On the right the distribution of hits forming the track are shown for both data and simulation.

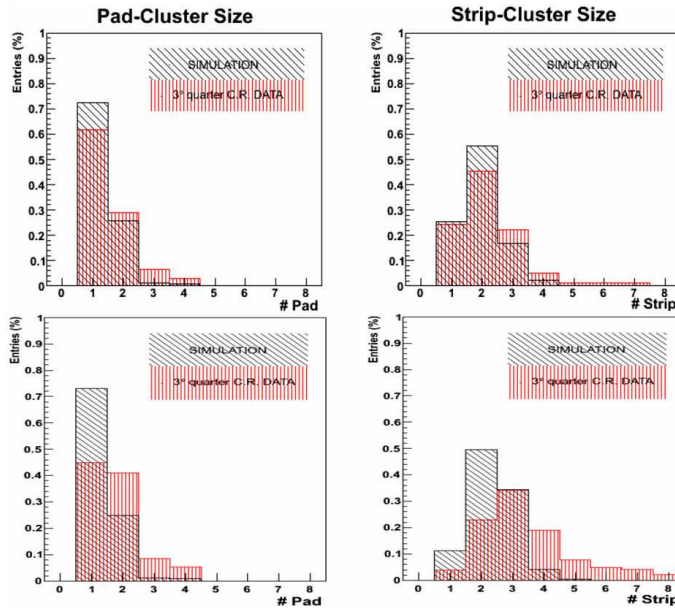


Figure 5.33: Pad and strip cluster size comparison between cosmic rays data and simulation. One of the best (top) and one of the worse (bottom) results are shown.

accurate analysis of the noise status of each detector for these data has not been done. When it will be done and after new data taking, with higher statistics and dedicated scan on the parameter involved (high voltage and threshold), the actual results will be improved a lot. We are nevertheless confident that what it has been found is encouraging in the sense that the simulation seems to be enough versatile to be properly tuned to satisfy its aim, even if extremely simple in its structure.

The efforts to find the right values of the parameters in the simulation have to be done only once. After they have been fixed, if something on the hardware setting is changed, the same has to be done on the simulation (i.e. changing the parameter that is affected by this change, without touching the other). If the gain for instance is increased, only the G in the digitization has to be properly scaled and the simulated output has to be automatically in agreement with data. As previously reported we don't have data that can be used to tune and test this capability on following (predicting) the behavior of the detectors when one parameter is changed. Nevertheless, some (even more than preliminary) tests have been done and they are reported in fig. 5.34. The digitization parameters have been considered free only initially, when the right values has been searched. After that, any modification has been done in a deterministic way (i.e. the threshold has been changed according to the one written in the chip register and the gain has

been properly scaled according to the high voltage applied on the divider board of the Triple GEM.).

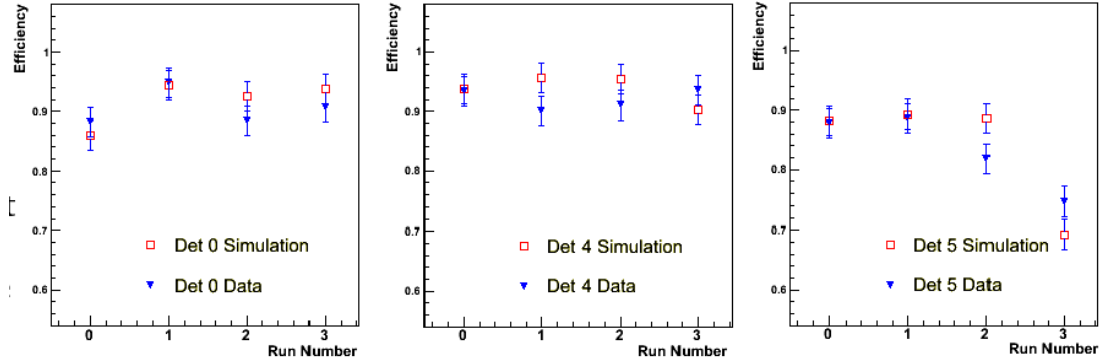


Figure 5.34: Efficiency comparison between data and simulation for three detectors and for four different runs with different detector settings (high voltage and VFAT2 thresholds).

5.3 LHC pp Collisions Data

In this section the T2 first measurements of LHC pp collisions will be shown. A brief introduction of the installation in the *CMS* experimental area of the Interaction Point five (IP5) will be given.

The T2 installation, where I've been involved, can be subdivided in three main periods:

- i. BEFORE the first LHC run (2008).
- ii. BETWEEN the the first and the second LHC run (2009).
- iii. AFTER the second LHC run (present).

The T2 status in IP5 at the end of these installation (BEFORE, BETWEEN, AFTER) periods was:

- i. The T2 Minus Far⁸ quarter installed, but not fully operative. All the services, for the full telescope, ready and tested (i.e. cooling, gas, high voltage, low voltage). Readout Electronics partially ready. DCS (detector control system) partially operative.
- ii. All the T2 quarters installed. The T2 Plus side fully operative. The T2 Minus side partially operative. Readout Electronics completely installed and partially ready. Trigger available. DCS (detector control system) operative.
- iii. All the T2 quarters installed and fully operative. Readout electronics completely installed and ready. DCS and DSS (detector safety system) operative.

The BEFORE period was characterized by the installation of all the cables and of the gas and cooling pipes. The power supplies (Low Voltage and High Voltage) have been installed with their control modules and tested. The data transmission lines and modules have been partially installed. One of the quarter of T2 has been assembled in laboratory and moved completely equipped in the cavern, where the connections to services, control loops and optical transmitters (data and triggers lines) have been done. During the LHC run it was not possible to acquire data for missing modules on the readout chain. We

⁸Minus and Plus End are relative to the interaction point, while Far and Near Side to the center of the LHC ring.

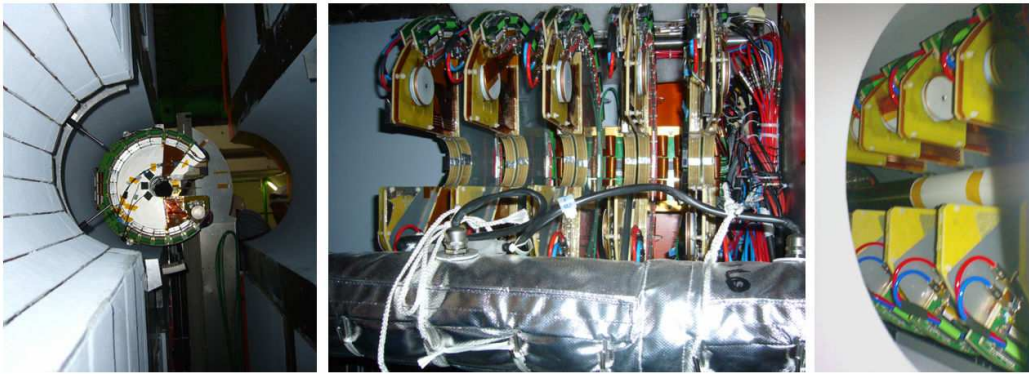


Figure 5.35: The Minus Far T2 quarter installed before the 2008 LHC run. Left: The quarter seen from the inside of the *CMS* Hadron Forward (i.e. from the IP side). The quarter is already inserted in HF. Middle: The quarter in proximity of the beam pipe. Right: The quarter in the final position with a clearance of about 10mm between the detectors inner radius and the beam pipe.



Figure 5.36: First closure of the forward region of *CMS* where T2 is installed. Left: The platform on two risers, ready to be moved on final position. The Hadron Forward, T2 and Castor are in open position to allow the insertion of the beam pipe. Middle: First Closure of the rotating shielding. Right: Forward region completely closed.

therefore used this run simply as an operational test, monitoring the status of the detectors and playing with the control loop of our electronics.

In the BETWEEN period we have completed the T2 installation. All the quarters have been equipped and moved in the CMS cavern. Part of the cabling have been modified and optimized according to the problems encountered in the previous closure. At the end of this period we encountered a problem in the control loop of the Minus Near quarter. The problem was related to the 11th card (see ch.2) installed on the detectors. There was no time to replace the card and the quarter has been excluded by the control loop (i.e. that quarter was unusable). The Minus Near quarter had instead a problem in the data output, due to a disconnection of one cable used to control the transmission. This accident happens during the final closure and there was no possibility to reconnect it, once the problem was found. The Plus End was instead completely working.

Half of T2 (in the Plus End) was therefore ready for the pp collisions of the second LHC run and the data will be partially shown in this section.

The AFTER period has been used to solve the problems encountered at the end of the BETWEEN period and during the second run. The full telescope has been left fully operative for the third run of the LHC.

The second run of the LHC started at the end of 2009. pp collisions with a cms energy of 0.9TeV and 2.36TeV have been obtained. The latter ones correspond to the higher collision energy ever achieved. The T2 Plus End was operative during this period and data from both energies have been collected. In

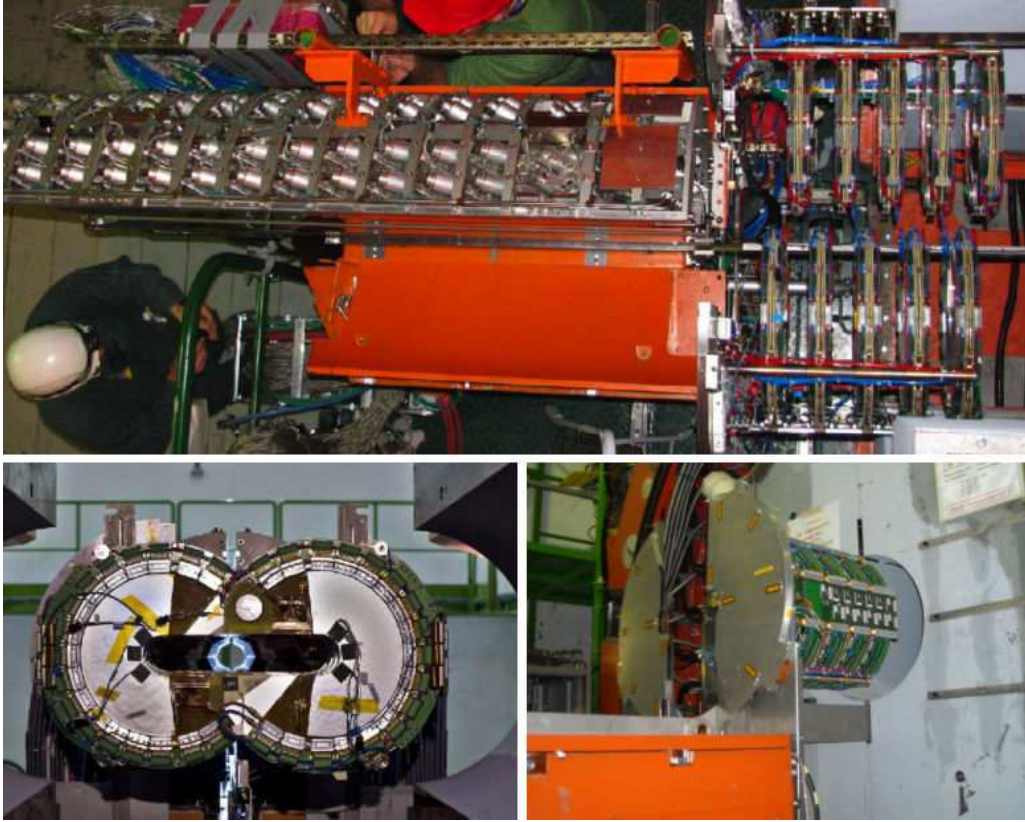


Figure 5.37: All the T2 quarter are installed on both end, minus and plus. Top: Minus End with the T2 quarters and half of CASTOR installed. Bottom: Plus End T2 quarters seen from inside HF and from the platform. The quarters in these picture are still outside the HF.

total a number of about 20k events has been stored. The second LHC run was the first one with part of the T2 telescope ready for the acquisition. The first issue was the research of the right latency for the VFAT2 chip. A minimum bias trigger has been generated by the triggering roads of our telescope (see ch.2 for details). Once that a rude estimate (transmission line delays) has been done, the search of the right latency has been done with a latency scan and looking at the plots in fig. 5.38, where the number of detectors involved per event is shown. Pads, strips and hits of both types belonging to different detectors in coincidence have been required. No collinearity has been however requested. The noise involves few detectors per event, while when a particle passes through the ten aligned detectors of each quarter this number increases drastically (as shown in the plots). The right latency has been found using these simple plots and looking at the their mean values. After that, a more accurate latency scan and a serious data taking has been done. The plots in fig. 5.38 have been the first evidence that we were revealing the passage of particles produced at the LHC.

The first collisions have been used to find an acceptable configuration for our detectors and for the triggering system. It was actually the first time that the full triggering chain of T2 has been tested in its completeness and the results were fine. After that, all the available "stable beams" have been used for data taking. A preliminary analysis has been started by our group [44]. I will show here only the most significative plots, based on raw data. They are all related to the tracking capability of our detector. I will add few qualitative consideration about these plots and their possible interpretation. It will be nothing else than a qualitative analysis based on raw data, waiting for the accurate one that is ongoing.

In fig. 5.39 the multiplicity of tracks per event and the η^9 distribution of tracks are shown for two

⁹ η is the pseudo-rapidity and it is defined as $\eta = -\ln \tan(\theta/2)$, where θ is the scattering angle. The T2 θ coverage is approximately between $3mrad$ and $10mrad$ from the beam center.

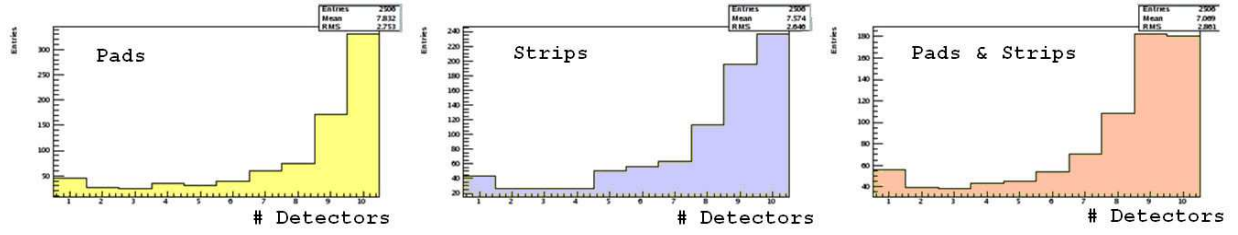


Figure 5.38: LHC Collisions: First T2 Particles signature. Distributions of the number of detectors with at least one pad (left), one strip (middle) and one strip and one pad (right) over threshold. No collinearity required. A cut on the maximum number of hits per detector fixed to 60.

sets of data acquired at $\sqrt{s} = 2.36\text{TeV}$. These information are particularly interesting because of the

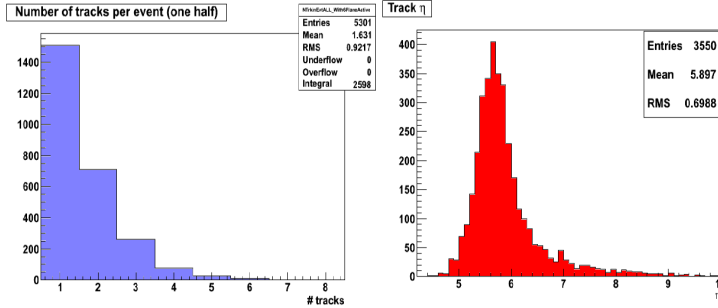


Figure 5.39: Track multiplicity per event (left) and their η distribution (right) for two set of data taken at $\sqrt{s} = 2.36\text{TeV}$ [45].

forward region covered by our telescope (i.e. an acceptance for particles arising from the interaction point between 5.3 and 6.5). In this region indeed, not enough data are available to compare with and validate the existing models. When a proper analysis of the found η distribution will be completed for instance, (i.e. the extraction of the signals from the background and right definition of the η scale) a direct comparison with MC could be done.

Once the tracks coming roughly from the interaction point have been selected, a first reconstruction of the primary vertex has been done. Due to its geometry and position, the accurateness of T2 is very good in the transversal plane (XY), while it is expected a rms of about $2 - 2.5m$ for the reconstructed Z of the primary vertex when the CMS magnetic field is on. The longitudinal coordinate of the vertex can be known more precisely with the measures of the $T1$ telescope. The transversal coordinates of the vertices reconstructed with the T2 tracker are shown in fig. 5.40 .

All these plots are obtained from raw data. No correction has been done at the moment. The analysis

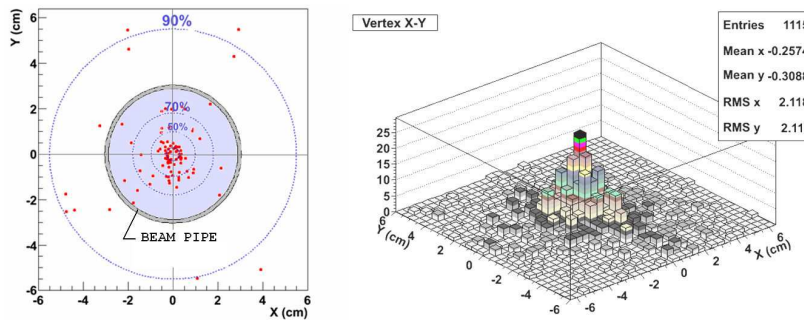


Figure 5.40: Primary vertex XY reconstruction [47] [46]. The beam pipe is shown as a reference.

of the detector alignment is ongoing and it will affect the previous results. One clear example of needed corrections is given in fig. 5.41 where the distribution of the longitudinal coordinate measured by one quarter is shown. The found distribution is centered at about $3.5m$ from the interaction point (the axis origin). This displacement can be explained by a quarter misalignment as the one drawn in the figure.

The vertex is between the interaction point and the detector. A common tilt of the quarter shifts the reconstructed vertex toward T2 if the quarter is less closer to the beam pipe in the interaction point direction than in the other. Displacement on a millimeter scale are enough to move the reconstructed Z of meters.

The knowledge of the quarters alignment is one important issue because it affects our measurement(as

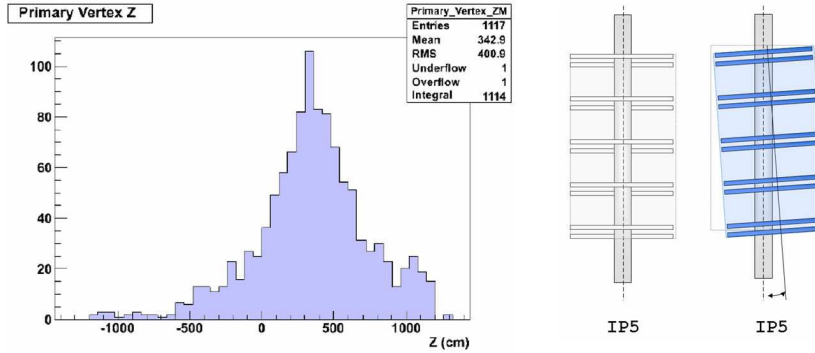


Figure 5.41: Left: Distribution of the primary vertex Z reconstruction [45]. Right: Quarter tilt that could explain the displaced distribution, moving the vertex in the quarter direction. The IP side is shown.

shown before). Unfortunately, the position of the quarters cannot be measured when they are in the final position (i.e. inside the hadron forward of *CMS*). The nominal distance between the inner radius of our detector and the beam pipe is of about 10mm. Displacement of few millimeters are accepted. For beam pipe safety, these distances are measured through position sensors mounted on few detectors (fig. 5.42). These sensors, even if very accurate for relative displacement, are not indicated for an absolute measurement. The data analysis will be therefore the only way to know with the proper accuracy the misalignment of the quarters. Once their absolute position with respect to the beam is found however, the alignment provided by the analysis will be checked looking at any relative displacement measured by the position sensors. The analysis alignment will be done therefore inside each quarter looking at the

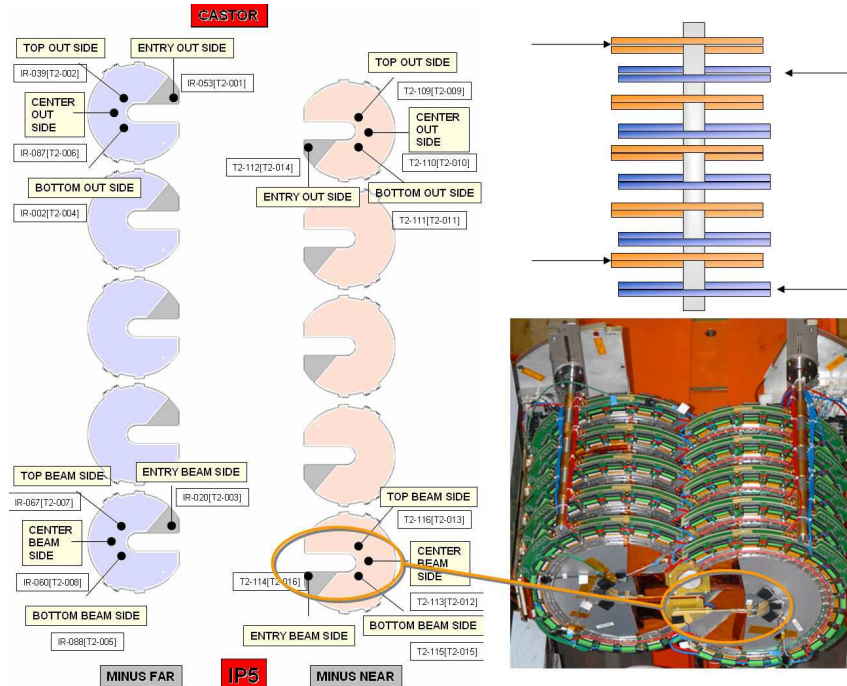


Figure 5.42: Position sensors installed on the T2 telescope. The mapping for the minus side (left). The sensors have been mounted on the detectors indicated with the arrows (top-right) and few of them are shown in the picture.

hit residuals with respect to the fitted tracks. Once the detectors inside the quarter have been aligned, the position with respect to the beam can be done looking at the reconstructed vertex distributions. A first check of the results can be done using events that happen in the overlapping region between two

quarters, whose positioning has been found independently. In fig. 5.43 one example of an event measured during the LHC run with one candidate in the overlapping region is shown. If the alignment is correct and if the hits in this regions are produced by a single particle, also an aligned single track has to be found by the two quarters.

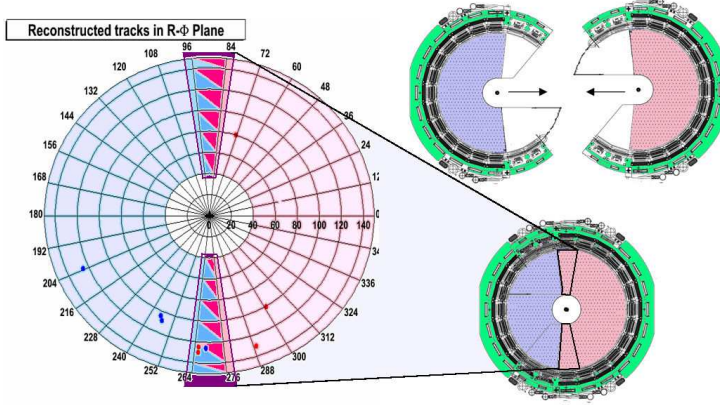


Figure 5.43: One event involving the two quarters with one track in the overlapping region.

It is very important to stress that the previous plots are based on raw data and that they are a very preliminary version. As an example, I will report a rude and qualitative analysis of effects that the displacement found in the Z coordinate of the primary vertex could have on the η distribution of fig. 5.39. Assuming that the tilt on the quarter is the explanation of the shift in Z, the η scale used in the previous plot has to be properly modified. Before doing that, it is useful to describe the triggering configuration used. As explained in ch.2, the trigger of our detector is generated by a triggering road (fig. 5.44).

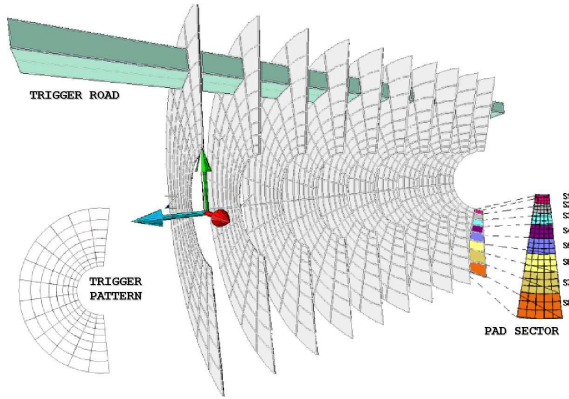


Figure 5.44: T2 Triggering Pattern. One triggering road is shown. Along R , the readout plane are subdivided in eight different sector $S1, \dots, S8$, with roughly the same $\Delta\eta$ for particles arriving from the interaction point. The data acquired during the 2009 LHC runs have been triggered on the internal $S1, \dots, S6$ sectors.

For each quarter 104 roads are foreseen (8 in R times 13 in Φ). Each road is done with ten aligned *Super-pads*, one for each detector of the quarter. The *super pads* are made of 15 pads of our readout plane (3 in R times 5 in Φ). The triggering pattern has, along R , eight different sectors $S1, \dots, S8$. During the 2009 LHC runs have been used. This will bias obviously the η distribution previously shown.

If the triggers acceptance is superimposed to the η distribution, few useful information can be extracted. The η coverage is obviously different if the triggering particles arrive from the interaction point or not. A particle created in a pp collision will pass through one of the triggering sectors depending on its η . The coverage of each sector $S1, \dots, S8$ is therefore limited by η_{min} and η_{max} . This is due to the fact that the origin of the track is fixed in the IP. The acceptance of the $S1, \dots, S8$ triggering sector, for background particles¹⁰ has instead only a limit on the η_{min} . This is due to the fact that now, the origin of the track is unknown. The η_{max} doesn't exist because a particle that is parallel to the beam ($\theta = 0$) can pass through any sectors, depending on its radial position. The η_{min} can be evaluated from the required number of detectors in coincidence and on the geometrical parameters of the road (i.e dimensions of *Super-pads* and their spatial separation). Its value can be different from the η_{min} used for particles emerging from the IP

¹⁰Particles emerging from secondary or beam gas interactions or belonging to the beam halo

and it becomes smaller when the required number of detectors in coincidence is reduced.

In fig. 5.45 the two cases are shown. The sectors coverage in the background case has been computed considering a coincidence of at least five detector over ten. The intervals superimposed are approximately right, as much as needed for this extremely qualitative approach.

I want to focus the attention on the acceptance interval for particles arising from the IP. In fig. 5.46 the

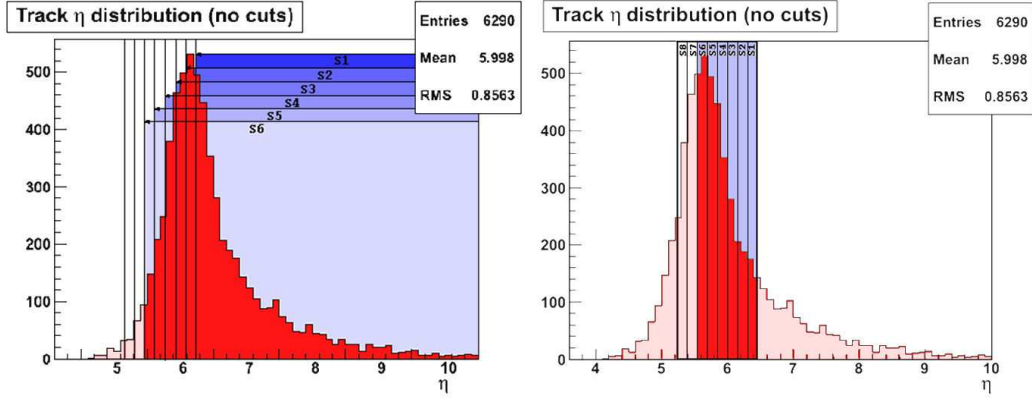


Figure 5.45: Nominal η acceptance of the trigger sectors S1,..S8 for background events (left) and particles arising from primary vertex in the interaction point (right). These intervals are superimposed to the η distribution of fig. 5.39. The darker part of each distribution represents events inside the nominal trigger acceptance. During the 2009 LHC runs, only the internal sectors S1,..S6 have been used.

same distribution is shown together with two subset of data, relative to the two available cms energy collisions, where a χ^2 cut on the tracks reconstruction has been applied. The nominal triggering acceptance of the S1,..S6 sectors is drawn.

The coverage of the nominal triggering interval for IP particles is displaced from the main peak of the η

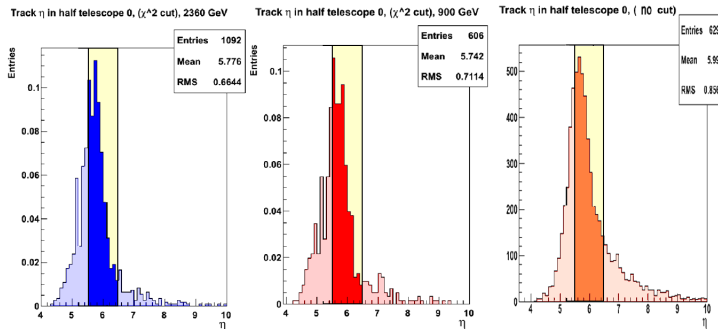


Figure 5.46: η Distributions with and without cuts for different cms energies. Nominal η trigger sectors (S1-S6) acceptance for particles arising from the interaction point.

distribution. If the previous hypothesis of a quarter misalignment with respect to its nominal position (a possible explanation for the Z primary vertex distribution of fig. 5.41) would be true, the η used for the distribution should be changed.

In fig. 5.47 the same plots are shown, having computed for simplicity the new triggering interval instead of changing the η scale of the distributions. The new range has been obtained considering the 3.5m shift in Z of the vertex. The main peak of all the η distribution fits quite well within the η acceptance of the trigger sectors S1,..S6. When a more accurate analysis will be done and the quarter misalignment will be taken into account, the proper η will be used in the previous distributions and at that time, they could be compared with the nominal trigger acceptance. This will be another useful check of the correction done. Even if approximative and qualitative, the result seems to be more than fortuitus and the peaked structure seems to describe the particles of our interest. The next step will be the understanding of the background. The trigger acceptance interval could be useful also from this point of view. In fig. 5.48 I've reported two of the previous distributions where there was a more evident distinction between the peaked structure in the IP trigger acceptance and the rest of the distribution. The counts for η corresponding to particles arising from the IP (the previous shift of the interval has been used) have been masked and a

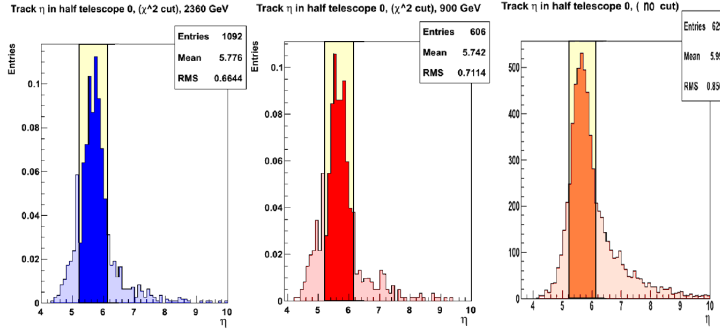


Figure 5.47: Same distributions as in fig. 5.46. In these plots the η trigger sectors (S1-S6) acceptance for particles arising from the interaction point has been calculate taking into account the Z displacement found in the vertex reconstruction(fig. 5.41).

“Landau” distribution has been superimposed as a very rude and approximative description of the background. In fig. 5.49, the same “Landau” distribution, properly scaled, has been therefore superimposed

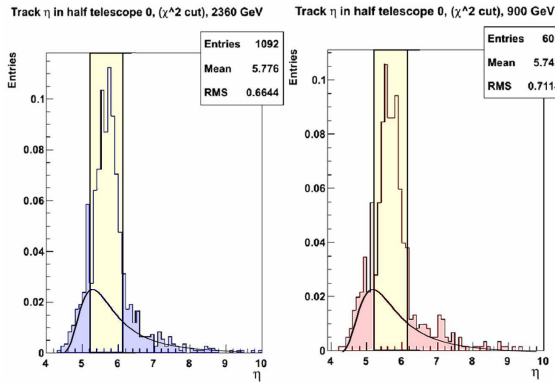


Figure 5.48: Qualitative background evaluation from the η distributions without the events belonging to the η range covered by particles arising from the IP. Landau distributions have been superimposed.

to the distribution of fig. 5.39 where the peaked structure was less evident. A background subtraction has been done and the remaining distribution (the signal) has been superimposed with the nominal trigger acceptance for particles arising from the IP. The new distribution found could be the real η distribution that we are looking for, without background contamination. If the assumed tilt of the quarter is applied

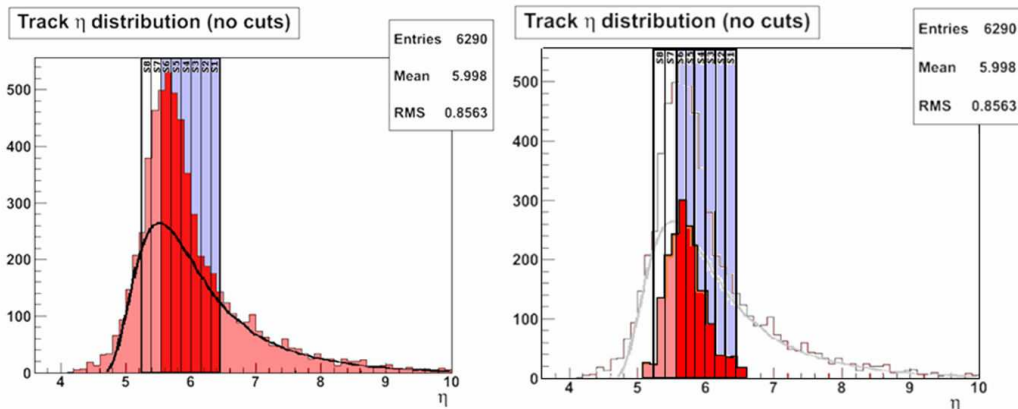


Figure 5.49: Qualitative signal extraction from the background. The nominal triggering windows for particles arising from the IP is shown.

on the signal (without background) of this distribution, the inclosure of the signal inside the trigger window is also more evident than before. In fig. 5.50 the η of the hypothetical signal is changed using the η corresponding to the raw data, the η corresponding to a Z shift of the primary vertex of 5m and finally

the η corresponding to a Z shift of the primary vertex of 3.5m (i.e. the one found experimentally). In the last case the maximum overlapping is found.

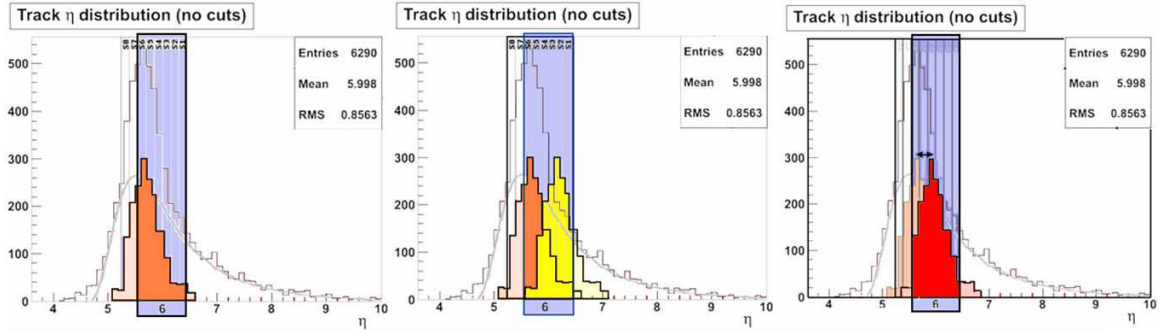


Figure 5.50: Different overlapping between the signal extracted from an hypothetical background and the nominal triggering window, considering different tilts of the T2 quarter. Left: raw data. Middle: displacement obtained by a detector tilt that moves the reconstructed Z of the primary vertex at 5m from the IP. Left: displacement obtained by a detector tilt that moves the reconstructed Z of the primary vertex at 3.5m from the IP (i.e. the shift observed).

All the data collected have been monitored and analyzed with the TOTEM software. With the available tools it is possible to check the status of the detectors event per event and to retrieve useful distributions related of interesting quantities (as the ones previously shown). After the 2009 LHC runs however, a visualizer of the status of the complete telescope in single display was missing. Due to the quantity of channels and the complexity of the system, these plots are very useful to have an idea of the detectors status before the analysis is performed. A 3D visualizer is in progress and it will be developed in the FROG¹¹ environment. In the meanwhile, a simple event display has been developed in the ROOT¹² environment. It is based in the ROOT TGEO class. Even if there are a lot of useful methods in this class that can be used to create an event display with very interesting features, only the simple visualization of 3D objects (detectors, strip, pads and tracks) has been used.

In fig. 5.51 four different events are shown as an example. The tracks are reconstructed requiring few

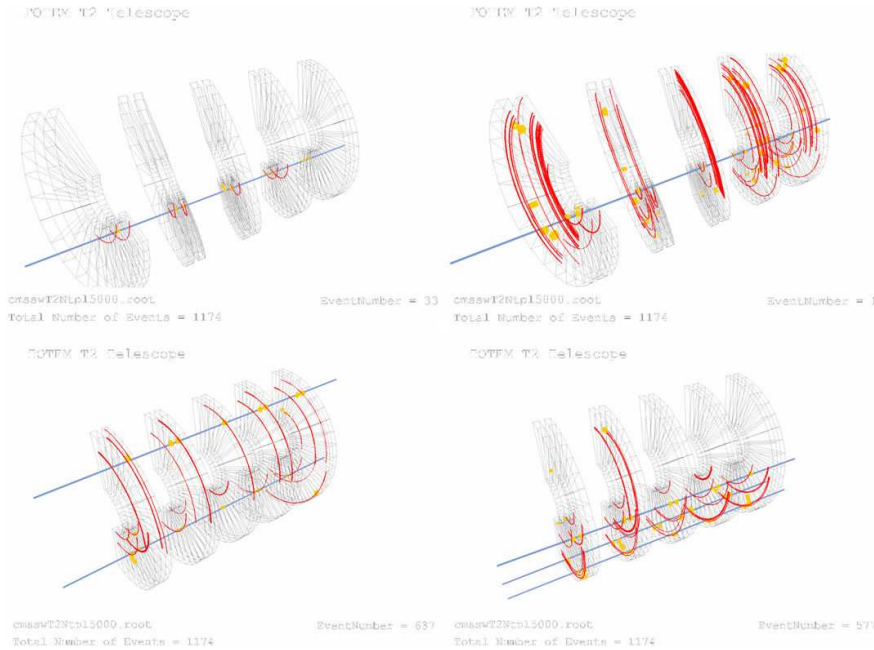


Figure 5.51: Four event of the 2009 LHC runs. One quarter is displayed. Detectors shape, hit pads (yellow blocks), hit strips (red lines) and reconstructed tracks (blue lines) are shown.

¹¹<http://projects.hepforge.org/frog/>

¹²<http://root.cern.ch/root/>

collinear detectors with a cluster of pads underlying a cluster of strips in its readout plane. The first display (top-left) shows a good example of very clean event. Strip clusters overlying pad clusters are found in seven detectors out of ten and they are collinear. In the second display (top-right) however, a more noisy condition has been found. The reconstruction of the tracks is capable anyway to isolate hits related to the passage of the particle from the noise. Examples of multi-tracks events are shown in the two bottom plots.

It is clear that a single event display is not as powerful as a distribution or a cumulative plot related to large number of events. It is however useful to monitor the current status of the detectors in real time and it can show in a much intuitive way which kind of problem the telescope could have, without mixing the information coming from different events. It is moreover useful for debugging purposes. In clean events like the ones shown for example, the tracks reconstruction can be directly checked with the corresponding real hits. Fig. 5.52 shows one event with a clear noise pattern in the sixth detector. The same pattern has been found quite frequently in other events. The external strips show a burst of noise (all the channels over threshold) that is reflected on few groups of underlying pads. Each trapezoidal shape in the pads pattern is associated to one VFAT2 chip. This suggests that the threshold used for those channels has to be probably increased. In fig. 5.53 an other typical event recorded during the 2009

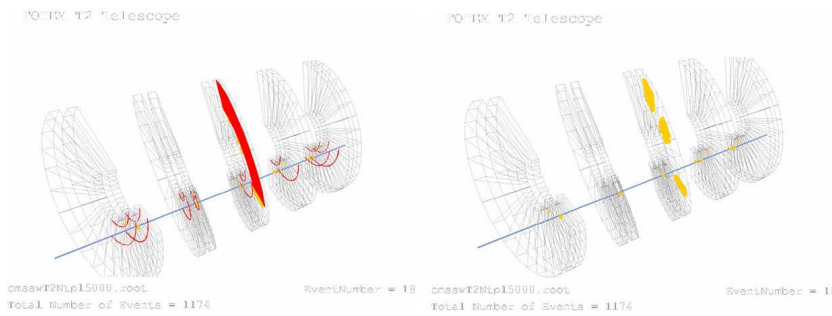


Figure 5.52: Typical noise pattern due to strips and pads cross talk. The same event is displayed with pads and strips (left) and only pads (right).

LHC runs is shown. Only the pads are shown for clearness. In this kind of event, a large number of pads and strips hits has been found in all the detectors. The tracking procedure normally reconstructs only few tracks (i.e. there are only few collinear hits between different detectors, with for each readout plane a cluster of strips overlying the pad cluster). They happen quite often. Pattern of noise, that involves so many channels per detector, normally doesn't show the found pads granularity. Single blocks of many channels are normally found in those cases. The same for strips. At the same time the passage of particles should induce more collinear hits of strips over pad cluster. This remains an open question that has to be faced off by the analysis and by the next data taking. We have to understand if the problem is in the hardware point of view (detectors) or if we are doing something wrong in the tracks reconstruction or if they are real secondary events with a vertex so close to the detector that, for their anomalous η , they cannot be recognized by the tracks reconstruction methods.

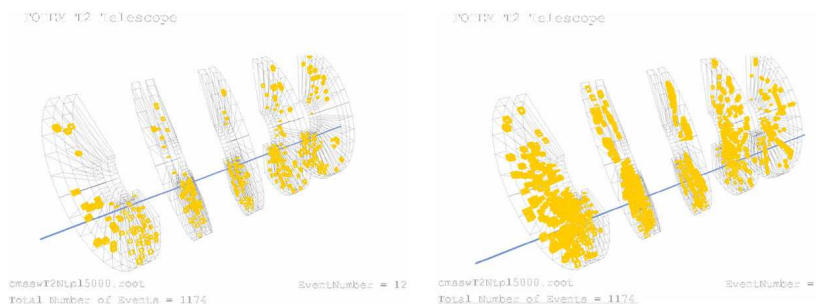


Figure 5.53: Two different events showing a large number of hits per detector. Only pads are shown for clearness.

In this ROOT-based real time event display, histograms can be easily added to improve the monitoring capabilities of the visualizer. In fig. 5.54 two different types of histograms have been added as

example: a cumulative plot of the number of hits (pads and strips) per detector and a cumulative plot of the number of hit detectors per event. The first one is useful to monitor the status of the 40 Triple GEMs installed in the T2 telescope. A wrong noise level or efficiency of one detector can be immediately seen looking to its counts in the histogram. The second plot is the same as the one used for the initial latency scan and it can be used for the same purpose and for monitoring the efficiency of each quarter. In the shown visualizer panel, the quarters have been drawn together in the full telescope and independently. This allows the possibility of plotting different quantities in the various windows. In the single quarter for instance, all the hits (pads and/or strips) could be shown, while in the full telescope only tracks and hits whit strip clusters overlaying the pad ones.

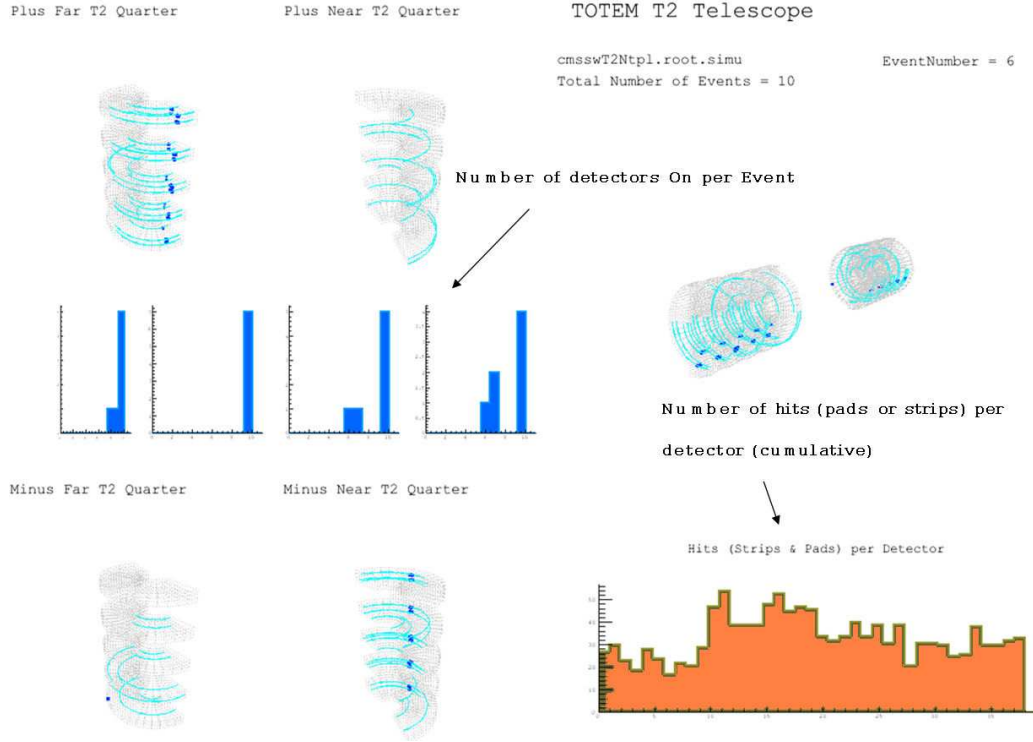


Figure 5.54: Example of integration of histograms in the Visualizer. Different quarter view can be done to show different objects (hits, tracks,...)

Conclusions

At the beginning of my Ph.D. program in the TOTEM experiment, the first prototypes of Triple GEM detectors for the T2 telescope were still under test, while the readout chip was still to come and a major part of the front end electronics had to be designed.

Today (March 2010) all the four quarters of the T2 telescope, forty Triple GEM detectors fully equipped with their readout, controls and services are installed in the experimental area of the Interaction Point 5 (IP5) of the LHC and they are successfully taking data.

A lot of time and resources were necessary in order to realize and test all the components of the complex system needed to have the telescope completely operative. Even if each part still requires further improvements, the current overall T2 performance is very much in good agreement with the requirements, and the first data taking has been successful. It was extremely important to be ready for the first collider runs because these runs have been very useful from several different points of view. First of all it was the first time that the complete data and trigger readout chains were tested for particles' detection in the experimental area. The horizontal orientation of our small telescope indeed didn't allow any cosmic ray run before LHC collisions. Moreover a preliminary analysis has been possible and a first test of the detectors alignment and background rejection have been done.

The TOTEM requirements on the pp beams to perform the Total Cross Section measurement are particular, complementary but also different (small beam divergence and low luminosity) from those necessary for the physics programs of the other experiments. In the preliminary phase of the LHC it is still hard to plan when we would have the ideal condition to achieve our measurement. The first period of the LHC restart, with beams that are not suitable for the Total Cross Section measurement, will be extremely important for us and it will be used for the commissioning of the experimental apparatus. Systematic and exhaustive tests will be performed in order to provide an accurate characterization of the detectors and front end status. A complete optimization of the full T2 system will result from these measurements and, in the meanwhile, the simulation will be tuned and validated together with the developed analysis.

My Ph.D. thesis work can be divided in three main phases: activities on the triple GEM detector, on the integration of the final electronics and on the final installation in the experimental area at the LHC. During the first period, basic tests have been performed on the detectors with standard electronics. The results have been compared with the ones of other experiment (COMPASS, LHCb) and a good agreement was found. In addition, a fine simulation has been written to understand their characteristics for further optimization and the results have been used to provide information for the digitization that has been developed.

After that, we started working with our final electronics. The integration between the digital VFAT2 chip and the detector was the first issue. Since the beginning we discovered an injection of noise correlated to the level of the internal digital activity of the chip. An activity induced by input signals that overcome the programmable internal threshold of the chip. The cross talk between channels, due to the coupling in the readout plane, propagates this noise that increases with a positive feedback mechanism and affects the whole detector. Being the inter-electrodes coupling unavoidable, this problem has been faced off by working on the chip grounding and on the shielding of the detector. The threshold level finally reached was compatible with the efficiency requirement (at least 80%) of the single detector.

Once the noise was sufficiently under control, test beams and cosmic ray tests have been performed. Among others, the most critical requirements for the T2 telescope were time and spatial resolutions and the track reconstruction efficiency. The time resolution of the integrated system Triple GEM-VFAT2 chip has been measured during a test beam run. It was in the range of $15 - 20\text{ ns}$ depending on the particular configuration used. The data stored in one memory cell of the chip has been fixed therefore to periods of

75ns – 100ns in order to guarantee an acceptable efficiency. This is compatible with the occupancy rate expected during the LHC runs. Despite this, we modified the internal field configuration of the detector to improve a little bit the obtained resolution. Other improvements however, if needed, could be done by changing the gas mixture. The possibility to add CF_4 has been foreseen in the experimental area.

The spatial resolution has been checked during test beam and cosmic ray data taking. The estimate has been done by looking at the distribution of radial and azimuth coordinate residuals between reconstructed hits and their expected positions as obtained by tracks' fitting. The radial resolution of about $100\mu m$ and the angular of half degree are close to what one would expect for a digital readout (pitch divided by square root of 12).

The track reconstruction efficiency of the full telescope, measured during the cosmic ray data taking, was about 99%, well within the expectations. This result, thanks to the redundancy of aligned detectors in the telescope, has been obtained without reaching the efficiency plateau of each triple GEM. Even if we are already in agreement with the requirement, a single detector optimization will be done during the commissioning, working on the programmed threshold and latency of the VFAT2 and on the triple GEM internal fields.

The previous phases have been followed by the final installation in the Interaction Point 5, completed by May 2009. Each quarter, fully equipped, has been tested in the TOTEM test beam area in Preveessin and then moved to the IP5 experimental area at the LHC, where we previously installed and tested all the needed services. During the LHC run of December 2009, two quarters were fully operational and data taking during the first pp collision has been performed. Now, for the 2010 LHC restart all four quarters are ready and fully operative.

All these works, done in the TOTEM experiment, have to be considered very precious taking into account that a big part of our system has been realized with prototypes or custom devices. Starting from the TOTEM Triple GEM detector or from the VFAT2 chip realized for Totem up to the control and readout system components developed for the CMS experiment. The experience with the triple GEM and the chip is surely the more promising resource for future activities.

The TOTEM collaboration, as an example, has decided to take the advantages of the radiation hardness of the triple GEM detectors in the upgrade of the T1 telescope, that could be replaced with Large Area Triple GEMs. The degrees of freedom in the shape, in the covered area, in the internal structure, in the gas choice and in the type of readout plane, make this kind of detector extremely versatile for any application where particles (photons) detection is needed.

The VFAT2 chip itself could be moreover the subject of future research lines. When analog information of the signal are not needed, this chip is an interesting choice for its readout rate (100kHz) and trigger capability. TOTEM results have shown that, even if optimized for silicon detectors, it can also be used with gaseous detectors, where signals are different (bigger but slower) and where the readout electrodes have bigger capacitance. Obviously, while the performances are excellent for silicon detectors, in the gaseous ones the integration is a little bit more complicated (problems have been found for both T2 Triple GEMs and T1 CSCs). Nevertheless the job done by the T1 and T2 groups has allowed to satisfy the initial requirements, justifying the choice of having a common chip for the three different sub-detectors of TOTEM (a choice that has greatly simplified the amount of work needed for the electronics development and data integration). Tests that we have done in other detectors, based on triple GEM technology, with the same chip have shown however strong dependencies on the readout pattern type. In certain cases, the cross talk among the various channels through the capacitive coupling in the readout plane, affects the level of noise in a not acceptable way. An upgrade of the chip, with changes in the input stage (coupling and analog stage parameters), would be then desirable to reduce or eliminate this kind of problems, providing a more reliable chip also for gaseous detectors.

Acknowledgments

I offer my regards and blessings first of all to the reader that is directly jumped to the page of the acknowledgments without having given a glance to the rest of the thesis before. I would have done the same, since this page is perhaps the only extract of humanity that can be found in these pages. My hesitation to express feelings, will unfortunately make dissatisfied such reader. Premised this, I now turn me to all those that have worked with me during the completion of the project. Your names, even if not expressly written, are privately guarded in the sincere gratitude that I feel toward each of you. I would like to offer a particular and explicit blessing to all the problems that we have met in these years. Without them, I would have learned very less and I would not perhaps have amused me so much.

Appendix Chapter 1

.1 Total Cross Section

.1.1 From the Optical Theorem to the Luminosity-Independent Total Cross Section Measurement [14]

In eq.6 is shown the Optical Theorem, derived using *Classical-Conservation of Energy* or *Q.M.-Conservation of Probability*:

$$\sigma_{TOT} = \frac{4\pi}{k} \Im F(t=0) \quad (6)$$

where $F(t=0)$ is the forward scattering amplitude and k is the particle momentum in the cms. The wave amplitude at great distance from the interaction point can be expressed as:

$$\psi(\mathbf{r}) \approx e^{ikz} + F(\theta) \frac{e^{ikr}}{r} \quad (7)$$

For large value of z and small angles:

$$r = \sqrt{x^2 + y^2 + z^2} \approx z + \frac{x^2 + y^2}{2z} \quad (8)$$

The Intensity is:

$$|\psi|^2 = \left| e^{ikz} + \frac{F(\theta)}{r} e^{ikz} e^{ik \frac{x^2+y^2}{2z}} \right|^2 = 1 + \frac{F(\theta)}{z} e^{ik \frac{x^2+y^2}{2z}} + \frac{F^*(\theta)}{z} e^{-ik \frac{x^2+y^2}{2z}} + \frac{|F(\theta)|^2}{z^2} \quad (9)$$

Dropping the term $\frac{1}{z^2}$ and considering that $C + C^* = 2\Re C$:

$$|\psi|^2 \approx 1 + 2\Re \frac{F(\theta)}{z} e^{ik \frac{x^2+y^2}{2z}} \quad (10)$$

In the approximation $F(\theta) = F(0)$:

$$\int da |\psi|^2 \approx A + 2\Re \frac{F(0)}{z} \int_{-\infty}^{\infty} dx e^{ik \frac{x^2}{2z}} \int_{-\infty}^{\infty} dy e^{ik \frac{y^2}{2z}} \quad (11)$$

The exponential can be treated as Gaussian and so:

$$\int da |\psi|^2 \approx A + 2\Re \frac{F(0)}{z} \frac{2z i \pi}{k} = A - \frac{4\pi}{k} \Im F(0) \quad (12)$$

Which is just the wave intensity if none is scattered, lessened by an amount:

$$\sigma_{TOT} = \frac{4\pi}{k} \Im F(t=0) \quad (13)$$

The differential cross section of the elastic scattering in the center of mass at zero degree can be written as:

$$\left(\frac{d\sigma_{el}}{d\Omega} \right)_{\theta=0} = |F(0)|^2 = (\Im F(0))^2 + (\Re F(0))^2 \quad (14)$$

where $\Im F(0)$ and $\Re F(0)$ are the imaginary and the real part of the forward elastic scattering amplitude. The ratio of these amplitudes is defined as:

$$\rho = \frac{\Re F(0)}{\Im F(0)} \quad (15)$$

Using the optical theorem expressed in eq.6 and the definition of ρ as in eq.15, we can rewrite eq.14:

$$\left(\frac{d\sigma_{el}}{d\Omega}\right)_{\theta=0} = (1 + \rho^2)(\Im F(0))^2 = (1 + \rho^2)\left(\frac{p\sigma_{TOT}}{4\pi}\right)^2 \quad (16)$$

where p is the momentum in the center of mass.

It is possible to describe the differential cross section in terms of the transfer three-momentum t :

$$\left(\frac{d\sigma_{el}}{dt}\right)_{t=0} = \frac{\pi}{p^2} \left(\frac{d\sigma_{el}}{d\Omega}\right)_{\theta=0} = (1 + \rho^2) \frac{\sigma_{TOT}^2}{16\pi} \quad (17)$$

In colliding experiments the cross section σ of a process is related to its event rate N through a process-independent quantity L , called luminosity that depends exclusively on beam parameters. We can explicit this relation between the σ_{TOT} and the elastic and inelastic processes observed rates:

$$N_{el} + N_{inel} = L(\sigma_{el} + \sigma_{inel}) = L\sigma_{TOT} \quad (18)$$

From the first part of eq.18 it is possible to write:

$$\left(\frac{N_{el}}{dt}\right)_{t=0} = L \left(\frac{d\sigma_{el}}{dt}\right)_{t=0} \quad (19)$$

Combining eq.17, eq.18 and eq.19 we obtain:

$$\left(\frac{dN_{el}}{dt}\right)_{t=0} = (1 + \rho^2)\sigma_{TOT} \left(\frac{N_{el} + N_{inel}}{16\pi}\right) \quad (20)$$

The total cross section σ_{TOT} can be written as:

$$\sigma_{TOT} = \frac{16}{(1 + \rho^2)} \frac{(dN_{el}/dt)_{t=0}}{N_{el} + N_{inel}} \quad (21)$$

1.2 Dispersion Relations [10]

According to Söding [15], for high energy pp and $p\bar{p}$ scattering, the dispersion relations (once subtracted) can be written as:

$$\rho_{p,\bar{p}}(E)\sigma_{p,\bar{p}}(E) = \frac{A}{p} + \frac{E}{\pi p} \int_m^\infty dE' p' \left[\frac{\sigma_{p,\bar{p}}(E')}{E'(E' - E)} - \frac{\sigma_{\bar{p},p}(E')}{E'(E' + E)} \right] \quad (22)$$

where E and p are the laboratory energy and momentum of the incoming particle and A is the subtraction constant. The parameter ρ is defined as the ratio of the real to imaginary part of the forward amplitude, $\rho = \Re F(t=0)/\Im F(t=0)$. Using the even and odd signature amplitude:

$$F_+ = (F_{p\bar{p}} + F_{pp})/2 \quad F_- = (F_{p\bar{p}} - F_{pp})/2 \quad (23)$$

we can define the quantities σ_+, σ_- and ρ_+, ρ_- . If we assume that the odd-signature amplitude becomes negligible at high energy, for the odd signature we will have:

$$\rho_+(E)\sigma_+(E) = \frac{2E}{\pi} \int_m^\infty dE' \frac{\sigma_+(E')}{E'^2(E' - E)^2}. \quad (24)$$

If the dependence of the total cross section on energy is smooth, as in high energy regime, eq.24 can be expressed in a simplified form, the *Derivative Dispersion Relations*:

$$\rho_+ = \frac{1}{\sigma_+} \tan\left(\frac{\pi}{2}\right) \frac{d\sigma_+}{d \log s} \approx \frac{\pi}{\sigma_+} \frac{d\sigma_+}{d \log s} \quad (25)$$

.1.3 Accelerator physics of colliders [27]

This part has been extracted from the "Accelerator physics of colliders" section in [27].

Luminosity

The event rate R is proportional to the interaction cross section σ_{int} . The factor of proportionality is the luminosity L .

$$\text{Event Rate} \quad R = L\sigma_{int} \quad (26)$$

For two colliding bunches with n_1 and n_2 particles, with a collision frequency f and with a Gaussian transverse profile in the interaction point characterized by σ_x (horizontal) and σ_y (vertical), the luminosity can be approximated with:

$$\text{Luminosity} \quad L = f \frac{n_1 n_2}{4\pi\sigma_x\sigma_y} \quad (27)$$

Locally, the beam size σ_i can be expressed in terms of:

- the transverse emittance ϵ . It is a beam quality concept reflecting the process of bunch preparation, extending all the way back to the source for hadrons and, in the case of electrons, mostly dependent on synchrotron radiation.
- the amplitude function β . It is a beam optics quantity and is determined by the accelerator magnet configuration.

$$\text{Emittance} \quad \epsilon = \pi\sigma^2/\beta \quad (28)$$

Using the value of the amplitude function β in the interaction point, identified with β^* , the luminosity in eq. 27 can be written as:

$$\text{Luminosity} \quad L = f \frac{n_1 n_2}{4\sqrt{\epsilon_x \beta_x^* \epsilon_y \beta_y^*}} \quad (29)$$

Betatron oscillation

The betatron oscillation is the motion of a particle that undergoes oscillation with respect to the designed trajectory, for the action of the alternating gradient focussing of the quadrupole magnetic fields. In one plane this motion can be written as:

$$x(s) = A\sqrt{\beta(s)}\cos(\phi(s) + \delta) \quad (30)$$

where A and δ are constants of integration and the phase advances according to $d\phi/ds = 1/\beta$. The dimension of A is the square root of length, reflecting the fact that the oscillation amplitude is modulated by the square root of the amplitude function. In addition to describing the envelope of the oscillation, β also plays the role of an "instantaneous" λ . The wavelength of a betatron oscillation may be some tens of meters, and so typically values of the amplitude function are of the order of meters rather than on the order of the beam size. The beam optics arrangement generally has some periodicity and the amplitude function is chosen to reflect that periodicity.

.1.4 Elastic Scattering: Optics Condition [5]

The trajectory of a particle at nominal momentum through the accelerator lattice can be described as [5]:

$$\begin{pmatrix} u(s) \\ u'(s) \end{pmatrix} = T_u(s) \begin{pmatrix} u^*(s) \\ u'^*(s) \end{pmatrix} \quad (31)$$

where $u(s)$ is one of the two transverse coordinates $x(s)$ and $y(s)$, $u'(s)$ is the projected scattering angle, $T(s)$ is the transfer matrix between the IP and the detector location and u^* and u'^* are the coordinate and angle at the IP. From eq.31, the angle at the IP, that is the physically relevant quantity, can be written as:

$$u'^*(s) = \frac{u(s) - T_{11}u^*}{T_{12}} \quad (32)$$

. The best strategy for the design of the optics is then:

- i. $T_{11} = 0$. The transverse position of the proton at the detector will be independent from its transverse position at the collision point (parallel-to-point focusing optics).
- ii. T_{12} as large as possible. In this way, if $T_{11} = 0$, the transverse position of the proton at the detector $u(s) = T_{12}u'^*(s)$ will be larger.

In this special case, the element $T_{12} = u(s)/u'^*(s)$, with small angle $u'^*(s)$, represents the effective distance of the detector from the IP. For an optics symmetric around the IP, the two elements can be expressed as follows:

$$\begin{aligned}
 \text{Magnification} \quad T_{11} &= v_u = \sqrt{\frac{\beta_u(s)}{\beta^*}} \cos \delta\mu_u(s) \\
 \text{Effective Length} \quad T_{12} &= L_u^{eff} = \sqrt{\beta_u(s)\beta^*} \sin \delta\mu_u(s) \\
 \text{Phase Advance} \quad \delta\mu_u(s) &= \int ds \frac{1}{\beta_u(s)}
 \end{aligned} \tag{33}$$

According to these definitions, the transverse displacement ($u_x(s), u_y(s)$) of a proton at a distance s from the IP is related to its transverse origin (u_x^*, u_y^*) and its momentum vector (expressed by the horizontal and vertical scattering angles (u_x', u_y') and by $\xi = \Delta p/p$) at the IP via the above optical functions and the dispersion $D(s)$ of the machine:

$$\begin{aligned}
 u_y(s) &= v_y(s) \cdot u_y^* + L_y^{eff}(s) \cdot u_y' \\
 u_x(s) &= v_x(s) \cdot u_x^* + L_x^{eff}(s) \cdot u_x' + \xi \cdot D(s)
 \end{aligned} \tag{34}$$

The condition of parallel-to-point focusing is achieved by requiring the betatron phase advance $\delta\mu_u(s) = \pi/2$ at the detector location. This condition will maximizes also L_u^{eff} , i.e. T_{12} as previously suggested. The minimum distance of a detector from the beam is proportional to the beam size:

$$y_{min} = K\sigma_y^{beam} = K\sqrt{\epsilon\beta_y(s)} \tag{35}$$

where ϵ is the transverse beam emittance and K is around $10 \div 15$. Combining this with Eqs.32-35 and assuming the parallel-to-point focusing condition, the smallest detectable angle is:

$$y_{min}'^* = \theta_{y_{min}}^* = K\sqrt{\epsilon/\beta_y^*} \tag{36}$$

Assuming the nominal value for ϵ, β^* has to be larger than 1000 m, if scattering angles of a few μrad are to be detected. The beam divergence $\sigma_\Theta^* = \sqrt{\epsilon/\beta^*}$ is then K times smaller than the minimum scattering angle.

Bibliography

- [1] Large Hadron Collider in the LEP Tunnel, Vol.I & II, Proceeding of the ECFA-CERN Workshop. 5 September 1984. ECFA 84/85, CERN 84-10 v1 & v2.
- [2] Total Cross Section, Elastic Measurement and Diffraction Dissociation at LHC. Expression of Interest, 2 March 1992.
- [3] TOTEM. Total Cross Section, Elastic Measurement and Diffraction Dissociation at LHC. Letter of Intent, CERN\ LHCC 97-49, LHCC/I 11, 15 August 1997.
- [4] Total Cross Section, Elastic Measurement and Diffraction Dissociation at LHC. Technical Proposal, CERN\ LHCC 99-7, LHCC/ P5, 15 March 1999.
- [5] TOTEM. Total Cross Section, Elastic Measurement and Diffraction Dissociation at the Large Hadron Collider at CERN. Technical Design Report, CERN-LHCC 2004-002, TOTEM-TDR-001, 7 January 2004.
- [6] TOTEM. Total Cross Section, Elastic Measurement and Diffraction Dissociation at the Large Hadron Collider at CERN. Technical Design Report - Addendum, CERN-LHCC 2004-020, TOTEM-TDR-001-Add1, 18 June 2004. addendum CERN-LHCC-2004-020.
- [7] The TOTEM Experiment at the CERN Large Hadron Collider. The TOTEM Collaboration, G Anelli et al 2008, JINST 3 S08007.
- [8] C.Augier et al., Phys.Lett. B315 (1993) 503.
- [9] The LHC study group, LHC Conceptual Design, CERN/AC/95-05 (LHC), p5-6, p43, 1995.
- [10] G.Matthie, Proton and antiproton cross section at high energies, Rep. Prog. Phys. 57 (1994) 743-790.
- [11] Khuri N N and Kinoshita T, Phys. Rev. B 137 (1965) 720.
- [12] Khuri N N On Testing Local QFT at LHC, RU94-5-B, hep-th/9405406.
- [13] J.Wess, International Workshop on Mathematical Physics, Kiev, May 1997.
- [14] Cervelli, PHD Lectures Notes. Department of Physics, Siena 2006.
- [15] Söding P., Phys. Lett. 8 (1964) 285.
- [16] H.Cheng and T.T.Wu, Phys.Rev.Lett. 24 (1970) 1456 and "Expanding Protons - Scattering at High Energies", Cambridge 1987.
- [17] C.Bourrely, J.Soffer, T.T.Wu. Eur.Phys.J. C 28 (2003) 97-105.
- [18] P.Gauron, B.Nicolescu, O.V.Selyugin. Phys.Lett. B 397 (1997) 305-310.
- [19] A.Donnachie, P.V.Landshoff. DIFFRACTION, CERN-TH, 4105/85 (1985).
- [20] P.Desgrolard, M.Giffon and E.Predazzi, Z.Phys. C63 (1994) 241.
- [21] A. Brandt et al. (UA8-Collaboration) Phys. Lett. B297 (1992) 417.

-
- [22] M. Gallinaro for the CDF collaboration, FERNILAB-CONF-03-403-E, presented at 23rd International Symposium on Multiparticle Dynamics (ISMD 2003), Cracow, Poland, 5-11 Sep. 2003.
F. Abe et al., Phys. Rev. Lett. 80 (1998) 1156.
D. Acosta et al., hep-ex/0311023
B. Abbot et al., Phys. Lett. B531 (2002) 52.
 - [23] M.L. Good and W.D. Walker. Phys. Rev. Lett. 31 (1960) 63.
 - [24] The CMS and TOTEM diffractive and forward physics working group. Prospects fo Diffractive and Forward Physics at the LHC. CERN/LHCC 2006-039/G-124, CMS Note-2007/002, TOTEM Note 06-5, 21 December 2006.
 - [25] M.M. Islam, R.J. Luddy and A.V. Prokudin, Near forward pp elastic scattering at LHC and nucleon structure, Int. J. Mod. Phys. A 21 (2006) 1
C. Bourrely, J. Soffer and T.-T. Wu, Impact-picture phenomenology for $\pi^\pm p$, $K^\pm p$ and pp , $\bar{p}p$ elastic scattering at high energies, Eur. Phys. J. C 28 (2003) 97;
V.A. Petrov, E. Predazzi and A. Prokudinl, Coulomb interference in high-energy pp and $\bar{p}p$ scattering, Eur. Phys. J. C 28 (2003) 525;
M.M. Block, E.M. Gregores, F. Halzen and G. Pancheri, Photon proton and photon photon scattering from nucleon nucleon forward amplitudes, Phys. Rev. D 60 (1999) 054024.
 - [26] C. Augier et al, Phys. Lett.B 316 (1993) 448.
 - [27] Review of Particle Physics, PDG. Journal of Physics G, Nuclear and Particle Physics 33 (July 2006) 252.
 - [28] TOTEM, Technical Design Report - Addendum. CERN-LHCC-2004-020 TOTEM-TDR-001-Add1 18 June 2004.
 - [29] F. Murtas, *Development of a gaseous detector based on Gas Electron Multiplier (GEM) Technology*, Frascati 28 November 2002.
 - [30] C. Altunbas et al., Nucl. Instr. Meth. A 490 (2002), p. 177
 - [31] P. Aspell, G. Anelli, P. Chalmet, J. Kaplon, K. Kloukinas, H. Mugnier b, W. Snoeys. VFAT2 : *A front-end system on chip providing fast trigger information, digitized data storage and formatting for the charge sensitive readout of multi- channel silicon and gas particle detectors*. 2007 TWEPP Workshop. September 3 - 7, 2007 Prague, Czech Republic.
 - [32] P.Aspell. VFAT2 *Operating Manual*, July 2006.
 - [33] P.Aspell. *Simulation of VFAT2 dynamic range*. TOTEM Electronics Meeting. CERN.
 - [34] J. Kaplon and W. Dabrowski, IEEE TRANSACTIONS ON NUCLEAR SCIENCE, VOL. 52, NO. 6, DECEMBER 2005.
 - [35] E.Noschis et al. Nucl. Instr. Meth. A 572 (2007) 378381.
 - [36] SauliLect) F.Sauli. *Principles of Operation of Multitwire Proportional and Drift Chambers*, CERN 77-09, 3 May 1997.
 - [37] G&A Engineering, Carsoli (Aquila) - Italy, www.gaengineering.com .
 - [38] *Testing procedure for the COMPASS triple-GEM detectors*, Internal Note, April 2001.
 - [39] P.Aspell, V.Avati, J.Kaspar, J.Petäjäjärvi. VFAT2 - *Operational Routines & Testing Procedures.*, TOTEM Internal Note 07-03, July 2007.
 - [40] S.Lami *First Test Results for the TOTEM T2 Telescope* ,IEEE 2009, N44-4.
 - [41] G.Croci *October RD51 Test Beam: CERN GDD and CMS preliminary data analysis* ,4th RD51 Collaboration Meeting, CERN.
 - [42] M.Berretti (G.Latino, E.Oliveri) *T2 Cosmic Ray Data Analysis* ,Sw group Meeting, CERN,15/05/2009.

-
- [43] M.Berretti (G.Latino, E.Oliveri) *T2 Cosmic Ray Data Analysis and Detector Simulation* ,TOTEM Collaboration Meeting, CERN,10/06/2009.
 - [44] M.Berretti et al. 2009 LHC run Data Analysis TOTEM Internal NOTE (in progress).
 - [45] K. Osterberg, *Data Analysis and Physics*, TOTEM referees Meeting, CERN, 16/02/2010.
 - [46] N. Turini, *T2 and Trigger Status*, TOTEM referees Meeting, CERN, 16/02/2010.
 - [47] K. Eggert, *First results from the LHC Runs 2009*, CERN Council, December 2009.