



Etude de la production resonante de graviton de Kaluza-Klein dans ses desintegrations en paires de muons dans le modele de Ramdall-Sundrum aupres de l'experience D0 au Tevatron

Nadia Lahrichi

► To cite this version:

Nadia Lahrichi. Etude de la production resonante de graviton de Kaluza-Klein dans ses desintegrations en paires de muons dans le modele de Ramdall-Sundrum aupres de l'experience D0 au Tevatron. Physique des Hautes Energies - Expérience [hep-ex]. Université Pierre et Marie Curie - Paris VI, 2004. Français. NNT : . tel-00007869

HAL Id: tel-00007869

<https://theses.hal.science/tel-00007869>

Submitted on 30 Dec 2004

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.



**THESE DE DOCTORAT DE L'UNIVERSITE PARIS 6
PIERRE ET MARIE CURIE**

Spécialité: Champs, Particules, Matières

présentée par

Nadia LAHRICHI

pour obtenir le grade de

DOCTEUR DE L'UNIVERSITE PARIS 6

Sujet:

**Etude de la production de graviton de Kaluza-Klein
dans ses désintégrations en paires de muons dans le modèle
de Randall-Sundrum auprès de l'expérience DØ au Tevatron**

soutenue le 2 juillet 2004

devant le jury composé de:

MM.	P. Brax	rapporteur
	J. Dumarchez	rapporteur
	P. Billoir	président
	J.-F. Grivaz	
	M. Besançon	
	D. Vilanova	

Table des matières

1	Les dimensions supplémentaires	9
1.1	Le Modèle Standard	11
1.1.1	Le mécanisme de Higgs	12
1.1.2	Les limites du modèle standard	13
1.2	Extensions et alternatives au modèle standard	15
1.2.1	Les particules composites	15
1.2.2	Les leptiquarks	16
1.2.3	La supersymétrie	16
1.2.4	La supergravité	16
1.2.5	Les théories des cordes	16
1.3	Les dimensions supplémentaires	20
1.3.1	Phénoménologie des dimensions supplémentaires	20
1.3.2	Des dimensions supplémentaires de l'ordre du TeV	21
1.3.3	Une unique dimension supplémentaire à l'échelle de Planck	23
1.4	Le graviton de Kaluza-Klein dans le modèle RS1	24
1.5	La recherche de gravitons auprès des collisionneurs	25
2	Le graviton de Kaluza-Klein	29
2.1	Règles de Feynman	31
2.1.1	Le propagateur du graviton	31
2.1.2	Couplages du graviton aux champs du modèle standard	32
2.1.2.1	Couplage aux fermions	33
2.1.2.2	Couplage aux bosons de jauge	34
2.1.2.3	Largeur de désintégration du graviton de Kaluza-Klein	34
2.2	Etude du signal attendu	37
2.2.1	Calcul de l'amplitude de probabilité	39
2.2.2	Section efficace $p\bar{p} \rightarrow G \rightarrow \mu\mu$	41
2.2.2.1	Section efficace différentielle	41
2.2.2.2	Les densités de partons	42
2.2.3	Génération aléatoire d'événements	43
2.2.4	Résultat de la génération	46
2.3	Terme d'interférence dans le processus de fusion des quarks	49
2.4	Influence du choix des fonctions de distribution des partons	52
2.5	Conclusion	53

3	Le détecteur DØ	57
3.1	Le système d'accélération	59
3.1.1	La création du faisceau de protons	59
3.1.2	La création du faisceau d'antiprotons	61
3.1.3	La luminosité	63
3.1.3.1	La luminosité délivrée	63
3.1.3.2	La luminosité reçue	64
3.2	Le détecteur DØ	65
3.2.1	Les modifications par rapport au Run I	65
3.2.2	Le repère géométrique	67
3.2.3	Les variables cinématiques	67
3.2.4	Le détecteur de traces central	68
3.2.4.1	Le détecteur de vertex SMT	68
3.2.4.2	Le détecteur de traces central CFT	70
3.2.4.3	Le solénoïde	72
3.2.4.4	Performances du détecteur de traces central	73
3.2.5	Le système calorimétrique	73
3.2.5.1	Le détecteur de pieds de gerbes	74
3.2.5.2	Le calorimètre	74
3.2.5.3	Les détecteurs inter-cryostats	77
3.2.5.4	Résolutions et performances du calorimètre	77
3.2.6	Le système à muons	78
3.2.6.1	L'aimant toroïdal	79
3.2.6.2	Les chambres à dérive	79
3.2.6.3	Les scintillateurs	81
3.3	Le traitement des données	82
3.3.1	Système de déclenchement et d'acquisition des données	82
3.3.2	Format de stockage des données	84
4	Identification et reconstruction des muons	89
4.1	Identification des muons	91
4.2	Reconstruction des muons	92
4.2.1	Dans le système à muons	92
4.2.2	Dans le détecteur de traces central	97
4.2.3	Algorithme d'appariement des traces locales et centrales	98
4.2.4	Les muons dits <i>calorimétriques</i>	101
4.2.5	Comparaison entre le Monte Carlo et les données	102
4.3	Critères de qualité des muons	103
4.4	Efficacités de reconstruction	105
5	Calcul de $\sigma(Z).\text{Br}(\mu\mu)$	111
5.1	Sélection des muons de qualité moyenne	114
5.2	Coupure sur l'impulsion transverse des muons	120
5.3	Déclenchement de l'acquisition	122
5.3.1	Déclenchement au niveau 1	123
5.3.2	Déclenchement au niveau 2	125
5.4	Appariement local-central	126

5.5	Isolation des muons	128
5.6	Sélection du lot de données analysées	132
5.7	Contamination des données en bruit de fond	134
5.7.1	Contamination en événements $b\bar{b}$	134
5.7.2	Contamination en événements $\tau^+\tau^-$	135
5.7.3	Proportion d'événements Z purs et Drell-Yan	136
5.7.4	Contamination en événements ZZ, ZW et WW	136
5.8	Section efficace de production de Z	137
6	Contraintes sur le graviton de Kaluza-Klein	143
6.1	Méthode du rapport de vraisemblance	145
6.2	Calcul des limites à 95% de niveau de confiance	146
6.2.1	Limites sur les paramètres du modèle	150
6.3	Conclusions et perspectives	150
6.3.1	Améliorations possibles de l'analyse	151
6.3.2	Canal électromagnétique	152
6.3.3	Autres études sur le modèle de Randall-Sundrum	152
	Annexe A - Calcul d'amplitudes de probabilités avec FORM	155

Et ne s'en suit-il pas que, dans cette région-ci, lorsque je vois une Figure Plane et que j'infère un Solide, je contemple en réalité une Quatrième Dimension inconnue de moi, autre que la couleur, mais qui existe bien qu'elle soit infinitésimale et ne puisse être mesurée?

[...]

Parvenus dans cette région bénie des Quatre Dimensions, hésiterons-nous au seuil de la Cinquième, sans oser y entrer? Ah, non. Décidons plutôt que notre ambition s'élèvera encore à mesure de notre ascension corporelle. Alors, cédant à notre assaut intellectuel, les portes de la Sixième Dimension s'ouvriront toutes grandes; puis ce sera au tour de la Septième, de la Huitième ...

Edwin A. ABBOTT, FLATLAND.

Introduction

Ce travail de thèse a pour but l'étude d'un signal physique qui mettrait en évidence l'existence de dimensions supplémentaires. Les modèles de dimensions supplémentaires sont nombreux mais deux des modèles phénoménologiques que l'on peut mettre à l'épreuve auprès de nos accélérateurs actuels sont le modèle d'Arkani-Hamed, Dimopoulos et Dvali (ADD) qui décrit de deux à sept dimensions supplémentaires à l'échelle du millimètre et le modèle de Randall et Sundrum qui décrit une seule dimension supplémentaire de petite taille.

Auprès du TeVatron, dans la collaboration DØ au sein de laquelle ce travail a été effectué, le modèle ADD de grandes dimensions supplémentaires a déjà été testé et contraint. Mais le travail développé dans cette thèse est le premier à s'intéresser au modèle de Randall-Sundrum. Dans ce modèle, seul le graviton peut se propager dans la dimension supplémentaire. Celle-ci étant compactifiée, elle est à l'origine des états de Kaluza-Klein qui apparaissent massifs dans nos quatre dimensions spatio-temporelles habituelles. Les premiers états excités ont des masses de l'ordre de quelques centaines de GeV ou de quelques TeV.

Le signal étudié est la production résonante de graviton de Kaluza-Klein. Nous recherchons le premier état excité de ce graviton à travers ses produits de désintégration. L'énergie disponible dans le centre de masse des collisions hadroniques que nous étudions est de 1960 GeV et par conséquent nous nous limiterons à la recherche du premier état excité du graviton à des masses comprises entre 100 GeV et 1 TeV. Le graviton peut se coupler à tous les champs du modèle standard et peut donc se désintégrer en une paire de particules. Le canal leptonique est le canal le plus propre pour rechercher des signaux exotiques étant donné qu'il y a peu de bruit de fond non hadronique dans les collisions hadroniques. Afin de bénéficier de l'expertise dans la reconstruction des muons présente dans le groupe DØ-Saclay, le canal de désintégration étudié est le canal muonique. Le but est donc de reconstruire une paire de muons ayant une très grande masse invariante.

Dans le premier chapitre, nous posons le cadre général des dimensions supplémentaires dans le contexte actuel des connaissances sur la physique des particules. Nous expliquons un peu plus en détail le modèle de Randall-Sundrum que nous étudions.

Dans le second chapitre, nous abordons l'aspect phénoménologique de ce modèle en explicitant les couplages et largeurs de désintégration ainsi que tous les termes qui contribuent au signal. Nous expliquons plus en détail comment nous avons dû étudier de près les programmes de génération du signal pour obtenir une description fiable du signal attendu.

Dans le troisième chapitre, nous décrivons le cadre expérimental ainsi que le détecteur qui nous fournit les données que nous analysons dans les chapitres qui suivent.

Le quatrième chapitre a pour but de comprendre la reconstruction des muons dans le

détecteur aussi bien qualitativement que quantitativement.

Etant donné que la reconstruction de la masse invariante des muons met en évidence le boson Z , nous effectuons une mesure de sa section efficace dans le chapitre 5, et nous en profitons pour relever les efficacités de sélection du signal sur des échantillons de Monte-Carlo simulant la production de bosons Z ainsi que la production de graviton de Kaluza-Klein pour différentes masses de graviton.

Le sixième et dernier chapitre est consacré à l'interprétation des résultats finals en termes de limites sur les paramètres du modèle de Randall-Sundrum, à savoir la masse du premier état excité du graviton et la valeur du couplage de ce dernier aux champs du modèle standard.

Chapitre 1

Les dimensions supplémentaires

Introduction

Au début du XXe siècle, on ne connaissait que la force gravitationnelle et la force électromagnétique à travers les lois de Newton et de Maxwell, mais très tôt, les tentatives d'unification de ces deux forces se sont multipliées. En effet, dès 1914 Nordström [1] publia ses travaux décrivant une théorie unifiée de l'électromagnétisme et de la gravitation. D'autres l'ont suivi, mais ceux qui ont laissé leur empreinte dans l'histoire de la physique des particules sont Théodor Kaluza qui publia ses travaux en 1921 [2] et montra que la dimension supplémentaire qu'il avait introduite devait être une dimension spatiale pour sauvegarder la positivité de l'énergie et enroulée (compactifiée) pour retrouver les équations de Newton ; puis Oskar Klein [3] qui, en 1926, publia ses travaux qui développent et approfondissent les travaux de Kaluza, notamment concernant la taille de la dimension supplémentaire. Ce sont les travaux de T.Kaluza complétés par ceux de O. Klein qui ont conduit à la théorie connue sous le nom de *théorie de Kaluza-Klein*.

Mais depuis la découverte des forces faible et forte et l'émergence des théories de jauge, les dimensions supplémentaires et la force gravitationnelle furent reléguées en arrière-plan, notamment à cause des échecs pour décrire la gravitation avec des théories de jauge, laissant place au modèle le plus puissant sur le plan prédictif et descriptif de la physique des particules, le *modèle standard*.

Dans ce chapitre nous allons décrire brièvement le modèle standard en nous intéressant notamment à ses limites. Nous mentionnerons alors les modèles au-delà du modèle standard en nous attardant quelque peu sur les théories des cordes car elles furent à l'origine de la réintroduction des dimensions supplémentaires dans les modèles alternatifs au modèle standard. Nous parlerons enfin de deux des modèles phénoménologiques actuellement mis à l'épreuve auprès de nos accélérateurs en insistant sur le modèle de Randall-Sundrum.

1.1 Le Modèle Standard

A basse énergie, le modèle standard décrit avec une grande précision et un formalisme simple les trois forces électromagnétique, faible et forte. C'est une théorie quantique des champs renormalisable et localement invariante sous le groupe de jauge $U(1)_Y \times SU(2)_L \times SU(3)_C$.

Le modèle standard se base sur des principes fondateurs dont entre autres :

- l'existence de *particules élémentaires fermioniques ponctuelles* qui interagissent entre elles par l'échange de *particules élémentaires bosoniques ponctuelles* ;
- les bosons sont les quanta de champ de jauge correspondant à une symétrie locale. Ils véhiculent la force dont ils sont la matérialisation ;
- l'existence d'un *champ scalaire fondamental* appelé *champ de Higgs* qui brise la symétrie électrofaible $SU(2)_L \times U(1)_Y$ et donne les masses aux particules grâce à leurs interactions.

L'interaction forte est véhiculée par huit gluons de masse nulle (correspondant aux huit générateurs du groupe $SU(3)_C$). Les interactions électromagnétique et faible sont unifiées dans la théorie de Weinberg-Salam-Glashow [4] sous le groupe de jauge $U(1)_Y \times SU(2)_L$.

Cette unification est réalisée à haute énergie mais il subsiste à basse énergie un couplage différent pour chacune de ces interactions. Les quatre champs de la force électrofaible sont : le photon non massif qui véhicule la force électromagnétique, le boson massif neutre Z^0 et les deux bosons massifs chargés W^+ et W^- qui véhiculent la force faible.

Les particules fermioniques sont réparties en deux catégories, les leptons et les quarks, contenant chacune trois familles ou trois générations.

Les fermions sont décrits par des spineurs de Dirac Ψ ayant deux projections chirales possibles : la chiralité gauche Ψ_L appartenant à un doublet d'isospin et la chiralité droite Ψ_R représentant un état singulet d'isospin. L'isospin est la charge conservée associée à la symétrie $SU(2)_L$ et est noté T . La charge conservée associée à la symétrie du groupe $U(1)_Y$ est l'hypercharge notée Y , avec $Q = T_3 + Y/2$ où Q représente la charge électrique de la particule. Les neutrinos appartiennent à des doublets gauches d'isospin, et ont une masse très petite (strictement nulle dans le modèle standard) . La seule chiralité qu'on leur a mesurée est gauche.

Les champs du modèle standard sont décrits dans le tableau 1.1.

1.1.1 Le mécanisme de Higgs

Le lagrangien du modèle standard doit être invariant sous les trois groupes de jauge $SU(3)_C$, $SU(2)_L$ et $U(1)_Y$. Il ne peut donc y avoir de terme de masse ni pour les fermions du type $m\bar{\Psi}\Psi$ ni pour les bosons de jauge pour lesquels $m^2 A^\mu A_\mu$ briserait l'invariance de jauge. Un moyen d'écrire un lagrangien décrivant des particules massives est d'introduire un champ scalaire appelé champ de Higgs qui engendre des termes de masse invariants de jauge au prix d'une brisure spontanée de la symétrie électrofaible.

En effet, le champ de Higgs est un doublet Φ_H de champs scalaires complexes où :

$$\Phi_H = \begin{pmatrix} \Phi^+ \\ \Phi^0 \end{pmatrix} \quad (1.1)$$

Ce champ introduit des termes additionnels dans le lagrangien :

- un terme cinétique

$$D_\mu \Phi_H (D^\mu \Phi_H)^\dagger \quad (1.2)$$

où D_μ est la dérivée covariante et vaut :

$$D_\mu = \partial_\mu - ig_3 \frac{\lambda_a}{2} G_\mu^a - ig_2 \frac{\sigma_b}{2} A_\mu^b - ig_1 \frac{Y}{2} B_\mu,$$

où g_1 , g_2 et g_3 sont les couplages respectifs des groupes $U(1)_Y$, $SU(2)_L$ et $SU(3)_C$, dont les champs respectifs sont notés B_μ , A_μ^b et G_μ^a . Les termes λ_a sont les matrices 3×3 de Gell-Mann, générateurs du groupe $SU(3)_C$ et σ_b sont les matrices 2×2 de Pauli, générateurs du groupe $SU(2)_L$.

- un terme décrivant le potentiel scalaire :

$$\mu^2 \Phi_H^\dagger \Phi_H + \lambda (\Phi_H^\dagger \Phi_H)^2 \quad (1.3)$$

dans lequel λ est un réel positif.

Ce potentiel possède une valeur minimale non nulle et a une infinité de solutions possibles vérifiant $|\Phi|^2 = -\mu^2/2\lambda$. Le choix d'un état parmi ce continuum d'états dégénérés brise la symétrie électrofaible. C'est une brisure spontanée de symétrie. L'absorption de 3 bosons de Goldstone (parmi les 4 degrés de liberté du champ de Higgs) par les bosons électrofaibles confère des masses aux bosons $W^\pm$ et Z par le mélange des champs et l'angle de mélange θ_W mesuré expérimentalement donne la relation entre la masse des bosons $W^\pm$ et la masse du boson Z :

$$m_W = m_Z / \cos(\theta_W) \quad (1.4)$$

Par contre, le photon conserve une masse nulle. Le quatrième degré de liberté correspond au boson de Higgs neutre massif.

- des termes de couplage aux fermions appelés aussi termes de Yukawa λ_f

Les masses acquises par les fermions sont données par :

$$m_f = \frac{1}{\sqrt{2}} \lambda_f v \quad (1.5)$$

où $v = \sqrt{-\mu^2/(2\lambda)}$ est la valeur moyenne dans le vide de la masse du boson de Higgs.

Le boson de Higgs donne très élégamment des masses aux particules en respectant les contraintes théoriques du modèle standard. Les recherches directes ont exclu une masse du boson de higgs inférieure à 115 GeV à 95% de niveau de confiance [5].

1.1.2 Les limites du modèle standard

Malgré le très bon accord entre le modèle standard et l'expérience, ce modèle souffre de lacunes et de faiblesses dont :

- le nombre de paramètres libres :
en effet, un premier point en défaveur du modèle standard est que les masses des particules ainsi que d'autres quantités sont des paramètres libres que l'on voudrait fixés par la théorie sous-jacente :
 - 9 masses de fermions : $m_e, m_\mu, m_\tau, m_u, m_d, m_c, m_s, m_b, m_t$;
 - 3 angles de mélange entre les quarks et une phase de violation de CP provenant de la matrice CKM (de Cabibbo, Kobayashi et Maskawa) ;
 - 3 couplages de jauge α_s, α_{em} et α_W (ou θ_W) ;
 - la masse d'un des bosons W ou Z ;
 - la masse du boson de Higgs ;
 - les masses des neutrinos $m_{\nu_e}, m_{\nu_\mu}, m_{\nu_\tau}$. Le modèle standard considère en toute rigueur que les masses des neutrinos sont nulles.

bosons de jauge	spin	groupe de jauge		nb de bosons
G^a	1	$SU(3)_C$		$a = 1, \dots, 8$
A^b	1	$SU(2)_L$		$b = 1, \dots, 3$
B	1	$U(1)_Y$		1
champ de matière	spin	représentation		hypercharge Y
		$SU(3)_C$	$SU(2)_L$	
$\psi_{qL} = \begin{pmatrix} u \\ d \end{pmatrix}_L$	1/2	3	2	1/3
$\psi_{qR} = u_R$	1/2	3	1	4/3
$\psi_{qR} = d_R$	1/2	3	1	- 2/3
$\psi_{lL} = \begin{pmatrix} e \\ \nu_e \end{pmatrix}_L$	1/2	1	2	-1
$\psi_{lR} = e_R$	1/2	1	1	- 2
$\Phi_H = \begin{pmatrix} \Phi^+ \\ \Phi^0 \end{pmatrix}$	0	1	2	1

TAB. 1.1 – *Particules élémentaires du modèle standard. Pour les fermions, seules les premières familles sont représentées, les deux autres familles étant décrites par les mêmes paramètres.*

- les charges fractionnaires :
on ne sait toujours pas pourquoi les quarks ont des charges électriques proportionnelles à 1/3 de la charge de l'électron ;
- la structure en famille de fermions :
rien n'explique le fait que les fermions peuvent être groupés en familles ou générations et surtout pourquoi trois générations de quarks et trois générations de leptons ayant une hiérarchie en masse telle que celle observée ;
- le problème de naturalité :
le problème de naturalité est lié aux corrections radiatives à la masse du boson de Higgs. En effet, dans le calcul des corrections à la masse du boson de Higgs, la contribution des diagrammes en boucle :

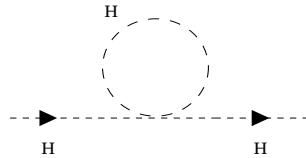


FIG. 1.1 – *Diagramme en boucle avec un boson de Higgs.*

est proportionnelle à $\int d^4k/k^2$ qui diverge quadratiquement.

Pour se débarrasser de cette divergence, on introduit un “cut-off” Λ comme borne supérieure de l’intégrale qui est la limite de validité du modèle standard ($\Lambda \sim M_{Pl}$), mais alors la masse m_H du boson de Higgs devient dépendante de Λ :

$$m_H^2 = m_{nue}^2 + \frac{\lambda \Lambda^2}{16\pi^2}. \quad (1.6)$$

Dans ce cas, pour une masse de Higgs de 100 GeV, les corrections radiatives doivent compenser la masse nue du boson de Higgs à 17 ordres de grandeur près. C'est le problème du *fine tuning*. Cet ajustement d'un des paramètres du modèle n'a rien de naturel ;

- la hiérarchie des échelles et des masses :
elle est liée à la différence entre les échelles du modèle standard et l'échelle de la gravitation. En effet, l'échelle de la gravitation donnée par la masse de Planck est bien supérieure à celle des autres forces. Le rapport entre l'échelle de Planck et l'échelle électrofaible M_{PL}/M_W est de l'ordre de $\simeq 10^{16}$!
La hiérarchie des échelles se retrouve aussi entre les masses des particules du modèle standard: la masse du lepton μ est 200 fois supérieure à celle de l'électron, la masse du quark top est environ 200 fois supérieure à celle du quark c , et aucune explication à ces différences de masses n'est apportée par le modèle standard;
- l'absence de la gravitation :
la force gravitationnelle n'est pas incluse dans le modèle standard. C'est pourtant la première force de la nature que l'homme ait décrite, mais du fait de la dimension en $1/M_{Pl}$ du couplage gravitationnel, il apparaît impossible d'élaborer une théorie de jauge renormalisable de la gravitation.

Tous ces problèmes posés par le modèle standard pourraient être dus au fait que ce modèle est une théorie effective à basse énergie d'une théorie plus générale qui n'est toujours pas établie. A l'heure actuelle, quelques modèles plus ou moins liés entre eux arrivent à résoudre quelques-uns des problèmes posés tout en gardant une cohérence interne et une consistance avec les résultats expérimentaux parfaitement établis. Nous donnons quelques exemples de modèles dans ce qui suit.

1.2 Extensions et alternatives au modèle standard

1.2.1 Les particules composites

Pour résoudre quelques uns des problèmes posés par le modèle standard, certains modèles font l'hypothèse que les particules du modèle standard supposées fondamentales ne le sont pas. Il existe trois types de modèles, chacun supposant que les fermions, les bosons de jauge ou le boson de Higgs sont composés de particules plus élémentaires [6].

Les modèles de fermions excités répondent à la question de la structure en familles des fermions. Ainsi, chaque famille est décrite comme un état excité des fermions eux-mêmes constitués de deux particules, une fermionique et une scalaire appelées "Préons". Ce modèle simplifie la "zoologie" du modèle standard et diminue le nombre de paramètres libres.

Les modèles de technicouleur permettent de résoudre le problème de naturalité en n'incluant aucun champ scalaire fondamental. Le boson de Higgs, selon ces modèles, n'est plus considéré comme une particule scalaire fondamentale mais est supposé être formé

d'un état lié de deux fermions appelés technifermions.

1.2.2 Les leptoquarks

Les modèles de leptoquarks sont issus des modèles de grande unification (GUT). Les quarks et les leptons appartiennent à un multiplet commun et l'interaction entre ces particules se fait par échange de nouveaux bosons de jauge, les leptoquarks. Une des motivations physiques de ces modèles revient à l'hypothèse que, aux énergies de la GUT, les leptons et les hadrons seraient représentés par un seul groupe de jauge brisé à l'échelle électrofaible.

1.2.3 La supersymétrie

La supersymétrie constitue une nouvelle symétrie globale de l'espace-temps [7]. La supersymétrie résout entre autres le problème de naturalité par l'introduction d'une symétrie entre les bosons et les fermions. A chaque boson du modèle standard est associé un fermion qui est son partenaire supersymétrique et inversement.

Dans ce modèle, les partenaires supersymétriques bosoniques des fermions du modèle standard contribuent aux diagrammes en boucles qui interviennent dans les corrections radiatives à la masse du boson de Higgs et ont pour effet d'annuler les contributions des boucles fermioniques ce qui permet d'éliminer les divergences quadratiques.

La supersymétrie doit être elle-même brisée afin de conférer aux partenaires supersymétriques des particules du modèle standard des masses très élevées pour expliquer leur discrétion dans les expériences menées jusqu'à nos jours.

1.2.4 La supergravité

C'est la version locale de la supersymétrie [8]. En effet, le commutateur de deux transformations supersymétriques locales est une translation dans l'espace-temps. On s'attend donc naturellement à voir apparaître la relativité générale, d'où le nom de supergravité. Dans le cas de transformations locales, on redonne au lagrangien de la supersymétrie son invariance en introduisant un nouveau champ de jauge de spin $3/2$ qui représente le gravitino, partenaire supersymétrique du graviton de spin 2.

En partant d'une description effective à basse énergie des théories des cordes on peut retrouver la supergravité.

1.2.5 Les théories des cordes

Dans le cas classique, la force gravitationnelle est décrite par la relativité générale. L'intégration de la force gravitationnelle aux autres forces en utilisant les théories de jauge pose quelques problèmes notamment en raison de la non-renormalisabilité de la relativité générale. En effet, le couplage en $1/M_{Pl}$ du graviton aux champs du modèle standard crée des divergences non renormalisables dans les corrections radiatives, contrairement au cas des champs du modèle standard dont les couplages non-dimensionnés permettent

de renormaliser les divergences. Par ailleurs, l'unification des quatre forces suppose que la force gravitationnelle soit décrite même au niveau quantique. Alors que les champs électromagnétique, faible et fort sont quantifiés dans un environnement spatio-temporel classique, pour la force gravitationnelle il s'agit de quantifier l'espace-temps lui-même. Cela rajoute une difficulté technique et même conceptuelle à la tentative d'unification de la force gravitationnelle avec les trois autres forces décrites par le modèle standard.

Les théories de cordes [12][13]

Une des pistes permettant de résoudre les problèmes du modèle standard et d'obtenir une description quantique de l'interaction gravitationnelle repose sur une description non plus en terme d'objets ponctuels comme les particules dans les théories des champs mais en terme d'objets avec une extension spatiale à savoir les cordes décrites dans les théories de cordes et les théories branaires.

En 1974 Scherk et Schwarz [9] ainsi que Yoneya [10] ont montré que la corde fermée contient toujours un état de spin 2 de masse nulle dans son spectre qu'on peut identifier au graviton. Scherk et Schwarz ont alors proposé les théories de cordes comme théories candidates à l'unification des interactions de jauge avec l'interaction gravitationnelle.

Les travaux sur la quantification de la corde relativiste bosonique ont montré [11] l'existence d'un nombre critique égal à 26 pour le nombre de dimensions de l'espace-temps afin d'éviter les anomalies à l'invariance de Lorentz. Toutefois cette corde bosonique présente l'inconvénient de ne pas inclure d'états fermioniques dans son spectre et de posséder un état fondamental tachyonique.

L'introduction de la supersymétrie permet d'obtenir une théorie de cordes sans tachyons, avec des états fermioniques dans son spectre. Ces théories sont connues sous le nom de théories de supercordes. Les arguments pour éviter les anomalies à l'invariance de Lorentz peuvent s'étendre pour montrer l'existence d'un nombre critique égal à 10 pour le nombre de dimensions de l'espace-temps dans les théories de supercordes.

Les premières théories de supercordes [12][13] connues sont d'une part les théories de supercordes de type I qui sont des théories de supercordes ouvertes et fermées et d'autre part les théories de supercordes de type II qui sont des théories de cordes fermées seulement.

On distingue deux sortes de théories de supercordes de type II suivant la chiralité opposée (théorie de type IIA) ou égale (théorie de type IIB) des deux gravitinos présents dans le spectre des états de masse nulle.

Les théories de supercordes hétérotiques sont des théories de cordes fermées dont la construction procède de la combinaison des modes gauches d'une corde bosonique fermée à 26 dimensions d'espace-temps évoquée plus haut et des modes droits d'une supercorde fermée à 10 dimensions d'espace-temps.

Pour effectuer cette combinaison, 16 des dimensions de la corde bosonique fermée à 26 dimensions d'espace-temps sont compactifiées selon une procédure dont les propriétés per-

mettent d'obtenir deux groupes de jauge à savoir $SO(32)$ et $E_8 \times E_8$. De plus, les théories de supercordes hétérotiques ont des anomalies de jauge et des anomalies gravitationnelles qui s'annulent si et seulement si les groupe de jauge sont $SO(32)$ et $E_8 \times E_8$. C'est pour cette raison qu'il existe deux théories de supercordes hétérotiques à savoir la théorie de supercordes hétérotiques dite $SO(32)$ et la théorie de supercordes hétérotiques dite $E_8 \times E_8$ [14].

La compactification de 6 des dimensions des deux théories de cordes hétérotiques, en particulier la théorie $E_8 \times E_8$, sur un certain type d'espaces compacts, a permis d'aboutir dans la limite "basse énergie" à des théories de jauge avec une supersymétrie locale $N=1$ à 4 dimensions d'espace-temps à une échelle de l'ordre de 10^{18} GeV proche de la l'échelle de Planck et avec un groupe de jauge pouvant contenir le groupe de jauge du modèle standard. Ce développement a permis l'étude d'une phénoménologie dite inspirée des cordes à partir de la seconde moitié des années 80 basée sur des théories de supergravité $N=1$ à quatre dimensions d'espace-temps avec le groupe exceptionnel E_6 , sous-groupe de E_8 comme groupe de jauge [15].

Il existe donc cinq théories de supercordes consistantes à savoir la théorie de type I, celle de type IIA et celle de type IIB, la théorie hétérotique $SO(32)$ et la théorie hétérotique $E_8 \times E_8$.

Les symétries de dualité [12][18]

En utilisant les théories effectives basse énergie correspondantes, il est possible de montrer que le régime de couplage fort ou encore le régime non perturbatif d'une théorie de supercordes de type I peut s'identifier au régime de couplage faible (le régime perturbatif) de la théorie de supercordes hétérotiques $SO(32)$, et vice versa qui sont par ailleurs perturbativement différentes. Ceci permet de définir une symétrie de dualité qui porte le nom de dualité S.

Les théories de supercordes fermées de type II à dix dimensions d'espace-temps peuvent être compactifiées en théories à 9 dimensions d'espace-temps où une dimension d'espace a été compactifiée sur un cercle de rayon R . Dans ce cas simple, les théories à 9 dimensions possèdent d'une part un spectre d'états dits de Kaluza-Klein où chaque état correspond à un mode k et d'autre part un spectre d'états dits d'enroulement où chaque état correspond à un mode m . Le nombre entier m correspondant au nombre d'enroulements de la dimension compactifiée sur le cercle de rayon R . Ce spectre est invariant par la transformation $R \leftrightarrow 2/R$ et la transformation de m en k . Cette symétrie est une autre symétrie de dualité (invariance du spectre) appelée dualité T.

Ainsi les théories de supercordes hétérotiques $SO(32)$ et $E_8 \times E_8$ sont T-duales entre elles de même que les théories de supercordes de type IIA et IIB.

Il est également possible de montrer que la théorie de supercordes de type I peut être obtenue à partir de la théorie de supercordes fermées de type IIB par une opération de projection (dite opération Ω) qui ne conserve que la somme diagonale des deux gravitinos présents dans le spectre de la théorie de type IIB.

Finalement, dans la limite des couplages forts (régime non perturbatif), la théorie de supercordes de type IIA se “décompactifie” vers une théorie à 11 dimensions d’espace-temps appelée M-théorie encore inconnue dont la supergravité à 11 dimensions d’espace-temps serait la limite “basse énergie” en théorie ponctuelle des champs. De même, dans la limite des couplages forts, la théorie de supercordes hétérotiques $E_8 \times E_8$ peut être reliée à la théorie à 11 dimensions non plus compactifiée sur un espace particulier appelé orbifold (en l’occurrence S^1/Z_2).

Les cinq théories de supercordes consistantes connues peuvent donc être reliées entre elles par des opérations de projection et des symétries de dualité dans diverses dimensions d’espace-temps.

Witten a montré en outre que ces symétries de dualité entraînent également que l’échelle caractéristique des théories de supercordes ne reste plus forcément fixée à des valeurs proches de l’échelle de la masse de Planck mais peut prendre des valeurs arbitraires. Lykken a poussé l’argument à son extrême en suggérant que cette échelle devenue arbitraire peut prendre alors des valeurs relativement petites, plus petites que des valeurs de l’ordre du TeV [16][17].

Les branes [12][18]

Chaque théorie de supercordes contient dans son spectre non seulement des états de spin 2 et des états de spin 0 comme nous l’avons évoqué plus haut mais également des tenseurs antisymétriques de masse nulle. Ces tenseurs antisymétriques permettent de définir des champs invariants par transformations de jauge généralisées particulières et donc de définir des charges conservées sous une forme spécifique particulière. Les objets naturels portant ces charges généralisées sont des objets étendus à p dimensions connus sous le nom de p -branes. De plus, étant donné que les états perturbatifs de chaque théorie de supercordes sont neutres du point de vue de ces charges généralisées, ce sont les états du régime non perturbatif de chaque théorie qui portent ces charges. Les p -branes sont donc des objets décrivant des états étendus du régime non perturbatif des théories de supercordes.

Comme nous l’avons vu plus haut, la théorie de supercordes de type I est une théorie de cordes fermées et ouvertes dans un espace-temps à 10 dimensions avec $SO(32)$ pour groupe de jauge. Dans cette théorie, les charges de jauge de $SO(32)$ sont portées par les extrémités des cordes ouvertes. Ces extrémités décrivent également une D p -brane (ou brane de Dirichlet ou D-branes), les champs de jauge sont contraints de se mouvoir sur cette D p -brane.

Les D p -branes ne sont pas des hyperplans rigides dans l’espace-temps mais sont des objets dynamiques qui admettent des fluctuations de position et de forme. Ces fluctuations sont décrites par certains états de cordes ouvertes correspondant à un champ de jauge. Les champs de jauge dynamiques vivent dans un volume d’univers défini par la D p -brane et sont associés aux configurations de cette D p -brane [18].

A cette étape du développement, un certain nombre de résultats ont facilité l’émer-

gence de modèles phénoménologiques de dimensions supplémentaires même si le lien avec une théorie fondamentale en terme de théorie de supercordes est parfois un peu lointain. Ce lien, même si rétabli d'une certaine manière dans plusieurs descriptions spécifiques (que nous n'aborderons pas ici) fait encore l'objet de travaux approfondis.

Ainsi les 5 théories de supercordes destinées à une description quantique unifiée de toutes les particules et interactions peuvent être définies de manière quantiquement consistante dans un espace-temps à 10 dimensions. Ainsi apparaît une notion de dimensions supplémentaires d'espace-temps associée à la notion d'unification de l'interaction gravitationnelle et des autres interactions connues ainsi qu'à la notion de quantification de l'interaction gravitationnelle. Les symétries de dualité permettent de relier les 5 théories de supercordes entre elles comme autant de secteurs (perturbativement) différents d'une seule et même théorie à 11-dimensions d'espace-temps encore inconnue. Une des conséquences de ces symétries de dualité est de rendre arbitraire l'échelle fondamentale des théories de supercordes qui n'est plus alors fixée à des valeurs proches de l'échelle de Planck. Puis, rapidement, la proposition a été faite que cette échelle peut prendre alors des valeurs aussi petites que le TeV ce qui entraîne la perspective d'accéder aux échelles de dimensions supplémentaires dans les expériences ainsi qu'aux effets quantiques de l'interaction gravitationnelle.

Finalement, les Dp-branes décrivant certains états étendus non-perturbatifs de la théorie de supercordes fermées de type II, sur lesquels les extrémités des cordes ouvertes de la théorie des supercordes de type I viennent s'attacher, permettent une description en terme de champs de jauge contraints de vivre dans un sous-espace à $p+1$ dimensions et d'interaction gravitationnelle existant dans l'espace-temps complet à 10 dimensions.

A partir de là, l'émergence de modèles où les champs du modèle standard sont confinés sur une brane à 4 dimensions d'espace-temps et l'interaction gravitationnelle peut vivre dans l'espace-temps complet à des échelles aussi petites que le TeV bénéficie d'un cadre conceptuel relativement naturel.

A noter cependant que la possibilité de l'existence de dimensions supplémentaires au TeV a été discutée dans le contexte des théories de supercordes dès 1990 par Ignatios Antoniadis [19] dans des travaux entrepris pour une meilleure compréhension de la brisure de la supersymétrie dans le contexte des théories de supercordes.

1.3 Les dimensions supplémentaires

1.3.1 Phénoménologie des dimensions supplémentaires

Nous décrirons brièvement dans ce qui suit les deux modèles phénoménologiques qui sont à la base des recherches expérimentales actuelles de dimensions supplémentaires auprès des accélérateurs, à savoir d'une part le modèle de *Large Extra-Dimensions* [20] publié en 1998, connu aussi sous l'appellation *modèle ADD* du nom de ses auteurs Nima Arkani-Hamed, Savas Dimopoulos, Gia R. Dvali et d'autre part le modèle contenant une

Nombre de dimensions supplémentaires	1	2	3	4	5	6	7
Rayon de compactification (en m)	10^{13}	10^{-3}	10^{-8}	10^{-11}	10^{-12}	10^{-14}	10^{-15}
Masse du premier état excité (en eV)	10^{-20}	10^{-4}	10^1	10^4	10^5	10^7	10^8

TAB. 1.2 – *Rayon de compactification et masse du premier mode de Kaluza-Klein en fonction du nombre de dimensions supplémentaires.*

seule petite dimension spatiale supplémentaire, le *modèle de RS* [21] publié en 1999 par Lisa Randall et Raman Sundrum.

La première solution apportée par ces deux modèles aux faiblesses du modèle standard est relative au problème de hiérarchie. En effet, l'introduction de dimensions supplémentaires oppose la masse de Planck perçue dans l'espace à 4 dimensions (1 temporelle et 3 spatiales) à la masse de Planck perçue dans l'espace à 4+n dimensions (1 temporelle et 3+n spatiales).

1.3.2 Des dimensions supplémentaires de l'ordre du TeV

Le modèle ADD est un modèle à plusieurs dimensions supplémentaires décrit par une géométrie plate factorisable. Il contient une brane correspondant à un sous-espace identifié à l'espace-temps que l'on connaît. Le nombre de dimensions supplémentaires de grande taille ($R \gg L_{Pl}$, la longueur de Planck $L_{Pl} \sim 10^{-34}$ m) peut aller jusqu'à 6 (ou 7 qui correspond au nombre de dimensions supplémentaires de *la théorie M*). Le tableau 1.2 montre la taille des dimensions supplémentaires et la masse du premier état massif du graviton de Kaluza-Klein en fonction du nombre de dimensions supplémentaires de grande taille considérées pour une masse de Planck dans toutes les dimensions $M_D=1$ TeV.

Cependant, la possibilité de n'avoir qu'une seule dimension supplémentaire est exclue, autrement les effets¹ auraient été observables à l'échelle du système solaire. C'est pourquoi les développements phénoménologiques et les études expérimentales se font sur des espaces contenant de 2 à 7 dimensions supplémentaires à l'échelle du TeV.

La géométrie est dite factorisable car les éléments de la restriction de la métrique à quatre dimensions ne dépendent pas des coordonnées dans les dimensions supplémentaires.

Ce modèle propose que ces dimensions soient de l'ordre du millimètre, l'échelle de Planck à 4+n dimensions étant ramenée à des énergies de l'ordre du TeV. L'expression qui relie les masses de Planck à 4 et à 4+n dimensions peut être déduite à partir de la loi de Gauss [20] :

1. Les dimensions supplémentaires apportent une correction à la loi en $1/r^2$ de la force gravitationnelle. Cet écart à la loi en $1/r^2$ n'a pas été observé pour des distances supérieures au millimètre. Or dans le modèle ADD, la correction apportée à cette loi par l'existence d'une seule dimension supplémentaire aurait été visible à l'échelle du système solaire (voir tableau 1.2).

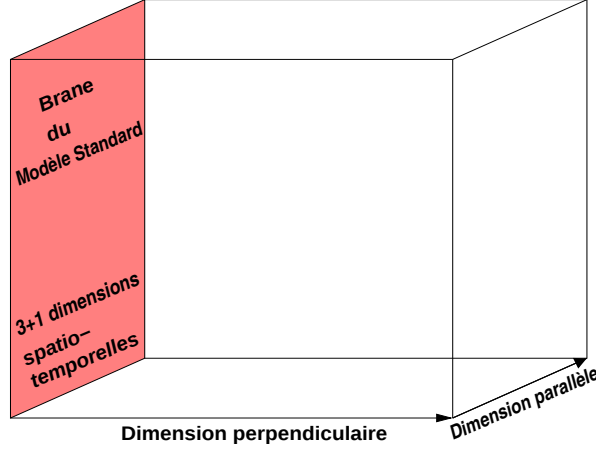


FIG. 1.2 – Représentation des dimensions supplémentaires dans le modèle ADD. La brane grisée correspond à la brane contenant les champs du modèle standard.

$$\bar{M}_{PL}^2 = M_D^{2+n} \mathcal{V}_{(ED)} \quad (1.7)$$

ou encore :

$$\bar{M}_{PL}^2 = M_D^{2+n} R^n \quad (1.8)$$

où $\mathcal{V}_{(ED)}$ est le volume des dimensions supplémentaires et R le rayon de compactification de ces dimensions inférieur au millimètre. Ainsi, dans les 4 dimensions spatio-temporelles, la masse de Planck $\bar{M}_{PL}$ ou $\bar{M}_{PL}^{(4)}$ vaut $\bar{M}_{PL} = [hc/G]^{1/2} = 1.22 \cdot 10^{19} \text{ GeV}/c^2$ tandis que la masse de Planck dans l'espace à $4+n$ dimensions, $M_{PL}^{(4+n)}$ ou M_D est de l'ordre du TeV.

Dans ce modèle, les gravitons sont les seules particules à pouvoir se propager dans les dimensions supplémentaires. Les champs du modèle standard sont confinés sur une brane. Etant donné que les dimensions supplémentaires sont compactifiées, la fonction d'onde du graviton peut être développée en série de Fourier dans les dimensions supplémentaires. Chaque terme de la série est un état de Kaluza-Klein du graviton dont la masse vaut k/R où k est le mode d'excitation du terme considéré. Deux états de Kaluza-Klein successifs sont donc espacés de $1/R$. Cela implique que le spectre des états de Kaluza-Klein est un quasi-continuum d'états excités, chacun ayant un couplage en $1/\bar{M}_{PL}$ à la matière.

Le graviton de Kaluza-Klein dans le modèle ADD est activement recherché auprès des accélérateurs. Les canaux explorés sont soit la production directe où le graviton n'est pas détecté et a pour effet de l'énergie manquante, soit la production en voie s , ce qui devrait engendrer un excès de paires particule-antiparticule par rapport aux prédictions du modèle standard. Bien que chaque état soit très supprimé du fait du très faible couplage à la matière, la section efficace totale de production de gravitons est compensée par le très grand nombre d'états finals possibles, pouvant prendre des valeurs de l'ordre du picobarn.

1.3.3 Une unique dimension supplémentaire à l'échelle de Planck

Deux modèles ont été développés par L.Randall et R.Sundrum, les modèles RS1 [21] contenant une seule dimension supplémentaire compactifiée et RS2 [22] décrivant une alternative à la compactification. Dans ces deux modèles, la métrique est une solution de l'équation d'Einstein et conserve l'invariance de Poincaré. Ils sont décrits dans une géométrie non factorisable.

Dans le modèle RS1 que nous étudions, la dimension supplémentaire est limitée par deux branes (voir figure 1.3). L'espace entre les deux branes est le *bulk*, et la constante cosmologique dans ce *bulk* est négative, d'où l'espace anti-de Sitter (AdS). La brane située à $\phi = 0$, de tension positive, contient des champs dont la masse est naturellement de l'ordre de l'échelle de Planck, c'est pourquoi on appelle cette brane la brane de Planck. La brane située à $\phi = \pi$ est de tension négative. Elle contient les champs du modèle standard et est à l'échelle du TeV.

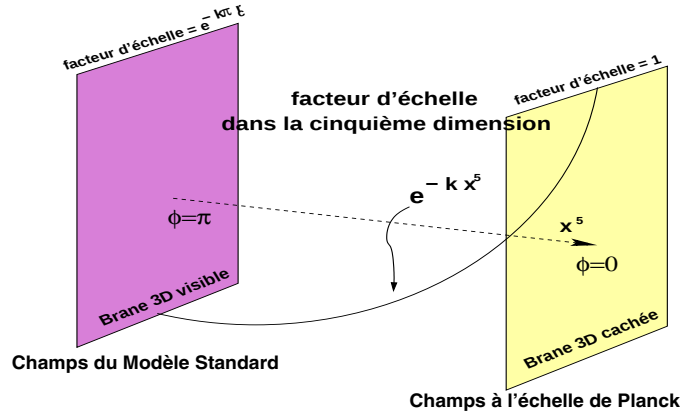


FIG. 1.3 – Représentation des dimensions supplémentaires dans le modèle RS1. L'espace entre les deux branes est le *bulk*. La constante cosmologique Λ_B dans le *bulk* est négative.

La position dans la cinquième direction est donnée par $x^5 = \phi r_c$ où r_c est le rayon de compactification de la dimension supplémentaire et $\phi \in [-\pi, \pi]^2$.

L'élément d'espace-temps s'écrit :

$$ds^2 = e^{-\phi k r_c} dx^2 - r_c^2 d\phi^2 \quad (1.9)$$

La relation entre la masse de Planck M_D dans toutes les dimensions et la masse de Planck M_{Pl} sur la brane visible est donnée par :

$$\bar{M}_{Pl}^2 = M_D^3 r_c \int_{-\pi}^{+\pi} d\phi e^{-2kr_c|\phi|} = \frac{M_D^3}{k} [1 - e^{-2kr_c\pi}] \quad (1.10)$$

2. L'introduction d'un champ scalaire dans le *bulk* permet de stabiliser le radion, champ scalaire quadridimensionnel associé à la position relative des deux branes dans le *bulk*. Cette stabilisation est connue comme le mécanisme de Goldberger et Wise [23][24]. Après stabilisation le radion devient massif mais léger et possède une phénoménologie qui ressemble à la phénoménologie du boson de Higgs [28]. La question se pose également de comprendre le comportement de la masse du radion sous l'effet des corrections radiatives. L'introduction de la supersymétrie permet d'appréhender cette question [25][26][27][28].

D'une part on peut noter que la masse de Planck fondamentale M_D est du même ordre de grandeur que la masse de Planck M_{Pl} perçue sur la brane du modèle standard, k étant lui-aussi de l'ordre de la masse de Planck M_{Pl} . D'autre part, cette relation montre qu'en fait la masse de Planck fondamentale M_D dépend peu du rayon de compactification³ car l'exponentielle est très petite ($\sim 10^{-16}$) .

1.4 Le graviton de Kaluza-Klein dans le modèle RS1

Le modèle RS1 décrit des champs localisés sur la brane du modèle standard et seul le graviton se propage dans la cinquième dimension. Sur la brane visible, les masses des particules sont à l'échelle électrofaible. Le paramètre k du modèle est de l'ordre de la masse de Planck. On ne peut demander $k \ll M_{Pl}$, autrement cela reviendrait à introduire une échelle de masse supplémentaire. La relation entre les masses des particules sur la brane visible ($\phi = \pi$) et sur la brane cachée ($\phi = 0$) est donnée par :

$$m_{(\pi)} = e^{-k \pi r_c} m_{(0)} \quad (1.11)$$

Lorsqu'on veut étudier les interactions du graviton, on doit alors intégrer sur tous les états possibles sur la cinquième dimension. La métrique prise dans une tranche Minkowskienne d'espace-temps à quatre dimensions peut être paramétrisée comme suit :

$$g_{\mu\nu} = e^{-2k\phi r_c} (\eta_{\mu\nu} + 2M_D^{-3/2} h_{\mu\nu}) \quad (1.12)$$

où $\mu, \nu = 0, 1, 2, 3$, $\eta_{\mu\nu}$ étant la métrique de Minkowski.

Le terme décrivant la gravitation est $h_{\mu\nu}(x, \phi)$. Du fait de la compactification de la dimension supplémentaire, $h_{\mu\nu}(x, \phi)$ contient un terme périodique qui ne dépend que de la cinquième dimension. On peut alors le développer en harmoniques :

$$h_{\mu\nu}(x, \phi) = \sum_{n=0}^{+\infty} h_{\mu\nu}^{(n)}(x) \frac{\chi_G^{(n)}(\phi)}{\sqrt{r_c}} \quad (1.13)$$

où les $h_{\mu\nu}^{(n)}(x)$ représentent les modes de Kaluza-Klein du graviton. Les $\chi_G^{(n)}(\phi)$ sont les fonctions d'ondes du graviton dans la cinquième dimension et ne dépendent par conséquent que de la variable ϕ .

La fonction d'onde du graviton dans le *bulk* a pour expression [29] :

$$\chi_G^{(n)}(\phi) \simeq e^{k r_c (2\phi - \pi)} \sqrt{k r_c} \frac{J_2(z_n^G)}{J_2(x_n^G)} \quad (1.14)$$

où J_q sont les fonctions de Bessel d'ordre q et $z_n^G(\phi) = m_n e^{2k r_c \phi} / k$ où m_n est la masse du mode n du graviton de Kaluza-Klein.

Les solutions $\chi_G^{(n)}(\phi)$ doivent être paires afin de permettre au mode zéro de représenter un graviton de masse nulle qui est responsable de l'interaction gravitationnelle habituelle,

3. Cette propriété fut à l'origine du deuxième papier de L. Randall et R. Sundrum présentant une alternative à la compactification en faisant tendre le rayon de compactification r_c vers l'infini.

observée dans l'espace à 4 dimensions. Les modes non nuls du graviton contribuent aussi à l'interaction gravitationnelle mais, étant très massifs, la portée de leur interaction est très courte ($\lambda = \hbar/m_n c \simeq 10^{-4} fm$!).

Les conditions de continuité sur les branes ($\phi = 0$ et $\phi = \pm\pi$) des fonctions d'onde ainsi que de leurs dérivées premières impliquent que

$$J_1(x_n^G) = 0 \quad (1.15)$$

avec $x_n \equiv z_n^G(\phi = \pi)$.

En supposant que la masse du premier état excité du graviton est très inférieure à k ($m_n/k \ll 1$) et étant donné que $e^{-2k\pi r_c} \gg 1$, on peut montrer que l'expression de la masse de l'état d'excitation n d'un graviton [3] s'écrit :

$$m_n = x_n k e^{-k\pi r_c} \quad (1.16)$$

Le mode zéro du graviton de Kaluza-Klein correspondant au graviton de masse nulle est donné par :

$$\chi_G^{(0)} = \sqrt{kr_c}. \quad (1.17)$$

Il est important de remarquer que les masses des états d'excitation du graviton sont proportionnelles aux zéros de la fonction J_1 de Bessel (éq.(1.16)), ce qui implique que les états ne sont pas équidistants en masse.

Par ailleurs, $x_0 \simeq 3.83$ et $k e^{-k\pi r_c}$ est de l'ordre de la centaine de GeV ou du TeV. On peut en déduire alors que les masses des premiers états excités sont de l'ordre de quelques centaines de GeV, et espacées d'environ autant. Chaque état apparaît alors comme une résonance massive bien séparée de la résonance suivante.

1.5 La recherche de gravitons auprès des collisionneurs

La production de graviton de Kaluza-Klein dans le modèle de Randall Sundrum est totalement différente de celle dans le modèle ADD. En effet, dans le modèle ADD les états excités sont exactement équidistants, chaque état a une masse de n/R^δ , où R est le rayon de compactification, et est pris comme étant le même pour chacune des δ dimensions supplémentaires. Les états apparaissent alors comme un quasi-continuum d'états excités ayant tous un couplage à la matière de l'ordre de $1/M_{Pl}$ donc très supprimé. Cependant la section efficace totale intégrée sur tous les états de Kaluza-Klein dans ce modèle est importante (de l'ordre du picobarn) du fait du très grand nombre d'états accessibles. Dans le modèle de Randall-Sundrum que nous étudions, cet ordre de grandeur correspond à la section efficace de production du premier état excité du graviton. Ainsi les deux modèles sont complètement différents aussi bien du point de vue théorique que du point de vue de leurs conséquences expérimentales.

Bibliographie

- [1] G. Nordström, *On the possibility of a unification of the electromagnetic and gravitation field*,
Z. Phys. 15 (1914) 504.
- [2] T. Kaluza, *On the problem of unity in Physics*,
Sitzungsber. Preuss. Akad. Wiss. Berlin (Math. Phys.) 1921 266-972.
- [3] O. Klein, *Quantum theory and five-dimensional theory of relativity*,
Z. Phys. 37 (1926) 895-906.
- [4] S. L. Glashow, Nucl. Phys. A22 (1961) 579 ;
A. Salam et J.C. Ward, Phys. Rev. Lett. 13 (1964) 168 ;
S. Weinberg, Phys. Rev. Lett. 19 (1967) 1264.
- [5] A. Heister *et al.*, ALEPH Collaboration, *Final results of the searches of neutral Higgs in e^+e^- collisions at $\sqrt{s}$ up to 209 GeV*,
Phys. Lett., **B526** (2002).
- [6] I. A. D'Souza, C. S. Kalman, *Preons : Models of leptons, quarks and gauge bosons as composite objects*,
World Scientific, Singapore 1992.
N. Delerue, thèse de doctorat de l'université Aix-Marseille, 2002.
- [7] P. Fayet et S. Ferrara, *Supersymmetry*, Phys. Rept. 32 (1977) 249.
J. Wess et B. Zumino, *Supergauge transformations in four dimensions*,
Nucl. Phys. **B70** p39-50, 1974.
- [8] P. van Nieuwenhuizen, *Supergravity*, Phys. Rep. **68** (1981) 189 ;
H. P. Nilles, *Supersymmetry, Supergravity and particle physics*,
Phys. Rept. **100** (1984) 1.
- [9] J. Scherk et J. Schwarz, Nucl. Phys. **B81** (1974) 155.
- [10] Yoneya, Prog. Theor. Phys. 51 (1974) 1907.
- [11] P. Goddard, J. Goldstone, C. Rebbi, C. B. Thorn, Nucl. Phys. **B56** (1973) 109.
- [12] J. Polchinski, *String Theory, Vol I et II*.
- [13] M. B. Green, J. H. Schwarz, E. Witten, *Superstring Theory, Vol I et II*.
- [14] J. Gross, J. Harvey, E. Martinec et R. Rohm, *The heterotic string*,
Phys. Rev. Lett. 54 (1985) 502-505 ;
Heterotic string theory 1. The free heterotic string,
Nucl. Phys. B256 (1986) 253-310 ;
Heterotic string theory 2. The interacting heterotic string,
Nucl. Phys. B267 (1986) 75-166.
- [15] Hewett et Rizzo, Phys Rev. **D33** (1986) 1476.
- [16] Witten, hep-th/9602070 Nucl. Phys. **B471** (1996) 121.

- [17] Lykken, Phys. Rev. **D54** (1996) 3693.
- [18] Polchinski, Phys. Rev. **D50** (1994) 6041;
Phys. Rev. Lett. 75 (1995) 4724.
- [19] I. Antoniadis *A possible new dimension at a few TeV*,
Phys. Lett. **B 246** (1990) 377-384.
- [20] N. Arkani-Hamed, S. Dimopoulos, G. R. Dvali, *The hierarchy problem and new dimensions at a millimeter*,
Phys. Lett. B429 (1998) 263-272
hep-ph/9803315.
- [21] L. Randall, R. Sundrum, *A large mass hierarchy from small extra dimension*,
Phys. Rev. Lett. 83 (1999) 3370-3373;
hep-ph/9905221.
- [22] L. Randall, R. Sundrum, *An Alternative to compactification*, Phys. Rev. Lett. 83:
4690-4693,1999 ;
hep-th/9906064.
- [23] W. D. Goldberger, M. B. Wise, *Modulus Stabilization With Bulk Fields*,
Phys. Rev. Lett. 83: 4922-4925,1999 ;
hep-ph/9911457.
- [24] W. D. Goldberger, M. B. Wise, *Phenomenology of Stabilized Modulus*,
Phys. Lett. **B475**: 275-279 ,2000 ;
hep-ph/9907447.
- [25] R. Altendorfer, J. Bagger et D. Nemeschansky, *Supersymmetric Randall-Sundrum Scenario*,
Phys. Rev. **D63** 125025, 2001.
- [26] J. Bagger et D. Nemeschansky, R.-J. Zhang, *Supersymmetric Radion in the Randall-Sundrum Scenario*, JHEP 0108 057, 2001;
hep-th/0003117.
- [27] T. Gherghetta et A. Pomarol Nucl. Phys. **B586** (2000) 141.
- [28] A. Falkowski, Z. Lalak, S. Pokorski, *Supersymmetrizing Branes with Bulk in Five-Dimensionnal Supergravity*, Phys. Lett. **B491** (2000) 172,
hep-th/0004093;
Five-Dimensional Gauged Supergravities with Universal Hypermultiplet and Warped Brane Worlds, Phys. Lett. *B509*:337-345,2001 ;
hep-th/0009167.
- [29] H. Davoudiasl, J. L. Hewett, T. G. Rizzo, *Experimental Probes of Localized Gravity: On and Off the Wall*,
Phys. Rev. **D63**: 075004 (2001),
hep-ph/0006041 ;
Warped phenomenology,
hep-ph/9909255,
Phenomenology of the Randall-Sundrum Gauge Hierarchy Model,
Phys. Rev. Lett. **84**, 2080 (2000).

Chapitre 2

Le graviton de Kaluza-Klein

Introduction

Dans ce chapitre nous allons décrire la forme du signal attendu. Notre canal d'analyse est la production résonante du premier état excité du graviton de Kaluza-Klein dans sa désintégration muonique avec des quarks ou des gluons dans l'état initial. Pour cela, il est nécessaire de comprendre ce que sont les états de Kaluza-Klein et les couplages qu'ils ont à la matière. Comme l'énergie dans le centre de masse du TeVatron est de 1.96 TeV, notre étude se limitera à la recherche du premier état excité du graviton. Nous rechercherons donc un graviton (particule de spin 2) dont la masse est comprise entre 100 GeV et 1 TeV. Nous donnerons tout d'abord les règles de Feynman qui permettent de calculer les diagrammes d'interaction du graviton avec la matière, puis nous détaillerons dans la section 2.2 les étapes de l'analyse de la forme du signal qui nous ont permis de rectifier une erreur dans le générateur PYTHIA que nous utilisons pour notre analyse. Dans la section 2.3 nous vérifierons que le terme d'interférence entre la production de graviton et celle de Drell-Yan et de Z est négligeable ce qui nous permet d'utiliser les générateurs habituels qui n'intègrent pas le terme d'interférence dans leurs calculs. Finalement nous verrons brièvement l'influence du choix des fonctions de distribution des partons dans le calcul des sections efficaces.

2.1 Règles de Feynman

Les expressions détaillées du propagateur et des vertex faisant intervenir un graviton de Kaluza-Klein n'ont jamais été explicitement publiées dans le cadre du modèle de Randall-Sundrum dans lequel nous nous plaçons. Ces calculs ont en effet été développés uniquement dans le cadre du modèle ADD [1] [2] dans une géométrie plate complètement factorisable, ou dans le cadre de modèles de Randall-Sundrum dans lesquels les champs du modèle standard peuvent se propager dans le bulk [3]. Néanmoins, en restreignant le calcul à la brane contenant les champs du modèle standard, on se retrouve dans un espace quadridimensionnel à courbure nulle dans les deux modèles ADD et RS1, et les règles de Feynman sont identiques moyennant des facteurs multiplicatifs aux vertex qui révèlent la différence entre les deux modèles [4].

2.1.1 Le propagateur du graviton

Pour calculer le propagateur du graviton, on peut partir de l'expression du tenseur d'Einstein, en supposant la métrique plate sur la brane visible et en se plaçant à basse énergie¹ :

$$\mathcal{G}_{\mu\nu} = \mathcal{R}_{\mu\nu} - 1/2 g_{\mu\nu} \mathcal{R} = -M_D^{-3} T_{\mu\nu} \quad (2.1)$$

$g_{\mu\nu} = \eta_{\mu\nu} + h_{\mu\nu}$ est la métrique explicitée précédemment.

1. Cette condition nous affranchit des éventuelles interactions entre les gravitons et la brane ainsi que les fluctuations de la position de la brane dans la cinquième dimension.

En développant l'expression de $\mathcal{G}_{\mu\nu}$ et ne gardant que les termes d'ordre 1 en $\hbar$, on peut écrire :

$$\begin{aligned} M_D^{3/2} \mathcal{G}_{\mu\nu} &= \square h_{\mu\nu} - \partial_\mu \partial^\rho h_{\rho\nu} - \partial_\nu \partial^\rho h_{\rho\mu} + \partial_\mu \partial_\nu h^\rho_\rho \\ &\quad - \eta_{\mu\nu} \square h^\rho_\rho + \eta_{\mu\nu} \partial^\rho \partial^\sigma h_{\rho\sigma} \\ &= -M_D^{-3/2} T_{\mu\nu} \end{aligned} \quad (2.2)$$

Partant de l'équation du mouvement à quatre dimensions pour un graviton de masse m_n :

$$(\square + m_n^2) \left(h_{\mu\nu}^{(n)} - \frac{1}{2} \eta_{\mu\nu} h^{(n)} \right) = 0 \quad (2.3)$$

nous pouvons extraire la forme du propagateur en se basant sur le développement donné par l'équation (2.2) :

FIG. 2.1 – Propagateur du graviton

$$\begin{aligned} P_{\mu\nu\alpha\beta}^{(n)} &= \frac{1}{2} (\eta_{\mu\alpha} \eta_{\nu\beta} + \eta_{\mu\beta} \eta_{\nu\alpha} - \eta_{\mu\nu} \eta_{\alpha\beta}) \\ &\quad - \frac{1}{2m_n^2} (\eta_{\mu\alpha} k_\nu k_\beta + \eta_{\nu\beta} k_\mu k_\alpha + \eta_{\mu\beta} k_\nu k_\alpha + \eta_{\nu\alpha} k_\mu k_\beta) \\ &\quad + \frac{1}{6} \left(\eta_{\mu\nu} + \frac{2}{m_n^2} k_\mu k_\nu \right) \left(\eta_{\alpha\beta} + \frac{2}{m_n^2} k_\alpha k_\beta \right) \end{aligned} \quad (2.4)$$

2.1.2 Couplages du graviton aux champs du modèle standard

Pour trouver la forme des couplages, nous revenons sur l'expression du lagrangien d'interaction entre le graviton $h_{\mu\nu}$ et la matière $T_{\mu\nu}$ [3] :

$$\mathcal{L}_{int}^{(4D \text{ eff})} = -\frac{1}{M_D^{3/2}} h_{\mu\nu}(x, \phi = \pi) T^{\mu\nu}(x) \quad (2.5)$$

On développe $h_{\mu\nu}$ en modes de Kaluza-Klein et on se place sur la brane du modèle standard, c'est-à-dire à $\phi = \pi$.

Le terme de mode zéro qui correspond au graviton non massif n'a pas le même couplage que les termes de modes $n > 0$ à cause de la différence entre les expressions de $\chi_G^{(0)}$ (équ. (1.17)) et $\chi_G^{(n)}$ (équ. (1.14)), nous séparerons donc le mode nul des modes non nuls dans le lagrangien d'interaction :

$$\mathcal{L}_{int}^{(4)} = -\frac{\sqrt{k}}{M_D^{3/2}} T^{\mu\nu}(x) h_{\mu\nu}^{(0)}(x) - \frac{\sqrt{k} e^{kr_c\pi}}{M_D^{3/2}} \sum_{n=1}^{\infty} T^{\mu\nu}(x) h_{\mu\nu}^{(n)}(x) \quad (2.6)$$

Par ailleurs, l'équation (1.10) donne que $\sqrt{k}/M_D^{3/2} = 1/\bar{M}_{Pl}$, ce qui nous permet d'écrire :

$$\mathcal{L}_{int}^{(4)} = -\frac{1}{\bar{M}_{PL}} T^{\mu\nu}(x) h_{\mu\nu}^{(0)}(x) - \frac{1}{\Lambda_\pi} \sum_{n=1}^{\infty} T^{\mu\nu}(x) h_{\mu\nu}^{(n)}(x) \quad (2.7)$$

où $\Lambda_\pi = e^{-k\pi r_c} \bar{M}_{PL}$ qui est de l'ordre de l'échelle électrofaible. Notons que les modes non nuls sont donc faiblement supprimés en comparaison avec le mode zéro ainsi qu'avec les modèles de larges dimensions supplémentaires dont nous avons discuté dans le chapitre précédent. Par ailleurs, nous retrouvons bien le couplage en $1/\bar{M}_{Pl}$ du mode zéro.

2.1.2.1 Couplage aux fermions

Le tenseur énergie-impulsion dans le cas des couplages aux fermions est donné par la relation :

$$T_{\mu\nu} = \frac{i}{4} \bar{\psi} (\gamma_\mu \partial_\nu + \gamma_\nu \partial_\mu) \psi - \frac{i}{4} (\partial_\mu \bar{\psi} \gamma_\nu + \partial_\nu \bar{\psi} \gamma_\mu) \psi \quad (2.8)$$

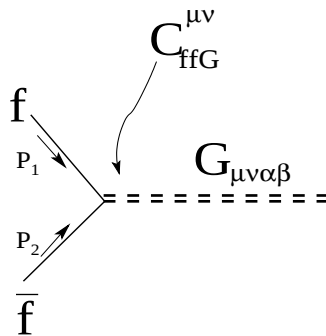


FIG. 2.2 – Couplage des fermions au graviton.

De cette expression on déduit les règles de Feynman du vertex décrivant ce couplage :

$$C_{ffG}^{\mu\nu} = \frac{-i}{4\Lambda_\pi} [(p_1 - p_2)^\mu \gamma^\nu + (p_1 - p_2)^\nu \gamma^\mu] \quad (2.9)$$

2.1.2.2 Couplage aux bosons de jauge

Le tenseur énergie-impulsion dans le cas des couplages aux bosons est donné par la relation :

$$T_{\mu\nu} = \eta_{\mu\nu} \left(\frac{1}{4} F^{\rho\sigma} F_{\rho\sigma} - \frac{m_A^2}{2} A^\rho A_\rho \right) - (F_\mu^\rho F_{\nu\rho} - m_A^2 A_\mu A_\nu) \quad (2.10)$$

Dans le cas des gluons, $m_A = 0$ et $F_{\mu\nu} = \partial_\mu A_\nu^a - \partial_\nu A_\mu^a - g_s f^{abc} A_\mu^b A_\nu^c$, a, b et c sont les indices de groupe, t^a sont les générateurs du groupe $SU(3)_C$ et f^{abc} sont les constantes de structure de ce groupe.

Les règles de Feynman pour ce couplage donnent, dans le cas d'un vertex à trois pattes faisant intervenir deux gluons :

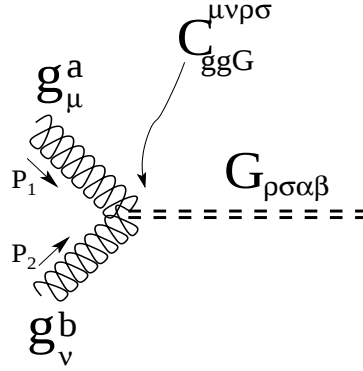


FIG. 2.3 – Couplage des gluons au graviton.

$$C_{ggG}^{\mu\nu\rho\sigma} = \frac{-i}{\Lambda_\pi} \delta^{ab} [W_{\mu\nu\rho\sigma} + W_{\nu\mu\rho\sigma}] \quad (2.11)$$

où on a posé :

$$W_{\mu\nu\rho\sigma} = \frac{1}{2} \eta_{\mu\nu} (p_{1\sigma} p_{2\rho} - p_1 \cdot p_2 \eta_{\rho\sigma}) + \eta_{\rho\sigma} p_{1\mu} p_{2\nu} \quad (2.12)$$

$$+ \eta_{\mu\rho} (p_1 \cdot p_2 \eta_{\nu\sigma} - p_{1\sigma} p_{2\nu}) - \eta_{\mu\sigma} p_{1\nu} p_{2\rho} \quad (2.13)$$

Des couplages entre un graviton et trois gluons ou plus sont aussi possibles mais très supprimés comparés au couplage à trois pattes. Nous n'en parlerons pas puisque notre étude se limite au couplage à deux gluons seulement.

2.1.2.3 Largeur de désintégration du graviton de Kaluza-Klein

La largeur partielle de désintégration d'un graviton $G^{(n)}$ en deux corps se calcule à partir de l'expression :

$$\frac{d\Gamma_n}{d\Omega} = \frac{1}{32\pi^2} |\mathcal{M}|^2 \frac{p^*}{m_n^2} \quad (2.14)$$

où p^* est l'impulsion des particules sortantes dans le centre de masse et vaut $m_n/2$ pour des particules ultra-relativistes, $\mathcal{M}$ est l'amplitude de probabilité de la désintégration du graviton en deux particules finales, et m_n est la masse du graviton.

Dans tous les canaux de désintégration du graviton en des champs du modèle standard, $|\mathcal{M}|^2 \propto 1/\Lambda_\pi^2$ car on a un couplage en $1/\Lambda_\pi$ au vertex. La largeur est donc :

$$\Gamma_n \propto m_n^3 / \Lambda_\pi^2. \quad (2.15)$$

En utilisant la relation $\Lambda_\pi = e^{-kr_c\pi} \bar{M}_{Pl}$ et $m_n = x_n k e^{-kr_c\pi}$, Γ_n se met sous la forme:

$$\Gamma_n \propto m_n x_n^2 \left(\frac{k}{\bar{M}_{Pl}} \right)^2. \quad (2.16)$$

La largeur est de l'ordre de quelques MeV pour de faibles couplages ($k/\bar{M}_{Pl} \simeq 0.01$), et quelques GeV pour les forts couplages ($k/\bar{M}_{Pl} \simeq 0.1$), pour des masses inférieures à 1 TeV. Retenons simplement que la largeur a une dépendance en $(k/\bar{M}_{Pl})^2$, et que le couplage a une dépendance en $k/\bar{M}_{Pl}$.

Pour un graviton produit hors couche de masse, sa largeur de désintégration varie selon les canaux de désintégrations possibles à l'énergie $\sqrt{\hat{s}}$ donnée. Pour une généralisation plus rigoureuse de la largeur, nous donnons ci-dessous les expressions des largeurs partielles du graviton à une énergie dans le centre de masse égale à $\sqrt{\hat{s}}$. Nous ne donnons ici que les contributions dues aux désintégrations du graviton en des particules du modèle standard [1] [13].

$$\Gamma(G \rightarrow gg, \gamma\gamma) = N_c \frac{\kappa^2 m_n^3}{160\pi} \quad (2.17)$$

$$\Gamma(G \rightarrow ZZ, WW) = \delta \frac{\kappa^2 m_n^3}{80\pi} (1 - 4r_A)^{1/2} \left(\frac{13}{12} + \frac{14}{39} r_A + \frac{4}{13} r_A^2 \right) \quad (2.18)$$

$$\Gamma(G \rightarrow f\bar{f}) = N_c \frac{\kappa^2 m_n^3}{320\pi} (1 - 4r_f)^{1/2} \left(1 + \frac{8}{13} r_f \right) \quad (2.19)$$

$$\Gamma(G \rightarrow H\bar{H}) = N_c \frac{\kappa^2 m_n^3}{960\pi} (1 - 4r_H)^{5/2} \quad (2.20)$$

où :

- N_c représente les états de couleur accessibles et vaut 8 pour les gluons, 3 pour les quarks et 1 pour les autres particules ;
- δ est le facteur de symétrie qui vaut 1/2 pour les désintégrations en deux particules identiques à savoir les bosons Z ;
- $r_X = m_X^2/\hat{s}$ est le rapport entre le carré de la masse de la particule X et le carré de l'énergie dans le centre de masse $\hat{s}$. Il va de soi que les canaux inaccessibles énergétiquement (pour $m_X > \sqrt{\hat{s}}/2$) sont "éteints" ;

$$- \kappa = 8\sqrt{\pi}/\Lambda_\pi .$$

Ce qui est mesurable auprès des accélérateurs, c’est la masse de la résonance ainsi que sa section efficace qui est proportionnelle au carré du couplage, donc proportionnelle à la largeur de la résonance. Ainsi, plutôt que d’utiliser les paramètres théoriques du modèle, nous utilisons la masse m_1 du premier état excité ainsi que $k/\bar{M}_{Pl}$ que nous appellerons par abus de langage le *couplage* dans les chapitres suivants.

Nous montrons sur la figure 2.4 les rapports d’embranchement des canaux de désintégration du graviton en des particules du modèle standard. La courbe désignée par “ $q\bar{q}$ ” somme toutes les contributions des quarks à l’exception du quark top dont le rapport d’embranchement est montré séparément.

La contribution des six paires leptons-antileptons est montrée sur une seule courbe nommée “ $l\bar{l}$ ”. La contribution des muons dans le rapport d’embranchement vaut un sixième de la valeur indiquée par cette courbe.

Le tableau 2.1 donne quelques exemples chiffrés de valeurs de rapports d’embranchements pour les gluons, les quarks (excepté le quark top), les muons et les photons.

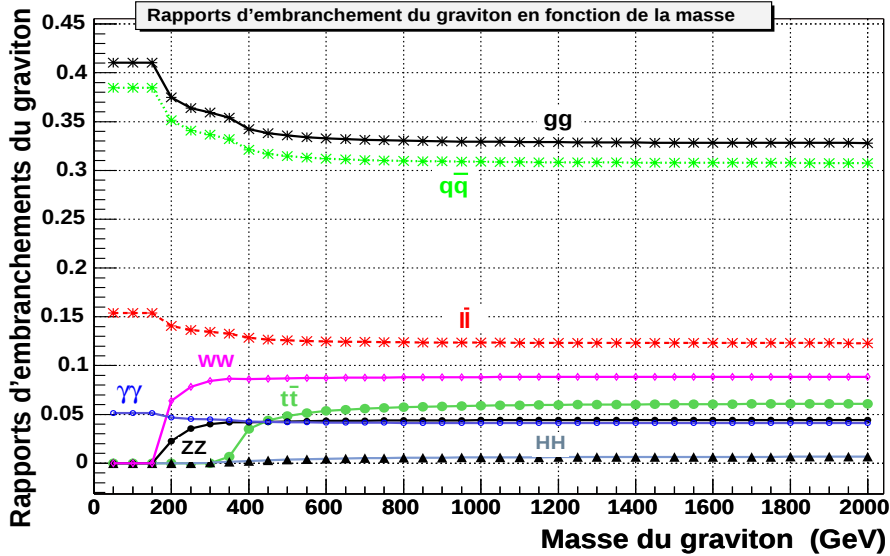


FIG. 2.4 – Rapports d’embranchement du graviton dans ses désintégrations en particules du modèle standard. Les valeurs dépendent de la masse du graviton et sont indépendantes de la valeur du couplage.

Nous pouvons noter que les canaux dijets qui sont produits par les désintégrations du graviton en deux gluons ou deux quarks sont les plus favorisés par le modèle. Cependant le bruit de fond pour ces canaux est aussi très élevé, principalement à cause des processus de QCD durs. Une étude préliminaire réalisée par la collaboration CDF montre que l’on a une meilleure sensibilité dans les canaux leptoniques [5] que dans le canal dijets [6].

Canal de désintégration	M_G en GeV				
	100	400	600	800	1000
$\text{BR}(G \rightarrow gg)$	$4,10 \cdot 10^{-1}$	$3,42 \cdot 10^{-1}$	$3,33 \cdot 10^{-1}$	$3,30 \cdot 10^{-1}$	$3,29 \cdot 10^{-1}$
$\text{BR}(G \rightarrow q\bar{q})$	$3,85 \cdot 10^{-1}$	$3,21 \cdot 10^{-1}$	$3,12 \cdot 10^{-1}$	$3,10 \cdot 10^{-1}$	$3,09 \cdot 10^{-1}$
$\text{BR}(G \rightarrow \mu^+ \mu^-)$	$2,57 \cdot 10^{-2}$	$2,14 \cdot 10^{-2}$	$2,08 \cdot 10^{-2}$	$2,07 \cdot 10^{-2}$	$2,07 \cdot 10^{-2}$
$\text{BR}(G \rightarrow \gamma\gamma)$	$5,13 \cdot 10^{-2}$	$4,28 \cdot 10^{-2}$	$4,16 \cdot 10^{-2}$	$4,13 \cdot 10^{-2}$	$4,12 \cdot 10^{-2}$
$\Gamma_{\text{tot}}(\text{GeV}) \text{ (k}/M_{Pl}=0,1)$	1,148	5,463	8,402	11,289	14,153

TAB. 2.1 – *Rapports d'embranchements pour les canaux de désintégration du graviton en une paire de gluons, de quarks, de muons et de photons, pour cinq valeurs de masses de graviton M_G indicatives. Les valeurs de la largeur totale sont indiquées pour chaque masse en dernière ligne du tableau.*

2.2 Etude du signal attendu

Les règles de Feynman que nous avons données précédemment nous permettent de calculer les sections efficaces de production de gravitons de Kaluza-Klein. Ce processus est implémenté dans le générateur PYTHIA [7] qui est interfacé avec l'environnement informatique de la collaboration. C'est pourquoi notre étude sera principalement basée sur ce générateur.

Nous étudions la production résonante de graviton, donc en voie s , à l'ordre des arbres. Les processus étudiés sont montrés sur la figure 2.5. Ils font intervenir soit des gluons soit des quarks dans l'état initial. La production de graviton par fusion de quark $q\bar{q} \rightarrow G \rightarrow \mu\mu$ interfère avec la production standard de paires de muons via un boson γ ou Z . Ce n'est bien-sûr pas le cas du processus $gg \rightarrow G \rightarrow \mu\mu$.

Nous étudierons dans cette section les deux processus. Dans un premier temps nous expliquerons les études que nous avons menées pour vérifier la forme du signal ; puis dans un deuxième temps nous montrerons que le terme d'interférence entre la production de gravitons et de photons ou bosons Z est négligeable dans le processus impliquant des quarks dans l'état initial.

$$\left| \begin{array}{c} g \\ g \end{array} \right\} \text{---} G \text{---} \left\{ \begin{array}{c} \mu^+ \\ \mu^- \end{array} \right\} \right|^2 + \left| \begin{array}{c} q \\ \bar{q} \end{array} \right\} \text{---} G, Z, \gamma \text{---} \left\{ \begin{array}{c} \mu^+ \\ \mu^- \end{array} \right\} \right|^2$$

FIG. 2.5 – *Diagrammes de Feynman des processus étudiés pour la production résonante de graviton de Kaluza-Klein. Le processus faisant intervenir des quarks dans l'état initial interfère avec la production dimuons via un boson γ ou Z , ce qui n'est pas le cas pour le processus dû à la fusion de gluons.*

Pour l'étude du signal Monte Carlo, les événements ont d'abord été générés avec la version 6.2 de PYTHIA.

Comme le montrent les expressions des largeurs du graviton données précédemment, la largeur de la résonance dépend de la valeur du couplage du graviton aux champs du modèle standard. Ainsi, plus le couplage est grand et plus la largeur de la résonance est grande.

PYTHIA prend en compte *a priori* les effets dus à la largeur finie du graviton. Bien que dans la version 6.5 du générateur HERWIG [8] (la version interfacée à l'environnement informatique de la collaboration) le traitement de la résonance de graviton ne soit pas tout à fait rigoureux, nous avons comparé les signaux donnés par ces deux générateurs, sachant qu'*a priori* nous devrions retrouver globalement la même forme.

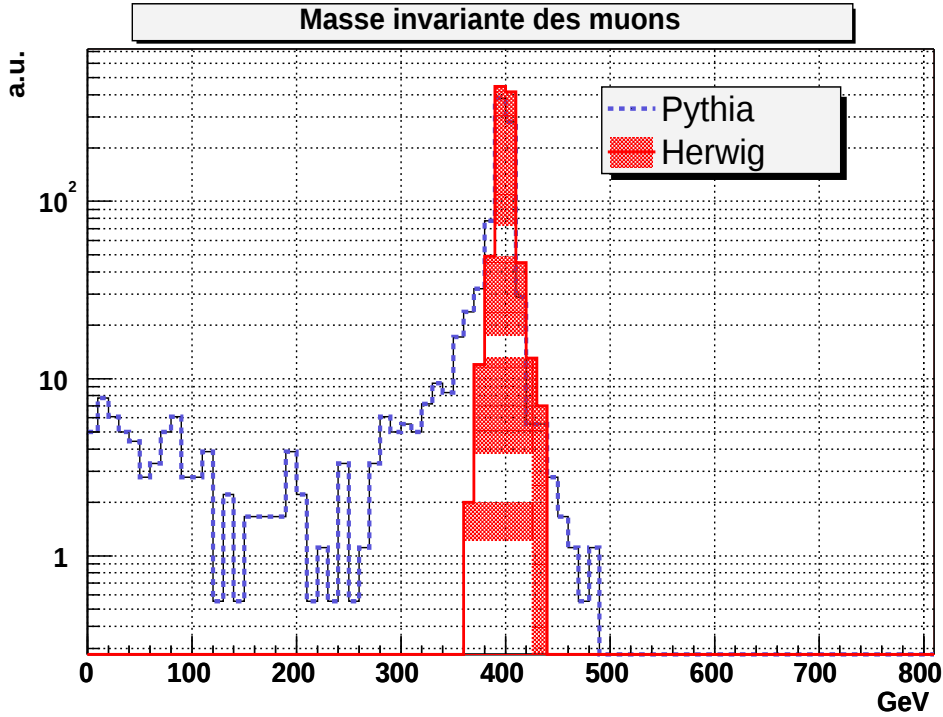


FIG. 2.6 – Distribution de masse invariante de paires de muons issus de la désintégration d'un graviton de masse 400 GeV pour une valeur de $k/\bar{M}_{Pl} = 0,1$. L'histogramme grisé montre la distribution donnée par le générateur HERWIG, l'autre histogramme est donné par le générateur PYTHIA.

Pour cela nous avons généré 1000 événements avec PYTHIA et HERWIG pour un graviton de masse 400 GeV et un couplage correspondant à $k/\bar{M}_{Pl} = 0,1$. La figure 2.6 montre les distributions de masse invariante des muons données par HERWIG et par PYTHIA. Nous observons que dans le cas de PYTHIA, la masse invariante montre beaucoup d'événements à basse masse, tandis que le pic généré par HERWIG est étroit et centré sur la masse de la résonance.

Nous avons ensuite regardé ce que PYTHIA donne pour des valeurs de largeur différentes pour une masse donnée de graviton, afin de voir l'influence de la largeur, c'est-à-dire du couplage, sur la forme du signal. La figure 2.7 montre la distribution de la masse invariante des produits de désintégration du graviton pour une masse de 400 GeV et pour des valeurs de largeurs indicatives de 1,4 GeV, 5,5 GeV et 22,6 GeV. Nous voyons que

pour de grandes largeurs, c'est-à-dire de forts couplages, la queue à basses masses devient très importante devant le pic de résonance, ce qui a une influence directe importante sur l'acceptance d'une analyse recherchant le signal à grande masse. C'est pour cela que nous avons tenu à vérifier la forme réelle que le modèle prédit pour la résonance.

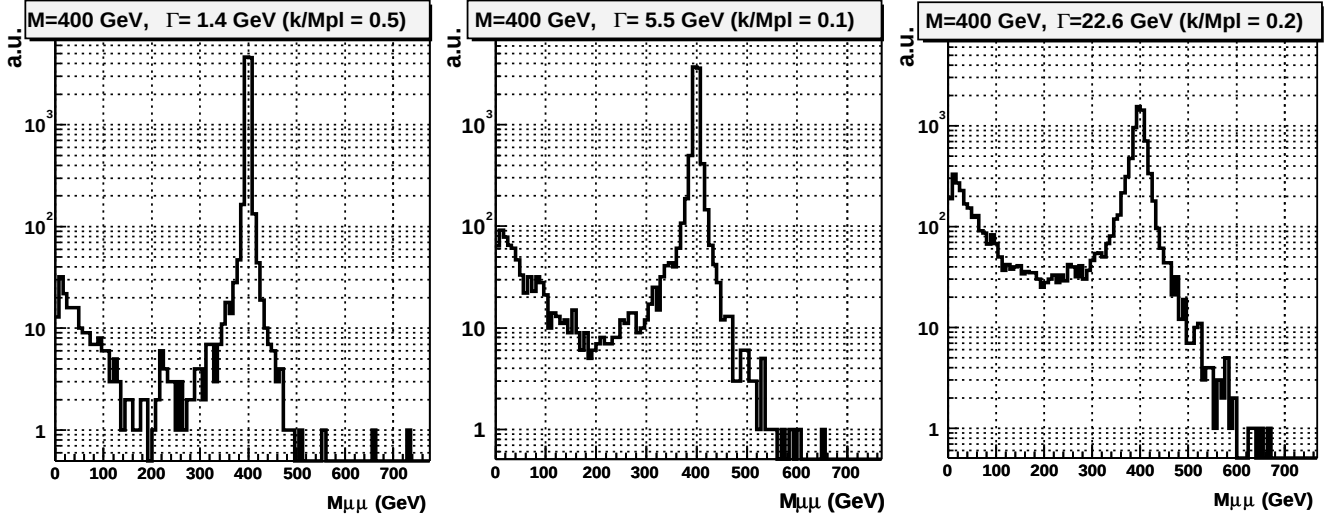


FIG. 2.7 – Masse invariante des muons issus de la désintégration du graviton généré par la version 6.2 de PYTHIA, pour une masse de graviton de 400 GeV et trois valeurs indicatives pour les largeurs (valeurs de $k/\bar{M}_{Pl}$) : 1,4 GeV (0,05), 5,5 GeV (0,1) et 22,6 GeV (0,2).

Pour pouvoir vérifier la forme de la résonance, nous avons développé un générateur d'événements indépendant de PYTHIA et HERWIG dans lequel nous contrôlons les amplitudes de probabilité des processus et les paramètres du problème. Ce générateur doit permettre de :

- calculer la section efficace différentielle des processus $gg \rightarrow G \rightarrow \mu\mu$ et $q\bar{q} \rightarrow G \rightarrow \mu\mu$ étudiés en fonction des variables cinématiques du problème ;
- intégrer la section efficace différentielle ainsi calculée. Cette intégration se fait numériquement par une méthode de Monte Carlo ;
- puis finalement générer des événements selon cette section efficace.

Nous allons décrire ces étapes dans ce qui suit.

Notre programme aura avant tout besoin de calculer la largeur Γ_G du graviton pour des valeurs données de masse et de couplage. Pour cela nous avons codé les expressions données par les équations (2.17) à (2.20). Nos résultats sont en très bon accord avec PYTHIA (voir tableau 2.2).

2.2.1 Calcul de l'amplitude de probabilité

La première étape est de calculer l'amplitude de probabilité des processus étudiés (voir figure 2.5).

$M_G(\text{GeV})$	$\Gamma_G(\text{PYTHIA}) (\text{GeV})$	$\Gamma_G(\text{Notre g�n�rateur}) (\text{GeV})$
100	1,14778	1,14417
400	5,46301	5,48472
600	8,40239	8,46059
800	11,28872	11,36436
1000	14,15276	14,25057

TAB. 2.2 – Comparaison des largeurs totales du graviton obtenues par notre programme   celles donn es par PYTHIA pour diff rentes masses, et pour une valeur de $k/\bar{M}_{Pl}=0,1$.

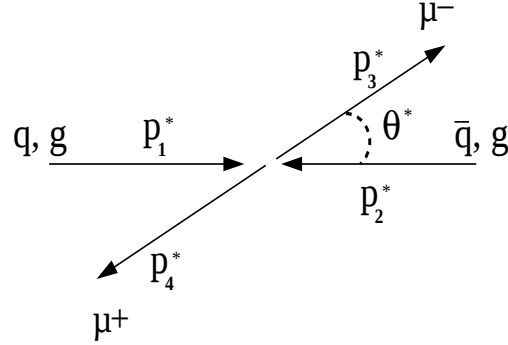


FIG. 2.8 – Sch ma de production r sonante de graviton repr sent  dans le centre de masse. Le processus fait intervenir des gluons ou des quarks dans l' tat initial, et des muons dans l' tat final. L'angle θ^* est par convention l'angle form  entre le proton et le muon.

Pour cela, on utilise l'outil FORM [9]. C'est un langage de programmation qui permet de simplifier des expressions math matiques litt rales en manipulant des vecteurs, tenseurs   plusieurs indices etc. , et qui contient d j  quelques outils math matiques comme le nombre complexe i , les tenseurs ϵ_{ijk} , les fonctions δ de Kr necker, les matrices γ^μ de Dirac ainsi que leur alg bre, et d'autres fonctions math matiques tr s utilis es en physique.

Le type de variables trait es comprend entre autres les symboles (nombres multiplicatifs), les indices (de sommation), les fonctions, les vecteurs (dont on peut sp cifier la dimension qui vaut 4 par d faut) et les tenseurs. Il est possible de pr ciser les r gles de commutation des fonctions ou tenseurs et pr ciser s'ils sont sym triques, cycliques par rapport aux variables que l'on doit pr ciser, etc. FORM effectue la sommation sur les indices r p t s (il utilise la convention de sommation d'Einstein).

L'une des fonctionnalit s les plus importantes pour le calcul des amplitudes carr es de probabilit  des processus physiques est la possibilit  de sp cifier la ligne de spins sur laquelle il faut calculer la trace des matrices lors du calcul.

D'une mani re g n rale, cet outil se r v le tr s puissant pour calculer les amplitudes de probabilit  de processus multiples ou des termes d'interf rence compliqu s o  le calcul   la main s'av re tr s p rilleux compte tenu des risques d'erreurs de signes ou de calcul de traces de produits de matrices de Dirac. G n ralement, le r sultat est impl ment  dans

des programmes sous forme de fonctions dont on calcule la valeur pour une ou plusieurs variables. Dans ce but, l'auteur a ajouté la possibilité de transformer l'expression littérale du résultat final dans un format compatible avec le langage Fortran.

Un exemple de programme permettant le calcul des amplitudes de probabilité des processus de production résonante de Z , γ et graviton se désintégrant en deux muons est commenté dans l'annexe A. C'est le programme qui est utilisé pour calculer les amplitudes dont on se sert dans cette section.

Nous avons donc donné tous les paramètres dont notre programme a besoin pour calculer l'amplitude de probabilité des processus $q\bar{q} \rightarrow G \rightarrow \mu\mu$ et $gg \rightarrow G \rightarrow \mu\mu$ à une énergie dans le centre de masse donnée $\sqrt{\hat{s}}$, et nous obtenons en retour les expressions suivantes :

$$|\mathcal{A}_q|^2 = |T_{fi}(q\bar{q}) T_{fi}^*(q\bar{q})| = \frac{\hat{s}^4}{16 \Lambda_\pi^4} \frac{(1 - 3 \cos^2(\theta^*) + 4 \cos^4(\theta^*))}{((\hat{s} - m_G^2)^2 + m_G^2 \Gamma_G^2)} \quad (2.21)$$

$$|\mathcal{A}_g|^2 = |T_{fi}(gg) T_{fi}^*(gg)| = \frac{\hat{s}^4}{16 \Lambda_\pi^4} \frac{1 - \cos^4(\theta^*)}{((\hat{s} - m_G^2)^2 + m_G^2 \Gamma_G^2)} \quad (2.22)$$

Dans cette expression nous retrouvons :

- le terme $1/((\hat{s} - m_G^2)^2 + m_G^2 \Gamma_G^2)$ décrivant le propagateur de la résonance, où m_G et Γ_G sont les masse et largeur du graviton ;
- le terme en $1/\Lambda_\pi^4$ provenant du vertex ;
- le terme en $(1 - \cos^4(\theta^*))$ dans la contribution des gluons qui décrit la distribution angulaire des particules de spin 1/2 issues de la désintégration d'une particule de spin 2. La distribution angulaire donnée par les quarks dans l'état initial est donnée par le terme en $1 - 3 \cos^2(\theta^*) + 4 \cos^4(\theta^*)$ dans l'expression de l'amplitude $|\mathcal{A}_q|^2$. Ces dépendances angulaires peuvent être retrouvées grâce aux matrices $\delta_{mm'}^J$ de Wigner donnant la distribution angulaire en fonction des spins et de l'hélicité des particules mises en jeu dans le processus².

Le plus important à noter est que la dépendance des amplitudes est en $\hat{s}^4$, contrairement à tous les processus "2→2" du modèle standard qui sont en $\hat{s}^2$. Ceci est dû au fait que tous les couplages du modèle standard sont adimensionnés, alors que le couplage du graviton a une dimension en $1/\bar{M}_{Pl}$ à chaque vertex.

2.2.2 Section efficace $p\bar{p} \rightarrow G \rightarrow \mu\mu$

2.2.2.1 Section efficace différentielle

Nous avons calculé les amplitudes de probabilité des sous-processus "durs" :

$$q\bar{q} \rightarrow G \rightarrow \mu^+ \mu^- \quad (2.23)$$

2. Dans les désintégrations d'une particule de spin J, les matrices de Wigner $\delta_{mm'}^J$ donnent la distribution angulaire des produits de désintégration pour un moment angulaire de l'état initial égal à m et un moment angulaire de l'état final égal à m'.

Dans le processus $q\bar{q} \rightarrow G \rightarrow \mu\mu$, les hélicités possibles des états initial et final sont LL et RR (L pour gauche et R pour droite). La distribution angulaire est donnée par $(\delta_{11}^2)^2 + (\delta_{1-1}^2)^2 + (\delta_{-11}^2)^2 + (\delta_{-1-1}^2)^2$. Dans le processus $gg \rightarrow G \rightarrow \mu^+ \mu^-$, les hélicités possibles de l'état final sont LL et RR, les gluons dans l'état initial sont réels donc de masse nulle et ont une polarisation transverse. La distribution angulaire est donnée par $(\delta_{21}^2)^2 + (\delta_{2-1}^2)^2 + (\delta_{-21}^2)^2 + (\delta_{-2-1}^2)^2$.

$$gg \rightarrow G \rightarrow \mu^+ \mu^-, \quad (2.24)$$

c'est-à-dire l'amplitude de probabilité étant donnés deux partons (deux gluons ou deux quarks) pour une énergie dans le centre de masse égale à $\sqrt{\hat{s}}$.

Pour des collisions hadroniques, les énergies des partons incidents ne sont pas connues. En effet, chaque parton emporte une fraction de l'énergie du hadron dont il est issu. Ainsi, pour le calcul de la section efficace totale du processus hadronique il faut tenir compte de la fraction x_1 emportée par le parton issu du proton et de la fraction x_2 emportée par le parton issu de l'antiproton.

Pour calculer la section efficace hadronique, il faut multiplier la section efficace différentielle précédente par la probabilité $q(x_1)$ et $q(x_2)$ d'avoir ces valeurs pour x_1 et x_2 . Pour une énergie dans le centre de masse hadronique $\sqrt{s}$, nous pouvons écrire la section efficace différentielle du processus dur $d\hat{\sigma}$:

$$d\hat{\sigma} = \Phi \times |\mathcal{A}(\hat{s} = x_1 x_2 s)|^2 \times \Delta P \quad (2.25)$$

où :

- Φ est le facteur de flux et vaut $1/2\hat{s}$;
- ΔP est l'espace des phases intégré sur l'angle azimutal ϕ et vaut $\frac{1}{16\pi} d\cos(\theta^*)$;

La section efficace différentielle du processus hadronique $d^3\sigma/dx_1 dx_2 d\cos(\theta^*)$ se déduit de l'expression précédente en prenant :

$$\frac{d^3\sigma}{dx_1 dx_2 d\cos(\theta^*)} = \frac{d\hat{\sigma}}{d\cos(\theta^*)} \delta(\hat{s} - x_1 x_2 s) q(x_1) q(x_2) \quad (2.26)$$

où $q(x_1)$ et $q(x_2)$ sont les probabilités d'avoir les valeurs x_1 et x_2 pour les partons en collision.

2.2.2.2 Les densités de partons

Les valeurs de x_1 ou de x_2 ne sont pas équiprobables pour un parton donné, et ne sont pas les mêmes selon le type de parton. Leur probabilité $q(x, Q^2)$ est donnée par une des fonctions de distribution des partons appelées *PDF* (*Parton Distribution Fonction*), Q étant l'échelle d'énergie à laquelle on "sonde" le proton. C'est une échelle d'énergie caractéristique du processus (elle peut être prise égale à l'énergie dans le centre de masse du processus, à l'impulsion transverse des particules sortantes, ou encore à la masse de la particule résonante). Les fonctions de distribution des partons donnent donc la probabilité de trouver un parton emportant une fraction x de l'énergie du proton (ou antiproton) à une échelle Q^2 donnée.

Les *PDF* ne peuvent pas être calculées à partir de la QCD perturbative. Elles sont donc le résultat d'ajustements des paramètres de la fonction avec de nombreuses données expérimentales obtenues à partir des processus de diffusion profondément inélastiques (souvent sur cible fixe).

Les différences majeures entre les *PDF* résident dans le choix des lots de données analysées ainsi que de l'échelle Λ_{QCD} utilisée pour les ajustements. Ainsi chaque groupe de travail a ses hypothèses de base, ses contraintes de calcul et de mesure ainsi que son choix de lot de données. La fonction peut varier plus ou moins sensiblement si l'on change par exemple la proportion des quarks de la mer, la densité de gluons selon les échelles d'énergie, ou le couplage fort α_s . Les fonctions ainsi paramétrées portent souvent les initiales de leurs auteurs ainsi que des suffixes indiquant si l'ajustement a été fait en tenant compte des effets au premier ordre (*Leading Order*) ou au second ordre (*Next to leading order*), ainsi que de la version de la mise à jour de la fonction. Pour plus de détails nous référons à la description de PDFLIB [11] qui dresse une liste exhaustive des *PDF* disponibles ainsi que de leurs caractéristiques.

Parmi les quelques groupes de travail qui réalisent ces ajustements, on compte trois groupes principaux :

- CTEQ (Coordinated Theoretical-Experimental Project on QCD) ;
- MRS (du nom de leurs auteurs Martin, Robert et Stirling) ;
- et GRV (de Glück, Reya et Vogt).

Dans la suite, pour le calcul des sections efficaces, nous ferons appel à la librairie PDFLIB développée au CERN qui donne les valeurs de $q(x, Q^2)$ pour un choix de *PDF*. Nous choisissons la paramétrisation CTEQ5L [14]. L'échelle d'énergie Q à laquelle nous nous plaçons est $\sqrt{\hat{s}}$.

2.2.3 Génération aléatoire d'événements

Dans ce qui suit nous expliquons la méthode employée pour intégrer la section efficace différentielle des processus étudiés (que l'on a calculée précédemment) puis générer des événements selon cette loi de probabilité. Les sections efficaces que l'on a calculées dépendent de trois paramètres principaux qui déterminent complètement le résultat. Les paramètres mis en jeu sont $\cos(\theta^*)$ qui est compris entre -1 et 1, x_1 et x_2 qui sont compris entre 0 et 1. Ces variables x_1 et x_2 sont les fractions d'énergie emportées par chaque parton en collision dans le proton ou l'antiproton dont le parton est issu.

Il nous faut donc tirer des valeurs aléatoires de x_1 , x_2 et $\cos(\theta^*)$ selon la distribution de la section efficace différentielle du processus que l'on notera $f(x_i)$ et dans les bornes de validité des variables.

Nous pouvons utiliser la méthode du double tirage. Elle permet de tirer selon une loi uniforme trois nombres y_1 , y_2 et y_3 compris entre 0 et 1 de telle sorte que :

$$x_1 = y_1 \times (x_{max} - x_{min}) + x_{min},$$

$$x_2 = y_2 \times (x_{max} - x_{min}) + x_{min},$$

où $x_{min}=0,0001$ et $x_{max}=0,9999$

$$\text{et } \cos(\theta^*) = 2 y_3 - 1.$$

Notons f_{max} la valeur maximale que peut prendre la section efficace différentielle du processus étudié dans l'intervalle considéré. On tire ensuite alors un triplet de variables y_1 , y_2 et y_3 (premier tirage) et on calcule la valeur $f(y_1, y_2, y_3)$ de la section efficace différentielle pour ces trois variables. On tire alors une valeur α quelconque comprise entre 0 et 1

(second tirage), et on compare cette valeur au rapport $\mathcal{R} = f(y_1, y_2, y_3)/f_{max}$. Si $\mathcal{R} > \alpha$ alors on garde le triplet et on enregistre l'événement correspondant, sinon on refait un double tirage en prenant un nouveau triplet de variables ainsi qu'une nouvelle valeur de α .

Cette méthode permet de garder davantage de points dans la région où f est grand que dans celle où f est très petit, ce qui permet de reproduire le comportement de f .

Si f présente des pics très marqués, cette méthode se révèle alors très inefficace. En effet tous les triplets (y_1, y_2, y_3) ont la même probabilité d'être tirés au départ, mais dans les régions où f est très faible, le taux de réjection des triplets est très élevé étant donné qu'il est peu probable que $f(y_1, y_2, y_3)/f_{max} > \alpha$. Dans notre cas, nous savons que la résonance que nous voulons observer est très étroite et très piquée à la masse du graviton, la valeur de l'amplitude de probabilité dans les régions éloignées de la résonance étant de plusieurs ordres de grandeur inférieure à la valeur maximale au pic.

Un changement de variables adéquat permet de s'affranchir du pic étroit de la fonction. En effet, le pic que l'on observe dans la section efficace différentielle est dû au dénominateur de la Breit-Wigner $(\hat{s} - m^2)^2 + m^2\Gamma^2$. En opérant le changement de variable $t = (\hat{s} - m^2)/m\Gamma$, on fait apparaître la forme $1/(1 + t^2)$ qui donne un arctan(t) après intégration. La fonction est plus aplanie, et l'on peut utiliser la méthode du double tirage en ayant diminué l'inefficacité du tirage.

Mais nous avons préféré utiliser une méthode qui fait appel à un ensemble de programmes *BASES/SPRING* [10] qui permet de réaliser l'intégration et la génération des événements. La partie *BASES* a pour but d'intégrer l'amplitude de probabilité que l'on donne en entrée après optimisation d'une grille d'échantillonnage du domaine d'intégration, et la partie *SPRING* se charge de générer des événements.

Nous avons fait appel à cet outil car il offre plusieurs possibilités de vérification des étapes de l'optimisation et de calcul, ainsi que des histogrammes de contrôle.

L'idée de cette méthode est qu'il faut trouver une fonction g reliée à f et dont la variation est adoucie par rapport à f , afin d'utiliser la méthode du double tirage à partir de cette fonction g et non plus de f .

Le principe d'intégration utilisé par *BASES* est le suivant : par changement de variable³ on se ramène au cas où les variables x_i du problème sont dans l'intervalle $[0, 1]$. Chaque intervalle $[0, 1]$ est découpé en sous-intervalles. On note n_i le nombre de subdivisions de l'intervalle d'intégration associé à la i ème variable, et $I_k^i = [a_{k-1}^i, a_k^i]$ les sous-intervalles correspondants. Ceci revient à découper l'espace des variables aléatoires en régions appelées *hypercubes*.

On commence par prendre les I_k^i d'égale longueur $\Delta a_k^i = a_k^i - a_{k-1}^i$. On tire alors un grand nombre N_p de points $\vec{x} = (x_1, \dots, x_n)$, puis on calcule une première estimation de F_1

3. Ce changement de variable consiste à redéfinir les variables dans un intervalle compris entre 0 et 1. Ce n'est pas le même changement de variable que dans le cas du double tirage où il faut trouver une expression qui donne une arctangente.

l'intégrale de f et de sa variance σ_1 avec :

$$F_1 = \frac{1}{N_p} \sum_{k=1}^{N_p} f(\vec{x}_k) \omega_k \quad (2.27)$$

où $\omega_k = \left(\prod_{i=1}^n n_i\right) v_k$
et $v_k = \prod_{i=1}^n \Delta a_k^i$ est le volume de l'hypercube contenant le point $\vec{x}_k$,

$$\sigma_1^2 = \frac{1}{N_p (N_p - 1)} \left[N_p \sum_{k=1}^{N_p} (f(\vec{x}_k) \omega_k)^2 - \left(\sum_{k=1}^{N_p} f(\vec{x}_k) \omega_k \right)^2 \right] \quad (2.28)$$

On fait varier la largeur Δa_k^i des intervalles I_k afin de minimiser la variance. On tire alors un nouvel ensemble de N_p points $\vec{x}$ dans l'intervalle $[0,1]$ puis on calcule avec les nouvelles valeurs des Δa_k^i une deuxième estimation F_2 de l'intégrale de f , ainsi que la nouvelle variance σ_2^2 . On calcule la valeur moyenne $\tilde{F}$ à partir de F_1 et F_2 et la variance de cette moyenne $\tilde{\sigma}^2$ avec σ_1 et σ_2 :

$$\frac{1}{\tilde{\sigma}^2} = \frac{1}{\sigma_1^2} + \frac{1}{\sigma_2^2} \quad (2.29)$$

$$\tilde{F} = \tilde{\sigma}^2 \left(\frac{F_1}{\sigma_1^2} + \frac{F_2}{\sigma_2^2} \right) \quad (2.30)$$

On remodifie alors les Δa_k^i pour minimiser la variance moyenne $\tilde{\sigma}$, puis on retire une troisième fois un ensemble de N_p points. On calcule une troisième estimation de l'intégrale F_3 et de la variance σ_3 , puis on fait une moyenne comme précédemment en tenant compte de $F_1, F_2, F_3, \sigma_1, \sigma_2$ et σ_3 .

Cette opération itérative d'optimisation de la largeur des intervalles Δa_k^i se poursuit jusqu'à ce que le nombre d'itérations maximal imposé soit atteint, ou que la précision de calcul souhaitée soit atteinte.

On part donc d'un espace découpé régulièrement et on finit avec un espace découpé selon des hypercubes de tailles très différentes. Ceux-ci sont très petits au voisinage des pics de la fonction f de sorte que l'on tirera davantage de points dans cette zone.

L'optimisation de la grille étant faite, on doit alors se baser sur ce "découpage" de l'espace en hypercubes pour intégrer la fonction. On tire alors N_T ensembles de N_p points pour lesquels on calcule une valeur moyenne $\tilde{F}$ ainsi qu'une variance $\tilde{\sigma}^2$ telles que :

$$\tilde{F} = \tilde{\sigma}^2 \left(\sum_{j=1}^{N_T} \frac{F_j}{\sigma_j^2} \right) \quad (2.31)$$

$$\frac{1}{\tilde{\sigma}^2} = \sum_{j=1}^{N_T} \frac{1}{\sigma_j^2} \quad (2.32)$$

Plus N_T est grand, plus grande est la précision du calcul. On augmente alors N_T et on réeffectue le calcul jusqu'à ce que le nombre maximal imposé pour N_T soit atteint ou que la précision de calcul maximale souhaitée soit atteinte.

Ce résultat final est utilisé par SPRING qui prend en charge la génération des événements.

En effet, SPRING utilise les résultats de l'optimisation de la grille (les Δa_k^i) qui sont enregistrés dans un fichier. Un hypercube est d'abord choisi selon une probabilité inversement proportionnelle à son volume. On tire au hasard un point dans cet hypercube de volume v , et on calcule $g(\vec{x}) = f(\vec{x}) \times v$. Le point $\vec{x}$ est gardé si $g(\vec{x})/g_{max} > \alpha$ où α est un nombre aléatoire.

Ainsi la fonction que l'on utilise pour le double tirage est la fonction g , par construction très aplaniée par rapport à f .

2.2.4 Résultat de la génération

A partir de l'expression de la section efficace différentielle de production de graviton nous avons généré des événements pour chaque contribution (2.23) et (2.24) pour une même luminosité. Le signal total obtenu est représenté par l'histogramme en trait plein sur la figure 2.9.

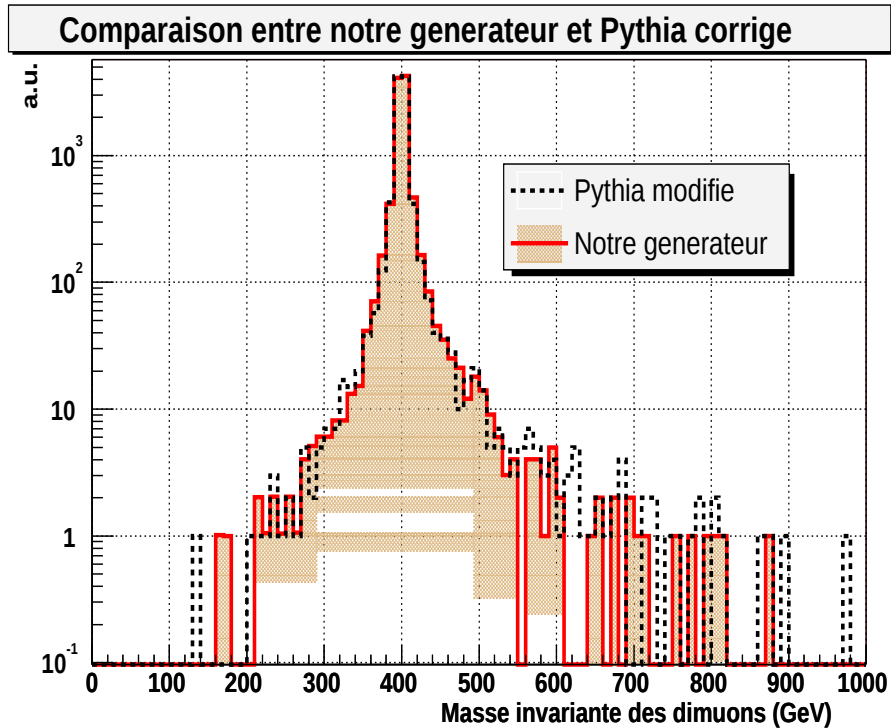


FIG. 2.9 – Distribution de masse invariante des événements générés par notre générateur (histogramme en trait plein) pour une masse de graviton de 400 GeV et un couplage $k/\bar{M}_{Pl}=0,1$. L'histogramme en pointillés montre ce qui est obtenu en modifiant PYTHIA (voir texte) pour les mêmes valeurs de masse et de couplage.

L'histogramme montré dans la figure 2.6 représente la distribution initialement donnée par la version 6.2 de PYTHIA tandis que l'histogramme en trait plein sur la figure 2.9 concerne la distribution obtenue par notre générateur. On constate que ce dernier ne prédit pas un nombre important d'événements à basse masse, contrairement à PYTHIA.

Après investigations et interactions avec l'auteur principal de PYTHIA, il s'est avéré [12] que dans l'expression de la section efficace différentielle codée dans PYTHIA (reposant sur [13]) un facteur $\hat{s}^2$ a été remplacé par m_G^4 .

Si l'on modifie le programme de PYTHIA en introduisant le facteur restaurant la dépendance de l'amplitude carrée en $\hat{s}^4$, on obtient l'histogramme en traits pointillés sur la figure 2.9. Nous observons un très bel accord entre ce dernier histogramme et celui obtenu par notre générateur.

Notons que la queue prédite par PYTHIA dans sa version de base est due principalement à la contribution des gluons et non des quarks dans l'état initial (voir figure 2.10). La cause revient à la convolution de la Breit-Wigner avec les PDF. En effet, la densité de gluons diminue beaucoup plus vite que celle des quarks à grand x . Pour des résonances larges⁴, les gluons contribueront fortement à très faibles valeurs de x , ce qui a pour effet de favoriser les basses masses et créer de ce fait une queue⁵.

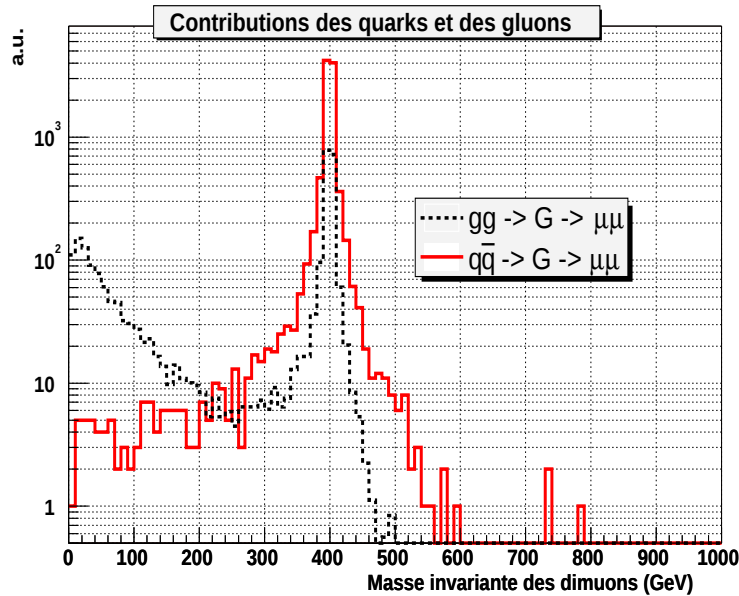


FIG. 2.10 – Contribution des gluons (en pointillés) et des quarks (en trait plein) dans l'état initial dans la production en voie s d'un graviton de 400 GeV de masse pour une valeur de $k/\bar{M}_{Pl}=0,1$ prédit par la version 6.202 de PYTHIA. Les deux contributions sont normalisées à la même luminosité.

Afin de vérifier cet argument, nous avons généré 10000 événements “Z” avec des quarks

4. ce terme est un peu vague et dépend de la masse de la résonance, et de plusieurs autres facteurs.

5. Notons que cet effet physique est aussi présent *a priori* dans notre générateur, mais son impact sur la distribution en masse est très fortement diminué car la dépendance en $\hat{s}^4$ de l'amplitude carrée supprime fortement les bas $\hat{s}$.

de la mer (nous avons pris des quarks c) et nous avons forcé la largeur du Z à 20 GeV et nous lui avons donné une masse arbitraire de 400 GeV afin de nous placer dans des conditions favorables à l'apparition d'une queue à basse masse. Ceci revient finalement à étudier l'effet des partons de la mer sur la forme d'une Breit-Wigner dans le cas d'un couplage électro-faible neutre. Le résultat est montré sur la figure 2.11. Ce résultat est très similaire à ce que donne PYTHIA dans sa version de base pour le signal graviton pour des grandes masses et des grands couplages.

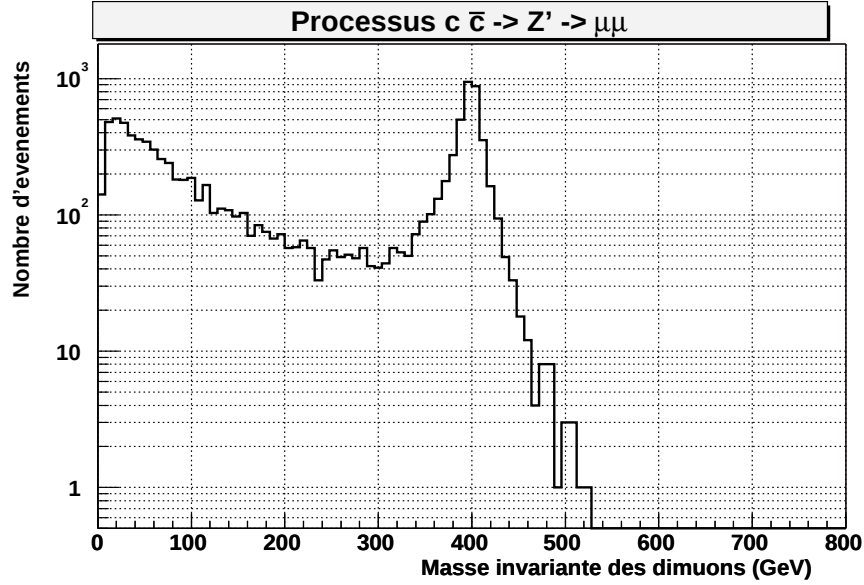


FIG. 2.11 – Distribution de masse invariante de dimuons dans l'hypothèse d'un boson Z de 400 GeV de masse et de largeur de 20 GeV produit par des quarks c de la mer dans des collisions $p\bar{p}$ à 1,96 TeV. La queue à basse masse est due à la convolution de la Breit-Wigner avec la densité de quarks c dans le proton qui tombe très rapidement à grandes valeurs de x .

Ainsi cette étude nous a permis de rectifier une coquille dans le code de PYTHIA qui, au final, change complètement la forme de la résonance ainsi que la valeur de la section efficace de production du graviton (tableau 5.1)! La section efficace est plus faible que ce qui avait été trouvé avec la version de base de PYTHIA. Cela affecte toutes les études effectuées sur la production résonnante de graviton de Kaluza-Klein dans le modèle de Randall-Sundrum, quel que soit le canal de désintégration, basées sur le générateur PYTHIA.

$M_G(\text{GeV})$	400		600		800	
contributions	avant	après	avant	après	avant	après
$\sigma(q\bar{q} \rightarrow G \rightarrow \mu\mu)$	$3,27 \cdot 10^{-1}$	$3,27 \cdot 10^{-1}$	$4,16 \cdot 10^{-2}$	$4,04 \cdot 10^{-2}$	$5,74 \cdot 10^{-3}$	$5,21 \cdot 10^{-3}$
$\sigma(gg \rightarrow G \rightarrow \mu\mu)$	$8,27 \cdot 10^{-2}$	$5,41 \cdot 10^{-2}$	$6,51 \cdot 10^{-3}$	$1,35 \cdot 10^{-3}$	$1,78 \cdot 10^{-3}$	$4,19 \cdot 10^{-5}$

TAB. 2.3 – Comparaison entre les sections efficaces (en pb) dans le canal muonique données par PYTHIA avant et après correction pour une valeur de $k/\bar{M}_{Pl}=0,1$.

2.3 Terme d'interférence dans le processus de fusion des quarks

Dans PYTHIA et HERWIG qui sont interfacés à l'environnement informatique de la collaboration, le processus de production de graviton dans le modèle de Randall-Sundrum est traité seul et aucun terme d'interférence n'est ajouté. Si l'on veut pouvoir utiliser ces générateurs sans omettre une contribution importante, il faut vérifier que le terme d'interférence entre la production de paires de muons médiée par un graviton et celle médiée par un Z ou un photon est négligeable.

Pour cela, nous avons recours au générateur indépendant décrit ci-dessus auquel on ajoute les amplitudes $q\bar{q} \rightarrow \gamma, Z \rightarrow \mu\mu$.

Nous utilisons les notations suivantes pour les amplitudes de probabilité :

processus	matrice de transition
$q\bar{q} \rightarrow Z^* \rightarrow \mu^+ \mu^-$	$T_{fi}^{(Z)}$
$q\bar{q} \rightarrow \gamma^* \rightarrow \mu^+ \mu^-$	$T_{fi}^{(\gamma)}$
$q\bar{q} \rightarrow G^* \rightarrow \mu^+ \mu^-$	$T_{fi}^{(G)}$

Nous calculons les matrices de transition T_{fi} pour chacun des trois processus, et nous en déduisons les termes suivants :

$$|\mathcal{M}_I|^2 = \left[\left(T_{fi}^{(Z)} + T_{fi}^{(\gamma)} \right) \times T_{fi}^{*(G)} \right] + \left[T_{fi}^{(G)} \times \left(T_{fi}^{*(Z)} + T_{fi}^{*(\gamma)} \right) \right] \quad (2.33)$$

$$|\mathcal{M}_{\gamma, Z+G}|^2 = \left(T_{fi}^{(Z)} + T_{fi}^{(\gamma)} \right) \times \left(T_{fi}^{*(Z)} + T_{fi}^{*(\gamma)} \right) + \left(T_{fi}^{(G)} \times T_{fi}^{*(G)} \right) \quad (2.34)$$

L'amplitude totale du processus étant $|\mathcal{M}_{tot}|^2 = |\mathcal{M}_I|^2 + |\mathcal{M}_{\gamma, Z+G}|^2$ où $|\mathcal{M}_I|^2$ représente le terme d'interférence entre le graviton et les champs du modèle standard et $|\mathcal{M}_{\gamma, Z+G}|^2$ correspond à la somme des processus individuels en incluant le terme d'interférence entre le Z et le photon.

La section efficace différentielle totale due à ce processus peut s'écrire :

$$\frac{d^3\sigma_X}{dx_1 dx_2 d\cos(\theta^*)} = \Phi \times |\mathcal{M}_{tot}|^2 \times \Delta P \times q(x_1)q(x_2) \quad (2.35)$$

$$\sigma_X = \sigma_{\gamma, Z} + \sigma_G + \sigma_I \quad (2.36)$$

Le terme d'interférence présente des valeurs négatives car c'est un terme qui n'est pas physique et ne peut être pris seul indépendamment des autres termes du processus, à savoir les contributions du modèle standard et du graviton pures. Pour cette étude, nous ne pouvons utiliser la partie *BASES* car ce programme ne sait pas intégrer des fonctions négatives, son but étant à l'origine d'intégrer des densités de probabilité qui sont positives. Nous utilisons cependant la partie de notre programme qui, pour des valeurs de x_1, x_2 et $\cos(\theta^*)$ tirées aléatoirement, calcule la section efficace différentielle $\sigma_{\gamma, Z} + \sigma_G$ et le terme σ_I seul. Ces valeurs ainsi que celles de x_1, x_2 et $\cos(\theta^*)$ sont écrites dans deux fichiers que

k/M_{Pl}	σ_I (pb)	σ_G (pb)
0,01	$8,42 \cdot 10^{-6}$	$4,358 \cdot 10^{-4} \pm 7,5 \cdot 10^{-5}$
0,05	$-3,71 \cdot 10^{-5}$	$9,995 \cdot 10^{-2} \pm 3,7 \cdot 10^{-3}$
0,1	$9,01 \cdot 10^{-5}$	$4,101 \cdot 10^{-1} \pm 6,5 \cdot 10^{-3}$

TAB. 2.4 – Contributions du terme d'interférence et du terme graviton pure à la section efficace de production de graviton de masse 400 GeV par fusion de quarks.

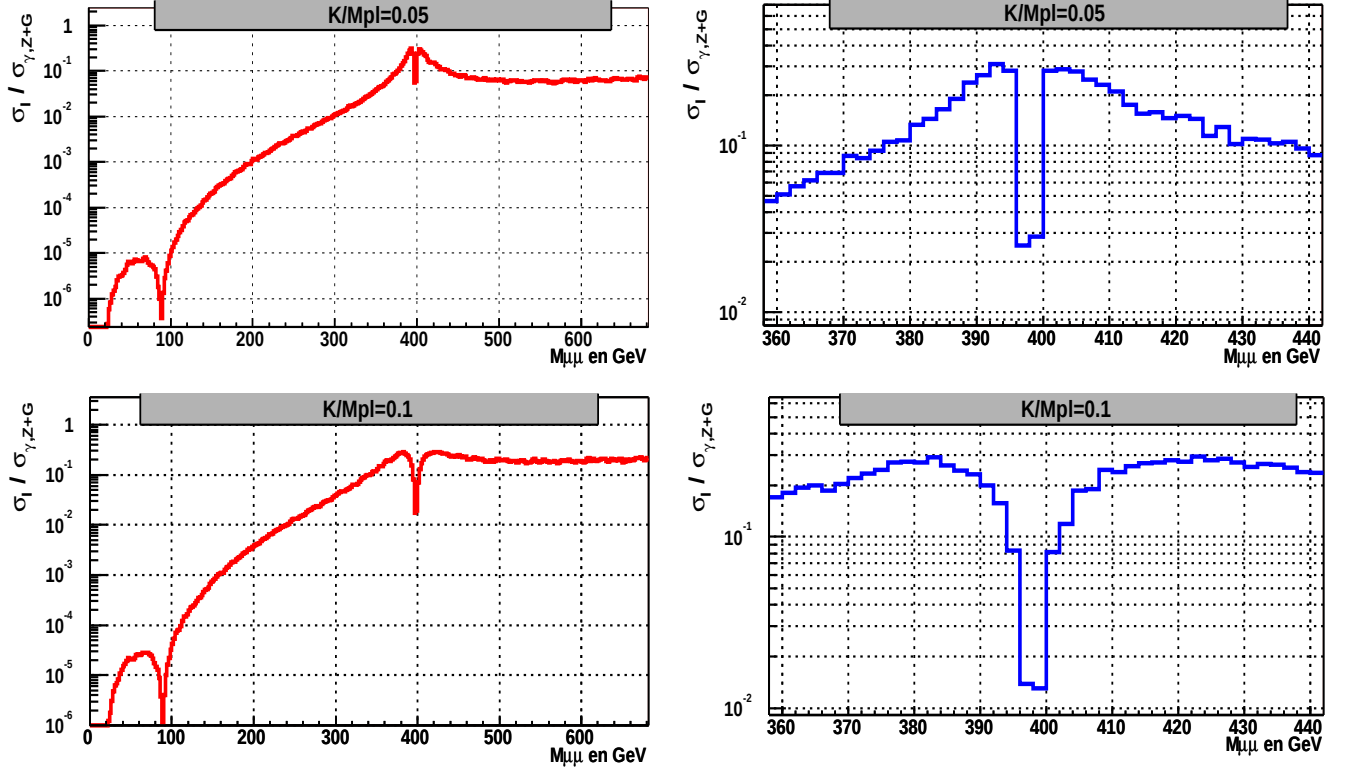


FIG. 2.12 – Rapport entre la contribution du terme d'interférence (pris en valeur absolue) et celle du graviton et des champs du modèle standard pour un graviton de masse de 400 GeV en fonction de la masse invariante des dimuons finals. De haut en bas les distributions correspondent à $k/\bar{M}_{Pl} = 0,05$ et $0,1$. La moitié gauche de la figure montre les distributions sur une large gamme d'énergie comprise entre 0 GeV et 700 GeV, la moitié droite est un agrandissement sur une gamme en énergie plus restreinte autour de 400 GeV. Les structures observées au voisinage de 400 GeV sont dues au fait que le terme d'interférence change de signe à la masse du graviton.

nous relisons par la suite.

L'intégration des différents termes est obtenue par une simple moyenne. Le tableau 2.4 compare la contribution du terme d'interférence seul à celui du processus de production de graviton seul par fusion de quarks. Notons que la contribution du terme d'interférence est bien en-deçà des erreurs dues à l'intégration numérique sur la valeur de la section efficace du graviton seul! On peut donc conclure que le terme d'interférence ne contribue pas de manière significative à la section efficace totale de production de graviton.

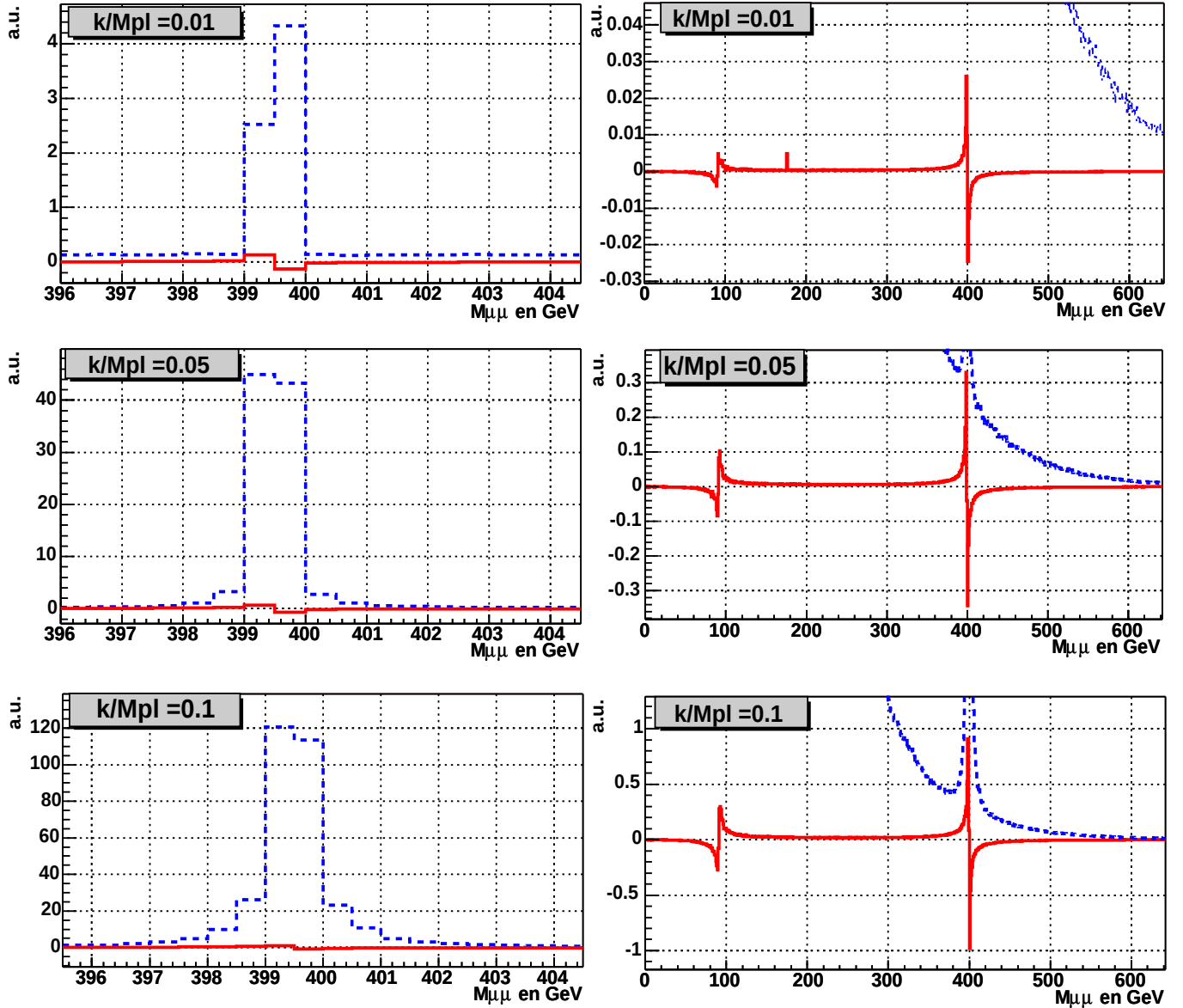


FIG. 2.13 – Contributions de l'interférence (en trait plein) et du processus total sans interférence (en pointillés) à la section efficace différentielle de production résonante de graviton par fusion de quarks. Les distributions sont réalisées dans le cas d'un graviton de masse de 400 GeV et pour trois valeurs de $k/\bar{M}_{Pl} = 0,01$, $0,05$ et $0,1$ en fonction de la masse invariante des dimuons finals. Les figures de droite diffèrent de celle de gauche seulement par la gamme de masse considérée.

Maintenant nous moyennons uniquement sur la variable $\cos(\theta^*)$. Ce que nous étudions à proprement parler est en fait $d\sigma/d\hat{s}$ où $\hat{s} = x_1 x_2 s$, $\sqrt{s}=1960$ GeV est l'énergie des protons dans le centre de masse. Nous comparons σ_I à $\sigma_{\gamma,Z+G}$. Les figures 2.12 et 2.13 montrent le résultat de cette étude pour un graviton de masse 400 GeV et pour deux couplages indicatifs $k/\bar{M}_{Pl}=0,05$ et $0,1$.

Sur la figure 2.12, nous voyons que le terme d'interférence est complètement négli-

geable devant la contribution des champs de modèle standard tant que nous restons hors de la “zone du pic” de la résonance de graviton. La partie droite de cette figure est un agrandissement de la “zone du pic” montrée sur la partie gauche.

La figure 2.13 montre que l’interférence change de signe avant et après la masse nominale du graviton. En dehors du pic, le terme d’interférence est complètement négligeable. Au voisinage du pic, l’interférence est non négligeable sur un intervalle de quelques GeV seulement, très petit devant les résolutions expérimentales.

En fait, on pourrait dire de manière qualitative que l’effet aurait pu être perçu dans le cas de résonances larges dans lesquelles le terme d’interférence aurait eu un effet de distortion marqué sur la distribution. C’est le cas pour l’interférence $Z\text{-}\gamma$, cela aurait pu l’être pour le graviton dans l’hypothèse de couplages du graviton aux champs du modèle standard environ 100 fois supérieures aux couplages permis par le modèle.

Ainsi nous avons montré que nous pouvons ignorer dans la suite de notre analyse la contribution du terme d’interférence.

2.4 Influence du choix des fonctions de distribution des partons

Les fonctions de distribution des partons ont été introduites dans la section 2.2.2.2. Le choix d’une fonction plutôt qu’une autre pourrait avoir une influence sur la forme de la résonance aussi bien que sur la valeur de la section efficace. Nous allons comparer la section efficace et la forme du signal de graviton en prenant trois autres *PDF* typiques représentatives des trois groupes CTEQ, MRS et GRV afin d’estimer la marge d’erreur que l’on commet en choisissant la paramétrisation CTEQ5L, c’est-à-dire celle qui est utilisée pour la génération de nos événements pour l’analyse dans les chapitres suivants.

PDF	$\sigma_q(pb)$	$\sigma_g(pb)$
CTEQ5L	$(3,404 \pm 0,11) 10^{-3}$	$(2,306 \pm 0,055) 10^{-5}$
MRST(c-g)LO	$(3,514 \pm 0,12) 10^{-3}$	$(2,906 \pm 0,121) 10^{-5}$
MRST(larger d/u) NLO	$(3,282 \pm 0,08) 10^{-3}$	$(2,691 \pm 0,084) 10^{-5}$
GRV98	$(3,029 \pm 0,12) 10^{-3}$	$(4,868 \pm 0,165) 10^{-5}$

TAB. 2.5 – Section efficace de production de graviton de masse 400 GeV pour un couplage de $k/\bar{M}_{Pl}=0,01$ pour chaque contribution de quarks et de gluons dans l’état initial.

Il est important de noter que ces erreurs sont les erreurs dues au choix des *PDF* et non l’erreur réelle sur l’estimation de la section efficace, ni même l’erreur théorique due à notre connaissance de la structure du proton⁶. Nous tiendrons compte toutefois de cette

6. Comparer les résultats obtenus avec différentes *PDF* ne prend pas en compte le fait que les différentes paramétrisations utilisent en grande partie les mêmes données ou font des hypothèses similaires quant à la forme analytique des fonctions de distribution. Depuis peu des outils existent pour évaluer les densités de partons et leurs erreurs [15]. Ils n’ont pas été implémentés dans cette étude basée, comme nous l’avons dit, sur la PDFLIB.

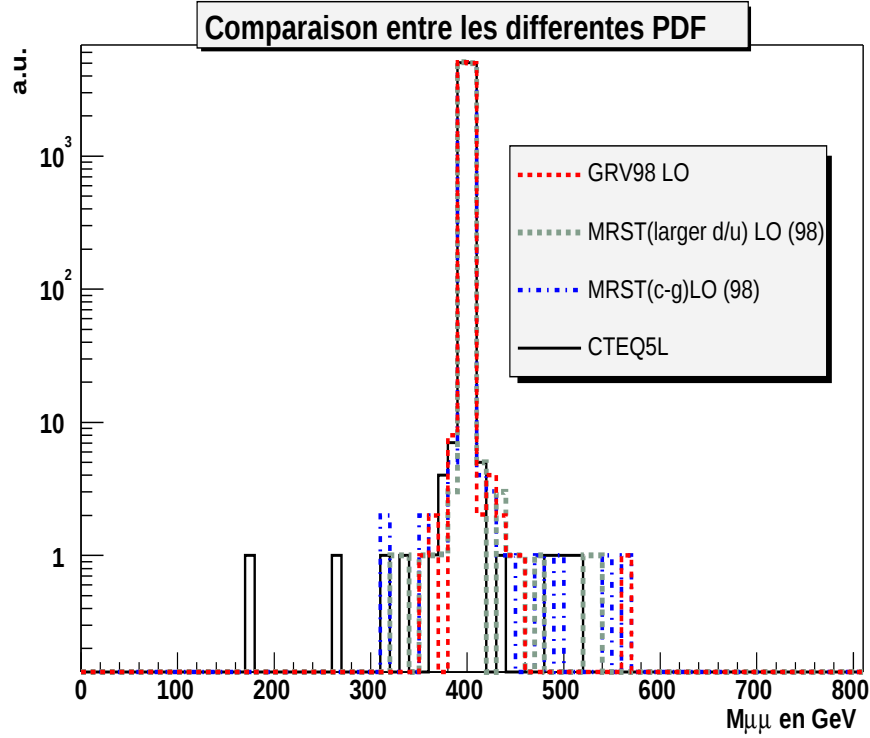


FIG. 2.14 – Comparaison entre les formes du signal données par les différentes PDF : CTEQ5L, MRST(c-g)LO, MRST(larger d/u)NLO et GRV98. Les distributions correspondent à un graviton de masse 400 GeV et un couplage $k/\bar{M}_{Pl}=0,01$ en fonction de la masse invariante des produits de désintégration.

erreur dans le calcul des erreurs systématiques lors de l'interprétation des résultats en termes de limites sur les paramètres du modèle.

Sur la figure 2.14, nous montrons les différentes distributions de la section efficace différentielle de production du graviton pour chaque PDF choisie. Elles sont toutes très similaires entre elles. Nous pouvons aussi noter que les sections efficaces sont compatibles entre chaque PDF. Nous montrons dans le tableau 2.5 les valeurs des sections efficaces de production de graviton en séparant les termes dus aux gluons et ceux dus aux quarks. Alors que la section efficace du processus $q\bar{q} \rightarrow G \rightarrow \mu\mu$ varie de moins de 10%, la section efficace du processus gluonique varie de manière significative. On s'y attend car le gluon à grand x est assez mal contraint. Mais ce processus contribue très peu à la section efficace totale.

2.5 Conclusion

Dans ce chapitre nous avons donné les règles de Feynman qui permettent de calculer les sections efficaces de production de graviton par fusion de quarks ou de gluons. Grâce à un générateur que nous avons développé, nous avons pu mettre en évidence une erreur dans le générateur PYTHIA. Nous avons donc corrigé le code de PYTHIA afin qu'il décrive

correctement la section efficace différentielle du graviton. Par la suite, nous utilisons cette nouvelle version de PYTHIA pour générer tous nos événements de signal qui servent dans l'analyse. La version 6.224 de PYTHIA, rendue publique début mai 2004, prend en compte cette modification.

Bibliographie

- [1] T.Han,J.D.Lykken,R.-J.Zhang, *Kaluza-Klein States from Large Extra Dimensions*
Phys.Rev. **D59**, 105006 (1999),
hep-ph/9811350
- [2] G.Giudice, R.Rattazzi, J.D.Wells, *Quantum Gravity and Extra Dimensions at High Energy Colliders*,
hep-ph/9811291
- [3] H.Davoudiasl, J.L. Hewett, T.G. Rizzo, *Experimental Probes of Localized Gravity: On and Off the Wall*,
Phys. Rev. **D63**: 075004 (2001),
hep-ph/0006041
- [4] H.Davoudiasl, J.L. Hewett, T.G. Rizzo, *Warped phenomenology*,
hep-ph/9909255,
Phenomenology of the Randall-Sundrum Gauge Hierarchy Model,
Phys.Rev.Lett. **84**, 2080 (2000).
- [5] www-cdf.fnal.gov/physics/exotic/run2/highmass-emucomb-2004/ee_mm_RS_Limit.eps
- [6] www-cdf.fnal.gov/physics/exotic/run2/highmass-jets-2003/limit_graviton_color.eps
- [7] T. Sjöstrand, *The PYTHIA (And JETSET) Web Page*,
<http://www.thep.lu.se/~torbjorn/PYTHIA.html> .
- [8] G.Marchesini,B.R.Webber,G.Abbiendi,I.G.Knowles,M.H.Seymour,L.Stanco,CAVENDISH-HEP-91-26, *HERWIG: A Monte Carlo Event Generator for Simulating Hadron Emission Reactions With Interfering Gluons*, VERSION 5.1 - APRIL 1991,
Comput.Phys.Commun.67 :465-508,1992
- [9] Jos Vermaseren, <http://www.nikhef.nl/~form> .
- [10] S.Kawabata, *A New Monte Carlo Event Generator for High Energy Physics*,
Comput.Phys.Commun. 41, 127-153 (1986).
S.Kawabata, *A New Version of the Multidimensional Integration and Event Generation Package BASES/SPRING*,
Comput.Phys.Commun. 88, 309-326 (1995).
- [11] H.Plochow-Besch, Computer Physics Commun. **75** (1993) 396, Int. J. Mod. Phys. **A10** (1995) 2901,
<http://consult.cern.ch/writeup/pdflib>.
- [12] T. Sjöstrand, communication privée.
- [13] J.Bijnens, P.Eerola,M.Maul, A.Månsson, T.Sjöstrand, *QCD Signature of Narrow Graviton Resonance in Hadron Colliders*,
Phys. Lett. **B503**,(2001) 341,
hep-ph/0101316

- [14] CTEQ web page
<http://www.phys.psu.edu/~cteq>
- [15] A.D. Martin, R.G. Roberts, W.J. Stirling, R.S. Thorne, *Uncertainties of predictions from parton distributions*,
Eur. Phys. Jour. **C28** (2003) 455-473,
hep-ph/0211080;
Uncertainties of predictions from parton distributions,
hep-ph/0308087;
The H1 Collaboration, Eur. Phys. Jour. **C30** (2003) 1-32) .

Chapitre 3

Le détecteur DØ

Introduction

Le Tevatron est un collisionneur proton/antiproton qui se situe à FNAL, Fermi National Laboratory situé à 80 km de Chicago. L'intérêt du concept de collisionneur hadronique est que l'on s'affranchit du problème de rayonnement synchrotron qui existe dans les collisionneurs électron-positron. En 1984 le Tevatron accéléra son premier faisceau test, mais ce n'est que huit ans plus tard qu'il devint opérationnel. Entre 1992 et 1996, période de prise de données appelée 'Run I', les détecteurs DØ et CDF, situés aux points de collision DØ et BØ du Tevatron, ont permis de faire des découvertes fondamentales pour la physique des particules. Ainsi, le 3 mars 1995 les collaborations CDF et DØ ont annoncé conjointement la découverte du quark top avec une masse de $174.3 \pm 5.1 \text{ GeV}/c^2$ [1], la masse du boson W fut mesurée avec une grande précision à $80.456 \pm 0.059 \text{ GeV}$ [2].

Afin d'opérer quelques modifications sur le système d'accélération, le collisionneur a été arrêté en 1996 pour redémarrer ses activités en mars 2001, entamant ainsi la période de prise de données 'Run II'.

Les modifications visent à augmenter la luminosité et l'énergie des faisceaux, ainsi que d'améliorer les détecteurs et leur électronique pour s'adapter à la plus grande luminosité instantanée et obtenir un meilleur rendement ainsi qu'une meilleure identification des particules.

Actuellement, l'expérience DØ réunit 73 instituts de 18 pays et plus de 600 physiciens.

Dans ce chapitre nous allons décrire d'une part le système d'accélération, et d'autre part le détecteur DØ grâce auquel nous avons pu mener une étude sur la recherche de graviton de Kaluza-Klein dans le modèle de Randall-Sundrum.

3.1 Le système d'accélération

Les protons et antiprotons sont créés et accélérés dans un complexe à plusieurs étages. La figure 3.1 en montre les principaux éléments. Durant la phase qui sépare le Run I du Run II, le système d'accélération a subi plusieurs modifications :

- la construction d'un nouvel accélérateur appelé *Main Injector* qui remplace le *Main Ring* ;
- la construction d'un nouvel anneau d'accumulation des antiprotons appelé *Recycler* ;
- l'augmentation de l'énergie des faisceaux d'environ 10% passant de 900 GeV par faisceau à 980 GeV ;
- l'augmentation du nombre de paquets de protons et d'antiprotons, passant de 6 à 36 paquets, ce qui induit une diminution du temps entre deux paquets, passant de 3,4 μs à 396 ns.

Nous allons décrire comment chacun des faisceaux de protons et d'antiprotons sont créés, stockés, puis injectés dans le collisionneur *Tevatron* et nous résumerons les aspects importants des différentes sous-parties de la chaîne d'accélération.

3.1.1 La création du faisceau de protons

Le faisceau de protons [3] est créé à partir d'une source d'hydrogène liquide. Le dispositif de Cockroft-Walton [4] extrait et accélère des ions H^+ jusqu'à une énergie cinétique de 750 keV. Ce faisceau est ensuite injecté dans le *Linac* [5], un accélérateur linéaire de

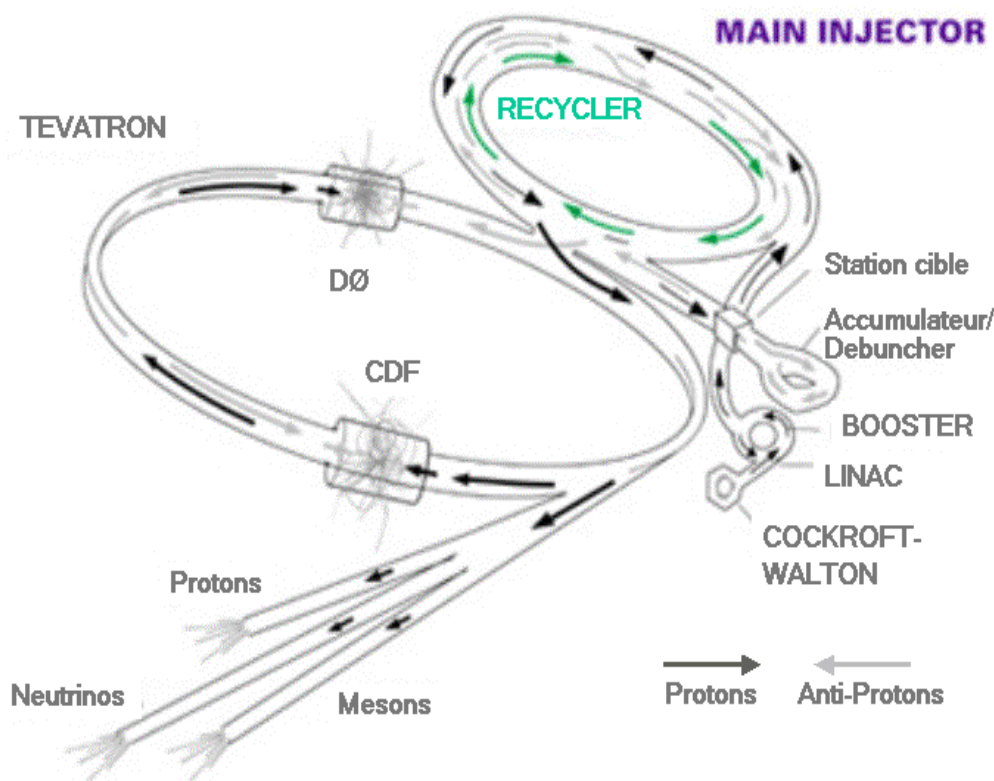


FIG. 3.1 – Chaîne de création et d'accélération des faisceaux de collision.

130 m de long, où ils atteignent une énergie cinétique de 400 MeV.

Les ions sont ensuite injectés dans le *Booster* [6] et sont épluchés à leur traversée d'une feuille de carbone située sur la première portion du Booster. Les protons résultants sont stockés dans l'anneau le temps que tous les ions H^- soient épluchés, c'est-à-dire environ $20 \mu s$, et traversent plusieurs fois la feuille de carbone sans être perturbés. A la fin de la phase d'épluchage, le paquet de protons ainsi stockés dans le Booster contient environ 10^{12} protons, accélérés à une énergie de 8 GeV avant d'être introduits dans l'injecteur principal (*Main Injector*) [7].

L'injecteur principal remplace le *Main Ring* qui existait au Run I. Il joue plusieurs rôles dans la constitution des faisceaux de collision :

- il accélère les protons à une énergie de 120 GeV soit pour des expériences sur cible fixe, soit pour constituer le faisceau source d'antiprotons ;
- il récupère en retour le faisceau d'antiprotons venant de l'accumulateur et le porte à une énergie de 150 GeV avant de l'envoyer vers le Tevatron ;
- il accélère les protons à une énergie de 150 GeV avant de les injecter dans le Tevatron où ils sont accélérés jusqu'à une énergie de 980 GeV .

Il peut en retour recevoir des faisceaux du *Booster*, du *Recycler*, du *Tevatron* ou de la source d'antiprotons.

A la sortie de l'injecteur principal, le faisceau de protons, d'énergie 150 GeV, est

injecté dans le *Tevatron* pour subir sa dernière accélération avant d'entamer la phase de collisions.

3.1.2 La création du faisceau d'antiprotons

La source d'antiprotons est constituée d'une cible de nickel bombardée par un faisceau discontinu constitué des paquets de protons issus du *Main Injector*. Les antiprotons sont produits selon la réaction $p p \rightarrow 3 p + \bar{p}$ dont le rendement est de l'ordre de un ou deux antiprotons pour 10^5 protons incidents.

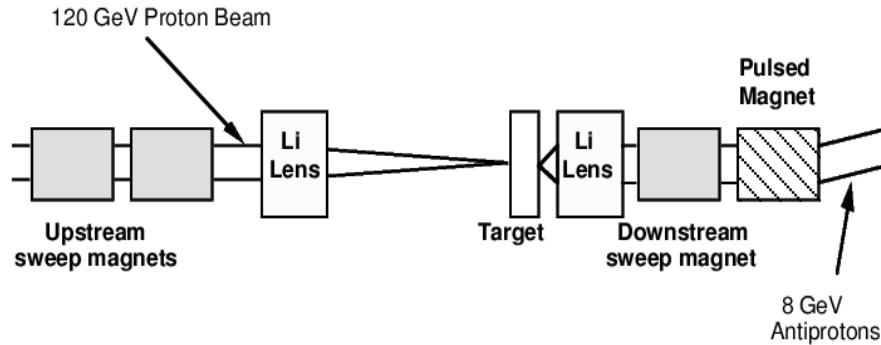


FIG. 3.2 – Dispositif de création d'antiprotons.

Lors de l'interaction des protons sur la cible, plusieurs particules secondaires sont produites, dont les antiprotons, avec une grande dispersion angulaire qui est une source importante d'inefficacité de production d'antiprotons. Les particules secondaires sont donc focalisées grâce à une lentille magnétique. La charge négative des antiprotons et leur énergie d'environ 8 GeV permet de les séparer des autres particules lors de leur passage dans un champ magnétique dipolaire. Leur trajectoire est donc incurvée, ce qui permet de les recueillir dans le *Debuncher* [8]. Les antiprotons sont recueillis par paquets (un paquet d'antiprotons par paquet de protons incident) dont l'énergie est distribuée selon une loi gaussienne autour de 8 GeV, et une cavité radiofréquence permet de ralentir les antiprotons très énergétiques et d'accélérer les plus lents, afin d'uniformiser l'énergie des protons à 8 GeV, c'est-à-dire de réduire la dispersion en énergie du faisceau, et de créer un faisceau continu d'antiprotons qui se substitue aux paquets d'antiprotons injectés initialement.

C'est ensuite que le faisceau est envoyé dans l'accumulateur où les paquets d'antiprotons sont formés, stockés, puis envoyés dans l'injecteur principal où ils sont accélérés à une énergie de 150 GeV. Une partie de ces antiprotons est envoyée vers le *Recycler*. C'est un anneau situé juste au-dessus de l'anneau du *Main Injector* dont il reçoit les protons, et qui a été initialement conçu pour remplir deux rôles :

- recevoir les antiprotons qui n'ont pas interagi dans le *Tevatron* afin de les stocker et les réutiliser pour un nouveau faisceau de collision. Cette opération devait permettre d'augmenter le nombre de particules par faisceau de 20%, et donc d'augmenter la luminosité instantanée (voir section 3.5) d'autant, mais cette opération perd de son intérêt à cause de la longue durée du recyclage et les difficultés techniques mises en oeuvre, c'est

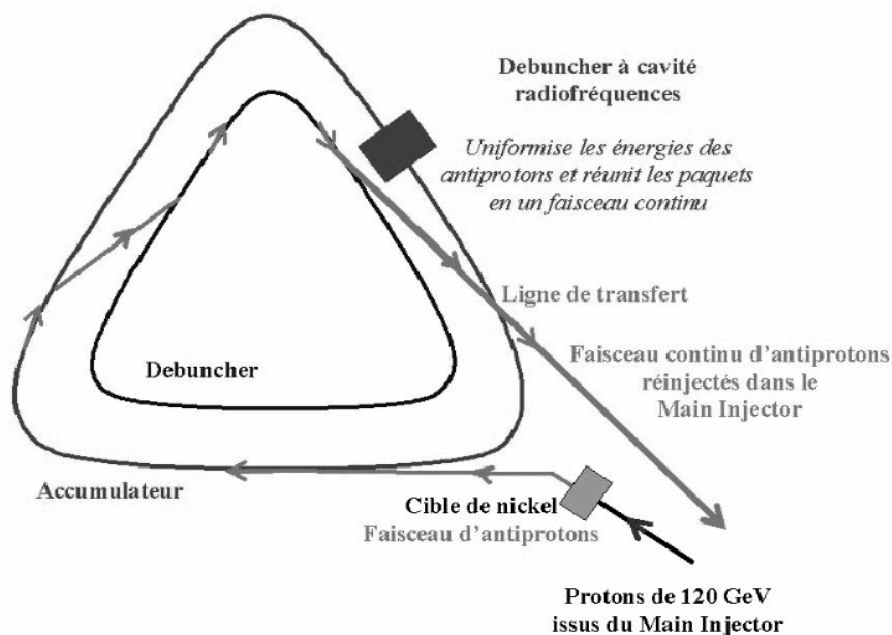


FIG. 3.3 – Schéma descriptif de l'Accumulateur/debuncher.

pourquoi ce rôle n'est finalement pas rempli par le *Recycler* ;

- stocker les antiprotons provenant du *Main Injector*. Plus ce dernier contient d'antiprotons et plus l'efficacité d'accumulation de nouveaux antiprotons diminue. C'est pourquoi le *Recycler* en prend une partie qu'il stocke pour une durée indéterminée, puis réinjecte dans le *MainRing*. Ce rôle est finalement le seul rempli par le *Recycler*.

A la sortie du *Main Injector*, le faisceau d'antiprotons est injecté dans le Tevatron qui les accélère à une énergie de 980 GeV.

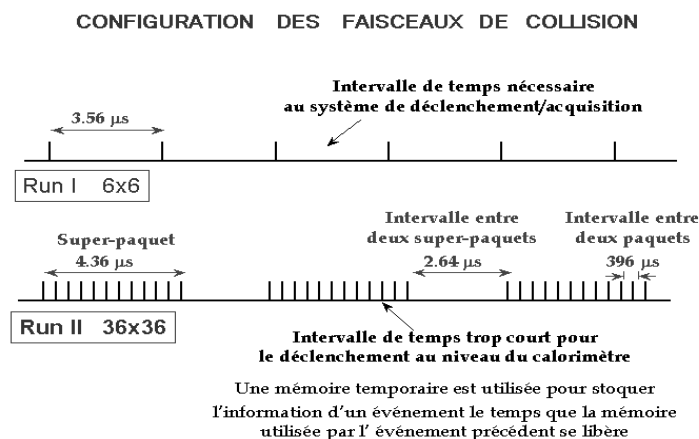


FIG. 3.4 – Description des paquets de particules dans le Tevatron.

3.1.3 La luminosité

3.1.3.1 La luminosité délivrée

La luminosité est un paramètre capital dans l'étude des interactions des particules à haute énergie auprès des accélérateurs. En effet, étant donné que le nombre d'événements observables N_{obs} est proportionnel à la luminosité intégrée $\mathcal{L}$, $N_{obs} = \epsilon \sigma \mathcal{L}$ où σ est la section efficace du processus étudié (en cm^2) et $\mathcal{L}$ est en cm^{-2} , plus la luminosité est grande, plus la probabilité d'observer des événements recherchés augmente. L'erreur statistique liée à ce nombre d'événements varie proportionnellement à $\sqrt{N_{obs}}$, c'est pourquoi il est très important de disposer d'un maximum de luminosité dans le choix de l'échantillon analysé.

Le Tevatron délivre une luminosité instantanée nominale $\mathcal{L}$ qui dépend des paramètres des faisceaux (voir tableau 3.1). La structure des faisceaux de collision est donnée sur la figure 3.4. La luminosité instantanée a pour expression :

$$\mathcal{L} = \frac{N_p \nu n_p n_{\bar{p}}}{2 \pi \beta^* (\epsilon_p + \epsilon_{\bar{p}})}$$

- N_p est le nombre de paquets ;
- ν est la fréquence de rotation des paquets et vaut 47,71 kHz ;
- n_p et $n_{\bar{p}}$ sont les nombres de protons et d'antiprotons par paquet ;
- β^* est liée à la longueur du paquet et vaut 35 cm ;
- ϵ_p et $\epsilon_{\bar{p}}$ sont les émittances transverses des paquets de protons et d'antiprotons, et sont liées à l'extension spatiale des faisceaux dans le plan perpendiculaire à l'axe des faisceaux ;
- La luminosité instantanée est de l'ordre de $\mathcal{L} \simeq 5 \cdot 10^{31} \text{ cm}^{-2} \text{ s}^{-1}$. Au Run I, la luminosité était de l'ordre de $\mathcal{L}_{Run\ I} \simeq 1 \cdot 10^{31} \text{ cm}^{-2} \text{ s}^{-1}$.

Paramètres des faisceaux de collision	Run I (1992-1996)	Run II (2001-2003) prévisions	Run II (atteint en 2003)
Energie de chaque faisceau (en GeV)	900	960	960
Nombre de paquets $p \times \bar{p}$	6×6	36×36	36×36
Nombre de protons par paquet	$2.3 \cdot 10^{11}$	$2.7 \cdot 10^{11}$	$2.14 \cdot 10^{11}$
Nombre d'antiprotons par paquet	$5.5 \cdot 10^{10}$	$3.0 \cdot 10^{10}$	$1.12 \cdot 10^{10}$
Taux de production d'antiprotons (h^{-1})	$0.6 \cdot 10^{11}$	$1.0 \cdot 10^{11}$	$1.0 \cdot 10^{11}$
Espacement entre les paquets (en ns)	3500	396	396
Longueur des paquets (dispersion en cm)	60	37	37
Emittance des protons (en mm mrad)	23π	20π	20π
Emittance des antiprotons (en mm mrad)	13π	15π	15π
Nombre d'interactions par croisement	2.5	2.3	2.3
Temps entre chaque croisement (ns)	3500	396	396
Luminosité instantanée (en $\text{cm}^{-2} \text{ s}^{-1}$)	$1.6 \cdot 10^{31}$	$8.6 \cdot 10^{31}$	$4.5 \cdot 10^{31}$

TAB. 3.1 – Paramètres des faisceaux de collision.

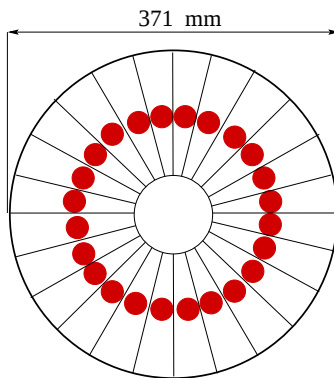


FIG. 3.5 – Schéma représentant les 24 scintillateurs triangulaires constituant un des deux disques du luminomètre, les points au centre de chaque triangle représentent les photomultiplicateurs.

3.1.3.2 La luminosité reçue

Pour la mesure de la luminosité reçue par l'expérience DØ des détecteurs de luminosité sont positionnés très près des faisceaux de collision, aux deux extrémités nord et sud du calorimètre du détecteur DØ, à une distance de ± 140 cm du centre de la zone de collision des faisceaux. Ces luminomètres consistent en deux scintillateurs en plastique constitués chacun de 24 scintillateurs triangulaires disposés en disques (figure 3.5) associés à des photomultiplicateurs. Ils détectent les particules légèrement déviées de l'axe des faisceaux après interaction. La couverture angulaire des scintillateurs s'étend à des valeurs de η (θ) (section 3.2.2) comprises entre 2,7 ($15,4^\circ$) et 4,4 ($2,1^\circ$).

Pour calculer la luminosité reçue, on mesure le nombre d'événements pour les processus de sections efficaces connues (voir tableau 3.2), ces processus étant de plusieurs ordres de grandeur dominants, le nombre d'événements observés divisé par leur section efficace totale σ_{eff} donne la luminosité reçue.

$$L_{re\acute{c}ue} = N_{\text{evt observés}} / \sigma_{eff}$$

où $\sigma_{eff} = \epsilon (A_{CD} \sigma_{CD} + A_{SD} \sigma_{SD} + A_{DD} \sigma_{DD}) = 43,26 \pm 2,07$

et $\epsilon = 0,907 \pm 0,02$ est l'efficacité des luminomètres.

Les effets d'acceptance géométrique des détecteurs de luminosité sont pris en compte pour chaque processus grâce aux paramètres A_{SD} , A_{DD} et A_{CD} donnés dans tableau 3.2.

Processus	section efficace (en mb)	Acceptance
Collision dure (CD)	$46,69 \pm 1,63$	$0,97 \pm 0,02$
Simple diffractif (SD)	$9,57 \pm 0,43$	$0,15 \pm 0,05$
Double diffractif (DD)	$1,29 \pm 0,20$	$0,72 \pm 0,03$

TAB. 3.2 – Processus pris en compte pour le calcul de la luminosité.

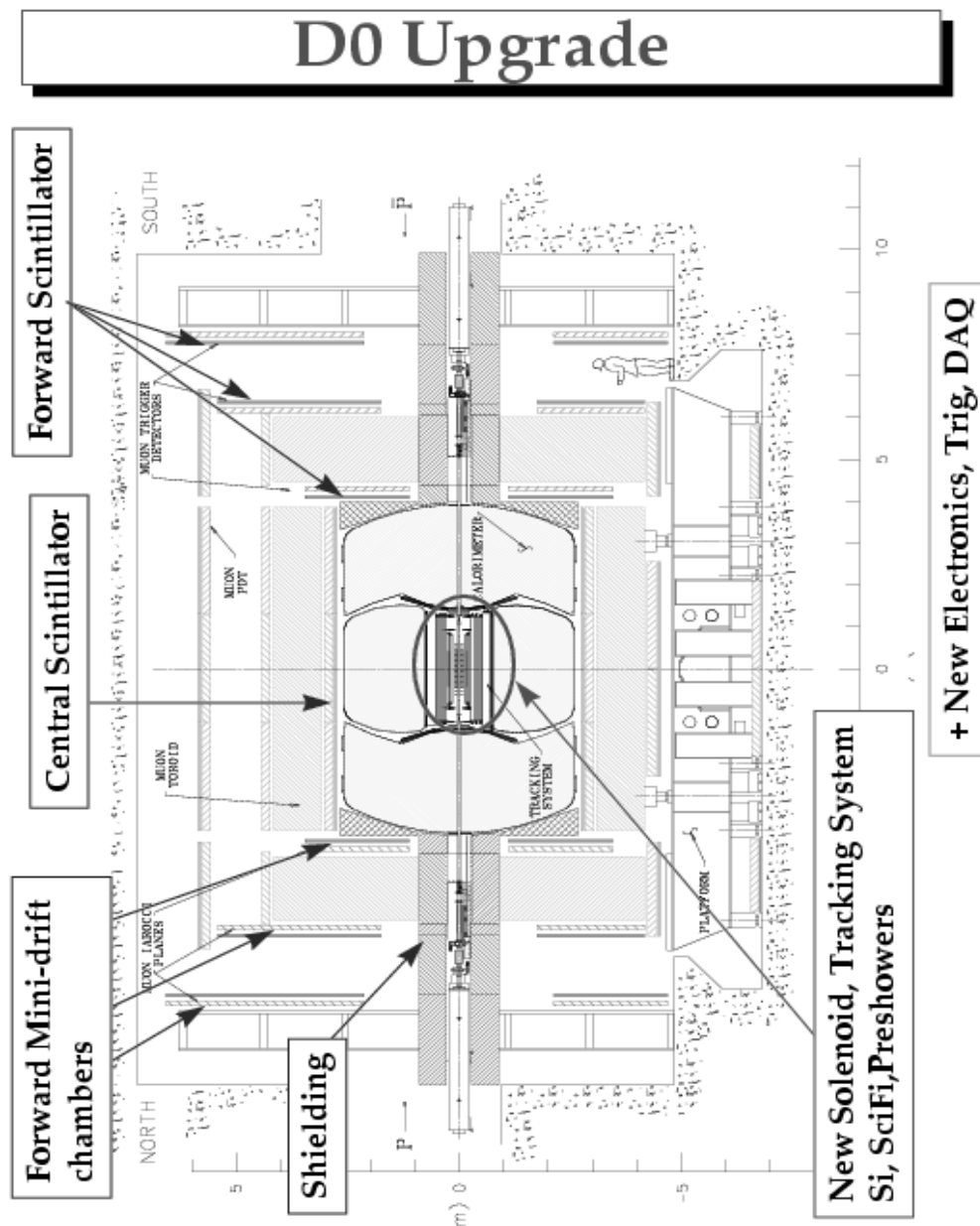


FIG. 3.6 – Vue générale du détecteur.

3.2 Le détecteur DØ

3.2.1 Les modifications par rapport au Run I

Durant les cinq années qui ont suivi la fin du Run I, le détecteur DØ a subi plusieurs modifications, aussi bien au niveau de l'électronique que sur des éléments du détecteur,

certaines sous-détecteurs étant tout à fait nouveaux, le but étant de permettre au détecteur de supporter une plus forte luminosité et d'avoir de meilleures performances de détection et de reconstruction des événements.

Le taux d'interaction étant beaucoup plus élevé et le rayonnement plus fort qu'au Run I, il a fallu modifier le détecteur en conséquence par l'ajout de blindages, notamment autour du tube à vide (voir section 3.2.6). D'autres modifications aussi bien sur le détecteur que sur son électronique se sont ajoutées afin de permettre une meilleure reconstruction des événements.

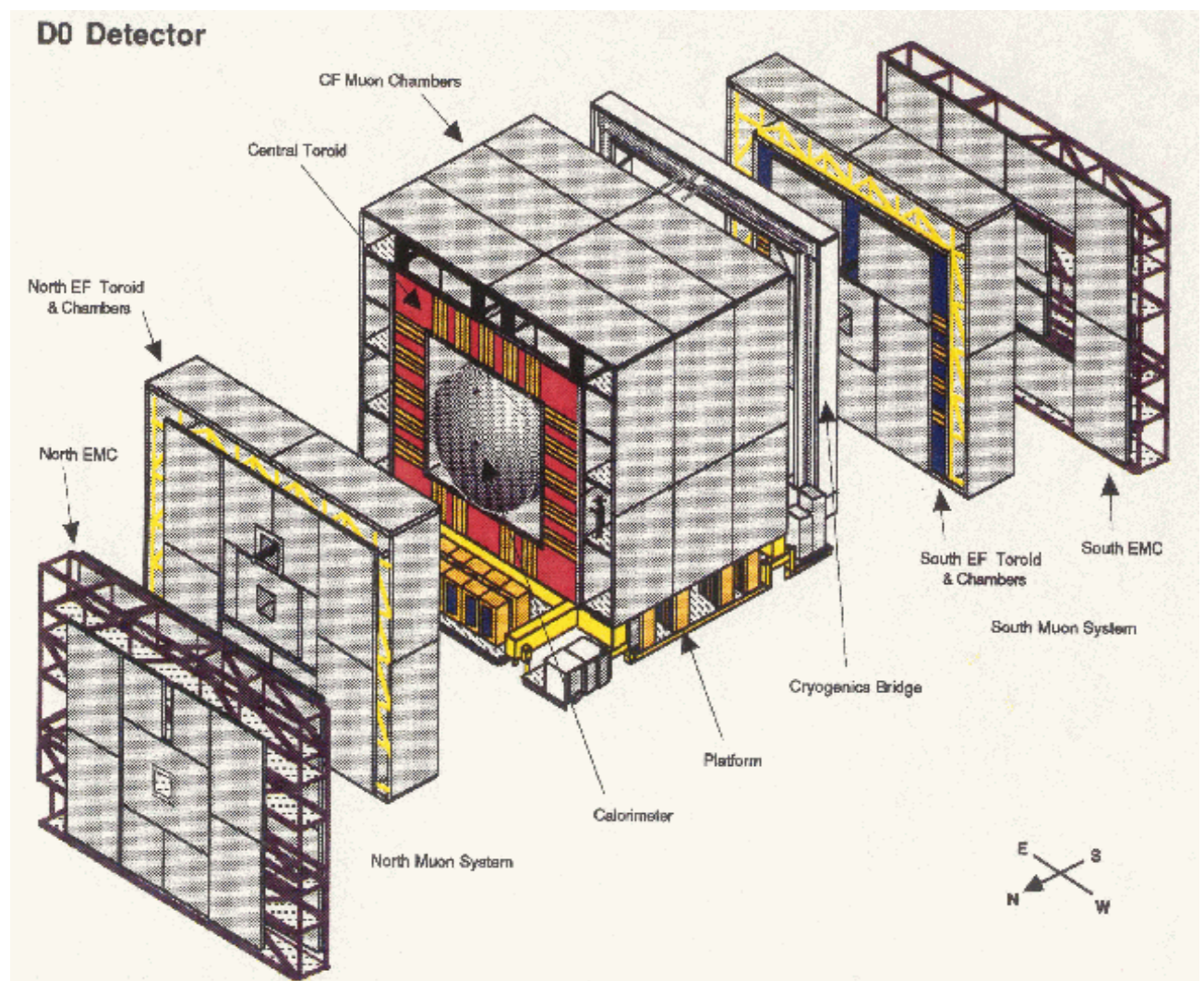


FIG. 3.7 – Schéma éclaté du détecteur. On distingue le système à muons (partie grisée), le toroïde et le calorimètre.

Les modifications du détecteur pour la nouvelle période de prise de données comprennent :

- le changement du détecteur de traces central ; il est remplacé par un détecteur cylindrique composé d'un détecteur de traces au silicium, entouré d'un détecteur à fibres scintillantes, puis d'un solénoïde qui crée un champ magnétique de 2 teslas.

- le changement de l'électronique du calorimètre ainsi que l'addition de détecteurs de pieds de gerbes ; le calorimètre lui-même n'a pas été modifié ;
- les chambres à muons avant/arrière sont remplacées par des mini-chambres à dérive ainsi que 3 couches de scintillateurs ;
- le système de déclenchement est complètement modifié, compte tenu des contraintes temporelles imposées par le plus grand taux d'interaction.

3.2.2 Le repère géométrique

Le repère géométrique utilisé dans la collaboration DØ est le repère cartésien direct, le centre O du repère coïncide avec le centre géométrique du détecteur, c'est-à-dire le centre du détecteur de vertex, le SMT.

L'axe (Oz) est orienté le long des faisceaux de collision, dans le sens du faisceau de protons, et l'axe y est orienté vers le haut. Les conventions concernant les variables sphériques sont les suivantes :

- la distance r à l'axe (Oz) du point considéré ;
- l'angle azimuthal ϕ orienté dans le sens positif dans le plan (Oxy) et vaut zéro le long de l'axe (Ox) ;
- l'angle polaire θ qui vaut zéro le long de l'axe (Oz).

3.2.3 Les variables cinématiques

Les collisions étant frontales, la composante de l'impulsion totale dans le plan transverse est nulle :

$$p_{Ti} = \sqrt{\left(\sum_i p_{xi}\right)^2 + \left(\sum_i p_{yi}\right)^2}$$

Cette variable est très importante pour reconstruire l'énergie manquante qui correspondrait à une particule non chargée comme le neutrino ou autre nouvelle particule.

En collision hadronique, l'impulsion longitudinale dans le laboratoire n'est pas nulle, à cause de la fraction d'impulsion emportée par chacun des quarks en collision, ce qui rend l'angle θ non invariant par une transformation de Lorentz, et complique considérablement le calcul des quantités qui dépendent de θ , notamment le calcul des sections efficaces de différents processus physiques dont les quantités physiques sont toujours exprimées dans le référentiel du centre de masse. La *rapidité* y qui apparaît plus naturellement dans les expressions mathématiques des sections efficaces est une fonction de l'angle θ , et est invariante de Lorentz. Elle a pour expression :

$$y = \frac{1}{2} \ln \frac{E + p_z}{E - p_z}$$

dont l'expression simplifiée dans le cas de particules ultra-relativistes est donnée par :

$$\eta = -\ln \tan(\theta/2) \quad (3.1)$$

où η est la *pseudo-rapacité*. Etant donné que les particules dans l'état final sont quasiment toujours ultra-relativistes, c'est donc cette dernière variable η qui est utilisée en référence à la coordonnée polaire.

Dans ce manuscrit nous parlerons tantôt de η_{phys} , tantôt de $\eta_{dét}$.

La variable η_{phys} est reliée par la relation (3.1) à l'angle θ physique de la particule, c'est-à-dire l'angle entre l'axe orienté (Oz) et l'impulsion de la particule.

La variable $\eta_{dét}$ est calculée à partir de l'angle θ donné par la position de la particule dans le repère du détecteur. Elle sert à décrire la couverture angulaire de chaque détecteur ainsi qu'à localiser la particule sur le détecteur et diffère selon que l'on traite le CFT (section suivante) ou les couches du système à muons, ou même des cellules du calorimètre.

Cette distinction est due au fait que la position du vertex principal de l'interaction suit une loi gaussienne le long de l'axe (Oz), centrée en O et dont la dispersion à 1σ est d'environ 25 cm. Dans la suite, nous parlerons toujours de η comme étant $\eta_{détect}$, sauf si l'on précise que c'est la variable physique qui sera utilisée.

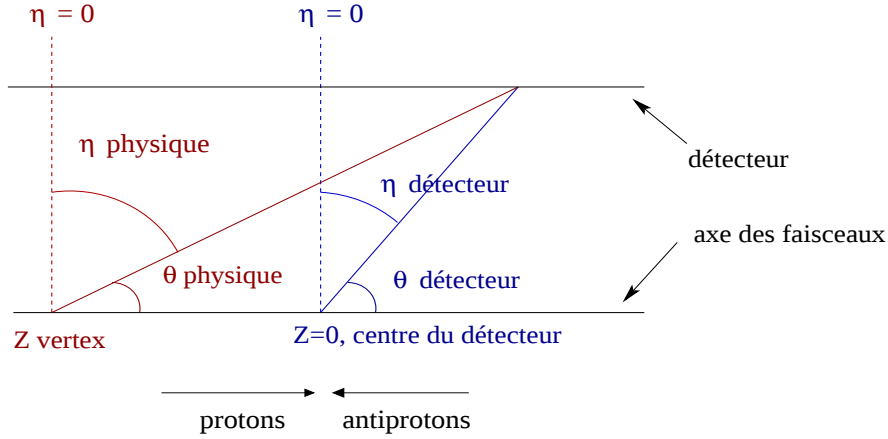


FIG. 3.8 – Différence entre $\eta_{physique}$ et $\eta_{détecteur}$.

3.2.4 Le détecteur de traces central

Il est composé de deux parties : un détecteur de traces au silicium, le SMT [12], qui joue aussi le rôle de détecteur de vertex, et un détecteur à fibres scintillantes, le CFT [13], qui permet de reconstruire les traces des particules chargées. Ces deux détecteurs sont plongés dans le champ magnétique $\vec{B} = B_Z \vec{e}_Z$ créé par le solénoïde qui les entoure, ce qui permet de reconstruire les impulsions des particules d'après la courbure de la trace.

3.2.4.1 Le détecteur de vertex SMT

Il est composé de trois types de détecteurs :

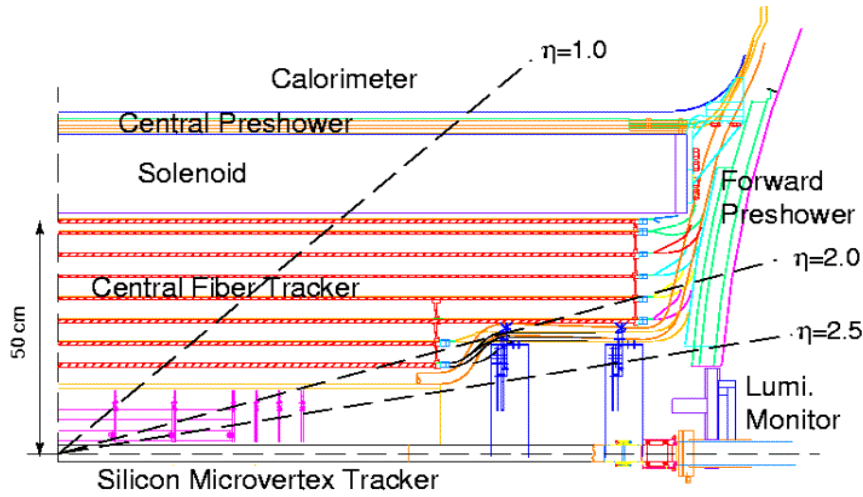


FIG. 3.9 – Schéma descriptif du détecteur de traces central. Autour de l'axe du faisceau on distingue de bas en haut : le SMT, le CFT, le solénoïde.

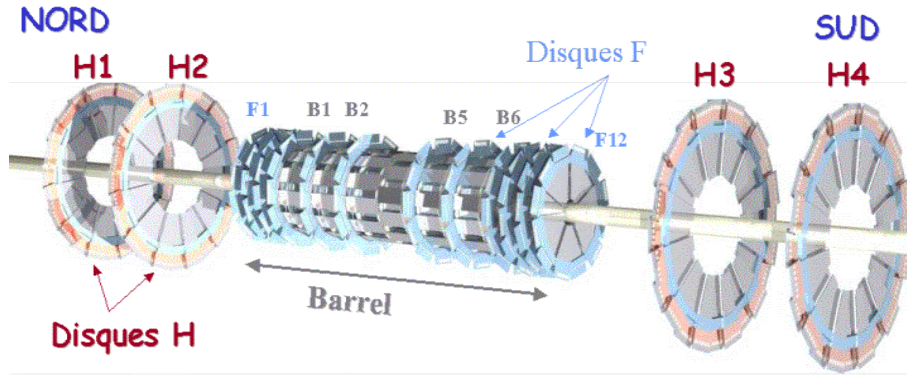
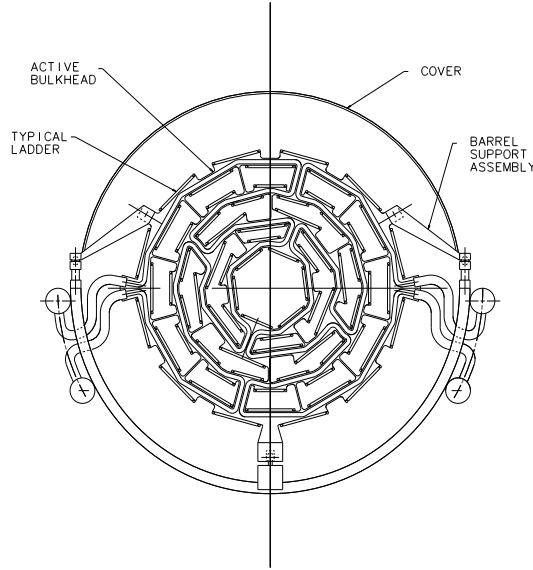
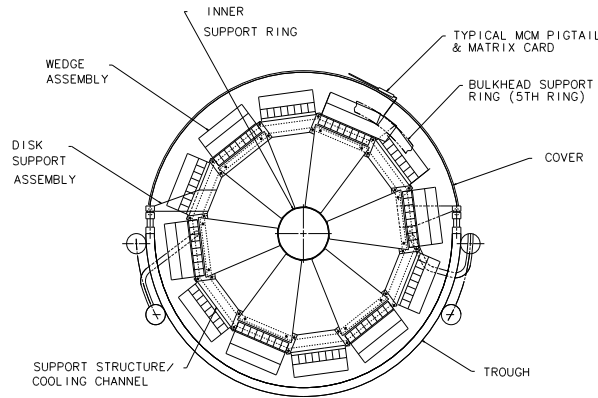


FIG. 3.10 – Schéma descriptif des différents modules du SMT.

- six systèmes cylindriques alignés appelés *barrels*. Ils sont numérotés de 1 à 6, du nord du détecteur au sud, formant un tube dont le centre se situe entre les *barrels* B_3 et B_4 . Chaque *barrel* mesure 12 cm de long, a un rayon intérieur de 2.7 cm et un rayon extérieur de 9,4 cm, et est constitué de 4 couches concentriques double-face de détecteurs en silicium parallèles à l'axe des faisceaux, totalisant 72 micro-pistes alignées à $10\ \mu\text{m}$ près et fixées sur un support en Beryllium (Fig.3.11). Chaque trace laissée par une particule chargée contient 4 à 8 coups dans les *barrels* du SMT. La géométrie des *barrels* permet de reconstruire des traces de $|\eta| < 1,5$ avec une grande précision ;

- douze disques F disposés de la façon suivante : quatre à chaque extrémité Sud et Nord du tube de *barrels*, et quatre disques entre chaque *barrel*, sauf au centre, entre les *barrels* B_3 et B_4 , où l'interstice est laissé libre. Ces disques ont un rayon intérieur de 2.6 cm et un rayon extérieur de 10.5 cm, et sont composés chacun de 12 trapèzes de silicium double-face se chevauchant. Leur intérêt est de renforcer la détection des traces à grands η (à faibles angles), typiquement $|\eta| < 3$.

- deux disques H de chaque côté des disques F terminaux. Chaque disque H est formé

FIG. 3.11 – *Vue en coupe d'un barrel double-face du SMT.*FIG. 3.12 – *Schéma d'un disque F.*

de 24 plaques trapézoïdales doubles-faces en silicium intercalées entre elles et se chevauchant. Ils sont très similaires aux disques F, mais de taille, d'orientation et d'électronique différentes. L'angle de 7.5° entre deux plaques trapézoïdales successives permet de reconstruire la position d'un coup en trois dimensions. Les disques H viennent compléter le système des barrells et des disques F, leur position éloignée par rapport au centre du détecteur permet d'augmenter la précision de mesure pour des traces à grand η (voir figure 3.13).

3.2.4.2 Le détecteur de traces central CFT

Il est constitué de huit cylindres concentriques. Les deux cylindres intérieurs mesurent 1,8 m de long et les six autres cylindres mesurent 2,6 m. Le rayon interne est de 20 cm et le rayon externe est de 50 cm (voir figure 3.14) . Sur chaque cylindre sont montées

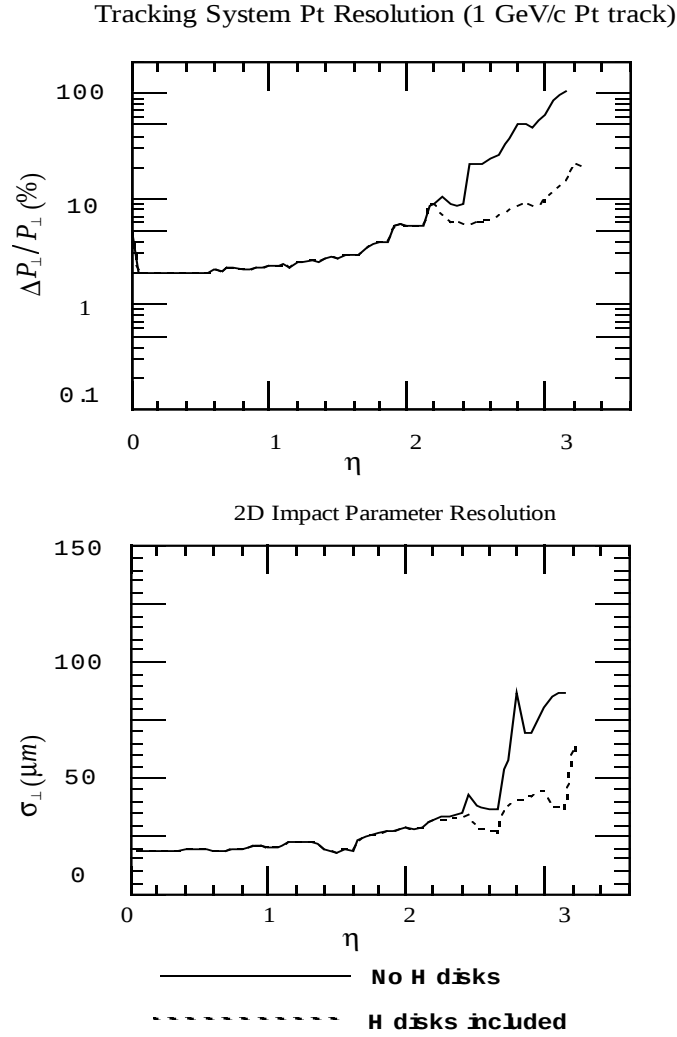


FIG. 3.13 – Résolution sur l'impulsion et sur le paramètre d'impact reconstruits par le SMT. Comparaison entre le SMT sans inclure les disques H (en pointillés) et avec les disques H (en traits pleins).

des fibres scintillantes de $835 \mu\text{m}$ de diamètre assemblées par paires : une double couche de fibres parallèles à l'axe du faisceau, et une double couche de fibres inclinées de ± 2 degrés par rapport à cet axe. Cette disposition en (u,v) permet une reconstruction en 3 dimensions des coups laissés par le passage des particules chargées dans le CFT. La précision de mesure obtenue est de $100 \mu\text{m}$ dans le plan (r, ϕ). Afin de préserver cette résolution, la position des fibres est mesurée à $50 \mu\text{m}$ près.

La lumière produite par la scintillation des fibres après le passage d'une particule est récoltée par des guides d'ondes connectés aux VLPC (*Visible Light Photon Counter*). Le CFT compte au total 71 680 canaux électroniques. La haute granularité obtenue grâce au petit diamètre des fibres est importante à la fois pour la reconstruction des traces et pour le système de déclenchement qui nécessite une réponse rapide (efficacité atteinte supérieure à 99,5% par fibre).

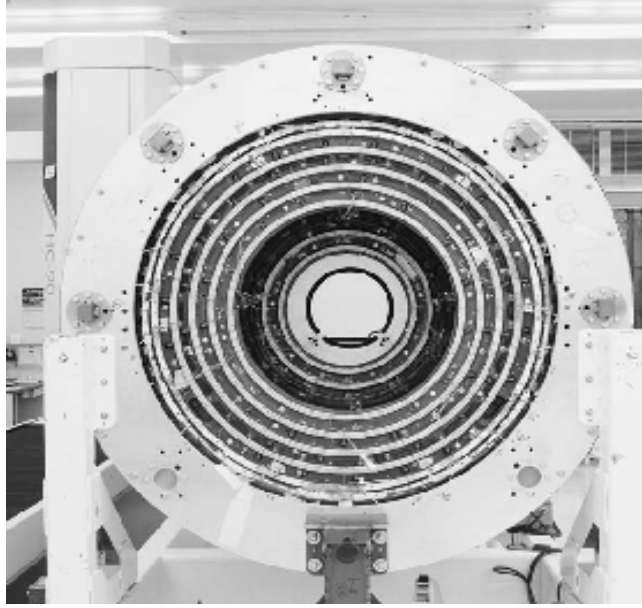


FIG. 3.14 – Photographie du CFT prise en bout. On distingue huit cercles concentriques correspondant aux huit couches cylindriques de doublets de fibres scintillantes.

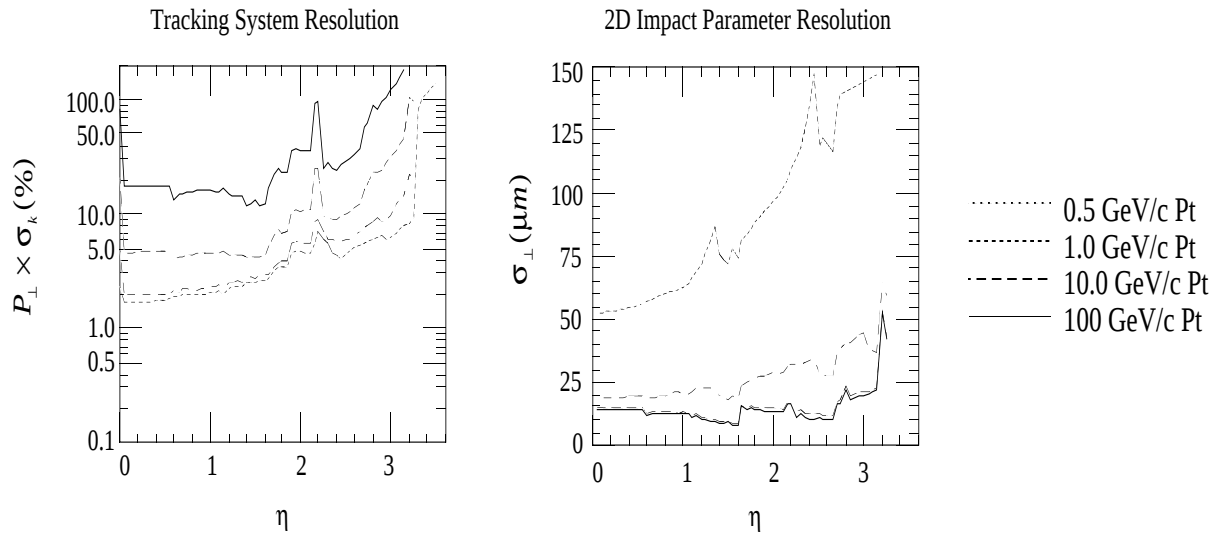


FIG. 3.15 – Résolution sur l'impulsion transverse des traces et sur le paramètre d'impact dans le plan transverse reconstruits par le système SMT, CFT et solénoïde en fonction de la pseudorapacité η et de P_T .

3.2.4.3 Le solénoïde

C'est un aimant supraconducteur qui a été conçu pour optimiser la précision de mesure en impulsion du détecteur de traces centrales, en fournissant un champ magnétique de 2 T uniforme à 0,5% près. Il entoure le détecteur de traces central, et mesure 2,73 m de long sur 1,42 m de diamètre.

3.2.4.4 Performances du détecteur de traces central

La résolution sur l'impulsion des traces est représentée sur la figure 3.15. Pour $|\eta| = 0$, elle peut être paramétrée comme suit :

$$\sigma_{p_T}/p_T = (0,015^2 + (0,014 p_T)^2)^{1/2}$$

L'alignement entre le CFT et le SMT est connu à $40 \mu\text{m}$ près. La position d'une trace est connue à $100 \mu\text{m}$ près et la position du vertex sur l'axe (Oz) à $35 \mu\text{m}$.

La reconstruction du vertex primaire est montrée sur la figure 3.16. Nous pouvons remarquer que les positions des vertex en x et en y ne sont pas centrées en zéro. Ceci est dû au fait que le faisceau peut être excentré par rapport au centre du repère géométrique.

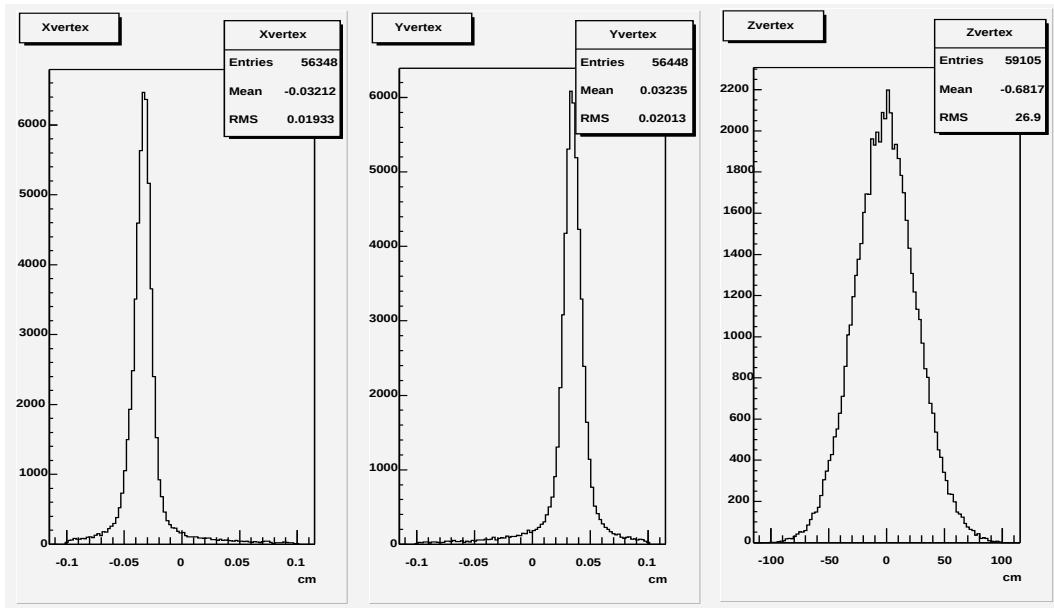


FIG. 3.16 – Distribution spatiale de la position du vertex primaire (en cm). Ces courbes ont été réalisées avec des données comprenant deux muons.

3.2.5 Le système calorimétrique

Le système calorimétrique compte trois parties complémentaires :

- le détecteur de pieds de gerbes qui est un élément nouveau apporté en complément du système calorimétrique du Run I, pour compenser les effets de la perte d'énergie des électrons et photons dans le solénoïde ;
- le calorimètre à argon liquide (c'est un calorimètre à compensation) qui est le même qu'au Run I ; seule l'électronique d'acquisition a changé pour tenir compte de la plus forte luminosité et du temps de croisement des faisceaux plus faible qu'au Run I, nécessitant tout une phase de réétalonnage du détecteur, réajustement des portes électroniques et des seuils de déclenchement ;
- les détecteurs intercyostat.

Nous allons brièvement décrire ces différentes parties.

3.2.5.1 Le détecteur de pieds de gerbes

Le rôle principal du détecteur de pieds de gerbes (*PreShower*) est de calculer avec une précision meilleure que celle du calorimètre la position des électrons et des photons.

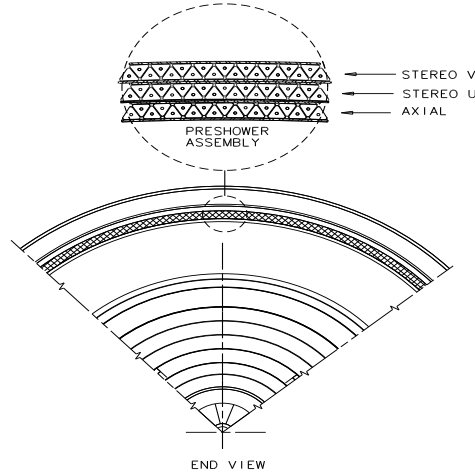


FIG. 3.17 – Schéma du détecteur de pieds de gerbes.

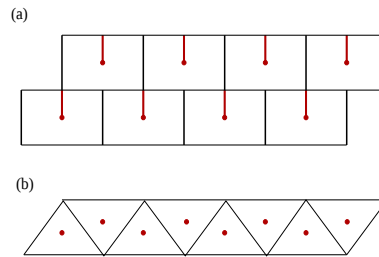


FIG. 3.18 – Schéma du détecteur de pieds de gerbes. La figure a) montre une vue longitudinale où les traits épais représentent les fibres scintillantes. La figure b) montre une vue en coupe transversale où l'on distingue la forme triangulaire de chaque tube, et les points représentant les fibres scintillantes.

Il est constitué d'un *PreShower* central, le CPS, qui couvre la région $|\eta| < 1,3$, et un *PreShower* avant, le FPS, qui couvre la région $1,4 < |\eta| < 2,5$. Le CPS et le FPS sont placés contre la première couche du calorimètre, juste après le solénoïde. Ils sont formés de deux couches de fibres scintillantes triangulaires (figures 3.17 et 3.18) qui s'imbriquent les unes contre les autres sans laisser d'espace vide. Ainsi, une particule traverse au moins deux fibres, ce qui permet d'avoir une très bonne précision sur la mesure de sa position, typiquement $600 \mu\text{m}$ pour un muon et 1.4 mm pour un électron.

3.2.5.2 Le calorimètre

Le calorimètre à argon liquide est un calorimètre à échantillonnage, composé d'une succession de plaques d'uranium apauvri (ou de cuivre pour la régions les plus externes), qui constitue le milieu absorbeur, et de couches d'argon liquide. Le tout est plongé dans

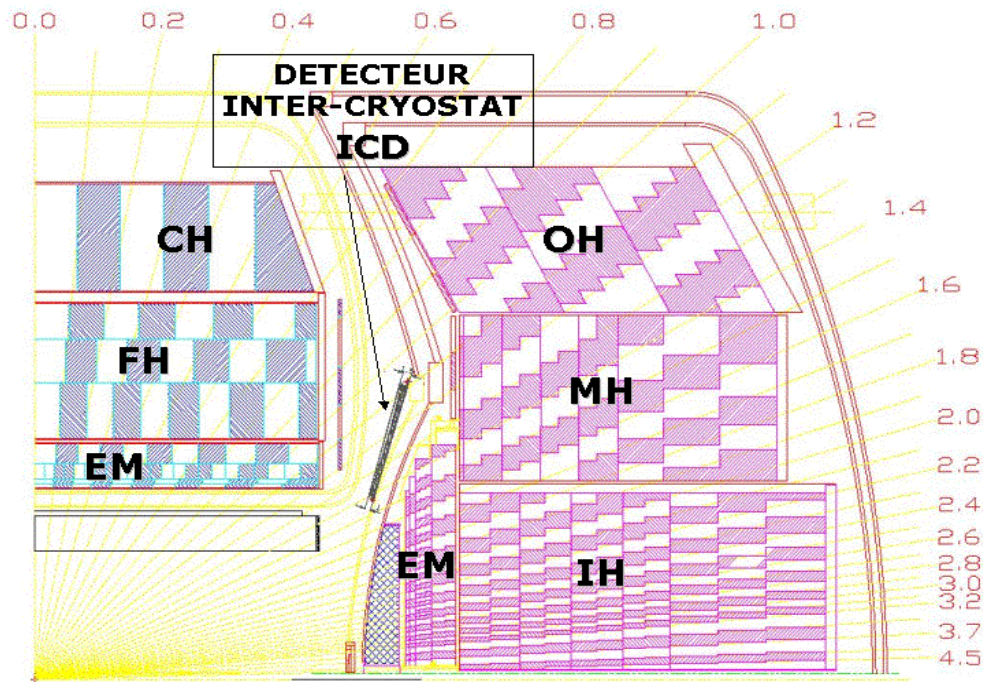


FIG. 3.19 – Schéma descriptif du calorimètre. On distingue le bloc central (CH, FH et EM central) du bloc avant (OH, MH, IH et EM à droite). Les parties nommées IH, MH et OH sont respectivement les parties intérieure, médiane, extérieure du calorimètre hadronique du bloc avant.

un cryostat maintenu à une température de 78 K. L'argon liquide a été choisi parce qu'il est stable, très pur, insensible aux radiations et son étalonnage est simple.

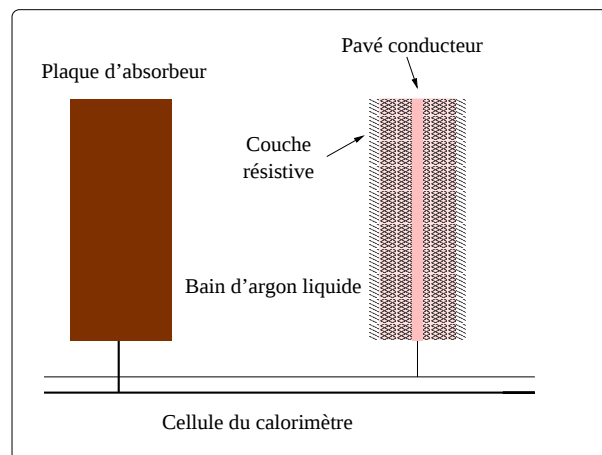


FIG. 3.20 – Schéma d'une cellule du calorimètre. Une succession de cellules reproduit l'alternance de couches d'absorbeur en uranium et d'argon liquide.

Il est constitué de trois blocs : un bloc central, noté CC, couvrant la région $|\eta| < 1,0$, et deux blocs vers l'avant, notés EC, couvrant la région $0,7 < |\eta| < 4,5$. Chaque bloc se compose d'une partie électromagnétique (Electro-Magnetic ou EM), une partie hadronique fine (Fine Hadronic, ou FH), et une partie hadronique "grossière" (Coarse Hadronic ou CH). Les caractéristiques des couches EM, FH et CH sont montrées dans le tableau

3.3. X_0 est la longueur de radiation électromagnétique, c'est-à-dire la longueur moyenne au bout de laquelle un électron a perdu $1-1/e$ de son énergie (soit environ 63% de son énergie initiale), l'électron perd son énergie principalement par rayonnement de freinage (*bremsstrahlung*) et les photons par création de paires électron-positron. λ_0 est la longueur d'interaction hadronique, c'est-à-dire la longueur moyenne pendant laquelle une particule hadronique ne subit aucune interaction. Dans l'uranium, X_0 vaut typiquement 0,32 cm et λ_0 vaut typiquement 10,5 cm.

Le calorimètre électromagnétique (EM) est constitué de quatre couches cylindriques concentriques dans la partie centrale ainsi que dans les parties avant (Nord et Sud), notées EM1, EM2, EM3 et EM4. Les cellules sont regroupées en tours semi-projectives par rapport au centre du détecteur. La taille d'une cellule en $\eta \times \phi$ est approximativement de $0,1 \times 0,1$, sauf pour EM3 qui a une granularité de $0,05 \times 0,05$ soit quatre fois plus fine que les autres couches car, au Run I, c'est dans cette couche que la gerbe électromagnétique était la plus développée, nécessitant une meilleure granularité. Par l'ajout du *PreShower*, la couche E3 ne correspond plus au maximum du développement de la gerbe.

Dans le bloc central, le calorimètre hadronique est constitué de trois couches cylindriques concentriques notées FH1, FH2 et FH3, le calorimètre hadronique "grossier" est constitué d'une seule couche CH. La granularité de ces couches est de $0,1 \times 0,1$ en $\eta \times \phi$. La couche CH contient un trou vertical de bout en bout dans sa partie supérieure par lequel passait l'ancien anneau *Main Ring* du Run I, remplacé par le *Main Injector* au Run II. Dans les blocs avant (Nord et Sud), le calorimètre hadronique est constitué de trois cylindres concentriques interne (noté IH), médian (MH) et externe (OH), chacun contenant une partie hadronique fine intérieure et une partie hadronique "grossier" extérieure. Pour des valeurs de η inférieures à 3,2, la granularité des cellules est de $0,2 \times 0,2$ en $\eta \times \phi$.

Couche	CC	EC
EM 1	$X_0 = 2$	$X_0 = 0,3$
EM 2	$X_0 = 2$	$X_0 = 3$
EM 3	$X_0 = 7$	$X_0 = 8$
EM 4	$X_0 = 10$	$X_0 = 9$
FH 1	$\lambda_0 = 1,3$	$\lambda_0 = 1,3$
FH 2	$\lambda_0 = 1$	$\lambda_0 = 1,2$
FH 3	$\lambda_0 = 0,9$	$\lambda_0 = 1,2$
FH 4		$\lambda_0 = 1,2$
CH 1	$\lambda_0 = 3$	$\lambda_0 = 3$
CH 2	$\lambda_0 = 3$	$\lambda_0 = 3$
CH 3	$\lambda_0 = 3$	$\lambda_0 = 3$

TAB. 3.3 – Caractéristiques du calorimètre.

La figure 3.20 montre une cellule du calorimètre. La plaque d'absorbeur est constituée d'uranium apauvri, sauf pour les couches extérieures du calorimètre hadronique "grossier" qui est constitué d'absorbeur en cuivre ou en acier inoxydable. L'épaisseur des absorbeurs varie selon la région du calorimètre, et l'écart entre l'absorbeur et l'électrode de lecture est de 2,3 mm. L'électrode est placée à un potentiel positif, la plaque d'absorbeur est placée à

la masse. Ainsi, lorsqu'une particule traverse une cellule du calorimètre, les électrons issus de l'ionisation du milieu sont attirés par la cathode, avec un temps de dérive typiquement de 400 ns. Le nombre d'électrons issus de l'ionisation est proportionnel à l'énergie de la particule incidente. Ainsi en mesurant l'intensité du signal récolté à la cathode, on en déduit par un facteur de proportionnalité l'énergie de la particule. Ces coefficients de proportionnalité sont mesurés pour chaque cellule du calorimètre, et constituent une partie de la calibration du détecteur.

3.2.5.3 Les détecteurs inter-cryostats

Entre la partie centrale et les parties avant (Nord et Sud) du calorimètre (soit pour $0,8 \leq |\eta| \leq 1,4$) sont placés des détecteurs inter-cryostats (ICD) qui ont pour but de combler ces zones très peu instrumentées et par lesquelles sont acheminés les câbles. Ce sont des scintillateurs trapézoïdaux de granularité $0,1 \times 0,1$ en $\eta \times \phi$, placés dans le prolongement des tours pseudo-projectives. Des cellules ne contenant pas de plaques d'absorbeur (dites cellules sans masse ou *massless gaps*) sont placées entre la couche externe du calorimètre et la paroi du cryostat, aussi bien dans la partie du côté extérieur du calorimètre central que du côté intérieur du calorimètre avant (Nord et Sud). Elles sont constituées uniquement de plaques de lecture plongées dans l'argon liquide, et utilisent les parois du calorimètre comme absorbeur.

3.2.5.4 Résolutions et performances du calorimètre

La précision de mesure σ sur l'énergie déposée dans le calorimètre dépend des incertitudes sur l'étalonnage du détecteur, du bruit de l'électronique ainsi que d'autres paramètres. Cette précision est quantifiée de la manière suivante :

$$\left(\frac{\sigma}{E}\right)^2 = \left(\frac{N}{E}\right)^2 + \frac{S^2}{E} + C^2$$

où :

- N paramétrise le bruit dû à l'électronique ainsi qu'à la contamination en éléments radioactifs présents dans les composants des cellules du calorimètre ;
- S est un terme dû aux fluctuations statistiques lors de la reconstitution des gerbes électromagnétiques ou hadroniques ;
- C est un terme constant qui englobe toutes les incertitudes sur la calibration du détecteur ainsi que sur l'épaisseur des plaques d'absorbeurs.

Les valeurs de ces trois paramètres sont déduites des résultats de calibration et d'intercalibration des cellules du calorimètre. Pour le Run II, les valeurs de ces paramètres sont :

- $N(\text{GeV}) = 0.202 \pm 0.002$;
- $S(\text{GeV}^{1/2}) = 0.23 \pm 0.10$;
- $C = 0.04 \pm 0.002$.

3.2.6 Le système à muons

Le système à muons est de forme parallélépipédique, et a une section approximativement carrée dans le plan (Oxy) d'environ 10 m (ou 6 m) de côté, tandis que sa dimension le long de l'axe (Oz) est d'environ 20 m (ou 9 m) en ce qui concerne les dimensions extérieures (intérieures).

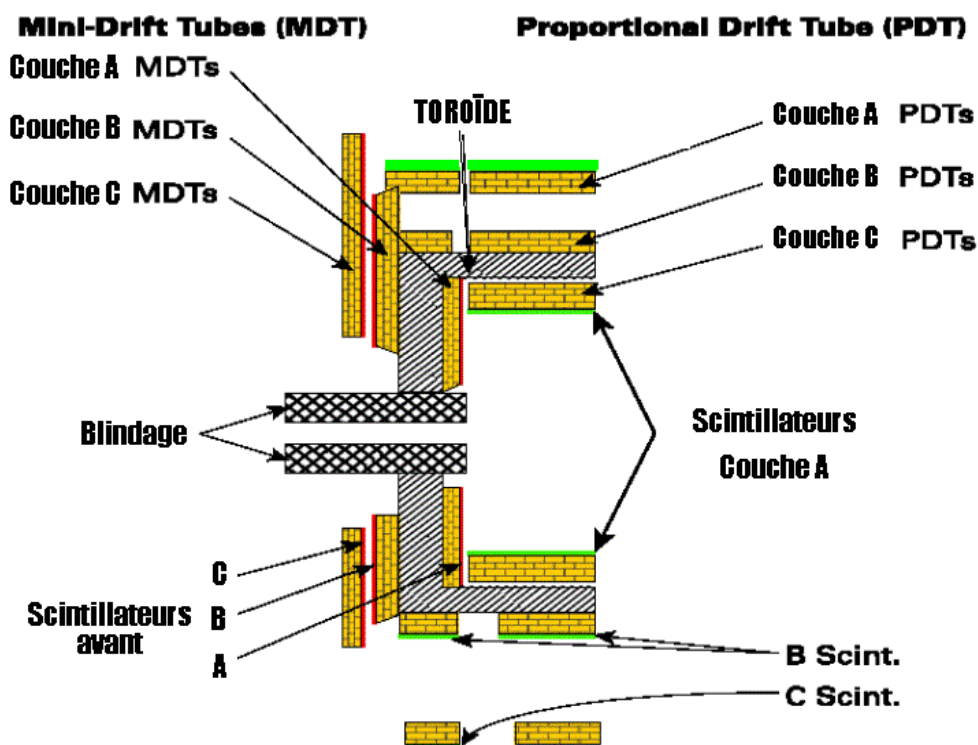


FIG. 3.21 – Le système à muons.

Il est divisé en huit sous-blocs appelés *octants*, chacun couvrant une ouverture angulaire de $\phi = \pi/4$.

Il est constitué de trois parties complémentaires (voir figure 3.21) :

- les chambres à dérive : elles comprennent les chambres à dérive proportionnelles (PDT : Proportional Drift Tubes) qui sont situées dans la partie centrale du détecteur, et les mini-chambres à dérive (MDT : Mini Drift Tubes) qui sont situées à l'avant du détecteur, dans les parties Nord/Sud. Ce sont les détecteurs qui permettent de reconstruire une trace laissée par les muons ; il y a trois couches de chambres à dérive : une couche située avant le toroïde (couche A), et deux couches situées après le toroïde (couches B et C). Les fils des chambres à dérive centrales sont perpendiculaires à l'axe du faisceau.
- les scintillateurs : ils servent principalement à déclencher le système d'acquisition ;
- le toroïde : il crée un champ magnétique qui courbe les muons à leur traversée du toroïde. Ceci permet de reconstruire l'impulsion des muons à partir de l'angle de déviation. Le champ magnétique est toujours perpendiculaire aux fils des chambres

à dérive.

Le détecteur à muons est le sous-détecteur le plus externe et ne reçoit quasiment que des muons. Toutes les particules (sauf évidemment les neutrinos) s'arrêtent dans le calorimètre, excepté les muons qui ont un libre parcours moyen dans la matière beaucoup plus long que celui des autres particules du fait de sa longue durée de vie et de sa section efficace d'interaction électromagnétique bien plus faible que celle de l'électron.

3.2.6.1 L'aimant toroïdal

C'est un aimant en fer de section carrée, situé entre la couche A et la couche B du système à muons. Il est formé d'une partie centrale et d'une partie avant (nord/sud). Il a une épaisseur de 109 cm et pèse 1973 tonnes. L'aimant génère un champ magnétique de 1,8 tesla à l'intérieur du toroïde perpendiculairement au faisceau. Il sert aussi à fermer les lignes de champs électromagnétiques qui sortent du solénoïde.

3.2.6.2 Les chambres à dérive

Les chambres à dérives sont constituées de tubes d'aluminium contenant un fil d'anode au centre et des plaques cathodiques le long de deux faces du tube, le tout plongé dans un mélange gazeux facilement ionisable. Une différence de potentiel est appliquée entre la cathode et l'anode (de l'ordre de 3000 V) de sorte que, lorsqu'un muon traverse une chambre à dérive, les électrons issus de l'ionisation du gaz sont attirés par l'anode. Les fils des chambres à dérive sont orientés le long des lignes de champs créées par l'aimant toroïdal pour permettre la mesure de l'angle de déviation de la trajectoire des muons dans le champ, et donc la mesure de l'impulsion.

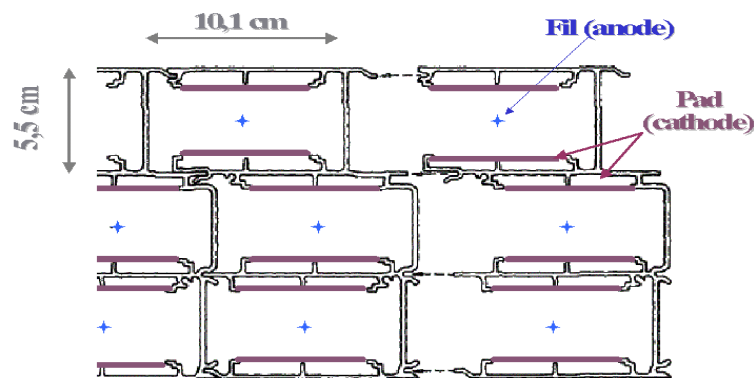


FIG. 3.22 – Vue transversale des PDT. On distingue les fils d'anode au centre des tubes de section rectangulaire et les plaques de cathode sur les côtés long des sections des tubes.

Les chambres à dérive centrales PDT

Une cellule de PDT est formée d'un long tube en aluminium long d'environ 2 m, et de section rectangulaire de 5,5 cm×10,1 cm (voir figure 3.22). Au centre de la cellule se trouve le fil d'anode d'un diamètre de 50 μm . La cathode est constituée de deux plaques fixées sur les côtés longs de la cellule. L'anode et la cathode baignent dans un mélange

gazeux constitué de 80% d'argon, 10% de méthane et 10% de tétrafluorure de carbone (CF_4). Ce mélange gazeux, modifié par rapport au Run I, permet de limiter le temps de dérive des électrons d'ionisation (voir ci-dessous) à un maximum de 500 ns (la vitesse de dérive étant de l'ordre de 10 cm/ μs).

La couverture angulaire des chambres à dérive centrales s'étend jusqu'à $|\eta| < 1,1$. La couche A est formée de quatre plans de cellules (sauf pour la partie inférieure du détecteur où l'on ne compte que trois plans parallèles au plan (xOz) pour $y < 0$). Les couches B et C sont formées de trois plans de chambres à dérive.

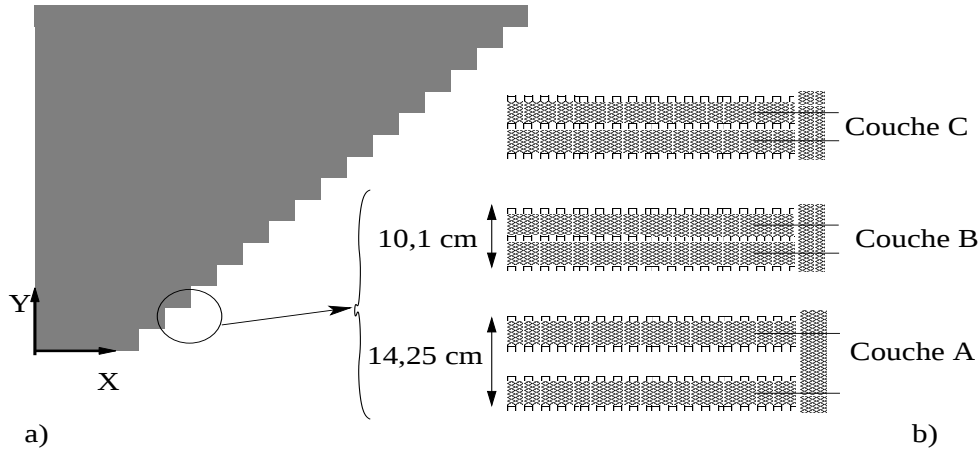


FIG. 3.23 – Schéma des MDT. La figure a) représente le huitième d'une couche de MDT la figure b) montre la disposition des MDT dans chacune des couches A, B et C.

Les chambres à dérive avant MDT

Une cellule de MDT est formée de tubes de longueur variable (voir le schéma a) de la figure 3.23) et de section carrée de 1 cm². Le fil d'anode est le même que pour les PDT, mais le mélange gazeux est différent : il est constitué de 90% de CF_4 et de 10% de méthane, ce qui permet d'obtenir un temps de dérive des électrons de 60 ns au maximum. Ce temps est très inférieur au temps de croisement des faisceaux (actuellement de 396 ns), et permet donc d'éviter l'empilement des signaux issus de deux événements différents. Les MDT ont d'autres avantages : ils ont un faible taux d'occupation, et sont notamment très résistants aux radiations, la partie avant du détecteur étant beaucoup plus soumise aux radiations que la partie centrale. La couverture angulaire des MDT s'étend entre des valeurs de $|\eta|$ allant de 1,1 à 2. Dans les régions très proches du faisceau, un blindage a été ajouté pour diminuer trois principales sources de bruit de fond à l'avant :

- les débris de protons et d'antiprotons après interaction qui créent du bruit de fond dans les couches A des chambres à dérive ;
- les débris de protons et d'antiprotons qui interagissent avec les aimants quadrupôles du Tevatron et qui créent du bruit de fond dans les couches extérieures des MDT ;
- les interactions du faisceau dans le tunnel.

La précision de mesure sur la position des muons dans le plan de dérive est d'environ 0,7 mm.

3.2.6.3 Les scintillateurs

Les scintillateurs jouent deux rôles principaux dans le système à muons :

- leur réponse rapide au passage d'une particule permet de déclencher rapidement le système d'acquisition ;
- une partie de l'électronique permet d'éliminer les muons cosmiques.

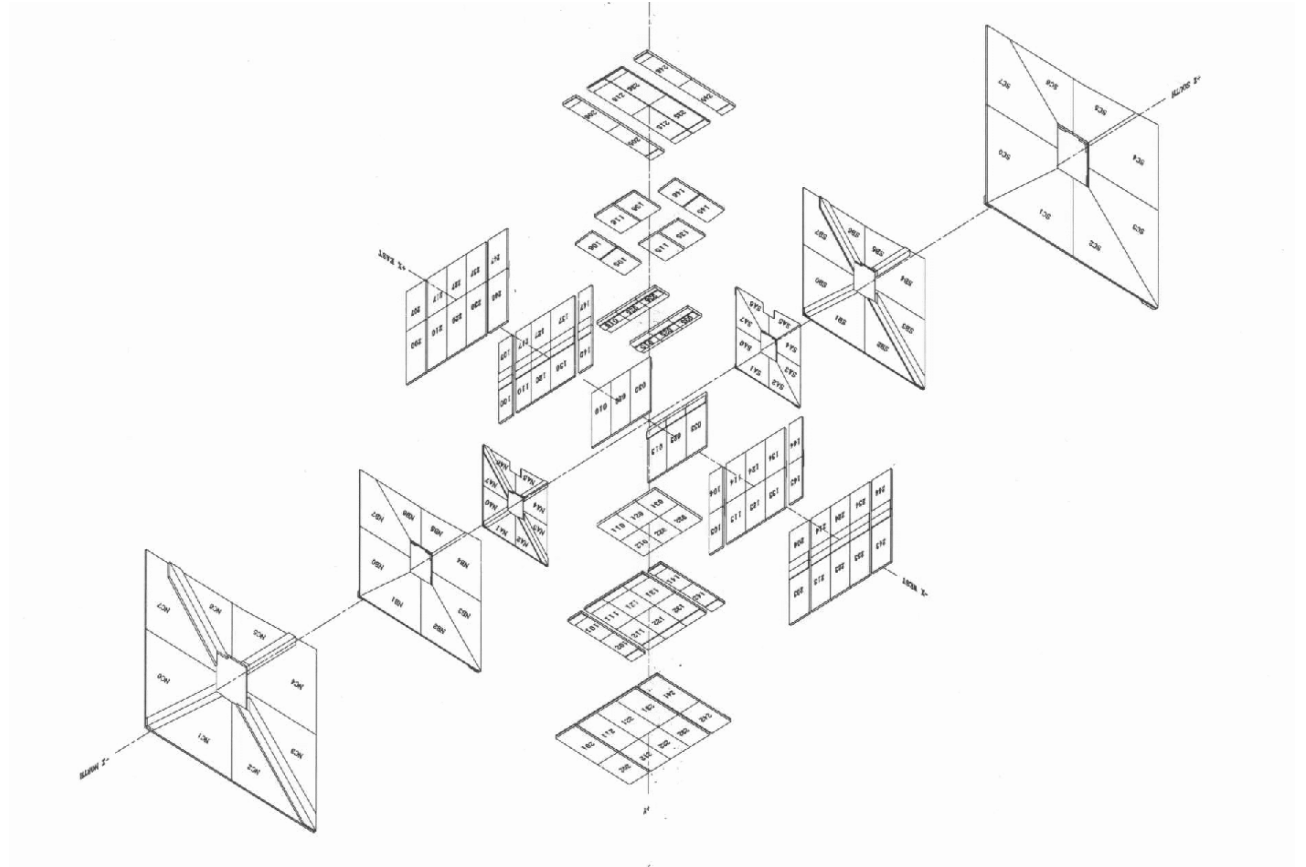


FIG. 3.24 – Représentation en 3 dimensions des chambres à dérives.

Dans la partie centrale sont montés deux types de scintillateurs :

- les *compteurs $A-\phi$* sont placés après le calorimètre, contre la couche A des PDT, et sont formés de lames scintillantes de longueur environ 85 cm le long de l'axe des faisceaux et de largeur angulaire en ϕ d'environ 4,5 degrés. Des photomultiplicateurs sont placés au bout des lames scintillantes pour récolter la lumière émise au passage des muons. La précision temporelle des compteurs $A-\phi$ est d'environ 4 ns, ce qui permet au système de déclenchement de bénéficier d'un temps de réponse très bref, et de rejeter ainsi les muons cosmiques. La segmentation en phi permet d'augmenter la précision sur la position en ϕ des muons ;

- les "*caps cosmiques*" sont placés juste au-dessus de la couche C pour la partie supérieure du détecteur, et entre la couche B et la couche C ainsi qu'au-dessous de la couche C pour la partie inférieure du détecteur. Ils forment avec les compteurs $A-\phi$ un système de

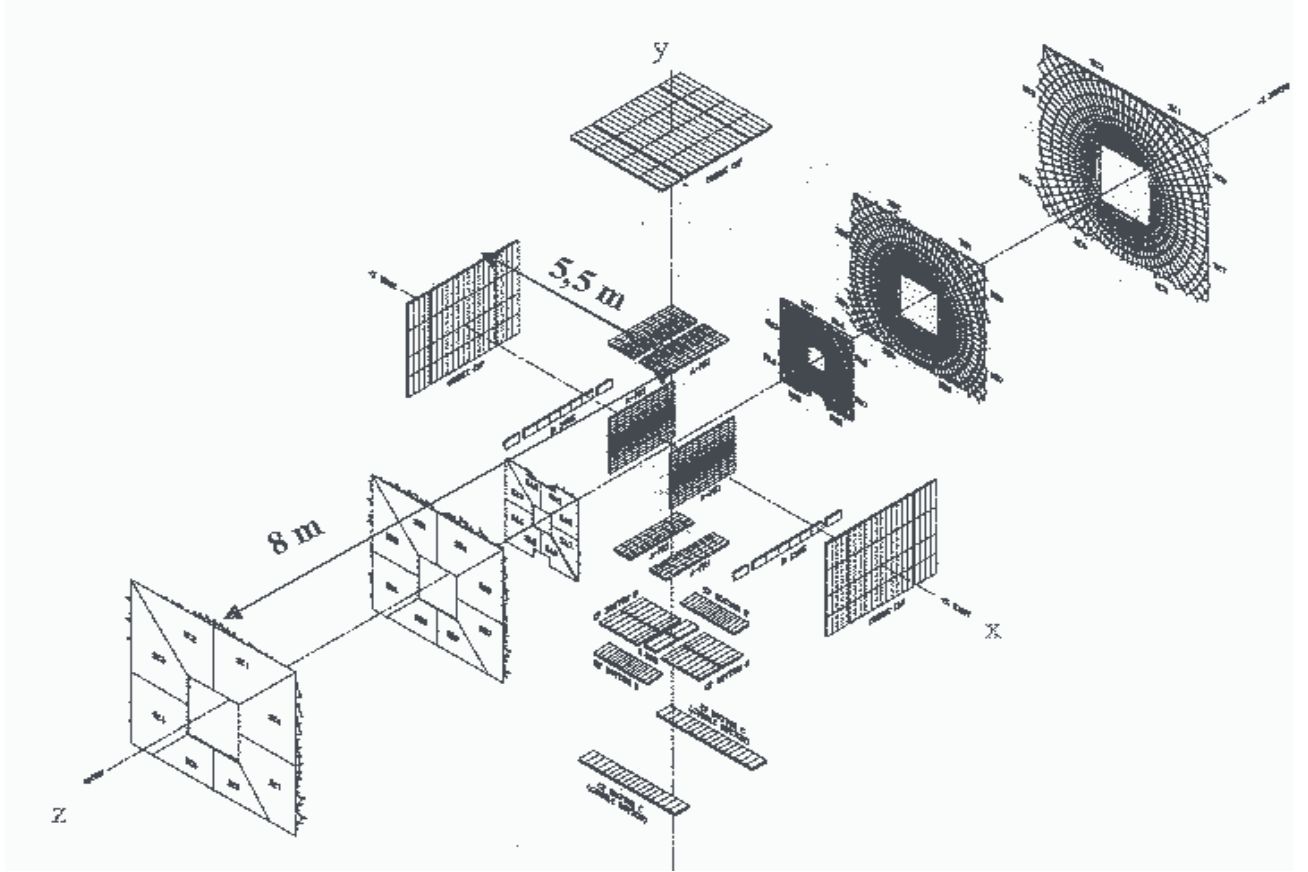


FIG. 3.25 – Représentation en 3 dimensions des scintillateurs du système à muons.

déclenchement performant et efficace (voir chapitre 5 section 5.3) . Leur précision temporelle est de 5 ns, améliorée par le programme de reconstruction jusqu'à une valeur de 2,5 ns.

Dans la partie avant du détecteur sont placés des scintillateurs à segmentation fine, contre le premier plan des couches A et C des MDT et contre le dernier plan de la couche B, offrant au total trois plans de scintillateurs dans la partie nord et trois plans dans la partie sud. Chaque octant est constitué de 96 lamelles scintillantes de segmentation en $\eta \times \phi$ de $0,07 \times 4,5$ pour les trois étages les plus près du faisceau, et de segmentation de $0,12 \times 4,5$ pour les neuf autres. La précision temporelle des scintillateurs pixels est de l'ordre de la nanoseconde. L'efficacité de détection d'une particule est de 99,9% pour chaque plan de scintillateur.

3.3 Le traitement des données

3.3.1 Système de déclenchement et d'acquisition des données

Le taux d'interactions ($\simeq 5 \cdot 10^6/s$) dû au faible temps de croisement des faisceaux de collision (396 ns entre chaque croisement) ainsi qu'au grand nombre de protons ($\simeq 10^{11}$) et d'antiprotons ($\simeq 10^{10}$) par paquet est beaucoup trop élevé pour que toutes les données soient enregistrées. La plupart des événements sont dus à des processus QCD à très faibles impulsions transverses. C'est pourquoi une présélection des événements s'impose. Ainsi

on sélectionne une petite partie de ces collisions grâce à un système de déclenchement (appelé aussi *trigger*) échelonné en trois niveaux.

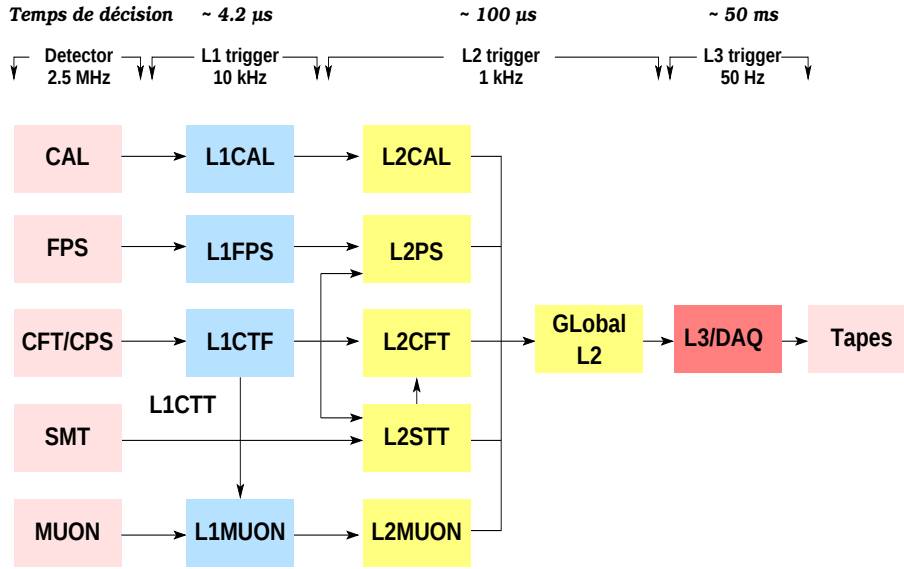


FIG. 3.26 – Schéma représentant les principaux éléments du système de déclenchement. Les flèches indiquent les éléments qui sont pris en compte pour les triggers dans chaque niveau de déclenchement (L1, L2 et L3/acquisition).

Le niveau 1 (L1)

Le niveau 1 du déclenchement utilise l'information provenant du CFT, des détecteurs de pieds de gerbe (CPS et FPS), des détecteurs de protons à l'avant (Forward Proton Detector: FPD) du calorimètre et du système à muons.

Un ensemble de termes logiques appelés *termes "ET/OU"* représentent les briques élémentaires du système de déclenchement. Le niveau de déclenchement L1 peut traiter 128 combinaisons de termes "ET/OU". L'association de plusieurs termes "ET/OU" représente un type de déclenchement.

Si un *trigger* a un taux de déclenchement trop élevé par rapport à un taux acceptable, un précomptage (ou *prescale p*) est appliqué de sorte que sur p événements qui ont déclenché seulement un événement est enregistré et soumis à la sélection du niveau 2.

Le niveau 2 (L2)

Un *trigger* au niveau L2 peut être basé sur un seul terme "ET/OU" ou sur une combinaison de ces termes. Le niveau L2 nécessite un temps de décision pour traiter l'information venant du L1 et vérifier les termes "ET/OU" qui ont été déclenchés. Pendant ce temps, les collisions continuent de solliciter le niveau L1, et pour éviter l'empilement des signaux de différentes collisions, une mémoire de stockage temporaire est utilisée pour stocker l'information du L1 pour chaque événement en attente d'être traitée par le niveau L2. Cette mémoire temporaire peut ainsi stocker l'information de 16 collisions.

Le niveau 2 du déclenchement fonctionne en deux étapes. La première étape consiste en l'acheminement de l'information issue du L1 vers des préprocesseurs afin de reconstruire,

à partir des sous-détecteurs, les premiers éléments “d’objets physiques” :

- des traces reconstruites à partir de l’information provenant du CTT (Central Track Trigger, utilisant le CFT et le CPS) [14] ou du STT (Silicon Track Trigger, utilisant le SMT et le CFT) [17] ;
- des dépôts d’énergie dans les détecteurs de pieds de gerbes ; des objets électromagnétiques, jets et autres quantités liées aux données du calorimètre ;
- des muons reconstruits à partir de l’information donnée par le système à muons et combinée avec le SMT et CFT.

Les algorithmes implémentés dans les processeurs ont un temps de calcul qui doit être inférieur à $50\ \mu\text{s}$ par événement, afin de permettre au système de déclenchement de niveau 2 de fonctionner à un taux de déclenchement maximal de 1000/s.

Le niveau 3 (*L3*)

Le troisième niveau de déclenchement L3 a un taux de déclenchement maximal de 50/s, ce qui pose des contraintes temporelles sur la décision prise pour chaque événement. Les informations de tous les sous-détecteurs sont envoyées vers une ferme d’ordinateurs qui reconstruisent partiellement les objets et quantités physiques. Plusieurs critères sont combinés afin de sélectionner les événements, le but étant de réduire le flux de données de 95%, ne sélectionnant donc que 50 événements par seconde de données qui doivent être enregistrés sur des bandes magnétiques sous forme de *raw data* (voir section suivante 3.3.2).

3.3.2 Format de stockage des données

Les données collectées durant l’acquisition sont enregistrées dans des *raw data* qui contiennent l’information brute de tous les détecteurs. Dans le programme de reconstruction [18] sont reconstruits les objets et sont enregistrés dans des DST (*Data Summary Tapes*).

Le problème avec ce format de stockage est que la fréquence d’acquisition des données est de 50 événements par seconde, ce qui porte le nombre d’événements enregistrés à environ 800 millions par an occupant un espace disque d’environ 120 Tbytes pour un format non condensée. Cette taille est très élevée et des problèmes de stockage se posent. C’est pour cela qu’un nouveau format de données a été construit, appelé *Thumbnail* [19]. Ce format contient les informations sous forme condensées. Il permet alors un gain de place de plus de 90%, mais n’est pas lisible directement pour l’analyse. il faut d’abord le décompresser pour le convertir en DST.

Pour l’analyse des données, il est possible d’utiliser directement les DST ou de les convertir en un format (les *HEP – tuples*) totalement indépendants de l’environnement informatique complexe de DØ.

Les HEP-tuples créés sont compatibles avec l’outil d’analyse et de visualisation graphique ROOT [20], largement employé dans la collaboration DØ. Leur inconvénient majeur est qu’ils occupent un espace considérable (environ 150 kbytes par événement) et ne

permettent pas de relier les différents blocks d'information entre eux.

Nous avons participé au développement d'un ensemble de programmes permettant l'écriture des données dans un format ROOT, baptisé TMBTree pour ThuMBnail Tree, beaucoup plus condensé et conservant les relations entre objets physiques reconstruits. Le stockage des données sous ce format peut se faire aussi bien à partir des DST que des thumbnails décompressés.

Les avantages du TMBTree par rapport au ROOT HEP-tuple sont multiples :

- l'espace utilisé par chaque événement est considérablement inférieur à celui utilisé par le ROOT HEP-tuple : ~ 20 kbytes au lieu de ~ 150 kbytes ;
- alors que le ROOT HEP-tuple était construit sur la base d'une structure en tableau, la construction TMBTree est orientée objet, utilisant les avantages du C++, langage officiel de la collaboration DØ ;
- La structure du TMBTree permet une grande efficacité de stockage des événements, aussi bien en espace utilisé (comme vu précédemment) qu'en durée d'enregistrement et en temps d'accès à l'information ;
- le TMBTree contient toutes les informations utiles sur les relations entre objets reconstruits alors que le ROOT HEP-tuple en connaît très peu.

Le TMBTree a une structure ramifiée, comme un arbre. Il contient des branches qui sont des blocs indépendants les uns des autres. Une branche peut représenter une particule identifiée comme les muons ou les jets, ou représenter un objet physique comme l'information sur les *triggers* ou les vertex. Chaque branche contient des feuilles (*leaf*) qui sont les variables à analyser. Une branche spéciale permet d'établir des liens entre les différentes branches.

Pour les détails techniques sur cette partie, voir les références [19][21][22][23].

Actuellement, les données sont stockées sous forme de *thumbnails* créés à partir des *raw data*. Chaque groupe de travail crée ensuite ses propres fichiers d'analyse à partir de ces *thumbnails* en les convertissant soit en *DST* par décompression, soit en *HEP-tuple*, soit en TMBTree.

Bibliographie

- [1] F.Abe *et al.*, Phys. Rev. Lett. 74, 2626 (1995);
S.Abachi *et al.*, Phys. Rev. Lett. 74, 2632 (1995).
- [2] e-Print Archive: hep-ex/0212027,
A paraître dans les *proceedings* des 37^{ièmes} rencontres de Moriond sur la QCD et les
interaction hadroniques, Les Arcs, France, 16-23 Mar 2002.
- [3] <http://www-bdnew.fnal.gov/pbar>
- [4] *L'accélérateur de Cockroft-Walton*,
<http://www.fnal.gov/pub/inquiring/physics/accelerators/chainaccel.html>
<http://lbl.gov/abc/wallchart/chapters/11/1.html>
- [5] *Le Linac*,
<http://linac.fnal.gov>
- [6] *Le Booster*,
<http://www-bd.fnal.gov/proton/boobster/booster.html>
- [7] *L'injecteur principal* Main Injector,
<http://www-fmi.fnal.gov>
- [8] *L'accumulateur/debuncher*
<http://www-bdnew.fnal.gov/pbar>
- [9] *Le Tevatron*
<http://bd-new.fnal.gov/Tevatron/>
- [10] *Run II luminosity goals* [http://www-bd.fnal.gov/Run II/wc_run2b.pdf](http://www-bd.fnal.gov/Run_II/wc_run2b.pdf)
DØ luminosity web page
<http://www.hep.brown.edu/lm/detector.htm>
- [11] *Run II Handbook*
[http://www-bd.fnal.gov/Run II/index.html](http://www-bd.fnal.gov/Run_II/index.html)
- [12] *Documentation sur le SMT*
http://d0server1.fnal.gov/projects/Silicon/www/Writeups/m_nimnote.ps
http://d0server1.fnal.gov/projects/Silicon/www/Writeups/fw_nim.ps
http://d0server1.fnal.gov/projects/Silicon/www/Writeups/m_proj_90.ps
<http://d0server1.fnal.gov/projects/silicon/www/SMTcoord.JPG>
http://www-d0.fnal.gov/meena/d0_private/silicon/coords/smt-numb
- [13] http://d0server1.fnal.gov/projects/SciFi/cft_home.html
<http://www-d0.fnal.gov/hardware/upgrade>
- [14] L. Babukhadia et M. Martin, *Track and preshower digital trigger in DØ*,
DØ note 3980.
- [15] H. Evans *et al.*, *A silicon track trigger for the DØ experiment in Run II*,
<http://www-d0.fnal.gov/trigger/stt/sttdesign.html>.

- [16] DØ Collaboration, *Proposal for a forward proton detector at DØ* ,
http://www-d0.fnal.gov/brandta/fpg/d0_private/fpd_pac3.ps.
A.Brandt, *Triggering with the forward proton detector* ,
http://www-d0.fnal.gov/brandta/fpg/d0_private/trigger/fpd_trigger.ps.
- [17] L. Babukhadia et M. Martin, *Track and preshower digital trigger in DØ* ,DØ note 3980.
- [18] http://d0server1.fnal.gov/projects/Computing/Web/Meeting/RAC/d0finance_0702_jq.pdf
- [19] S.Baffioni, E.Nagy, S.Protopopescu, *Thumbnail: a compact data format*,
DØ note 3979.
- [20] <http://root.cern.ch>
- [21] <http://www-d0.fnal.gov/~serban/thumbnail/>
Proposal for Thumbnail Contents.
- [22] http://marwww.in2p3.fr/~nagy/tmb_tutorial/
- [23] S.Baffioni, N.Lahrichi, E.Nagy, S.Protopopescu, E.Thomas, *The tmb_tree package*,
DØ note 3978.

Chapitre 4

Identification et reconstruction des muons

Introduction

Le signal étudié étant la masse invariante de deux muons dans l'état final, l'accent sera mis sur la reconstruction de l'impulsion des muons. Dans ce chapitre nous allons décrire les étapes de la reconstruction et de la sélection des muons.

Le canal physique étudié ne requiert aucune information particulière donnée par le calorimètre. En effet, nous n'utilisons *a priori* aucune particule électromagnétique ni aucun jet hadronique mis à part les jets issus de radiation de gluons dans l'état initial et la radiation de photons par les muons dans l'état final. Ces radiations auront un effet sur le critère d'isolation des muons, nous en parlerons dans le chapitre 5, dans la section 5.5.

4.1 Identification des muons

Dans un premier temps, l'identification des muons se fait au niveau du déclenchement parmi les *triggers* liés aux muons :

Niveau 1 (L1)

Le niveau 1 [1] prend en compte les informations données par les scintillateurs du système à muons ainsi que les chambres à dérive et prend en compte éventuellement les informations du détecteur de traces central L1CFT selon le type de *trigger*.

Ces *triggers* peuvent ensuite exiger d'autres caractéristiques liées à l'impulsion transverse du muon (P_T), la région (en η) dans laquelle se situe le muon, le nombre de scintillateurs touchés ainsi que le nombre de coups dans les couches de chambres à dérive.

Le nombre de muons identifiés au L1 est important dans la sélection de l'échantillon de données. En effet, lorsque deux muons sont identifiés par leur passage dans les scintillateurs, un système de coïncidences permet de détecter leur passage quasi-simultané et enregistre un événement *dimuon*. Cette condition sera exigée dans la sélection des données que nous analyserons dans le chapitre suivant.

Niveau 2 (L2)

Au niveau 2 [2], le *trigger* muons contient toutes les informations primaires concernant les muons récoltées par le L1 ainsi que celles des autres détecteurs. Ces données permettent d'ajouter des conditions supplémentaires à la sélection des événements effectuée au niveau 1, comme la qualité des muons identifiés, la présence de jets d'énergie minimale donnée ou d'énergie manquante dans le calorimètre, ou la présence de traces centrales. Ces critères peuvent être combinés entre eux, constituant ainsi différents éléments de déclenchement.

Niveau 3 (L3)

Le niveau 3 [3] est réalisé dans une ferme d'ordinateurs : il dispose des informations de tous les détecteurs, et utilise les outils des programmes de reconstruction. A ce niveau l'information est plus précise que celle collectée par le L1 et le L2, ce qui permet d'obtenir des précisions sur les énergies et impulsions comparables à celles de la reconstruction

différée.

Les *triggers* utilisés

Pour notre échantillon d’analyse final, nous sélectionnons les événements qui ont déclenché le *trigger* “dimuon” 2MU_A_L2M0 : il demande au moins deux muons au niveau L1 et est sensible à toute la couverture angulaire du détecteur de muons, c’est-à-dire $|\eta| < 2$. Au niveau 2 du déclenchement il requiert qu’au moins un des deux muons qui ont déclenché le *trigger* soit de qualité moyenne¹ (voir section 4.3). Aucun critère de sélection n’est appliqué sur l’impulsion transverse des muons.

Une étude sur les efficacités de déclenchement de ce *trigger* est détaillée dans la section 5.3. Pour cette étude nous utilisons aussi des *triggers* “monomuon” :

- le *trigger* monomuon MU_A_L2M0 pour l’étude du L1. Il demande un muon au niveau L1 qui ait laissé un coup dans les scintillateurs à muons dans la région $|\eta| < 2$, un muon de qualité moyenne au niveau L2 et aucune condition particulière au niveau L3;
- le *trigger* monomuon MU_JT20_L2M0 pour l’étude du L2. Il requiert un muon au L1 dans la région $|\eta| < 2$ ainsi qu’une tour dans le calorimètre contenant une énergie transverse supérieure à 5 GeV, un muon de qualité moyenne au L2, et un jet reconstruit avec une énergie transverse supérieure à 20 GeV au L3.

Concernant l’étude de la coupure sur le critère de qualité des muons (voir la section 5.1), nous sélectionnons les événements qui ont déclenché un des deux *triggers* monomuons :

- le *trigger* monomuon MU_W_L2M0. Il requiert un muon au niveau 1 se trouvant dans la région $|\eta| < 1,5$ et un muon de qualité moyenne au L2.
- le *trigger* monomuon MU_W_L2M5-TRK10. Il requiert un muon au L1 se trouvant dans la région $|\eta| < 1,5$. Au L2, il requiert un muon de qualité moyenne et d’impulsion transverse supérieure à 5 GeV. Au L3, une trace dans le CFT d’impulsion transverse supérieure à 10 GeV est requise.

4.2 Reconstruction des muons

Dans le détecteur de muons, la reconstruction des muons se fait principalement en trois étapes : la reconstruction de points correspondant aux “coups” laissés par le passage des muons dans les chambres à dérive ou dans les scintillateurs, la reconstruction de segments de droite joignant ces points, puis la reconstruction de traces joignant ces segments dans les trois couches de chambres à dérive. Une description détaillée de ces étapes est réalisée dans [4].

4.2.1 Dans le système à muons

Reconstruction des coups

Les premiers objets dont dispose le programme de reconstruction sont les coups laissés

1. Ce critère de qualité au niveau du trigger n’est pas rigoureusement identique à celui qui est décrit dans la section 4.3.

par les muons dans les chambres à dérives dans chaque plan des couches de MDT ou PDT ainsi que dans les scintillateurs.

La position du coup dans les scintillateurs est prise au centre du scintillateur touché. Dans la partie centrale, les scintillateurs sont des plaques rectangulaires pouvant atteindre plusieurs mètres de long, ce qui donne des précisions sur la position du coup allant de 25 cm dans la couche A à 2 m pour les couches B et C dans la moitié haute du détecteur (figure 4.1). Par contre dans la partie avant, les plaques de scintillateurs étant plus petites, la précision est de l'ordre de quelques centimètres.

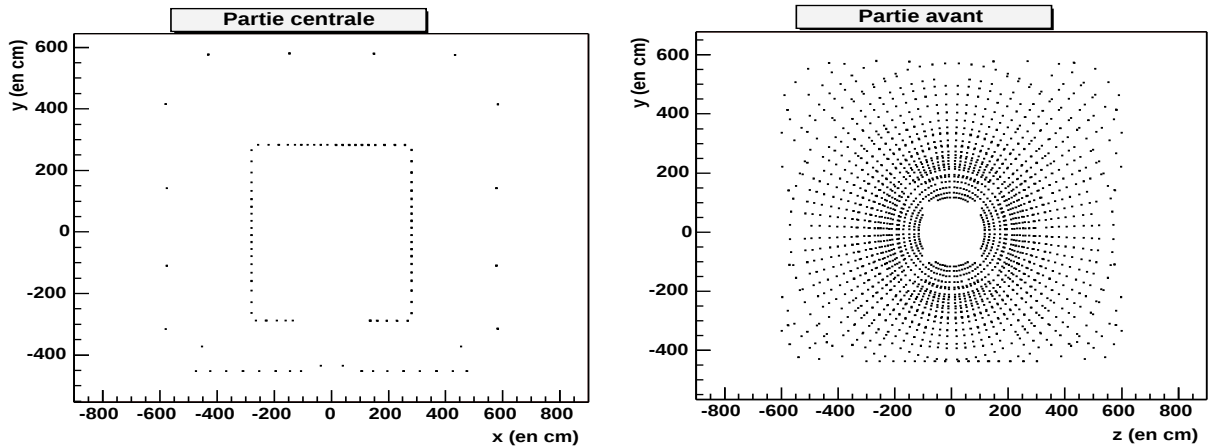


FIG. 4.1 – Carte des coups relevés dans les scintillateurs. La figure de gauche concerne la partie centrale des scintillateurs, la figure de droite concerne la partie avant.

Dans les PDT, les cellules ont chacune un système de lecture et sont reliées deux à deux par un câble à ligne de retard de 20 ns. L'information temporelle est donnée par deux systèmes de lecture. Connaissant le temps d'arrivée de chaque signal ainsi que la différence de 20 ns entre les deux signaux, une soustraction des deux temps permet de déduire le temps de parcours dans chacun des fils d'anode, et ainsi reconstruire la position le long de chaque fil.

La mesure de la position du coup laissé par le muon dans le plan perpendiculaire à l'axe du fil n'est pas simple étant donné que la relation entre le temps de dérives (information disponible) et la position n'est pas linéaire (le champ électrique est central, mais la section rectangulaire des cellules cause une distorsion des lignes de champs [5]). Pour cela les positions des coups sont calculées en supposant que le muon a traversé la chambre dans le plan perpendiculaire à l'axe du fil, et sont amenées à être modifiées lors de l'ajustement des segments de droite reconstruits à partir de ces coups.

Dans les MDT, chaque fil d'anode est indépendant de ses voisins. On suppose dans un premier temps que la particule est passée au milieu du fil (cf Fig 4.3). Un temps de dérives est alors grossièrement calculé. La position du coup le long du fil est calculée avec plus de précision grâce à l'information donnée par les scintillateurs. Ce calcul se fait lors

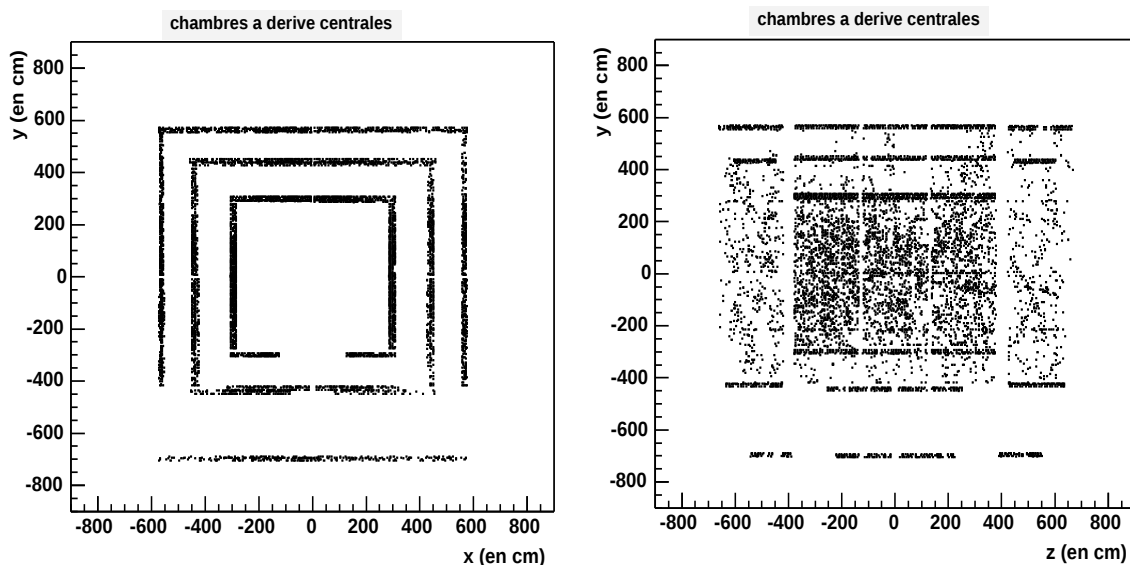


FIG. 4.2 – Carte des coups relevés dans les chambres à dérive centrales PDT. La figure de gauche est une projection dans le plan $(x0y)$. La figure de droite est une projection dans le plan $(z0y)$.

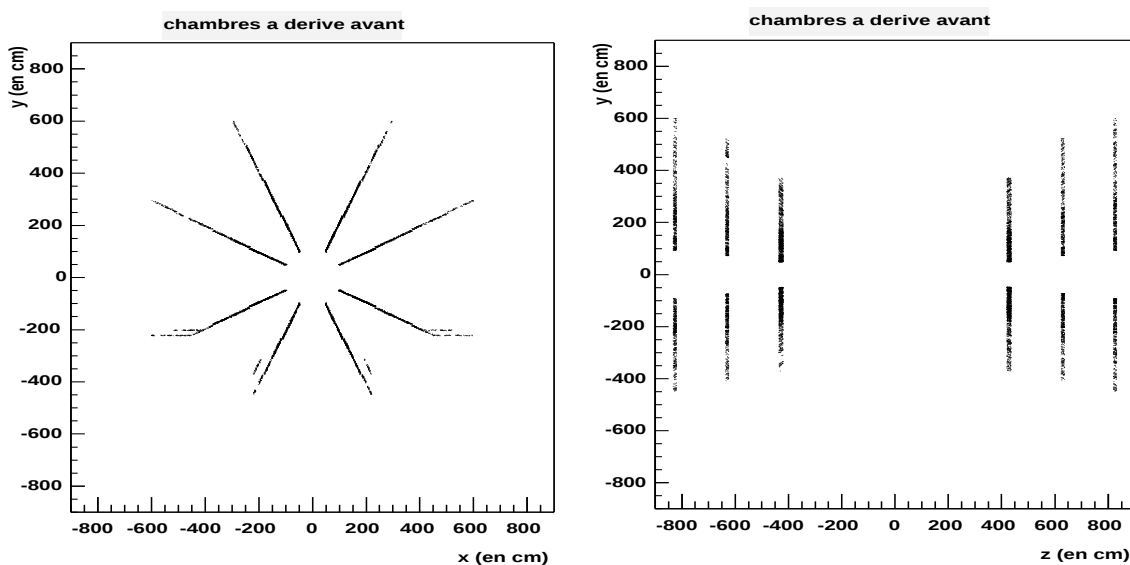


FIG. 4.3 – Carte des coups relevés dans les chambres à dérive avant MDT. La figure de gauche est une projection dans le plan $(x0y)$, les coups sont par défaut situés au milieu du fil de chaque tube de MDT. La figure de droite est une projection dans le plan $(z0y)$, les trois couches de MDT sont bien visibles.

de la reconstruction de segments.

Reconstruction de segments

En général, un muon peut traverser une ou deux cellules dans chaque plan de cellules suivant l'angle d'incidence, ce qui crée entre 4 et 8 coups pour les chambres A, et 3 à 6

coups dans chacune des chambres B ou C. Chaque couche fournit donc un ensemble de coups à partir desquel on reconstruit un ou plusieurs segments, selon les combinaisons possibles entre les coups. Les segments relient entre eux les coups laissés dans chacune des chambres à dérive. Dans chaque couche de PDT ou de MDT, les segments sont d'abord reconstruits dans le plan perpendiculaire aux fils. Chaque coup est déterminé par sa distance de dérive (très approximative dans un premier temps). Les points correspondant aux coups sont d'abord joints deux à deux, ce qui donne un premier ensemble de segments à deux points. Les segments qui partagent un point commun sont ensuite combinés deux à deux. Les paires de segments formant un angle supérieur à $0,3$ rad sont rejetées (figure 4.4). La courbure due au champ magnétique peut être négligée au sein de chaque couche.

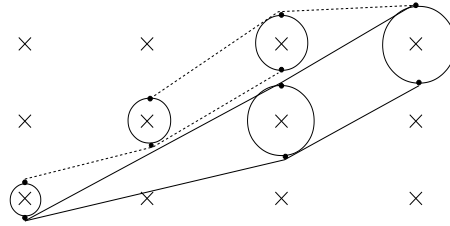


FIG. 4.4 – Exemple de paires de segments acceptées (en ligne continue) et rejetée (en pointillés) joignant deux points sur les cercles matérialisant la position du coup. Les croix représentent la position des fils.

La même opération se renouvelle afin d'associer ces paires de segments avec un quatrième point dans la couche A ou avec un coup dans les scintillateurs et le même critère de sélection est appliqué. Un ajustement est réalisé sur les segments les plus longs trouvés afin de recalculer la position finale des coups. L'ajustement est appliqué à la fois dans le plan perpendiculaire aux fils et le long du fil indépendamment. Pour les PDT, la position le long du fil est connue, mais pour les MDT c'est la position du coup des scintillateurs qui permet de déterminer plus précisément la position le long du fil.

Il n'y a pas de matière (hormis les scintillateurs dans certaines zones) ni de champ magnétique entre les couches B et C. Nous parlerons donc plus tard de “la couche BC”. Les couches B et C sont séparées d'environ 1 m, ce qui permet d'obtenir une meilleure précision de mesure sur la position des segments reconstruits. A ce niveau, tout un ensemble de segments est gardé, et aucune sélection a priori n'est faite. Les informations sont conservées et la sélection se fait par la suite au niveau de l'appariement des segments entre les couches A et BC ainsi que des coups laissés dans les scintillateurs.

La précision de mesure sur la position du muon donnée par le segment est d'environ 40 mrad sur ϕ . La précision sur θ est d'environ 10 mrad pour des segments de la couche A et 0.6 mrad dans le cas d'un segment de la couche BC [6].

Reconstruction de traces

Lorsque des segments sont reconstruits à la fois dans la couche A et la couche BC, le programme de reconstruction combine ces segments pour en faire une seule trace appelée

trace locale.

Le champ magnétique (toujours parallèle aux fils des chambres à dérivation et perpendiculaire à l'axe du faisceau) courbe la trajectoire des muons qui définit alors un plan de déviation.

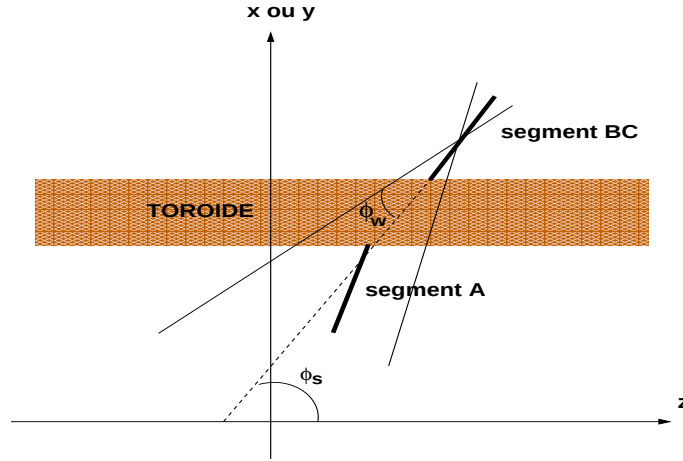


FIG. 4.5 – Compatibilité d'un segment de la couche A avec un segment de la couche BC.

On calcule une première estimation de l'impulsion connaissant l'angle de déviation θ_d formé par les deux segments, sachant que $P = 0.3 \times B \times L / \theta_d$ où P est l'impulsion reconstruite dans le plan de déviation et s'exprime en GeV, B est le champ magnétique pris en Teslas et L est la distance en mètres parcourue par les muons dans le toroïde. Cette impulsion sera modifiée au cours de l'ajustement des paramètres. La valeur de P au niveau de la couche A est corrigée de la perte d'énergie dans le toroïde qui est d'environ 15.7 MeV/cm.

La trace est propagée pas à pas depuis le centre de gravité du segment BC jusqu'au centre de gravité du segment A en prenant en compte à chaque étape la perte d'énergie dans le toroïde ainsi que la diffusion multiple.

La position de la trace, les angles, l'impulsion dans le plan de déviation associée à cette trace ainsi que toutes les erreurs liées à ces paramètres sont donnés par le résultat de l'ajustement.

L'impulsion des muons dans le système local est donnée par le résultat de l'ajustement entre le segment A et le segment BC ainsi trouvés. La résolution sur l'impulsion est montrée sur la figure 4.6 dans le cas de muons issus de la désintégration d'un Z et sur la figure 4.7 dans le cas de muons issus de la désintégration d'un graviton de masse 400 GeV. Sur ces deux figures réalisées sur un échantillon de données simulées par une méthode de Monte Carlo, nous montrons à gauche la distribution de l'impulsion transverse des muons (piquée à la moitié de la masse de la résonance) et nous traçons dans la partie droite des figures la résolution sur l'impulsion transverse.

On voit que pour des muons d'impulsion transverse distribuée autour de 45 GeV, la précision de mesure du système local est d'environ 42%. Pour des muons d'impulsion transverse distribuée autour de 200 GeV, l'erreur est supérieure à 90%. Cela montre que, pour la reconstruction de la masse invariante des dimuons étudiés, nous ne pouvons plus

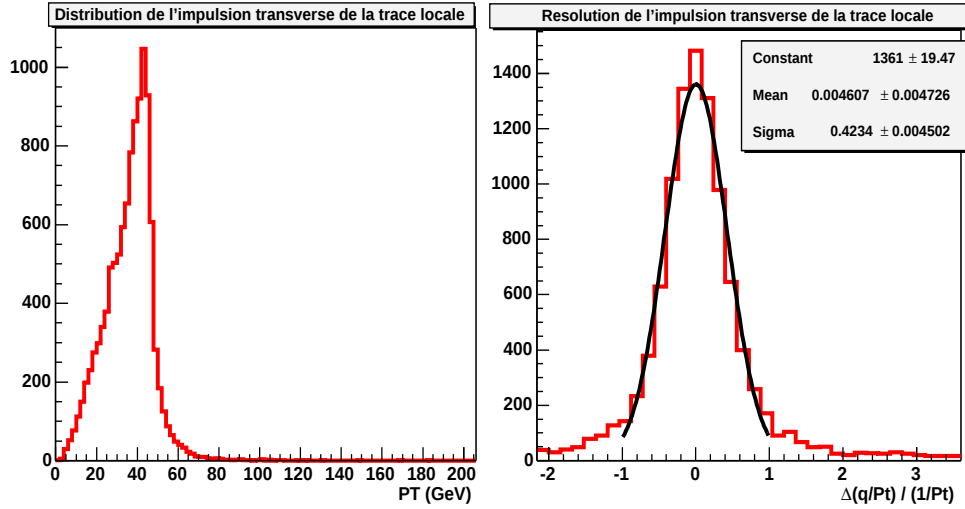


FIG. 4.6 – Distribution de l'impulsion transverse des muons (à gauche) et précision de mesure de l'impulsion des muons dans le système local issus de la désintégration d'un Z (à droite) : en abscisse est montrée $(q_g/P_{T_g} - q_r/P_{T_r}) * P_{T_g}$ où q_g (q_r) et P_{T_g} (P_{T_r}) sont la charge et l'impulsion transverse générées (reconstruites).

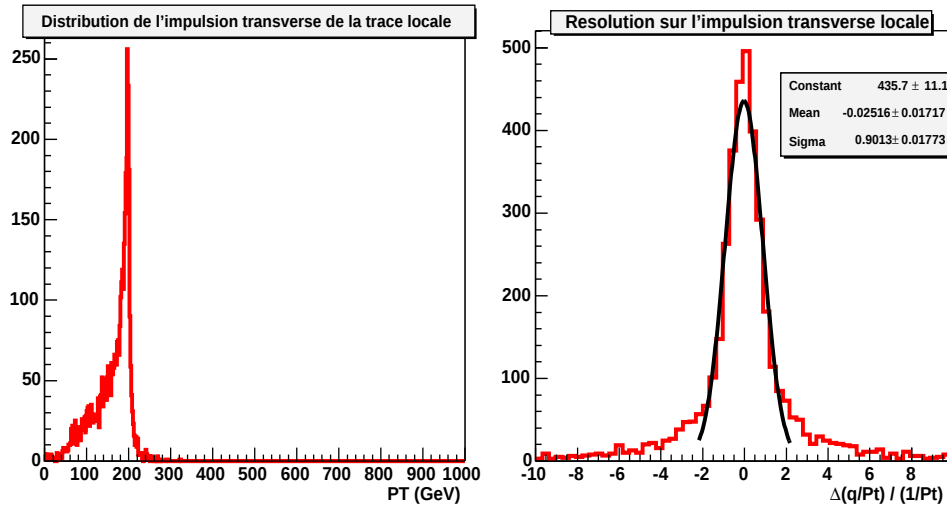


FIG. 4.7 – Distribution de l'impulsion transverse des muons (à gauche) et précision de mesure de l'impulsion des muons dans le système local issus de la désintégration d'un graviton de 400 GeV de masse (à droite) : en abscisse est montrée $(q_g/P_{T_g} - q_r/P_{T_r}) * P_{T_g}$ où q_g (q_r) et P_{T_g} (P_{T_r}) sont la charge et l'impulsion transverse générées (reconstruites).

tenir compte du système local seul pour la mesure de l'impulsion et nous devons aussi prendre en considération l'impulsion transverse donnée par le système central.

4.2.2 Dans le détecteur de traces central

Dans le chapitre précédent nous avons vu que le détecteur de traces central est composé de 2 sous-détecteurs : le SMT et le CFT. La reconstruction des traces centrales se

fait en tenant compte des deux détecteurs.

Il existe quatre algorithmes [8][9] qui reconstruisent des traces à partir des coups laissés par les particules chargées dans le SMT et le CFT. De ces différentes méthodes nous décrirons les deux principales qui donnent les meilleurs résultats :

“Global Track Finder” [10]

Dans cet algorithme, on divise l’espace en quatre régions en η comprenant le CFT et la partie du SMT la plus centrale, la partie avant du SMT comprenant les disques F, une partie couvrant les couches extérieures du CFT et la zone du SMT qu’elles recouvrent, puis finalement l’espace qui est compris entre la zone de recouvrement et la partie avant du SMT. Dans chaque région, l’algorithme reconstruit des “routes” en reliant les coups laissés dans le SMT et le CFT, en partant des couches externes et en extrapolant les trajectoires reconstruites jusqu’aux couches internes des détecteurs. A chaque association d’un nouveau point à la trace, les positions des coups et les erreurs sont recalculées tenant compte de l’ajustement des paramètres.

“Histogramming Track Finder” [11]

Il effectue une présélection des coups dans les détecteurs en utilisant une technique d’histogrammes (basée sur la méthode dite de Hough) . C’est une méthode qui réalise une transformation de l’espace des paramètres appliquée au plan transverse (xOy) puis au plan (rOz). Cette méthode permet de réduire considérablement l’espace combinatoire et réduit donc le temps de reconstruction des traces. La sortie de cette première étape est un ensemble de traces plus ou moins réalistes, qui va être ensuite “écrémé” en utilisant un filtre de Kalman. La reconstruction des traces se fait aussi bien à partir du SMT vers le CFT que du CFT vers le SMT. Il est possible de passer outre la première étape comprenant la mise en histogrammes des données en partant directement de quadruplets de coups dans le SMT. Mais cette étape n’est valable que pour des traces ayant laissé au moins 4 coups dans le SMT.

Dans la version de reconstruction que nous utilisons, l’algorithme de reconstruction des traces centrales est une combinaison de ces deux algorithmes [12]. Nous montrons sur les figures 4.8 et 4.9 la distribution de l’impulsion transverse (partie droite des figures) et la précision de mesure de l’impulsion transverse (partie gauche des figures) des muons issus de la désintégration d’un Z et d’un graviton de 400 GeV de masse.

La résolution en énergie est de l’ordre de 8% pour des valeurs d’impulsion transverse autour de 45 GeV et de 48% pour une distribution de l’impulsion transverse piquée à 200 GeV. Cette résolution est bien meilleure que celle obtenue avec le système à muons seul.

4.2.3 Algorithme d’appariement des traces locales et centrales

L’appariement de traces locales avec les traces centrales a un intérêt majeur dans l’amélioration de l’estimation de l’impulsion du muon étant donné que la précision de mesure de l’impulsion est meilleure dans le détecteur central. Il permet aussi de reconstruire des candidats muons qui n’ont pas de coups dans les couches BC (un seul segment dans

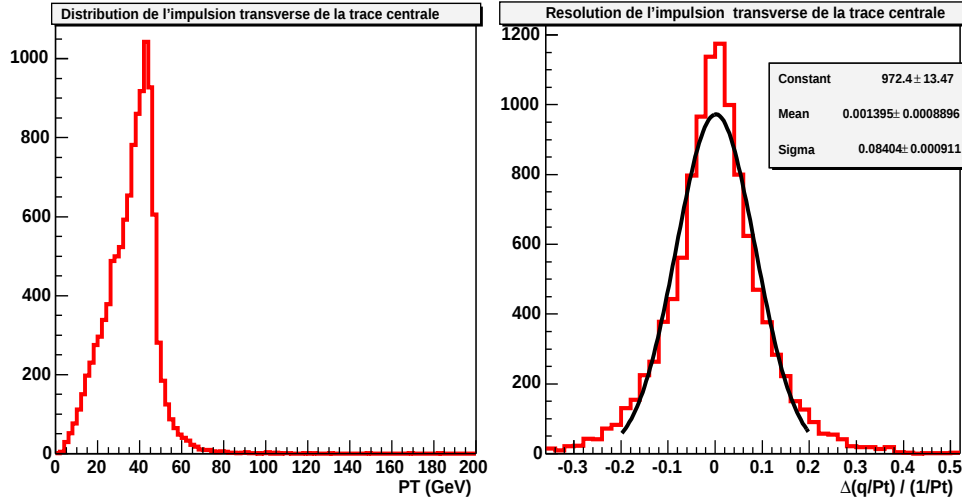


FIG. 4.8 – Distribution de l'impulsion transverse des muons (à gauche) et précision de mesure de l'impulsion des muons dans le système central issus de la désintégration d'un Z (à droite) : en abscisse est montrée $(q_g/P_{Tg} - q_r/P_{Tr}) * P_{Tg}$ où q_g (q_r) et P_{Tg} (P_{Tr}) sont la charge et l'impulsion transverse générées (reconstruites).

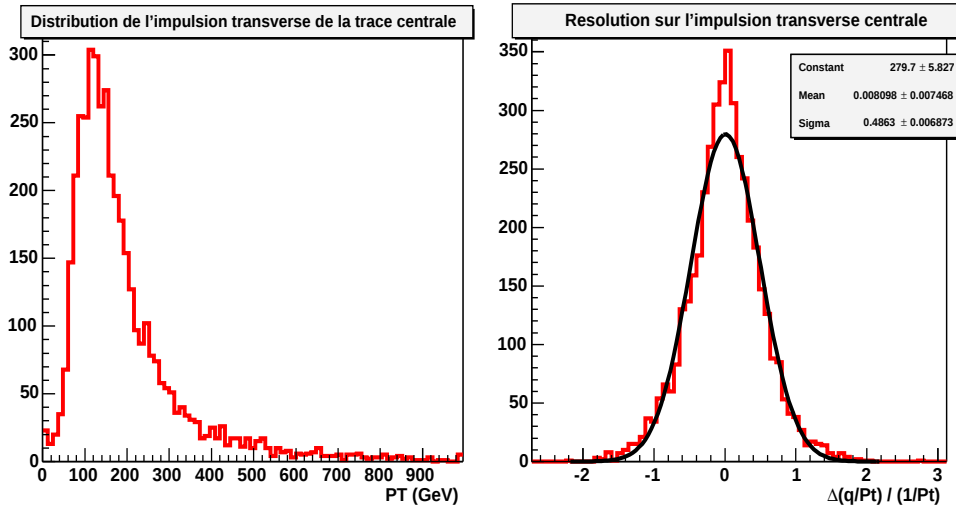


FIG. 4.9 – Distribution de l'impulsion transverse des muons (à gauche) et précision de mesure de l'impulsion des muons dans le système central issus de la désintégration d'un graviton de 400 GeV de masse (à droite) : en abscisse est montrée $(q_g/P_{Tg} - q_r/P_{Tr}) * P_{Tg}$ où q_g (q_r) et P_{Tg} (P_{Tr}) sont la charge et l'impulsion transverse générées (reconstruites).

la couche A).

L'idée est de partir d'abord de la trace du muon. On propage alors celle-ci du système à muon jusqu'au détecteur de traces centrales.

Si au moins une trace dans le détecteur central est trouvée dans un cône d'acceptance (dont l'ouverture dépend de la précision de mesure de l'impulsion dans le système à muons), un ajustement est effectué sur l'appariement de ces traces en tenant compte

des matrices d'erreur des deux traces.

Les traces centrales d'impulsion inférieure à 1.5 GeV sont rejetées dès le départ car un muon de si basse énergie ne peut sortir du calorimètre et atteindre la couche A du détecteur à muons.

Si à l'issue de cette procédure aucune trace centrale n'est associée à la trace locale du muon, alors la procédure inverse est suivie. Il s'agit de boucler sur toutes les traces trouvées dans le détecteur de traces central et voir si pour chaque trace centrale on peut associer une trace locale.

Chaque trace centrale est extrapolée à partir de son point le plus proche de la ligne de faisceau jusqu'à la couche externe du CFT (à un rayon de 51 cm), en tenant compte de la courbure de la trace due au champ magnétique créé par le solénoïde. La trace est ensuite propagée jusqu'à la couche A du détecteur à muons, en combinant les matrices d'erreurs.

Les paires ainsi formées sont classées par meilleure convergence.

Cet algorithme est aussi applicable à des segments de la couche A qui ne sont associés à aucun segment BC et aux segments BC seuls.

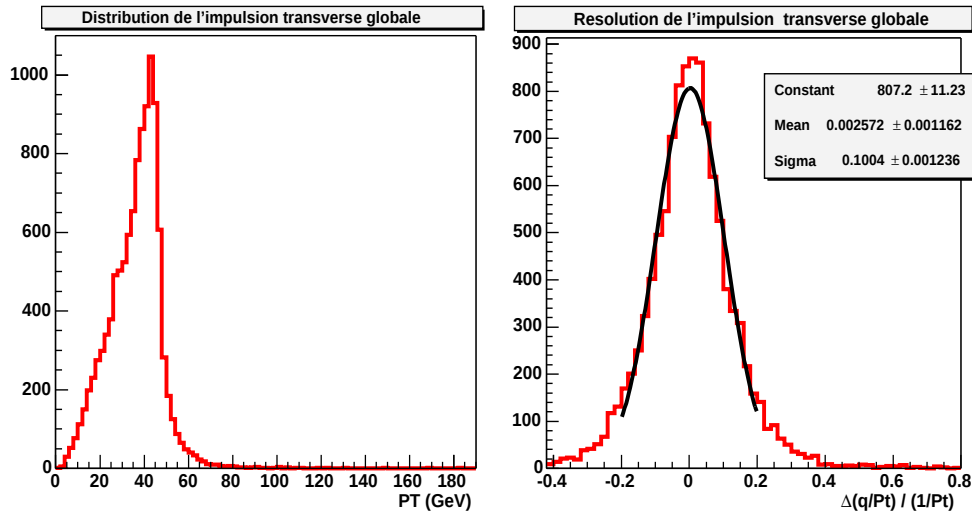


FIG. 4.10 – Distribution de l'impulsion transverse des muons (à gauche) et précision de mesure de l'impulsion des muons appariés à une trace centrale issus de la désintégration d'un Z (à droite) : en abscisse est montrée $(q_g/P_{Tg} - q_r/P_{Tr}) * P_{Tg}$ où q_g (q_r) et P_{Tg} (P_{Tr}) sont la charge et l'impulsion transverse générées (reconstruites).

La précision de mesure de l'impulsion transverse des muons après association de la trace centrale avec la trace locale est montrée sur la figure 4.10 pour des muons issus de la désintégration du Z et sur la figure 4.11 pour des muons issus de la désintégration d'un graviton de masse 400 GeV. Les résolutions données par l'ajustement global des traces des muons sont très proches de celles données par les traces centrales. En effet, nous obtenons 10% de résolution pour des muons issus de la désintégration du Z et 48% pour des muons issus de la désintégration d'un graviton.

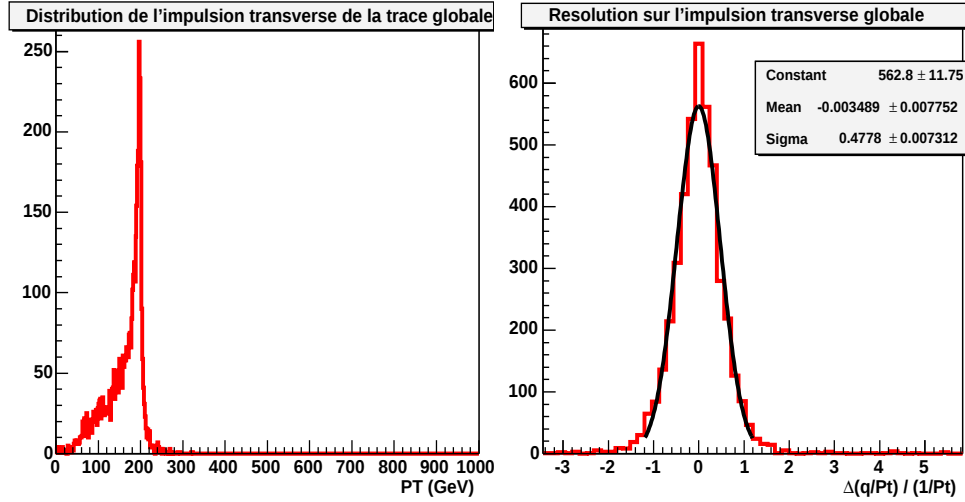


FIG. 4.11 – *Distribution de l'impulsion transverse des muons (à gauche) et précision de mesure de l'impulsion des muons appariés à une trace centrale issus de la désintégration d'un graviton de 400 GeV de masse (à droite): en abscisse est montrée $(q_g/P_{Tg} - q_r/P_{Tr}) * P_{Tg}$ où q_g (q_r) et P_{Tg} (P_{Tr}) sont la charge et l'impulsion transverse générées (reconstruites).*

L'impulsion P du muon global résulte de la combinaison des matrices d'erreurs. Si aucune valeur "correcte" de l'impulsion locale n'est donnée (cas de segments A ou BC seuls), alors l'impulsion du muon sera celle de la trace centrale.

Si aucun appariement n'est réalisé, le muon ne contient alors que les informations données par le système local. Dans ce cas, ne sont gardés que les muons dont la trace locale est reconstruite avec un segment de la couche A et un segment de la couche BC.

4.2.4 Les muons dits *calorimétriques*

Certains muons ont une impulsion assez faible (inférieure à 3-4 GeV) et ne sont pas assez énergétiques pour traverser le toroïde. De plus, s'ils n'ont pas laissé suffisamment de coups dans la couche A du système à muons, on ne peut pas leur associer un segment dans le système local.

Dans le calorimètre, ils laissent cependant une trace linéaire de basse énergie (environ 3-4 GeV) qui correspond au dépôt moyen d'énergie par le muon à sa traversée de la matière présente dans le calorimètre. Ces traces sont appelées "muons calorimétriques".

Dans notre sélection, nous ne prendrons pas en compte ces muons puisque les événements que nous étudions sont à très grande impulsion transverse. Ils seront automatiquement rejetés lors de l'application des coupures que nous détaillerons dans le chapitre suivant. Par contre ils seront utiles pour le calcul de certaines efficacités de sélection.

4.2.5 Comparaison entre le Monte Carlo et les données

La distribution de l'impulsion transverse des muons globaux est montrée sur la partie gauche de la figure 4.12. Sur la partie droite de cette figure nous montrons la distribution de la masse invariante des deux muons pour les données et pour des événements simulés.

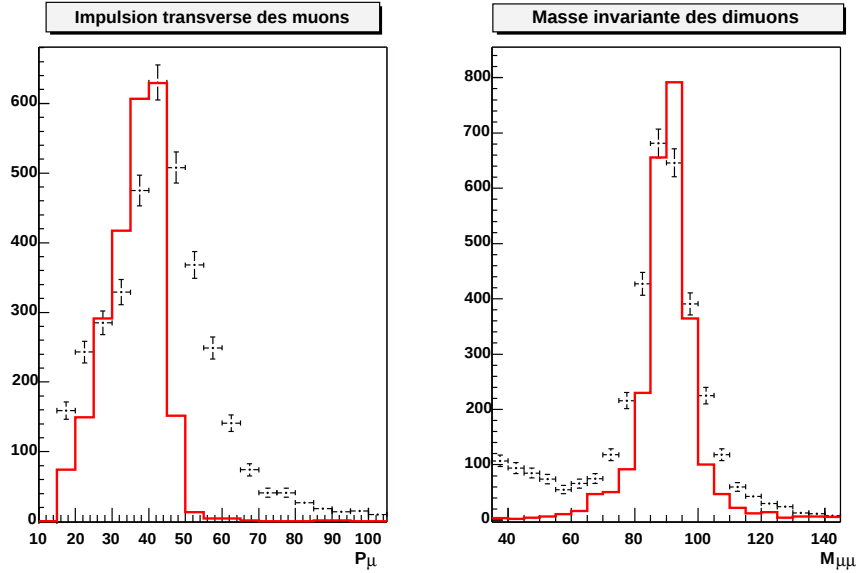


FIG. 4.12 – Comparaison entre les données (indiquées par des croix) et le MC (en trait plein) avant correction du MC. La figure de gauche montre la comparaison entre la distribution de l'impulsion transverse des muons dans les événements Z dans les données et dans le Monte-Carlo. La figure de droite montre la comparaison entre la distribution de masse invariante des dimuons donnée par le Monte-Carlo et les données.

Nous observons une différence assez nette entre ce que donne le Monte Carlo (que nous noterons MC dans la suite) et ce qui est observé sur les données. Ceci est dû au fait que le MC surestime la précision de mesure de l'impulsion transverse des muons. Par ailleurs, les données ne tiennent pas compte du décalage dans le plan transverse du SMT par rapport au CFT. Il faut donc appliquer deux corrections, une première appliquée aux données et qui tient compte de ce décalage entre le SMT et le CFT, la deuxième appliquée sur le MC et qui tient compte de l'optimisme de la simulation.

Pour tenir compte du décalage au sein du détecteur de traces central, il faut corriger l'impulsion transverse des traces centrales de la façon suivante [14]. Comme nous nous intéressons à l'impulsion des muons, nous n'avons aucun moyen d'effectuer la correction adéquate à l'impulsion des muons directement sur les données étant donné que l'information sur l'impulsion des muons qui est fournie tient compte des ajustements effectués lors de l'appariement des traces locales et centrales. Sachant que l'impulsion des muons est principalement donnée par la trace centrale, nous avons alors appliqué la correction valable pour les traces centrales directement aux muons, sachant que cette information n'est pas tout à fait correcte. Toutefois, cette correction est plus proche de l'impulsion

réelle des muons que la valeur de l'impulsion avant correction.

L'autre correction concerne le décalage entre le MC et les données. On veut appliquer une correction sur la valeur de l'impulsion transverse des muons afin de retrouver la largeur du pic du Z dans les données. Les centres des pics trouvés par le MC et par les données étant très proches, on peut considérer qu'il n'y a pas de décalage du pic, mais juste une sous-estimation de la largeur. Pour cela plusieurs façons d'appliquer une correction ont été testées et ont montré qu'elles étaient équivalentes [15]:

- corriger l'impulsion reconstruite:

$$\frac{1}{P_T^{reco}} = \frac{1}{P_T^{reco}} + f \times G \quad (4.1)$$

où f est un paramètre à définir et G est une gaussienne de largeur 1.

- corriger la différence entre l'impulsion générée et l'impulsion reconstruite:

$$\frac{1}{P_T^{reco}} = \frac{1}{P_T^{reco}} + f \left(\frac{1}{P_T^{reco}} - \frac{1}{P_T^{gen}} \right) \quad (4.2)$$

où P_T^{reco} (P_T^{gen}) est l'impulsion transverse reconstruite (générée) et f est à déterminer.

- corriger directement l'impulsion générée avant reconstruction:

$$\frac{1}{P_T^{gen}} = \frac{1}{P_T^{gen}} \left(1 + G \sqrt{A^2 P_T^2 + \frac{B}{\sin \theta}} \right) \quad (4.3)$$

Nous utilisons la première méthode qui est plus simple à intégrer dans le code d'analyse et donne d'aussi bons résultats que les deux autres méthodes. La valeur de f qui permet un bon accord entre les données et le MC est $f = 0.023$.

Les distributions obtenues après correction pour l'impulsion transverse des muons et pour la masse invariante sont montrées sur la figure 4.13 . Nous voyons que l'accord entre les données et le MC est raisonnable. Pour nos études sur le MC, nous nous baserons sur ces résultats.

4.3 Critères de qualité des muons

Selon le nombre de coups utilisés pour la reconstruction des muons dans le système local, des critères de qualité sont appliqués afin de rejeter les événements produits par du bruit de fond électronique et de purifier au mieux l'échantillon de données analysées. Ces critères de qualité s'appliquent au muon reconstruit dans le système local, et ne reposent pas sur des critères liés à la qualité des traces du détecteur de traces central.

Ainsi, la qualité d'un muon local est répartie en trois niveaux [16]: un muon de qualité excellente, un muon de qualité moyenne et un muon de qualité acceptable.

Critère de qualité excellente

Un muon est de qualité excellente s'il a :

- au moins 2 fils touchés dans la couche A ;

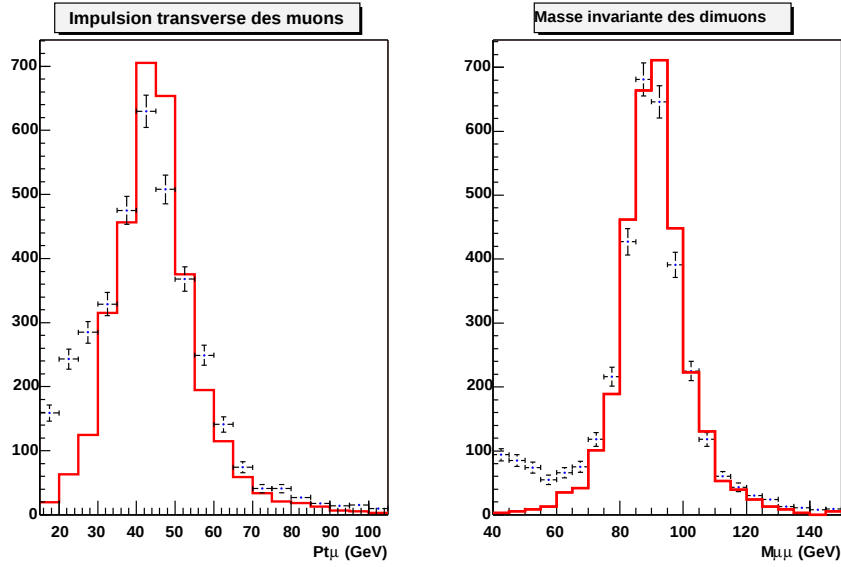


FIG. 4.13 – Comparaison entre les données (indiquées par des croix) et le MC (en trait plein) après correction du MC. La figure de gauche montre la comparaison entre la distribution de l'impulsion transverse des muons dans les événements Z dans les données et dans le Monte-Carlo. La figure de droite montre la comparaison entre la distribution de masse invariante des dimuons donnée par le Monte-Carlo et les données. L'excès d'événements à basse masse contenus dans les données par rapport au MC s'explique par la présence de bruits de fond QCD et le DY qui ne sont pas simulés par le MC utilisé ici.

- un scintillateur de la couche A touché ;
- au moins 3 coups dans les couches B et C (PDT ou MDT) ;
- au moins 1 scintillateur B ou C touché ;
- une trace locale reconstruite et dont l'ajustement est convergent.

Critère de qualité moyenne

Un muon est de qualité moyenne s'il a :

- au moins 2 fils touchés dans la couche A ;
- au moins 2 fils touchés dans les couches B et C (PDT ou MDT) ;
- au moins 1 coup enregistré dans les scintillateurs de la couche A ;
- au moins 1 coup enregistré dans les scintillateurs des couches B ou C.

Critère de qualité acceptable

Un muon est de qualité acceptable s'il satisfait à au moins trois des quatre critères requis pour les muons de qualité moyenne. Ce critère n'est pas appliqué pour les muons qui se trouvent dans la zone correspondant au trou dans la couche A du détecteur (octants 5 et 6).

La sélection de muons de qualité au minimum acceptable permet de s'affranchir d'une partie du bruit de fond non muonique. Ce dernier provient principalement d'éléments des gerbes hadroniques très énergétiques qui traversent le calorimètre et laissent des coups dans la couche A des chambres à muons. Il est par contre impossible pour ces jets de

traverser le toroïde en fer, et donc aucun coup n'est laissé par ces jets dans les couches B et C. On peut se débarrasser de ces faux muons en posant une condition sur la qualité des muons en demandant qu'ils satisfassent à au moins un des critères de qualité. Nous pouvons aussi limiter ces événements de fond en imposant un critère d'isolation du muon, ce que nous détaillons dans la section 5.5 .

4.4 Efficacités de reconstruction

Etant donné que l'on requiert des événements ayant déclenché le *trigger* dimuon (voir détails dans la section 5.3), il nous faut déterminer l'efficacité de reconstruction des muons due à l'acceptance géométrique du *trigger* dimu. L'efficacité de déclenchement étant calculée par ailleurs et pour éviter d'introduire des biais dans les calculs d'efficacité, on ne sélectionne que les événements dans lesquels au moins deux muons se trouvent dans l'acceptance géométrique du système de déclenchement des muons, c'est-à-dire à $|\eta| \leq 2^2$.

Par ailleurs, la partie inférieure du système à muon ne couvre pas la région en $|\eta| < 1.1$ et $4.25 \leq \phi \leq 5.15$. Ce trou est une cause de perte d'efficacité de reconstruction des muons qui apparaît clairement dans les distributions angulaires des muons dans les événements sélectionnés dans les données.

Dans les zones de détecteur correspondant à la jonction des régions avant/central, lorsqu'un muon laisse des coups, il peut traverser alternativement des couches de MDT et PDT, ce qui rend la reconstruction des muons dans ces zones très compliquée (voir figure 4.14). Même si l'effort est mis sur l'appariement des segments dans ces zones mixtes, les muons ne sont pas toujours reconstruits avec tous les coups disponibles, et souvent seul le segment dans la couche A est reconstruit (si deux couches A sont touchées, les deux segments sont gardés, mais un seul sera associé à une trace centrale).

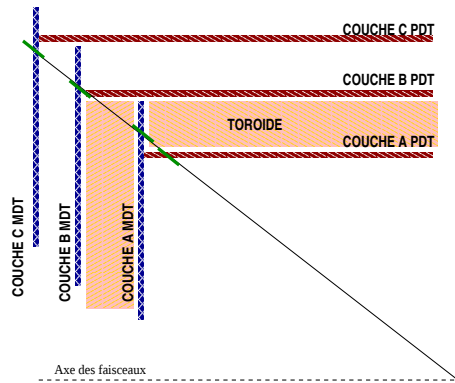


FIG. 4.14 – Schéma représentant un muon traversant des couches de PDT et MDT. L'information laissée par le muon dans chacun des deux types de chambres n'est pas regroupée.

De ce fait, l'efficacité de reconstruction des muons dans les zones mixtes sera diminuée du fait des muons perdus dans ces zones.

2. Dans ce cas, η est le η_{det} pris au niveau du système à muons.

D'autre part, on requiert aussi que les muons soient associés à une trace du détecteur de traces central. De la même façon que précédemment, on devra donc sélectionner les événements qui sont dans l'acceptance géométrique du CFT, c'est-à-dire imposer que $|\eta| \leq 1.6^3$. Cette condition n'est pas une simple restriction à de plus petites valeurs de η de la condition précédente. En effet, selon la position du vertex de l'interaction suivant l'axe (Oz), un muon peut se situer dans l'acceptance géométrique du CFT et pas dans l'acceptance du système à muons, et inversement.

Dans les événements que nous avons sélectionnés pour notre analyse, les muons doivent se trouver tous les deux dans l'acceptance géométrique du système à muons ainsi que du détecteur de traces central que nous appellerons globalement par la suite "acceptance géométrique du détecteur", sauf si nous spécifions un sous-détecteur particulièrement. Les corrélations angulaires entre les deux muons doivent être prises en compte (voir figure 4.15 et figure 4.16 réalisées pour du MC et qui représentent les variables générées).

On doit tenir compte de ces corrélations angulaires lors du calcul de l'efficacité due à l'acceptance géométrique du détecteur en calculant directement le rapport entre le nombre de paires de muons dans l'acceptance géométrique du détecteur et le nombre de paires de muons initiales. Bien sûr, cette efficacité dépend de la particule (Z ou graviton) et de la masse de la résonance. En effet, une particule très massive, un graviton de 800 GeV par exemple est produite presque au repos. Les particules sortantes sont pratiquement dos-à-dos (voir figure 4.15 pour les masses de graviton ≥ 400 GeV et qui représentent les variables générées).

signal	Nombre de paires générées	Nombre de paires dans l'acceptance	efficacité
Z	12000	4102	0.342 ± 0.004
G M=100 GeV	4000	2070	0.518 ± 0.008
G M=200 GeV	4000	2276	0.569 ± 0.008
G M=400 GeV	4000	2236	0.559 ± 0.008
G M=600 GeV	4000	2219	0.557 ± 0.008
G M=800 GeV	4000	2255	0.564 ± 0.008

TAB. 4.1 – *Efficacité due à l'acceptance géométrique du détecteur.*

Nous disposons pour cela d'échantillons de 4000 événements pour chaque valeur de masse de graviton (M=200 GeV, 400 GeV, 600 GeV, 800 GeV) et d'un échantillon de 12000 événements Z simulés. Nous calculons ensuite le nombre de paires de muons qui ont passé la coupure de l'acceptance géométrique.

L'efficacité ainsi trouvée est reportée dans le tableau 4.1. Elle correspond aux valeurs générées. Nous voyons que pour le graviton, l'efficacité due à l'acceptance géométrique du détecteur est quasiment la même quelle que soit la masse du graviton, alors que pour le Z l'efficacité est particulièrement faible. Cela est dû principalement à la distribution angulaire intrinsèque des muons qui diffère entre le graviton et le Z. En effet, les muons

3. Dans ce cas, η est le η_{det} . pris au niveau du CFT.

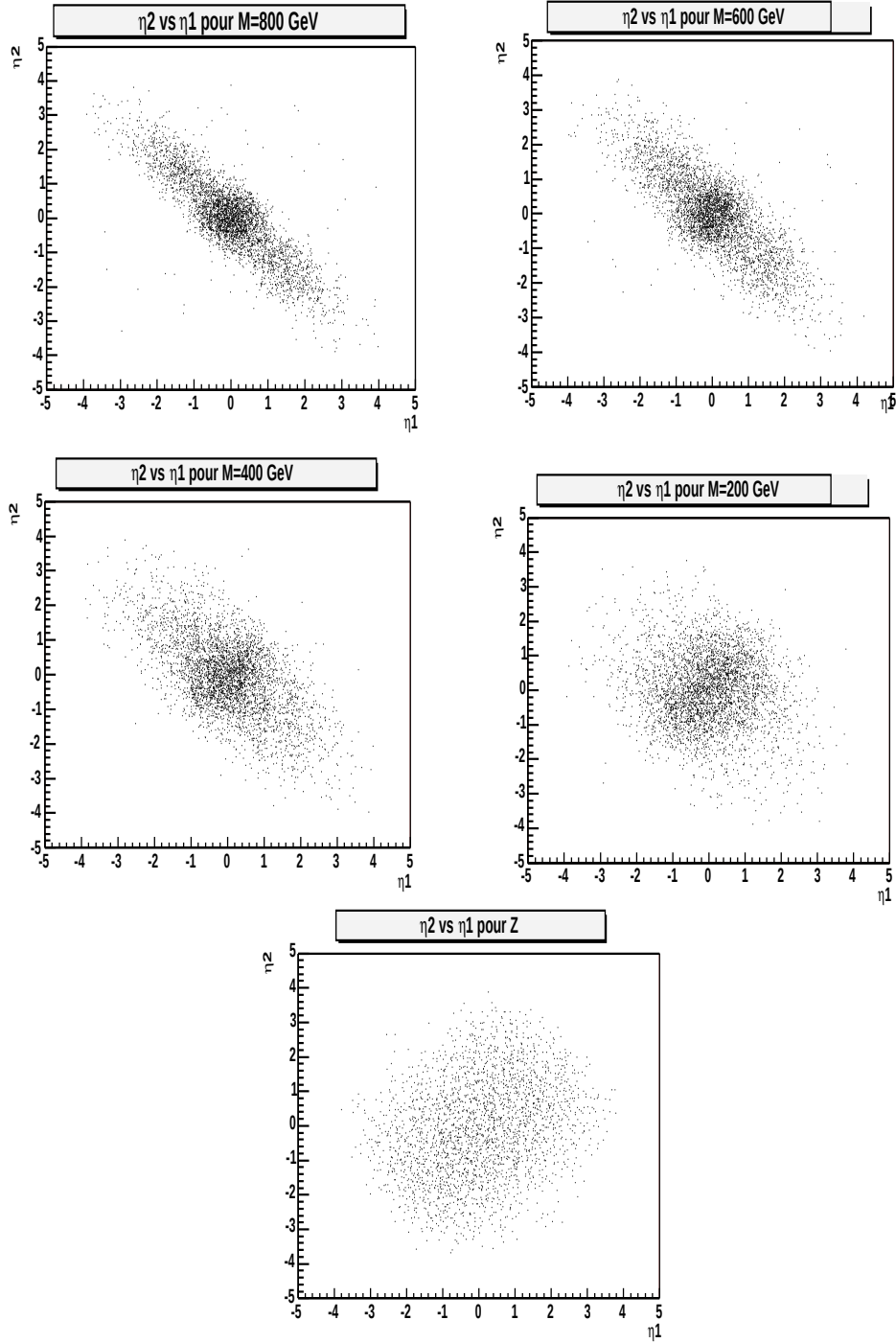


FIG. 4.15 – Distributions angulaires entre les deux muons issus de la désintégration d'un Z et de gravitons de masse 200 GeV, 400 GeV, 600 GeV et 800 GeV. Chaque figure montre la distribution de η_2 de l'antimuon en fonction de η_1 du muon .

issus de la désintégration d'une particule de spin 2 ont une distribution angulaire plus piquée à $|\eta| = 0$ que celle des muons issus de la désintégration d'une particule de spin 1 (voir figure 4.17) . La coupure à $|\eta| < 1.6$ due à l'acceptance géométrique du CFT élimine davantage d'événements "Z" que d'événements "gravitons".

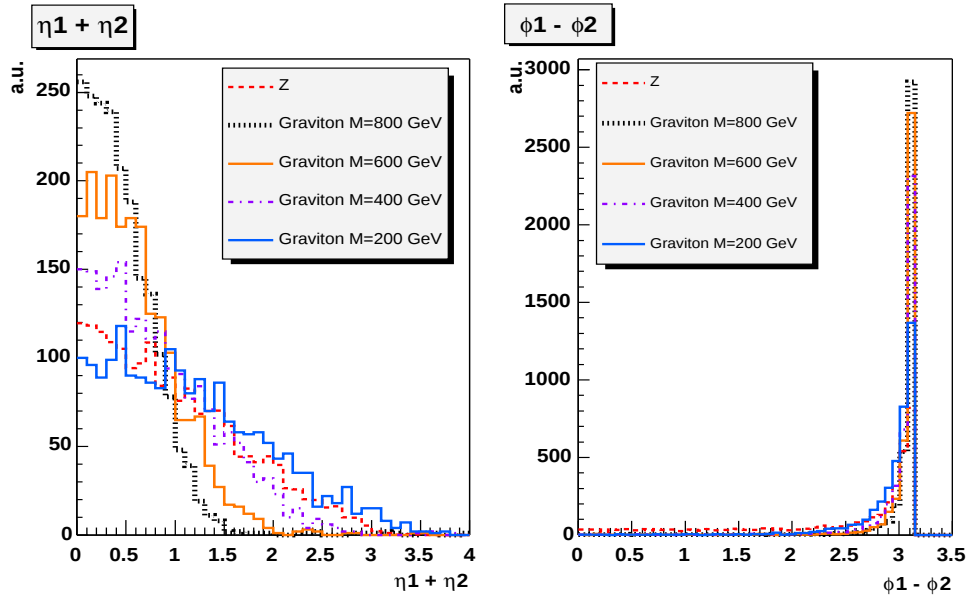


FIG. 4.16 – Distributions angulaires entre les deux muons issus de la désintégration d'un Z et de gravitons de masse 200 GeV, 400 GeV, 600 GeV et 800 GeV. La figure de gauche montre la somme des valeurs de η des deux muons, la figure de droite montre la différence en ϕ entre les deux muons.

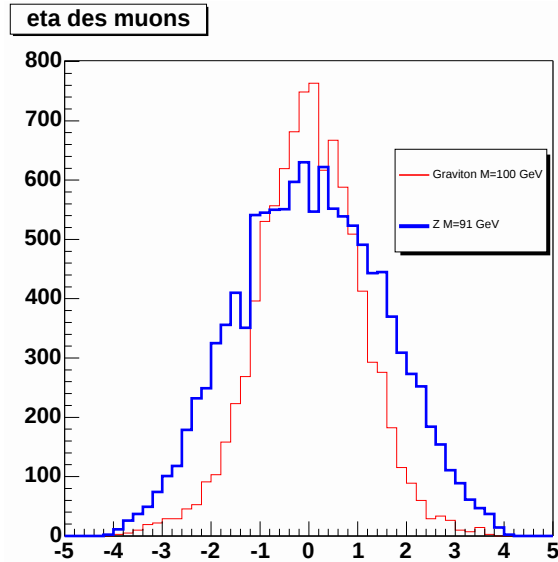


FIG. 4.17 – Distribution de η physique pour les deux muons issus de la désintégration d'un Z (histogramme en gras) et ceux issus de la désintégration d'un graviton de 100 GeV de masse (en trait fin).

Bibliographie

- [1] http://www-d0.fnal.gov/phys_id/muon_id/d0_private/muonid_primer.html
- [2] Muon Level-2 trigger Algorithmes,
<http://www-d0.fnal.gov/~maciel/l2alg/index.html>
- [3] www-d0.fnal.gov/phys_id/muon_id/d0_private/l3_muon.html
- [4] F. Déliot, thèse de doctorat de l'université Paris VII, avril 2002.
- [5] B.Gomez, *Time To Distance Relationship for the DØ Proportionnal Drift Tubes PDTs*,
DØ note 3892 (2001).
- [6] O.Peters, DØ note 3901 (2001),
www-d0.fnal.gov/d0pub/d0_private/3901/m_MuonSegmentReconstruction.doc
- [7] The DØ Collaboration, *Content of the p13 Muon Thumbnail*,
DØ note 4091 (2003).
- [8] http://www-d0.fnal.gov/global_tracking/algorithms.html
- [9] http://www-d0.fnal.gov/global_tracking/results/p14.00.00_index.html
- [10] G.Hesket, *Centrak Track Extrapolation Through the DØ Detector*,
DØ note 4079.
- [11] A.Khanov, *HTF: Histogramming method for finding tracks. The algorithm description*,
DØ note 3778 (2000).
- [12] <http://www-d0.fnal.gov/computing/www/review/tarc02/TARCRecom.pdf>
- [13] P.Balm, DØ note 3934 (2002),
www-d0.fnal.gov/d0pub/d0_private/3934/L3CentralTrackMatching_v30.doc
- [14] http://www-d0.fnal.gov/~kazu/d0_private/emid/tracking_Apr24_03.pdf;
http://www-d0.fnal.gov/~zdrazil/d0_private/zmm0501.html;
http://www-d0.fnal.gov/d0_private/4161/m_align.pdf.
- [15] M. Agelou, A. Askew, M. Besançon, F. Déliot, N. Lahrichi, Y. Maravin, E. Perez, B. Tuchming, M. Verzocchi, *Measurement of the inclusive $W \rightarrow \mu\nu$ Cross-Section in $p\bar{p}$ collisions at $\sqrt{s} = 1.96$ TeV*,
DØ note 4128, (2003).
- [16] http://www-d0.fnal.gov/d0dist/dist/releases/development/muo_cand/src/MuoCandidate.cpp;
http://www-d0.fnal.gov/phys_id/muon_id/d0_private/talks/adm151102.ppt

Chapitre 5

Calcul de $\sigma(Z).\text{Br}(\mu\mu)$

Introduction

L'état final du signal de graviton recherché est une paire de muons de signes opposés dont la masse invariante est très élevée (masse du graviton de l'ordre de quelques centaines de GeV) . Pour la sélection des données nous pouvons reconstruire *dans un premier temps* des événements Z à partir desquels nous pouvons calculer la section efficace de production de Z dans ses désintégrations muoniques.

Cette première étape est intéressante et constructive car elle permet d'une part de contrôler la sélection des données et de s'assurer de la cohérence de cette première partie d'analyse en comparant la mesure à la section efficace théorique bien connue, et d'autre part elle est une étape vers la recherche de graviton, étant donné que les critères de sélection des événements Z en deux muons sont repris pour la recherche d'événements graviton en deux muons.

Dans ce chapitre nous décrivons la sélection des événements de notre lot de données en appliquant des coupures que nous détaillons dans ce qui suit. Nous disposons des données prises entre novembre 2002 et juillet 2003, ce qui correspond à une luminosité de $107,8 \text{ pb}^{-1}$ (voir section 5.6) .

Certaines des efficacités de coupure sont spécifiques au Z, comme l'efficacité de coupure sur l'écart angulaire entre les muons ($|\Delta\phi| \geq 2$ radians) ou la coupure sur l'impulsion transverse du muon. Mais d'autres efficacités sont plus générales ou tout simplement transposables au cas du graviton, comme l'efficacité d'appariement des muons locaux avec des traces centrales, de sélection de muons de qualité intermédiaire, ou les coupures d'isolation.

Nous détaillerons dans ce qui suit la sélection des données analysées, la justification des coupures appliquées ainsi que la méthode de calcul de leurs efficacités. Cette étude nous permettra de calculer la valeur de la section efficace de production du boson Z dans le canal muonique. Les événements ainsi sélectionnés serviront de base pour l'étude du signal graviton que nous rechercherons dans la queue de la distribution de la masse invariante des muons.

La valeur de la section efficace de production du boson Z dans le canal muonique ($\sigma(Z) \times Br(Z \rightarrow \mu\mu)$) donnée par PYTHIA est de 184 pb à l'ordre des arbres. Un facteur de 30% [1] doit être ajouté à cette valeur afin de tenir compte des contributions des diagrammes aux ordres supérieurs, ce qui donne une valeur de la section efficace d'environ 240 pb.

Pour notre mesure de la section efficace de production du Z dans le canal muonique, nous utilisons la relation suivante :

$$\sigma(Z).Br(\mu\mu) = \frac{N \times (1 - f)}{\epsilon \times \mathcal{L}}$$

où :

- N est le nombre final de candidats Z en 2 muons après toutes les coupures,
- f est la proportion de faux événements Z qui contaminent le signal,
- ϵ est l'efficacité totale des coupures appliquées sur le Z,

- $\mathcal{L}$ est la luminosité intégrée de l'échantillon de données analysé.

Nous avons déjà vu dans le chapitre précédent les efficacités de reconstruction des muons dans l'acceptance géométrique du détecteur. Nous allons à présent étudier les efficacités de sélection des événements qui permettent de purifier l'échantillon de dimuons.

Dans tous les événements, nous appliquons dans l'ordre suivant les coupures permettant de sélectionner les paires de muons “opposés” en ϕ ($\Delta\phi > 2$ radians) :

- les deux muons sont dans l'acceptance géométrique du détecteur, c'est-à-dire dans l'acceptance géométrique du CFT ainsi que celle du système à muons ;
- les deux muons sont au moins de qualité moyenne ;
- les deux muons ont un $P_T > 15 \text{ GeV}$ (pour s'affranchir de la majeure partie du fond Drell-Yan et J/Ψ) ;
- les événements ont déclenché le L1 et le L2 du *trigger* 2MU_A_L2M0 ;
- les deux muons sont appariés à une trace centrale (muons globaux). Chaque trace doit être reconstruite avec au moins 3 coups dans le SMT et 3 coups dans le CFT ;
- au moins un muon est isolé énergétiquement¹ dans le calorimètre et dans le détecteur de traces central (pour s'affranchir du fond QCD contenant des muons dans des jets hadroniques) .

Pour notre étude, nous sélectionnerons les événements qui ont déclenché le *trigger* “dimuon” 2MU_A_L2M0 : il demande au moins deux muons au niveau L1 et est sensible à toute la couverture angulaire du détecteur de muons, c'est-à-dire pour $|\eta| < 2$. Au niveau 2 du déclenchement il requiert qu'au moins un muon soit de qualité moyenne (voir section 4.3). Aucun critère de sélection n'est appliqué sur l'impulsion transverse des muons.

Pour le calcul de l'efficacité de sélection sur le critère de qualité des muons détaillé dans la section 5.1, nous utilisons un échantillon sélectionné selon des *triggers* monomuons dont la couverture angulaire s'étend jusqu'à $|\eta| < 1,3$.

5.1 Sélection des muons de qualité moyenne

Le critère de qualité appliqué aux muons est le critère de qualité moyenne décrit dans le chapitre précédent. La raison de ce choix plutôt que des muons de qualité acceptable revient, à la base, au choix du lot de données disponibles pour l'analyse. En effet, les données ont été filtrées selon différents critères pour chaque type d'analyse réalisée au sein de la collaboration, et celles que nous utilisons pour l'analyse ont été sélectionnées en demandant deux muons de qualité moyenne pour seul critère.

Nous devons calculer l'efficacité du critère de sélection des muons de qualité moyenne sur les événements Z en dimuons ou graviton en dimuons. Pour cela nous avons pris un échantillon de données contenant principalement des événements Z en dimuons.

1. Le critère d'isolation consiste à demander d'une part que la somme des énergies transverses dans le calorimètre dans un “halo” autour de la trajectoire du muon soit inférieure à 2,5 GeV et d'autre part que la somme des impulsions transverses des traces présentes dans un cône autour de la trace du muon soit inférieure à 2,5 GeV. Le “halo” et le cône sont définis dans la section 5.5.

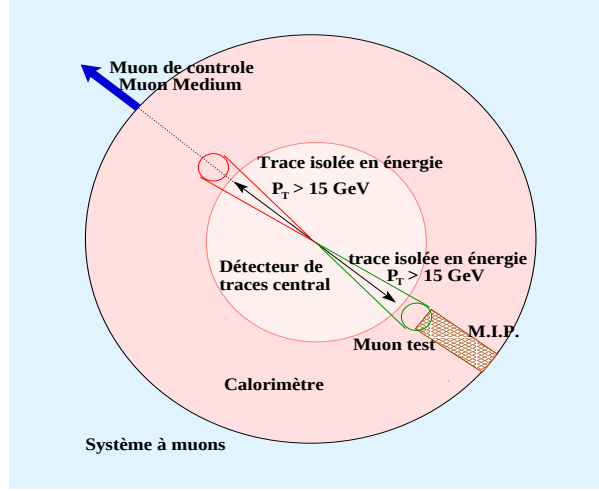


FIG. 5.1 – Schéma illustrant la sélection du muon de contrôle et du muon test.

Pour être sûr que ce que l'on teste est effectivement un muon, on sélectionne une paire de muons compatible avec un événement Z, c'est-à-dire deux muons dont la masse invariante est comprise entre 40 GeV et 140 GeV, auxquels on applique des coupures afin d'éliminer le bruit de fond dû aux paires de muons qui ne viennent pas du Z. La paire de muons est choisie en sélectionnant “un muon de contrôle” et “un muon test” définis comme suit :

un muon de “contrôle” :

- de qualité moyenne ;
- apparié à une trace centrale d'impulsion transverse supérieure à 20 GeV ;
- isolé énergétiquement à la fois dans le détecteur de traces central et dans le calorimètre ;
- situé dans l'acceptance géométrique du détecteur (nous rappelons que le η considéré est le η_{det} qui a été défini dans la section 3.2.2²) ;
- correspondant au muon qui a déclenché un *trigger* lié aux muons au niveau L2 et qui requiert au moins un muon de qualité moyenne³ ;
- la distance d'approche minimale entre la trace centrale et le point d'interaction (que l'on appellera DCA pour simplifier) doit être inférieure à 0,2 cm afin d'éliminer la majorité des événements cosmiques.

un muon “test” :

- une trace centrale d'impulsion transverse supérieure à 20 GeV et associée à une trace dans le calorimètre compatible avec le dépôt d'énergie d'un muon dans le calorimètre (voir section 4.2.3) ;
- la trace centrale doit pointer vers une zone où se trouvent déjà des coups dans les chambres à dérive de la couche A du système à muons ;

2. Dans ce cas nous demandons d'une part que $|\eta_{det}|$ au niveau du CFT soit inférieur à 1.6 qui correspond à l'acceptance géométrique du CFT, et d'autre part que $|\eta_{det}|$ calculé au niveau du système à muons soit inférieur à 2 et hors du trou de la couche A.

3. En toute rigueur nous devrions nous assurer que le muon de contrôle correspond aussi au muon qui a déclenché le *trigger* au L1, mais cette information n'est pas disponible. Nous disposons cependant de l'information concernant le muon qui a déclenché le *trigger* au niveau L2 et nous l'utilisons afin de minimiser le biais introduit par la corrélation entre la qualité du muon et le déclenchement du *trigger*.

- elle doit être isolée énergétiquement dans le détecteur de traces central ;
- son extrapolation doit être dans l'acceptance géométrique du détecteur à muons ;
- la DCA entre la trace centrale et le point d'interaction doit être inférieure à 0,2 cm afin d'éliminer la majorité des événements cosmiques ;
- la charge électrique du muon test doit être opposée à celle du muon de contrôle ;
- l'angle azimutal défini par la trace du muon de contrôle et la trace du muon test doit être supérieur à 2 radians .

Pour ce calcul, on demande que l'impulsion transverse de la trace centrale du muon test et du muon de contrôle soit supérieure à 20 GeV afin d'éliminer la majeure partie des événements J/ψ , Drell-Yan et $b\bar{b}$ à basse énergie⁴.

L'échantillon utilisé pour le calcul de l'efficacité de sélection des muons de qualité moyenne a déclenché un *trigger* à un seul muon qui n'est sensible, au niveau L1, qu'à la région $|\eta| < 1,3$, et requiert un muon de qualité moyenne au niveau L2 du déclenchement. Une condition additionnelle est appliquée au niveau L3, demandant qu'au moins une trace soit reconstruite au niveau 3 avec une impulsion transverse supérieure à 10 GeV, mais cette condition n'a pas de répercussion sur le résultat de notre calcul étant donné que les énergies des traces centrales associées aux muons que nous sélectionnons sont supérieures à 20 GeV et que cette information est indépendante de ce que l'on teste.

La figure 5.2 montre la distribution de masse invariante des paires de muons reconstruites pour tous les événements (histogrammes en blanc) et pour les événements dans lesquels le muon test n'est pas de qualité moyenne (histogrammes grisés). Pour calculer les efficacités, on utilise deux méthodes:

- sans soustraction du bruit de fond. C'est une méthode de simple comptage qui sert de contrôle. Pour calculer les efficacités sans soustraction de bruit de fond, on se place sous le pic du Z, c'est-à-dire à une masse invariante de la paire de muons comprise entre 75 GeV et 105 GeV. Les événements à prendre en compte pour le calcul des efficacités sont :
 - le nombre total N_T d'événements sous le pic ;
 - le nombre d'événements N_0 sous le pic dans lesquels le muon test n'est pas de qualité moyenne.

L'efficacité vaut dans ce cas $\epsilon_f = 1 - N_0/N_T$;

- avec soustraction de bruit de fond. On réalise un ajustement de la distribution de la masse invariante des deux muons et ce pour une masse comprise entre 40 GeV et 140 GeV pour tous les événements sélectionnés d'une part et pour les événements dans lesquels le muon test n'est pas de qualité moyenne d'autre part. La fonction d'ajustement est la somme d'une gaussienne modélisant le pic du Z et d'un polynôme de degré 1 modélisant le bruit de fond. Les nombres d'événements pris en compte dans ce cas sont :
 - N_T le nombre total d'événements donné par l'intégrale de la fonction d'ajuste-

4. Notons que pour la sélection des événements qui serviront pour notre analyse finale, nous demandons $P_T > 15 \text{ GeV}$.

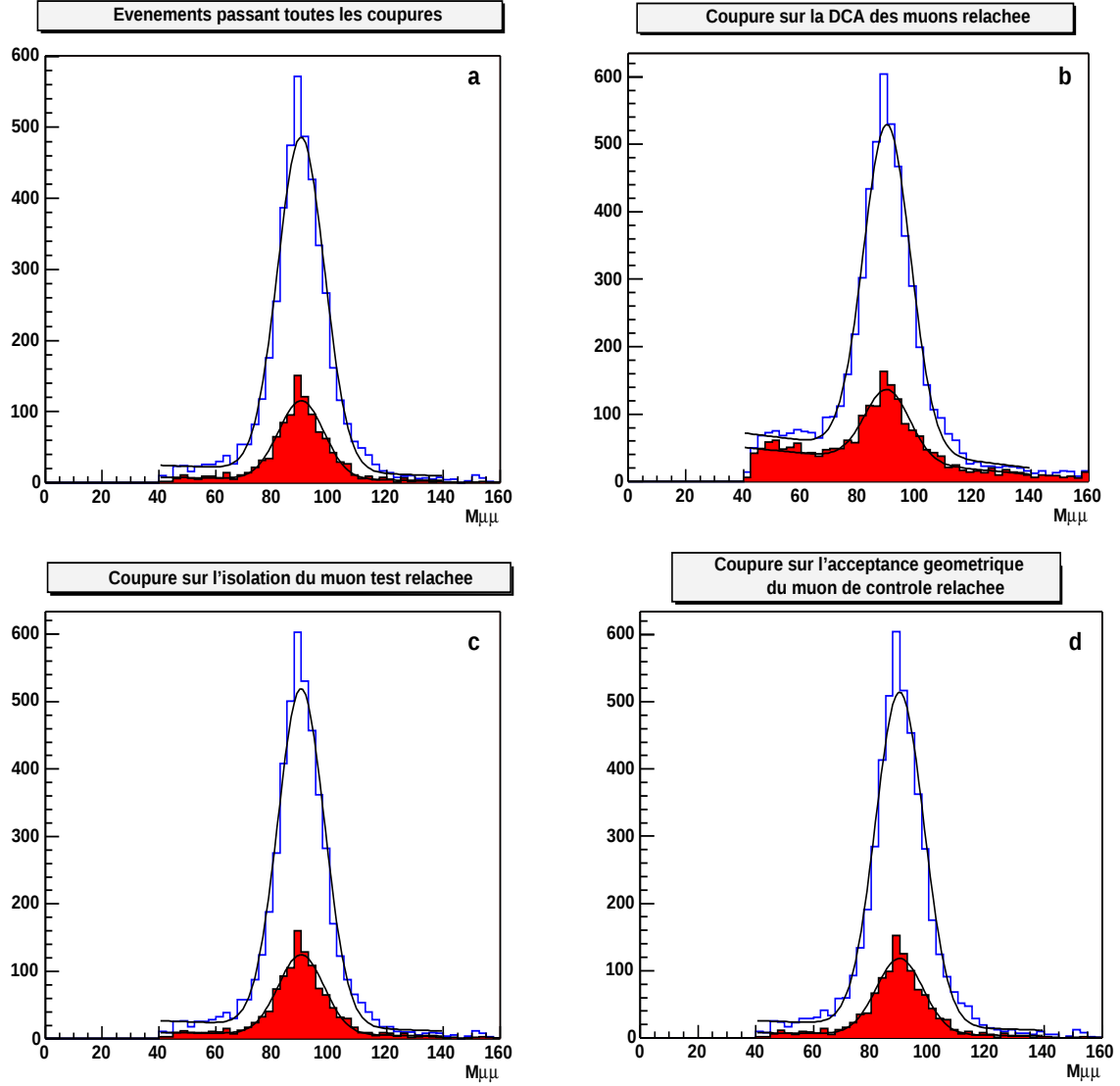


FIG. 5.2 – Masse invariante des dimuons pour tous les événements (histogrammes blancs) et pour les événements dans lesquels le muon test n'est pas de qualité moyenne (histogrammes grisés). Les distributions sont réalisées dans le cas où les événements passent toutes les coupures (a), le cas où on a relâché la coupure sur la distance minimale d'approche (DCA) des traces centrales des deux muons au point d'interaction (b), le cas où l'on n'exige pas l'isolation énergétique du muon test (c) et le cas où le muon de contrôle n'est pas obligatoirement dans l'acceptance géométrique du trigger (d).

ment ;

- N_{Tf} le nombre d'événements de bruit de fond contenus dans N_T . Ce nombre est donné par l'intégrale du polynôme de degré 1 de la fonction d'ajustement ;
- N_0 le nombre d'événements donné par l'intégrale de la fonction d'ajustement dans lesquels le muon test n'est pas de qualité moyenne ;
- N_{0f} est le nombre d'événements de bruit de fond contenus dans N_0 . Ce nombre est donné par l'intégrale du polynôme de degré 1 de la fonction d'ajustement.

L'efficacité dans ce cas est donnée par $\epsilon = 1 - (N_0 - N_{0f}) / (N_T - N_{Tf})$.

Les erreurs sur les efficacités calculées sont les erreurs statistiques.

Sur la figure 5.2-a, réalisée pour tous les événements ayant passé toutes les coupures concernant le muon de contrôle et le muon test, il y a très peu de bruit de fond, ce qui montre que nos coupures sont très sélectives. Les efficacités sont les suivantes :

- sans soustraction du bruit de fond : $\epsilon_f^a = 76,4\% \pm 0,68\%$;
- avec soustraction du bruit de fond : $\epsilon^a = 76,8\% \pm 0,69\%$.

L'efficacité ϵ^a est notre résultat de référence.

Afin de vérifier les effets de nos coupures sur le muon test et le muon de contrôle, nous relâchons quelques coupures.

Sur la figure 5.2-b nous avons relâché la coupure sur la distance de moindre approche des traces centrales des deux muons au point d'interaction. Cette coupure vise à l'origine à éliminer les événements cosmiques qui peuvent traverser le détecteur loin de la ligne des faisceaux. Les efficacités avec et sans soustraction du bruit de fond sont les suivantes :

- sans soustraction du bruit de fond : $\epsilon_f^b = 72,4\% \pm 0,68\%$;
- avec soustraction du bruit de fond : $\epsilon^b = 78,2\% \pm 0,67\%$.

Sans coupure sur la DCA l'efficacité est plus faible que dans le premier cas si l'on ne soustrait pas le bruit de fond. Ceci est dû à la contamination en événements cosmiques et, dans une moindre proportion, en événements QCD. Si l'on soustrait le bruit de fond, on retrouve une efficacité quelque peu supérieure à ϵ^a . Ceci indique que nous avons soustrait trop de bruit de fond contenant aussi des événements non cosmiques.

Sur la figure 5.2-c la coupure sur l'isolation énergétique de la trace du muon test est relâchée. Les distributions de masse montrent que la contamination en événements $b\bar{b}$ reste négligeable, que l'on demande ou non que le muon test soit de qualité moyenne. Les efficacités sont :

- sans soustraction du bruit de fond : $\epsilon_f^c = 75,9\% \pm 0,67\%$;
- avec soustraction du bruit de fond : $\epsilon^c = 76,6\% \pm 0,68\%$.

Ces résultats restent tout à fait compatibles avec la valeur de ϵ^a dans les marges d'erreurs. Nous pouvons constater que la méthode avec soustraction de bruit de fond donne de meilleurs résultats que la méthode sans soustraction de bruit de fond.

Finalement, la figure 5.2-d correspond aux événements dans lesquels aucune condition n'est requise sur la distribution angulaire du muon de contrôle qui peut alors se trouver hors de l'acceptance géométrique du *trigger* au niveau L1. Nous relâchons cette coupure pour quantifier le biais introduit du fait que l'on ne peut pas savoir si c'est le muon de contrôle qui a déclenché le *trigger* au L1 ou non. Concrètement, si ce n'est pas le muon de contrôle qui a déclenché le L1, c'est le muon que l'on teste qui l'a déclenché. Autrement dit, le muon test est de qualité suffisamment bonne pour déclencher le L1, augmentant ainsi la probabilité qu'il soit de qualité au moins moyenne. C'est ce que l'on constate en calculant l'efficacité dans le cas de la figure 5.2-d. Les efficacités sont :

- sans soustraction du bruit de fond : $\epsilon_f^d = 77,3\% \pm 0,65\%$

- avec soustraction du bruit de fond: $\epsilon^d = 77,7\% \pm 0,66\%$

Comme nous n'avons pas la possibilité de vérifier que c'est bien le muon de contrôle qui a déclenché le niveau L1 du *trigger*, nous contournons le problème en nous plaçant alors dans le cas précis où le muon de contrôle ne peut avoir déclenché le L1 (voir figure 5.3) .

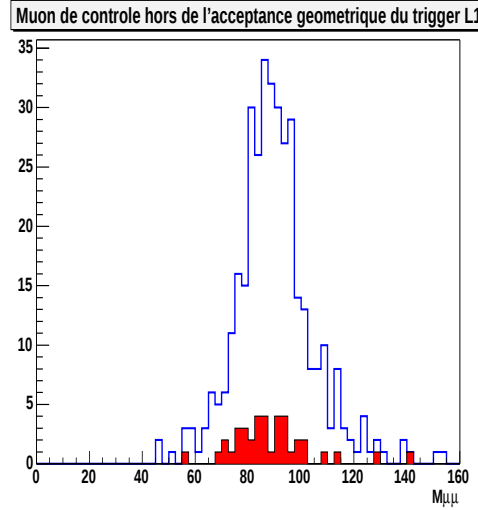


FIG. 5.3 – Distribution de masse invariante de paires de muons pour un muon de contrôle explicitement hors de l'acceptance géométrique du système de déclenchement au niveau L1 pour les triggers à un muon requis pour cette étude.

On sait que la couverture angulaire du L1 s'étend à des valeurs de $|\eta| < 1,3$, on demande alors explicitement que le muon de contrôle soit à $|\eta| > 1,5$. Les événements ainsi sélectionnés sont des événements dans lesquels le niveau L1 du *trigger* a été déclenché par le muon test. Il est donc naturel que l'on trouve une efficacité particulièrement élevée sans soustraction du bruit de fond: $\epsilon_f^{out} = 89,4\% \pm 0,18$. Cette valeur est de 13,6% supérieure à ϵ^a .

Il est certes rare qu'un muon ne déclenche pas le L1 mais déclenche le L2, de sorte que le nombre d'événements est très faible dans ce cas, comme le montre la figure 5.3 . Mais cela introduit un biais dans le calcul de l'efficacité que l'on peut estimer à partir de l'efficacité de déclenchement du *trigger* au L1 pour un muon de qualité moyenne qui est de 93,1% pour la période de prise de données la plus inefficace (voir section 5.3) . Ainsi dans 6,9% des cas un muon de qualité moyenne ne déclenche pas le *trigger* au L1, et la probabilité pour que ce même muon déclenche le *trigger* au L2 est bien plus faible. Nous pouvons estimer la proportion maximale d'événements dans lesquels le muon de contrôle a déclenché le L2 sans avoir déclenché le L1, ce qui correspond à $6,9\% \times 13,6\% = 0,94\%$ des cas. Nous pouvons prendre comme erreur systématique la moitié de cette valeur, ce qui donne une efficacité totale de sélection des muons sur le critère de qualité moyenne de : $\epsilon_{moy} = 76,8\% \pm 0,69(stat)\% \pm 0,47(syst)\%$.

La figure 5.4 montre l'efficacité de sélection des muons de qualité moyenne avec toutes les coupures appliquées sur le muon de contrôle et sur le muon test (le cas de la figure

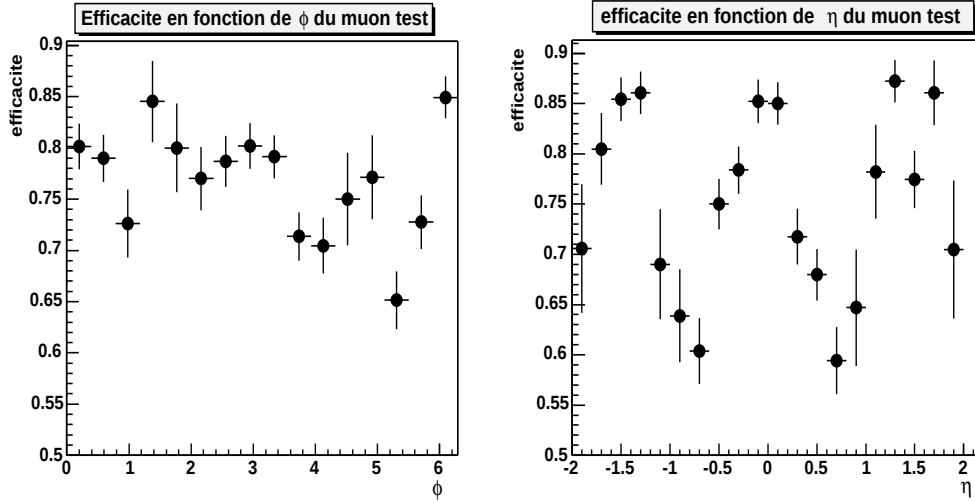


FIG. 5.4 – Efficacité de sélection de muons de qualité moyenne en fonction de la valeur de l’angle azimutal ϕ du muon test (à gauche) et de la valeur de η du muon test (à droite).

5.2-a) en fonction de la valeur de η et de la valeur de ϕ du muon test.

Sur la partie gauche de la figure 5.4, nous pouvons constater que l’efficacité de reconstruction d’un muon de qualité moyenne en fonction de l’angle ϕ du muon est globalement plate, moyennant les erreurs statistiques, sauf dans la partie $4 < \phi < 5,5$ où l’efficacité semble plus faible du fait de la présence du “trou” dans la couche A du système à muons dans la moitié basse du détecteur, correspondant à $|\eta| < 1,1$.

Sur la partie droite de la figure 5.4, nous voyons que l’efficacité dépend très fortement du η du muon test. La chute d’efficacité dans les régions en η autour de ± 1 est due à la mauvaise reconstruction de muons dans les régions mixtes décrites dans le chapitre précédent⁵. La chute d’efficacité pour des régions à grandes valeurs de η (typiquement $|\eta| > 1,8$) s’explique simplement par le fait que la couverture angulaire du système à muons s’étend jusqu’à $|\eta| < 2$.

5.2 Coupure sur l’impulsion transverse des muons

Pour nous affranchir des événements J/ψ ainsi que d’une grande partie du Drell-Yan, on doit effectuer une coupure sur l’impulsion transverse des muons de qualité moyenne. Nous avons vu par ailleurs dans le chapitre précédent que la précision de mesure de l’impulsion était bien meilleure pour des muons globaux que pour des muons locaux. C’est pour cela que nous calculerons l’efficacité de la coupure sur l’impulsion transverse en uti-

5. Pour remédier à cette perte d’efficacité conséquente les versions ultérieures à la version du programme de reconstruction que nous utilisons pour notre analyse permettent d’une part d’améliorer la reconstruction des traces locales dans ces régions et d’autre part de redéfinir les critères de qualité dans les zones particulières (incluant aussi le trou dans la couche A).

lisant l'impulsion transverse globale plutôt que l'impulsion transverse locale.

Nous disposons des échantillons d'événements Monte Carlo (nous utiliserons dorénavant la notation MC) $Z \rightarrow \mu\mu$ utilisés pour le calcul de l'efficacité de l'acceptance géométrique. Nous sélectionnons des muons de qualité moyenne appariés à une trace centrale et nous traçons l'impulsion transverse des muons globaux (voir partie gauche de la figure 5.5).

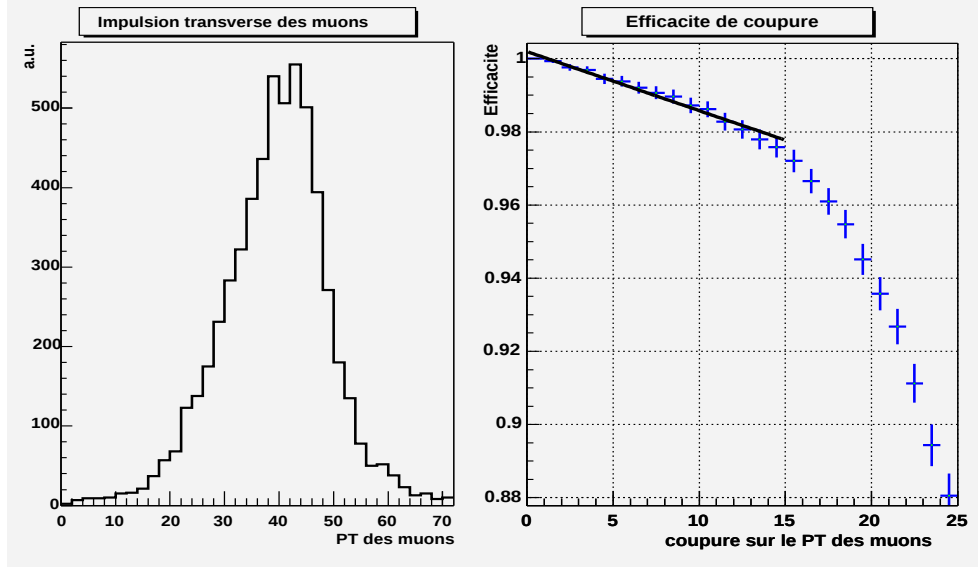


FIG. 5.5 – Distribution de l'impulsion transverse des muons (à gauche) et de l'efficacité de la coupure sur l'impulsion transverse en fonction de la valeur de la coupure (à droite). Cette figure utilise des événements MC $Z \rightarrow \mu\mu$ reconstruits, l'impulsion transverse prise en compte est l'impulsion transverse reconstruite dans le système global.

La partie droite de la figure 5.5 montre en abscisse la valeur de la coupure sur l'impulsion transverse, en ordonnée est reportée l'efficacité de cette coupure.

Nous pouvons constater que l'efficacité de la coupure décroît à peu près linéairement jusqu'à une coupure $P_T > 15$ GeV. Au delà de cette valeur, l'efficacité décroît très rapidement. Nous pouvons choisir cette valeur pour notre coupure. Ainsi, pour $P_T \geq 15$ GeV, 2823 événements passent la coupure sur les 2896 événements sélectionnés. L'efficacité est donc de $97,5\% \pm 0,3\%$.

L'efficacité de cette coupure est très satisfaisante et permet de se placer loin du pic du J/ψ (dont la masse est d'environ 3 GeV) et de s'affranchir d'une grande partie du Drell-Yan.

La paramétrisation de la correction de l'impulsion donnée par le MC (voir section 4.2.5) n'est pas parfaite. Pour évaluer le biais introduit par cette valeur sur notre estimation de l'efficacité de la coupure sur l'impulsion transverse des muons, nous pouvons choisir d'étudier :

- l'effet de la modification du facteur f introduit dans l'équation 4.1. Pour cela nous pouvons supposer que l'erreur sur ce paramètre est de $\pm 50\%$, on calcule alors l'efficacité pour une valeur de f prise à $f \times 1,5$ et $f \times 0,5$. Nous trouvons alors :
 - pour $f \times 1,5$: l'efficacité est $\epsilon_{1,5 f} = 2822/2896 = 97,4\%$;

- pour $f \times 0,5$: l’efficacité est $\epsilon_{0,5\ f} = 2822/2896 = 97,6\%$.
- ou l’effet de la modification de l’impulsion transverse elle-même après correction de $\pm 10\%$ qui correspond à la résolution sur l’impulsion globale des muons (voir section 4.2.3) . Nous trouvons alors :
 - pour $P_T \times 1,1$: l’efficacité est $\epsilon_{1,1\ P_T} = 2809/2896 = 97,0\%$;
 - pour $P_T \times 0,9$: l’efficacité est $\epsilon_{0,9\ P_T} = 2834/2896 = 97,9\%$.

La marge d’erreur introduite par le choix du facteur f et la résolution sur l’impulsion des muons peut être estimée à partir des erreurs précédentes, ce qui donne une erreur systématique maximale de $\pm 0,5\%$.

5.3 Déclenchement de l’acquisition

Le *trigger* pris en compte dans cette analyse est le *trigger* dimuon le moins restrictif. En effet, il requiert deux muons au niveau 1 du déclenchement en coïncidence avec le déclenchement des luminomètres (cette dernière condition n’est présente que sur une partie des données), et un muon de qualité moyenne au niveau 2, aucune restriction n’est faite sur l’impulsion minimale ni sur la région angulaire des muons.

L’échantillon de données utilisé pour notre analyse finale s’étale sur plusieurs mois d’acquisition. Durant cette période la définition des *triggers* a changé. Pour notre *trigger* dimuon, la condition de coïncidence entre l’information des scintillateurs à muons avec les luminomètres a été supprimée au niveau L1 au mois de février 2003, et le système de déclenchement du niveau L2 a été amélioré. Les efficacités correspondant aux périodes avant et après ces changements sont donc différentes. Pour tenir compte de ce changement nous séparons en deux parties les échantillons qui servent à l’étude de l’efficacité de déclenchement du *trigger* et l’efficacité finale est la moyenne des deux efficacités pondérées par la luminosité de chaque période.

Pour le calcul des efficacités de déclenchement, l’échantillon de données que nous utilisons a été présélectionné sur un critère de déclenchement de *triggers* purement lié à l’information du calorimètre (demandant que l’énergie dans une tour du calorimètre soit supérieure à un certain seuil) , donc totalement indépendant d’un quelconque critère relatif aux muons. Ceci nous permet de nous affranchir de tout biais de calcul lié au choix de l’échantillon.

Pour calculer l’efficacité de déclenchement du *trigger*, nous pouvons sélectionner des paires de muons et voir si ces événements déclenchent notre *trigger* dimuon au L1 et au L2. Cette méthode, que nous appellerons “méthode avec le *trigger* dimuon”, est rigoureuse mais souffre d’une très faible statistique étant donné qu’il est peu probable que des événements contenant deux muons dont un est isolé énergétiquement dans le calorimètre puisse déclencher les *triggers* calorimétriques. Il est alors possible de calculer ces efficacités avec une plus grande statistique en étudiant l’efficacité de déclenchement d’un “*trigger* monomuon” qui requiert les mêmes conditions que notre *trigger* dimuon au L1, et un autre *trigger* monomuon qui requiert les mêmes conditions que le *trigger* dimuon au L2,

à la seule différence près que les *triggers* monomuon utilisés ne requièrent qu'un muon au lieu de deux.

Dans ce cas, nous pouvons extraire très simplement des efficacités liées au *trigger* monomuon les valeurs des efficacités du *trigger* dimuon. Nos résultats seront donnés par cette deuxième méthode. La méthode avec le *trigger* dimuon servira de résultat de contrôle. L'écart entre les deux résultats donnera l'erreur systématique sur la mesure de l'efficacité de déclenchement au L1 ainsi qu'au L2.

Nous détaillons dans ce qui suit ces méthodes pour le L1 et le L2.

5.3.1 Déclenchement au niveau 1

Méthode avec le *trigger* monomuon

Nous calculons l'efficacité de déclenchement du *trigger* pour des événements ayant déclenché un *trigger* calorimétrique contenant un seul muon de qualité moyenne et d'impulsion transverse supérieure à 15 GeV (c'est-à-dire passant les deux premières coupures de sélection décrites précédemment). Ce muon doit se trouver dans l'acceptance géométrique du détecteur (c'est-à-dire du système à muons, du CFT et hors du trou de la couche A). Des coupures sur la DCA et les temps relevés par les scintillateurs sont appliquées afin de s'affranchir des événements cosmiques⁶.

Parmi les N_{tot} événements sélectionnés nous calculons le nombre d'événements N_{L1} qui ont déclenché le *trigger* monomuon au niveau 1.

La figure 5.6 montre l'efficacité du *trigger* L1 en fonction de η des muons. La valeur du plateau correspond à $\epsilon_{1\mu L1} = 98,4\% \pm 0,2\%$.

On calcule le nombre d'événements qui déclenchent le *trigger* monomuon ainsi que le nombre d'événements qui ont déclenché le *trigger* au niveau 1 sans tenir compte de la condition de coïncidence avec les luminomètres. Les résultats sont montrés sur la figure 5.6. Nous constatons que la condition de coïncidence a une efficacité d'environ $\epsilon_c = 96,2 \pm 0,3\%$. Pour la période où la condition de coïncidence était requise, l'efficacité de déclenchement du niveau 1 pour le *trigger* dimuon est donc :

$$\epsilon_{2\mu L1}^{av} = \epsilon_c \epsilon_{1\mu L1}^2 = 93,1 \pm 0,4\%.$$

On en déduit par la même occasion l'efficacité de déclenchement des muons pour la période après changement de *triggers* :

$$\epsilon_{2\mu L1}^{ap} = \epsilon_{1\mu L1}^2 = 96,8 \pm 0,3\%.$$

L'efficacité totale pour le L1 est donnée par la moyenne des efficacités avant et après changement, pondérée par les luminosités dans chaque période.

Ainsi, on trouve $\epsilon_{2\mu L1} = 94,4 \pm 0,5\%$.

6. Ces coupures sont appliquées dans toutes les études sur les efficacités du système de déclenchement.

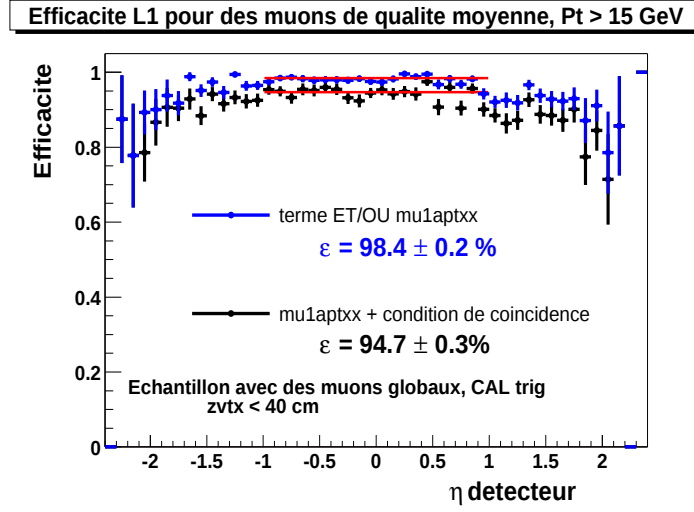


FIG. 5.6 – Efficacité du déclenchement au L1 sur des événements contenant un seul muon de qualité moyenne. Le trigger monomuon étudié est le “mu1aptxx”.

Méthode avec le *trigger* dimuon

Pour comparaison, on refait la même étude que pour le *trigger* monomuon mais cette fois avec le *trigger* dimuon qui a servi à sélectionner les événements de notre échantillon d’analyse. On demande alors deux muons passant les mêmes coupures de sélection que dans le cas précédent.

Pour la première période où le *trigger* requiert une condition de coïncidence avec les luminomètres, on trouve une efficacité $\epsilon_{2\mu L1}^{*av} = 89,1 \pm 2,1\%$. Cette valeur est assez faible par rapport à ce qui est trouvé dans le cas du trigger monomuon. Ceci est dû principalement à la faible statistique de cet échantillon, même si les marges d’erreur ne compensent pas totalement cet écart⁷. L’erreur systématique sur ce résultat peut être prise comme étant la moitié de l’écart entre les efficacités calculées avec les deux méthodes, ce qui correspond à $\pm 2\%$.

On calcule le nombre d’événements qui ont déclenché le *trigger* au niveau 1 sans tenir compte de la condition de coïncidence.

Pour la deuxième période, on trouve une efficacité de $\epsilon_{2\mu L1}^{*ap} = 93,5 \pm 1,6\%$. Ce résultat est encore une fois plus faible que $\epsilon_{2\mu L1}^{ap}$. La moitié de l’écart entre ces deux résultats peut être pris comme erreur systématique, ce qui correspond à $\pm 1,6\%$.

Nous prenons pour référence les résultats donnés par la méthode avec le *trigger* monomuon et les erreurs systématiques sont estimées à partir de l’écart entre les résultats des deux méthodes. En pondérant ces efficacités par les luminosités pour les périodes avant

7. Du fait de la faible statistique de l’échantillon, il n’est pas possible de réaliser une soustraction du bruit de fond. Il est donc probable que l’échantillon contienne un nombre important d’événements cosmiques.

et après le changement de *trigger*, on trouve :

$$\epsilon_{L1} = 94,4\% \pm 0,5\% (stat.) \pm 2\% (syst.) .$$

5.3.2 Déclenchement au niveau 2

Pour le calcul de l'efficacité de déclenchement du *trigger* dimuon au L2, nous procédons de la même façon que dans le cas du calcul de l'efficacité du déclenchement du L1 puisque le L2 a subi aussi des modifications qui lui ont permis de gagner en efficacité de déclenchement.

Nous voulons calculer l'efficacité de déclenchement du *trigger* dimuon au L2 sachant que l'on a dans nos événements deux muons de qualité moyenne et d'impulsion transverse supérieure à 15 GeV. Les événements sélectionnés ont déclenché le niveau L1 du *trigger* étudié.

D'une part nous calculons l'efficacité de déclenchement du *trigger* monomuon au L2 $\epsilon_{1\mu L2}$ (la condition est qu'un muon de qualité moyenne soit détecté au L2) lorsque dans les événements étudiés nous avons à chaque fois un et un seul muon de qualité moyenne.

Nous en déduisons l'efficacité de déclenchement du *trigger* dimuon au L2 $\epsilon_{2\mu L2}$ (la condition étant que –au moins– un muon de qualité moyenne soit détecté au L2) lorsque dans les événements étudiés se trouvent déjà deux muons de qualité moyenne.

Le lien entre $\epsilon_{1\mu L2}$ et $\epsilon_{2\mu L2}$ est donné par $\epsilon_{2\mu L2} = 1 - (1 - \epsilon_{1\mu L2})^2$.

Il est donc clair que dans notre cas, l'efficacité $\epsilon_{2\mu L2}$ que nous cherchons à évaluer sera particulièrement bonne étant donné que $\epsilon_{1\mu L2}$ est déjà élevée.

D'autre part, comme dans la section précédente, nous calculons aussi $\epsilon_{2\mu L2}$ directement en utilisant le *trigger* dimuon.

Méthode avec le *trigger* monomuon

On utilise le *trigger* monomuon MU_JT20_L2M0 décrit dans la section 4.1. Dans l'échantillon analysé on ne prend en compte que les événements dans lesquels un seul muon est de qualité moyenne et d'impulsion transverse supérieure à 15 GeV sachant que l'événement a déclenché le L1. La distribution de l'efficacité de déclenchement au L2 est montrée sur la figure 5.7 en fonction des variables η et ϕ du muon pour les périodes avant et après le changement de *trigger*.

A partir des résultats du *trigger* monomuon on déduit l'efficacité de déclenchement dans le cas où l'on a deux muons de qualité moyenne et d'impulsion transverse supérieure à 15 GeV.

Ainsi, on trouve :

$\epsilon_{1\mu L2}^{av} = 90,34 \pm 0,28\%$, ce qui donne un efficacité L2 calculée de $\epsilon_{2\mu L2}^{av} = 99,06 \pm 0,2\%$ pour la période avant le changement

et $\epsilon_{1\mu L2}^{ap} = 97,51 \pm 0,42\%$, ce qui donne un efficacité L2 calculée de $\epsilon_{2\mu L2}^{ap} = 99,93 \pm 0,3\%$

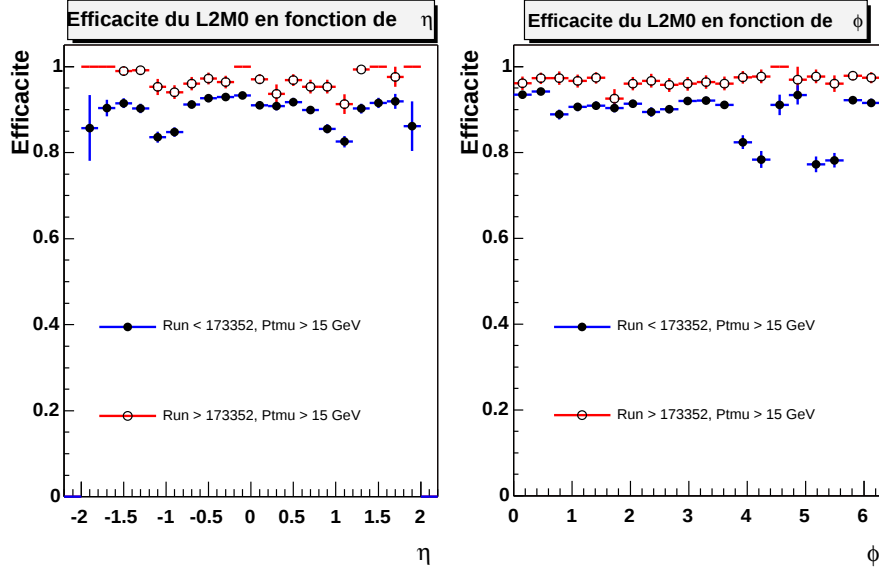


FIG. 5.7 – Efficacité du déclenchement au L2 sur des événements contenant un seul muon de qualité moyenne. Les irrégularités de l'efficacité en fonction de ϕ concernant la période avant le changement de trigger n'apparaissent plus pour la période après le changement où des corrections ont été apportées au système de déclenchement.

pour la période après le changement.

Méthode avec le *trigger* dimuon

En calculant avec le *trigger* dimuon, on trouve une efficacité de $99,50 \pm 0,5\%$ (sur les 204 événements sélectionnés un seul n'a pas passé la condition de *trigger* au L2) pour la période avant le changement et 100% d'efficacité pour la période après le changement, les 94 événements qui ont passé le *trigger* au L1 ayant aussi passé le *trigger* au L2. La moyenne de ces deux efficacités pondérées par la luminosité de chaque échantillon est alors de 99,8%. Les erreurs systématiques sont calculées comme dans le cas du L1. Ainsi on trouve une erreur systématique de 0,1%.

En pondérant ces efficacités par les luminosités pour la période avant le changement de *trigger* et la période après le changement, on trouve :
 $\epsilon_{L2} = 99,7\% \pm 0,3\% \text{ (stat.)} \pm 0,1\% \text{ (syst.)}$.

5.4 Appariement local-central

Nous avons sélectionné les événements afin de ne garder que les muons qui se trouvent dans l'acceptance géométrique du détecteur de traces central. Nous en avons tenu compte dans le chapitre précédent lors du calcul de l'efficacité due à l'acceptance géométrique des détecteurs. Tous les muons sont donc supposés laisser une trace dans le détecteur de traces central, c'est-à-dire que, dans le cas idéal, tous les muons que nous sélectionnons devraient être appariés à une trace centrale. Ce n'est pas le cas, et ce à cause de plusieurs

raisons :

- l'inefficacité du détecteur de traces central ;
- l'inefficacité de reconstruction d'une trace centrale ;
- l'inefficacité d'appariement d'une trace locale avec une trace centrale.

Ces trois points seront pris en compte dans une seule efficacité d'appariement d'un muon local avec une trace centrale. Pour ce calcul, nous reprenons la même idée que pour le cas du critère de qualité des muons, en sélectionnant un muon de contrôle et un muon test selon les critères suivants :

un muon de “contrôle” :

- un muon local de qualité moyenne ;
- apparié à une trace centrale ;
- il doit être isolé énergétiquement dans le calorimètre ;
- l'impulsion transverse de la trace centrale doit être supérieure à 15 GeV ;
- la distance de la trace centrale au point d'interaction doit être inférieure à 2 mm.

un muon “test” :

- un muon local de qualité moyenne ;
- l'impulsion transverse de la trace locale doit être supérieure à 15 GeV.

Les deux muons doivent être presque opposés en ϕ : on demande que la différence entre les angles ϕ des deux muons soit supérieure à 2,5 radians. Le muon test est considéré comme étant apparié à une trace centrale si cette dernière a les propriétés suivantes :

- elle est reconstruite avec plus de 3 coups dans le SMT et plus de 3 coups dans le CFT ;
- son impulsion transverse est supérieure à 15 GeV ;
- sa distance au point d'interaction est inférieure à 2 mm.

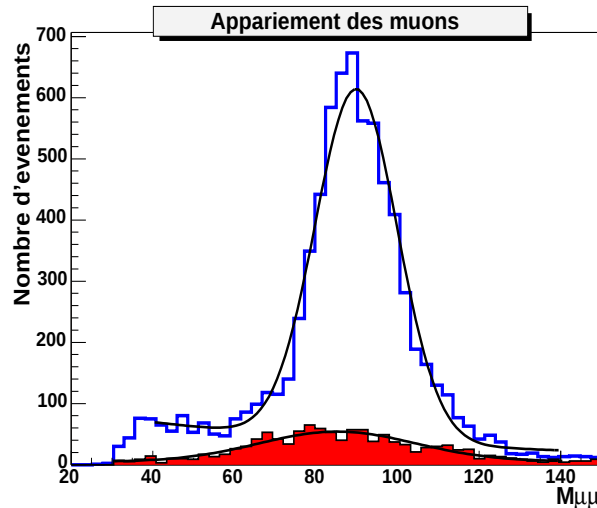


FIG. 5.8 – *Distribution de la masse invariante des paires de muons pour tous les événements (histogramme blanc) et pour les événements dans lesquels le muon test n'est pas apparié à une trace centrale (histogramme grisé) .*

De la même façon que dans la section 5.1, on étudie les distributions de la masse invariante des paires de muons pour tous les événements et pour les événements dans lesquels le muon test n'est pas apparié à une trace centrale (voir figure 5.8) et on calcule les efficacités de l'appariement avec et sans soustraction de bruit de fond pour un seul muon :

- sans soustraction du bruit de fond : $\epsilon_f = 84,7\% \pm 0,5\%$
- avec soustraction du bruit de fond : $\epsilon = 84,8\% \pm 0,5\%$

La distribution en η et en ϕ du muon test est montrée sur la figure 5.9. Le faible nombre d'événements à $4 < \phi < 5,5$ et à $1 < \phi < 2,5$ est une conséquence du trou dans la couche A du système à muons. En effet, dans ces régions soit on reconstruit moins de muons de contrôle (d'où le faible nombre de muons test en $1 < \phi < 2,5$ qui est la zone opposée au trou de la couche A) soit on reconstruit moins de muons test (d'où le faible nombre de muons test dans le trou à $4 < \phi < 5,5$).

Les efficacités en fonction de ces variables sont montrées sur la figure 5.10.

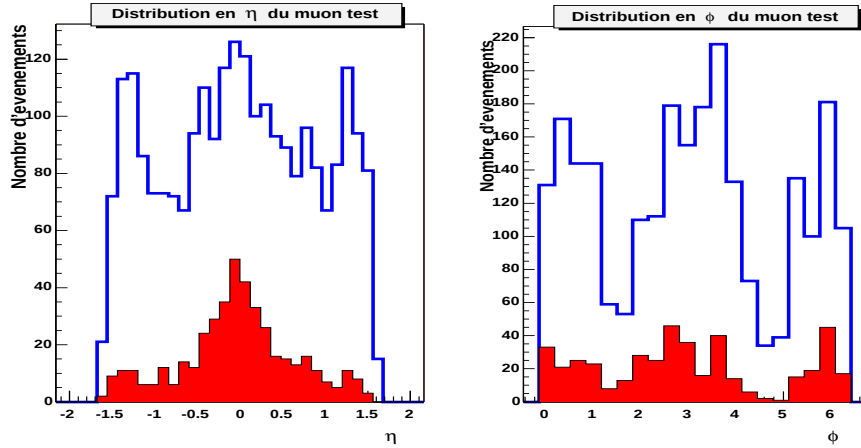


FIG. 5.9 – Distribution en η (à gauche) et en ϕ (à droite) du muon test pour tous les événements sélectionnés (histogramme blanc) et pour les événements dans lesquels le muon test n'est pas apparié à une trace centrale (histogramme grisé) .

Nous voyons que l'efficacité de l'appariement dépend très fortement de la distribution en η du muon test, mais reste presque plate en ϕ . Il faut en tenir compte dans le calcul de l'efficacité d'appariement.

5.5 Isolation des muons

L'isolation énergétique des muons est liée à la quantité d'énergie qui se trouve autour du muon. Un muon dans un jet apparaîtra dans le détecteur de traces central entouré de plusieurs traces dont l'énergie peut être très élevée, et dans le calorimètre l'énergie des cellules dans un cône autour du muon sera supérieure à l'énergie laissée par le muon seul.

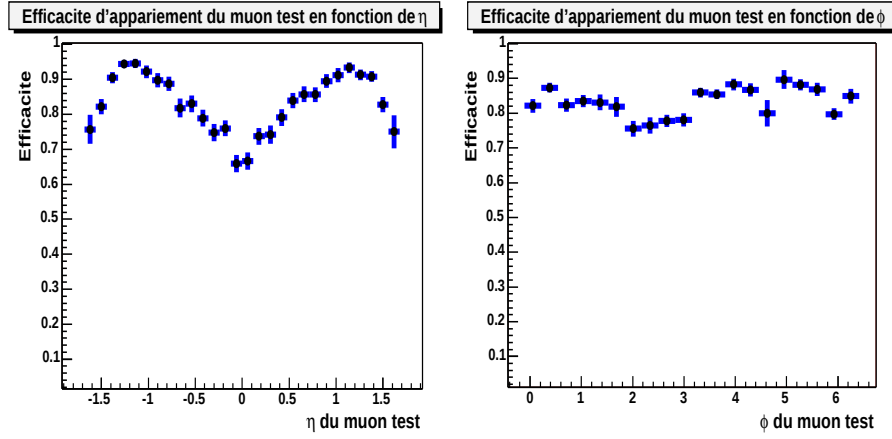


FIG. 5.10 – Efficacité d'appariement du muon test en fonction de η (à gauche) et de ϕ (à droite) du muon test.

Des muons dans un jet sont issus soit de la désintégration leptonique de quarks lourds (principalement des b ou c), ce qui devrait donner une paire de muons non isolés énergétiquement, soit d'un ISR (*Initial State Radiation*, c'est-à-dire la production d'un quark ou gluon associé à la production de Z) qui donne une gerbe hadronique qui peut se trouver près d'un des muons, ce qui donnerait un seul muon non isolé. Comme nous calculons la section efficace *inclusive* de production du Z , nous devons tenir compte de la possibilité d'avoir un jet de particules près d'un muon.

Le critère d'isolation consiste à demander que :

- la somme des énergies transverses⁸ de toutes les traces centrales se trouvant dans un cône d'ouverture $\Delta R^9 < 0,5$ soit inférieure à 2,5 GeV ;
- la somme des énergies transverses dans le calorimètre se trouvant dans un cône d'ouverture $0,1 < \Delta R < 0,4$ (que l'on appelle aussi "halo") soit inférieure à 2,5 GeV.

Nous voyons sur la figure 5.11 que dans la majorité des événements dimuons l'énergie transverse est inférieure à 2,5 GeV. Cette valeur est très discriminante par rapport aux événements contenant des jets.

Sur la figure 5.12 nous montrons la distribution de la masse invariante des dimuons dans le cas où les deux muons sont isolés (histogramme blanc) , dans le cas où un seul muon est isolé (histogramme grisé) et dans le cas où aucun muon n'est isolé (histogramme en ligne pointillée). Dans ce dernier cas les événements sont principalement dus au bruit de fond QCD dans lequel les muons sont issus de la désintégration de mésons et se trouvent donc dans des jets.

La figure 5.13-a montre la distribution en η du muon test. Le faible nombre d'événe-

8. Le plan transverse ici est le plan (xOy), c'est le plan transverse défini dans le repère du détecteur. Ce n'est donc pas le plan perpendiculaire à la trace du muon.

9. $\Delta R = \sqrt{(\Delta\phi)^2 + (\Delta\eta)^2}$ est la variable généralement utilisée pour décrire l'ouverture d'un cône ou la "distance angulaire" en fonction de ϕ et de η . $\Delta\phi = |\phi_1 - \phi_2|$ à 2π près et $\Delta\eta = |\eta_1 - \eta_2|$ où ϕ_1 (η_1) et ϕ_2 (η_2) sont les paramètres angulaires de la particule 1 (2).

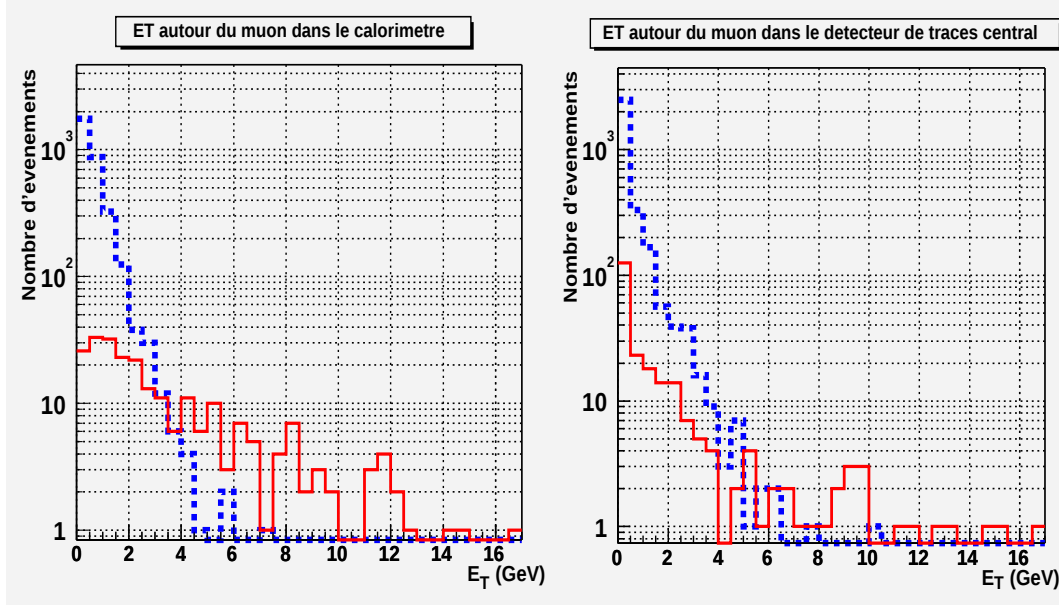


FIG. 5.11 – Distribution de l'énergie transverse autour des muons des événements sélectionnés sous le pic du Z dans le calorimètre (figure de gauche) et dans le détecteur de traces central (figure de droite). L'histogramme en pointillé correspond aux événements pour lesquels aucun jet n'a été reconstruit, l'histogramme en ligne continue correspond aux événements pour lesquels au moins un jet a été reconstruit près d'un des deux muons, c'est-à-dire à une distance en ϕ inférieure à 0.5 radians. Ces distributions ont été réalisées avec les données de notre analyse.

ments dans la région où $|\eta| < 1,1$ est un effet dû au trou dans la couche A du système à muons. Pour des valeurs de $|\eta| > 1,5$ les muons sont hors de l'acceptance géométrique du CFT, d'où le faible nombre d'événements dans ces régions.

Nous ne pouvons observer ces effets sur la distribution correspondant au cas où le muon test n'est pas isolé (histogramme grisé) du fait de la faible statistique de cette distribution.

De même nous constatons sur la figure 5.13-b que le nombre d'événements est faible dans les régions où $4 < \phi < 5,5$ à cause du trou dans la couche A ainsi que dans les régions où $1 < \phi < 2,5$ à cause de la corrélation en ϕ entre les deux muons.

La figure 5.14 montre que l'efficacité d'isolation des muons est très stable en fonction des variables angulaires η et ϕ du muon test. Nous calculons alors une efficacité indépendamment de la position angulaire des muons.

Pour calculer l'efficacité du critère d'isolation des muons, nous utilisons un muon de contrôle et un muon test qui répondent à tous les critères de sélection des muons sauf le critère d'isolation. Le muon de contrôle doit être isolé énergétiquement dans le calorimètre et dans le détecteur de traces central. Nous comptons le nombre d'événements pour lesquels le muon test est isolé énergétiquement ainsi que le nombre total d'événements. Nous calculons l'efficacité avec et sans soustraction de bruit de fond, comme cela est expliqué dans la section 5.1. Nous trouvons :

- sans soustraction du bruit de fond : $\epsilon_f = 91,4\% \pm 0,42\%$;
- avec soustraction du bruit de fond : $\epsilon = 92,3\% \pm 0,39\%$.

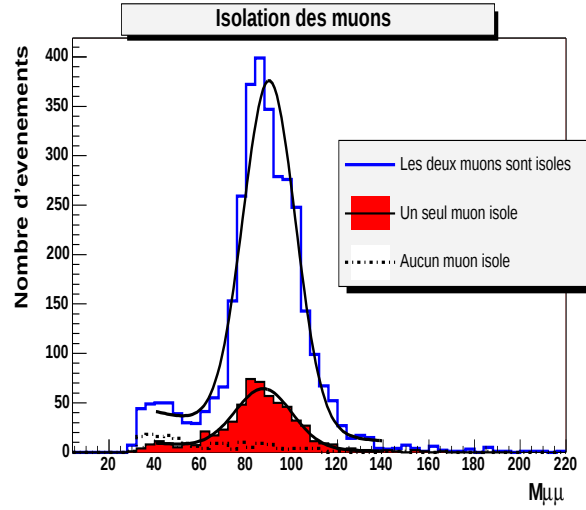


FIG. 5.12 – Distribution de la masse invariante des dimuons dans le cas où les deux muons sont isolés (histogramme gras), dans le cas où un seul muon est isolé (histogramme grisé) et dans le cas où aucun muon n'est isolé (histogramme en ligne pointillée) .

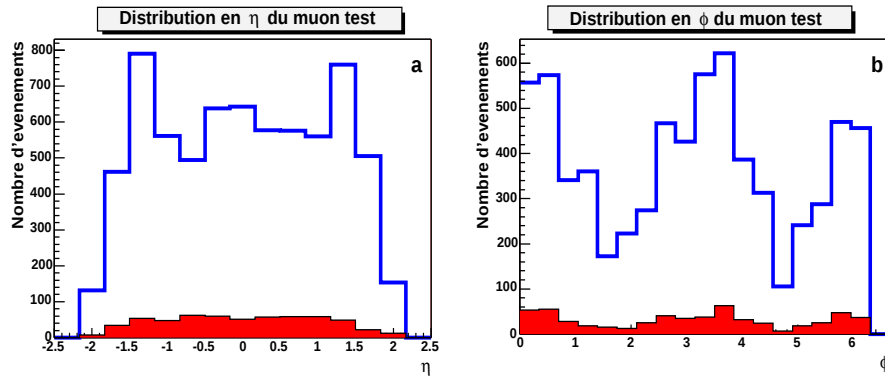


FIG. 5.13 – Distribution des variables angulaires η (à gauche) et ϕ (à droite) du muon test pour tous les événements (histogramme blanc) et pour les événements dans lesquels le muon test est isolé (histogramme grisé) .

Ces efficacités concernent la sélection d'un *seul muon isolé*.

De même, le critère d'isolation ne semble pas dépendre de manière sensible de la valeur de l'impulsion transverse du muon test, bien que l'efficacité augmente très légèrement avec l'impulsion transverse, comme le montre la figure 5.16.

La figure 5.15 montre que les distributions de masse invariante des dimuons dans le cas où les deux muons sont isolés et dans le cas où un seul muon est isolé sont similaires. Au lieu de se restreindre au cas où les deux muons sont isolés, nous pouvons gagner en statistique en incluant aussi dans notre sélection les événements dans lesquels un seul muon est isolé.

Ainsi, contrairement aux autres critères de sélection, au lieu de demander que les deux

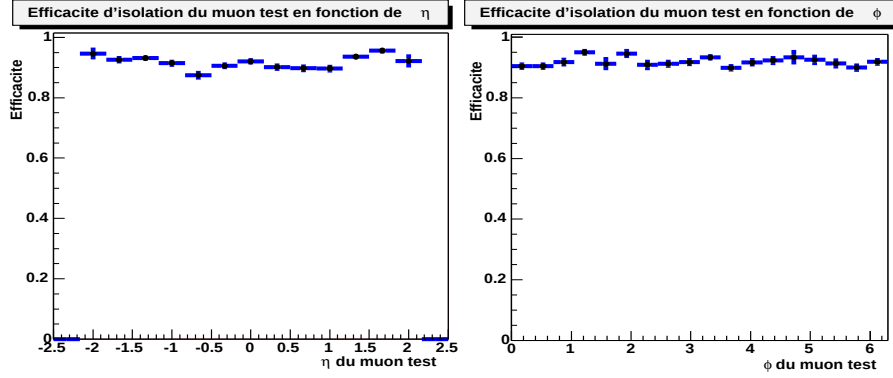


FIG. 5.14 – Efficacité d'isolation du muon test en fonction de η (à gauche) et de ϕ (à droite) .

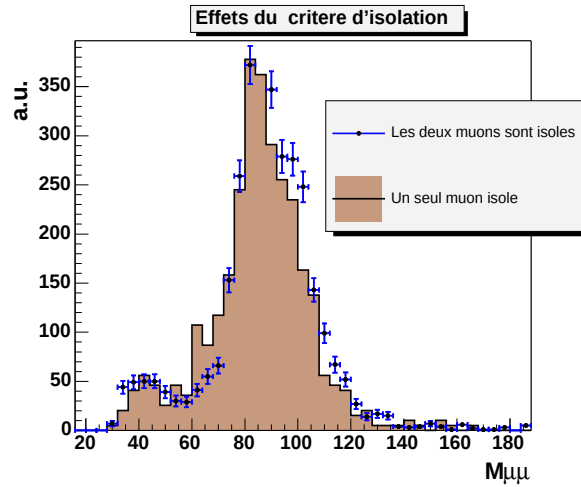


FIG. 5.15 – Distribution de la masse invariante des dimuons dans le cas où les deux muons sont isolés (points), dans le cas où un seul muon est isolé (histogramme grisé) . Les deux histogrammes sont normalisés au même nombre d'événements.

muons passent le critère d'isolation, nous exigeons qu'*au moins un* des deux muons passe ce critère.

Dans ce cas, l'efficacité de sélection des événements dans lesquels *au moins un* des deux muons est isolé ϵ_{isol} est reliée à l'efficacité de sélection des événements dans lesquels *un seul muon* est isolé ϵ calculée précédemment par $\epsilon_{isol} = 1 - (1 - \epsilon)^2$, ce qui donne $\epsilon_{isol} = 99,4\% \pm 0,5\%$.

5.6 Sélection du lot de données analysées

Nous avons sélectionné les données acquises durant la période couvrant de novembre 2002 à juillet 2003. Ces données ont été triées parmi toutes les données enregistrées durant cette période selon un seul critère à savoir que tous les événements contiennent deux muons de qualité moyenne.

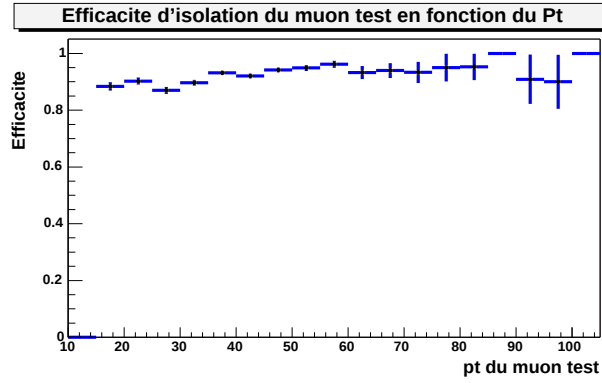


FIG. 5.16 – Efficacité d'isolation du muon test en fonction de son impulsion transverse.

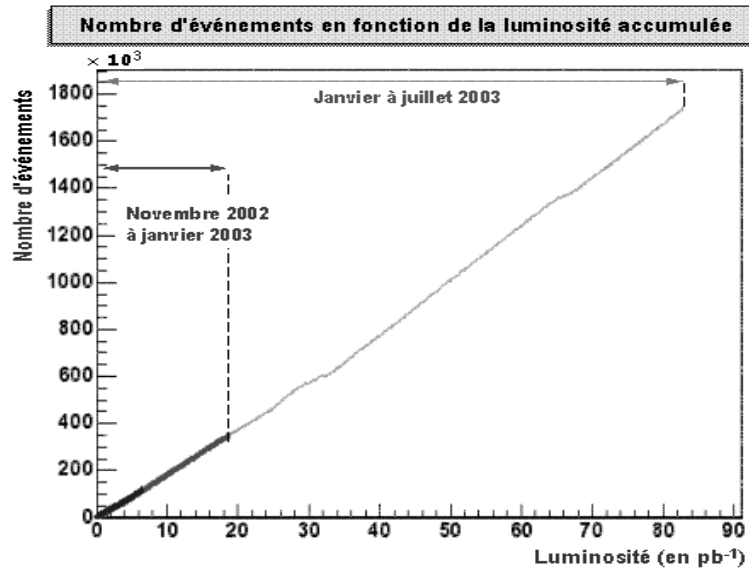


FIG. 5.17 – Nombre d'événements, ayant déclenché le trigger dimuon de notre analyse, intégré dans l'échantillon de présélection en fonction de la luminosité intégrée.

La figure 5.17 montre le nombre d'événements présents dans l'échantillon et ayant déclenché le *trigger* dimuon 2_MU_A_L2M0 en fonction de la luminosité intégrée liée à ce *trigger*. Cette distribution doit être linéaire car pour une luminosité donnée, le nombre d'événements doit être constant. Mais la pente de la courbe grise (correspondant à la période de janvier à juillet) change à partir d'une luminosité intégrée d'environ 36 pb^{-1} , ce qui correspond à la période de prise de données coïncidant avec le changement de *trigger*.

Nous remarquons plus clairement ce décrochement sur la figure 5.18. La durée d'acquisition varie d'un lot à l'autre (mais ne dépasse pas 4 heures), c'est pourquoi les erreurs statistiques sont variables. Nous distinguons sur cette figure deux périodes qui présentent un nombre moyen d'événements différent. Cette différence correspond aux deux périodes utilisant une condition de *trigger* différente.

Après avoir enlevé les lots dans lesquels des problèmes ont eu lieu durant l'acquisition et des lots présentant un taux anormalement élevé ou bas (voir figure 5.18), on calcule la luminosité totale des données restantes. Pour toute la période considérée, l'échantillon

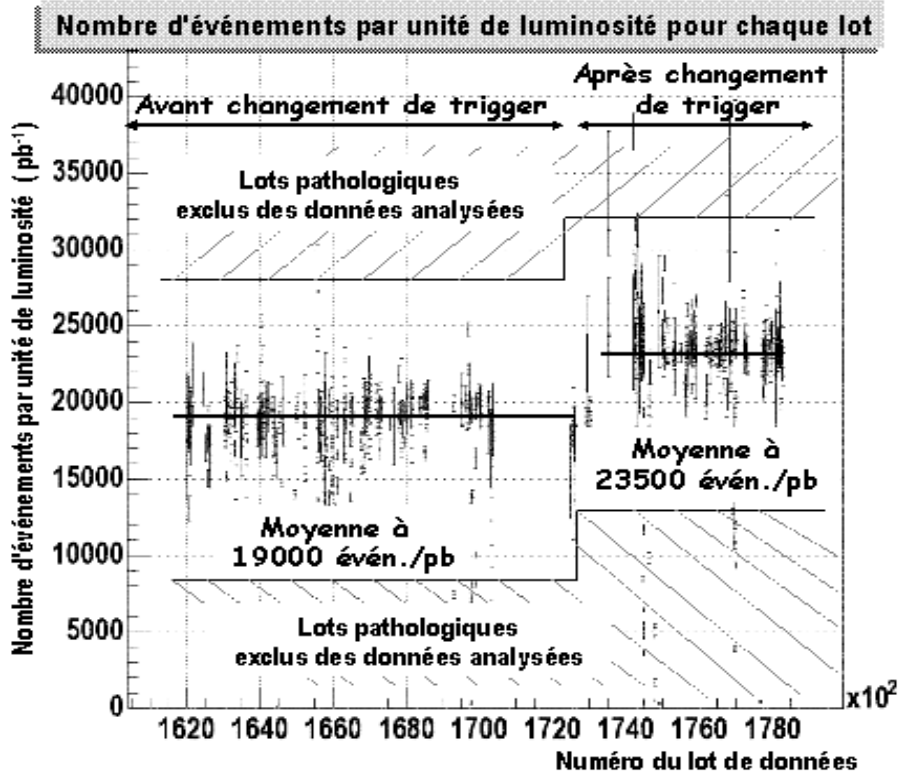


FIG. 5.18 – Stabilité de la luminosité en fonction du numéro de lot de données. En ordonnée est reporté le nombre d'événements par unité de luminosité intégrée prise en picobarns pour chaque lot de données (run). Les événements ont déclenché le trigger dimuon au L1 et au L2.

analysé représente une luminosité intégrée de $107,8 \text{ pb}^{-1}$. L'incertitude liée à la luminosité est une incertitude systématique de 10% liée au système de comptage de la luminosité.

5.7 Contamination des données en bruit de fond

Dans cette section nous allons estimer les contributions au bruit de fond muonique présent dans l'échantillon final sélectionné. Nous nous attendons à ce que les bruits de fond dominants soient le Drell-Yan non résonnant et les événements $b\bar{b}$. Dans une proportion beaucoup plus faible les événements $Z, \gamma \rightarrow \tau\tau$ peuvent ajouter quelques événements. Nous verrons que les bruits de fond tels que la production associée de bosons de jauge (ZZ , ZW et WW) sont négligeables.

5.7.1 Contamination en événements $b\bar{b}$

Dans les événements $b\bar{b}$, les muons sont généralement dans des jets et apparaissent donc en général comme non isolés énergétiquement.

A haute énergie, les jets de b sont très collimés, ce qui laisse supposer que des muons issus de jets de b de haute énergie apparaîtront plus souvent non isolés que dans le cas de muons issus de jets de b de basse énergie.

Nous avons sélectionné des paires de muons d'impulsion transverse supérieure à 15 GeV dans lesquels au moins un muon est isolé. Dans ce cas, on peut s'attendre à ce que la proportion d'événements $b\bar{b}$ dans nos événements soit faible. Si l'on note ϵ_b la probabilité pour qu'un muon issu d'un jet de b apparaisse comme étant isolé, l'efficacité pour que l'on retrouve des événements $b\bar{b}$ avec un seul muon isolé est $2\epsilon_b(1 - \epsilon_b)$ et l'efficacité pour sélectionner des événements $b\bar{b}$ avec deux muons isolés est ϵ_b^2 .

Pour estimer le nombre d'événements $b\bar{b}$ dans notre échantillon, nous relevons dans les données les nombres d'événements suivants :

- $N_{tot} = N_{Z/\gamma/\tau\tau} + N_b$, N_{tot} étant le nombre total de paires de muons qui satisfont à tous les critères de sélection mis à part le critère d'isolation;
- $N_0 = (1 - \epsilon_{isol})^2 N_{Z/\gamma/\tau\tau} + (1 - \epsilon_b)^2 N_b$, où N_0 est le nombre de paires de muons dans lesquelles aucun des deux muons n'est isolé;
- $N_1 = (1 - \epsilon_{isol})\epsilon_{isol} N_{Z/\gamma/\tau\tau} + (1 - \epsilon_b)\epsilon_b N_b$. C'est le nombre total de paires de muons dans lesquelles un seul muon est isolé;
- $N_2 = \epsilon_{isol}^2 N_{Z/\gamma/\tau\tau} + \epsilon_b^2 N_b$. C'est le nombre total de paires de muons dans lesquelles les deux muons sont isolés.

Nous avons calculé dans la section 5.5 l'efficacité ϵ_{isol} qu'un muon issu de la désintégration d'un $Z/\gamma/\tau\tau$ apparaisse comme étant isolé, avec $\epsilon_{isol} = 92,3\%$.

Nous ne connaissons *a priori* ni $N_{Z/\gamma/\tau\tau}$, ni N_b , ni ϵ_b . Nous pouvons calculer ces trois paramètres inconnus à partir des trois équations indépendantes données par N_{tot} , N_0 et N_2 .

On trouve alors une valeur de l'efficacité $\epsilon_b = 13,2\%$ et un nombre total d'événements $b\bar{b}$ dans tout notre échantillon de 247 événements. Le nombre d'événements $b\bar{b}$ se trouvant alors dans les données que nous avons sélectionnées sur le critère d'isolation (au moins un muon isolé) contient donc : $N_{b\bar{b}} = (1 - (1 - \epsilon_b)^2)N_b = 61$ événements sur les 4040 événements sélectionnés après application de toutes les coupures, soit une contamination f_b de notre échantillon final en événements $b\bar{b}$ de $f_b = N_{b\bar{b}}/N_2 + N_1 = 1,5\% \pm 0,2\%$.

5.7.2 Contamination en événements $\tau^+\tau^-$

Pour la mesure de la contamination en événements $\tau^+\tau^-$ dans les données finales on dispose de deux échantillons MC de $Z, \gamma \rightarrow \mu\mu$ et de $Z, \gamma \rightarrow \tau\tau$ reconstruits dans le détecteur.

Nous appliquons toutes les coupures de sélection des événements sur ces deux échantillons de MC. La distribution de la masse invariante des dimuons dans les événements restants est montrée sur la figure 5.19 .

Nous constatons que la proportion d'événements $\tau^+\tau^-$ est naturellement très faible et localisée principalement à basse masse invariante. La fraction d'événements $\tau^+\tau^-$ dans les événements est estimée à $f_\tau = N_{\tau\tau}/N_{\mu\mu} = 0,17\% \pm 0,05\%$.

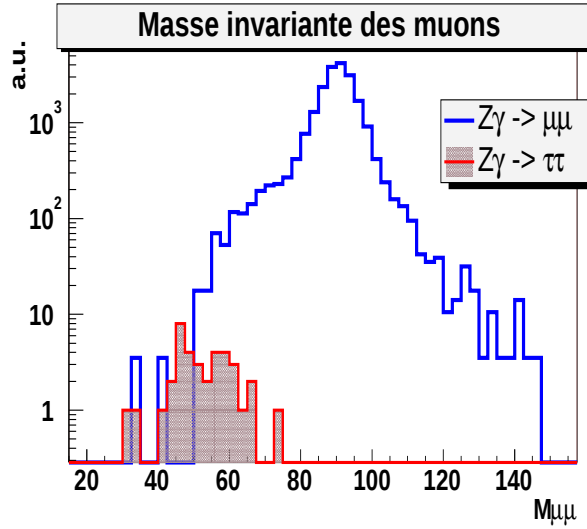


FIG. 5.19 – Distribution de la masse invariante des dimuons issu du processus $q\bar{q} \rightarrow Z, \gamma \rightarrow \mu\mu$ (histogramme blanc) et du processus $q\bar{q} \rightarrow Z, \gamma \rightarrow \tau\tau$ (histogramme en grisé) pour des muons passant toutes les coupures de sélection. Ces distributions ont été réalisées avec des échantillons MC normalisés à la même luminosité.

5.7.3 Proportion d'événements Z purs et Drell-Yan

Pour ce calcul on se base sur des événements MC générés par PYTHIA. On produit séparément des événements Z purs et des événements Z/ γ (où les contributions dues au Z et au photon sont considérées). Pour une même luminosité, on compte le nombre d'événements Z purs et Z/ γ dans lesquels les muons ont une impulsion transverse générée supérieure à 15 GeV (ce qui correspond à peu près à la coupure effectuée sur l'impulsion reconstruite des muons). Nous montrons sur la figure 5.20 la distribution de la masse invariante (générée) des dimuons pour les deux cas. Nous pouvons en déduire la proportion d'événements Z en dimuons qui restent dans les événements finals des données sélectionnées.

A l'ordre des arbres et pour des muons d'impulsion transverse minimale de 15 GeV, PYTHIA donne une valeur de la section efficace du processus non physique Z seul de 166 pb et une section efficace du processus physique Z/ γ de 195 pb. Nous prendrons donc $f_{DY}=14,8\%\pm 0,5\%$ des événements sélectionnés sont du Drell-Yan¹⁰. L'erreur est déduite des incertitudes des intégrations numériques des sections efficaces par PYTHIA.

5.7.4 Contamination en événements ZZ, ZW et WW

D'autres bruits de fond tels que la production associée de deux bosons de jauge ZZ, ZW et WW ont une section efficace totale de production respectivement de 1 pb, 2,5 pb et 8,5 pb pour tous les canaux de désintégration. Tenant compte du rapport d'embranchement de 3,37% du Z en muons et de 10,57% du W en muons, la contribution de ces bruits de fond n'excède pas 0,067 pb, 0,084 pb et 0,095 pb respectivement. Les bruits de fond ZZ

10. Dans cette estimation nous comptons à la fois les événements $q\bar{q} \rightarrow \gamma \rightarrow \mu\mu$ et la contribution du terme d'interférence entre les processus faisant intervenir un photon et ceux faisant intervenir un Z.

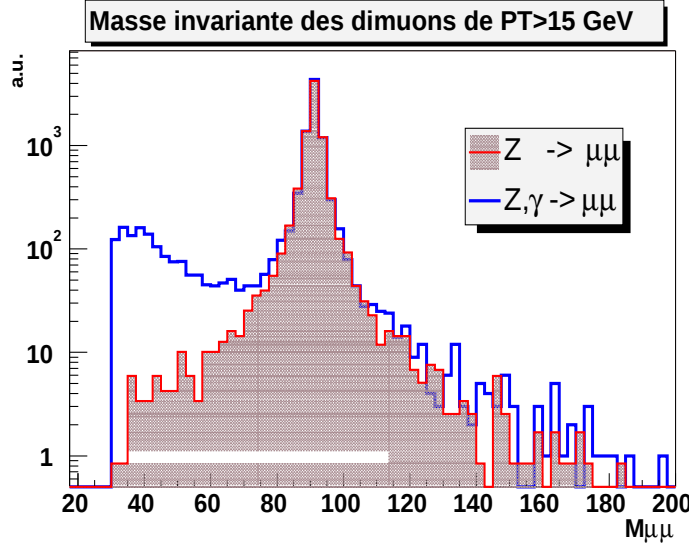


FIG. 5.20 – Distribution de la masse invariante des dimuons issus du processus $q\bar{q} \rightarrow Z \rightarrow \mu\mu$ (histogramme grisé) et du processus $q\bar{q} \rightarrow Z, \gamma \rightarrow \mu\mu$ (histogramme blanc) pour des muons d'impulsion transverse supérieure à 15 GeV. Les variables (masse et impulsions transverses) sont les valeurs générées. Ces histogrammes ne tiennent pas compte de l'effet de la reconstruction des événements dans le détecteur.

et ZW suivent la même sélection que celle de notre signal. Avec une efficacité de sélection de 13,5% (voir tableau des résultats finals 5.1), ces deux bruits de fond contribuent au maximum pour 2,6 événements.

Dans le cas du bruit de fond WW, les deux muons issus de la désintégration des bosons W ne sont pas forcément opposés en ϕ . En effet, seulement 20% des paires de muons satisfont la condition $\Delta\phi > 2,5$. On s'attend donc à avoir tout au plus 0,3 événement dans les données dû au processus de production de paires de W, ce qui est complètement négligeable.

Ainsi, moins de 3 événements dans les données sélectionnées pourraient provenir de la production de paires de bosons ZZ, ZW et WW.

5.8 Section efficace de production de Z

Après toutes les coupures de sélection de notre signal, il reste $N_{obs}=4040$ événements dont une partie f_b provient des événements $b\bar{b}$ et une partie f_τ provient des événements $\tau^+\tau^-$; de ce qui reste une partie f_{DY} provient d'événements de Drell-Yan.

Ainsi, N_{obs} peut s'écrire :

$$N_{obs} = N_Z + N_{DY} + N_{\tau\tau} + N_{b\bar{b}} \quad (5.1)$$

avec

- $N_{b\bar{b}} = f_b N_{obs}$;
- $N_{\tau\tau} = f_\tau (N_Z + N_{DY})$;
- et $N_{DY} = f_{DY} N_Z$.

On réécrit alors :

$$N_{obs} (1 - f_b) = (N_Z + N_{DY}) \times (1 + f_\tau) = N_Z \times (1 + f_{DY}) \times (1 + f_\tau) \quad (5.2)$$

Ainsi, la section efficace de production du Z dans le canal muonique peut s'écrire sous la forme :

$$\sigma = \frac{N_Z}{\epsilon \times \mathcal{L}} \quad (5.3)$$

C'est-à-dire :

$$\sigma = \frac{N_{obs} \times (1 - f_b)}{\epsilon \times \mathcal{L} \times (1 + f_{DY}) \times (1 + f_\tau)} \quad (5.4)$$

Le tableau 5.1 résume les efficacités sur les coupures de sélection ainsi que les contaminations en bruits de fond.

Coupures de sélection	Efficacité (%)	Erreur (%)
Acceptance géométrique	34,2	0,4
Qualité moyenne	58,9	1,1
$P_T > 15$ GeV	97,5	1,0
<i>Trigger</i> L1	94,4	2
<i>Trigger</i> L2	99,7	0,7
Appariement	71,9	0,7
Isolation	99,4	0,5
Efficacité totale	13,5	0,4
Événements	Taux de contamination (%)	Erreur (%)
$b\bar{b}$	1,5	0,2
$\tau^+\tau^-$	0,17	0,05
DY	14,87	1
Nombre total d'événements	4040	$\pm 1,60\%$
Luminosité (pb^{-1})	107,8	$\pm 10\%$
Section Efficace (pb)	$237,6 \pm 3,8$ (stat) $\pm 6,0$ (syst) $\pm 23,8$ (lumi)	

TAB. 5.1 – *Résumé des coupures de sélection et contamination des différents bruits de fond des événements sélectionnés. La section efficace de production de Z dans le canal muonique est reportée en dernière ligne du tableau.*

La section efficace de production de boson Z dans le canal muonique est donc de 237,6 pb. Cette valeur est tout à fait conforme à ce que l'on attend, sachant que PYTHIA prédit une valeur de 240 pb (initialement 184 pb multipliée par un facteur K de 1,3 pour tenir compte des contributions des diagrammes aux ordres supérieurs) .

Nous pouvons constater que l'erreur statistique est raisonnablement faible (1,6%) . Par contre, l'erreur systématique est beaucoup plus grande que l'erreur statistique. Cette valeur importante est principalement due à la faible statistique dont on dispose pour calculer les efficacités du système de déclenchement, ce qui introduit des erreurs importantes

dans les calculs, ainsi qu'au manque d'information pertinente concernant le déclenchement sur les données (lors du calcul de l'efficacité de sélection de muons de qualité moyenne) .

Les erreurs systématiques peuvent être améliorées avec un lot de données plus important ainsi qu'avec une sélection des événements avec d'autres *triggers* qui permettraient d'avoir un plus grand nombre d'événements pour le calcul de l'efficacité de déclenchement.

Nous pouvons par ailleurs améliorer la statistique des données analysées en choisissant un échantillon qui n'a pas été sélectionné sur le critère de qualité des muons, ce qui permettrait de sélectionner des muons de qualité acceptable au lieu de muons de qualité moyenne. Ce choix améliorerait l'efficacité de sélection d'environ 50%.

Finalement, l'erreur due à la luminosité est une erreur incontournable pour le moment puisque l'étalonnage des luminomètres n'a pas été effectué. Les luminomètres sont les mêmes que ceux du Run I, mais l'électronique a changé. On s'attend naturellement à ce que leur efficacité de détection ainsi que leur sensibilité soient différentes de celles du Run I. De ce fait, une erreur de 10% sur la mesure de la luminosité est estimée par le groupe de travail concerné.

Bibliographie

- [1] R. Hamberg, W.L. Neerven, T. Matsuura, *A Complete Calculation of The Order α_s^2 Correction to The Drell-Yan K-Factor*,
Nucl. Phys. **B359** (1991) 343-405.
- [2] D. Whiteson, M. Kado, *Muon Isolation Studies*,
DØ note 4070 (2003).
- [3] E. Nurse, P. Telford, *Measurement of $\sigma.Br$ for $Z \rightarrow \mu\mu$ in $p\bar{p}$ collisions at $\sqrt{s}=1.96$ TeV*,
DØ note 4231 (2003).

Chapitre 6

Contraintes sur le graviton de Kaluza-Klein

Introduction

Nous avons sélectionné les événements qui ont passé toutes les coupures détaillées dans le chapitre précédent. Nous allons maintenant étudier la compatibilité du spectre de masse des muons avec la présence d'un signal de graviton. Pour cela nous utilisons la méthode du rapport de vraisemblance développée par T. Junk [3] et largement utilisée au LEP. Nous en déduisons ensuite la courbe d'exclusion du signal à 95% de niveau de confiance relative à la section efficace de production du graviton en fonction de sa masse. Nous en déduisons la zone d'exclusion à 95% de niveau de confiance dans le plan des paramètres du modèle, c'est-à-dire dans le plan $(M_{grav} , k/M_{Pl})$.

6.1 Méthode du rapport de vraisemblance

Dans nos données (figure 6.1) , nous observons un léger excès d'événements par rapport aux prédictions du modèle standard. Nous voulons tester la compatibilité de cet excès avec un hypothèse de bruit de fond seul (B) ou de bruit de fond avec du signal de graviton (B+S). Le nombre d'événements est relativement faible dans les queues des distributions de masse invariante. Ces événements suivent une statistique poissonnienne.

Notre but est d'estimer le nombre limite s_{lim} d'événements de signal de graviton qui reste compatible avec les observations des données sachant que l'on observe d événements de données et qu'on attend b événements de bruit de fond. Ce nombre s_{lim} doit correspondre à la limite à 95% de niveau de confiance de validité de l'hypothèse (S+B), c'est-à-dire que l'on cherche s_{lim} tel que la probabilité $P(d | (b + s_{lim}))$ pour que l'excès observé soit dû à du signal doit être inférieure à 5%.

De ce nombre s_{lim} nous déduisons la section efficace limite σ_{lim} de production du graviton donnée par :

$$\sigma_{lim} = \frac{s_{lim}}{\epsilon \mathcal{L}} . \quad (6.1)$$

Plusieurs méthodes permettent d'évaluer cette limite. Une méthode simple consiste à prendre pour niveau de confiance $CL = CL_{s+b}$. Dans le cas de distributions qui suivent une loi de probabilité de Poisson et pour n canaux différents,

$$CL_{s+b} = P(d < d_{obs} | (b + s)) = \sum_{d=0}^{d_{obs}} \prod_{i=0}^n \frac{(s_i + b_i)^{d_i} e^{-(s_i + b_i)}}{d_i !} . \quad (6.2)$$

CL_{s+b} mesure le niveau de confiance pour que les d_{obs} événements des données soient compatibles avec un nombre $b = \sum_{i=0}^n b_i$ de bruit de fond attendu et un nombre $s = \sum_{i=0}^n s_i$ de signal.

Si l'on se place dans le cas où aucun événement de données n'est observé, le niveau de confiance du signal est donné par $CL = e^{-(b+s)}$. Si $b \leq 3$, alors $CL = CL_{s+b} \leq 0,05$ quelle que soit la valeur de s , ce qui veut dire que l'on exclut le signal à 95% de niveau de confiance quelle que soit sa section efficace. Ce résultat est évidemment incorrect, ce

qui restreint l'utilisation de la méthode aux cas où le nombre d'événements observés d est supérieur au nombre d'événements attendus b .

Par ailleurs, observer un petit CL_{s+b} ne prouve pas que le signal est absent car le fond lui-même peut avoir fluctué vers le bas, auquel cas CL_b est petit. C'est pour gérer de tels cas que nous décrivons CL_s en renormalisant CL_{s+b} , ce qui se traduit par la prescription $CL_s = CL_{s+b}/CL_b$.

Afin de garder une évaluation du niveau de confiance acceptable dans tous les cas de figure, le théorème de Neyman-Pearson¹ affirme que la méthode du rapport de vraisemblance [1] est la mieux adaptée pour estimer les niveaux de confiance des hypothèses $H_0 = B$ supposant l'absence de signal et $H_1 = S + B$ supposant la présence de signal dans les données observées.

Notons CL_{s+b} le niveau de confiance pour l'hypothèse S+B donné par l'équation (6.2). Pour un nombre d'événements n observé, le niveau de confiance CL_b pour le bruit seul est donné par :

$$CL_b = P(d < d_{obs} \mid b) = \sum_{d=0}^{d_{obs}} \prod_{i=0}^n \frac{b_i^{d_i} e^{-b_i}}{d_i!} . \quad (6.3)$$

Le niveau de confiance du signal CL_s est alors donné par le rapport :

$$CL_s = \frac{CL_{s+b}}{CL_b} \quad (6.4)$$

Cette méthode présente l'avantage d'être plus sensible dans les canaux à faible nombre d'événements et permet de combiner simplement les canaux entres eux.

Par convention, on parle de découverte (à 5σ) si $CL_b < 5,7 \cdot 10^{-7}$. De même, par convention, on rejette l'hypothèse d'un signal si $CL_{s+b}/CL_b < 5 \cdot 10^{-2}$. C'est la limite à 95% de niveau de confiance.

Notons que, à statistique de test donnée et analyses identiques, le résultat ne peut s'améliorer que si l'on augmente la luminosité, c'est-à-dire b et $s+b$. Cela vient du fait que l'écart standard d'une loi de Poisson de paramètre N est $\sqrt{N}$. Ainsi, les événements se "resserrent" autour de la valeur centrale à mesure que l'on augmente N .

Afin de tenir compte des erreurs systématiques sur le bruit de fond σ_b et sur le signal σ_s , la méthode utilisée se base sur la prescription de Cousins et Highland [2].

6.2 Calcul des limites à 95% de niveau de confiance

Cette méthode, basée sur la méthode du rapport de vraisemblance, a été développée par T. Junk. Contrairement à la méthode du simple comptage, elle présente l'avantage de

1. Le théorème stipule que dans le cas d'un test d'une hypothèse simple contre une autre hypothèse simple complètement spécifiées, si les densités de probabilité sous les deux hypothèses ont la même forme analytique (toutes deux des lois de Poisson par exemple), alors il existe un test plus puissant que tous les autres, à savoir le rapport de vraisemblance des deux hypothèses.

tenir compte de tout le spectre de masse, ce qui lui permet de tenir compte d'un maximum d'informations fournies par les spectres de masse des données, du signal et du bruit de fond.

L'idée est de partir des distributions des signaux de données (D), de bruit de fond (B) et de la somme du bruit de fond et du signal (B+S), où la normalisation de la distribution du signal S correspond à une section efficace σ_0 de production du graviton. On calcule ensuite le niveau de confiance CL_b de la distribution du bruit de fond et le niveau de confiance CL_{s+b} de la distribution de la somme du bruit de fond et du signal. Chaque intervalle des histogrammes de spectres de masses est traité séparément des autres intervalles, le contenu de chacun est sujet à des fluctuations statistiques qui suivent une loi poissonnienne.

Le niveau de confiance du signal est défini par $CL_s = CL_{s+b}/CL_b$. Ces valeurs reflètent la compatibilité des données avec la distribution considérée. Ainsi, si CL_s est inférieur à 5%, les données excluent l'hypothèse S+B à 95% de niveau de confiance et donc la section efficace du signal est inférieure à σ_0 .

Nous utilisons les outils fournis par le logiciel ROOT² pour calculer les limites à 95% de niveau de confiance [4]. Pour cela nous disposons d'histogrammes des données, du signal et du bruit de fond. Deux histogrammes additionnels, contenant les erreurs systématiques pour le signal et pour le bruit de fond, sont nécessaires à la prise en compte des erreurs systématiques lors du calcul des limites. Ces erreurs sont supposées indépendantes de la masse des dimuons ainsi que de la masse du graviton pour le signal considéré.

Sur la figure 6.1 sont reportés les spectres de masse des dimuons pour les données (croix), le bruit de fond du modèle standard (histogramme en ligne continue), le signal de graviton (tirets épais) et la somme de ces deux derniers histogrammes (pointillés).

La distribution observée présente un accord raisonnable avec celle prédite par le modèle standard. Pour des masses supérieures à 150 GeV, 49 événements sont observés dans les données et $33 \pm 3,5$ sont attendus dans les processus standards. Avec une erreur systématique de 10%, la probabilité que le fond attendu fluctue à 49 événements ou plus est de l'ordre de 1%, ce qui correspond à une déviation de l'ordre de $1,8 \sigma$. L'événement de plus grande masse observé dans les données est mesuré à une masse de 380 GeV.

Pour une masse de graviton donnée, la largeur de la résonance dépend du couplage, c'est-à-dire de k/M_{Pl} , mais n'excède pas quelques dizaines de GeV même pour les résonances les plus larges. La figure 6.2 montre les valeurs de la largeur de la résonance d'un graviton pour différentes valeurs de masse et de k/M_{Pl} . La forme du signal de graviton reconstruit dans le détecteur dépend entièrement de la résolution sur l'impulsion des muons. Ainsi, pour une masse de graviton donnée, la forme de la résonance est la même pour toutes les valeurs de couplage permises par le modèle (c'est-à-dire k/M_{Pl} compris entre 0,01 et 0,1).

2. Les classes auxquelles on fait appel sont TConfidenceLevel, TLimit, TLimitDataSource.

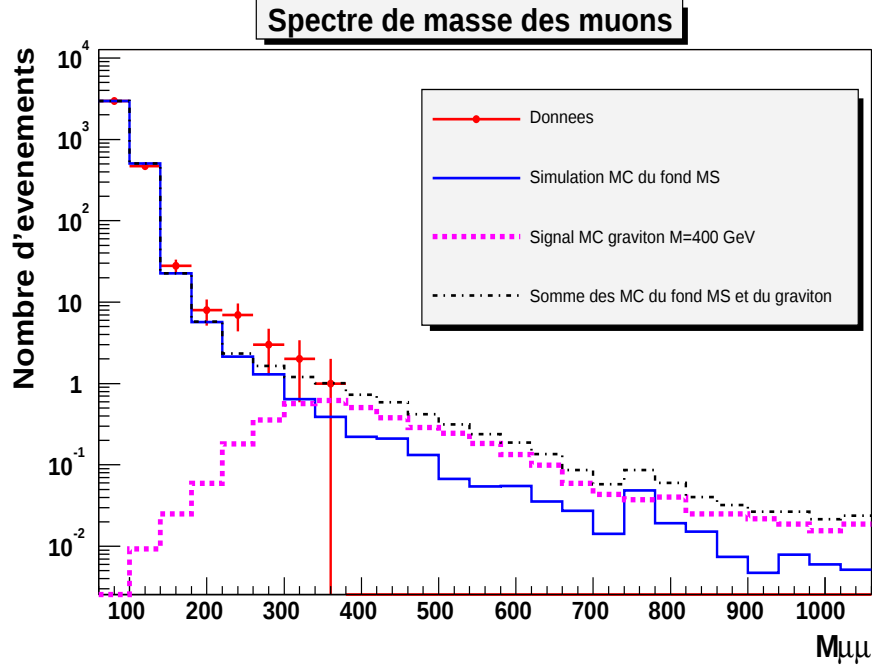


FIG. 6.1 – Spectres de masse invariante de paires de muons pour les données (croix), la simulation MC du bruit de fond du modèle standard (ligne continue), pour la simulation MC du signal de graviton de masse 400 GeV (tirets épais) et pour la somme du bruit de fond et du signal graviton (ligne fine pointillée). Le signal de graviton est tracé pour une section efficace de production de graviton de 1 pb. Tous les histogrammes sont normalisés à la luminosité des données, c'est-à-dire $107,8 \text{ pb}^{-1}$.

Ainsi, pour calculer les limites sur la section efficace de production du graviton pour une masse donnée, nous prenons un signal de graviton de couplage *a priori* quelconque³. Nous modifions de manière itérative la valeur de la section efficace de production du signal (et donc la normalisation de la distribution du signal) jusqu'à l'obtention d'une valeur de $CL_s = 0,05 \pm 5 \cdot 10^{-4}$ qui correspond à l'exclusion du signal à 95% de niveau de confiance pour la valeur de la section efficace de production considérée.

Dans le tableau 6.1 nous reportons quelques valeurs de sections efficaces limites obtenues par cette méthode en fonction de la masse du graviton. A titre comparatif et pour une estimation de l'ordre de grandeur, nous utilisons une méthode de simple comptage tenant compte des erreurs systématiques avec une coupure (non optimisée) sur la masse invariante des dimuons afin de confronter les résultats de cette méthode avec les résultats obtenus avec la méthode du rapport de vraisemblance. Nous voyons que pour des grandes masses la méthode utilisant le rapport de vraisemblance a une sensibilité meilleure que la méthode de simple comptage.

Sur la figure 6.3 est reportée la courbe d'exclusion du signal à 95% de niveau de confiance. On a tracé également les courbes théoriques donnant la section efficace de production de graviton en fonction de la masse de ce dernier pour quelques valeurs indicatives

3. Nous avons généré les signaux de graviton avec $k/M_{Pl} = 0,05$.

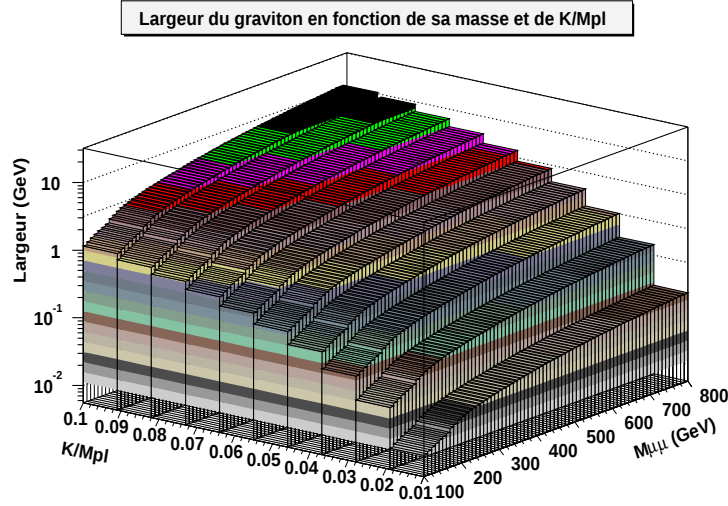


FIG. 6.2 – Largeurs des résonances de graviton pour différentes valeurs de masse et de k/M_{Pl} .

Masse du graviton(GeV)	200	300	400	500	600	800
Rapport de vraisemblance σ limite (pb)	1,44	0,751	0,365	0,245	0,214	0,228
comptage des événements σ limite (pb)	1,36	0,655	0,432	0,408	0,419	0,487
Coupure en masse: M(GeV) >	180	260	300	300	300	300

TAB. 6.1 – Section efficace limite à 95% de niveau de confiance calculée pour 6 masses différentes de graviton en utilisant d’une part la méthode du rapport de vraisemblance et d’autre part la méthode de simple comptage utilisée après une coupure non optimisée sur la masse invariante des paires de muons. La coupure sur la masse utilisée pour ce calcul est reportée en dernière ligne du tableau.

de k/M_{Pl} .

Les sections efficaces de ces courbes théoriques ne sont pas multipliées par un “ facteur $K = 1,3$ ” communément utilisé. En effet, ce “ facteur K ” est employé dans les diagrammes des processus du modèle standard et a été évalué par des calculs complexes pour tenir compte de la contribution à la section efficace des diagrammes aux ordres supérieurs. Il dépend entre autres du type de couplage au vertex, de la masse de la particule produite et des énergies mises en jeu, de sorte qu’il prend une valeur différente selon le processus étudié.

Nous n’avons donc *a priori* aucune raison de prendre la même valeur de ce facteur pour le cas du graviton, d’autant plus que dans ce cas les couplages sont dimensionnés et, de surcroît, que la théorie n’est pas renormalisable.

Nous sous-estimons probablement la section efficace en ne tenant pas compte de la contribution des ordres supérieurs (non évalués à ce jour) , mais le résultat que nous obtenons est conservatif.

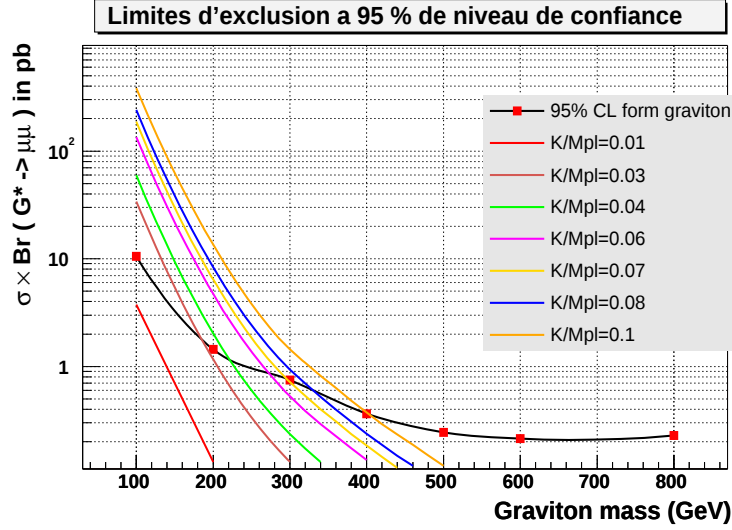


FIG. 6.3 – Limite à 95% de niveau de confiance sur la section efficace de production du premier état excité du graviton en fonction de la masse de ce dernier. Les courbes théoriques donnant la section efficace de production en fonction de la masse sont reportées pour quelques valeurs de k/M_{Pl} .

6.2.1 Limites sur les paramètres du modèle

Nous pouvons traduire la limite supérieure à 95% de niveau de confiance sur la section efficace de production du graviton en termes de contraintes sur les paramètres du modèle de Randall-Sundrum. Pour chaque valeur de couplage et de masse on calcule la section efficace limite par une méthode d'interpolation exponentielle⁴ à partir des points calculés explicitement (voir tableau 6.1).

Les résultats de cette interpolation définissent une zone d'exclusion du signal représentée sur la figure 6.4 dans le plan $(M_{grav}, k/M_{Pl})$.

Nous constatons que la masse la plus élevée exclue par notre analyse est de 400 GeV. Cette valeur est tout à fait satisfaisante compte tenu du fait que les sections efficaces données par la version correcte de PYTHIA sont inférieures à celles données par la version erronée utilisée jusqu'alors dans les études sur ce modèle. Notons que si nous avions appliqué un K-facteur de 1,3, la plus grande masse que nous aurions exclue serait de 438 GeV.

6.3 Conclusions et perspectives

Ce travail a permis de poser des contraintes sur les paramètres du modèle de Randall-Sundrum à travers la recherche du premier état massif graviton de Kaluza-Klein. Cette analyse est la première réalisée au sein de la collaboration DØ. Les résultats obtenus

4. L'interpolation exponentielle est préférée à l'interpolation linéaire utilisée dans des résultats préliminaires antérieurs car elle se rapproche davantage du comportement des courbes théoriques.

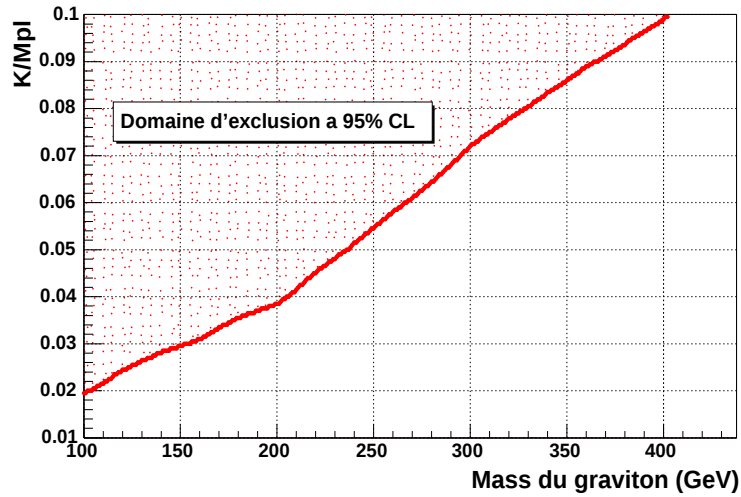


FIG. 6.4 – *Domaine d'exclusion à 95% de niveau de confiance des paramètres du modèle de Randall-Sundrum dans le plan k/M_{Pl} en fonction de la masse du premier état excité du graviton.*

sont comparables à ceux obtenus par la collaboration CDF, tenant compte de la différence entre les choix des deux collaborations DØ et CDF pour cette étude. En effet, la collaboration CDF a choisi de multiplier les sections efficaces de production de graviton par un K facteur de 1,3, donnant alors une plus grande sensibilité de l'analyse à un signal de graviton. Ce choix n'est pour l'instant pas complètement justifié sur le plan théorique, comme cela a été mentionné dans la section 6.2. Par ailleurs, la collaboration CDF utilise pour son analyse une version erronée de PYTHIA qui surestime les section efficaces théoriques (nous en avons parlé au chapitre 2).

6.3.1 Améliorations possibles de l'analyse

Nous avons vu que l'efficacité sur la production du Z suivie d'une désintégration muonique est d'environ 13,5%, principalement affectée par la couverture angulaire du détecteur. Dans le canal muonique, la couverture angulaire du système à muons ne peut être améliorée.

Nous pouvons alors améliorer l'efficacité de sélection des données en choisissant un lot de données non pas présélectionné sur des critères de qualité des muons mais sur un critère de déclenchement, en prenant le *trigger* le moins restrictif ou un ensemble de *triggers* muoniques qui permettent d'augmenter la statistique de l'échantillon de données analysées.

La sélection sur un critère de *trigger* permettrait en plus de sélectionner des muons de qualité acceptable, ce qui augmenterait l'efficacité de sélection des données sur le critère de qualité des muons de 50%, et permettrait aussi d'améliorer la statistique des lots de données utilisés pour calculer les efficacités de déclenchement.

6.3.2 Canal électromagnétique

Il est possible d'élargir les canaux d'étude de ce signal au canal électromagnétique par la reconstruction d'événements à haute masse avec des paires d'électrons ou de photons reconstruits dans le calorimètre.

La précision de mesure de l'énergie dans le calorimètre est meilleure que celle du système à muons pour des particules de haute énergie⁵, ce qui permettrait d'avoir un signal de graviton plus étroit et limiterait l'étalement du signal sur une trop grande gamme de masses.

De plus, ce choix permettrait d'avoir accès à un plus grand nombre total d'événements (electrons et photons confondus) puisque le rapport d'embranchement du graviton en électrons et photons est d'environ 6% tandis que le rapport d'embranchement du graviton en deux muons est seulement de 2%.

6.3.3 Autres études sur le modèle de Randall-Sundrum

Nous avons établi des contraintes sur le modèle de RS de base. Comme nous l'avons vu dans le premier chapitre, des extensions du modèle de Randall-Sundrum permettent une description plus complète du modèle de base, incluant de ce fait un champ scalaire dans le *bulk* appelé *radion* permettant de stabiliser le rayon de compactification. D'un point de vue expérimental, ce radion présente un grand intérêt phénoménologique et des possibilités de découverte auprès des accélérateurs très intéressantes, ce qui a mené les collaborations ATLAS et CMS à s'intéresser à ce canal.

Dans le cadre de l'extension supersymétrique du modèle de Randall Sundrum, les masses des états excités du graviton de Kaluza-Klein ne sont pas données par $m_n = x_n k e^{-k\pi r_c}$ mais sont proportionnelles à $(n+1/4)x_n k e^{-b_0 k\pi r_c}$ [8] (où b_0 est une constante introduite par la supersymétrisation du *bulk*). Cela a pour conséquence de décaler le spectre de masse des états excités de graviton d'environ un quart de la masse du premier état excité.

D'un point de vue expérimental, si nous n'observons que le premier état excité, nous ne pouvons rien conclure sur ce décalage. Par contre, si nous observons deux résonances successives de spin 2, l'écart entre ces deux résonances ainsi que la masse de la première résonance donneront des indications sur la pertinence de la description supersymétrique.

5. La résolution en énergie pour des particules électromagnétiques d'impulsion transverse de 45 GeV (respectivement 200 GeV) est d'environ 3% (respectivement 20%) tandis que celle du système à muons pour des muons de même impulsion transverse est de 8% (respectivement 48%)[6][7].

Bibliographie

- [1] W. T. Eadie, D. Drijard, F. E. James, M. Roos, B. Sadoulet, *Statistical Methods in Experimental Physics*, North-Holland, 1971, ISBN 0 7204 02395.
- [2] R. D. Cousins, V. L. Highland, *Incorporating Systematic Uncertainties into an Upper Limit*,
Nucl. Instr. Meth. **A320** (1992), 331-335.
- [3] T. Junk, Nucl. Instr. Meth. **A434** (1999) 435.
- [4] C. Delaere, <http://aleph-proj-alphapp.web.cern.ch/aleph-proj-alphapp/doc/tlimit.html>
- [5] R. Hamberg, W.L. Neerven, T. Matsuura, *A Complete Calculation of The Order α_s^2 Correction to The Drell-Yan K-Factor*,
Nucl. Phys. **B359** (1991) 343-405.
- [6] V. Buescher, J. Zhu, *em_cert: EM Certification Tools*,
DØ-note 4171.
- [7] D. Chapin, H. Fox, J. Gardner, R. Illingworth, A. Lyon, J. Zhu, *Measurement of $Z \rightarrow e^+e^-$ and $W \rightarrow e^\pm \nu$ Production Cross Sections with $\eta < 2, 3$* ,
DØ-note 4403.
- [8] A. Falkowski, Z. Lalak, S. Pokorski, *Five-Dimensional Gauge Supergravities with Universal Hypermultiplet and Warped Brane Worlds*,
Phys. Lett. **B509** (2001) 337-345.

Annexe A

Calcul d'amplitudes de probabilités avec FORM

Cette annexe a pour but de donner un exemple illustrant l'utilisation de FORM pour le calcul des amplitudes de probabilité de processus physiques. Ce code est celui que nous avons utilisé pour calculer les amplitudes des processus décrits dans le chapitre 2. Les notations reprennent celles de ce chapitre.

Il faut au préalable trouver l'expression littérale des termes de transition entre l'état initial et l'état final T_{fi} et T_{fi}^* en utilisant les règles de Feynman ainsi que les propagateurs et les couplages aux vertex.

Nous utilisons les notations décrivant:

- un fermion entrant : u et un fermion sortant : $\bar{u}$
- un anti-fermion entrant : $\bar{v}$ et un anti-fermion sortant : v
- un gluon de couleur a : g^a

ainsi que les propriétés suivantes:

$$\begin{aligned}\sum_{s_i} u(p) \bar{u}(p) &= \not{p} + m, \\ \sum_{s_i} v(p) \bar{v}(p) &= \not{p} - m, \\ \sum_{s_i} g_\mu^a(p) g_\nu^b(p) &= g_{\mu\nu} \delta^{ab}\end{aligned}$$

où l'on a noté les spins possibles s_i . En fait, nous considérerons que les quarks et les muons sont de masse nulle car leur masse est négligeable devant les énergies mises en jeu.

L'expression de T_{fi} s'écrit donc:

$$T_{fi} = \bar{u}(q1) C_{ffG}^{\alpha\beta} v(q2) P_{\mu\nu\alpha\beta} g_\rho^a(p1) C_{ggG}^{\mu\nu\rho\sigma} g_\sigma^b(p2) \quad (5)$$

où on a pris l'expression du propagateur $P_{\mu\nu\alpha\beta}$ de l'expression (2.4), et les $C_{ffG}^{\alpha\beta}$ et $C_{ggG}^{\mu\nu\rho\sigma}$ sont les expressions des couplages au graviton respectivement des fermions (ici les muons) et des gluons. Ces couplages s'écrivent:

$$C_{ffG}^{\alpha\beta} = \frac{-i}{4\Lambda_\pi} [(p_1 - p_2)_\alpha \gamma_\beta + (p_1 - p_2)_\beta \gamma_\alpha] \quad (6)$$

$$\text{et } C_{ggG}^{\mu\nu\rho\sigma} = \frac{-i}{\Lambda_\pi} \delta^{ab} [W_{\mu\nu\rho\sigma} + W_{\nu\mu\rho\sigma}] \quad (7)$$

où l'expression de $W_{\mu\nu\rho\sigma}$ est donnée par la relation (2.12) .

Par ailleurs, nous pouvons faire apparaître la dépendance en θ^* de la matrice de transition T_{fi} en faisant intervenir le produit scalaire entre les quantités de mouvement des particules en jeu. Nous prenons θ^* comme étant l'angle formé entre la particule incidente et la particule sortante, comme le montre la figure 2.8, et on définit les produits scalaires suivants:

$$p1.p3 = \hat{s} * (1 - \cos(\theta^*)) / 4$$

$$p2.p3 = \hat{s} * (1 + \cos(\theta^*)) / 4$$

$$p1.p4 = \hat{s} * (1 + \cos(\theta^*)) / 4$$

$$p2.p4 = \hat{s} * (1 - \cos(\theta^*)) / 4$$

où $\hat{s}$ est le carré de l'énergie dans le centre de masse de la collision et vaut $x_1 x_2 s$ avec $\sqrt{s}=1960$ GeV . $\sqrt{s}$ est l'énergie disponible lors de la collision proton-antiproton, x_1 et x_2 sont les fractions d'énergie des partons en collision dans le proton ou l'antiproton (ce sont les $x_{\text{Björken}}$).

Dans le programme ci-dessous, nous avons noté:

- **s** pour $\hat{s}$;
- **mz, mg, Mpl** respectivement pour les masses du Z, du graviton et de Planck;
- **wz, wg** respectivement pour les largeurs du Z et du graviton;
- **eq, emu** pour les charges électriques du quark et du muon;
- **e** pour le couplage électrofaible;
- **gvmu, gamu, gvq, gaq** sont les termes pour les muons et pour les quarks intervenant dans le couplage (V-A) avec le Z;
- **ct** pour $\cos(\theta^*)$;
- **costhw** pour $\cos(\theta_W)$, θ (Weinberg) qui est un des paramètres électrofaible.

Les types de variables peuvent se mettre au singulier comme au pluriel. Dans les versions précédentes de FORM, les types au singulier ne pouvaient déclarer qu'une seule variable à la fois, les types au pluriel ne pouvaient pas déclarer une seule variable mais au moins deux. Cette restriction grammaticale a été levée dans les versions récentes afin de simplifier l'écriture du code. Nous pouvons utiliser indifféremment le singulier ou le pluriel.

Nous donnons quelques notes d'explication afin de faciliter la compréhension du code.

Concernant le type de variables:

Le type *Symbol* a les propriétés mathématiques des scalaires.

Le type *Indice* est un indice de sommation. Dès que le même indice se retrouve deux fois dans la même expression, FORM effectue la sommation dessus.

Le type *Vector* est traité comme un quadri-vecteur, sauf si l'on spécifie au début du programme "dimension 3" (ou une autre dimension), la dimension sera alors changée. Le produit scalaire entre deux vecteurs V1 et V2 est noté V1.V2, FORM transforme en produit scalaire toutes les paires de vecteur lorsqu'ils ont le même indice V1(mu)*V2(mu). Par ailleurs, dès que FORM trouve V(mu), ce terme est considéré comme un scalaire et peut alors commuter avec tout autre terme de l'expression.

Le type *Function* peut avoir des indices ou des vecteurs en argument. *CFunction* est une fonction qui commute avec tout autre objet.

Le type *Tensor* que nous n'avons pas utilisé dans notre programme a les propriétés des tenseurs mathématiques, mais aucune différence n'est faite quant aux indices covariants ou contravariants. Les règles de sommation sur les indices doivent être spécifiées dans le code. *CTensor* sont des tenseurs qui commutent avec tout autre objet.

Le type *Local* spécifie une expression à calculer, ce n'est pas une déclaration de variable mais plutôt une définition de variable qui s'exprime en fonction de tous les paramètres précédemment déclarés.

Notons que la déclaration de variables se fait en tout début de programme et nous ne pouvons pas déclarer de variable à n'importe quel moment comme cela est possible en C++.

Concernant les outils mathématiques que FORM connaît:

Le nombre complexe i est noté i_- dans le code.

La métrique $\eta_{\mu\nu}$ est notée $d_-(mu, nu)$. En fait, $d_-(mu, nu)$ correspond au symbole δ de Kröneckers. Mais dans les calculs avec les matrices γ de Dirac, cela ne fait pas de différence. Les matrices γ^μ de Dirac sont notées $g_-(mu, a)$ où $a=1,2,\dots$ est un entier qui sert à désigner une ligne de spin. Dans notre code, on a noté 1 la ligne de spin concernant les muons et 2 celle concernant les quarks ou les gluons. Lorsque l'on écrit "trace4,1", cela indique que l'on effectue la trace en quatre dimensions de toutes les expressions définies sous le type "Local" en ne tenant compte que de la ligne de spin désignée par 1. Comme nous avons deux lignes de spin, nous devons calculer deux traces différentes: "trace4,1" et "trace4,2". En trois dimensions, nous aurions écrit "trace3,1".

FORM connaît d'autres outils mathématiques très utiles dont nous ne donnerons pas la liste. Nous référons au site web de FORM <http://www.nikhef.nl/~form> qui contient les exécutable ainsi qu'un "tutorial" et un manuel d'utilisation.

Concernant les opérateurs et les commandes:

L'opérateur "*id*" permet de remplacer l'expression qui le suit par celle qui se trouve à droite du signe "=". L'opérateur "?" généralise la variable qui se trouve à sa gauche à toutes les variables du même type. Nous l'avons utilisé pour les vecteurs lorsqu'on a écrit dans le code ci-dessous: "p1?.p1?=0" ce qui signifie que l'on veut remplacer le produit scalaire de tout vecteur avec lui-même par zéro (ce produit scalaire donne en réalité la masse carrée de la particule concernée, mais nous avons considéré que leurs masses étaient nulles). Nous avons bien sûr pris garde à d'abord demander "**id** k.k = s" où s est le carré de l'énergie dans le centre de masse et k est l'impulsion de la particule produite en voie s , $k=p_1+p_2=q_1+q_2$.

La commande "format fortran" permet de convertir le résultat en un format compatible avec le fortran. Le résultat peut être copié tel quel et collé dans un programme en fortran.

La commande ".sort" permet au préprocesseur d'effectuer une compilation préliminaire du programme.

La commande "print" permet d'écrire le résultat de la variable déclarée (ou plutôt définie) sous le type "Local". L'option "+s" qui la suit permet d'obtenir une meilleure clarté et lisibilité du résultat puisque son rôle est d'écrire l'expression en revenant à la ligne à chaque nouveau terme lorsque l'expression est une somme de plusieurs termes.

La commande ".end" désigne la fin du programme.

Dans ce qui suit nous donnons le programme que nous avons utilisé.

```
*****
*****                               *****
*****      CODE DE CALCUL DES      *****
*****      AMPLITUDES DE PROBABILITE      *****
*****
```

***** Déclaration des variables

Symbol s, mz, wz, mg, wg, Mpl, eq, emu, e, gvmu, gamu, gvq, gaq, g;
Symbols ct, costhw;

CFunctions ubarg,vbarg,ug,vg;
Function umu, ubarmu, vmu, vbarmu;
Functions uq, ubarq, vq, vbarq;

Indices mu, nu, alp, bet;
Index rho, sig, eps, del, phi, khi, psi, tau, a, b, c, d;

Vectors p1, p2, q1, q2, k, k1, k2;

```
*****
*****      Expression des termes à calculer      *****
*****
```

```
*****
*****      Termes pour le graviton      *****
```

***** Propagateur

Local Gpropag = $\mathbf{i}_- * ((s - m_g^2) - (\mathbf{i}_- * m_g * w_g))$
 $* \left(\frac{1}{2} * (\mathbf{d}_-(\mu, \alpha p) * \mathbf{d}_-(\nu, \beta t) + \mathbf{d}_-(\mu, \beta t) * \mathbf{d}_-(\nu, \alpha p) - \mathbf{d}_-(\mu, \nu) * \mathbf{d}_-(\alpha p, \beta t)) \right.$
 $- \frac{1}{(2 * m_g^2)} * (\mathbf{d}_-(\mu, \alpha p) * k(\nu) * k(\beta t) + \mathbf{d}_-(\nu, \beta t) * k(\mu) * k(\alpha p)$
 $\quad + \mathbf{d}_-(\mu, \beta t) * k(\nu) * k(\alpha p) + \mathbf{d}_-(\nu, \alpha p) * k(\mu) * k(\beta t))$
 $\left. + \frac{1}{6} * (\mathbf{d}_-(\mu, \nu) + (2 / m_g^2 * k(\mu) * k(\nu))) * (\mathbf{d}_-(\alpha p, \beta t) + 2 / m_g^2 * k(\alpha p) * k(\beta t)) \right);$

Local Gpropagstar = $-\mathbf{i}_- * ((s - m_g^2) + (\mathbf{i}_- * m_g * w_g))$
 $* \left(\frac{1}{2} * (\mathbf{d}_-(\rho, \epsilon s) * \mathbf{d}_-(\sigma, \delta l) + \mathbf{d}_-(\rho, \delta l) * \mathbf{d}_-(\sigma, \epsilon s) - \mathbf{d}_-(\rho, \sigma) * \mathbf{d}_-(\epsilon s, \delta l)) \right.$
 $- \frac{1}{(2 * m_g^2)} * (\mathbf{d}_-(\rho, \epsilon s) * k(\sigma) * k(\delta l) + \mathbf{d}_-(\sigma, \delta l) * k(\rho) * k(\epsilon s)$
 $\quad + \mathbf{d}_-(\rho, \delta l) * k(\sigma) * k(\epsilon s) + \mathbf{d}_-(\sigma, \epsilon s) * k(\rho) * k(\delta l))$
 $\left. + \frac{1}{6} * (\mathbf{d}_-(\rho, \sigma) + (2 / m_g^2 * k(\rho) * k(\sigma))) * (\mathbf{d}_-(\epsilon s, \delta l) + 2 / m_g^2 * k(\epsilon s) * k(\delta l)) \right);$

***** Vertex $q\bar{q} \rightarrow G G \rightarrow m m$

Local VtxGq = -i₋/(4*Mpl) *(k1(mu)*g₋(2,nu)+k1(nu)*g₋(2,mu));
Local VtxGqstar = i₋/(4*Mpl)*(k1(rho)*g₋(2,sig)+k1(sig)*g₋(2,rho));

Local VtxGmu = -i₋/(4*Mpl)*(k2(alp)*g₋(1,bet)+k2(bet)*g₋(1,alp));
Local VtxGmustar = i₋/(4*Mpl)*(k2(eps)*g₋(1,del)+k2(del)*g₋(1,eps));

***** Vertex gg- > GG

Local VtxGg = -i₋/Mpl*(k1(mu)*g₋(2,nu)+k1(nu)*g₋(2,mu));
Local VtxGgstar = i₋/Mpl*(k1(rho)*g₋(2,sig)+k1(sig)*g₋(2,rho));

***** Termes pour le Z

Local Zpropag = -i₋*(d₋(phi,khi)-k(phi)*k(khi)/mz^2)
 *((s-mz^2)-(i₋*mz*wz))/((s-mz^2)^2+(mz*wz)^2);

Local Zpropagstar = i₋*(d₋(psi,tau)-k(psi)*k(tau)/mz^2)
 *((s-mz^2)+(i₋*mz*wz))/((s-mz^2)^2+(mz*wz)^2);

Local VtxZq = -i₋*g/(2*costhw)*g₋(2,phi)*(gvq-gaq*g5₋(2));
Local VtxZqstar = i₋*g/(2*costhw)*(gvq+gaq*g5₋(2))*g₋(2,psi);

Local VtxZmu = -i₋*g/(2*costhw)*g₋(1,khi)*(gvmu-gamu*g5₋(1));
Local VtxZmustar= i₋*g/(2*costhw)*(gvmu+gamu*g5₋(1))*g₋(1,tau);

***** Termes pour le DY

Local DYpropag = -i₋*(1/s)*d₋(a,b);
Local DYpropagstar= i₋*(1/s)*d₋(c,d);

Local VtxDYq = -i₋*(e^2*eq)*g₋(2,a);
Local VtxDYqstar = i₋*(e^2*eq)*g₋(2,c);

Local VtxDYmu = -i₋*(e^2*emu)*g₋(1,b);
Local VtxDYmustar= i₋*(e^2*emu)*g₋(1,d);

***** Calcul de Tfi et Tfistar pour chaque processus *****

Local Gfiqq = ubarmu*VtxGmu*vmu* Gpropag *vbarq*VtxGq*uq;
Local Gfiqqstar= ubarq*VtxGqstar*vq * Gpropagstar *vbarmu*VtxGmustar*umu;

Local Gfigg = ubarmu*VtxGmu*vmu* Gpropag *vbarg*VtxGg*ug;
Local Gfiggstar= ubarg*VtxGgstar*vg * Gpropagstar *vbarmu*VtxGmustar*umu;

Local Zfi = ubarmu*VtxZmu*vmu* Zpropag *vbarq*VtxZq*uq;
Local Zfistar= ubarq*VtxZqstar*vq * Zpropagstar *vbarmu*VtxZmustar*umu;

Local DYfi = ubarmu*VtxDYmu*vmu* DYpropag *vbarq*VtxDYq*uq;
Local DYfistar= ubarq*VtxDYqstar*vq * DYpropagstar *vbarmu*VtxDYmustar*umu;

 ***** AMPLITUDES *****

Local Gampgg =Gfigg *Gfiggstar;
Local Gampqq =Gfiqq *Gfistarqq;
Local Zamp =Zfi *Zfistar;
Local DYamp =DYfi*DYfistar;

Local DYZinterf =((Zfi*DYfistar)+ (DYfi*Zfistar));

Local GZint = ((Zfi*Gfiqqstar) + (Gfiqq* Zfistar));
Local GDYint = (DYfi*Gfiqqstar) + (Gfiqq*DYfistar);

Local GDYZamp = Gampqq + Zamp + DYamp + DYZinterf + GZint + GDYint;
Local DYZamp = Zamp + DYamp+ DYZinterf;

 ***** REPLACEMENT : id *****

id uq=1;
id ubarq=g_(2,p1);
id vq=1;
id vbarq=g_(2,p2);

id umu=1;
id ubarmu=g_(1,q1);
id vmu=1;
id vbarmu=g_(1,q2);

id ubarg(mu?)*ug(nu?)=d_(mu,nu);

```
id vbarg(nu?)*vg(mu?)=d_(mu,nu);
```

```
trace4,1;
```

```
trace4,2;
```

```
.sort
```

```
id k.k = s;
```

```
id k1=p1-p2;
```

```
id k2=q1-q2;
```

```
id k=q1+q2;
```

```
id p1.q1 = s*(1-ct)/4;
```

```
id p2.q1 = s*(1+ct)/4;
```

```
id p1.q2 = s*(1+ct)/4;
```

```
id p2.q2 = s*(1-ct)/4;
```

```
id p1?.p1? = 0;
```

```
id q1.q2 = s/2;
```

```
id p1.p2 = s/2;
```

```
*****
```

```
format fortran;
```

```
print +s Gampgg;
```

```
print +s Gampqq;
```

```
print +s DYamp;
```

```
print +s Zamp;
```

```
print +s DYZinterf;
```

```
print +s GDYint;
```

```
print +s GZint;
```

```
.end
```


Remerciements

Je voudrais tout d'abord remercier les membres du jury Pierre Billoir, Jean-François Grivaz et Didier Vilanova qui ont accepté de se pencher sur mon travail, et plus particulièrement les rapporteurs Philippe Brax et Jacques Dumarchez pour le temps qu'ils ont consacré à relire le manuscrit.

Je tiens à remercier Pascal Debu qui m'a accueillie dans ce laboratoire et qui, même après avoir cédé son rôle de chef de service, a montré de l'intérêt à l'avancement de mes travaux et mes projets d'avenir.

Je remercie Marc Besançon qui a proposé un sujet de recherche dans un domaine de physique très intéressant et m'a laissé entièrement libre.

Je remercie aussi tous les membres du groupe DØ-Saclay: Mathieu Agelou, Jiri Bystriky, Frédéric Déliot, Pavel Demine, Antoine Kouchner, Alexander Kupco, Pierre Lutz, Marine Michaut, Emmanuelle Perez, Christophe Royon, Boris Tuchming, Didier Vilanova et Armand Zylberstejn.

Je tiens à remercier en particulier les chefs du groupe DØ-Saclay qui se sont impliqués dans le déroulement de ma thèse: Armand Zylberstejn qui m'a bien prise en charge lors de mes premières missions à Fermilab, je garderai toujours un très bon souvenir de mes missions avec lui et son côté paternel avec tous les membres du groupe; je remercie aussi Emmanuelle Perez qui, malgré toutes ses responsabilités et son emploi du temps très chargé, a su trouver du temps à consacrer à l'avancement de ce travail! Elle y a beaucoup apporté en qualité, en rigueur et en outils d'analyse. Elle a joué un rôle très important et je lui en suis reconnaissante.

Je remercie particulièrement Didier Vilanova qui a accepté de relire ma thèse avec beaucoup de sérieux ainsi que Boris Tuchming qui a indéniablement contribué à la qualité de ce travail. Ils m'ont tous deux accordé de leur temps avec beaucoup de gentillesse et j'ai particulièrement apprécié leur rigueur scientifique, leur grande disponibilité et leur humeur joviale et encourageante.

Je tiens aussi à exprimer ma gratitude envers le groupe DØ-Marseille au CPPM. J'ai été très bien accueillie dans leur laboratoire pendant la semaine de mission au sein de leur équipe; j'ai notamment apprécié leur simplicité et leur très grande gentillesse. Ils m'ont tout de suite intégrée et m'ont comptée parmi les leurs. Au delà de leurs compétences et leur grande implication dans la collaboration, ils sont restés très humains et au contact très agréables. Pour votre savoir-vivre, votre humanité et votre gentillesse, Elemér Nagy, Marie-Claude Cousinou, Talby Mossaddek, Eric Kajfasz, Arnaud Duperrin, Smaïn Kermich, Eric Thomas, Frédéric Villeneuve, Stéphanie Baffioni et Alexis Cothenet, je vous présente mes sincères remerciements.

Je ne remercierai jamais assez Michel Cribier, mon parrain de thèse, qui m'a toujours défendue, soutenue, conseillée, aidée. Sa présence, sa disponibilité et son authenticité m'ont été d'un très grand secours. Il a marqué positivement ma thèse et je lui dois beau-

coup. Je tiens à lui exprimer toute mon admiration, ma gratitude et ma reconnaissance pour ce qu'il est et ce qu'il a fait pour ma thèse.

Je ne pourrai oublier le temps que Nicolas Châtillon a consacré à partager sa compréhension des subtilités de la relativité générale et des dimensions supplémentaires. Il a toujours répondu promptement à mes innombrables questions avec un remarquable sens de la pédagogie. Il m'a tout simplement redonné goût à la physique et à la phénoménologie. Pour tout cela, Nico, je m'incline bien bas et te remercie de m'avoir accompagnée tout ce temps.

J'ai beaucoup apprécié les discussions très enrichissantes sur des sujets très variés avec Steve Muanza qui m'a beaucoup apporté tant par ses connaissances en physique que pour son soutien moral, ainsi que Oleg Kouznetsov qui m'a toujours impressionnée par son ouverture d'esprit et la diversité de ses expériences humaines et socio-culturelles. Ce fut un immense plaisir de les côtoyer et j'espère les croiser plus souvent dans le futur.

Je remercie Greg Landsberg qui m'a encouragée et a contribué à la qualité de cette thèse. Il a réorienté mon sujet de thèse vers l'étude d'un modèle non exploré dans la collaboration DØ. Ce manuscrit contient son empreinte et je lui suis reconnaissante pour son implication, sa lucidité et surtout ses encouragements.

Je remercie aussi Laurent Chevalier pour m'avoir rassurée sur plusieurs points et pour m'avoir encouragée, Vanina Rhulmann-Kleider pour ses conseils et orientations (j'aurai dû la consulter plus souvent), Philippe Brax pour ses références théoriques très intéressantes et sa disponibilité, Marie-Claude Lemaire et Caroline Collard pour les interactions constructives que nous avons eues, ainsi qu'Anne-Isabelle Etievre.

Je suis très reconnaissante envers Gautier Hamel qui m'a accordé de son temps avec beaucoup de gentillesse et de simplicité. Il a été présent dans des moments critiques, et pour passer à travers mes transparents.

Je remercie tout particulièrement Nathalie Besson qui a toujours été disponible avec beaucoup d'enthousiasme et de bonne humeur. Elle a été de très bon conseil sur tous les plans concernant ma thèse et m'a apporté son aide dans de nombreux domaines. Sa présence, ainsi que celle de Guillaume Gouge et Nicolas Kerschen, a donné une touche agréable à mes trois années de thèse. Grâce à eux j'en garderai de bons souvenirs.

Je voudrais particulièrement remercier les personnes qui ont toujours cru en moi, qui m'ont toujours soutenue durant toute ma thèse, principalement parce que beaucoup d'entre eux, étant thésards eux-mêmes, étaient très bien placés pour me comprendre et m'aider à relativiser et à essayer de voir un côté positif aux choses, même quand il n'y en avait pas! Pour toutes nos discussions très encourageantes, mes remerciements sincères et ma grande gratitude vont donc à Nicolas Châtillon, Daphné Denans, Guillaume Gouge, Samira Hassani, Fabrice Jouvenot, Nicolas Kerschen, Olivier Lenoble, Orianna Perru, Grégory Schott, Flore Skaza, mais aussi Jérémy Argiriades, Antoine Cazes, Simon Fiorucci, Clarice Hamadach, Muriel Pivk, Sébastien Saouter et Nazim Djeghloul. J'ai une pensée toute particulière pour Nicolas Châtillon, Nicolas Kerschen, Guillaume Gouge et Samira Hassani : ils ont été tout simplement merveilleux, je leur souhaite un grand avenir,

à la mesure de la grandeur de coeur et de leur grande compétence.

Je garderai un excellent souvenir de tous les moments passés avec les thésards du groupe DØ-France, aussi bien à Fermilab qu'à Beaune: Jean-Laurent Agram, Stéphanie Baffioni, Emmanuel Busato (pour la Hancock Tower et les dérapages sur la neige), Jérôme Coss (le polisson du T.), Alexis Cothenet (il est meilleur au naturel!), Pierre-Antoine Delsart, Sébastien Gréder (un être exceptionnel), Anne-Catherine Le Bihan, Anne-Marie Magnan, Marine Michaut, Tuan Vu Ahn (monsieur Pourquoi), Alexandre Zabi, j'espère n'oublier personne!

Pour finir, je ne saurais quantifier ni qualifier à sa juste valeur ce que m'a apporté la présence et le soutien de ma famille et particulièrement de mon fiancé, Sébastien Troquereau. Sébastien a été un soutien sans faille. Il a su me rassurer et m'aider à mener à terme ce travail et me rappeler chaque jour que la recherche a toujours été le métier que je désire exercer. Sans sa présence je n'aurais assurément jamais trouvé le courage ni la force de passer outre les difficultés qui ont jalonné les trois années de thèse. Ce manuscrit est en grande partie son œuvre, c'est à lui que je le dédie.

Merci à tous,

N. S. Lahrichi

Résumé

Dans cette thèse nous avons posé les premières contraintes sur les paramètres du modèle de dimensions supplémentaires de Randall-Sundrum, à savoir k/M_{Pl} qui est proportionnel au couplage du graviton aux champs du modèle standard et M_G qui est la masse du premier état excité du graviton.

L'analyse du signal Monte Carlo est basée sur le générateur PYTHIA. Cette étude a permis de mettre en évidence et de rectifier une erreur dans le générateur PYTHIA grâce à l'élaboration d'un générateur indépendant dédié à cette analyse.

Le lot de données utilisées pour l'analyse correspond à la période de prise des données effectuée par la collaboration Dzero au Tevatron allant de novembre 2002 à juillet 2003, qui correspond à une luminosité totale de $107,8 \text{ pb}^{-1}$.

La recherche du graviton dans ses désintégrations en paires de muons a permis de mesurer dans un premier temps la section efficace de production du boson Z multiplié par le rapport d'embranchement du Z en deux muons.

Mots-clés

Graviton, dimension supplémentaire, Kaluza-Klein, Randall-Sundrum, muon, Dzero, Tevatron.

Summary

In this thesis we have put the first constraints on the fundamental parameters of the Randall-Sundrum model of extra dimensions, k/M_{Pl} which is proportional to the coupling of the graviton to the standard model fields and M_G which is the mass of the first excited state of the Kaluza-Klein graviton.

The analysis performed on Monte carlo sample of the signal allowed to find an error in the PYTHIA generator. The elaboration of an independent generator dedicated for this special analysis helped to find out and correct the error.

The data sample used for the analysis covers the period running from november 2002 up to july 2003 taken by the Dzero collaboration at Tevatron, which corresponds to an accumulated luminosity of $107,8 \text{ pb}^{-1}$.

The search for the graviton in the dimuon channel allowed to measure the Z production cross-section times the branching ratio in dimuons.

Key words

Graviton, extra dimension, Kaluza-Klein, Randall-Sundrum, muon, Dzero, Tevatron.