



Recherche de gluinos dans la topologie à jets de quarks b et énergie transverse manquante avec le détecteur D0 au TeVatron

Thomas Millet

► To cite this version:

Thomas Millet. Recherche de gluinos dans la topologie à jets de quarks b et énergie transverse manquante avec le détecteur D0 au TeVatron. Physique des Hautes Energies - Expérience [hep-ex]. Université Claude Bernard - Lyon I, 2007. Français. NNT: . tel-00175297

HAL Id: tel-00175297

<https://theses.hal.science/tel-00175297>

Submitted on 27 Sep 2007

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.



N° d'ordre 62-2007
LYCEN – T 2007-09

Thèse

présentée devant

l'Université Claude Bernard Lyon-I

pour l'obtention du

DIPLOME de DOCTORAT
Spécialité : Physique des Particules

(arrêté du 7 août 2006)

par

Thomas MILLET

**Recherche de gluinos dans la topologie à jets de quarks b et
énergie transverse manquante avec le détecteur D0
au TeVatron**

Soutenue le 11 mai 2007
devant la Commission d'Examen

Jury :	M.	B.	Ille	Président du jury
	M.	A.	Deandrea	Directeur de thèse
	M.	G.-S.	Muanza	
	M.	D.	Froidevaux	Rapporteur
	M.	G.	Brooijmans	Rapporteur
	M.	A.	Duperrin	
	M.	P.	Verdier	

UNIVERSITÉ CLAUDE BERNARD - LYON 1
Institut de Physique Nucléaire de Lyon

Mémoire de thèse
pour l'obtention du grade de
Docteur de l'Université Claude Bernard - Lyon 1
Spécialité : Physique de particules
au titre de l'Ecole doctorale de Physique et Astrophysique fondamentale Rhône-Alpes

présentée et soutenue publiquement le 11 mai 2007
par M. Thomas MILLET

Recherche de gluinos dans la topologie à jets de quark b et énergie transverse manquante avec le détecteur DØ au TeVatron

Devant la commission d'examen formée de :
M. Bernard ILLE
M. Aldo DEANDREA (directeur de thèse)
M. Steve MUANZA
M. Daniel FROIDEVAUX (rapporteur)
M. Gustaaf BROOIJMANS (rapporteur)
M. Arnaud DUPERRIN
M. Patrice VERDIER

Table des matières

Remerciements	7
Introduction	11
1 Cadre théorique	13
1.1 Introduction	14
1.2 Les symétries en physique des particules	15
1.2.1 Notions de bases	15
1.2.2 Les symétries externes	15
1.2.3 Les symétries internes	16
1.2.4 Les brisures des symétries	17
1.3 Le modèle standard de la physique des particules	18
1.3.1 Introduction	18
1.3.2 L'électrodynamique quantique : point de départ d'une théorie aboutie . . .	18
1.3.3 La force faible et l'interaction électrofaible	20
1.3.4 Le mécanisme de Higgs	23
1.3.5 La chromodynamique quantique	25
1.3.6 Quelques mots sur la renormalisation d'une théorie	29
1.3.7 Les réussites du modèle standard	30
1.3.8 Le contenu en champs	31
1.3.9 Les limites du modèles standard et quelques exemples d'extension	32
1.4 La supersymétrie	35
1.4.1 Introduction et motivations	35
1.4.2 L'algèbre supersymétrique	36
1.4.3 Le lagrangien supersymétrique	38
1.4.4 Le modèle standard supersymétrique minimal (MSSM)	42
1.4.5 Conclusion	48
2 Le dispositif expérimental	51
2.1 Introduction	52
2.2 La chaîne de production, d'accélération et de collision des protons et des antiprotons	53
2.2.1 Généralité	53
2.2.2 Principes de fonctionnement	56
2.3 Le détecteur DØ	62
2.3.1 Généralités	63
2.3.2 Le trajectographe interne	64
2.3.3 Le calorimètre	69

2.3.4	Le spectromètre à muons	77
2.3.5	La mesure de la luminosité	80
3	Des collisions aux données reconstruites	85
3.1	Introduction	86
3.2	Le système de déclenchement de $D\bar{O}$ au <i>Run IIa</i>	86
3.2.1	Un système de sélection à trois niveaux	87
3.2.2	Les outils utilisés pour l'étude et la conception de nouvelles conditions de déclenchement	91
3.2.3	Les conditions de déclenchement spécifiques aux signaux à jets de hautes énergies transverses et à haute énergie transverse manquante	92
3.3	Améliorations apportées au <i>Run IIb</i>	104
3.3.1	Description détaillée des modifications apportées pour la partie calorimétrique	106
3.3.2	Etalonnage du système de déclenchement du calorimètre au niveau 1	111
3.3.3	Conception des conditions de déclenchement spécifiques aux signaux à jets et $\cancel{E}_T$ pour le <i>Run IIb</i> : la liste <i>V15</i>	125
3.4	Reconstruction des objets physiques	136
3.4.1	Les traces et les vertex	136
3.4.2	Les muons	138
3.4.3	Les particules électromagnétiques	141
3.4.4	Les jets	144
3.4.5	L'énergie transverse manquante	162
3.4.6	L'étiquetage des jets de hadrons beaux	163
3.4.7	Conclusion	170
4	Recherche d'un signal supersymétrique	171
4.1	Introduction et motivations	172
4.2	Caractéristiques du signal	172
4.2.1	Le cadre théorique	172
4.2.2	Production de paires de gluinos au TeVatron	173
4.2.3	Désintégration en un état final $4b + mE_T$	174
4.2.4	Simulation du signal	175
4.3	Les lots de données utilisés	177
4.3.1	Le lot de données du détecteur	178
4.3.2	Les lots de données <i>Monte Carlo</i>	179
4.4	Sélection des événements et traitement des données <i>Monte Carlo</i>	188
4.4.1	La stratégie de l'analyse	188
4.4.2	Le bruit fond instrumental <i>QCD</i>	195
4.4.3	La confirmation des jets par le trajectographe interne	201
4.4.4	Coupsures sur le nombre de leptons et zones de contrôle du bruit de fond physique	211
4.4.5	L'étiquetage des jets des quarks b	216
4.4.6	Coupsures finales avant l'optimisation	226
4.5	Optimisation finale	229
4.5.1	Méthode statistique	229
4.5.2	Incertitudes systématiques	230

4.5.3	Détermination des masses limites	230
4.5.4	Contours d'exclusion dans le plan $(m_{\tilde{g}}, m_{\tilde{b}_1})$	239
Conclusion		243
Bibliographie		253

Remerciements

La thèse est souvent synonyme de travail, mais elle rime également avec rencontres heureuses et relations humaines. J'aimerais remercier tous ceux qui ont été là tout au long de ces trois années, ces soutiens multiples qui rendent le travail plus facile, ces nombreuses sources d'énergie sans qui rien n'aurait été possible.

Je tiens tout d'abord à remercier Bernard Ille, directeur de l'I.P.N.L., d'une part pour son accueil au sein du laboratoire lyonnais, d'autre part pour avoir accepté de présider mon jury.

Merci à l'équipe D0 de Lyon pour son accompagnement tout au long de cette thèse. Je remercie tout particulièrement Steve Muanza pour son encadrement, qui après un court stage de D.E.A., a bien voulu poursuivre l'aventure avec moi et ainsi m'a permis de continuer mon apprentissage du monde de la physique des particules. Je tiens également à remercier Patrice Verdier pour son aide, ses conseils et sa rigueur au travail. Même si j'ai parfois buté sur cette dernière, je dois bien reconnaître qu'elle facilite le travail et le rend meilleur. J'espère ne pas l'oublier. Merci à Catherine Biscarat, ta venue dans l'équipe a été un véritable plaisir pour ma dernière année de thèse. Je remercie aussi Gérard Grenier, Tibor Kurca, Patrice Lebrun, Jean-Paul Martin qui complètent cette équipe. J'ai eu la chance de trouver auprès d'eux les conseils avisés indispensables à mon travail. J'ajoute à cette liste de physiciennes et physiciens, Aldo Deandrea mon directeur de thèse. Je te remercie pour tes éclaircissements théoriques, le premier chapitre de ce manuscrit en a grandement bénéficié.

J'aimerais remercier une autre équipe de l'I.P.N.L. sans qui cette thèse n'aurait pas été possible : l'équipe du service missions composée de Marie Berthier, Andrée Ducloux et Narcisse Katy. Grâce à elles, j'ai pu m'envoler si souvent à Fermilab.

Je remercie également Gustaaf Brooijmans et Daniel Froidevaux, rapporteurs de thèse, pour la lecture de ce manuscrit et leurs commentaires.

Une thèse dans l'expérience D0, c'est l'avantage du travail en groupe et de ces fameuses rencontres. Les réunions D0-France à Lyon, Grenoble, Paris et Marseille, et les meetings de collab' à Fermilab ou même Vancouver sont à chaque fois de nouvelles occasions de retrouver tous ceux que j'aimerais remercier. La liste des gens que j'ai eu la chance de rencontrer est longue et je m'excuse d'avance de ne pouvoir tous vous citer. Merci à tous !

Je tiens à remercier leur digne représentant : Eric Kajfasz (comme ça se prononce). J'ai pu constater qu'à tes côtés convivialité, accessibilité et sourire n'ont rien d'incompatible avec le travail.

Je remercie également l'équipe "Commissionning" pour ces mois passés avec vous à Fermilab, ces parties de Uno, de tarot, ces matchs de coupe du monde, ces soirées à refaire le monde autour de quelques verres. Merci à mes colocos du moment : Sam et Bertrand. Sans vous, jamais je n'aurais pu pousser la voiture à 7h30 du matin. Bon courage à toi, Sam, pour la dernière ligne droite. Merci

aussi à Christophe et Florent, qui complètent à merveille cette équipe. Je n'oublierai pas les longs débats nocturnes indispensables à l'accomplissement de notre travail. Les talents d'orateur de Christophe se sont alors avérés bien utiles. Je remercie bien sûr Vincent, qui sans faire partie de l'équipe, a grandement contribué à intellectualiser ces discussions.

Cette équipe reste incomplète si j'oublie de citer nos encadrants. Je remercie Arnaud Duperrin sans qui je n'aurais jamais pu connaître les joies des "triggers". Je remercie également Jan Stark. Il faudra que tu m'expliques un jour comment tu fais pour réfléchir aussi vite.

Je tiens également à remercier tous les doctorants de D0-France. Bon courage à tous pour la suite. Je remercie Vincent L., Marine que j'ai eu le plaisir de mieux connaître en Autriche, Aurélien, Benoît, Nikola avec qui j'ai eu la chance de partager une chambre de l'autre côté du globe au doux son du yukulélé, Fabrice, Anne-Fleur à la bonne humeur communicative, Jérémy. Je remercie très chaleureusement Marion, qui a eu la "chance" d'entendre mes coups de gueules journaliers dans les derniers moments de thèse. Merci beaucoup pour cet agréable soutien. J'en profite pour remercier ses coloc pour leur accueil lors de mes passages à Paris : Iro, Marine, Rémy et William.

La thèse, c'est aussi l'entourage, l'échappatoire, le défouloir, celui qui permet de relativiser les soucis et doutes du doctorant catastrophé devant un calcul incompréhensif ou un graphique qu'il ne trouve pas à son goût. Je voudrais remercier tous ses gens qui m'ont si souvent remis les pieds sur terre.

Je tiens à remercier l'équipe de hand de Lyon9, mes partenaires d'entraînement, de victoires, de défaites et surtout de pots en tous genres. Vous me manquerez à la capitale. Je remercie donc Sylvain, Alex, Hervé, François le métronome et Clarisse (je pense qu'on peut considérer que tu fais partie de l'équipe), Maxime le droitier, Maxime le gaucher clairvoyant, coach Jeson, ex-coach Bruno, Seb, Onofre qui a mis la barre très haute pour mon futur voisin, Ludo et son éternel sourire, David la mobylette, Guigui, Tomtom, Thomas, Pierre, Philippe, Bernard, Camille, Boris et capitaine Alain (merci aussi à Morgane pour ses belles oeuvres d'art, elles trouveront une place de choix dans mon futur appart). S'il vous plaît les gars, faites en sorte de monter avant la retraite de notre capitaine...

Je remercie également ceux qui ont fait que la pause repas du midi a été un des moments clé de la journée, ceux qui m'ont fait aimer le café. Merci à Pad l'homme à tout faire, Micho, Michel, Manu à qui le costume va si bien et David que je suis bien content de rejoindre à Paris.

Merci aussi à Mme Manassès, qui m'a fait aimer la physique.

Merci aux potes connus plus tôt et toujours présents. Merci à toi Hélène. Je tiens à remercier Ludo et Clo, que je pourrais venir voir plus souvent à présent. Je remercie également Pete, Elise, Gregory, Anne, Bertrand, Samuel, John, Steven, Céline, Delphine. Je tiens à remercier Oliv toujours là quand il faut et Gillian, Yo, Steph et Lies, Aurèl et Anne, et Fabien. Je remercie également Florian qui m'a fait découvrir la physique des particules et Camille. J'aimerais également remercier les bretons que je retrouve toujours avec autant de plaisir : Nico et Sosso, Stan, Micky, Erwan, Florian, François, Ronan, Onenn, Adeline. Merci aussi à Manu pour son amitié.

Pour finir, j'aimerais remercier de tout mon coeur ceux qui ont toujours été là, les soutiens inconditionnels, en un mot la famille. Je remercie sincèrement mes parents. Merci à ma mère pour son soutien téléphonique, rendu parfois difficile au vue de mon manque de discussion. Remerciements spéciaux à mon père qui a réussi l'incroyable performance de lire le manuscrit en un jour et

de corriger toutes les fautes d'orthographe que je n'ai pas ajoutées depuis. Merci enfin à Manon, François et Elise.

Introduction

La physique des particules s'appuie depuis des décennies sur le modèle standard pour expliquer les phénomènes intervenant dans le monde subatomique. Ce modèle s'attache à décrire les constituants élémentaires de la matière, ainsi que les interactions fondamentales de la nature. Fort de ses confirmations expérimentales, il offre ainsi une description parfaite de la physique jusqu'aux énergies actuellement accessibles dans les collisionneurs de particules. Jamais directement mis en défaut, il doit une grande partie de son succès à son côté prédictif. Néanmoins, des considérations théoriques et quelques observations expérimentales inexpliquées à ce jour laissent à penser que le modèle standard n'est qu'une théorie effective valable, en première approximation, aux énergies actuellement accessibles. Il existe de nombreux modèles théoriques qui proposent une extension au modèle standard. On regroupe une partie de ces modèles dans la famille, dite des modèles supersymétriques, qui séduisent notamment par leur élégance théorique. Ces modèles prédisent une augmentation du nombre de particules élémentaires et, de ce fait, offrent de nombreuses possibilités de recherche pour les expériences de physique des particules. Le modèle standard et son extension supersymétrique sont décrits dans le premier chapitre. Ce chapitre motive alors la recherche de particules supersymétriques auprès des collisionneurs.

L'objet de cette thèse est la recherche de particules supersymétriques, appelées gluino ($\tilde{g}$), sbottom ($\tilde{b}_1$) et neutralino ($\tilde{\chi}_1^0$), au TeVatron grâce au détecteur DØ. Le TeVatron est un collisionneur protons-antriprotons qui fonctionne avec une énergie de 1.96 TeV dans le centre de masse. Le deuxième chapitre donne les principaux détails du dispositif expérimental permettant la production de faisceaux de protons et d'antiprotons ainsi que leur accélération à l'énergie voulue. Ce chapitre décrit également le fonctionnement du détecteur DØ, qui enregistre les collisions ensuite utilisées pour les analyses de physique.

Le troisième chapitre présente les principaux ingrédients nécessaires à la recherche des particules supersymétriques citées précédemment. Il s'applique notamment à décrire le système de déclenchement. Ce dispositif est utilisé pour trier les collisions en ligne, de façon à n'enregistrer que les collisions considérées comme intéressantes pour les analyses de physique. Il est indispensable d'un point de vue pratique et technologique, tant il est impossible d'enregistrer la totalité des collisions produites par le TeVatron. La conception des critères de sélection, souvent amenés à être rapidement modifiés, est une étape primordiale en amont des recherches, telles les recherches de nouvelle physique. Ce troisième chapitre détaille également le traitement des données, simulées ou réelles, qui permet une construction des objets physiques indispensables à la compréhension des collisions enregistrées par le détecteur.

Une fois ce travail préliminaire effectué, il est possible d'analyser les données collectées et de rechercher la présence éventuelle de déviations par rapport aux prédictions théoriques du modèle standard. Le dernier chapitre s'intéresse à la recherche de telles déviations afin de mettre en évidence la présence du signal supersymétrique recherché. Ce signal est, en fait, une cascade de désintégration conduisant à un état final particulier et rare pour les processus prédits par le modèle

standard :

$$p\bar{p} \rightarrow \tilde{g} + \tilde{g} \rightarrow \tilde{b}_1 + \bar{b} + \bar{\tilde{b}}_1 + b \rightarrow 2b + 2\bar{b} + 2\tilde{\chi}_1^0$$

Aucune analyse de ce type n'a actuellement été menée par la collaboration DØ, une analyse similaire a été entreprise par la collaboration CDF. La totalité des données utilisées pour l'analyse présentée dans cette thèse couvre la période s'étalant d'avril 2002 au printemps 2006, ce qui représente une luminosité intégrée d'environ 1 fb^{-1} (cette notion est expliquée dans le deuxième chapitre). L'ensemble des techniques employées pour isoler le signal des bruits de fond du modèle standard et le résultat final de cette recherche sont détaillés dans ce dernier chapitre, qui conclura le travail de thèse.

Chapitre
1

Cadre théorique

Sommaire

1.1	Introduction	14
1.2	Les symétries en physique des particules	15
1.2.1	Notions de bases	15
1.2.2	Les symétries externes	15
1.2.3	Les symétries internes	16
1.2.4	Les brisures des symétries	17
1.3	Le modèle standard de la physique des particules	18
1.3.1	Introduction	18
1.3.2	L'électrodynamique quantique : point de départ d'une théorie aboutie . .	18
1.3.3	La force faible et l'interaction électrofaible	20
1.3.4	Le mécanisme de Higgs	23
1.3.5	La chromodynamique quantique	25
1.3.6	Quelques mots sur la renormalisation d'une théorie	29
1.3.7	Les réussites du modèle standard	30
1.3.8	Le contenu en champs	31
1.3.9	Les limites du modèles standard et quelques exemples d'extension	32
1.4	La supersymétrie	35
1.4.1	Introduction et motivations	35
1.4.2	L'algèbre supersymétrique	36
1.4.3	Le lagrangien supersymétrique	38
1.4.4	Le modèle standard supersymétrique minimal (MSSM)	42
1.4.5	Conclusion	48

1.1 Introduction

La physique des particules et le cadre théorique qui l'accompagne ont vu leur avènement tout au long du XX^e siècle. Néanmoins, les motivations premières de cette quête de l'infiniment petit remontent à l'antiquité. Les penseurs grecs, déjà, cherchaient à comprendre la diversité des formes et des couleurs qui nous entourent. Au IV^e siècle av. J.C., les philosophes Leucippe et Démocrite émettent l'hypothèse que la matière est constituée de particules minuscules : les atomes (du grec *atomos* qui signifie indivisible). Ils ont ainsi voulu définir des briques élémentaires de la matière. C'est finalement le modèle d'Aristote basé sur quatre éléments : "l'eau, l'air, la terre et le feu", qui avait été préféré. Leur objectif était de déterminer quels constituants élémentaires pourraient être à la base de tout le reste. Cette vision de briques élémentaires se retrouve dans de nombreuses théories physiques qui ont suivi jusqu'au modèle standard de la physique des particules. En plus de ces briques, il faut être à même de comprendre les forces qui interviennent entre ces constituants élémentaires et ceci est motivé, encore une fois, par la diversité des structures proposées par la Nature. Ces motivations ont conduit Boyle à une définition moderne des éléments en 1661, et Avogadro et Dalton à proposer un modèle atomique en 1860. Auparavant, Newton apporta une première théorie de la gravitation en 1687 qu'Einstein bouleversa avec sa théorie de la relativité générale en 1915 [1]. Maxwell a quant à lui permis la compréhension des interactions électromagnétiques en proposant des équations les décrivant en 1860. Le modèle de l'atome est alors modifié à la suite d'une série de découvertes : l'électron en 1897 par Thomson, la mise en évidence du noyau atomique par Rutherford en 1911 et la compréhension de l'existence du proton les années suivantes, l'hypothèse du neutrino par Pauli en 1930 et, comme dernier exemple, la découverte du neutron par Chadwick en 1932 [2]. L'ensemble des découvertes de ces nouvelles particules et des forces régissant leur comportement a ainsi amélioré notre perception des phénomènes microscopiques. C'est finalement la naissance de deux grandes théories qui a permis de comprendre et de donner un début d'explication à ces phénomènes. La relativité restreinte, proposée par Einstein en 1905 [3], et la mécanique quantique dont les fondements ont été posés entre 1924 et 1927, ont ouvert les portes de la physique de l'infiniment petit et marqué le début de la physique des particules. La physique des particules s'attache à décrire les constituants élémentaires de la matière et les interactions fondamentales qui sont par ordre d'intensité croissante :

- les interactions gravitationnelles, responsables de la chute des objets et du mouvement des planètes ;
- les interactions faibles à l'origine des désintégrations radioactives ou des processus de fusion ;
- les interactions électromagnétiques, responsables de la structure des atomes et des molécules ;
- les interactions fortes, à l'origine de la cohésion de la matière nucléaire.

La description de ces interactions est intimement liée à la découverte des particules élémentaires. Les trois dernières de la précédente liste sont décrites avec une grande précision dans le même cadre théorique : le modèle standard détaillé dans la Section 1.3. Ce cadre théorique basé sur les symétries (voir la Section 1.2) fut développé autour de 1970. Les prédictions théoriques et les observations expérimentales s'accordent de façon spectaculaire. C'est ce caractère prédictif qui a contribué au succès de ce modèle. Aucun résultat expérimental ne l'a remis en cause malgré les nombreuses questions ouvertes, insolubles à ce jour (voir la Section 1.3.9). Elles conduisent à penser que le modèle standard est une théorie effective à basse énergie d'une théorie plus globale. De nombreuses théories ont été proposées dans ce sens ; la Supersymétrie (Section 1.4) est l'une d'entre elles. Activement recherchée dans les collisionneurs actuels et passés, aucune trace de cette dernière n'a pour l'instant été décelée. Le *TeVatron*, décrit dans le Chapitre 2, contribue à

repousser encore plus loin les frontières en énergie de cette éventuelle théorie plus globale. Les espoirs se portent alors vers le futur collisionneur basé au *CERN* à Genève : le *Large Hadron Collider* (*LHC*), pour découvrir ses nouvelles particules tant attendues et continuer à améliorer la description de la matière.

1.2 Les symétries en physique des particules

Cette section [4, 5] s'attache à décrire l'importance des symétries dans la théorie de la physique des particules. Les détails mathématiques, essentiellement reliés à la théorie des groupes ou à la théorie des champs, seront volontairement limités pour ne garder que les bases nécessaires à la compréhension des sections suivantes. Les symétries peuvent être classées en deux catégories. Les symétries externes sont des transformations de l'espace-temps qui affectent le système de coordonnées, alors que les symétries internes sont des transformations abstraites des champs de particules. En règle générale, une symétrie est une transformation mathématiques des objets de la théorie ne modifiant pas ses prévisions physiques.

1.2.1 Notions de bases

L'ensemble des symétries d'une théorie forme un groupe : groupe de symétrie de la théorie. Les propriétés mathématiques de ce groupe permettent de décrire directement la physique prédite par la théorie. La physique des particules s'appuie sur une théorie quantique des champs. Les particules sont ainsi associées à des champs d'opérateurs sur un espace vectoriel abstrait (espace de Hilbert). L'évolution du système physique, c'est-à-dire d'un vecteur de l'espace vectoriel, est décrit par un lagrangien contenant ces champs d'opérateurs. A chaque symétrie de la théorie, on associe un opérateur linéaire.

1.2.2 Les symétries externes

La physique des particules se place actuellement dans le cadre de la relativité restreinte [3]. L'espace-temps est donc l'espace à quatre dimensions de Minkowski (des théories de grandes unifications, non-développées dans cette thèse, utilise plutôt le cadre de la relativité générale [1]) et le groupe de symétrie associé est le groupe de Poincaré. Ce groupe contient les transformations suivantes :

- les translations dans l'espace-temps ;
- les rotations dans l'espace à trois dimensions ;
- les transformations ou *boost* de Lorentz ;
- les inversions spatiales ou temporelles.

La théorie doit donc être invariante sous l'action de chacune de ces transformations. De plus, le théorème de Noether [6] montre qu'à toute symétrie on peut associer une quantité conservée. Par exemple, l'invariance suivant les translations se traduit par une conservation de l'impulsion et l'invariance par les rotations se traduit par une conservation du moment cinétique. On peut définir un ensemble minimal de vecteurs (une représentation irréductible du groupe) stables sous les transformations de symétrie. Ces vecteurs correspondent aux états physiques du système étudié, les transformations représentent des changements de référentiels. L'étude des représentations irréductibles d'un groupe permet alors de définir les propriétés physiques des différents états du

système. A la suite de cette étude, on peut montrer qu'il existe deux cas physiques intéressants selon la valeur de leur impulsion au carré :

- si $p^2 > 0$, ce sont des vecteurs de type temps caractérisés par leur masse, leur spin (demi-entier) et leur parité (1 ou -1) ;
- si $p^2 = 0$, ce sont des vecteurs de type lumière de masse nulle caractérisés par leur hélicité λ ($\lambda = \frac{n}{2}$ avec n un nombre relatif).

L'étude des symétries externes à travers les représentations irréductibles du groupe de Poincaré a permis de proposer une première description des particules massives ou non, qui correspond parfaitement aux observations expérimentales. De plus, la théorie des champs devant également vérifiée les symétries du groupe de Lorentz (composé des trois dernières transformations de la liste précédente), trois types de champs peuvent intervenir : scalaire, vectoriel ou spinoriel. Cette description peut ensuite être complétée par l'étude des symétries internes.

1.2.3 Les symétries internes

Plusieurs types de champs, décrits dans la section précédente, peuvent intervenir dans le lagrangien. Des symétries plus abstraites peuvent transformer ces champs sans modifier le lagrangien, ce sont les symétries internes. Par exemple, si un champ ψ est transformé selon :

$$\psi \rightarrow \psi' = U\psi \quad (1.1)$$

L'invariance selon la symétrie U s'exprime de la façon suivante :

$$\mathcal{L}'(\psi', \partial_\mu \psi') = \mathcal{L}(\psi, \partial_\mu \psi) \quad (1.2)$$

D'après le théorème de Noether, on peut associer à chacune des symétries des courants conservés et une charge conservée. Des exemples de ces quantités conservées dans le cadre du modèle standard seront détaillés dans la Section 1.3.

Invariance de jauge globale

Si les paramètres de la transformation ne dépendent pas de l'espace-temps, on dit qu'on a une transformation de jauge globale. Dans le cas de l'électrodynamique quantique, l'invariance de jauge globale par transformation de la phase conduit, par exemple, à montrer la conservation de la charge électrique (voir la Section 1.3.2). Cette propriété n'avait pas été mise en évidence par l'étude des symétries externes.

Invariance de jauge locale

L'invariance de jauge locale, développée par Weyl en 1929 [7], est plus complexe et intervient quand les paramètres de la transformation dépendent de l'espace-temps. Les termes dérivatifs du lagrangien ne permettent généralement pas d'avoir un lagrangien invariant selon ce type de transformation. Ce problème peut être résolu en introduisant de nouveaux champs, champ de jauge, en remplaçant la dérivée ∂_μ par la dérivée covariante $D_\mu = \partial_\mu - igX_\mu$ et en ajoutant des termes invariants (de jauge et de Lorentz) pour décrire la dynamique des nouveaux champs de jauge. Des détails seront donnés dans la Section 1.3.2 en développant l'invariance de jauge locale dans le cas de l'électrodynamique quantique. Les conséquences physiques de ces développements mathématiques sont :

- la construction d’une théorie avec des champs non-uniformes dans l’espace-temps ;
- la prédiction de nouvelles particules, appelées bosons de jauge, qui engendrent des interactions entre les particules ;
- la description de la dynamique de ces bosons de jauge au sein du même lagrangien.

Les symétries internes permettent donc une description des interactions entre les particules. Elles jouent un rôle fondamental dans les théories de jauge.

1.2.4 Les brisures des symétries

Finalement, les symétries peuvent être brisées. On différencie alors deux types de brisure :

- les brisures explicites quand le lagrangien n’est pas invariant après la transformation. C’est le cas pour la symétrie globale $SU(2)$ décrite dans la Section 1.3.5. Les différences de masses entre les quarks brisent cette symétrie.
- Les brisures spontanées correspondent à des lagrangiens invariants sous la transformation mais dont l’état fondamental n’est pas symétrique. Un exemple d’une telle brisure est explicité dans la Section 1.3.4. La Figure 1.1 illustre ce cas en représentant un potentiel symétrique selon l’axe $Im(\Phi) = 0$, $Re(\Phi) = 0$ alors que la position au repos décrit par ce potentiel ne respecte pas cette symétrie.

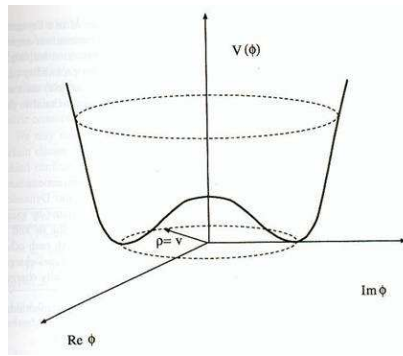


FIG. 1.1 – Potentiel en forme de chapeau mexicain brisant spontanément la symétrie selon l’axe $Im(\Phi) = 0$, $Re(\Phi) = 0$.

Conclusion

La notion de symétrie, transformation mathématique complexe, explique de nombreux phénomènes physiques. Elle peut décrire, par exemple, des concepts fondamentaux comme : la charge, le spin, la parité ou la masse d’une particule. Elle donne également une description des interactions entre les particules. Elle peut même avoir un caractère prédictif dans le cas des bosons de jauge comme il sera expliqué dans la suite. Finalement, les différents concepts qui ont été évoqués, sans être détaillés dans cette section, sont à l’origine de la construction du modèle standard, qui permet d’expliquer, à lui seul, la quasi-totalité des observations expérimentales en physique des particules.

1.3 Le modèle standard de la physique des particules

1.3.1 Introduction

Le modèle standard repose sur des principes d'invariance de jauge locale, c'est-à-dire des transformations dont les paramètres dépendent des coordonnées de l'espace-temps (voir la Section 1.2). En décrivant le lagrangien d'un électron libre, Dirac posa en 1927 les fondements de ce futur modèle [8, 9]. A partir de ce lagrangien, Tomonaga, Schwinger [10] et Feynman [11] utilisèrent le principe de jauge locale pour développer l'électrodynamique quantique entre 1948 et 1949. Les détails seront donnés dans la Section 1.3.2 ; ils permettront d'illustrer les propos de la Section 1.2. Cette description de l'électron grâce à la théorie quantique des champs a constitué une première étape importante. L'invariance selon une transformation d'un élément de la symétrie $U(1)$ fut ensuite généralisée par Yang et Mills en 1954 [12] en utilisant le groupe de symétrie $SU(2)$. Cette théorie de Yang-Mills permet de décrire l'interaction faible et de l'unifier à l'électromagnétisme sous le nom d'interaction électrofaible (voir la Section 1.3.3). La symétrie utilisée pour cette interaction est brisée à l'échelle électrofaible ; le mécanisme de Higgs [13] apporte une solution à ce problème comme il sera explicité dans la Section 1.3.4. Finalement, le développement de la chromodynamique quantique (voir la Section 1.3.5) conclut la description du modèle standard comme il est connu aujourd'hui. Toutes les interactions connues, exceptée la gravitation, sont décrites dans le même cadre théorique. Ce modèle permet d'identifier par la même occasion les constituants élémentaires de la matière. Leurs caractéristiques seront données dans la Section 1.3.8.

1.3.2 L'électrodynamique quantique : point de départ d'une théorie aboutie

Le point de départ de cette théorie est le lagrangien de Dirac qui décrit le mouvement d'un électron libre. Ainsi, si $\psi(x)$ est le champ (ou bi-spineur de Dirac) correspondant à un électron de masse m et dépendant de $x = (t, \vec{r})$, on a le lagrangien suivant :

$$\mathcal{L}_D = \bar{\psi}(x)(i\gamma^\mu \partial_\mu - m)\psi(x) \quad (1.3)$$

Les termes γ^μ correspondent aux matrices de Dirac, c'est-à-dire :

$$\gamma^0 = \begin{pmatrix} I_2 & 0 \\ 0 & -I_2 \end{pmatrix}, \gamma^i = \begin{pmatrix} 0 & \sigma^i \\ -\sigma^i & 0 \end{pmatrix} \quad (1.4)$$

Et les termes σ^i sont les matrices de Pauli, c'est-à-dire :

$$\sigma_0 = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} = I_2, \sigma_x = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}, \sigma_y = \begin{pmatrix} 0 & -i \\ i & 0 \end{pmatrix}, \sigma_z = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} \quad (1.5)$$

Les matrices de Dirac ont été obtenues pour trouver une équation linéaire comme celle de Klein-Gordon. La différence entre ces deux équations est le spin des particules décrites : celle de Klein-Gordon est utilisée pour les particules de spin 0 alors que celle de Dirac (voir l'Equation 1.3) est utile aux particules de spin 1/2. Ce dernier lagrangien est invariant sous transformation d'un élément de la symétrie $U(1)$ globale :

$$\psi(x) \rightarrow \psi'(x) = e^{-i\alpha}\psi(x) \quad (1.6)$$

$$\bar{\psi}(x) \rightarrow \bar{\psi}'(x) = e^{i\alpha} \bar{\psi}(x) \quad (1.7)$$

Avec α un réel quelconque.

Si on veut maintenant que le lagrangien reste invariant sous la transformation locale ou de jauge, i.e. si le paramètre α dépend de l'espace-temps : $\alpha(x)$, il va falloir introduire des termes supplémentaires comme il a été évoqué dans la Section 1.2. En effet, si le terme de masse est invariant sous la transformation de jauge, le terme avec la dérivée ne l'est pas. Pour remédier à ce problème, on introduit une dérivée covariante possédant la propriété suivante sous une transformation locale :

$$D_\mu \psi(x) \rightarrow e^{-i\alpha(x)} D_\mu \psi(x) \quad (1.8)$$

Cette dérivée covariante est obtenue en introduisant un champ vecteur (un champ de jauge) $a_\mu(x)$ de la façon suivante :

$$D_\mu \psi(x) = (\partial_\mu + ie a_\mu(x)) \psi(x) \quad (1.9)$$

Avec e une constante.

Ce champ de jauge doit se transformer sous une transformation du groupe $U(1)$ selon :

$$a_\mu(x) \rightarrow a'_\mu(x) = a_\mu(x) + \frac{1}{e} \partial_\mu \alpha(x) \quad (1.10)$$

Le champ de jauge n'étant pas un champ dynamique, il faut ajouter un terme cinétique pour rendre ce champ physique. Il faut, de plus, que le terme cinétique soit invariant de jauge pour conserver la symétrie locale du lagrangien. On introduit alors le terme :

$$f_{\mu\nu}(x) = \partial_\mu a_\nu(x) - \partial_\nu a_\mu(x) \quad (1.11)$$

On obtient finalement le lagrangien de l'électrodynamique quantique (QED) :

$$\mathcal{L}_{QED} = \bar{\psi}(x)(i\gamma^\mu D_\mu - m)\psi(x) - \frac{1}{4} f_{\mu\nu}(x) f^{\mu\nu}(x) \quad (1.12)$$

On constate qu'il n'y a pas de terme de masse pour le champ a_μ . En effet, Un terme de masse briserait l'invariance de jauge. En fait, le champ a_μ correspond au potentiel que l'on trouve dans les équations de Maxwell, qui n'est autre que le photon de l'interaction électromagnétique. Le photon est donc sans masse. Si on développe le lagrangien de l'Equation 1.12, on obtient un terme cinétique pour l'électron, un terme cinétique pour le photon, un terme de masse pour l'électron et enfin un terme décrivant l'interaction entre un photon et un électron : $e\bar{\psi}(x)\gamma^\mu\psi(x)a_\mu(x)$. L'interprétation proposée par Feynman pour chacun de ces termes est la propagation d'une particule chargée, la propagation d'un photon et l'absorption ou l'émission d'un photon par une particule chargée. Ces différents phénomènes physiques sont résumés dans les diagrammes de Feynman (1947), qui sont bien plus que de simples schémas. En fait, ils correspondent au développement mathématique d'un processus physique donné et sont des aides visuelles pour le calcul des termes du développement perturbatif. Un exemple d'un tel calcul est représenté sur la Figure 1.2. Les lignes externes sont les particules initiales et finales du processus étudié, alors que les intersections entre les lignes sont les vertex de l'interaction. Dans le cas de la QED , le seul vertex possible fournit par le lagrangien de l'Equation 1.12 est un vertex à trois pattes avec deux électrons (ou fermions si on généralise à toutes les particules de spin 1/2) et un photon. Pour avoir la probabilité totale d'un tel processus, il faut sommer les probabilités de toutes les configurations possibles donnant le même résultat final en partant du même état initial. L'ordre de grandeur des différentes contributions est directement

lié au nombre de vertex. Ainsi, le calcul est souvent effectué grâce à la théorie des perturbations. On calcule alors la quantité voulue à l'ordre le plus bas (contribution due aux diagrammes de Feynman avec le minimum de vertex) et on corrige le calcul avec les ordres suivants en fonction de la précision souhaitée.

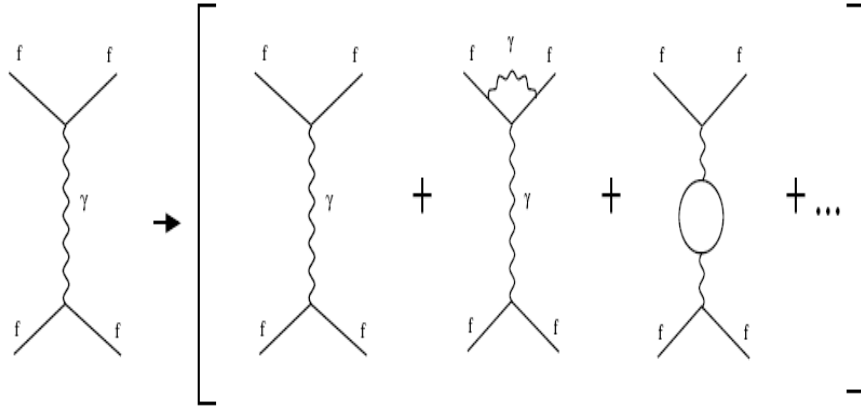


FIG. 1.2 – Exemple de développement perturbatif d’une interaction électromagnétique avec dans l’état initial deux fermions.

Cette théorie a rencontré des problèmes pour le calcul de certains diagrammes de Feynman, qui comportent des boucles fermioniques (voir par exemple le quatrième diagramme de la Figure 1.2). La résolution de ce problème ne sera pas détaillée ici. En effet, une procédure complexe a été mise au point en 1948 par Feynman pour corriger les divergences du calcul : la renormalisation.

En détaillant le modèle de l’électrodynamique quantique, on montre qu’en rendant locale l’invariance par transformation de phase, on couple l’électron au photon, c’est-à-dire à un champ de spin 1 sans masse. Cette théorie reproduit parfaitement les observations expérimentales comme la diffusion Compton ou le *Bremsstrahlung*, par exemple. Cette théorie a permis également de mettre au point des techniques de calcul très poussées schématisées par les diagrammes de Feynman. Finalement, deux mesures de précision ont accentué le succès de la *QED* : une description de la structure fine de l’atome d’hydrogène (*shift* de Lamb) et la mesure du moment magnétique anomal de l’électron par Kusch. Dans ce dernier cas, la théorie est vérifiée à dix décimales près !

1.3.3 La force faible et l’interaction électrofaible

Historique

L’interaction faible est une interaction à faible portée (de l’ordre de 10^{-18} m) dont l’action est confinée à l’intérieur des noyaux atomiques, d’où sa découverte tardive. C’est Becquerel en 1896 qui mit le premier en évidence la radioactivité. Plusieurs types de radioactivité sont connues, l’une d’entre elles a longtemps posé problème : la radioactivité β . Ce processus concerne l’expulsion d’un électron hors d’un noyau atomique (ce qui correspond en fait à la désintégration d’un neutron). Or, contrairement aux prédictions théoriques du modèle de Rutherford, le spectre d’énergie de l’électron est continu, ce qui suggère la présence d’une autre particule indétectée. La solution fut trouvée par Pauli en 1930 qui postula l’existence d’une particule neutre interagissant très peu

avec la matière : le neutrino. La radioactivité β correspond donc au processus de désintégration : $n \rightarrow p^+ e^- \nu_e$. Deux découvertes confirmèrent l'hypothèse de Pauli : la découverte du neutron en 1932 [2] et la découverte de la radioactivité β^+ avec émission du positron en 1933 par I. Curie et F. Joliot. En 1934, Fermi proposa une première théorie pour expliquer ces deux types de désintégration [14] : la théorie $V - A$. Cette théorie explique correctement les observations expérimentales mais n'est qu'une théorie effective à basse énergie. En 1958, la découverte des neutrinos par Cowan et Reines [15] confirme l'hypothèse de Pauli et la théorie de Fermi à basse énergie. Ce sont finalement Glashow [16], Weinberg et Salam [17] qui proposèrent en 1961 une théorie basée sur la théorie de Yang-Mills. Cette théorie permet alors d'expliquer à la fois l'interaction électromagnétique et l'interaction faible par l'échange des bosons de jauge Z^0 et $W^{+/-}$.

Théorie de Yang-Mills

La théorie de l'électrodynamique quantique peut se généraliser en utilisant un autre groupe de symétries qui comportent l'ensemble des rotations de phase où la phase est une matrice :

$$\psi(x) \rightarrow \psi'(x) = S\psi(x) \quad (1.13)$$

Avec S une matrice spéciale unitaire appartenant au groupe $SU(2)$ et $\psi(x)$ un doublet (champ à deux composantes).

Comme dans le cas de la QED , on demande l'invariance de jauge par rapport aux rotations locales du groupe $SU(2)$: $S(x) = e^{-i\theta^a(x)\frac{\sigma^a}{2}}$ avec σ^a les matrices de Pauli définies précédemment et $\theta^a(x)$ des nombres réels dépendant des coordonnées de l'espace-temps. L'invariance de jauge impose l'introduction d'une dérivée covariante (qui ne sera pas explicitée ici) comprenant trois bosons de jauge. Contrairement au cas de la QED , les bosons de jauge sont chargés par rapport à la charge de $SU(2)$, alors que le photon est neutre par rapport à la charge de $U(1)$ (la charge électrique). Ceci s'explique mathématiquement par le fait que le groupe $SU(2)$ n'est pas comme le groupe $U(1)$ un groupe abélien (commutatif). Ces trois bosons de jauge sont nommés W^+ , W^- et W^0 . Il faut finalement ajouter des termes cinétiques invariants sous une transformation de $SU(2)$ pour rendre physique ces bosons. Une autre différence par rapport à une théorie abélienne est la présence de termes d'auto-interaction entre les bosons d'échange. En effet, cette théorie prédit des vertex composés exclusivement des trois bosons W^+ , W^- et W^0 . Le lagrangien de Yang-Mills, ainsi défini, permet de décrire une théorie invariante de jauge sous les transformations du groupe $SU(2)$. C'est ce lagrangien qui est le point de départ de la description de l'interaction faible.

Modèle de Glashow, Weinberg et Salam

Pour expliquer la désintégration du neutron observée dans la radioactivité β , Glashow dans les années 50 puis Weinberg et Salam dans les années 60 proposent un modèle unifiant la force électromagnétique et la force faible : le modèle standard électrofaible.

Plusieurs études dédiées aux désintégrations des hadrons étranges (ensemble de quarks contenant un quark s) ont permis de mettre en évidence que contrairement à l'interaction électromagnétique (ou l'interaction forte décrite dans la Section 1.3.5), l'interaction faible ne conserve pas la parité P (symétrie par rapport à $\vec{r} \rightarrow -\vec{r}$) et ne conserve pas non plus la conjugaison de charge C . La violation de la symétrie CP est de plus largement étudiée dans la physique des mésons B (ensemble de quarks contenant des quarks b) comme dans l'expérience $BaBar$ par exemple. On

distingue alors les particules d'hélicité droite des particules d'hélicité gauche sous l'interaction faible.

Le groupe utilisé pour décrire l'interaction faible est le groupe $SU(2)$. Les particules interagissant sous cette interaction sont donc classées en doublets ou en singlets selon la valeur du nombre quantique associé à cette interaction : l'isospin T_3 . Les particules d'hélicité gauche sont classées en doublet d'isospin $T_3 = \pm\frac{1}{2}$. Et les particules d'hélicité droite sont classées en singlet d'isospin $T_3 = 0$. On trouve par exemple le doublet $(e^-, \nu_e)_L$ et les singlets $e_R^-, \nu_{eR}, \dots$. Les dernières expériences concernant la physique du neutrino prédisent une masse pour cette particule, le neutrino droit, initialement non décrit par le modèle standard, a volontairement été ajouté ici.

Glashow, Weinberg et Salam s'appuient sur les travaux de Yang et Mills et sur le modèle des quarks décrit dans la Section 1.3.5 pour écrire le lagrangien de leur modèle. Le groupe de symétries choisi pour unifier ces deux interactions est le groupe $SU(2)_L \otimes U(1)_Y$. L'indice L indique le couplage des bosons de jauge avec les particules d'hélicité gauche ou d'isospin $T_3 = \pm\frac{1}{2}$, l'indice Y correspond au nombre quantique désignant l'hypercharge Y . Si Q est le signe de la charge électrique, on peut montrer que les trois nombres quantiques vérifient : $Q = T_3 + \frac{Y}{2}$. La théorie de Yang-Mills prédit donc la présence de trois bosons de jauge dont un neutre pour l'invariance de jauge selon le groupe $SU(2)$. L'invariance de jauge selon le groupe $U(1)$ fournit un boson neutre supplémentaire, que l'on appelle B^0 dans le cas du modèle standard électrofaible (le photon est en fait une combinaison des deux bosons neutres, c.f. Section 1.3.4).

Finalement, le modèle standard électrofaible prédit quatre bosons de jauge, deux chargés et deux neutres. Les bosons W^+ , W^- et W^0 sont couplés au courant faible d'isospin avec une constante de couplage appelée g_2 , le boson B^0 est couplé au courant faible d'hypercharge avec une constante g_1 . Ce modèle décrit les interactions électromagnétique et faible des fermions dans un même cadre théorique. Le Tableau 1.1 donne les nombres quantiques associés à deux doublets de fermions. On peut y vérifier la relation $Q = T_3 + \frac{Y}{2}$.

Fermions	T_3	Y	Q
e^-	$-\frac{1}{2}$	-1	-1
ν_e	$\frac{1}{2}$	-1	0
d	$-\frac{1}{2}$	$\frac{1}{3}$	$-\frac{1}{3}$
u	$\frac{1}{2}$	$\frac{1}{3}$	$\frac{2}{3}$

TAB. 1.1 – Nombres quantiques associés aux doublets $(e^-, \nu_e)_L$ et $(u, d)_L$.

La désintégration du neutron qui a lieu dans les processus de radioactivité β correspond donc à l'émission d'un boson W^- comme indiqué sur le schéma de la Figure 1.3. Le neutron est en effet formé d'un triplet de quarks (u,d,d) alors que le proton est formé du triplet (u,u,d).

A ce niveau de description de l'interaction électrofaible, le lagrangien ne contient pas de terme de masse pour les bosons de jauge ou pour les fermions contrairement aux observations expérimentales. De tels termes brisent en effet l'invariance du lagrangien selon les transformations du groupe $SU(2)_L \otimes U(1)_Y$. Il faut donc être en mesure d'expliquer la masse des particules impliquées dans les interactions faibles. Le mécanisme de Higgs fournit une solution pour résoudre cette incompatibilité entre la théorie et l'expérience.

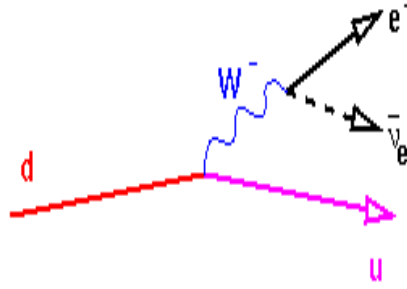


FIG. 1.3 – Désintégration du quark d en quark u par interaction faible avec émission d'un boson W^- .

1.3.4 Le mécanisme de Higgs

En 1967, Higgs [13] propose un mécanisme de brisure spontanée de symétrie (voir la Section 1.2) pour compléter la description de l'interaction électrofaible.

L'idée est d'introduire dans le lagrangien de l'interaction électrofaible un potentiel V , qui est une fonction des doublets décrivant les fermions :

$$V(\Psi) = \mu^2 \Psi^\dagger \Psi + \lambda (\Psi^\dagger \Psi)^2 \quad (1.14)$$

Avec $\mu^2 < 0$ et $\lambda > 0$, et Ψ , un doublet de $SU(2)$ composé de deux champs scalaires :

$$\Psi = \begin{pmatrix} \psi^+ \\ \psi^0 \end{pmatrix} \quad (1.15)$$

Ce potentiel est invariant sous les transformations du groupe $SU(2)_L \otimes U(1)_Y$. On peut schématiser la forme de ce potentiel comme celle montrée sur la Figure 1.1.

La dérivée covariante du modèle standard électrofaible est la suivante :

$$D_\mu = \partial_\mu + \frac{ig_1}{2} B_\mu + \frac{ig_2}{2} W_\mu \quad (1.16)$$

Avec B_μ le champ de jauge de $U(1)$ et $W_\mu = \frac{1}{2} W_\mu^i \sigma_i$ sont les trois champs de jauge correspondant à $SU(2)$.

Le lagrangien, ainsi construit, est invariant sous toute transformation du groupe $SU(2)_L \otimes U(1)_Y$, ce qui n'est pas le cas des états de vide prédits par ce lagrangien. En effet, on peut démontrer que dans une théorie de jauge si les états de vide s'annulent sous l'action des générateurs de la symétrie alors ces états de vide sont eux-mêmes invariants sous les transformations du groupe. Or, le minimum du potentiel est atteint pour :

$$\langle \Psi \rangle_0 = \frac{1}{\sqrt{2}} \begin{pmatrix} 0 \\ v \end{pmatrix}, v = \sqrt{\frac{-\mu^2}{\lambda}} \quad (1.17)$$

Et ce minimum ne s'annule pas sous l'action de $T_i = \frac{1}{2} \sigma_i$ pour $SU(2)$, il ne s'annule pas non plus sous l'action de $Y = I_2$ pour $U(1)$. On parle alors de brisure spontanée de la symétrie. On

constate, par contre, que ce minimum s'annule sous l'action de $T_3 + \frac{1}{2}Y$, qui correspond à un générateur de la symétrie $U(1)$. Les états de vide sont donc invariants par changement de phase.

La symétrie étant brisée par l'ajout d'un potentiel, il s'agit maintenant de développer la théorie au voisinage du minimum pour examiner l'effet de ce potentiel sur les bosons de jauge et les fermions impliqués dans le lagrangien :

$$\Psi = \frac{v + h(x)}{\sqrt{2}} \begin{pmatrix} 0 \\ 1 \end{pmatrix} \quad (1.18)$$

Avec $h(x)$ un scalaire fonction des coordonnées de l'espace-temps. On effectue le développement tout en associant au générateur $T_3 + \frac{1}{2}Y$ un champ noté A_μ , qui est une combinaison linéaire des champs de jauge. Une fois le développement effectué et les combinaisons linéaires des champs de jauge bien choisies, on obtient les combinaisons de champs suivantes :

$$W_\mu^\pm = \frac{1}{2}(W_\mu^1 \pm iW_\mu^2) \quad (1.19)$$

$$A_\mu = W_\mu^3 \sin(\theta_w) + B_\mu \cos(\theta_w) \quad (1.20)$$

$$Z_\mu = W_\mu^3 \cos(\theta_w) - B_\mu \sin(\theta_w) \quad (1.21)$$

Avec θ_w , l'angle de Weinberg défini par :

$$\cos(\theta_w) = \frac{g_2}{\sqrt{g_1^2 + g_2^2}} \quad (1.22)$$

Les termes de masse correspondant aux combinaisons linéaires de champs précédentes sont les suivants :

- $\frac{1}{2}g_2v$ pour le boson W ;
- $\sqrt{\frac{g_1^2 + g_2^2}{2}}v$ pour le boson Z ;
- 0 pour le photon ;
- $(2\lambda)^{\frac{1}{2}}v$ pour la particule scalaire appelée boson de Higgs.

Ce modèle a l'avantage de donner une masse aux bosons de jauge et confirme que le photon est une particule non-massive. Mathématiquement, ceci est dû au fait que lorsqu'une symétrie est conservée (ici la symétrie selon les transformations du groupe $U(1)$) alors les champs associés aux générateurs de cette symétrie n'ont pas de masse (le champ A_μ). Néanmoins, il introduit une nouvelle particule, le boson de Higgs, et un nouveau paramètre libre au modèle : le paramètre v . De plus, les fermions acquièrent également une masse par couplage avec le boson de Higgs.

Le modèle standard électrofaible et le mécanisme de Higgs prédisent l'existence de courant neutre, l'existence de bosons de jauge massifs et l'existence d'une particule scalaire massive : le boson de Higgs. Si les deux premières prédictions sont parfaitement vérifiées par l'ensemble des observations expérimentales connues à ce jour, la troisième reste en suspens. Aucune expérience n'a pour l'instant montré de preuve directe de l'existence d'une telle particule scalaire, malgré les contraintes expérimentales et théoriques sur sa masse.

Confirmation du modèle standard électrofaible

Le boson de Higgs est aujourd'hui activement recherché par les collisionneurs comme le *TeVatron*. Les prédictions qui découlent du modèle standard électrofaible ont en effet été remarquablement vérifiées par les expériences de physique des particules.

En 1973, l'expérience *GARGAMELLE*, expérience de chambre à bulles, au *CERN* a découvert les courants neutres. La découverte des bosons Z et W au *CERN* également et la mesure de leur masse ont confirmé ensuite les prédictions. Finalement, cette théorie a valu à Glashow, Weinberg et Salam le prix Nobel de physique en 1979.

1.3.5 La chromodynamique quantique

Introduction

La chromodynamique quantique ou *QCD en anglais* pour *Quantum ChromoDynamics* est la théorie qui décrit les quarks et l'interaction forte. La découverte des quarks, particules élémentaires, constituants du proton ou du neutron par exemple, est le résultat d'une série d'expériences et d'observations physiques vers le milieu du *XIX^{ème}* siècle. La découverte du neutron en 1932 est une première étape importante dans la découverte de la sous-structure du noyau atomique. Ce noyau est formé de nucléons : les protons ou les neutrons. Les protons et les neutrons ont des masses très proches, mais tandis que le premier est chargé positivement, le second est neutre. Un premier modèle, construit par analogie avec la *QED*, postule qu'il existe une symétrie entre les nucléons et qu'ils interagissent en échangeant des mésons, particules encore hypothétiques à ce stade. Ce modèle est le fruit du travail de Yukawa [18]. Une découverte va ensuite soulever des interrogations sur ce modèle : la découverte du pion en 1947. S'ensuit alors une série de découvertes de particules hadroniques, c'est-à-dire de particules subissant l'interaction forte (voir la Figure 1.4 pour l'historique des découvertes et la Figure 1.5 pour la classification des particules élémentaires avant le modèle des quarks).

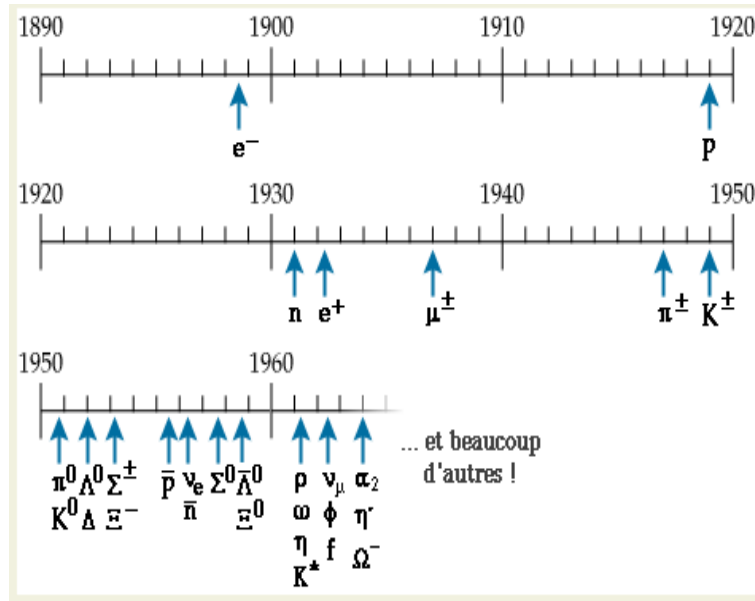


FIG. 1.4 – Historique des découvertes de particules jusqu'en 1964.

Photon (γ)		
Leptons		Electron (e^-), Neutrino (ν)
Hadrons	Mésons	Pi (π^+, π^0, π^-), K ($K^+, K^0, \bar{K}^0, K^-$), $\eta^0, \dots$
	Nucléons	Proton (p), Neutron (n)
	Baryons	Hypérons
		Delta ($\Delta^{++}, \Delta^+, \Delta^0, \Delta^-$), Sigma ($\Sigma^+, \Sigma^0, \Sigma^-$), Ksi (Ξ^0, Ξ^-), $\Lambda^0, \dots$

FIG. 1.5 – Classification des particules élémentaires dans les années 50.

Le modèle des quarks

Pour expliquer le nombre important de particules hadroniques, Gell-Mann et Nishijima proposent un nouveau nombre quantique, l'étrangeté. De la même manière que pour le modèle standard électrofaible, une théorie de jauge est alors progressivement mise au point. La symétrie $SU(3)$ est choisie en 1961 pour rendre compte de tous les hadrons observés. D'après Gell-Mann et Zweig [19], les hadrons ne sont plus des particules élémentaires mais plutôt des particules composites formées de quarks. Ils postulent l'existence de trois quarks se différenciant par leur saveur : u , d et s . On distingue alors deux familles de hadrons : les baryons composés de trois quarks et les mésons composés d'un quark et d'un anti-quark. Par exemple, un proton est composé de la combinaison de quarks uud alors qu'un neutron est composé de udd . Comme toutes les théories de jauge, la théorie basée sur l'invariance selon les transformations du groupe $SU(3)$ prédit l'existence de boson de jauge, vecteur de l'interaction forte. Ces bosons sont appelés les gluons et sont mis en évidence dans une expérience de diffusion électron-proton à Stanford. Cette théorie laisse malgré tout quelques interrogations :

- Pourquoi les baryons sont composés de trois quarks et les mésons d'un quark et d'un anti-quark ? Pourquoi n'est-il pas possible de trouver des quarks libres ou alors des baryons formés de quarks et d'anti-quarks ?
- Comment expliquer en 1951 la découverte de la particule Δ^{++} formée de trois quarks u ? Cette particule est en effet de spin $3/2$. Or, le principe de Pauli exclue la possibilité pour deux fermions de se retrouver dans le même état quantique.
- Comment expliquer enfin la différence entre les sections efficaces de production de paires $q\bar{q}$ et les sections efficaces de production de paires de muons ? Le rapport $R_{had} = \frac{\sigma(e^+e^- \rightarrow hadrons)}{\sigma(e^+e^- \rightarrow \mu^+\mu^-)}$ des sections efficaces évolue en effet en fonction de l'énergie comme le montre la Figure 1.6. Ce rapport devrait être constant et égal à la somme quadratique des charges électriques de quarks.

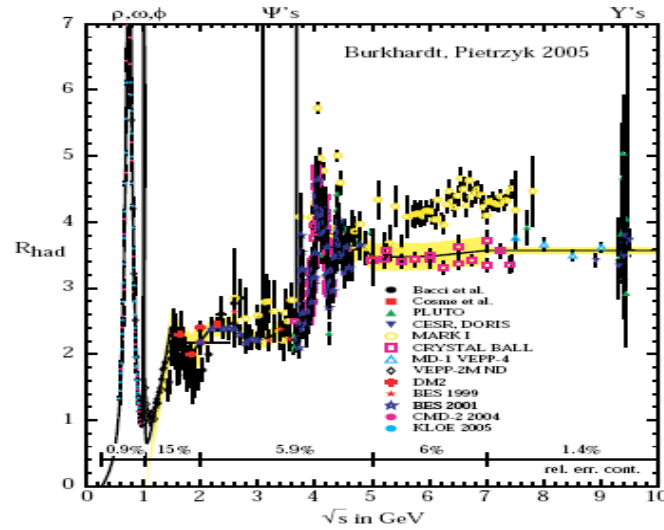


FIG. 1.6 – Le rapport R en fonction de l'énergie dans le centre de masse pour différentes mesures.

La QCD

La solution à ces trois problèmes est l'introduction d'un nouveau nombre quantique pouvant prendre trois valeurs pour les quarks : la couleur. Ainsi, la particule Δ^{++} ne contredit pas le principe d'exclusion de Pauli. La valeur du rapport R est comprise jusqu'à 3 GeV. Et, les hadrons sont des sommes "nulles" des trois couleurs, c'est-à-dire la somme d'une couleur et d'une anti-couleur ou la somme des trois couleurs. Par analogie avec les couleurs rouge, vert et bleu, les hadrons sont dits "blancs" de couleur.

Le modèle des quarks, ainsi défini, ne décrit pas l'interaction forte. La chromodynamique quantique, qui doit son nom à la charge de couleurs des quarks, doit vérifier les conditions suivantes pour répondre aux contraintes expérimentales :

- Contrairement au changement de saveur qui modifie la masse des quarks, le changement de couleurs ne modifie pas les quarks. La symétrie utilisée pour la QCD doit donc être exacte.
- La dimension de la représentation fondamentale de la symétrie doit être de trois comme le nombre de couleurs.
- La représentation fondamentale utilise des matrices complexes car les charges et anti-charges de couleur sont différentes.
- Le modèle doit enfin expliquer la liberté asymptotique, c'est-à-dire le fait que la constante de couplage diminue lorsqu'on s'approche des constituants des hadrons, c'est-à-dire à très haute énergie. Cette liberté asymptotique assure de plus qu'un quark libre n'est pas stable à basse énergie.

Ces contraintes fixent le groupe de symétrie à utiliser pour décrire la force de couleur : $SU(3)_c$ (le groupe $SU(3)$ "couleur" remplace le groupe $SU(3)$ "saveur"). Ce groupe est non-abélien, non-commutatif, et possède huit générateurs qui sont des matrices complexes, auxquels on associe huit champs de jauge. Ces champs correspondent aux gluons (Ce ne sont pas les mêmes objets mathématiques que ceux prédits par Gell-Mann et Zweig en 1961, néanmoins ce sont les mêmes qui avaient été mis en évidence expérimentalement à Standford). Les gluons ont la possibilité de se coupler entre eux, i.e. il existe des vertex à trois ou quatre gluons. Cette caractéristique est mise en

évidence par les termes cinétiques du lagrangien de la QCD . Les gluons sont bicolores et porteurs de la charge de couleur. On peut trouver par exemple un gluon $R\bar{V}$ pour les couleurs rouge et anti-vert. Ils sont finalement de masse et de charge nulle. Ce sont les vecteurs de l'interaction forte à l'intérieur du noyau, ils agissent sur la couleur sans modifier la saveur des quarks.

Le spectre des quarks s'est ensuite étendu jusqu'à la découverte du quark top, le plus lourd, en 1995 par les expériences $D\bar{O}$ et CDF [20, 21]. On compte désormais six quarks : *up* (u), *down* (d), *charm* (c), *strange* (s), *bottom* (b) et *top* (t). Les caractéristiques de ces quarks seront données dans la Section 1.3.8. La matrice CKM pour Cabbibo-Kobayashi-Maskawa définit les couplages entre ces quarks. C'est une matrice 3×3 unitaire, qui est construite à partir de trois angles et d'une phase complexe.

Une caractéristique remarquable de la QCD est l'explication de la liberté asymptotique, explication qui a d'ailleurs valu le prix Nobel à Gross, Politzer et Wilczek en 1994 [22, 23]. La constante de couplage diminue fortement pour les hautes énergies (voir la Figure 1.7), ce qui permet d'utiliser, comme pour la QED , la théorie des perturbations pour calculer les sections efficaces des processus de QCD . A plus basses énergies, c'est le contraire, la constante augmente très fortement. Des modèles effectifs complexes permettent alors d'estimer ces sections efficaces (lagrangien effectif, QCD sur réseau...). L'interaction qui lie deux quarks entre eux est en effet très forte, ce qui explique d'ailleurs la stabilité du noyau atomique. Elle est due à un échange permanent de gluons entre les quarks. Le proton est donc formé de quarks et d'un nuage gluonique. Cette structure empêche tout calcul de la masse du proton à partir de la masse des quarks. De la même manière, il est très difficile d'estimer la masse des quarks u et d . A l'échelle des quarks, on peut dire que l'interaction forte augmente fortement en fonction de la distance entre les quarks : c'est ce qu'on appelle le confinement. Si cette distance devient trop importante, les gluons continuellement échangés par les quarks sont assez énergétiques pour pouvoir recréer des paires quarks-antiquarks, qui peuvent à leur tour se combiner aux quarks initiaux pour former des hadrons. C'est le phénomène de création de jets ou d'hadronisation observé dans les collisionneurs comme illustré sur la Figure 1.8.

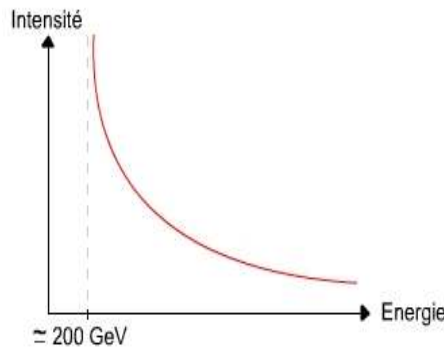


FIG. 1.7 – Evolution de l'intensité de l'interaction forte en fonction de l'énergie.

Conclusion

La chromodynamique quantique permet donc d'expliquer l'interaction forte, qui garantit la stabilité du noyau. Cette théorie, basée sur l'invariance de jauge selon les transformations du

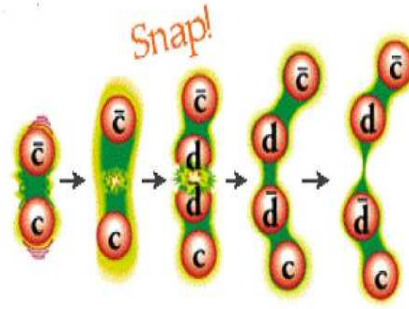


FIG. 1.8 – Phénomène d’hadronisation ou de fragmentation.

groupe $SU(3)$, ajoute à la description des particules élémentaires un nouveau nombre quantique, la couleur. Les champs utilisés comme triplet de couleur sont les quarks et les champ de jauge sont les gluons au nombre de huit. Ces dernières particules élémentaires complètent le contenu en champs donné dans les sections précédentes, qui forment ainsi les constituants de base du modèle standard (voir la Section 1.3.8)

1.3.6 Quelques mots sur la renormalisation d’une théorie

Les théories quantiques des champs, comme la QED détaillée en Section 1.3.2, le modèle standard électrofaible expliqué en 1.3.3 ou encore la QCD vue en Section 1.3.5, doivent être renormalisables. La renormalisabilité d’une théorie indique qu’après des manipulations mathématiques bien choisies le calcul de toutes les quantités physiques donne des valeurs finies. Ce sont Schwinger [10], Feynman [24], Tomonaga et Dyson [25] qui ont développé en 1949 ces manipulations mathématiques. En effet, le calcul des diagrammes de Feynman peut faire apparaître des divergences dues aux contributions des diagrammes à boucles comme il est illustré sur la Figure 1.2. Les divergences apparaissent parce que ces théories sont systématiquement construites à partir d’un découpage arbitraire entre la description des champs libres et la description des champs sous interactions. Les divergences introduites dans les quantités calculées sont non-physiques, elles peuvent être éliminées en utilisant des artefacts mathématiques. Les manipulations mathématiques, qui ont été mises au point, permettent donc de s’affranchir de cet arbitraire de la théorie.

Physiquement, on peut interpréter ces divergences grâce à la polarisation du vide. On considère par exemple un électron, dont on veut déterminer la charge. L’électron interagit continuellement avec des photons virtuels par interaction électromagnétique (par virtuel, on entend qu’ils ont une masse et une très courte durée de vie). Il en absorbe ou en émet. Les photons ont tendance à créer des paires d’électrons-positrons, qui s’alignent avec l’électron considéré. Toutes ces paires entourant l’électron masquent sa charge réelle. En fait, cette charge, comme sa masse, n’est finalement pas observable. Elles correspondent au cas d’un électron libre, c’est-à-dire à un cas purement mathématique. Seule la masse et la charge effective, c’est-à-dire la charge et la masse à une distance donnée de l’électron, sont mesurables. Le calcul de ces quantités effectives fait donc apparaître une distance, qui peut modifier ces deux grandeurs caractéristiques de l’électron. Cette distance arbitraire est l’artefact mathématique utilisé pour renormaliser la théorie. Finalement, les quantités physiques doivent être calculées en utilisant une masse et une charge effectives, qui prend

en compte la particule considérée et le "nuage" de particules l'entourant. La renormalisabilité n'est donc pas seulement une astuce mathématique pour s'affranchir de divergences indésirables ; elle rétablit en fait la réalité physique du processus en considérant indiscernables la particule et les interactions qu'elle subit. Dans le cas de la *QED*, seuls deux paramètres sont nécessaires pour traiter tous les diagrammes de Feynman : la masse et la charge. Une théorie est dite renormalisable si le nombre de paramètres effectifs à utiliser est fini. Mathématiquement, on dit qu'un terme du lagrangien est renormalisable s'il vérifie :

$$\Sigma = 4 - d - \sum_i n_i(s_i + 1) \geq 0 \quad (1.23)$$

Avec d le nombre de dérivées, n_i le nombre de champ de type i , s_i leur spin.

Dans l'exemple précédent, la charge électrique de l'électron augmente si l'énergie augmente. D'une autre manière, on peut dire que la constante de couplage de la *QED* augmente avec l'énergie. Pour la charge de couleur, deux effets doivent être pris en compte :

- la création de paires quark-antiquark conséquence au niveau des diagrammes de Feynman des boucles de quarks. Ces paires de quarks écrantent de la même manière qu'en *QED* la charge de couleur centrale.
- La contribution des boucles de gluons, possibles en *QCD* car les gluons peuvent interagir entre eux, a plutôt un effet anti-écran qui a tendance à augmenter la charge de couleur si on s'éloigne du quark considéré.

L'effet d'anti-écran est dominant si le nombre de saveurs n'est pas trop élevé. La Formule 1.24 donne l'évolution de la constante de couplage en fonction de l'énergie. On constate que si N_f , le nombre de saveurs, est inférieur à 17, la constante augmente avec l'énergie et l'effet des boucles de gluons est par conséquent dominant.

$$\alpha_s = \frac{g_3^2}{\hbar c} = \frac{1}{b \times \ln(\frac{\mu^2}{\Lambda^2})}, b = \frac{33 - 2N_f}{12\pi} \quad (1.24)$$

Avec μ l'énergie, Λ est l'énergie correspondant au pôle de Landau.

On retrouve la liberté asymptotique pour $\mu \gg \Lambda$, régime permettant le calcul perturbatif, et le confinement pour $\mu < \Lambda$.

Finalement, les théories de jauge développées pour expliquer les interactions fondamentales sont des théories renormalisables. Elles découplent systématiquement la matière et les interactions de façon arbitraire sachant que l'un ne peut aller sans l'autre (une particule libre n'existe pas et ne peut être qu'un objet mathématique). L'arbitraire, ainsi instauré, impose l'utilisation d'artefacts mathématiques dans le calcul des grandeurs physiques. Les paramètres utilisés pour calculer ces grandeurs sont souvent des paramètres effectifs dépendant de l'énergie à laquelle on se place, i.e. de la force de l'interaction que subit la particule. Ceci a pour répercussion une évolution des constantes de couplage, qui ne sont plus à proprement parlé des constantes. Dans un modèle ultime permettant une description exacte de la nature à partir d'une unique interaction, ces constantes seraient égales. C'est l'objectif des théories au-delà du modèle standard, des théories de grande unification qui seront développées dans la suite.

1.3.7 Les réussites du modèle standard

Le modèle standard a des imperfections théoriques et phénoménologiques (voir la Section 1.3.9 pour plus de détails). Néanmoins, aucune mesure expérimentale n'a permis de mettre à mal ce

modèle. Son côté prédictif a plusieurs fois contribué à son succès : prévision puis observation des courants neutres, prévision puis découverte du quark top... De plus, l'accord entre les valeurs des observables physiques prévues par le modèle et mesurées expérimentalement donne encore plus de force à ce modèle théorique. En exemple, on peut donner l'accord obtenu pour différentes observables électrofaibles. Les valeurs mesurées O^{meas} et les valeurs prédites O^{fit} sont données sur la Figure 1.9 qui récapitule les résultats à l'été 2006. Le *pull* représente l'écart entre ces deux valeurs de la façon suivante : $pull = \frac{O^{meas} - O^{fit}}{\sigma^{meas}}$ où σ^{meas} est l'incertitude sur la valeur mesurée.

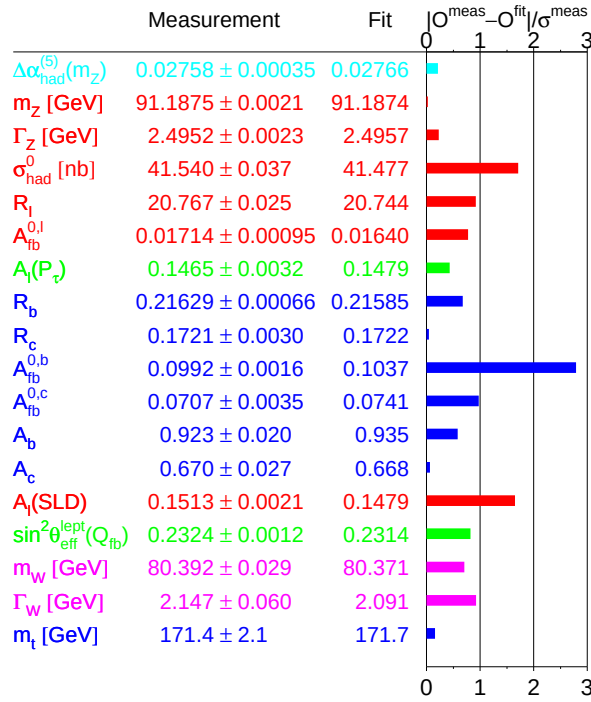


FIG. 1.9 – Valeurs mesurées et ajustées des observables électrofaibles du modèle standard.

1.3.8 Le contenu en champs

Le modèle standard décrit les particules élémentaires et les interactions intervenant entre ces particules. Les particules de matière, les fermions, sont composées de deux familles : les leptons et les quarks. Les seconds sont sensibles à l'interaction forte et portent par conséquent une charge de couleur. Chacune des familles est elle-même constituée de trois générations de particules classées en doublets d'isospin. Les vecteurs des interactions fondamentales sont les bosons de jauge et on distingue le photon pour l'interaction électromagnétique, les bosons Z^0 et $W^\pm$ pour l'interaction faible et les gluons pour l'interaction forte. Le mécanisme de Higgs, introduit pour rendre compte

de la brisure dynamique de la symétrie électrofaible, apporte une dernière particule. Le boson de Higgs complète ainsi le contenu en champs du modèle standard. Toutes ces particules sont décrites dans le même cadre théorique, qui intègre à la fois l'unification électrofaible et la chromodynamique quantique. Le groupe de symétries de cette théorie de jauge est le groupe : $SU(3)_c \otimes SU(2)_L \otimes U(1)_Y$. Le tableau 1.2 donne les principales caractéristiques de chacune de ces particules [26].

Section des fermions		Charge électrique	Spin	Masse (GeV)
Quarks	u	2/3	1/2	$1.5 \text{ à } 3 \cdot 10^{-3}$
	d	-1/3	1/2	$3 \text{ à } 7 \cdot 10^{-3}$
	c	2/3	1/2	1.25 ± 0.09
	s	-1/3	1/2	$95 \pm 25 \cdot 10^{-3}$
	t	2/3	1/2	174.2 ± 3.3
	b	-1/3	1/2	4.70 ± 0.07
Leptons	e^-	-1	1/2	$0.511 \pm 0.000 \cdot 10^{-3}$
	ν_e	0	1/2	$< 2 \cdot 10^{-6}$
	μ^-	-1	1/2	0.106 ± 0.000
	ν_μ	0	1/2	$< 0.19 \cdot 10^{-3}$
	τ^-	-1	1/2	1.777 ± 0.000
	ν_τ	0	1/2	$< 18 \cdot 10^{-3}$
Section des jauge		Charge électrique	Spin	Masse (GeV)
	photon (γ)	0	0	0
	$W^\pm$	± 1	1	80.403 ± 0.029
	Z^0	0	1	91.188 ± 0.002
	gluon (g)	0	1	0
Section de Higgs		Charge électrique	Spin	Masse (GeV)
	H	0	0	$> 114.4 \text{ @ } 95 \% \text{ C.L.}$

TAB. 1.2 – Principales caractéristiques des particules du modèle standard.

Le modèle standard compte finalement au moins 18 paramètres libres, c'est-à-dire non contraints par la théorie et obtenus directement des observations expérimentales. On compte ainsi 6 masses de quarks, 4 paramètres pour la matrice CKM , au moins 3 masses de leptons, 3 constantes de couplage et 2 paramètres pour le secteur de Higgs (dont la masse du boson de Higgs et par exemple la masse du boson Z^0 ou du boson $W^\pm$). On ajoute éventuellement un paramètre supplémentaire pour caractériser la violation de CP. Si on considère que les neutrinos ont une masse, il faut alors ajouter 3 paramètres correspondant aux masses, 4 paramètres pour définir une nouvelle matrice CKM , nécessaire pour les couplages entre leptons, et 2 nouvelles phases pour la violation de CP. Au total, le modèle standard peut compter jusqu'à 25 paramètres.

1.3.9 Les limites du modèles standard et quelques exemples d'extension

Le modèle standard constitue une description remarquable des phénomènes physiques jusqu'à des énergies de l'ordre de 100 GeV. Néanmoins, quelques interrogations théoriques restent en suspens. De plus, à plus haute énergie où les effets de la gravitation deviennent importants, les

corrections quantiques au secteur scalaire, pour la masse du boson de Higgs, deviennent trop importantes.

Les interrogations théoriques insolubles à ce jour sont à première vue d'ordre esthétique mais elles soulèvent la volonté permanente des physiciens de découvrir les briques élémentaires du monde qui nous entoure. En effet, on peut se poser les questions suivantes :

- Pourquoi y a t'il trois groupes de symétries ?
- Pourquoi y a t'il trois générations de quarks et de leptons ?

Ce qui revient au même que de se demander, s'il n'existe pas une interaction unique à la base des trois autres et des sous-structures aux fermions. L'interaction unique, socle d'une théorie du "tout", suggérerait d'ailleurs que toutes les constantes de couplages prennent la même valeur à une énergie, dite de grande unification. Les équations de renormalisation, qui permettent de calculer l'évolution de ces constantes, ne permettent pas par extrapolation cette convergence à hautes énergies, ce qui va dans le sens d'une imperfection du modèle standard et/ou qui peut confirmer le caractère effectif de la théorie. On est également en droit de se demander d'où viennent les nombres quantiques des particules, pourquoi la symétrie CP est-elle violée ou encore comment expliquer les différences d'ordre de grandeur entre les masses des particules de matière : des neutrinos aux masses non mesurées jusqu'au quark top. Le boson de Higgs reste encore introuvable, alors qu'il apporterait une nouvelle confirmation importante. Le mécanisme de masse reste en effet un mécanisme hypothétique purement théorique. Enfin, à très haute énergie, de nombreux problèmes théoriques viennent s'ajouter, qui confirment le caractère effectif de la théorie. Les corrections à la masse du Higgs sont de plusieurs ordres de grandeur supérieurs à la masse du boson elle-même. Il faut donc effectuer des réglages très fins sur les paramètres du lagrangien pour réduire cet écart d'ordre de grandeur. La masse m du boson de Higgs, corrigée de la masse "nue" m_0 en l'absence d'interaction, peut en effet être donnée par la Formule suivante :

$$m^2 = m_0^2 - \frac{y_f^2 \Lambda^2}{8\pi^2} \quad (1.25)$$

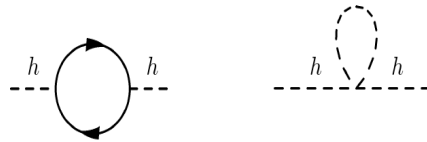


FIG. 1.10 – Corrections à une boucle à la masse du Higgs, boucle fermionique à gauche et boucle avec un champ scalaire à droite.

La Formule 1.25 tient compte uniquement de la divergence quadratique introduite par la correction à une boucle fermionique (voir le schéma de gauche de la Figure 1.10) de la masse du boson de Higgs. y_f est la valeur du couplage usuel du champ scalaire aux fermions (le couplage de Yukawa). Et, Λ est une coupure qui représente la limite de validité du modèle standard, i.e. le modèle standard est une théorie effective applicable jusqu'à une énergie Λ (échelle de grande unification). Comme il a été évoqué dans la Section 1.3.6, on utilise plutôt des quantités effectives ajustées sur les valeurs expérimentales pour s'affranchir des divergences dues aux diagrammes de Feynman à boucles. On utilise alors une masse "nue" susceptible de vérifier les contraintes

expérimentales sur la masse du Higgs. Par exemple, si on impose $m \approx 100$ GeV (la contrainte expérimentale) et $\Lambda = 10^{19}$ GeV, il faut ajuster m_0 et y_f à 36 décimales près.

$$m_0^2 = 10^{38} \left(\frac{y_f^2}{8\pi^2} + 10^{-36} \right) \quad (1.26)$$

On appelle ce problème le *fine tuning* en anglais.

On parle également d'un problème de hiérarchie pour soulever la différence entre l'échelle d'énergie de la brisure électrofaible ($m_{Z^0} \approx 90$ GeV) et l'échelle de grande unification, qui correspond à la masse de Planck ($m_{GUT} = \frac{1}{\sqrt{8\pi G_{Newton}}} = 2.4 \times 10^{18}$ GeV). L'échelle de Planck est en fait l'échelle à laquelle les effets quantiques de la gravitation ne sont plus négligeables. Comment donc concilier l'interaction gravitationnelle décrit par la gravitation générale et le formalisme quantique des champs, qui décrit actuellement les trois autres interactions fondamentales. A ces questions, on peut ajouter celles posées par les observations de la cosmologie et de l'astrophysique. D'où vient l'asymétrie matière-antimatière ? De quoi est constituée la matière noire qui semble composer une majeure partie de l'Univers ?

De nombreux modèles théoriques tentent de répondre aux interrogations et de résoudre les problèmes. Actuellement aucun modèle n'a toutes les réponses tout en redonnant à basse énergie la physique du modèle standard. Toutes les solutions proposées sont des solutions partielles ou des modèles effectifs à une énergie intermédiaire entre m_{Z^0} et m_{GUT} . On trouve par exemple les modèles suivants :

- La technicouleur [27] plutôt que d'utiliser un champ scalaire (le boson de Higgs), dont les corrections radiatives divergent à hautes énergies, utilise des champs fermioniques pour expliquer les masses des bosons de jauge et des fermions. Ces champs fermioniques, appelé technifermions, forment un état lié. L'avantage des fermions est qu'il n'y a plus le problème des divergences quadratiques à haute énergie. Par contre, ce modèle n'explique pas la masse des fermions et semble incompatible avec les observations expérimentales.
- Les théories de grande unification [28] cherchent à décrire les quatre interactions fondamentales comme une unique interaction à haute énergie. Les quatre interactions seraient alors le fruit de brisures de symétrie comme l'électromagnétisme et l'interaction faible sont les fruits de l'interaction électrofaible. Ces modèles cherchent alors un groupe de symétries commun pour toutes les interactions. Le choix le plus simple est le groupe de symétrie $SU(5)$. Ce modèle minimal a l'avantage de donner naturellement l'unification des constantes de couplage et d'expliquer les charges fractionnaires des quarks. Par contre, il n'explique ni le nombre de familles, ni le *fine tuning* et ne permet pas d'avoir des prédictions phénoménologiques correctes.
- La supersymétrie proposée par Wess et Zumino en 1974 [29] postule l'existence d'une nouvelle symétrie entre les fermions et les bosons pour régler le problème du *fine tuning*. Une description détaillée de cette théorie et des avantages qu'elle comporte est donnée dans la Section 1.4.
- Les dimensions supplémentaires permettent de fournir un modèle élégant pour expliquer le problème de hiérarchie en diminuant la masse de Planck. Elles décrivent la masse de Planck comme une masse effective à notre échelle, qui dépendrait du nombre et de la taille des dimensions supplémentaires.
- La théorie des supercordes [30–32] a pour but initial de donner une description quantique de la gravitation. Le problème de cette interaction est sa non-renormalisabilité dans le cadre

d'une théorie de jauge, comme celles détaillées plus tôt. L'idée est donc d'utiliser des objets de taille finie à plusieurs dimensions, des cordes, plutôt que des objets ponctuels pour décrire les particules élémentaires. Les modes d'oscillations de la corde sont alors associés à un type de particule d'une masse donnée. Si on associe à ce modèle des cordes la supersymétrie qui corrige des imperfections théoriques (l'apparition du tachyon, particule de masse carrée négative), on obtient le modèle des supercordes. Il existe plusieurs théories des supercordes qui proviennent d'une unique théorie appelée M-théorie (M pour magique, mystérieuse ou Mère), qui reste encore à mettre au point.

L'un des plus prometteurs reste la supersymétrie. Ce modèle théorique est détaillé dans la Section 1.4. La recherche expérimentale de particules supersymétriques est l'objet de cette thèse. Les résultats de cette recherche sont donnés dans le Chapitre 4.

1.4 La supersymétrie

1.4.1 Introduction et motivations

La supersymétrie [5, 33] est une nouvelle symétrie introduite par Weiss et Zumino en 1974 [29]. L'objectif initial était purement mathématique. Ils voulaient généraliser le groupe de Poincaré des symétries d'espace-temps. Cette nouvelle théorie permet de corriger le problème des divergences quadratiques introduites par les corrections radiatives à la masse des scalaires, c'est-à-dire à la masse du boson de Higgs dans le cas du modèle standard.

La correction à la masse du Higgs donnée par la Formule 1.25 peut être développée à l'ordre suivant et donne :

$$\Delta m_h^2 = -\frac{y_f^2 \Lambda^2}{8\pi^2} - \frac{3m_f^2}{8} \ln\left(\frac{\Lambda}{m_f}\right) - \dots \quad (1.27)$$

Si on considère maintenant, la contribution d'une boucle avec un champ scalaire à la correction de la masse du Higgs, on obtient le schéma de droite de la Figure 1.10 et la Formule suivante :

$$\Delta m_h^2 = +\frac{\lambda_s^2 \Lambda^2}{16\pi^2} - \frac{m_s^2}{8} \ln\left(\frac{\Lambda}{m_s}\right) - \dots \quad (1.28)$$

Avec λ_s , le couplage au champ scalaire et m_s la masse du scalaire qui se trouve dans la boucle.

On constate que les boucles fermioniques donnent un signe négatif alors que les boucles avec un champ scalaire donnent un signe positif. L'idée de la supersymétrie est donc de compenser systématiquement les divergences quadratiques en introduisant une nouvelle symétrie de telle façon à ce que chaque fermion ait deux partenaires scalaires avec $\lambda_s = y_f$. Les corrections radiatives ne varient plus comme le carré de l'énergie Λ mais comme le logarithme de cette énergie, comme pour les corrections à la masse des fermions. Ainsi, la supersymétrie permettrait de régler le problème de l'ajustement fin (*fine tuning*). Initialement, la symétrie de superjauge, renommée ensuite supersymétrie, devait être une symétrie entre les fermions et les bosons existants. Cette approche, qui avait l'avantage de ne pas introduire des nouvelles particules, s'est avérée impossible à cause des caractéristiques physiques trop différentes entre les fermions et les bosons connus.

En 1976, Fayet [34] postule l'existence de nouvelles particules, partenaires supersymétriques ou superpartenaires des particules du modèle standard. Il introduit également de nouveaux champs de Higgs, ainsi que leurs superpartenaires et définit la R-parité qui sera décrite ultérieurement. Aucun superpartenaire n'ayant été découvert à ce jour, on en déduit qu'ils n'ont pas la même

masse que les particules du modèle standard. Il faut donc introduire un mécanisme de brisure de la supersymétrie qui sera expliqué en Section 1.4.4 dans le cadre d'un modèle particulier de supersymétrie appelé le *MSSM* pour modèle standard supersymétrique minimal.

Les conséquences positives de ce modèle basé sur le postulat d'une nouvelle symétrie sont nombreuses. On peut citer en exemple la résolution du *fine tuning*. Une autre conséquence obtenue indirectement est la modification des équations de renormalisation, qui permet une convergence des constantes de couplage au même point comme le montre la Figure 1.11. Sur cette figure, α_1 , α_2 , et α_3 sont respectivement les constantes de couplage des interactions électromagnétique, faible et forte. Cette convergence laisse à penser que la supersymétrie pourrait être une théorie effective d'une théorie du "tout" et donc un nouveau pas en avant vers les modèles de grande unification. Un dernier avantage de la supersymétrie en guise d'exemple est qu'elle repousse le temps de vie du proton de 10^{31} années, estimé trop court, à 10^{32} années. D'autres avantages seront donnés au cours de la description détaillée de ce nouveau modèle théorique.

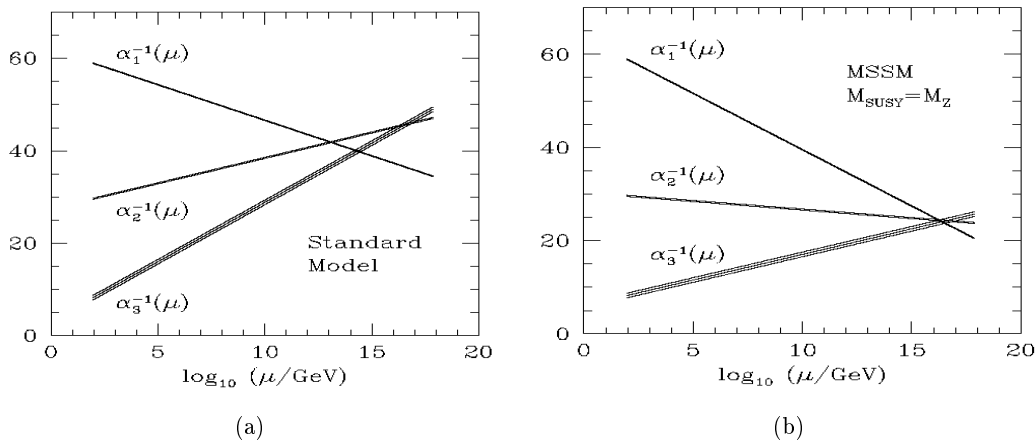


FIG. 1.11 – Evolution de l'inverse des constantes de couplage dans le cas du modèle standard (a) et dans le cas du MSSM (b).

1.4.2 L'algèbre supersymétrique

Introduction

Physiquement, la supersymétrie est possible grâce à l'introduction d'un nouvel espace complexe dont les coordonnées anticommutent contrairement à l'espace-temps de Minkowski. On appelle ce nouvel espace un super-espace. Les transformations de supersymétrie sont en fait des transformations correspondant à des rotations dans le super-espace. Elles permettent de transformer un fermion $|f\rangle$ en un boson $|s\rangle$ et inversement, comme indiqué dans les Formules 1.29 avec Q un générateur des transformations de supersymétrie. La supersymétrie n'est donc pas une symétrie interne mais une symétrie de l'espace-temps. L'algèbre de la supersymétrie fera donc intervenir les générateurs de l'algèbre de Poincaré.

$$Q|s\rangle = |f\rangle, Q|f\rangle = |s\rangle \quad (1.29)$$

Mathématiquement, cela s'explique par le fait que ces générateurs sont des opérateurs fermioniques de spin $1/2$. Ils permettent en effet de modifier le spin d'une particule d'une demi-unité. Or, le théorème de Coleman et Mandula [35] interdit les générateurs de spin différent de 0, exceptés les générateurs du groupe de Poincaré, pour former un groupe de Lie avec des paramètres réels. Ainsi, afin de décrire de telles symétries de l'espace-temps, la seule possibilité est de généraliser l'algèbre de Poincaré en utilisant des relations anticommutation. On parle alors de superalgèbre ou d'algèbre de Lie graduée.

Quelques mots sur les notations

Un spineur de Weyl est un multiplet à deux composantes qui permet de décrire une particule, ou une transformation, de spin $1/2$ et de parité donnée. Ainsi, on note Q_α un spineur de Weyl de chiralité gauche et $\bar{Q}_{\dot{\alpha}}$ un spineur de Weyl de chiralité droite. On peut créer un spineur à quatre composantes englobant les deux chiralités d'un même spineur : $\begin{pmatrix} Q_\alpha \\ \bar{Q}_{\dot{\alpha}} \end{pmatrix}$. On parle dans ce cas de spineur de Majorana. La propriété caractéristique des spineurs de Majorana est qu'ils sont égaux à leur conjugué de charge. Si on se restreint à un nombre $N = 1$ de générateurs pour la supersymétrie, on peut construire une charge conservée dans la symétrie au sens du théorème de Noether : $Q = \begin{pmatrix} Q_\alpha \\ \bar{Q}_{\dot{\alpha}} \end{pmatrix}$. Cette charge est un spineur de Majorana qui regroupe deux générateurs de la supersymétrie. Deux générateurs sont en effet nécessaires pour décrire les deux chiralités d'un spineur de Weyl. On utilise également les générateurs du groupe de Poincaré pour construire l'algèbre de Lie graduée : $M_{\mu\nu}$ pour les rotations et P_μ pour les translations.

Définition de l'algèbre

Les relations de commutation avec l'algèbre de Poincaré sont les suivantes :

$$[M_{\mu\nu}; Q_\alpha] = \frac{1}{2}(\sigma_{\mu\nu})_\alpha^\beta \quad (1.30)$$

$$[M_{\mu\nu}; \bar{Q}_{\dot{\alpha}}] = -\frac{1}{2}(\bar{\sigma}_{\mu\nu})_{\dot{\alpha}}^{\dot{\beta}} \quad (1.31)$$

$$[Q_\alpha; P_\mu] = 0 \quad (1.32)$$

$$[\bar{Q}_{\dot{\alpha}}; P_\mu] = 0 \quad (1.33)$$

Les relations d'anticommutation pour les générateurs fermioniques sont :

$$\{Q_\alpha; \bar{Q}_{\dot{\beta}}\} = 2(\sigma^\mu)_{\alpha\dot{\beta}} P_\mu \quad (1.34)$$

$$\{Q_\alpha; Q_\beta\} = 0 \quad (1.35)$$

Dans les équations précédentes, les matrices σ^μ et $\bar{\sigma}^\mu$ sont définies à l'aide des matrices de Pauli : $\sigma^\mu = (1, \sigma^1, \sigma^2, \sigma^3)$ et $\bar{\sigma}^\mu = (1, -\sigma^1, -\sigma^2, -\sigma^3)$. La matrice $\sigma_{\mu\nu}$ est finalement définie de la façon suivante : $\sigma_{\mu\nu} = (\bar{\sigma}^\mu \sigma^\nu - \bar{\sigma}^\nu \sigma^\mu)$.

Conséquences physiques

Les prédictions de ce superalgèbre sur la physique introduite par la supersymétrie sont les suivantes (les justifications mathématiques de ces prédictions ne seront pas détaillées ici) :

- Les superpartenaires ont la même masse et un spin différent de leur partenaire du modèle standard, ce qui impose que la supersymétrie est brisée.
- Un état de vide avec une énergie non nulle équivaut à une brisure de la supersymétrie globale. De plus, l'état fondamental est forcément d'énergie positive ou nulle. Cette propriété a des conséquences cosmologiques importantes, car l'énergie du vide contribue à la densité d'énergie de l'Univers
- Il y a autant d'opérateurs bosoniques que d'opérateurs fermioniques.
- La Formule 1.34 indique que Q^2 est une translation de l'espace-temps. Dans ce cas, si la supersymétrie est locale, alors elle peut être utilisée pour une quantification de la gravitation en terme de théorie des champs.

Le nombre N de générateurs de la supersymétrie est indifférent. Pour chaque N donné, on a un modèle de supersymétrie. Pour des raisons de simplicité, le cas $N = 1$ est souvent préféré, c'est ce modèle, appelé *MSSM*, qui sera décrit dans la suite.

1.4.3 Le lagrangien supersymétrique

Les supermultiplets

Les supermultiplets, c'est-à-dire les représentations irréductibles du superalgèbre, sont constitués de particules ayant les mêmes nombres quantiques excepté le spin. Les opérateurs supersymétriques commutent en effet avec toutes les symétries internes car ils correspondent à des symétries de l'espace-temps (c'est une propriété générale à toutes les symétries externes). On peut donc former un supermultiplet à partir de particules ayant la même masse (avant toute brisure de la supersymétrie) et les mêmes nombres quantiques excepté le spin. Il est néanmoins possible de construire un nouveau nombre quantique à partir des spins des particules : le superspin. Les valeurs du superspin au sein d'un même multiplet vérifient : $M_S = M_j, M_j + \frac{1}{2}, M_j - \frac{1}{2}, M_j$, où M_S est la valeur du spin et M_j la valeur du superspin. Finalement, un supermultiplet est constitué de particules de même masse et de même superspin.

Dans le cas $M_j = 0$, on parle de supermultiplet chiral ou scalaire. Ce supermultiplet comporte deux champs scalaires réels ou un champ scalaire complexe, Φ , et un champ fermionique de spin $1/2$, ψ . Le nombre de degrés de liberté de spin passe de deux à quatre pour un fermion quand il est hors couche de masse (*off-shell*). Si on veut conserver l'équilibre entre degrés de liberté fermionique et bosonique pour ce supermultiplet, on est amené à ajouter un champ auxiliaire scalaire complexe, F , qui ne se propage pas, c'est-à-dire sans terme cinétique. Sur couche de masse (*on-shell*), on utilise alors l'équation du mouvement $F = F^* = 0$. Dans ce cas, que ce soit sur couche de masse ou hors couche de masse, le nombre de degrés de liberté fermionique et bosonique est constamment 4. Le supermultiplet chiral peut donc être écrit : $\Psi = (\Phi, \psi_\alpha, F)$.

Le supermultiplet vecteur ou de jauge correspond à $M_j = \frac{1}{2}$. Il comporte un boson jauge de masse nulle, A_a^μ , donc trois degrés de liberté bosonique hors couche de masse ou deux polarisations sur couche de masse. Il est également constitué d'un fermion de Weyl de masse nulle λ^a , si on considère que la symétrie est conservée, qui a deux ou quatre degrés de liberté. De la même façon que pour le supermultiplet chiral, on introduit un champ scalaire réel, D^a pour corriger le déficit

d'un degré de liberté quand on est hors couche de masse. a est un indice utilisé pour la description de tous les bosons de jauge du modèle standard.

Le champ scalaire complexe du multiplet chirale est appelé le sfermion et le champ fermionique du supermultiplet vecteur est appelé jaugino. De la même façon, le nom des partenaires supersymétriques des fermions est construit en ajoutant un "s". On parle par exemple de selectron, de sneutrino, de sbottom... Et le nom des superpartenaires des bosons de jauge est formé du nom du boson de jauge auquel on ajoute le suffixe "ino". On a par exemple : le gluino, le zino, le wino...

Le lagrangien libre

Comme point de départ, on considère le lagrangien libre, c'est-à-dire sans introduire de boson de jauge, sur couche de masse. Ce lagrangien se décompose en deux parties. La première permet de caractériser le mouvement d'un champ scalaire, $\mathcal{L}_{\text{scalaire}}$, la seconde le mouvement d'un fermion de spin $1/2$, $\mathcal{L}_{\text{fermion}}$. On a alors le lagrangien suivant :

$$\mathcal{L}_{\text{libre}} = \mathcal{L}_{\text{fermion}} + \mathcal{L}_{\text{scalaire}} = -i\psi^\dagger \bar{\sigma}^\mu \partial_\mu \psi - \partial^\mu \Phi \partial_\mu \Phi^* \quad (1.36)$$

Pour une transformation infinitésimale supersymétrique, le champ scalaire est transformé en champ fermionique, et on peut écrire :

$$\delta\Phi = \epsilon^\alpha \psi_\alpha, \delta\Phi^* = \epsilon^\dagger \psi^\dagger = \bar{\epsilon}_{\dot{\alpha}} \bar{\psi}^{\dot{\alpha}} \quad (1.37)$$

Avec ϵ^α un fermion de Weyl infinitésimal indépendant des coordonnées de l'espace-temps : $\partial^\mu \epsilon_\alpha = 0$.

Un champ scalaire a la dimension d'une masse, alors qu'un champ de Weyl qui décrit une particule fermionique a la dimension d'une masse à la puissance $\frac{3}{2}$. On en déduit que ϵ^α a la dimension d'une masse à la puissance $-\frac{1}{2}$. Une transformation infinitésimale de $\mathcal{L}_{\text{scalaire}}$ est donc :

$$\delta\mathcal{L}_{\text{scalaire}} = -\epsilon^\alpha (\partial^\mu \psi_\alpha) \partial_\mu \Phi^* - \partial^\mu \Phi \bar{\epsilon}_{\dot{\alpha}} (\partial_\mu \bar{\psi}^{\dot{\alpha}}) \quad (1.38)$$

Pour que le modèle sans interaction soit invariant sous les transformations de supersymétrie, il faut, d'après le principe de moindre action, que l'action soit également invariante sous les transformations de supersymétrie (L'action S vérifie : $S = \int d^4x \mathcal{L}_{\text{libre}}$). Les transformations de $\mathcal{L}_{\text{scalaire}}$ et $\mathcal{L}_{\text{fermion}}$ doivent donc se compenser ou différer d'une dérivée totale. Par conséquent, la transformation d'un fermion $\delta\psi$ doit être linéaire en ϕ et $\epsilon^\dagger$. De plus, on veut que la transformation d'un fermion ait la dimension d'une masse à la puissance $\frac{3}{2}$. Les champs fermioniques se transforment donc de la façon suivante :

$$\delta\psi_\alpha = i(\sigma^\mu \epsilon^\dagger)_\alpha \partial_\mu \Phi, \delta\bar{\psi}^{\dot{\alpha}} = -i(\epsilon \sigma^\mu)^{\dot{\alpha}} \partial_\mu \Phi^* \quad (1.39)$$

En utilisant les identités pour les matrices de Pauli, on peut finalement montrer que la somme de $\delta\mathcal{L}_{\text{scalaire}}$ et $\delta\mathcal{L}_{\text{fermion}}$ est égale à une dérivée totale et donc que la variation infinitésimale de l'action est nulle. Ce modèle est donc invariant sous les transformations supersymétriques. Pour se convaincre que les transformations des champs fermioniques et des champs scalaires sont bien des transformations supersymétriques, il suffit de les appliquer deux fois et ainsi constater que : $\Phi \rightarrow \psi \rightarrow \partial\Phi$.

Pour vérifier que les transformations de supersymétrie, qui viennent d'être définies, appartiennent effectivement à la superalgèbre, il faut que l'algèbre de ces transformations soit fermée. Autrement dit, il faut que le commutateur de deux transformations infinitésimales de supersymétrie appartiennent à la superalgèbre. On peut montrer que le commutateur appliqué à un champ scalaire est proportionnel à la dérivée de ce champ et donc à l'impulsion. Par contre, si on applique ce même commutateur sur un champ fermionique, deux termes apparaissent. Le premier est également proportionnel à une dérivée, alors que le second n'appartient pas à la superalgèbre mais s'annule sur la couche de masse. C'est la raison pour laquelle on introduit un champ scalaire auxiliaire F et le lagrangien correspondant : $\mathcal{L}_{\text{auxiliaire}} = FF^*$. Les champs scalaires obéissent aux lois de transformations suivantes :

$$\delta F = i\bar{\epsilon}^{\dot{\alpha}}(\bar{\sigma}^{\mu})_{\dot{\alpha}}^{\beta}\partial_{\mu}\psi_{\beta}, \delta F^* = -i\partial_{\mu}\bar{\psi}^{\dot{\beta}}(\bar{\sigma}^{\mu})_{\dot{\beta}}^{\alpha}\epsilon_{\alpha} \quad (1.40)$$

On remarque que ces transformations ne font intervenir que le champ fermionique étant donné que l'action du commutateur sur un champ scalaire appartient déjà au superalgèbre. Ces transformations, ainsi définies, permettent de construire la variation infinitésimale $\delta\mathcal{L}_{\text{auxiliaire}}$. Cette variation est nulle sur couche de masse grâce aux équations du mouvement $F = F^* = 0$. Pour conserver l'invariance de l'action hors couche de masse malgré cette nouvelle variation du lagrangien, il faut modifier la transformation infinitésimale du champ fermionique (en effet, $\delta\mathcal{L}_{\text{auxiliaire}}$ dépend uniquement de ψ et non de Φ) sous l'action de la supersymétrie. On a alors :

$$\delta\psi_{\alpha} = i(\sigma^{\mu}\epsilon^{\dagger})_{\alpha}\partial_{\mu}\Phi + \epsilon_{\alpha}F, \delta\bar{\psi}^{\dot{\alpha}} = -i(\epsilon\sigma^{\mu})^{\dot{\alpha}}\partial_{\mu}\Phi^* + \bar{\epsilon}^{\dot{\alpha}}F^* \quad (1.41)$$

Finalement, on peut vérifier que l'action du commutateur de deux transformations infinitésimales de supersymétrie sur un champ quelconque ($\Phi, \Phi^*, \psi, \psi^{\dagger}, F$ ou F^*) est proportionnelle à la dérivée de ce champ, donc à l'impulsion. Autrement dit, sans faire référence aux équations de mouvement, on a construit des transformations qui appartiennent à la superalgèbre et qui transforment un fermion en un boson et vice versa. Le champ scalaire auxiliaire a permis de généraliser la théorie hors couche de masse et de fermer l'algèbre des transformations. Le champ F , artefact mathématique, prend toute sa signification physique et son importance dans la partie décrivant la brisure de la supersymétrie de la Section 1.4.4.

Les interactions

Interactions du supermultiplet chiral Dans la section précédente, on a décrit le lagrangien sans interaction du supermultiplet chiral. On peut généraliser le lagrangien obtenu pour plusieurs supermultiplets chiraux $\Psi_i = (\Phi_i, \psi_i, F_i)$. L'indice i correspond aux degrés de liberté de jauge et de saveurs. Les transformations des champs d'un supermultiplet chiral ont été également données dans la section précédente. Dans la suite, on donne, sans les détails des démonstrations mathématiques, les résultats concernant la construction d'un terme d'interaction invariant sous les transformations de supersymétrie. La forme la plus générale pour un terme d'interaction entre multiplets chiraux est la suivante :

$$\mathcal{L}_{\text{int}} = -\frac{1}{2}W^{ij}\psi_i\psi_j + W^iF_i + c.c. \quad (1.42)$$

c.c. désigne le complexe conjugué.

La Formule 1.23 donne les conditions de renormalisabilité de chacun des termes. Elle impose que W^{ij} et W^i soient des fonctions des champs scalaires Φ_i et Φ_i^* uniquement : W^{ij} est une

fonction linéaire et W^i est au plus un polynôme de degré deux de ces champs. Des champs bosoniques auraient pu également respecter la condition de renormalisabilité, néanmoins il aurait été impossible de les compenser par d'autres termes dans la variation infinitésimale du lagrangien. Les contraintes sur l'invariance selon les transformations de supersymétrie permettent de conclure (après plusieurs manipulations mathématiques) aux définitions suivantes pour W^{ij} et W^i :

$$W^{ij} = \frac{\partial^2 W}{\partial \Phi_i \partial \Phi_j} \quad (1.43)$$

$$W^i = \frac{\partial W}{\partial \Phi_i} \quad (1.44)$$

Avec $W = \frac{1}{2}m^{ij}\Phi_i\Phi_j + \frac{1}{6}y^{ijk}\Phi_i\Phi_j\Phi_k$, où m^{ij} est la matrice symétrique de masse et y^{ijk} la matrice totalement symétrique des couplages de Yukawa entre un scalaire et deux fermions.

Les équations du mouvement pour le lagrangien avec ces nouveaux termes d'interaction deviennent : $F_i = -W_i^*$ et $F^{*i} = -W^i$. Ainsi, on peut écrire le lagrangien sans y faire apparaître explicitement le champ scalaire auxiliaire F :

$$\mathcal{L} = -i\psi^{i\dagger}\bar{\sigma}^\mu\partial_\mu\psi_i - \partial^\mu\Phi_i\partial_\mu\Phi^{i*} - \frac{1}{2}m^{ij}\psi_i\psi_j - \frac{1}{2}m_{ij}^*\psi^{i\dagger}\psi^{j\dagger} - V - \frac{1}{2}y^{ijk}\Phi_i\psi_j\psi_k - \frac{1}{2}y_{ijk}^*\Phi^{i*}\psi^{j\dagger}\psi^{k\dagger} \quad (1.45)$$

Avec V le potentiel scalaire de la théorie : $V = W^i W_i^* = F_i F^{i*}$

Physiquement, le terme $\mathcal{L}_{int}$ décrit toutes les interactions exceptées celle de jauge, i.e. celles qui font intervenir un boson de jauge. Les bosons de jauge sont indispensables pour décrire complètement la théorie supersymétrique, qui n'est autre qu'une extension du modèle standard.

Interactions de jauge supersymétriques Les bosons de jauge apparaissent dans le supermultiplet vecteur. Il contient en effet un champ de jauge A_μ^a sans masse, un fermion de Weyl λ^a également sans masse car la supersymétrie n'est pas brisée à ce stade, et un champ scalaire auxiliaire réel D^a . On rappelle que l'indice a est utile à la description de tous les bosons de jauge du modèle standard. Les conditions d'invariance de jauge et la renormalisabilité de la théorie donnent le lagrangien du supermultiplet vecteur sans interaction :

$$\mathcal{L}_{jauge} = -\frac{1}{4}F_{\mu\nu}^a F^{a\mu\nu} - i\lambda^{a\dagger}\bar{\sigma}^\mu D_\mu\lambda^a + \frac{1}{2}D^a D^a \quad (1.46)$$

avec les définitions usuelles à toutes les théories de jauge :

$$F_{\mu\nu}^a = \partial_\mu A_\nu^a - \partial_\nu A_\mu^a - gf^{abc}A_\mu^b A_\nu^c \quad (1.47)$$

$$D_\mu\lambda^a = \partial_\mu\lambda^a - gf^{abc}A_\mu^b\lambda^c \quad (1.48)$$

où f^{abc} sont les constantes de structure antisymétriques du groupe de jauge considéré, D_μ les dérivées covariantes associées au champ fermionique et g le couplage de jauge.

Ce lagrangien est déjà supersymétrique et on peut le montrer comme dans le cas du lagrangien pour le multiplet chiral. A ce stade, c'est-à-dire sans interaction entre le multiplet chiral et le multiplet vecteur, les équations du mouvement donnent : $D^a = 0$. Le champ scalaire auxiliaire ne se propage donc pas.

Pour poursuivre la généralisation du modèle standard et décrire les interactions de jauge, il faut relier le lagrangien de jauge et le lagrangien du multiplet chiral, de telle façon à ce que ce dernier soit invariant sous les transformations de jauge et que l'ensemble soit invariant sous les transformations supersymétriques. Ainsi, on remplace les dérivées du lagrangien chiral par des dérivées covariantes et on introduit de nouveaux termes pour rendre compte des interactions entre les champs du supermultiplet chiral et ceux du supermultiplet vecteur. Les variations sous les transformations supersymétriques des nouveaux termes entraînent des modifications des variations infinitésimales des champs du supermultiplet chiral afin de maintenir l'invariance supersymétrique. Une fois, ces opérations mathématiques effectuées, on peut réévaluer les équations du mouvement et trouver : $D^a = -g(\Phi^* T^a \Phi)$, avec T^a les générateurs des transformations de jauge. Le potentiel scalaire devient finalement : $V = F_i F^{i*} + \frac{1}{2} D^a D^a$. Une propriété importante de ce potentiel scalaire qui sera développée dans la suite est qu'il est toujours positif ou nul. Ce potentiel scalaire est, de plus, complètement déterminé par les autres de la théorie, c'est-à-dire par les couplages de Yukawa et les couplages de jauge.

Conclusion

L'introduction d'une nouvelle symétrie entre bosons et fermions a permis de généraliser l'algèbre de Poincaré. Les supermultiplets vecteur et les supermultiplets chiraux sont regroupés au sein du même formalisme, qui décrit une généralisation du modèle standard. Le potentiel scalaire de la théorie ne souffre pas d'un ajustement fin pour éviter les divergences quadratiques puisqu'il est directement lié aux autres interactions de la théorie. La physique et la phénoménologie de cette théorie supersymétrique sont détaillées dans la Section 1.4.4, où l'on trouve une description du modèle standard supersymétrique minimal qui introduit une brisure de la supersymétrie.

1.4.4 Le modèle standard supersymétrique minimal (MSSM)

Aucune particule supersymétrique n'a encore été découverte. La supersymétrie est par conséquent une symétrie brisée. Pour avoir une description phénoménologique et pour étudier les modèles supersymétriques dans les collisionneurs, il faut être en mesure de modéliser ou du moins de décrire de façon effective cette brisure. De plus, le modèle supersymétrique décrit précédemment ne modifie pas les nombres quantiques des particules. On rappelle qu'un supermultiplet est composé de particules de mêmes nombres quantiques excepté le spin. Il est donc théoriquement impossible d'associer les fermions et les bosons existants, de telle façon à ce que les premiers soient les partenaires supersymétriques des seconds. Il faut finalement ajouter de nouvelles particules pour rendre le modèle consistant.

La brisure souple de la supersymétrie

Le mécanisme de la brisure de la supersymétrie est encore inconnu. La solution communément adoptée est une description effective de cette brisure par l'ajout de termes de brisures explicites. De cette façon, les modèles supersymétriques étudiés actuellement sont des extensions au-delà du modèle standard mais restent des théories effectives. Concrètement, la supersymétrie permet généralement d'étudier les phénomènes physiques à l'échelle du TeV alors que le modèle standard fournit une description d'une très grande précision jusqu'à des échelles d'énergie de l'ordre de 100 GeV.

Les termes ajoutés explicitement doivent vérifier plusieurs contraintes pour conserver les propriétés et les avantages des modèles supersymétriques. Tout d'abord, la supersymétrie a été motivée par la suppression des divergences quadratiques qui obligent l'ajustement fin du modèle standard. Par conséquent, les termes de brisure ne doivent pas diverger quadratiquement. On parle alors de brisure souple. Cette brisure souple se matérialise par des termes ayant des couplages avec une dimension de masse positive, des termes de masse par exemple. Ce sont ces termes de brisure souple qui vont donner des masses différentes aux particules. Dans le cas d'une théorie renormalisable, on peut écrire ces termes de la façon suivante :

$$\mathcal{L}_{\text{souple}} = -\frac{1}{2}m_\lambda\lambda^a\lambda^a - \frac{1}{6}a^{ijk}\Phi_i\Phi_j\Phi_k - \frac{1}{2}b^{ij}\Phi_i\Phi_j + c.c. - (m^2)_j^i\Phi_j^*\Phi_i \quad (1.49)$$

Les termes a^{ijk} et b^{ij} sont de la même forme que y^{ijk} et m^{ij} . Les champs chiraux n'interviennent pas dans la Formule 1.49, car il est toujours possible de redéfinir le superpotentiel et les constantes pour les éliminer. Cette formule est la plus générale pour une brisure souple explicite. Malheureusement, elle contient 109 paramètres libres dans le cas du *MSSM* si on ne fait aucune hypothèse ou si on ne suppose pas l'existence de nouvelles symétries (comme la R-parité qui sera évoquée plus loin). Un tel modèle ne permet donc pas une description phénoménologique simple étudiable dans les collisionneurs.

De plus, on a vu précédemment que la supersymétrie est conservée si et seulement si le vide a l'énergie zéro. Dans le cas contraire, si la supersymétrie est brisée spontanément, le vide a une énergie strictement positive. Par conséquent, une brisure spontanée de la supersymétrie peut être obtenue si les champs F_i et D^a ne sont pas nuls dans le vide. On en déduit que les équations du mouvement $F_i = 0$ et $D^a = 0$ ne peuvent être simultanément satisfaites. Pour ceci, il suffit d'ajouter les termes de brisure dans le potentiel scalaire afin de modifier les équations du mouvement. Deux mécanismes existent pour décrire une brisure souple et spontanée de la supersymétrie :

- Le mécanisme de Fayet-Iliopoulos [36] introduit un terme proportionnel à D^a dans le potentiel scalaire, ce qui modifie les équations de mouvement pour le champ D^a et conduit à une valeur non nulle dans le vide pour ce champ. Ce modèle souffre néanmoins de difficulté dans le cas du *MSSM*. Il induit ou bien une brisure de la couleur et de l'électromagnétisme pour certains superpartenaires ou alors des masses incompatibles avec les limites expérimentales.
- Le mécanisme de O'Raifeartaigh [37] permet la brisure grâce à un terme de type F (contenant seulement les champs F_i). Ce terme modifie le potentiel V de façon à ce que les équations $F_i = 0$ et $\frac{\delta W^*}{\delta \Phi^{i*}} = 0$ ne soient pas simultanément satisfaites.

Une dernière contrainte à considérer est l'existence de règles de somme pour les masses, que la supersymétrie soit brisée ou non. Ces règles s'expriment à partir des traces des matrices de masses au carré :

$$Tr[m^2(\text{scalaires complexes})] = Tr[m^2(\text{fermions de Weyl})] \quad (1.50)$$

Cette règle est automatiquement vérifiée avec le mécanisme de O'Raifeartaigh. Néanmoins, elle ne se vérifie pas expérimentalement. En effet, les fermions de Weyl ont tous des masses petites exceptés le top et les higgsinos. Or, aucun scalaire supersymétrique léger n'a encore été observé. Par conséquent, deux solutions pourraient permettre d'expliquer de telles différences de masse : des corrections radiatives importantes ou le postulat d'un secteur caché. Par secteur caché, on entend un secteur de très haute énergie inaccessible aux expériences de physique des particules. Cette dernière solution est souvent préférée. Elle possède de plus l'avantage de proposer des modèles phénoménologiques avec peu de paramètres libres et donc directement étudiables expérimentale-

ment. On distingue alors deux secteurs : le secteur visible et le secteur caché. L'ensemble de ces deux secteurs vérifient la Formule 1.50, alors que le secteur visible ne la vérifie pas. Les deux secteurs interagissent pour communiquer la brisure de la supersymétrie. Plusieurs modèles existent pour décrire l'interaction entre les deux secteurs selon que l'interaction soit de jauge (on trouve par exemple les modèles *AMSB* et *GMSB*) ou soit gravitationnelle (le modèle *mSUGRA*).

Sans entrer dans les détails mathématiques de ces modèles de brisure de la supersymétrie, on peut en citer les conséquences fondamentales. Ils comportent peu de paramètres libres, de 109 on passe par exemple à 5 pour le modèle *mSUGRA*. Les masses des squarks et des sleptons ne dépendent que de leurs nombres quantiques de jauge. De plus, les particules qui interagissent par interactions fortes, squarks et gluinos, sont systématiquement plus lourdes que les particules interagissant par interactions faibles, sleptons et higgsinos.

Les superpartenaires des particules du modèle standard

Les particules du modèle standard possèdent toute un partenaire supersymétrique : un superpartenaire scalaire pour les fermions et un superpartenaire fermionique pour les bosons de jauge ou les scalaires. Les principales différences ou particularité du *MSSM* par rapport au modèle standard sont les suivantes :

- Les superpartenaires fermioniques des bosons de jauge sont des spineurs de Weyl, ou de Majorana, alors que le modèle standard ne comporte que des spineurs de Dirac (multiplet à quatre composantes indépendantes).
- On distingue deux types de superpartenaires scalaires, $\tilde{q}_L$ et $\tilde{q}_R$, selon que le fermion du modèle standard associé soit d'hélicité droite ou d'hélicité gauche. Il faut néanmoins garder à l'esprit que les squarks sont des scalaires et que les notations L et R font référence à leur partenaire du modèle standard et non à une propriété d'hélicité.
- Le secteur de Higgs est étendu et est constitué de deux doublets au lieu d'un seul. On a alors deux bosons de Higgs neutres H_u^0 et H_d^0 pour le couplage aux quarks de type "up" et "down" respectivement. On a également deux bosons de Higgs chargés H_u^+ et H_d^- .
- Les jauginos et les higgsinos sont des fermions qui possèdent les mêmes nombres quantiques. On ne peut pas les distinguer. Leurs champs se recombinaient donc pour former des états propres de masse, différents après la brisure de la symétrie électrofaible. Les neutralinos, $\tilde{\chi}_{1,2,3,4}^0$, correspondent au mélange des champs associés aux particules $\tilde{B}^0$, $\tilde{W}^0$, $\tilde{H}_u^0$ et $\tilde{H}_d^0$. Les charginos, $\tilde{\chi}_{1,2}^\pm$, sont les états propres de masse des $\tilde{W}^\pm$, H_u^+ et H_d^- .

L'avant-dernier point est nécessaire pour deux raisons. La première est qu'un deuxième doublet de Higgs est nécessaire pour que les diagrammes d'anomalies triangulaires se compensent entre eux (voir la Figure 1.12), ce qui évite ainsi tout problème de divergence. La seconde est que le terme F du potentiel ne fait pas intervenir de terme conjugué, comme le montre la Formule 1.44, alors que le couplage du boson de Higgs aux quarks u, c et fait intervenir le champ scalaire conjugué du boson de Higgs. La création d'un second doublet avec un nouveau champ de Higgs d'hypercharge égale à -1 règle ce problème. Elle nécessite par contre l'apparition des deux bosons de Higgs chargés. De plus, on peut montrer que le vide est électriquement neutre, c'est-à-dire qu'il ne brise pas la symétrie électromagnétique. En développant le potentiel scalaire, fonction des champs scalaires auxiliaires, on peut également contraindre le secteur de Higgs et ainsi montrer que la valeur attendue dans le vide des deux bosons de Higgs est : $\langle H_u \rangle = \begin{pmatrix} v_1 \\ 0 \end{pmatrix}$ et $\langle H_d \rangle = \begin{pmatrix} 0 \\ v_2 \end{pmatrix}$. Il n'est pas possible d'avoir les mêmes masses pour les deux bosons de Higgs. On définit donc la valeur

β à partir de ces valeurs dans le vide : $\tan(\beta) = \frac{v_1}{v_2}$. Le paramètre $\tan(\beta)$ est un paramètre libre du *MSSM*. Les recherches directes de la supersymétrie aux collisionneurs ont exclu une partie des petites valeurs pour ce paramètre. Ce paramètre est commun à tous les modèles supersymétriques, y compris ceux avec un nombre restreint de paramètres comme le modèle *mSUGRA*.

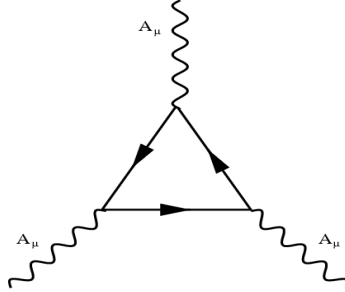


FIG. 1.12 – Diagrammes de Feynman d'anomalies triangulaires.

Le tableau 1.3 résume la composition en champs du modèle standard supersymétrique avant le mélange des higgsinos et des jauginos. Dans ce tableau, on remarque que le *MSSM* est une extension du modèle standard sans considérer la possibilité d'une masse pour les neutrinos. Des modèles, qui incluent le neutrino droit, existent mais ils ne seront pas évoqués dans cette thèse.

Les supermultiplets chiraux				
Composition du multiplet	Notations	Spin 0	Spin 1/2	Spin 1
Leptons et sleptons ($\times 3$ familles)	L	$(\tilde{\nu}, \tilde{e}_L)$	(ν, e_L)	
	$\bar{e}$	$\tilde{e}_R^*$	$e_R^\dagger$	
Quarks et squarks ($\times 3$ familles)	Q	$(\tilde{u}_L, \tilde{d}_L)$	(u_L, d_L)	
	$\bar{u}$	$\tilde{u}_R^*$	$u_R^\dagger$	
	$\bar{d}$	$\tilde{d}_R^*$	$d_R^\dagger$	
Higgs et higgsinos	H_u	(H_u^+, H_u^0)	$(\tilde{H}_u^+, \tilde{H}_u^0)$	
	H_d	(H_d^0, H_d^-)	$(\tilde{H}_d^0, \tilde{H}_d^-)$	
Les supermultiplets de jauge				
Bosons W et winos	W		$(\tilde{W}^\pm, \tilde{W}^0)$	$(W^\pm, W^0)$
Boson B et bino	B		$\tilde{B}^0$	B^0
Gluon et gluino	G		$\tilde{g}$	g

TAB. 1.3 – Les supermultiplets du *MSSM*.

Le superpotentiel.

Avec les notations du Tableau 1.3, on peut écrire le superpotentiel du *MSSM*, W_{MSSM} . En simplifiant l'écriture, c'est-à-dire en omettant les indices de jauge et de famille, on obtient :

$$W_{MSSM} = y_u \bar{u} Q H_u - y_d \bar{d} Q H_d - y_e \bar{e} L H_d + \mu H_u H_d \quad (1.51)$$

Avec y_u , y_d et y_e les matrices de Yukawa, 3×3 , dans l'espace des saveurs.

Elles permettent de donner les masses aux quarks et leptons, ainsi que les angles et phases de la matrice CKM après la brisure électrofaible. Le terme μ représente le couplage entre les deux doublets de Higgs, on dit que μ est relié à l'échelle de la brisure de la supersymétrie. A partir de ce potentiel et de la Formule 1.42, on peut réécrire le lagrangien décrivant toutes les interactions exceptées celles de jauge.

Une hypothèse importante a été faite sur ce potentiel. La forme du potentiel général comporte deux termes : le terme de la Formule 1.51 et un second terme qui brise la R-parité. La R-parité est une symétrie discrète usuellement utilisée qui est définie à partir du nombre baryonique B , du nombre leptonique L et du spin S d'une particule : $R = (-1)^{3B-3L+2S}$. Le nombre baryonique d'un système donné est défini comme le tiers du nombre de quarks moins le tiers du nombre d'antiquarks (pour un méson $B = 0$ et pour un hadron $B = \pm 1$) ou plus généralement le nombre B traduit la sensibilité à l'interaction forte. Et le nombre leptonique traduit la sensibilité à l'interaction faible : 1 pour un lepton, -1 pour un anti-lepton et 0 sinon. On peut montrer alors que pour une particule du modèle standard $R = 1$ et $R = -1$ pour leurs partenaires supersymétriques. Par exemple, le quark u vérifie : $B = 1/3$, $L = 0$ et $S = 1/2$, alors que le squark $\tilde{u}$ appartenant au même supermultiplet possède les mêmes nombres quantiques excepté le spin donc : $B = 1/3$, $L = 0$ et $S = 0$. La R-parité est supposée conservée dans l'expression du superpotentiel et plus généralement dans la plupart des modèles supersymétriques. Il existe néanmoins des études pour des modèles avec R-parité brisée [38, 39]. Les conséquences de cette supersymétrie sont importantes et font partie des motivations d'une théorie supersymétriques. Les particules supersymétriques sont en effet toutes produites par paires, autrement dit la particule supersymétrique la plus légère est stable. Elle est communément appelée *LSP* pour *Lightest Supersymmetric particle*. Celle-ci étant neutre électriquement et interagissant très peu avec la matière, elle échappe à la détection et fournit ainsi un bon candidat pour expliquer la matière noire de l'Univers.

Le spectre de masse

Pour déterminer les masses des particules supersymétriques après la brisure de la supersymétrie, il faut faire des hypothèses sur la valeur des couplages à hautes énergies et étudier leur évolution grâce aux équations du groupe de renormalisation. L'évolution des constantes de couplages $\sqrt{\frac{5}{3}}g_1$, g_2 et g_3 est donnée par la formule suivante :

$$\frac{dg_i}{dt} = \frac{1}{16\pi^2} b_i g_i^3 \quad (1.52)$$

Avec $t = \ln(Q/Q_0)$ où Q est l'énergie et Q_0 une énergie de référence.

La normalisation de $\sqrt{\frac{5}{3}}g_1$ plutôt que de g_1 a été choisie pour être en accord avec celle des théories de grandes unifications. C'est une simple convention qui peut être modifiée en remplaçant b_1 par $\sqrt{\frac{3}{5}}b_1$. Les valeurs des coefficients b_i sont (41/10, -19/6, -7) pour le modèle standard. Ces valeurs augmentent pour le *MSSM* avec la contribution des diagrammes de Feynman avec des particules supersymétriques dans les boucles et deviennent : (33/5, 1, -3). Ces valeurs permettent naturellement la convergence des constantes à une énergie de l'ordre 2×10^{16} GeV comme le montre la Figure 1.11.

Les masses des jauginos varient également en fonction des coefficients b_i et des constantes g_i :

$$\frac{dm_i}{dt} = \frac{1}{8\pi^2} b_i g_i^2 m_i \quad (1.53)$$

A partir des Formules 1.52 et 1.53, on montre que le rapport $\frac{m_i}{g_i^2}$ est constant et donc que : $\frac{m_1}{g_1^2} = \frac{m_2}{g_2^2} = \frac{m_3}{g_3^2}$. A l'échelle de Planck, c'est-à-dire à une énergie supérieure à l'unification des couplages, ces masses sont égales à une seule masse que l'on note $m_{1/2}$ et qui est la masse de tous les fermions supersymétriques. A l'échelle électrofaible, on peut calculer le rapport relatif entre ces masses grâce aux valeurs connues des constantes de couplages. Ainsi, $m_2 = 2m_1$ et $m_3 = 7m_1$ pour $Q \approx m_Z$.

Les neutralinos et les charginos Les neutralinos et charginos sont les états propres de masse des binos, winos et higgsinos après brisure de la symétrie électrofaible. Ils dépendent des paramètres m_1, m_2 et des paramètres des matrices 4×4 des termes de masse. Les masses des neutralinos et des charginos sont obtenues en diagonalisant ces deux matrices :

$$\begin{pmatrix} m_1 & 0 & -\cos(\beta)\sin(\theta_w)m_Z & \sin(\beta)\sin(\theta_w)m_Z \\ 0 & m_2 & \cos(\beta)\cos(\theta_w)m_Z & -\sin(\beta)\cos(\theta_w)m_Z \\ -\cos(\beta)\sin(\theta_w)m_Z & \cos(\beta)\cos(\theta_w)m_Z & 0 & -\mu \\ \sin(\beta)\sin(\theta_w)m_Z & -\sin(\beta)\cos(\theta_w)m_Z & -\mu & 0 \end{pmatrix} \quad (1.54)$$

$$\begin{pmatrix} 0 & 0 & m_2 & \sqrt{2}\cos(\beta)m_W \\ 0 & 0 & \sqrt{2}\sin(\beta)m_W & \mu \\ m_2 & \sqrt{2}\sin(\beta)m_W & 0 & 0 \\ \sqrt{2}\cos(\beta)m_W & \mu & 0 & 0 \end{pmatrix} \quad (1.55)$$

La matrice donnée dans la Formule 1.54 est exprimée dans la base $(\tilde{B}^0, \tilde{W}^0, \tilde{H}_d^0, \tilde{H}_u^0)$ alors que celle de la Formule 1.55 est exprimée dans la base $(\tilde{W}^+, \tilde{H}_u^+, \tilde{W}^-, \tilde{H}_d^-)$. La diagonalisation de ces matrices amène à des formules complexes qui ne seront pas détaillées.

Les fermions On distingue généralement deux cas pour les fermions [40]. Les fermions de première et deuxième génération ont des couplages de Yukawa qui peuvent être négligés par rapport aux couplages de Yukawa de la troisième famille. Si on ne néglige pas ces couplages, cela entraîne un mélange des états propres électrofaibles et cela donne des états propres de masse différents notés $\tilde{t}_{1,2}$, $\tilde{b}_{1,2}$ et $\tilde{\tau}_{1,2}$. Pour les squarks et sleptons de premières générations, on garde les notations liées à la chiralité de leur partenaire du modèle standard : $\tilde{q}_{R,L}$ avec $q = u, d, s, c$ et $\tilde{l}_{R,L}$ avec $l = e, \mu$.

Le sbottom est l'objet de l'analyse détaillée dans le Chapitre 4. En effet, dans certains cas, si le mélange entre les états propres électrofaibles $\tilde{b}_R$ et $\tilde{b}_L$ est suffisamment important alors le sbottom le plus léger, toujours noté $\tilde{b}_1$, est plus léger que le gluino. Ainsi, la désintégration du gluino en sbottom le plus léger est autorisée. C'est ce canal de recherche qui sera étudié dans le Chapitre 4. La matrice de mélange du sbottom est la suivante :

$$\begin{pmatrix} \tilde{m}_{bL}^2 & m_b(A_b - \mu \cdot \tan(\beta)) \\ m_b(A_b - \mu \cdot \tan(\beta)) & \tilde{m}_{bR}^2 \end{pmatrix} \quad (1.56)$$

Avec $\tilde{m}_{bL}^2$ et $\tilde{m}_{bR}^2$ les masses des états propres électrofaibles donnés par :

$$\tilde{m}_{bL}^2 = \tilde{m}_Q^2 + m_b^2 - \frac{1}{6}(2m_W^2 + m_Z^2)\cos(2\beta)\tilde{m}_{bR}^2 = \tilde{m}_d^2 + m_b^2 + \frac{1}{3}(m_W^2 - m_Z^2)\cos(2\beta) \quad (1.57)$$

Les termes $\tilde{m}_Q^2$ et $\tilde{m}_d^2$ sont les termes de masses utilisés dans la brisure douce de la supersymétrie et sont calculés à partir des équations du groupe de renormalisation (pour le calcul des états propres de masse du stop, on trouve : $\tilde{m}_Q^2$ et $\tilde{m}_u^2$, et pour le stau, on a : $\tilde{m}_L^2$ et $\tilde{m}_e^2$). A_b est le couplage trilinéaire du sbottom avec le boson de Higgs H_u^0 , qui apparaît également dans les termes de brisures explicites de la supersymétrie (on définit de la même façon A_t et A_τ).

Les termes non diagonaux sont proportionnels à la masse du quark b, c'est pourquoi les sbottoms sont dans la plupart des cas plus légers que les squarks des deux premières familles. Pour favoriser un sbottom léger, il faut augmenter les termes de mélange par rapport aux termes diagonaux. Le stau possède également une contribution en $\mu \cdot \tan(\beta)$ à son terme de mélange, alors que le stop possède une contribution en $\frac{\mu}{\tan(\beta)}$.

Les gluinos Ils possèdent un nombre de couleur. Ils ne peuvent par conséquent être associés aux autres jauginos qui possèdent tous les mêmes nombres quantiques. Leur masse dépend donc uniquement de $m_{1/2}$ et de la valeur de la constante unifiée, ou autrement dit de la constante de couplage à l'échelle de Planck : Q_{Planck} .

Exemple d'évolution des masses des particules supersymétriques La Figure 1.13 donne un exemple de l'évolution des masses calculée à partir des équations du groupe de renormalisation. Pour cette figure, on a pris $Q_0 = 2.5 \times 10^{16}$ GeV. On constate que la valeur de $m_{H_u^0} + \mu^2$ devient négative à l'échelle électrofaible, c'est ce qui implique la brisure de la symétrie électrofaible. Cette brisure devient alors une conséquence directe du modèle supersymétrique sans avoir été provoquée par l'ajout de terme de brisure explicite. Les lignes pontillées nommées H_u et H_d représentent respectivement les quantités $\sqrt{m_{H_u^0} + \mu^2}$ et $\sqrt{m_{H_d^0} + \mu^2}$. L'évolution de la masse des sleptons et des squarks est représentée par les lignes nommées sleptons et squarks, avec en lignes pointillées l'évolution des masses de la troisième famille.

La Figure 1.14 donne un exemple de spectre de masse. Sur cette figure, les charginos sont notés $\tilde{C}_{1,2}$ et les neutralinos $\tilde{N}_{1,2,3,4}$. Ces deux figures montrent que les sleptons sont toujours plus légers que les squarks et les gluinos, comme il a été mentionné plus tôt. Dans l'exemple de la Figure 1.14, on remarque également que la cascade de désintégration $\tilde{g} \rightarrow \tilde{b}_1 \rightarrow \tilde{\chi}_1^0$ est possible. Ce sont des cascades de ce type qui seront étudiés en détails dans le Chapitre 4.

1.4.5 Conclusion

La supersymétrie est actuellement un modèle théorique hypothétique qui n'a pas encore été observé expérimentalement. De nombreuses analyses au *LEP* puis au TeVatron ont permis de repousser les limites du spectre de masse sans jamais observer de preuves de l'existence de particules supersymétriques. L'échelle du TeV n'est néanmoins pas encore atteinte et tous les espoirs sont encore permis pour cette extension très élégante du modèle standard. Les motivations principales pour cette théorie sont avant tout théoriques. La supersymétrie est une théorie qui généralise l'algèbre de Poincaré en rendant possible la transformation de fermions en bosons et inversement. Elle est indispensable à la théorie des cordes, qui est un modèle intéressant pour une possible

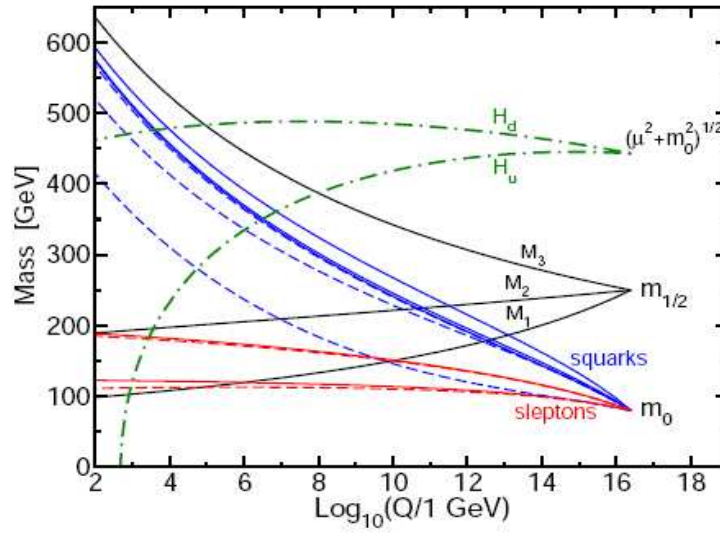
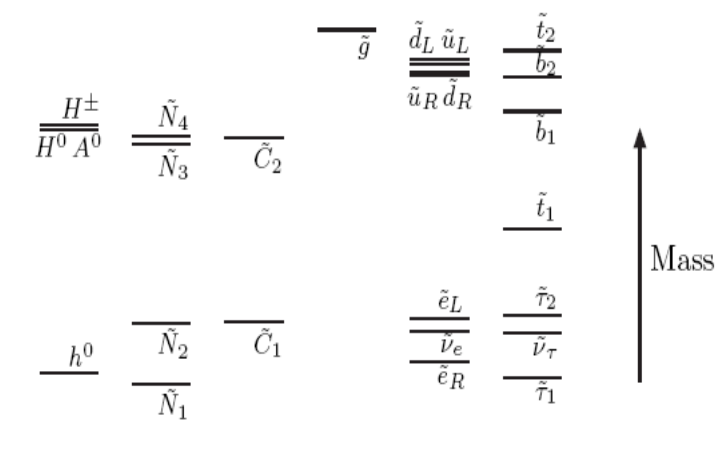


FIG. 1.13 – Evolution des masses des superpartenaires de l'échelle du GUT jusqu'à l'échelle électrofaible.

Grande Unification des quatre interactions fondamentales dans le même formalisme théorique. Elle permet de plus de régler le problème de l'ajustement fin introduit par le mécanisme de Higgs, ce dernier étant un complément indispensable au modèle standard pour expliquer les masses des particules. Enfin, elle induit directement la convergence des constantes de couplage qui laisse à penser qu'une unification des trois interactions du modèle standard est possible. Expérimentalement, la supersymétrie permet d'apporter une explication plausible à la matière noire de l'Univers. Elle favorise de plus l'existence d'un boson de Higgs léger, ce dernier étant toujours actuellement recherché auprès des collisionneurs.

FIG. 1.14 – Un exemple de spectre de masse pour un jeu de paramètres du *MSSM*.

Chapitre 2

Le dispositif expérimental

Sommaire

2.1	Introduction	52
2.2	La chaîne de production, d'accélération et de collision des protons et des antiprotons	53
2.2.1	Généralité	53
2.2.2	Principes de fonctionnement	56
2.3	Le détecteur DØ	62
2.3.1	Généralités	63
2.3.2	Le trajectographe interne	64
2.3.3	Le calorimètre	69
2.3.4	Le spectromètre à muons	77
2.3.5	La mesure de la luminosité	80

2.1 Introduction



FIG. 2.1 – Vue aérienne du site de Fermilab avec au premier plan, le *Main Injector* et au second plan le TeVatron

Le laboratoire Fermilab [41], pour *Fermi National Accelerator Laboratory*, regroupe de nombreuses expériences de physique des hautes énergies. On peut citer, entre autres, les expériences MINOS [42], MiniBooNE [43], CDF [44] et DØ [45] pour le champ de la physique des particules. Ces deux dernières, CDF pour *Collider Detector at Fermilab* et DØ pour le nom d'un des points de croisement des faisceaux, sont les noms des deux détecteurs situés sur l'accélérateur TeVatron. Cet accélérateur de protons p et d'antiprotons $\bar{p}$ sera décrit dans la Section 2.2.

Ce vaste laboratoire d'environ 28 km^2 se trouve à 60 km de Chicago (Illinois, Etats-Unis). Il a été le siège de grandes découvertes pour la physique des hautes énergies. Il a notamment permis de découvrir les dernières briques du modèle standard : la quark bottom en mai-juin 1977, le quark top en février 1995 [20, 21] et finalement la première observation directe du neutrino tau en juillet 2000. Ces ultimes vérifications du modèle standard ont véritablement posé les objectifs du TeVatron pour la seconde période de prise de données, appelée *Run II*. Le *Run II* correspond à la période de prise de données de mars 2001 à aujourd'hui, par opposition à la période de prise de données entre 1992 et 1996, appelée *Run I*. Le laps de temps entre le *Run I* et le *Run II* a permis d'améliorer la chaîne d'accélération et les deux détecteurs CDF et DØ. Les principaux objectifs du *Run II* sont la physique du top, la physique des hadrons beaux, l'étude des processus de *QCD* et d'interaction électrofaible, la recherche du boson de Higgs et de nouvelle physique... Pour satisfaire au mieux ces objectifs, deux paramètres importants (voir la Section 2.2.1) ont été améliorés du *Run I* au *Run II*. L'énergie dans le centre de masse est passée de 1.8 TeV à 1.96

TeV. La luminosité instantanée a également été augmentée, ainsi 125 pb^{-1} de luminosité intégrée a été collectée au *Run I* alors que plus de 1 fb^{-1} de données sont exploitables au *Run IIa* de 2001 à avril 2006 (*Run IIa* et *Run IIb* sont deux périodes de prise de données qui seront décrites plus en détails dans le Chapitre 3). Les notions de luminosité instantanée et intégrée seront détaillées dans la Section 2.2.1.

Les données utilisées pour les études de cette thèse correspondent à la période de prise de données appelée *Run IIa* (mars 2001 à avril 2006) et ont été enregistrées par le détecteur DØ, qui sera plus largement décrit dans la Section 2.3. Ce détecteur permet de mesurer les coordonnées du quadri-vecteur énergie-impulsion des différentes particules qui peuvent être produites lors d'une collision $p\bar{p}$: électrons, photons, particules hadroniques et muons. Il est constitué de trois sous-détecteurs. Du plus proche du faisceau vers l'extérieur, on trouve le détecteur de traces, le calorimètre uranium/argon liquide et un spectromètre pour détecter les muons. DØ a subi plusieurs modifications entre chaque période de prise de données afin de profiter au mieux de l'augmentation de l'énergie et de la luminosité.

2.2 La chaîne de production, d'accélération et de collision des protons et des antiprotons

2.2.1 Généralité

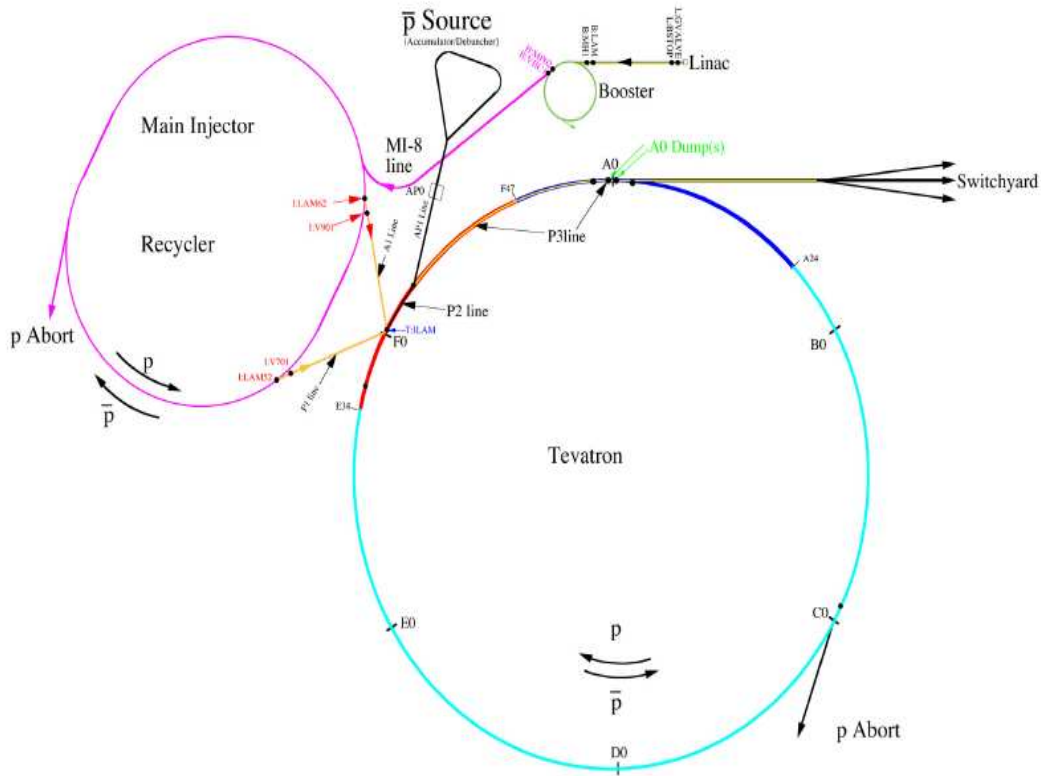


FIG. 2.2 – Schéma de la chaîne d'accélération des protons et des antiprotons.

Le TeVatron [46, 47], pour *tera electron volt*, est un accélérateur circulaire de protons et d'antiprotons entrant en collision en deux endroits précis de l'anneau où se trouvent les détecteurs CDF et DØ. Plusieurs types d'accélérateurs ainsi que de particules accélérées sont envisageables pour étudier les constituants élémentaires de la matière. Le choix effectué à Fermilab sera discuté dans la Section 2.2.1. Ce dispositif de création et d'accélération de particules est complexe et comporte de nombreuses étapes, qui seront détaillées dans la Section 2.2.2. L'ensemble de cette chaîne d'accélération permet aujourd'hui au TeVatron d'être l'accélérateur avec la plus haute énergie dans le centre de masse. Cette énergie dans le centre de masse est de 1.96 TeV (voir la Section 2.2.1).

Choix du dispositif expérimental

Le but d'un accélérateur, ou plus précisément d'un collisionneur, de particules est de produire un maximum d'interactions avec la plus haute énergie possible. Ainsi, de nouvelles particules élémentaires, différentes des particules incidentes, peuvent être produites. L'accélérateur devient de cette façon un excellent outil pour sonder la matière [48].

Pour obtenir la plus haute énergie possible, il faut accélérer les particules initiales et produire une collision. Deux types d'accélérateurs sont possibles : un accélérateur linéaire ou circulaire. Le TeVatron fait partie de la seconde famille. Le principal avantage d'un tel accélérateur par rapport à un accélérateur linéaire est la dimension de celui-ci. Son principal désavantage est l'énergie perdue par rayonnement synchrotron. L'énergie perdue est en effet inversement proportionnelle à la masse au carré de la particule accélérée. Une particule relativiste de masse m et d'énergie E , subissant une accélération centripète dans un anneau de rayon de courbure R , rayonne à chaque tour une énergie ΔE :

$$\Delta E = \frac{E^4}{m^4 R} \quad (2.1)$$

Pour cette raison, le TeVatron utilise des protons et des antiprotons contrairement au LEP (*Large Electron Positron*), qui utilisait des électrons et des positrons à plus basse énergie. L'ordre de grandeur du rapport des énergies perdues dans le cas d'un collisionneur $p\bar{p}$ par rapport à un collisionneur e^+e^- est donnée par la Formule 2.2.

$$\frac{\Delta E(p)}{\Delta E(e)} = \left(\frac{m_e}{m_p}\right)^4 \approx 10^{-13} \quad (2.2)$$

Le TeVatron a une seconde particularité par rapport au LEP. En effet, l'utilisation de protons, particules composites, plutôt que des électrons ne permet pas de connaître avec précision l'énergie de la collision partonique mais seulement l'énergie transverse au faisceau. L'énergie de l'interaction dure n'a donc pas une valeur fixée par le dispositif expérimental, mais balaie plutôt un intervalle d'énergie. C'est la particularité des machines hadroniques, qui sont par conséquent orientées "découvertes", par opposition aux machines leptoniques (électrons-positrons) orientées mesures de précision.

De plus, le TeVatron utilise des collisions protons-antiprotons. L'avantage d'utiliser l'antiparticule correspondante au proton est le coût de l'accélérateur et l'augmentation de la section efficace de certains processus (ceci est valable jusqu'à une énergie dans le centre de masse d'environ 3 TeV [49]). Le même dispositif d'aimants est en effet utilisé pour accélérer dans un sens les protons et dans l'autre les antiprotons (seule la charge différenciant les deux particules), et ceci contrairement au futur collisionneur LHC (*Large Hadron Collider*) dont les deux faisceaux seront composés

de protons. Ce dispositif plus complexe a néanmoins l'avantage d'obtenir une luminosité (voir la Section 2.2.1) beaucoup plus importante, et à 14 TeV (énergie dans le centre de masse prévue pour le LHC) une collision protons-antiprotons a une section efficace différentielle du même ordre de grandeur qu'une collision protons-protons. En effet, à très haute énergie la section efficace différentielle est dominée par les processus gluon-gluon et gluon-quark, alors qu'à plus basse énergie ce sont les processus quark-anti-quarks qui prédominent.

Finalement, le TeVatron n'est pas un collisionneur sur cible fixe. Les faisceaux de protons et d'antiprotons possèdent tous les deux une énergie de 980 GeV. Cette méthode permet d'obtenir une énergie dans le centre de masse plus importante que pour les expériences sur cible fixe. L'inconvénient d'un tel dispositif est la section efficace de collision beaucoup plus faible.

Quelques grandeurs caractéristiques

Deux grandeurs précédemment évoquées sont primordiales pour définir les performances d'un collisionneur. La première est l'énergie dans le centre de masse. Plus elle est importante, plus la création de particules massives est aisée. Cette énergie dans le centre de masse est donnée par la formule $\sqrt{s} = \sqrt{4E_1E_2}$, où E_1 et E_2 sont les énergies des faisceaux de protons et d'antiprotons. Au TeVatron, on a $E_1 = E_2 = 980$ GeV, donc $\sqrt{s} = 1.96$ TeV. L'énergie réellement utilisée pour l'interaction n'est en fait que la fraction d'énergie emportée par chaque parton participant à cette interaction, c'est-à-dire : $\sqrt{s} = \sqrt{4x_1E_1x_2E_2}$, où x_1 et x_2 sont les fractions d'énergie emportées par le parton issu du proton et celui issu de l'antiproton.

La seconde grandeur caractéristique est la luminosité instantanée qui permet de définir le nombre d'interactions en un temps donné et pour une section efficace donnée :

$$\mathcal{L} = \frac{1}{\sigma} \frac{dN}{dt} \quad (2.3)$$

où dN est le nombre d'interactions dans l'intervalle de temps dt et σ est la section efficace du processus étudié. Cette quantité s'exprime en $cm^{-2}s^{-1}$.

Ces deux quantités permettent de définir les objectifs de recherche du TeVatron dans le temps. En effet, dans le cas de la nouvelle physique, plus l'énergie est importante, plus la section efficace de production de nouvelles particules est importante. Et, plus la luminosité est grande, plus le nombre de particules produites en un temps donné est grand. On peut définir ensuite la luminosité intégrée, qui définit la statistique disponible pour le programme de physique :

$$L = \int_T \mathcal{L}(t) dt \quad (2.4)$$

où T est la période d'acquisition ou plus précisément la durée pendant laquelle des collisions sont enregistrées. Cette quantité s'exprime en cm^{-2} , mais on emploie plus souvent l'unité pb^{-1} ou fb^{-1} avec $1 \text{ barn} = 10^{-24} cm^2$.

Dans le cas du TeVatron, c'est-à-dire d'un collisionneur protons-antiprotons, la luminosité instantanée est fonction de la façon dont interagissent les deux faisceaux, c'est-à-dire de leur structure et de la fréquence des croisements. La formule suivante exprime cette luminosité instantanée :

$$\mathcal{L} = \frac{fnN_pN_{\bar{p}}}{A} \quad (2.5)$$

où N_p et $N_{\bar{p}}$ sont le nombre de protons et d'antiprotons par paquets (un faisceau est constitué de plusieurs paquets, 36×36 au *Run II*). f est la fréquence de croisement des faisceaux, n est le nombre

de paquets par faisceau et A est l'aire de la section orthogonale commune aux deux faisceaux. Les deux faisceaux étant différents, on peut définir deux sections orthogonales différentes pour le faisceau de protons et d'antiprotons. De plus, on peut représenter la répartition des protons ou des antiprotons dans chaque faisceau par une gaussienne. Ainsi, A s'exprime en fonction de σ_p et de $\sigma_{\bar{p}}$ les largeurs des deux gaussiennes. La luminosité du TeVatron est finalement donnée par la formule suivante :

$$\mathcal{L} = \frac{fnN_pN_{\bar{p}}}{2\pi(\sigma_p^2 + \sigma_{\bar{p}}^2)} F\left(\frac{\sigma_l}{\beta^*}\right) \quad (2.6)$$

où F est un facteur de forme qui est fonction de la longueur du paquet, σ_l , et de la focalisation longitudinale au point de collision, caractérisée par β^* . Ce facteur de forme est en fait un rapport sans dimension.

Le Tableau 2.1 donne des valeurs caractéristiques pour les variables évoquées plus haut ainsi que pour t_{paquet} , le temps entre deux paquets et $\mathcal{L}$ est une valeur caractéristique de la luminosité instantanée sous ces conditions.

N_p	$N_{\bar{p}}$	$E_1 = E_2$	$t_{\text{croisement}}$	n	σ_l	β^*	F	$\mathcal{L}$
$2,7 \times 10^{11}$	7×10^{10}	980 GeV	396 ns	36	37 cm	35 cm	0,7	$2 \times 10^{32} \text{ cm}^{-2} \text{ s}^{-1}$

TAB. 2.1 – Grandeurs caractéristiques pour le calcul de la luminosité au TeVatron.

Le paramètre limitant de cette formule est le nombre d'antiprotons par paquet. La production des antiprotons sera détaillée dans la Section 2.2.2. Toutes les étapes de la chaîne d'accélération seront également détaillées dans la Section 2.2.2, le but de cette chaîne étant d'optimiser tous ces paramètres afin d'obtenir la meilleure luminosité possible. L'évolution de cette luminosité instantanée au cours du *Run II* jusqu'à août 2006 est représentée sur la Figure 2.3. Chaque période d'arrêt du détecteur ou de l'accélérateur est marquée par l'absence de point. On remarque par exemple la transition du *Run IIa* au *Run IIb* d'avril 2006 à juin 2006.

2.2.2 Principes de fonctionnement

Comme on peut le voir sur la Figure 2.2, Fermilab possède toute une chaîne d'accélérateurs qui permettent à la fois de créer les particules et ensuite de les accélérer. De nombreuses étapes sont nécessaires avant la collision des faisceaux de protons et d'antiprotons aux points $D\bar{O}$ et $B\bar{O}$ du TeVatron. Les étapes intervenant dans la création des protons et leur accélération avant l'entrée dans le TeVatron seront détaillées dans un premier temps. La chaîne de production des antiprotons, plus complexe, sera décrite dans un deuxième temps. Pour finir, on donnera les détails de la dernière accélération des faisceaux en vue des collisions au sein du TeVatron [50].

Création et accélération des paquets de protons

La production de protons (et d'antiprotons) à une énergie de 980 GeV commence par une simple bouteille d'hydrogène (6,9 litres à une pression d'environ 138 bar, l'équivalent de 5×10^{25} atomes d'hydrogène ou un an de production de protons). A partir de cette bouteille, trois étapes permettent de créer un premier faisceau de protons à 8 GeV.

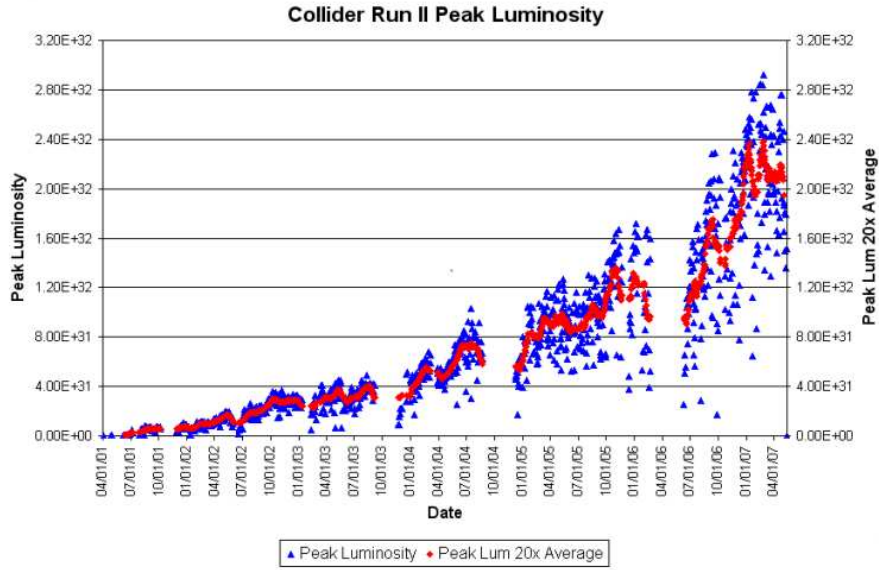


FIG. 2.3 – Evolution de la luminosité instantanée initiale en fonction du temps.

La première étape consiste en la préparation d'un faisceau d'ions H^- en ionisant le gaz d'hydrogène. Ce faisceau est ensuite accéléré à une énergie de 750 keV par un accélérateur de type Cockroft-Walton en utilisant simplement un champ électrostatique. La Figure 2.4 donne le schéma de l'ionisation à partir d'une plaque de césium et d'un champ électrique, ainsi qu'une photographie du pré-accelérateur (l'accélérateur de type Cockroft-Walton). Les ions H^- pénètrent ensuite dans le *LINAC*.

Le *LINAC*, pour *Linear Accelerator*, est un accélérateur linéaire de 130 mètres de long qui accélère le faisceau d'ions H^- de 750 KeV à 400 MeV à l'aide de cavités radiofréquences accélératrices. Un schéma des cavités accélératrices est montré sur la Figure 2.5 et une illustration animée du principe de fonctionnement de tel accélérateur peut être trouvée en [51]. Le faisceau d'ions est accéléré entre les cavités et garde une vitesse constante dans chacune des cavités sous vide. Les cavités sont de plus en plus longues car la vitesse des particules augmente au fur et à mesure. Le faisceau d'ions H^- entre finalement dans le *BOOSTER* avec une énergie de 400 MeV pour la dernière étape de production de protons à une énergie de 8 GeV.

Le *BOOSTER* est le premier synchrotron des six que compte Fermilab. En entrant dans ce synchrotron, les ions H^- traversent une feuille de carbone et perdent leurs électrons. On passe ainsi à un faisceau de protons. Ce faisceau est accéléré dans le tunnel d'une circonférence de 475 mètres jusqu'à une énergie de 8 GeV en moins de 67 ms. Cet accélérateur est constitué de 96 aimants le long de l'anneau, afin de courber et focalisé le faisceau, ainsi que de cavités radiofréquences accélératrices pour augmenter l'énergie des protons. Une fois l'énergie de 8 GeV atteinte, le faisceau peut être injecté dans le *Main Injector* pour la dernière accélération avant l'entrée dans le TeVatron, mais ce faisceau peut également être orienté vers d'autres expériences comme les expériences *MiniBooNE* ou *MINOS*. Une photographie du *BOOSTER* proche du bâtiment principal de Fermilab ainsi qu'une autre du *Main Injector* au premier plan sont montrées sur la Figure 2.6.

Le *Main Injector* ou injecteur principal [52] est l'avancée majeure du *Run II* par rapport au

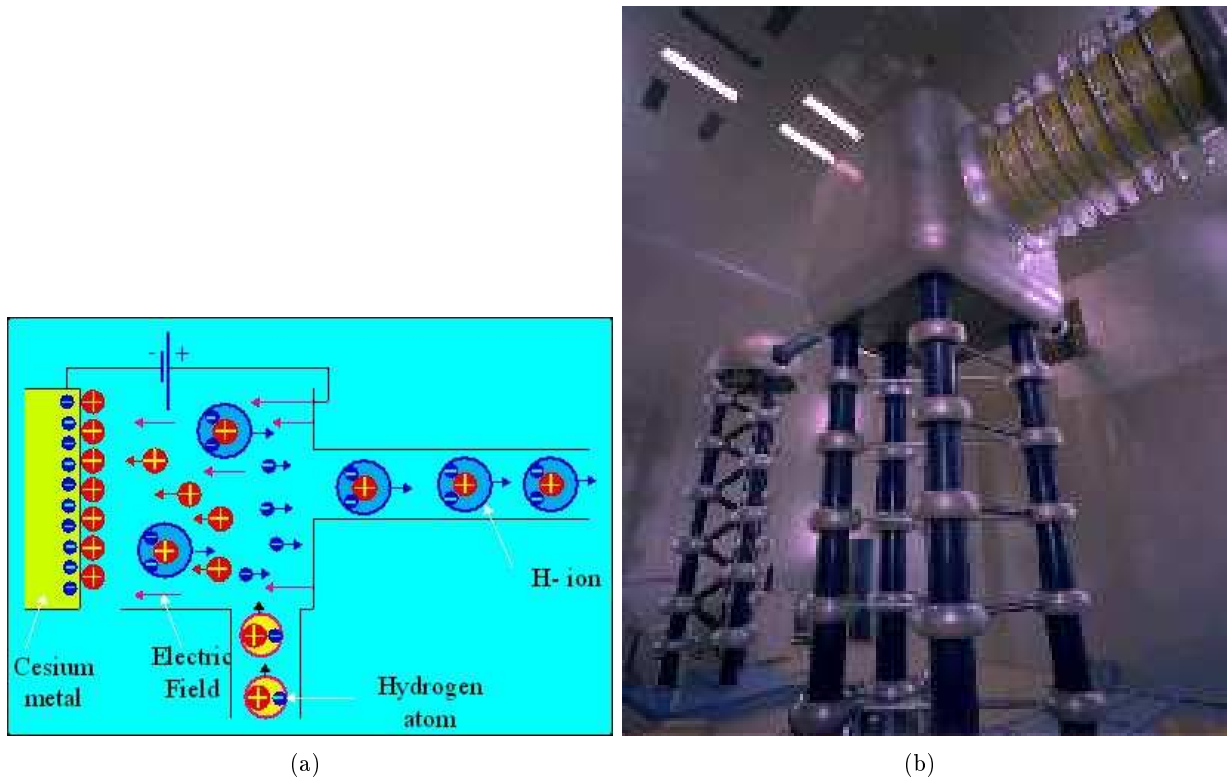


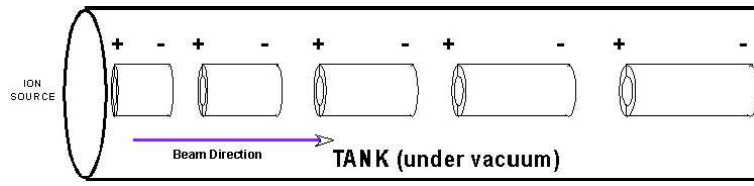
FIG. 2.4 – Production des ions H^- (a) et l'accélérateur de type Cockcroft-Walton (b).

Run I. Il a été conçu en 1998 afin d'augmenter la luminosité du TeVatron et d'éviter de dégrader les données enregistrées par les deux détecteurs du TeVatron. Au *Run I*, l'anneau principal ou *Main Ring* jouait en effet le rôle joué actuellement par l'injecteur principal, c'est-à-dire accélérer les protons de 8 GeV à 150 GeV. De plus, l'anneau principal joue également un rôle important dans la production des antiprotons (voir le paragraphe suivant). Le *Main Injector* est un synchrotron de forme elliptique (voir le premier plan de la Figure 2.1) qui peut accepter les protons (ou antiprotons) à une énergie de 8 GeV. Un système d'aimants permet de stabiliser le faisceau par paquets de protons, créés par résonance avec les cavités radiofréquences du synchrotron. Le faisceau n'est donc pas continu et plusieurs paquets circulent dans le même temps à l'intérieur de l'anneau. Une fois les paquets stabilisés et correctement focalisés, l'accélération peut commencer grâce à l'énergie apportée par les radiofréquences pour atteindre finalement une énergie de 150 GeV. On obtient finalement un faisceau compact de dimension réduite dans le plan transverse pour maximiser la luminosité. Ce faisceau est injecté dans le TeVatron pour la dernière étape avant les collisions (voir le dernier paragraphe décrivant l'accélération finale).

Création et accélération des paquets d'antiprotons

La production d'un faisceau d'antiprotons est beaucoup plus complexe et c'est l'étape qui limite essentiellement la luminosité.

La production et la stabilisation d'un faisceau d'antiprotons à une énergie de 8 GeV néces-

FIG. 2.5 – Schéma de principe du *LINAC*.FIG. 2.6 – Le *BOOSTER*.

site essentiellement trois étapes. La première utilise un faisceau de protons de 120 GeV issu de l'injecteur principal toutes les 1.5 secondes. Ce faisceau est envoyé sur une cible fixe de nickel (voir la Figure 2.7). L'énergie de la collision permet de créer de nombreuses particules différentes. En effet, pour un million de protons heurtant la cible, seulement 20 antiprotons sont finalement recueillis dans l'accumulateur (*accumulator*). Ce flux de particules s'échappe de la cible fixe avec une large dispersion angulaire et nécessite une focalisation. La lentille de lithium représentée sur la Figure 2.7 permet de créer un faisceau de particules. La plupart de ces particules est ensuite filtrée à l'aide d'un champ magnétique fourni par un aimant dipolaire de 1.5 T agissant comme un spectromètre de masse et sélectionnant les charges positives afin d'isoler les antiprotons des autres particules.

La structure du faisceau de protons est une structure en paquets, ce qui implique que la structure du faisceau d'antiprotons résultant est également découpée en paquets. Une fois le faisceau d'antiprotons créé, celui-ci est envoyé pour la deuxième étape du processus vers le *debuncher* (*bunch* signifiant paquet en anglais, le terme *debuncher* désigne la destruction de la structure en paquets du faisceau). L'objectif de cette deuxième étape est alors de réduire la dispersion en énergie du faisceau, tout en augmentant la dispersion en temps. On obtient ainsi un faisceau continu et homogène en énergie. Comme indiqué sur le schéma de la Figure 2.7, les particules les plus énergétiques (orbite la plus large) sont ralenties par le *debuncher* alors que les particules les moins énergétiques (orbite la plus petite) sont accélérées. Ce processus prend 100 millisecondes. Le reste du temps imparti, le faisceau restant 1.5 secondes dans le *debuncher*, est utilisé pour un processus

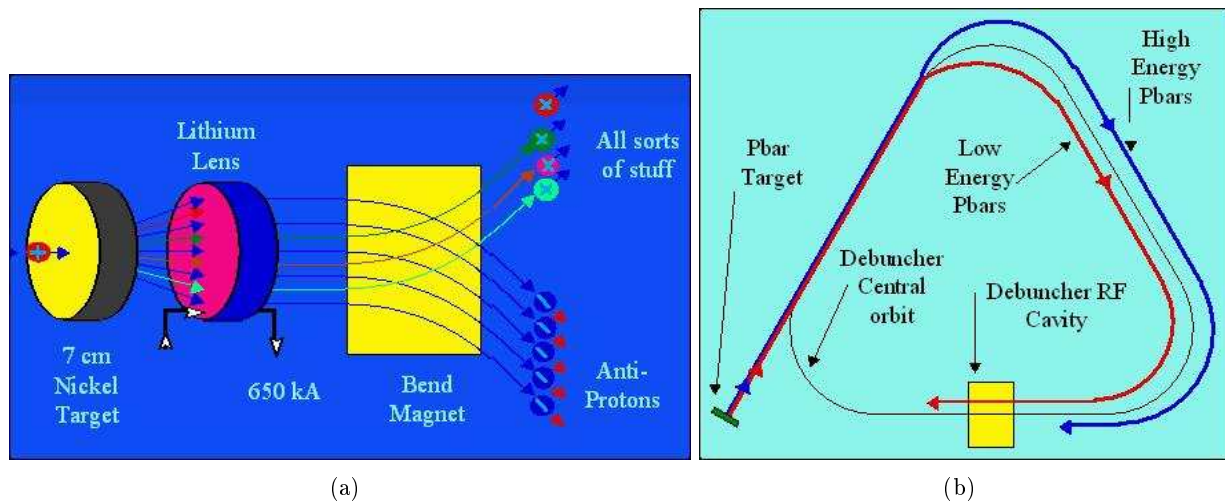


FIG. 2.7 – Collision des protons sur une cible fixe de nickel (a) et schéma du *debuncher* (b).

de refroidissement stochastique. Ce processus, qui commence dans le *debuncher* pour se terminer dans l'accumulateur, permet finalement de réduire la répartition spatiale et angulaire du faisceau. Cette étape est une nouvelle fois primordiale pour augmenter la luminosité des futures collisions.

Finalement, le faisceau d'antiprotons est envoyé dans l'accumulateur. Ce synchrotron triangulaire permet de stocker les antiprotons jusqu'à en obtenir un nombre suffisant pour les collisions. Cette étape peut durer plusieurs heures et a lieu entre deux séries de prises de données (appelées *store*). De plus, un nouvel accélérateur a été créé depuis décembre 2004 pour la gestion des antiprotons : le recycleur ou *recycler*. Il se trouve juste au dessus de l'anneau de l'injecteur principal. Son objectif initial était de stocker les antiprotons en excès dans l'accumulateur pour éviter des instabilités et de récupérer des antiprotons disponibles en fin de *store*. Seul le premier objectif est en fait respecté. Il permet ainsi de pallier la difficulté et la lenteur du processus de production d'antiprotons. Son utilisation a permis d'augmenter la luminosité.

Avant d'être injecté dans le TeVatron pour la dernière étape, le faisceau d'antiprotons est accéléré à une énergie de 150 GeV dans l'injecteur principal de la même façon que le faisceau de protons.

Accélération finale et collisions au sein du TeVatron

Caractéristiques générales du TeVatron Le TeVatron est un synchrotron quasiment circulaire de 1 km de rayon. C'est actuellement l'accélérateur qui permet d'avoir l'énergie dans le centre de masse la plus haute du monde 1.96 TeV, en attendant le début de l'expérience LHC (14 TeV). Mis en service en 1983, il est également le premier synchrotron à utiliser la supraconductivité pour faire fonctionner ses aimants. Ainsi à 1.96 TeV, les 722 dipôles sont refroidis à une température de 4.3 K par un système cryogénique à l'hélium liquide. Les dipôles, parcourus par un courant de 4350 A, produisent un champ de 4.2 T (8.4 T pour les dipôles supraconducteurs du LHC). Ils permettent de courber la trajectoire des faisceaux de particules à 980 GeV. L'avantage de cette technique est qu'à très basse température, le courant ne rencontre quasiment aucune résistance. Par conséquent, les bobines des aimants dissipent très peu de puissance électrique. De plus, 180

quadripôles contrôlent la taille transverse des faisceaux et 8 cavités radiofréquences accélèrent le faisceau de protons dans un sens et le faisceau d'antiprotons dans le sens opposé de 150 GeV à 980 GeV.

Collisions et prises de données Une fois, cette énergie atteinte, les faisceaux se rencontrent en deux points précis pour la détection des particules, qui sera détaillée dans la Section 2.3. Les deux faisceaux ne sont pas continus mais regroupés en trois "super-paquets" séparés de $2.6 \mu\text{s}$. Ces "super-paquets" sont composés eux-mêmes de 12 paquets de protons ou d'antiprotons séparés de 396 ns. Chacun des paquets de protons et d'antiprotons se rencontre donc précisément toutes les 396 ns aux points DØ et BØ (voir le schéma de la Figure 2.8).

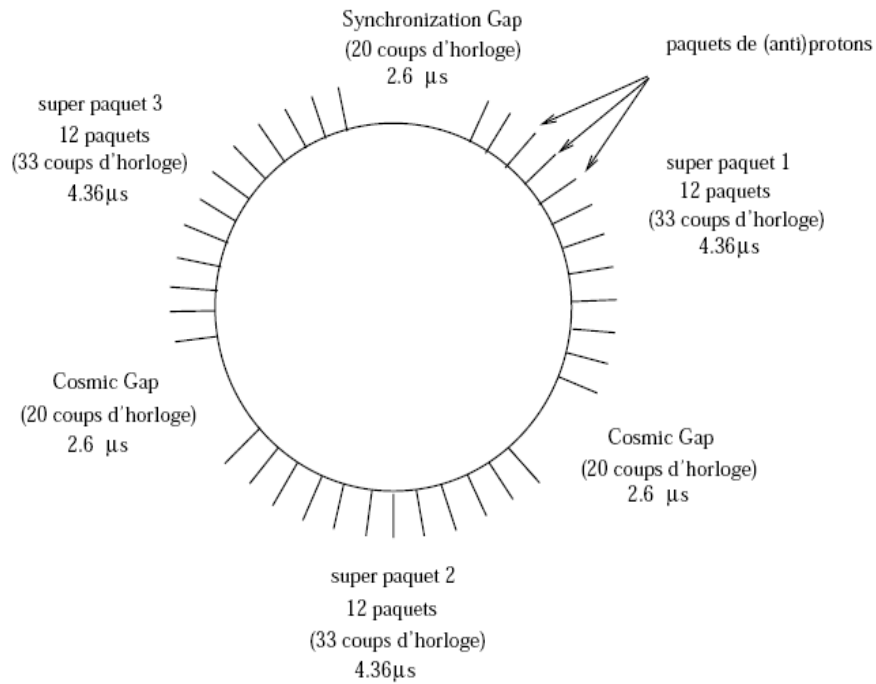


FIG. 2.8 – Schéma de la structure en paquets des faisceaux de protons et d'antiprotons.

Les interactions avec le gaz résiduel contenu dans le tube à vide diminuent le temps de vie des faisceaux et limitent les prises de données. Il est donc nécessaire de reproduire régulièrement des faisceaux et de procéder à de nouvelles injections. Les périodes de prise de données entre deux injections de faisceaux sont appelées des *stores* et durent de 12 à 16h environ. Ces périodes de temps sont découpées en *runs*, qui correspondent à des périodes de prise de données variant de 2h à haute luminosité jusqu'à 4h à faible luminosité. Entre chaque *run*, les menus de déclenchement peuvent être modifiés pour s'adapter à l'évolution de la luminosité.

Résumé des différentes étapes Les Tableaux 2.2 et 2.3 résument les différentes étapes de production, d'accélération et finalement de collisions des faisceaux de protons et d'antiprotons, ainsi que les différents accélérateurs utilisés. Deux soucis ont toujours été pris en compte pour chacune de ces étapes : l'augmentation de l'énergie et l'obtention des meilleurs faisceaux possibles pour une luminosité maximale.

Accélérateur	Fonction	Energie en sortie
<i>Preacc</i>	Création du faisceau de H^-	750 KeV
<i>LINAC</i>	Accélération des H^-	400 MeV
<i>BOOSTER</i>	Production des protons	8 GeV
<i>Main Injector</i>	Accélération des protons	150 GeV
TeVatron	Accélération finale et collisions	980 GeV

TAB. 2.2 – De la production des protons à la collision.

Accélérateur	Fonction	Energie en sortie
Cible de nickel	Création des antiprotons	< 120 GeV
<i>Debuncher</i>	Uniformisation du faisceau	8 GeV
<i>Accumulator</i>	Uniformisation et stockage	8 GeV
<i>Main Injector</i>	Accélération des antiprotons	150 GeV
TeVatron	Accélération finale et collisions	980 GeV

TAB. 2.3 – De la production des antiprotons à la collision.

2.3 Le détecteur DØ

L'expérience DØ a été proposée en 1983. L'objectif de ce détecteur [53] est d'étudier les collisions protons-antiprotons du TeVatron. Les premières collisions enregistrées par le détecteur ont eu lieu en 1992 à une énergie dans le centre de masse de 1,8 TeV. Depuis 2001, l'énergie dans le centre de passe a augmenté jusqu'à 1,96 TeV. Cette augmentation d'énergie, ainsi que les améliorations de la chaîne d'accélération évoquées dans la Section 2.2, ont nécessité plusieurs modifications du détecteur pour s'adapter au mieux aux nouvelles conditions et pour élargir les ambitions du programme de physique. C'est précisément le détecteur utilisé pendant cette période (*Run IIa* : 2001-2006) qui sera décrit dans les sections suivantes. Le détecteur DØ a été conçu pour satisfaire aux conditions suivantes :

- Mesurer précisément l'énergie des jets (ils sont la conséquence de l'hadronisation des partons. Ce phénomène sera décrit dans la Section 2.3.3) à l'aide d'un calorimètre suffisamment segmenté et uniforme.
- Détermination de l'énergie transverse manquante pour accéder à l'information sur les particules interagissant faiblement avec le détecteur.
- Bonne couverture angulaire pour l'identification des électrons et des muons

L'expérience DØ a en effet pour objectif principal l'étude des états finaux à haute masse invariante et à haute énergie transverse (quark top ou supersymétrie par exemple). La multiplicité en leptons et en jets peut être importante pour de telles topologies.

Quelques généralités sur le détecteur seront données dans la Section 2.3.1 avant de détailler les trois sous-détecteurs : le détecteur de traces chargées (Section 2.3.2), le calorimètre électromagnétique et hadronique (2.3.3) et le spectromètre à muons (Section 2.3.4).

2.3.1 Généralités

Création de nouvelles particules

Le TeVatron permet d'avoir une collision d'un faisceau de protons contre un faisceau d'anti-protons au point DØ de l'anneau. Chacun des faisceaux a une énergie de 980 GeV. Cette énergie est en fait l'énergie maximale que peut emporter le parton (ou éventuellement les partons) interagissant au moment de la collision. En effet, le proton (ou l'antiproton) est une structure complexe formée essentiellement de quarks et d'un nuage gluonique.

L'énergie du faisceau est la somme de l'énergie cinématique (T) et de la "masse au repos" de la particule ($E_0 = mc^2$). Cette relation est donnée par la formule 2.7.

$$E_T = T + E_0 \quad (2.7)$$

On peut également exprimer cette énergie comme le produit de γ , "facteur relativiste", par "la masse au repos".

$$E_T = \gamma \times E_0 \quad (2.8)$$

Avec :

$$\gamma = \sqrt{\frac{1}{1 - \frac{v^2}{c^2}}} \quad (2.9)$$

Ainsi, en connaissant l'énergie totale de la particule et sa masse (la masse du proton est de 938 MeV), on peut connaître γ et finalement la vitesse de la particule. Par exemple, un proton ayant une énergie de 980 GeV possède une vitesse de 99,99995% de la vitesse de la lumière. Par conséquent même si les partons interagissant lors de la collision emportent seulement une fraction des 980 GeV, ces particules sont considérées comme relativistes. A de telles énergies, les effets les plus spectaculaires de la relativité restreinte ne sont plus négligeables. Les interactions inélastiques sont notamment possibles. Ces interactions, dont l'état final est différent de l'état initial, sont précisément celles étudiées par les collisionneurs.

Les collisionneurs comme le TeVatron permettent par ce principe de créer de nouvelles particules. Le détecteur est alors l'outil pour étudier ces nouveaux états de la matière instables à basse énergie. En détectant les produits de désintégrations de ces particules instables et en leur associant un quadrivecteur énergie-impulsion, on peut remonter aux produits initiaux de la collision et ainsi caractériser l'événement étudié. Plusieurs sous-détecteurs sont nécessaires pour identifier et mesurer les caractéristiques des particules stables (ou à long temps de vie). Le schéma de la Figure 2.9 montre ces particules et les sous-détecteurs utilisés pour accéder aux informations propres à ces particules.

Le système de détection des traces, décrit dans la Section 2.3.2, permet de mesurer l'impulsion ou l'énergie des électrons, des muons ou des hadrons chargés. Le calorimètre électromagnétique, deuxième couche en partant du tube à vide, mesure le dépôt d'énergie des électrons et des photons, ainsi qu'une fraction de l'énergie des jets. L'autre fraction de l'énergie est finalement mesurée par le calorimètre hadronique (voir la Section 2.3.3). Enfin, les muons déposent très peu d'énergie dans ces sous-détecteurs. Cette énergie est essentiellement mesurée par les chambres à muons situées complètement à l'extérieur du détecteur (voir la Section 2.3.4).

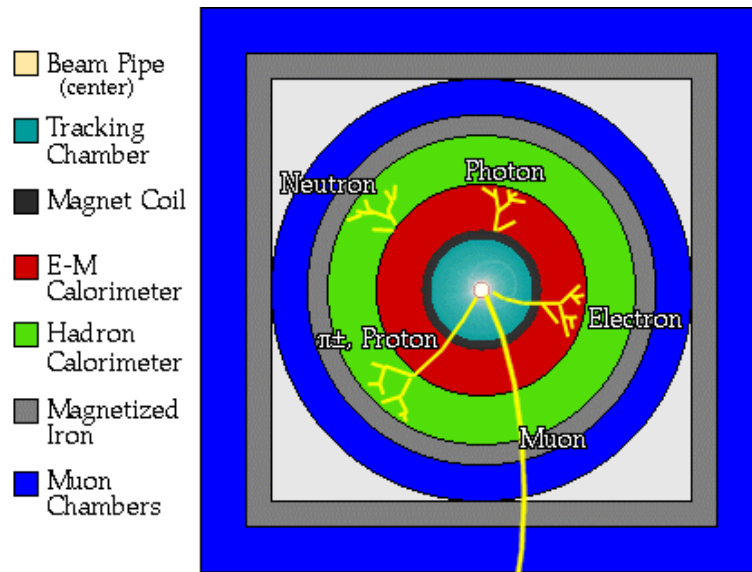


FIG. 2.9 – Schéma de la détection des particules dans une coupe transverse du détecteur.

Système de coordonnées utilisé dans l'expérience DØ

L'expérience DØ utilise un référentiel cylindrique pour décrire son détecteur. La direction du faisceau de protons définit le sens positif de l'axe z (du Nord au Sud) et l'axe y est la verticale ascendante. L'angle azimutal ϕ donne la position dans le plan transverse à l'axe z alors que l'angle polaire θ est utilisé pour le plan longitudinal. Le schéma de la Figure 2.10 récapitule ces conventions.

Les particules interagissant au moment de la collision, initiales et finales, sont relativistes (voir la Section 2.3.1). Elles subissent donc une poussée de Lorentz non négligeables. La distribution de ces particules n'est, par conséquent, pas uniforme par rapport à l'angle polaire θ . Par contre, la rapidité $y = \frac{1}{2} \ln \left(\frac{E+p_z c}{E-p_z c} \right)$ est un invariant de Lorentz. Dans le cas ultra-relativiste $\frac{mc^2}{E} \rightarrow 0$, la rapidité tend vers la pseudo-rapidity définie comme $\eta = -\ln(\tan(\frac{\theta}{2}))$. En utilisant cette nouvelle coordonnée, la distribution des particules est la même au repos ou dans un référentiel se déplaçant à une vitesse proche de celle de la lumière. Le système de coordonnées finalement utilisé est (r, ϕ, η) .

2.3.2 Le trajectographe interne

Le détecteur DØ est schématisé dans son ensemble sur la Figure 2.11 avec deux échelles graduées en mètres pour apprécier ses dimensions. La partie centrale représente le système de traces, appelé aussi le trajectographe.

Le but du trajectographe est de mesurer avec la meilleure résolution possible la trajectoire des traces des particules chargées. L'information sur la position et la courbure des traces fournissent alors la position du vertex (point de production), la quantité de mouvement et la charge de la particule détectée. Le détecteur DØ est composé de deux sous-détecteurs de traces chargées : un détecteur au silicium (le *SMT* : *Silicon Microstrip Tracker*) entouré d'un détecteur à base de fibres scintillantes (le *CFT* : *Central Fiber Tracker*). L'ensemble de ces deux sous-détecteurs est plongé dans un champ magnétique uniforme de 2 T produit par une bobine solénoïdale supraconduc-

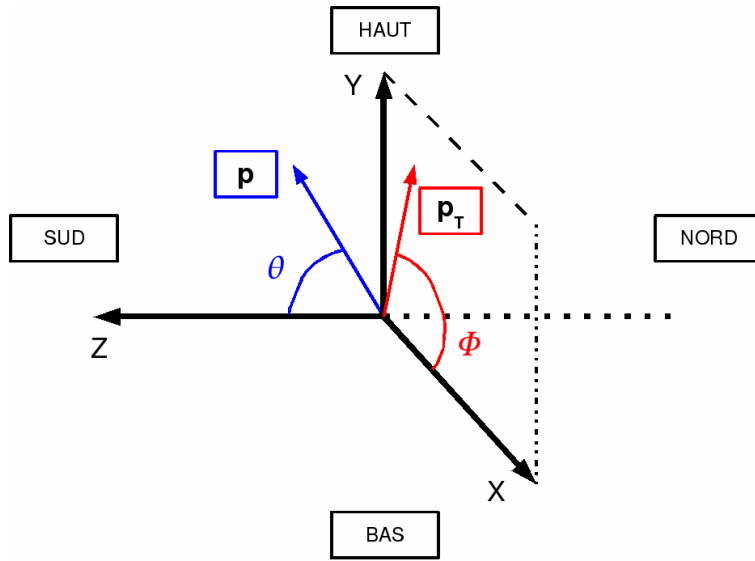


FIG. 2.10 – Système de coordonnées de l'expérience DØ

trice. Cette bobine a été ajoutée au *Run II*. Deux détecteurs de pieds de gerbe (voir Section 2.3.2) complètent cet ensemble plongé à l'intérieur de la cavité calorimétrique et permettent de compenser la perte en résolution du calorimètre due à l'ajout de la bobine. L'ajout de matière avant le calorimètre est en effet de 0.8 à 2 X_0 en fonction de η . La distance X_0 est appelée longueur de radiation et correspond à la distance moyenne entre deux interactions avec la matière. Elle dépend du matériau choisi.

Le détecteur au silicium

Le *SMT* [54] est un détecteur constitué de silicium. Le principe de détection [55] est l'utilisation de semi-conducteurs dopés n (excès d'électrons), par exemple avec du phosphore, et dopés p (carence en électrons), par exemple avec du bore, juxtaposés. La juxtaposition de ces deux matériaux crée une zone d'équilibre neutre, appelées zone de déplétion, et de part et d'autre de cette zone une accumulation de charges positives et négatives. On obtient ainsi une jonction *p-n* (pour plus de détails : [56]). La taille de cette zone peut être amplifiée par un champ électrique pour atteindre jusqu'à 300 μm . Une particule ionisante qui pénètre dans cette zone va libérer des paires électrons-trous entraînant la formation d'un signal électrique. Ce signal est finalement collecté par des bandes conductrices parallèles.

Deux types de détecteurs au silicium existent : les détecteurs à simple face de lecture, qui donnent accès à seulement deux des trois coordonnées et les détecteurs à double faces de lecture, qui permettent une mesure précise tridimensionnelle du passage des particules chargées (utilisés pour la première fois au LEP en 1992). Le *SMT* est composé de ces deux types de détecteur.

Le *SMT* a pour objectif de mesurer les traces proches de la collision et ainsi de déterminer avec précision la position du vertex dans la zone d'interaction. La mesure des traces doit avoir une couverture angulaire proche de celle du spectromètre à muons et du calorimètre. Ces contraintes fixent la géométrie du détecteur. La Figure 2.12 montre une vue tridimensionnelle de sa géométrie.

La zone d'interaction ($\sigma_Z \approx 25 \text{ cm}$) fixe la longueur du détecteur. La volonté de mesurer la

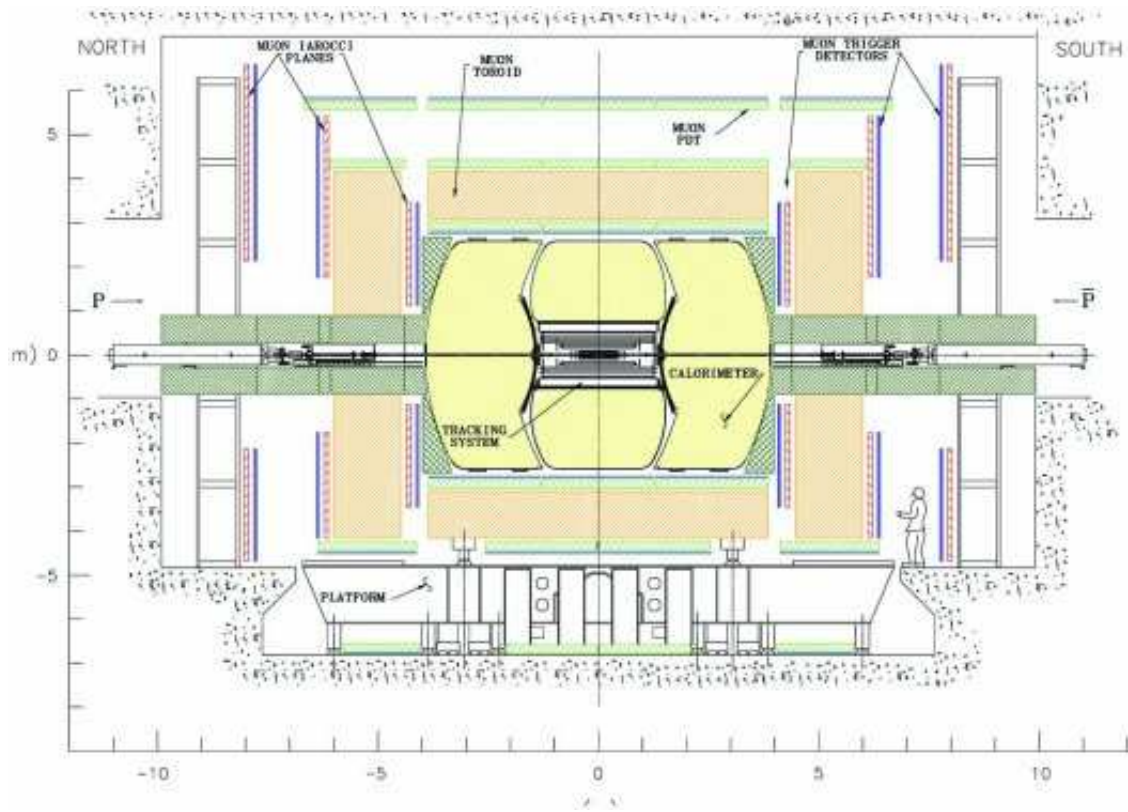


FIG. 2.11 – Schéma du détecteur DØ

trajectoire des traces sur une large couverture angulaire nécessite deux formes de détecteur au silicium : les tonneaux et les disques.

Les tonneaux, longs de 12 cm, sont au nombre de six dans la partie centrale. Ils déterminent la position longitudinale du vertex et ont une acceptance $|\eta| < 1,5$. Chacun d'eux se termine par un disque F. On trouve trois nouveaux disques F de chaque côté du dispositif. Cet ensemble constitue la partie centrale du *SMT*. Finalement, deux derniers disques de rayon plus important complètent l'édifice de chaque côté de la partie centrale pour atteindre une couverture angulaire de $|\eta| < 3$. Avec cette géométrie, le vertex est reconstruit en trois dimensions grâce aux disques pour les grandes valeurs de η et grâce aux tonneaux et au détecteur de traces à fibres pour les petites valeurs de η . Au final, la position du vertex primaire de l'interaction est déterminée avec une résolution de $35 \mu\text{m}$ le long de l'axe z .

Les tonneaux possèdent quatre couches successives de détecteurs au silicium, constituées elles-mêmes de sous-couches comme l'illustre la Figure 2.13. Les deux premières couches en possèdent 12 alors que les deux dernières en possèdent 24. Au total, les tonneaux comptent 432 détecteurs : à double face pour les quatre tonneaux du centre et pour les deux derniers tonneaux à simple face une couche sur deux. Ces couches sont distantes du tube, où se trouve le faisceau de protons, de 2,7 à 9,4 cm.

Les disques F, positionnés de $|z| = 12,5 \text{ cm}$ à $|z| = 48,1$, possèdent six "pétales" avec deux capteurs à double face dos à dos sur chacun des "pétales". Les disques H sont positionnés deux fois plus loin par rapport au centre du détecteur afin d'obtenir la couverture angulaire désirée :

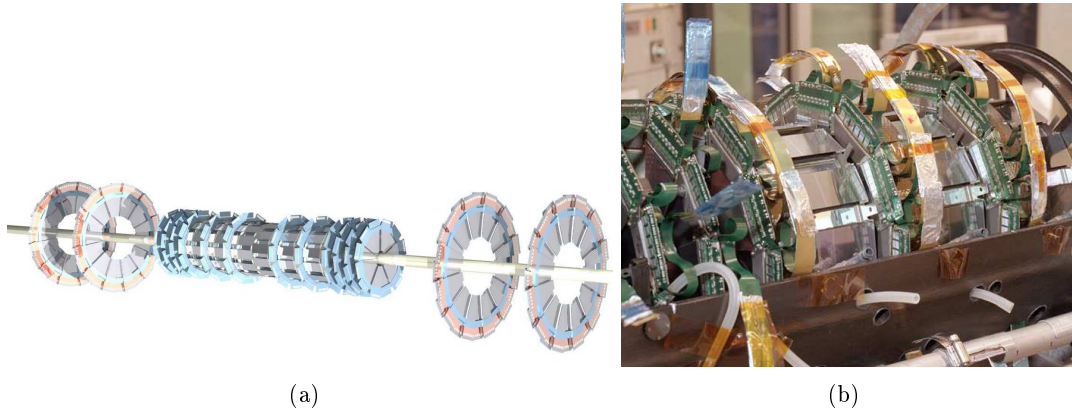


FIG. 2.12 – Schéma tridimensionnel du *SMT* (a) et photographie des tonneaux du *SMT* (b).

$|z| = 100,4$ cm et $|z| = 121,0$ cm. On trouve 24 "pétales" avec des détecteurs simple face sur ces disques disposés de la même façon que sur les disques F.

Au total, le *SMT* compte 912 détecteurs et 792576 canaux de lecture, suffisamment rapides pour être utilisés au niveau 2 du système de déclenchement (voir le chapitre 3 pour plus de détails). L'ensemble du dispositif, refroidi à une température de -5°C pour éviter la détérioration due aux radiations, permet finalement d'obtenir une résolution de l'ordre de $10\ \mu\text{m}$ dans le plan $r-\phi$ et de $40\ \mu\text{m}$ dans le plan longitudinal. Ces performances permettent un étiquetage performant des quarks b , comme il sera détaillé dans la Section 3.4.6.

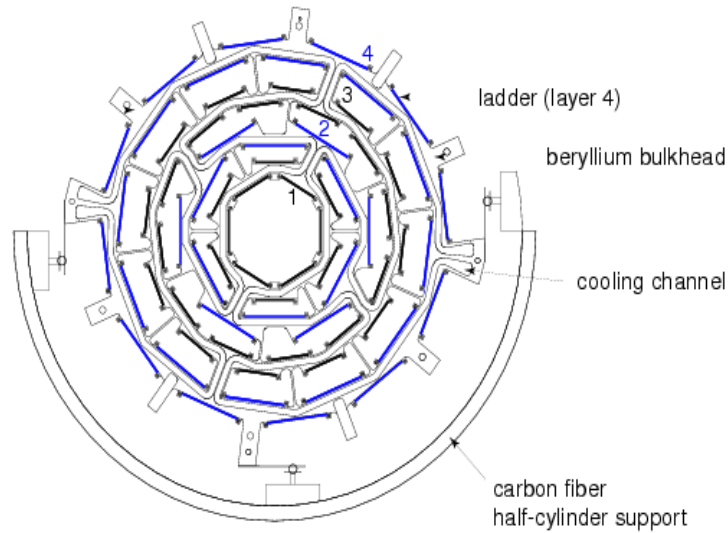
Le détecteur de traces à fibres

Le détecteur à fibres scintillantes ou *CFT* [57], amélioration apportée du *Run I* au *Run II*, entoure le *SMT*. Il est composé de huit cylindres situés de 20 à 52 cm du tube contenant le faisceau. Les deux premiers cylindres sont moins longs en raison du diamètre des disques H du *SMT*. Ces deux premiers cylindres ont une longueur de 1,66 m contre 2,52 m pour les six cylindres externes comme le montre la Figure 2.14. Cette géométrie permet d'obtenir une couverture angulaire jusqu'à $|\eta| \approx 1,7$.

Les fibres scintillantes sont utilisées pour la détection des photons émis quand une particule chargée traverse la fibre. Les fibres employées pour la reconstruction des traces sont des fibres à base de scintillateur organique en couche de faible épaisseur. On mesure ainsi le temps de passage de la particule tout en minimisant la perte d'énergie de la particule (fraction de X_0 très faible).

Sur chaque cylindre, on trouve deux doubles couches de fibres scintillantes. La première double couche est parallèle à l'axe z , chaque sous-couche étant séparée de l'autre d'un rayon de fibre ($\sim 417\ \mu\text{m}$). Les deux autres sous-couches ne sont pas parallèles à l'axe z afin de déterminer la trajectoire des traces en trois dimensions. La première d'entre elles fait un angle stéréo suivant le cylindre de $+3^\circ$ par rapport à l'axe z et la seconde un angle stéréo de -3° .

Le faible rayon des fibres et la précision sur leur positionnement ($< 50\ \mu\text{m}$) donnent une précision sur l'impact de la particule de $100\ \mu\text{m}$ dans le plan $r-\phi$. La résolution d'un tel dispositif est bien moins importante que celle d'un détecteur au silicium. Néanmoins ce dispositif permet à moindre coût de couvrir un volume beaucoup plus important. Le choix de l'expérience DØ a

FIG. 2.13 – Coupe du *SMT* et géométrie des tonneaux

donc été de combiner l'information du *SMT* et du *CFT* afin d'avoir une bonne précision sur la position des vertex (très utile pour l'identification des jets de quarks b , le *b-tagging*), tout en ayant une bonne résolution sur l'impulsion des particules chargées, c'est-à-dire sur la courbure de la trajectoire, et en déterminant le signe de la charge de la particule. Le trajectographe de l'expérience *CMS* (*The Compact Muon Solenoid*) sera en revanche entièrement en silicium.

Une des deux extrémités de la fibre est ensuite prolongée par des fibres optiques utilisées comme des guides d'onde de 7,8 à 11,9 m, qui transforment la lumière de scintillation (430 nm de longueur d'onde) en lumière visible (530 nm). Les guides d'onde transmettent ce flux de photons jusqu'aux *VLPC* (pour *Visible Light Photon Counters*), qui permettent de convertir ce flux de lumière en flux électrique. L'autre extrémité de la fibre est fermée par une couche d'aluminium qui joue le rôle de miroir avec une réflectivité de l'ordre de 90 %. Les *VLPC* sont suffisamment rapides et efficaces (efficacité quantique supérieure à 75 %) pour donner une information au niveau 1 du système de déclenchement. Pour maintenir leur gain élevé, ils doivent être néanmoins refroidis constamment dans un cryostat à l'hélium liquide à 9 K. Au total, le *CFT* utilise approximativement 200 km de fibres scintillantes et près de 800 km de guides d'onde.

Les détecteur de pieds de gerbe

Les détecteurs de pieds de gerbes, l'un central (*CPS* pour *Central PreShower*) et les deux autres à l'avant et à l'arrière du détecteur $D\phi$ (*FPS* pour *Forward PreShower*), ont été ajoutés du *Run I* au *Run II*. Leur objectif principal est l'identification des électrons et le rejet du bruit de fond tant au niveau du système de déclenchement qu'une fois la collision définitivement enregistrée.

Ils ont été conçus pour pallier la perte d'énergie entre le trajectographe et le calorimètre, notamment dans la bobine supraconductrice. Le nombre de longueurs de radiation de la bobine varie en effet de $1 X_0$ à $2 X_0$ en fonction de η . Dans la partie centrale du détecteur, une couche de plomb a été ajoutée pour uniformiser cette longueur de radiation et pour avoir approximativement $2 X_0$ quelque soit la valeur de η . Les détecteurs de pied de gerbe permettent d'introduire une

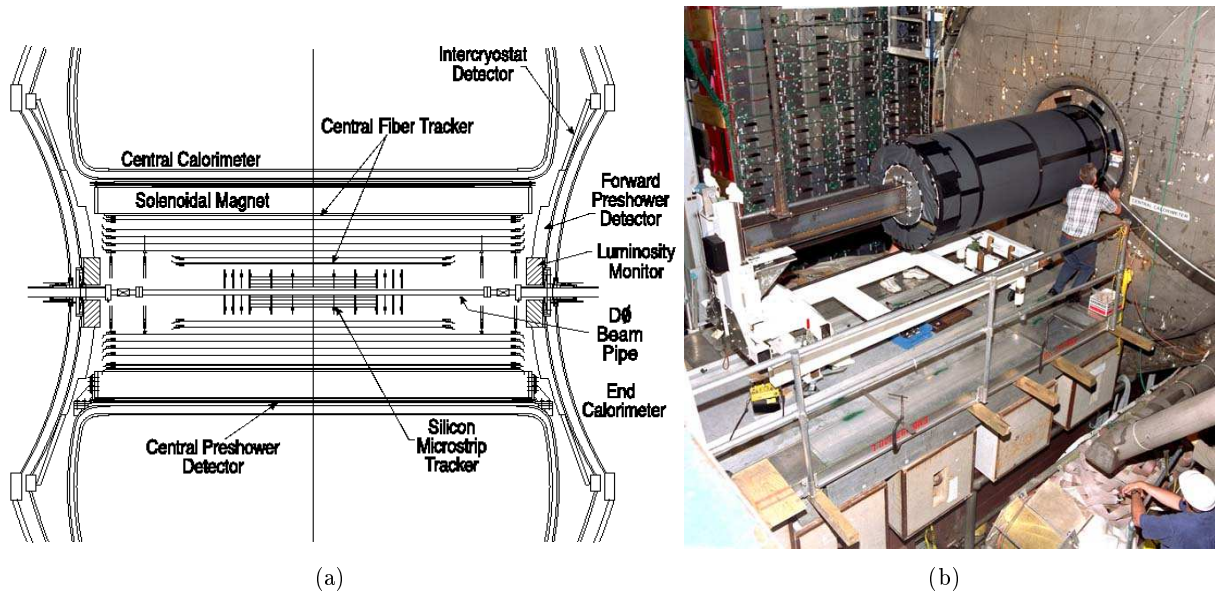


FIG. 2.14 – Schéma d'ensemble du détecteur de traces (a) et photographie de l'entrée du *CFT* dans l'aimant central (b).

correction à ce développement précoce de la gerbe électromagnétique. Ils sont également utiles pour mesurer plus précisément la position des particules et aident finalement à l'identification de ces dernières. Les détecteurs de pieds de gerbe jouent donc un double rôle :

- Détecteur de traces afin de faire coïncider les traces des électrons, mesurées par le trajectographe, et les amas électromagnétiques, mesurés par la calorimètre.
- Calorimètre en mesurant le début du développement des gerbes électromagnétiques.

Le *CPS* [58], formé d'une couche de plomb et du matériau actif, se trouve entre le solénoïde et le calorimètre. Il couvre la région $|\eta| < 1.3$, alors que les deux *FPS* ([59]) sont fixés sur la paroi interne des bouchons du calorimètre dans la région $1.5 < |\eta| < 2.5$. De la même façon que le *CFT*, les détecteurs de pieds de gerbes sont constitués de fibres scintillantes reliées à des *VLPC*.

Le *CPS* compte trois couches concentriques de fibres scintillantes de forme triangulaire afin de minimiser les zones mortes (voir la Figure 2.15). La première couche de fibres est parallèle à l'axe z , la seconde fait un angle stéréo de $+23,8^\circ$ par rapport à cet axe et la dernière de $-24,0^\circ$. Cette disposition des trois couches permet une nouvelle fois d'avoir une reconstruction tridimensionnelle. Chacune des couches est formée de 1280 fibres.

Les deux *FPS* (Nord et Sud) sont formés de deux couches sous-couches de fibres scintillantes séparées par une plaque de plomb ($2 X_0$). Un schéma d'un *FPS* est montré sur la Figure 2.15.

2.3.3 Le calorimètre

Introduction

Les différents sous-détecteurs du trajectographe ne donnent pas toute l'information nécessaire pour reconstruire un événement. Le calorimètre [60] apporte un supplément d'information conséquent. Il permet l'identification et la mesure des propriétés des particules neutres, en général

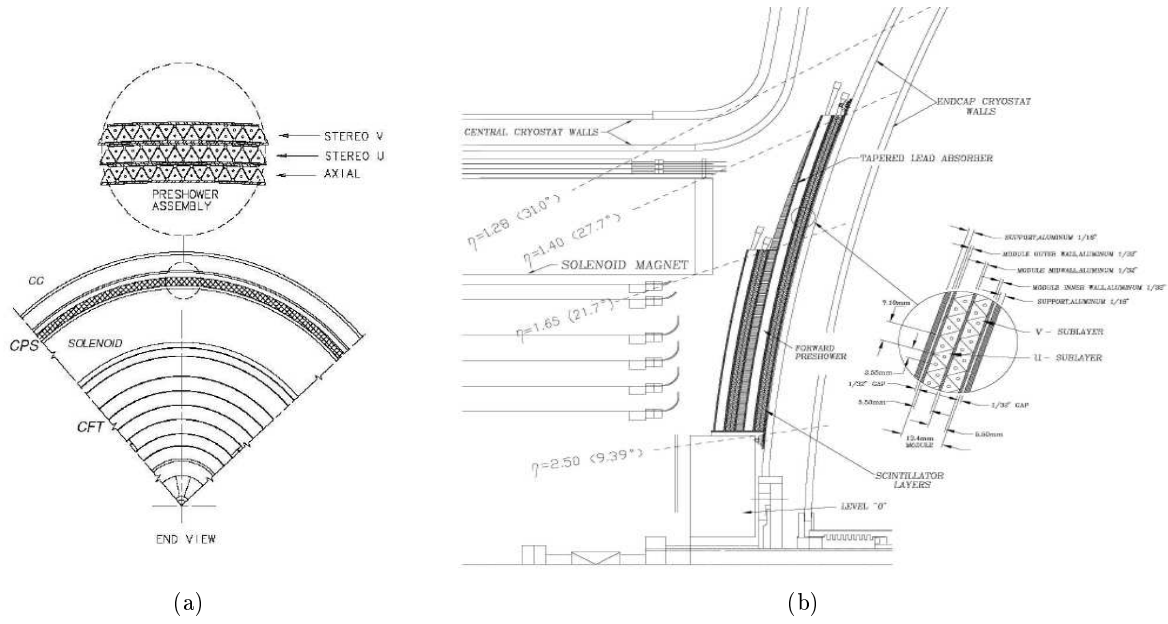


FIG. 2.15 – Schéma du CPS (a) et du FPS (b).

impossible avec le trajectographe. Il aide à l'identification et à la mesure de l'énergie des jets, des photons et des électrons. La mesure de l'énergie des jets est primordiale dans le cas d'un collisionneur hadronique pour lequel la multiplicité en jets peut être très importante. Le calorimètre peut également mesurer l'impulsion des particules à grandes quantités de mouvement, dont la trajectoire est faiblement courbée par le champ magnétique. Enfin, une bonne herméticité du calorimètre est indispensable pour la mesure de l'énergie manquante (en fait l'énergie transverse manquante) produite par les particules agissant faiblement avec le détecteur et par les imperfections de la détection.

La structure du calorimètre de DØ n'a pas été changée du *Run I* au *Run II*. Ce détecteur a été conçu pour mesurer l'énergie des photons, des électrons et des jets en l'absence de champ magnétique central. Le calorimètre de l'expérience DØ est un calorimètre à échantillonnage à l'argon liquide. Seule l'électronique de lecture a été modifiée afin de faire face à la diminution de l'intervalle de temps entre deux croisements de faisceaux. Le temps entre deux collisions passent en effet de $3.2 \mu\text{s}$ à 396 ns . De même, l'électronique de lecture du système de déclenchement a connu quelques changements du *Run IIa* au *Run IIb* pour pallier l'augmentation importante de la luminosité. Plus de détails sur ces derniers changements seront donnés dans la Section 3.3. Le principe de fonctionnement d'un calorimètre et plus précisément d'un calorimètre à échantillonnage sera décrit dans le paragraphe suivant. Suivra ensuite la description des différentes caractéristiques du détecteur.

Le principe de fonctionnement d'un calorimètre

Un calorimètre recueille et mesure l'énergie des particules qui interagissent avec les matériaux qui le composent. Plusieurs types d'interaction peuvent avoir lieu selon le type de particules, dont on cherche à mesurer l'énergie.

Pour les photons et les électrons, une gerbe électromagnétique se forme dans le calorimètre par cascades successives d'interactions avec la matière. Les photons de hautes énergies ($E_\gamma > 1 \text{ GeV}$) se matérialisent en une paire $e^+ - e^-$, alors que les électrons ($E_e > 100 \text{ MeV}$) vont interagir par émission de rayonnement de freinage ou *Bremsstrahlung* : $e^\pm \rightarrow e^\pm + \gamma$. Le nombre de particules se multiplie au fur et à mesure que la gerbe se développe. Dans le même temps, l'énergie des particules diminue, c'est ce qui arrête finalement le développement de la gerbe. La longueur de la gerbe dépend donc de l'énergie de la particule initiale (on peut montrer qu'elle est proportionnelle à cette énergie) et de la distance entre deux interactions avec le matériau, c'est-à-dire la longueur de radiation X_0 . Dans le cas du calorimètre de DØ, on utilise notamment l'uranium ($X_0 = 3.2 \text{ mm}$). Cette distance fixe finalement les dimensions du calorimètre électromagnétique.

Pour les particules hadroniques, le principe est le même, c'est-à-dire formation d'une gerbe, mais les interactions avec l'absorbeur sont plus complexes et variées. On peut trouver deux contributions distinctes dans une gerbe hadronique :

- Une composante électromagnétique avec comme interaction initiale : $\pi^0 \rightarrow \gamma\gamma$
- Une composante hadronique avec les interactions fortes des pions chargés, des kaons, des protons, des neutrons,...

Un schéma d'une gerbe hadronique est représentée sur la Figure 2.16. Dans le cas d'une telle gerbe, ce n'est plus la longueur de radiation qui caractérise son développement mais la longueur d'absorption nucléaire, notée λ_A . Cette longueur est plus longue que la longueur de radiation, pour l'uranium $\lambda_A = 10 \text{ cm}$. Une gerbe hadronique est en effet plus longue à initier, de plus les processus mis en jeu produisent moins de particules qui interagissent avec l'absorbeur. Ainsi, les dimensions longitudinales du calorimètre hadronique sont beaucoup plus grandes que celles du calorimètre électromagnétique (une particule de 100 GeV est arrêtée dans le calorimètre après une longueur de $9 \lambda_A$).

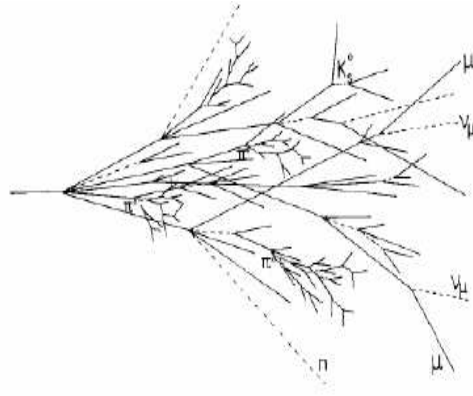


FIG. 2.16 – Schéma d'une gerbe hadronique

L'absorbeur, nécessairement dense pour limiter la taille du détecteur, permet donc de former des gerbes de particules et de recueillir une fraction de l'énergie de la particule détectée. Il faut ensuite un milieu sensible pour mesurer l'énergie des particules chargées créées. Deux types de calorimètres sont alors possibles :

- Les calorimètres à milieu homogène (scintillateurs par exemple). Dans ce cas, l'absorbeur est également le détecteur. Ces calorimètres donnent une bonne résolution en énergie mais une résolution spatiale limitée.

- Les calorimètres à échantillonnage : L'absorbeur et le détecteur sont séparés. Seule une fraction de l'énergie est collectée mais la résolution spatiale est améliorée. Un autre avantage d'un tel dispositif est le coût et la compacité du détecteur.

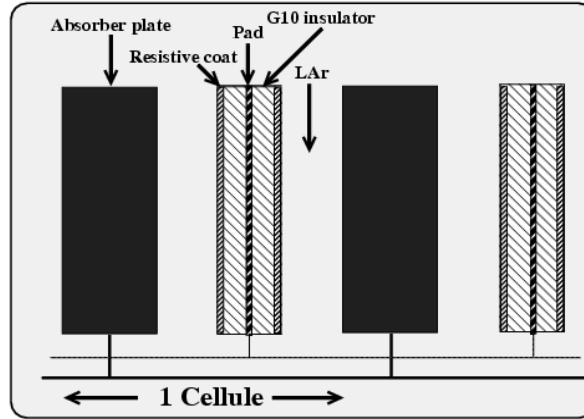


FIG. 2.17 – Schéma de principe d'une cellule du calorimètre

Les calorimètres électromagnétiques et hadroniques de DØ sont des calorimètres à échantillonnage comme le montre le schéma d'une cellule calorimétrique sur la Figure 2.17. Chaque cellule est formée de l'alternance d'un absorbeur et d'une électrode plongée au centre d'un espace d'argon liquide, le milieu actif. L'absorbeur, qui initie et entretient la gerbe, est reliée à la masse alors que l'électrode est soumise à un potentiel positif. Le milieu actif est un milieu ionisant qui recueille une fraction de l'énergie de la particule incidente. De cette façon, le champ électrique créé fait dériver les électrons de l'argon liquide vers l'électrode en un temps proche des 450 ns. La mesure du courant électrique donne l'information sur l'énergie déposée dans la cellule (seule l'énergie déposée dans l'argon liquide est réellement mesurée). Le choix de l'alternance d'un absorbeur et de l'argon liquide a plusieurs avantages : la simplicité de calibration, la flexibilité de segmentation du calorimètre et une réponse stable dans le temps. Par contre, il nécessite un système de refroidissement à 78 K qui introduit des zones non instrumentées dans le détecteur.

Les différents absorbeurs utilisés et leur caractéristique sont décrits dans la Section 2.3.3.

Caractéristique du calorimètre de DØ

Le calorimètre de DØ est découpé en trois parties complètement séparées les unes des autres pour permettre l'accessibilité aux différentes parties du détecteur : une partie centrale qui couvre la zone $|\eta| < 1$ et deux bouchons à l'avant et à l'arrière du détecteur qui étendent la couverture angulaire jusqu'à $|\eta| \approx 4$. La séparation entre chacune des parties est perpendiculaire à l'axe du faisceau (voir sur la Figure 2.18) afin de minimiser la dégradation sur la résolution de l'énergie transverse manquante.

Les dimensions du calorimètre sont fixées par des critères pratiques, comme le volume du hall de l'expérience DØ, et par des critères physiques, comme la profondeur des gerbes de particules et la nécessité de la mesure de l'impulsion du muon en dehors du calorimètre (voir la Section 2.3.4). Pour satisfaire chacune de ces contraintes, plusieurs types distincts de modules sont utilisés. Les modules électromagnétiques constitués de plaques fines d'uranium pour l'absorbeur permettent

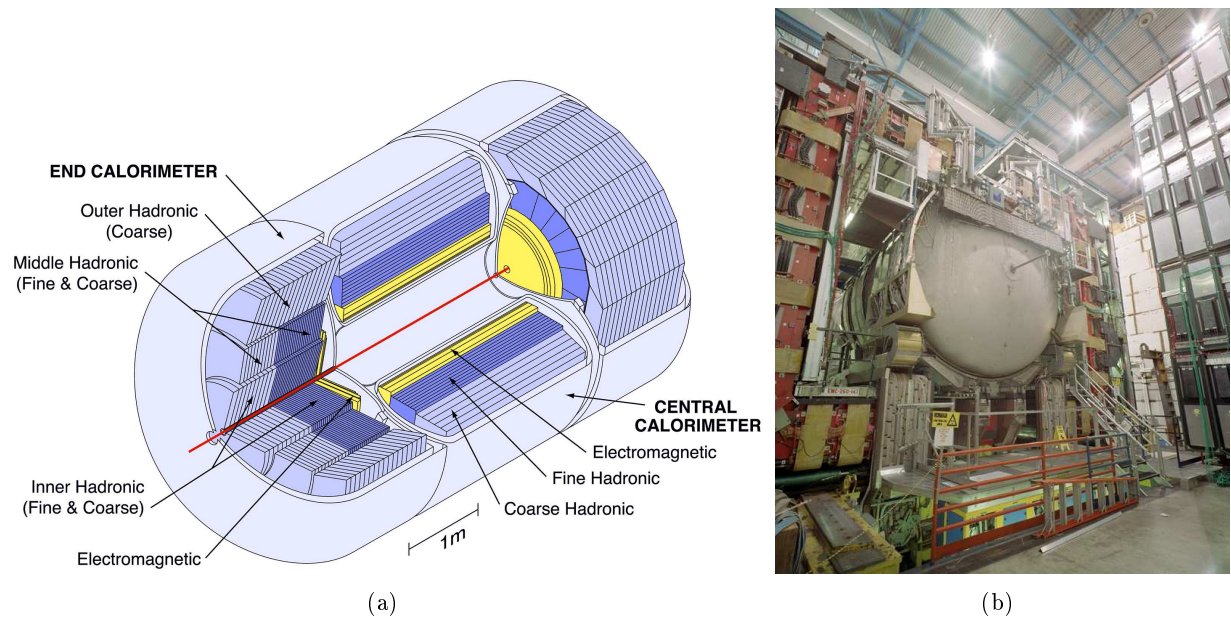


FIG. 2.18 – Vue isométrique du calorimètre de DØ (a) et photographie de l'extérieur d'un bouchon (b).

le déploiement longitudinal de la gerbe électromagnétique. Les modules hadroniques fins (*Fine Hadronic*) avec des plaques d'uranium plus épaisses permettent généralement de contenir toute la gerbe hadronique. Ils sont donc utiles à la mesure de l'énergie et de la position des hadrons incidents. Enfin, les modules hadroniques de faible granularité (*Coarse Hadronic*) sont constitués de plaques de cuivre ou d'acier inoxydables. Cette partie hadronique, dite grossière, protège en fait les chambres à muons contre le développement de gerbes tardives. Le détail de ces modules est donné dans les paragraphes suivants et résumé dans le Tableau 2.4.

Le calorimètre central Le calorimètre central est segmenté selon l'angle ϕ . Il est formé de 32 modules pour sa partie électromagnétique (*CCEM*), 16 pour la partie hadronique fine (*CCFH*) et 16 pour la partie hadronique grossière (*CCCH*). Chaque module est de plus segmenté en cellules : l'unité de mesure de base du calorimètre. La taille des cellules contraint la résolution du détecteur. Ces cellules, de forme trapézoïdales, sont organisées en tours pseudo-projectives pointant vers le centre du détecteur (voir la Figure 2.19). Une tour est un empilement de cellules.

La longueur des cellules est imposée par la taille des gerbes. Le calorimètre électromagnétique compte quatre couches cylindriques concentriques de distance 85, 87, 92 et 99 cm par rapport à l'axe des faisceaux (*EM1*, *EM2*, *EM3* et *EM4*). Chacune de ces couches correspond respectivement à une longueur de radiation de 2, 2, 6.8 et 9.8 X_0 , ce qui correspond à une longueur totale de 20.6 X_0 et 0.76 λ_A . Il faut ajouter, à toutes ces longueurs de radiation, 2 X_0 correspondant au trajectographe, au solénoïde et à la couche de plomb en amont (voir la Section 2.3.2). Les deux premières couches, de faibles longueurs de radiation, sont utiles à la différenciation des photons et des π^0 , qui n'est possible qu'au début du développement de la gerbe (cette différenciation peut être améliorée en utilisant conjointement les détecteurs de pieds de gerbe). Les cellules de ces

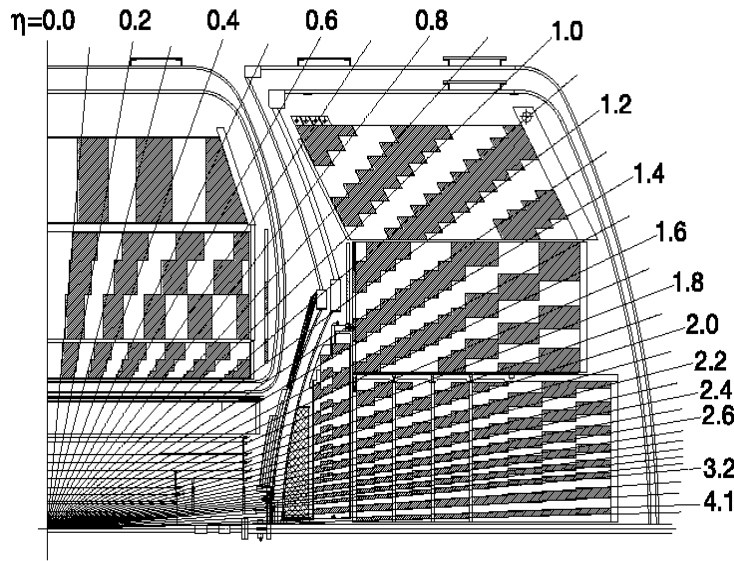


FIG. 2.19 – Schéma d'une portion du calorimètre et agencement projectif des modules.

couches sont formées d'une fine plaque d'uranium (3 mm).

CCFH est constitué de trois couches de longueur d'absorption nucléaire 1.3, 1 et 0.9 λ_A . L'absorbeur pour les cellules de ces couches est une plaque d'uranium plus épaisse (6 mm) que celles utilisées pour les cellules du *CCEM*. Finalement, le *CCCH* n'est constitué que d'une couche de 3.2 λ_A avec pour absorbeur une plaque de cuivre. Ces dernières couches sont situées à des distances de 119, 141, 158 et 195 cm de l'axe z .

La section des cellules, caractérisée par les variables physiques ϕ et η décrites dans la Section 2.3.1, est imposée par la taille caractéristique d'une gerbe hadronique : $\Delta R = \sqrt{\Delta\phi^2 + \Delta\eta^2} \sim 0.5$. Toutes les cellules ont une section de $\Delta\phi \times \Delta\eta = 0.1 \times \frac{2\pi}{64}$, exceptée pour la troisième couche du calorimètre électromagnétique (*EM3*) où les électrons déposent le plus d'énergie et où la gerbe est la plus large. Celle-ci est deux fois plus segmentée : $\Delta\phi \times \Delta\eta \approx 0.05 \times 0.05$.

Les bouchons Le calorimètre compte deux bouchons : un au Nord du détecteur et l'autre au Sud. Les bouchons sont formés de trois cylindres concentriques pour la partie hadronique (interne, du milieu et externe), la partie électromagnétique (*ECEM*) s'appuyant sur les deux cylindres les plus proches de l'axe z (voir le schéma de la Figure 2.18).

Le calorimètre électromagnétique des bouchons possède la même géométrie pour ses cellules que le calorimètre électromagnétique central. On compte quatre couches de cellules avec un absorbeur en uranium appauvri. Les longueurs de radiation de chacune des couches sont 0.3, 2.6, 7.9 et 9.3 X_0 . Il faut, en fait, ajouter deux longueurs supplémentaires à la première couche en raison de la présence du cryostat, de la plaque de plomb et des détecteurs de pied de gerbe. La couche *EM1* est située à 1.7 m du centre de détecteur et est longue de 2 cm. Les couches *EM2*, *EM3* et *EM4* font respectivement 2 cm, 8 cm et 9 cm d'épaisseur. De plus, une plaque d'acier inoxydable est placée entre les couches *EM3* et *EM4* par souci de construction. La segmentation de ces cellules est identique à la segmentation du calorimètre électromagnétique central jusqu'à $|\eta| = 2.6$. Au-delà, la segmentation de *EM3* redevient $\Delta\phi \times \Delta\eta \approx 0.1 \times 0.1$, car à grand η les

cellules sont de plus en plus petites (voir le schéma de la Figure 2.19).

Le calorimètre hadronique interne est constitué d'une partie hadronique fine de quatre couches ($1.2 \lambda_A$ chacune) et d'une couche grossière de $3.6 \lambda_A$. Les couches hadroniques fines ont pour absorbeur une plaque d'uranium alors que les couches grossières sont formées d'une plaque d'acier inoxydable. La seconde partie du calorimètre hadronique (du milieu) possède les mêmes caractéristiques que la partie interne avec des longueurs d'absorption nucléaire de $1 \lambda_A$ pour la partie fine et de $4.1 \lambda_A$ pour la partie de faible granularité. La dernière partie (externe) n'est constitué que d'une couche grossière avec des plaques d'acier inoxydable pour absorbeur ($7 \lambda_A$). La segmentation de toutes les cellules est $\Delta\phi \times \Delta\eta \approx 0.1 \times 0.1$. Elle devient quatre fois plus importante à partir de $|\eta| = 3.2$.

Les détecteurs inter-cryostat La zone entre les bouchons et le calorimètre central est peu instrumentée car elle contient notamment les cryostats des différents modules du calorimètre. Cette région en pseudo-rapacité $0.8 < |\eta| < 1.4$ nécessite une infrastructure supplémentaire pour mesurer approximativement l'énergie des particules qui la traversent et améliorer la précision sur l'énergie transverse manquante. On appelle cette région la région inter-cryostat (*ICR* pour *InterCryostat Region*).

L'ajout d'une bobine supraconductrice centrale du *Run I* au *Run II* a pour conséquence de nombreuses modifications pour les détecteurs de cette région (voir [61] pour plus de détails). Deux types de détecteurs spécifiques permettent de compléter les informations données par le calorimètre : le détecteur inter-cryostat (*ICD* pour *InterCryostat Detector* et les calorimètres sans absorbeur (*massless gaps*). L'*ICD* est un détecteur formé de scintillateurs qui prolongent les tours du calorimètre. Le signal est collecté par des fibres reliées à des guides d'onde, qui transmettent ce signal lumineux à des photo-multiplicateurs situés en dehors du champ magnétique. Le détecteur inter-cryostat est fixé sur les parois des cryostats des bouchons ($1.1 < |\eta| < 1.4$). Les *massless gaps* sont des cellules du même type que celles du calorimètre avec l'argon liquide comme milieu actif. Le rôle de l'absorbeur est joué par les parois du calorimètre et/ou les parois des cryostats. Ils sont situés entre les parois du calorimètre et les cryostats. On les trouve dans la partie centrale et dans les bouchons.

L'électronique de lecture Le schéma électrique de l'électronique de lecture du calorimètre est détaillée sur la Figure 2.20. On y trouve trois parties importantes [62] :

- Acquisition et amplification du signal électrique en provenance de la cellule calorimétrique.
- Transport du signal pour un échantillonnage, un tri et une soustraction du bruit.
- Conversion du signal analogique en un signal numérique.

Les deux premières parties sont complètement nouvelles pour le *Run II*. Elles étaient indispensables pour gérer la diminution du temps de croisement entre deux paquets de protons et d'antiprotons.

La première étape de l'acquisition pour une cellule donnée consiste à amplifier le signal électrique mesuré par cette cellule. La principale difficulté de cette étape est l'obtention de signaux électriques identiques en forme et en amplitude pour toutes les canaux d'acquisition. Les câbles qui transportent le signal de la cellule aux pré-amplificateurs ont été remplacés pour le *Run II* afin d'avoir une impédance et une longueur proche pour chacun des câbles. Contrairement aux pré-amplificateurs, les cartes *BLS* (pour *Base Line Subtractor*) sont faciles d'accès. Elles sont situées sous le cryostat. Ces cartes remplissent de nombreuses fonctions. Elles échantillonnent d'abord le signal toutes les 132 ns. Cette étape est effectuée par un "formeur" (pour *shaper* en

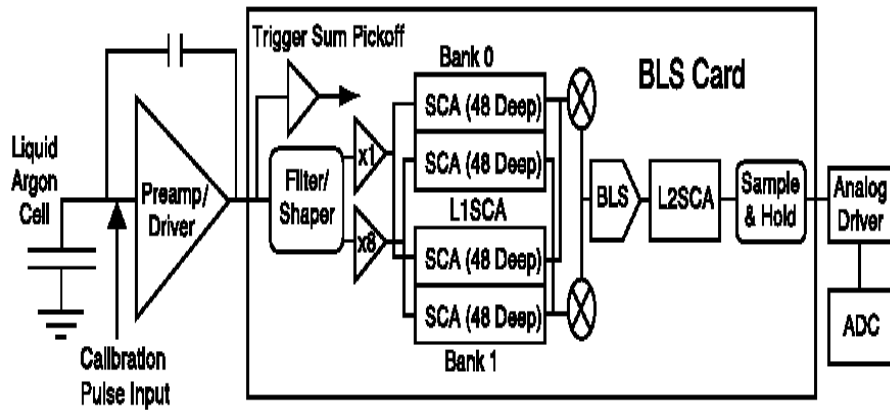


FIG. 2.20 – Chaîne de l'électronique de lecture du calorimètre.

anglais) bi-gain ($\times 1$ et $\times 8$). On rappelle que le signal provenant de la cellule a un temps d'acquisition d'environ 450 ns, de plus la période de relaxation est d'environ 15 μ s. Les cartes *BLS* sont capables ensuite de stocker le signal pendant 4 μ s en attendant la décision du niveau 1 du système de déclenchement. Cette mémoire analogique s'appelle un *SCA* (pour *Switch Capacitor Array*). Si l'événement passe ce premier niveau du système de déclenchement, la carte soustrait le bruit résiduel dû aux collisions précédentes (une collision toutes les 396 ns) et au bruit de basses fréquences. Le résultat est de nouveau stocké dans un *SCA* pour environ 2 ms en attendant la décision du niveau 2 du système de déclenchement. Le signal est finalement envoyé aux cartes *ADC* (*Analog-to-Digital Convertor*) qui le convertissent en signal numérique, qui peut être utilisé pour la décision finale du niveau 3. Ce signal est multiplié par 8 si le gain du "formeur" était $\times 1$. Le "formeur" bi-gain permet ainsi d'étendre la gamme dynamique des valeurs enregistrées par les cellules du calorimètre.

L'électronique du système de déclenchement sera détaillée plus en détails dans le chapitre 3. Il existe en effet deux voies de lecture différentes et indépendantes pour la mesure de l'énergie déposée dans le calorimètre (voir la Figure 2.20 : bifurcation juste avant le "formeur"). La première est celle décrite plus tôt avec deux gains possibles, la seconde est l'électronique de lecture du système déclenchement qui utilise seulement le gain $\times 1$. Par conséquent, cette dernière n'est pas aussi précise que la voie de lecture utilisée hors-ligne (appelée également *precision readout* en anglais). On distinguera alors dans la suite l'électronique de déclenchement de l'électronique de précision.

Les performances générales du calorimètre

On caractérise les performances d'un calorimètre en estimant l'incertitude sur la mesure de l'énergie. Cette incertitude peut être due à deux effets évoqués dans les paragraphes précédents : la résolution intrinsèque du calorimètre et la non-linéarité de la réponse électronique. La résolution en énergie d'un calorimètre est généralement paramétrée par la formule 2.10.

$$\frac{\sigma(E)}{E} = \sqrt{\left(\frac{A}{E}\right)^2 + \left(\frac{B}{\sqrt{E}}\right)^2 + C^2} \quad (2.10)$$

A est le terme de bruit lié à l'électronique et à l'activité de l'uranium. Ce terme est indépendant de l'énergie de la particule détectée. B est le terme stochastique dû à la non-uniformité de l'échantillonnage. Il dépend du nombre de particules chargées susceptibles de produire un signal. On peut réduire ce terme en augmentant la fraction d'échantillonnage. Et C est le terme constant relié à la non-uniformité de l'électronique et de la mécanique.

Au *Run I*, des tests ont été effectués avec des faisceaux de pions et d'électrons [60]. Ces mesures ont donné les termes suivants :

- Pour les électrons : $A = 0.29 \pm 0.03 \text{ GeV}$, $B = 0.157 \pm 0.006 \text{ GeV}^{\frac{1}{2}}$, $C = 0.003 \pm 0.003$.
- Pour les pions : $A \approx 0.29 \text{ GeV}$, $B = 0.41 \pm 0.04 \text{ GeV}^{\frac{1}{2}}$, $C = 0.032 \pm 0.004$.

Pour le *Run II*, le temps n'a pas permis d'effectuer de tels tests. Les mesures des paramètres A, B et C ont donc été effectuées en conditions normales de collision. Les résultats sont les suivants :

- Pour les électrons (la méthode est décrite dans [63]), on montre qu'il y a une forte dépendance de la résolution en fonction de l'angle d'incidence. Il n'est donc plus possible d'utiliser simplement la fonction de la Formule 2.10. Il est nécessaire de paramétrer la dépendance en $|\eta|$ [64].
- Pour l'énergie transverse des jets, les résultats sont dépendants de $|\eta|$ et sont détaillés dans la DØ note donnée dans la référence [65]. Un exemple est donné sur la Figure 2.21.

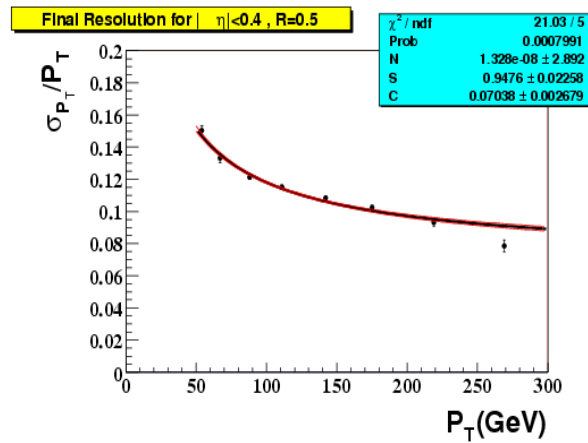


FIG. 2.21 – Résolution sur l'énergie transverse d'un jet de cône $R = 0.5$ pour $|\eta| < 0.4$.

2.3.4 Le spectromètre à muons

Introduction

La couche la plus externe du détecteur DØ est le spectromètre à muons [66]. Les muons sont en effet les seules particules leptoniques qui traversent le calorimètre en y déposant très peu d'énergie grâce à leur masse élevée (0.11 GeV) et leur temps de vie suffisamment long (voir la Figure 2.9).

Deux types de détecteurs sont utilisés pour le spectromètre à muons. Les chambres proportionnelles à dérive, dont le fonctionnement est décrit dans la la suite, permettent de mesurer la position et l'énergie des muons ionisant le milieu. Ces détecteurs n'étant pas assez rapides pour le système de déclenchement, des scintillateurs complètent le dispositif.

L'ensemble du spectromètre à muons comporte trois parties. Un aimant toroïdal permet de courber la trajectoire des muons et donne une information sur leur impulsion et leur charge. Une partie centrale couvre la région $|\eta| < 1$ et une partie à l'avant et à l'arrière, complètement modifiée du *Run I* au *Run II*, étend la couverture angulaire ($1 < |\eta| < 2$). Les deux systèmes de détection des muons utilisent le système de coordonnées (x,y,z). Les axes x et y appartiennent à la partie centrale et l'axe z aux parties avant et arrière. Le schéma des chambres à dérive est représenté sur la Figure 2.22, alors que celui des scintillateurs est représenté sur la Figure 2.23.

Principe de fonctionnement des chambres proportionnelles à dérive

Les chambres proportionnelles à dérive sont des détecteurs à ionisation. Ils mesurent la charge déposée par une particule dans un milieu ionisable. Une particule très énergétique peut arracher les électrons du milieu qu'elle traverse, c'est le processus d'ionisation. Dans de tel détecteur, le milieu utilisé est un gaz, facilement ionisable, soumis à un champ électrique très fort (7.4 kV pour la partie centrale et 3.1 kV pour les parties avant et arrière) généré par une paire d'électrodes. Chaque tube des chambres proportionnelles à dérive est traversé par un fil, qui joue le rôle d'anode, et les cathodes sont situées sur les parois des tubes. Le champ électrique est suffisamment fort pour que les électrons, isolés par le processus d'ionisation, arrachent à leur tour des électrons des atomes du milieu. C'est le phénomène d'avalanche. Les électrons suivent les lignes de champ magnétique et sont rapidement capturés par l'anode pendant que les ions dérivent vers la cathode. Ce processus induit un signal électrique. Trois informations sont enregistrées par les chambres à dérives à chaque passage d'un muon :

- Le temps de dérive des électrons,
- La différence de temps d'arrivée des charges sur le fil et sur le fil voisin,
- La quantité de charge recueillie.

Les chambres proportionnelles à dérive sont de bons détecteurs de position, mais la résolution sur l'énergie en utilisant seulement l'information de ces détecteurs est mauvaise. Pour une meilleure détermination de l'énergie et de la position des muons, on utilise la complémentarité entre le spectromètre à muons et les informations du trajectographe et du calorimètre.

La vitesse de dérive dans le gaz des chambres proportionnelles à dérive de l'expérience DØ est d'environ 10 cm/ μ s (pour la partie centrale), ce qui implique un temps de dérive maximum de 500 ns (étant donnée la taille des tubes 10 cm $\times$ 5 cm). Ce temps étant supérieur au temps de croisement entre les faisceaux, ces détecteurs ne peuvent pas être utilisés pour le système de déclenchement. Pour cette raison, le spectromètre à muons utilise conjointement des scintillateurs rapides (voir les Sections 2.3.4 et 2.3.4).

L'aimant toroïdal

Les aimants toroïdaux sont situés à l'extérieur du calorimètre comme on peut le voir sur la Figure 2.11. Le toroïde central est un aimant en forme de carré situé à 318 cm de l'axe du faisceau et faisant une épaisseur de 109 cm. Les toroïdes avant et arrière sont situés dans l'intervalle $454 \leq |z| \leq 610$ cm. Ces aimants sont constitués d'enroulement de conducteurs et créent un champ magnétique de 1.8 T.

Les toroïdes permettent une mesure de l'impulsion des muons indépendantes des autres sous-détecteurs. L'avantage est de pouvoir utiliser cette information, même imprécise, dès le niveau

1 du système de déclenchement. Cette mesure peut, de plus, être complétée par la mesure du trajectographe pour une plus grande précision.

La partie centrale : WAMUS

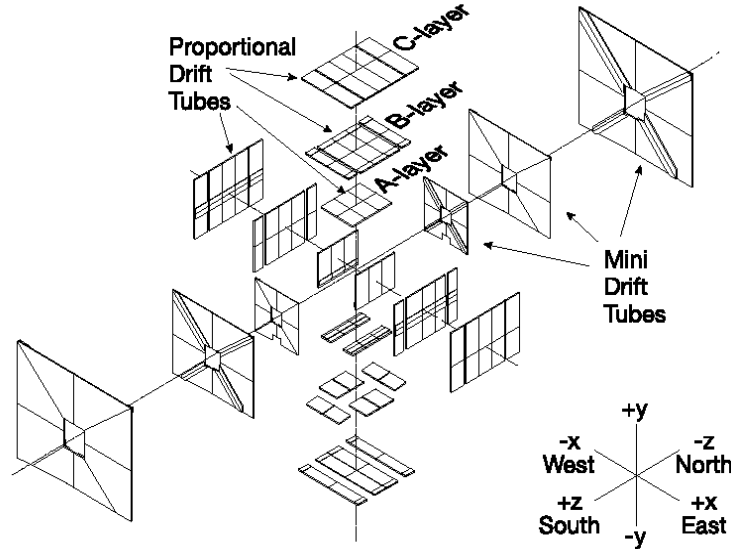


FIG. 2.22 – Schéma du système de chambres à fils du spectromètre à muons.

La partie centrale du spectromètre à muons est également appelée *WAMUS* pour *Wide Angle MUon Spectrometer*. Elle est composée de trois couches de chambres à dérive constituées de *PDT* (*Proportionnal Drift Tube*). La première couche, dénommée A, est située avant le toroïde central et comporte quatre plans de *PDT*. Les couches, dites B et C, ne comportent que trois plans de *PDT* et se trouvent à l'extérieur du toroïde central (voir la Figure 2.22).

Les *PDT* sont remplis d'un mélange gazeux, jouant le rôle de milieu ionisable : 84 % d'argon, 8% de CH_4 et 8% de CF_4 . Ils sont de forme rectangulaire de section $10 \text{ cm} \times 5 \text{ cm}$, de longueur 2 m, de largeur 1 m et de profondeur 20 cm.

Le dispositif central est complété par des scintillateurs pour le système de déclenchement. Une couche de scintillateurs, appelée $A\Phi$, est située entre le calorimètre et les *PDT* de la couche A (voir la Figure 2.23). Ces scintillateurs ont une segmentation en ϕ de 4.5 degrés. Ils permettent à la fois d'associer les muons détectés par les chambres à dérive avec un croisement de faisceau et de rejeter les cosmiques (muons produits par des rayons cosmiques. Ils sont nombreux en raison du faible enterrement du détecteur DØ). Pour le rejet des cosmiques, la principale méthode utilisée est la mesure de la coïncidence en temps entre un impact détecté et un croisement de faisceau. La segmentation des couches $A\Phi$ correspond à la géométrie du *CFT*, ce qui permet l'association des traces reconstruites avec un impact dans le spectromètre à muons. Des couches de scintillateurs sont également installées sur la couche B. On trouve approximativement 396 scintillateurs entre la couche B et C. Ces couches de scintillateurs associées à l'information du *CFT* permettent de construire des conditions de déclenchement sur les muons. Finalement, deux couches de scintillateurs entourent la couches C des chambres à dérive. Ils sont essentiellement utilisés pour la détection des cosmiques, de la même façon que les scintillateurs des couches $A\Phi$.

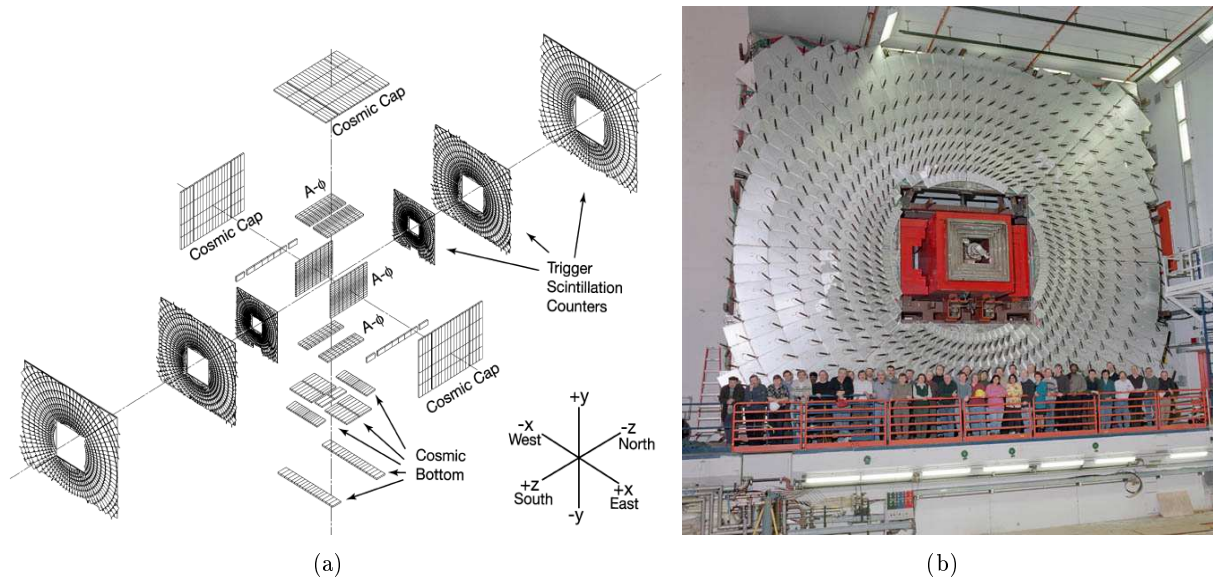


FIG. 2.23 – Vue explosée du système de scintillateurs du spectromètre à muons (a) et photographie de l'extérieur des scintillateurs (b).

Finalement, un système de blindage en fer et en polyéthylène, installé sur les parois du calorimètre, protège les premières couches du spectromètre à muons contre les gerbes tardives ou les débris de protons et d'antiprotons émis à l'avant. Une couche de plomb entoure ce système de blindage pour absorber les photons éventuellement émis par cette protection.

Le système avant et arrière : *FAMUS*

Le *FAMUS* pour *Forward Angle MUon Spectrometer* couvre la partie Nord et Sud du détecteur. Les chambres à dérive sont beaucoup plus près du faisceau que le spectromètre de la partie centrale et donc beaucoup plus sensibles aux radiations. Pour cette raison, un blindage supplémentaire a été ajouté du *Run I* au *Run II* le long de l'axe z pour réduire un déclenchement inopiné de la couche la plus proche du faisceau. De plus, les *PDT* ont été remplacés par les *MDT* : *Mini Drift Tube*. Le mélange gazeux est différent et est plus stable vis-à-vis des radiations : 90% de CF_4 et 10% de CH_4 . La section de ces tubes perpendiculaires au faisceau est de $1\text{ cm} \times 1\text{ cm}$, et le temps de dérive est réduit à 60 ns pour les électrons d'ionisation. De la même façon que pour la partie centrale, des scintillateurs sont ajoutés pour un veto des muons cosmiques et un déclenchement sur les muons grâce à la complémentarité avec le *CFT*. Trois couches de scintillateurs sont installées sur les chambres à dérives : sur la face interne des couches A et C et sur la face externe de la couche B. Ils permettent de plus de déterminer la position de la trajectoire en ϕ .

2.3.5 La mesure de la luminosité

Les moniteurs de luminosité [67, 68] sont placés à l'avant et à l'arrière du détecteur sur les faces internes des bouchons à approximativement 135 cm du centre du détecteur. Ils couvrent la région $2.7 < |\eta| < 4.4$. Ils sont formés de 24 pétales identiques de scintillateurs reliés à des

photo-multiplificateurs (ronds sur la Figure 2.24).

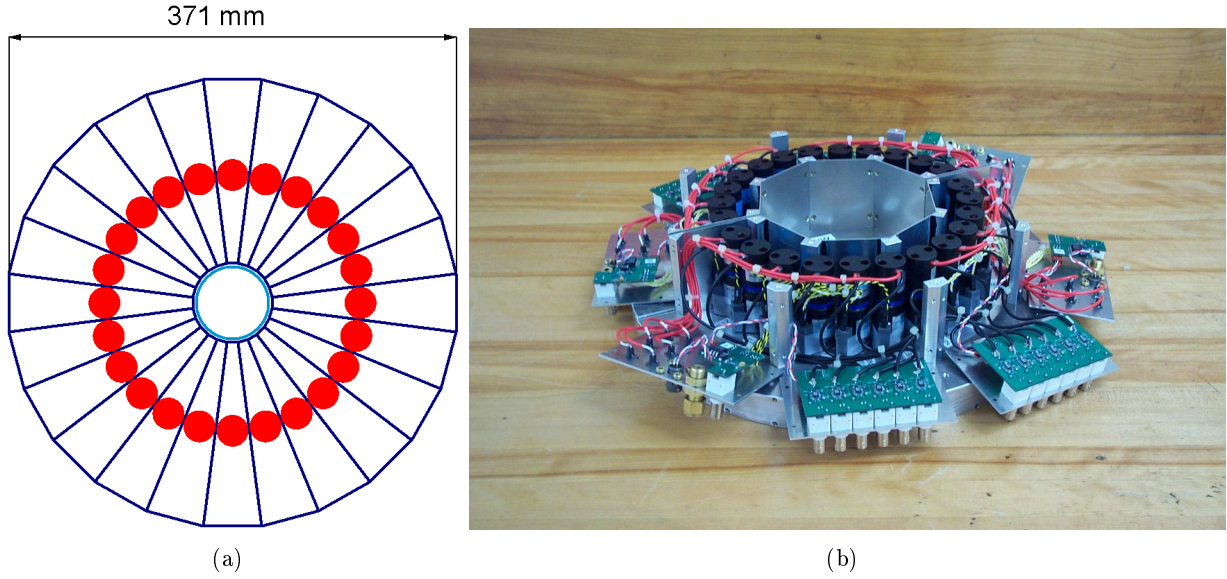


FIG. 2.24 – Schéma du moniteur de luminosité (a) et photographie de ce moniteur (b).

Les scintillateurs sont chargés d'identifier les particules qui les traversent lors d'une interaction inélastique. Pour déterminer si une particule détectée provient bien d'une collision inélastique, on utilise la différence de temps entre les moniteurs de luminosité avant et arrière. On considère en effet qu'une collision inélastique à une position du vertex $|z| < 100$ cm (mesurée avec une précision de 6 cm) et que les temps de vol des particules provenant de la collision sont identiques. La coïncidence en temps des deux moniteurs de luminosité permet ainsi de rejeter l'essentiel des particules solitaires du "halo".

Le détecteur de luminosité compte alors les interactions inélastiques $p\bar{p}$ par unité de temps : $\frac{dN}{dt}$. Pour accéder à la luminosité, il suffit de normaliser cette quantité par la section efficace effective des collisions inélastiques (σ_{eff}), définie comme le produit de la section efficace théorique ($\sigma_{p\bar{p}}$) par l'efficacité des moniteurs (ϵ) et par leur acceptance (A). La formule 2.11 résume ce calcul :

$$\mathcal{L} = \frac{1}{A \times \epsilon \times \sigma_{p\bar{p}}} \frac{dN}{dt} \quad (2.11)$$

Le nombre de collisions inélastiques par unité de temps est paramétré par le produit de la fréquence de croisement des faisceaux (f) et du nombre moyen d'interactions par croisement (μ). Les moniteurs de luminosité sont incapables de distinguer plusieurs interactions par croisement. La mesure de μ n'est pas directe. La probabilité d'avoir n interactions par croisement suit une loi de Poisson :

$$P(n) = \frac{\mu^n}{n!} e^{-\mu} \quad (2.12)$$

Par conséquent, la probabilité d'avoir au moins une interaction par croisement est donnée par $P(n > 0) = 1 - P(0)$. Ainsi, on estime μ en fonction de $P(n > 0)$:

$$\mu = -\ln(1 - P(n > 0)) \quad (2.13)$$

La probabilité d'avoir au moins une collision par croisement de faisceaux peut être estimée par le rapport du nombre de fois où les moniteurs ont détecté une collision inélastiques (N_{inel}) sur le nombre de *ticks* (N_{ticks}), c'est-à-dire le nombre de croisements de faisceaux différents (il y a 159 *ticks* au total, ils correspondent aux coups d'horloge de la Figure 2.8). Cette quantité est calculée pendant un intervalle de temps suffisamment court pour que la luminosité ne diminue pas significativement pendant cet intervalle. Cet intervalle de temps est appelé bloc de luminosité ou en anglais *Luminosity Block Number* (*LBN*), et est de 60 s. De plus, $P(n > 0)$ peut varier en fonction des caractéristiques des paquets de protons et d'antiprotons, qui entrent en collisions. On estime donc $P(n > 0)$ pour chacun des 159 croisement potentiels. Finalement, la luminosité mesurée par l'expérience DØ est donnée par la formule 2.14.

$$\mathcal{L} = -\frac{1}{A \times \epsilon \times \sigma_{p\bar{p}}} \frac{1}{159} \sum_{i=1}^{159} \ln\left(1 - \frac{N_{inel}}{N_{ticks}/159}\right) \quad (2.14)$$

La section efficace théorique est bien connue et la valeur de $\sigma_{p\bar{p}} = 60.7 \pm 2.4$ mb est utilisée par les deux collaborations CDF et DØ. La détermination de l'acceptance et de l'efficacité des moniteurs de luminosité ont récemment été réajustées [69] pour les valeurs suivantes :

- $A = 79 \pm 3\%$
- $\epsilon = 85 \pm 2\%$

Calorimètre	<i>CCEM</i>	<i>CCFH</i>	<i>CCCH</i>	<i>ECEM</i>	<i>ECIH</i>		<i>ECMH</i>		<i>ECHOH</i>
					FH	CH	FH	CH	
# de modules	32	16	16	1	1		16		16
# de couches	4	3	1	4	4	1	4	1	1
Epaisseur d'absorbeur (mm)	3 (U)	6 (U)	46.5 (Cu)	4 (U)	6 (U)	46.5 (Inox)	6 (U)	46.5 (Inox)	46.5 (Inox)
Epaisseur des niveaux	2, 2, 6.8 et 9.8 X_0	1, 1.3 et 0.9 λ_A	3.2 λ_A	0.3, 2.6, 7.9 et 9.3 X_0	1.2 λ_A chacun	3.6 λ_A	1 λ_A chacun	4.1 λ_A	7 λ_A
X_0 total	20.6	96	32.9	20.6	121.8	32.8	115.5	37.9	65.1
λ_A total	0.76	3.2	3.2	0.95	4.8	3.6	4	4.1	7
Fraction d'échantillonnage	11.8 %	6.8 %	1.5 %	11.9 %	5.7 %	1.5 %	6.7 %	1.6 %	1.6 %
Couverture en $ \eta $	≤ 1.2	≤ 1	≤ 0.6	$\geq 1.4, \leq 4$	$\geq 1.6, \leq 4.5$	$\geq 2, \leq 4.5$	$\geq 1, \leq 1.7$	$\geq 1.3, \leq 1.9$	$\geq 0.7, \leq 1.4$
# de cellules	≈ 11370	≈ 3000	≈ 1120	≈ 7490	≈ 4290	≈ 930	≈ 1430	≈ 1340	

TAB. 2.4 – Tableau récapitulatif des caractéristiques du calorimètre.

Chapitre 3

Des collisions aux données reconstruites

Sommaire

3.1	Introduction	86
3.2	Le système de déclenchement de DØ au <i>Run IIa</i>	86
3.2.1	Un système de sélection à trois niveaux	87
3.2.2	Les outils utilisés pour l'étude et la conception de nouvelles conditions de déclenchement	91
3.2.3	Les conditions de déclenchement spécifiques aux signaux à jets de hautes énergies transverses et à haute énergie transverse manquante	92
3.3	Améliorations apportées au <i>Run IIb</i>	104
3.3.1	Description détaillée des modifications apportées pour la partie calorimétrique	106
3.3.2	Etalonnage du système de déclenchement du calorimètre au niveau 1	111
3.3.3	Conception des conditions de déclenchement spécifiques aux signaux à jets et $\cancel{E}_T$ pour le <i>Run IIb</i> : la liste <i>V15</i>	125
3.4	Reconstruction des objets physiques	136
3.4.1	Les traces et les vertex	136
3.4.2	Les muons	138
3.4.3	Les particules électromagnétiques	141
3.4.4	Les jets	144
3.4.5	L'énergie transverse manquante	162
3.4.6	L'étiquetage des jets de hadrons beaux	163
3.4.7	Conclusion	170

3.1 Introduction

La compréhension des données issues des collisions nécessite un long et complexe travail préliminaire avant toute analyse de physique. Deux chaînes parallèles sont nécessaires pour comparer les données aux prédictions théoriques (voir la Figure 3.1). La première consiste en un tri des données en ligne (voir la Section 3.2), suivi d'une reconstruction des objets physiques tels les jets, les électrons, les muons (voir la Section 3.4). La seconde chaîne permet d'obtenir les mêmes objets physiques à partir des prédictions théoriques. Ces prédictions permettent une simulation des événements du modèle standard, de la même façon les signaux supersymétriques recherchés peuvent être simulés.

Une fois la reconstruction des objets physiques effectuée pour ces deux chaînes, la comparaison des données aux simulations est possible. On compare ainsi les distributions des variables associées aux objets physiques détaillés en Section 3.4 pour différentes sélections. Ces sélections permettent d'isoler certains processus physiques pour des études plus spécifiques. Elles permettent également de détecter d'éventuelles déviations entre les prédictions théoriques du modèle standard et les données enregistrées par le détecteur DØ. Les déviations peuvent être dues aussi bien à des problèmes isolés dans l'acquisition des données qu'à la découverte de nouvelle physique. Dans le premier cas, on corrige les problèmes. Des exemples de problèmes dans la qualité des données seront donnés dans le Chapitre 4. Dans ce même chapitre, on trouvera également une recherche de déviations dues à de la nouvelle physique. Quelle que soit l'origine de la déviation, il est primordial de maîtriser les différentes étapes des deux chaînes de traitement des données, c'est-à-dire de l'acquisition et le traitement hors-ligne pour les données du détecteur ou de la simulation pour les données simulées jusqu'à la reconstruction des objets physiques.

Dans ce chapitre, je décris les étapes successives de ces deux chaînes parallèles en insistant particulièrement sur le système de déclenchement. Mes travaux sur cette étape de la chaîne d'acquisition ont en effet constitué la majeure partie de mes activités sur DØ en amont de l'analyse de physique. Je développerai le suivi et la conception des conditions de déclenchement propres aux états finaux avec des jets de hautes énergies et de l'énergie transverse manquante, $\cancel{E}_T$ (voir la Section 3.2.3). Cette activité concerne le *Run IIa*. Pour le *Run IIb*, les modifications apportées au système de déclenchement ont nécessité un nouvel étalonnage au niveau 1 de la partie calorimétrique. Les améliorations apportées et cet étalonnage sont détaillées en Section 3.3.

3.2 Le système de déclenchement de DØ au *Run IIa*

Le but du système de déclenchement [70] est de trier en ligne les collisions dites intéressantes. Les collisions générant des processus physiques déjà étudiés et/ou bien connus sont jugées "inintéressantes". La plupart de ces processus physiques ont d'ailleurs des sections efficaces de plusieurs ordres de grandeurs supérieurs aux processus que l'on souhaite étudier au TeVatron comme la production de paires de quark top ou bien de possibles événements supersymétriques. De plus, le TeVatron produisant un croisement de faisceaux toutes les 396 ns correspondant à une fréquence d'environ 2.5 MHz, il est technologiquement et financièrement impossible d'enregistrer tous les événements correspondant à ces collisions. Ainsi, le système de déclenchement est un outil indispensable pour enregistrer sur bandes les événements, qui seront ensuite étudiés pour les analyses de physique s'inscrivant dans le programme de physique du TeVatron. Il est important de noter que le choix des événements enregistrés est irréversible. Par conséquent, la conception des condi-

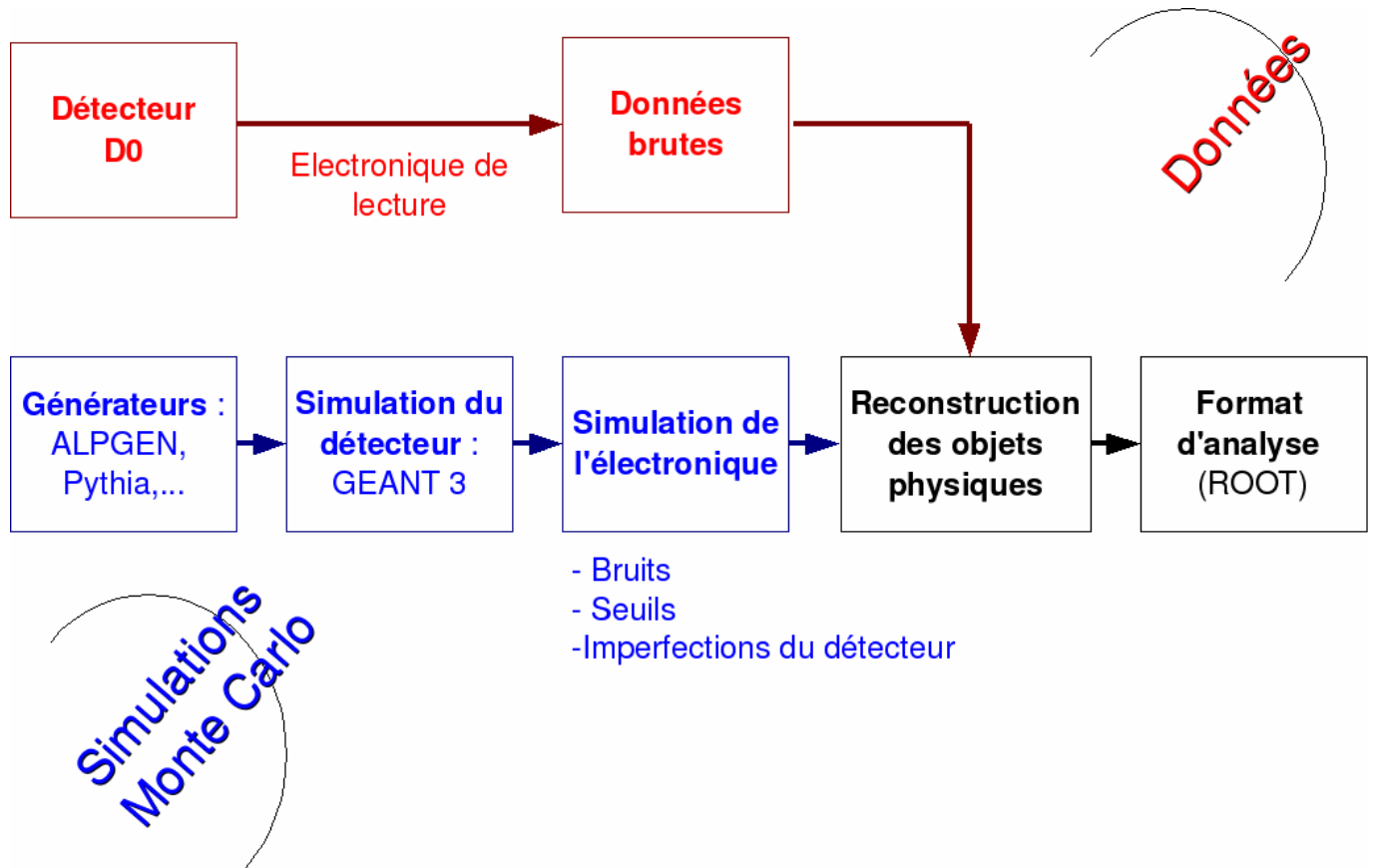


FIG. 3.1 – Schéma des deux chaînes d'acquisition parallèles.

tions de déclenchement, ensemble de critères à respecter pour conserver un événement, doit être bien réfléchi. Les détails d'une conception de conditions de déclenchement sont donnés en Section 3.2.3 pour le *Run IIa* et en Section 3.3 pour le *Run IIb*. Quelques généralités sur le système de déclenchement de DØ sont fournies en Section 3.2.1.

3.2.1 Un système de sélection à trois niveaux

L'architecture du système de déclenchement de DØ est découpée en trois niveaux. A chaque niveau le temps de calcul nécessaire au traitement d'un événement augmente. Par voie de conséquence, la complexité des conditions de déclenchement s'accroît. D'un niveau à l'autre, la proportion d'événements choisis diminue alors que la sélection s'affine (voir la Figure 3.2). A noter que ce système de déclenchement est conçu pour limiter au maximum les temps morts. En effet, un système complexe de communication entre les niveaux permet d'enregistrer plusieurs événements alors que la condition de déclenchement du niveau supérieur est en cours. C'est ce même système de communication qui centralise les décisions de chacun des niveaux. Un menu de déclenchement est composé d'une liste de conditions pour chacun des trois niveaux. On appelle termes de déclenchement les critères pour satisfaire une condition de déclenchement. Chaque condition de niveau 3 est associée à une condition de niveau 2, qui est elle-même associée à une condition de niveau

1. L'événement est finalement enregistré s'il a déclenché au moins une des conditions de niveau 3, ainsi que les conditions de niveaux inférieurs qui lui sont associées. Dans la suite, on parle de conditions de déclenchement pour désigner l'ensemble des conditions associées aux trois niveaux.

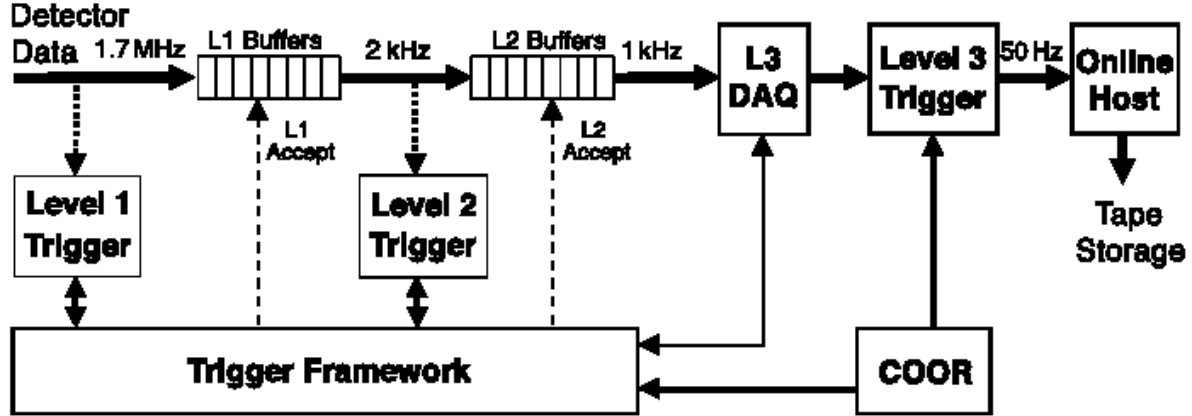


FIG. 3.2 – Schéma du système de déclenchement et bandes passantes.

Niveau 1 (L1)

Le niveau 1 est limité à un taux de sortie d'environ 1.6 kHz pour un taux d'entrée approximativement mille fois supérieur. Le détecteur de luminosité permet la sélection des collisions inélastiques à chaque croisement de faisceaux, cette sélection réduit alors le taux d'entrée des 2.5 MHz initial au 1.7 MHz de la Figure 3.2. L'information utilisée pour déclencher à ce niveau n'est pas numérisée, ce qui permet un temps de décision très rapide d'environ 9 μ s. En contrepartie, les termes de déclenchement sont relativement simples et les communications entre les sous-détecteurs sont limitées comme le montre la Figure 3.3. De plus, il est possible de définir un système de déclenchement pour chacun des sous-détecteurs excepté le détecteur de vertex qui ne dispose pas d'une électronique de lecture suffisamment rapide.

Le calorimètre utilise ainsi, comme objet de base pour le déclenchement, la tour de déclenchement, ensemble semi-projectif de cellules recouvrant une surface $\Delta\eta \times \Delta\phi = 0.2 \times 0.2$ (voir l'illustration de la Figure 3.4). Deux types de tours sont ainsi définis : les tours dites *EM*, pour électromagnétiques, et les tours dites *HD*, pour hadroniques. Les tours hadroniques n'utilisent pas l'information de la partie grossière du calorimètre (*CH*), ni l'information des parties *ICD* et *MG* du calorimètre, en raison de bruits électroniques trop importants. Finalement, la couverture angulaire de ces tours est réduite à $|\Delta\eta| < 3.2$. Les termes de déclenchement utilisant les tours sont très simples pour le *Run IIa*. Un événement passe au deuxième niveau s'il possède un certain nombre de tours au-dessus d'un seuil en énergie (voir exemples en Section 3.2.3). Les algorithmes de sélection sont plus complexes depuis le printemps 2006, c'est-à-dire depuis le début du *Run IIb* (voir la Section 3.3).

Le système de déclenchement propre aux traces utilise l'information issue du *CFT* et des détecteurs de pieds de gerbes. Les détecteurs de pieds de gerbes permettent d'étendre la couverture angulaire de $|\Delta\eta| < 1.6$ à $|\Delta\eta| < 2.5$. Pour qu'un événement satisfasse les termes définis par ce système de déclenchement, il doit dépasser les seuils sur l'impulsion transverse calculée à partir

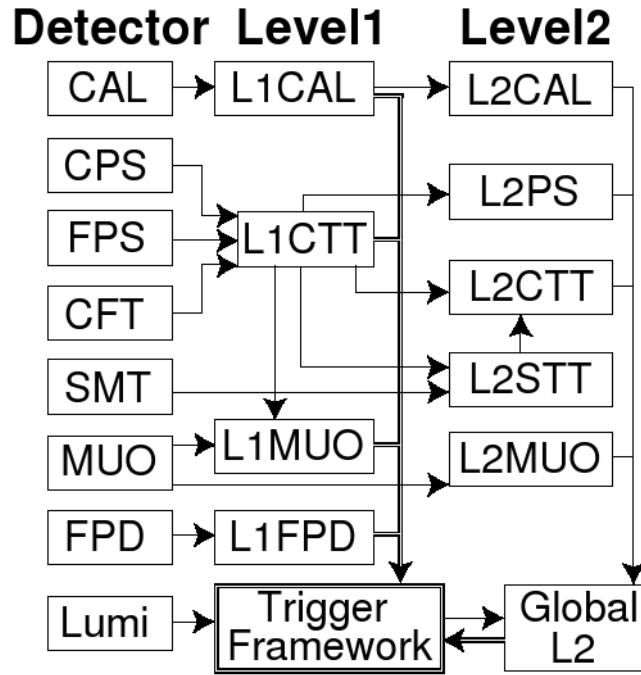


FIG. 3.3 – Schéma du niveau 1 et 2 du système de déclenchement.

des coups dans les 8 couches axiales du détecteur de traces. On remarquera que le *SMT* n'est pas utilisé à ce niveau. Néanmoins, sa lecture est commandée par le déclenchement du niveau 1, ce qui réduit la bande passante disponible au niveau 1.

Enfin, le déclenchement sur les muons est très précis et efficace dès le niveau 1. Les termes sont en effet définis à partir des scintillateurs, des chambres à dérives et de l'information combinée provenant du détecteur de traces. Le système déclenche finalement sur des muons ayant une impulsion transverse suffisante. Ce système de déclenchement permet de plus le rejet des muons provenant de rayons cosmiques.

Une fois, qu'un événement passe le niveau 1, son traitement à l'aide d'algorithmes plus complexes est possible au niveau 2 du système de déclenchement (voir la Figure 3.3).

Niveau 2 (L2)

Le niveau 2 [71, 72] permet de réduire le taux de données d'un facteur 2 approximativement. Il passe ainsi de 1.6 kHz à 800 Hz en introduisant moins de 5% de temps mort. Ce niveau 2 est composé de deux sous-parties. La première est formée de pré-processeurs qui traitent l'information de chacun des sous-détecteurs. Le terme pré-processeur est une traduction directe du terme anglais *pre-processor*. En réalité, ce sont de véritables processeurs informatiques (Pentium III ou IV). On utilise le préfixe "pré" pour désigner un premier traitement de l'information en amont de la sélection des événements. Chacun de ces pré-processeurs forme en effet une liste d'objets qui seront ensuite utilisés par le processeur global pour la sélection des événements. Cette seconde sous-partie permet par conséquent de faire jouer les coïncidences et les corrélations entre les différents sous-détecteurs pour une sélection plus efficace des événements jugés intéressants. Les pré-processeurs

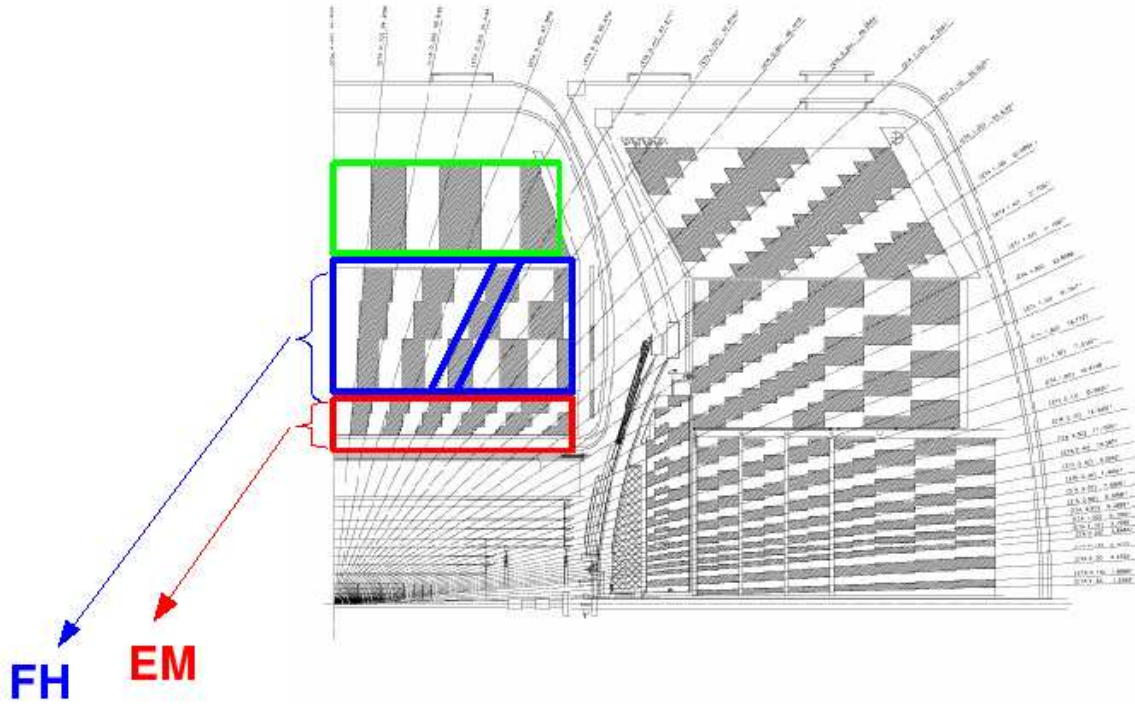


FIG. 3.4 – Schéma de la partie calorimétrique du niveau 1.

ont environ $50 \mu\text{s}$ pour former la liste des objets et le processeur global a le même temps pour prendre sa décision.

Le temps de calcul étant plus long pour ce niveau de déclenchement, les objets formés par les pré-processeurs sont plus complexes et plus élaborés qu'au niveau 1. Par exemple, le pré-processeur, traitant l'information provenant du calorimètre, permet la formation de jets. Ces jets sont construits à l'aide d'une surface 5×5 de tours de déclenchement. La direction du jet pointe du centre du détecteur vers le centre de la surface composée des tours de déclenchement. De même, les pré-processeurs fournissent une reconstruction des vertex, des muons, etc. Cette reconstruction reste néanmoins limitée par le temps de calcul disponible au niveau 2, ce qui implique une résolution sur l'énergie ou l'impulsion plus faible que celle des objets finaux. La position de ces objets dans le détecteur est pour les mêmes raisons moins précise que la position des objets reconstruits hors-ligne.

Une nouvelle partie du détecteur $D\bar{O}$ est utilisée par le niveau 2 pour sélectionner les événements en ligne : le *SMT*. L'information de ce détecteur combinée à celle du *CFT* forme le système de déclenchement appelé *L2STT* (voir la Figure 3.3). Le *L2STT* permet de déclencher sur la présence de vertex déplacés, donc sur la présence de particules à long temps de vie.

Niveau 3 (L3)

Le niveau 3 de déclenchement [73] est constitué d'une ferme d'ordinateurs. Ce stade ultime de la sélection en ligne permet une reconstruction des objets physiques, qui se rapprochent des objets

utilisés pour les analyses de physique [74]. L'information numérisée de tous les sous-détecteurs est envoyée à cette ferme de calcul qui a 50 à 70 ms pour prendre une décision et qui doit réduire le flux des données d'un facteur 10 environ. En effet, de 50 à 80 Hz de données peuvent être enregistrés sur bande magnétique (au format appelé *raw data*) pour être ensuite analysés par les algorithmes de reconstruction hors-ligne (*offline*). La sélection des événements n'est pas linéaire pour réduire au maximum les temps morts. Ce dispositif est légèrement plus complexe. Tout d'abord, l'information du niveau 1 et 2 n'est pas utilisée par le niveau 3. Ensuite, si un événement passe une condition du niveau 1, les données brutes de chaque sous-détecteur sont numérisées et en attente d'une décision positive du niveau 2. L'électronique de lecture peut enregistrer jusqu'à 16 événements ce qui laisse le temps au niveau 2 de prendre une décision. Finalement, dans le cas favorable, ce sont ces données numérisées qui sont envoyées vers la ferme d'ordinateurs. Pour ce dernier niveau de déclenchement, on parle de filtres pour désigner les termes de déclenchement. Et contrairement au niveau 1 et 2, limités à 128 termes chacun (au *Run IIb* 4096 termes sont disponibles au niveau 2), le niveau 3 ne possède pas un nombre de filtres fixés. Deux exemples de filtres seront détaillés en Section 3.2.3.

3.2.2 Les outils utilisés pour l'étude et la conception de nouvelles conditions de déclenchement

Pour étudier des conditions de déclenchement, il faut être en mesure d'estimer ou de mesurer le taux d'acquisition et les efficacités des coupures de sélection. Le but pour la conception de nouvelles conditions de déclenchement est en effet d'avoir la meilleure efficacité pour le signal recherché tout en ayant un taux réduit.

Le taux d'acquisition dépend des conditions fixées et de la luminosité instantanée. Pour une étude, il faut alors prévoir l'évolution de ce taux en fonction des coupures de sélection sur les variables disponibles à chaque niveau du système de déclenchement. Mais il faut aussi estimer l'évolution en fonction de la luminosité instantanée pour prévoir la luminosité instantanée maximale sous laquelle le système de déclenchement peut fonctionner. En cas de taux trop hauts, plusieurs solutions sont possibles : le durcissement des coupures, une nouvelle conception à base de nouveaux algorithmes, l'arrêt temporaire de certaines conditions de déclenchement au-dessus d'une certaine luminosité ou le rejet aléatoire d'événements passant une condition de déclenchement (on parle alors de *prescale* en anglais, la valeur de ce *prescale* est le facteur de réduction appliqué sur une condition de déclenchement.). Le dernier cas est bien sûr le moins souhaitable et à éviter pour une condition de déclenchement dédiée à la recherche de nouvelle physique. Le premier cas peut être une solution en cas d'augmentations isolées de la luminosité instantanée, enfin le deuxième cas est préféré pour une augmentation importante telle celle attendue pour le *Run IIb*. Ces deux derniers cas seront détaillés dans la suite. Pour prévoir ces dépendances du taux d'acquisition, un outil a été mis au point par la collaboration DØ : *trigger_rate_tool* [75]. Cet outil donne le taux d'acquisition propre à chacune des conditions de déclenchement, ainsi que le taux total de toutes les conditions de déclenchement en tenant compte des recouvrements entre les différentes conditions. Pour estimer l'évolution avec la luminosité instantanée, *trigger_rate_tool* est utilisé sur des périodes de prise de données spéciales avec des conditions très lâches au niveau 1 pour enregistrer le maximum de données et éviter de biaiser les estimations. A partir de ces prises de données, on peut reconstruire des termes de déclenchement et par extrapolation linéaire estimer le taux d'acquisition à une luminosité instantanée donnée.

Le deuxième point important pour l'étude d'un système de déclenchement est la détermination

de l'efficacité d'une coupure. Pour les données des collisions, l'information des variables définies aux trois niveaux est en général directement stockée dans les *raw data*. Par contre, pour les données issues des simulations *Monte Carlo*, cette information doit être simulée. Il est nécessaire d'estimer ces variables à partir de l'information hors-ligne pour déterminer si oui ou non un événement respecte une condition de déclenchement. C'est le rôle de l'outil appelé *d0trigsim* [76]. Cet outil permet par exemple d'estimer et d'optimiser l'efficacité d'une condition de déclenchement sur un signal supersymétrique.

Ces deux outils sont utilisés conjointement pour les études présentées par la suite.

3.2.3 Les conditions de déclenchement spécifiques aux signaux à jets de hautes énergies transverses et à haute énergie transverse manquante

Présentation et historique

Les conditions de déclenchement propres aux topologies à jets et $\cancel{E}_T$ sont au nombre de trois depuis le début du *Run IIb*, juin 2006. Ces trois conditions de déclenchement, se différenciant essentiellement par le nombre de jets requis dans l'état final, sont détaillées dans la Section 3.3. La mise en place d'une méthode déclenchant sur ces topologies particulières est utilisée par la collaboration DØ depuis mars 2003. Des détails sur la mise au point de cette méthode et sur la conception de la première condition de déclenchement, appelé *MHT30*, peuvent être trouvés en [77, 78]. Depuis, deux types de méthode et donc de conditions de déclenchement ont été conçus et ont évolué jusqu'à la fin du *Run IIa* [79, 80]. En tant que responsable de ces deux conditions de déclenchement pendant un an, j'ai été amené à suivre leur comportement avec la luminosité instantanée et finalement à apporter des modifications pour le niveau 1, comme il sera détaillé dans les paragraphes suivants. Ainsi depuis mars 2003, la liste des conditions de déclenchement spécifiques aux topologies à jets et $\cancel{E}_T$ a évolué de une à trois conditions comme le montre le Tableau 3.1. De cette façon, chacune des conditions possède ses propres spécificités et des coupures plus dures que l'unique condition conçue initialement. Cette évolution a été guidée par la nécessité de réduire le taux d'acquisition tout en se rapprochant au maximum des topologies des signaux recherchés.

Nom du menu	Date de début	Nom des conditions			Luminosité intégrée (pb^{-1})
V11	25/03/03	MHT30			≈ 63
V12	16/07/03	MHT30			≈ 227
V13	25/06/04	JT1_ACO_MHT_HT	JT2_MHT25_HT		≈ 372
V14	13/07/05	JT1_ACO_MHT_HT	JT2_MHT25_HT		≈ 330
V15	09/06/06	<i>monojet</i>	<i>dijet</i>	<i>multijet</i>	en cours

TAB. 3.1 – Historique des conditions de déclenchement spécifiques aux topologies à jets et énergie transverse manquante du *Run IIa* au *Run IIb*.

Les variables utilisées pour définir les conditions de chaque niveau

Comme expliqué dans la Section 3.2.1, le système de déclenchement possède trois niveaux et d'un niveau à l'autre les variables utilisées pour la sélection des événements sont de plus en plus

complexes.

Le premier niveau propose des termes standard de déclenchements sur les jets : les termes $CJT(n, X, Y)$. La lettre "C" désigne le calorimètre, alors que les lettres "JT" est la nomenclature utilisée pour les jets. De la même façon, il existe des termes $CEM(n, X, Y)$ pour les objets électromagnétiques. Le n correspond au nombre de tours de déclenchement définies en Section 3.2.1, au-dessus d'un seuil en énergie transverse fixée à X GeV et vérifiant $|\eta| < Y$. Ces tours peuvent être électromagnétiques pour les termes de type CEM ou totales (somme de la partie électromagnétique et de la partie hadronique) pour les termes de type CJT . Ces termes sont communs à de nombreuses conditions de déclenchement, par exemple ils sont utilisés pour les études sur le fond multijet, dit *QCD*. Par exemple, la première condition de déclenchement utilisait le terme $CJT(3, 5, 3.2)$, c'est-à-dire : au moins trois tours totales au-dessus de 5 GeV étaient requises pour passer au niveau 2. Finalement, une condition au niveau 1 peut être une combinaison logique de plusieurs de ces termes. C'est la raison pour laquelle ces termes font partie de ce qui est communément appelé par la collaboration DØ les *And/Or Terms*.

Au niveau 2, les tours totales du niveau 1 sont utilisées pour reconstruire des jets [81]. Les algorithmes de reconstruction sont simplifiés pour limiter le temps de calcul. Ainsi, un jet de niveau 2 est constitué de 5×5 tours, alors qu'un jet *offline* est en première approximation constitué d'un cône de rayon $\Delta R = \sqrt{(\Delta\eta)^2 + (\Delta\phi)^2} = 0.5$ (voir les détails dans la Section 3.4). Le barycentre de ces 25 tours permet de déterminer les coordonnées η et ϕ en prenant comme origine le centre du détecteur au lieu du vertex primaire de l'interaction pour un jet *offline*. A partir des jets de niveau 2, on peut construire les différentes variables qui sont utilisées pour définir les termes de déclenchement. Les variables suivantes sont celles utilisées par les conditions de déclenchement décrites dans cette section :

- $L2JET(n, X)$: avec X le seuil en énergie transverse et n le nombre de jets au-dessus de ce seuil.
- $L2MJT(X)$: ce terme correspond à la somme vectorielle des jets au-dessus d'un seuil en énergie transverse (fixé à 10 GeV dans notre cas). Ce terme est satisfait si cette somme est supérieure à X GeV. Il permet de sélectionner les événements possédant des jets et une forte énergie transverse manquante. Au *Run IIa*, il n'existait pas d'énergie transverse manquante définie au niveau 2 et une étude sur cette variable est en cours pour une future utilisation au *Run IIb*.
- $L2HT(X)$: ce terme est défini à partir de la somme scalaire de l'énergie transverse des jets et est satisfait si cette somme est supérieure à X GeV. Initialement, ce terme n'utilisait que les jets au-dessus de 10 GeV sans restriction en η pour le calcul de la somme. A partir de juin 2005 (menu de déclenchement *V14*) et donc pour les études menées dans les paragraphes suivants, cette somme n'est calculée qu'avec les jets au-dessus de 6 GeV et vérifiant $|\eta| < 2.6$.
- $L2ACOP(X)$: ce terme est satisfait si l'événement possède un jet uniquement ou si la distance angulaire entre les deux jets de plus haute énergie transverse est inférieure à X degrés, c'est-à-dire si $\Delta\phi(jet_1, jet_2) < X$ (angle azimutal entre ces deux jets). De même que pour les deux variables définies au-dessus, $L2ACOP(X)$ n'utilisait initialement que les jets de plus de 10 GeV. Ce seuil a été ramené à 5 GeV par la suite et c'est ce dernier qui est utilisé pour les résultats détaillés dans les paragraphes suivants. Le plan transverse au niveau 2 est divisé en 160 secteurs de 2.25 degrés, il existe donc 80 valeurs de X possibles. La valeur utilisée pour les conditions de déclenchement décrites dans la suite est 168.75.

Finalement, les jets de niveau 3 sont des objets plus complexes et très proches de ceux recons-

truits *offline*. Plutôt que d'utiliser l'information des tours de déclenchement, le niveau 3 utilise l'information des cellules du calorimètre pour reconstruire les jets. La résolution sur l'énergie des jets de niveau 3 est donc meilleure que celle des jets de niveau 2. De plus, l'information des traces provenant du *SMT* et du *CFT* permet de déterminer la position selon la variable z (de direction parallèle à l'axe du faisceau) du vertex primaire de l'interaction. A partir de cette information, la direction ϕ et η du jet est calculée. Cette information sur la position du vertex primaire n'existait pas initialement et de la même façon qu'au niveau 2 l'origine du détecteur était choisie. Finalement, le jet est défini à partir d'un algorithme de simple cône de rayon 0.5 (initialement ce rayon était de 0.7). Seuls les jets de plus de 9 GeV sont ensuite utilisés pour définir les variables caractérisant les filtres du niveau 3. Ces variables sont les mêmes que celles utilisées au niveau 2, ainsi qu'une nouvelle variable topologique introduite en juillet 2005, c'est-à-dire :

- $L3JET(n, X, Y)$: pour chacune des variables la lettre n correspond au nombre minimum de jets requis, X au seuil en énergie transverse imposée sur ces jets et vérifiant $|\eta| < Y$.
- $L3MHT(n, X, Y)$: ce terme désigne la somme vectorielle de l'énergie transverse des jets. L'énergie transverse manquante n'est pas utilisée pour les conditions de déclenchement décrites ici, néanmoins des études sur cette variable ont été menées par la collaboration DØ [82].
- $L3HT(n, X, Y)$ pour la somme scalaire de l'énergie transverse des jets.
- $L3ACOP(n, X, Y)$ pour l'angle azimutal entre les deux jets de plus haute énergie transverse. Ce filtre est satisfait s'il n'y a qu'un seul jet au-dessus de 9 GeV dans l'événement.
- $L3AngleMHTJET(n, X, Y)$: ce dernier filtre calcule l'angle minimal entre tous les jets et la variable $L3MHT(n, X, Y)$. Il a été introduit pour pallier l'augmentation de la luminosité au cours du *Run IIa* comme il est détaillé dans un paragraphe ultérieur.

La nomenclature définie ici sera utilisée par la suite pour décrire les coupures sur ces variables pour les conditions de déclenchement spécifiques aux topologies à jets et $\cancel{E}_T$.

Les signaux de physiques utilisés pour l'étude des conditions de déclenchement

Les conditions de déclenchement décrites dans cette partie ont été conçues par le groupe de recherche de Nouvelle Physique pour sélectionner efficacement les processus impliquant dans l'état final des jets et de l'énergie transverse manquante. Ces processus couvrent en effet un large domaine de recherche pour l'étude de la physique au-delà du modèle standard ou pour la recherche du boson de Higgs. Il est indispensable d'avoir un système de déclenchement bien maîtrisé et efficace pour ces processus quand on sait que le fond dominant multijet peut avoir des sections efficaces de plusieurs ordres de grandeur supérieures à celles des processus supersymétriques par exemple. La première condition de déclenchement avait été conçue en utilisant la production $HZ \rightarrow \bar{b}b\nu\nu$ ($m_H = 115$ GeV) comme signal de référence. Ce signal possède dans l'état final deux jets acoplanaires plutôt mous. A ce signal, on ajoute les processus suivants pour les travaux décrits dans les paragraphes suivants :

- Production de paires de sbottoms [83] avec pour masses $m(\tilde{b}_1^0) = 155$ GeV et $m(\tilde{\chi}_1^0) = 75$ GeV. Ce signal a une topologie très proche de celle de la production $HZ \rightarrow \bar{b}b\nu\nu$. Il sera utilisé comme complément ou point de vérification pour l'étude des conditions de déclenchement spécifiques aux jets acoplanaires (voir le paragraphe suivant pour des détails sur ces conditions pour la liste *V14* ou la Section 3.3 pour la liste *V15*, première liste du *Run IIb*).
- Production de paires de squarks [84] dans un modèle *mSUGRA* avec les paramètres du

modèle suivants : $m_0 = 25$ GeV, $m_{1/2} = 140$ GeV, $A_0 = 0$ GeV, $\mu > 0$ et $\tan(\beta) = 3$. Ce processus présente une topologie similaire aux deux précédents avec comme particularité des jets de hautes énergies transverses. Il est donc plus aisé de déclencher sur un tel signal, néanmoins il sera utilisé en complément des deux signaux précédents.

- Production de paires de gluinos [84] dans le même cadre que le point précédent mais avec un autre jeu de paramètres $mSUGRA$: $m_0 = 500$ GeV, $m_{1/2} = 90$ GeV, $A_0 = 0$ GeV, $\mu > 0$ et $\tan(\beta) = 3$. Ce processus possède plus de deux jets dans l'état final. L'acoplanarité n'est donc pas une variable efficace pour sélectionner un signal de ce type. Il sera plutôt utilisé pour les études sur la condition de déclenchement avec au moins trois jets dans l'état final et une coupure importante sur la variable H_T . De même, des détails seront donnés dans le paragraphe suivant pour la liste *V14* ou la Section 3.3 pour la liste *V15*.
- Un signal dans un modèle de dimensions supplémentaires [85] avec pour paramètres du modèle $N = 4$ et $M = 800$ GeV. Ce signal possède un unique jet de très haute énergie et une forte énergie transverse manquante contrebalançant ce jet. Ce signal utilisait préférentiellement la condition dédiée aux signaux acoplanaires au cours du *Run IIa*. Il a finalement été décidé de concevoir une condition de déclenchement propre à cette topologie particulière pour la première liste du *Run IIb*).

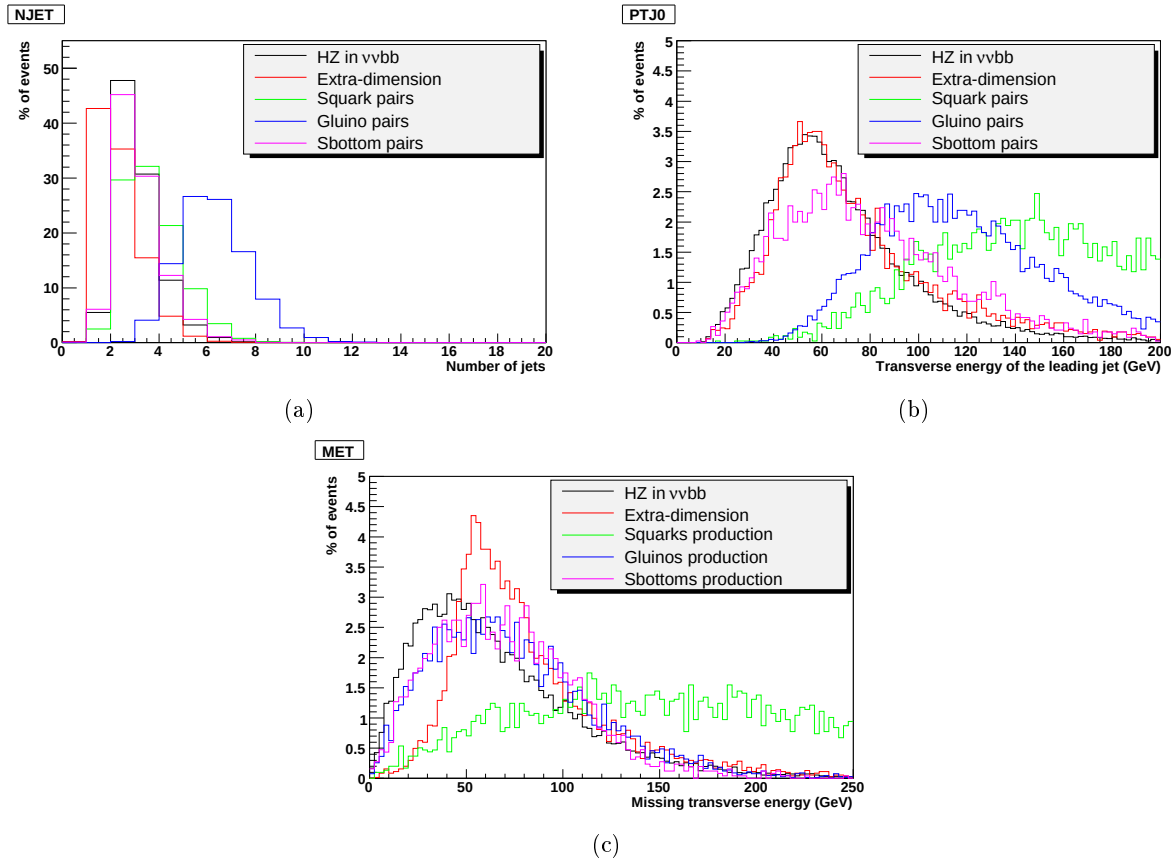


FIG. 3.5 – Distribution du nombre de jets (a), de l'énergie transverse du jet le plus dur (b) et de l'énergie transverse manquante pour les cinq signaux utilisés (c).

La Figure 3.5 montre les caractéristiques de ces cinq signaux utilisés pour la recherche du boson de Higgs et pour la recherche de physique au-delà du modèle standard. On constate notamment que les signaux correspondant à la recherche du boson de Higgs et à la production de sbottoms ont en majorité deux jets avec une énergie transverse minimum de 20 GeV pour le jet le plus dur. Le signal correspondant à la production de squarks a plutôt un état final à deux ou trois jets et avec des jets de plus hautes énergies. Le signal correspondant à la production de gluinos se caractérise par la présence de plus de quatre jets dans l'état final. Enfin le signal utilisé pour la recherche de dimensions supplémentaires possède un seul jet dans plus de 40 % des événements et une énergie transverse manquante de plus de 50 GeV en moyenne.

Ces cinq signaux seront donc utilisés par la suite pour vérifier de nouvelles coupures de sélection, pour en ajouter ou pour concevoir de nouvelles conditions de déclenchement.

Présentation des conditions de déclenchement spécifiques au processus à jets et $\cancel{E}_T$ de la liste *V14*

La liste des conditions de déclenchement, appelée liste *V14*, est la liste qui a été utilisée de l'été 2005 au printemps 2006. Cette liste comporte deux conditions de déclenchement spécialement conçus pour les analyses du groupe de Nouvelle Physique recherchant des processus à jets et $\cancel{E}_T$. La première condition de déclenchement, appelé *JT1_ACO_MHT_HT*, est utilisé pour les topologies finales à un jet ou deux jets acoplanaires comme la recherche de la production $HZ \rightarrow \bar{b}b\bar{\nu}\nu$ ou la recherche de dimensions supplémentaires par exemple. Le second, nommé *JT2_MHT25_HT*, est plus spécifique aux analyses multijet, c'est-à-dire aux analyses avec au moins trois jets dans l'état final. On peut citer en exemple la recherche de production de paires de gluinos.

Ces deux conditions de déclenchement ont évolué de leur conception initiale jusqu'à la fin du *Run IIa* pour être plus largement modifiés au *Run IIb*. Je ne décrirai pas en détails les différentes modifications qui ont eu lieu. Je présenterai un état donné de ces deux conditions, l'été 2005, c'est-à-dire le moment à partir duquel j'ai commencé mon étude sur les conditions de déclenchement propres aux signaux évoqués précédemment. Et je décrirai plus largement les modifications sur lesquelles j'ai travaillé.

Les conditions de déclenchement au démarrage de la liste *V14* sont résumées dans le Tableau 3.2.

Conditions	JT1_ACO_MHT_HT	JT2_MHT25_HT
Level 1	CJT(3,5)CJT(3,4,2.6)	CJT(3,5)CJT(3,4,2.6)
Level 2	L2MJT(20)_L2ACOP(168.75)	L2JET(3,6)_L2HT(75)
Level 3	L3HT(1,9,3.6) > 50 GeV L3MHT(1,9,3.6) > 30 GeV L3ACOP(1,9,3.6) < 170° L3AngleMHTJET(1,9,3.6) > 25°	L3HT(1,9,3.6) > 125 GeV L3MHT(1,9,3.6) > 25 GeV L3JET(3,20,3.6)

TAB. 3.2 – L1,L2 et L3 pour JT1_ACO_MHT_HT et JT2_MHT25_HT.

Le niveau 1 est identique aux deux conditions de déclenchement, ainsi d'ailleurs qu'à toutes les conditions de déclenchement spécifiques aux jets exceptés ceux visant à l'étude de la *QCD*. On constate que, par rapport à la condition *MHT30*, un critère a été ajouté pour réduire le taux du niveau 1 : au moins trois tours comprises dans la région $|\eta| < 2.6$ doivent avoir une énergie

transverse supérieure à 4 GeV. Les deux conditions commencent à se différencier au niveau 2. Celle spécifique aux signaux monojet ou dijet acoplanaire sélectionne plutôt sur l'acoplanarité et sur la variable $\cancel{H}_T$, alors que la seconde demande au moins trois jets et une somme scalaire de l'énergie des jets importante. C'est dans cette logique de différenciation qu'a été conçu le niveau 3. Pour $JT1_ACO_MHT_HT$, les coupures sur l'acoplanarité et $\cancel{H}_T$ sont maintenues. On y ajoute une coupure sur H_T , ainsi que sur l'angle minimum entre tous les jets et le vecteur $\cancel{H}_T$. Cette dernière est particulièrement efficace pour réduire les événements du fond QCD et donc réduire le taux d'acquisition. En effet, pour ces événements, la topologie finale ressemble souvent à deux jets dos à dos. Il suffit qu'un de ces jets soit mal reconstruit pour créer de l'énergie transverse manquante de même sens qu'un des deux jets. Finalement, pour $JT2_MHT25_HT$, une coupure supplémentaire sur $\cancel{H}_T$ permet de réduire le taux d'acquisition. La distribution de l'énergie transverse manquante des événements QCD peut, en effet, en première approximation être assimilée à une exponentielle décroissante. Les autres coupures sont durcies.

Vérifications des modifications apportées à la liste $V14$ pour les signaux squarks et gluinos

La liste $V14$ a permis d'enregistrer des données de juillet 2005 jusqu'à la fin du *Run IIa*. Elle a été conçue pour maintenir un taux d'acquisition raisonnable malgré l'augmentation de la luminosité instantanée et ainsi remplacée la liste $V13$ qui prenait des données depuis janvier 2004. Plusieurs modifications par rapport à cette dernière ont été proposées pour les deux conditions de déclenchement décrites précédemment. L'objectif affiché pour ces modifications étaient alors la réduction du taux d'acquisition de l'ordre de 40% au niveau 3. Les signaux correspondant à la production de squarks et à la production de gluinos ont alors été utilisés pour vérifier que ces modifications n'introduisaient pas une diminution conséquente de l'efficacité. Il aurait également été judicieux d'utiliser les autres signaux évoqués pour compléter cette étude. Cependant, la nécessité de réagir rapidement à l'augmentation de la luminosité instantanée a poussé à agir au plus vite et à utiliser les signaux disponibles au moment de l'étude. En effet, il était préférable de durcir les coupures plutôt que d'arrêter l'utilisation de ces conditions de déclenchement à hautes luminosités.

Le niveau 1 a donc été modifié par l'ajout du terme $CJT(3,4,2.6)$ pour les conditions $JT1_ACO_MHT_HT$ et $JT2_MHT25_HT$. Au niveau 2, deux algorithmes ont été changés. L'algorithme $L2HT$ n'utilise plus que les jets vérifiant $|\eta| < 2.6$ contre $|\eta| < 3.2$ pour la liste $V13$. De plus la coupure sur cette variable a été augmentée de 70 à 75 GeV. Et l'algorithme $L2ACOP$ utilise les jets au-dessus de 10 GeV contre 5 GeV pour la liste précédente. Les principales modifications ont eu lieu au niveau 3. Premièrement, l'algorithme reconstruisant les jets de niveau 3 a été modifié. Au lieu d'utiliser les traces au-dessus de 1 GeV pour déterminer la position du vertex primaire de l'interaction, le nouvel algorithme de jet utilise les traces au-dessus de 3 GeV. Cette modification a permis de réduire les temps de calcul, qui augmentent rapidement avec la luminosité. Deuxièmement, deux coupures importantes et fortement discriminantes contre le fond instrumental QCD ont été ajoutées : la coupure sur $L3AngleMHTJET$ à 25 degrés pour la condition $JT1_ACO_MHT_HT$ et trois jets au-dessus de 20 GeV pour la condition $JT2_MHT25_HT$. Ce sont ces deux dernières coupures qui nécessitent plus particulièrement des vérifications sur la stabilité de l'efficacité du signal.

Avant de présenter les résultats de cette étude, il est nécessaire de bien définir l'efficacité. En effet, deux types d'efficacité peuvent être définis :

- l’efficacité absolue, c’est-à-dire le rapport du nombre d’événements attendus après les trois niveaux sur le nombre d’événement initial,
- l’efficacité relative, qui correspond au rapport du nombre d’événements passant à la fois des coupures de présélection et la condition de déclenchement sur le nombre d’événements passant ces mêmes coupures de présélection.

Cette dernière efficacité est celle sur laquelle on conclura les résultats de cette étude. En effet, ce qui importe est le nombre d’événements éventuellement perdus après les coupures de présélection, sachant que ces coupures étant le niveau de sélection le plus lâche de chacune des deux analyses. Elles sont donc indispensables et systématiquement appliquées.

Les coupures de présélection utilisées pour cette étude sont les suivantes :

- Pour la recherche de la production de paires de squarks :
 - $E_T(j_0) > 60 \text{ GeV}$, $E_T(j_1) > 40 \text{ GeV}$
 - $\cancel{E}_T > 60 \text{ GeV}$
 - $\Delta\phi(j_0, j_1) < 165^\circ$
 - $\Delta\phi_{min}(jets, \cancel{H}_T) > 30^\circ$
- Pour la recherche de la production de gluinos :
 - $E_T(j_0) > 60 \text{ GeV}$, $E_T(j_1) > 40 \text{ GeV}$, $E_T(j_2) > 30 \text{ GeV}$
 - $\cancel{E}_T > 60 \text{ GeV}$

Avec j_0 , j_1 et j_2 , les jets *offline* classés par ordre décroissant d’énergies transverses. La variable E_T représente leur énergie transverse. $\cancel{E}_T$ est l’énergie transverse manquante et $\cancel{H}_T$ la somme vectorielle de l’énergie transverse des jets.

Après avoir appliqué la simulation du système de déclenchement, *d0trigsim*, et les coupures de présélection pour chacun des deux signaux, on obtient les efficacités au niveau 3 correspondant à la production de gluinos pour la condition *JT2_MHT25_HT* dans le Tableau 3.3 et correspondant à la production de squarks pour la condition *JT1_ACO_MHT_HT* dans le Tableau 3.4. Les erreurs sur les efficacités sont les erreurs statistiques. Ces efficacités indiquent que les modifications de la liste *V13* à la liste *V14* n’entraînent pas de pertes significatives de signal pour la condition *JT2_MHT25_HT*. Par contre, l’efficacité absolue calculée pour la condition de déclenchement *JT1_ACO_MHT_HT* chute approximativement de 20 %. Cependant, la véritable perte de signal est d’approximativement 3 %, ce qui reste une perte raisonnable sachant l’objectif, qui était fixé. En effet, ces deux coupures ont pour but de réduire le taux d’acquisition au niveau 3 de l’ordre de 40 %.

Efficacités	V13	V14
Absolues	$87.9 \pm 0.3\%$	$85.0 \pm 0.4\%$
Relatives	$99.7 \pm 0.1\%$	$99.5 \pm 0.1\%$

TAB. 3.3 – Efficacités des coupures appliquées par la condition *JT2_MHT25_HT* sur le signal correspondant à la production de gluinos

Deux méthodes ont été utilisées pour estimer ou vérifier la diminution du taux d’acquisition. La première est une méthode a priori. En utilisant, une période de prise de données de référence antérieure aux modifications des conditions de déclenchement (appelée *Run 203351*, avec une luminosité instantanée initiale de $84.53 \times 10^{30} \text{ cm}^{-2} \text{ s}^{-1}$ et finale de $58.05 \times 10^{30} \text{ cm}^{-2} \text{ s}^{-1}$) et l’outil décrit dans la Section *trigger_rate_tool*, on estime par extrapolation linéaire le taux d’acquisition

Efficacités	V13	V14
Absolues	$86.8 \pm 0.5\%$	$69.9 \pm 0.7\%$
Relatives	$99.0 \pm 0.1\%$	$96.0 \pm 0.3\%$

TAB. 3.4 – Efficacités des coupures appliquées par la condition JT1_ACO_MHT_HT sur le signal correspondant à la production de squarks

qu'aurait le niveau 3 à la luminosité instantanée $100 \times 10^{30} \text{cm}^{-2} \text{s}^{-1}$. La seconde méthode est une méthode a posteriori. Deux tests ont alors été effectués. On utilise tout d'abord différentes périodes de prise de données (*runs*) antérieures et postérieures aux modifications. On moyenne la luminosité intégrée au cours de la période de données et en modélisant le comportement du taux d'acquisition en fonction de la luminosité par une droite, on obtient une estimation à $100 \times 10^{30} \text{cm}^{-2} \text{s}^{-1}$ (voir la Figure 3.6). Ensuite, on peut également utiliser une période de prise de données à haute luminosité instantanée initiale et mesurer la luminosité au cours de cette période en comptabilisant le nombre d'événements enregistrés toutes les minutes (période appelée *LBN* comme l'explique la Section 2.3.5). La Figure 3.7 montre l'évolution du taux d'acquisition le long d'une période de prise de données pour les conditions JT1_ACO_MHT_HT (en rouge), JT2_MHT25_HT (en bleu) et pour le recouvrement entre ces deux conditions (en vert). Les deux méthodes sont en bon accord et donnent une bonne estimation du taux d'acquisition à $100 \times 10^{30} \text{cm}^{-2} \text{s}^{-1}$. La seconde figure montre de plus que l'hypothèse de la linéarité du taux d'acquisition est raisonnable à partir de $40 \times 10^{30} \text{cm}^{-2} \text{s}^{-1}$.

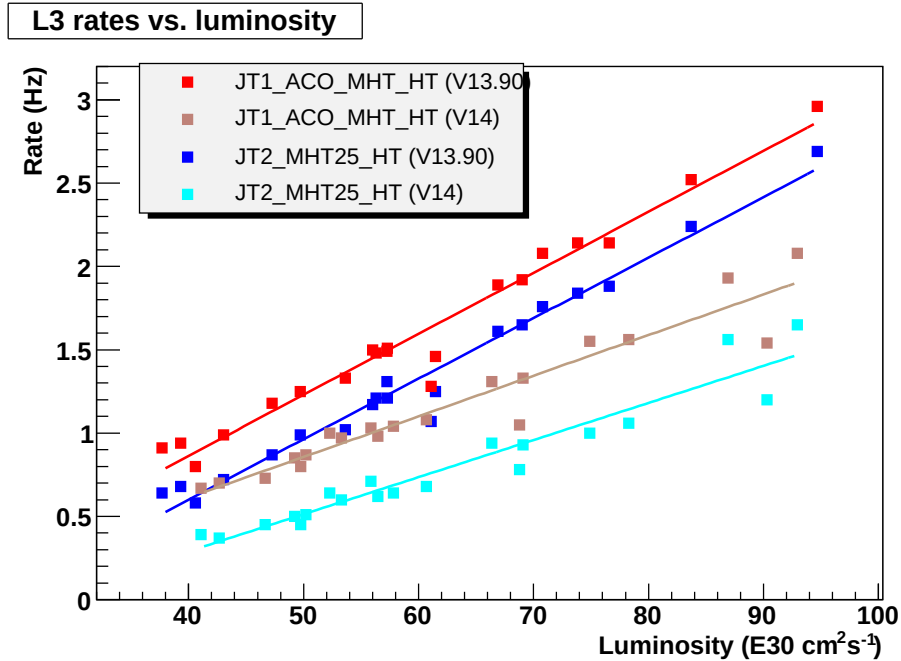


FIG. 3.6 – Taux d'acquisition du niveau 3 pour les conditions de déclenchement JT1_ACO_MHT_HT et JT2_MHT25_HT en fonction de la luminosité instantanée moyenne pour des périodes de prise de données avec les listes V13.90 et V14

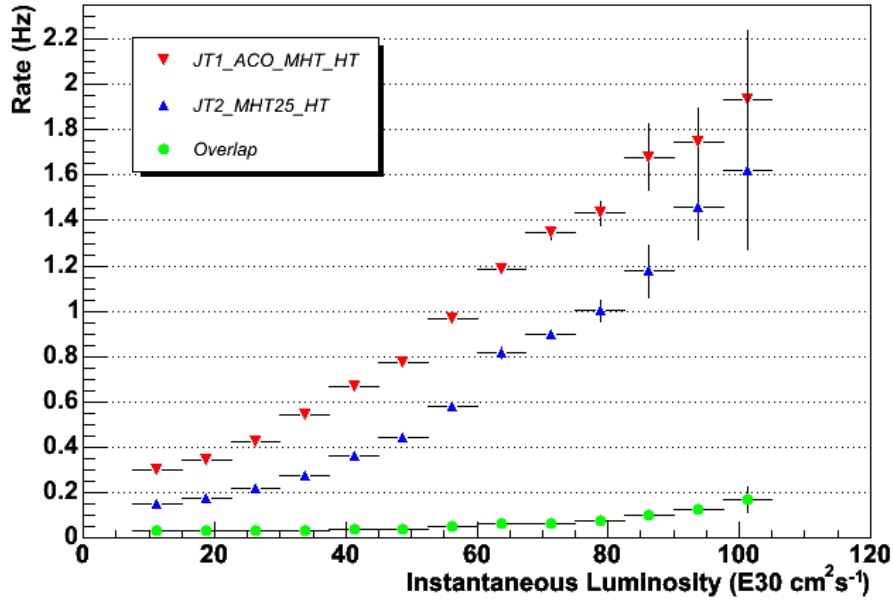


FIG. 3.7 – Taux d’acquisition du niveau 3 pour les conditions JT1_ACO_MHT_HT et JT2_MHT25_HT pour une période de prise de données avec la liste *V14*

Les résultats sont résumés dans le Tableau 3.5 avec en italique les estimations de la méthode a priori. Les erreurs sur ces taux sont une nouvelle fois statistiques. Les deux méthodes donnent des résultats comparables malgré les nombreuses approximations introduites par la méthode dite a priori. Les résultats du tableau indiquent finalement que le taux d’acquisition de la condition de déclenchement *JT1_ACO_MHT_HT* a diminué de plus de 30 %, alors que celui de *JT2_MHT25_HT* a diminué de plus de 40 %. Ces diminutions ont permis d’utiliser ces mêmes conditions de niveaux 3 de juillet 2005 jusqu’à la fin du *Run IIa*, malgré une nouvelle augmentation de la luminosité instantanée due aux bonnes performances de l’accélérateur en fin d’année 2005 comme il sera détaillé dans le paragraphe suivant.

Menus	JT1_ACO_MHT_HT	JT2_MHT25_HT
<i>V13</i>	3.1 Hz	2.8 Hz
<i>V13</i>	<i>3.3 ± 0.2 Hz</i>	<i>2.9 ± 0.2 Hz</i>
<i>V14</i>	2.1 Hz	1.6 Hz
<i>V14</i>	<i>1.6 ± 0.2 Hz</i>	<i>1.6 ± 0.2 Hz</i>

TAB. 3.5 – Taux d’acquisitions estimés à une luminosité instantanée de $100 \times 10^{30} \text{ cm}^{-2} \text{ s}^{-1}$ avec les taux obtenus par extrapolation linéaire pour les lignes 2 et 4, et en italique, les estimations de *trigger_rate_tool*, pour les lignes 3 et 5.

Conception d'un nouveau niveau 1 pour pallier l'augmentation brutale de la luminosité

A la fin de l'année 2005, plusieurs périodes de prise de données (en anglais *stores*) ont commencé avec une luminosité instantanée proche de $150 \times 10^{30} \text{cm}^{-2} \text{s}^{-1}$ alors que la liste *V14* avait été conçue initialement pour des luminosités maximales de $120 \times 10^{30} \text{cm}^{-2} \text{s}^{-1}$. Ces excellentes performances du *TeVatron* ont donc causé plusieurs temps morts au niveau 1 (en anglais *Front End Busy*), perturbant la prise de données et entraînant une perte approximative de 5 % de données au début de ces périodes. Des changements rapides des niveaux 1 des différentes conditions de déclenchement ont donc été requis. Ces changements se font toujours dans le même esprit. Il faut être rapide pour éviter une perte importante de données, mais maintenir les efficacités des signaux recherchés à celles initiales ou du moins limiter une chute éventuellement trop importante de ces efficacités. De même que pour l'étude précédente, il est préférable de comparer l'évolution des efficacités relatives plutôt que celle des efficacités absolues. Les responsables du système de déclenchement ont finalement demandé une modification du niveau 1 des conditions de déclenchement spécifiques aux jets (exceptés ceux propres aux études du bruit instrumental *QCD*). Ce niveau 1 concerne en effet une longue liste de conditions de déclenchement et à un taux d'acquisition relativement élevé. La diminution de ce taux nécessaire pour continuer à prendre des données dans de bonnes conditions a été estimée autour de 20 %. Pour concevoir un nouveau niveau 1, il faut prendre en compte plusieurs contraintes liées au fait que ce niveau traite l'information des sous-détecteurs au niveau matériel, *hardware* en anglais, par opposition au niveau logiciel, *software* en anglais (voir la Section 3.2.1). Le nombre de termes maximum utilisables au niveau 1 est de 128. Autrement dit, l'information du niveau 1 est une série de 128 bits (0 ou 1) à partir desquels on forme des combinaisons logiques, c'est-à-dire les termes de déclenchement. Ce nombre de termes (ou *refset* en anglais) est fixé à quatre pour le déclenchement sur les jets :

1. *jt3* : nombre de tours vérifiant $E_T > 3\text{GeV}$
2. *jt4eta26* : nombre de tours vérifiant $|\eta| < 2.6$ et $E_T > 4\text{GeV}$
3. *jt5* : nombre de tours vérifiant $E_T > 5\text{GeV}$
4. *jt7eta18* : nombre de tours vérifiant $|\eta| < 1.8$ et $E_T > 7\text{GeV}$

Par exemple, le terme de déclenchement conçu initialement pour la liste *V14*, *CJT(3,5,3.2)CJT(3,4,2.6)*, utilise 3 fois le deuxième terme et 3 fois le troisième. Ainsi, il faut utiliser une combinaison de ces quatre termes ou modifier un de ces termes, en vérifiant au préalable qu'il n'est pas utilisé par une autre condition de déclenchement. Le deuxième terme, *jt4eta26*, n'étant justement pas utilisé par une autre condition, il est possible de le modifier. Dans les combinaisons testées pour trouver un niveau 1 plus sélectif, on a donc essayé de remplacer ce deuxième terme par l'introduction éventuelle des termes *jt4eta16* et *jt4eta12*. Les signaux utilisés pour cette étude et pour le choix final d'un nouveau niveau 1 sont ceux correspondant à la production $HZ \rightarrow b\bar{b}\nu$ et à la production de paires de sbottoms. En effet, ces deux signaux ont dans l'état final des jets de plus basse énergie transverse que les autres signaux et déclenchent donc plus difficilement le niveau 1. Ensuite, les résultats obtenus ont été vérifiés avec les signaux correspondant à la production de gluinos et à la production de squarks. De la même façon que pour l'étude précédente, les taux d'acquisition sont estimés avec *trigger_rate_tool*, alors que les efficacités des signaux sont calculés après simulation du système de déclenchement grâce à *d0trigsim* et après application de coupures de présélection. Les résultats obtenus sont résumés dans les Tableaux 3.6 et 3.7. La première ligne correspond au niveau 1 utilisé par la liste *V13*,

la seconde correspond à celui initialement choisi pour la liste *V14* et les autres lignes sont les propositions faites pour la version *V14.80* de la liste *V14*. Les taux donnés sont ceux correspondant à une luminosité instantanée de $100 \times 10^{30} \text{cm}^{-2} \text{s}^{-1}$.

Niveaux 1 proposés	HZ in $\nu\nu b\bar{b}$		Sbottoms	
	Efficacités relatives (L1)	Efficacités relatives (L3)	Efficacités relatives (L1)	Efficacités relatives (L3)
CJT(3,5,3.2)	91.1±0.3	89.1±0.3	92.4±0.8	90.6±0.8
CJT(3,5,3.2)CJT(3,4,2.6)	91.0±0.3	89.1±0.3	92.4±0.8	90.6±0.8
CJT(3,5,3.2)CJT(2,4,1.6)	89.9±0.3	88.1±0.3	92.2±0.8	90.4±0.8
CJT(3,5,3.2)CJT(3,4,1.6)	87.4±0.3	85.5±0.3	91.0±0.8	89.2±0.8
CJT(3,5,3.2)CJT(2,4,1.2)	86.8±0.3	84.9±0.3	90.0±0.8	88.3±0.8
CJT(3,5,3.2)CJT(3,4,2.6)CJT(1,7,1.8)	90.3±0.3	88.4±0.3	92.1±0.8	90.3±0.8
CJT(3,5,3.2)CJT(2,4,1.6)CJT(1,7,1.8)	89.5±0.3	87.7±0.3	91.9±0.8	90.0±0.8
CJT(3,5,3.2)CJT(3,4,1.6)CJT(1,7,1.8)	87.0±0.3	83.5±0.3	90.7±0.8	88.9±0.8

TAB. 3.6 – Différentes propositions de conditions au niveau 1 pour la liste *V14.80* et efficacités relatives correspondantes.

Niveaux 1 proposés	Taux d'acquisition (L1) (Hz)	Evolution relative du taux (L1)
CJT(3,5,3.2)	448 ± 10	+24 %
CJT(3,5,3.2)CJT(3,4,2.6)	386 ± 9	0
CJT(3,5,3.2)CJT(2,4,1.6)	345 ± 9	-11 %
CJT(3,5,3.2)CJT(3,4,1.6)	249 ± 8	-35 %
CJT(3,5,3.2)CJT(2,4,1.2)	305 ± 8	-21 %
CJT(3,5,3.2)CJT(3,4,2.6)CJT(1,7,1.8)	290 ± 8	-25 %
CJT(3,5,3.2)CJT(2,4,1.6)CJT(1,7,1.8)	283 ± 8	-27 %
CJT(3,5,3.2)CJT(3,4,1.6)CJT(1,7,1.8)	218 ± 7	-44 %

TAB. 3.7 – Différentes propositions de conditions au niveau 1 pour la liste *V14.80* et taux d'acquisition correspondants.

Le meilleur choix, c'est-à-dire celui qui limite les baisses relatives d'efficacités tout en permettant une diminution prévisionnelle de plus de 20 % pour le taux de niveau 1, est la condition *CJT(3,5,3.2)CJT(3,4,2.6)CJT(1,7,1.8)* (en gras dans les Tableaux 3.6 et 3.7). De plus, cette condition a l'avantage supplémentaire de ne pas nécessiter la modification d'un des quatre termes. Une étude du groupe de recherche du boson de Higgs a confirmé ce résultat et les chiffres obtenus pour d'autres types de signaux. Cette nouvelle condition au niveau 1 a été acceptée et implémentée dans la liste *V14.80*. Cette liste a commencé à enregistrer des données le 01/11/2005. Les taux d'acquisition du niveau 1 prévus par *trigger_rate_tool* ont ensuite été vérifiés grâce aux données prises à partir du 01/11/2005 par la même méthode que celle employée pour la Figure 3.6. Une extrapolation linéaire permet d'estimer ce taux à $100 \times 10^{30} \text{cm}^{-2} \text{s}^{-1}$ et de le comparer

aux prévisions du Tableau 3.7. La Figure 3.8 montre l'évolution du taux d'acquisition de niveau 1 en fonction de la luminosité moyenne à partir de données enregistrées avec la liste *V13* (en rouge), la liste *V14* (en bleu) initiale et la nouvelle liste (en vert). Le Tableau 3.8 récapitule les résultats obtenus avec *trigger_rate_tool* et les taux effectivement enregistrés ensuite. Ces derniers confirment la baisse de plus de 20 % du taux d'acquisition de niveau 1.

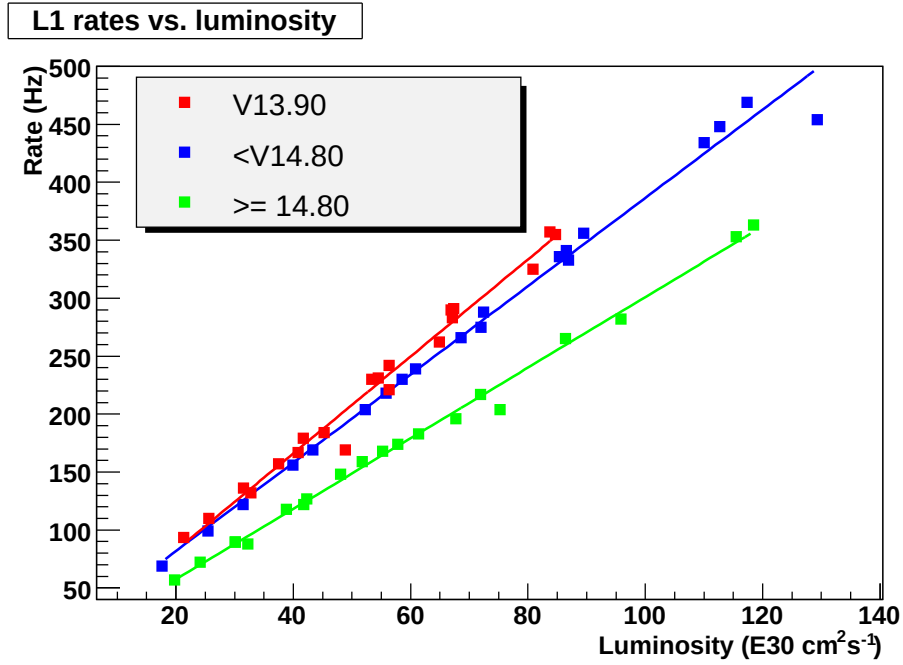


FIG. 3.8 – Taux d'acquisition du niveau 1 en fonction de la luminosité moyenne pour différentes listes de conditions de déclenchement.

	V13.90	< V14.80	≥ V14.80
Résultats de l'extrapolation linéaire	417 ± 2 Hz	387 ± 2 Hz	302 ± 2 Hz
Résultats de <code>trigger_rate_tool</code>	448 ± 10 Hz	386 ± 9 Hz	290 ± 8 Hz

TAB. 3.8 – Taux d'acquisition à une luminosité instantanée de $100 \times 10^{30} \text{ cm}^{-2} \text{ s}^{-1}$

Finalement, cette modification du niveau 1 a permis de prendre des données à des luminosités parfois supérieures à $150 \times 10^{30} \text{ cm}^{-2} \text{ s}^{-1}$ jusqu'à la fin du *Run IIa*. Le défi suivant est la conception d'une nouvelle liste permettant de prendre des données jusqu'à $300 \times 10^{30} \text{ cm}^{-2} \text{ s}^{-1}$, c'est-à-dire la valeur de la luminosité de référence pour le *Run IIb*. Pour réussir à maintenir les efficacités à leur valeur du *Run IIa* tout en prenant des données à cette luminosité, plusieurs nouveaux outils et différentes modifications du système de déclenchement ont été prévus. Celles concernant la partie calorimétrique du système de déclenchement sont décrites dans la Section 3.3.

3.3 Améliorations apportées au *Run IIb*

Pour augmenter le potentiel de découverte du TeVatron, il est important d'enregistrer une statistique de plus en plus importante de données. Les améliorations apportées à la chaîne d'accélération [86] ont pour but d'augmenter la luminosité intégrée tout en gardant la même énergie dans le centre de masse. Ainsi, cette luminosité permettra de repousser au maximum les limites de la physique des hautes énergies avant le démarrage du *LHC*, c'est-à-dire soit découvrir de nouvelles particules, tel le boson de Higgs (voir la Figure 3.9), soit préparer au mieux les recherches du futur collisionneur du *CERN*, en limitant l'espace des paramètres supersymétriques par exemple. Le *Run IIb* devrait permettre d'enregistrer une luminosité intégrée de 4 à 8 fb^{-1} d'ici 2009, alors que le *Run IIa* a collecté environ 1 fb^{-1} . La Figure 3.10 montre l'évolution des pics de luminosité instantanée au cours du *Run II* et la luminosité intégrée correspondante. On constate que depuis le début du *Run IIb* les pics de luminosité instantanée atteignent des records jamais obtenus lors du *Run IIa* et la luminosité intégrée augmente donc d'autant plus vite.

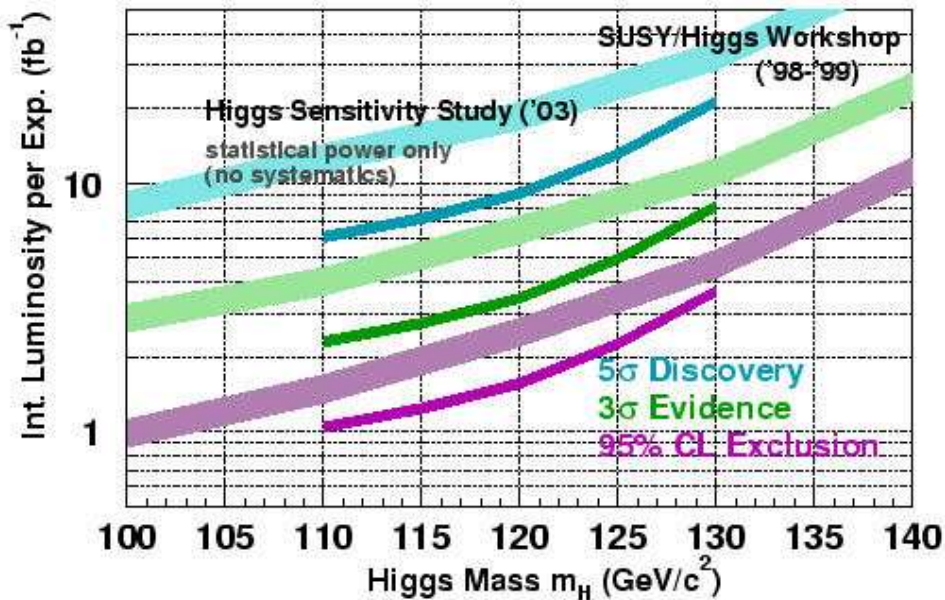
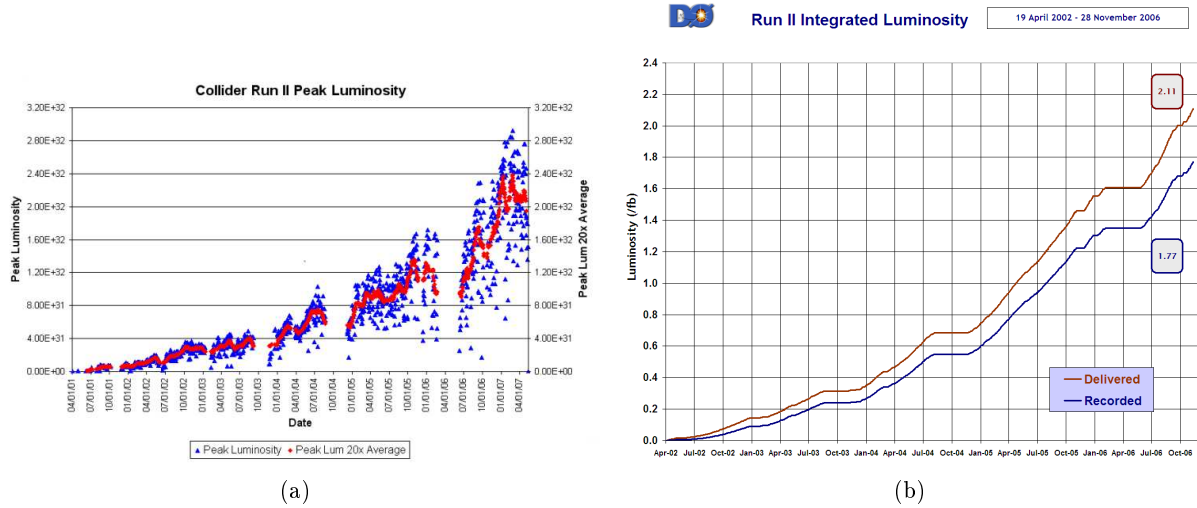


FIG. 3.9 – Prédiction du potentiel de découverte du boson de Higgs au TeVatron en fonction de sa masse et la luminosité intégrée.

L'augmentation de la luminosité instantanée et du nombre d'interactions de biais minimum (interactions dues à des collisions supplémentaires lors d'un même croisement de faisceau) entraînent par conséquent d'importantes modifications pour le système de déclenchement [87, 88], si on veut être en mesure de maintenir de hautes efficacités pour les signaux recherchés. L'architecture de base à trois niveaux et les taux d'acquisition maximum ne changent pas du *Run IIa* au *Run IIb*. Le temps entre deux croisements de faisceaux prévus initialement à 132 ns reste finalement à sa valeur initiale de 396 ns. Les contraintes de temps ainsi que la volonté de garder au maximum l'architecture du système de *Run IIa* ont limité les améliorations à cinq sous-systèmes :

- Le niveau 1 du système de déclenchement de la partie calorimétrique (*L1CAL*),

FIG. 3.10 – Evolution de la luminosité instantanée (a) et intégrée (b) lors du *Run II*.

qui sera décrit dans le paragraphe suivant, a été complètement modifié avec l'introduction de nouvelles cartes d'acquisition et de nouveaux algorithmes.

- **Le niveau 1 déclenchant sur les traces centrales (*L1CTT*)** a vu son algorithme de reconstruction des traces modifié. L'idée de base de cette modification est illustrée sur la Figure 3.12. Plutôt que d'utiliser systématiquement des doublons de fibres scintillantes du *CFT* (voir la Partie 2.3.2) pour reconstruire une trace, chaque couche de fibres est vue comme une entité (*Singlet equations* en anglais). Des simulations ont estimé que le facteur de rejet des traces identifiées à tort est approximativement de 10. Par contre, ces nouvelles équations demandent beaucoup plus de ressources, l'électronique de lecture associée aux fibres scintillantes a donc été en partie remplacée.
- **L'association de l'information sur les traces du *L1CTT* avec l'information sur les objets électromagnétiques et hadroniques du *L1CAL* (*L1CALTRACK*)** est une nouvelle fonctionnalité ajoutée au *Run IIb* pour le niveau 1. Conçue de la même façon que le système de déclenchement sur les muons (*L1MUO*), le *L1CALTRACK* permet d'associer notamment la position azimutale des traces reconstruites avec la position des objets reconstruits par le *L1CAL*. Cette association permet de réduire fortement les taux d'acquisition de niveau 1 et procure un avantage important par rapport aux expériences *CMS* et *ATLAS*, qui ne disposent pas de cette fonctionnalité (voir la Figure 3.11). Les détecteurs de pied de gerbes *CPS* et *FPS* (voir la Section 2.3.2) sont également utilisés.
- **Les cartes électroniques permettant de gérer les informations et algorithmes du niveau 2 (*L2β Processors*)** ont été partiellement améliorées pour permettre l'intégration d'algorithmes plus complexes.
- **Le système de déclenchement de niveau 2 permettant l'association des traces reconstruites au niveau 1 par le *CFT* avec les informations provenant du détecteur *SMT* (*L2STT*)** utilise désormais la couche supplémentaire de silicium ajoutée au plus proche du détecteur (*Layer 0*, [89]).

Toutes ces modifications permettent d'avoir un système de déclenchement efficace en mesure

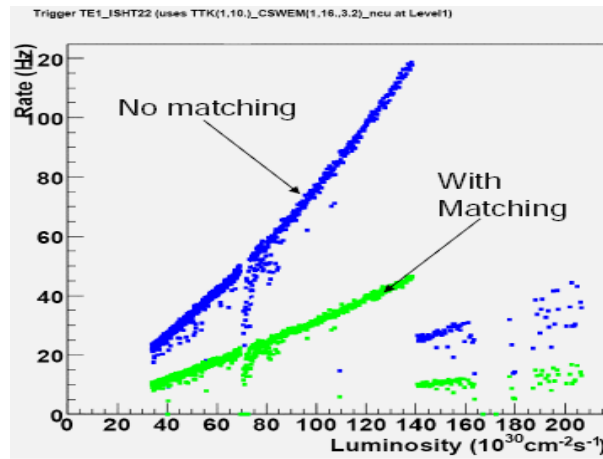


FIG. 3.11 – Taux d’acquisition au niveau 1 en fonction de la luminosité instantanée pour un déclenchement sur les électrons de 16 GeV ($CSWEM(16)$) avec $L1CALTRACK$ en vert et sans en bleu. Le gain estimé est de l’ordre de 2.5. A partir d’une luminosité instantanée de $140 \times 10^{30} \text{ cm}^{-2} \text{ s}^{-1}$, on ajoute un *prescale*.

de fonctionner avec une luminosité instantanée deux fois plus importante que celle enregistrée lors du *Run IIa*. Des algorithmes plus complexes et le remplacement de certaines parties du système d’acquisition nécessitent cependant un travail préliminaire important afin de bénéficier au mieux des nouvelles fonctionnalités dès le départ du *Run IIb*. Ces travaux préliminaires sont détaillés dans les paragraphes suivants pour la partie calorimétrique, qui en effet a subi des changements très importants.

3.3.1 Description détaillée des modifications apportées pour la partie calorimétrique

La partie calorimétrique ($L1CAL$) du *Run IIa* a été décrite dans la Section 3.2.1 et de plus amples détails peuvent être trouvés en [90]. Les principales modifications apportées au *Run IIb* sont le remplacement des cartes électroniques d’acquisition [91] et l’implémentation de nouveaux algorithmes de reconstruction d’objets électromagnétiques, hadroniques et de l’énergie transverse manquante [92]. On appellera par la suite ces objets : les électrons, les jets et l’énergie transverse manquante de niveau 1.

Modifications de l’électronique de lecture

L’électronique et l’architecture du système de déclenchement utilisé au *Run IIa* sont vieilles de 15 ans. Elles n’intègrent pas les progrès effectués au cours des dernières années dans le domaine de l’électronique. De plus, le système est devenu obsolète ce qui rend difficile la maintenance. Les progrès dans le domaine de l’électronique permettent d’obtenir un système plus rapide et plus précis dans l’estimation de l’énergie déposée dans les tours de déclenchement. Ils permettent également d’obtenir un système beaucoup plus compact. Enfin, la conception initiale des nouvelles cartes d’acquisition prévoyait une utilisation avec un croisement de faisceau toutes les 132 ns. Le fait de garder les 396 ns du *Run IIa* permet d’assurer une sécurité et une fiabilité renforcée. C’est

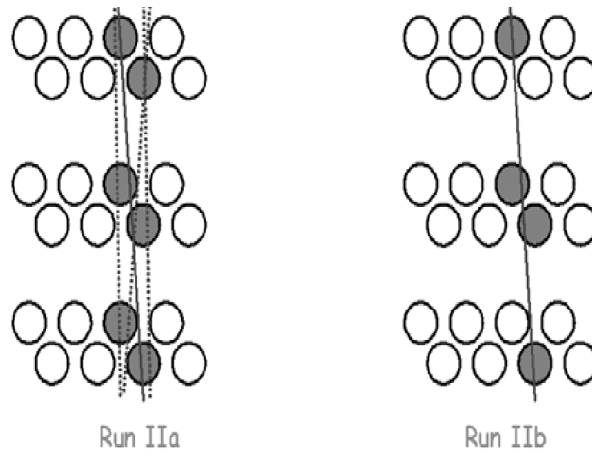


FIG. 3.12 – Schéma des algorithmes utilisés au niveau 1 par le système de déclenchement sur les traces *L1CTT* pour le *Run IIa* et le *Run IIb*.

pour toutes ces raisons qu'une partie du système électronique du *Run IIa* a été grandement améliorée.

Au *Run IIa*, les cartes, appelées *BLS* (décrites dans la Section 2.3.3), effectuent des sommes analogiques pour chacune des tours de déclenchement du calorimètre (2560 au total) et délivrent les 2560 signaux électroniques aux cartes électroniques appelées *CTFE* (pour *Calorimeter Trigger Front End*). Ces cartes digitalisent les signaux (voir la Figure 3.13) et calculent partiellement des sommes d'énergie et comptent le nombre de tours au-dessus de certains seuils en énergie. Ces résultats sont ensuite envoyés au système qui construit les termes de niveau 1 et prend en charge de garder ou non l'événement (ce système est appelé *Level 1 Trigger Framework*). Ce dernier ainsi que les cartes *BLS* restent inchangées au *Run IIb*. Par contre, toute la partie responsable de la digitalisation du signal et de l'attribution d'une énergie par tour et par croisement des faisceaux est modifiée. Les cartes *CTFE* sont donc remplacées par les cartes dites *ADF* (pour *Analog to Digital converter and Filter board*) [93, 94]. Le schéma de la Figure 3.14 récapitule toutes les étapes du détecteur à la prise de décision.

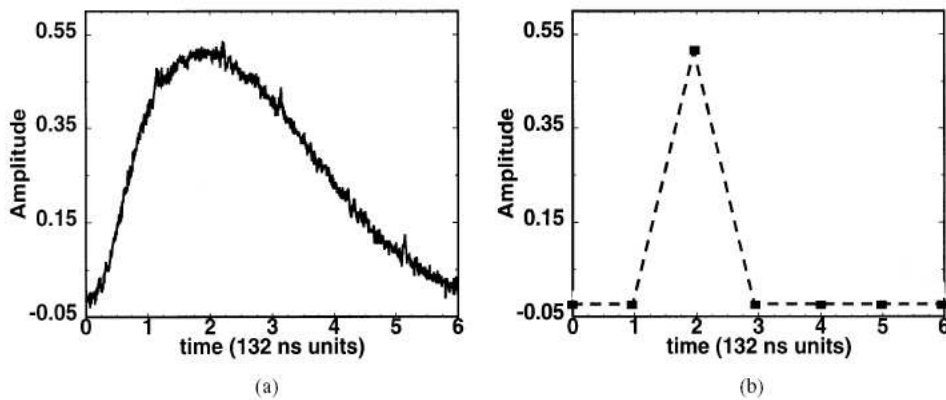


FIG. 3.13 – Signal analogique (a) et forme du même signal après digitalisation (b).

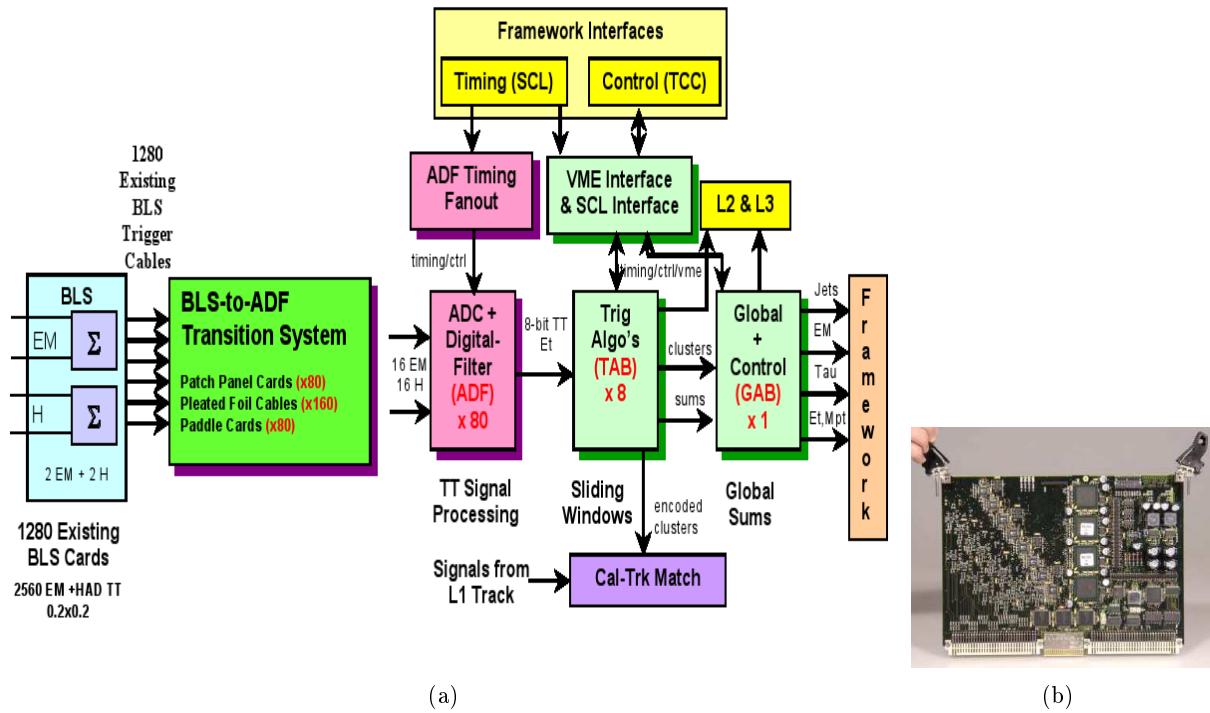


FIG. 3.14 – Flux des données pour la partie électronique du système de déclenchement (a) et photographie d'une carte *ADF* (b).

Ces cartes ont été conçues pour fonctionner avec un croisement de faisceaux toutes les 132 ns. Or, la Figure 3.13 montre qu'après 132 ns, l'énergie résiduelle dans une tour peut atteindre approximativement 75 % de l'énergie déposée lors du croisement précédent. Il faut donc être en mesure de soustraire cette énergie résiduelle pour affecter une énergie correcte à chaque tour. Cette énergie résiduelle ne posait pas réellement de problèmes au *Run IIa*. En effet, celui-ci fonctionnait avec un croisement toutes les 396 ns et comptait seulement le nombre de tours individuelles au-dessus d'un seuil en énergie. De même, au *Run IIb*, le temps entre deux croisements de faisceaux est finalement resté à 396 ns, néanmoins l'utilisation d'algorithmes plus complexes nécessite une précision plus importante sur la mesure de l'énergie et donc sur la soustraction de l'énergie résiduelle. La solution choisie pour pallier ce problème est de digitaliser ce signal en utilisant une fenêtre de temps adaptée mesurant l'énergie maximum du pic représenté sur la Figure 3.13. De nombreux tests ont été effectués afin d'optimiser le choix de cette fenêtre de temps. La modification de celle-ci peut d'ailleurs avoir des conséquences sur l'étalonnage de ce nouveau système. Finalement, les cartes *ADF* fournissent, aux algorithmes de niveau 1, l'énergie de chaque tour en coups *ADC* (pour *Analog-to-Digital Convertor*). Un coup *ADC* représente 0.25 GeV et est l'unité également utilisée au *Run IIa*.

Les nouveaux algorithmes de sélection implémentés

Le *Run IIb* utilise toujours les 2560 tours de déclenchement décrits en Section 3.2.3 pour trier les événements. Mais plutôt que de comptabiliser les tours au-dessus d'un certain seuil indépen-

damment les unes des autres, on les associe entre elles grâce à un algorithme dit de "fenêtres glissantes" (*Sliding Windows* en anglais) développé par la collaboration *ATLAS*. Cet algorithme se déploie aussi bien sur les tours *EM* pour reconstruire des électrons de niveau 1, que sur les tours *HD* pour reconstruire des jets de niveau 1. Seule la reconstruction des jets, utiles pour les conditions de déclenchement décrites dans cette partie, sera détaillée. Le principe de base de cet algorithme est la recherche de maxima locaux en déplaçant une grille composée d'un nombre fixé de tours sur tout le calorimètre. Plusieurs algorithmes ont été testés afin de fournir la meilleure précision sur l'énergie des jets en se rapprochant au plus proche des jets *offline*. La solution suivante a finalement été adoptée :

- On définit une région *R* composée de 2×2 tours *HD* (en bleu sur la Figure 3.15) ou la Figure 3.16). Cette région est utilisée pour trouver les maxima locaux. A cette région, on associe la tour en bas à gauche, appelée tour *T* et on affecte à cette tour la somme en énergie de la région *R*. L'ensemble des tours *T* constituent l'espace *R* (ou *R-space* sur la Figure 3.16).
- L'algorithme des fenêtres glissantes compare ensuite les tours *T* dans l'espace *R* et recherche les maxima dans un espace 5×5 tours selon les critères fixés par la Figure 3.15.
- Un jet de niveau 1 est finalement constitué de l'énergie déposée dans une région $\Delta\eta \times \Delta\phi = 0.8 \times 0.8$ (4×4 tours) autour de la région *R* (en orange sur les Figures 3.15 et 3.16).

La Figure 3.16 donne un exemple de l'application de cet algorithme dans le cas où deux jets de 5 et 9 GeV sont reconstruits.

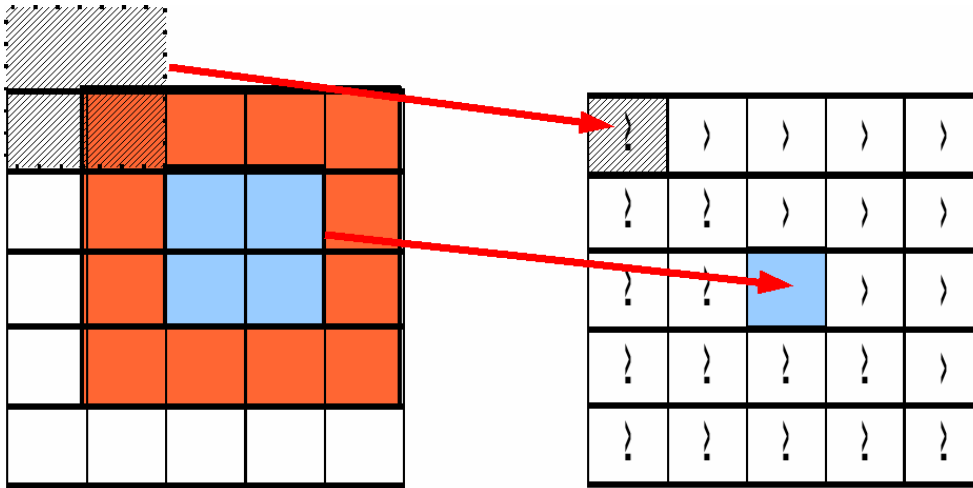
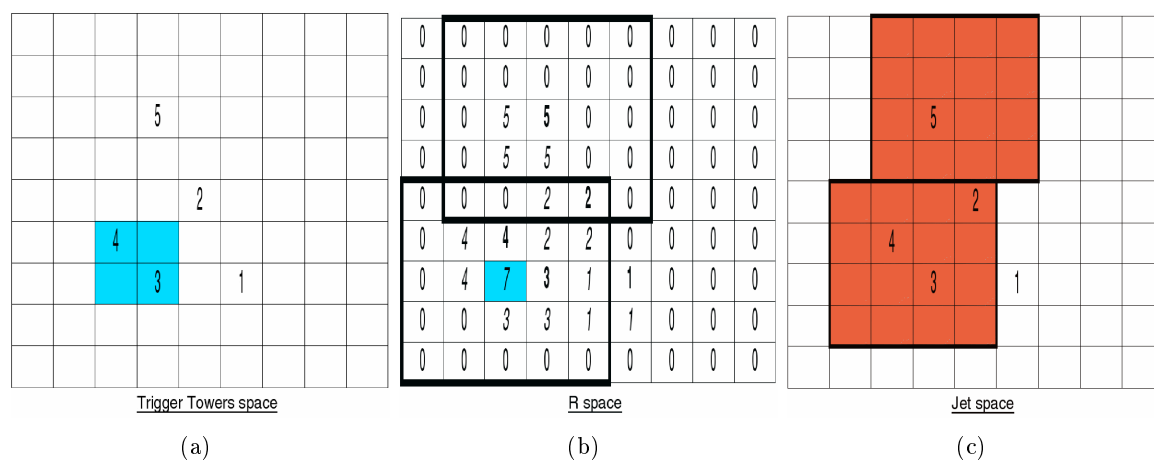


FIG. 3.15 – Schéma de l'algorithme de *Sliding Window*

Un autre algorithme très utile pour la conception de conditions de déclenchement sur les topologies à jets et $\cancel{E}_T$ est l'algorithme permettant d'estimer l'énergie transverse manquante au niveau 1. Cet algorithme est simplement la somme vectorielle de toutes les tours. Des études, qui seront détaillées par la suite, ont montré qu'il est plus efficace de reconstruire cette variable en utilisant seulement les tours au-dessus d'un seuil en énergie. Les conventions finalement utilisées pour ces deux nouvelles variables au niveau 1 sont les suivantes :

- CSWMET(*X*) : un événement passe si $\cancel{E}_T > X$ GeV.
- CSWJT(*n,X,Y*) : un événement déclenche ce terme s'il y a *n* jets de niveau 1 vérifiant $E_T > X$ GeV et $|\eta| < Y$.

FIG. 3.16 – Illustration de l’algorithme de *Sliding Window* pour un événement à deux jets

Ces deux termes vont permettre de réduire fortement le taux d’acquisition tout en maintenant une efficacité proche, voire meilleure, que celle obtenue lors du *Run IIa*. Le détail de ces résultats est décrit dans un paragraphe ultérieur.

3.3.2 Etalonnage du système de déclenchement du calorimètre au niveau 1

Les cartes électroniques *CTFE* ayant été remplacées par les cartes *ADF*, il faut donc étalonner ces dernières pour que les coups *ADC* mesurés correspondent effectivement à la valeur de 0.25 GeV. Seul le système d'acquisition du niveau 1 a été modifié, le système électronique *offline* permettant de mesurer les énergies recueillies par les cellules du calorimètre reste inchangé (voir la description dans la Section 2.3.3). L'étalonnage de ce dernier est largement documenté pour la partie électromagnétique en [95] et pour la partie hadronique en [96, 97]. Ces deux systèmes électroniques sont de plus totalement indépendants. L'étalonnage du système électronique du *L1CAL* s'effectue par conséquent par rapport au système *offline*, dit de précision. Ainsi, ces deux systèmes deviennent complètement consistants malgré les différences évoquées plus loin entre ces deux systèmes.

Ce travail a été effectué au printemps 2006. Il a débuté quelques semaines avant le début de la prise de données du *Run IIb* afin de mettre en place les outils et de les tester sur des données du *Run IIa*. Il s'est poursuivi quelques semaines après le début de la prise de données, temps nécessaire à la bonne compréhension des données et à la correction de nouveaux problèmes introduits par les modifications du système de déclenchement. Durant cette période, une procédure reproductible d'étalonnage du *L1CAL* a pu être mise en place et les outils appliquant cette procédure sont documentés en [98]. L'avantage de cette procédure par rapport aux étalonnages effectués dans le passé est qu'elle tient compte des non-uniformités du calorimètre. En effet, contrairement au *Run IIa* où les constantes d'étalonnage étaient calculées pour un anneau en ϕ , c'est-à-dire à η constant, cette procédure donne une constante par tour tant que la statistique le permet. De plus, des outils étant développés et documentés, il est aisé de reproduire rapidement cette procédure à n'importe quel moment, ce qui n'avait pas été fait au *Run IIa*. Les outils ne seront pas décrits ici, néanmoins la procédure ainsi que les résultats obtenus sont détaillés dans la suite.

Lots de données utilisés

Deux types d'échantillon de données sont utilisés dans la suite pour étalonner le nouveau système *L1CAL*. Un échantillon de données sans collision est utilisé pour mesurer le bruit dû essentiellement à l'électronique de lecture et ainsi fixer le piédestal de chacune des tours, qui correspond à l'énergie moyenne mesurée par la tour sans collision dans le détecteur en fonctionnement. Cet échantillon compte approximativement 135 000 événements et a été enregistré en juin 2006. Un deuxième échantillon qui correspond à plusieurs périodes de prise de données (*stores*) est utilisé pour ajuster la valeur donnée par le *L1CAL* avec celle obtenue *offline* et ainsi donner les constantes d'étalonnage par tour. Ce deuxième lot de données compte près de 4 650 000 événements passant les critères sur la qualité des données [99]. Le choix des données utilisées pour ce deuxième lot et la méthode pour calculer les constantes sont plus largement détaillés ensuite. Avant de pouvoir appliquer la procédure d'étalonnage, plusieurs problèmes ont été observés et rapidement réglés au démarrage du TeVatron. Un problème mineur dû à des mauvais branchements de fils causait l'absence de certaines tours de déclenchement. Il a suffi de repérer ces tours pour solutionner ce problème. De plus, l'utilisation de la couche A du système à muons a créé pendant plusieurs jours un bruit important pour la partie *CH* du calorimètre, essentiellement dans le bouchon au nord. Les conséquences de ce bruit sont visibles sur la Figure 3.17. Sur cette figure, les variables suivantes sont tracées :

- $METE$ représente l'énergie transverse manquante calculée à partir de la partie EM et FH du calorimètre.
- $METD$ est la somme de $METE$ et de la contribution de la région inter-cryostat (ICR).
- $METC$ correspond à la somme de $METD$ et de la contribution de la partie CH du calorimètre.

La modification de l'horloge des PDT (voir description du système à muons en Section 2.3.4) a finalement résolu ce problème, qui empêchait tout étalonnage propre de la partie $L1CAL$ sans risque de biais introduit par ce bruit.

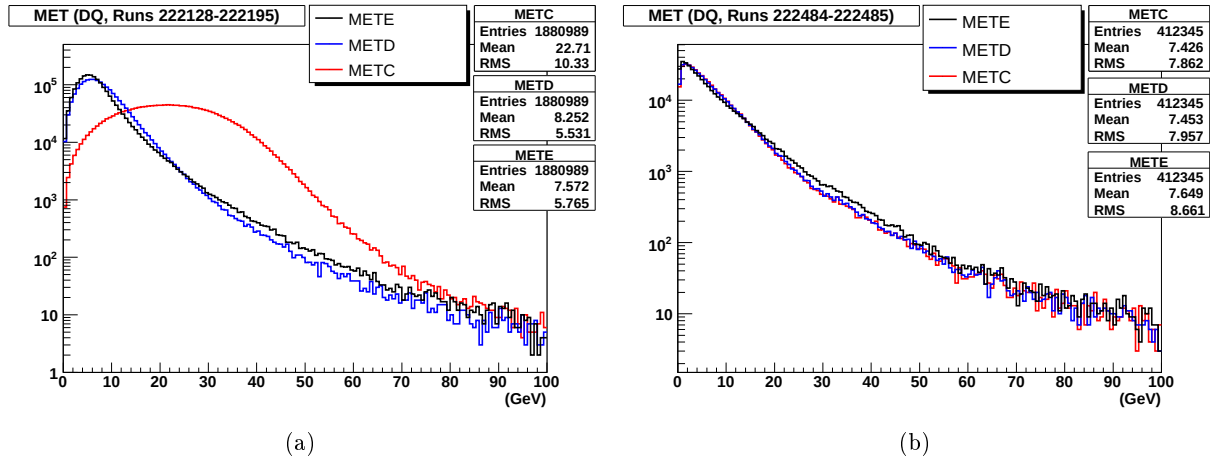


FIG. 3.17 – Distribution de $\cancel{E}_T$ (a) avec le bruit dû à l'utilisation de couche A du système à muons et (b) sans le bruit.

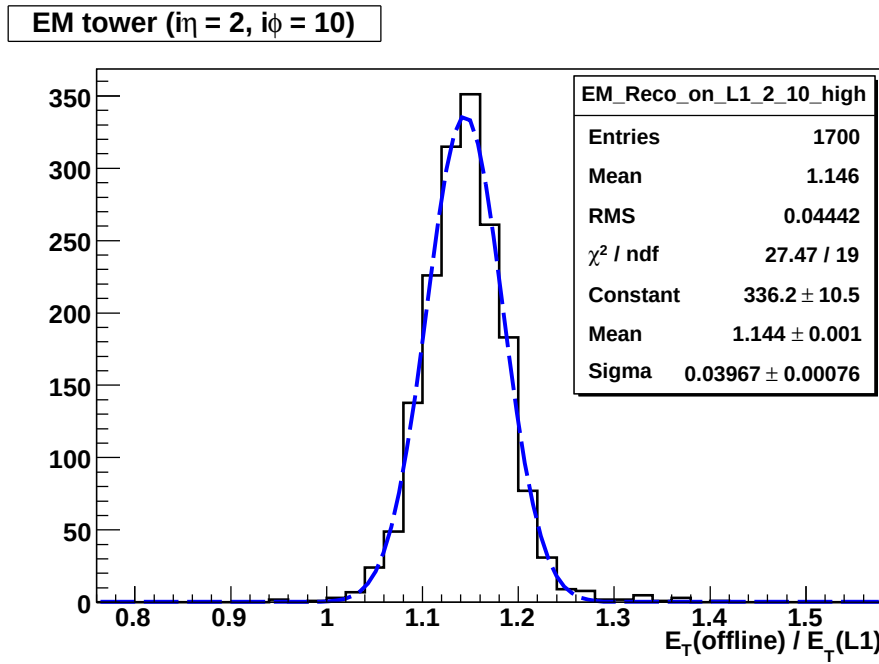
Choix de la méthode

Une fois ces problèmes réglés et l'enregistrement d'une quantité suffisante de données (approximativement 2 heures ou un *run*), la procédure d'étalonnage a pu être mise au point. Deux méthodes ont été testées pour calculer les constantes d'étalonnage :

- L'ajustement (*fit* en anglais) gaussien du rapport $R = \frac{E_T(offline)}{E_T(L1)}$, où $E_T(offline)$ est la somme projective de l'énergie des cellules du calorimètre mesurée par le système *offline* et correspondant à la tour de déclenchement dont l'énergie transverse est donnée par $E_T(L1)$ (voir la Figure 3.18).
- L'ajustement linéaire de $E_T(offline)$ exprimée en fonction de $E_T(L1)$.

La constante d'étalonnage est donnée par le centre de la gaussienne dans le premier cas et par la pente de la droite dans le second cas.

Le choix entre les deux méthodes a été fait de sorte à obtenir les incertitudes les plus faibles sur les constantes. La valeur de ces incertitudes en fonction de la coordonnée $i\eta$ est représentée sur les Figures 3.19 et 3.20. La coordonnée $i\eta$ correspond à un nombre entier tel que $i\eta = \eta \times 10$ excepté pour les tours se trouvant dans la région inter-cryostat (ICR et MG) qui sont aux coordonnées $|i\eta| = 19, 20$. De même, la coordonnée $i\phi$ correspond à une valeur entière vérifiant $i\phi = \phi \times 10$. Les plus larges incertitudes pour les coordonnées $|i\eta| = 7$ s'expliquent par l'absence des couches $EM1$ et $EM2$ du calorimètre électromagnétique. Et les incertitudes importantes observées pour $|i\eta| > 14$ sont dues à la faible statistique dans cette région du calorimètre, correspondant aux

FIG. 3.18 – Exemple du rapport R pour une tour EM .

petits angles proches du faisceau, et à la moins bonne précision des cellules se trouvant dans la région inter-cryostat.

Les Figures 3.19 et 3.20 montrent que la méthode utilisant l'ajustement gaussien est beaucoup plus précise. Les incertitudes sur la constante d'étalonnage sont effectives dans la plupart des cas inférieures à 1 %. De plus, l'incertitude minimum de la seconde méthode est imposée, d'une part, par les paramètres de la procédure d'étalonnage et, d'autre part, par la granularité du *L1CAL* (0.25 GeV). On peut estimer cette valeur minimale par la Formule 3.1 où Thr_{max} et Thr_{min} représentent les seuils maximum que l'on a imposés pour l'utilisation ou non d'une tour. Le choix de ces seuils sera décrit plus loin, sachant que $Thr_{max} < 60$ GeV, seuil de saturation d'une tour de déclenchement. De plus, S est la valeur de la pente et ΔS , l'incertitude sur cette pente.

$$\Delta S = S \times \left(1 - \left(\frac{Thr_{max} - Thr_{min} + 0.25}{Thr_{max} - Thr_{min} - 0.25}\right)\right) \quad (3.1)$$

Cette formule approximative est estimée en supposant que la pente est fixée par les deux points extrême $E_T = Thr_{min}$ et $E_T = Thr_{max}$. Les coordonnées de ces points sont donc :

- $(Thr_{min} \pm 0.125, S \times Thr_{min} \pm 0.125)$
- $(Thr_{max} \pm 0.125, S \times Thr_{max} \pm 0.125)$

L'incertitude irréductible sur la pente est finalement le maximum entre S et les pentes calculées à partir des points précédents. Dans le cas, où les deux points précédents ne sont pas parfaitement sur la droite, l'incertitude est encore plus importante. C'est pourquoi la Formule 3.1 est une bonne estimation de l'incertitude irréductible sur la pente. Cette incertitude ne peut donc pas être réduite en dessous de 1% et atteindre celle donnée par la méthode utilisant l'ajustement gaussien. C'est pour cette raison, que cette dernière méthode sera utilisée dans la suite. Par contre, il est important de garder à l'esprit que cette seconde méthode est un rapport avec $E_T(L1)$ comme dénominateur.

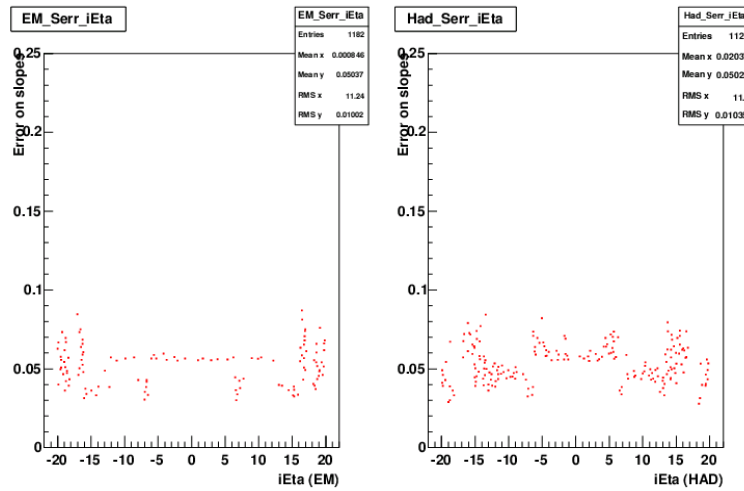


FIG. 3.19 – Incertitude sur la pente pour les tours *EM* (gauche) et pour les tours *HD* (droite).

La valeur de la constante dépend donc de celle du piedestal, ce qui implique rigoureusement que les constantes doivent être réévaluées après chaque modification des piedestaux.

La granularité de 0.25 GeV imposée par construction pour le *L1CAL* introduit un biais systématique sur le calcul de la constante d'étalonnage. Le meilleur moyen de réduire le biais introduit par cette granularité est l'application d'un seuil minimum sur les tours utilisées pour le calcul. Par la même occasion, ce seuil permet de réduire le biais introduit par un éventuel bruit dans la tour. Pour choisir le seuil à appliquer à une tour, deux effets sont à tenir en compte. D'une part, plus ce seuil est haut plus le biais est réduit, mais d'une autre part la statistique disponible pour calculer la constante se trouve amoindrie. Il faut donc trouver un compromis entre ces deux effets. Pour bien comprendre l'importance du biais en fonction du seuil fixé sur la tour de déclenchement, on trace la valeur du rapport R en fonction de $E_T(L1)$ (voir la Figure 3.21). Dans le cas de la tour *EM*, on constate que le ratio R diminue avec $E_T(L1)$. Ainsi, si on suppose que le seuil fixé sur cette tour est $E_T(L1)^{Thr}$ et sachant que $E_T(L1)$ décroît exponentiellement, cela implique que la constante est calculée pour des énergies proches de $E_T(L1)^{Thr}$. La conséquence est donc que si R augmente trop au-delà de $E_T(L1)^{Thr}$ alors la constante d'étalonnage sera sous-évaluée pour des énergies importantes. Le choix du seuil est donc capital dans le calcul de la constante.

On peut modéliser cette évolution par la Formule 3.2 où C et B sont des constantes. C est le maximum que peut atteindre la constante d'étalonnage. C'est également la constante théorique que l'on aurait si les valeurs des énergies du *L1CAL* formaient un spectre continu. B permet d'avoir une première information sur la valeur du piedestal de la tour. En effet, dans le cas de la tour *EM* de la Figure 3.21, la valeur de B est négative, ce qui implique que $E_T^{mes}(L1) > E_T^{real}(L1)$, et donc que le piedestal de cette tour est anormalement bas. Cette observation a été vérifiée par ailleurs en utilisant le lot de données sans collision ; le piedestal était en effet inférieur de 18 % à la valeur attendue.

$$R = \frac{E_T^{mes}(L1)}{E_T(of\ line)} = \frac{C}{1 + \frac{B}{E_T^{real}(L1)}} \quad (3.2)$$

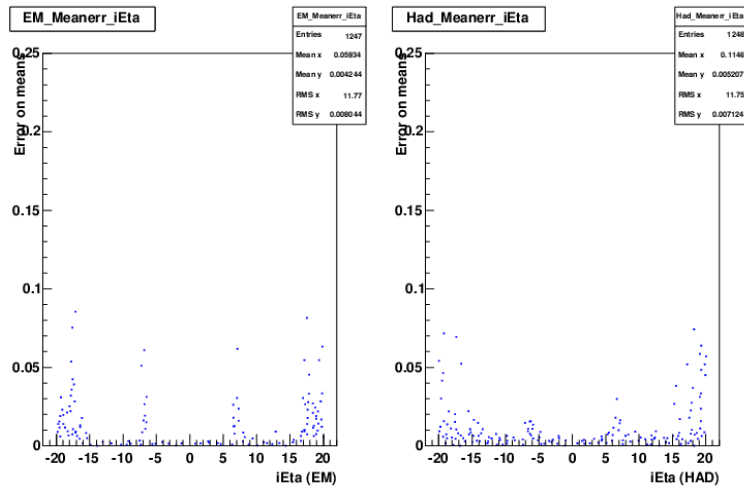


FIG. 3.20 – Incertitude sur le centre de la gaussienne pour les tours *EM* (gauche) et pour les tours *HD* (droite).

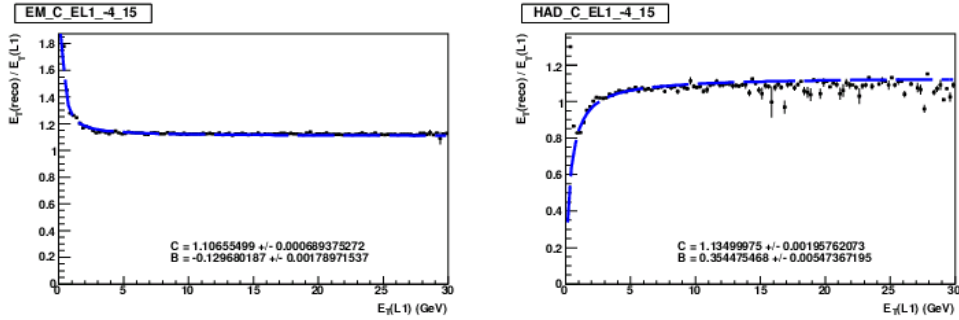


FIG. 3.21 – Exemples du rapport R en fonction de $E_T(L1)$ pour une tour *EM* (gauche) et la tour *HD* ayant les mêmes coordonnées (droite).

Un résumé des valeurs C obtenues en fonction de $i\eta$ est tracé sur les Figures 3.22. Ces valeurs C sont proches des constantes d'étalonnage que l'on obtiendra par la suite.

Finalement, on détermine le seuil optimal par valeur de $i\eta$ en ayant pris soin de regarder chaque tour une par une. Un seuil conservatif sur $E_T(\text{offline})$ a été ajouté pour éviter d'éventuels bruits enregistrés par le système électronique *offline* absente du côté du *L1CAL*. Ces seuils sont résumés dans le Tableau 3.9. On remarque que celui-ci diminue en fonction de $i\eta$ pour avoir suffisamment de statistique pour déterminer la constante. Un dernier paramètre est appliqué pour clore la description de cette méthode. En effet, un seuil maximal sur l'énergie des tours permet d'éviter des problèmes particuliers que l'on rencontre sur certaines tours dus à la géométrie du calorimètre. Par exemple, une tour *HD* correspondant à la valeur $|i\eta| = 5$ se situe dans une zone bien particulière du calorimètre. La moitié de cette tour est composée de cellules des couches *FH1* et *FH2*, l'autre moitié ne possède aucune cellule dans la couche *FH2*. Le poids attribué à chacune des couches étant le même dans tout le calorimètre central, il n'y a aucun moyen de compenser la différence entre un événement très énergétique déposant préférentiellement dans la première

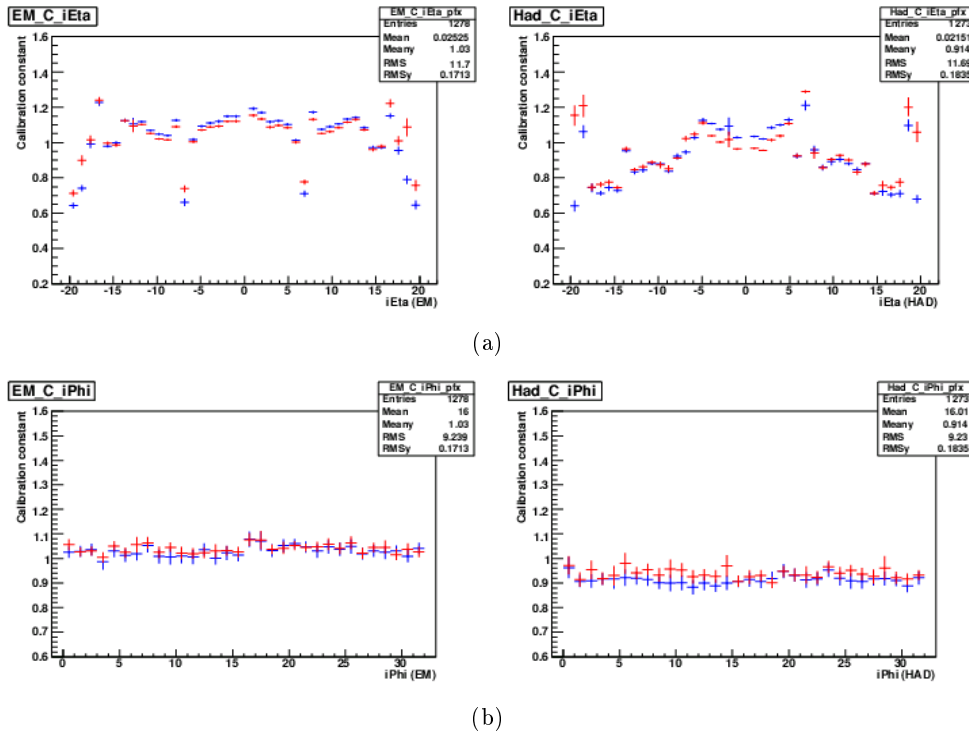


FIG. 3.22 – Paramètres C obtenus pour les tours EM (gauche) et pour les tours HD (droite) en fonction de $i\eta$ (a) et en fonction de $i\phi$ (b).

moitié d'un autre déposant dans la seconde moitié. Ainsi, le calcul de la constante fait apparaître deux pics (voir la Figure 3.23). L'ajout d'un seuil maximal en favorisant les événements déposant moins de 20 GeV dans une tour a été adopté, ce qui permet de faire disparaître le second pic (voir la Figure 3.23). Ce seuil de 20 GeV a été appliqué de façon systématique à toutes les tours.

Biais éventuellement introduits par le choix des conditions de déclenchement

Les biais dus aux bruits, à la discrétisation, aux particularités du calorimètre étant contrôlés, les incertitudes sur la détermination de la constante d'étalonnage étant réduites, il faut être en mesure d'estimer les biais introduits par le lot de données utilisé. Il faut en effet estimer si l'utilisation de conditions de déclenchement utilisant des termes basés sur le calorimètre introduisent un nouveau biais à la méthode utilisée. Pour cette étude, on utilise le lot de données précédent, ainsi qu'une sous-partie de celui-ci ne contenant que les événements ayant déclenchés des conditions de déclenchement dont les termes sont basés exclusivement sur les détecteurs de traces et le spectromètre à muons. Les constantes sont calculées pour la partie EM et FH de chacun des lots de données. On trace ensuite la distribution du rapport des constantes calculées avec toutes les conditions de déclenchement par les constantes calculées avec la sous-partie du lot de données. Les distributions de ces rapports sont tracées sur la Figure 3.24, la distribution de ces rapports est de plus paramétrée par un ajustement gaussien. Pour éviter de tenir compte des constantes n'ayant pas été estimées avec une précision raisonnable, on impose que l'incertitude sur ces dernières soit

$i \eta $	Seuils sur les tours <i>EM</i> (GeV)	Seuils sur les tours <i>HD</i> (GeV)
20	4	5
19	4	4
18	3	2
17	4	3
16	5	4
15	6	4
14	7	5
13	7	5
12	10	7
11	10	8
10	10	8
9	10	8
8	8	8
7	5	5
≤ 6	10	10

TAB. 3.9 – Seuils sur $E_T(L1)$.

inférieure à 1%. Pour la partie *EM*, on constate que le centre de la gaussienne est proche de 1 et la déviation par rapport à 1 est négligeable par rapport au biais que peut introduire la discrétisation vue précédemment. De plus la largeur de la gaussienne est consistante avec la valeur de 1 % imposée sur l'incertitude. Cette étude a été répétée avec une valeur maximale sur l'incertitude de 3 % et 5% pour obtenir les mêmes conclusions pour les tours *EM* et pour les tours *HD*, c'est-à-dire une distribution gaussienne centrée sur 1 et une largeur consistante avec l'incertitude.

La conclusion de cette étude est que l'utilisation de toutes les conditions de déclenchement n'introduit aucun biais sur la méthode mais a l'avantage de fournir une statistique beaucoup plus importante.

Finalement, la procédure finale adoptée est l'utilisation du rapport $R = \frac{E_T(\text{offline})}{E_T(L1)}$ avec un seuil sur $E_T(L1)$ dépendant de $i|\eta|$ et un seuil maximal de 20 GeV. L'utilisation de toutes les conditions de déclenchement est préférée pour cette procédure.

Estimation des piédestaux

Comme il a été décrit auparavant, avant de calculer les constantes d'étalonnage, il faut au préalable estimer la valeur du piédestal pour chacune des tours. Ces valeurs sont estimées en coups *ADC* et la valeur correspondant à 0 GeV est :

- 0 pour $i|\eta| = 19, 20$,
- 8 autrement.

Pour estimer ces piédestaux, il suffit de mesurer l'énergie mesurée par une tour en l'absence de collision et un ajustement gaussien de la distribution de l'énergie donne la valeur recherchée. Deux exemples sont illustrés sur les Figures 3.25 et 3.26.

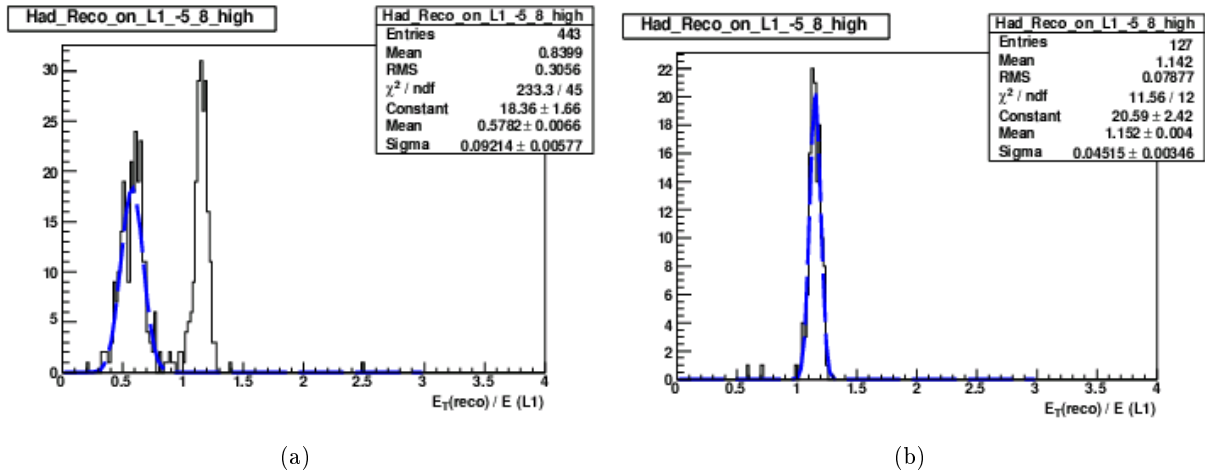


FIG. 3.23 – Valeur de R sans le seuil (a) et avec le seuil (b).

Deux paramètres sont importants pour cette mesure : le centre de la gaussienne fournit la valeur du piédestal et permet d'effectuer un réajustement par rapport à 8 (ou 0) si besoin et la largeur de la gaussienne donne un ordre de grandeur des bruits enregistrés par une tour. Ces bruits sont constants par rapport à la variable $i\phi$, par contre ils fluctuent largement par rapport à $i\eta$. On peut modéliser ces fluctuations en supposant qu'il y a deux types de bruit. Le premier a lieu avant la projection sur l'espace transverse et est donc dépendant de $i\eta$, le second a lieu après la projection mais avant la digitalisation et est donc indépendant de $i\eta$. Si on suppose que ces deux bruits sont indépendants, on peut les sommer en quadrature et obtenir une modélisation selon la Formule 3.3, où A et B sont des constantes tenant compte de l'importance relative des deux bruits et C est la constante d'étalonnage.

$$\sigma = \frac{\sqrt{(A \times \cos(\frac{\pi}{2} - \text{atan}(e^{-\eta})))^2 + B^2}}{C} \quad (3.3)$$

Les valeurs des largeurs obtenues (en noir) et la paramétrisation des bruits (en bleu) est représentée sur la Figure 3.27 en prenant comme valeurs :

- Pour les tours *EM* : A = 0.65 et B = 0.4
- Pour les tours *HD* : A = 0.65 and B = 0.3

On constate que pour la majorité des tours la paramétrisation des bruits est proche des mesures effectuées.

calcul des constantes d'étalonnage

Une fois, les piédestaux fixés à leur valeur de fonctionnement (8 ou 0), les constantes d'étalonnage peuvent être évaluées pour le nouveau système *L1CAL* du *Run IIb* à partir de la procédure largement décrite précédemment. Les constantes ont donc été calculées à partir du lot de données utilisé pour les études précédentes. Les constantes de 1236 tours *EM* sur 1280 et de 1231 tours *HD* sur 1280 ont pu être calculées avec une précision inférieure à 5 %. Pour les autres tours, la moyenne sur tout $i\eta$, à la manière du *Run IIa*, a été affectée pour la valeur de la constante. Les

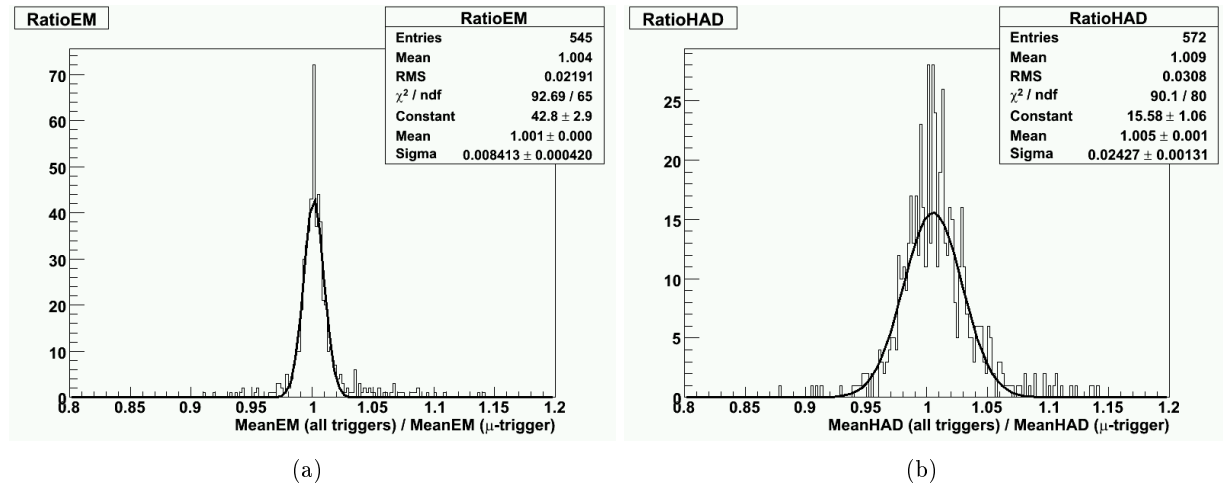


FIG. 3.24 – Comparaison entre les constantes calculées avec toutes les conditions de déclenchement et celles calculées avec les conditions déclenchant exclusivement sur les traces et les muons, pour la partie *EM* (a) et la partie *HD* (b).

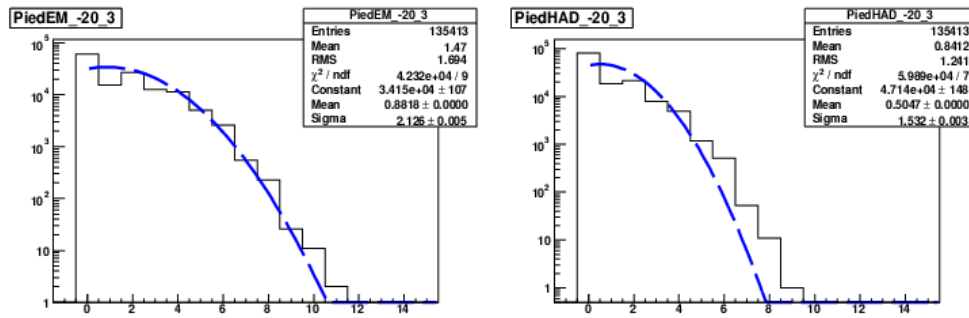


FIG. 3.25 – Distribution de l'énergie transverse pour une tour ayant pour coordonnées $i\eta = -20$ et $i\phi = 3$.

variations par rapport à 1 ont été propagées sur les cartes *ADF* afin d'avoir une énergie mesurée par le système *LICAL* consistante avec celle mesurée *offline*. A partir des données enregistrées après cette modification, la procédure a de nouveau été appliquée pour vérifier l'efficacité de la méthode. Les résultats pour la mesure des pedestaux et l'impact sur le taux d'acquisition au niveau 1 sont représentés sur la Figure 3.28. On constate que le réajustement des pedestaux réduit brutalement le taux d'acquisition.

Les constantes d'étalonnage estimées avant la procédure décrite préalablement et après cette procédure sont représentées en fonction de $i\eta$ sur la Figure 3.29. On constate que la dispersion est fortement réduite, la valeur du *RMSy* (*RMSy* désigne l'amplitude de la dispersion des valeurs selon l'axe y, *RMS* signifie *Root Mean Square* en anglais) est divisée par un facteur 10 approximativement, et que la valeur moyenne des constantes est centrée sur 1. La légère diminution observée dans le calorimètre central pour les tours *EM* est due au fait que pour cette itération les ajustements des pedestaux et des constantes d'étalonnage ont été effectués simultanément au lieu

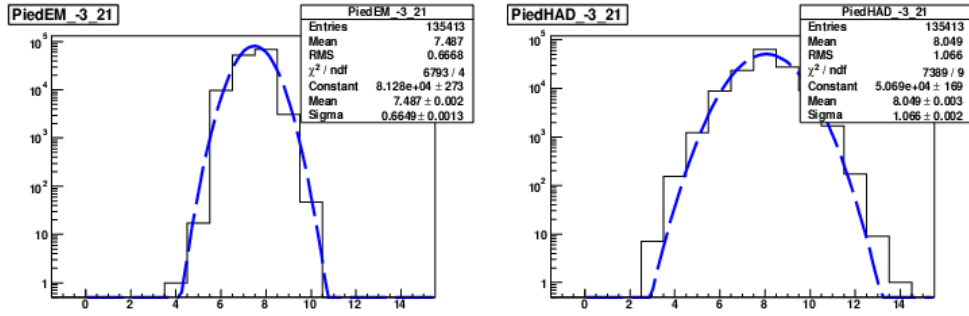


FIG. 3.26 – Distribution de l'énergie transverse pour une tour ayant pour coordonnées $i\eta = -3$ et $i\phi = 21$.

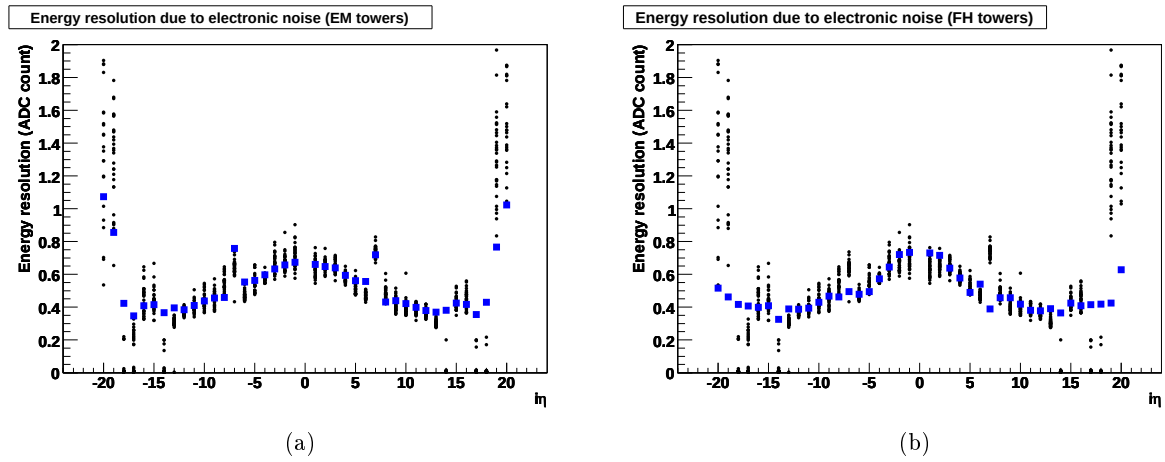


FIG. 3.27 – Largeur de la gaussienne en fonction de $i\eta$ pour toutes les tours *EM* (a) et pour toutes les tours *HD* (b).

de séquentiellement. La différence par rapport à 1 de moins de 2 % a été jugée raisonnable en l'attente d'une nouvelle itération.

Conséquences positives de l'étalonnage

Finalement, la procédure mise au point a permis d'étalonner le nouveau système *L1CAL* par rapport à l'information recueillie par l'électronique de précision. La première conséquence positive est la diminution des taux d'acquisition du niveau 1 sans aucune modification des conditions de déclenchement et ce dès le réajustement des piédestaux. La seconde conséquence positive peut être observée directement sur l'efficacité des nouveaux algorithmes mis en place dès le niveau 1 (voir [100, 101] pour plus de détails). En effet, la Figure 3.30 montre l'efficacité d'une coupure à 20 GeV sur l'énergie transverse manquante de niveau 1 par rapport à l'énergie transverse manquante *offline* (variable calculée avec un algorithme effectuant les mêmes opérations que celles effectuées au niveau 1 mais à partir de l'électronique de précision), alors que la Figure 3.31 montre l'efficacité d'une coupure à 20 GeV sur l'énergie d'un jet de niveau 1 (cette courbe a été obtenue par le même

procédé que la première). On constate que dans le premier cas que la point d'inflexion passe de 21.7 à 20.8 GeV et la dispersion autour de ce point est réduite de 4.9 à 3.8 GeV. De même, la seconde figure montre une dispersion réduite de 6.1 à 5.9 GeV. Les algorithmes deviennent donc plus efficaces après l'étalonnage, ce qui permet d'obtenir des efficacités de sélection équivalentes pour des coupures plus dures et donc des taux d'acquisition plus bas. Le choix de la valeur de ces coupures et des termes développés pour les conditions de déclenchement propres aux topologies à jets et $\cancel{E}_T$ est développé dans la Section 3.3.3.

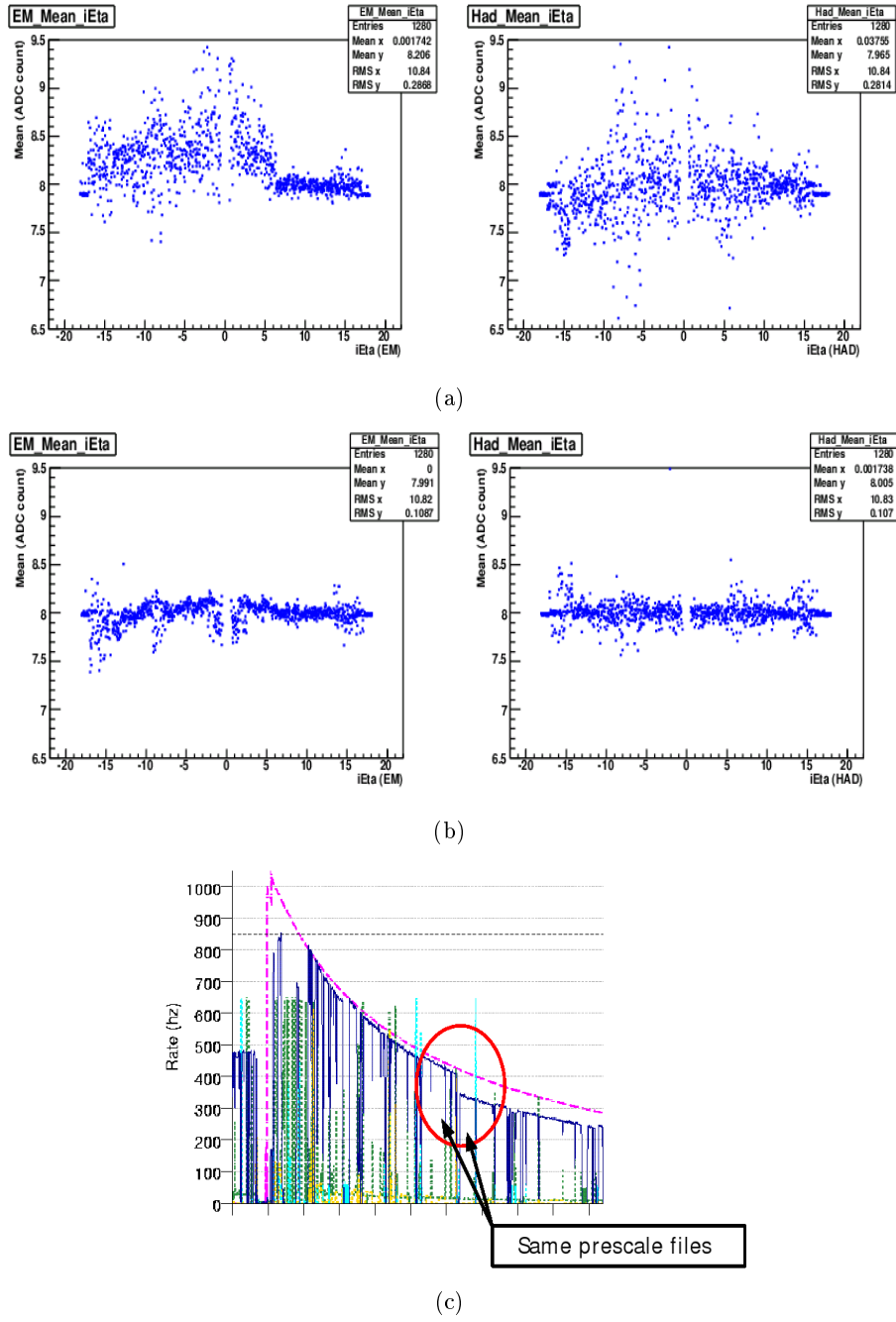
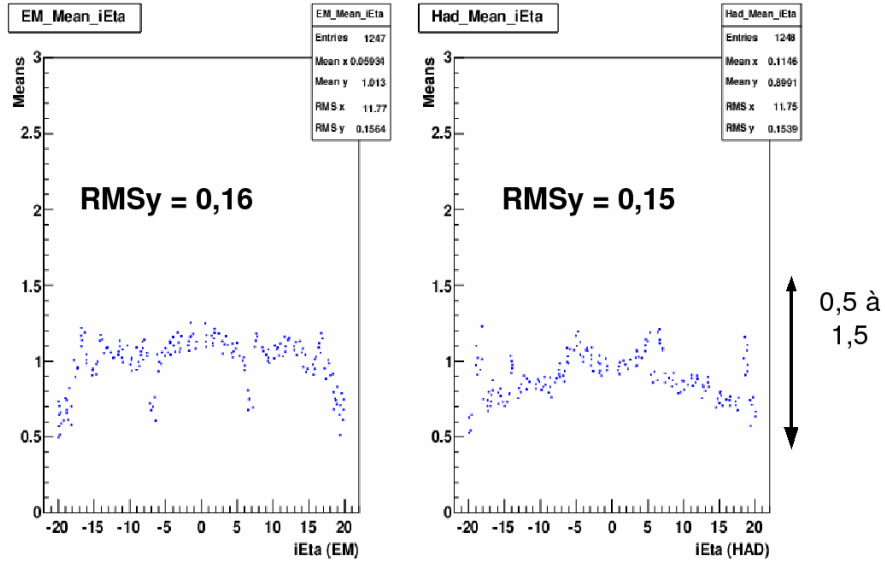
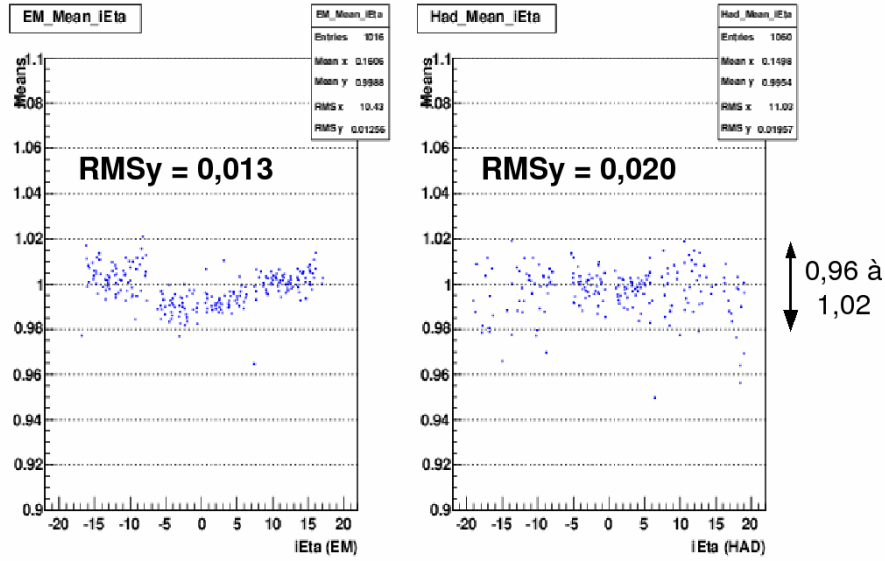


FIG. 3.28 – Mesure des piédestaux avant (a) et après réajustement (b) par rapport à 8 *ADC* (ou 0) et conséquence sur le taux d'acquisition au niveau 1 (décrochage observé sur la courbe bleue) (c). Les graphiques de gauche de (a) et (b) montre les piédestaux pour les tours *EM* alors que les graphiques de droite de (a) et (b) les tours *HD*.



(a)



(b)

FIG. 3.29 – Constantes d'étalonnage en fonction de $i\eta$ avant la procédure (a) et après la procédure (b) pour des tours EM à gauche et HD à droite.

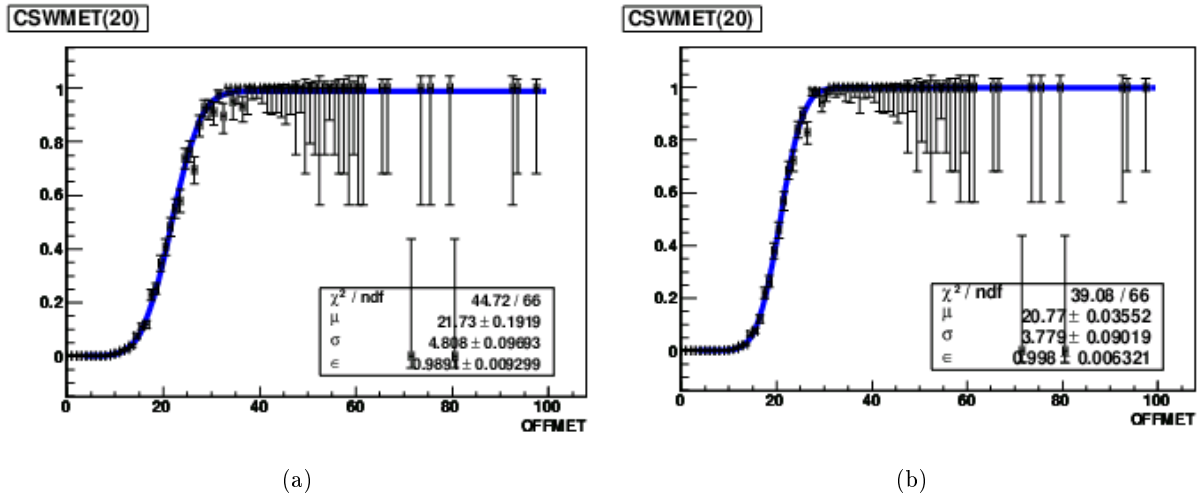


FIG. 3.30 – Courbes représentant l’efficacité de la coupure $CSWMET(20)$ en fonction de l’énergie transverse manquante calculée *offline* (a) avant l’étalonnage et (b) après l’étalonnage.

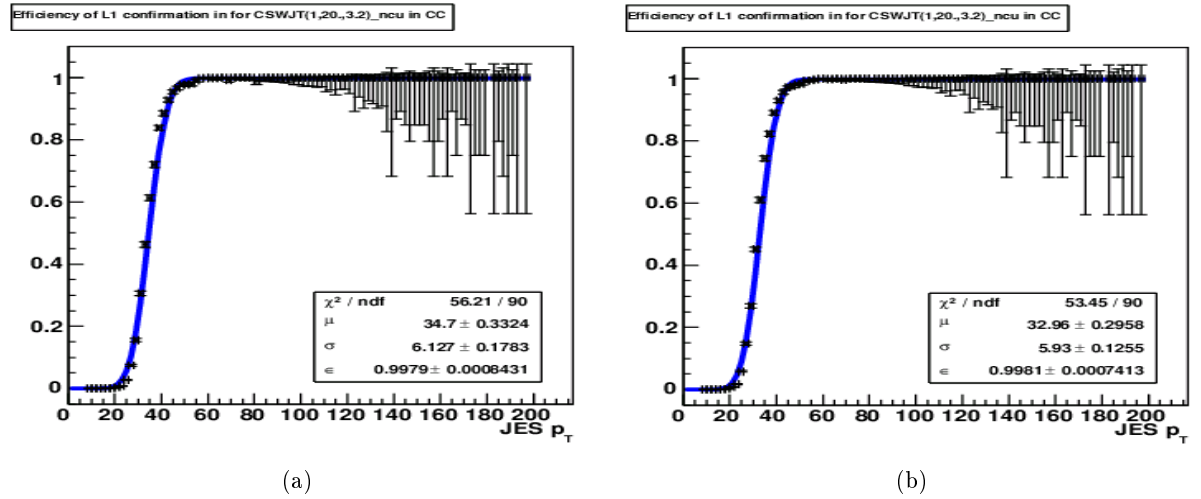


FIG. 3.31 – Courbes représentant l’efficacité de la coupure $CSWJT(1,20,3.2)$ en fonction de l’énergie transverse manquante calculée *offline* (a) avant l’étalonnage et (b) après l’étalonnage.

3.3.3 Conception des conditions de déclenchement spécifiques aux signaux à jets et E_T pour le *Run IIb* : la liste *V15*

La liste *V15.00* est le premier menu de déclenchement qui a pris des données au démarrage du *Run IIb*. Cette liste a été conçue pour pouvoir enregistrer des données jusqu'à des luminosités instantanées de $300 \times 10^{30} \text{cm}^{-2} \text{s}^{-1}$, c'est-à-dire des pics de luminosité deux fois plus hauts que ceux obtenus dans les meilleures conditions du *Run IIa*. Les nouveaux algorithmes utilisables pour cette liste ont été développés pour être les plus efficaces du point de vue de la physique recherchée tout en étant les plus discriminants pour le fond instrumental comme il a été détaillé auparavant. Pour les conditions de déclenchement propres aux signaux décrits en Section 3.2.3, seuls les algorithmes de niveau 1 détaillés dans la Section 3.3.1 sont utilisés. Dans un premier temps, c'est-à-dire plusieurs mois avant le début du *Run IIb*, les deux premiers niveaux du système de déclenchement ont été mis au point. Les nouveaux algorithmes, comme les algorithmes de "fenêtres glissantes" par exemple, n'étant pas encore en ligne, une émulation de ces derniers reproduisant exactement le même code a été utilisée pour simuler les termes de niveau 1. Dans un deuxième temps, c'est-à-dire un mois après le démarrage du *Run IIb* et une fois une quantité suffisante de données enregistrées, la conception du niveau 3 a commencé ainsi que la vérification du bon fonctionnement des deux premiers niveaux. N'ayant participé qu'à la première étape, je m'attacherai à décrire celle-ci en détails [80]. Je ne donnerai que les résultats et les conclusions obtenus pour la deuxième étape, c'est-à-dire la liste complète des termes de déclenchement pour les trois niveaux [102]. Au démarrage du *Run IIb*, les niveaux 3 des conditions de déclenchement propres aux topologies à jets et E_T étaient ceux des conditions de la liste *V14*.

Pour la conception des niveaux 1 et 2, les signaux et les outils décrits en Section 3.2.3 ont permis de mettre au point trois conditions de déclenchement, un de plus que pour le *Run IIa*. La première condition, dite condition *Monojet*, est dédiée aux signaux particuliers ayant au moins un jet de très haute énergie transverse (E_T) et une forte énergie transverse manquante contrebalançant ce jet. Cette condition est une nouveauté de la liste *V15*. La deuxième condition, dite condition *Dijet*, est la continuité de *JT1_ACO_MHT_HT*, alors que celle appelée condition *Multijet* est la continuité de *JT2_MHT25_HT*.

Conception du niveau 1

Historiquement, les variables et les coupures choisies au niveau 1 ont été développées avant l'optimisation de l'algorithme sur l'énergie transverse manquante de niveau 1. En effet, il a été décidé d'ajouter un seuil sur les tours de déclenchement utilisées dans le calcul de cette variable. Cette étude sur l'optimisation du choix du seuil sera décrite ultérieurement. La conséquence de cette optimisation tardive est que la conception du niveau 1 a été faite sans cette optimisation mais a été vérifiée par la suite. De plus, comme il sera détaillé plus loin, cette optimisation n'a pas de conséquence négative sur les efficacités du signal mais réduit par contre les taux d'acquisition.

Les variables utilisées pour la conception du niveau 1 sont les suivantes :

- $CSWJT(n, X, Y)$
- $CSWMET(X)$
- $CSWAKL(X, Y, n)$: ce terme a été développé pour rejeter le fond instrumental *QCD* par le groupe de recherche du boson de Higgs. Contrairement à tous les autres termes de déclenchement, c'est en partie un terme de veto. Ce terme considère toutes les paires de jets vérifiant $5 < E_T < X$ GeV pour le premier et $5 < E_T < Y$ GeV pour le second (le seuil

de 5 GeV a été poussé à 8 GeV ensuite). L'événement passe en effet par défaut s'il existe au moins une paire vérifiant ces critères, **sauf** si une de ces paires est dites dos à dos. Le nombre entier n représente le nombre de secteurs contigus qui définissent le terme dos à dos. Il y a en effet 32 valeurs possibles pour la coordonnées ϕ au niveau 1. Par exemple, sur la Figure 3.32, la valeur de n est 2.

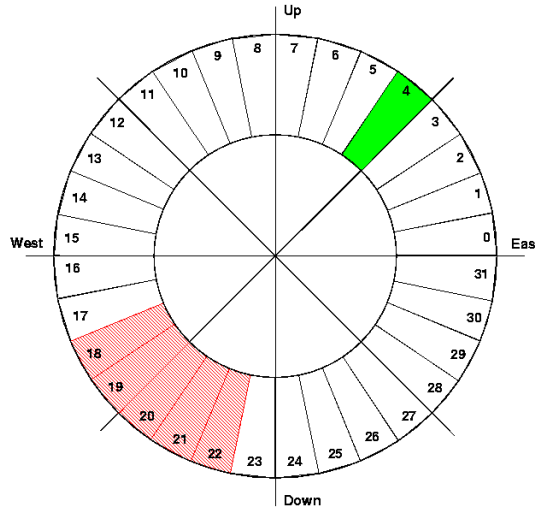


FIG. 3.32 – Les 32 secteurs en ϕ au niveau 1.

La Figure 3.33 montre les distributions de niveau 1 des signaux qui ont été utilisés pour la conception des deux premiers niveaux de ces conditions de déclenchement. C'est à partir de distributions de ce type que sont construites les termes de déclenchement de chacune des conditions. De plus, comme il a été évoqué précédemment, il est plus utile d'étudier l'efficacité relative à une présélection conservatrice que l'efficacité absolue d'un terme de déclenchement. C'est pour cette raison que l'on trace, pour toutes les variables pertinentes, l'évolution de ces deux efficacités en fonction de la valeur de la coupure appliquée sur cette variable. Ainsi, on est en mesure de déterminer une valeur minimale, qui est conservatrice en terme d'efficacité pour les signaux étudiés. Des exemples de ce procédé sont montrés sur la Figure 3.34.

De plus, les conditions de déclenchements propres aux topologies à jets et $\cancel{E}_T$ ont été étudiées par le groupe de recherche du boson de Higgs, et en particulier les équipes en charge du canal $HZ \rightarrow b\bar{b}\nu\nu$ et de la production Hbb avec $H \rightarrow b\bar{b}$. La première analyse, également étudiée ici, a un état final avec deux jets dos à dos, alors que la seconde a une topologie finale multijet comparable à la production de gluinos. Les niveaux 1 conçus pour ces deux analyses, appelés par commodité HZ et HBB , ont donc été testés sur les signaux évoqués en Section 3.2.3. En complément de ces deux conditions de déclenchement, une condition *Monojet*, appelée *MONO*, a été proposée. Finalement, le niveau 1 d'une des conditions de déclenchement proposé pour les études sur le fond QCD a également été étudié. Il sera simplement nommé QCD . Les conditions de chacun de ces niveaux 1 sont les suivantes :

- QCD : CSWJT(1,45,3.2)
- $MONO$: CSWJT(1,30,3.2)CSWMET(25)
- HZ : CSWAKL(20,20,0)CSWJT(1,20,2.4)CSWJT(2,8,3.2)CSWMET(25)
- HBB : CSWJT(1,30,2.4)CSWJT(2,15,2.4)CSWJT(3,8,3.2)

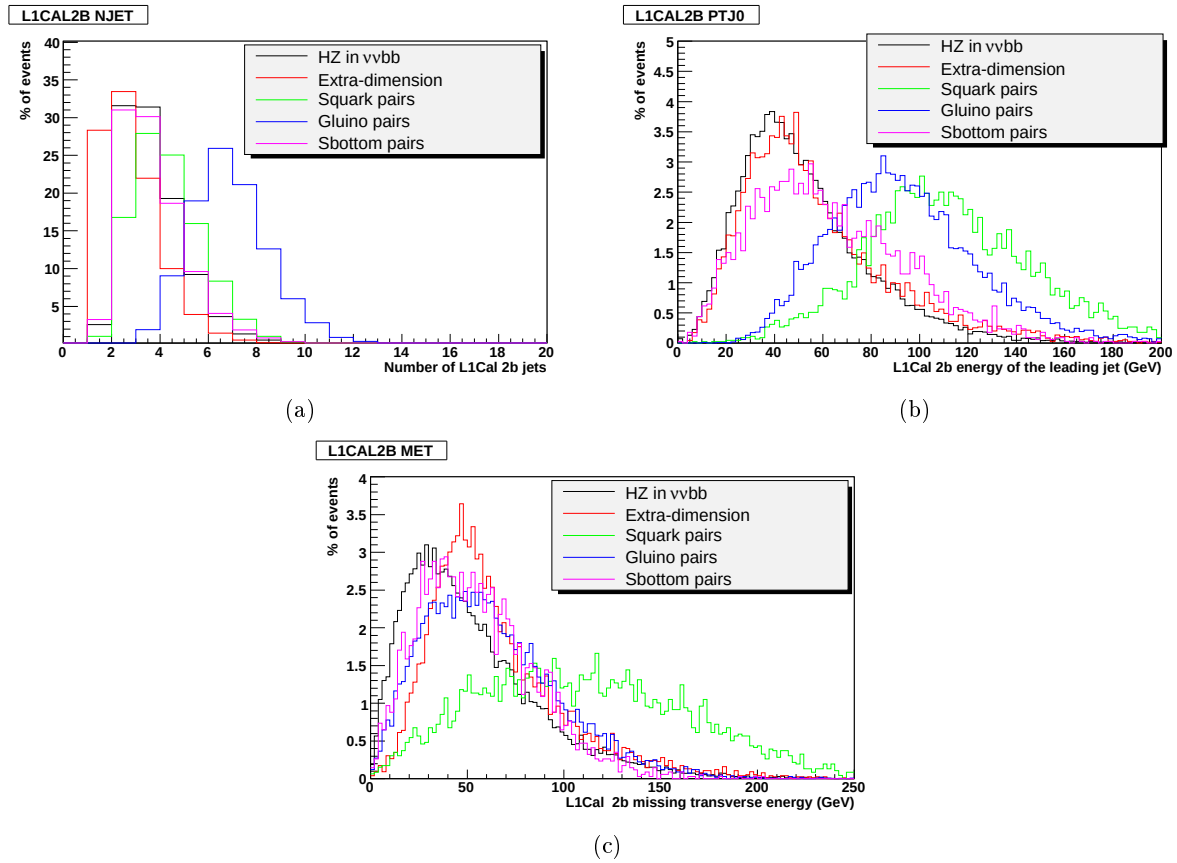


FIG. 3.33 – Distribution au niveau 1 du nombre de jets reconstruits (a), de l'énergie transverse du jet le plus dur (b) et de l'énergie transverse manquante pour les cinq signaux utilisés (c).

Les efficacités absolues et relatives pour les signaux étudiés sont résumées dans le Tableau 3.10. Ce tableau montre que les niveaux 1 proposés par le groupe de recherche du boson de Higgs sont exploitables pour les signaux du groupe de recherche de nouvelle physique. Il n'est donc pas nécessaire ni de les modifier, ni d'en ajouter. On constate par exemple que la recherche de gluinos peut parfaitement utiliser la condition conçue initialement pour la production Hbb . Afin d'augmenter les efficacités relatives des signaux correspondant à la production $HZ \rightarrow \bar{b}b\nu\nu$ et à la production de sbottoms, il a été décidé d'utiliser un "ou" des niveaux 1 *MONO*, *HZ* et *HBB*. Les efficacités relatives de ces deux signaux bénéficient de ce "ou" entre plusieurs termes et passent ainsi à 97.2 ± 0.5 pour la production de sbottoms et à 92.8 ± 0.3 pour la production HZ . Enfin, le niveau 1 appelé *MONO* est préféré au terme *QCD*, car la coupure de présélection sur $\cancel{E}_T$ est plus en adéquation avec ce terme et qu'il est préférable d'avoir une condition propre aux topologies dites monojet que de partager le même niveau 1 que celui conçu pour les études du fond *QCD*.

Finalement, les niveaux 1 adoptés pour la première liste *V15* sont les suivants :

- CSWJT(1,30,3.2)CSWMET(25) "Ou"
 CSWAKL(20,20,0)CSWJT(1,20,2.4)CSWJT(2,8,3.2)CSWMET(25) "Ou"
 CSWJT(1,30,2.4)CSWJT(2,15,2.4)CSWJT(3,8,3.2) : 446.5 ± 17.7 Hz

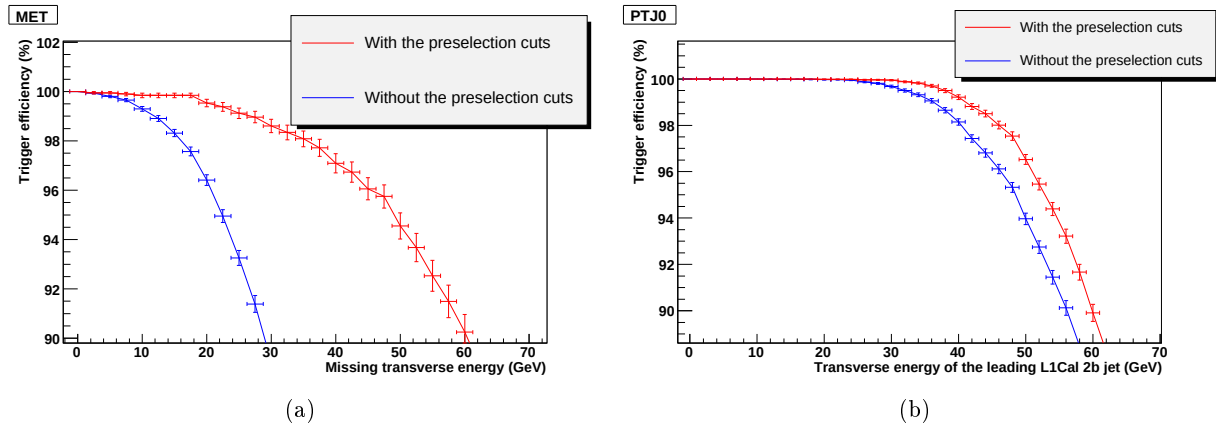


FIG. 3.34 – Efficacités relatives et absolues pour une coupure sur E_T de niveau 1 pour le signal correspondant à la recherche de dimensions supplémentaires (a) et pour une coupure sur le E_T du jet le plus dur de niveau 1 pour le signal correspondant à la recherche de gluinos (b).

	sbottoms	hZ	squarks	monojet	gluinos
<i>QCD</i>	67.1 ± 0.9	53.1 ± 0.3	97.9 ± 0.2	59.2 ± 0.6	96.4 ± 0.2
<i>MONO</i>	77.4 ± 0.8	68.9 ± 0.3	96.1 ± 0.3	81.9 ± 0.5	87.4 ± 0.4
<i>ZH</i>	71.7 ± 0.8	64.7 ± 0.3	92.0 ± 0.4	40.6 ± 0.6	80.8 ± 0.4
<i>HBB</i>	28.4 ± 0.8	23.8 ± 0.3	65.4 ± 0.8	9.9 ± 0.4	98.9 ± 0.1
<i>QCD</i>	78.2 ± 1.1	65.6 ± 0.5	99.2 ± 0.2	97.2 ± 0.4	98.2 ± 0.2
<i>MONO</i>	92.9 ± 0.7	86.2 ± 0.3	98.9 ± 0.2	98.9 ± 0.3	91.6 ± 0.3
<i>ZH</i>	89.2 ± 0.9	87.7 ± 0.3	96.5 ± 0.3	44.0 ± 1.2	85.5 ± 0.4
<i>HBB</i>	31.3 ± 1.3	29.2 ± 0.4	68.0 ± 0.8	12.5 ± 0.8	99.8 ± 0.1

TAB. 3.10 – Résumé des efficacités du niveau 1 pour la liste *V15* (absolues pour les quatre premières lignes et relatives pour les cinq suivantes).

- CSWJT(1,30,3.2)CSWMET(25) : 125.0 ± 9.4 Hz
- CSWJT(1,30,2.4)CSWJT(2,15,2.4)CSWJT(3,8,3.2) : 229.6 ± 12.7 Hz

Ces niveaux 1 permettent d'obtenir des efficacités relatives supérieures ou égales à celles obtenues avec les **mêmes outils et les mêmes signaux** de la liste *V14* (voir le Tableau 3.6) pour un taux d'acquisition total de 446.5 ± 17.7 Hz, c'est-à-dire approximativement deux fois plus faible que celui qu'aurait la liste *V14* à $300 \times 10^{30} \text{ cm}^{-2} \text{ s}^{-1}$ (résultat obtenu par une simple règle de trois). On en conclut donc que l'algorithme de 'fenêtres glissantes' s'avère très efficace et très discriminant pour le fond instrumental.

Afin de réduire le taux d'acquisition du niveau 1 conçu précédemment, différentes configurations ont été essayées pour l'algorithme de reconstruction de l'énergie transverse manquante de niveau 1. Les possibilités suivantes ont été étudiées pour calculer cette quantité :

- Utilisation ou non de la région inter-cryostat ($ICR : 1.0 < |\eta| < 1.4$).
- Appliquer une coupure sur la variable η .
- Appliquer un seuil sur l'énergie des tours (exprimée en *ADC*).

Le but de ces optimisations est de déclencher sur une énergie transverse manquante due essentiellement à la physique de l'événement plutôt qu'aux bruits et imperfections éventuels du calorimètre. Ainsi, ces modifications ne doivent pas avoir une incidence importante sur les signaux étudiés. Il a rapidement été décidé de ne pas utiliser la région inter-cryostat car elle peut être très sensible au bruit électronique. La coupure sur la variable η a également été rejetée, car elle modifiait très peu la forme des distributions et n'avait qu'un effet minime sur le taux d'acquisition. Le seul paramètre capable de modifier fortement ce taux est le seuil sur l'énergie transverse des tours de déclenchement. Cependant, si on augmente trop ce seuil, on interfère sur les énergies réelles des objets physiques et l'on crée une $\cancel{E}_T$ artificielle. Supposons par exemple qu'un événement possède seulement deux jets acoplanaires de même énergie, son énergie transverse manquante est nulle. Si l'énergie de ces deux jets n'a pas été déposée de la même façon et qu'on impose un seuil trop haut sur l'énergie des tours, alors l'énergie des deux jets ne se compensera plus et une $\cancel{E}_T$ artificielle sera créée. Il faut donc trouver le seuil optimal permettant de réduire la $\cancel{E}_T$ instrumentale sans augmenter trop fortement la $\cancel{E}_T$ artificiellement créée. La Figure 3.35 montre ces deux effets sur les taux d'acquisition des deux niveaux 1 utilisant la coupure sur $CSWMET$. On constate qu'un seuil de 4 ADC donne une valeur minimale pour le taux d'acquisition. La Figure 3.36 montre ensuite l'effet de ce seuil sur les efficacités relatives du signal, chaque point correspond à une valeur différente de la coupure sur $CSWMET$ et les cercles correspondent à la valeur 25 GeV finalement adoptée. La diminution des efficacités relatives est moindre, voire nulle, alors que le taux d'acquisition est diminuée d'un facteur 2 approximativement.

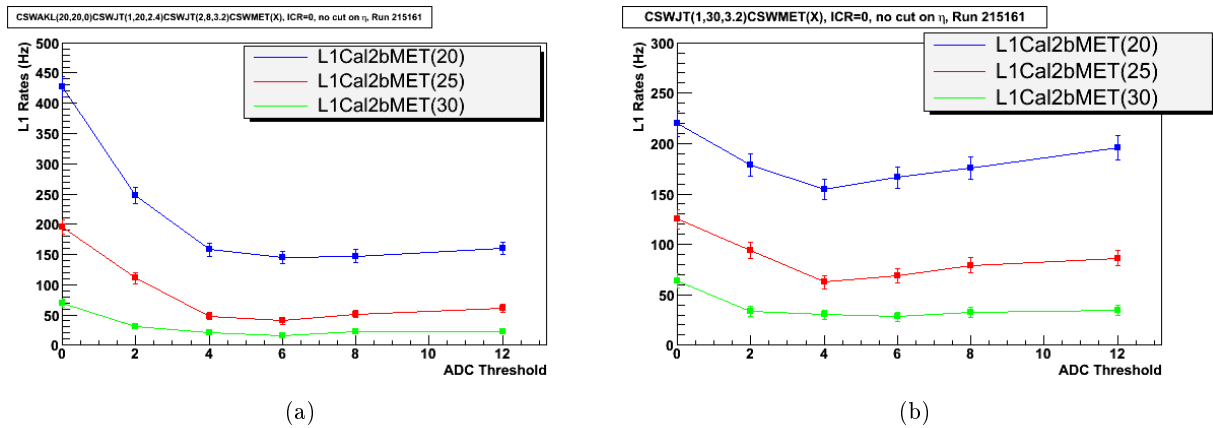


FIG. 3.35 – Taux d'acquisition du niveau 1 en fonction du seuil en coups ADC et de la valeur de la coupure sur $\cancel{E}_T$ pour la condition de déclenchement *dijet* (a) et pour la condition *monojet* (b).

Le seuil de 4 ADC a finalement été adopté et implémenté dans les algorithmes de reconstruction en ligne de l'énergie transverse manquante de niveau 1. Les taux d'acquisition des trois niveaux 1 conçus pour les conditions de déclenchement spécifiques aux jets et $\cancel{E}_T$ sont résumés dans le Tableau 3.11.

Conception du niveau 2

Les outils et termes permettant de concevoir les conditions de déclenchement sont les mêmes que ceux de la liste *V14* (voir la Section 3.2.3). Le niveau 2 a donc été mis au point avec ces

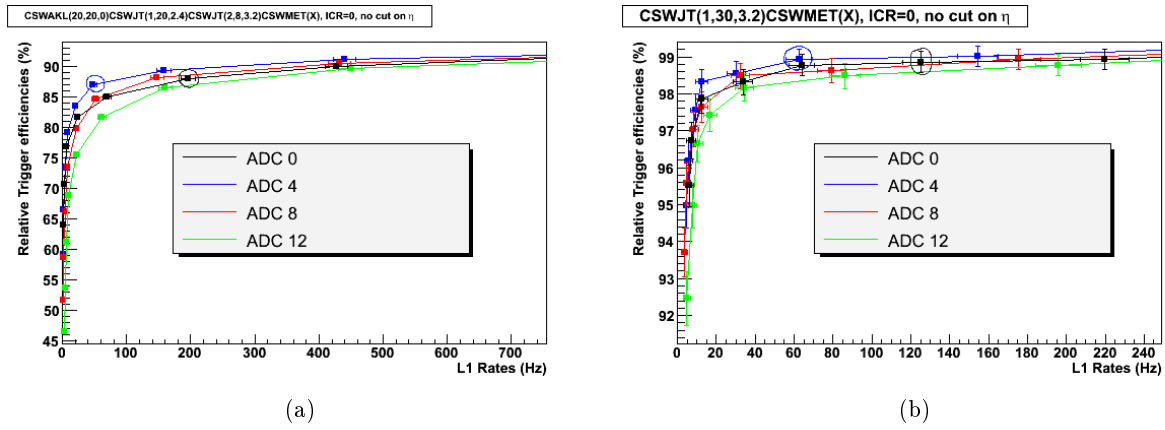


FIG. 3.36 – Effet du seuil et de la coupure sur E_T sur les efficacités relatives des signaux correspondant à la production HZ (a) et correspondant à la recherche de dimensions supplémentaires (b).

Niveau 1	Taux d'acquisition (Hz)
CSWAKL(20, 20, 0)CSWJT(1, 20, 2.4)CSWJT(2, 8, 3.2)CSWMET(25)	47.7 ± 5.8
CSWJT(1, 30, 3.2)CSWMET(25)	62.5 ± 6.6
CSWJT(1, 30, 2.4)CSWJT(2, 15, 2.4)CSWJT(3, 8, 3.2)	229.6 ± 12.7
Total	294.8 ± 14.4

TAB. 3.11 – Résumé des taux d'acquisition du niveau 1 de la liste *V15* après optimisation de l'algorithme calculant *CSWMET*

termes dans le même esprit que le niveau 1, c'est-à-dire :

- Une coupure sur le E_T du jet le plus dur et sur $\cancel{E}_T$ pour la condition de déclenchement dite *monojet*. Les variables utilisées sont donc $L2JET(1, X, Y)$, où Y est une coupure supplémentaire sur η et $L2MJT(X)$.
- Une sélection de deux jets acoplanaires avec de l'énergie transverse manquante est requise pour la condition dite *dijet*. On déclenche alors le niveau 2 grâce aux variables $L2JET(1, X, Y)$, $L2JET(2, X, Y)$ éventuellement, $L2MJT(X)$ et $L2ACOP(X)$.
- La condition dite *multijet* demande au moins trois jets dans l'état final avec une faible énergie transverse manquante. La sélection s'effectue sur $L2JET(1, X, Y)$, $L2JET(2, X, Y)$, $L2JET(3, X, Y)$, $L2HT(X)$ et éventuellement $L2MJT(X)$.

Pour estimer la valeur des coupures de chacune des variables (les lettres X dans la liste précédente), on trace l'efficacité relative du signal en fonction de la valeur prise par X en tenant compte seulement des événements ayant déclenché le niveau 1. La Figure 3.37 donne deux exemples permettant d'estimer une valeur conservative de la coupure, cette valeur est représentée par une flèche.

Pour la conception finale du niveau 2, on trace l'efficacité relative en fonction du taux d'événements déclenchant ce niveau 2 à une luminosité instantanée de $300 \times 10^{30} \text{ cm}^{-2} \text{ s}^{-1}$. La Figure 3.38 montre deux exemples de ce type de courbes.

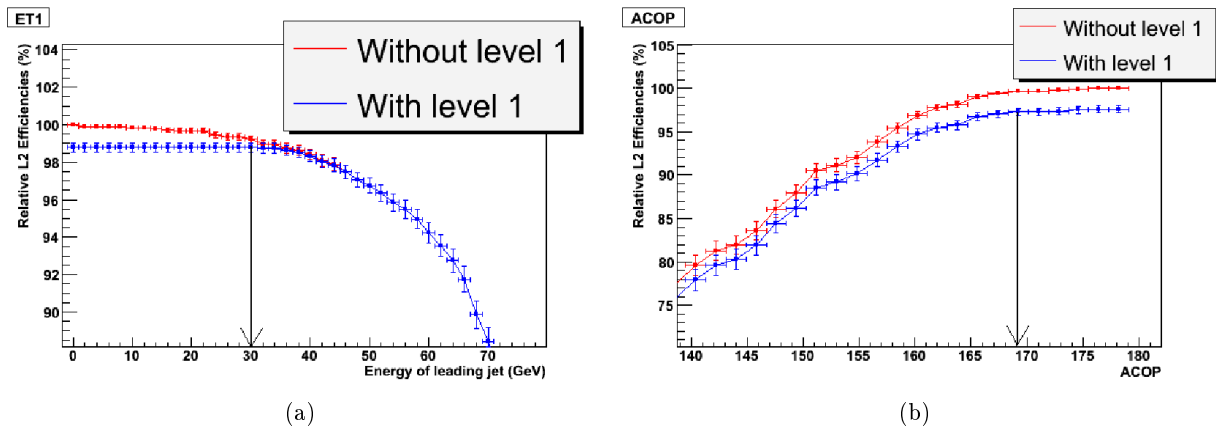


FIG. 3.37 – Efficacités relatives du signal correspondant à la recherche de dimensions supplémentaires en fonction d’une coupure sur $L2JET(1,X)$ (a) et correspondant à la production de sbottoms en fonction de $L2ACOP(X)$ (b).

Pour la condition de déclenchement *multijet*, il a été décidé de limiter la coupure sur $L2HT(X)$ à 75 GeV même si le signal correspondant à la production de gluinos permet d’augmenter encore cette coupure. En effet, cette condition est partagée avec des signaux, tels la production de Hbb , qui ne possèdent pas des jets de E_T suffisant pour déclencher ce terme au-delà de 75 GeV. Une coupure à 10 GeV sur $L2MJT(X)$ a été ajoutée pour réduire le taux d’acquisition. Pour la condition dite *dijet*, la Figure 3.38 montre qu’il est préférable d’augmenter la coupure sur $L2ACOP(X)$ plutôt que d’introduire une coupure sur l’énergie du deuxième jet le plus dur. La valeur de 168.75 identique à la liste *V14* a été choisie. Les niveaux 2 finalement choisis sont récapitulés dans le Tableau 3.12. Les taux d’acquisition ainsi que les efficacités relatives à ces niveaux 2 sont donnés dans le Tableau 3.13.

Type de signal	Nom du niveau 2
<i>Monojet</i>	L2JET(1,35,2.4)L2MJT(25)
<i>Dijet</i>	L2JET(1,20,2.4)L2MJT(20)L2HT(35)L2ACOP(168.75)
<i>Multijet</i>	L2JET(1,20,2.4)L2JET(2,15,2.6)L2JET(3,8,3.6)L2MJT(10)L2HT(75)

TAB. 3.12 – Résumé des niveaux 2 choisis.

Les niveaux 1 et 2 détaillés précédemment ont été conçus à partir d’estimations du taux d’événements passant ces deux niveaux à une luminosité instantanée de $300 \times 10^{30} \text{cm}^{-2} \text{s}^{-1}$, sachant qu’aucune période de prise de données du *Run IIa* n’atteint cette valeur. L’outil utilisé a été *trigger_rate_tool* en supposant que les taux d’acquisition sont linéaires avec la luminosité, ce qui était difficilement vérifiable a priori. De plus, les efficacités et taux obtenus ont été estimés à partir d’une émulation des nouveaux outils proposés pour la liste *V15*. Les résultats précédents, obtenus avant le début du *Run IIb*, doivent donc être confirmés avec d’une part les premières données du *Run IIb* et d’autre part une nouvelle estimation des efficacités des signaux. Ces deux points sont détaillés en [102]. En [102], les signaux sont identiques. Ils correspondent aux mêmes processus physiques. Par contre, les versions des outils utilisés sont différentes. Elles tiennent compte

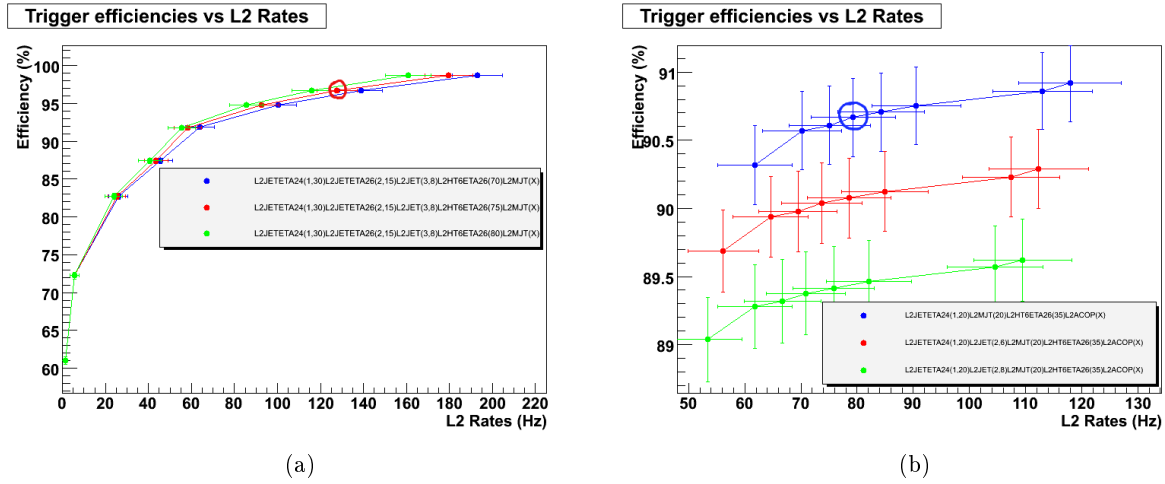


FIG. 3.38 – Efficacités relatives du signal correspondant à la production de gluinos en fonction d’une coupure sur $L2MJT(X)$ et sur $L2HT(X)$ (a) et correspondant à la production HZ en fonction de $L2ACOP(X)$ et de $L2JET(2,X,Y)$ (b).

Type de signal	Signaux	Taux d’acquisition de niveau 2 (Hz)	Efficacités absolues (%)	Efficacités relatives (%)
<i>Monojet</i>	Dimensions supplémentaires squarks	43.5 ± 5.5	78.9 ± 0.5	98.3 ± 0.3
			95.2 ± 0.3	98.6 ± 0.2
<i>Dijet</i>	HZ sbottoms	79.3 ± 7.5	66.6 ± 0.3	90.6 ± 0.3
			73.7 ± 0.8	95.3 ± 0.6
<i>Multijet</i>	gluinos	127.8 ± 9.5	96.7 ± 0.2	98.4 ± 0.2

TAB. 3.13 – Résumé des taux d’acquisition et des efficacités obtenus pour le niveau 2.

des améliorations de la simulation du système de déclenchement (*d0trigsim*) et de la simulation *Monte Carlo* des signaux. Ces versions proposent de plus les algorithmes de ‘fenêtres glissantes’ identiques à ceux présents en ligne. Une autre amélioration importante est le fait que les événements de données ajoutés aux événements *Monte Carlo* des signaux générés reproduisent mieux le profil de luminosité du *Run Iib* (ajout d’événements dits de *zero bias* en anglais). En effet, à hautes luminosités le nombre de collisions moyennes par croisement de faisceaux est augmenté. Ce paramètre peut s’avérer très important pour le terme $CSWAKL(n,X,Y)$. C’est d’ailleurs pour éviter une perte d’efficacité trop grande pour ce terme due aux événements des collisions supplémentaires que le seuil minimal imposé sur les jets est passé de 5 à 8 GeV. Finalement, la coupure sur $CSWMET(X)$ est de 24 GeV au lieu de 25 GeV en ligne (voir [100] pour plus d’explications à ce sujet). Malgré toutes ces différences et améliorations, les efficacités et taux obtenus avant le démarrage du *Run Iib* sont dans la plupart des cas en bon accord avec ceux obtenus a posteriori, dans les autres cas les différences sont bien comprises et expliquées en [102]. Les niveaux 1 et 2 détaillés ici sont donc confirmés avec les deux différences sur $CSWMET(X)$ et $CSWAKL(n,X,Y)$ évoqués au-dessus.

Finalement, les études antérieures au *Run IIb* ont permis de montrer qu'il était possible de concevoir trois conditions de déclenchement pour les signaux à jets et $\cancel{E}_T$ avec des efficacités relatives à des coupures de présélection supérieures à celles obtenues avec la liste *V14* et des taux d'acquisition autorisant un fonctionnement à une luminosité instantanée de $300 \times 10^{30} \text{ cm}^{-2} \text{ s}^{-1}$. Ces comparaisons sont résumées dans les Tableaux 3.14 et 3.15.

Type de signal	Taux <i>V14</i> de niveau 1 (Hz)	Taux <i>V15</i> de niveau 1 (Hz)	Taux <i>V14</i> de niveau 2 (Hz)	Taux <i>V15</i> de niveau 2 (Hz)
<i>Monojet</i>	817.1 ± 24.0	62.5 ± 6.6	115.1 ± 9.0	43.5 ± 5.5
<i>Dijet</i>	817.1 ± 24.0	294.8 ± 14.4	115.1 ± 9.0	79.3 ± 7.5
<i>Multijet</i>	817.1 ± 24.0	229.6 ± 12.7	249.2 ± 13.2	128.5 ± 9.5
Total	817.1 ± 24.0	294.8 ± 14.4	315.9 ± 14.9	183.9 ± 11.4

TAB. 3.14 – Comparaisons entre les listes *V14* et *V15* des taux d'acquisition estimés à une luminosité instantanée de $300 \times 10^{30} \text{ cm}^{-2} \text{ s}^{-1}$

Signaux	Niveau 1 (<i>V14</i>) (%)	Niveau 1 (<i>V15</i>) (%)	Niveau 2 (<i>V14</i>) (%)	Niveau 2 (<i>V15</i>) (%)
Dimensions supplémentaires	89.3 ± 1.0	98.6 ± 0.3	85.9 ± 1.1	98.3 ± 0.3
Squarks	99.9 ± 0.0	98.7 ± 0.2	91.3 ± 0.5	98.6 ± 0.2
<i>HZ</i>	88.5 ± 0.3	92.0 ± 0.3	86.9 ± 0.3	90.6 ± 0.3
Sbottoms	90.9 ± 0.7	97.2 ± 0.5	90.0 ± 0.8	95.3 ± 0.6
Gluinos	100.0 ± 0.0	99.8 ± 0.1	99.8 ± 0.1	98.4 ± 0.2

TAB. 3.15 – Comparaisons entre les listes *V14* et *V15* des efficacités relatives

Conception du niveau 3

Au démarrage du *Run IIb*, c'est-à-dire pour la liste *V15.00*, les niveaux 3 des conditions de déclenchement spécifiques aux topologies à jets et $\cancel{E}_T$ étaient identiques à ceux de la liste *V14*. Ces niveaux 3 ont par conséquent dû être modifiés ou complètement reconçus. Les détails de ces études sur le choix de chacune des conditions peuvent être trouvés en [102]. Le choix final des termes choisis pour les niveaux 3 est résumé dans le Tableau 3.16. On y trouve également le rappel des niveaux 1 et 2.

La particularité de cette conception est l'introduction d'une nouvelle variable disponible au niveau 3 : *L3MET*. Cette variable s'avère beaucoup plus efficace que la variable *L3MHT* pour les signaux étudiés. Par mesure de précaution, une condition de déclenchement dite *dijet alternatif* a été conservée afin d'être sûr de bien comprendre ce nouveau terme *L3MET* et de pouvoir l'appliquer correctement sur les simulations *Monte Carlo*. En effet, n'étant pas utilisée auparavant, des études complémentaires sont nécessaires. La comparaison des taux d'acquisition et des efficacités entre la liste *V15.00* et la liste *V15.20*, première liste intégrant ces nouveaux niveaux 3, est donnée par le Tableau 3.17.

	<i>dijet avec $\cancel{E}_T$</i>
Niveau 1	CSWMET(24)CSWJT(1,20,2.4)CSWJT(2,8,2.4)CSWAKL(0,20,20) OR Niveau 1 <i>monojet</i> OR Niveau 1 <i>multijet</i>
Niveau 2	L2JET(1,20,2.4)L2MJT(20)L2HT(35)L2ACOP(168.75)
Niveau 3	L3JET(2,9,3.6), L3MHT(1,9,3.6) > 25, L3ACOP(1,9,3.6) > 170 L3AngleMHTJET(1,9,3.6) > 25, L3MET > 25
	<i>dijet alternatif</i>
Niveau 1	CSWMET(24)CSWJT(1,20,2.4)CSWJT(2,8,2.4)CSWAKL(0,20,20) OR Niveau 1 <i>monojet</i> OR Niveau 1 <i>multijet</i>
Niveau 2	L2JET(1,20,2.4)L2MJT(20)L2HT(35)L2ACOP(168.75)
Niveau 3	L3JET(2,9,3.6), L3MHT(1,9,3.6) > 35 L3ACOP(1,9,3.6) > 170, L3AngleMHTJET(1,9,3.6) > 25
	<i>monojet</i>
Niveau 1	CSWMET(24)CSWJT(1,30,3.2)
Niveau 2	L2JET(1,35,2.4)L2MJT(25)
Niveau 3	L3MHT(1,9,3.6) > 45, L3HT(1,9,3.6) > 70, L3JET(1,9,3.6)
	<i>multijet</i>
Niveau 1	CSWJT(1,30,2.4)CSWJT(2,15,2.4)CSWJT(3,8,3.2)
Niveau 2	L2JET(1,30,2.6)L2JET(2,15,2.6)L2JET(3,8,3.2)L2HT(75)L2MJT(10)
Niveau 3	L3MHT(1,9,3.6) > 40, L3HT(1,9,3.6) > 140, L3JET(3,20,3.6)

TAB. 3.16 – Description des termes de niveau 1,2 et 3 pour les conditions spécifiques aux jets et $\cancel{E}_T$ de la liste *V15.20*

	Taux d'acquisition (Hz)				Gain (%)	Efficacités relatives (%)		Gain (%)
	<i>V15.00</i>		<i>V15.20</i>			<i>V15.00</i>	<i>V15.20</i>	
	Inclusifs	Exclusifs	Inclusifs	Exclusifs				
<i>dijet</i>	6.9±0.4	3.3±0.3	3.9±0.3	1.4±0.2	-44%			
<i>HZ</i>						87.9±0.6	91.3±0.5	+ 3%
Sbottoms						88.0±0.8	89.3±0.7	+ 1%
Squarks						98.3±0.2	98.3±0.2	0
<i>dijet alternatif</i>	6.9±0.4	3.3±0.3	4.6±0.3	2.1±0.2	-33%			
<i>HZ</i>						87.9±0.6	83.2±0.6	- 5%
Sbottoms						88.0±0.8	87.7±0.8	- 1%
Squarks						98.3±0.2	98.2±0.2	0
<i>monojet</i>	3.5±0.3	0.2±0.1	2.0±0.2	1.2±0.2	-43%			
Dimensions supplémentaires						92.1±0.7	98.9±0.3	+7%
<i>multijet</i>	4.7±0.3	4.3±0.3	1.9±0.2	1.4±0.2	-60%			
Gluinos						97.4±0.3	96.2±0.3	-1%

TAB. 3.17 – Résumé des résultats et comparaison entre la liste *V15.00* et la liste *V15.20*

Le Tableau 3.17 montre que pour chacun des signaux dijets et monojet, les efficacités augmentent (de 3 % à 7 %) alors que le taux d'acquisition diminue de près de 45 %. Pour la condition de déclenchement *multijet*, ce taux descend même de 60 % approximativement alors que l'efficacité relative ne subit qu'une diminution de 1 %. Cette proposition de niveaux 3 a donc été acceptée pour compléter le travail fait au préalable pour les niveaux 1 et 2. Les signaux ayant des topologies à jets et E_T bénéficient donc de conditions de déclenchement efficaces en mesure des prendre des données jusqu'à des luminosités instantanées proches de $300 \times 10^{30} \text{ cm}^{-2} \text{ s}^{-1}$. Les estimations des taux prévus dans le Tableau 3.17 ont finalement été vérifiées une fois la liste *V15.20* mise en ligne. La Figure 3.39 montre ces taux pour la condition de déclenchement *dijet* avant la nouvelle conception du niveau 3 (liste *V15.17*) et après l'introduction des nouveaux niveaux 3 (liste *V15.20*). Le taux d'acquisition est effectivement diminué d'un facteur 2 approximativement comme les estimations le prévoyaient. Toutes les vérifications effectuées ont finalement conclu au bon fonctionnement des conditions de déclenchement de la liste *V15* en bon accord avec les études menées et les estimations proposées.

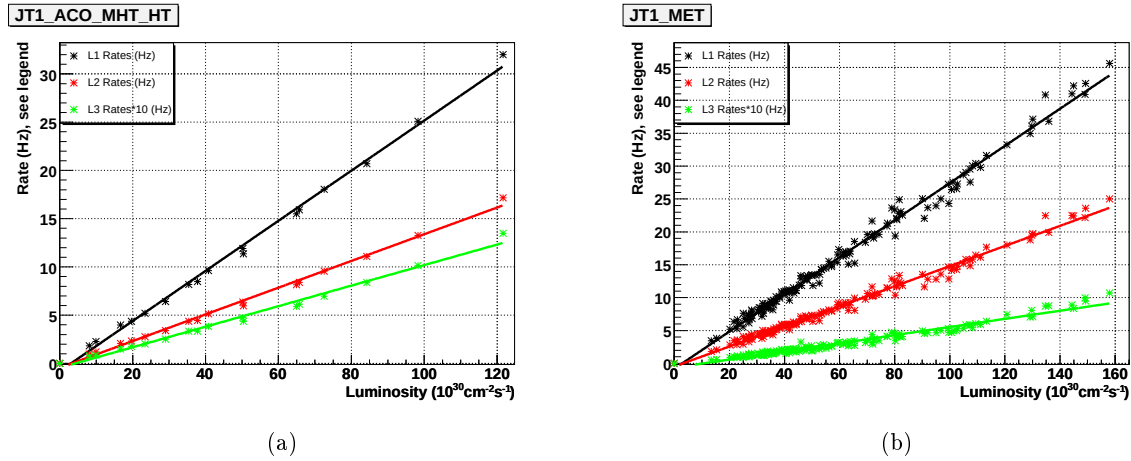


FIG. 3.39 – Comparaison entre les taux de niveaux 1,2 et 3 pour les listes *V15.17* (a) et *V15.20* (b).

3.4 Reconstruction des objets physiques

Une fois qu'un événement a satisfait au moins une condition de déclenchement, il est enregistré sur bande magnétique dans un format brut avec l'ensemble des informations des sous-détecteurs renvoyé par l'électronique de lecture. Ce format est appelé *raw data*. A ce stade de la chaîne d'acquisition et du traitement de données, aucun algorithme de précision n'a agi, les contraintes de temps du système de déclenchement limitant en effet tout traitement plus approfondi des données. Il est donc nécessaire de convertir cette information en terme d'objets physiques tels les muons, les électrons ou les jets ou en terme de grandeurs comme les traces, les vertex ou l'énergie transverse manquante. Ces quantités plus élaborées seront ensuite utilisées par l'ensemble des analyses de physique. Elles représentent une traduction de la physique de l'événement par le détecteur et le traitement de l'information qu'il fournit. Cette opération est effectuée par le programme *d0reco*, qui regroupe l'ensemble des algorithmes nécessaires. Le programme *d0reco* tourne sur les données du détecteur et les données des simulations *Monte Carlo*.

Les algorithmes et méthodes de reconstruction employés par le programme *d0reco* sont décrits dans les sections suivantes. Dans la Section 3.4.1, la reconstruction des traces et des vertex seront détaillées conjointement, l'un n'allant pas sans l'autre. Les Sections 3.4.2 et 3.4.3 sont consacrées aux deux leptons les plus légers. Je mettrai finalement l'accent sur les Sections 3.4.4, 3.4.5 et 3.4.6, qui correspondent à une description plus détaillée des objets physiques ou grandeurs globales utilisés dans le Chapitre 4 dédié à l'analyse de données et à la recherche des gluinos.

3.4.1 Les traces et les vertex

Reconstruction des traces

La reconstruction des traces et des vertex utilise les impacts détectés dans le *CFT* et le *SMT* décrits dans la Section 2.3.2. Les algorithmes permettant cette reconstruction sont ceux qui consomment le plus de temps lorsque l'exécutable du programme *d0reco* tourne. Ce temps de calcul augmente avec le nombre d'événements empilés comme illustré sur la Figure 3.40, donc avec la luminosité instantanée, de façon exponentielle. La recherche de performance et d'optimisation de ces algorithmes est un véritable défi [103], et ce d'autant plus depuis le début du *Run IIb*.

Plusieurs régions des deux détecteurs utilisés sont à distinguer pour la reconstruction des traces :

- Dans la région centrale, $|\eta| < 1.7$, un maximum d'informations est disponible et la reconstruction s'en trouve facilitée. Les algorithmes de reconstruction utilisent les impacts détectés à la fois dans le *CFT* et le *SMT*. Une trace possède jusqu'à 16 impacts dans le *CFT*.
- Pour la région $1.7 < |\eta| < 2$, le nombre d'impacts dans le *CFT* est réduit de 8 à 15.
- Dans la région à faibles angles, $|\eta| > 2$, seul le *SMT* permet la reconstruction des traces.

Cette reconstruction s'effectue en deux temps. Les algorithmes sont basés sur un ajustement de Kalman [104]. Dans un premier temps, deux algorithmes, le *Histogram Track Finding* décrit en [105, 105] et le *AATrack Finder* [106], effectuent une reconstruction des traces en calculant les paramètres de ces traces par rapport au centre géométrique du détecteur. Ensuite, un troisième algorithme les ajuste en tenant compte de la courbure magnétique et des interactions avec les matériaux du détecteur. Ce troisième algorithme, appelé *Global Track Finder* [107], fournit en fait une description très précise du mouvement d'une particule chargée dans le détecteur DØ. Cette première étape permet de positionner le vertex associé à ces traces. Dans un second temps, les

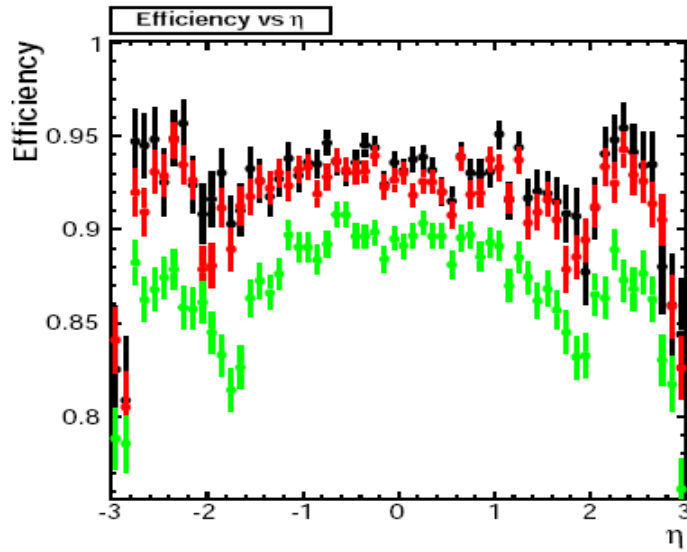


FIG. 3.40 – Efficacités de reconstruction des traces en fonction de η avec 0 événement empilé (noir), 4 (rouge) et 8 (vert) pour des simulations *Monte Carlo* de production de $t\bar{t}$.

paramètres des traces sont réévalués par rapport au vertex primaire ainsi déterminé. La reconstruction des traces se limitent aux traces d’une impulsion minimum de 183 MeV.

Identification des vertex de l’événement

L’identification des vertex primaires d’un événement est une étape primordiale et nécessaire à la reconstruction de tous les objets physiques. Le vertex primaire associé à l’interaction dure, en opposition aux interactions de biais minimum, est en effet utilisé pour calculer les composantes de l’énergie des cellules du calorimètre et permet donc de mesurer l’énergie transverse des jets et des objets électromagnétiques. Une mauvaise identification du vertex ou une imprécision sur sa position entraîne donc une détérioration de l’énergie transverse manquante (voir les Sections 3.4.3, 3.4.4 et 3.4.5). L’impulsion transverse des muons est également mesurée à partir de la position de ce vertex primaire (voir la Section 3.4.2). Enfin, la précision sur sa position a un impact direct sur la précision et l’efficacité de l’étiquetage des hadrons beaux (voir la Section 3.4.6). Cette position est limitée par les dimensions du faisceau. Elle peut varier de l’ordre de 30 μm dans le plan transverse et de quelques centimètres le long de l’axe z . La précision des algorithmes d’identification des vertex a donc des conséquences importantes sur l’ensemble du traitement *offline* des données.

Plusieurs étapes sont requises pour déterminer la liste des vertex de l’événement [108] et leurs positions. Tout d’abord, il faut sélectionner les traces qui vont être utilisées en entrée des algorithmes décrits ensuite. Ces traces doivent avoir une impulsion calculée à partir du centre du détecteur suffisante, $p_T > 0.5$ GeV, et au moins deux impacts dans le *SMT* si les traces sont dans l’acceptance du détecteur. Un algorithme permet alors de former des pseudo-vertex le long de l’axe z séparés au minimum de 2 cm à partir des traces sélectionnées. Une première itération à partir de ces traces et des pseudo-vertex permet d’affiner la position des vertex et de caractériser la géométrie du faisceau (position, largeur, etc.). Pour chacun des pseudo-vertex, on calcule à

l'aide d'un ajustement de Kalman [104] la contribution au χ^2 des traces et on rejette celles de plus grande contribution jusqu'à ce que le χ^2 par degré de liberté pour un pseudo-vertex donné soit inférieur à 10. La procédure reprend ensuite avec les traces rejetées pour les autres pseudo-vertex. Les positions trouvées correspondent alors à la liste des vertex primaires. Une deuxième itération utilise les informations sur la géométrie du détecteur (l'algorithme précédent utilisait une valeur fixe pour la largeur du faisceau de $30 \mu\text{m}$) et la position des vertex pour déterminer plus précisément les caractéristiques des traces et ainsi réappliquer la même procédure. On a ainsi une détermination précise de la position des vertex primaires et les valeurs des impulsions des traces provenant de ces vertex.

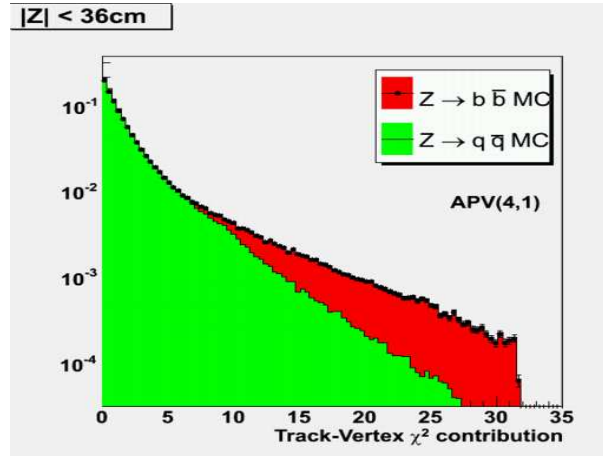


FIG. 3.41 – Contribution au χ^2 des traces pour des événements avec ou sans particules massives (quark b dans cet exemple).

Les deux principales difficultés rencontrées dans la reconstruction des vertex est d'une part la haute luminosité instantanée qui augmente le nombre de vertex primaires. Ce problème est correctement pris en compte par un simple ajustement de Kalman. La seconde difficulté est la présence de vertex secondaires dus à des désintégrations de particules massives, comme les hadrons beaux par exemple [108, 109]. La Figure 3.41 montre qu'un événement avec des saveurs lourdes possède plus de traces avec une contribution importante au χ^2 de l'ajustement de Kalman. Un algorithme a récemment été développé pour réduire la contribution des traces provenant de vertex secondaires dans l'ajustement de Kalman : *Adaptive Vertex Fitting*. Cet algorithme augmente la précision sur la position des vertex primaires comme le montre le Tableau 3.18 qui donne la résolution sur la position du vertex à partir de simulations *Monte Carlo*. De plus, ce nouvel algorithme augmente considérablement l'efficacité de reconstruction du vertex primaire si sa position sur l'axe z est éloignée du centre du détecteur.

Une sélection probabiliste est finalement effectuée pour déterminer quel est le vertex de l'interaction dure. Cette sélection est détaillée en [110].

3.4.2 Les muons

La reconstruction des muons dans DØ [111] peut utiliser les informations de trois sous-détecteurs. Comme l'illustre la Figure 2.9, les muons vont interagir dans le détecteur de traces, puis déposer de l'énergie dans le calorimètre, pour finalement être identifiés dans le spectromètre

Algorithmes	$Z \rightarrow b\bar{b}$ Résolution (μm)	$Z \rightarrow q\bar{q}$ Résolution (μm)
Ajustement de Kalman	14.1	9.1
<i>Adaptive Vertex Fitting</i>	12.8	9.3

TAB. 3.18 – Comparaison entre l'ajustement de Kalman simplement et l'*Adaptive Vertex Fitting* à partir de données *Monte Carlo*.

à muons. Pour l'instant, des études sont menées pour l'identification des muons à l'aide du calorimètre, mais l'algorithme existant appelé *Muon Tracking in the Calorimeter (MTC)* ne permet pas d'avoir des efficacités d'identification suffisante pour être utilisé (approximativement 50 %). Un muon est tout d'abord reconstruit à partir de l'information des trois couches du spectromètre à muons (voir la Section 2.3.4 pour une description plus complète de ce sous-détecteur), qui couvre près de 90 % de l'acceptance angulaire totale du détecteur, c'est-à-dire jusqu'à une valeur de la pseudo-rapacité $|\eta| < 2$. Les muons reconstruits grâce à cette information seulement sont dits "locaux". En addition du spectromètre à muons, le détecteur de traces, constitué du *SMT* et du *CFT*, fournit une estimation de l'impulsion du muon beaucoup plus précise et est très efficace pour trouver des traces dans toutes l'acceptance du spectromètre à muons. Un muon "local" dont les traces ont été reconstruites dans le détecteur de traces centrales est appelé *Central track-matched muon* en anglais.

L'algorithme de reconstruction des muons commence donc par construire des segments dans les couches du spectromètre à muons à partir des impacts mesurés dans les chambres à dérive ou les scintillateurs. La taille de ces segments varie en fonction du nombre de couches : un segment avec des impacts dans les trois couches ($|nseg| = 3$), un segment avec des impacts dans les couches B et C seulement ($|nseg| = 2$), c'est-à-dire en dehors du toroïde ou bien un segment avec des impacts dans la couche A seulement ($|nseg| = 1$). Ces segments sont ensuite associés dans la mesure du possible aux traces centrales.

Les muons sont alors classifiés en plusieurs catégories selon leur "type" ou leur "qualité". Sept "types" de muon sont définis. Les trois premiers "types" de muons correspondent à ceux qui ont été associés à au moins une trace centrale, $nseg > 0$. Les trois derniers sont les muons "locaux", $nseg < 0$. Le "type", $nseg = 0$, est un "type" particulier qui décrit un muon avec des traces centrales reconstruites et un impact dans le spectromètre ou une identification à l'aide du calorimètre. Le critère de "qualité" des muons est ensuite plus restrictif en ajoutant des conditions supplémentaires à celles imposées par le "type" du muon. Trois "qualités" sont définies *Loose*, *Medium* ou *Tight*.

Efficacités de reconstruction

L'efficacité de reconstruction est calculée pour les données du détecteur et les données *Monte Carlo* par l'étude des processus $Z \rightarrow \mu^+ \mu^-$. Cette efficacité est estimée en fonction des coordonnées spatiales η et ϕ . La Figure 3.42 montre dans un plan à deux dimensions l'évolution de l'efficacité d'identification dans l'espace (η, ϕ) en fonction de critères de plus en plus restrictifs sur la "qualité" du muon. On remarque une zone où l'efficacité est très faible voire nulle. Cette région du détecteur est en fait une région non instrumentée pour le spectromètre à muons située sous le détecteur DØ.

L'efficacité moyenne sur les données du détecteur varie en moyenne de 94.8 % à 75.9 % du

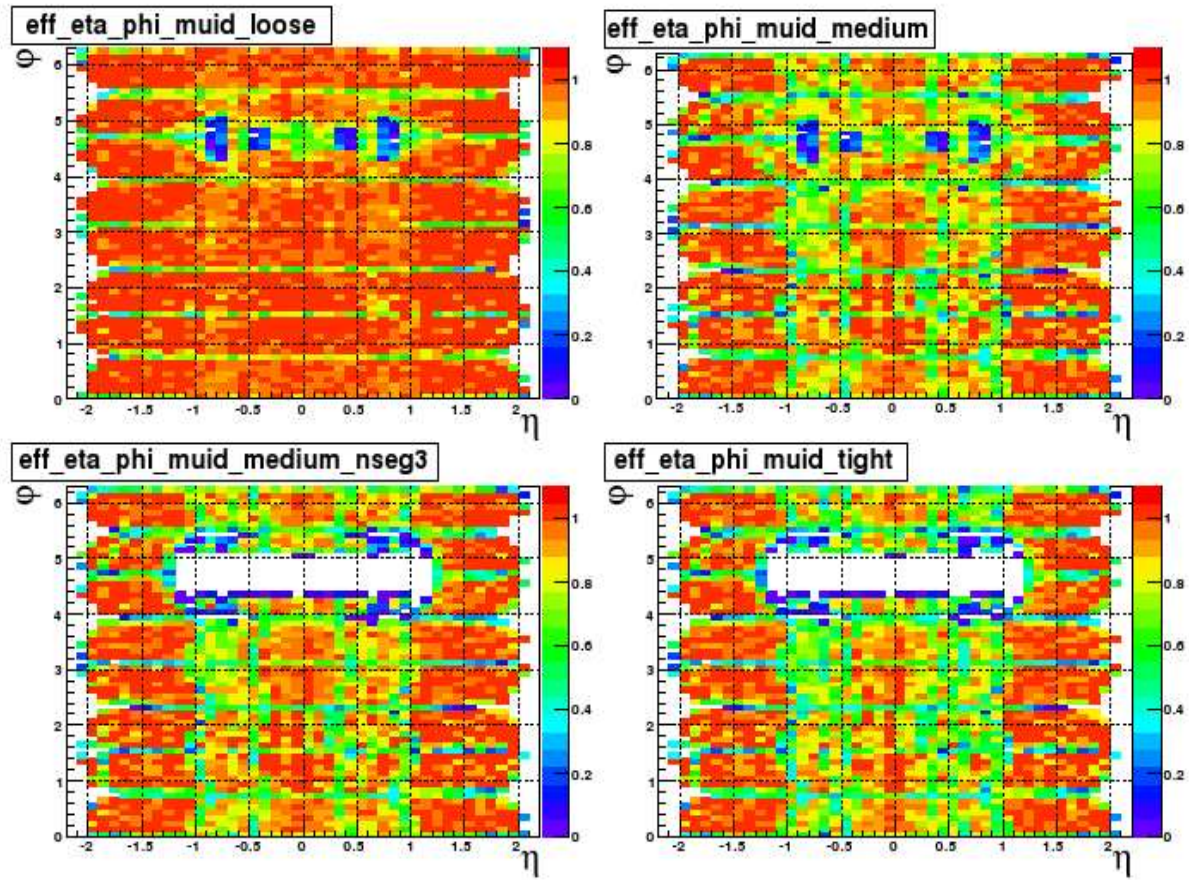


FIG. 3.42 – Efficacité d'identification d'un muon dans l'espace (η, ϕ) en fonction de différents critères de "qualité".

critère *Loose* au critère *Tight*. Cette efficacité est légèrement supérieure pour les données *Monte Carlo*, il faut alors introduire un facteur correctif allant en moyenne de 0.995 pour les muons *Loose* à 0.970 pour les muons *Tight*.

Critères d'isolation

Les critères d'isolation permettent de séparer les muons provenant des processus $W \rightarrow \mu\nu_\mu$ des muons provenant des désintégrations de hadrons beaux ou charmés. En effet, dans le deuxième cas, les traces associées au jet ou bien l'énergie déposée par le jet sont proches géographiquement de la trace centrale associée au muon. Ces critères peuvent s'avérer très utiles dans de nombreuses analyses pour discriminer le fond instrumental avec des muons dans l'état final. Ces critères seront notamment utilisés dans le Chapitre 4. On peut définir par exemple les variables :

- *TrackHalo* : Somme vectorielle des traces associées aux jets contenus dans un cône défini par $\Delta R(\text{traces des jets}, \text{trace du muon}) < 0.5$ (voir la Figure 3.43).
- *CalorimeterHalo* : Somme vectorielle de l'énergie transverse des cellules associées aux jets contenus dans un cône défini par $0.1 < \Delta R(\text{cellules associées aux jets}, \text{trace du muon}) <$

0.4.

- $\Delta R(\mu, jet)$: distance entre le muon et le jet le plus proche (voir la Figure 3.43).
- *ScaledCalorimeterHalo* qui correspond au ratio de la variable *CalorimeterHalo* par l'impulsion transverse du muon.
- *ScaledTrackHalo* qui correspond au ratio de la variable *TrackHalo* par l'impulsion transverse du muon.

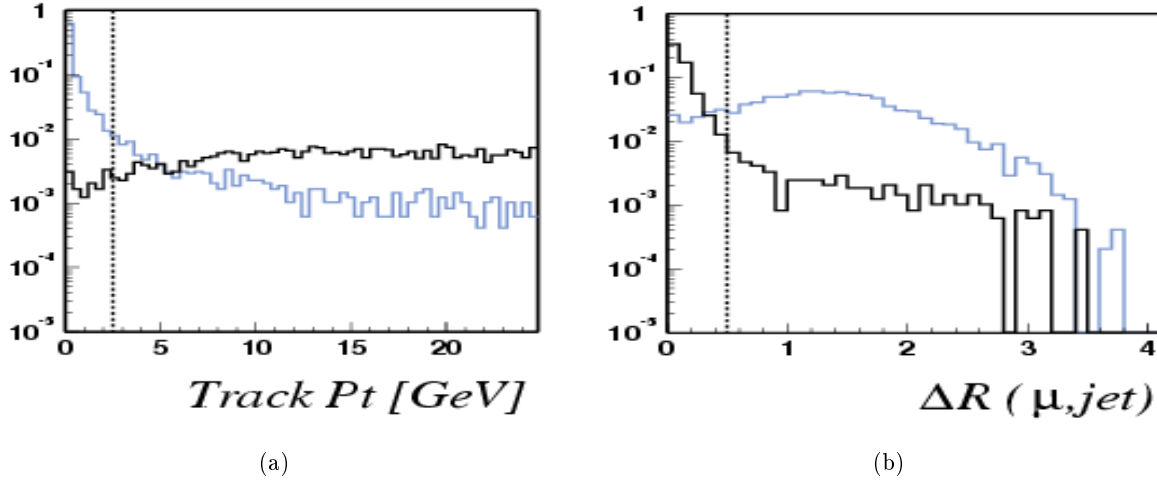


FIG. 3.43 – Distributions des variables $\Delta R(\mu, jet)$ (a) et *TrackHalo* (b) pour des données simulées correspondant à la désintégration du boson W en muons (bleu) et correspondant à la désintégration semi-leptonique en muon d'une production de paires de quarks top (noir).

Résolution sur l'impulsion et correction de l'impulsion des données *Monte Carlo*

La résolution sur l'impulsion des muons est estimée en utilisant la largeur du pic de la résonance du boson Z calculée pour les processus $Z \rightarrow \mu^+ \mu^-$. Cette résolution est déterminée à la fois pour les données du détecteur et les données *Monte Carlo*. La différence sur la position et la largeur du pic nécessitent des corrections à appliquer aux données *Monte Carlo* en recalculant l'impulsion du muon comme donnée par la Formule 3.4 où q est la charge du muon, Rnd est un nombre aléatoire gaussien (largeur 1 et centré sur 0), et A et B sont deux paramètres à déterminer.

$$\frac{q}{p_T} = \frac{q}{\alpha \times p_T} + \left(A + \frac{B}{\alpha \times p_T} \right) \times Rnd \quad (3.4)$$

Les résultats de cette correction sont illustrés sur la Figure 3.44.

3.4.3 Les particules électromagnétiques

La reconstruction des particules électromagnétiques est basée essentiellement sur les informations fournies par le calorimètre et plus particulièrement celles du calorimètre électromagnétique et celles de la couche *FH1* du calorimètre hadronique.

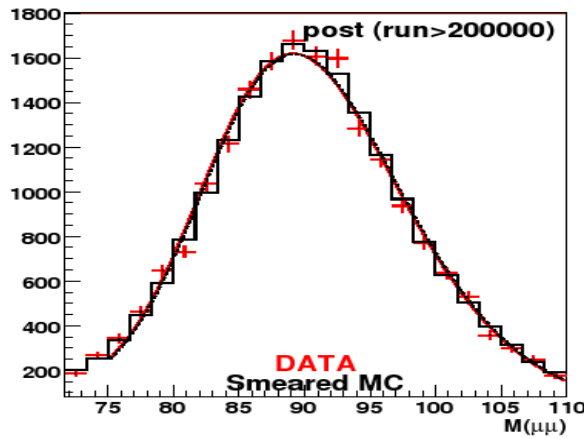


FIG. 3.44 – Pic du Z pour les données du détecteur (en rouge) et les données *Monte Carlo* corrigées (histogrammes en noirs).

Introduction sur l'algorithme de *Simple Cone*

L'objet de base pour le calorimètre est la cellule (voir la Section 2.3.3). Pour une cellule, on peut mesurer son énergie E_{cell} et son énergie $\mathbf{p}_{cell}$, dont la norme est E_{cell} et la direction est définie à partir de la position du vertex primaire de l'interaction. Pour reconstruire les objets calorimétriques et mesurer leur énergie, on utilise plutôt la tour comme objet de base. Elle est constituée de l'ensemble des cellules de toutes les couches du calorimètre ayant pour coordonnées $(\eta_{tour}, \phi_{tour})$. Une cellule peut avoir une énergie négative due aux fluctuations électroniques autour de sa valeur nominale. Par conséquent, certaines tours peuvent également avoir une énergie négative. Ces dernières sont supprimées par un algorithme appelé *T42* pour les algorithmes de reconstruction décrits dans la suite.

L'algorithme dit de *Simple Cone* est le point de départ de toute reconstruction d'un objet calorimétrique, que ce soit un électron, un photon ou bien encore un jet. L'information de départ de cet algorithme est l'ensemble des tours ayant une énergie supérieure à 0.5 GeV. L'algorithme de *Simple Cone* commence alors par ajouter la tour la plus énergétique à ce qu'on appelle un "pré-amas" et la retire de la liste ordonnée des tours. Il cherche alors dans un cône de rayon R ($R = \sqrt{(\Delta\eta)^2 + (\Delta\phi)^2}$) autour de cette tour une autre tour de la liste et ajoute la plus énergétique au "pré-amas" tout en recalculant le barycentre du "pré-amas". Le barycentre devient alors le centre du "pré-amas" considéré, ce qui implique qu'au final une tour peut être éloignée du centre du "pré-amas" d'une distance angulaire supérieure à R . Ce processus continue jusqu'à ce que toutes les tours de la liste aient été utilisées. On a alors une liste de "pré-amas", point de départ à tous les algorithmes de reconstructions des objets calorimétriques.

Identification des photons et des électrons

Pour les objets électromagnétiques, le rayon de l'algorithme de *Simple Cone* est 0.4. Pour chacun des "pré-amas" précédemment définis, on construit un cône de rayon 0.2 dans le calorimètre central et de 10 cm dans les bouchons autour de la tour la plus énergétique. L'ensemble des tours à l'intérieur de ce nouveau cône forme un candidat électromagnétique. Jusqu'à présent toutes

les couches du calorimètre avaient été utilisées pour définir les candidats électromagnétiques, ce n'est finalement qu'au moment du calcul des variables associées à ces candidats que l'on retire les cellules du calorimètre hadronique exceptée la couche *FH1*. Cet algorithme ainsi défini est bien adapté à la reconstruction des électrons isolés de plus de 15 GeV. Pour les autres cas, c'est-à-dire les électrons de plus basse énergie ou présents à l'intérieur d'un jet, on utilise plutôt la "Méthode de la Route" [112] qui ne sera pas décrite ici.

L'énergie des candidats électromagnétiques doit de plus être corrigée afin de tenir compte de la géométrie du détecteur [113]. En effet, pour remonter à l'énergie initiale du candidat, il faut tenir compte des différents matériaux traversés selon l'angle d'incidence. Finalement, l'échelle d'énergie absolue des candidats électrons est estimée à partir d'événements $Z \rightarrow e^+e^-$ ou J/Ψ pour obtenir un pic de la résonance du boson Z centré sur la valeur connue expérimentalement [114].

Les principales variables choisies pour définir les différents candidats électromagnétiques et qui permettent de discriminer les gerbes produites par des jets ou des taus sont les suivantes :

- La fraction électromagnétique, emf , est la fraction de l'énergie de la gerbe déposée dans les couches électromagnétiques de la gerbe. Le critère utilisé par toutes les définitions d'électrons et pour le photon est $emf > 0.9$. La gerbe d'un jet est en effet beaucoup plus grande et son initiation est plus tardive. Ce critère est donc très discriminant pour les jets.
- L'isolation, iso , est définie par la formule :

$$iso = \frac{E_{tot}(\Delta R < 0.4) - E_{EM}(\Delta R < 0.2)}{E_{EM}(\Delta R < 0.2)} \quad (3.5)$$

où E_{tot} (E_{EM}) est l'énergie totale (électromagnétique) des tours se trouvant à une distance maximale de 0.4 (0.2) du candidat. Le critère utilisé pour un photon est $iso < 0.15$ et varie de $iso < 0.15$ à $iso < 0.2$ pour les différentes définitions des électrons.

- La matrice $H8$, $Hmx8$, est le χ^2 d'une matrice à huit variables qui compare la forme de la gerbe à la forme attendue pour un candidat électromagnétique. Les variables utilisées pour définir cette matrice sont les fractions d'énergie dans les quatre couches, la largeur des gerbes suivant η et ϕ , la position en z du vertex et le logarithme de l'énergie. Il existe également une matrice $H7$ qui est formée des variables précédentes exceptée la largeur selon la coordonnée η . En fonction du type d'électron choisi pour l'analyse, cette coupure peut varier de 0 à 50 sur $Hmx7$ et de 0 à 75 sur $Hmx8$. Elles peuvent également varier du calorimètre central aux bouchons.

En plus de ces variables, on peut ajouter une variable basée sur les informations du trajectographe central afin de séparer les candidats électron des candidats photon. On définit alors la probabilité qu'un objet électromagnétique soit associé spatialement à une trace : $P(\chi^2_{spatial})$. Les photons n'étant pas associés à une trace centrale, on peut imposer une coupure maximum sur cette variable pour leur identification. Finalement, l'information du *CPS* décrit dans la Section 2.3.2 apporte depuis peu [115] un complément pour l'identification des photons [116]. Ce détecteur permet notamment une meilleure description de la forme de la gerbe.

Pour étudier l'efficacité d'identification et de reconstruction des objets électromagnétiques, on utilise les processus $Z \rightarrow e^+e^-$ [117]. De la même façon que pour les muons, on utilise la méthode dite de *Tag And Probe* en anglais. On demande deux électrons l'un avec des critères très lâches (le *Probe*) et l'autre avec des critères beaucoup plus sélectifs (le *Tag*) dont la masse invariante est proche de celle du boson Z . Le premier sert d'électron test, il permet de mesurer les efficacités, le second est nécessaire pour purifier le lot de données. On détermine par cette méthode les efficacités

pour chaque type de coupure sur les variables évoquées précédemment, et ensuite les différences entre les données du détecteur et les données *Monte Carlo*. Ces études permettent alors de définir des critères standard dépendant de la topologie et de la qualité des objets électromagnétiques requise par l'analyse.

3.4.4 Les jets

Un jet n'est pas un objet physique à proprement parlé (voir la Section 2.3.3), comme le sont les muons, les électrons ou les photons. C'est un objet complexe formé du regroupement de particules, angulairement proches, issues de l'hadronisation d'un quark ou d'un gluon dans une gerbe hadronique comme l'illustre la Figure 2.16. Les hadrons, ainsi formés, sont essentiellement des pions ou des kaons. Ils interagissent dans le calorimètre et déposent de l'énergie dans les tours décrites au préalable. Ainsi, un jet est la signature expérimentale des partons produits initialement au moment de la collision.

Les différentes étapes de la formation d'un jet et les différents processus physiques mis en jeu permettent de définir plusieurs niveaux de reconstruction pour un jet. L'énergie déposée dans le calorimètre est l'information directement mesurée. On reconstruit alors les gerbes hadroniques à partir de l'énergie déposée dans les tours. On obtient ce qu'on appelle un jet de calorimètre ou un jet au niveau du détecteur. On peut ensuite déconvoluer les effets du détecteur, comme l'énergie sous-jacente ou bien les non-uniformités du détecteur, pour reconstruire un jet dit de particules. On obtient en fait l'énergie des hadrons formés par l'hadronisation, appelée également la fragmentation. C'est aussi l'énergie que fournit un générateur d'événements *Monte Carlo*. On pourrait enfin remonter à l'énergie du parton mais ce niveau de reconstruction n'est pas utilisé dans l'analyse présentée dans le Chapitre 4. Accéder à l'énergie du parton n'a d'intérêt que dans le cas de reconstruction de masses invariantes. Les études pour ce dernier niveau de reconstruction ne seront pas détaillées dans la suite mais accessibles en [118]. Ces différents niveaux de reconstruction sont schématisés sur la Figure 3.45.

Les algorithmes utilisés par la collaboration DØ

Deux grandes familles d'algorithmes sont disponibles pour reconstruire les jets de calorimètre dans la collaboration DØ [119] : un algorithme de cône [120] décrit en détail ensuite et un algorithme de k_T [121]. Dans le premier, on reconstruit un jet à partir de cônes de tailles "fixes" alors que dans le second la forme de la gerbe n'est pas définie a priori. Le problème du second algorithme est sa forte sensibilité au bruit. Il n'est quasiment plus utilisé au *Run II* à cause de l'augmentation du bruit dans le calorimètre et ne sera donc pas décrite dans la suite. Une comparaison entre ces deux méthodes est disponible en [122]. Une autre modification importante du *Run I* au *Run II* est la façon de calculer les variables du jet une fois celui-ci reconstruit à partir des tours. Au *Run I*, la méthode utilisée était appelée le *Snowmass sheme*. Cette méthode calculait le barycentre scalaire des variables en utilisant l'énergie transverse comme poids pour les entités constituant le jet (les tours). On obtenait par exemple les Formules 3.6 avec i l'indice pour une tour.

$$E_T^{jet} = \sum_i E_T^i, \eta^{jet} = \frac{\sum_i \eta^i E_T^i}{\sum_i E_T^i}, \phi^{jet} = \frac{\sum_i \phi^i E_T^i}{\sum_i E_T^i}. \quad (3.6)$$

Cette procédure avait l'avantage de conserver facilement l'invariance de Lorentz, prérequis pour un algorithme de jet correct. L'inconvénient est qu'elle suppose que tous les jets ont une

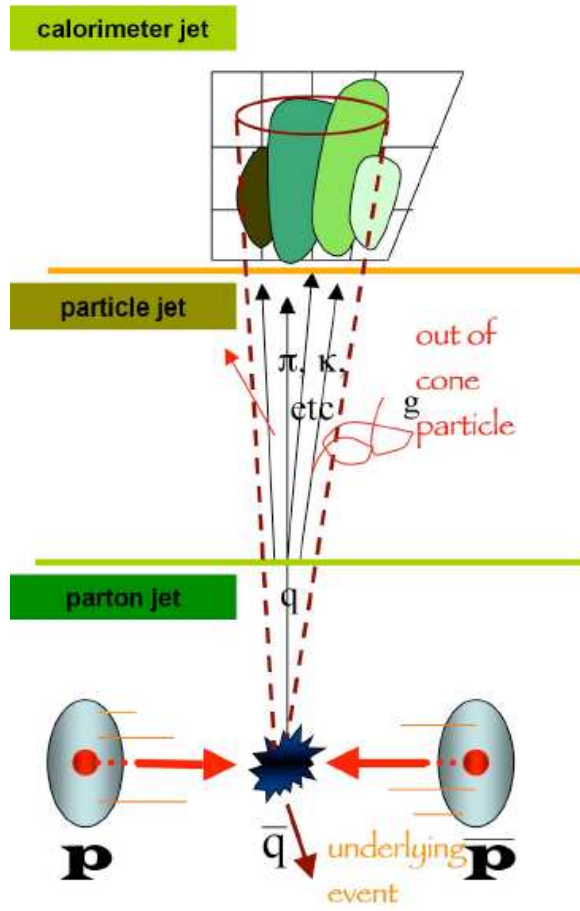


FIG. 3.45 – Schéma représentant les différentes définitions d'un jet selon le niveau de reconstruction.

masse nulle ($M^{jet} \ll E_T^{jet}$). A partir du *Run II*, il a été décidé d'utiliser plutôt la procédure nommée *E-scheme* qui se sert de l'information complète du quadri-vecteur et qui recalcule ensuite chacune des variables à partir de ce quadri-vecteur comme l'illustre les Formules 3.7.

$$P^{jet} = (E^{jet}, \mathbf{p}^{jet}) = \sum_i (E^i, \mathbf{p}^i), y^{jet} = \frac{1}{2} \frac{E^{jet} + p_z^{jet}}{E^{jet} - p_z^{jet}}, \phi^{jet} = \tan^{-1} \left(\frac{p_y^{jet}}{p_x^{jet}} \right). \quad (3.7)$$

Dans ce schéma la rapidité a été préférée à la pseudo-rapidité, la pseudo-rapidité peut être également calculée à partir de l'angle polaire dans l'hypothèse des masses nulles.

Dans la suite, c'est ce deuxième schéma qui sera implicitement utilisé pour toute description d'une variable liée au jet (de calorimètre ou de particules).

Reconstruction et identification des jets

Reconstruction La reconstruction des jets [123–125] décrite ici et utilisée ensuite dans le Chapitre 4 est basée sur plusieurs algorithmes de cône successifs.

Le point de départ de la reconstruction est la formation des "pré-amas" à partir de l'algorithme de *Simple Cone*, détaillé dans la Section 3.4.3, avec un rayon de 0.3. Pour réduire le nombre de faux jets dûs au bruit, un traitement particulier est appliqué pour les tours susceptibles d'être fortement affectées par les bruits du calorimètre. Si la cellule la plus énergétique se trouve dans les *massless gap* des bouchons ou dans la partie grossière (*CH*) du calorimètre, la condition d'utilisation de la tour devient :

$$p_T^i - p_T^c > 0.5 \text{ GeV} \quad (3.8)$$

Avec i l'indice de la tour et c celui de la cellule la plus énergétique.

On ne garde finalement que les "pré-amas" ayant une énergie totale supérieure à 1 GeV et composés d'au moins deux tours pour former la liste ordonnée utilisée en entrée d'un nouvel algorithme de cône, appelé *ILCA* pour *Improved Legacy Cone Algorithm*. Cette deuxième itération fonctionne alors en trois étapes.

La première étape consiste en la formation de nouveaux agrégats de tours appelés "proto-jets" ($Proto_k$ avec k l'indice du "proto-jet" considéré) à partir de la liste ordonnée des "pré-amas" décrites précédemment et des tours du calorimètre. Les "proto-jets" sont des cônes stables de rayon R_{cone} . Deux valeurs pour le rayon sont utilisées par la collaboration DØ : 0.5 et 0.7 (0.5 dans le cas de l'analyse traitée dans le Chapitre 4). Il est à noter que la distance angulaire ΔR est exprimée non plus en fonction de la pseudo-rapacité mais plutôt de la rapidité pour cet algorithme. Tout d'abord, l'algorithme considère un "pré-amas" (A_i avec i son indice) de la liste en partant du "pré-amas" du plus haut E_T . Deux cas se présentent :

- $\Delta R(Proto_k, A_i) > 0.5$ avec $Proto_k$ correspondant à tous les "proto-jets" déjà reconstruits. Dans ce cas, le "pré-amas" i devient à son tour un "proto-jet".
- $\Delta R(Proto_k, A_i) < 0.5$. Pour cette deuxième possibilité, il y a formation d'un nouveau "proto-jet" à partir des tours du "pré-amas" d'indice i et centré initialement sur $Proto_k$. L'algorithme procède alors de nouveau par itérations en ajoutant les tours distantes de R_{cone} au "proto-jet" et en recalculant le centre du nouveau "proto-jet" ainsi formé et ainsi de suite jusqu'à stabilisation de la direction comme l'illustre le schéma de la Figure 3.46. Les critères de stabilisation ou de fin d'itérations sont les suivants :
 - Si l'énergie transverse du candidat "proto-jet" calculée à partir du schéma *E-scheme* est inférieure à 3 GeV, c'est-à-dire la moitié de l'énergie minimum requise pour un jet final, ce candidat est rejeté.
 - Si la distance angulaire entre un candidat "proto-jet" et un candidat "proto-jet" nouvellement construit est inférieure à 0.001. On considère qu'il y a stabilisation et le "proto-jet" est conservé.
 - Si le nombre d'itérations est égal à 50, le "proto-jet" est également conservé.

Une fois un "proto-jet" $Proto_k$ conservé, une dernière vérification est nécessaire pour l'ajouter à la liste finale des "proto-jets". Il faut en effet que celui-ci n'ait pas été trouvé précédemment, c'est-à-dire qu'il n'existe aucun "proto-jet" $Proto_j$ vérifiant :

$$\left| \frac{p_T^k - p_T^j}{p_T^j} \right| < 0.01, \Delta R(Proto_k, Proto_j) < 0.005. \quad (3.9)$$

On obtient finalement à la suite de cette première étape, plus complexe que l'algorithme de *Simple Cone*, une liste de "proto-jets".

La deuxième étape est très proche de cette dernière dans le procédé avec comme liste d'entrée pour l'algorithme non plus les "pré-amas" mais les barycentres en η et ϕ associés à chaque paires

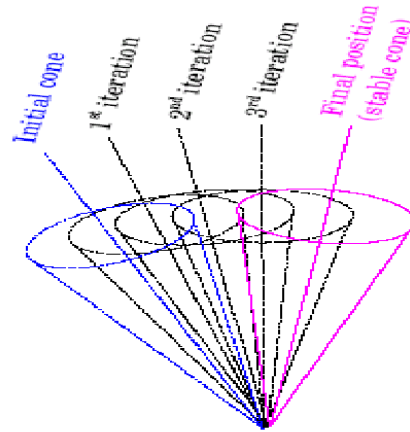


FIG. 3.46 – Schéma illustrant la stabilisation de la direction d'un "proto-jet".

de "proto-jets". Ils sont appelés *Midpoints* en anglais. Ne sont considérées pour constituer ces *Midpoints* que les paires séparées d'au moins R_{cone} et au plus de $2 \times R_{cone}$. Deux conditions sont de plus retirées par rapport à l'algorithme précédent : le critère de séparation entre "proto-jets" les plus proches et la dernière étape retirant les "proto-jets" dupliqués. Cette deuxième étape permet de rendre l'algorithme insensible aux radiations, infrarouges par exemple, c'est une autre condition pour la conception d'un bon algorithme de reconstruction de jets.

Après cette deuxième étape, on a une nouvelle liste de "proto-jets". Il faut alors vérifier que les "proto-jets" n'ont pas de tours en commun, c'est l'étape dite du *Merging and Splitting*. Un dernier algorithme est donc appliqué pour aboutir aux jets finaux utilisés ensuite dans les analyses de physique. La liste de "proto-jets" est ordonnée par ordre décroissant en énergie transverse. Deux cas sont à distinguer :

- Si le "proto-jet" ne partage aucune tour avec ses voisins, il devient un jet final.
- s'il partage des tours avec un "proto-jet" voisin, l'énergie transverse de la partie commune est calculée, ainsi que le rapport entre cette énergie transverse et l'énergie transverse du voisin. Deux nouveaux cas se présentent, comme le montrent les deux schémas de la Figure 3.47 :
 - Si ce rapport est supérieur à 50 %, les deux "proto-jets" fusionnent pour n'en former qu'un et le voisin est retiré de la liste.
 - Si ce rapport est inférieur à 50 %, les tours communes sont réparties à l'un et l'autre des "proto-jets" en fonction de leur distance angulaire avec les centres des deux "proto-jets". On obtient alors deux "proto-jets" ajoutés à la liste.

La liste est alors réordonnée et le processus recommence jusqu'à la fin de la liste.

Après ce processus, seuls les jets ayant une énergie transverse supérieure à 6 GeV sont conservés (ce seuil était auparavant de 8 GeV mais il pénalisait les analyses recherchant de jets de basse énergie transverse dans l'état final).

Identification Les jets reconstruits par les algorithmes de cône précédents peuvent également être des jets de bruit par opposition aux jets dits physiques qui correspondent effectivement à l'hadronisation d'un parton sous forme de gerbe dans le calorimètre. Il est donc nécessaire

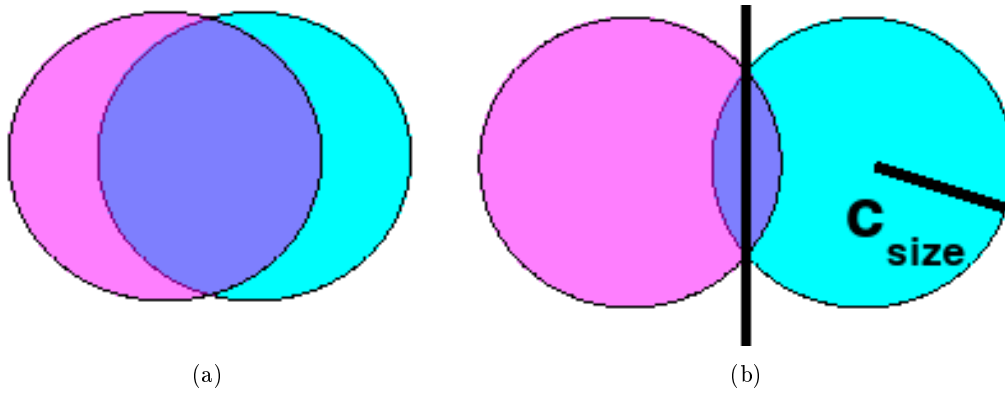


FIG. 3.47 – Fusion des deux "proto-jets" (a) ou séparation (b).

de discriminer ces faux jets ou plutôt de sélectionner les jets physiques à partir de critères de qualité [126]. Ces critères évoluent d'études en études, le but étant d'augmenter l'efficacité de reconstruction des jets physiques tout en diminuant l'efficacité de reconstruction d'un faux jet.

Les variables utilisées pour définir ces critères sont les suivantes :

- $n90$: nombre de tours portant 90 % de l'énergie transverse du jet. Cette variable permet de limiter les jets formés d'une seule tour chaude. La coupure était de $n90 > 1$ et est supprimée depuis des études récentes menées sur l'efficacité d'identification [126].
- $hotf$: rapport entre les énergies transverses des deux cellules les plus énergétiques du jet. Cette variable réduit le nombre de jets constitués de cellules chaudes. La coupure choisie était : $hotf < 10$, pour les mêmes raisons que $n90$, la coupure a récemment été supprimée.
- chf : la fraction du E_T se trouvant dans la partie grossière du calorimètre (CH). La partie CH peut potentiellement contenir des cellules bruyantes, alors qu'un jet dépose principalement son énergie dans les parties EM et FH du calorimètre. La coupure imposée utilisée jusqu'en 2006 était de : $chf < 0.4$. La volonté d'avoir une efficacité d'identification constante en fonction de E_T et de η autour de 98-99 % a conduit à une optimisation de ces constantes.

Un jet passe ce critère de sélection si :

- * $chf < 0.4$, **ou**
- * $chf < 0.6$, $0.85 < |\eta_{det}| < 1.25$ (c'est-à-dire dans la partie $ECMH$) et $n90 < 20$, **ou**
- * $chf < 0.46$, $|\eta| < 0.8$ (le calorimètre central), **ou**
- * $chf < 0.33$, $1.5 < |\eta| < 2.5$ (les bouchons).

La différence entre η_{det} et η est que le premier correspond à la localisation géographique du jet en prenant le centre du détecteur comme origine, alors que le second utilise le vertex primaire de l'interaction principale comme origine.

- emf : la fraction du E_T se trouvant dans la partie électromagnétique du calorimètre (EM). Une coupure minimale sur cette variable est motivée également par la suppression des jets dominés par le bruit hadronique. Une coupure maximale permet de différencier les jets des électrons et des photons. De même que pour la variable précédente, des études récentes, au cours de l'année 2005, ont permis d'optimiser la condition initiale fixée à : $0.05 < emf < 0.95$. La coupure minimale a été modifiée de telle sorte qu'un jet passe ce critère de qualité si :

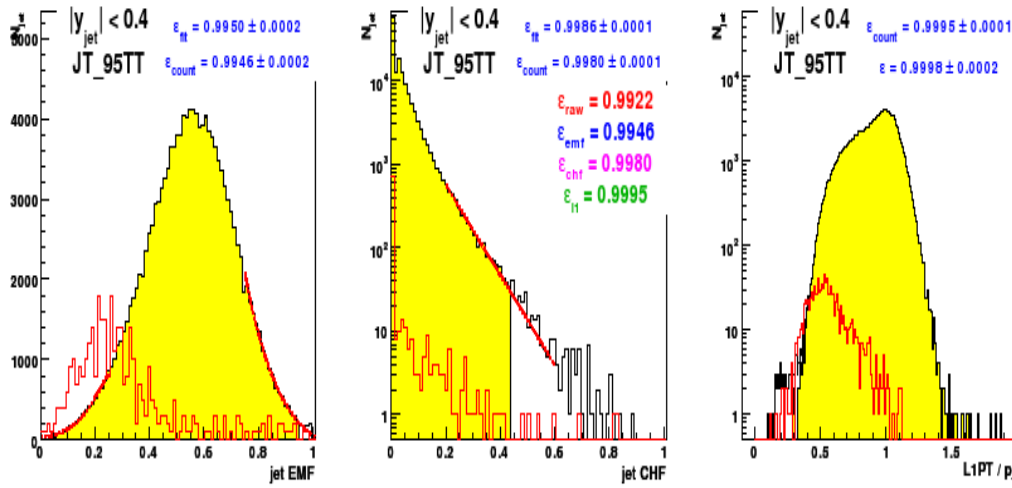
- * $emf > 0.05$, **ou**
 - * $1.3 > ||\eta_{det}| - 1.25| + \max(0.4 \times (\sigma_n - 0.1))$, avec σ_n la largeur du jet selon η , **ou**
 - * $emf > 0.03$, et $1.1 < |\eta_{det}| < 1.4$, **ou**
 - * $emf > 0.04$, et $2.5 < |\eta|$.
- $L1_{ratio}$: cette dernière variable permet de comparer l'énergie transverse du jet reconstruite à partir de deux électroniques de lecture indépendantes, celle de précision et celle de déclenchement à partir des tours du niveau 1 du système de déclenchement. Les bruits de ces deux canaux n'étant pas corrélés, $L1_{ratio}$ est une vérification supplémentaire pour rejeter les jets de bruit. Comme il a été décrit dans la Section 3.2.1, les tours de niveau 1 n'utilisent pas la partie grossière du calorimètre. On définit donc une nouvelle énergie transverse du jet compatible avec cette définition : $E_T^{sansCH} = E_T \times (1 - chf)$. Le ratio est défini comme $L1_{ratio} = \frac{E_T^{sansCH}}{E_T^{L1}}$. Les dernières critères de sélection sur cette variable sont les suivants :
- * $L1_{ratio} > 0.5$, **ou**
 - * $L1_{ratio} > 0.35$, $E_T < 15$ GeV et $1.4 < |\eta|$, **ou**
 - * $L1_{ratio} > 0.1$, $E_T < 15$ GeV et $3 < |\eta|$, **ou**
 - * $L1_{ratio} > 0.2$, $E_T > 15$ GeV et $3 < |\eta|$.

Les efficacités des critères de qualité ont été étudiées en [127]. Les résultats ont été obtenus en utilisant les distributions des trois variables définissant les critères de qualité chf , emf et $L1_{ratio}$ que l'on trouve sur la Figure 3.48 pour un lot de données QCD avec un ou deux jets dans l'état final. Ces résultats ont été obtenus par la méthode dite des "distributions". Sur chacun des graphes représentés sur la Figure 3.48, on a un histogramme jaune représentant les données passant ce critère et en blanc celles ne passant pas ce critère. En traits pleins rouges, la fraction des événements ne passant pas les deux autres critères (multiplié par 10 pour le graphe de gauche représentant la distribution de emf). Et en traits pointillés rouges, la fonction d'ajustement extraite des données dans la zone où le critère de sélection est appliqué. Deux façons de calculer l'efficacité d'une coupure sachant les deux autres passées : en utilisant directement les données et on obtient ϵ_{count} ou en utilisant la fonction d'ajustement pour l'efficacité ϵ_{fit} . L'efficacité finale pour une coupure donnée est ϵ_{fit} si $\epsilon_{fit} > \epsilon_{count}$, ϵ_{count} sinon. Finalement, la valeur ϵ_{raw} correspond au produit des trois efficacités.

Les résultats sont ensuite exprimés en fonction de E_T pour différents intervalles en η . La méthode dite "des distributions" pour les lots monojet et dijet est comparée à la méthode dite de *Tag And Probe* décrite précédemment et appliquée à des lots de données dijet ou γ +jet. A partir des courbes obtenues (voir un exemple sur la Figure 3.49), on peut dériver des efficacités pour les données du détecteur et les données *Monte Carlo* et ainsi corriger les données *Monte Carlo*, si l'efficacité obtenue dans les simulations est mal estimée. La courbe noire de la Figure 3.49 représente la paramétrisation estimée dans cette exemple avec la fonction d'ajustement : $\epsilon(E_T) = \epsilon(0) + a.exp(-b.E_T)$.

Correction de l'énergie des jets : JES

Les jets reconstruits par la méthode décrite dans la Section précédente sont des jets de calorimètre (voir la Figure 3.45). A partir de la gerbe hadronique, on reconstruit un objet et on mesure les propriétés de ce jet en calculant les différentes variables utiles avec le schéma *E-scheme*. L'énergie ainsi mesurée E^{mes} ne tient pas compte des effets de réponse du calorimètre (la réponse du calorimètre est différente pour un pion, un kaon, ...), du bruit de l'uranium, de l'effet des interac-

FIG. 3.48 – Distribution de emf , chf et $L1_{ratio}$.

tions spectatrices et sa mesure peut être biaisée par les caractéristiques de la méthode employée, comme le cône de taille fixe par exemple. Pour tenir compte de tous ces effets et mesurer l'énergie d'un jet de particules, on corrige l'échelle d'énergie des jets appelée *JES* pour *Jet Energy Scale* en anglais [125, 128, 129].

L'énergie d'un jet de particules E^{part} est donnée en fonction de E^{mes} par la Formule 3.12 :

$$E^{part} = \frac{E^{mes} - E^{off}(R_{cone}, \eta, \mathcal{L})}{\mathcal{R}(R_{cone}, \eta, E^{mes}) \mathcal{S}(R_{cone}, \eta, E^{mes})} \quad (3.10)$$

où :

- E^{off} est l'énergie sous-jacente. Elle est due aux bruits de l'uranium et de l'électronique, à l'empilement dû aux croisements précédents et aux interactions spectatrices. Cette énergie est liée à la taille du cône R_{cone} (0.5 ou 0.7), à la position du jet dans le détecteur η et à la luminosité $\mathcal{L}$ dont dépend le nombre d'interactions spectatrices.
- $\mathcal{R}$ est la réponse du calorimètre. Il prend en compte l'énergie perdue dans les zones non instrumentées et la différence de réponse entre les calorimètres électromagnétique et hadronique. Ce terme est un nombre sans dimension inférieur à 1.
- $\mathcal{S}$ est la fraction d'énergie se trouvant à l'intérieur du cône. En effet, une particule se trouvant à l'intérieur de la gerbe hadronique peut avoir déposé de l'énergie hors du cône. A contrario, une particule n'appartenant pas à la gerbe peut avoir déposé de l'énergie dans le cône. Ce nombre est également sans dimension et est inférieur à 1, ce qui signifie que le flux d'énergie se fait principalement de l'intérieur vers l'extérieur du cône.

Energie sous-jacente L'énergie sous-jacente est l'énergie qui est déposée dans le calorimètre ne provenant pas de l'interaction dure mais se trouvant à l'intérieur du cône. Pour estimer, l'énergie à l'intérieur du cône provenant réellement du processus produit par l'interaction dure, il faut estimer puis soustraire cette énergie sous-jacente. On peut diviser les contributions à l'énergie sous-jacente en trois parties :

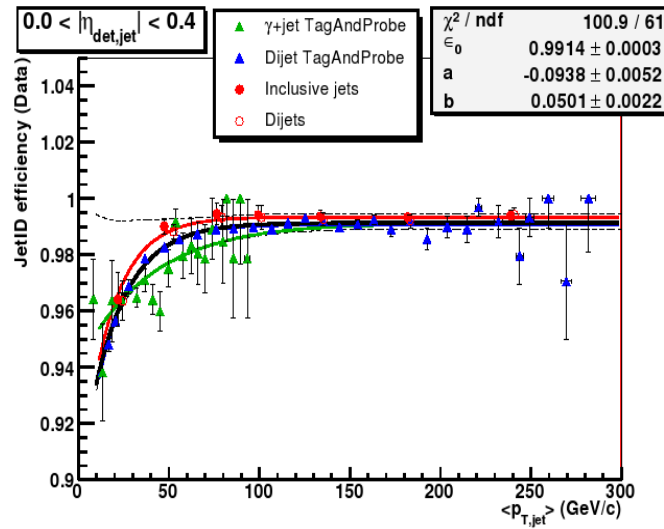


FIG. 3.49 – Evolution de l'efficacité d'identification en fonction de E_T pour $0 < |\eta| < 0.4$ et pour les différentes méthodes de mesure d'efficacité.

- *UE* pour *Underlying Events*. Lors d'une collision proton-antiproton, deux partons interagissent lors de l'interaction dure. Les autres partons ou éventuellement les radiations de gluons peuvent également interagir, c'est ce qu'on appelle l'*Underlying Events* en anglais.
- *NP* désigne le bruit et les effets d'empilement.
- *MI*, pour *Multiple Interaction*, est le terme pour les interactions spectatrices dues aux collisions entre les autres protons et antiprotons. Ce terme est directement lié à la luminosité instantanée.

Dans la collaboration DØ, il a été récemment choisi de ne pas soustraire la contribution *UE*, car elle correspond à un véritable processus physique.

On utilise alors les événements collectés par deux conditions de déclenchement appelés *minimum bias* et *zero bias*. Le premier enregistre les événements ayant déclenché le compteur de luminosité mais n'ayant déclenché aucune autre condition de déclenchement. Ils correspondent par conséquent à une collision inélastique enregistrée par le détecteur DØ. Le second correspond seulement à un croisement de faisceaux sans vertex primaire et sans déclenchement du compteur de luminosité. A partir de la condition *zero bias*, on mesure la contribution *NP*. Ensuite, la différence entre la condition *minimum bias* avec un unique vertex primaire et cette contribution *NP* donne la contribution *UE*. Enfin, pour déterminer la contribution due aux interactions spectatrices, on mesure la différence entre les événements enregistrés par la condition *minimum bias* avec $n + 1$ vertex primaires et les événements avec n vertex primaires ($n > 0$) collectés par cette même condition.

On estime alors la densité d'énergie transverse due à la contribution *NP* et *MI* déposée dans le calorimètre en fonction du nombre de vertex, qui donne une bonne indication de la dépendance en luminosité, et de la variable η : $D_T(n_{vertex}, \eta)$. On exprime alors l'énergie sous-jacente par la relation donnée en Section 3.11.

$$E^{off} = \pi \cdot R_{cone}^2 \times D_T(n_{vertex}, \eta) \times \cosh(\eta) \quad (3.11)$$

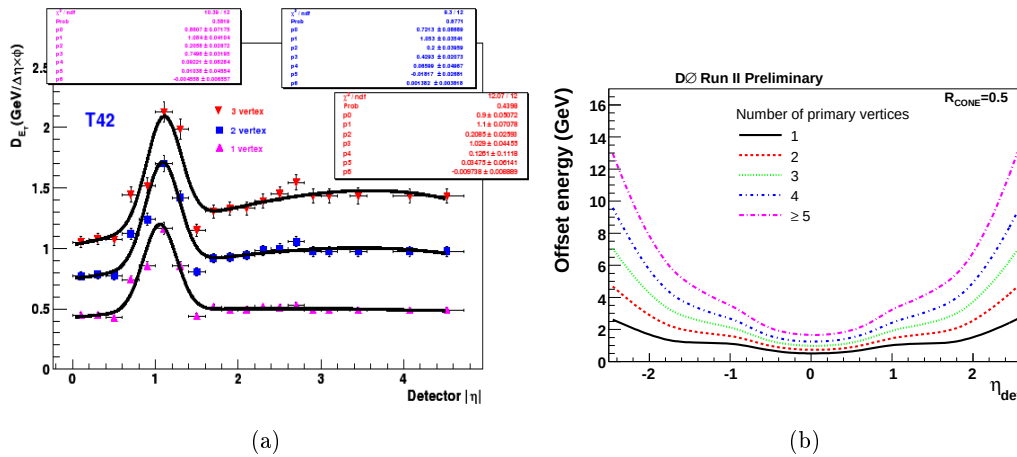


FIG. 3.50 – Densité d'énergie transverse en fonction de η_{det} pour différents nombres de vertex (a) et énergie sous-jacente pour un cône de 0.5 pour différents nombres de vertex également (b).

La densité d'énergie mesurée et l'énergie sous-jacente résultante sont données par les graphes de la Figure 3.50. On constate que la dépendance avec le nombre de vertex est forte à cause des interactions spectatrices. La densité d'énergie transverse devient importante pour une région en η_{det} correspondant à la région inter-cryostat (*ICR*). Cet effet est dû aux poids trop importants attribués aux détecteurs inter-cryostatiques.

Réponse du calorimètre La réponse du calorimètre est la conséquence directe de ses propriétés. Le détecteur a été calibré sous faisceaux test d'électrons et de pions, néanmoins le rapport entre la réponse électromagnétique et la réponse hadronique varie d'une particule à l'autre. On dit que le calorimètre est non-compensé. De plus, les particules produites peuvent également perdre de l'énergie dans les parties non instrumentées comme le solénoïde et le cryostat par exemple. Et finalement, la réponse du détecteur peut varier d'une cellule à l'autre car par construction leur direction varie en fonction de leur position dans le détecteur. La réponse doit donc tenir compte de tous ces effets pour estimer l'énergie d'un jet de particule.

La méthode d'estimation de cette réponse se base sur le principe de la conservation de l'énergie transverse lors d'un processus γ +jet. Si l'énergie électromagnétique du photon est bien connue (l'échelle d'énergie électromagnétique est étalonnée grâce au pic du boson Z, voir la Section 3.4.3), on utilise ainsi l'équilibre en énergie transverse pour déterminer celle du jet. Malheureusement, cet équilibre peut être dégradé, ce qui peut créer de l'énergie transverse manquante. On utilise alors la méthode dite de la fraction projetée de masse transverse manquante (*MPF*), qui utilise la relation 3.12 et qui est illustrée sur la Figure 3.51.

$$\vec{E}_T^\gamma + \mathcal{R} \cdot \vec{E}_T^{jet} = -\vec{\cancel{E}}_T \quad (3.12)$$

$\vec{E}_T^\gamma$ est l'énergie transverse du photon (après application de la réponse électromagnétique). $\vec{E}_T^{jet}$ est l'énergie transverse du jet après soustraction de l'énergie sous-jacente. Et $\vec{\cancel{E}}_T$ est le vecteur énergie transverse manquante calculée à partir des tours et corrigée comme indiquée sur la Formule

3.13 de la réponse électromagnétique.

$$\vec{E}_T = \left(\sum_{tours} \vec{E}_T^{tours} \right) - (1 - R_{em}) \vec{E}_T^{\gamma, \text{ avant correction}} \quad (3.13)$$

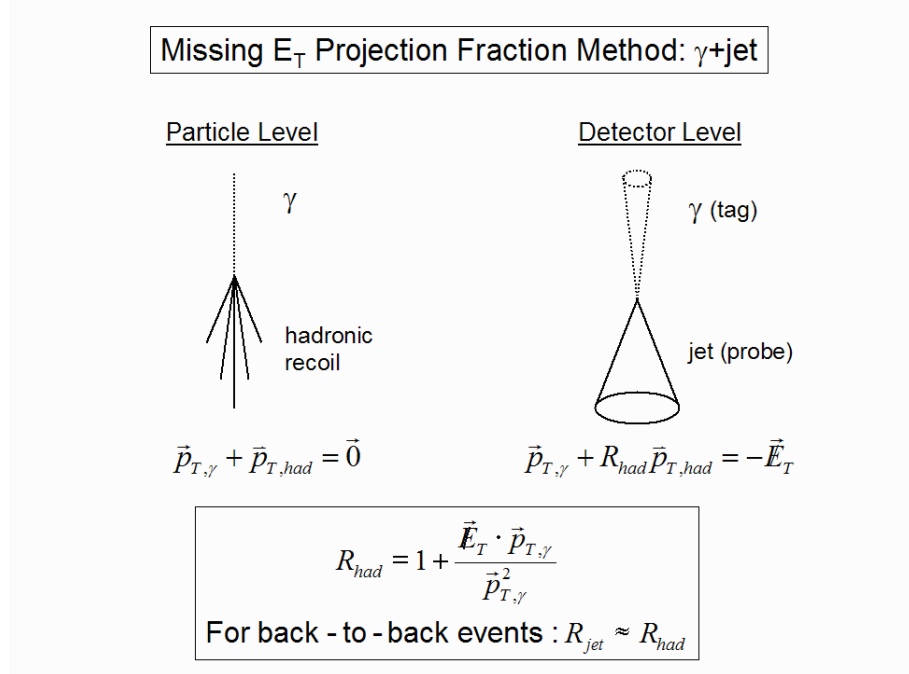


FIG. 3.51 – Schéma de la méthode *MPF*.

On en déduit une Formule pour la réponse du calorimètre :

$$\mathcal{R} = 1 + \frac{\vec{E}_T \cdot \vec{n}_T^\gamma}{E_T^\gamma} \quad (3.14)$$

où $\vec{n}_T^\gamma = \vec{E}_T^\gamma / E_T^\gamma$ est la direction de l'énergie transverse du photon.

Dans la Formule 3.14, aucune grandeur liée à l'algorithme de reconstruction n'est utilisée et on utilise la $\vec{E}_T$ et le E_T du photon déjà corrigés de la réponse électromagnétique.

L'avantage de la méthode est qu'elle ne dépend pas de l'algorithme de reconstruction du jet. Par contre, les photons ayant souvent des énergies transverses faibles, une extrapolation à haute énergie est donc nécessaire pour $\mathcal{R}$. Une autre méthode basée sur des événements dijet peut être utilisée pour pallier cette imperfection.

La première étape de la détermination de la réponse est ce qui est appelé l' *η -intercalibration*. Le but est d'égaliser la réponse en fonction de la pseudo-rapidité et ainsi obtenir une courbe qui paramètre la réponse en fonction de l'énergie du jet uniquement. Pour cela, on utilise la méthode *MPF* avec des lots de données dijets, qui permettent d'avoir une plus grande statistique, et les données de données γ +jet en demandant que le photon ou le jet servant de référence soit dans le calorimètre central ($|\eta| < 0.5$). Les deux lots de données donnant des résultats en très bon accord, ils sont donc combinés pour limiter les incertitudes statistiques et également pour leur

complémentarité. En effet, le lot γ +jet domine à basse énergie transverse alors que le lot dijet est prépondérant à haute énergie transverse. La correction pour égaliser la réponse est finalement donnée en fonction de l'énergie transverse du photon ou jet de référence en utilisant des intervalles fins en pseudo-rapidité (typiquement $\Delta\eta$ 0.1). La Figure 3.52 montre les corrections à appliquer en fonction de η pour différentes valeurs de l'énergie transverse. Ces corrections sont estimées de nouveau après toute la procédure de correction de l'échelle d'énergie des jets, dont les dernières étapes sont détaillées plus loin, afin de vérifier la cohérence de cet étalonnage en fonction η (voir graphe b de la Figure 3.52). Les variations sont alors inférieures à 2 % et la réponse est centrée sur 1.

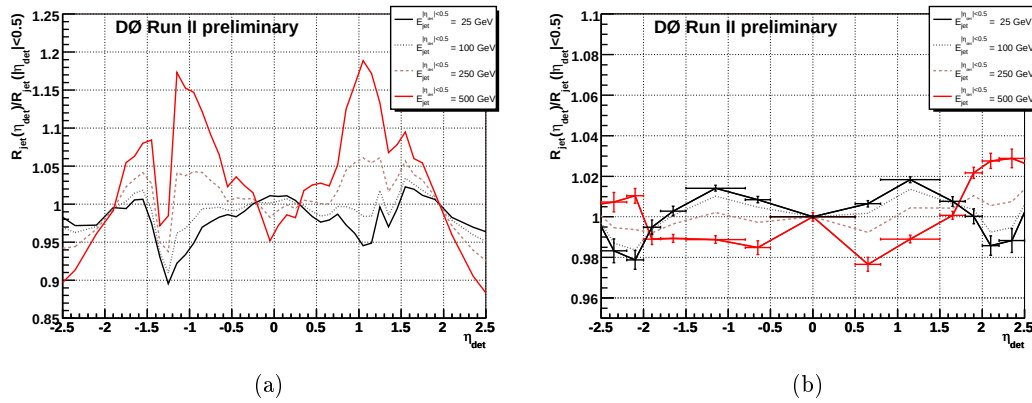


FIG. 3.52 – Correction relative de la réponse en fonction de η avant l' η -intercalibration (a) et après toutes les corrections *JES* (b).

La seconde étape pour estimer la réponse en fonction de l'énergie du jet est effectuée après soustraction de l'énergie sous-jacente et après l' η -intercalibration à partir du lot de données γ +jet. Les critères de sélection du photon sont très stricts, de plus un jet seulement est requis. Ce jet doit être dos à dos avec le photon. La réponse est donnée pour différentes régions en η , ce qui permet de vérifier l'effet positif de l' η -intercalibration. La valeur de la réponse est finalement donnée sur la Figure 3.53. On constate que cette réponse est inférieure à 1 et augmente avec l'énergie. De plus, on observe que quelle que soit la région en η étudiée la réponse est cohérente avec la région centrale.

Fraction de l'énergie dans le cône Deux effets peuvent intervenir pour cette dernière correction. Le premier effet est physique, il correspond par exemple aux radiations de gluons déposant leur énergie dans le cône mais ne provenant pas de la fragmentation du parton. Le deuxième est dû au détecteur. Il corrige les effets instrumentaux qui affectent le développement de la gerbe, qui n'est donc plus complètement contenue dans un cône de taille fixe. C'est seulement cet effet dû au détecteur qui est corrigé par le terme $\mathcal{S}$ de la Formule 3.12. L'étude est effectuée sur des événements γ +jet collectés par le détecteur et également simulés par des générateurs. Le premier lot de données permet d'évaluer l'effet physique et l'effet du détecteur, alors que le second peut être utilisé pour évaluer l'effet physique uniquement. Une soustraction donne finalement l'effet dû au détecteur qui est utilisé pour estimer le nombre $\mathcal{S}$. Le nombre $\mathcal{S}$ est en fait le rapport de

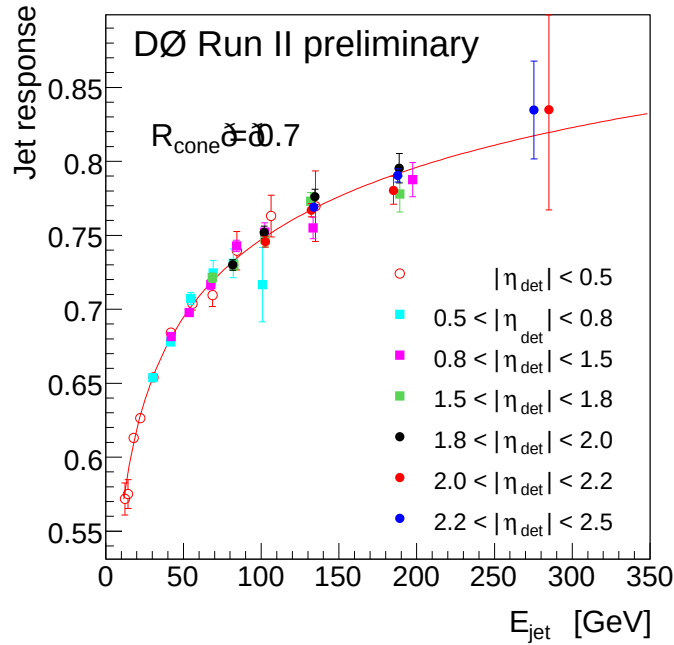


FIG. 3.53 – Réponse pour un cône de 0.7 en fonction de l'énergie du jet partiellement corrigée.

l'énergie du jet dans le cône par l'énergie du jet dans un cône de rayon R_{limite} . Ce deuxième rayon dépend de la position dans le calorimètre et donc de η_{det} . Il varie de 1.0 à 1.6 et a été déterminé lors du *Run I* [130]. Les énergies qui sont utilisées dans ce rapport ont au préalable été retranchées de la ligne de base, qui est une énergie résiduelle due aux bruits et aux effets d'empilement. Le schéma de la Figure 3.54 illustre le calcul du rapport $\mathcal{S}$ (pour plus de détails sur l'estimation de la ligne de base et du facteur correctif se référer à [123]).

Un exemple de la correction due à la fuite d'énergie hors du cône est représentée sur la Figure 3.55.

Les incertitudes systématiques sur la JES Les incertitudes systématiques dues à la correction de l'échelle d'énergie des jets sont des paramètres très importants pour les analyses de physique. La plus grande partie du travail du groupe responsable de ces corrections est la réduction de ces incertitudes systématiques qui sont souvent les plus importantes dans les analyses de DØ, bien plus grandes que les incertitudes statistiques. Ces erreurs sont de plus propagées à toutes les variables calculées à partir de l'énergie des jets, comme les sommes scalaire et vectorielle des E_T , les variables topologiques entre jets et même l'énergie transverse manquante par exemple. Dans une analyse comme celle présentée dans le Chapitre 4 avec une topologie multijet et de l'énergie transverse manquante, la systématique sur la JES a des répercussions directes sur les résultats finaux. La contribution majoritaire à cette incertitude est l'incertitude due au calcul de la réponse. Cette incertitude est importante à bas E_T à cause de l'incompréhension des queues non gaussiennes de la distribution de la réponse. Elle est également importante à haut E_T pour des raisons statistiques essentiellement. C'est pour cette dernière raison que l'incertitude augmente

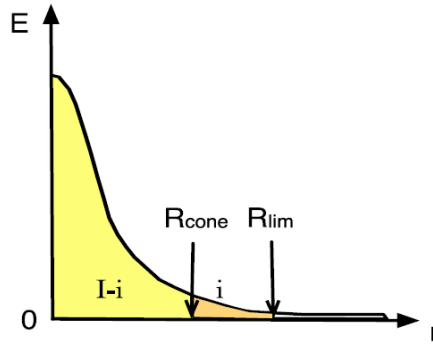


FIG. 3.54 – Représentation schématique de la distribution de l'énergie utilisée pour le calcul de $\mathcal{S}$.

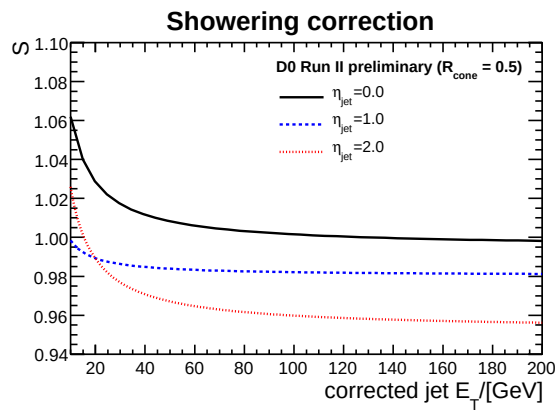


FIG. 3.55 – Facteur correctif $\mathcal{S}$ pour un cône de rayon 0.5 en fonction de l'énergie corrigée du jet pour différente valeur de η .

également à haut E_T pour la correction de la fuite d'énergie. Le résumé de ces incertitudes en fonction de l'énergie non corrigée des jets pour différentes valeurs de η est montré sur les graphes de la Figure 3.56 pour un cône de rayon 0.5.

Corrections spécifiques aux données Monte Carlo

Les données du détecteur et les données *Monte Carlo* ne sont pas parfaitement compatibles pour la reconstruction et l'identification des jets. La complexité des processus physiques mis en jeu lors d'une collision, la difficulté de simuler parfaitement la réponse du détecteur à une interaction $p\text{-}\bar{p}$ et finalement la difficulté de simuler les imperfections du détecteur rendent impossible une description parfaite des données enregistrées par le détecteur DØ. Il est donc nécessaire d'appliquer des corrections aux données *Monte Carlo* après la simulation des générateurs pour approcher cette description complète. Ces corrections sont déterminées par comparaison entre les données du détecteur et les données *Monte Carlo*. Pour les jets, plusieurs corrections sont indispensables. Il faut corriger la différence de résolution de l'énergie, l'échelle d'énergie (*JES*) elle-même et l'efficacité de reconstruction puis d'identification, qui est généralement plus faible dans les données du détecteur. Toutes ces corrections sont appliquées en même temps de façon consistante et en

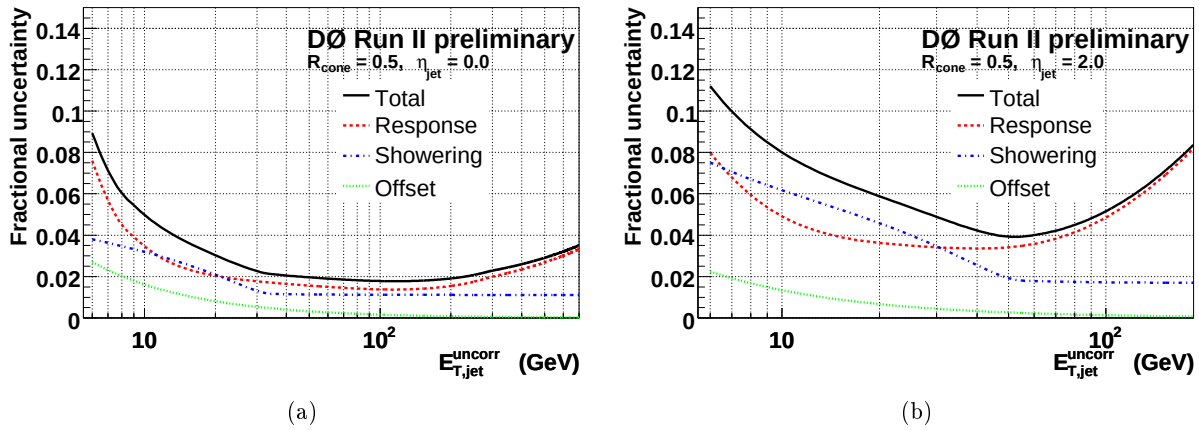


FIG. 3.56 – Incertitudes sur la *JES* en fonction de l'énergie non corrigée des jets de cône de rayon 0.5 pour $|\eta_{det}| = 0.0$ (a) et pour $|\eta_{det}| = 2.0$ (b).

tenant compte des corrélations entre elles grâce à la méthode *SSR* pour *Smearing Shifting and Removing* en anglais [131, 132]. Le terme *Smearing* désigne une dégradation de la résolution des jets simulés pour la faire correspondre à celles des jets effectivement collectés par le détecteur. Le mot *Shifting* correspond à la modification de l'échelle d'énergie des jets par un décalage de l'énergie transverse des jets simulés. Et le terme *Removing* est utilisé pour décrire la suppression de certains jets simulés pour diminuer l'efficacité de reconstruction et d'identification.

Une fois toutes les corrections sur l'échelle absolue d'énergie des jets (*JES*) appliquées aux données du détecteur et aux données *Monte Carlo*, on utilise encore une fois la conservation de l'énergie transverse dans le cas de lot de données $Z + \text{jet}$ et $\gamma + \text{jet}$ pour étudier les différences entre les données réelles et les simulations. Pour les lots de données utilisées, on demande de plus que le jet et le photon ou le boson Z soient dos à dos. On définit alors une variable ΔS donnée par la Formule 3.15.

$$\Delta S = \frac{E_T^{jet} - E_T^\gamma}{E_T^\gamma} \quad (3.15)$$

L'énergie transverse du jet E_T^{jet} est corrigée et on suppose connue l'énergie transverse du photon E_T^γ (ou du boson Z dans le cas du lot de données $Z + \text{jet}$). L'avantage de cette méthode est qu'elle estime en une fois l'incertitude systématique totale des corrections à effectuer sur les données simulées plutôt que de les ajouter quadratiquement. Cette méthode a donc le grand avantage de réduire les incertitudes systématiques, ce qui a des conséquences immédiates sur toutes les analyses de physique ayant des jets dans l'état final.

La stratégie employée est la suivante :

1. On calcule ΔS pour différents intervalles en E_T du photon ou du boson Z .
2. On ajuste les distributions de ΔS trouvées par une fonction d'ajustement judicieusement choisie.
3. On extrait des paramètres des fonctions d'ajustement l'information nécessaire pour corriger les données simulées.

Deux exemples de distributions de ΔS sont représentés sur les graphes de la Figure 3.57 pour un lot de données réelles γ +jet. En trait noir, on distingue les fonctions d'ajustement qui correspondent à ces deux distributions.

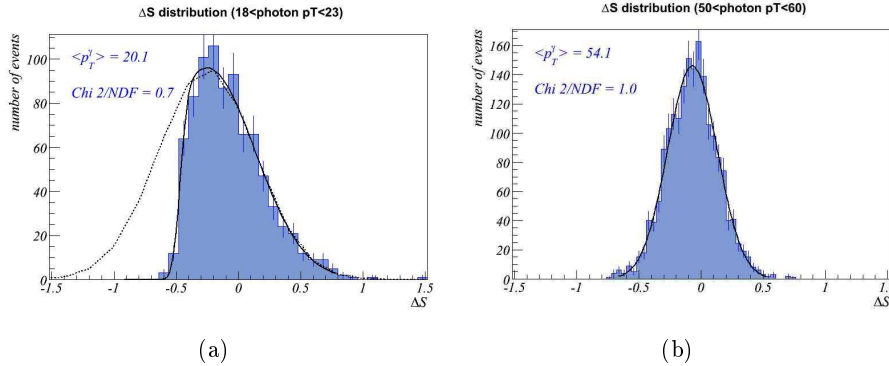


FIG. 3.57 – Distribution de ΔS pour un intervalle $13 < E_T^\gamma < 26$ (a) et pour un intervalle $50 < E_T^\gamma < 60$ (b).

En fonction du E_T du photon ou du boson Z, deux types de fonctions d'ajustement sont utilisées. La première correspond aux hautes énergies transverses, c'est-à-dire supérieures à 26 GeV pour les données simulées et supérieures à 45 GeV pour les données réelles, ce qui correspond au graphe de droite de la Figure 3.57 :

$$g(< \Delta S >) = N \times e^{-\frac{(< \Delta S > - < \Delta S >_0)^2}{2\sigma^2}} \quad (3.16)$$

La seconde est utilisée pour les énergies inférieures aux seuils donnés plus tôt et donc au graphe de gauche de la Figure 3.57 :

$$f(< \Delta S >) = N \times (1 + \text{erf}(\frac{< \Delta S > - \alpha}{\beta\sqrt{2}})) \times e^{-\frac{(< \Delta S > - < \Delta S >_0)^2}{2\sigma^2}} \quad (3.17)$$

Le nombre entier N des Formules 3.16 et 3.17 permet de normaliser les fonctions d'ajustement. Le terme $< \Delta S >_0$, la valeur moyenne des gaussiennes, donne l'information sur l'échelle relative de l'énergie des jets. Ce terme doit, idéalement, être à 0. La largeur des gaussiennes, σ , permet d'accéder à l'information sur la résolution du jet. Enfin, les deux paramètres α et β de la fonction erreur fournissent l'information sur le produit des efficacités de reconstruction et d'identification.

Cette stratégie est employée dans un premier temps pour le lot de données γ +jet pour trois régions en η (calorimètre central, région inter-cryostat et les bouchons du calorimètre). On extrait alors les paramètres des fonctions d'ajustement pour chacune des trois régions. Le lot de données γ +jet, qui a une statistique suffisante pour une étude dans chacune des trois régions, souffre de problèmes d'impureté dus à des mauvaises identifications des photons (notamment un jet identifié comme étant un photon). Les bosons Z des lots de données Z+jet ne rencontrent pas ce problème, néanmoins la section efficace du processus Z+jet est trop faible pour une étude par régions de η . On utilise donc une comparaison systématique des deux lots sur tout le calorimètre et on en déduit une correction pour les paramètres déterminés avec le premier lot. Dans les faits, la correction n'est utile que pour l'échelle relative d'énergie des jets. Les paramètres de dégradation de la résolution

sont en effet identiques dans les deux lots de données. On a finalement un lot de paramètres pour chacune des régions en η .

Les graphes de la Figure 3.58 montre l'exploitation des paramètres des fonctions d'ajustement pour appliquer ensuite les corrections nécessaires aux données simulées. Le premier graphe à gauche donne la différence entre les largeurs des gaussiennes, c'est-à-dire la différence de résolution entre les données du détecteur et les données *Monte Carlo*, en fonction du E_T du photon. On utilise cette différence pour dégrader les résolutions des données simulées en modifiant le E_T corrigé du jet. On multiplie cette énergie transverse par un nombre tiré aléatoirement dans une gaussienne centrée sur 1 et de largeur $\sqrt{(\sigma_{data}^2 - \sigma_{MC}^2)}$. Une fois cette dégradation effectuée, le graphe du centre de la Figure 3.58 permet d'accéder à l'échelle relative d'énergie des jets.

On calcule alors la différence entre les valeurs moyennes de ΔS pour les données réelles et les simulations pour accéder à l'échelle relative d'énergie des jets $\mathcal{D}$:

$$\mathcal{D} = \langle \Delta S \rangle_{data} - \langle \Delta S \rangle_{MC} \quad (3.18)$$

où $\langle \Delta S \rangle_{data}$ est la valeur moyenne de ΔS pour les données réelles et $\langle \Delta S \rangle_{MC}$ pour les données simulées. Ce nombre $\mathcal{D}$ permet de plus de vérifier que l'échelle d'énergie absolue est bien appliquée. En effet, si la *JES* est correctement estimée et appliquée aux données (réelles ou simulées), ce nombre doit être nul.

Pour corriger la différence relative, on modifie donc le E_T du jet à l'aide d'un facteur multiplicatif dérivé grâce à la paramétrisation représentée en bleu sur ce graphe. Enfin, le dernier graphe à droite montre que l'efficacité atteint approximativement un plateau à partir de 13 GeV. Il a donc été décidé d'appliquer la procédure *SSR* pour des jets de $E_T > 13$ GeV pour avoir une efficacité de reconstruction et d'identification constante en fonction du E_T du jet corrigé. La différence d'efficacité est corrigée ensuite en retirant aléatoirement des jets dans les données *Monte Carlo*, ces efficacités ayant été étudiées par ailleurs comme l'illustre par exemple la Figure 3.49.

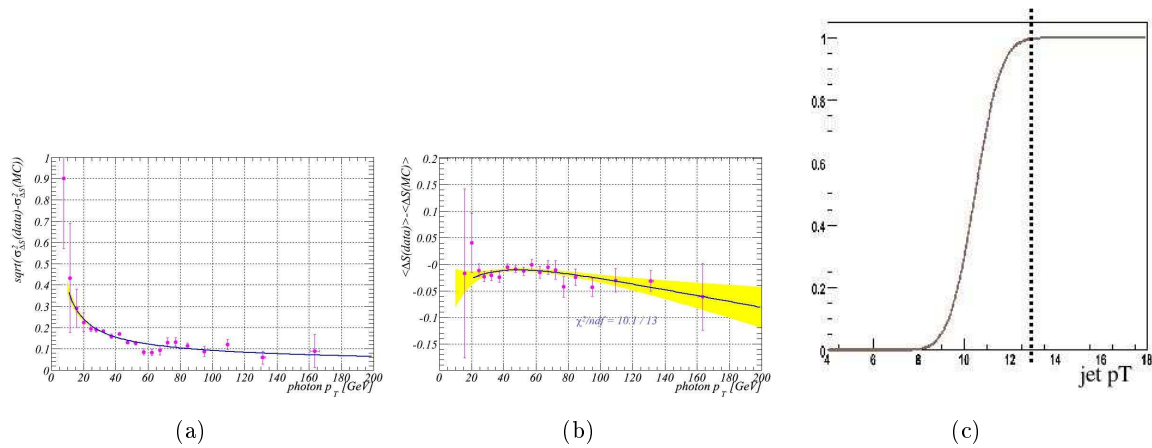


FIG. 3.58 – Différence entre les résolutions : $\sqrt{(\sigma_{data}^2 - \sigma_{MC}^2)}$ (a) , différence entre les centres des gaussiennes (b) et exemple de fonction erreur (c).

On peut finalement vérifier la consistance de ces corrections en traçant la variable $\mathcal{D}$ définie dans la Formule 3.18. Les résultats avant et après les corrections sont résumés sur la Figure 3.59.

La procédure de correction de l'échelle relative réduit donc fortement les différences entre les données simulées et les données réelles. De plus, les graphes de la Figure 3.60 montrent que ces corrections permettent d'obtenir un meilleur accord entre la simulation et les données enregistrées par le détecteur.

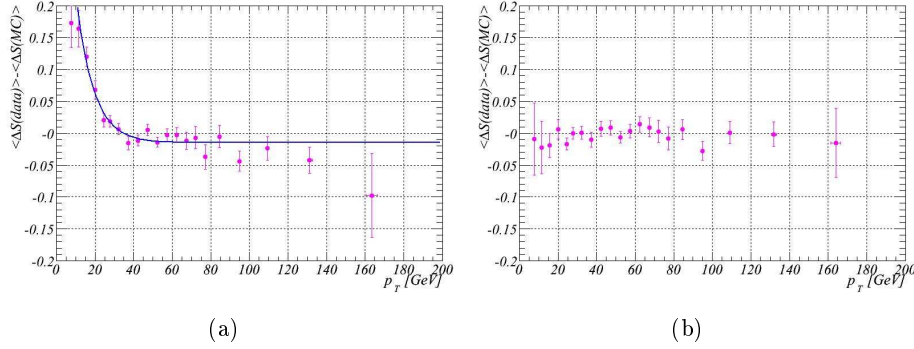


FIG. 3.59 – Variable $\mathcal{D}$ en fonction du E_T corrigé avant (a) et après les corrections de la procédure SSR (b).

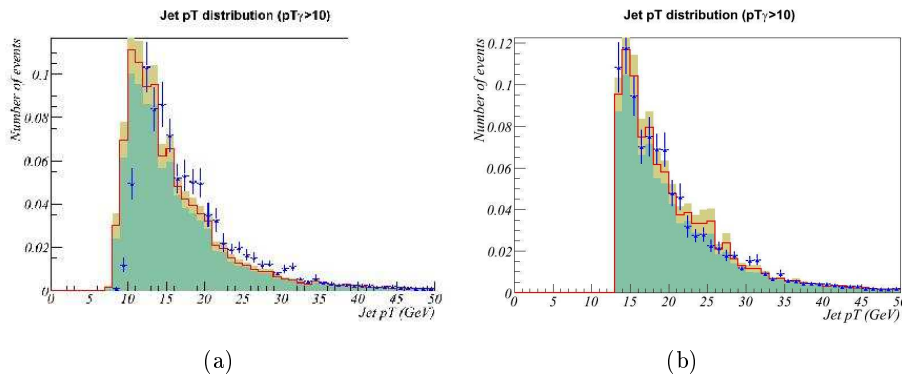


FIG. 3.60 – Energie transverse des jets avant (a) et après les corrections de la procédure SSR (b) pour les données réelles (points avec les incertitudes statistiques) et les données simulées (histogrammes avec les incertitudes systématiques).

Les incertitudes systématiques introduites par cette méthode peuvent être dues à la sélection des événements utilisée pour la procédure SSR , dues aux incertitudes sur l'énergie transverse du photon, aux dépendances en η et finalement aux différences entre les lots de données γ +jet et Z +jet. Ces incertitudes combinées avec les incertitudes statistiques dépendent de l'énergie transverse du jet et augmentent fortement aux hautes valeurs à cause du manque de données comme le montre la Figure 3.61. De 40 à 100 GeV, les incertitudes sont inférieures à 2 % dans le calorimètre central et peuvent atteindre 4 % dans les bouchons pour ce même intervalle en énergie transverse. Les jets, que l'on trouve typiquement dans l'analyse du Chapitre 4, se trouvent dans cet intervalle.

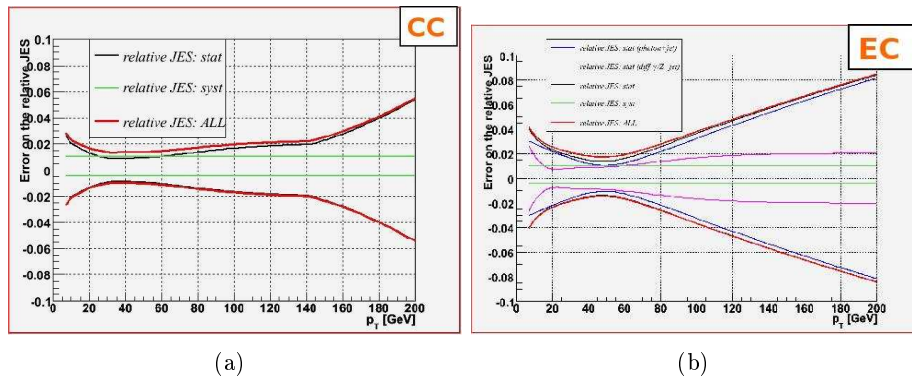


FIG. 3.61 – Incertitudes systématiques et statistiques pour la procédure *SSR* appliquée dans le calorimètre central (a) ou appliquée dans les bouchons (b).

Cette procédure est appliquée à toutes les analyses de physique recherchant des jets dans l'état final. De plus, elle a des effets immédiats sur toutes les variables ayant pour ingrédients les jets, comme par exemple l'énergie transverse manquante décrite ensuite dans la Section 3.4.5. Elle permet donc d'avoir un bon accord entre les jets simulés et les jets réellement observés dans les données et ce malgré la complexité physique d'une gerbe hadronique. Cette complexité a été évoquée tout au long de cette section : de la reconstruction de la gerbe, en passant par l'identification des jets pour finir sur l'évaluation de leur énergie. Enfin, cette procédure, dernière étape dans la compréhension des jets, permet une estimation des systématiques en tenant compte des corrélations, ce qui permet de ne pas surévaluer ces dernières et ce qui a un impact positif non négligeable sur tous les résultats de physique.

3.4.5 L'énergie transverse manquante

L'énergie transverse manquante est la variable qui permet de signer les particules échappant à la détection. Ces particules sont neutres électriquement et n'ont pas de charge de couleur. Elles ne laissent donc aucune trace dans le spectrographe central, aucun dépôt d'énergie dans le calorimètre et échappent au détecteur à muons. Finalement, elles n'interagissent que par interaction faible avec la matière. On peut citer par exemple les neutrinos pour le modèle standard. On peut citer également l'exemple du neutralino le plus léger qui est la particule supersymétrique la plus légère et qui est stable dans de nombreux modèles supersymétriques (voir le Chapitre 1 pour les détails théoriques et le Chapitre 4 pour des détails expérimentaux sur cette particule).

On utilise l'énergie transverse manquante plutôt que l'énergie manquante, car seule la partie transverse de l'énergie totale dans le centre de masse est connue lors d'une collision proton-antiproton. Celle-ci est nulle. La composante selon l'axe z des deux partons interagissant ne peut en effet être connue. Pour calculer cette variable, il faut utiliser l'énergie transverse mesurée par l'ensemble des sous-détecteurs sachant que idéalement si toutes les particules étaient détectées cette énergie serait nulle. On mesure alors deux composantes : la composante x $\cancel{E}_{Tx}$ et la composante y $\cancel{E}_{Ty}$. L'énergie transverse manquante est donnée par la Formule :

$$\cancel{E}_T = \sqrt{(\cancel{E}_{Tx})^2 + (\cancel{E}_{Ty})^2} \quad (3.19)$$

La première étape du calcul est la somme de toutes les cellules du calorimètre exceptées celles des couches hadroniques grossières trop bruyantes pour être directement prises en compte. Dans la Formule 3.20, i représente à la fois la composante x et la composante y .

$$\cancel{E}_{T_i}^{cal} = - \sum_{\text{cellules sauf CH}} p_i \quad (3.20)$$

Ensuite, chacune des corrections des objets physiques évoqués dans les sections précédentes est propagée au calcul de $\cancel{E}_T$. On propage ainsi les corrections des objets électromagnétiques, les photons et les électrons, et on définit : $d\cancel{E}_{T_i}^{EM} = p_i(\text{corrige}) - p_i(\text{non corrige})$. De la même façon, on définit la correction $d\cancel{E}_{T_i}^{JES}$, où l'on calcule la différence entre les jets corrigés par la JES et les jets non corrigés en ne prenant en compte que la correction liée à la réponse définie dans la Formule 3.12 :

$$d\cancel{E}_{T_i}^{JES} = -p_i(\text{non corrige}) \times (1 - \mathcal{R}) \quad (3.21)$$

Les cellules de la partie CH du calorimètre n'étant pas comptabilisées dans la Formule 3.20, il faut ajouter une nouvelle correction à $\cancel{E}_T$ en utilisant uniquement les cellules de la partie CH utilisées pour la reconstruction des jets : $d\cancel{E}_{T_i}^{CH} = -\sum_{\text{jets}} p_i^{CH}$. Cette procédure évite d'introduire une trop forte influence du bruit dans le calcul. Après toutes ces corrections on obtient :

$$\cancel{E}_{T_i}^{corrCALO} = \cancel{E}_{T_i}^{cal} + d\cancel{E}_{T_i}^{EM} + d\cancel{E}_{T_i}^{JES} + d\cancel{E}_{T_i}^{CH} \quad (3.22)$$

Pour les données simulées, une étape supplémentaire tenant compte de la correction liée à l'échelle relative de l'énergie des jets (SSR) est nécessaire. Elle est calculée de la même façon que les corrections précédentes par simple différence entre l'énergie transverse des jets corrigés et l'énergie transverse des jets non corrigés.

Cette définition de l'énergie transverse manquante est celle qui sera utilisée dans le Chapitre 4. Elle correspond en bonne approximation à l'énergie transverse des particules non détectées si

la topologie recherchée ne contient pas de muons isolés dans l'état final. Dans le cas contraire, il faut ajouter une dernière correction pour prendre en compte l'énergie déposée par les muons (de qualité *medium* associés à une trace) dans le calorimètre, qui a été comptabilisée dans la Formule 3.20. Cette correction, $d\cancel{E}_{T_i}^{MU}$, est alors ajoutée à la Formule 3.22.

L'énergie transverse manquante est une variable très importante pour les analyses de physique recherchant la supersymétrie (dans les modèles où la R-parité est conservée, voir le Chapitre 1). Elle fait partie des variables discriminantes pour le bruit de fond, elle permet également d'estimer le bruit de fond multijet (*QCD*) comme il sera détaillé dans le Chapitre 4. Enfin, l'énergie transverse manquante étant calculée à partir de toutes les corrections des objets physiques, sa distribution permet de bien contrôler les algorithmes de reconstruction d'une part et la compréhension des données en les comparant aux simulations *Monte Carlo* d'autre part.

3.4.6 L'étiquetage des jets de hadrons beaux

Introduction

L'identification des quarks *b* [133, 134] est une partie importante pour de nombreuses analyses de physique. On peut citer en exemple les analyses étudiant les caractéristiques du quark top, la recherche du boson de Higgs si sa masse n'est pas trop élevée (c'est-à-dire si la désintégration en deux bosons *W* n'est pas prépondérante) ou encore l'analyse présentée dans le Chapitre 4. La différenciation des jets issus de l'hadronisation d'un quark *b* des autres jets peut donc permettre de mieux isoler les processus précédents pour une recherche et une étude plus efficace.

Une propriété importante d'un hadron beau, utilisée pour son identification, est sa durée de vie τ de l'ordre de 1.6 ps ($c\tau \approx 400\mu\text{m}$). Cette durée de vie est suffisamment élevée pour permettre à un quark *b* d'énergie 40 GeV le parcours d'une distance d'environ 3 mm dans le détecteur avant de se désintégrer. Cette distance n'est pas négligeable par rapport aux distances caractéristiques du détecteur de traces et implique la possibilité d'identification d'un vertex secondaire déplacé par rapport au vertex primaire de l'interaction $p\bar{p}$. Autrement dit, les traces associées au jet de quark *b* ne convergent pas vers le vertex primaire de l'événement, i.e. elles possèdent un paramètre d'impact *IP* élevé par rapport aux traces issues du vertex primaire. On définit le paramètre d'impact comme le périégée de la trajectoire d'une particule avec le vertex primaire comme le montre le schéma de la Figure 3.62. Les caractéristiques du paramètre d'impact sont utilisés par l'algorithme d'étiquetage

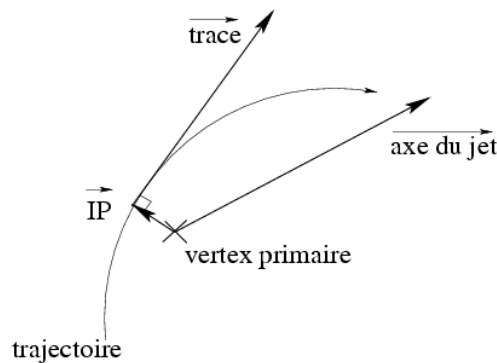


FIG. 3.62 – Schéma du paramètre d'impact *IP*.

des hadrons beaux. La masse du quark b , $m_b \approx 4.7$ GeV, permet également de le différencier des autres quarks plus légers (excepté le quark top). En effet, cette masse implique une masse invariante du vertex secondaire plus élevée que la masse d'un vertex primaire. Elle implique également une fragmentation plus dure, ce qui a pour conséquence un cône de jet plus ouvert et une plus grande impulsion transverse relative à l'axe du jet pour les particules associées à un tel jet : un plus grand p_T^{rel} (voir l'exemple de la Figure 3.63). Une dernière propriété utile à la

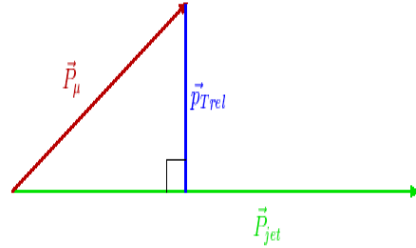


FIG. 3.63 – Schéma du E_T relatif à l'axe du jet pour un muon issu de la fragmentation.

différenciation d'un quark b par rapport aux autres quarks est la désintégration semi-leptonique plus élevée. Dans 20 % des cas, le quark b se désintègre en produisant un muon ou un électron et comme il a été évoqué plus tôt une coupure sur le p_T^{rel} du muon ou de l'électron peut améliorer l'étiquetage. Cette technique est essentiellement employée pour la désintégration en muon, car l'identification des électrons dans le cône d'un jet n'est pas aisée. L'utilisation de ces propriétés particulières aux hadrons beaux est détaillée dans la suite pour l'algorithme d'étiquetage. Il est important de noter que le temps de vol d'un quark c ($c\tau \approx 100\mu m$) peut également engendrer dans une moindre mesure un vertex déplacé, de même qu'un lepton τ se désintégrant hadroniquement ($c\tau \approx 87\mu m$). Ces deux particules peuvent donc contribuer à diminuer l'efficacité d'étiquetage.

Deux étapes sont nécessaires pour l'étiquetage d'un jet. La première consiste à sélectionner les jets sur lesquels peut s'appliquer l'algorithme : les jets *taggable* ou étiquetable. Ensuite, l'algorithme d'étiquetage à proprement parlé est appliqué sur ces jets uniquement.

Définition et paramétrisation de la taggabilité

Un jet doit satisfaire un minimum de conditions pour pouvoir lui appliquer l'algorithme d'étiquetage décrit dans la section suivante. s'il satisfait ces conditions, il est dit *taggable* ou étiquetable [135]. Les jets considérés sont des jets reconstruits avec l'algorithme de cône décrit dans la Section 3.4.4 avec un rayon de cône de 0.5 corrigé de la *JES*. On demande de plus que ces jets satisfassent les conditions suivantes :

- $E_T > 15$ GeV
- $|\eta| < 2.5$

Pour pouvoir appliquer l'algorithme, il faut pouvoir reconstruire un vertex déplacé et donc utiliser l'information des traces associées au jet. Ainsi, on demande qu'un jet de traces chargées de rayon de cône 0.5 soit reconstruit à une distance $\Delta R < 0.5$ du jet en question. Le jet de traces chargées est un objet physique reconstruit de la façon suivante :

- On utilise le même principe que pour tous les algorithmes de cône en partant des traces d'impulsion $p_T > 1$ GeV. Ces traces sont classés par ordre décroissant et constituent ce qu'on appelle les *graines*.
- On ajoute ensuite itérativement aux graines l'ensemble des traces vérifiant $|\Delta z(\text{trace} - \text{graine})| < 0.2$ cm et $\Delta R(\text{trace}, \text{graine}) < 0.5$. A chaque itération, la direction de l'objet, appelé *graine*, est recalculée.
- Si la direction du jet de traces chargées se stabilise, on demande finalement qu'il soit constitué d'au moins deux traces.

De plus, les traces utilisées par l'algorithme de reconstruction du jet de traces chargées doivent vérifier :

- $p_T > 0.5$ GeV,
- Au minimum un impact dans le détecteur *SMT*, c'est-à-dire dans les tonneaux ou les disques F (voir la Section 2.3.2),
- $|IP_z| < 0.4$ cm (voir la Figure 3.62 pour la définition du vecteur $\vec{IP}$),
- $|IP_{xy}| < 0.2$ cm,

Un jet pouvant être associé avec un tel objet physique est dit étiquetable. On peut donc lui affecter une probabilité d'être étiqueté.

Pour les données *Monte Carlo*, la difficulté de bien simuler les traces implique l'impossibilité de reconstruire correctement un jet de traces chargées. Autrement dit, la différence d'efficacité de reconstruction des traces pour un événement donné est trop importante entre les données simulées et les données réelles pour appliquer un tel algorithme. En effet, il est très délicat de simuler le bruit dans le *SMT*. On préfère donc une paramétrisation de la *taggabilité* à partir des données réelles et en fonction de variables bien simulées dans les événements *Monte Carlo*, c'est-à-dire le E_T du jet, sa position dans le détecteur (η, ϕ) et la position selon l'axe z du vertex primaire. La paramétrisation peut dépendre de la topologie des événements, il est donc souhaitable de la réévaluer pour chaque type analyse une fois la topologie bien définie. Deux exemples de paramétrisation seront détaillés dans le Chapitre 4 pour deux topologies différentes.

Description de l'algorithme

L'algorithme d'étiquetage appliqué pour cette thèse est basé sur un réseau de neurones [136–138] construit à partir de trois algorithmes précédemment utilisés par la collaboration $D\bar{O}$:

1. *JLIP* [139] pour *Jet Lifetime Probability* en anglais.
2. *CSIP* [140] pour *Counting Signed Impact Parameters*.
3. *SVT* [141] pour *Secondary Vertex Tagging*.

Un réseau de neurones est une combinaison de variables auxquelles on affecte un poids. Les poids sont définis de telle façon à ce que la variable finale du réseau de neurones soit égale à 1 dans le cas de l'objet recherché (un signal physique, une particule ou dans le cas présent un jet de b) et égale à 0 dans le cas contraire (le fond instrumental ou physique du signal recherché ou un jet de saveur légère dans le cas de l'étiquetage des hadrons beaux). Pour étiqueter un jet comme étant un jet issu d'un quark b , il suffit finalement d'imposer une coupure inférieure sur cette variable de sortie. Plus la coupure est élevée plus l'efficacité d'identification diminue mais dans le même temps plus l'efficacité d'étiqueter un jet de saveur légère comme étant un jet de quark b est faible. L'augmentation de la valeur de la coupure permet donc d'améliorer la sélectivité. Différentes valeurs de la coupure définissent finalement différents points de fonctionnement, qui

seront détaillés dans la suite. Une description plus approfondie sur la construction d'un réseau de neurones peut être trouvée en [142]. Les détails de l'implémentation de cette méthode et de la construction de l'algorithme sont résumés en [143].

Les trois algorithmes précédemment nommés utilisent l'information sur les traces chargées pour étiqueter les hadrons beaux. Le premier, *JLIP*, combine les valeurs des paramètres d'impact de toutes les traces associées au jet étudié pour former la probabilité *JLIP Prob* que ces traces proviennent toutes du vertex primaire de l'interaction principale. Plus cette probabilité est faible, plus le jet a de chance d'être issu de l'hadronisation d'un quark *b*. Différents points de fonctionnement sont définis en fonction de la coupure sur cette probabilité. Le second algorithme, *CSIP*, compte le nombre de traces associées à un jet ($\Delta R < 0.5$) qui ont une signification du paramètre d'impact, $\mathcal{R}_{IP} = \frac{IP}{\sigma_{IP}}$, élevée par rapport au vertex primaire. σ_{IP} est l'erreur sur le paramètre d'impact. Pour qu'un jet soit étiqueté, il faut qu'il soit associé à 2 traces de signification supérieure à 3 ou trois traces de signification supérieure à 2. Le dernier algorithme, *SVT*, se sert des traces suffisamment éloignées du vertex primaire pour reconstruire un vertex secondaire. Un jet est considéré comme étiqueté si un vertex secondaire est proche ($\Delta R < 0.5$) de celui-ci.

Deux lots de données simulées ont permis d'entraîner l'algorithme de réseau de neurones. Par entraînement d'un réseau de neurones, on entend définir l'algorithme de telle façon à ce que la valeur finale approche 1 pour le lot contenant des quarks *b*. Ce lot correspond à des processus de *QCD* avec une production de paires de quarks *b* dans l'état final. Le second lot correspondant à des processus de *QCD* avec des saveurs légères (*u,d,s*) dans l'état final permet dans le même temps d'obtenir une valeur de sortie proche de 0 pour les jets le constituant. La première étape consiste à sélectionner un jeu de variables qui différencie les deux lots [137]. La seconde étape consiste à entraîner le réseau de neurones construits à partir du jeu de variables grâce aux deux lots de données simulées. Sept variables de ces trois algorithmes ont été choisies pour construire le réseau de neurones :

- *SVT DLS* est la signification de la distance de vol, c'est-à-dire la distance entre le vertex secondaire et le vertex primaire divisée par l'erreur sur cette distance. Si aucun vertex secondaire est trouvé, cette valeur est fixée à 0. C'est la variable la plus discriminante de l'algorithme *SVT*.
- *CSIP Comb* est une combinaison linéaire de quatre variables de l'algorithme *CSIP*. Chacune de ces variables est un nombre entier correspondant aux nombres de traces vérifiant des conditions sur le paramètre d'impact. L'avantage d'une combinaison linéaire entre quatre variables différentes est de fournir plus de possibilités de valeur pour la variable utilisée. Une variable continue est en effet préférable pour un réseau de neurones. Un deuxième avantage est la réduction du nombre de variables utilisées pour le réseau de neurones qui gagne ainsi en simplicité.
- *JLIP Prob* est l'unique variable de l'algorithme *JLIP* retenue.
- *SVT χ^2_{dof}* est le χ^2 par nombre de degré de liberté du vertex secondaire.
- *SVT N_{tracks}* est le nombre de traces utilisées pour reconstruire le vertex secondaire.
- *SVT Mass* est la masse du vertex secondaire.
- *SVT Num* est le nombre de vertex secondaire autour du jet.

Ces sept variables sont représentées sur les graphes de la Figure 3.64 pour les deux lots de données utilisés pour l'entraînement du réseau de neurones, ainsi que pour un lot de données réelles correspondant à des événements de *QCD* de saveurs légères (majoritairement car la section efficace de production est plus importante pour les quarks légers) ou lourdes.

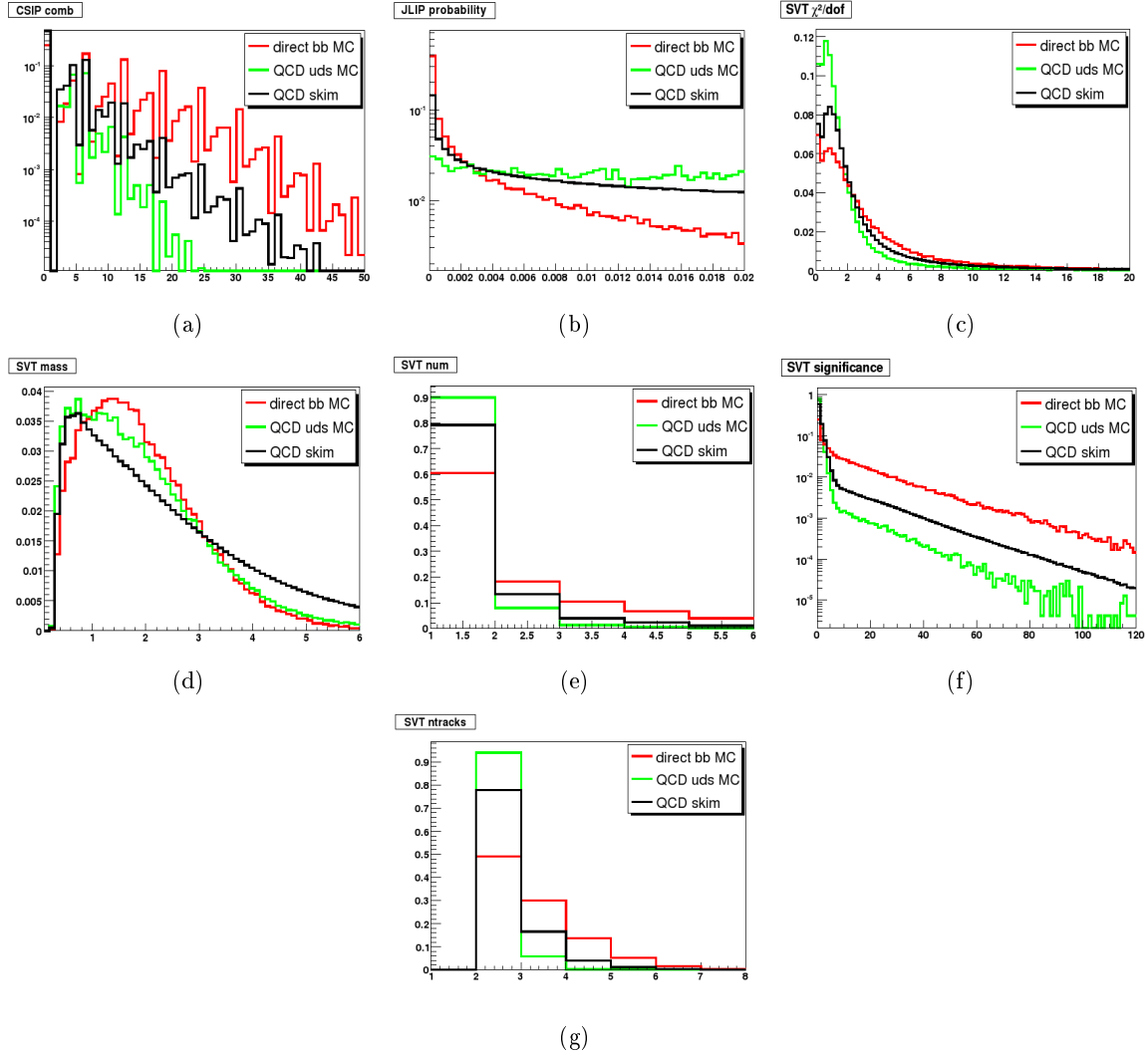


FIG. 3.64 – Distribution des variables *CSIP Comb* (a), *JLIP Prob* (b), *SVT χ^2_{dof}* (c), *SVT Mass* (d), *SVT Num* (e), *SVT DLS* (f) et *SVT N_{tracks}* (g) pour deux lots de données *Monte Carlo QCD* correspondant à la production de saveurs légères ou de paires de quark b et pour un lot de données réelles correspondant à des processus de *QCD*.

La variable de sortie du réseau de neurones pour les deux lots de données simulées est donnée sur la Figure 3.65. On constate que cette variable est bien comprise entre 0 et 1 et qu'elle permet surtout d'apporter une différenciation entre les deux de données simulées meilleure que n'importe laquelle des variables utilisées par les trois algorithmes décrits plus tôt. Cette distribution montre toute la force d'un réseau de neurones et justifie son utilisation pour augmenter l'efficacité d'étiquetage.

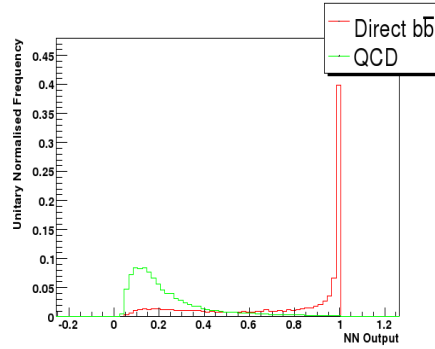


FIG. 3.65 – Variable de sortie du réseau de neurones pour deux lots de données *Monte Carlo QCD* correspondant à la production de saveurs légères ou de paires de quark b .

Pour des analyses comme celle présentée dans le Chapitre 4 ou encore celle détaillée en [83], l'algorithme *JLIP* était préféré aux deux autres. La mise en place du réseau de neurone a largement amélioré les performances de l'étiquetage des hadrons beaux par la collaboration $DØ$. La comparaison de l'ancien algorithme et de l'algorithme actuellement utilisé est représentée sur la Figure 3.66 pour plusieurs points de fonctionnement.

On constate par exemple que pour une efficacité donnée d'étiqueter un jet léger comme étant un jet de quark b , l'efficacité d'identifier un jet de quark b a augmenté de 29 %. La sélectivité de l'algorithme s'en trouve grandement améliorée et par conséquent la sensibilité des analyses utilisant l'étiquetage des hadrons beaux. De plus amples détails sur les efficacités du réseau de neurones sont donnés en [138].

Finalement, 12 points de fonctionnement ont été définis avec une coupure sur la variable de sortie représentée sur la Figure 3.65 allant de 0.1 à 0.925. L'analyse présentée dans le Chapitre 4 utilise trois d'entre eux : le point dit *LOOSE* avec une coupure à 0.45, le point *MEDIUM* avec une coupure à 0.65 et le point *TIGHT* avec une coupure à 0.775.

Application de la méthode d'étiquetage

Plusieurs méthodes sont possibles pour étiqueter les données simulées, en revanche une seule méthode existe pour les données réelles. Pour ces dernières, on utilise directement les informations des détecteurs de traces et du calorimètre pour construire pour chaque jet un booléen correspondant à la *taggabilité* et pour calculer la variable de sortie du réseau de neurones pour les jets dits *taggable* (booléen égal à 1). En fonction du point de fonctionnement choisi, on coupe sur la variable de sortie et on sait si le jet est étiqueté ou non. Ainsi, pour un point de fonctionnement donné, on est en mesure de déterminer quels jets de l'événement sont étiquetés et lesquels ne le sont pas. Pour les données *Monte Carlo*, les imprécisions dans la reconstruction des traces rendent

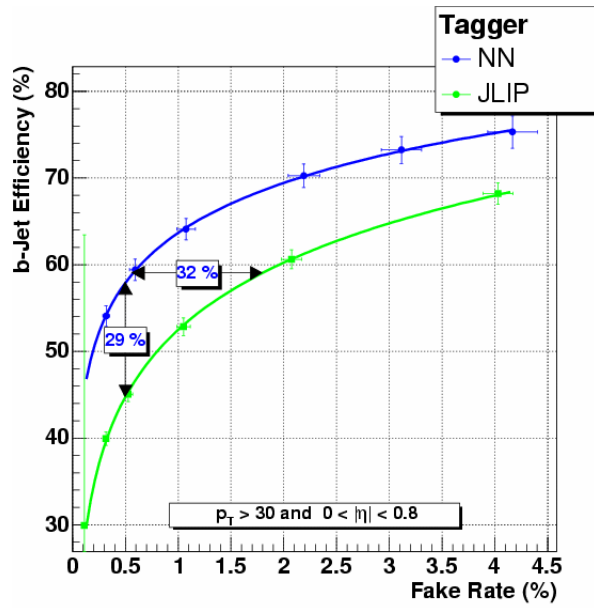


FIG. 3.66 – Comparaison entre l’algorithme *JLIP* et le réseau de neurones pour plusieurs points de fonctionnement pour un jet de $E_T > 30$ GeV dans le calorimètre central.

plus complexe la méthode. Tout d’abord, on ne sait pas si un jet est *taggable* ou non, on connaît uniquement la probabilité qu’il le soit à partir de ces caractéristiques (calculées à partir du calorimètre uniquement). On définit alors une probabilité par jet : $\mathcal{P}_{taggable}$. La méthode utilisée ensuite dans le cadre de cette thèse pour étiqueter un jet est basée également sur une probabilité estimée en fonction des caractéristiques du jet (E_T et η). Cette probabilité, appelée *TRF* pour *Tag Rate Function* en anglais, est donnée par le groupe responsable des algorithmes d’étiquetage et dépend du point de fonctionnement. Le même groupe fournit aussi les incertitudes systématiques sur ces paramétrisations. Finalement, on définit une probabilité pour que le jet provienne de l’hadronisation d’un quark b : $\mathcal{P}_{tagged} = \mathcal{P}_{taggable} \times TRF$. Il existe également d’autres méthodes où le réseau de neurone est appliqué directement sur les données *Monte Carlo* et corrigé ensuite par des facteurs d’échelle obtenus en comparant les résultats des données réelles et des données simulées. Cette méthode ne sera pas utilisée dans le cadre de cette thèse. Ces facteurs d’échelle sont également calculés en fonction des caractéristiques du jet. Pour un événement simulé, il n’est donc pas possible de savoir quel jet est étiqueté. Il n’est pas possible non plus de construire simplement des variables topologiques utilisant des jets de quark b . On peut seulement affecter par événement un poids correspondant au nombre de jet étiqueté requis.

L’application de cette méthode pour les données réelles et les données simulées, ainsi que des comparaisons après étiquetage seront données dans le Chapitre 4.

3.4.7 Conclusion

Cette partie décrit les différents algorithmes et méthodes utilisés par la collaboration DØ pour étudier les particules ou objets physiques, constitués eux-mêmes de particules, détectés par l'ensemble des sous-détecteurs. Elle décrit également la variable utilisée, $\cancel{E}_T$, pour caractériser celles qui échappent à la détection. Les algorithmes permettent de reconstruire des objets physiques, de les identifier ou les étiqueter, puis de leur affecter une énergie corrigée de tous les biais éventuellement introduits par les interactions entre les particules et le détecteur. Cette étape est indispensable à l'étude ou la recherche de processus physique. Les méthodes ou algorithmes utilisés dans le chapitre suivant ont été plus largement détaillés, c'est-à-dire les jets, l'énergie transverse manquante et l'étiquetage des hadrons beaux. Ces trois points sont les points fondamentaux d'une topologie finale comprenant quatre jets de b et de l'énergie transverse manquante.

Chapitre 4 Recherche de paires de gluinos se désintégrant dans un état final $4b + mE_T$

Sommaire

4.1	Introduction et motivations	172
4.2	Caractéristiques du signal	172
4.2.1	Le cadre théorique	172
4.2.2	Production de paires de gluinos au TeVatron	173
4.2.3	Désintégration en un état final $4b + mE_T$	174
4.2.4	Simulation du signal	175
4.3	Les lots de données utilisés	177
4.3.1	Le lot de données du détecteur	178
4.3.2	Les lots de données <i>Monte Carlo</i>	179
4.4	Sélection des événements et traitement des données <i>Monte Carlo</i>	188
4.4.1	La stratégie de l'analyse	188
4.4.2	Le bruit fond instrumental <i>QCD</i>	195
4.4.3	La confirmation des jets par le trajectographe interne	201
4.4.4	Coupures sur le nombre de leptons et zones de contrôle du bruit de fond physique	211
4.4.5	L'étiquetage des jets des quarks <i>b</i>	216
4.4.6	Coupures finales avant l'optimisation	226
4.5	Optimisation finale	229
4.5.1	Méthode statistique	229
4.5.2	Incertitudes systématiques	230
4.5.3	Détermination des masses limites	230
4.5.4	Contours d'exclusion dans le plan $(m_{\tilde{g}}, m_{\tilde{b}_1})$	239

4.1 Introduction et motivations

Ce chapitre est consacré à la recherche d'un signal supersymétrique au TeVatron à l'aide du détecteur DØ, dispositif expérimental qui a été décrit dans le Chapitre 2. Le modèle théorique sur lequel s'appuie cette étude est le modèle standard supersymétrique minimal (*MSSM*) avec R-parité conservée dont les principaux fondements théoriques ont été donnés dans la Section 1.4.4. De ce modèle, on apprend que dans un collisionneur hadronique les particules supersymétriques sont toujours produites par paires par deux voies différentes :

1. les partons issus des protons et antiprotons interagissent par interaction faible et donnent des charginos, des neutralinos ou des sleptons ;
2. les partons interagissent par interaction forte pour engendrer des squarks et des gluinos.

Dans un collisionneur hadronique, à masses égales, la section efficace de production de squarks et gluinos est toujours supérieure à celle de production de charginos, neutralinos ou sleptons. Néanmoins, les signatures des états finaux sont toujours plus complexes à isoler dans un environnement riche en jets.

Dans cette analyse, on s'intéressera à la production de gluinos par paires. De précédentes analyses ont cherché des gluinos soit produits par paires soit avec un squark [84, 144, 145]. Ces analyses sont toutes basées sur des coupures topologiques pour isoler le signal du bruit de fond du modèle standard et n'utilisent pas la possibilité d'étiqueter des saveurs lourdes dans l'état final. Pour l'analyse présentée dans ce chapitre, on étudie, pour compléter les résultats précédents, la possibilité de désintégration du gluino en cascade dans un état final constitué de quarks *b* :

$$p\bar{p} \rightarrow \tilde{g} + \tilde{g} \rightarrow \tilde{b}_1 + \bar{\tilde{b}}_1 + b \rightarrow 2b + 2\bar{b} + 2\tilde{\chi}_1^0$$

L'état final comporte quatre jets de quark *b* et deux neutralinos $\tilde{\chi}_1^0$, c'est-à-dire la particule supersymétrique la plus légère dans ce modèle (*LSP*). Cette dernière est stable et échappe à la détection. Les événements possédant des neutralinos $\tilde{\chi}_1^0$ seront caractérisés par une forte énergie transverse manquante mesurée par le détecteur. On utilise ainsi une signature très discriminante par rapport au bruit de fond pour rechercher des paires de gluinos. De cette façon, on bénéficie d'une part des sections efficaces relativement élevées de production de particules supersymétriques par interaction forte, et d'autre part d'une signature très particulière et rare dans les bruits de fond du modèle standard.

Cette analyse utilise l'ensemble des données collectées lors du *Run IIa*, c'est-à-dire une luminosité intégrée approximativement de 1 fb^{-1} .

Dans un premier temps, les caractéristiques du signal seront détaillées de la production des gluinos à l'obtention de l'état final. La description des lots de données utilisés pour cette analyse sera donnée dans un second temps. La succession des coupures de rejet du fond instrumental *QCD* et le traitement des données *Monte Carlo* permet ensuite de définir un lot de données bien compris en vue d'une optimisation des coupures finales de sélection.

4.2 Caractéristiques du signal

4.2.1 Le cadre théorique

Cette analyse se place dans le cadre du modèle standard supersymétrique minimal. Le but est d'avoir un modèle suffisamment prédictif tout en limitant le nombre de paramètres contraints.

Ainsi, les résultats sont donc facilement interprétables dans un modèle très contraint comme le modèle *mSUGRA* par exemple.

4.2.2 Production de paires de gluinos au TeVatron

Les deux modes de production de gluinos par paires au TeVatron [40, 146, 147] sont la fusion de gluons (voir les diagrammes de Feynman de la Figure 4.2) et l'annihilation quark-antiquark (voir les diagrammes de Feynman de la Figure 4.1). Ce deuxième mode est dominant car les énergies des faisceaux de protons et d'antiprotons ne favorisent pas suffisamment les interactions entre gluons de la mer (appelée également nuage gluonique par ailleurs). La situation est inversée dans le cas du LHC comme il a été évoqué dans la Section 2.2.1.

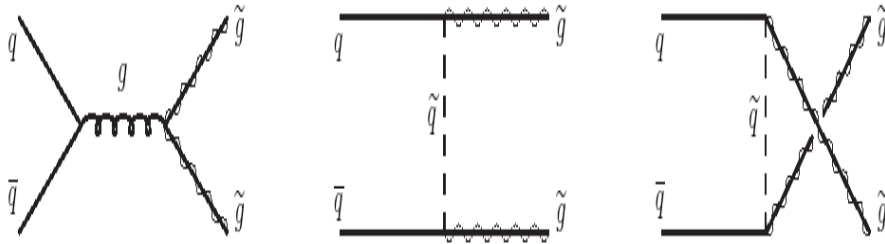


FIG. 4.1 – Diagrammes de Feynman à l'arbre pour la production de paires de gluinos avec un quark et un anti-quark dans l'état initial.

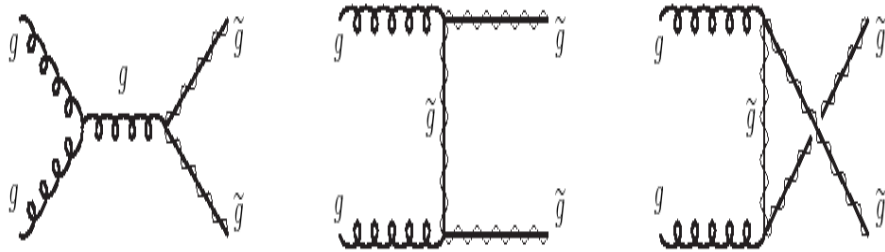


FIG. 4.2 – Diagrammes de Feynman à l'arbre pour la production de paires de gluinos avec deux gluons dans l'état initial.

La section efficace de production de gluinos par paire dépend essentiellement de deux paramètres libres du modèle *MSSM* : la masse du gluino et la masse des squarks de premières générations ($\tilde{u}$, $\tilde{d}$, $\tilde{s}$, $\tilde{c}$). A masse du squark fixé, la section efficace de production décroît exponentiellement en fonction de la masse du gluino. L'effet de la masse du squark intervient dans les diagrammes de Feynman d'annihilation quark-antiquark avec échange d'un squark virtuel. Ces diagrammes sont en interférence destructrice avec le diagramme de gauche de la Figure 4.1. Les squarks virtuels sont essentiellement de premières générations, car ils se couplent avec les quarks du modèle standard largement dominés par les quarks légers. L'interférence étant destructrice, plus la masse du squark est élevée, plus la section efficace de production de paires de gluinos est

importante comme l'illustre la Figure 4.3. Sur cette Figure sont tracées les sections efficaces à l'arbre (LO pour *Leading Order*) de production pour différentes masses de squarks de première génération. Dans la suite, on supposera que leurs masses sont dégénérées. Les résultats de cette analyse, c'est-à-dire les limites sur les masses du gluino et du sbottom le plus léger, seront données pour différentes masses de squarks de premières générations. Ces masses varient de 400 GeV, masse minimale exclue en référence [84], jusqu'à 1 TeV.

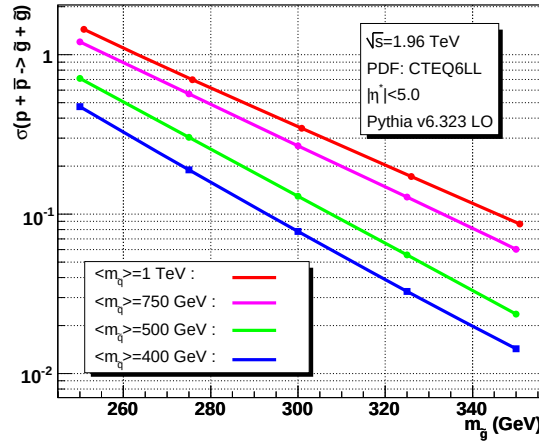


FIG. 4.3 – Sections efficaces (LO) de production de paires de gluinos pour différentes masses de squarks de première génération.

4.2.3 Désintégration en un état final $4b + mE_T$

Désintégration des gluinos

La seule possibilité de désintégration pour le gluino est une désintégration due à un couplage fort. Le gluino ne peut donc se désintégrer qu'en un squark sur couche de masse ou virtuel. Si la désintégration à deux corps est cinématiquement autorisée, c'est-à-dire si la masse du squark est inférieure à la masse du gluino, alors ce mode de désintégration est le seul mode significatif : $\tilde{g} \rightarrow q\tilde{q}$. Comme il a été mentionné dans la Section 1.4.4, il est possible d'avoir un sbottom léger. Dans cette analyse, on choisit des masses de squarks de première génération de 400 à 1000 GeV, c'est-à-dire systématiquement supérieures à celles des gluinos recherchés. De plus, les paramètres du $MSSM$ sont choisis de telle sorte que la masse du sbottom le plus léger $\tilde{b}_1$ soit inférieure à celle du gluino. Ainsi, la désintégration $\tilde{g} \rightarrow \tilde{b}\tilde{b}_1$ est possible, on suppose son taux de branchement égal à 100 % dans la suite. De cette façon, on obtient deux sbottoms $\tilde{b}_1$. Cette stratégie de production de sbottoms a l'avantage par rapport à la production directe [83] d'avoir une section efficace d'un ordre de grandeur supérieure et un état final plus discriminant (quatre quarks b contre deux quarks b). La contrepartie est que l'étude d'une production directe de paire de sbottom est une analyse moins dépendante du modèle car elle ne nécessite aucune hypothèse sur la masse du gluino.

Désintégration des sbottoms

Le mode dominant est la désintégration en gluino : $\tilde{b}_1 \rightarrow b\tilde{g}$, si celle-ci est cinématiquement autorisée. Dans le cas contraire, c'est-à-dire celui de cette étude, les squarks se désintègrent préférentiellement en quark et chargino ou neutralino. Le mode dominant de désintégration est la désintégration en neutralino $\tilde{\chi}_1^0$. On supposera le taux de branchement de ce mode égal à 100 %.

La cascade de désintégration du gluino est schématisée sur la Figure 4.4.

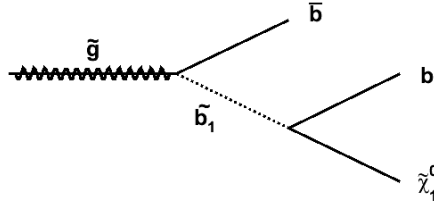


FIG. 4.4 – Chaîne de désintégration du gluino en $2b + mE_T$.

Les deux hypothèses faites sur les taux de branchement égaux à 100 % peuvent paraître optimistes et erronées. Elles permettent néanmoins de fournir des résultats généraux exploitables dans n'importe lequel des modèles supersymétriques. Ainsi, si ces taux de branchement sont différents dans un modèle particulier, il suffit de déterminer de nouveau la limite sur les masses du sbottom et du gluino obtenues en extrapolant simplement les résultats finaux montrés en Section 4.5. De cette façon, plutôt que de traiter un cas particulier sans hypothèse sur les taux de branchement, on préfère se mettre dans un cadre très général en émettant des hypothèses aisément modifiables.

4.2.4 Simulation du signal

Outils utilisés pour la génération

Les différents lots de données *Monte Carlo*, correspondant aux différents points de l'espace des paramètres *MSSM*, sont générés avec le générateur d'événements Pythia v6.323 [148]. Les paramètres d'entrée pour ce générateur, c'est-à-dire les masses des particules et les taux de branchement, sont calculés avec le programme Sdecay v1.1.a [149]. Ce programme donne le spectre de masse des particules supersymétriques à partir des paramètres du modèle théorique choisi. Dans le cas de l'analyse effectuée dans ce chapitre, les paramètres utiles sont essentiellement (voir la Section 1.4) :

- m_3 pour fixer la masse du gluino ;
- $m_{\tilde{b}_R}$, A_b , μ et $\tan(\beta)$ pour fixer la masse du sbottom $\tilde{b}_1$;
- m_1 , μ et $\tan(\beta)$ permettent de fixer la masse du neutralino $\tilde{\chi}_1^0$;
- $m_{\tilde{q}_R}$ choisi suffisamment élevée pour interdire la désintégration du gluino en squarks de première ou deuxième génération.

Sdecay fournit alors au programme Pythia l'ensemble des paramètres nécessaires dans un format standard qui respecte la convention appelée *The SUSY Les Houches Accord* [150]. Cette convention définit les architectures des fichiers d'entrée et de sortie de la plupart des outils utilisés pour les études de signaux supersymétriques. Elle facilite ainsi l'interopérabilité entre ces différents outils.

Une fois la génération des événements supersymétriques effectuée, la chaîne de reconstruction des objets physiques, détaillée dans le Chapitre 3, permet l'obtention de lots de données d'événements supersymétriques dans le format d'analyse usuel de DØ, c'est-à-dire un format similaire et directement comparable au format utilisé pour les données du détecteur. Pour tenir compte des collisions multiples $p\bar{p}$ à haute luminosité, on ajoute à ces événements *Monte Carlo* des collisions supplémentaires obtenues à partir de la condition de déclenchement dite *zero_bias* décrite dans le Chapitre 3. En effet, ces collisions fournissent une bonne modélisation des effets d'empilement et du bruit dans le détecteur. De cette façon, on reproduit ces effets d'empilement et ces bruits, dépendants de la luminosité, dans les événements *Monte Carlo*. Cette modélisation est néanmoins dépendante des collisions multiples $p\bar{p}$ ajoutées et notamment du profil de luminosité de l'ensemble de ces collisions. Pour reproduire parfaitement, le profil de luminosité du lots de données réelles analysées, il est parfois nécessaires de corriger les événements *Monte Carlo*. Cette correction sera détaillée dans la Section 4.4.3.

Caractéristiques des points de l'espace des paramètres *MSSM* générés

Les masses des trois particules supersymétriques étudiées fixent en majeure partie les caractéristiques des événements à analyser.

$m_{\tilde{g}}$ (GeV)	$m_{\tilde{b}_1}$ (GeV)	$m_{\tilde{\chi}_1^0}$ (GeV)	N'_{gen}	σ_{LO} (pb) Pythia v6.323	$K - factor$ Prospino v2
200.	95.	75.	8746	6.860	1.57
200.	150.	75.	9467		
200.	180.	75.	9249		
250.	95.	75.	8973	1.552	1.51
250.	150.	75.	9220		
250.	230.	75.	8972		
300.	150.	75.	9215	0.332	1.48
300.	250.	75.	7999		
300.	280.	75.	9470		
325.	150.	75.	9227	0.185	1.47
325.	200.	75.	9003		
325.	250.	75.	7102		
325.	305.	75.	8990		
350.	105.	75.	8515	0.089	1.46
350.	150.	75.	9014		
350.	200.	75.	8999		
350.	330.	75.	9474		
375.	150.	75.	9214	0.045	1.46
375.	250.	75.	9235		

TAB. 4.1 – Les principales propriétés des signaux supersymétriques générés pour une masse moyenne des squarks de première et deuxième génération égale à 1 TeV.

La section efficace dépend essentiellement de la masse du gluino. Cette section efficace est donnée à l'ordre dominant (*LO* pour *Leading Order*) par le générateur Pythia. Afin d'obtenir

la section efficace à l'ordre supérieur de QCD (NLO pour *Next-to-Leading Order*), et ainsi tenir compte des diagrammes de Feynman à une boucle, on utilise l'outil Prospino [151]. A partir des sections efficaces NLO et des sections efficaces LO calculées par cet outil, on détermine un facteur multiplicatif, appelé $K-factor$ défini de la manière suivante : $K-factor = \frac{\sigma_{NLO}(\tilde{g}\tilde{g})}{\sigma_{LO}(\tilde{g}\tilde{g})}$. Ce facteur est ensuite directement appliqué sur la section efficace LO fournie par Pythia. Les paramètres utilisés pour le calcul des sections efficaces par Prospino sont exactement les mêmes que ceux utilisés pour la génération des événements *Monte Carlo* : $m_t = 175$ GeV, $m_b = 4.75$ GeV et $m_c = 1.55$ GeV, $|\hat{\eta}| < 5$ et la fonction de densité de partons (PDF pour *Parton Density Functions*) CTEQ6M1 (resp. CTEQ6L1)[152] permettent le calcul NLO (resp. LO). Le symbole $\sqrt{s}$ réfère au centre de masse de la collision des deux partons. Les fonctions de densité de partons décrivent le contenu en partons des faisceaux de protons et d'antiprotons à une énergie donnée (980 GeV dans le cas de cette thèse). Ces fonctions prédisent notamment qu'à hautes énergies, c'est-à-dire au-delà de 3 TeV, les densités de gluons sont plus importantes que les densités des quarks composant les protons et antiprotons ($u, d, \bar{u}, \bar{d}$). Ce point a été abordé dans la Section 2.2.1. L'unique différence entre les outils Pythia et Prospino réside dans le choix de l'échelle de renormalisation $\mathcal{Q}_R$ et l'échelle de factorisation $\mathcal{Q}_F$. La première est l'échelle en énergie à laquelle les sections efficaces partoniques sont calculées (voir la Section 1.3.6 pour plus de détails sur le principe de la renormalisation). La seconde est l'échelle en énergie à laquelle les fonctions de densité de partons sont estimées, elle correspond, en ordre de grandeur, à la limite supérieure du confinement. En deçà de cette limite, il est impossible d'utiliser un développement perturbatif pour estimer les sections efficaces de QCD . Prospino utilise plutôt : $\mathcal{Q}_R = \mathcal{Q}_F = m_{\tilde{g}}$; alors que Pythia utilise : $\mathcal{Q}_R = \mathcal{Q}_F = \sqrt{\hat{m}^2(\tilde{g}) + \hat{p}_T^2(\tilde{g})}$. Cette différence explique que la section efficace LO de Prospino soit systématiquement supérieure à celle donnée par Pythia.

Les sections efficaces LO de Pythia et les $K-factor$ de Prospino sont données, pour une partie des points générés pour une masse moyenne des squarks de première et deuxième génération égale à 1 TeV, dans le Tableau 4.1. Dans ce tableau, N'_{gen} représente le nombre d'événements générés après l'application des critères de qualité des données qui seront détaillés dans la Section 4.4.1 (de plus amples détails peuvent être trouvés en [99] pour les définitions des critères de qualité). Pour tenir compte des variations de sections efficaces et de $K-factor$ en fonction de cette masse moyenne, comme il a été évoqué précédemment, on calcule ces deux paramètres pour différentes valeurs de la masse moyenne : 400, 500, 750 et 1000 GeV.

Dans un second temps, les différences de masses entre le gluino et le sbottom $m_{\tilde{b}_1}$ et entre ce sbottom et le neutralino $m_{\tilde{\chi}_1^0}$ donnent des informations sur les distributions des variables caractéristiques de l'événement. En d'autres termes ces différences de masse fixent la topologie du signal et les efficacités des coupures de sélection par la suite. Ces différences de masse permettent également de définir la stratégie de l'analyse qui sera détaillée dans la Section 4.4.1.

4.3 Les lots de données utilisés

Comme il a été vu dans le Chapitre 3, deux lots de données sont nécessaires pour chercher d'éventuelles déviations dues à de la physique au-delà du modèle standard. Dans ce chapitre, on cherche une déviation qui pourrait être due à un signal supersymétrique dont la signature a été donnée dans la Section 4.2. Le premier lot de données comporte l'ensemble des événements ayant une topologie proche de ce signal. Ce lot est décrit dans la Section 4.3.1. Le second est constitué d'événements *Monte Carlo* simulant les processus du modèle standard. Les différents processus

du modèle standard simulés pour cette analyse sont donnés dans la Section 4.3.2.

4.3.1 Le lot de données du détecteur

Le lot de données collectées par le détecteur DØ comporte l'ensemble des collisions enregistrées lors du *Run IIa* d'avril 2003 à février 2006. Les collisions antérieures à avril 2003 n'ont pas été retenues car aucune condition de déclenchement propre aux topologies à jets et énergie transverse manquante n'avait été conçue.

Le système de déclenchement

Trois conditions de déclenchement sont utilisées pour définir le lot de données analysé dans la suite de ce chapitre :

- *MHT30_3CJT5* pour les menus de déclenchement V11 et V12 ;
- *JT1_ACO_MHT_HT* pour les menus V13 et V14 ;
- *JT2_MHT_25_HT* pour les menus V13 et V14.

Les termes de déclenchement de ces trois conditions ainsi que leur évolution tout au long du *Run IIa* sont détaillés dans la Section 3.2.3.

La luminosité intégrée

La luminosité intégrée totale enregistrée par le détecteur DØ d'avril 2003 à février 2006 est de 1.14 fb^{-1} . Cette luminosité se répartit à 6.3 % pour le menu de déclenchement V11, à 22.9 % pour le menu V12, à 37.6 % pour le menu V13 et enfin à 33.2 % pour le menu V14. Quand on parle de luminosité intégrée, il faut différencier trois types de valeurs :

- la luminosité intégrée délivrée par le TeVatron $\mathcal{L}_{del}$;
- la luminosité effectivement enregistrée par le détecteur DØ $\mathcal{L}_{rec}$, qui tient compte des temps morts lors de la prise de données ;
- la luminosité intégrée après application des critères de qualité des données $\mathcal{L}_{good}$, détaillés dans la Section 4.4.1.

Cette dernière est la luminosité réellement utilisée dans la suite. En effet, les lots de données *Monte Carlo* sont normalisés avec cette luminosité intégrée de la façon suivante :

$$N_{equi} = \frac{\mathcal{L}_{good} \times \sigma_{LO} \times K - factor}{N'_{gen}} \quad (4.1)$$

Avec les N'_{gen} , le nombre d'événements générés après application des critères de qualité des données, N_{equi} le nombre d'événements équivalent à une luminosité $\mathcal{L}_{good}$ sachant la section efficace *NLO* du processus en question de $\sigma_{LO} \times K - factor$ (voir détails dans la Section 4.3.2).

Un facteur correctif est calculé pour chacune des conditions de déclenchement pour tenir compte des temps morts dans l'enregistrement des données [153]. Ce facteur correctif permet de passer de $\mathcal{L}_{del}$ à $\mathcal{L}_{rec}$. On définit alors l'efficacité de prises de données, c'est-à-dire le rapport entre $\mathcal{L}_{rec}$ et $\mathcal{L}_{del}$. L'évolution de cette efficacité du début du *Run IIa* jusqu'à aujourd'hui est donnée sur la Figure 4.5. Cette efficacité tient compte de toutes les conditions de déclenchement.

Finalement, les différentes luminosités intégrées correspondant aux trois conditions de déclenchement utilisées pour l'analyse présentée dans ce chapitre sont résumées dans le Tableau 4.2.

La luminosité totale obtenue après les coupures standard de qualité des données est de 0.99 fb^{-1} pour la combinaison des trois conditions de déclenchement.

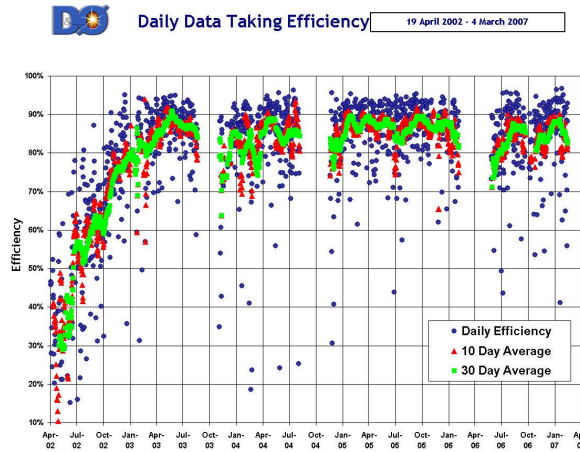


FIG. 4.5 – Efficacités dans la prise de données du détecteur DØ.

Menus de déclenchement	<i>MHT30_3CJT5</i>			<i>JT1_ACO_MHT_HT</i>			<i>JT2_MHT_25_HT</i>		
	$\mathcal{L}_{del}$	$\mathcal{L}_{rec}$	$\mathcal{L}_{good}$	$\mathcal{L}_{del}$	$\mathcal{L}_{rec}$	$\mathcal{L}_{good}$	$\mathcal{L}_{del}$	$\mathcal{L}_{rec}$	$\mathcal{L}_{good}$
V11	79.3	79.1	62.5	-	-	-	-	-	-
V12	277.1	249.9	226.8	-	-	-	-	-	-
V13	-	-	-	464.0	425.5	373.3	464.0	422.0	371.2
V14	-	-	-	416.8	388.8	329.5	416.8	388.8	329.5
Total	356.4	321.0	289.3	880.8	814.3	702.8	880.8	810.8	700.7

TAB. 4.2 – Luminosité intégrée délivrée (en pb^{-1}), enregistrée et après application des critères de qualité des données pour les trois conditions de déclenchement propres aux topologies à jets et énergie transverse manquante.

4.3.2 Les lots de données *Monte Carlo*

Génération des événements de bruit de fond

Générateurs d'événements utilisés Les principaux bruits de fond ont été simulés en utilisant le générateur Alpgen v2.05 [154], exceptés les processus de production de paires de bosons simulés avec le générateur Pythia v6.323 et les processus de quark top célibataire produits avec le générateur Comphep v4.1.10 [155]. Ces trois générateurs permettent de produire les interactions inélastiques prédites par le modèle standard à partir de la modélisation des protons et des antiprotons donnée par les fonction de densité de partons CTEQ6L1.

Dans le cas de la génération d'événements par Alpgen, on utilise une interface entre ce générateur et Pythia. De cette façon, on profite de la complémentarité des deux générateurs. Pythia simule des processus de type $2 \rightarrow 2$ pour deux particules dans l'état initial et l'état final. Il complète ensuite l'état final en produisant des radiations de gluons supplémentaires dans l'état initial et l'état final. Ces dernières engendrent alors des gerbes partoniques, c'est-à-dire des quarks et gluons de basses énergies colinéaires en général au gluon créé, ce dernier étant également coli-

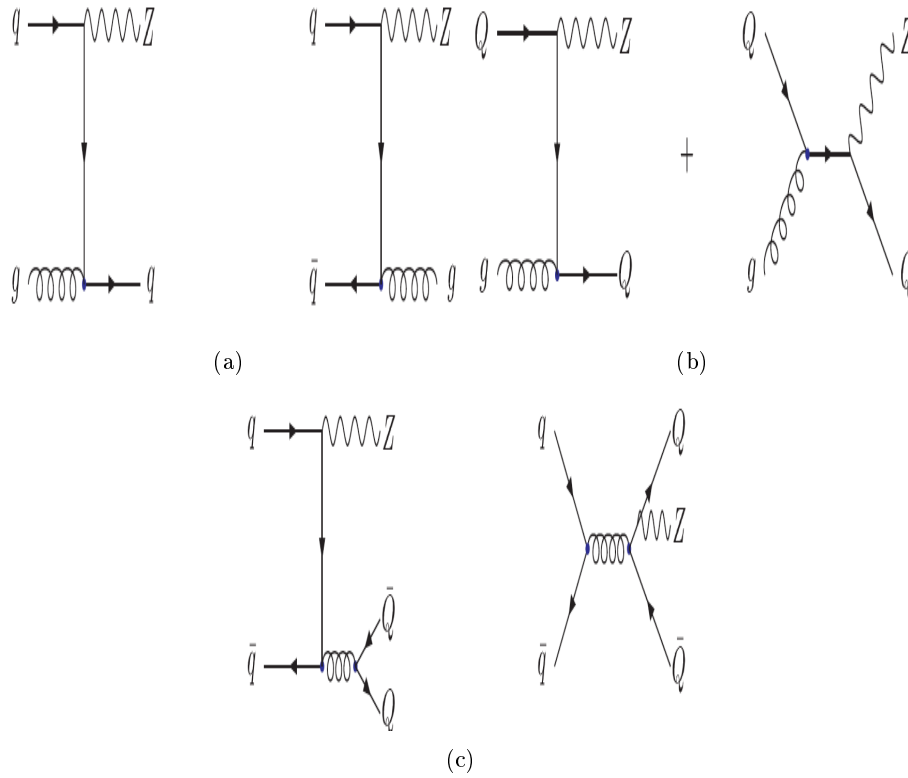


FIG. 4.6 – Exemples de diagrammes de Feynman des modes de productions associées $Z + 1$ jet léger (a), $Z + 1$ jet de saveur lourde (b) et $Z + 2$ jets de saveur lourde (c).

néaire en première approximation au parton dont il est issu (on parle de *parton shower* en anglais). Finalement, il simule l'hadronisation des partons pour une énergie de l'ordre de 200 MeV. Cette manière de procéder a l'avantage de produire des processus avec un nombre quelconque de jets dans l'état final de façon inclusive. Le désavantage, que l'on constate essentiellement pour les topologies multijets, est que Pythia effectue des approximations dans le calcul des sections efficaces $2 \rightarrow n$ (avec n un entier quelconque supérieur à 2). Le générateur AlpGen simule plutôt des processus $2 \rightarrow n$ sans approximation dans le calcul des sections efficaces et en tenant compte des corrélations entre diagrammes de Feynman. Il permet ainsi la production de partons de hautes énergies et/ou à grands angles dans l'état final (les variables topologiques entre les jets sont donc mieux simulées), même si ces derniers sont plus rares. Ce type de production est exclusive (pour un nombre donné de partons dans l'état final) et a le désavantage d'être une génération d'événements au niveau partonique uniquement. L'interface entre les deux générateurs permet alors de produire des processus $2 \rightarrow n$ de façon exclusive, complétés par la gerbe partonique et l'hadronisation de Pythia.

Néanmoins, il peut y avoir des redondances d'états finaux. Il faut alors définir une procédure pour pallier un éventuel double comptage. Cette procédure est couramment appelé procédure de *matching*. La génération des événements par le générateur AlpGen interfacé avec le générateur Pythia se fait pour un nombre fixé de partons dans l'état final. On parle, par exemple, de production

de boson Z avec 0, 1, 2, ... jets supplémentaires comme le montrent les exemples de la Figure 4.6. Ces jets supplémentaires peuvent également provenir de la gerbe partonique. On distingue alors les partons simulés par le générateur Alpgen des partons additionnels produits par Pythia. On définit également deux niveaux de reconstruction : le niveau partonique et les jets de parton. Ces derniers sont reconstruits à partir d'un algorithme de simple cône similaire à celui détaillé dans la Section 3.4.4 (L'énergie transverse minimale est fixée à 8 GeV). Au niveau partonique, on a alors $n_{partons}$ partons dans l'état final produits par le générateur Alpgen. Au niveau des jets de partons, on a n_{jets} jets de partons provenant aussi bien des $n_{partons}$ partons produits par Alpgen que de la gerbe partonique de Pythia. Si on veut générer un lot de données *Monte Carlo* correspondant à une production d'un boson Z et exclusivement n jets supplémentaires, il faut vérifier : $n = n_{partons} = n_{jets}$, sinon l'événement est rejeté. Dans le cas d'une production dite inclusive, c'est-à-dire une production d'un boson Z avec au moins n jets supplémentaires, il faut seulement vérifier : $n_{jets} \geq n_{partons} = n$. Cette procédure évite le double comptage. On appelle dans la suite lp , pour *light parton* (u,d,s et c), les jets de partons supplémentaires. On note *excl.* et *incl.* les productions exclusives et inclusives respectivement.

Par ailleurs, cette procédure de production a également l'avantage d'augmenter la statistique des processus contribuant majoritairement au bruit de fond de l'analyse détaillée dans ce chapitre. En effet, plus le nombre de jets produits en association avec un boson Z est important et plus la section efficace est faible (car plus le nombre de vertex est important dans les diagrammes de Feynman).

Processus avec des saveurs lourdes dans l'état final L'état final du signal recherché possède plusieurs quarks b . On procédera alors à un étiquetage de ces quarks b . Pour plus de précision dans la description des données réelles, il faut donc simuler les processus ayant dans l'état final des saveurs lourdes. Outre les processus produisant un quark top résumés dans le Tableau 4.3, il faut également produire des processus du type $V + jets$, avec V un boson de jauge de l'interaction électrofaible (Z , $W^\pm$ ou γ hors couche de masse), avec dans les jets des saveurs lourdes, i.e. des quarks b ou c . L'ensemble des processus générés avec des saveurs lourdes dans l'état final est donné dans les Tableaux 4.7 et 4.4. Ce dernier tableau donne également les processus $W + jets$ sans saveur lourde dans l'état final. Deux exemples de diagramme de Feynman avec des saveurs lourdes dans l'état final sont donnés sur la Figure 4.6.

Pour pouvoir utiliser à la fois les processus avec des saveurs lourdes et les processus sans production forcée de saveur lourde, il faut filtrer certains événements des processus $W + jets$ et $\gamma^*/Z + jets$ des Tableaux 4.4, 4.5 et 4.6. De cette façon, on évite toute possibilité de double comptage [156]. Les événements à filtrer sont ceux qui contiennent :

- des paires $c\bar{c}$ au niveau de la génération des partons, ceci est valable pour les processus $W + jets$ et $\gamma^*/Z + jets$;
- des paires $c\bar{c}$ ou $b\bar{b}$ au niveau de la radiation de gluons effectuée par Pythia, ceci est valable pour les processus $W + jets$ et $\gamma^*/Z + jets$;
- des paires $c\bar{c}$ au niveau de la radiation de gluons, ceci est valable pour les processus $W + b\bar{b} + jets$ et $\gamma^*/Z + b\bar{b} + jets$.

Une fois ce filtrage terminé, on calcule de nouveau les sections efficaces des processus concernés selon la proportion d'événements rejetés. Les sections efficaces données dans tous les tableaux suivants tiennent compte de ce filtrage.

Nombre d'événements et sections efficaces Les Tableaux 4.3, 4.4, 4.5, 4.6, 4.7 et 4.8 fournissent l'ensemble des processus du modèle standard générés pour l'analyse de ce chapitre. Pour chacun des processus, on donne les différents mode de désintégration qui ont été simulés. Le nombre d'événements après application des critères de qualité des données N'_{gen} , la section efficace $\sigma_{LL/LO}$ et le facteur multiplicatif, appelé *k-factor*, permettent d'obtenir la luminosité intégrée équivalente aux lots de données *Monte Carlo* générées. Ce nombre est à comparer à la luminosité intégrée totale utilisée pour cette analyse, i.e. 0.99 fb^{-1} . Il est important que les processus du modèle standard, qui ont des signatures proches du signal recherché, aient une luminosité intégrée équivalente d'un ordre de grandeur supérieur à 0.99 fb^{-1} . De cette façon, on limite grandement les incertitudes statistiques. Le facteur multiplicatif *k-factor* permet, comme il a été évoqué pour le signal supersymétrique, d'obtenir la section efficace *NLO* d'un processus donné à partir de la section efficace *LO* ou *LL* dans le cas du générateur Alpgen. *LL*, pour *Leading Log* en anglais, signifie que la précision sur la section efficace est à l'ordre logarithmique le plus haut. La procédure choisie [157] pour déterminer ce facteur est la suivante :

$$k - \text{factor} = \frac{[\sigma^{incl}(X + 0lp)]_{NLO}}{[\sigma(X + 0lp)]_{LO}} \quad (4.2)$$

Où $\sigma^{incl}(X + 0lp)$ désigne la section efficace inclusive du processus.

Pour déterminer la luminosité intégrée équivalente $\mathcal{L}_{equi}$, on applique la formule suivante :

$$\mathcal{L}_{equi} = \frac{N'_{gen}}{\sigma_{LL/LO} \times k - \text{factor}} \quad (4.3)$$

Bruits de fond physique

Les bruits de fond physique sont les processus du modèle standard dont la topologie s'approche de celle du signal recherché. Tous ces bruits de fond ont été simulés par les outils décrits précédemment. Les bruits de fond principaux sont la production de paires de quark top, les processus $W + jet$ avec des saveurs lourdes et les processus $Z(\rightarrow \nu\bar{\nu}) + jets$ avec des saveurs lourdes également.

Le bruit de fond instrumental

Les processus de *QCD* sont également appelés bruit de fond instrumental. Ils contribuent fortement au bruit de fond total de l'analyse à cause de leurs sections efficaces particulièrement élevées. Les figures de la Section 4.4.2 illustrent cette différence d'ordre de grandeur entre le nombre d'événements des processus *QCD* et le nombre d'événements des autres processus du modèle standard. Il n'est pas possible d'avoir une simulation précise de ce fond instrumental. D'une part, les sections efficaces de ces processus étant très importantes, il faudrait simuler une statistique énorme d'événements pour décrire correctement ce fond (voir la Formule 4.3). D'autre part, la simulation du détecteur ne permet pas une estimation assez précise des mauvaises mesures des jets. Or, le fond instrumental est très sensible à ces mauvaises mesures, car ces dernières sont à l'origine de l'énergie transverse manquante pour ce bruit de fond. Enfin, ce bruit de fond n'est pas simulé par le générateur Alpgen et Pythia est trop imprécis pour les processus multijets, comme il a été décrit précédemment.

Processus	Modes de désintégration	N'_{gen}	$\sigma_{LL/LO}$ (pb)	K-factor	Luminosité équivalente (fb^{-1})
$t + \bar{t} + 0lp$ (excl.)	$2b + 4lp$	96734	1.300	1.374	54.16
	$2\ell + \nu\bar{\nu} + 2b$	223768	0.324	1.374	502.65
	$\ell\nu + 2b + 2lp$	283765	1.300	1.374	158.87
$t + \bar{t} + 1lp$ (excl.)	$2b + 5lp$	90094	0.538	1.374	121.88
	$2\ell + \nu\bar{\nu} + 2b + 1lp$	96381	0.135	1.374	130.38
	$\ell\nu + 2b + 3lp$	98424	0.540	1.374	132.65
$t + \bar{t} + 2lp$ (incl.)	$2b + 6lp$	23723	0.254	1.374	67.97
	$2\ell + \nu\bar{\nu} + 2b + 2lp$	148099	0.063	1.374	1710.90
	$\ell\nu + 2b + 4lp$	92684	0.254	1.374	265.57
$t + q + b$	$e\nu b + qb$	130068	0.220	1.	591.22
$t + q + b$	$\mu\nu b + qb$	539838	0.220	1.	2453.81
$t + q + b$	$\tau\nu b + qb$	273306	0.220	1.	1242.30
$t + q + b$	$q'\bar{q}b + qb$	140708	1.518	1.	92.69
$t + b$	$e\nu b + b$	109303	0.098	1.	1115.34
$t + b$	$\mu\nu b + b$	112348	0.098	1.	1146.41
$t + b$	$\tau\nu b + b$	141197	0.098	1.	1440.79
$t + b$	$q'\bar{q}b + b$	137379	0.675	1.	203.52

TAB. 4.3 – Les principales propriétés des processus du modèle standard produisant un quark top ou une paire de quarks top.

Par conséquent, le bruit de fond instrumental QCD n'est pas estimé dans cette recherche d'un signal supersymétrique. L'estimation du nombre d'événements issus de ce bruit de fond est directement extraite des données du détecteur comme il est détaillé dans la Section 4.4.2.

Processus	Modes de désintégration	N'_{gen}	σ_{LL} (pb)	K-factor	Luminosité équivalente (fb^{-1})
$W^\pm + 0lp$ (excl.)	$\ell\nu + 0lp$	2257330	4574.36	1.318	0.37
$W^\pm + 1lp$ (excl.)	$\ell\nu + 1lp$	2754645	1273.94	1.318	1.64
$W^\pm + 2lp$ (excl.)	$\ell\nu + 2lp$	1564990	298.56	1.318	3.98
$W^\pm + 3lp$ (excl.)	$\ell\nu + 3lp$	789121	70.56	1.318	8.49
$W^\pm + 4lp$ (excl.)	$\ell\nu + 4lp$	778692	15.831	1.318	37.32
$W^\pm + 5lp$ (incl.)	$\ell\nu + 5lp$	57973	8.064	1.318	5.45
$W^\pm + c + 0lp$ (excl.)	$\ell\nu + c + 0lp$	291790	46.0	1.318	4.81
$W^\pm + c + 1lp$ (excl.)	$\ell\nu + c + 1lp$	379995	17.05	1.318	16.91
$W^\pm + c + 2lp$ (excl.)	$\ell\nu + c + 2lp$	190068	4.53	1.318	31.83
$W^\pm + c + 3lp$ (incl.)	$\ell\nu + c + 3lp$	249288	1.0	1.318	189.14
$W^\pm + c\bar{c} + 0lp$ (excl.)	$\ell\nu + c\bar{c} + 0lp$	481859	71.145	1.318	5.14
$W^\pm + c\bar{c} + 1lp$ (excl.)	$\ell\nu + c\bar{c} + 1lp$	337862	29.854	1.318	8.59
$W^\pm + c\bar{c} + 2lp$ (excl.)	$\ell\nu + c\bar{c} + 2lp$	332547	10.251	1.318	24.61
$W^\pm + c\bar{c} + 3lp$ (incl.)	$\ell\nu + c\bar{c} + 3lp$	372977	13.135	1.318	21.54
$W^\pm + b\bar{b} + 0lp$ (excl.)	$\ell\nu + b\bar{b} + 0lp$	739617	19.182	1.318	29.25
$W^\pm + b\bar{b} + 1lp$ (excl.)	$\ell\nu + b\bar{b} + 1lp$	261515	7.939	1.318	24.99
$W^\pm + b\bar{b} + 2lp$ (excl.)	$\ell\nu + b\bar{b} + 2lp$	171486	2.637	1.318	49.34
$W^\pm + b\bar{b} + 3lp$ (incl.)	$\ell\nu + b\bar{b} + 3lp$	163674	1.244	1.318	99.83

TAB. 4.4 – Les principales propriétés des processus du modèle standard correspondant à la production par fusion d'un boson W.

Processus	Modes de désintégration	N'_{gen}	σ_{LL} (pb)	K-factor	Luminosité équivalente (fb^{-1})
$\gamma^*/Z + 0lp$ (excl., $15 < \hat{m} < 60$ GeV)	$e^\pm e^\mp + 0lp$	561955	336.212	1.329	1.26
	$\mu^\pm \mu^\mp + 0lp$	552242	336.212	1.329	1.24
	$\tau^\pm \tau^\mp + 0lp$	534902	336.212	1.329	1.20
$\gamma^*/Z + 0lp$ (excl., $60 < \hat{m} < 130$ GeV)	$e^\pm e^\mp + 0lp$	2034300	140.291	1.329	10.91
	$\mu^\pm \mu^\mp + 0lp$	2277864	140.291	1.329	12.22
	$\tau^\pm \tau^\mp + 0lp$	2151362	140.291	1.329	11.54
$\gamma^*/Z + 0lp$ (excl., $130 < \hat{m} < 250$ GeV)	$e^\pm e^\mp + 0lp$	95535	0.906	1.329	79.34
	$\mu^\pm \mu^\mp + 0lp$	100603	0.906	1.329	83.55
	$\tau^\pm \tau^\mp + 0lp$	99596	0.906	1.329	82.72
$\gamma^*/Z + 0lp$ (excl., $250 < \hat{m} < 1960$ GeV)	$e^\pm e^\mp + 0lp$	97702	0.0701057	1.329	1048.64
	$\mu^\pm \mu^\mp + 0lp$	92722	0.0701057	1.329	995.19
	$\tau^\pm \tau^\mp + 0lp$	93027	0.0701057	1.329	998.46
$\gamma^*/Z + 1lp$ (excl., $15 < \hat{m} < 60$ GeV)	$e^\pm e^\mp + 1lp$	427237	39.4202	1.329	8.15
	$\mu^\pm \mu^\mp + 1lp$	423180	39.4202	1.329	8.08
	$\tau^\pm \tau^\mp + 1lp$	430829	39.4202	1.329	8.22
$\gamma^*/Z + 1lp$ (excl., $60 < \hat{m} < 130$ GeV)	$e^\pm e^\mp + 1lp$	363201	42.270	1.329	6.47
	$\mu^\pm \mu^\mp + 1lp$	375077	42.269	1.329	6.68
	$\tau^\pm \tau^\mp + 1lp$	556291	42.269	1.329	9.90
$\gamma^*/Z + 1lp$ (excl., $130 < \hat{m} < 250$ GeV)	$e^\pm e^\mp + 1lp$	83992	0.365395	1.329	172.96
	$\mu^\pm \mu^\mp + 1lp$	90775	0.365395	1.329	186.93
	$\tau^\pm \tau^\mp + 1lp$	89944	0.365395	1.329	185.22
$\gamma^*/Z + 1lp$ (excl., $250 < \hat{m} < 1960$ GeV)	$e^\pm e^\mp + 1lp$	88427	0.0369237	1.329	1802.00
	$\mu^\pm \mu^\mp + 1lp$	88021	0.0369237	1.329	1793.73
	$\tau^\pm \tau^\mp + 1lp$	87978	0.0369237	1.329	1792.85
$\gamma^*/Z + 2lp$ (excl., $15 < \hat{m} < 60$ GeV)	$e^\pm e^\mp + 2lp$	164417	10.297	1.329	12.01
	$\mu^\pm \mu^\mp + 2lp$	163310	10.297	1.329	11.93
$\gamma^*/Z + 2lp$ (excl., $60 < \hat{m} < 130$ GeV)	$e^\pm e^\mp + 2lp$	250703	10.4663	1.329	18.02
	$\mu^\pm \mu^\mp + 2lp$	178071	10.4663	1.329	12.80
	$\tau^\pm \tau^\mp + 2lp$	177619	10.4663	1.329	12.77
$\gamma^*/Z + 2lp$ (excl., $130 < \hat{m} < 250$ GeV)	$e^\pm e^\mp + 2lp$	86647	0.0986418	1.329	660.95
	$\mu^\pm \mu^\mp + 2lp$	86490	0.0986418	1.329	659.75
	$\tau^\pm \tau^\mp + 2lp$	80372	0.0986418	1.329	613.08
$\gamma^*/Z + 2lp$ (excl., $250 < \hat{m} < 1960$ GeV)	$e^\pm e^\mp + 2lp$	88344	0.0118	1.329	5633.39
	$\mu^\pm \mu^\mp + 2lp$	82230	0.0118	1.329	5243.52
	$\tau^\pm \tau^\mp + 2lp$	81915	0.0118	1.329	5223.44
$\gamma^*/Z + 3lp$ (incl., $15 < \hat{m} < 60$ GeV)	$e^\pm e^\mp + 3lp$	77676	3.08402	1.329	18.95
	$\mu^\pm \mu^\mp + 3lp$	76757	3.08402	1.329	18.73
	$\tau^\pm \tau^\mp + 3lp$	76321	3.08402	1.329	18.62
$\gamma^*/Z + 3lp$ (incl., $60 < \hat{m} < 130$ GeV)	$e^\pm e^\mp + 3lp$	225478	3.90446	1.329	43.45
	$\mu^\pm \mu^\mp + 3lp$	242080	3.90446	1.329	46.65
	$\tau^\pm \tau^\mp + 3lp$	155664	3.90446	1.329	30.00
$\gamma^*/Z + 3lp$ (incl., $130 < \hat{m} < 250$ GeV)	$e^\pm e^\mp + 3lp$	74660	0.040569	1.329	1384.74
	$\mu^\pm \mu^\mp + 3lp$	73215	0.040569	1.329	1357.94
	$\tau^\pm \tau^\mp + 3lp$	71102	0.040569	1.329	1318.75
$\gamma^*/Z + 3lp$ (incl., $250 < \hat{m} < 1960$ GeV)	$e^\pm e^\mp + 3lp$	74062	0.0048711	1.329	11440.45
	$\mu^\pm \mu^\mp + 3lp$	77377	0.0048711	1.329	11952.53
	$\tau^\pm \tau^\mp + 3lp$	75722	0.0048711	1.329	11696.88

TAB. 4.5 – Principales propriétés des processus $Z(\rightarrow \ell^\pm \ell^\mp) + jets$.

Processus	Modes de désintégration	N'_{gen}	σ_{LL} (pb)	K-factor	Luminosité équivalente (fb^{-1})
$Z + 0lp$ (excl.)	$\nu\bar{\nu} + 0lp$	739759	818.043	1.322	0.68
$Z + 1lp$ (excl.)	$\nu\bar{\nu} + 1lp$	1228440	245.504	1.322	3.78
$Z + 2lp$ (excl.)	$\nu\bar{\nu} + 2lp$	241912	61.184	1.322	2.99
$Z + 3lp$ (excl.)	$\nu\bar{\nu} + 3lp$	76159	14.605	1.322	3.94
$Z + 4lp$ (excl.)	$\nu\bar{\nu} + 4lp$	90435	3.368	1.322	20.31
$Z + 5lp$ (incl.)	$\nu\bar{\nu} + 5lp$	73942	1.857	1.322	30.12

TAB. 4.6 – Principales propriétés des processus $Z(\rightarrow \nu\bar{\nu}) + jets$.

Processus	Modes de désintégration	N'_{gen}	σ_{LL} (pb)	K-factor	Luminosité équivalente (fb^{-1})
$Z + c\bar{c} + 0lp$ (excl.)	$\nu\bar{\nu} + c\bar{c} + 0lp$	170383	17.7439	1.322	7.26
$Z + c\bar{c} + 1lp$ (excl.)	$\nu\bar{\nu} + c\bar{c} + 1lp$	89270	6.26213	1.322	10.78
$Z + c\bar{c} + 2lp$ (incl.)	$\nu\bar{\nu} + c\bar{c} + 2lp$	90739	2.51	1.322	27.35
$\gamma^*/Z + c\bar{c} + 0lp$ (excl.)	$e^{\pm}e^{\mp} + c\bar{c} + 0lp$	47252	3.04996	1.329	11.66
$\gamma^*/Z + c\bar{c} + 0lp$ (excl.)	$\mu^{\pm}\mu^{\mp} + c\bar{c} + 0lp$	47125	3.04996	1.329	11.63
$\gamma^*/Z + c\bar{c} + 0lp$ (excl.)	$\tau^{\pm}\tau^{\mp} + c\bar{c} + 0lp$	39295	3.04996	1.329	9.69
$\gamma^*/Z + c\bar{c} + 1lp$ (excl.)	$e^{\pm}e^{\mp} + c\bar{c} + 1lp$	42901	1.07254	1.329	30.17
$\gamma^*/Z + c\bar{c} + 1lp$ (excl.)	$\mu^{\pm}\mu^{\mp} + c\bar{c} + 1lp$	42776	1.07254	1.329	30.08
$\gamma^*/Z + c\bar{c} + 1lp$ (excl.)	$\tau^{\pm}\tau^{\mp} + c\bar{c} + 1lp$	43491	1.07254	1.329	30.58
$\gamma^*/Z + c\bar{c} + 2lp$ (incl.)	$e^{\pm}e^{\mp} + c\bar{c} + 2lp$	21921	0.438884	1.329	37.57
$\gamma^*/Z + c\bar{c} + 2lp$ (incl.)	$\mu^{\pm}\mu^{\mp} + c\bar{c} + 2lp$	23317	0.438884	1.329	39.97
$\gamma^*/Z + c\bar{c} + 2lp$ (incl.)	$\tau^{\pm}\tau^{\mp} + c\bar{c} + 2lp$	20930	0.438884	1.329	35.87
$Z + b\bar{b} + 0lp$ (excl.)	$\nu\bar{\nu} + b\bar{b} + 0lp$	189986	5.8059	1.322	24.74
$Z + b\bar{b} + 1lp$ (excl.)	$\nu\bar{\nu} + b\bar{b} + 1lp$	88251	2.1492	1.322	31.06
$Z + b\bar{b} + 2lp$ (incl.)	$\nu\bar{\nu} + b\bar{b} + 2lp$	87349	0.8869	1.322	74.50
$\gamma^*/Z + b\bar{b} + 0lp$ (excl.)	$e^{\pm}e^{\mp} + b\bar{b} + 0lp$	230272	0.987807	1.329	175.41
$\gamma^*/Z + b\bar{b} + 0lp$ (excl.)	$\mu^{\pm}\mu^{\mp} + b\bar{b} + 0lp$	266953	0.987807	1.329	203.35
$\gamma^*/Z + b\bar{b} + 0lp$ (excl.)	$\tau^{\pm}\tau^{\mp} + b\bar{b} + 0lp$	93013	0.987807	1.329	70.85
$\gamma^*/Z + b\bar{b} + 1lp$ (excl.)	$e^{\pm}e^{\mp} + b\bar{b} + 1lp$	47738	0.367608	1.329	97.72
$\gamma^*/Z + b\bar{b} + 1lp$ (excl.)	$\mu^{\pm}\mu^{\mp} + b\bar{b} + 1lp$	48052	0.367608	1.329	98.36
$\gamma^*/Z + b\bar{b} + 1lp$ (excl.)	$\tau^{\pm}\tau^{\mp} + b\bar{b} + 1lp$	182386	0.367608	1.329	373.33
$\gamma^*/Z + b\bar{b} + 2lp$ (incl.)	$e^{\pm}e^{\mp} + b\bar{b} + 2lp$	20702	0.152731	1.329	102.01
$\gamma^*/Z + b\bar{b} + 2lp$ (incl.)	$\mu^{\pm}\mu^{\mp} + b\bar{b} + 2lp$	21627	0.152731	1.329	106.57
$\gamma^*/Z + b\bar{b} + 2lp$ (incl.)	$\tau^{\pm}\tau^{\mp} + b\bar{b} + 2lp$	86872	0.152731	1.329	428.07

TAB. 4.7 – Principales propriétés des processus $\gamma^*/Z + jets$ avec des saveurs lourdes. Les processus $Z(\rightarrow \ell^{\pm}\ell^{\mp})$ ont été générés dans l'intervalle de masse $60 < \hat{m} < 130$ GeV.

Processus	Modes de désintégration	N'_{gen}	σ_{LO} (pb)	K-factor	Luminosité équivalente (fb^{-1})
$W^{\pm} + W^{\mp}$	incl.	467777	11.54	1.	40.54
$W^{\pm} + Z$	incl.	276775	3.49	1.	79.31
$\gamma^*/Z + \gamma^*/Z$	incl.	191420	1.56	1.	122.71

TAB. 4.8 – Principales propriétés des processus produisant une paire de bosons.

4.4 Sélection des événements et traitement des données *Monte Carlo*

4.4.1 La stratégie de l'analyse

La topologie du signal recherché est fortement liée aux différences de masse entre les particules supersymétriques : $\Delta M = m(\tilde{g}) - m(\tilde{b}_1)$ et $\Delta m = m(\tilde{b}_1) - m(\tilde{\chi}_1^0)$. En effet, si les différences de masses ΔM et Δm sont suffisamment grandes (supérieures à 50 GeV), la reconstruction des jets et l'étiquetage des quarks b sont plus efficaces et on s'attend alors à quatre jets de hautes impulsions dans l'état final. En fonction du point de fonctionnement choisi, on peut dans ce cas étiqueter de deux à quatre jets. Un autre cas de figure est possible si ΔM ou Δm est relativement faible (un bon ordre de grandeur est entre 10 et 20 GeV). Dans ce cas, on s'attend plutôt à deux jets de hautes impulsions et à une forte énergie transverse manquante. On peut alors étiqueter de un à deux jets. Ces deux cas de figure sont illustrés sur la Figure 4.7. Afin de rechercher des signaux supersymétriques dans la plus grande partie du plan $(m(\tilde{g}), m(\tilde{b}_1))$, on est amené à définir deux canaux d'analyse avec des conditions de déclenchement et des coupures de sélection différentes. Le choix des conditions de déclenchement et des coupures de sélection appliquées est détaillé dans la suite.

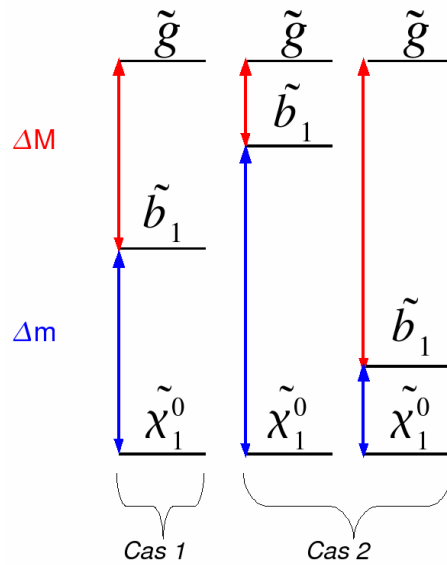


FIG. 4.7 – Deux types de signaux à étudier suivant les différences de masse entre les particules supersymétriques.

La qualité des données

Les deux canaux d'analyse ont en commun la sélection des événements satisfaisant les critères de qualité des données. En effet, seules les données prises dans de bonnes conditions sont conservées pour les analyses de physique. Il existe deux types de critère de qualité des données. En effet, on peut être amené à rejeter soit une période de prise de données, soit un événement particulier.

Ainsi, dans le premier cas, si pendant une période de prise de données, un des sous-détecteurs est absent ou ne fonctionne pas, alors le *Run* en question est rejeté. Les sous-détecteurs utilisés pour définir les mauvais *Run* sont les suivants : *SMT*, *CFT*, système à muons et calorimètre. De la même façon, on définit des mauvais *LBN*, correspondant à une durée d'une minute de prise de données. C'est l'unité utilisée pour la mesure de la luminosité instantanée comme il est expliqué dans la Section 2.3.5. Si pendant cet intervalle de temps, les moniteurs de luminosité ne fonctionnent pas correctement, on rejette le *LBN* considéré. De même, on rejette des *LBN*, si la valeur moyenne de variables, comme l'énergie transverse manquante par exemple, dévie de la valeur usuelle (on recherche en fait des déviations brutales de la valeur moyenne par *LBN* en fonction du temps). Dans la suite de l'analyse, les *Run* et les *LBN* considérés comme mauvais ne sont pas utilisés.

Dans le deuxième cas, on peut rejeter certains événements particuliers sujets à des problèmes dans le calorimètre. Ces problèmes sont classés en plusieurs catégories selon leur origine, soit interne pour les problèmes dus à l'électronique du calorimètre soit externe pour les problèmes provoqués par d'autres appareils électroniques, et selon leur façon de se manifester dans le calorimètre. On trouve les catégories de bruits ou problèmes suivants [158] :

- Caisse manquante (*empty crate*) : au moins un des châssis de lecture est vide de tout signal. Ce problème est dû à un mauvais fonctionnement de l'électronique de lecture du calorimètre.
- Bruit "cohérent" (*coherent noise*) : saut piédestal identique pour de nombreuses cartes d'acquisition *ADC*. Ce bruit électronique peut toucher 10 % des événements. La source de ce bruit est actuellement mal connue.
- "Anneau de feu" (*Ring Of Fire*) : anneau de cellules ayant un signal anormalement élevé. La source de ce bruit est externe et liée aux générateurs de hautes tensions des bouchons du calorimètre, qui sont des électrodes circulaires.
- "Bruit de midi" (*Noon Noise*) : ce bruit affecte plusieurs châssis de lecture et la somme scalaire de leur énergie transverse devient très élevée. Au moins un des châssis de lecture a une occupation supérieure à 35 %. La source est externe au calorimètre.
- "Brouillard violet" (*Purple Haze*) : c'est un bruit de midi avec une somme scalaire des énergies transverses des cellules supérieures à 2 TeV ! La source est inconnue. Ce bruit n'est pas réapparu au début du *Run IIb*.
- "Eventail espagnol" (*Spanish Fan*) : ce bruit a été découvert en juin 2006. Pour ce bruit, des anneaux en ϕ situés dans l'intervalle $0.7 < |\eta_{det}| < 0.8$ ont une occupation anormalement élevée. Les événements subissant ce bruit sont facilement identifiables, étant donné qu'il possède une fraction *CHF* supérieure à 0.3 et une fraction électromagnétique *EMF* inférieure à 0.3.

Des illustrations de ces bruits sont données sur les Figures 4.8 et 4.9. Sur ces figures, on trouve des schémas à deux ou trois dimensions de l'occupation des cellules du calorimètre en fonction de η et ϕ .

Les événements possédant au moins un de ces bruits, liés au calorimètre sont rejetés, pour les deux canaux d'analyse. Il est également possible que de tels bruits apparaissent dans les données *Monte Carlo* en raison des collisions multiples $p\bar{p}$ ajoutées pour simuler le profil de luminosité. Les événements *Monte Carlo* concernés sont également rejetés. L'efficacité des coupures de sélection permettant d'éliminer ces événements a été mesurée dans les données réelles [159] : 97.14 ± 0.003 %. On applique alors un poids global de 0.97 pour les données simulées afin de tenir compte de cette inefficacité observée dans les données réelles.

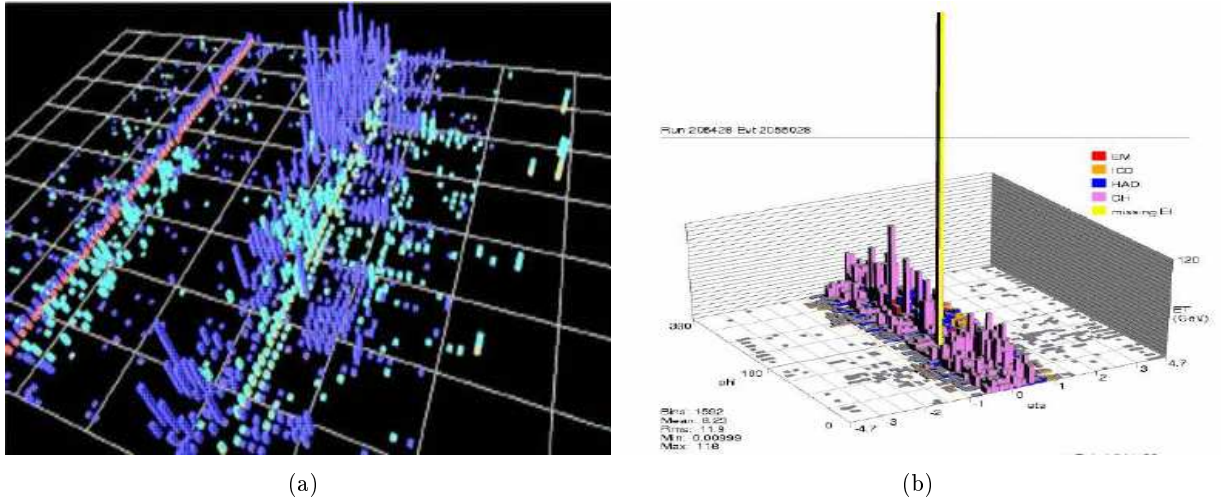


FIG. 4.8 – Exemples de bruits dans le calorimètre : *Ring Of Fire* (a) et *Purple Haze*, en jaune l'énergie transverse manquante (b).

Le calcul de la luminosité intégrée totale utilisée dans la recherche du signal supersymétrique n'utilise ni ces événements rejetés ni ces périodes de prise de données rejetées.

Les conditions de déclenchement

Les conditions de déclenchement *MHT30_3CJT5*, *JT1_ACO_MHT_HT* et *JT2_MHT25_HT* sont simulées pour les données *Monte Carlo* [160, 161]. À partir de lots de données réelles, une paramétrisation des objets physiques utilisés au niveau du système de déclenchement est déterminée. On paramètre alors le nombre de tours de déclenchement au dessus d'un seuil en énergie transverse, l'énergie transverse des jets de niveau 2 et des jets de niveau 3, en fonction des caractéristiques des jets *offline*. De cette façon, on est en mesure d'estimer si un événement *Monte Carlo* passe l'ensemble des termes de déclenchement.

Les incertitudes systématiques sur les simulations des conditions de déclenchement propres aux topologies à jets et énergie transverse manquante ont été étudiées dans la référence [162]. Une incertitude conservatrice de 2% sera utilisée dans la suite.

Recherche d'un état final avec deux jets et l'énergie transverse manquante

De façon à couvrir un maximum de l'espace des paramètres *MSSM* accessible au TeVatron, deux canaux d'analyse ont été conçus. Le premier, appelée analyse *JT1* dans la suite, permet préférentiellement une recherche d'événement à deux jets dans l'état final et une forte énergie transverse manquante (c'est-à-dire faible valeur de ΔM ou Δm). Comme il a été vu dans la Section 3.2.3, les conditions de déclenchement *MHT30_3CJT5* puis *JT1_ACO_MHT_HT* correspondent à cette topologie. À partir de ces deux conditions de déclenchement, on définit alors une présélection pour cette analyse. Les coupures ont été choisies de façon à sélectionner la topologie recherchée et à limiter les biais éventuellement introduits par les conditions de déclenchement. Cette présélection,

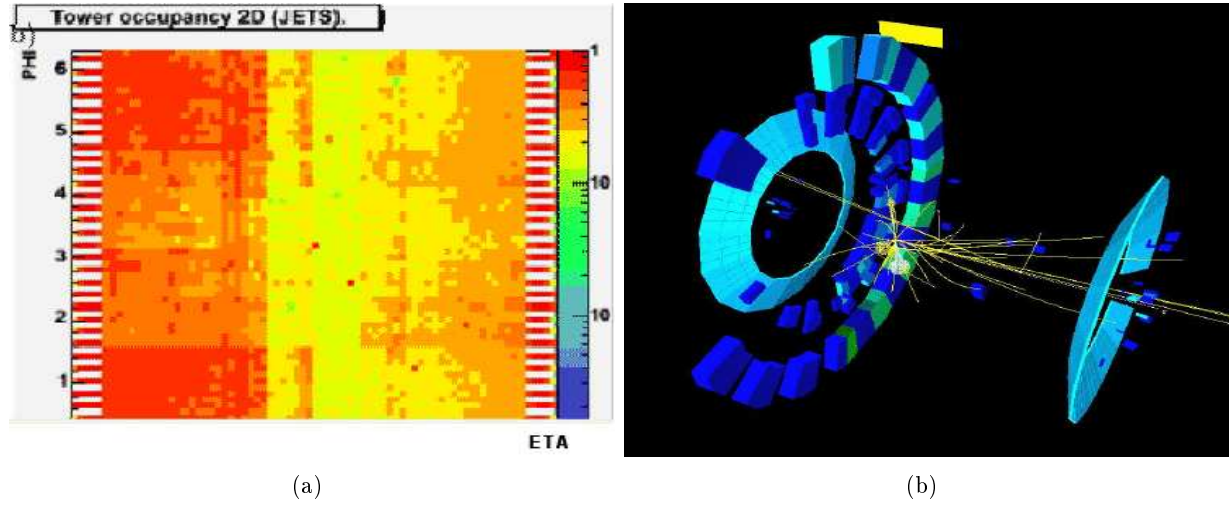


FIG. 4.9 – Exemples de bruits dans le calorimètre : *coherent noise* (a) et *Spanish Fan* (b).

notée $C1$, est la suivante :

$$\left\{ \begin{array}{l} \text{Conditions de déclenchement : } MHT30_3CJT5 \text{ ou } JT1_ACO_MHT_HT \\ \text{Critères de qualité des données} \\ \text{Au moins un vertex primaire reconstruit (PV), avec : } |z_{PV}| < 60 \text{ cm} \\ \text{Au moins 2 jets } (j_1 \text{ et } j_2) \text{ avec } E_T > 40 \text{ GeV et } |\eta_{\text{det}}(j_1, j_2)| < 0.8 \\ E_T > 50 \text{ GeV} \\ \cancel{E}_T > 50 \text{ GeV} \\ \Delta\phi(j_1, j_2) < 165^\circ \\ \Delta\phi_{\min}(\cancel{E}_T, jets) > 40^\circ \end{array} \right.$$

Recherche d'un état final multijet avec de l'énergie transverse manquante

La deuxième analyse, dite $JT2$, sélectionne plutôt les événements avec au moins trois jets dans l'état final. Les conditions de déclenchement conçues pour de telles topologies sont $MHT30_3CJT5$ puis $JT2_MHT25_HT$. De la même façon que pour l'analyse $JT1$, on définit les coupures de présélection $C1$:

$$\left\{ \begin{array}{l} \text{Conditions de déclenchement : } MHT30_3CJT5 \text{ ou } JT2_MHT25_HT \\ \text{Critères de qualité des données} \\ \text{Au moins un vertex primaire reconstruit (PV), avec : } |z_{PV}| < 60 \text{ cm} \\ \text{Au moins 3 jets } (j_1, j_2, \text{ et } j_3) \text{ avec } E_T > 40 \text{ GeV et } |\eta_{\text{det}}(j_1, j_2)| < 0.8, |\eta_{\text{det}}(j_3)| < 2.5 \\ E_T > 50 \text{ GeV} \\ \cancel{E}_T > 50 \text{ GeV} \\ H_T > 180 \text{ GeV} \end{array} \right.$$

Efficacités relatives des conditions de déclenchement

Une fois les coupures de présélection déterminées pour chacune des deux analyses, on estime les efficacités relatives des conditions de déclenchement $JT1_ACO_MHT_HT$ et $MHT30_3CJT5$ pour l'analyse $JT1$ et des conditions de déclenchement $JT2_MHT25_HT$ et $MHT30_3CJT5$ pour l'analyse $JT2$. Ces efficacités sont calculées uniquement pour les événements sélectionnés par les coupures détaillées précédemment. Elles sont données pour les processus du modèle standard qui ont un nombre d'événements générés supérieur à 1000 après les coupures de présélection, de cette manière on ne décrit que les processus qui contribuent significativement au bruit de fond des deux analyses. Le Tableau 4.9 donne les résultats pour l'analyse $JT1$, alors que le Tableau 4.10 donne ceux de l'analyse $JT2$. On retrouve dans ces deux tableaux les fonds du modèle standard dominants à ce stade des coupures de sélection.

On constate que plus l'énergie transverse manquante moyenne est importante, plus l'efficacité relative des conditions de déclenchement est importante. C'est pour cette raison que les efficacités relatives des signaux supersymétriques sont systématiquement supérieures. On constate également que les coupures de présélection, choisies de telle façon à limiter les biais introduits par les conditions de déclenchement, permettent des efficacités proches pour des topologies proches. En effet, les efficacités relatives des signaux supersymétriques, pour des $\Delta M, m$ du même ordre de grandeur, sont identiques à l'incertitude statistique près. De plus, plus cette énergie transverse manquante moyenne est importante, plus la contribution d'un processus est importante. L'énergie transverse manquante ne tenant pas compte des muons dans l'événement (voir la Section 3.4.5), les processus $Z + jets \rightarrow \mu^\pm \mu^\mp + jets$ peuvent contribuer au bruit de fond comme on le constate dans le Tableau 4.9. Ce n'est pas le cas des processus $Z + jets \rightarrow e^\pm e^\mp + jets$.

Processus du modèle standard	Mode de désintégration	Efficacités (%)
$t + \bar{t} + 0lp$ (excl.)	$2\ell + \nu\bar{\nu} + 2b$ $\ell\nu + 2b + 0lp$	91.47 ± 0.14 90.67 ± 0.10
$t + \bar{t} + 1lp$ (excl.)	$2\ell + \nu\bar{\nu} + 2b$ $\ell\nu + 2b + 1lp$	91.15 ± 0.20 89.99 ± 0.20
$t + \bar{t} + 2lp$ (incl.)	$2\ell + \nu\bar{\nu} + 2b$ $\ell\nu + 2b + 2lp$	89.97 ± 0.14 88.65 ± 0.20
$t + q + b$	$e\nu b + qb$ $\mu\nu b + qb$ $\tau\nu b + qb$	89.85 ± 0.43 92.06 ± 0.17 91.67 ± 0.22
$t + b$	$e\nu b + b$ $\mu\nu b + b$ $\tau\nu b + b$	89.52 ± 0.31 94.42 ± 0.17 92.26 ± 0.20
$W^\pm + 2lp$ (excl.)	$\ell\nu + 2lp$	87.57 ± 0.37
$W^\pm + 3lp$ (excl.)	$\ell\nu + 3lp$	89.98 ± 0.26
$W^\pm + 4lp$ (excl.)	$\ell\nu + 4lp$	90.16 ± 0.17
$W^\pm + 5lp$ (incl.)	$\ell\nu + 5lp$	91.44 ± 0.43
$W^\pm + c + 3lp$ (incl.)	$\ell\nu + c + 3lp$	89.36 ± 0.40
$W^\pm + c\bar{c} + 2lp$ (excl.)	$\ell\nu + c\bar{c} + 2lp$	91.37 ± 0.41
$W^\pm + c\bar{c} + 3lp$ (incl.)	$\ell\nu + c\bar{c} + 3lp$	90.24 ± 0.31
$W^\pm + b\bar{b} + 2lp$ (excl.)	$\ell\nu + b\bar{b} + 2lp$	90.87 ± 0.46
$W^\pm + b\bar{b} + 3lp$ (incl.)	$\ell\nu + b\bar{b} + 3lp$	90.59 ± 0.30
$\gamma^*/Z + 2lp$ (excl.)	$\tau^\pm\tau^\mp + 2lp$ ($250 < \hat{m} < 1960$ GeV)	80.46 ± 0.66
$\gamma^*/Z + 3lp$ (incl.)	$\mu^\pm\mu^\mp + 3lp$ ($60 < \hat{m} < 130$ GeV)	93.13 ± 0.26
	$\mu^\pm\mu^\mp + 3lp$ ($130 < \hat{m} < 250$ GeV)	92.45 ± 0.41
	$\mu^\pm\mu^\mp + 3lp$ ($250 < \hat{m} < 1960$ GeV)	93.82 ± 0.30
	$\tau^\pm\tau^\mp + 3lp$ ($130 < \hat{m} < 250$ GeV)	81.85 ± 0.64
	$\tau^\pm\tau^\mp + 3lp$ ($250 < \hat{m} < 1960$ GeV)	81.17 ± 0.45
$Z + 4lp$ (excl.)	$\nu\bar{\nu} + 4lp$	93.65 ± 0.40
$Z + 5lp$ (incl.)	$\nu\bar{\nu} + 5lp$	94.41 ± 0.28
$Z + b\bar{b} + 2lp$ (incl.)	$\nu\bar{\nu} + b\bar{b} + 2lp$	93.80 ± 0.40
$W^\pm + W^\mp$ (incl.)		87.78 ± 0.33
$W^\pm + Z$ (incl.)		88.78 ± 0.43
$\gamma^*/Z + \gamma^*/Z$ (incl.)		90.23 ± 0.42
Exemples de points supersymétriques ($m(\tilde{g}), m(\tilde{b}_1), m(\tilde{\chi}_1^0)$) GeV		Efficacités (%)
(350,330,75)		97.71 ± 0.26
(325,305,75)		97.53 ± 0.14
(300,280,75)		96.76 ± 0.31
(275,265,75)		97.57 ± 0.14
(250,230,75)		96.51 ± 0.36

TAB. 4.9 – Efficacités relatives des conditions de déclenchement $JT1_ACO_MHT_HT$ et $MHT30_3CJT5$ pour l'analyse $JT1$. Les incertitudes sont uniquement les incertitudes statistiques.

Processus du modèle standard	Mode de désintégration	Efficacités (%)
$t + \bar{t} + 0lp$ (excl.)	$2\ell + \nu\bar{\nu} + 2b$ $\ell\nu + 2b + 0lp$ $2b + 4lp$	92.13 ± 0.41 93.59 ± 0.10 74.72 ± 0.71
$t + \bar{t} + 1lp$ (excl.)	$2\ell + \nu\bar{\nu} + 2b$ $\ell\nu + 2b + 1lp$ $2b + 5lp$	93.44 ± 0.40 94.29 ± 0.17 75.74 ± 0.65
$t + \bar{t} + 2lp$ (incl.)	$2\ell + \nu\bar{\nu} + 2b$ $\ell\nu + 2b + 2lp$ $2b + 6lp$	94.56 ± 0.17 94.60 ± 0.14 78.10 ± 1.07
$t + q + b$	$\mu\nu b + qb$ $\tau\nu b + qb$	91.92 ± 0.46 91.46 ± 0.59
$t + b$	$e\nu b + b$ $\mu\nu b + b$ $\tau\nu b + b$ $q'\bar{q}b + b$	90.25 ± 0.76 92.84 ± 0.61 91.04 ± 0.54 77.31 ± 0.90
$W^\pm + 4lp$ (excl.)	$\ell\nu + 4lp$	90.98 ± 0.52
$W^\pm + 5lp$ (incl.)	$\ell\nu + 5lp$	92.38 ± 0.66
$W^\pm + b\bar{b} + 3lp$ (incl.)	$\ell\nu + b\bar{b} + 3lp$	91.55 ± 0.56
$\gamma^*/Z + 3lp$ (incl.)	$\tau^\pm\tau^\mp + 3lp$ ($250 < \hat{m} < 1960$ GeV)	90.12 ± 0.52
$Z + 5lp$ (incl.)	$\nu\bar{\nu} + 5lp$	93.77 ± 0.53
Exemples de points supersymétriques ($m(\tilde{g}), m(\tilde{b}_1), m(\tilde{\chi}_1^0)$) GeV		Efficacités (%)
(350, 200, 75)		96.82 ± 0.17
(350, 150, 75)		96.37 ± 0.31
(325, 250, 75)		96.56 ± 0.20
(325, 200, 75)		96.18 ± 0.17
(325, 150, 75)		95.96 ± 0.17
(300, 200, 75)		96.35 ± 0.20
(300, 150, 75)		95.52 ± 0.37

TAB. 4.10 – Efficacités relatives des conditions de déclenchement $JT2_MHT25_HT$ et $MHT30_3CJT5$ pour l'analyse $JT2$. Les incertitudes sont uniquement les incertitudes statistiques.

4.4.2 Le bruit fond instrumental QCD

Le bruit fond instrumental QCD n'est pas simulé. Il faut donc être en mesure d'évaluer la contribution de ce bruit de fond et, de cette façon, montrer qu'après des coupures de sélection, judicieusement choisies, ce bruit de fond devient négligeable. Pour estimer cette contribution, on utilise le fait que le bruit de fond QCD ne possède pas une énergie transverse manquante importante étant donné que les états finaux sont composés exclusivement de jets. La méthode d'estimation du nombre d'événements de QCD , détaillée dans la section suivante, est basée sur cette propriété.

De plus, il a été montré que, quel que soit le point de l'espace des paramètres recherché, l'optimisation finale des coupures donne une énergie transverse manquante systématiquement supérieure à 100 GeV. Cette coupure permet notamment de s'affranchir d'une majeure partie du bruit de fond instrumental. L'optimisation sera détaillée dans la Section 4.5, néanmoins on décide d'ajouter aux coupures de sélection $C1$ des analyses $JT1$ et $JT2$: $E_T > 100$ GeV. En effet, l'étude des événements ayant une énergie transverse manquante inférieure à cette valeur seuil ne présente pas d'intérêt pour le choix des coupures de sélection qui suivent. Par contre, on maintient la coupure initiale à 50 GeV pour les estimations du nombre d'événements du bruit de fond instrumental.

Méthode d'estimation du nombre d'événements

Le bruit de fond QCD est estimé en extrapolant la forme de la distribution d'énergie transverse manquante à partir des basses valeurs ($50 < E_T < 120$ GeV), où ce bruit de fond est dominant, jusqu'aux hautes valeurs, c'est-à-dire vers la région où l'on trouve le signal. Pour calculer la contribution de ce fond aux basses valeurs, on soustrait le nombre d'événements de données réelles par le nombre d'événements estimés du modèle standard. Cette soustraction suppose que les bruits de fond du modèle standard, excepté le bruit de fond instrumental, sont bien connus et bien simulés par les générateurs d'événements *Monte Carlo*. On ajuste alors la distribution du bruit de fond QCD à l'aide de deux fonctions :

$$f(E_T) = a_{pow} \times E_T^{b_{pow}} \quad (4.4)$$

$$g(E_T) = a_{exp} \times e^{-b_{exp} E_T} \quad (4.5)$$

Où a_{pow} , b_{pow} , a_{exp} et b_{exp} sont des constantes à déterminer.

Ces deux fonctions ont tendance à sous-estimer la contribution du bruit de fond QCD , ce qui conduira par la suite à des limites conservatives sur les masses des particules supersymétriques. La loi de puissance fournit toujours une estimation supérieure à la loi exponentielle et surestime même, dans certains cas, la contribution du bruit de fond instrumental. Comme il n'y a aucune raison physique qui conduit à préférer l'une ou l'autre de ces deux fonctions, on n'utilise la moyenne des deux pour notre estimation finale. Le nombre d'événements de QCD , N_{QCD} , est finalement donné par la formule suivante :

$$N_{QCD} = \left(\frac{\int_{E_{Tcut}}^{980} f(x) dx + \int_{E_{Tcut}}^{980} g(x) dx}{2} \right) \pm \left(\frac{\int_{E_{Tcut}}^{980} f(x) dx - \int_{E_{Tcut}}^{980} g(x) dx}{2} \right) \quad (4.6)$$

Où E_{Tcut} est la valeur de la coupure sur l'énergie transverse manquante.

Les estimations du nombre d'événements QCD après les coupures de présélection $C1$ sont données, pour les deux canaux d'analyse, sur la Figure 4.10. Sur cette figure, on trouve également les contributions des différents bruits de fond du modèle standard et les distributions de $\cancel{E}_T$ relatives à deux signaux caractéristiques ($m(\tilde{g}) = 350$, $m(\tilde{b}_1) = 330$ et $m(\tilde{\chi}_1^0) = 75$ pour un signal de type $JT1$, $m(\tilde{g}) = 350$, $m(\tilde{b}_1) = 150$ et $m(\tilde{\chi}_1^0) = 75$ pour un signal de type $JT2$).

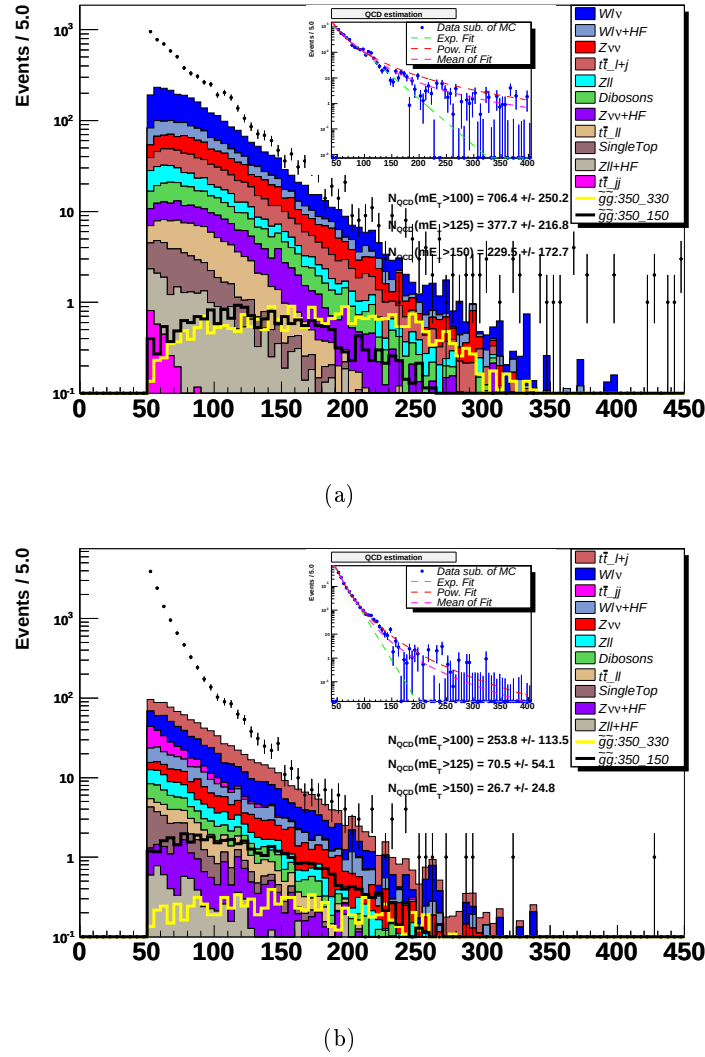


FIG. 4.10 – Distribution de $\cancel{E}_T$ après les coupures de présélection $C1$ pour l'analyse $JT1$ (a) et pour l'analyse $JT2$ (b). Le second cadre, inséré en haut à droite, représente la distribution de $\cancel{E}_T$ pour le bruit de fond QCD et les fonctions d'ajustement : en vert la loi exponentielle, en rouge la loi de puissance et en rose la moyenne des deux.

Après les coupures de présélection $C1$, les Tableaux 4.11 et 4.12 donnent le nombre d'événements attendus pour le bruit de fond, les estimations du bruit de fond QCD , le nombre d'événements effectivement observé, ainsi que le nombre d'événements de signal attendus (pour une

masse moyenne de squarks légers de 1 TeV) à ce stade et l'efficacité correspondante.

Données	$N_{obs} = 1562$
Processus	N_{exp}
$W \rightarrow \ell\nu + jets$	375.2 ± 0.0
$Z \rightarrow \nu\bar{\nu} + jets$	150.8 ± 0.0
$t + \bar{t} \rightarrow \ell\nu + 2b + jets$	109.9 ± 0.6
$W \rightarrow \ell\nu + HF + jets$	108.4 ± 0.4
VV	45.5 ± 0.9
$Z \rightarrow \ell^+\ell^- + jets$	39.6 ± 0.0
$Z \rightarrow \nu\bar{\nu} + HF + jets$	39.0 ± 0.3
$t + \bar{t} \rightarrow 2\ell + \nu\bar{\nu} + 2b + jets$	23.1 ± 0.2
$Z \rightarrow \ell^+\ell^- + HF + jets$	7.9 ± 0.0
$t + q + b, t + b$	6.3 ± 0.0
$t + \bar{t} \rightarrow 2b + jets$	0.0 ± 0.0
Bruit de fond QCD	706.4 ± 250.2
Total	1612.1 ± 252.6
(350,150,75)	15.9 ± 0.5 (12.0 %)
(350,330,75)	24.5 ± 0.5 (18.5 %)

TAB. 4.11 – Nombres d'événements restant après les coupures de présélection $C1$ de l'analyse $JT1$ avec la coupure supplémentaire à 100 GeV sur $\cancel{E}_T$. Seules les incertitudes statistiques sont indiquées, exceptée l'incertitude sur le bruit de fond instrumental. HF désigne les saveurs lourdes dans l'état final, V les bosons de l'interaction électrofaible.

On constate que l'estimation du nombre d'événements de QCD est très approximative mais reste néanmoins un bon indicateur du bruit de fond non simulé. Cet indicateur guide les coupures qui suivent ; elles visent en effet à réduire au maximum ce bruit de fond dominant tout en conservant de hautes efficacités de signal supersymétrique. Ces deux tableaux permettent également de vérifier que le signal avec de fortes différences de masses est plus sensible à l'analyse $JT2$, alors que le deuxième exemple de signal est plus sensible à l'analyse $JT1$.

Les jets mal reconstruits ou mal identifiés

On considère qu'un jet, reconstruit par les algorithmes détaillés dans la Section 3.4.4, est mauvais s'il ne vérifie pas les critères d'identification donnés dans cette même section. Ces jets peuvent être des jets de bruit, ils créent alors de l'énergie transverse manquante. Ils peuvent être aussi des vrais jets mais non reconnus comme tel. Ces derniers n'étant pas corrigés de l'échelle d'énergie des jets (JES), ils contribuent par conséquent à l'énergie transverse manquante. Les événements du bruit de fond instrumental possédant plusieurs jets dans l'état final, ils sont donc plus sensibles à cette inefficacité d'identification des jets et possèdent une plus grande fraction de mauvais jets.

Afin de réduire le bruit de fond instrumental, on applique des coupures de sélection sur les mauvais jets. On coupe d'une part sur l'énergie transverse des mauvais jets, dont la distribution est donnée sur la Figure 4.11, et d'autre part sur la fraction CH de ces mauvais jets, dont la

Données	$N_{obs} = 638$
Processus	N_{exp}
$t + \bar{t} \rightarrow \ell\nu + 2b + jets$	110.6 ± 0.6
$W \rightarrow \ell\nu + jets$	78.5 ± 0.0
$Z \rightarrow \nu\bar{\nu} + jets$	24.5 ± 0.0
$W \rightarrow \ell\nu + HF + jets$	23.0 ± 0.1
$Z \rightarrow \ell^+\ell^- + jets$	10.1 ± 0.0
VV	7.3 ± 0.4
$t + \bar{t} \rightarrow 2\ell + \nu\bar{\nu} + 2b + jets$	6.2 ± 0.1
$t + \bar{t} \rightarrow 2b + jets$	5.9 ± 0.2
$Z \rightarrow \nu\bar{\nu} + HF + jets$	5.9 ± 0.1
$Z \rightarrow \ell^+\ell^- + HF + jets$	2.7 ± 0.0
$t + q + b, t + b$	2.4 ± 0.0
Bruit de fond QCD	253.8 ± 113.5
Total	530.9 ± 115.0
(350,150,75)	25.1 ± 0.5 (19.1 %)
(350,330,75)	8.1 ± 0.3 (6.1 %)

TAB. 4.12 – Nombres d'événements restant après les coupures de présélection $C1$ de l'analyse $JT2$ avec la coupure supplémentaire à 100 GeV sur $\cancel{E}_T$. Seules les incertitudes statistiques sont indiquées, exceptée l'incertitude sur le bruit de fond instrumental. HF désigne les saveurs lourdes dans l'état final, V les bosons de l'interaction électrofaible.

distribution est donnée sur la Figure 4.12. Ainsi, on ne sélectionne pas les événements possédant un mauvais jet vérifiant : $E_T > 20$ GeV ou $CHF > 0.02$. Ces nouvelles coupures définissent le niveau de sélection $C2$ pour les deux canaux d'analyse.

Les coupures topologiques

Le bruit de fond instrumental ne possède pas d'énergie transverse manquante physique, c'est-à-dire de l'énergie transverse manquante due à des particules non-détectées comme les neutralinos. L'essentiel de la contribution à $\cancel{E}_T$, une fois les bruits dans le calorimètre nettoyés, provient d'erreur dans la mesure de l'énergie des jets. Par conséquent, le vecteur $\vec{\cancel{E}}_T$ est, dans la majorité des cas, aligné avec l'un des jets de l'événement. La coupure sur $\Delta\phi_{min}(\cancel{E}_T, jets)$, effectuée pour l'analyse $JT1$, enlève une grande partie des événements dont un des jets est aligné avec le vecteur $\vec{\cancel{E}}_T$. On ajoute uniquement une coupure minimale sur $\Delta\phi_{min}(\cancel{E}_T, j_1)$ à 100° pour ce canal d'analyse.

Pour le canal d'analyse $JT2$, trois coupures supplémentaires permettent de réduire le bruit de fond instrumental tout en gardant une haute efficacité de signal :

$$\begin{cases} \Delta\phi_{min}(\cancel{E}_T, j_1) > 80^\circ \\ \Delta\phi_{min}(\cancel{E}_T, j_2) > 20^\circ \\ \Delta\phi_{min}(\cancel{E}_T, j_3) > 20^\circ \end{cases}$$

Les distributions de $\Delta\phi_{min}(\cancel{E}_T, j_1)$, $\Delta\phi_{min}(\cancel{E}_T, j_2)$ et $\Delta\phi_{min}(\cancel{E}_T, j_3)$ avant les coupures précédentes sont données sur la Figure 4.13.

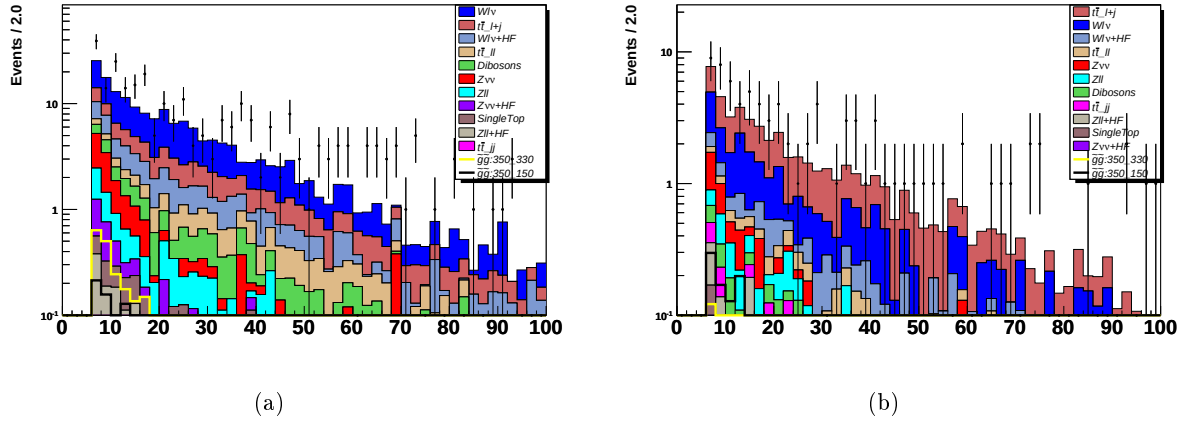


FIG. 4.11 – Distribution du E_T des mauvais jets après les coupures de présélection $C1$, avec la coupure supplémentaire à 100 GeV sur $\cancel{E}_T$, pour l'analyse $JT1$ (a) et pour l'analyse $JT2$ (b).

Des coupures maximales sur ces trois variables ne sont pas possibles, car elles correspondraient à la région du signal recherché. Ces coupures topologiques définissent le niveau de sélection $C3$.

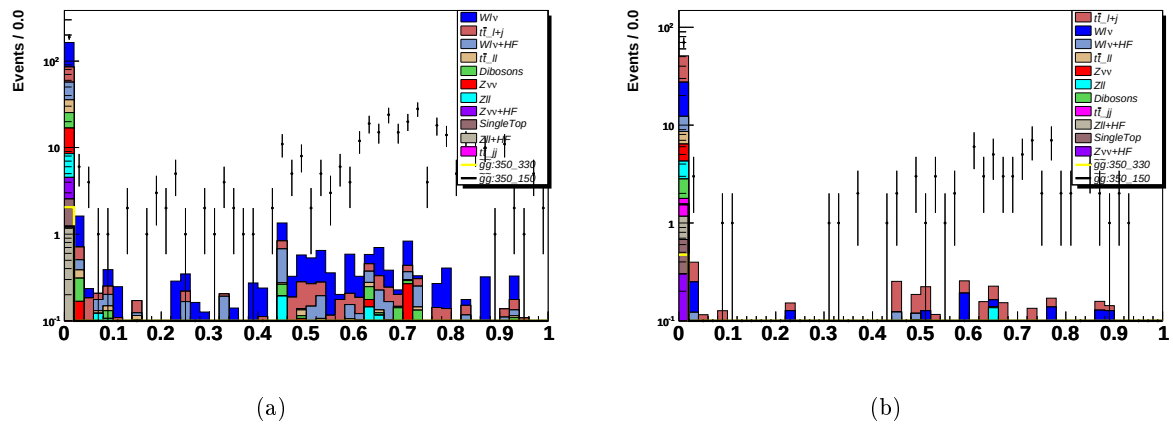


FIG. 4.12 – Distribution de la fraction CHF des mauvais jets après les coupures de présélection $C1$, avec la coupure supplémentaire à 100 GeV sur $\cancel{E}_T$, pour l'analyse $JT1$ (a) et pour l'analyse $JT2$ (b).

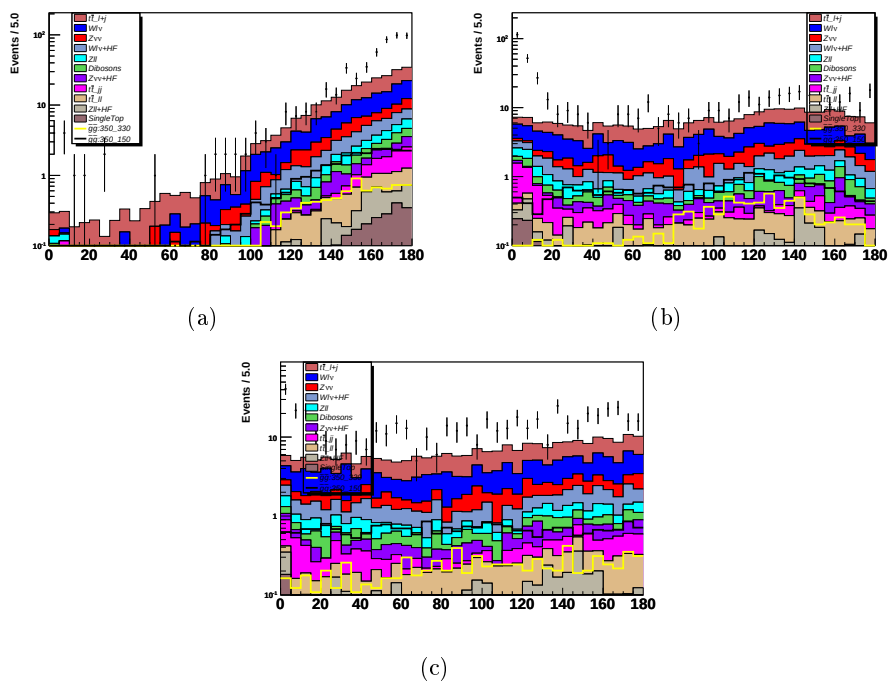


FIG. 4.13 – Distribution de $\Delta\phi_{min}(\not{E}_T, j_1)$ (a), $\Delta\phi_{min}(\not{E}_T, j_2)$ (b) et $\Delta\phi_{min}(\not{E}_T, j_3)$ (c) après les coupures de présélection $\mathcal{C}^2$, avec la coupure supplémentaire à 100 GeV sur $\not{E}_T$, pour l'analyse $JT2$.

4.4.3 La confirmation des jets par le trajectographe interne

Afin d'augmenter la qualité des lots de données utilisés par rapport aux bruits éventuels dans le calorimètre, par rapport aux jets mal mesurés ou par rapport aux rayons cosmiques déposant une grande quantité d'énergie dans le calorimètre, on impose une confirmation des jets identifiés dans le calorimètre par le trajectographe interne. Pour cette confirmation, on utilise une variable appelée CPF défini pour un jet et un vertex primaire donnés.

Définition de la variable CPF_0

Pour un jet i et un vertex primaire j , on définit la variable $CPF(j_i, PV_j)$ de la façon suivante :

$$CPF(j_i, PV_j) = \frac{\sum_{trk} p_T^{trk}(j_i, PV_j)}{\sum_{k=1}^{N_{PV}} \sum_{trk} p_T^{trk}(j_i, PV_k)} \quad (4.7)$$

Où $p_T^{trk}(j_i, PV_j)$ est une trace provenant du vertex primaire j et contenu dans un cône autour du jet i .

Schématiquement, on somme l'impulsion transverse de toutes les traces dans le cône provenant du vertex primaire considéré, trait bleu sur la Figure 4.14, et on divise cette somme par la somme de l'impulsion transverse de toutes les traces dans le cône, traits bleu et rouge.

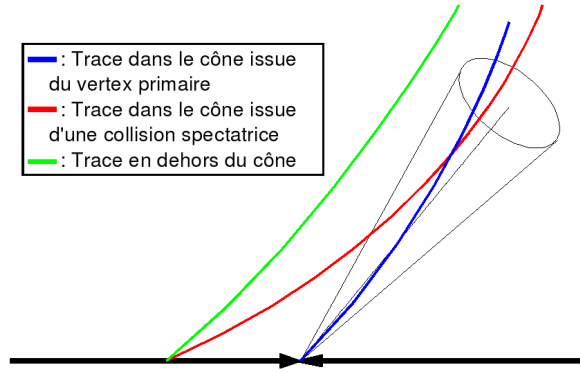


FIG. 4.14 – Les traces pour la définition de la variable CPF_0 : seules les traces représentées par les traits rouge et bleu sont utilisées.

Pour un jet donné, on définit CPF_0 pour le vertex primaire de l'interaction principale de la collision. Si le dénominateur de la Formule 4.7 est nul, alors $CPF_0 = -1$. Dans la suite, on considère qu'un jet est confirmé par le trajectographe interne, si $CPF_0 > 0.85$. Cette définition n'est possible que pour des jets centraux, ce qui est le cas des deux jets de plus hautes énergies sur lesquels cet algorithme sera appliqué pour les deux canaux d'analyse. En effet, si un jet est situé dans un des bouchons du calorimètre, il est grandement possible qu'aucune trace ne lui soit associé et donc que : $CPF_0 = -1$. L'événement ne doit pas être rejeté pour autant.

On définit alors trois probabilités distinctes pour un jet d'être confirmé par les traces :

- $Prob(-1)$: probabilité qu'il y ait au moins une trace dans le cône, i.e. $CPF_0 \neq -1$;
- $Prob(0)$: probabilité que le jet ne provienne pas d'une collision multiple $p\bar{p}$, i.e. $CPF_0 \neq 0$;
- $Prob(0.85)$: probabilité que le jet soit confirmé par les traces, i.e. $CPF_0 > 0.85$, sachant les deux précédentes conditions vérifiées.

Ces trois définitions sont indépendantes. La probabilité totale pour un jet donné d'être confirmé par les traces est le produit de ces trois probabilités : $Prob(-1) \times Prob(0) \times Prob(0.85)$. On étudie alors les dépendances de chacune de ces trois probabilités en fonction de variables caractéristiques de l'événement comme la position du vertex primaire ou le nombre de vertex primaire et en fonction de variables caractéristiques du jet comme son énergie transverse et sa position en η . Cette étude est menée aussi bien pour les données du détecteur que pour les données simulées.

Comparaisons entre données du détecteur et données *Monte Carlo*

La simulation des traces est très complexe et amène souvent à des différences d'efficacité entre les données réelles et les données simulées pour les coupures basées sur les informations fournies par le trajectographe interne. La différence est due essentiellement à la difficulté de simuler les bruits qui peuvent provoquer des impacts dans le *SMT* ou le *CFT*. Ces bruits étant mal simulés, on surestime, en général, l'efficacité des algorithmes basés sur les informations du trajectographe interne, tel l'algorithme *CPF₀*. Pour cette raison, on compare l'efficacité de la coupure sur *CPF₀* entre les données du détecteur et les données *Monte Carlo*.

Pour cette comparaison, on utilise un lot de données riche en jets. Pour les données réelles, on utilise les données enregistrées par les conditions de déclenchement propres aux études *QCD*. Ces conditions sont satisfaites si au moins un jet a une énergie transverse supérieur à un seuil donné. Pour les données *Monte Carlo*, on utilise une génération d'événements de *QCD* par le générateur Pythia.

Pour enrichir le lot de données réelles en événements *QCD* et ainsi le comparer au lot de données simulées, on impose également les coupures suivantes :

$$\left\{ \begin{array}{l} \text{Critères de qualité des données} \\ \text{Veto sur les mauvais jets avec } E_T > 20 \text{ GeV} \\ \text{Au moins un vertex primaire reconstruit (PV), avec : } |z_{PV}| < 60 \text{ cm} \\ \text{Exactement 2 jets } (j_1, j_2) \text{ avec } E_T > 40 \text{ GeV et } |\eta_{\text{det}}(j_1, j_2)| < 0.8 \\ \Delta\phi(j_1, j_2) > 170^\circ \end{array} \right.$$

La comparaison des deux lots est finalement possible en normalisant le lot de données simulées, de telle façon à avoir autant d'événements dans les deux lots. On obtient finalement les distributions de la Figure 4.15, qui montre un bon accord pour l'énergie transverse des deux jets de plus hautes énergies transverses, ainsi que pour l'énergie transverse manquante. L'étude de l'efficacité de la coupure sur *CPF₀* est effectuée sur ces deux lots de données pour les événements sélectionnés par les coupures précédentes.

Correction du nombre de vertex et application de la coupure

On étudie la dépendance des trois probabilités définies précédemment en fonction de l'énergie transverse du jet, E_T , de la position en η du jet, η_{det} ($= \eta \times 10$), de la position selon l'axe z du meilleur vertex primaire de l'événement, z_{vtx} , et du nombre de vertex primaires de l'événement.

Dépendance de la probabilité $Prob(-1)$ $Prob(-1)$ est définie comme la probabilité d'associer des traces aux jets. Cette probabilité dépend essentiellement de la position du jet dans le calorimètre et de la position du vertex primaire de l'événement. Ces deux dépendances sont

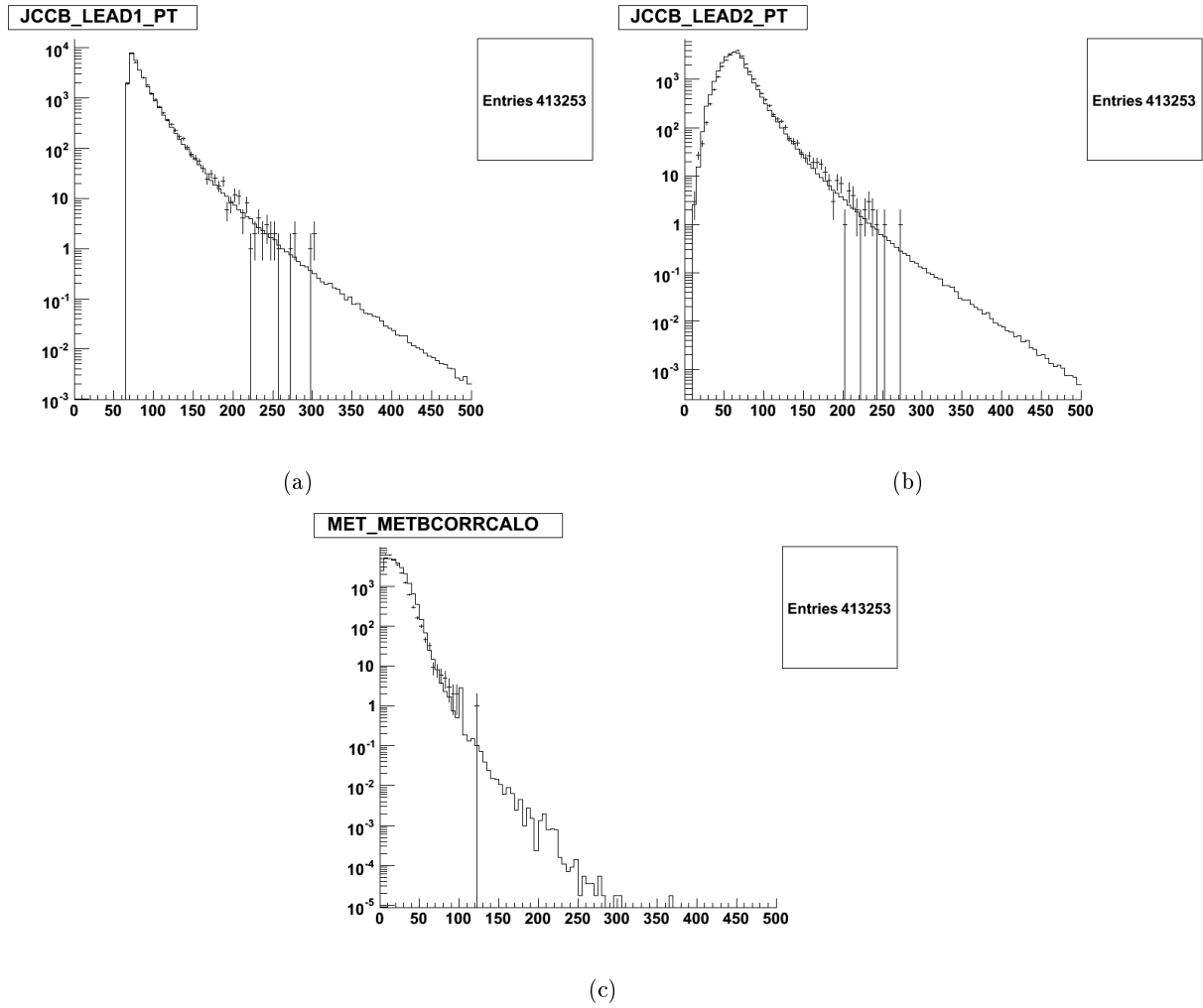


FIG. 4.15 – Comparaison entre les données réelles et les données *Monte Carlo* pour les distributions du p_T du premier jet (a), du second jet (b) et pour la distribution de $\cancel{E}_T$ (c).

données sur la Figure 4.16 pour les données réelles et les données simulées. On constate que la probabilité d'associer des traces au jet est stable dans la partie centrale du calorimètre et qu'elle est de l'ordre de 100 % pour les données simulées et entre 98 et 99 % pour les données réelles.

Dépendance de la probabilité $Prob(0)$ Cette probabilité définit la probabilité que le jet ne soit pas issu d'une collision multiple $p\bar{p}$, elle dépend donc fortement du nombre de vertex primaires dans l'événement, mais également de l'énergie transverse du jet. En effet, les jets issus de l'interaction principale de l'événement ont, dans la plupart des cas, une plus haute énergie transverse. Ces deux dépendances sont données sur la Figure 4.17. La dépendance avec le nombre de vertex primaires, ou avec la luminosité instantanée, diminue fortement pour les données réelles mais très légèrement pour les données simulées. Cet effet est donc à corriger. On constate également que la probabilité atteint un plateau pour des énergies transverses supérieures à 30-40 GeV.

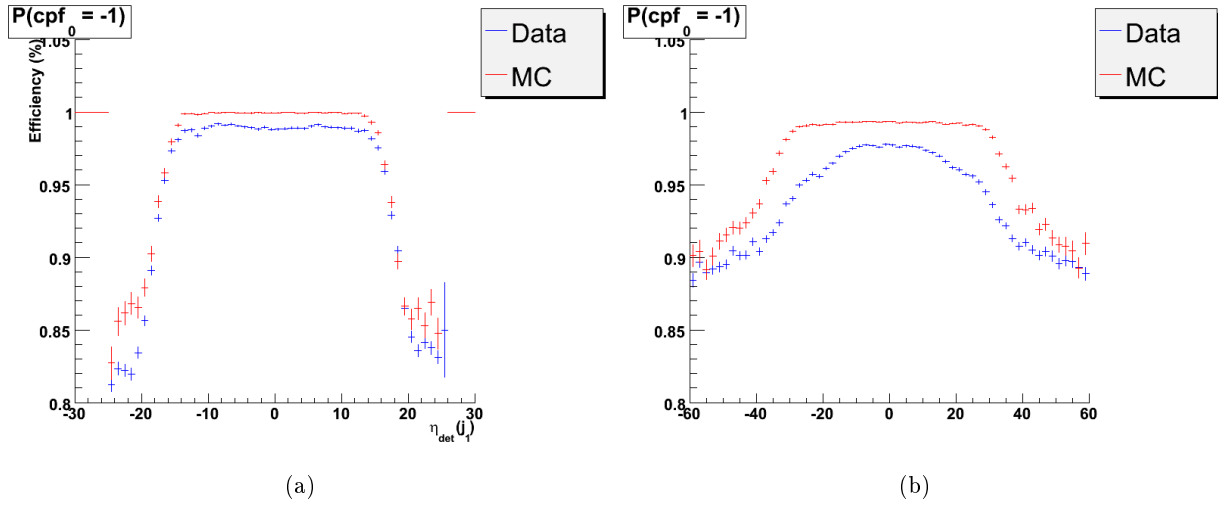


FIG. 4.16 – Dépendance de $Prob(-1)$ en fonction de η_{det} du jet le plus dur (a) et en fonction de z_{vtx} (b) pour les données réelles et les données simulées.

Dépendance de la probabilité $Prob(0.85)$ On donne deux exemples de dépendance pour la probabilité de confirmer un jet par le trajectographe interne, la dépendance en η_{det} et en E_T , sur la Figure 4.18. Pour des jets centraux et pour les données simulées, on observe que cette probabilité est quasiment constante quelle que soit l'énergie transverse du jet. Le comportement des données réelles est différent. Cette probabilité se stabilise à partir d'un seuil en énergie transverse de l'ordre de 40 GeV

On construit alors un facteur de correction qui est le rapport de la probabilité pour les données réelles sur la probabilité pour les données simulées. Pour des jets centraux, l'efficacité de confirmer un jet par les traces ne dépend pas significativement de η_{det} . La dépendance en E_T du facteur de correction peut être paramétrée comme le montre la Figure 4.19. On constate que pour des valeurs supérieures à 40-45 GeV ce facteur correctif est constant et de l'ordre de 99.2 %. Pour l'analyse présentée dans cette thèse, seuls les deux jets de plus hautes énergies transverses sont confirmés par le trajectographe interne. Ces jets vérifient : $E_T > 40$ GeV et $|\eta_{det}| < 0.8$. On considère alors qu'on est sur les plateaux d'efficacité pour ces deux variables caractéristiques du jet. De plus, on peut montrer que, dans ces conditions, le facteur correctif à appliquer aux données *Monte Carlo* varie très peu en fonction de la position du principal vertex primaire de l'événement.

Corrections à appliquer aux données *Monte Carlo* pour des jets centraux de $E_T > 40$ GeV Etant sur les plateaux d'efficacité des variables η_{det} , z_{vtx} et E_T , on peut considérer que les différences d'efficacité entre les données réelles et les données *Monte Carlo* ne dépendent que du nombre de vertex primaires dans l'événement, c'est-à-dire de la luminosité (voir le paragraphe qui décrit les collisions multiples $p\bar{p}$ ajoutées aux données simulées dans la Section 4.3).

On calcule alors le rapport de l'efficacité trouvée dans les données réelles et dans les données *Monte Carlo* pour déterminer les facteurs correctifs à appliquer aux données simulées. La dépendance de ce rapport pour chacune des trois probabilités en fonction du nombre de vertex primaire est ajustée par une fonction linéaire. On calcule alors trois paramétrisations pour le jet de plus

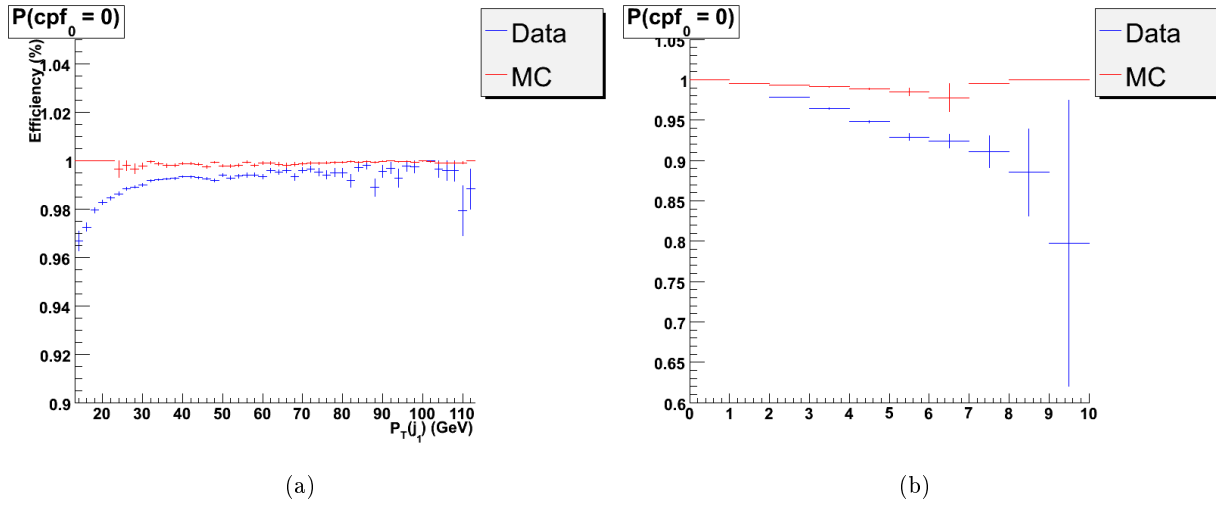


FIG. 4.17 – Dépendance de $Prob(0)$ en fonction de E_T du jet le plus dur (a) et en fonction du nombre de vertex (b) pour les données réelles et les données simulées.

haute énergie transverse, comme le montre la Figure 4.20. On calcule, de même, trois nouvelles paramétrisations pour le deuxième jet, sachant que le premier vérifie $CPF_0 > 0.85$. La Figure 4.21 donne les variations des facteurs correctifs en fonction du nombre de vertex primaire pour le deuxième jet.

Les fonctions d'ajustement donnent finalement les facteurs correctifs suivants, pour le jet de plus haute énergie transverse :

$$\begin{cases} CPF_0 \neq 0 : & 1. - 0.003467 \times (N_{PV} - 1) \\ CPF_0 \neq -1 : & 0.9958 - 0.0003442 \times (N_{PV} - 1) \\ CPF_0 > 0.85 : & 0.9959 - 0.02817 \times (N_{PV} - 1) \end{cases}$$

Pour le deuxième jet de plus haute énergie transverse, sachant le premier confirmé par le trajectographe interne, on obtient :

$$\begin{cases} CPF_0 \neq 0 : & 0.9975 - 0.0007508 \times (N_{PV} - 1) \\ CPF_0 \neq -1 : & 1. - 0.00004998 \times (N_{PV} - 1) \\ CPF_0 > 0.85 : & 0.9976 - 0.01426 \times (N_{PV} - 1) \end{cases}$$

La plus grosse inefficacité provient de la troisième probabilité, c'est-à-dire la probabilité que la condition $CPF_0 > 0.85$ soit vérifiée. Elle augmente approximativement de 3 % par nombre de vertex primaire dans l'événement. L'inefficacité pour le deuxième jet est deux fois moins importante. On constate que, pour un seul vertex primaire dans l'événement, l'inefficacité due à la coupure sur CPF_0 est quasiment nulle. Ceci est dû essentiellement au fait que les deux jets sélectionnés sont des jets centraux.

Correction du nombre de vertex La différence d'efficacité entre les données réelles et les données simulées dépend uniquement du nombre de vertex primaires pour les deux canaux d'analyse présentés dans ce chapitre, car l'algorithme est appliqué à des jets centraux de hautes énergies

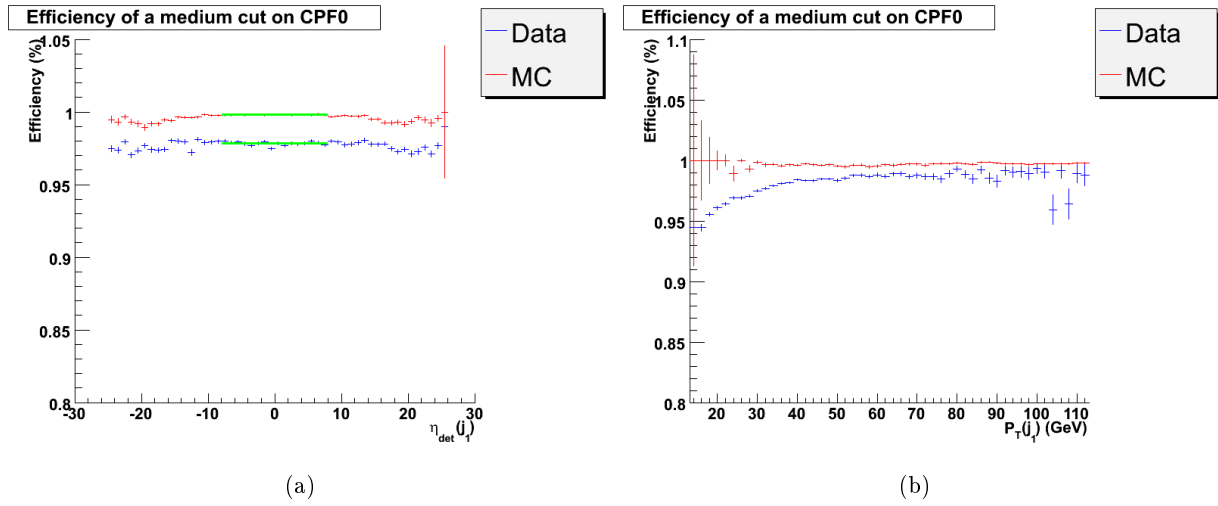


FIG. 4.18 – Dépendance de $Prob(0.85)$ en fonction de η_{det} (a) et du E_T (b) du jet le plus dur pour les données réelles et les données simulées.

transverses. Or, le profil de luminosité, fortement corrélé au nombre de vertex primaires, diffère dans les données *Monte Carlo* et les données du détecteur. Pour appliquer correctement la coupure sur la variable CPF_0 , il est nécessaire de corriger la distribution du nombre de vertex aussi bien pour le fond simulé du modèle standard que pour les signaux supersymétriques. Cette correction est effectuée au niveau *C3* des coupures de sélection.

On applique alors un facteur correctif fonction du nombre de vertex primaires de l'événement, de telle façon à ce que la forme des distributions, montrées sur les Figures 4.22 et 4.23, soit identique pour les données réelles et les données simulées tout en conservant le nombre total d'événements simulés.

Les facteurs correctifs obtenus, à partir des distributions des Figures 4.22 et 4.23, sont résumés dans le Tableau 4.13.

N_{vertex}	Analyse <i>JT1</i>		Analyse <i>JT2</i>	
	Bruits de fond	Signaux	Bruits de fond	Signaux
1	0.707	0.749	0.580	0.618
2	1.085	1.112	1.212	1.224
3	1.960	1.580	2.293	1.975
4	2.612	1.578	2.025	1.307
5	7.332	3.185	4.523	2.422
≥ 6	19.750	3.945	14.826	3.108

TAB. 4.13 – Facteurs correctifs à appliquer à la distribution du nombre de vertex primaires pour avoir une simulation correcte du profil de luminosité.

La coupure sur CPF_0 est finalement appliquée pour les deux jets de plus hautes énergies transverses après la correction du profil de luminosité des événements *Monte Carlo*. On atteint alors le niveau *C4* de sélection pour les deux canaux d'analyse.

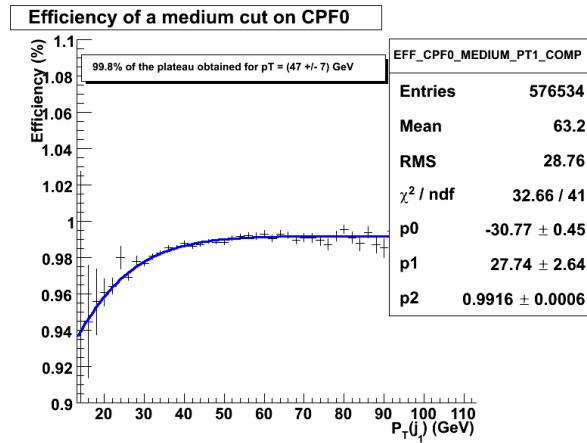


FIG. 4.19 – Facteur de corrections à appliquer aux données simulées en fonction du E_T du jet le plus dur, si ce jet est situé dans le calorimètre central..

La variation du nombre d'événements observés et du nombre d'événements attendus du niveau de sélection $C1$ au niveau $C4$ est résumée dans les Tableaux 4.14 et 4.15 pour les analyses $JT1$ et $JT2$ avec une coupure à 100 GeV sur l'énergie transverse manquante. Ces tableaux donnent plusieurs informations. La première est que les deux signaux ont des sensibilités différentes aux deux canaux d'analyse, c'est ce qui était recherché initialement. L'analyse $JT1$ est plus sensible à une faible différence de masse entre deux particules supersymétriques, alors que l'analyse $JT2$ sélectionne plutôt les signaux avec de fortes différences de masse. Deuxièmement, les différentes coupures de sélection ont permis de réduire fortement le fond instrumental non-simulé. Les coupures topologiques ont été très efficaces pour le canal d'analyse $JT2$, les coupures sur les mauvais jets et la confirmation des jets par le trajectographe interne ont permis de réduire ce bruit de fond pour l'analyse $JT1$. Finalement, à ce niveau de sélection, le bruit de fond instrumental n'est plus le bruit de fond dominant pour les deux analyses.

Niveau de sélection	$C1$	$C2$	$C3$	$C4$
Données réelles	1562	1179	1169	886
Fond physique	905.7 ± 2.4	812.0 ± 2.2	805.7 ± 2.2	746.2 ± 2.4
Fond instrumental	706.0 ± 250.2	147.3 ± 61.9	148.5 ± 61.7	35.3 ± 18.6
(350,150,75)	12.0 %	11.5 %	10.7 %	10.1 %
(350,330,75)	18.5 %	18.1 %	18.0 %	16.9 %

TAB. 4.14 – Récapitulatif des nombres d'événements observés et attendus du niveau de sélection $C1$ au niveau $C4$ pour l'analyse $JT1$.

Niveau de sélection	$C1$	$C2$	$C3$	$C4$
Données réelles	638	525	221	196
Fond physique	277.1 ± 1.5	247.3 ± 1.4	193.8 ± 1.1	177.8 ± 1.1
Fond instrumental	253.8 ± 113.5	204.3 ± 91.8	8.9 ± 5.0	4.4 ± 2.6
(350,150,75)	19.1 %	18.6 %	14.9 %	13.8 %
(350,330,75)	6.1 %	6.1 %	5.3 %	5.1 %

TAB. 4.15 – Récapitulatif des nombres d'événements observés et attendus du niveau de sélection $C1$ au niveau $C4$ pour l'analyse $JT2$.

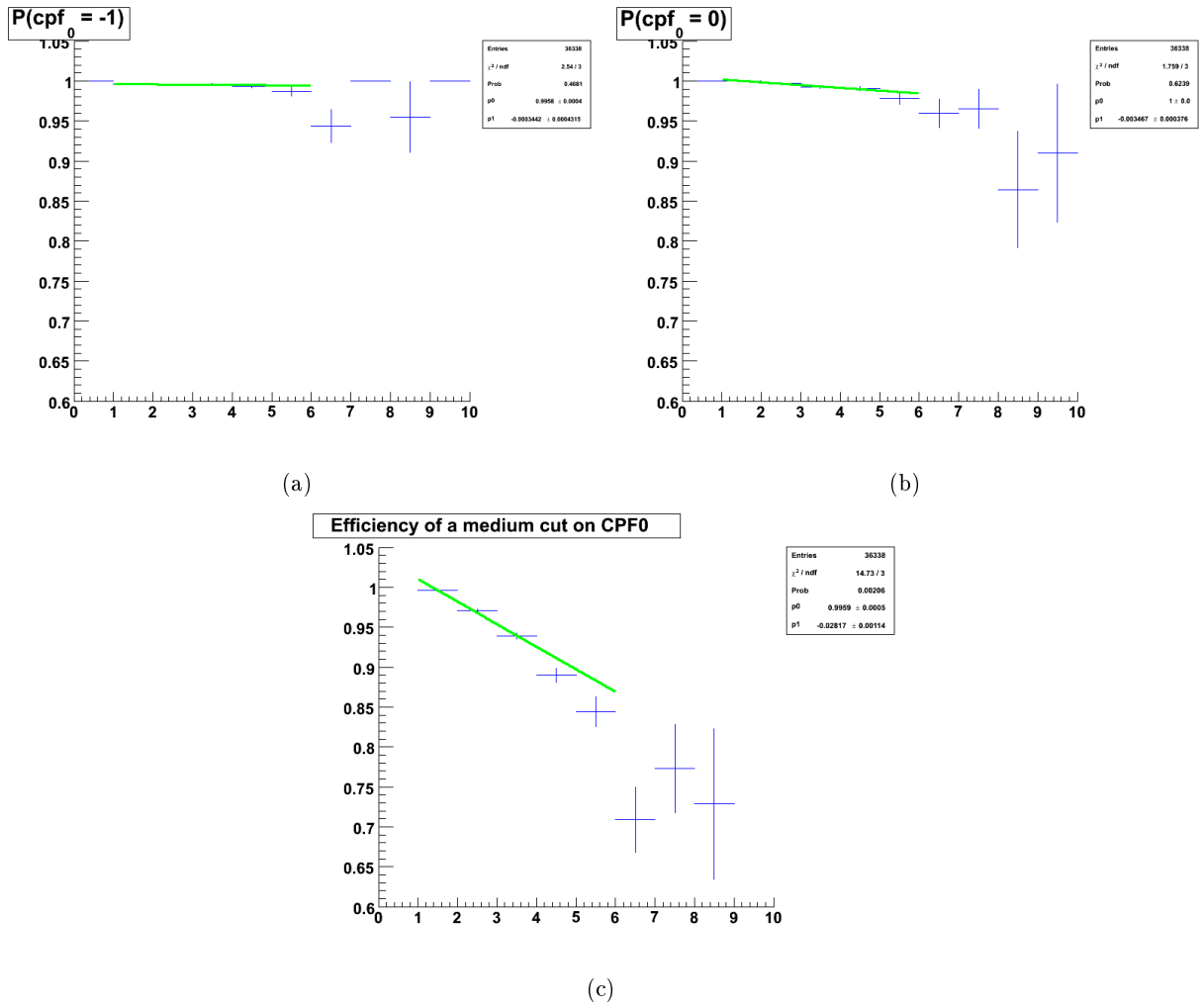


FIG. 4.20 – Facteurs de correction à appliquer aux données simulées en fonction du nombre de vertex pour les probabilités $Prob(-1)$ (a), $Prob(0)$ (b) et $Prob(0.85)$ (c). Ces facteurs sont valables pour une confirmation du jet le plus dur de l'événement.

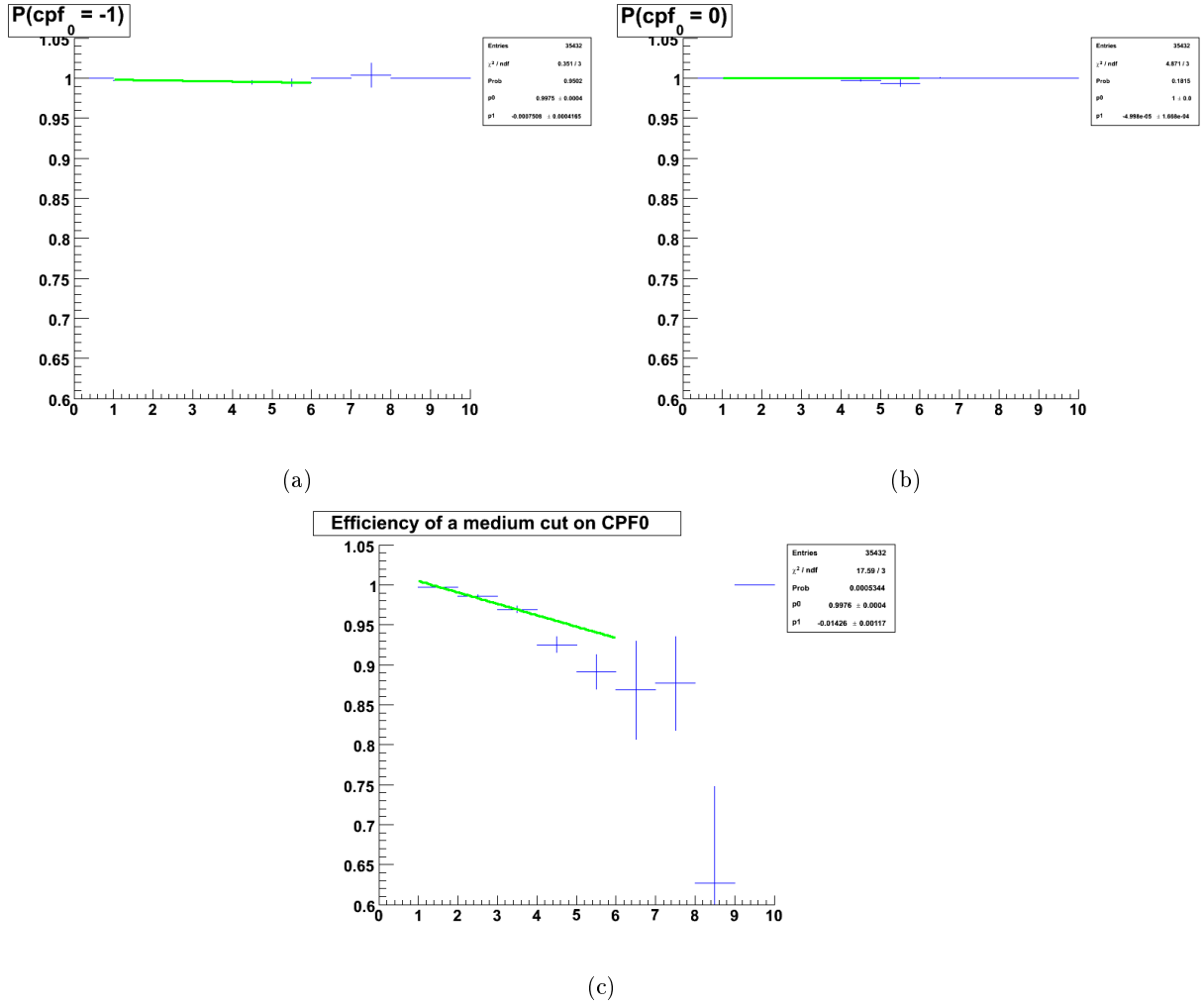


FIG. 4.21 – Facteurs de correction à appliquer aux données simulées en fonction du nombre de vertex pour les probabilités $Prob(-1)$ (a), $Prob(0)$ (b) et $Prob(0.85)$ (c). Ces facteurs sont valables pour une confirmation du second jet le plus dur de l'événement sachant que le jet le plus dur est également confirmé.

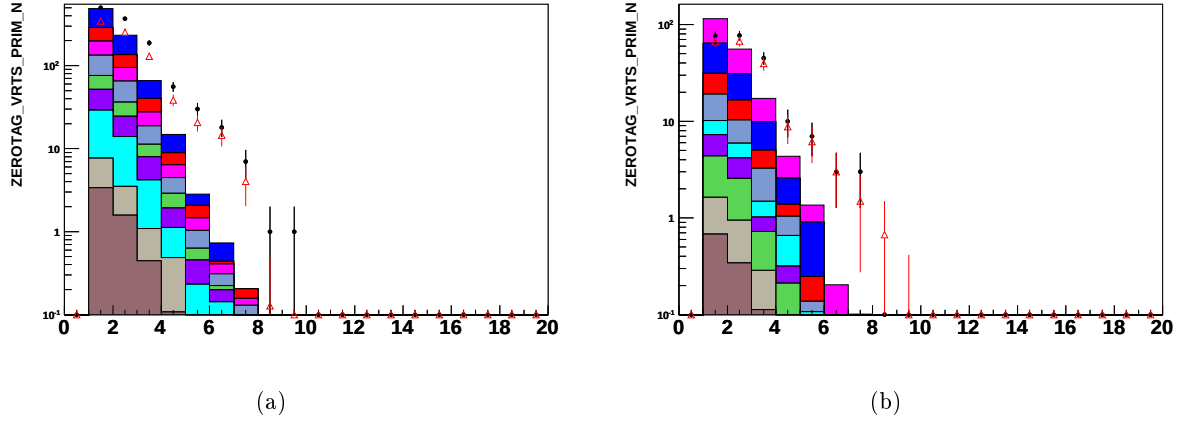


FIG. 4.22 – Distributions du nombre de vertex pour les données réelles (ronds noirs), les fonds simulés du modèle standard (histogrammes de couleur) et les fonds du modèle standard une fois la correction appliquée (triangles rouges) pour l'analyse *JT1* (a) et pour l'analyse *JT2* (b).

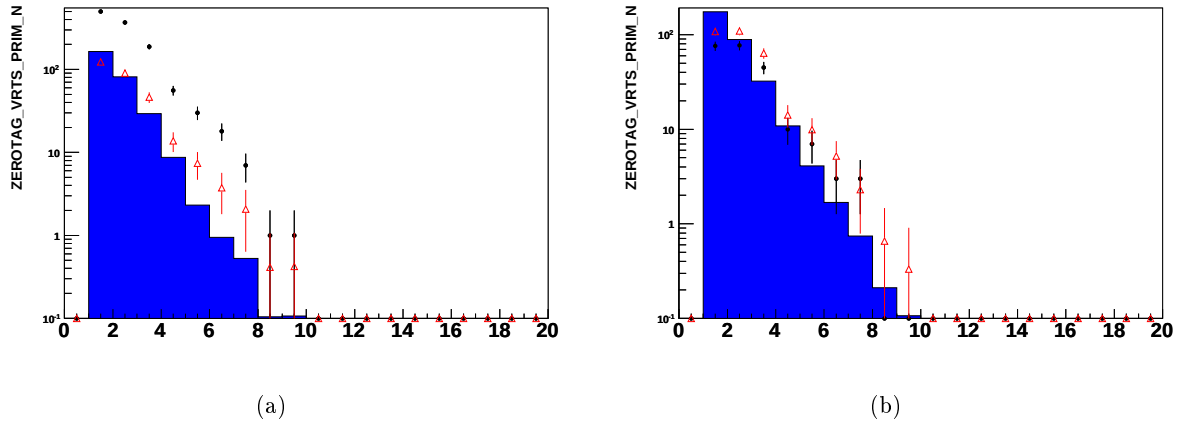


FIG. 4.23 – Distributions du nombre de vertex pour les données réelles (ronds noirs), les signaux simulés (histogramme bleu) et les signaux une fois la correction appliquée (triangles rouges) pour l'analyse *JT1* (a) et pour l'analyse *JT2* (b).

4.4.4 Coupures sur le nombre de leptons et zones de contrôle du bruit de fond physique

Au niveau $C4$ des coupures de sélection, quatre bruits de fond dominant pour les deux canaux d'analyses comme le montre le Tableau 4.16. Ce tableau donne la fraction d'événements attendus pour les quatre bruits de fond physiques dominants. Cette fraction est calculée par rapport aux nombres d'événements observés. Pour ce calcul, on fait l'hypothèse que la différence entre le nombre d'événements attendus et le nombre d'événements observés est exclusivement constituée du bruit de fond instrumental.

	Analyse $JT1$	Analyse $JT2$
Données réelles	886	196
Total des bruits de fond physique	746.2 ± 2.4	177.8 ± 1.1
$W \rightarrow \ell\nu + jets$	34 %	25 %
$Z \rightarrow \nu\bar{\nu} + jets$	15 %	10 %
$W \rightarrow \ell\nu + HF + jets$	10 %	8 %
$t + \bar{t} \rightarrow \ell\nu + 2b + jets$	9 %	37 %

TAB. 4.16 – Fraction du nombre d'événements attendus pour les quatre bruits de fond dominants du modèle standard.

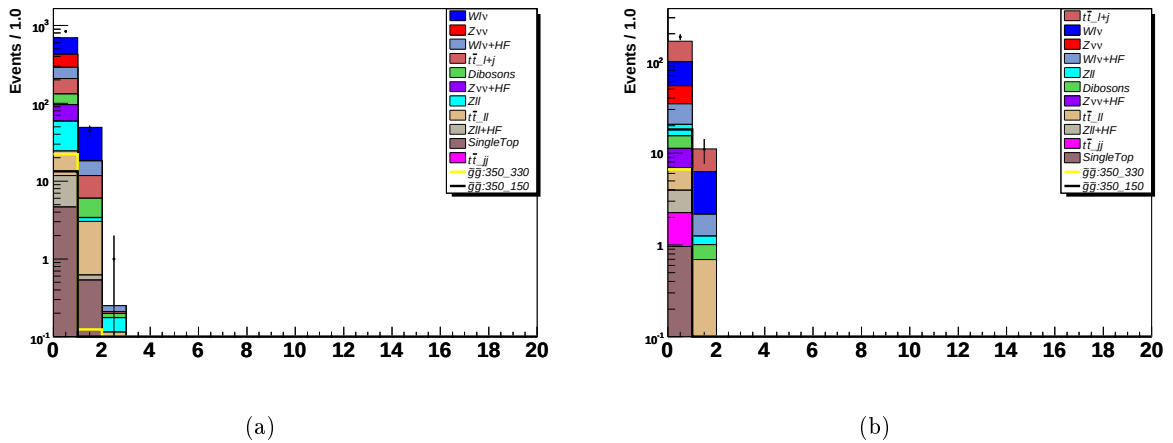


FIG. 4.24 – Distributions du nombre d'électrons isolés au niveau de sélection $C4$ pour l'analyse $JT1$ (a) et pour l'analyse $JT2$ (b).

On constate que trois de ces bruits de fond possèdent un lepton dans l'état final. On utilise alors un veto sur les électrons isolés et sur les muons isolés pour améliorer la sensibilité au signal. Ce veto constitue le niveau de sélection $C5$. Les électrons choisis pour cette nouvelle coupure de sélection ont les caractéristiques suivantes (voir la Section 3.4.3 pour plus de détails sur l'identification des électrons) : $emf > 0.9$, $iso < 0.2$, $HMx7 < 50$ et $E_T > 5$ GeV. Pour différencier les électrons des photons, on utilise un critère supplémentaire, qui demande l'association d'une trace avec le candidat électromagnétique sélectionné. On effectue également un veto sur les muons isolés. Les

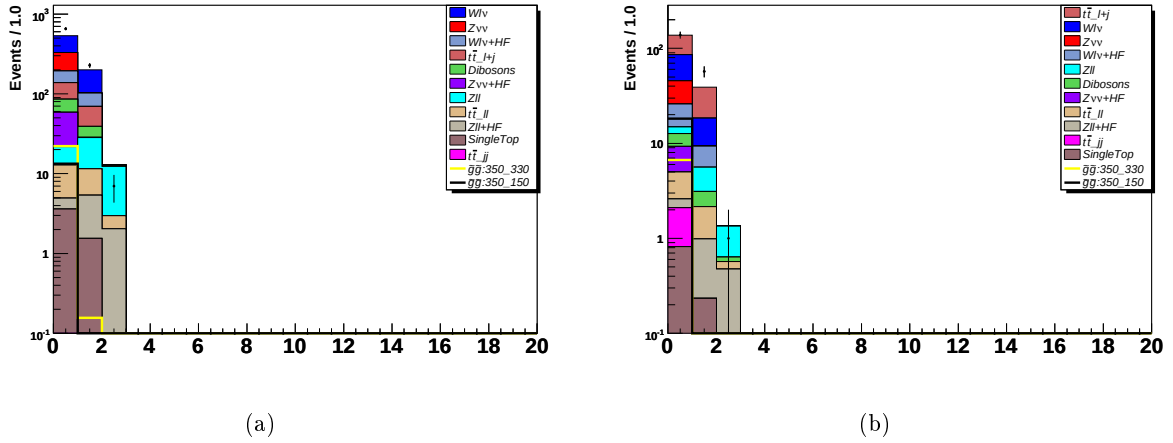


FIG. 4.25 – Distributions du nombre de muons isolés au niveau de sélection C_4 pour l'analyse $JT1$ (a) et pour l'analyse $JT2$ (b).

muons choisis pour ce veto sont des muons de qualité *Medium*, de type $|nseg| = 3$ et $p_T > 5$ GeV. Ces caractéristiques sont détaillées dans la Section 3.4.2.

Le bruit de fond instrumental ne possède quasiment pas de leptons dans l'état final comme le montrent les Figures 4.24 et 4.25. En inversant les coupures précédentes, c'est-à-dire en demandant explicitement un ou deux leptons, on peut d'une part étudier la composition en bruit de fond des événements rejetés et d'autre part s'assurer que les principaux bruits de fond physique sont correctement simulés et normalisés. On définit alors plusieurs zones de contrôle selon qu'on ait exactement un électron isolé, exactement un muon isolé ou au moins deux leptons isolés. Pour augmenter la statistique de cette étude, on diminue la coupure sur l'énergie transverse manquante à 75 GeV pour les études avec exactement un lepton isolé et à 50 GeV pour au moins deux leptons isolés. Les autres coupures de sélection du niveau C_4 restent inchangées. Des distributions dans les zones de contrôle, ainsi définies, sont données sur les Figures 4.26, 4.27 et 4.28.

Les graphes de la Figure 4.26 permettent de contrôler le bruit de fond $W \rightarrow \ell\nu + jets$. Les accords obtenus permettent d'être d'une part d'être confiant dans la simulation de ce bruit de fond et d'autre part d'être confiant dans la procédure appliquée pour reproduire le profil de luminosité. Les distributions du nombre de vertex primaires de l'événement attendu et observée sont, en effet, en très bon accord. Les graphes de la Figure 4.27 ajoute la possibilité de contrôler le bruit de fond $t + \bar{t} \rightarrow \ell\nu + 2b + jets$, sachant le bruit de $W \rightarrow \ell\nu + jets$ correctement simulé. Compte-tenu des incertitudes statistiques sur les données réelles, les graphes obtenus fournissent un accord raisonnable. Finalement, les graphes de la Figure 4.28 indique que le bruit de fond $Z \rightarrow \ell^+\ell^- + jets$ est également bien simulé.

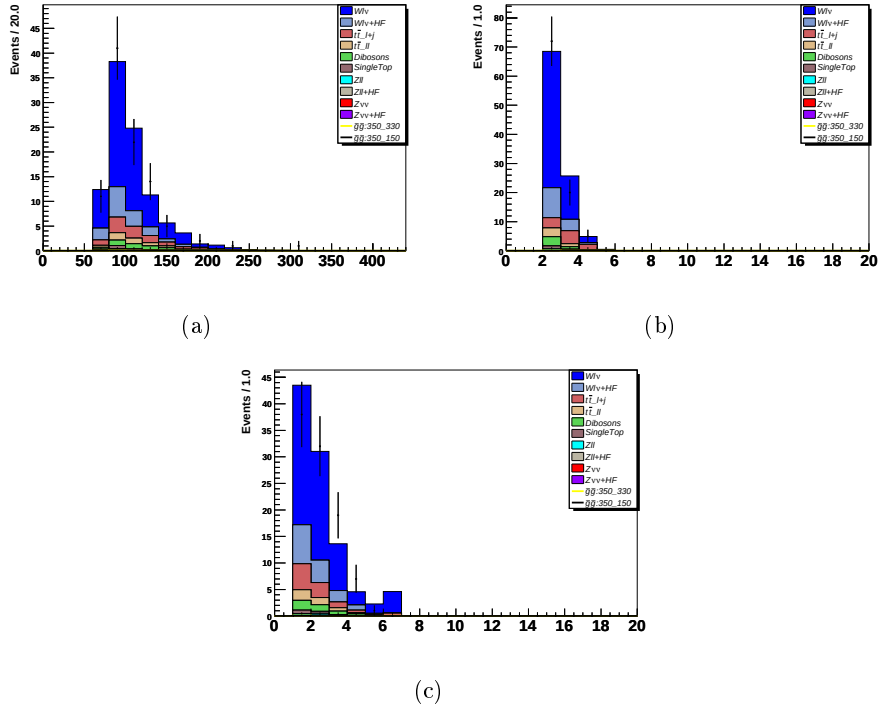


FIG. 4.26 – Distributions de E_T (a), du nombre de jets centraux (b) et du nombre de vertex primaires (c) au niveau de sélection C_4' en demandant exactement un électron isolé pour l'analyse $JT1$.

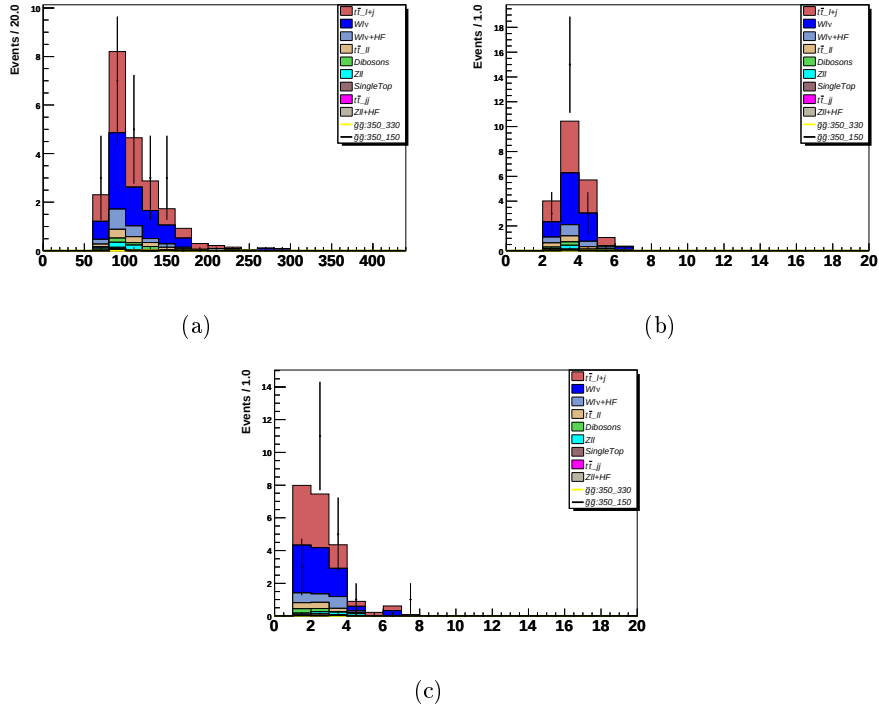


FIG. 4.27 – Distributions de E_T (a), du nombre de jets centraux (b) et du nombre de vertex primaires (c) au niveau de sélection C_4' en demandant exactement un électron isolé pour l'analyse $JT2$.

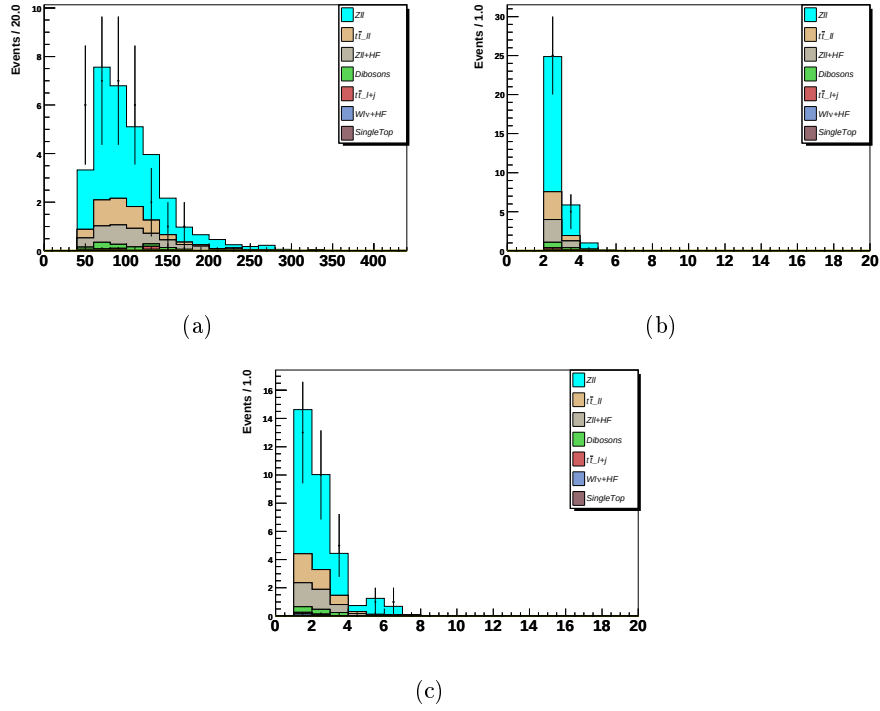


FIG. 4.28 – Distributions de E_T (a), du nombre de jets centraux (b) et du nombre de vertex primaires (c) au niveau de sélection C_4' en demandant au moins deux leptons isolés pour l'analyse $JT1$.

4.4.5 L'étiquetage des jets des quarks b

L'étiquetage des quarks b permet une réduction du bruit de fond instrumental, il permet également une réduction de deux des principaux bruits de fond : $W \rightarrow \ell\nu + jets$ et $Z \rightarrow \nu\bar{\nu} + jets$. De plus, les états finaux des signaux supersymétriques recherchés possèdent de deux à quatre quarks de b. L'étiquetage des quarks b fournit donc une façon très efficace d'isoler le signal en vue d'une optimisation finale des coupures de sélection. L'algorithme d'étiquetage des quarks b choisi pour cette analyse est basé sur un réseau de neurones, dont la construction a été donnée dans la Section 3.4.6.

Paramétrisation de la taggabilité

Comme il a été détaillé dans la Section 3.4.6, la simulation des traces ne permet pas d'estimer convenablement la *taggabilité* pour les données *Monte Carlo*. En effet, la *taggabilité* caractérise la possibilité d'appliquer les algorithmes d'étiquetage, qui sont tous basés sur les informations du trajectographe central. On préfère alors paramétrer la *taggabilité* à partir de lots de données réelles pour déterminer un poids à appliquer sur les données simulées.

La *taggabilité* dépend fortement de la topologie du signal recherché et donc des coupures de sélection qui sont appliquées pour cette recherche. Pour cette raison, on utilise deux lots de données réelles pour paramétrer cette *taggabilité*. Le premier, correspondant à l'analyse *JT1*, utilise les coupures de sélection suivantes :

$$\left\{ \begin{array}{l} \text{Conditions de déclenchement : } MHT30_3CJT5 \text{ ou } JT1_ACO_MHT_HT \\ \text{Critères de qualité des données} \\ \text{Au moins un vertex primaire reconstruit (PV), avec : } |z_{PV}| < 60 \text{ cm} \\ \text{Au moins 2 jets } (j_1 \text{ et } j_2) \text{ avec } E_T > 15 \text{ GeV et } |\eta_{\text{det}}(j_1, j_2)| < 2.5 \\ \cancel{E}_T > 40 \text{ GeV} \\ \Delta\phi(j_1, j_2) < 165^\circ \\ \Delta\phi_{\min}(\cancel{E}_T, jets) > 40^\circ \\ CPF_0(j_1, j_2) > 0.85 \end{array} \right.$$

Le second lot de données correspond à l'analyse *JT2* et est construit à partir des coupures de sélections suivantes :

$$\left\{ \begin{array}{l} \text{Conditions de déclenchement : } MHT30_3CJT5 \text{ ou } JT2_MHT25_HT \\ \text{Critères de qualité des données} \\ \text{Au moins un vertex primaire reconstruit (PV), avec : } |z_{PV}| < 60 \text{ cm} \\ \text{Au moins 3 jets } (j_1, j_2, \text{ et } j_3) \text{ avec } E_T > 15 \text{ GeV et } |\eta_{\text{det}}(j_1, j_2, j_3)| < 2.5 \\ \cancel{E}_T > 50 \text{ GeV} \\ CPF_0(j_1, j_2) > 0.85 \end{array} \right.$$

Pour chacun des deux lots de données, on étudie alors la probabilité pour un jet d'être étiquetable en fonction des quatre variables suivantes :

1. la position du vertex primaire suivant l'axe z, z_{vtx} ;
2. la position en η du jet considéré ;

3. la position en ϕ du jet considéré ;
4. l'énergie transverse, E_T , du jet ;

La première variable est une variable caractéristique de l'événement, les trois suivantes sont propres au jet considéré. Les deux premières variables caractérisent la capacité du trajectographe interne à associer des traces au jet considéré. En fonction de la position du jet et du vertex primaire de l'interaction considéré, il peut s'avérer difficile voire impossible de reconstruire des traces associées au jet. Pour rendre compte de la géométrie du trajectographe et de la trajectoire des traces associées au jet, on utilise plutôt la variable : $\eta \times \text{sign}(z_{vtx})$ dans la suite de cette étude.

Analyse dite JT1 Les dépendances de la *taggabilité*, en fonction des variables explicitées précédemment, sont données sur la Figure 4.29 pour le lot de données correspondant à l'analyse JT1. On utilise alors trois types de fonctions d'ajustement pour paramétrer ces trois dépendances.

Pour la dépendance en $\eta_{\text{sign}} = \eta \times \text{sign}(z_{vtx})$, on utilise un polynôme d'ordre 8, qui diffère selon l'intervalle en $|z_{vtx}|$. Trois intervalles ont ainsi été définis, ainsi que les trois barycentres et les trois fonctions d'ajustement correspondant :

1. $0 \leq |z_{vtx}| < 20$ cm, $B20$, $f20(\eta_{\text{sign}})$;
2. $20 \leq |z_{vtx}| < 36$ cm, $B36$, $f36(\eta_{\text{sign}})$;
3. $36 \leq |z_{vtx}| < 60$ cm, $B60$, $f60(\eta_{\text{sign}})$.

Où $B20$, $B36$ et $B60$ sont les trois barycentres correspondant aux trois intervalles et $f20$, $f36$ et $f60$ sont les trois fonctions d'ajustement, polynômes d'ordre 8, associées.

La dépendance en E_T est ajustée à l'aide d'une fonction du type :

$$g(X) = A + \frac{\sum_{i=0}^5 B_i X^i}{\sum_{i=0}^5 C_i X^i}$$

Avec A , B_i et C_i des constantes à déterminer.

Finalement, la fonction d'ajustement choisie pour la dépendance en ϕ est donnée par la formule :

$$h(X) = A + \sum_{i=1}^3 (B_i + C_i X) \times \sin(D_i + i \times X)$$

Où A , B_i , C_i et D_i sont des constantes à déterminer.

Les résultats de ces ajustements pour l'analyse JT1 sont donnés sur la Figure 4.29.

Les fonctions d'ajustement, une fois déterminées, fournissent trois paramétrisations. Les fonctions g et h sont deux paramétrisations à une dimension, qui donne deux poids notés w_{ET} et w_ϕ . Ces deux poids correspondent simplement à : $w_{ET} = g(E_T)$ et $w_\phi = h(\phi)$. Les trois polynômes d'ordre 8 et les trois barycentres définis précédemment permettent de construire une paramétrisation à deux dimensions en fonction de η_{sign} et de $|z_{vtx}|$. Cette paramétrisation fournit le poids w_η et est construite de la façon suivante :

- Si $|z_{vtx}| < B20$, alors w_η est le barycentre des fonctions $f20(\eta_{\text{sign}})$ et de $f20(-\eta_{\text{sign}})$, c'est-à-dire :

$$w_\eta = \frac{(|z_{vtx}| + B20) \times f20(\eta_{\text{sign}})}{2B20} + \frac{(B20 - |z_{vtx}|) \times f20(-\eta_{\text{sign}})}{2B20}$$

- Si $B20 < |z_{vtx}| < B36$, alors w_η est le barycentre des fonctions $f20(\eta_{sign})$ et de $f36(\eta_{sign})$:

$$w_\eta = \frac{(|z_{vtx}| - B20) \times f36(\eta_{sign})}{B36 - B20} + \frac{(B26 - |z_{vtx}|) \times f20(\eta_{sign})}{B36 - B20}$$

- Si $B36 < |z_{vtx}|$, alors on distingue trois cas suivant la valeur de $|\eta_{sign}|$:
 - Si $|\eta_{sign}| < 1.1$, alors w_η est le barycentre de $f36(\eta_{sign})$ et de $f60(\eta_{sign})$, construit comme les barycentres précédents.
 - Si $1.1 < |\eta_{sign}| < 2.1$, on calcule le barycentre de la même façon mais en considérant que la fonction d'ajustement $f60(\eta_{sign})$ est nulle. Cette approximation est en accord avec la forme de la distribution de la Figure 4.29 et avec la géométrie du trajectographe interne.
 - Si $2.1 < |\eta_{sign}|$, alors $w_\eta = 0$. En effet, aucune trace ne peut être associée à un jet dans un des deux bouchons si le vertex primaire est situé du côté de ce même bouchon.

Ces trois paramétrisations sont résumées sur les graphes de la Figure 4.30 pour l'analyse *JT1*.

L'efficacité globale obtenue pour le lots de données correspondant à l'analyse *JT1* est de 75.3 %. On note, w_{glob} , cette efficacité. Finalement, le poids final à appliquer à un événement *Monte Carlo* est donné par la formule suivante :

$$w_{tot} = \frac{w_{ET} \times w_\phi \times w_\eta}{w_{glob}^2}$$

Pour valider cette procédure, on applique ce poids aux événements de données réelles du lot de données utilisés pour déterminer cette paramétrisation. Et on compare les distributions obtenues aux distributions obtenues en appliquant directement les critères de *taggabilité*. Les résultats de cette comparaison sont donnés sur la Figure 4.31 pour les distributions caractéristiques de l'événement. On constate que l'accord est très bon et par conséquent que la paramétrisation reproduit bien la *taggabilité* pour ce lot de données.

Analyse dite JT2 La procédure, décrite pour l'analyse *JT1*, est appliquée pour le lot de données correspondant à l'analyse *JT2*. De la même façon, on détermine des fonctions d'ajustement comme le montre la Figure 4.32. A partir de ces fonctions, on détermine la paramétrisation donnée sur les graphes de la Figure 4.33. Cette paramétrisation et l'efficacité globale de 78.5 % permettent de définir un poids à appliquer aux événements *Monte Carlo*. La vérification de ce poids global est donnée par les graphes de la Figure 4.34.

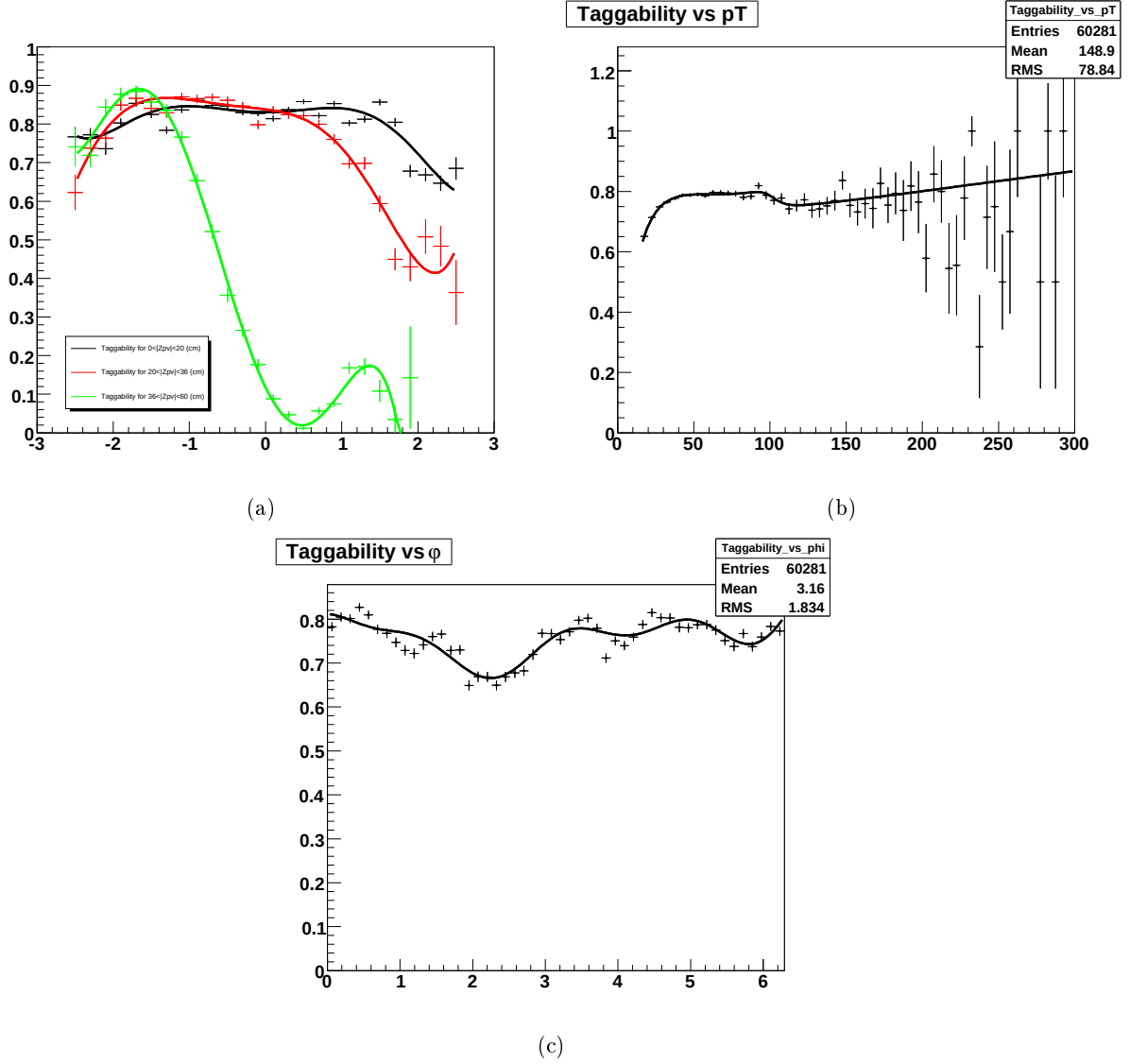


FIG. 4.29 – Probabilité pour un jet d'être *taggable* en fonction du $\eta \times \text{sign}(z_{vtx})$ pour différentes valeurs de $|z_{vtx}|$ (a), en fonction du E_T (b) et en fonction du ϕ (c) du jet considéré. Cette paramétrisation est valable pour l'analyse dite *JT1*.

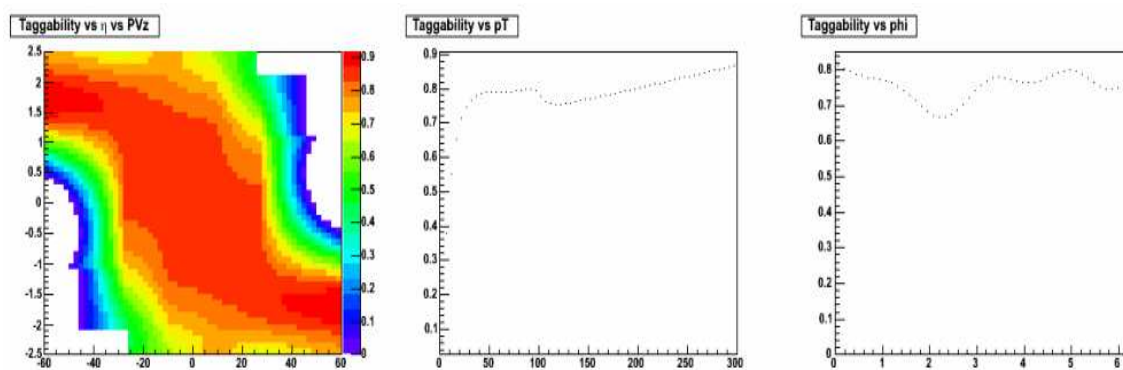


FIG. 4.30 – Paramétrisation de la *taggabilité* en fonction de η , z_{vtx} , p_T et ϕ du jet pour l'analyse JT1.

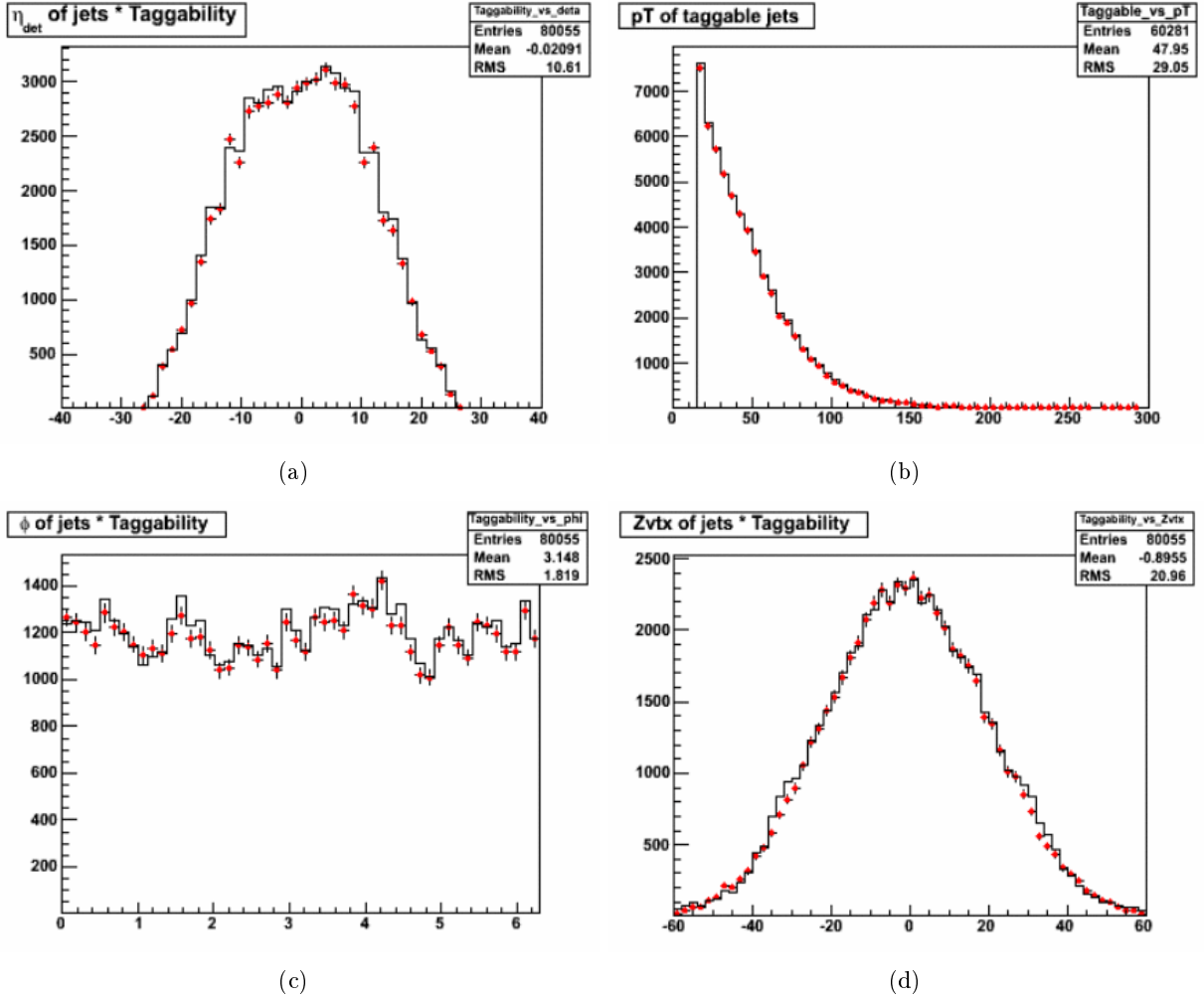


FIG. 4.31 – Vérification de la paramétrisation en fonction de η_{det} (a), de E_T (b), de ϕ (c) du jet considéré et de z_{vtx} (d). Les histogrammes représentent les distributions pour les jets *taggables* en utilisant la définition de la *taggabilité*, les points rouges correspondent aux distributions des jets *taggables* en utilisant la paramétrisation. Ces vérifications sont valables pour l'analyse *JT1*.

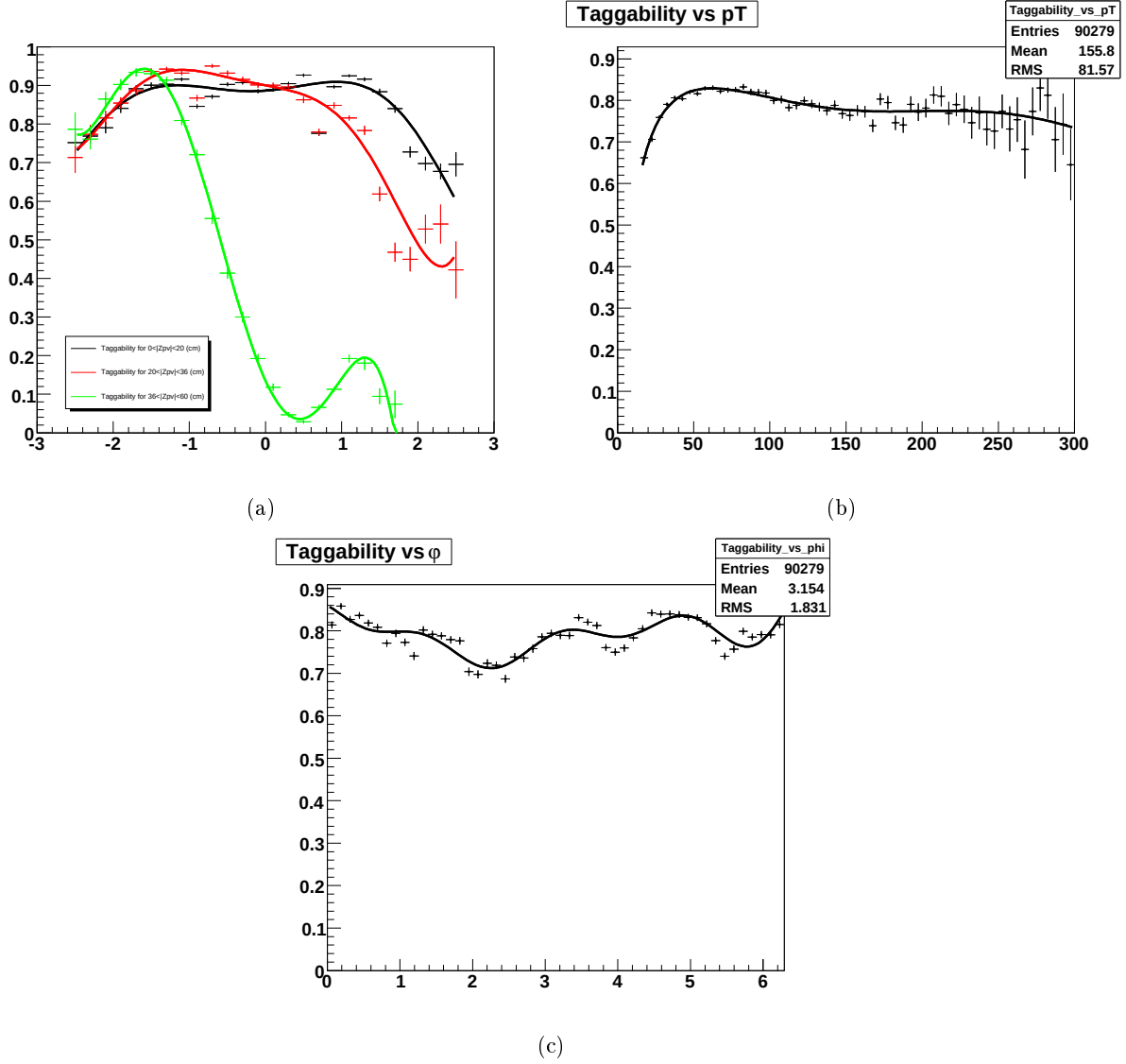


FIG. 4.32 – Probabilité pour un jet d'être *taggable* en fonction du $\eta \times \text{sign}(z_{vtx})$ pour différentes valeurs de $|z_{vtx}|$ (a), en fonction du E_T (b) et en fonction du ϕ (c) du jet considéré. Cette paramétrisation est valable pour l'analyse dite *JT2*.

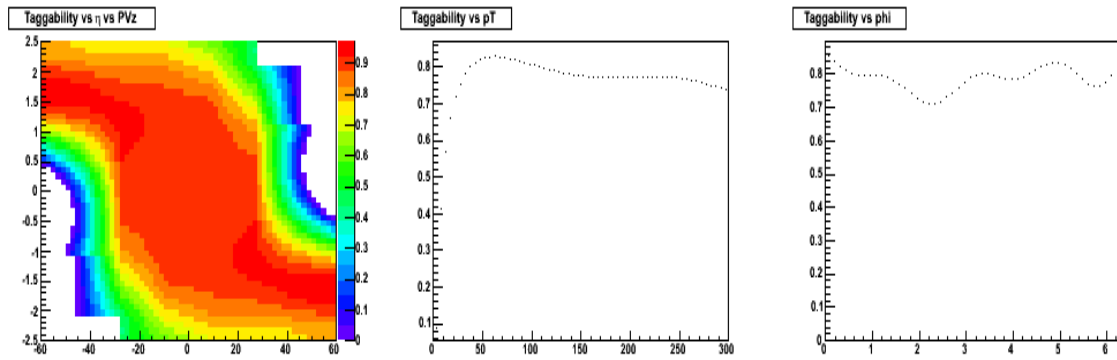


FIG. 4.33 – Paramétrisation de la *taggabilité* en fonction de η , z_{vtx} , E_T et ϕ du jet pour l'analyse JT1.

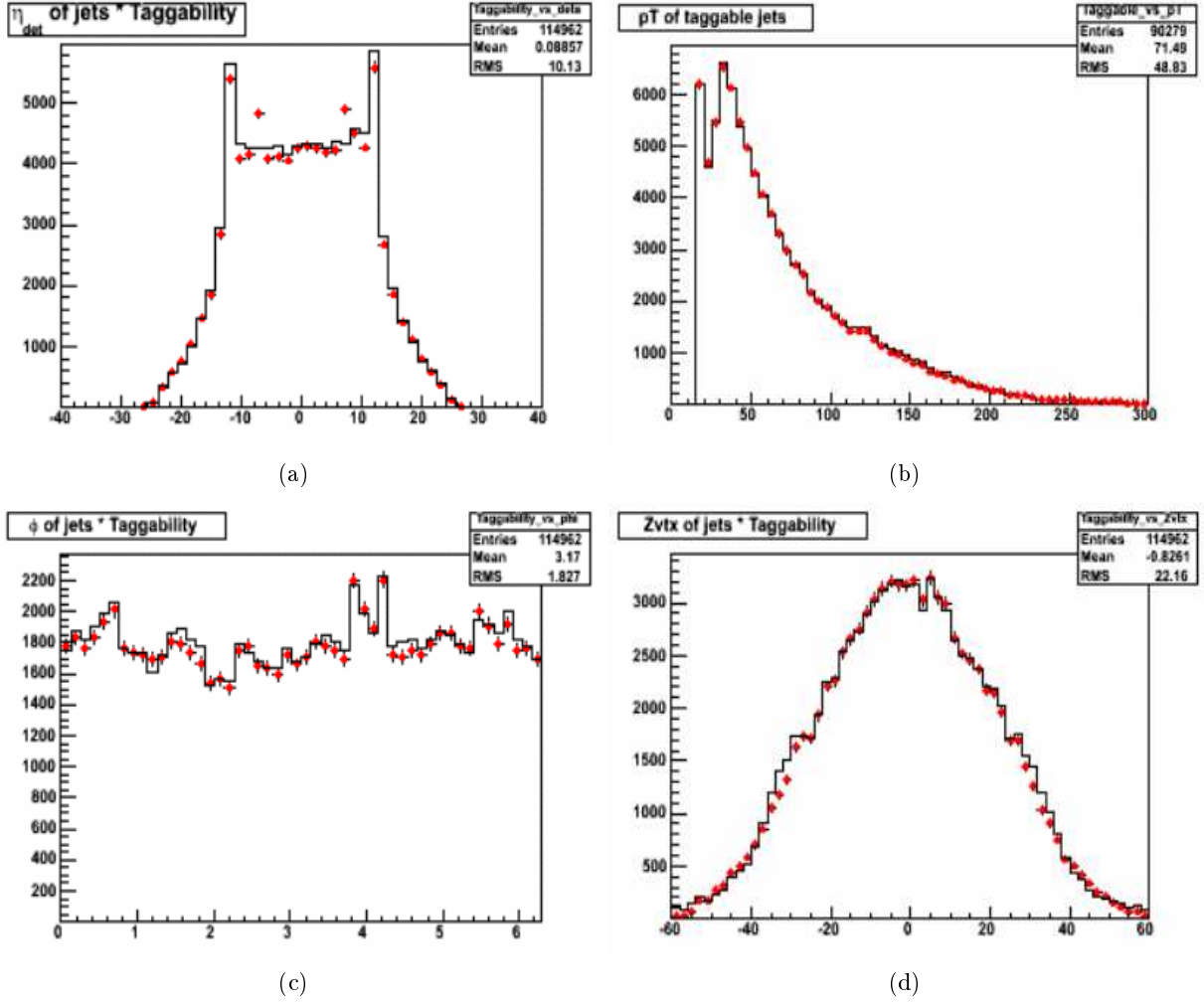


FIG. 4.34 – Vérification de la paramétrisation en fonction de η_{det} (a), de E_T (b), de ϕ (c) du jet considéré et de z_{vtx} (d). Les histogrammes représentent les distributions pour les jets *taggables* en utilisant la définition de la *taggabilité*, les points rouges correspondent aux distributions des jets *taggables* en utilisant la paramétrisation. Ces vérifications sont valables pour l'analyse JT2.

L'étiquetage des jets taggables

Les points de fonctionnement Trois points de fonctionnement ont été choisis pour cette recherche d'un signal supersymétrique : *Loose*, *Medium* et *Tight*. Les caractéristiques de ces points de fonctionnement sont données dans la Section 3.4.6. Les étiquetages croisés, du type un étiquetage *Loose* et un étiquetage *Tight*, n'ont pas été utilisés. On permet alors de 0 à 3 étiquetages de quarks b pour isoler le signal recherché.

Description de l'algorithme L'algorithme d'étiquetage utilise la valeur de sortie du réseau de neurones pour les données du détecteur. Si celle-ci est supérieure au seuil fixé par le point de fonctionnement le jet est considéré comme étiqueté. Cet algorithme s'applique uniquement sur les jets *taggables*. Pour les données simulées, on utilise la probabilité appelée *TRF* décrite dans la Section 3.4.6. Elle donne la probabilité pour un jet étiquetable d'être étiqueté suivant sa saveur. On distingue trois types de *TRF* (voir la Figure 4.35) :

1. les *TRF* pour les quarks b, c'est-à-dire la probabilité pour un quark b d'être étiqueté ;
2. les *TRF* pour les quarks c, c'est-à-dire la probabilité pour un quark c d'être étiqueté comme un quark b ;
3. les *TRF* pour les quarks légers, c'est-à-dire la probabilité pour un quark u,d ou s d'être étiqueté comme un quark b. On parle aussi d'inefficacité d'étiquetage dans ce cas.

Pour un jet issu d'un événement simulé, on affecte un poids fonction de la saveur du jet et du point de fonctionnement. Ce poids est à multiplier au poids déterminé par les études de *taggabilité* pour donner le poids total, correspondant à la probabilité du jet d'être étiqueté comme un quark de b.

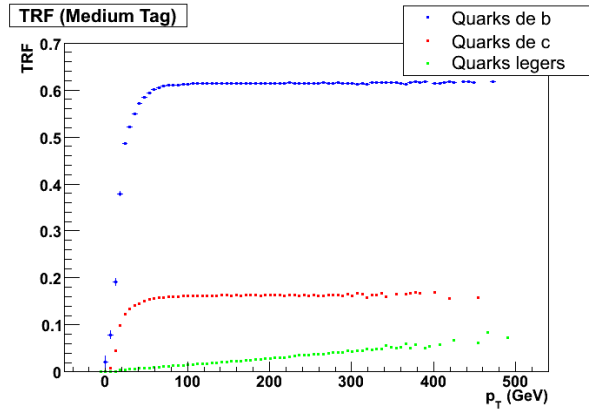


FIG. 4.35 – *TRF* en fonction de la saveur et de l'énergie transverse du jet pour un lot de données simulant un processus $t + \bar{t} \rightarrow \ell\nu + 2b + jets$.

A partir des poids obtenus pour chaque jet, on détermine un poids global pour l'événement selon qu'on demande 1, 2 ou 3 jets étiquetés dans l'état final. Cet algorithme est appliqué uniquement pour les quatre jets de plus hautes énergies transverses, confirmés par le trajectographe interne, avec $E_T > 15$ GeV et $|\eta| < 2.5$. Si un jet ne vérifie pas ces conditions, on lui affecte un poids de 0.

Comparaisons données réelles et données simulées Pour vérifier l'accord entre les données réelles et les données simulées après l'étiquetage d'un quark b, on utilise une inversion de coupures sur les leptons isolés. On se place ainsi au niveau $C4$ de l'analyse et on demande un électron isolé, avec la même définition que celle donnée précédemment, avec une coupure à 75 GeV pour l'énergie transverse manquante. De cette façon, on teste aussi bien la méthode d'étiquetage utilisé que la simulation des fonds dominants après ce type de sélection. Le bruit de fond physique est essentiellement constitué des processus $t\bar{t}$ avec un lepton dans l'état final et des processus $W \rightarrow \ell\nu + jets$ avec des saveurs lourdes dans l'état final.

	<i>JT1</i>			<i>JT2</i>		
	<i>L</i>	<i>M</i>	<i>T</i>	<i>L</i>	<i>M</i>	<i>T</i>
N_{obs}	22	17	16	10	9	9
N_{exp}	19.1 ± 0.5	15.7 ± 0.5	13.8 ± 0.4	9.5 ± 0.3	8.5 ± 0.2	7.8 ± 0.2

TAB. 4.17 – Nombre d'événements attendus et observés après un étiquetage de quark b (*L* pour *Loose*, *M* pour *Medium* et *T* pour *Tight*) pour le niveau de sélection $C4$ avec un électron isolé et $\cancel{E}_T > 75$ GeV. Seules les incertitudes statistiques sont indiquées.

Pour un étiquetage d'un quark b, on obtient les résultats du Tableau 4.17. Ces résultats sont raisonnablement en accord, compte-tenu des incertitudes importantes sur l'échelle d'énergie des jets, l'étiquetage des quarks b ou encore l'estimation des sections efficaces. Ces incertitudes seront détaillées dans la Section 4.5.2. Les mêmes accords sont également obtenus pour 2 ou 3 jets étiquetés comme étant un jet de quark b. Pour compléter cette comparaison, on donne les graphes de la Figure 4.36, qui montrent également un bon accord pour les formes de quelques distributions. Des exemples de distributions avec un veto sur les leptons isolés sont donnés dans la section suivante afin de compléter cette comparaison entre les données réelles et les données simulées.

4.4.6 Coupures finales avant l'optimisation

Le fond instrumental est négligé pour l'optimisation finale. Il faut donc être en mesure de déterminer la coupure minimale sur $\cancel{E}_T$ à appliquer pour que cette hypothèse soit cohérente. On estime alors le nombre d'événements de fond instrumental par la méthode décrite dans la Section 4.4.2 pour les deux canaux d'analyse, pour un nombre d'étiquetage de quarks b fixé et pour un point de fonctionnement donné. On choisit la coupure sur $\cancel{E}_T$ de façon à avoir moins de 10 % d'événements de fond *QCD* par rapport au nombre total d'événement observé. Ce critère ad-hoc permet d'avoir un bon accord entre les événements observés et les événements attendus et permet ainsi d'avoir une limite sur les sections efficaces attendue et observée proches. Cette hypothèse donne finalement une limite conservative sur la section efficace observée.

La Figure 4.37 donne les distributions de $\cancel{E}_T$ pour 0, 1, 2 ou 3 étiquetages *Medium* de quarks b, ainsi que les estimations du nombre d'événements de fond instrumental pour l'analyse dite *JT2*. De ces estimations et de ces distributions, on détermine les coupures minimales sur $\cancel{E}_T$ à appliquer pour la suite. Ces dernières sont résumées dans le Tableau 4.18.

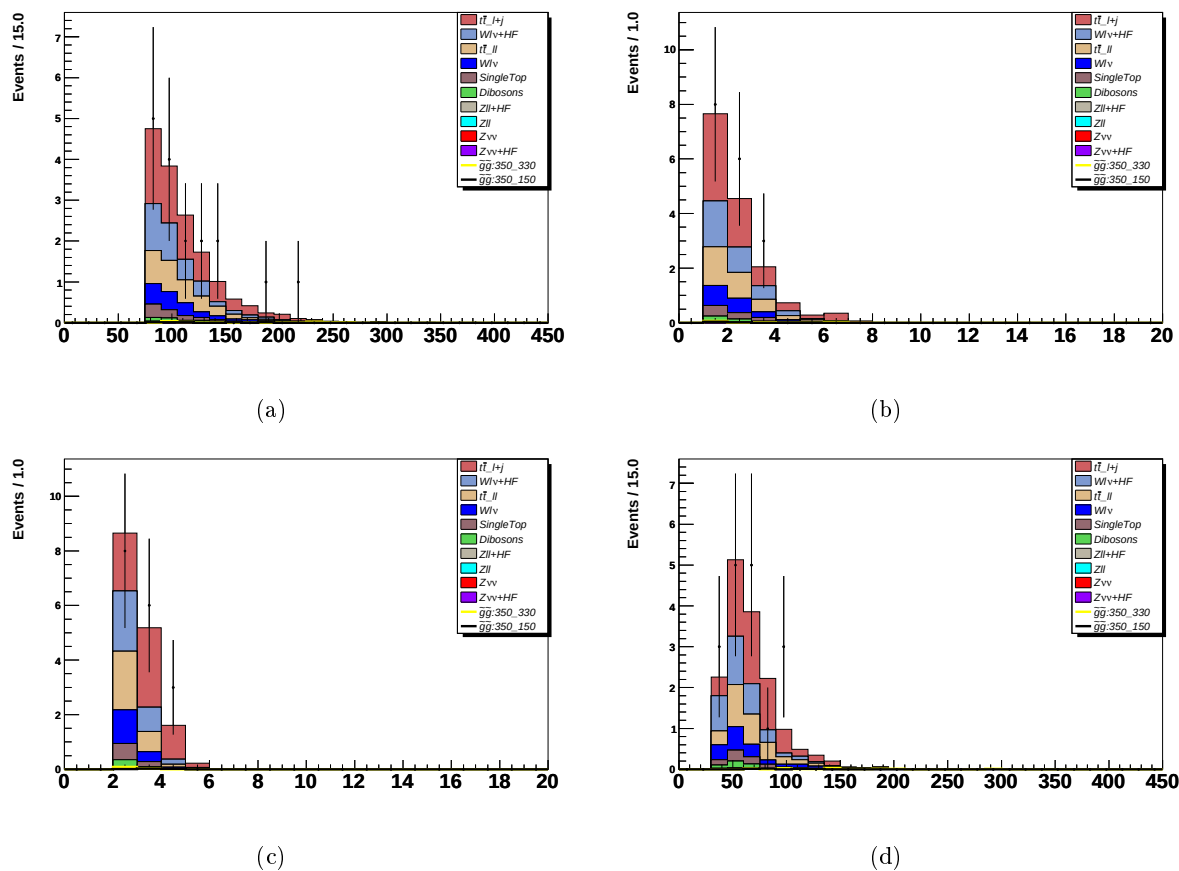


FIG. 4.36 – Distribution de E_T (a), du nombre de vertex primaires (b), du nombre de jets centraux (c) et de l'énergie transverse du second jet le plus dur (d) pour le niveau de sélection C_4 avec un électron isolé et $E_T > 75$ GeV.

Nombre de quarks étiquetés	0	L	M	T	$2L$	$2M$	$2T$	$3L$	$3M$	$3T$
$JT1, \not{E}_T > X \text{ GeV}$	125	100	100	100	100	100	100	75	75	75
$JT2, \not{E}_T > X \text{ GeV}$	125	150	150	150	125	125	100	100	100	100

TAB. 4.18 – Coupure minimale sur $\not{E}_T$ pour les deux canaux d’analyses et pour différentes sélections de quarks b. L désigne un étiquetage *Loose*, $2M$ deux étiquetages *Medium*, $2T$ deux étiquetages *Tight*,...

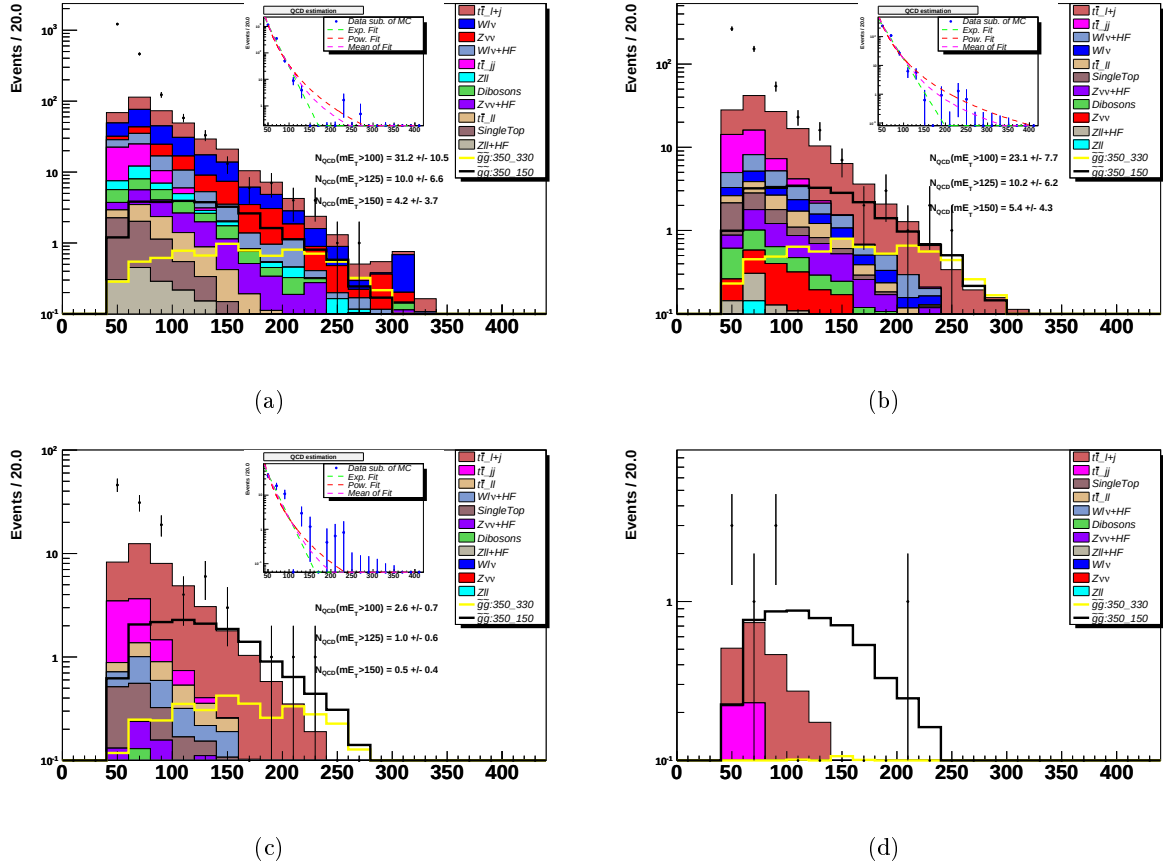


FIG. 4.37 – Distribution de E_T , pour l'analyse *JT2*, pour 0 (a), 1 (b), 2 (c) ou 3 (d) étiquetages *Medium* de quarks *b*. Les estimations du fond *QCD* sont données pour différentes coupures sur E_T .

4.5 Optimisation finale

L'objectif de l'optimisation finale est de trouver le meilleur jeu de coupures, qui minimise la valeur de la section efficace attendue du signal recherché. Deux variables sont utilisées pour cette minimisation : l'énergie transverse manquante, $\cancel{E}_T$, et la somme scalaire des jets, H_T , qui varient par pas de 25 GeV. L'optimisation est effectuée pour chacune des configurations d'étiquetage des quarks b étudiées dans ce chapitre.

4.5.1 Méthode statistique

Pour trouver le meilleur jeu de coupures et ainsi minimiser la section efficace, on utilise une méthode fréquentiste basée sur le calcul d'un niveau de confiance, noté CL_s [163]. Cette variable est calculée, d'une part, à partir du nombre d'événements de fond physique attendu et du nombre d'événements de signal prédit. On parle alors de niveau de confiance attendu CL_s^{exp} . D'autre part, on calcule le niveau de confiance observé CL_s^{obs} à partir du nombre d'événements de données réelles et du nombre d'événements de signal. L'intérêt de ce calcul est qu'il tient compte des incertitudes statistiques et systématiques sur les nombres d'événements attendus et observés. Il peut tenir compte également des corrélations entre le signal et le bruit de fond. En effet, deux approches sont possibles. La première prend en compte la forme des distributions pour calculer le niveau de confiance, la seconde est une expérience de comptage et ne considère que les intégrales des distributions (c'est-à-dire les nombres d'événements). Le choix, fait pour la recherche présentée dans ce chapitre, est d'utiliser une expérience de comptage en faisant varier la coupure sur $\cancel{E}_T$. Cette approche est peu différente, en première approximation, d'un calcul qui prendrait en compte la forme de la distribution de $\cancel{E}_T$.

La valeur du CL_s^{exp} caractérise la sensibilité de l'analyse. Plus ce chiffre est faible, plus il devient impossible de négliger le signal par rapport au bruit de fond physique. Le but de l'optimisation est la minimisation de cette valeur. La valeur minimale de CL_s^{exp} fixe les meilleures coupures sur $\cancel{E}_T$ et H_T à appliquer.

La valeur du CL_s^{obs} permet de savoir si le point de signal étudié est exclu ou non (dans le cas, bien sûr, où aucun excès n'a été découvert dans les données). En effet, si ce chiffre est inférieur à 5 %, on estime que le signal est exclu à 95 %. La valeur du CL_s^{obs} sera calculée pour les coupures déterminées par la minimisation de CL_s^{exp} . On détermine ainsi un jeu de coupures par point de signal et par nombre de jets étiquetés. Il est important de noter que cette façon de procéder nécessite d'avoir un bon accord entre les données réelles et les données simulées. Dans le cas de cette analyse, si le fond instrumental n'est pas négligeable, la valeur du CL_s^{obs} sera systématiquement bien supérieure à la valeur du CL_s^{exp} , et par conséquent la limite observée sur la section efficace moins bonne que la limite attendue. Les coupures finales sur $\cancel{E}_T$ détaillées précédemment limitent ce problème.

Pour déterminer les limites minimales sur les sections efficaces observées et attendues, on utilise une procédure itérative. On multiplie la section efficace de production du signal supersymétrique par un facteur k . On cherche alors le facteur k par dichotomie jusqu'à d'atteindre un CL_s égal à 0.05 à 0.001 près. Si σ_{signal} est la section efficace théorique du signal recherché, alors $\sigma_{signal} \times k$ est la valeur minimale de la section efficace pour laquelle on peut exclure ce signal à 95 % de niveau de confiance. Dans la suite, toutes les limites sur les sections efficaces ou sur les masses des particules supersymétriques seront données à 95 % de niveau de confiance.

4.5.2 Incertitudes systématiques

Les incertitudes systématiques sur les coupures de sélection sont utilisées pour le calcul du niveau de confiance. Plus ces incertitudes sont faibles, plus la section efficace exclue est faible. On trouve des incertitudes expérimentales qui influent sur la précision des coupures de sélection, on trouve également des incertitudes théoriques sur les sections efficaces qui influent sur la normalisation.

Les différentes sources d'incertitudes utilisées pour le calcul du niveau de confiance et leur amplitude sont résumées dans le Tableau 4.19. Il faut ajouter à ces incertitudes les incertitudes statistiques.

Les incertitudes sur l'échelle d'énergie des jets et sur les probabilités d'étiquetage, TRF , sont estimées en prenant la variation maximale que l'on trouve si on ré-effectue l'analyse avec $JES \pm 1\sigma$ et $TRF \pm 1\sigma$. En effet, ces incertitudes étant très sensibles au jeu de coupures choisis et au nombre de jets étiquetés comme étant un jet de quark b , il faut les estimer à chaque fois. Néanmoins, pour le signal, on prend la valeur conservative de 7 % pour l'incertitude sur l'échelle d'énergie des jets, qui correspond à l'incertitude maximale de l'analyse de la référence [84]. On procède de la même façon pour les incertitudes sur l'identification et la résolution des jets. La valeur utilisée pour CPF_0 correspond à l'incertitude sur la détermination des fonctions d'ajustement. Une valeur conservative de 1 % est, de plus, ajoutée pour tenir compte des incertitudes sur la détermination des fonctions d'ajustement de la paramétrisation de la *taggabilité*. Finalement, les incertitudes sur le veto des leptons isolés ont été estimées comme négligeables. En effet, le nombre d'événements sélectionnés, après le niveau de coupure $C5$, n'est pas modifié si on fait varier les efficacités d'identification et d'isolation d'un électron ou d'un muon (pour les définitions données précédemment) de $\pm 1\sigma$ pour les événements simulés.

Toutes ces incertitudes sont considérées comme indépendantes et l'incertitude totale pour le bruit de fond physique est donnée par la somme quadratique de chacune des valeurs du tableau. On rappelle que le fond instrumental a été négligé ici, ce qui donnera une valeur conservative pour la section efficace limite. Ceci est valable dans le cas où aucune déviation par rapport aux prédictions du modèle standard n'est observée.

Les incertitudes systématiques sur la section efficace du signal ne sont pas utilisées pour le calcul du niveau de confiance mais pour la détermination de la masse limite des particules supersymétriques. Ces dernières tiennent compte des incertitudes statistiques (nombre d'événement générés pour déterminer la section efficace), des incertitudes sur les fonctions de densité de partons, et des incertitudes dues aux choix des échelles de renormalisation et de factorisation. Le détail de chacune de ces contributions est donné dans le Tableau 4.20 pour une masse de squarks, de première et deuxième génération, de 1 TeV.

4.5.3 Détermination des masses limites

La détermination des limites sur les masses des particules supersymétriques se fait en plusieurs étapes. Tout d'abord, on cherche le nombre optimal de jets étiquetés par canal analyse. Ensuite, on étudie l'évolution des sections efficaces attendues en fonction des masses de particules ($m_{\tilde{g}}$ ou $m_{\tilde{b}_1}$). Cette étape est effectuée en fixant une condition sur la masse non-étudiée. On détermine ainsi le meilleur canal d'analyse avec le nombre de jets étiquetés correspondant pour une condition donnée. Finalement, on extrapole la section efficace observée, pour ce canal d'analyse, afin de trouver la masse limite de la particule étudiée sous la condition que l'on s'est fixée.

Sources	Effet sur le bruit de fond	Effet sur le signal
Echelle d'énergie des jets (<i>JES</i>)	$\pm 1\sigma$ (<i>JT1</i> : 4-13 %, <i>JT2</i> : 6-18%)	7% [84]
Identification des jets [84]	2%	2%
Résolution des jets [84]	5%	5%
Luminosité [69]	6.5%	6.5%
Conditions de déclenchement [162]	2%	2%
<i>CPF</i> ₀	5%	5%
<i>taggabilité</i>	1%	1%
Etiquetage des quarks b (<i>TRF</i>)	$\pm 1\sigma$ (4-6 %, 6-13 %, 15-21 % pour 1, 2 ou 3 étiquetages)	$\pm 1\sigma$ (1-5 %, 5-14 %, 11-23 %)
Sections efficaces du fond	15%	—
Acceptance due aux choix des <i>PDF</i> [84]	5%	5%

TAB. 4.19 – Les incertitudes systématiques sur le nombre d'événements attendus pour le bruit de fond physique et pour le signal.

Détermination du nombre de jets étiquetés par canal d'analyse

Dans une première étape, on cherche le meilleur jeu de coupures pour un point de signal et pour un nombre de jets étiquetés donnés. Cette étape s'effectue pour chacun des deux canaux d'analyse. Pour un point de signal donné, on a donc 8 optimisations possibles (2 analyses avec 0, 1, 2 ou 3 étiquetages). A ces 8 optimisations correspondent 8 CL_s^{exp} , 8 sections efficaces attendues et 8 sections efficaces observées.

On cherche alors, pour chaque point de signal, la combinaison qui donne la valeur minimale du CL_s^{exp} . Les résultats, obtenus pour une partie des points générés et dans le cas d'une masse moyenne des squarks légers de 1 TeV, sont résumés dans le Tableau 4.21. Ce tableau montre que le nombre optimal de jets étiquetés pour l'analyse *JT1* est 1, voire 2. Ce nombre est plutôt 2 ou 3 pour l'analyse *JT2*. Dans la suite, on ne considère que ces possibilités pour l'estimation des masses limites et on y ajoute la possibilité de quatre étiquetages pour l'analyse *JT2*. On constate par ailleurs que cette dernière possibilité n'est véritablement efficace que dans le cas d'extrapolation à hautes luminosités intégrées (voir les détails dans le paragraphe suivant). Le Tableau 4.21 permet également de vérifier que l'analyse *JT1* est plus sensible aux basses valeurs de $\Delta M, m$ ou aux sections efficaces les plus élevées ($m_{\tilde{g}} = 200, 250$ GeV).

$m_{\tilde{g}}$ (GeV)	σ_{LO} (pb)	$\Delta\sigma_{LO}$ (pb)	σ_{NLO} (pb)	$\Delta\sigma_{NLO}$ (pb)	K-factor
200	7.670	0.01% (stat) +19.65% (PDF) -13.25% +42.11% (scale) -27.12% +46.47% (total) -30.18%	12.000	0.01% (stat) +21.28% (PDF) -13.11% +18.33% (scale) -17.00% +28.09% (total) -21.47%	1.57^{+0.85}_{-0.58}
250	1.750	0.05% (stat) +18.50% (PDF) -12.90% +42.29% (scale) -27.43% +46.16% (total) -30.31%	2.640	0.07% (stat) +20.36% (PDF) -12.38% +18.94% (scale) -18.18% +27.81% (total) -22.00%	1.51^{+0.81}_{-0.56}
300	0.400	0.21% (stat) +19.54% (PDF) -11.81% +43.75% (scale) -28.25% +52.08% (total) -30.62%	0.590	0.29% (stat) +20.16% (PDF) -12.07% +21.36% (scale) -18.81% +29.37% (total) -22.35%	1.48^{+0.88}_{-0.56}
325	0.207	0.40% (stat) +20.29% (PDF) -11.74% +44.93% (scale) -28.50% +49.30% (total) -30.83%	0.304	0.57% (stat) +20.40% (PDF) -12.20% +22.37% (scale) -19.41% +30.28% (total) -22.93%	1.47^{+0.85}_{-0.56}
350	0.100	0.82% (stat) +20.67% (PDF) -12.91% +46.00% (scale) -29.00% +50.44% (total) -31.75%	0.147	1.18% (stat) +22.01% (PDF) -12.02% +23.13% (scale) -20.41% +31.95% (total) -23.72%	1.46^{+0.88}_{-0.58}
375	0.051	1.63% (stat) +21.98% (PDF) -13.13% +45.45% (scale) -29.64% +50.51% (total) -32.46%	0.074	2.34% (stat) +22.03% (PDF) -13.48% +24.36% (scale) -20.30% +32.93% (total) -24.48%	1.46^{+0.88}_{-0.59}
400	0.025	3.34% (stat) +23.01% (PDF) -14.41% +48.78% (scale) -28.46% +54.04% (total) -32.07%	0.036	4.82% (stat) +23.76% (PDF) -13.73% +25.07% (scale) -20.89% +34.88% (total) -25.46%	1.46^{+0.94}_{-0.60}

TAB. 4.20 – Les incertitudes théoriques sur les sections efficaces du signal calculées avec Prospino. La mention *scale* désigne les incertitudes dues aux choix des échelles, la mention *PDF* celles dues au fonction de densité de partons et *stat* celles dues à la statistique.

Points de signal	Canal d'analyse	Nombre de jets étiquetés	Coupures optimales (GeV)	Efficacités (%)
(375,250,75)	<i>JT2</i>	3 <i>Loose</i>	$\cancel{E}_T > 100, H_T > 300$	5.5
(375,150,75)	<i>JT2</i>	3 <i>Loose</i>	$\cancel{E}_T > 100, H_T > 300$	4.1
(350,330,75)	<i>JT1</i>	1 <i>Tight</i>	$\cancel{E}_T > 200, H_T > 250$	5.3
(350,200,75)	<i>JT2</i>	3 <i>Loose</i>	$\cancel{E}_T > 100, H_T > 300$	4.0
(350,150,75)	<i>JT2</i>	3 <i>Loose</i>	$\cancel{E}_T > 100, H_T > 300$	3.4
(350,105,75)	<i>JT2</i>	3 <i>Loose</i>	$\cancel{E}_T > 100, H_T > 300$	2.2
(325,305,75)	<i>JT1</i>	1 <i>Medium</i>	$\cancel{E}_T > 175, H_T > 225$	6.2
(325,250,75)	<i>JT2</i>	2 <i>Loose</i>	$\cancel{E}_T > 125, H_T > 180$	8.2
(325,200,75)	<i>JT2</i>	3 <i>Loose</i>	$\cancel{E}_T > 100, H_T > 180$	4.5
(325,150,75)	<i>JT2</i>	3 <i>Loose</i>	$\cancel{E}_T > 100, H_T > 300$	2.9
(300,280,75)	<i>JT1</i>	1 <i>Medium</i>	$\cancel{E}_T > 150, H_T > 200$	6.8
(300,250,75)	<i>JT1</i>	2 <i>Loose</i>	$\cancel{E}_T > 125, H_T > 175$	5.3
(300,200,75)	<i>JT2</i>	2 <i>Tight</i>	$\cancel{E}_T > 100, H_T > 180$	6.5
(300,150,75)	<i>JT2</i>	2 <i>Tight</i>	$\cancel{E}_T > 100, H_T > 200$	5.1
(300,105,75)	<i>JT2</i>	3 <i>Loose</i>	$\cancel{E}_T > 100, H_T > 180$	1.8
(290,260,60)	<i>JT1</i>	2 <i>Loose</i>	$\cancel{E}_T > 100, H_T > 225$	5.6
(275,265,75)	<i>JT1</i>	1 <i>Tight</i>	$\cancel{E}_T > 125, H_T > 100$	8.3
(275,250,75)	<i>JT1</i>	1 <i>Tight</i>	$\cancel{E}_T > 125, H_T > 150$	6.6
(275,200,75)	<i>JT2</i>	2 <i>Tight</i>	$\cancel{E}_T > 100, H_T > 180$	5.0
(275,150,75)	<i>JT2</i>	2 <i>Tight</i>	$\cancel{E}_T > 100, H_T > 200$	5.0
(275,105,75)	<i>JT1</i>	2 <i>Loose</i>	$\cancel{E}_T > 100, H_T > 100$	3.1
(250,230,75)	<i>JT1</i>	1 <i>Tight</i>	$\cancel{E}_T > 100, H_T > 125$	6.8
(250,150,75)	<i>JT1</i>	1 <i>Tight</i>	$\cancel{E}_T > 100, H_T > 125$	5.4
(250,105,75)	<i>JT1</i>	1 <i>Loose</i>	$\cancel{E}_T > 100, H_T > 150$	4.3
(250,95,75)	<i>JT1</i>	1 <i>Loose</i>	$\cancel{E}_T > 100, H_T > 150$	4.6
(200,180,75)	<i>JT1</i>	1 <i>Loose</i>	$\cancel{E}_T > 100, H_T > 125$	3.6
(200,150,75)	<i>JT1</i>	1 <i>Tight</i>	$\cancel{E}_T > 100, H_T > 125$	2.1
(200,95,75)	<i>JT1</i>	1 <i>Loose</i>	$\cancel{E}_T > 100, H_T > 125$	2.4

TAB. 4.21 – Résumé des meilleurs jeux de coupures par point de signal supersymétrique et efficacités correspondantes.

Evolution des sections efficaces observées

L'évolution des sections efficaces limites en fonction de la masse $\tilde{g}$ est donnée pour différents cas. Chacun des cas correspond à une masse $m_{\tilde{b}_1}$ fixée. On donne quelques exemples sur les Figures 4.38, 4.39, 4.40 et 4.41. Sur ces figures, on trouve les évolutions des sections efficaces limites, attendues et observées, correspondant aux canaux d'analyse les plus favorables. Les extrapolations des résultats, que l'on obtiendrait à une luminosité intégrée de $8fb^{-1}$, sont également données en pointillés. De plus, par extrapolation linéaire entre les points générés, on peut déterminer les masses limites pour chacun des cas considéré. La masse limite est alors l'intersection entre cette extrapolation et la bande jaune des figures précédemment citées. Cette bande jaune correspond, en fait, à la section efficace théorique plus ou moins l'incertitude sur cette section efficace. Les résultats et toutes les figures de cette section correspondent au cas $m_{\tilde{q}} = 1$ TeV et $m_{\tilde{\chi}_1^0} = 75$ GeV. Des courbes similaires avec une masse $m_{\tilde{q}} = 400$ GeV ou $m_{\tilde{q}} = 500$ GeV ont également permis de compléter les résultats finaux, qui seront donnés dans la suite.

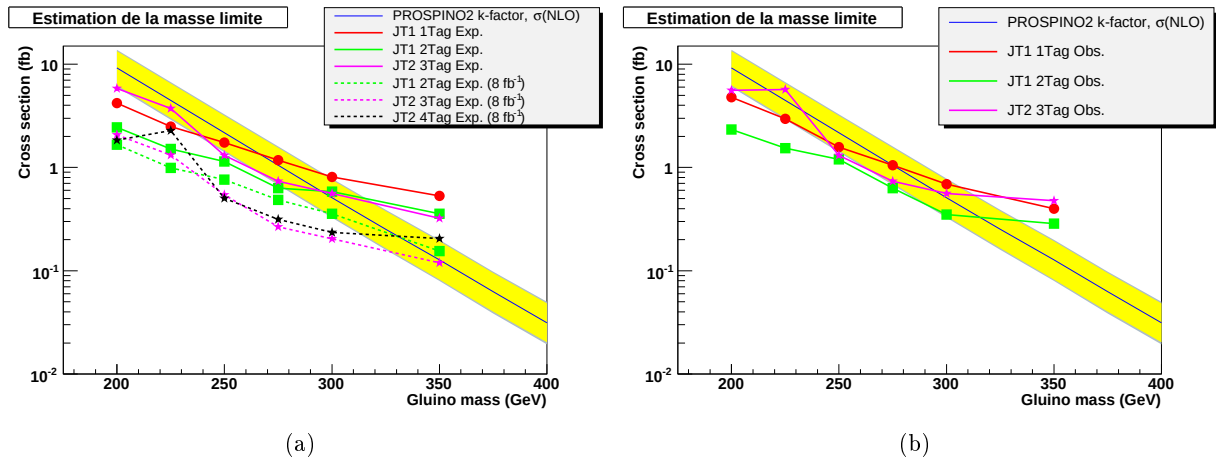


FIG. 4.38 – Evolution de la limite sur la section efficace attendue (a) et observée (b) pour $m_{\tilde{b}_1} = 105$ GeV pour les deux canaux d'analyse.

De chacune des figures, on extrait plusieurs informations. L'évolution de la section efficace attendue donne la combinaison (canal d'analyse et nombre de jets étiquetés) qui est la plus sensible au cas étudié. Par exemple, la Figure 4.38, correspondant au cas $m_{\tilde{b}_1} = 105$ GeV, nous indique que l'analyse *JT1* avec 2 jets étiquetés est la combinaison la plus sensible pour des basses masses de $\tilde{g}$. Aux hautes masses, c'est la combinaison composée de l'analyse *JT2* avec 3 jets étiquetés, qui devient la plus sensible. Dans la zone limite, c'est-à-dire proche de l'intersection avec la section efficace théorique minorée de son incertitude, c'est la première combinaison qui permet d'estimer la masse limite attendue minimale sur le $\tilde{g}$. On utilise, par conséquent, cette combinaison pour déterminer la masse limite observée. On montre alors que, pour $m_{\tilde{b}_1} = 105$ GeV, $m_{\tilde{g}} > 286$ GeV à 95 % de niveau de confiance. On reproduit cette stratégie pour différentes masses de $\tilde{b}_1$ fixées et également dans le cas où $m_{\tilde{g}} - m_{\tilde{b}_1} = 5, 10, 20, 25$ GeV (voir par exemple les Figure 4.42 et 4.43). Dans ces derniers cas, seuls les résultats de l'analyse *JT1* sont montrés sur les Figures 4.42 et 4.43. L'analyse *JT2* n'est pas assez efficace pour de si faibles différences de masse entre le gluino et sbottom. La complémentarité des deux analyses permet alors d'exclure de plus grandes masses

de gluino pour une masse donnée de sbottom.

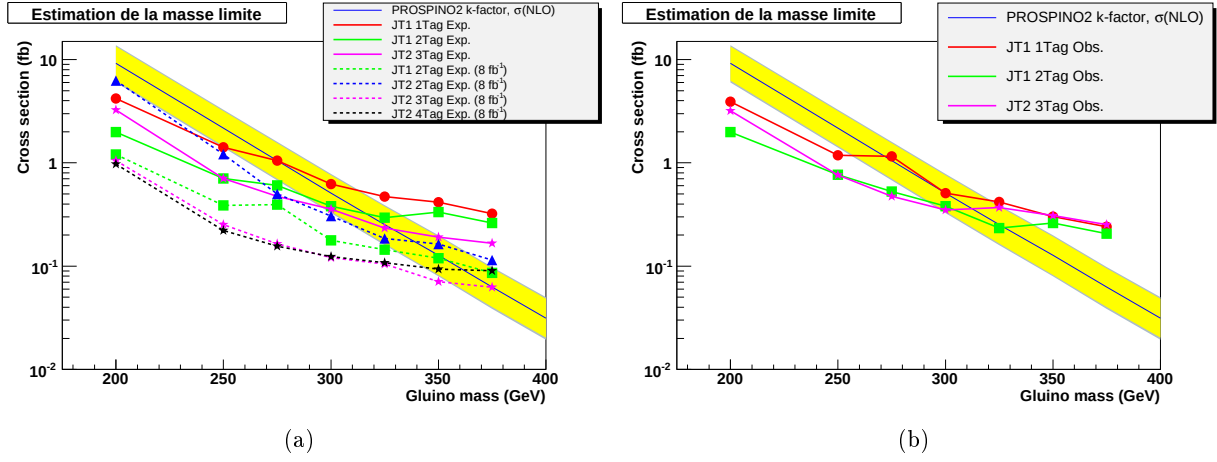


FIG. 4.39 – Evolution de la limite sur la section efficace attendue (a) et observée (b) pour $m_{\tilde{b}_1} = 150$ GeV pour les deux canaux d'analyse.

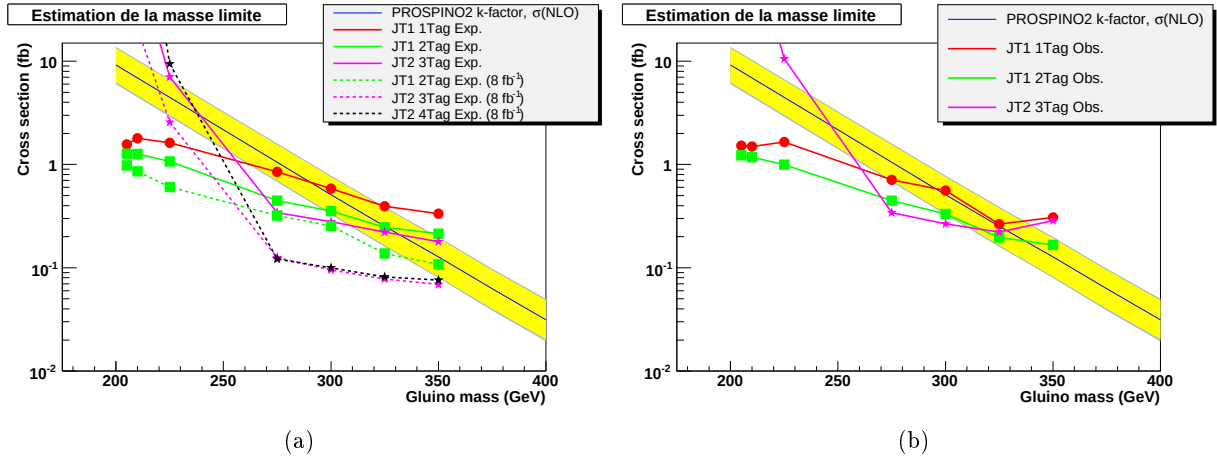


FIG. 4.40 – Evolution de la limite sur la section efficace attendue (a) et observée (b) pour $m_{\tilde{b}_1} = 200$ GeV pour les deux canaux d'analyse.

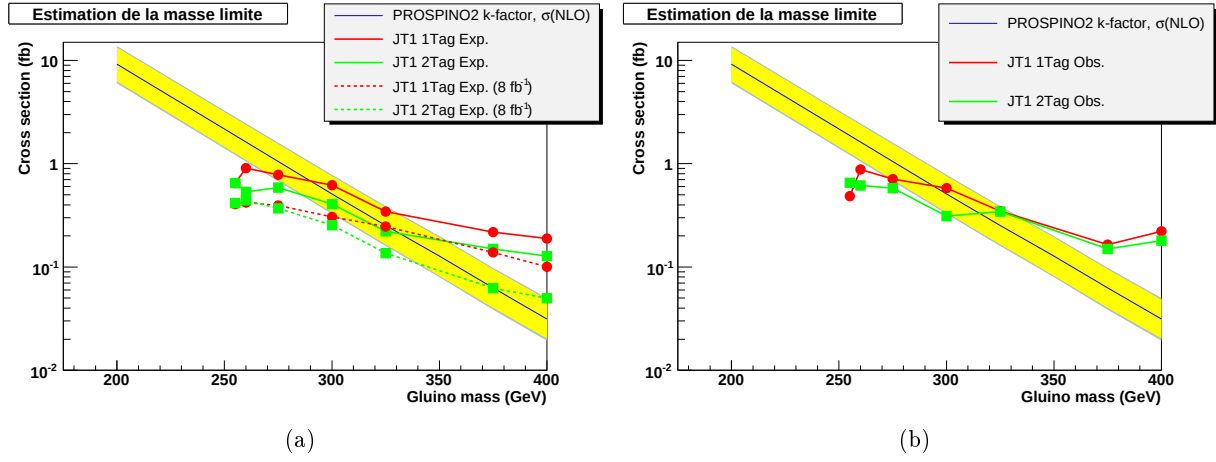


FIG. 4.41 – Evolution de la limite sur la section efficace attendue (a) et observée (b) pour $m_{\tilde{b}_1} = 250$ GeV pour les deux canaux d'analyse.

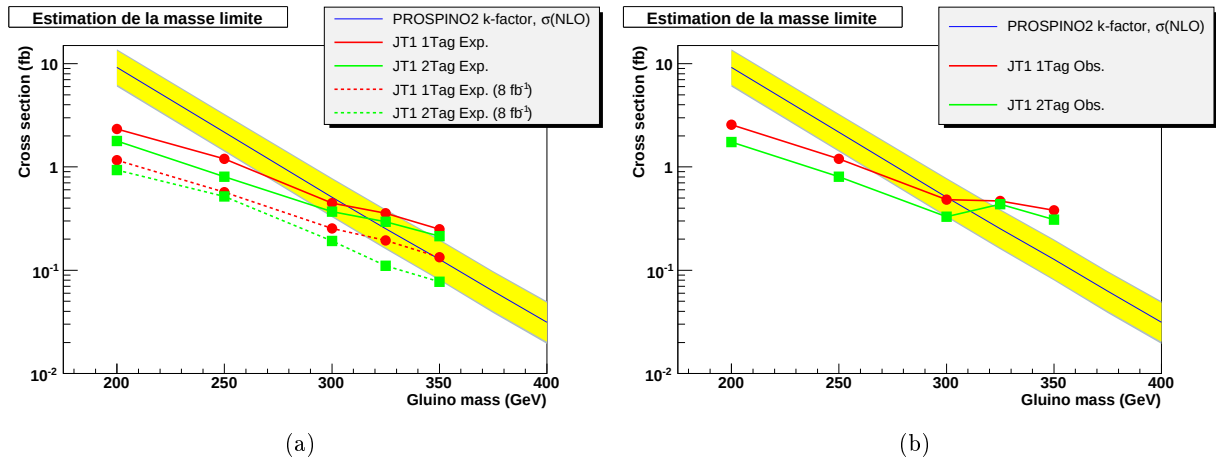


FIG. 4.42 – Evolution de la limite sur la section efficace attendue (a) et observée (b) pour $m_{\tilde{g}} - m_{\tilde{b}_1} = 20$ GeV pour les deux canaux d'analyse.

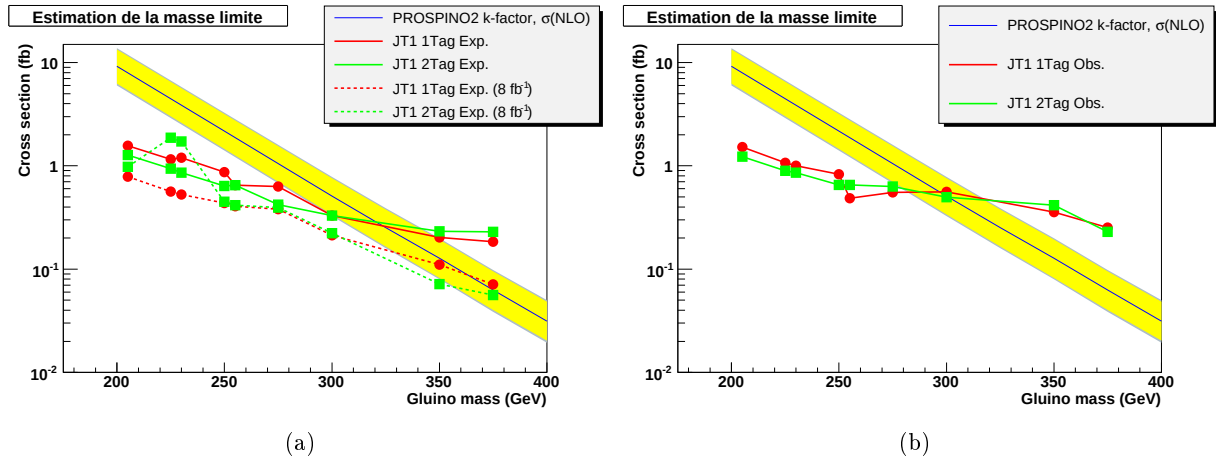


FIG. 4.43 – Evolution de la limite sur la section efficace attendue (a) et observée (b) pour $m_{\tilde{g}} - m_{\tilde{b}_1} = 5$ GeV pour les deux canaux d'analyse.

Masses exclues

L'étude de l'évolution des sections efficaces limites fournit les limites attendues et observées sur les masses du $\tilde{g}$ et du $\tilde{b}_1$. Ces limites sont résumées dans le Tableau 4.22 pour les différents cas étudiés et pour $m_{\tilde{q}} = 1$ TeV. Dans ce tableau, on indique également le meilleur canal d'analyse dans la zone limite, où l'exclusion devient impossible à 95 % de niveau de confiance.

Conditions	Limites attendues $m_{\tilde{g}}$ en GeV	Limites observées $m_{\tilde{g}}$ en GeV	Canal d'analyse et nombre de jets étiquetés
$m_{\tilde{b}_1} = 95$ GeV	< 261	< 266	<i>JT1</i> , 2 jets
$m_{\tilde{b}_1} = 105$ GeV	< 277	< 286	<i>JT1</i> , 2 jets
$m_{\tilde{b}_1} = 150$ GeV	< 295	< 296	<i>JT2</i> , 3 jets
$m_{\tilde{b}_1} = 175$ GeV	< 289	< 307	<i>JT2</i> , 3 jets
$m_{\tilde{b}_1} = 200$ GeV	< 308	< 309	<i>JT2</i> , 3 jets
$m_{\tilde{b}_1} = 225$ GeV	< 300	< 300	<i>JT1</i> , 2 jets
$m_{\tilde{b}_1} = 250$ GeV	< 285	< 301	<i>JT1</i> , 2 jets
$m_{\tilde{g}} - m_{\tilde{b}_1} = 5$ GeV	< 300	< 281	<i>JT1</i> , 2 jets
$m_{\tilde{g}} - m_{\tilde{b}_1} = 20$ GeV	< 290	< 299	<i>JT1</i> , 2 jets

TAB. 4.22 – Limites attendues et observées sur les masses du $\tilde{g}$ et du $\tilde{b}_1$ dans le cas où $m_{\tilde{q}} = 1$ TeV et $m_{\tilde{\chi}_1^0} = 75$ GeV.

Le Tableau 4.23 montre les résultats obtenus pour $m_{\tilde{q}} = 500$ GeV. La section efficace de production de paire de gluinos étant plus faibles, les limites sur les masses des particules supersymétriques sont plus faibles. Il en résulte également que les meilleurs canaux d'analyses changent par rapport au cas où $m_{\tilde{q}} = 1$ TeV.

Conditions	Limites attendues $m_{\tilde{g}}$ en GeV	Limites observées $m_{\tilde{g}}$ en GeV	Canal d'analyse et nombre de jets étiquetés
$m_{\tilde{b}_1} = 95$ GeV	< 229	< 236	<i>JT1</i> , 2 jets
$m_{\tilde{b}_1} = 105$ GeV	< 230	< 231	<i>JT1</i> , 2 jets
$m_{\tilde{b}_1} = 150$ GeV	< 253	< 254	<i>JT2</i> , 3 jets
$m_{\tilde{b}_1} = 175$ GeV	< 253	< 256	<i>JT1</i> , 2 jets
$m_{\tilde{b}_1} = 200$ GeV	< 257	< 261	<i>JT1</i> , 2 jets
$m_{\tilde{b}_1} = 225$ GeV	< 248	< 250	<i>JT1</i> , 2 jets
$m_{\tilde{b}_1} = 250$ GeV	< 258	< 256	<i>JT1</i> , 2 jets
$m_{\tilde{g}} - m_{\tilde{b}_1} = 5$ GeV	< 258	< 257	<i>JT1</i> , 2 jets
$m_{\tilde{g}} - m_{\tilde{b}_1} = 10$ GeV	< 254	< 252	<i>JT1</i> , 2 jets
$m_{\tilde{g}} - m_{\tilde{b}_1} = 20$ GeV	< 246	< 250	<i>JT1</i> , 2 jets
$m_{\tilde{g}} - m_{\tilde{b}_1} = 25$ GeV	< 246	< 251	<i>JT1</i> , 2 jets

TAB. 4.23 – Limites attendues et observées sur les masses du $\tilde{g}$ et du $\tilde{b}_1$ dans le cas où $m_{\tilde{q}} = 500$ GeV et $m_{\tilde{\chi}_1^0} = 75$ GeV.

4.5.4 Contours d'exclusion dans le plan $(m_{\tilde{g}}, m_{\tilde{b}_1})$

A partir des limites observées sur les masses des particules supersymétriques étudiées, on construit une zone d'exclusion dans le plan $(m_{\tilde{g}}, m_{\tilde{b}_1})$. A l'intérieur de cette zone, tous les points de signal sont exclus à 95 % de niveau de confiance dans le cadre des hypothèses fixées pour cette étude. Ces hypothèses sont rappelées sur la Figure 4.44, qui donne le contour d'exclusion obtenu pour différentes masses de squarks légers.

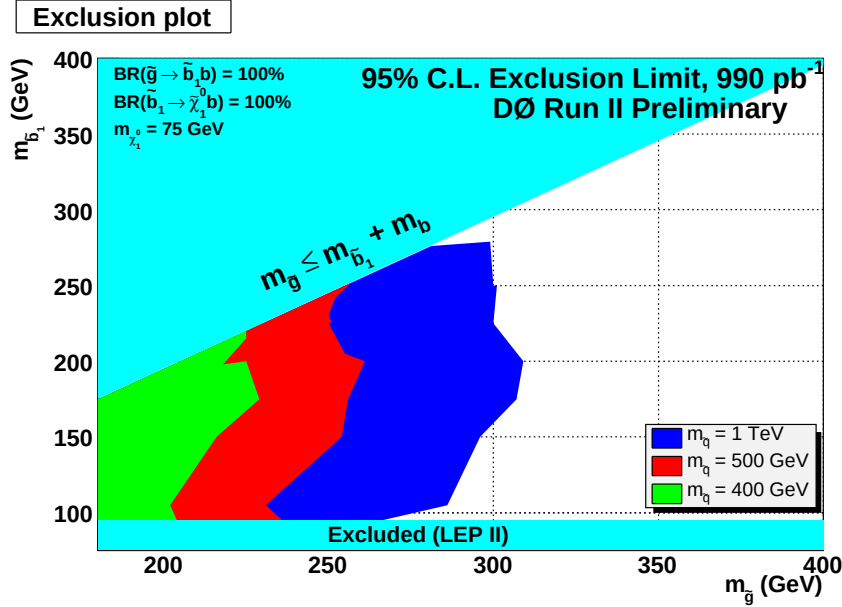


FIG. 4.44 – Contour d'exclusion dans le plan $(m_{\tilde{g}}, m_{\tilde{b}_1})$ pour $m_{\tilde{\chi}_1^0} = 75$ GeV.

Les zones d'exclusion représentées sur la Figure 4.44 correspondent aux masses limites obtenues. On peut également tracer des zones d'exclusion pour les masses limites attendues, c'est-à-dire celles déterminées à partir des données simulées uniquement. La comparaison des contours d'exclusion attendus et observés est donnée sur la Figure 4.45. Elle confirme les comparaisons entre les données réelles et les données simulées obtenues dans les précédents paragraphes, qui montrent un bon accord entre les événements collectés et les prédictions théoriques du modèle standard.

On peut également comparer l'ensemble de ces résultats, sur les masses limites des gluinos et des sbottoms, avec les résultats obtenus dans les références [84] et [83]. Pour mener à bien cette comparaison, il faut effectuer les mêmes hypothèses théoriques. Il faut, par conséquent, utiliser la même méthode pour normaliser les signaux supersymétriques. Les sections efficaces utilisées sont donc les sections efficaces *NLO* directement données par l'outil Prospino avec les incertitudes théoriques fournies également par cet outil. Cette méthode augmente la section efficace de production des particules supersymétriques, car la section efficace *LO* donnée par le générateur Pythia est systématiquement inférieure à celle donnée par Prospino. De plus, cette méthode réduit les incertitudes systématiques sur la section efficace de production du signal, car elle ne tient pas compte des incertitudes sur la section efficace *LO*. Il faut finalement utiliser les mêmes hypothèses sur les masses des particules supersymétriques, c'est-à-dire $m_{\tilde{\chi}_1^0} = 75$ GeV et $m_{\tilde{q}} > 400$ GeV. On obtient alors les résultats de la Figure 4.46. Cette comparaison permet de mettre en avant la

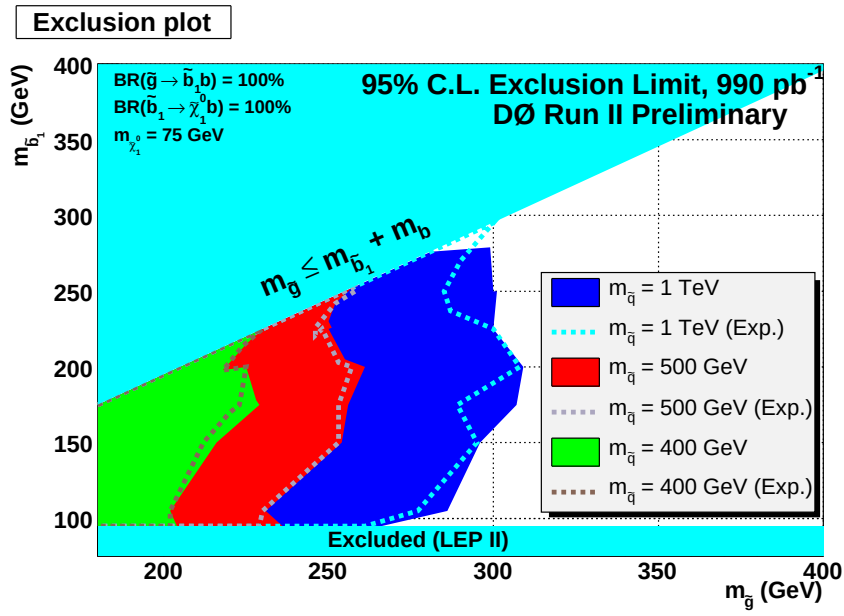


FIG. 4.45 – Contour d'exclusion dans le plan $(m_{\tilde{g}}, m_{\tilde{b}_1})$ pour $m_{\tilde{\chi}_1^0} = 75$ GeV et comparaison entre les limites attendues et observées.

complémentarité entre chacune des analyses.

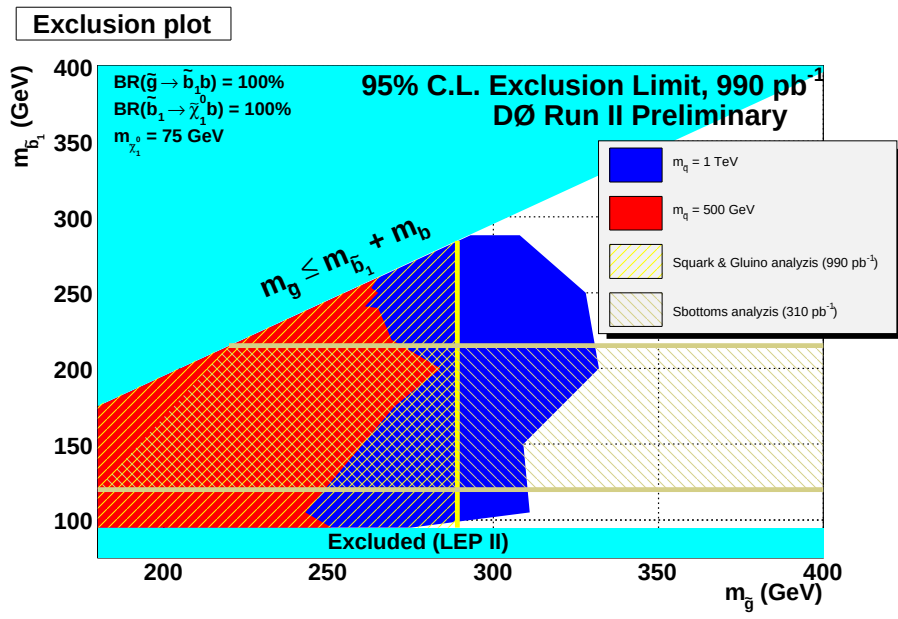


FIG. 4.46 – Contour d'exclusion dans le plan $(m_{\tilde{g}}, m_{\tilde{b}_1})$ pour $m_{\tilde{\chi}_1^0} = 75$ GeV et comparaison avec les analyses des références [84] et [83].

Conclusion

La supersymétrie définit un cadre théorique élégant qui ajoute une nouvelle symétrie fondamentale à celles déjà connues dans le modèle standard. De ce fait, le modèle standard deviendrait une théorie effective, valable jusqu'à des énergies de l'ordre de quelques centaines de GeV. Cette nouvelle symétrie, entre les bosons et les fermions, propose donc une possibilité d'extension du modèle standard valable à plus hautes énergies. Elle pourrait ainsi pallier les imperfections théoriques de ce dernier et elle permettrait de rendre compte d'observations expérimentales inexplicables à ce jour.

Ce nouveau cadre théorique prédit de nouvelles particules. La recherche de ces nouvelles particules au TeVatron est un enjeu important pour la physique des particules. C'est pour cette raison que de nombreux canaux de recherches avec des modes de productions de particules supersymétriques et des états finaux variés sont investis depuis le début du *Run IIa* en 2001 par la collaboration DØ. L'analyse proposée dans cette thèse est la recherche d'une production directe d'une paire de gluinos se désintégrant dans un état final comprenant quatre quarks b et deux χ_1^0 , particule supersymétrique stable dans le cadre théorique utilisé, i.e. le *MSSM*.

Cette analyse comporte de nombreux jets dans l'état final. Or le bruit de fond principal d'un collisionneur hadronique comme le TeVatron correspond aux processus de *QCD*, c'est-à-dire à des processus produisant de nombreux jets dans l'état final. L'intégralité des collisions ne pouvant être traitées et enregistrées, il est donc important de maîtriser le système de déclenchement afin d'isoler les topologies intéressantes tout en rejetant un maximum d'événements des processus de *QCD*. Des conditions de déclenchement spécialement conçues pour des topologies avec des jets et de l'énergie transverse manquante ont été étudiées et mise en place pour le *Run IIa* et le *Run IIb*, qui a commencé au printemps 2006. Ces conditions ont permis de mener à bien la recherche du signal supersymétrique évoqué précédemment. La compréhension de ce système de déclenchement, la connaissance de l'énergie des jets et la possibilité d'étiqueter ces jets comme étant des jets de quarks b permettent finalement d'isoler le signal recherché.

La recherche entreprise dans cette thèse ne montre finalement aucun excès dans les données du détecteur par rapport aux prédictions théoriques du modèle standard. En effet, on observe un bon accord entre les données réelles et les données simulées quel que soit le jeu de coupures choisi. Cette absence de déviation par rapport au modèle standard permet d'estimer des limites sur les sections efficaces de production de paires de gluinos au TeVatron et dans le même temps d'estimer les masses limites correspondant à ces sections efficaces. On exclut ainsi les gluinos jusqu'à une masse de 309 GeV.

Les résultats obtenus améliorent la sensibilité de l'expérience DØ à un signal supersymétrique produisant des gluinos. En effet, l'analyse générique sans étiquetage des jets de quark b [84] exclut des gluinos jusqu'à une masse de 289 GeV contre 334 GeV pour l'analyse présentée dans ce chapitre (si on se place dans les mêmes conditions de normalisation). Cette amélioration obtenue uniquement sous l'hypothèse de $\tilde{b}_1$ plus léger que les $\tilde{g}$, rend les résultats de ces deux analyses complémentaires. De plus, cette analyse est également complémentaire avec la recherche de production

directe de $\tilde{b}_1$ [83]. En effet, pour une masse $m_{\tilde{\chi}_1^0} = 75$ GeV, la recherche de production directe exclut des masses de $\tilde{b}_1$ entre 115 et 205 GeV, alors que la recherche indirecte exclut des masses entre 95 et 289 GeV. La différence entre ces deux analyses est que la première ne dépend pas de la masse du gluino en première approximation, alors que la seconde en dépend directement dans la détermination de la masse limite sur le $\tilde{b}_1$. Ces résultats ne sont pas directement comparables mais complémentaires.

Cette analyse a utilisé environ 1 fb^{-1} de luminosité intégrée, c'est-à-dire l'intégralité du *Run IIa*. Le début du *Run IIb* montre d'excellentes performances de l'accélérateur combinées avec de grandes efficacités dans la prise de données du détecteur DØ. Ces efficacités sont permises notamment grâce à un étalonnage plus précis du système de déclenchement propre au calorimètre et grâce également à l'évolution des conditions de déclenchement étudiée dans cette thèse. En effet, le détecteur DØ prend régulièrement des données avec des pics de luminosité instantanée autour de $250 \times 10^{30} \text{ cm}^{-2} \text{ s}^{-1}$. Ainsi, plus du double de luminosité intégrée est actuellement disponible pour poursuivre les recherches entreprises au *Run IIa*. Cette statistique importante devrait permettre d'améliorer les résultats présentés dans cette thèse, voire de mener à des découvertes de signaux supersymétriques. En effet, si on fait l'hypothèse très conservatrice que les incertitudes systématiques restent inchangées et qu'on se place à une luminosité intégrée de 8 fb^{-1} , on obtient les extrapolations données par la Figure 4.47. Si le *Run IIb* n'aboutit à aucune découverte de particules supersymétriques, le futur *LHC*, qui commencera à prendre des données au cours de l'année 2008, deviendra rapidement l'outil idéal pour poursuivre ces recherches.

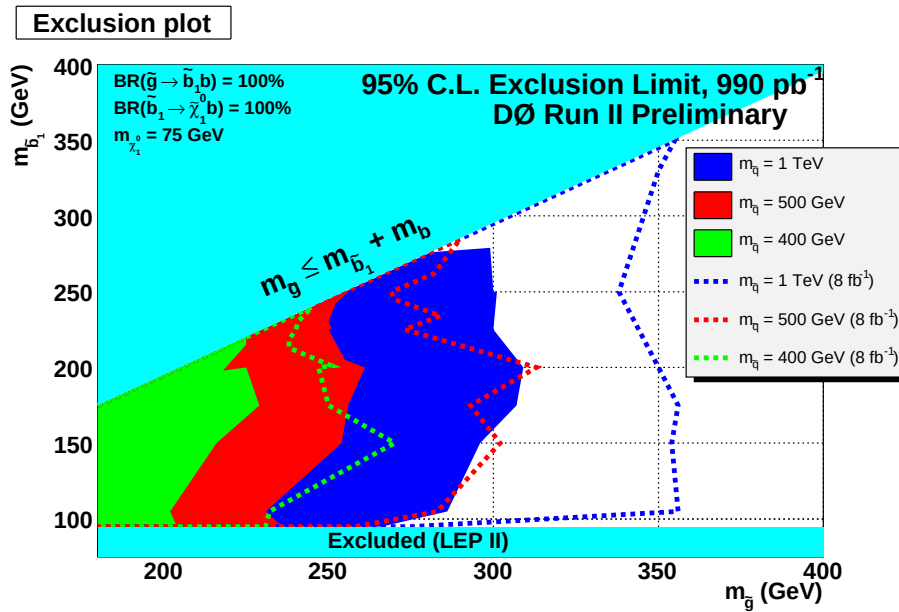


FIG. 4.47 – Contour d'exclusion dans le plan $(m_{\tilde{g}}, m_{\tilde{b}_1})$ pour $m_{\tilde{\chi}_1^0} = 75$ GeV et extrapolations à une luminosité intégrée de 8 fb^{-1} .

Bibliographie

- [1] Albert Einstein. The foundation of the general theory of relativity. *Annalen Phys.*, 49 : 769–822, 1916.
- [2] J. Chadwick. Possible existence of a neutron. *Nature*, 129 :312, 1932.
- [3] A. Einstein. On the electrodynamics of moving bodies. *Annalen Phys.*, 17 :891–921, 1905.
- [4] P.A. Delsart. *Etude du signal $H^0/A^0 \rightarrow \tau\mu$ aux collisionneurs hadroniques et intercalibration du calorimètre de DØ au Run II du TeVatron*. PhD thesis, Université Claude Bernard Lyon-1, 2003.
- [5] A. Deandrea. Interactions électrofaibles et introduction à la supersymétrie, <http://deandrea.cern.ch/deandrea>.
- [6] E. Noether. Invariante variationsprobleme. *Kgl. Ges. Wiss. Nachr. Math.-phys.*, 2 :235, 1918.
- [7] H. Weyl. Electron and gravitation. *Z. Phys.*, 56 :330–352, 1929.
- [8] Paul A. M. Dirac. The quantum theory of electron. *Proc. Roy. Soc. Lond.*, A117 :610–624, 1928.
- [9] P. A. M. Dirac. The quantum theory of electron. 2. *Proc. Roy. Soc. Lond.*, A118 :351, 1928.
- [10] Julian S. Schwinger. Quantum electrodynamics. iii : The electromagnetic properties of the electron : Radiative corrections to scattering. *Phys. Rev.*, 76 :790–817, 1949.
- [11] R. P. Feynman. Mathematical formulation of the quantum theory of electromagnetic interaction. *Phys. Rev.*, 80 :440–457, 1950.
- [12] Chen-Ning Yang and Robert L. Mills. Conservation of isotopic spin and isotopic gauge invariance. *Phys. Rev.*, 96 :191–195, 1954.
- [13] Peter W. Higgs. Spontaneous symmetry breakdown without massless bosons. *Phys. Rev.*, 145 :1156–1163, 1966.
- [14] E. Fermi. An attempt of a theory of beta radiation. 1. *Z. Phys.*, 88 :161–177, 1934.
- [15] F. Reines and C. L. Cowan. Free anti-neutrino absorption cross-section. 1 : Measurement of the free anti-neutrino absorption cross-section by protons. *Phys. Rev.*, 113 :273, 1959.
- [16] S. L. Glashow. Partial symmetries of weak interactions. *Nucl. Phys.*, 22 :579–588, 1961.

- [17] Jeffrey Goldstone, Abdus Salam, and Steven Weinberg. Broken symmetries. *Phys. Rev.*, 127 :965–970, 1962.
- [18] Hideki Yukawa. On the interaction of elementary particles. *Proc. Phys. Math. Soc. Jap.*, 17 :48–57, 1935.
- [19] Murray Gell-Mann. Symmetries of baryons and mesons. *Phys. Rev.*, 125 :1067–1084, 1962.
- [20] S. Abachi et al. Observation of the top quark. *Phys. Rev. Lett.*, 74 :2632–2637, 1995.
- [21] F. Abe et al. Observation of top quark production in $\bar{p}p$ collisions. *Phys. Rev. Lett.*, 74 :2626–2631, 1995.
- [22] D. J. Gross and Frank Wilczek. Ultraviolet behavior of non-abelian gauge theories. *Phys. Rev. Lett.*, 30 :1343–1346, 1973.
- [23] H. David Politzer. Reliable perturbative results for strong interactions? *Phys. Rev. Lett.*, 30 :1346–1349, 1973.
- [24] R. P. Feynman. Space-time approach to quantum electrodynamics. *Phys. Rev.*, 76 :769–789, 1949.
- [25] F. J. Dyson. The radiation theories of tomonaga, schwinger, and feynman. *Phys. Rev.*, 75 :486–502, 1949.
- [26] Page web du *particle data group*,
http://pdg.lbl.gov/2006/tables/contents_tables.html.
- [27] Kenneth Lane. Two lectures on technicolor,
<http://www.citebase.org/abstract?id=oai:arxiv.org:hep-ph/0202255>, 2002.
- [28] Dimitri V. Nanopoulos. Protons are not forever. Invited talk given at Seminar on Proton Stability, Madison, Wisc., Dec 8, 1978.
- [29] J. Wess and B. Zumino. Supergauge transformations in four-dimensions. *Nucl. Phys.*, B70 :39–50, 1974.
- [30] Michael B. Green and John H. Schwarz. Anomaly cancellation in supersymmetric d=10 gauge theory and superstring theory. *Phys. Lett.*, B149 :117–122, 1984.
- [31] Edward Witten. Noncommutative geometry and string field theory. *Nucl. Phys.*, B268 :253, 1986.
- [32] David J. Gross, Jeffrey A. Harvey, Emil J. Martinec, and Ryan Rohm. Heterotic string theory. 1. the free heterotic string. *Nucl. Phys.*, B256 :253, 1985.
- [33] Julien Welzel, David Gherson, and John R. Ellis. New particle physics. (in french). 2005.
- [34] P. Fayet and S. Ferrara. Supersymmetry. *Phys. Rept.*, 32 :249–334, 1977.
- [35] Sidney R. Coleman and J. Mandula. All possible symmetries of the s matrix. *Phys. Rev.*, 159 :1251–1256, 1967.

- [36] P. Fayet and J. Iliopoulos. Spontaneously broken supergauge symmetries and goldstone spinors. *Phys. Lett.*, B51 :461–464, 1974.
- [37] L. O’Raifeartaigh. Spontaneous symmetry breaking for chiral scalar superfields. *Nucl. Phys.*, B96 :331, 1975.
- [38] A.M. Magnan. *Recherche de particules supersymétriques se désintégrant en R-parité violée (couplage λ_{121}) dans un état final à trois leptons, avec les données du Run II de l’expérience DØ au Tevatron*. PhD thesis, Université Joseph-Fourier - Grenoble I, 2005.
- [39] A.C. Le Bihan. *Identification des leptons taus dans l’expérience DØ auprès du TeVatron et recherche de particules supersymétriques se désintégrant avec R-parité violée (couplage λ_{133})*. PhD thesis, Université Louis Pasteur de Strasbourg, 2005.
- [40] Stephen P. Martin. A supersymmetry primer. 1997.
- [41] Page web du laboratoire fermilab,
<http://www.fnal.gov/>, .
- [42] Page web de l’expérience minos,
<http://www-boone.fnal.gov/>, .
- [43] Page web de l’expérience miniboone,
<http://www-numi.fnal.gov:8875/>, .
- [44] Page web de l’expérience cdf,
<http://www-cdf.fnal.gov/>, .
- [45] Page web de l’expérience dØ
<http://www-d0.fnal.gov/>, .
- [46] Operation rookie book,
http://www-bdnew.fnal.gov/operations/rookie_books/rbooks.html, .
- [47] Run ii handbook,
<http://www-ad.fnal.gov/runii/index.html>, .
- [48] S. Laplace and al. La physique au lhc. *Proceedings du cours de G. Unal à l’école de Gif 2004 (CERN)*, 2006.
- [49] S. Eidelman et al. Review of particle physics. *Phys. Lett.*, B592 :1, 2004.
- [50] Quelques détails sur la chaîne d’accélération de fermilab,
<http://www-bd.fnal.gov/public/chain.html>, .
- [51] Illustration du principe de fonctionnement d’un accélérateur linéaire,
<http://www.sciences.univ-nantes.fr/physique/perso/gtulloue/meca/charges/linac.html>, .
- [52] Page web du l’injecteur principal,
<http://www-fmi.fnal.gov/>, .
- [53] V. M. Abazov et al. The upgraded dØ detector. 2005.

- [54] DØ silicon tracker technical design report,
http://d0server1.fnal.gov/projects/silicon/www/tdr_final.ps, .
- [55] Y. Arnoud. Cours de dea,
<http://lpsc.in2p3.fr/arnoud/master1/>.
- [56] Charles Kittel. *Introduction à la physique de l'état solide*. Dunod Université. Dunod, Paris, 3rd edition, 1972.
- [57] La page web du détecteur cft,
http://d0server1.fnal.gov/projects/scifi/cft_home.html, .
- [58] M. Adams et al. Design report of the central preshower detector for the dØ upgrade. *DØ Note 3014*, 1996.
- [59] A. Gordeev et al. Design report of the forward preshower detector for the dØ upgrade. *DØ Note 3445*, 1998.
- [60] S. Abachi et al. The dØ detector. *Nucl. Instr. and Methods.*, A338 :185, 1994.
- [61] K. De et al. Technical design report for the upgrade of the icd for dØ run ii. *DØ Note 2686*, 1997.
- [62] M. Ridet. *Reconstruction du flux d'énergie et recherche de squarks et gluinos dans l'expérience DØ*. PhD thesis, Université Paris XI Orsay, 2002.
- [63] J. Zhu et al. Determination of electron energy scale and energy resolution using p14 z->ee data. *DØ Note 4323*, 2003.
- [64] J. Stark. Présentation pour le groupe mesure de la masse du w sur la résolution de l'énergie des électrons,
it <http://www-d0.hef.kun.nl/askarchive.php?base=agenda&categ=a0651&id=a0651s1t0/transparencies>.
- [65] M. Agelou et al. 0,5 and 0,7 jet p_T resolution using jes v05-03-00. *DØ Note 4775*, 2005.
- [66] The muon system of the run ii dØ detector,
http://www-d0.fnal.gov/hardware/upgrade/muon_upgrade/d0_private/nim1.pdf, .
- [67] Page web du moniteur de luminosité,
<http://www-d0online.fnal.gov/www/groups/lum/index.html>, .
- [68] A. Lo et al. Luminosity monitor technical design report. *DØ Note 3320*, 1997.
- [69] B. Casey et al. Improved determination of the dØ luminosity. *DØ Note 5140*, 2006.
- [70] G. Blazey et al. The d0 run ii trigger. *FERMILAB-CONF-97-395-E*, 1997.
- [71] J.T. Linnenmann. The dØ level 2 trigger. *DØ Note 3334*, 1997.
- [72] D. Edmunds et al. Technical design report for the level 2 global processor. *DØ Note 3402*, 1998.

- [73] Robert D. Angstadt et al. The dzero level 3 data acquisition system. *IEEE Trans. Nucl. Sci.*, 51 :445–450, 2004.
- [74] A. Boehnlein et al. Description of d0 l3 trigger software components. *DØ Note 3630*, 1999.
- [75] A. Verzocchi et al. The trigger_rate_tool package : a tool to estimate rates and overlaps for the development of physics trigger lists. *DØ Note 4640*, 2004.
- [76] Page web du simulateur du système de déclenchement de d0 *d0trigsim*, <http://www-d0.fnal.gov/computing/d0trigsim.html>, .
- [77] A. Zabi. *Recherche de Leptoquarks dans la topologie à jets et énergie transverse manquante avec le détecteur DØ au TeVatron*. PhD thesis, Université Paris VII, 2004.
- [78] Zabi A. et al. A trigger for jets and missing et. *DØ Note 4315*, 2004.
- [79] Duperrin A. et al. The v13 physics trigger list and new phenomena triggers. *DØ Note 4641*, 2004.
- [80] Millet T. et al. Jets and met triggers for the new phenomena group in the v14 and v15 triggers lists. *DØ Note 5120*, 2006.
- [81] Adams M. et al. Level 2 calorimeter preprocessor technical design report, d0 *DØ Note 3651*, 1999.
- [82] Lee Sawyer M. K. Certification studies for the level 3 missing et tool. *DØ Note 4036*, 2002.
- [83] V. M. Abazov et al. Search for pair production of scalar bottom quarks in $p\bar{p}$ collisions at $\sqrt{s} = 1.96\text{-tev}$. *Phys. Rev. Lett.*, 97 :171806, 2006.
- [84] V. M. Abazov et al. Search for squarks and gluinos in events with jets and missing transverse energy in p anti-p collisions at $s^{*}(1/2) = 1.96\text{-tev}$. *Phys. Lett.*, B638 :119–127, 2006.
- [85] L. Duflot et al. Search for large extra spatial dimensions in jets + missing e_t topologies. *DØ Note 4400*, 2005.
- [86] T. Berenc et al. Plans for tevatron runiib. 2001.
- [87] V. M. Abazov et al. Run iib upgrade technical design report. FERMILAB-PUB-02-327-E.
- [88] M. Abolins et al. The run iib trigger upgrade for the d0 experiment. *IEEE Trans. Nucl. Sci.*, 51 :340–344, 2004.
- [89] Demarteau M. Dzero layer 0 conceptual design report. *DØ Note 4415*, 2004.
- [90] M. Abolins et al. The fast trigger for the d0 experiment. *Nucl. Instrum. Methods Phys. Res.*, A289 :543–560, 1990.
- [91] J. Bystricky et al. Algorithms and architecture for the l1 calorimeter trigger at d0 run iib. *IEEE Trans. Nucl. Sci.*, 51 :351–355, 2004.
- [92] S. Lammers et al. Simulation of run iib l1 cal trigger and em algorithm optimization. *DØ Note 4663*, 2004.

- [93] Page web, décrivant les cartes *adf*, du laboratoire du *cea* de saclay,
http://www-clued0.fnal.gov/perez/r2btrigger/adf_description/adf_descri.html.
- [94] Page web, décrivant les cartes *adf*, du laboratoire du *msu* : *michigan state univeristy*,
<http://www.pa.msu.edu/hep/d0/run2b/>.
- [95] Westein M. et al. Gain calibration for the em calorimeter in run ii. *DØ Note 5004*, 2006.
- [96] Cwiok M. et al. Run ii eta-intercalibration of the hadronic calorimeter. *DØ Note 5006*, 2006.
- [97] Kvita J. et al. Run ii phi-intercalibration of the fine hadronic calorimeter. *DØ Note 5005*, 2006.
- [98] Millet T. et al. Calibration of the d0 level 1 calorimeter trigger for runiib. *DØ Note 5199*, 2006.
- [99] Duflot L. et al. *cal_event_quality* package. *DØ Note 4614*, 2004.
- [100] Calvet S. et al. *l1cal2b_met_cert* package and level 1 missinget triggers certification in run iib. *DØ Note 5198*, 2006.
- [101] Martin B. et al. Validation of l1 jet and electron triggers in early runiib data and software tools for regular checks in the future. *DØ Note 5201*, 2006.
- [102] Lacroix F. et al. Higgs and new phenomena jets+met triggers : L3 design and commissioning in v15 runiib trigger list. *DØ Note 5282*, 2006.
- [103] Bloom K. et al. Primary vertex reconstruction by means of adaptative vertex fitting. *DØ Note 4918*, 2005.
- [104] Greenlee H. The dØ kalman track fit. *DØ Note 4303*, 2004.
- [105] Khanov A. et al. Htf : histogramming method for finding tracks. the algorithm description. *DØ Note 3778*, 2000.
- [106] Page web décrivant les algorithmes de reconstruction des traces,
http://www-d0.fnal.gov/global_tracking/.
- [107] Hesketh G. et al. Central track extrapolation through the dØ detector. *DØ Note 4079*, 2002.
- [108] Schwartzman A. et al. DØ tracking performance at high luminosity. *DØ Note 4980*, 2006.
- [109] D'Hondt J. et al. Sensitivity of robust vertex fitting algorithms. *IEEE Trans. Nucl. Sci.*, 51, 2004.
- [110] Schwartzman A. et al. Probabilistic primary vertex selection. *DØ Note 4042*, 2002.
- [111] Calfayan P. et al. Muon identification certification for p17 data. *DØ Note 5157*, 2007.
- [112] Beaudette F. et al. The road method (an algorithm for the identification of electrons in jet). *DØ Note 3976*, 2002.

- [113] Crepe-Renaudin S. Energy corrections for geometry effects for electrons in run ii. *DØ Note 4023*, 2002.
- [114] Kermiche S. et al. Energy scale studies and calibration of the d0 electromagnetic calorimeter using z0 and j/psi -> e+e- run ii events. *DØ Note 4945*, 2005.
- [115] Askew A. et al. Cps variables for photon identification. *DØ Note 4949*, 2005.
- [116] Atramentov O. et al. Photon identification in p17 data. *DØ Note 4976*, 2006.
- [117] Hays J. et al. Single electron efficiencies in p17 data and monte-carlo using d0correct from release p18.05.00. *DØ Note 5105*, 2006.
- [118] Magerkurth A. Parton level corrections for jetcorr 5.3. *DØ Note 4708*, 2005.
- [119] Blazey G.C. et al. Run ii jet physics. *DØ Note 3750*, 2000.
- [120] Abbott B. et al. Fixed cone jet definitions in dØ. *FERMILAB-PUB-97-242-E*, 1997.
- [121] Ellis S.D. et al. Successive combination jet algorithm for hadron collisions. *Phys. Rev.*, D48 : 3160–3166, 1993.
- [122] Stephens S. A comparison of inclusive jets reconstructed with the k_t and cone algorithms. *DØ Note 3010*, 1996.
- [123] J.L. Agram. *Mesure de la section efficace inclusive de production de jets en fonction de leur impulsion transverse dans l'expérience DØ au Fermilab*. PhD thesis, Université de Haute-Alsace, 2004.
- [124] E. Busato. *Recherche de la production électrofaible du quark top dans le canal électron+jets dans l'expérience DØ auprès du Tevatron*. PhD thesis, Université Paris VII, 2005.
- [125] N. Makovec. *Recherche de nouvelle physique dans la topologie à jets et énergie transverse manquante avec le détecteur DØ au Tevatron*. PhD thesis, Université Paris VII, 2006.
- [126] Harel A. Jetid optimization. *DØ Note 4919*, 2005.
- [127] Harel A. et al. Combined jetid efficiency for p17. *DØ Note 5218*, 2006.
- [128] Makovec N. et al. The relative data - monte carlo jet energy scale. *DØ Note 4807*, 2005.
- [129] Page web décrivant les corrections de l'énergie des jets,
it [http ://www-d0.fnal.gov/phys_id/jes/public/plots_v7.1/](http://www-d0.fnal.gov/phys_id/jes/public/plots_v7.1/).
- [130] Abott B. et al. Jet energy scale at dØ. *DØ Note 3287*, 1997.
- [131] Makovec N. et al. Shifting, smearing and removing simulated jets. *DØ Note 4914*, 2005.
- [132] C. Ochando. Présentation pour le groupe recherche du boson de higgs sur la procédure *ssr*,
it [http ://www-d0.hef.kun.nl/askarchive.php?base=agenda&categ=a06904&id=a06904s2t4/transparencie](http://www-d0.hef.kun.nl/askarchive.php?base=agenda&categ=a06904&id=a06904s2t4/transparencie)

- [133] S. Greder. *Etiquetage des quarks beaux et mesure de la section efficace de production de paires de quarks top à $\sqrt{s} = 1.96$ TeV dans l'expérience DØ* PhD thesis, Université Louis Pasteur de Strasbourg, 2004.
- [134] B. Clément. *Production électrofaible du quark top au Run II de l'expérience DØ* PhD thesis, Université Louis Pasteur de Strasbourg, 2006.
- [135] Anghel I.M. et al. Studies of taggability versus number of primary vertices for p17 data. *DØ Note 5240*, 2007.
- [136] Scanlon T. et al. Performance of the nn b-tagging tool on pass2 p14 data. *DØ Note 4890*, 2006.
- [137] Scanlon T. A neural network b-tagging tool. *DØ Note 4889*, 2005.
- [138] Anastasoiaie M. et al. Performance of the nn b-tagging tool on p17 data. *DØ Note 5213*, 2007.
- [139] Bloch D. et al. Update of the jlip b-tagger performance in p14/pass2 with jes 5.3. *DØ Note 4824*, 2005.
- [140] Demina R. et al. Measurement of b-tagging efficiency and mis-tagging rates with csip method. *DØ Note 4432*, 2004.
- [141] Boline D. et al. Update on b-quark jet identification with secondary vertex reconstruction using p14-pass2 samples. *DØ Note 4796*, 2005.
- [142] Description de la construction d'un réseau de neurones : Mlpfit, a tool for multi-layer perceptrons,
<http://schwind.web.cern.ch/schwind/talks/itseminar/>.
- [143] Implémentation dans root de la méthode tmultilayerperceptron,
<http://root.cern.ch/root/html/tmultilayerperceptron.html#tmultilayerperceptron> :description.
- [144] S. Abachi et al. Search for squarks and gluinos in $p\bar{p}$ collisions at $\sqrt{s} = 1.8$ tev. *Phys. Rev. Lett.*, 75 :618–623, 1995.
- [145] B. Abbott et al. Search for squarks and gluinos in events containing jets and a large imbalance in transverse energy. *Phys. Rev. Lett.*, 83 :4937–4942, 1999.
- [146] C. Rott. *Search for scalar bottom quarks from gluino decays in proton-antiproton collisions at a center-of-mass energy of 1.96 TeV*. PhD thesis, The faculty of Purdue University, 2004.
- [147] W. Beenakker, R. Hopker, M. Spira, and P. M. Zerwas. Squark and gluino production at hadron colliders. *Nucl. Phys.*, B492 :51–103, 1997.
- [148] Torbjorn Sjostrand, Leif Lonnblad, Stephen Mrenna, and Peter Skands. Pythia 6.3 : Physics and manual. 2003.
- [149] M. Muhlleitner, A. Djouadi, and Y. Mambrini. Sdecay : A fortran code for the decays of the supersymmetric particles in the mssm. *Comput. Phys. Commun.*, 168 :46–70, 2005.

- [150] P. Skands et al. Susy les houches accord : Interfacing susy spectrum calculators, decay packages, and event generators. *JHEP*, 07 :036, 2004.
- [151] W. Beenakker, R. Hopker, and M. Spira. Prospino : A program for the production of supersymmetric particles in next-to-leading order qcd. 1996.
- [152] H. L. Lai et al. Global qcd analysis and the ctq parton distributions. *Phys. Rev.*, D51 : 4763–4782, 1995.
- [153] Ahn S.H. et al. D0 luminosity in run 2 : Recorded. *DØ Note 3972*, 2002.
- [154] Mangano M.L. et al. Alpgen, a generator for hard multiparton processes in hadronic collisions. *JHEP*, 2003.
- [155] E. Boos et al. Comphep 4.4 : Automatic computations from lagrangians to events. *Nucl. Instrum. Meth.*, A534 :250–259, 2004.
- [156] Gillberg D. Heavy flavour removal and determination of weighting factors for alpgen w+jet monte carlo. *DØ Note 5129*, 2006.
- [157] Muanza S. A compilation of mcfm v5.1 cross sections. *DØ Note 5300*, 2006.
- [158] A.F. Barfuss. Compte-rendu des journées jeunes chercheurs à la rochelle, 2006.
- [159] Park S. et al. Efficiency of the data quality calorimeter flags. *DØ Note 5324*, 2007.
- [160] Haas A. Level 1 and level 2 calorimeter trigger performance in multi-jet events. *DØ Note 4073*, 2003.
- [161] Petteni M. et al. Search for $zh \rightarrow \nu\nu\bar{b}\bar{b}$ at dØ. *DØ Note 5352*, 2007.
- [162] A. Zabi et al. A trigger for jets and missing e_t . *DØ Note 4315*, 2004.
- [163] Alexander L. Read. Modified frequentist analysis of search results (the cl(s) method). Prepared for Workshop on Confidence Limits, Geneva, Switzerland, 17-18 Jan 2000.

NOM : MILLET (avec précision du nom de jeune fille, le cas échéant) Prénoms : Thomas		DATE de SOUTENANCE 11 mai 2007
TITRE : Recherche de gluinos dans la topologie à jets de quarks b et énergie transverse manquante avec le détecteur D0 au TeVatron		
Numéro d'ordre : 62-2007		
DIPLOME DE DOCT.	DOCTEUR- INGENIEUR	DOCTORAT D'ETAT
DOCTORAT DE 3e CYCLE	Spécialité : Physique des Particules	
X	☐	☐
☐	☐	
Cote B.I.U. – Lyon : T 50/210/19 / et bis		CLASSE :
RESUME : Le modèle standard est le modèle théorique de référence pour expliquer les phénomènes subatomiques. Jamais directement mis en défaut et couronné de succès depuis des dizaines d'années, il souffre néanmoins d'imperfections théoriques et peine à expliquer certains phénomènes physiques observés à hautes énergies. Pour ces raisons, plusieurs théories proposent des extensions de ce modèle. La supersymétrie, appréciée pour son élégance théorique, est l'une d'entre elles. Le travail présenté dans ce manuscrit concerne la recherche d'un signal supersymétrique, dont les caractéristiques principales sont la présence de plusieurs jets de quark b et une forte énergie transverse manquante. L'analyse de ce signal est effectuée sur les données du Run II collectées par le détecteur D0 d'avril 2003 à mars 2006 (~1 fb⁻¹). La signature d'une telle topologie implique une bonne résolution sur l'énergie des jets, ainsi qu'une mesure précise de l'énergie transverse manquante. Par conséquent, il faut au préalable bien comprendre les informations recueillies par le calorimètre, une des parties cruciales du détecteur pour cette analyse. La présence de plusieurs quarks de saveurs lourdes dans l'état final nécessite un étiquetage maîtrisé des jets de quark b. Cet étiquetage est une variable fortement discriminante entre le signal et le bruit de fond du modèle standard. Les collisionneurs hadroniques comme le TeVatron ont une section efficace de QCD inclusive des ordres de grandeurs supérieurs à celle du signal recherché dans cette analyse. Il faut donc être en mesure de trouver des coupures rapides et efficaces dès le système de déclenchement du détecteur. L'essentiel de mon travail de préparation à l'analyse porte sur la mise au point et l'étude de conditions de déclenchement basées sur le calorimètre et permettant de sélectionner efficacement les signaux caractérisés par la présence de jets et d'énergie transverse manquante.		
MOTS-CLES : Supersymétrie – gluinos – système de déclenchement – étiquetage des quarks b – jets – énergie transverse manquante – D0 - TeVatron		
Laboratoire de recherche : Institut de Physique Nucléaire de Lyon		
Directeurs de recherches : Aldo Deandrea		
Président du jury : B. Ille Composition du jury : G.-S. Muanza – D. Froidevaux – G. Brooijmans – A. Duperrin – P. Verdier		253 pages