

Thèse de doctorat de l'université Pierre & Marie Curie

Spécialité : physique des particules

présentée par

Sébastien TROQUEREAU

Pour obtenir le titre de

DOCTEUR EN SCIENCES PHYSIQUES

**Performances des chambres à dérive
dans l'expérience HARP au CERN.**

Soutenue le 18 juin 2003 devant le jury composé de :

Jean-Eudes	AUGUSTIN	Président
Murat	BORATAV	Examineur
Jacques	DUMARCHEZ	Co-Directeur
Daniel	SILLOU	Rapporteur
Thierry	STOLARCZYK	Rapporteur
François	VANNUCCI	Co-Directeur

Table des matières

I	Les motivations	3
1	Le neutrino et ses oscillations	5
1.1	Préliminaires sur les neutrinos	5
1.1.1	La découverte des neutrinos	5
1.1.2	Les neutrinos dans le modèle standard	6
1.1.3	Le neutrino au-delà du modèle standard	8
1.1.4	Les oscillations de neutrinos	11
1.1.4.1	Le modèle d'oscillation entre 2 familles	11
1.1.4.2	Le modèle d'oscillation entre 3 familles	12
1.1.4.3	Les oscillations dans la matière et l'effet MSW	13
1.1.5	Les sources de neutrinos	16
1.1.5.1	Les sources naturelles	16
1.1.5.2	Les sources artificielles	20
1.2	Les recherches d'oscillations actuelles et futures	20
1.2.1	Les recherches sur les neutrinos solaires	21
1.2.2	Les recherches avec les neutrinos atmosphériques.	21
1.2.3	Recherche auprès de réacteurs.	22
1.2.4	Les neutrinos d'accélérateur	22
1.2.4.1	Préliminaires	22
1.2.4.2	Les faisceaux conventionnels	23
1.2.4.3	Les usines à neutrinos.	25
1.2.4.4	Apport des usines à neutrinos	27
1.2.4.5	En attendant les usines... les "Super-Beams"	28
1.2.5	Les neutrinos : peut-être un casse-tête théorique.	28
1.2.6	Conclusion	28
II	L'expérience HARP	31
2	L'expérience HARP	33
2.1	Une expérience au bout de la ligne T9 du CERN.	33
2.2	Présentation de l'expérience HARP	35
2.2.1	La maîtrise du faisceau incident	35
2.2.1.1	L'identification des particules du faisceau	36
2.2.1.2	Position du faisceau	37
2.2.2	Système de déclenchement	38
2.2.3	Les cibles	41

2.2.4	Un détecteur en 2 parties	42
2.2.4.1	Spectromètre arrière : La TPC (Chambre à Projection Temporelle)	43
2.2.4.2	Les RPC (Resistive Plate Chambers)	47
2.2.4.3	Le Spectromètre avant	48
2.2.4.4	Le Time-Of-Flight couramment appelé TOF	50
2.2.4.5	Le Tcherenkov	51
2.2.4.6	Les identificateurs d'électrons et de muons	51
2.2.5	Programme de mesure	53
3	Les Chambres à dérive	59
3.1	Description et principe d'une cellule de dérive	61
3.2	Les panneaux	65
3.3	La position des fils	66
3.4	Conclusion	66
III	Développement et environnement logiciel	67
4	Développement logiciel... vers l'analyse physique	69
4.1	L'architecture informatique de la collaboration	69
4.1.1	GAUDI	70
4.1.2	Objectivity	72
4.2	Simulation	72
4.2.1	GEANT 4	73
4.2.2	La géométrie des chambres à dérive.	74
4.3	La simulation dans les chambres à dérive	75
4.3.1	La digitisation des chambres à dérive.	76
4.3.2	Les paramètres optionnels.	80
4.4	Reconstruction des traces dans les chambres à dérive.	81
4.4.1	Avant-propos	81
4.4.2	Le temps dans les chambres à dérive.	81
4.4.3	L'algorithme de reconstruction	83
4.4.4	Performances de l'algorithme de reconstruction.	90
4.4.4.1	Conversion des temps en distances.	91
4.4.4.2	Recherche des segments.	91
4.4.4.3	Efficacité de reconstruction	91
IV	Performances des chambres à dérive	97
5	Alignement	99
5.1	Pourquoi effectuer un alignement ?	99
5.2	Quelques notions	99
5.3	Procédure d'alignement	100
5.3.1	Position du problème.	100
5.3.2	Alignement interne	103

5.3.2.1	Effet des différents paramètres sur les distributions des résidus.	103
5.3.2.2	Comment corrige-t-on concrètement tous ces effets?	104
5.3.2.3	Corrections sur la position des fils.	105
5.3.2.4	Une relation temps-distance affinée pour chaque plan. . . .	107
5.3.2.5	Les résultats	109
5.3.2.6	Utilisation des fichiers de calibration.	109
5.3.3	Alignement inter-modulaire	111
5.3.3.1	Extrapolation de segments pour l'ajustement inter-modulaire	111
5.3.4	Quelques remarques de conclusion.	112
6	Efficacité des chambres.	115
6.1	Etude de l'efficacité avec les rayons cosmiques	115
6.1.1	Principe de l'étude.	115
6.1.2	Influence de l'efficacité des chambres sur l'efficacité de reconstruction.	116
6.1.3	Efficacité de reconstruction à partir des données.	118
6.1.4	Conclusion.	118
7	Conclusion	121
V	Annexes	123
A	Corrections pour l'alignement inter-modulaire	125
A.1	Position du problème	125
A.2	Calcul des paramètres	126
B	Petit historique sur la production hadronique ($p+A \rightarrow \pi, K, \dots$)	129
B.1	Activités des années 60-70	129
B.2	NA20 et SPY/NA56	130
B.3	E-802	131
B.4	E-910	131
C	Les membres de la collaboration HARP	135
D	Nomenclature pour décrire la géométrie	137

Introduction

Depuis la découverte expérimentale des neutrinos dans les années 50, la communauté scientifique ne cesse de vouloir étudier les propriétés de la particule la plus mystérieuse du modèle standard. Du fait de leurs très faibles sections efficaces d'interaction, l'étude des propriétés intrinsèques des neutrinos, repose sur une très bonne connaissance de leurs sources naturelles ou artificielles.

A la question primordiale actuelle, “Le neutrino, a-t-il une masse?”, deux indications expérimentales récentes apportent une première réponse : l'observation d'oscillations de neutrinos, faites sur les neutrinos solaires et sur les neutrinos atmosphériques, implique des masses non nulles.

Ces indications, qui reposent en particulier sur des calculs de flux, demandent une confirmation et une étude exhaustive pour des expériences sur accélérateur, où les conditions initiales sont mieux maîtrisées. Dans cette perspective, des projets existent, à moyen et long terme, basés sur des superfaisceaux et des usines à neutrinos. Leur sensibilité repose sur une bonne connaissance de la production des hadrons, sources de neutrinos (ou de muons pour les usines à neutrinos). L'incertitude sur ces sections efficaces de production (de l'ordre de 15%) provient d'expériences conduites dans les années 70 qui étaient limitées par la technique de l'époque.

L'expérience HARP au CERN s'est donc attachée à reprendre ces mesures avec l'objectif de limiter les incertitudes à 2%. Ceci est possible en accumulant assez de statistiques dans un détecteur couvrant 4π d'angle solide et doté d'identification de particules. On verra au cours de cette thèse la manière avec laquelle HARP prétend atteindre ces résultats en ce qui concerne les mesures de sections efficaces de production des hadrons issus de l'interaction des protons avec des noyaux cibles. Plus largement l'objectif de HARP est de faire ces mesures sur de nombreuses cibles, permettant une étude systématique de leur dépendance avec le numéro atomique, et une optimisation des choix de cible en amont des futures usines à neutrinos. HARP veut aussi mesurer la production hadronique dans des cibles utilisées par des expériences actuelles (K2K et MiniBooNE) qui visent à confirmer les indications d'oscillations. Enfin l'étude va s'étendre à des cibles cryogéniques d'azote et d'oxygène, ce qui permettra les calculs de flux de neutrinos atmosphériques (au moins dans la partie de basse énergie).

Cette thèse va d'abord commencer par présenter l'actualité de la recherche dans la physique du neutrino. Puis elle présentera l'expérience HARP et plus particulièrement les chambres à dérive sur lesquelles j'ai travaillé. Je parlerai ensuite des programmes de simulation et de reconstruction auxquels j'ai contribué. Finalement je consacrerai la dernière

partie aux performances des chambres à dérive obtenues avec un traitement des premières données prises lors des années 2001-2002. L'analyse des données proprement dites et l'extraction des sections efficaces ne sera possible qu'avec un programme de reconstruction complet, pas encore disponible au moment de la rédaction de cette thèse.

Première partie

Les motivations

Chapitre 1

Le neutrino et ses oscillations

Il s'agit dans ce chapitre de parler du neutrino et de faire une revue de ses caractéristiques en vue de cerner les interrogations de la physique du neutrino et comment on cherche à trouver des solutions aux problèmes posés.

Les oscillations de neutrinos représentent actuellement un des sujets prédominants de la physique des particules. De nombreuses recherches se déroulent en vue d'étudier ces particules mystérieuses que le modèle standard dit "sans masse". En effet la mise en évidence d'oscillation de neutrinos serait la preuve que ceux-ci ont une masse non nulle. Ce résultat permettrait d'aller au delà du modèle standard que les physiciens cherchent à tester.

1.1 Préliminaires sur les neutrinos

1.1.1 La découverte des neutrinos

Les neutrinos sont des particules dont beaucoup de secrets restent à percer. L'existence du neutrino a été tout d'abord postulée par W. Pauli en 1930, afin d'expliquer l'énergie manquante lors de la désintégration β des noyaux radioactifs. En même temps, le neutrino expliquait les variations de spin dans le processus. Selon Pauli, une particule neutre, de faible masse et de spin $1/2$ qui interagit faiblement avec la matière s'échapperait. Elle permettait de garder le principe de conservation d'énergie. C'est en 1933 que Fermi décide de l'appeler le "neutrino".

En 1956, la première détection d'un neutrino fut établie par Reines et Cowan. Cette recherche s'est faite auprès du réacteur nucléaire de Savannah River, source d'anti-neutrino électronique. La détection se faisait par l'utilisation d'un processus β inverse ou une interaction par courant chargé :

$$\bar{\nu}_e + p \rightarrow e^+ + n \quad (1.1)$$

Le résultat fut en accord avec les prédictions théoriques. L'hélicité du neutrino fut mesurée en 1958 par Goldhaber.

Les théoriciens avaient prédit en 1946 l'existence d'autres types de neutrinos (Inoué et Sakata, 1946) . Les neutrinos muoniques ont été découverts en 1962 à Brookhaven. Ceci

laissa penser qu’une association lepton chargé-neutrino existait. Le dernier neutrino fut en fait une conséquence directe de la découverte en 1975 à Stanford du troisième et dernier lepton : le Tau (τ). Son neutrino associé (ν_τ) fut observé très tardivement (DONUT 2001).

D’après la largeur du spectre de masse du Z^0 (figure 1.1, Collaboration (ALEPH, OPAL, DELPHI, L3) 1989 CERN , et par le détecteur Mark II à SLAC), il existe trois sortes différentes de neutrinos (six, si l’on considère leurs anti-particules). Cependant, comme on le verra plus tard, des physiciens émettent l’idée de l’existence d’une quatrième saveur de neutrino dit “neutrino stérile”, celui-ci n’interagissant pas avec la matière (peut être un neutrino d’hélicité droite) et n’étant pas visible dans la désintégration du Z^0 .

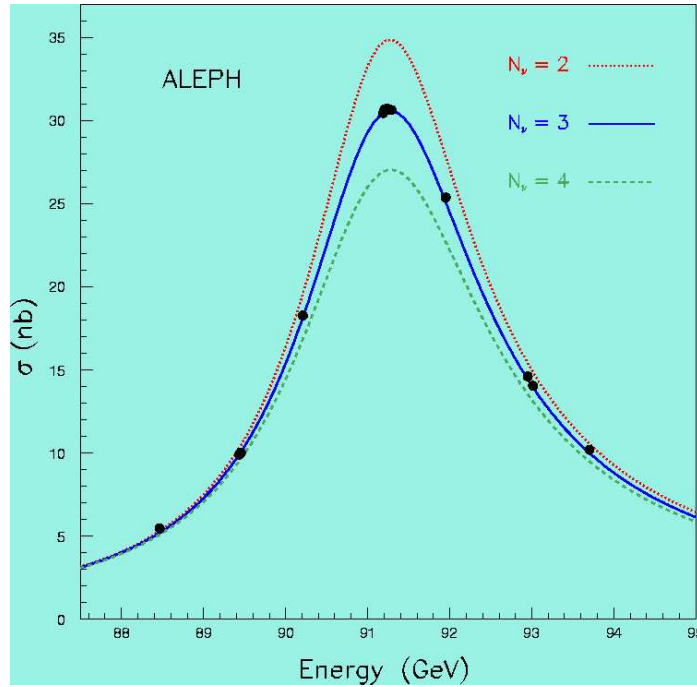


FIG. 1.1 – Spectre de masse du Z^0

1.1.2 Les neutrinos dans le modèle standard

Les neutrinos sont des leptons de charge nulle, de masse nulle et de spin $1/2$. Ils n’interagissent avec la matière que par l’intermédiaire de l’interaction faible. Il existe 3 saveurs différentes de neutrinos : les neutrinos électronique (ν_e), muonique (ν_μ) et tauique (ν_τ). Ces neutrinos sont associés aux leptons chargés formant des doublets.

$$\nu_e \longleftrightarrow e^- \quad \nu_\mu \longleftrightarrow \mu^- \quad \nu_\tau \longleftrightarrow \tau^- \quad (1.2)$$

Chaque doublet (e, ν_e ; μ, ν_μ ; τ, ν_τ) est associé à un nombre leptonique (L_e, L_μ, L_τ) qui se conserve lors d’une interaction (on prend $+1$ pour les leptons et -1 pour les anti-leptons). Néanmoins avoir trois nombres leptoniques traduit également qu’il n’y a pas de mélange considéré entre les différentes saveurs.

Mais d’une manière plus formalisée, les fermions de spin $1/2$ satisfont le Lagrangien de

Dirac défini par :

$$L_D = \bar{\Psi}(i\gamma^\mu\partial_\mu - m)\Psi \quad (1.3)$$

$$(1.4)$$

convention : μ allant de 0 à 3, et i allant de 1 à 3

Les fermions étant alors sous la forme d'un champ de matière Ψ qui est en fait un spineur à 4 composantes représentant sa particule et son anti-particule mais aussi leur projection de spin possible, $\pm 1/2$.

Sachant que :

$$\bar{\Psi} = \Psi^\dagger \gamma^0 \quad (1.5)$$

et les matrices de Dirac γ^μ , dans la représentation de Weyl (ou représentation chirale) :

$$\gamma^0 = \begin{pmatrix} 0 & I \\ I & 0 \end{pmatrix} \quad \gamma^i = \begin{pmatrix} 0 & \sigma^i \\ -\sigma^i & 0 \end{pmatrix} \quad \text{et} \quad \gamma^5 = i\gamma^0\gamma^1\gamma^2\gamma^3 = \begin{pmatrix} -I & 0 \\ 0 & I \end{pmatrix}$$

Avec les matrices de Pauli σ (portant le spin) :

$$\sigma^1 = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} \quad \sigma^2 = \begin{pmatrix} 0 & -i \\ i & 0 \end{pmatrix} \quad \sigma^3 = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}$$

En utilisant les projecteurs P_L et P_R définis par :

$$P_L = \frac{1 - \gamma^5}{2} \quad P_R = \frac{1 + \gamma^5}{2} \quad (1.6)$$

On obtient alors les composantes chirales de Ψ :

$$\Psi_R = P_R \cdot \Psi = \begin{pmatrix} 0 & 0 \\ 0 & I \end{pmatrix} \Psi \quad (1.7)$$

$$\Psi_L = P_L \cdot \Psi = \begin{pmatrix} I & 0 \\ 0 & 0 \end{pmatrix} \Psi \quad (1.8)$$

Et donc Ψ_R (Ψ_L) renferme les deux premières (dernières) composantes de Ψ .

L'utilisation de ces 2 composantes dans le Lagrangien de Dirac nous donne 2 termes :

– Un terme cinétique

$$L_c = i\bar{\Psi}\gamma^\mu\partial_\mu\Psi = i\bar{\Psi}_L\gamma^\mu\partial_\mu\Psi_L + \bar{\Psi}_R\gamma^\mu\partial_\mu\Psi_R \quad (1.9)$$

– un terme de masse

$$L_m = -m\bar{\Psi}\Psi = -m(\bar{\Psi}_R\Psi_L + \bar{\Psi}_L\Psi_R) \quad (1.10)$$

On a alors une situation très intéressante avec un terme cinétique L_c qui ne mélange pas les états droit et gauche, puis un terme de masse qui, au contraire nous donne un mélange de ces états. Or le modèle standard en nous imposant la masse d'un neutrino nulle permet d'éliminer en même temps ce mélange possible droite-gauche. Cela justifie la non-considération des oscillations de neutrinos dans le modèle standard. Réciproquement,

supposer une masse aux neutrinos signifie permettre de tels mélanges.

Donc d'après le modèle standard, les oscillations de neutrinos n'existent pas. Mais est-ce par simple commodité ou est-ce la réalité ? Nous verrons que des signes d'oscillation existent mais il reste à les confirmer.

1.1.3 Le neutrino au-delà du modèle standard

Le modèle standard suppose le neutrino comme étant une particule de masse nulle et d'hélicité gauche (droite pour l'antineutrino). De plus on a vu précédemment que le modèle standard ne permet pas les oscillations de neutrinos en éliminant le terme de masse dans le Lagrangien de Dirac. Cependant ce n'est pas une preuve. La masse du neutrino pourrait être non nulle et nous allons voir les conséquences que cela amène pour certains modèles.

Les neutrinos de Majorana et de Dirac

Dans ce passage, je vais expliquer deux modèles théoriques. Dans le monde des neutrinos, deux modèles restent incontournables : le neutrino de Majorana et le neutrino de Dirac. Dans le cas d'un neutrino gauche, par invariance de CPT, on obtient un antineutrino droit. Jusque là, tout reste conforme. Cependant si on considère que le neutrino est massif, celui ne peut plus aller à la vitesse de la lumière et donc on peut avoir un observateur qui peut aller plus vite que le neutrino et qui pourra alors le voir d'hélicité droite. Ainsi, un boost de Lorentz permet d'obtenir un neutrino "droit", ν_R .

Neutrino de Dirac

Si on considère que le ν_R n'est pas l'antineutrino $\overline{\nu_R}$, le neutrino ν_R a donc sa propre image par CPT, $\overline{\nu_L}$ et les deux antineutrinos sont reliés par le boost de Lorentz. On a donc quatre états de même masse non nulle. Ceci est l'approche pour expliquer ce que l'on appelle le neutrino de Dirac (figure 1.2).

Neutrino de Majorana

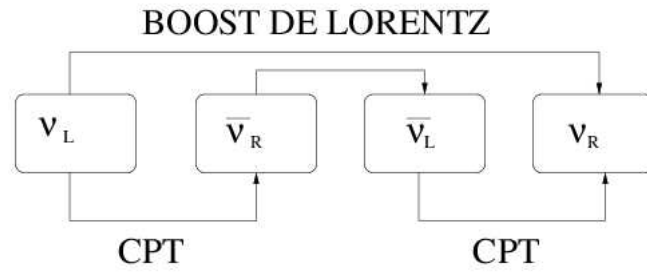
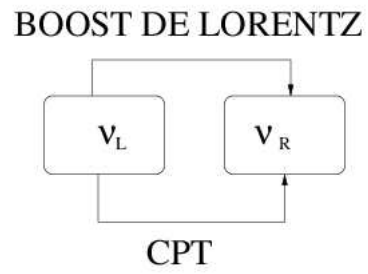
Si maintenant, on considère ν_R comme la même particule que l'antineutrino $\overline{\nu_R}$, on a donc les particules qui sont leur propres antiparticules et donc le boost de Lorentz a le même effet que la transformation par CPT. C'est le principe du neutrino de Majorana (figure 1.3).

On peut effectuer le petit exercice suivant :

On prend un ν_L pour lui appliquer CPT, on obtient neutrino droit en effectuant un boost de Lorentz à ce neutrino droit, on revient à l'état initial du neutrino gauche.

Des termes de masse dans le Lagrangien de Dirac.

Il est bon de rappeler que le terme de masse est un scalaire qui se doit d'être un invariant de Lorentz. Il est possible d'écrire d'autres termes de masse en couplant le champ Ψ avec son C-conjugué (changement de charge leptonique).

FIG. 1.2 – *Les quatre neutrinos de Dirac*FIG. 1.3 – *Les deux neutrinos de Majorana*

$$-m\bar{\Psi}\Psi^C - h.c \quad (1.11)$$

Ces 2 termes brisent la conservation de la charge (leptonique pour le neutrino) mais conservent la chiralité. Or le modèle standard ne considère pas la violation de la conservation de la charge sous une transformation U(1).

Si on utilise le C-conjugué pour écrire le Lagrangien de Dirac, on peut rajouter un terme au Lagrangien cinétique :

$$+i\bar{\Psi}^C\gamma^\mu\partial_\mu\Psi^C \quad (1.12)$$

Et le terme de masse devient pour sa part :

$$L_m = -\frac{1}{2}(\bar{\Psi}_L^C\bar{\Psi}_R)M\begin{pmatrix}\Psi_L \\ \Psi_R^C\end{pmatrix} - \frac{1}{2}(\bar{\Psi}_L\bar{\Psi}_R^C)M\begin{pmatrix}\Psi_L^C \\ \Psi_R\end{pmatrix} \quad (1.13)$$

avec la matrice de masse M :

$$M = \begin{pmatrix} m_L & m_D \\ m_D & m_R \end{pmatrix} \quad (1.14)$$

Soit si on développe :

$$L_m = -m_D\bar{\Psi}_L\Psi_R - \frac{1}{2}(m_L\bar{\Psi}_L^c\Psi_L + m_R\bar{\Psi}_R^c\Psi_R) + h.c \quad (1.15)$$

Le terme de masse m_D est le terme de Dirac qui est en fait le terme déjà présent dans le modèle standard car il conserve U(1) mais pas la chiralité. Les termes nouveaux par rapport au modèle standard sont les termes de Majorana droit (m_R) et gauche (m_L) qui ne conservent pas U(1) mais conservent la chiralité.

L'état gauche du neutrino est donné par :

$$\nu = \begin{pmatrix} \nu_R \\ \nu_R^c \end{pmatrix} \quad (1.16)$$

La matrice M n'étant pas diagonale, on peut alors calculer les valeurs propres de masse :

$$m_1 = \frac{m_L + m_R}{2} - \frac{1}{2}\sqrt{4m_D^2 + (m_R - m_L)^2} \quad (1.17)$$

$$m_2 = \frac{m_L + m_R}{2} + \frac{1}{2}\sqrt{4m_D^2 + (m_R - m_L)^2} \quad (1.18)$$

Les états de masse deviennent :

$$\Psi_m = \begin{pmatrix} \Psi_1 \\ \Psi_2 \end{pmatrix} = \begin{pmatrix} \cos\theta.\nu_L - \sin\theta.\nu_R^c \\ \sin\theta.\nu_L + \cos\theta.\nu_R^c \end{pmatrix} \quad (1.19)$$

avec un angle de mélange θ qui est défini par $\tan 2\theta = \frac{2m_D}{(m_R - m_L)}$.

Le mécanisme de la bascule ("seesaw mechanism")

Nous allons nous intéresser à cette matrice de masse auquel on a rajouté ces termes de Majorana (extension du modèle standard).

m_D représente les masses classiques des particules élémentaires dans le modèle standard tandis que pour les masses de Majorana, on peut faire les hypothèses suivantes : on prend une masse très faible pour l'une et une masse très importante pour l'autre d'où l'effet de la bascule.

$$m_R \gg m_D \gg m_L \quad (1.20)$$

si on néglige m_L , la matrice devient :

$$M = \begin{pmatrix} 0 & m_D \\ m_D & m_R \end{pmatrix} \quad (1.21)$$

dont les deux valeurs propres sont :

$$m_{lourd} \simeq m_R \quad (1.22)$$

$$m_{leger} \simeq \frac{m_D^2}{m_R} \quad (1.23)$$

m_D représente une masse de Dirac qu'on suppose de l'ordre de grandeur des masses des autres composants de la même génération. On a donc deux masses possibles pour le neutrino, une masse légère qui serait celle du neutrino gauche (état propre de l'interaction faible) et une masse lourde pour le neutrino droit non observé, et cela pour chaque famille de neutrino.

Pour trois familles de neutrino, on aurait donc trois doublets de valeurs propres.

$$m_{leger}^e, m_{leger}^\mu, m_{leger}^\tau, m_{lourd}^e, m_{lourd}^\mu, m_{lourd}^\tau \quad (1.24)$$

Il y aurait donc deux valeurs propres de masses pour chaque saveur et donc beaucoup de choses à découvrir. Ce modèle de balançoire est actuellement le modèle favori des théoriciens pour expliquer la petitesse des masses de neutrinos devant les autres masses des composants connus.

1.1.4 Les oscillations de neutrinos

1.1.4.1 Le modèle d'oscillation entre 2 familles

Nous allons voir le modèle simplifié d'oscillation dans le vide en ne considérant que deux familles de neutrinos, pour fixer les idées on ne considèrera dans ce paragraphe que ν_e et ν_μ .

Les Hamiltoniens H_V dans le vide sont diagonaux dans la base des états de masse et non de saveur.

$$H_V |\nu_1\rangle = E_1 |\nu_1\rangle \quad (1.25)$$

$$H_V |\nu_2\rangle = E_2 |\nu_2\rangle \quad (1.26)$$

Cependant les états propres d'interactions faibles (effectifs lors de la production ou de l'interaction des neutrinos) ne sont pas les états de masse mais les états de saveur. On a donc des états de saveur (états qui représentent nos observables) qui sont le résultat de combinaisons linéaires des états de masse.

$$\begin{pmatrix} \nu_e \\ \nu_\mu \end{pmatrix} = \begin{pmatrix} \cos \theta & \sin \theta \\ -\sin \theta & \cos \theta \end{pmatrix} \begin{pmatrix} \nu_1 \\ \nu_2 \end{pmatrix} \quad (1.27)$$

où θ est appelé l'angle de mélange¹.

Prenons un neutrino d'impulsion p à un temps t_0 dans un état initial ν_e (comme pour les neutrinos solaires) :

$$|\nu(t=0)\rangle = |\nu_e\rangle \quad (1.28)$$

En utilisant les opérateurs d'évolution dans le temps qui s'appliquent aux états ν_1 et ν_2 , on obtient pour un temps t ($> t_0$) :

$$|\nu(t)\rangle = \cos \theta e^{iE_1 t} |\nu_1\rangle + \sin \theta e^{iE_2 t} |\nu_2\rangle \quad (1.29)$$

sachant que $E = \sqrt{p^2 + m^2} \approx p \left(1 + \frac{m^2}{2p^2}\right)$ car $m \approx 0$.

La probabilité que ce neutrino qui était initialement un ν_e devienne un ν_μ est donnée par :

$$P(\nu_e \rightarrow \nu_\mu) = |\langle \nu_\mu | \nu(t) \rangle|^2 = 4 \sin^2 \theta \cos^2 \theta \sin^2 \frac{(E_1 - E_2)t}{2} \quad (1.30)$$

En considérant que $E_1 - E_2 \simeq \frac{m_1^2 - m_2^2}{2p}$ et en posant $\Delta m^2 = m_1^2 - m_2^2$, la probabilité devient :

$$P(\nu_e \rightarrow \nu_\mu)(t) = \sin^2 2\theta \sin^2 \left(\frac{\Delta m^2}{4p} t \right) \quad (1.31)$$

On voit immédiatement que la période temporelle de cette probabilité est $T = \frac{4\pi p}{\Delta m^2}$. On écrit souvent cette probabilité en terme de distance plutôt que de temps. Il est alors aisé de remplacer le contenu du dernier sinus carré par une composante en distance :

$$P(\nu_e \rightarrow \nu_\mu)(L) = \sin^2 2\theta \sin^2 \left(1,27 \frac{\Delta m^2 x}{E} \right) \quad (1.32)$$

où Δm^2 est en eV^2 , E en GeV, x en km.

1.1.4.2 Le modèle d'oscillation entre 3 familles

Nous passons maintenant au cas où l'on considère les trois neutrinos (ν_e , ν_μ , ν_τ) détectés jusqu'à présent. Dans ce cas, on doit avoir une matrice de mélange 3×3 .

$$\begin{pmatrix} \nu_e \\ \nu_\mu \\ \nu_\tau \end{pmatrix} = U \begin{pmatrix} \nu_1 \\ \nu_2 \\ \nu_3 \end{pmatrix} \quad (1.33)$$

La matrice de mélange est de la même forme que la matrice CKM (Cabibbo-Kobayashi-Maskawa) pour le mélange des quarks dans l'interaction faible. La matrice appliquée au neutrino est appelée la matrice MNS (Maki-Nakagawa-Sakata).

¹Ce n'est pas le même θ que dans le mécanisme de la balançoire.

La matrice unitaire MNS comporte trois angles de mélanges θ_{ij} une phase δ , et des termes $c_{ij} = \cos \theta_{ij}$, $s_{ij} = \sin \theta_{ij}$. Elle s'écrit :

$$U = \begin{pmatrix} U_{e1} & U_{e2} & U_{e3} \\ U_{\mu1} & U_{\mu2} & U_{\mu3} \\ U_{\tau1} & U_{\tau2} & U_{\tau3} \end{pmatrix} = \begin{pmatrix} c_{12}c_{13} & s_{12}c_{13} & s_{13}e^{-i\delta} \\ -s_{12}c_{23} - c_{12}s_{23}s_{13}e^{i\delta} & c_{12}c_{23} - s_{12}s_{23}s_{13}e^{i\delta} & s_{23}c_{13} \\ s_{12}s_{23} - c_{12}c_{23}s_{13}e^{i\delta} & -c_{12}s_{23} - s_{12}c_{23}s_{13}e^{i\delta} & c_{23}c_{13} \end{pmatrix} \quad (1.34)$$

L'évolution temporelle d'un état de saveur est d'une manière plus générale :

$$|\nu_\alpha > (t) = \sum_{i=1}^3 U_{\alpha i} e^{-iE_i t} |\nu_i > \quad (1.35)$$

et donc la probabilité de changement de saveur d'un ν_α vers un ν_β ($\alpha, \beta = e, \mu, \tau$) devient :

$$P_{\nu_\alpha \rightarrow \nu_\beta} = \sum_{i=1}^3 |U_{\alpha i}|^2 |U_{\beta i}|^2 \quad (1.36)$$

$$+ \sum_{j \neq i} \Re(U_{\alpha i} U_{\alpha j}^* U_{\beta j} U_{\beta i}^*) \cos\left(\frac{m_i^2 - m_j^2}{2p_\nu} x\right) \\ + \sum_{j \neq i} \Im(U_{\alpha i} U_{\alpha j}^* U_{\beta j} U_{\beta i}^*) \sin\left(\frac{m_i^2 - m_j^2}{2p_\nu} x\right) \quad (1.37)$$

On insère souvent dans les expressions en cosinus et sinus les longueurs d'oscillations :

$$L_{ij} = 2\pi \frac{2p_\nu}{|m_i^2 - m_j^2|} = 2\pi \frac{2p_\nu}{\delta m_{ij}^2} \quad (1.38)$$

L'expression de $\left(\frac{m_i^2 - m_j^2}{2p_\nu} x\right)$ dans la probabilité d'oscillation devient alors $2\pi \frac{x}{L_{ij}}$.

Il faut souligner que les éléments de matrice de U font l'objet d'une recherche intense, et l'objectif de toutes les expériences neutrinos en cours ou en préparation cherche à mesurer ces paramètres ou en cas d'échec à y mettre des limites. Chaque expérience est établie pour explorer un domaine de sensibilité bien précis dans le diagramme de type $\Delta m^2 - \sin^2 2\theta$.

1.1.4.3 Les oscillations dans la matière et l'effet MSW

On va voir dans ce paragraphe que l'influence de la matière sur le neutrino peut avoir des effets très intéressants. Ces effets sont différents pour les ν_e et les autres neutrinos. Les ν_e peuvent interagir par courant chargé et courant neutre alors que les autres neutrinos ne le peuvent que par courant neutre (voir figure 1.4). Cette dissymétrie fait émerger un terme supplémentaire dans l'hamiltonien pour le neutrino électronique que l'on décrit comme étant une énergie potentielle (V). Cependant l'hamiltonien H_V dans le vide est

diagonal dans la base des états de masse alors que cette énergie potentielle l'est dans la base des saveurs.

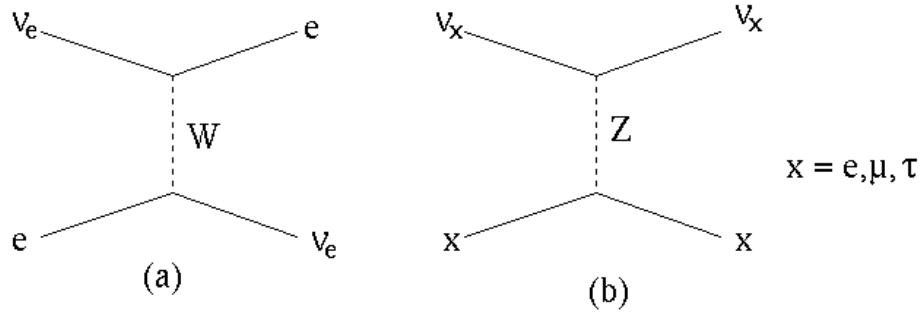


FIG. 1.4 – (a) interaction des neutrinos électroniques par courant chargé. (b) interaction des neutrinos e, μ, τ par courant neutre.

En ne considérant que les ν_e et les ν_μ , on a donc :

$$V|\nu_e\rangle = (C + \sqrt{2}G\rho_e)|\nu_e\rangle \quad (1.39)$$

$$V|\nu_\mu\rangle = C|\nu_\mu\rangle \quad (1.40)$$

où G est la constante de Fermi, ρ_e est la densité d'électrons. On appellera dans la suite θ_V l'angle de mélange dans le vide.

L'hamiltonien dans la matière $H_m (= H_V + V)$ n'étant diagonal ni dans les états de masses, ni dans les états de saveurs, il y aura donc de nouveaux états propres.

On peut établir l'hamiltonien dans le vide dans la base des saveurs :

$$H_V = \begin{pmatrix} H_{ee}^V & H_{e\mu}^V \\ H_{\mu e}^V & H_{\mu\mu}^V \end{pmatrix} \quad (1.41)$$

$$= \begin{pmatrix} E_1 \cos^2 \theta_V + E_2 \sin^2 \theta_V & -(E_2 - E_1) \sin \theta_V \cos \theta_V \\ -(E_2 - E_1) \sin \theta_V \cos \theta_V & E_2 \cos^2 \theta_V + E_1 \sin^2 \theta_V \end{pmatrix} \quad (1.42)$$

et donc H_m devient :

$$H_m = \begin{pmatrix} E_1 \cos^2 \theta_V + E_2 \sin^2 \theta_V + C + \sqrt{2}G\rho_e & -(E_2 - E_1) \sin \theta_V \cos \theta_V \\ -(E_2 - E_1) \sin \theta_V \cos \theta_V & E_2 \cos^2 \theta_V + E_1 \sin^2 \theta_V + C \end{pmatrix} \quad (1.43)$$

Lorsque l'on diagonalise H_m , nous avons deux valeurs propres λ_+ et λ_- qui sont les énergies propres. En posant pour simplifier :

$$H_m = \begin{pmatrix} a & c \\ c & b \end{pmatrix}$$

On obtient alors les valeurs propres $\lambda_\pm$ et l'angle de mélange θ_m dans la matière :

$$\lambda_{\pm} = \frac{1}{2}(a + b) \pm \frac{1}{2}\sqrt{((a - b)^2 + 4c^2)} \quad (1.44)$$

$$\tan 2\theta_m = \frac{2c}{a - b} \quad (1.45)$$

Il est important de souligner que l'on peut introduire une longueur d'oscillation dans la matière définie par :

$$L_m = \frac{2\pi}{|E_{+m} - E_{-m}|} \quad (1.46)$$

$$= \frac{L_V}{\sqrt{(\cos 2\theta_V - \sqrt{2}G\rho_e \times \frac{2p_\nu}{\Delta m^2})^2 + \sin^2 2\theta_V}} \quad (1.47)$$

et une amplitude d'oscillation dépendante de celle du vide :

$$\sin^2 2\theta_m = \frac{\sin^2 2\theta_V}{\left(\cos 2\theta_V - \frac{2p_\nu \sqrt{2}G\rho_e}{\Delta m^2}\right)^2 + \sin^2 2\theta_V} \quad (1.48)$$

L'influence de la matière pour les oscillations de neutrinos doit être prise en compte pour expliquer les résultats obtenus avec les neutrinos atmosphériques et surtout le déficit de neutrinos solaires mesuré. La matière a la faculté d'amplifier ces oscillations par un effet résonnant, ce qui est clair sur la formule (1.48) : c'est ce que l'on appelle l'effet MSW (Mikheyev-Smirnov-Wolfenstein)

L'effet MSW

L'effet MSW apparait lors du passage des neutrinos dans un milieu de densité variable. En effet on parle souvent de résonance et il existe une résonance différente suivant l'énergie du neutrino. Donc un milieu ayant une densité électronique variable permet d'avoir une résonance pour des neutrinos d'énergie différentes. Cette notion de résonance est importante et l'on va mieux voir ce qu'elle signifie vraiment dans l'explication suivante.

Nous allons voir ce qui se passe dans le cas de 2 familles (ν_e et ν_μ) de neutrinos lorsque l'on augmente la densité électronique en partant du vide.

On peut distinguer 3 zones distinctes sur la figure 1.5.

- Lorsque l'on part de ρ_e nulle, nous avons 2 états de propagation ν_+ et ν_- . L'angle de mélange est celui du vide et la formule 1.45 donne $\tan 2\theta_V = \frac{2c}{a-b}$. Ayant un angle de mélange du vide très faible, nous avons ν_- qui est proche de ν_e et ν_+ qui est proche de ν_μ
- Lorsque ρ_e augmente, l'angle de mélange a tendance à augmenter également pour atteindre finalement l'angle maximal de $\pm \frac{\pi}{4}$ (c'est-à-dire quand $a = b$) où les états de propagation $\nu_{\pm}$ ne privilégient plus un état d'interaction par rapport à un autre ν_e, ν_μ : c'est ce que l'on appelle la *résonance*.

- Puis lorsque l'on rencontre une densité d'électrons encore plus importante, l'angle de mélange θ_m diminue pour associer un état d'interaction avec un état de propagation. Cependant on s'aperçoit que la résonance a provoqué une inversion des états de propagation. En effet ν_+ se retrouve plus proche de ν_e et ν_- de ν_μ .

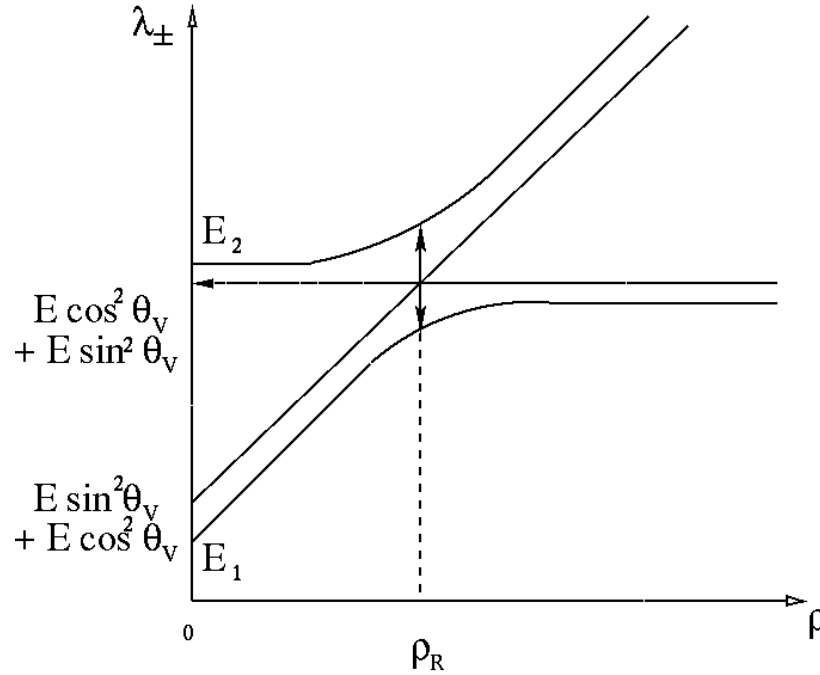


FIG. 1.5 – Evolution des valeurs propres $\lambda_{\pm}$ en fonction de la densité électronique ρ_e

Donc l'effet MSW permet d'expliquer l'influence de la matière sur les oscillations de neutrinos par résonance. Suivant le milieu qu'un neutrino traverse et son énergie, un neutrino peut changer de saveur plus ou moins facilement. Ceci est en particulier réalisé à l'intérieur du soleil où la densité varie de manière très importante entre le centre et la périphérie.

Une remarque importante à faire est que la hiérarchie des masses joue un rôle primordial en ce sens que cet effet ne peut exister que si le ν_μ est plus lourd que le ν_e car sinon on obtiendrait une densité ρ_e négative, et l'effet MSW ne pourrait alors se produire.

1.1.5 Les sources de neutrinos

Les sources de neutrinos naturelles sont abondantes mais l'homme sait aussi en fabriquer artificiellement afin d'avoir un contrôle plus important. Nous verrons que les sources naturelles détectables proviennent d'origines très différentes et qui ont l'avantage d'avoir des énergies diverses et donc permettent différents types d'étude potentiels.

1.1.5.1 Les sources naturelles

Les neutrinos "fossiles"

On oublie souvent de les évoquer mais il existe un fond cosmologique de neutrinos issus du Big-Bang. Néanmoins ceux-ci ont une si faible énergie que leur libre-parcours moyen est plus grand que la taille de l'univers. La température actuelle de ce fond est estimée à 1,9 K et quant à son origine, il provient du découplage des neutrinos qui s'est produit au moment où l'univers s'est suffisamment refroidi si bien qu'il n'y avait plus de réaction β inverse (absorption d'un neutrino par un proton), le neutrino n'étant plus assez énergétique. Il ne resta plus que la production de neutrino par désintégration du neutron. Le processus d'absorption étant quasi-impossible, seule l'émission restait en jeu d'où le découplage des neutrinos.

On estime que la densité pour chaque type de ces neutrinos fossiles est d'environ $110 \nu.cm^{-3}$ soit en considérant 3 familles $330 \nu.cm^{-3}$, ce qui est considérable. Mais ceux-ci étant très peu énergétiques, il est actuellement impossible de les détecter. Donc on peut les oublier pour le moment.

Le Soleil

Le Soleil est une source intense en neutrinos. Ceux-ci sont créés au coeur du soleil lors des réactions de fusions thermonucléaires. Le flux de neutrinos arrivant sur la Terre est de $\Phi \simeq 65 \times 10^9 cm^{-2} s^{-1}$ ce qui en fait la source la plus importante par son flux. Les neutrinos produits par le soleil sont à 100% des neutrinos électroniques (ν_e)

Le spectre du flux de neutrinos nous est donné sur la figure 1.6 suivant le modèle de Bahcall, Basu et Pinsonneault (1998).

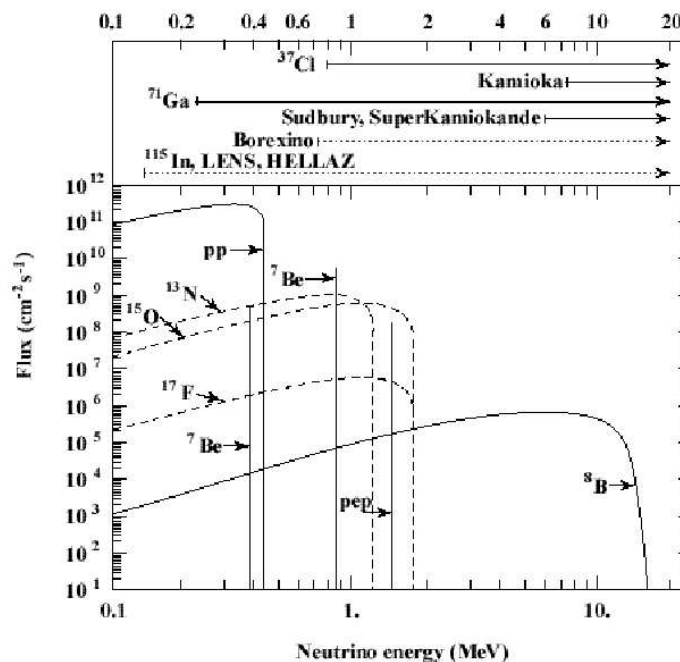


FIG. 1.6 – *Spectre des neutrinos solaires*

- Comme on peut le voir sur la figure les neutrinos dont le flux est le plus important sont les neutrinos dits “pp” (ils proviennent d’une fusion de 2 protons). Ce sont les moins énergétiques ayant un spectre continu jusqu’à 423 keV. Ils sont produits par

la réaction :

$$p + p \rightarrow e^+ + \nu_e + d$$

- Ensuite nous avons le deuxième flux le plus important : les neutrinos ${}^7\text{Be}$ qui ont la particularité d’avoir un spectre discret avec une raie à 861 keV. En fait une deuxième raie existe à 383 keV (10% des cas) mais elle est noyée dans le spectre ν_{pp} . Le ${}^7\text{Be}$ est obtenu à la suite de plusieurs réactions issues de la réaction précédente p-p.

$$\begin{aligned} p + d &\rightarrow \gamma + {}^3\text{He} \\ {}^3\text{He} + {}^3\text{He} &\rightarrow {}^4\text{He} + p + p \\ {}^3\text{He} + {}^4\text{He} &\rightarrow \gamma + {}^7\text{Be} \end{aligned}$$

et l’on obtient le neutrino ${}^7\text{Be}$ par la réaction suivante :

$$e^- + {}^7\text{Be} \rightarrow \nu_e + {}^7\text{Li}$$

La réaction produisant 2 corps dans l’état final explique l’aspect mono-énergétique des ν_{Be} .

- La troisième contribution spectrale est donnée par les neutrinos du Bore (${}^8\text{B}$). Ce sont les plus énergétiques avec un spectre continu allant jusqu’à 15 MeV. Ils proviennent d’une autre réaction avec le ${}^7\text{Be}$:

$$p + {}^7\text{Be} \rightarrow \gamma + {}^8\text{B}$$

puis

$${}^8\text{B} \rightarrow {}^8\text{Be} + e^+ + \nu_e$$

Ils ont certes un flux moins important que les classes précédentes avec toutefois une énergie très supérieure. C’est pourquoi il y a autant d’expériences pouvant les détecter alors qu’aux faibles énergies, les candidats se font rares.

Les supernovae

Les explosions de supernovae sont rares (du point de vue de la détection) mais lorsqu’elles se produisent, elles représentent une source ultra-intense de neutrinos. Ils sont produits au moment de l’effondrement du cœur de l’étoile où des électrons sont capturés par des protons formant ainsi des neutrons et des ν_e .

$$p + e^- \rightarrow n + \nu_e \quad (1.49)$$

Le flux de ces sources est très important mais bref. Le nombre de neutrinos produit lors d’une telle explosion est de l’ordre de $10^{58} \nu$. Mais il ne faut pas oublier que celles-ci se produisent loin et par conséquent, que le nombre de neutrinos arrivant jusqu’à nous est assez réduit. On se souvient de la supernova SN1987A qui se trouvait dans le nuage de Magellan à 50 kpc de la Terre. Le flux intégré a été calculé à partir des expériences IMB (Bratton et al., 1988) et Kamiokande (Hirata et al., 1988) pour finalement obtenir un flux sur terre d’environ $2 \times 10^{10} \text{cm}^{-2}$.

Les neutrinos atmosphériques

Ces neutrinos naturels sont les plus proches de notre Terre (mis à part la radioactivité naturelle de la Terre et des Hommes) car ceux-ci sont produits lors de l’interaction

de rayons cosmiques avec notre atmosphère. Les rayons cosmiques sont composés à 90% de protons, 9% de noyau d'hélium et pour le reste il s'agit de noyaux plus lourd (carbone, fer, ...). Les rayons cosmiques ont une origine qui reste inconnue mais dont la recherche est très active afin de percer leur mystère. Ils ont des énergies diverses allant de quelques MeV jusqu'à 10^{20} eV.

Le processus d'interaction des rayons cosmiques avec l'atmosphère (donc des noyaux d'azote et d'oxygène principalement) produit une gerbe "atmosphérique" suivant la figure 1.7.

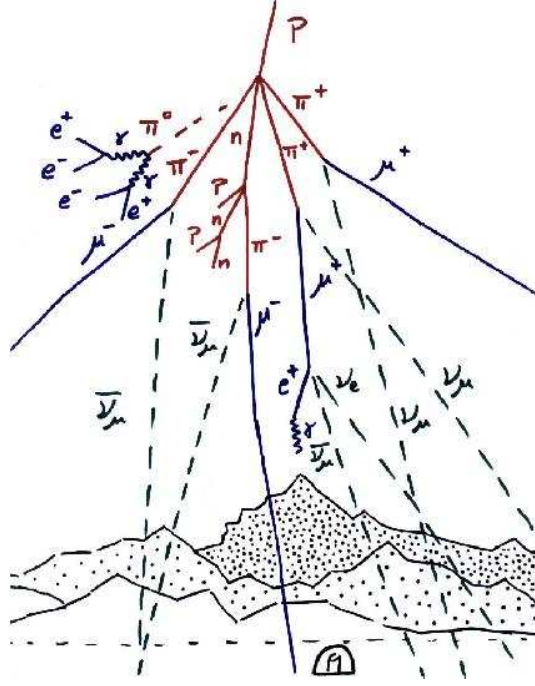


FIG. 1.7 – Gerbe électromagnétique issue d'un rayon cosmique

Lors de l'interaction d'un rayon cosmique avec un noyau de l'atmosphère, il est produit de nombreuses particules secondaires qui sont principalement des mésons π^0 qui donnent des électrons et des photons, développant la gerbe électromagnétique, et des mésons $\pi^\pm$ qui vont se désintégrer en couple $\mu - \nu_\mu$.

$$\pi^+ \rightarrow \mu^+ \nu_\mu \quad (1.50)$$

$$\pi^- \rightarrow \mu^- \bar{\nu}_\mu \quad (1.51)$$

Ensuite les muons se désintègrent à leur tour dans les canaux suivants :

$$\mu^+ \rightarrow e^+ + \nu_e + \bar{\nu}_\mu \quad (1.52)$$

$$\mu^- \rightarrow e^- + \bar{\nu}_e + \nu_\mu \quad (1.53)$$

Donc en faisant la somme des neutrinos produits à la suite de la désintégration d'un pion puis du muon résultant, on obtient un rapport (pour les neutrinos) muonique/électronique

de 2. Ce rapport est fiable malgré une baisse de celui-ci aux très hautes énergies car alors le muon a une durée de vie trop importante (effet relativiste) et donc atteint la Terre avant de se désintégrer.

1.1.5.2 Les sources artificielles

Les réacteurs nucléaires

Les neutrinos issus de réacteurs ont été les premiers à être détectés.

Les réacteurs produisent des neutrinos lors de la fission de noyaux (^{235}U par exemple) qui nous sert à produire de l'énergie en grande quantité.

Dans le cas de l'uranium, la fission s'accompagne de 6 désintégrations β et fournit donc un ensemble de 6 $\bar{\nu}_e$. Ce type de source ne produit que des $\bar{\nu}_e$ avec toutefois un taux de production qui varie suivant le produit de fission utilisé. De plus le flux de neutrinos varie suivant la puissance thermique du réacteur et aussi du nombre total de fission β qui s'enchaîne à la suite de la première fission.

Les $\bar{\nu}_e$ sont relativement de basse énergie allant de 1 à 10 MeV (la moyenne se situant vers 3 - 4 MeV) avec un flux intégré de l'ordre de 10^{20} - 10^{21} $\bar{\nu}_e.s^{-1}$ soit pour un détecteur situé à une centaine de mètres, 10^{11} $\bar{\nu}_e.s^{-1}$.

Les accélérateurs de particules

Les anti-neutrinos issus d'un accélérateur ont des intérêts très variés. L'aspect prépondérant est que c'est la manière d'obtenir des neutrinos la plus maniable par l'Homme, car en effet on peut choisir de produire ce que l'on veut aussi bien au niveau de la saveur qu'au niveau de l'énergie.

Ensuite ces neutrinos peuvent atteindre de très hautes énergies allant jusqu'à plusieurs dizaines de GeV, ce qui a pour effet d'être plus facilement détectable étant donné la section efficace plus importante.

Le principe est de former à l'aide d'un accélérateur, des particules instables qui se désintègrent en éjectant des neutrinos (comme pour les neutrinos atmosphériques). Les expériences NOMAD[1][2] et CHORUS ont utilisé un faisceau de protons de 450 GeV en l'envoyant sur une cible de béryllium. Cette interaction produit principalement des pions qui se désintègrent à leur tour en des muons et des neutrinos muoniques. On a donc un faisceau de ν_μ d'énergie moyenne de 23,5 GeV. Il reste une faible contamination en ν_e persistante dans ce type de faisceau du fait de la présence de K^- .

1.2 Les recherches d'oscillations actuelles et futures

L'étude des oscillations de neutrinos représente une recherche permettant de mesurer des paramètres d'oscillation dans le cas de résultats positifs, du moins de poser des limites en cas d'échec. Aussi diverses soient-elles, les expériences doivent s'adapter en fonction du type de source qu'elles cherchent à étudier. Cette adaptation reflète souvent la sensibilité d'une expérience à la mesure des paramètres-clés des oscillations.

1.2.1 Les recherches sur les neutrinos solaires

Les recherches sur les neutrinos solaires ont été très intenses ces dernières années. Nous pouvons compter de nombreuses expériences qui se sont intéressées à ce domaine d'étude : Homestake, GALLEX, SAGE, (Super)Kamiokande et SNO.

SNO s'est notamment illustrée par ses récents résultats quant à la mise en évidence d'oscillations de neutrino [3]. Celle-ci se révèle très sensible aux différentes saveurs de neutrinos et reste la seule expérience à retrouver une mesure de flux total² des neutrinos du 8B (toutes saveurs confondues), compatible avec le modèle solaire standard. Cette observation semble donc satisfaire à la réalisation d'oscillation $\nu_e \rightarrow \nu_x$, le flux reçu n'étant pas uniquement des ν_e . L'analyse des données de SNO dans [4] en complément des données déjà existantes dans le cadre des oscillations MSW, a sélectionné la solution aux grands angles de mélange dite LMA (Large Mixing Angle) pour les neutrinos solaires (figure 1.8).

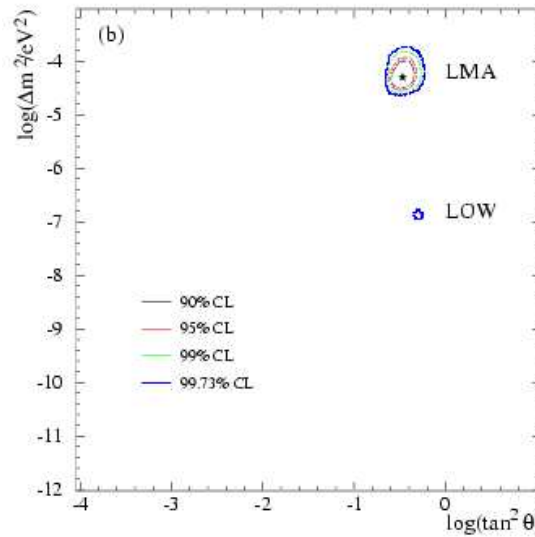


FIG. 1.8 – Solution autorisée pour les paramètres d'oscillation des neutrinos solaires [4].

1.2.2 Les recherches avec les neutrinos atmosphériques.

Dans le monde de nombreuses expériences ont effectué ou sont en train d'effectuer des recherches à propos des neutrinos atmosphériques. Super-Kamiokande (détecteur Cherenkov de 50 kT d'eau pure) affirme avoir mis en évidence des oscillations de neutrinos. En effet, elle a mesuré un rapport $\frac{\nu_\mu + \bar{\nu}_\mu}{\nu_e + \bar{\nu}_e}$ inférieur à 2, qui est normalement le rapport attendu s'il n'y a pas d'oscillation. On observe un déficit de ν_μ pour les neutrinos qui vont du bas vers le haut et donc ont traversé la Terre (figure 1.9). Au départ, on aurait pu penser à l'effet MSW mais les résultats ne mettent pas en évidence une influence de la Terre en tant que matière résonante. On ne voit en effet pas d'effet jour-nuit dans le cadre des neutrinos solaires. S'il y a effectivement des oscillations de neutrinos, on ne peut considérer que l'influence de la distance.

²Le flux total est en fait la somme des flux mesurés par SNO et SuperKamiokande.

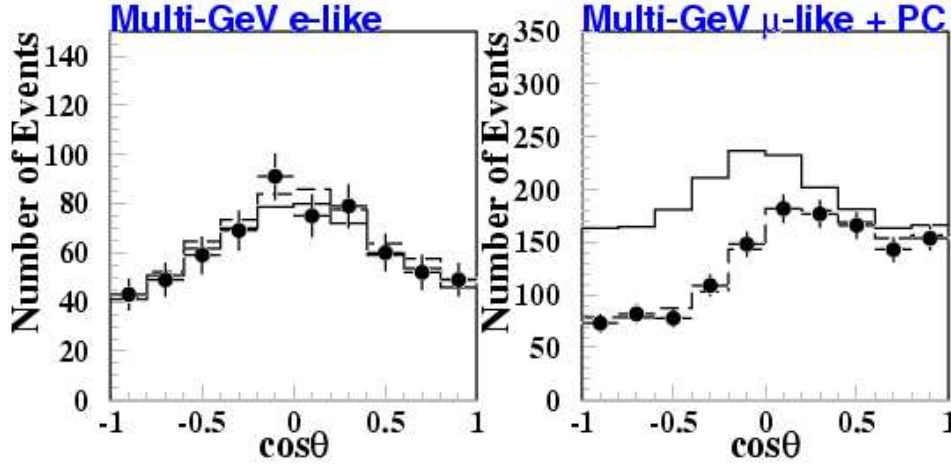


FIG. 1.9 – *Distribution angulaire (Zenith) de l'expérience SK pour les événements électroniques (a) et les événements muoniques. Les cercles noirs montrent les données, les histogrammes, le Monte-Carlo.*

Le problème de ce type de mesure est qu'il existe une incertitude de 5% pour le flux relatif, et encore c'est déjà bien car pour le flux absolu, les incertitudes sont d'environ 20% ce qui est énorme. La mise en évidence de ces oscillations ne se fera que lorsque l'on aura réduit ces incertitudes. En effectuant le rapport, on élimine les incertitudes liées à la section efficace d'interaction des neutrinos mais aussi au calcul de flux de neutrinos atmosphériques qui repose au départ sur une bonne connaissance des caractéristiques du processus de la production des mésons, particules primaires parentes de ces neutrinos. Dans l'annexe B, on peut se rendre compte qu'une grande partie de la production hadronique ouverte aux neutrinos atmosphériques de basse énergie est assez peu étudiée. De plus les données les plus utiles concernant les sections efficaces sont les plus anciennes et donc les moins précises du point de vue des incertitudes recherchées. C'est justement le point où l'on va voir plus tard que HARP va pouvoir améliorer en principe notre connaissance sur la production de ces mésons aux très basses énergies incidentes, vierges de données à l'heure actuelle.

1.2.3 Recherche auprès de réacteurs.

Une expérience étudiant les neutrinos issus de réacteurs nucléaires, KamLAND (au Japon) affirme avoir observé des oscillations par disparition de $\bar{\nu}_e$ avec une énergie supérieure à 3,5 MeV provenant de plusieurs réacteurs comptabilisant une distance moyenne de 180 km [5]. Le résultat de KamLAND semble être compatible uniquement avec la solution LMA et l'affirmerait donc comme étant la solution pour les oscillations des ν_e solaires.

1.2.4 Les neutrinos d'accélérateur

1.2.4.1 Préliminaires

L'étude d'oscillation auprès d'accélérateurs permet une grande flexibilité de par la possibilité de choisir la distance d'oscillation L à étudier entre la source d'émission du

neutrino et le détecteur. Par ailleurs, il est plus simple d'éliminer les bruits de fonds car on a l'information du temps de départ du neutrino. De plus en plus, on réalise que l'on a besoin d'aller dans les distances croissantes pour être sensible à des énergies plus faibles. On verra que les faisceaux conventionnels ont des limites ne permettant pas d'exploiter tout ce qui est théoriquement possible et que l'on envisage de mettre au point ce que l'on appelle des "usines à neutrinos".

1.2.4.2 Les faisceaux conventionnels

Il existe actuellement peu de faisceaux dans le monde. Comme il a été dit précédemment, il est important d'avoir des études d'oscillation auprès des accélérateurs sur des distances de plus en plus longues afin de faire varier le rapport L/E et améliorer la sensibilité à d'autres gammes d'énergie. Les premières lignes de faisceaux de neutrinos permettaient quelques centaines de mètres (NOMAD, CHORUS). Mais devant les résultats négatifs aux très faibles distances, une recherche sur des distances d'oscillations de plus en plus importantes est maintenant la tendance actuelle. Cependant nous verrons une exception dans le cas de MiniBooNE qui montre que les petites Baselines sont toujours d'actualité du fait que l'on cherche à confirmer les résultats de LSND (Liquid Scintillating Neutrino Detector)[6][7].

Je parlerai de K2K et de MiniBoone avec qui HARP a collaboré pour apporter une meilleure connaissance de leur faisceau (en réduisant les erreurs systématiques).

K2K

Il s'agit d'une expérience se faisant au Japon (K2K signifiant KEK To SuperKamiokande) et est actuellement la plus longue baseline du monde. L'objectif de K2K est de confirmer les résultats de Super-Kamiokande avec les neutrinos atmosphériques en envoyant un faisceau de neutrinos ν_μ partant du KEK (Tsukuba) vers le détecteur Cherenkov SuperKamiokande situé à environ 250 km (Figure 1.10).

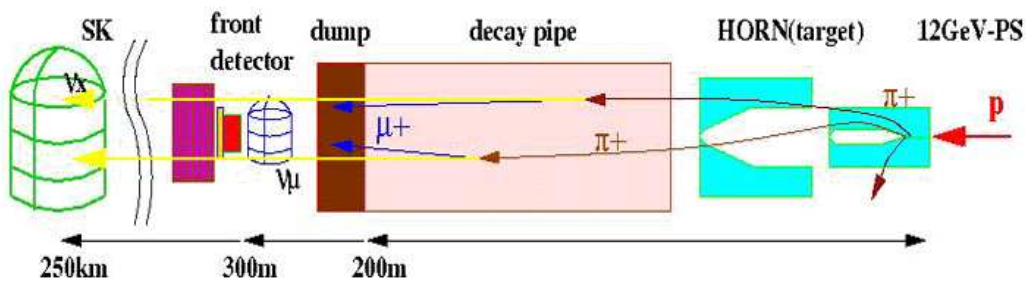


FIG. 1.10 – Schéma de production des neutrinos au KEK puis envoi vers Kamioka

Elle a commencé ses prises de données en 1999. Le faisceau de neutrinos est fabriqué selon la méthode classique actuelle : un faisceau de 12 GeV est envoyé sur une cible en aluminium ce qui produit des mésons (majoritairement des pions et kaons). Ces mésons sont collectés en charge par une corne. Cette sélection en charge est importante car elle permet de sélectionner les mésons suivant les neutrinos que l'on désire obtenir. Les pions chargés sont envoyés dans un tunnel de désintégration pour se désintégrer en muon et

neutrino (suivant les canaux montrés en 1.50 et 1.51). Les muons restant qui ne sont pas utilisés sont arrêtés par un “dump” ne laissant donc passer que les neutrinos.

$$\pi^+ \rightarrow \mu^+ + \nu_\mu \quad (1.54)$$

$$\pi^- \rightarrow \mu^- + \bar{\nu}_\mu \quad (1.55)$$

On obtient alors un faisceau de neutrinos dont on estime la composition à 99% en ν_μ et à 1% en ν_e , la composante en ν_e provenant de la désintégration de kaons.

Pour étudier ce faisceau de neutrinos en amont, un mini-détecteur Cherenkov (1 kilotonne) version réduite de SuperKamiokande (50 kT) (figure 1.11) est disposé à environ 300m après le tunnel de désintégration, permettant de le visualiser suivant le cône de lumière Tcherenkov produit la direction du faisceau au départ. Ce détecteur est secondé par un autre détecteur constitué de fibres scintillantes, verre au plomb et de chambres à muons afin de mesurer l'énergie des muons issus de la désintégration des pions mais aussi permettre de faire des mesures plus précises sur le profil du faisceau et sur les sections efficaces physiques afin d'apporter des éléments de données pour l'interprétation des résultats de SuperKamiokande.

K2K vise à étudier la région de $\Delta m^2 = 10^{-2} - 10^{-3} eV^2$ qui est la région où a été détectée l'anomalie pour les neutrinos atmosphériques avec Super-Kamiokande c'est-à-dire un signe positif d'oscillation.

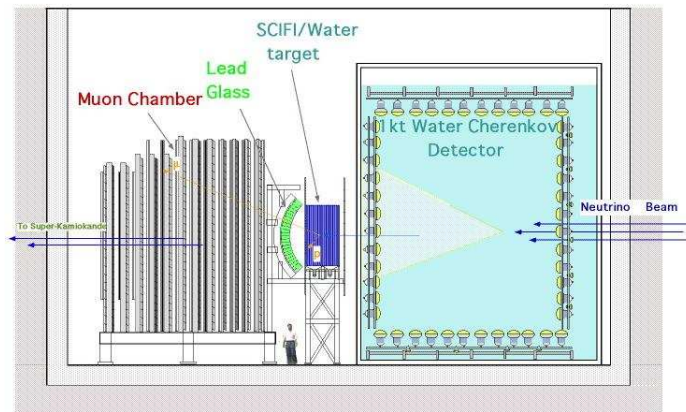


FIG. 1.11 – Réplique miniature de SuperKamiokande, fibres scintillantes et chambres à muons dans K2K.

MiniBooNE

L'expérience MiniBooNE (pour Mini-Booster Neutrino Experiment) s'effectue aux Etats-Unis au Fermilab (Chicago, Illinois) et a débuté ses prises de données en 2002. Elle vise également à confirmer les résultats des neutrinos atmosphériques mais également celui de LSND [7][8] en 1996. Le canal de recherche est donc $\nu_\mu \rightarrow \nu_e$ et donc vise à détecter l'apparition d'un ν_e [9].

Le détecteur est un imageur Cherenkov est à 500m de la ligne de faisceau de neutrinos, produits par un faisceau de protons de 8 GeV. Tout comme dans le cas de K2K, il s'agit de récupérer les mésons produits par l'intermédiaire d'une corne. La faible énergie du faisceau de protons limite la production de ν_e qui contamine le faisceau de ν_μ car on ne peut pas produire alors de kaons. Le faisceau est donc majoritairement peuplé de neutrinos ν_μ d'énergie entre 0,5 et 1 GeV.

La cible est une sphère de 11 mètres de diamètre contenant 445 tonnes d'huile minérale pure. Sur les bords de la sphère, 1220 PM couvrent 10% de la surface (voir figure 1.12). Dans une couche plus externe, 292 PM sont disposés en tant que "VETO" pour éliminer les signaux des rayons cosmiques. La détection se fait par la lumière émise par les particules produites issues de l'interaction quasi-élastique du neutrino (interaction courant chargé).

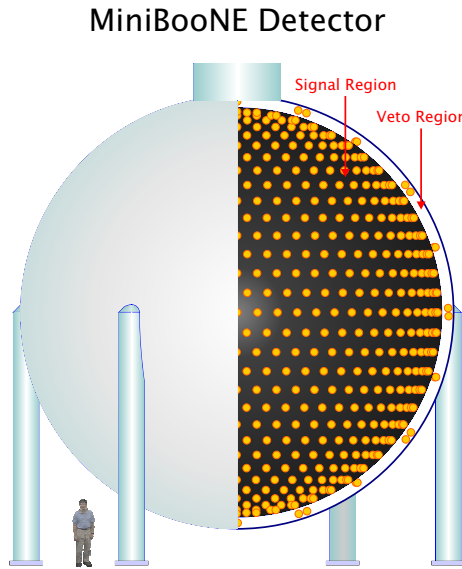


FIG. 1.12 – Une coupe de l'imageur Cherenkov de MiniBooNE avec ses PM

1.2.4.3 Les usines à neutrinos.

Les usines à neutrinos représentent l'avenir de la physique des neutrinos sur accélérateur. En effet pour le moment, celles-ci n'existent que sur le papier mais de nombreux physiciens cherchent actuellement à mettre en œuvre ce type d'expérience qui, comme l'on va voir dans la suite, permettra d'accéder à des régions d'étude jusqu'à présent inaccessibles.

Les Défis à relever

Les usines à neutrinos n'ont pas encore une architecture bien définie. Cependant on a déjà une idée de ce à quoi cela pourrait ressembler. La figure 1.13 est une possibilité de forme pour ce genre d'usine. Les incertitudes reposent surtout sur la forme de l'anneau de stockage à muon dont nous allons reparler.

Une usine à neutrinos est différente des faisceaux classiques actuels de par son objectif :

Obtenir des faisceaux de neutrinos par la désintégration du muon dont on connaît parfaitement le processus de désintégration en un neutrino et un antineutrino. Il s'agit donc d'avoir une parfaite maîtrise des muons.

Une telle performance passe par une étude approfondie de défis que l'on ne maîtrise pas bien encore. Nous allons voir précisément les obstacles sur lesquels la communauté scientifique doit se pencher au travers de la description des étapes pour la construction du faisceau de neutrinos.

Voici une approche qualitative de ces étapes :

- Tout commence comme dans les faisceaux conventionnels où l'on envoie un faisceau de protons (d'une énergie qui reste à déterminer) sur une cible (nature et dimensions à déterminer). L'interaction produit alors des mésons (pions et kaons).
- De ces mésons produits, on ne veut que les pions chargés dont on sait qu'ils se désintègrent en muons et ν_μ avec un rapport de branchement de 99,99% contre 63,51% pour les kaons chargés. On a donc une pureté plus importante avec les pions d'où cet intérêt. On cherche donc à construire une sorte de corne magnétique qui permettrait de ne récupérer que les pions, chose que l'on ne sait pas faire actuellement. Ces pions seraient envoyés dans un tunnel de désintégration à la suite de quoi, on oublie le neutrino ν_μ produit (qui est le centre d'intérêt des lignes classiques) pour se concentrer sur le muon. Néanmoins ce muon ne va pas être laissé dans un tunnel de désintégration comme pour les pions car la durée de vie du muon augmentée par l'effet relativiste demanderait un tunnel d'une dizaine de kilomètres. La solution serait alors de l'injecter dans un anneau de stockage.
- On cherche donc à faire un faisceau de muons extrêmement bien collimé et intense : pour cela, on pense les refroidir afin de réduire leur espace de phase puis les envoyer dans une unité d'accélération. Le faisceau est alors prêt pour être injecté dans un anneau de stockage.
- L'anneau de stockage devra alors comporter une ou deux sections de droites permettant aux muons qui s'y désintégreront d'émettre des neutrinos dans une direction très bien défini avec un minimum de divergence mais surtout très intense.

Les problèmes actuels

Les problèmes de conception sont multiples. On va commencer par le début en parlant du faisceau de protons. Le but est d'envoyer un faisceau de protons ultra-intense (avec une puissance de l'ordre de 4 MW ce qui est 2 ordres de grandeur supérieur à ce que l'on fait actuellement). L'envoi d'une telle puissance est colossale pour une cible. Cela peut poser des problèmes d'échauffement excessif mais aussi d'usure prématurée. L'étude de la production hadronique lors de l'interaction avait été effectuée dans les années 70 par diverses expériences (cf annexe B) cependant on verra que celles-ci ont obtenu des mesures de sections efficaces avec des incertitudes trop importantes (environ 15%) pour réaliser une machine qui pourrait capturer de manière précise les pions. Ces incertitudes sont issues d'un manque de recouvrement en angle solide forçant des interpolations et

extrapolations, sources d'incertitudes.

Il y a ensuite tout le processus de refroidissement par ionisation [10] (utilisation d'absorbeur d'énergie par $\frac{dE}{dX}$) qui reste délicat pour les muons car il s'agit de ne pas perdre trop de temps étant donné qu'à plus faible énergie, sa probabilité de désintégration augmente. Donc il faut traiter les muons afin d'en faire un paquet puis alors les accélérer (ce qui augmente leur durée de vie) pour l'injection dans l'anneau de stockage.

Il y a donc un nombre important de points à étudier, et l'on verra que HARP vise l'étude de la production hadronique aux faibles énergies.

Remarque : On peut noter que le principe d'usine à neutrino permet également d'alimenter abondamment des collisionneurs à muons, source d'un nouveau type de physique.

1.2.4.4 Apport des usines à neutrinos

Les premiers résultats potentiels d'oscillation [11] provenant des neutrinos atmosphériques (SuperKamiokande 1999) sont à considérer avec des réserves car ceux-ci possèdent des incertitudes trop importantes pour être concluants.

La physique des neutrinos auprès des accélérateurs cherche à confirmer les résultats des expériences à neutrino naturel. Des baselines existent déjà mais on va voir que les usines à neutrinos sont la solution aux limites des lignes de faisceaux conventionnelles.

Certes des études sont actuellement en cours pour essayer de simplement améliorer les faisceaux conventionnels [12] déjà existants. Ce type d'amélioration nécessite toutefois des connaissances plus approfondies sur des points dont l'usine à neutrinos a besoin (faisceau de proton très intense, production hadronique...).

On peut effectuer la comparaison entre un faisceau classique et un faisceau issu d'une usine à neutrino. Comme nous l'avons déjà vu, dans le premier cas, les neutrinos proviennent de la désintégration en deux corps des mésons chargés alors que pour le second, il s'agit de la désintégration en trois corps du muon (résultat de la désintégration en deux corps des pions chargés). Les faisceaux conventionnels n'ont pour but que de produire des faisceaux purs de ν_μ en vue de mesurer un signal de type $\nu_\mu \rightarrow \nu_e$. Il subsiste généralement une contamination en ν_e qui provient de la désintégration du kaon dans des canaux mineurs (mais gênants) et de muons qui se sont désintégrés avant d'être arrêtés donnant un neutrino électronique. Cette contamination est un problème sérieux pour la mesure d'un signal d'apparition ν_e . De plus ce taux de ν_e n'est pas suffisamment important pour obtenir un faisceau de ν_e pour l'étude de $\nu_e \rightarrow \nu_\mu$. Une usine à neutrino permet d'avoir une véritable production de ν_e pour ce type de recherche.

La désintégration du muon

Comme nous avons pu le voir au (1.52) et (1.53), les muons (anti-muons) se désintègrent en 3 corps permettant d'avoir une proportion égale de ν_e ($\bar{\nu}_\mu$) et de $\bar{\nu}_e$ (ν_μ) sans ambiguïté. Ainsi on peut obtenir 2 types de saveurs avec leurs anti-particules associées. On accède ainsi à une solution efficace pour les études de la violation de CP.

De plus il est possible de modifier le spectre en énergie des neutrinos en choisissant la polarisation des muons [13].

1.2.4.5 En attendant les usines... les “Super-Beams”

On vient de voir que la conception d’une usine à neutrino n’est pas encore immédiate. C’est pourquoi un autre projet a été élaboré : les *super-beams*. Ce projet en cours d’étude n’est en fait qu’une amélioration des faisceaux classiques. L’amélioration consiste comme pour les usines à neutrino à augmenter l’intensité de ces faisceaux par la montée en puissance des faisceaux de protons. Ces Super-Beams reposent toutefois sur la méthode de production classique des neutrinos en utilisant les neutrinos issus de la désintégration des pions et kaons. Ce type de faisceaux semble être à court terme la solution la plus abordable pour effectuer de nouvelles expériences neutrino.

1.2.5 Les neutrinos : peut-être un casse-tête théorique.

En considérant le schéma classique des états de saveurs composés des états de masse, on pense avoir résolu une conception théorique de ce qu’est le neutrino. Cependant les différents résultats obtenus sur les signes positifs d’oscillation, nous conduisent à de nouvelles interrogations quant à cette conception proprement dite. Récapitulons les valeurs de Δm^2 des divers cas d’expérimentation :

$$\Delta m_{sun}^2 = 10^{-12} - 10^{-11} eV^2 (vide) - 10^{-4} eV^2 (MSW) \quad (1.56)$$

$$\Delta m_{atm}^2 = 10^{-3} - 10^{-2} eV^2 \quad (1.57)$$

$$\Delta m_{LSND}^2 = 10^{-1} eV^2 \quad (1.58)$$

On se retrouve alors avec 3 différences de masses indépendantes impliquant l’existence d’un 4ème neutrino. Ce neutrino n’interagirait pas par interaction faible et ne serait donc pas impliqué dans une contribution à la masse de bosons faibles : il serait alors appelé le neutrino “stérile” (ν_s). En combinant les 3 différences de masse, on arrive à 2 types de modèles : d’une part, le modèle (3+1) considérant 3 états de masse proches qui sont séparés par Δm_{sun}^2 et Δm_{atm}^2 et le 4ème isolé à environ Δm_{LSND}^2 de ces derniers, et d’autre part, le modèle (2+2) considérant 2 ensembles de 2 états proches séparés de Δm_{LSND}^2 . Mais à l’heure actuelle, ces 2 scénarios ne sont pas favorisés et il se peut qu’une confirmation par MiniBooNE des résultats de LSND puisse amener à reconsidérer le cadre théorique des neutrinos.

1.2.6 Conclusion

La physique des neutrinos est un secteur en plein essor et obtiendra sans doute des résultats décisifs dans les 20 prochaines années pour trouver des réponses aux interrogations actuelles. Pour l’heure, des résultats encourageants restent à être confirmés. La mise en œuvre d’expériences sensibles aux différentes énergies cherche à être secondée par des expériences accélérateurs qui sont en évolution technologique et beaucoup plus souple en utilisation. Du concept de ligne de faisceau de neutrino classique au concept d’usine à neutrino, il n’y a en fait qu’un retour aux bases de la création du neutrino pour mieux rebondir vers des performances adéquates. Il s’agit alors d’avoir une parfaite maîtrise.

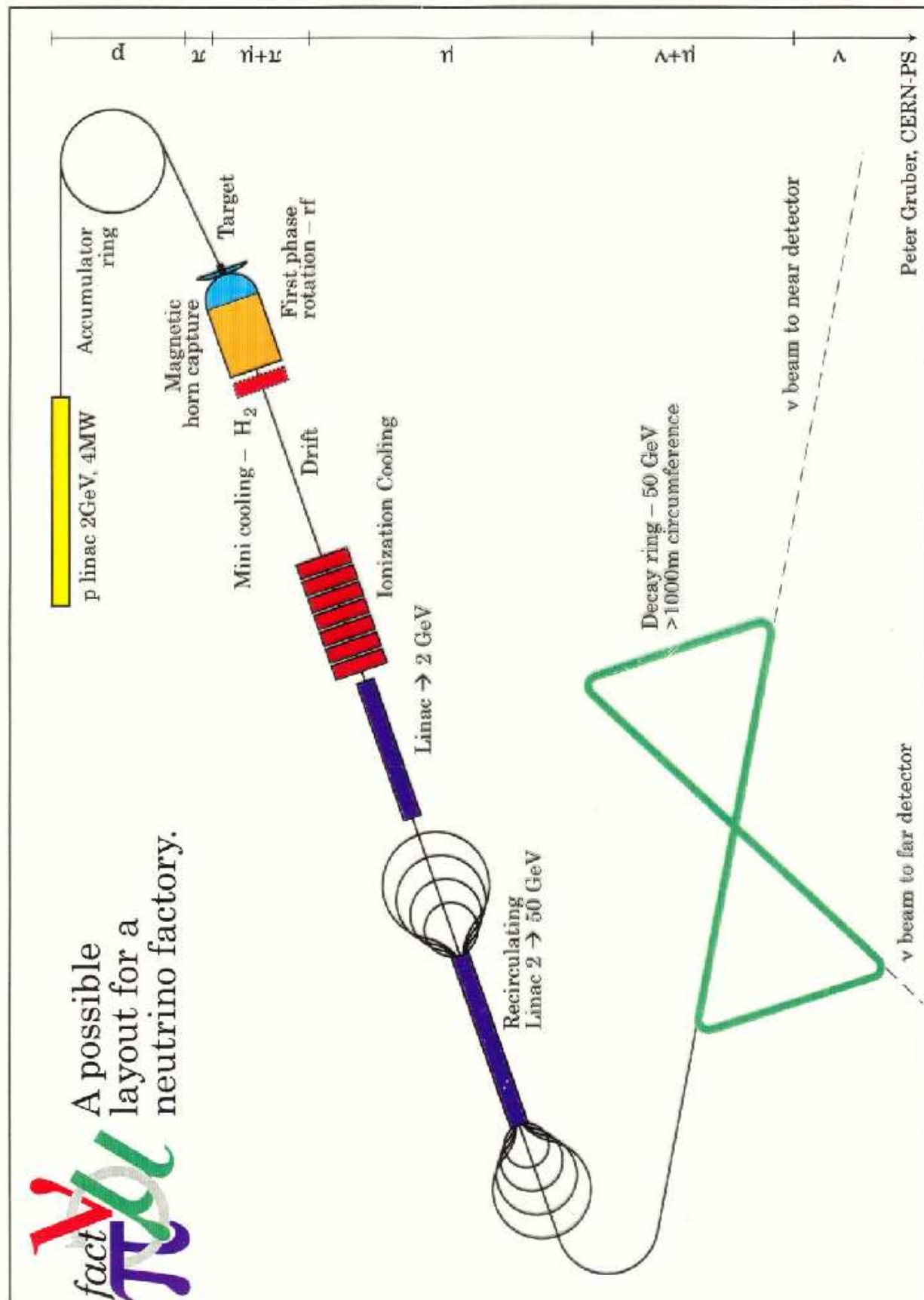


FIG. 1.13 – Schéma éventuel d'une usine à Neutrinos.

Deuxième partie

L'expérience HARP

Chapitre 2

L'expérience HARP

L'expérience HARP (signifiant HAdRon Production) est une collaboration internationale regroupant 24 instituts pour un total de 120 physiciens (détail dans l'annexe C).

Elle consiste à effectuer des mesures de la section efficace de la production hadronique dans les interactions proton-cible de faibles énergies (entre 2 et 15 GeV/c). La connaissance précise de ces sections efficaces est nécessaire à la conception des usines à neutrinos, au calcul des flux de neutrinos atmosphériques [14] à la compréhension des faisceaux classiques de neutrinos et à l'amélioration des générateurs hadroniques dans les codes Monte-Carlo aux basses énergies [15]. Ce type de mesure n'est certes pas nouveau et avait déjà été effectué dans diverses expériences des années 70 (voir annexe B). Alors pourquoi y revenir ? Tout simplement parce que les incertitudes sur ces mesures sont de l'ordre de 15%. La raison est principalement une faible acceptance et une mauvaise maîtrise des incertitudes systématiques dues à la composition du faisceau et à l'identification des particules secondaires produites. De plus ce type de recherche n'a pas été effectué de manière plus approfondie en variant les cas de figures. HARP cherche à réduire considérablement les sources d'incertitudes statistiques et systématiques. Sa sensibilité aux basses énergies et son étude sur des cibles multiples en font une expérience sans concurrence.

À l'heure actuelle, après 14 mois de prises de données réparties sur 2 ans (2001-2002), le détecteur est en cours de démantèlement et une phase d'étude de la qualité des données a été amorcée.

Dans ce chapitre, nous verrons une présentation globale de l'instrumentation utilisée dans cette expérience. Un premier paragraphe est consacré au faisceau incident et à son instrumentation de contrôle. Ensuite on abordera le système de déclenchement des événements, et les cibles utilisées. On finira par une présentation du détecteur avec une description personnalisée des sous-détecteurs.

2.1 Une expérience au bout de la ligne T9 du CERN.

L'accélérateur PS (Proton Synchrotron) maintient des protons à 24 GeV/c. Ces protons ne sont pas forcément directement injectés dans HARP. Le principe (figure 2.1) est d'envoyer une partie du faisceau de protons sur une cible et d'utiliser les particules secondaires pour constituer un nouveau faisceau. La sélection sur ces particules s'effectue sur l'énergie et la charge des particules. Concrètement on ne dispose pas alors d'un faisceau pur (on va voir par la suite que c'est notamment le cas aux très basses énergies).

D'autres lignes (figure 2.2), orientées autour du PS alimentent en faisceau de nombreuses expériences en se répartissant les paquets ou "spills" du PS. Pour notre expérience, on a utilisé en moyenne 2 des 6 "spills" de protons ou pions s'étalant sur 400 ms pour un cycle de 14,4 s. L'intensité acceptée par l'expérience varie de 10^4 à 10^5 particules incidentes suivant l'épaisseur de cible utilisée (voir plus loin).

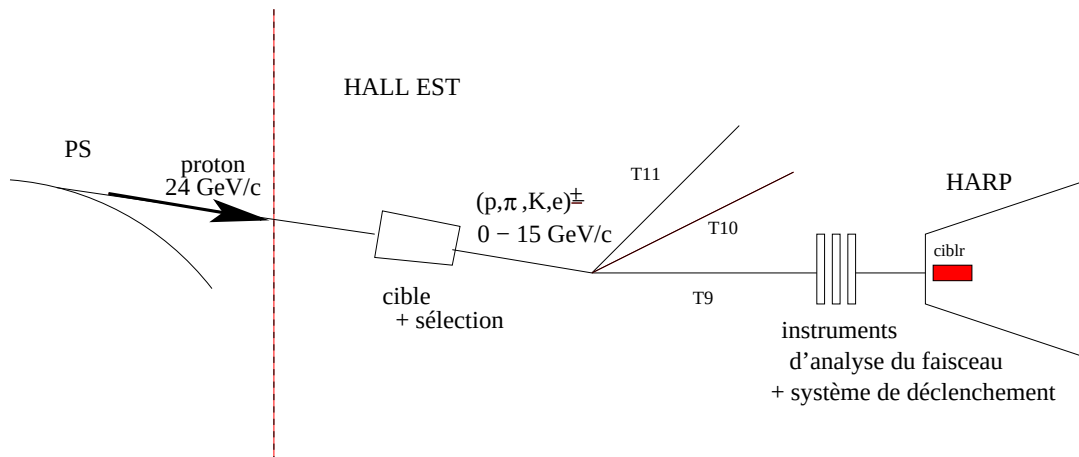


FIG. 2.1 – La ligne T9 provenant du PS : des protons de $24 \text{ GeV}/c$ sont envoyés sur une cible produisant des particules secondaires. Les particules de même charge et de même énergie sont sélectionnées pour former un faisceau vers l'expérience HARP.

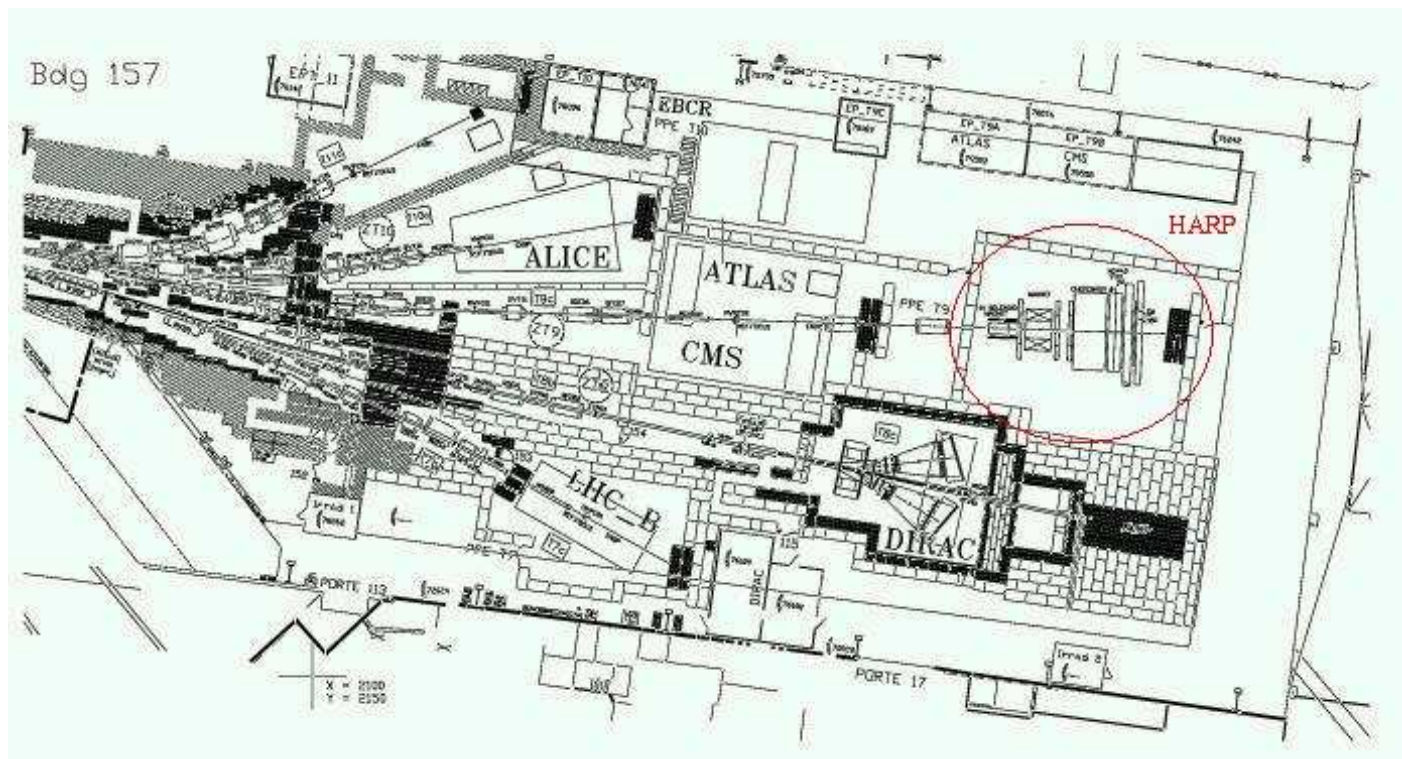


FIG. 2.2 – Vue du dessus du hall est : on aperçoit la ligne T9 avec à son bout l'expérience HARP (dans le rond).

Etant données les conditions expérimentales imposées par le faisceau incident, l'expérience HARP s'est dotée d'une instrumentation de faisceau en vue d'identifier chaque particule incidente afin d'optimiser l'enregistrement de bons événements. Pour le détecteur, l'objectif était de couvrir la plus importante acceptance possible. Aux basses énergies, beaucoup de particules secondaires empruntent des directions transverses, et l'on verra le rôle clé de la TPC (Chambre à Projection Temporelle, *en français*) dans ce cas précis. Cependant pour les particules orientées plus en avant, un spectromètre composé d'identificateurs (Time Of Flight, détecteur Tcherenkov et calorimètres) et de trajectographes (Chambres à dérive) permet de les analyser.

2.2 Présentation de l'expérience HARP

2.2.1 La maîtrise du faisceau incident

Dans ce type d'expérience, le faisceau incident est un élément clé. Il s'agit de bien le connaître sur plusieurs points :

- L'identification des particules incidentes. Les faisceaux disponibles ne sont jamais des faisceaux purs et il est donc important de pouvoir sélectionner, on-line ou off-line, les interactions.
- La structure en quantité de mouvement du faisceau.
- La géométrie de la ligne de faisceau, pour pouvoir prévoir et contrôler l'endroit où le faisceau frappe la cible.

En général, la section efficace d'interaction avec une cible fixe est donnée par la formule suivante :

$$\sigma = \frac{N_{obs}}{\epsilon N_{faisceau} T \rho_{cible}} \quad (2.1)$$

où :

ϵ est l'efficacité de sélection des particules incidentes. Elle comprend les efficacités de détection, de reconstruction des traces, d'identification des particules ;

$N_{faisceau}$ est le nombre de particules incidentes ayant la possibilité d'interagir ;

ρ_{cible} est la densité de la cible ;

N_{obs} est le nombre d'interactions c'est-à-dire le nombre d'événements observés après des coupures de sélection et la soustraction du bruit de fond ;

T est l'épaisseur de la cible.

Le numérateur dépend essentiellement de l'analyse des données et des choix de coupures de sélection. Par contre le dénominateur est sensible au faisceau incident mais également à la cible. Nous verrons dans le paragraphe consacré à la cible, les caractéristiques qui doivent être connues et comprises.

Les instruments nécessaires à l'étude du faisceau sont de 3 types :

- 2 Tcherenkov à seuil (BC1 et BC2).

- des scintillateurs : pour le calcul du temps de vol (TOF A et TOF B), pour la réjection du halo (Halo A et Halo B), pour le positionnement du faisceau (BS3 (Beam Scintillator) et TDS (Time Defining Scintillator)) (figure 2.7)
- 4 chambres à fils (MWPC = Multi-Wire Proportional Chamber)

Les Tcherenkov et les scintillateurs ont un rôle d'identificateur des particules du faisceau et peuvent être utilisés dans le système de déclenchement, alors que les chambres à fils nous permettent de mesurer la trajectoire du faisceau, y compris donc son angle d'incidence sur la cible.

2.2.1.1 L'identification des particules du faisceau

TOF

A basse énergie, la mesure du temps de vol peut permettre d'identifier les particules incidentes : on détermine leur masse à partir de la différence de temps Δt des signaux dans les plans de scintillateurs placés à une distance d :

$$m = \sqrt{\left(\frac{\Delta t \cdot p}{d}\right)^2 - p^2} \quad (2.2)$$

Les TOF A et B sont deux plans de scintillateurs séparés de 21 m. Leur résolution temporelle est de 170 ps.

Tcherenkov

L'identification des particules peut aussi être faite grâce à la détection de lumière dite "Tcherenkov". Cette lumière est émise lorsqu'une particule va plus vite que la vitesse de la lumière dans un milieu, vitesse définie pour un indice de réfraction n . La ligne de faisceau utilisée par HARP est équipée de deux Tcherenkov remplis de CO_2 . L'indice de réfraction dépend de la nature du gaz et de sa pression. Ainsi on augmente la capacité de détection des compteurs Tcherenkov en ajustant la pression selon la quantité de mouvement du faisceau.

HARP ayant pris des données avec différentes particules incidentes et un éventail d'impulsions allant de 1,5 à 15 GeV, c'est en combinant les conditions sur le temps de vol et les pressions dans BC1 et BC2 que l'identification et la sélection ont été possibles. Par exemple à 3 GeV/c, les électrons sont parfaitement identifiés dans BC1 avec une pression de 1 bar.

Il a été mesuré une quantité d'électrons, variable avec l'énergie. Comme on peut le voir dans le tableau 2.1, c'est aux faibles énergies que cette proportion est la plus importante. L'énergie est positive (négative) pour le faisceau lorsqu'il est composé de particules positives (négatives). On remarque qu'il y a des électrons même pour un faisceau chargé

positivement. La teneur en pions est comparable à celle des protons avec toutefois une relative stabilité vis-à-vis de l'énergie du faisceau. En ce qui concerne les kaons, la ligne T9 du CERN offre un ordre de grandeur inférieur à celui des pions.

Energie du faisceau (GeV)	quantité d'électron (%)
-1,5	75
1,5	23
2	14
3	7
5	6
8	1

TAB. 2.1 – *Proportion des électrons en fonction de l'énergie du faisceau.*

La stratégie adoptée pour l'identification des pions et des électrons superflus est montrée dans le tableau 2.2. On voit clairement qu'aux faibles énergies, on utilise les éléments de contrôle du faisceau alors que pour les pions, le TOF pour "Time of Flight" (dont nous parlerons dans la partie consacrée au détecteur) est jugé plus efficace.

Energie du faisceau en GeV	identification des électrons	identification des pions
inférieur à 5	BC2 à 1 bar	TOF
5	BC1 à 0,6 bar	TOF + BC2 à 2,5 bars
supérieur à 5	analyse Off-line	BC1 à 1,25 bar + BC2 à 1,5 bar

TAB. 2.2 – *stratégie pour l'identification des électrons et des pions.*

Le pouvoir de séparation des particules aux faibles énergies est largement effectif comme on peut le voir sur la figure 2.3 qui montre le temps de vol mesuré entre les deux éléments TOF A et TOF B pour un faisceau à 3 GeV/c. L'identification des particules aux plus hautes énergies est assurée par les compteurs Tcherenkov BC1 et BC2.

Cette séparation diminue lorsque l'énergie du faisceau augmente. La figure 2.4 montre qu'à 2 GeV/C, la séparation entre les protons et les pions est de 20σ et à 5 GeV/C, elle chute à $4,6\sigma$.

2.2.1.2 Position du faisceau

La trajectoire des particules du faisceau est reconstruite à l'aide de 4 chambres à fils (MWPC, figure 2.5). Chaque chambre est constituée de deux plans de fils orientés à 90 degrés l'un par rapport à l'autre. L'association des deux plans permet de reconstruire une position tridimensionnelle le long de la trajectoire. La combinaison des 4 points nous donne la direction des particules du faisceau [16].

La reconstruction de la trajectoire des particules incidentes lors d'une interaction est bien sûr importante mais c'est surtout on-line que ces chambres ont été essentielles. Leur

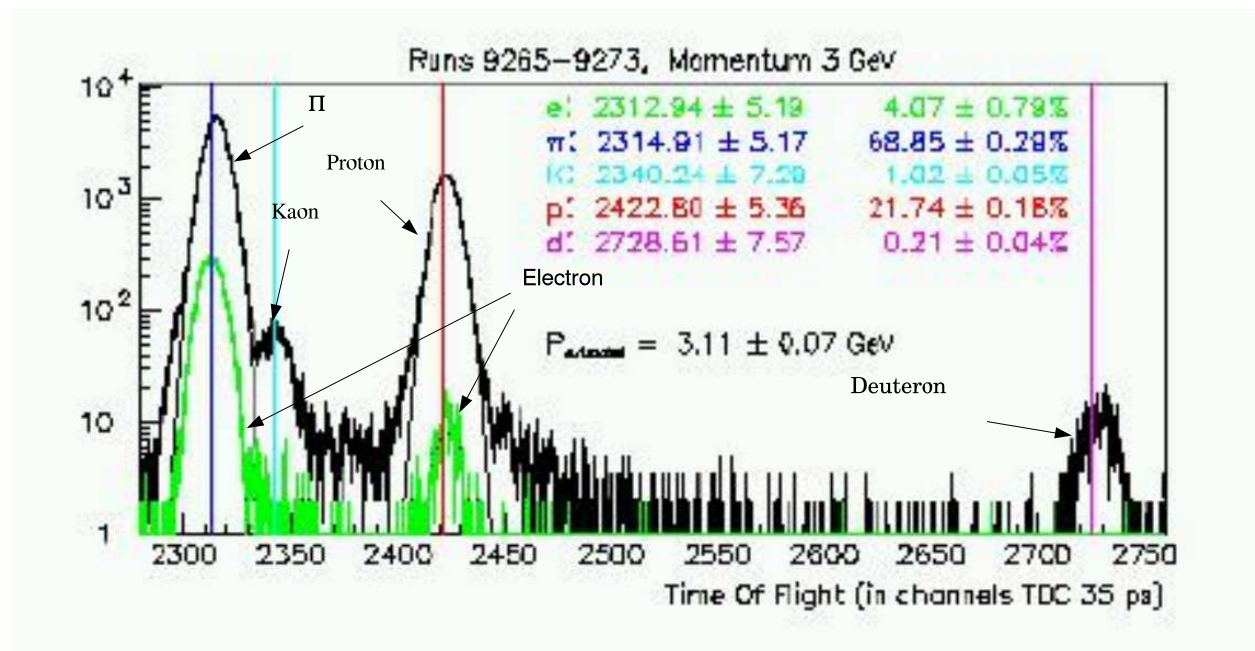


FIG. 2.3 – Séparation des particules du faisceau par mesure du temps de vol pour une quantité de mouvement de 3 GeV/c.

monitorage a constitué un élément de diagnostic permanent du faisceau : les profils enregistrés ont permis d'ajuster les quadrupoles et les collimateurs pour que la tache du faisceau au niveau de la cible soit plus petite que son diamètre, mais aussi que sa tache au niveau du FTP (voir la partie suivante : système de déclenchement) soit contenue dans son trou (figure 2.6).

En extrapolant des traces reconstruites (avec les MWPC) de particules du faisceau de quantité de mouvement de 3 GeV/c au niveau du FTP, on peut voir que les particules, n'ayant pas interagi dans la cible, passent à l'intérieur d'un trou de 6 cm de diamètre percé au centre du FTP (figure 2.6), et n'amènent pas de déclenchement.

2.2.2 Système de déclenchement

Le but de l'expérience HARP étant la mesure des sections efficaces de production de hadrons, le système de déclenchement doit permettre d'enregistrer tous les événements où une particule incidente donnée a produit dans la cible un ou des hadrons secondaires. Le trigger doit donc être capable à la fois d'identifier la particule incidente (comme décrit dans la partie précédente) et de réagir au passage de particules secondaires.

Le déclenchement sur les particules du faisceau s'établit par une chaîne logique (2.7) reposant sur les unités suivantes : TOF A et TOF B que l'on a déjà vu pour l'identification

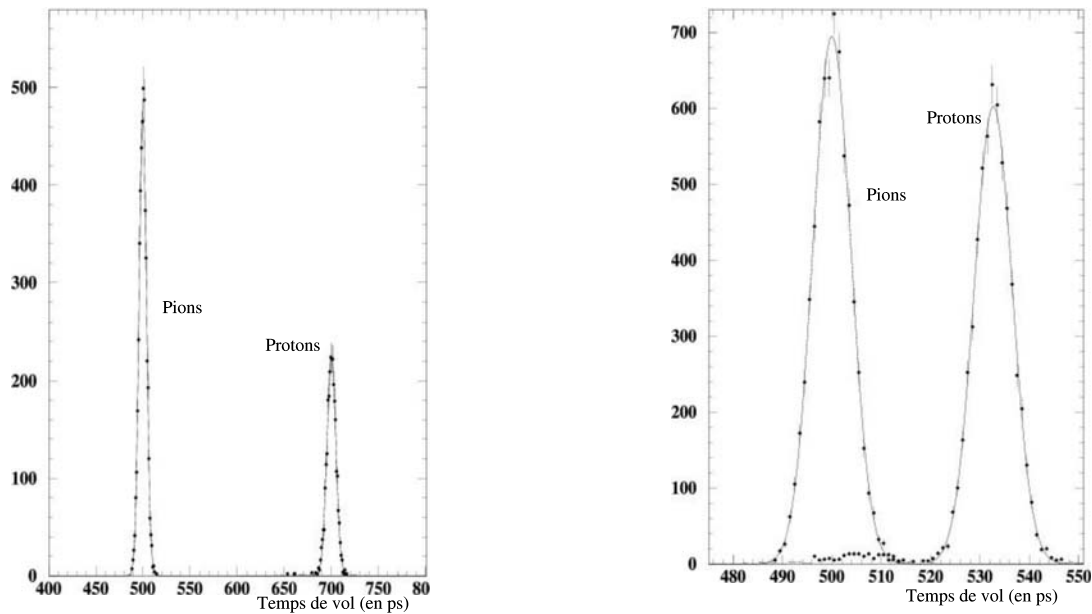


FIG. 2.4 – Séparation des protons et des pions dans les TOF de faisceau à 2 (*gauche*) et 5 GeV/c (*droite*) pour l'énergie du faisceau - abscisse : temps (en ps) - ordonnée : unité arbitraire.

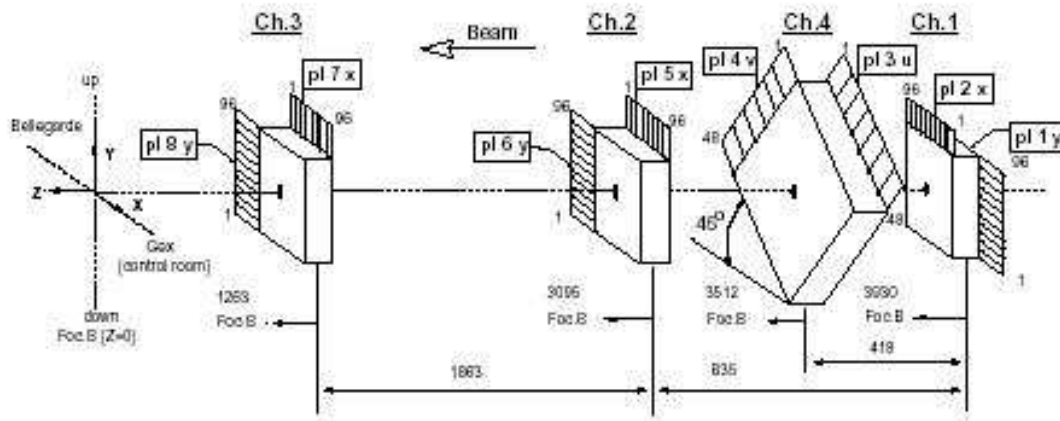


FIG. 2.5 – Les 4 chambres à fils (MWPC) permettant de reconstruire les trajectoires des particules du faisceau. Trois de ces chambres sont identiques avec une surface de détection de $10 \text{ cm} \times 10 \text{ cm}$. Une dernière chambre aux dimensions plus importantes est orientée à 45 degrés.

(voir 2.2.1), les HALO A et HALO B pour le rejet de particule ne se dirigeant pas vers la cible (ceux-ci ont un trou au centre ne déclenchant pas sur les particules de faisceaux), et le BS3 (Beam Scintillator) avec le TDS (Target Defining Scintillator) pour contrôler le contact avec la cible.

Quant au déclenchement sur les particules secondaires (ce qui est un signal d'interaction), il est assuré dans HARP par deux détecteurs : ITC et FTP (figure 2.7).

Les particules secondaires chargées produites dans la cible à grand angle sont détectées à l'aide de l'ITC (Inner Trigger Chamber) : ce détecteur est constitué de trois couches de fibres scintillantes formant un cylindre autour de la cible. Son efficacité pour détecter

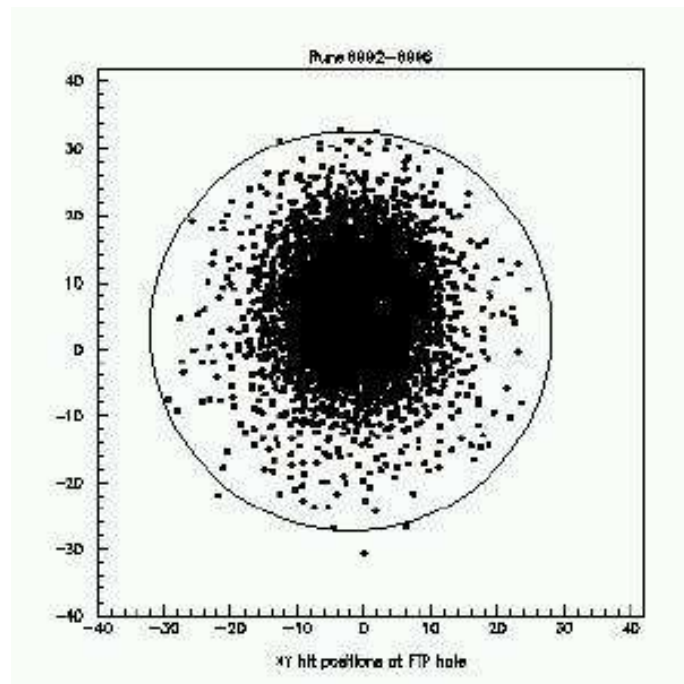


FIG. 2.6 – *Reconstruction de la position des particules du faisceaux au niveau du FTP. Le cercle noir représente le trou au centre du FTP.*

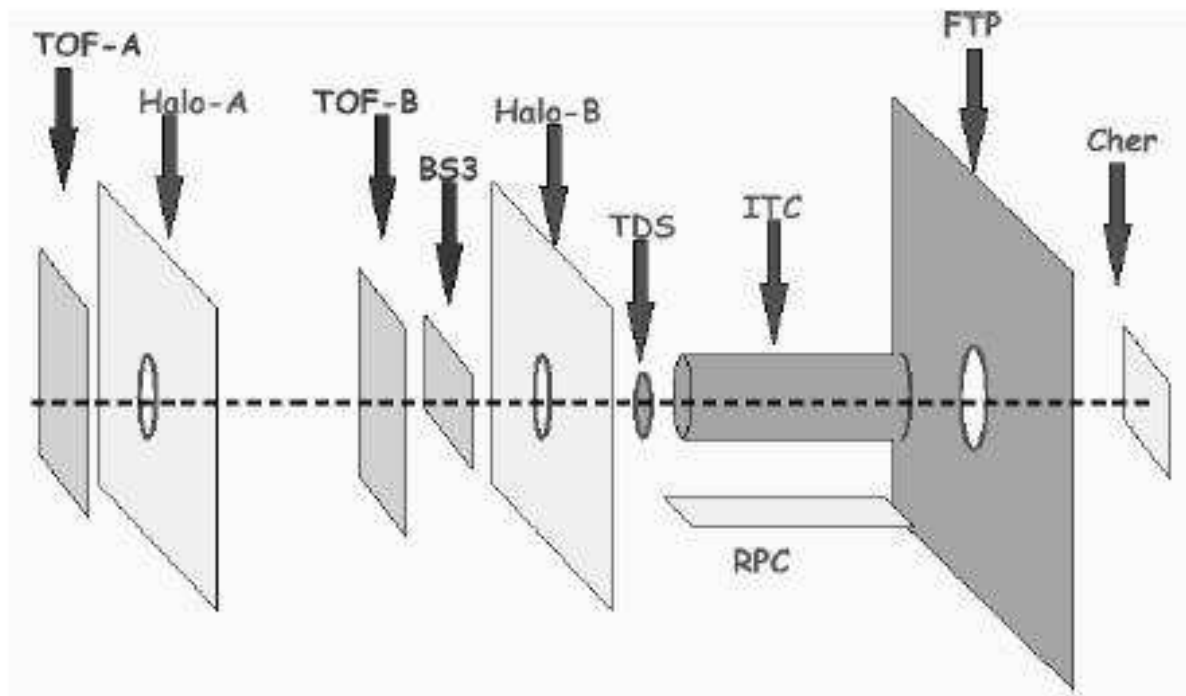


FIG. 2.7 – *Le système d'analyse du faisceau et le système de déclenchement. Le faisceau est représenté par la ligne centrale et se dirige de la gauche vers la droite. La cible est localisée à l'intérieur de l'ITC.*

les particules chargées dépasse les 99% [17]. Les particules secondaires produites à petit angle, hors de portée de l'ITC, sont détectées par le FTP (Forward Trigger Plane). Il se compose de deux plans de scintillateurs, 7 lattes horizontales et 7 lattes verticales, couvrant une surface de $1,4 \times 1,4 \text{ m}^2$. Au centre du FTP, un trou de 6 cm permet aux particules du faisceau n'ayant pas interagi avec la cible, de passer sans l'activer.

En plus de la condition de sélection par identification de la particule incidente (condition variant d'une prise de donnée à l'autre), la condition de déclenchement est donc :

$$\text{TRIG} = \text{TOF A} \cdot \overline{\text{HALOA}} \cdot \text{TOF B} \cdot \text{BS 3} \cdot \overline{\text{HALOB}} \cdot \text{TDS} \cdot (\text{ITC} \cup \text{FTP})$$

En restant dans cette configuration, nous ne pouvons obtenir, dans le cas des cibles minces, que des événements d'interaction biaisant ainsi notre mesure de section efficace par un flux incident apparent plus faible. Ainsi une logique a été ajoutée pour prendre un échantillon des particules incidentes qui n'interagissent pas (dans la cible). Pour cela, on utilise un facteur de réduction (DOWNSCALE) de 64 au niveau du trigger $\text{DOWNSCALE} = \text{BS 3} \cdot \text{TOF B}$. Ainsi dans le cas des prises de données sur cible mince, on rajoute des événements de non-interaction par la logique suivante :

$$\text{TRIG}_{\text{mince}} = \text{DOWNSCALE} \cup \text{TRIG}$$

2.2.3 Les cibles

La cible est placée au cœur de la TPC pour permettre la détection de toutes les particules chargées secondaires, y compris celles qui vont vers l'arrière. La TPC est équipée d'une cage de champ cylindrique de 20 cm de diamètre dans laquelle est fixé l'ITC : la cible est insérée au milieu de ce cylindre. (figure 2.10)



FIG. 2.8 – Photo du porte-cible (à gauche) et du tube contenant la cible choisie (à droite).

Le grand nombre de cibles prévues pour l'expérience a conduit la collaboration à concevoir un système de porte-cible (figure 2.8) facile à manipuler pour minimiser le temps d'interruption de la prise de données (inférieur à une demi-heure) lors des changements de cible. Ce support a été réalisé en fibre de carbone minimisant la matière devant la cible mais suffisamment résistant pour avoir une flèche négligeable.

Matériau	Z	épaisseur		Remarque
		$0,02\lambda$ (cm)	1λ (cm)	
Be	4	0,81		0,5 λ (cible MiniBooNE)
Be	4			
C	6	0,76	38,01	
Al	13	0,79	39,44	0,5 λ (cible K2K)
Al	13			
Cu	29	0,30	15,00	
Sn	50	0,45		
Ta	73	0,22	11,14	
Pb	82	0,34	17,05	
Matériau	Z	λ (%)	Température d'ébullition (K)	
H_2	1	0,84	20,4	
D_2	1	2,13	23,6	
N_2	7	5,52	77,4	
O_2	8	7,52	99,2	

TAB. 2.3 – Les cibles utilisées dans l'expérience HARP.

La liste des cibles [18] est donnée par le tableau 2.3. Les cibles solides ont été séparées en deux catégories :

- Les cibles minces. On utilise des cibles de 2 et/ou 5% de longueur d'interaction (λ) en vue d'étudier le processus fondamental d'une interaction proton-noyau cible. Dans le tableau 2.3, seules les cibles avec 2% de longueur d'interaction sont représentées.
- Les cibles épaisses (une longueur d'interaction). Elles sont destinées à étudier la situation réelle rencontrée dans les faisceaux classiques de neutrinos où la ré-interaction des particules secondaires au sein de la cible joue un rôle.

Des données ont également été prises avec le porte-cible mais sans cible, pour évaluer les systématiques du système de déclenchement.

Toutes ces cibles ont une face amont plane et une face aval en forme de calotte sphérique afin de minimiser les écarts de parcours des particules incidentes dans la cible selon leur angle.

Un jeu de cibles cryogéniques reflétant les composés de l'atmosphère terrestre, a également été utilisé principalement pour l'étude du flux de neutrinos atmosphériques. Ces cibles cryogéniques sont dotées d'un autre support équipé d'un doigt froid relié à un cryostat. Leur préparation demande beaucoup plus de temps avec en moyenne de 4 à 6 heures pour le remplissage de la cible contre seulement une heure pour la vider. Là encore des données ont été prises également avec le support et la cible vide. Le tableau 2.3 reporte les températures d'ébullition de ces cibles.

2.2.4 Un détecteur en 2 parties

Dans une interaction proton-cible, nous avons affaire à des traces produites dans toutes les directions mais avec une majorité de traces vers l'avant étant donnée l'asymétrie

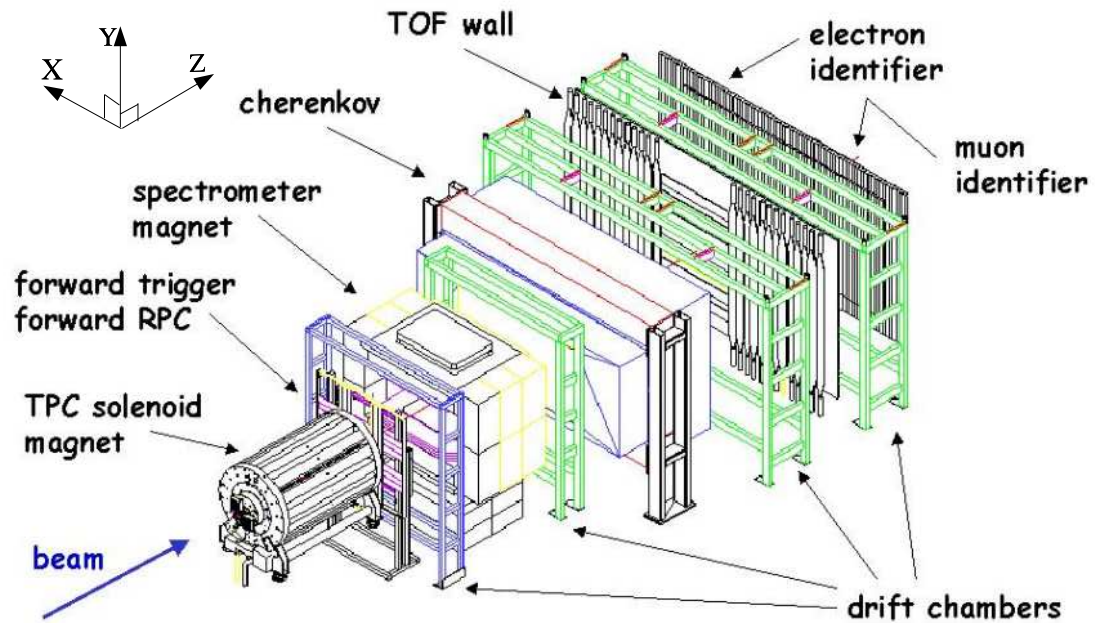


FIG. 2.9 – L'expérience HARP et ses sous-détecteurs.

de la réaction dans le référentiel de l'expérience.

La source principale d'incertitude pour les mesures anciennes réside dans le manque de couverture en angle solide, qui obligeait à mettre en oeuvre des extrapolations de données. HARP a donc voulu concevoir un détecteur autour de la cible de telle sorte que l'on ait une acceptance de pratiquement 4π stéradians. (Cf. figure 2.9). Le repère orthonormé permettant de s'orienter dans l'expérience est tel que la direction en Z est suivant le faisceau incident, l'axe Y est suivant la verticale et la direction X complète le trièdre de façon à être direct. L'origine est située au niveau de la cible.

Les éléments de détection se divisent en 2 groupes selon qu'ils s'attachent à mesurer les traces aux grands angles (spectromètre arrière) ou aux petits angles (spectromètre avant). Dans ces 2 cas, on combine détecteurs de traces et identificateurs de particules.

2.2.4.1 Spectromètre arrière : La TPC (Chambre à Projection Temporelle)

Présentation

L'élément clé du spectromètre arrière est une chambre à projection temporelle cylin-

drique (1541 mm de long et 784 mm de diamètre) à l'intérieur de laquelle est insérée la cible. La TPC a 2 fonctions :

- Elle joue le rôle de trajectographe pour les particules de faibles énergies (reconnaissance de trace et mesure de l'impulsion).
- Elle permet d'identifier les particules par le coefficient de perte d'énergie linéique $\frac{dE}{dX}$

2.11
La TPC est composée de 3972 pads de mesure (figure 2.10) répartis sur 20 rangées circulaires juxtaposées de manière concentrique.

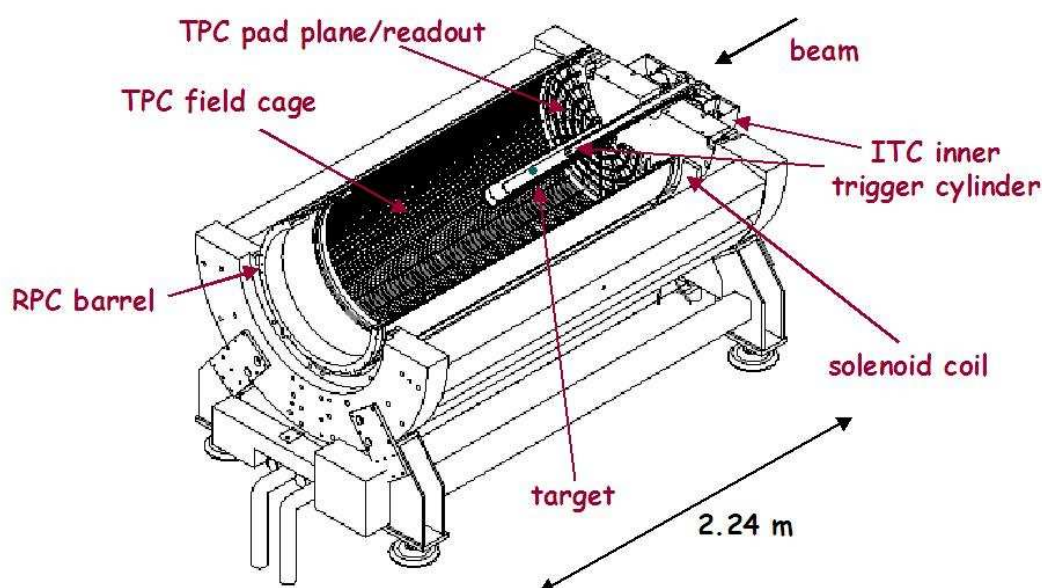


FIG. 2.10 – La TPC de l'expérience HARP : la cible est placée dans la TPC à l'intérieur de l'ITC à l'aide d'un support.

Une particule chargée traversant le volume de gaz (91% Ar - 9% CH_4) arrache des électrons d'ionisation le long de sa trajectoire. Ces électrons dérivent ensuite, grâce au champ électrique, parallèlement à l'axe de la chambre. Lorsqu'ils arrivent sur la face avant, ce signal est amplifié par les anodes et induit une charge sur les pads en regard.

Acceptance géométrique

Dans le programme de reconstruction (encore préliminaire), une trace est considérée quand sa projection sur le plan de pads touche au moins 7 pads dans la direction radiale.

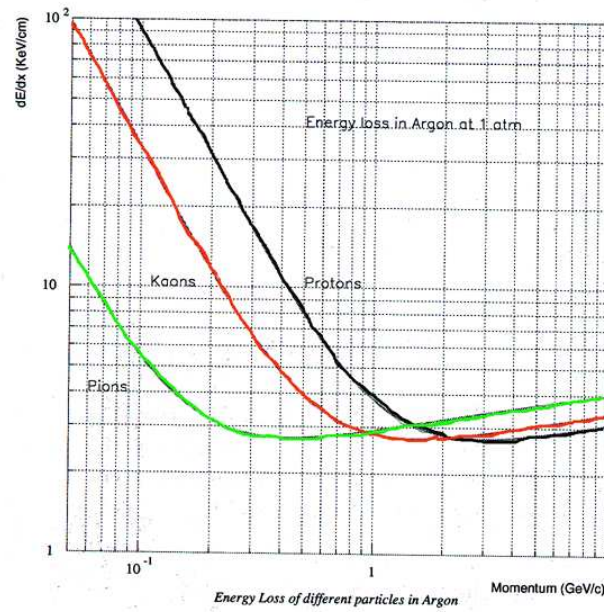
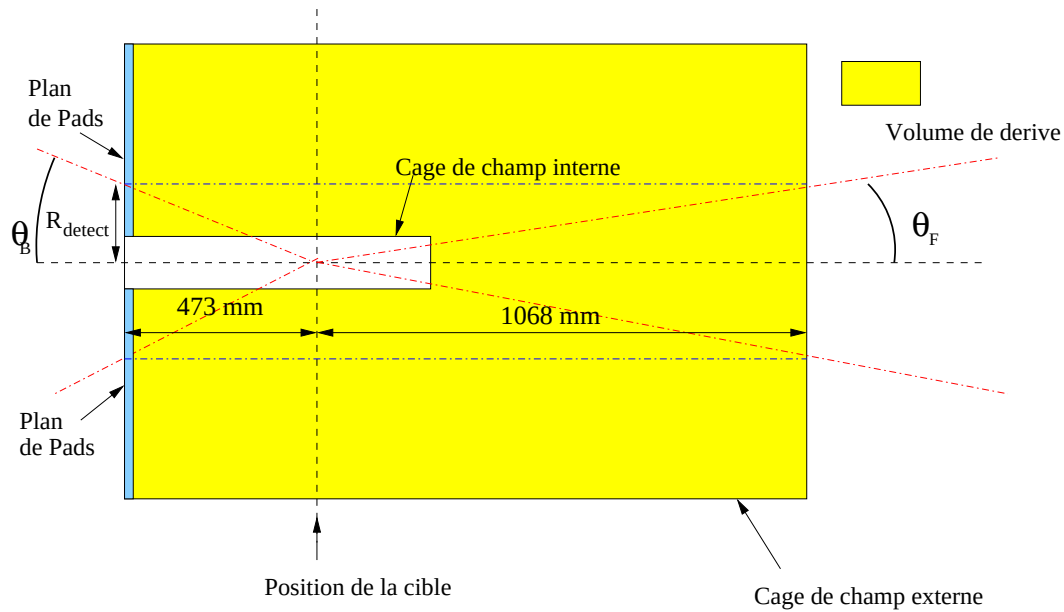
FIG. 2.11 – dE/dX en fonction de la quantité de mouvement dans l'Argon.

FIG. 2.12 – Angles critiques de détection dans la TPC.

On peut alors calculer l'acceptance angulaire de la TPC. On appelle R_{detect} la distance radiale minimale pour une trace :

$$R_{detect} = R_{RCI} + 7 \times L_{pad} = 77mm + 7 \times 15,5mm = 185,5mm \quad (2.3)$$

où R_{RCI} est le rayon central interne de la cage interne et L_{pad} la longueur d'un pad.

Donc les angles d'émission minimaux vers l'avant (θ_F) et vers l'arrière (θ_B) pour qu'une particule soit détectée sont (figure 2.12) :

$$\theta_F = \arctan\left(\frac{185,5}{1068}\right) \sim 9,8^\circ \quad (2.4)$$

$$\theta_B = \arctan\left(\frac{185,5}{473}\right) \sim 21,4^\circ \quad (2.5)$$

Mesure de l'impulsion.

La TPC est insérée dans un aimant solénoïdal fournissant un champ de 0,7 Tesla. Cet aimant utilisé dans les tests du prototype de la TPC d'ALEPH a dû être modifié : pour supprimer le bouchon aval de retour de fer, sans introduire de distorsions du champ dans la TPC, la bobine a été rallongée de 40 cm. L'angle d'émission minimal vers l'avant peut se traduire par une condition sur le P_T :

$$P_T = \cos \theta_F \cdot P_{tot} = \frac{1}{5,8} P_{tot} \quad (2.6)$$

Par ailleurs si on considère que la quantité de mouvement ne sera correctement mesurée que si le rayon de courbure est supérieur à $R_{detect}/2$, on peut estimer la quantité de mouvement transverse minimale à partir de

$$R(m) < \frac{p_T(GeV/c)}{0,3 \times B(T)} \quad (2.7)$$

pour obtenir

$$p_{Tmin} \sim 19,5 MeV/c \quad (2.8)$$

Aux grands angles, l'identification des électrons par rapport aux pions effectuée par dE/dX est limitée aux très basses énergies. Nous verrons plus loin comment compenser cette perte d'information par l'intermédiaire des RPC.

Identification par dE/dX

La figure 2.11 nous montre le dE/dX des pions, kaons et protons dans l'argon gazeux à pression atmosphérique en fonction de l'impulsion. Un champ électrique de $E = 120$ V/cm fait dériver les électrons dans le mélange gazeux à une vitesse de $5,2$ cm/ μs . La charge induite sur les pads est échantillonnée à une fréquence d'horloge de 10 MHz : on obtient alors un échantillonnage en coupe transverse du dépôt de charge primaire tous les $0,52$ cm, ce qui nous donne un nombre maximum de 205 échantillonnages pour une trace partant de la cible. Le nombre d'échantillon est minimal pour des traces à grand angle et tend vers 1 pour des traces perpendiculaires au faisceau (car alors parallèle au plan de mesure (pads)).

La résolution sur le dE/dX peut être obtenue par la formule empirique donnée par Lehraus [19].

$$\frac{\Delta dE/dx}{dE/dx} = \frac{13,5}{2,35} (N.l(\mu m))^{-0,37} \quad (2.9)$$

où N est le nombre de points par trace et l la longueur radiale d'un pad.

Pour une longueur de pad de 15,5 mm et une trace touchant 20 pads, nous obtenons une résolution d'environ 5,5%. Dans le cas minimal de 7 pads touchés, la résolution n'atteint que 8%.

Système d'étalonnage.

La qualité de l'identification repose sur la mesure de la charge collectée sur les pads, charge reliée à la perte d'énergie de la particule. Un étalonnage de la TPC est donc nécessaire et repose sur trois techniques : les rayons cosmiques, le laser, et le krypton.

- Les rayons cosmiques, majoritairement des muons dont le spectre en quantité de mouvement est connu, sont sélectionnés grâce à un système de trigger utilisant des RPC diamétralement opposés. La reconstruction des traces permet d'effectuer l'étalonnage des pads de lecture notamment en ce qui concerne l'égalisation du gain¹ d'un pad à l'autre.
- L'étalonnage au laser (figure 2.13) permet de dresser une carte détaillée des inhomogénéités du champ électrostatique dans l'espace de dérive. La lumière laser est distribuée sur 198 fibres optiques aux terminaisons aluminisées sur une épaisseur de 100 μm qui émettent des photons-électrons qui vont traverser tout l'espace de dérive. La mesure du temps d'arrivée permet donc de tester le champ électrostatique.
- l'utilisation du krypton (figure 2.14) vise à étalonner les signaux digitisés sur les pads en une échelle absolue d'énergie du dépôt par ionisation. Par périodes, du krypton radioactif est introduit dans le gaz de la TPC. Il se désexcite suivant un spectre bien connu de 9 à 42 keV (figure 2.15) [20] dû à des conversions internes d'électrons, des électrons Auger et des photons. Ces désexcitations se traduisent par des dépôts d'énergie très localisés qui permettent donc l'étalonnage de la TPC.

Lors du fonctionnement et de l'étalonnage de la TPC, des anomalies ont été détectées dont la plus importante est la présence d'un diaphone (cross-talk) dans la carte mère de la TPC. La conséquence pourrait être une limitation de la résolution en $r - \theta$ mais une correction logicielle est en cours d'étude.

2.2.4.2 Les RPC (Resistive Plate Chambers)

Pour remédier à la faiblesse d'identification de la TPC à grand angle (faible échantillonnage) et à la difficulté de séparer électrons et pions entre 100 et 250 MeV par $\frac{dE}{dX}$, des RPC ont été installées autour de la TPC pour mesurer le temps de vol des particules traversantes.

Les RPC (figure 2.16) sont présentées en 2 blocs distincts :

- une partie constitue un cylindre qui occupe l'espace restant de 30 mm entre la face externe de la TPC et la face interne de l'aimant solénoïdal : 24 compteurs RPC, de 11 mm d'épaisseur, 2178 mm de long et 140 mm de large couvrent presque complètement la TPC. Il reste cependant 2 zones non-couvertes à ± 35 degrés par rapport à l'axe vertical à cause des pieds de la TPC.

¹proportionnalité entre la charge induite sur les pads de lecture et les électrons primaires produits au passage d'une particule chargée

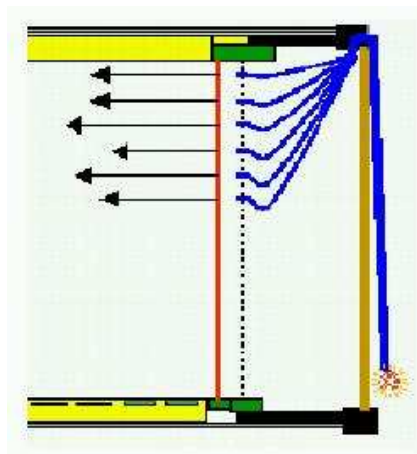


FIG. 2.13 – système d'étalonnage au laser de la TPC : des fibres optiques guident la lumière vers des couches d'aluminium situées après la membrane de haute tension (en pointillé).

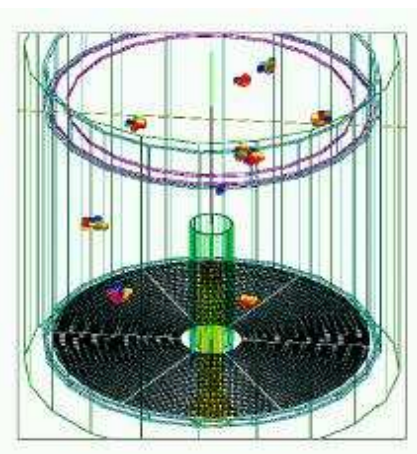


FIG. 2.14 – Visualisation du dépôt d'énergie consécutif à la désexcitation du krypton. Les dépôts d'énergie sont localisés en cluster couvrant généralement 5 pads.

- une autre partie est installée en aval de la TPC juste devant l'hodoscope de trigger (FTP) sous la forme de 2 ensembles de 4 compteurs disposés horizontalement et 2 autres ensembles de 3 compteurs disposés verticalement. Il s'agit là essentiellement de couvrir l'angle solide vers l'avant : les particules qui quittent la TPC par sa face aval.

La résolution temporelle atteinte par les RPC de HARP est de 150 ps, résolution mesurée sur les cosmiques.

2.2.4.3 Le Spectromètre avant

Pour les particules émises vers l'avant dans la cible (en général des particules de haute énergie), la TPC est peu efficace à la fois dans la mesure de la quantité de mouvement et pour l'identification. Le spectromètre avant est là pour combler cette lacune : il est constitué d'un ensemble de chambres à dérive pour reconstituer la trajectoire des parti-

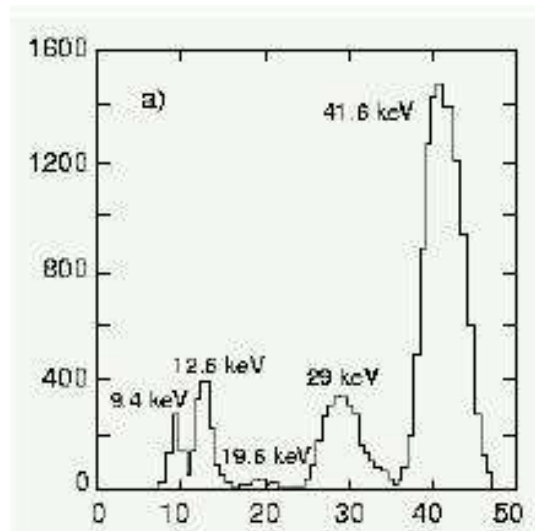


FIG. 2.15 – Spectre de la désexcitation du krypton (en keV).

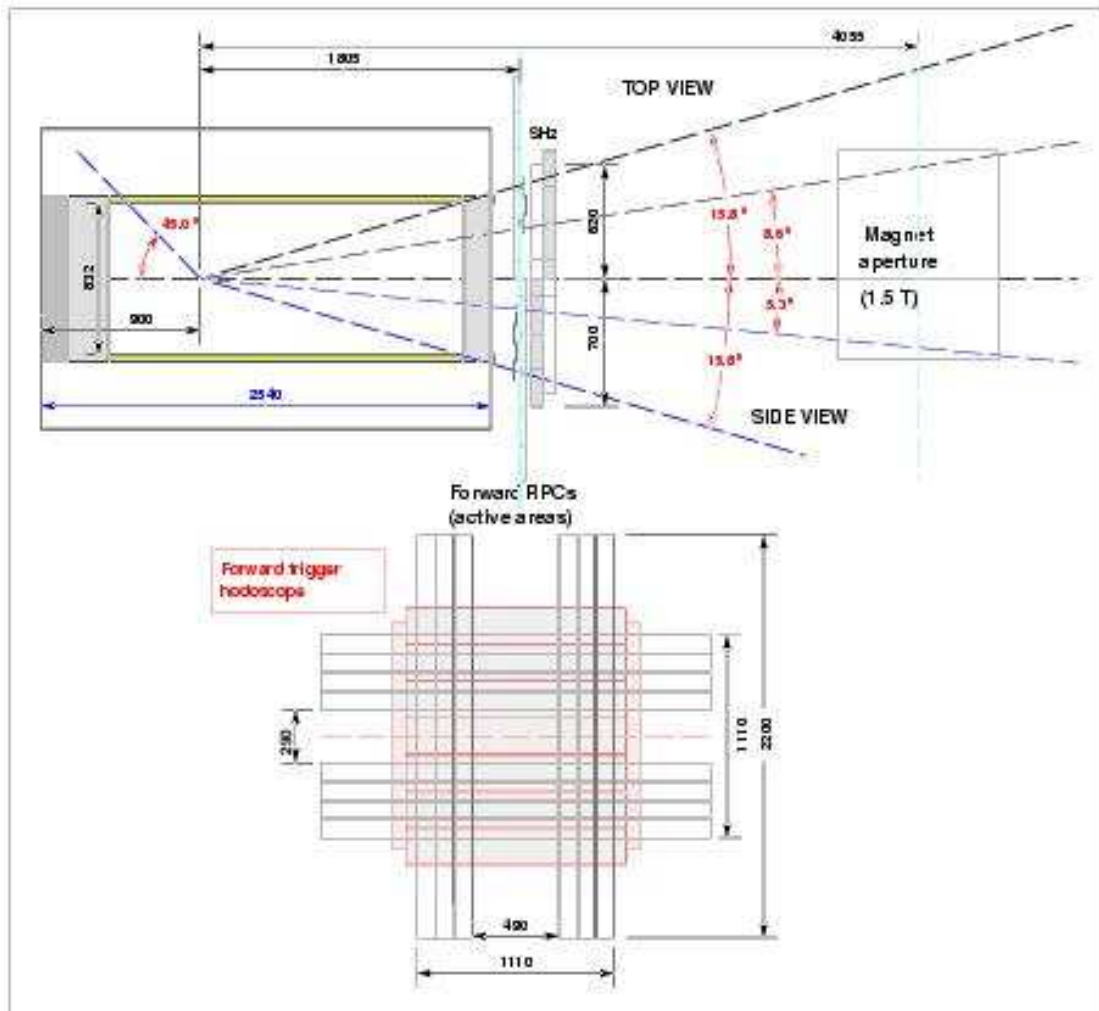


FIG. 2.16 – Disposition des RPC après la TPC suivant l'ouverture du dipôle.

cules chargées. Ces chambres encadrent un aimant dipolaire pour la mesure des impulsions

suivant une configuration (1+4) c'est-à-dire avec un module avant le dipôle et 4 après. Cette asymétrie s'explique par l'angle solide de la couverture avant (cet aspect est repris dans la partie consacrée aux chambres à dérive). L'identification est assurée par un Tcherenkov à seuil, un détecteur de temps de vol, un calorimètre électromagnétique et un identificateur de muons.

Je ne parlerai ici rapidement que des sous-détecteurs chargés de l'identification puis je consacrerai le chapitre suivant au sous-détecteur sur lequel j'ai travaillé, l'ensemble des chambres à dérive.

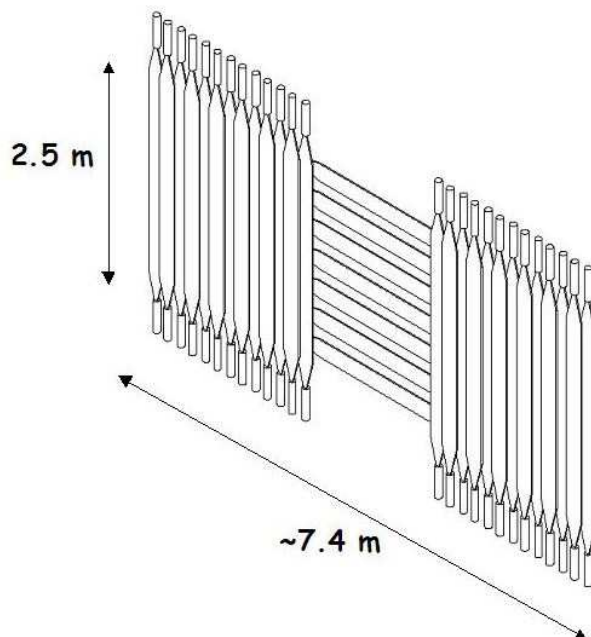


FIG. 2.17 – *Le mur du TOF.*

2.2.4.4 Le Time-Of-Flight couramment appelé TOF

Le TOF est un identificateur qui mesure le temps de parcours des particules depuis la cible et permet de les séparer suivant leur masse. Ce sous-détecteur est un mur de lattes de scintillateurs de $2,5 \text{ m} \times 20 \text{ cm} \times 1 \text{ cm}$, disposées comme indiqué sur la figure 2.17, à 10 m de la cible. Chaque latte est équipée de 2 photomultiplicateurs. Pour assurer une mesure précise du temps, le TOF utilise des TDC au pas de 35 ps. Il est doté d'un système d'étalonnage par laser qui permet de vérifier quotidiennement sa stabilité. La résolution temporelle du TOF est de 150 ps et permet une séparation proton-pion jusqu'à 4 GeV/C [21]. En effet, nous pouvons voir à titre d'exemple sur la figure 2.18 les différences de temps de vol concernant les couples protons-pions et pions-kaons pour la distance de vol de 10 mètres.

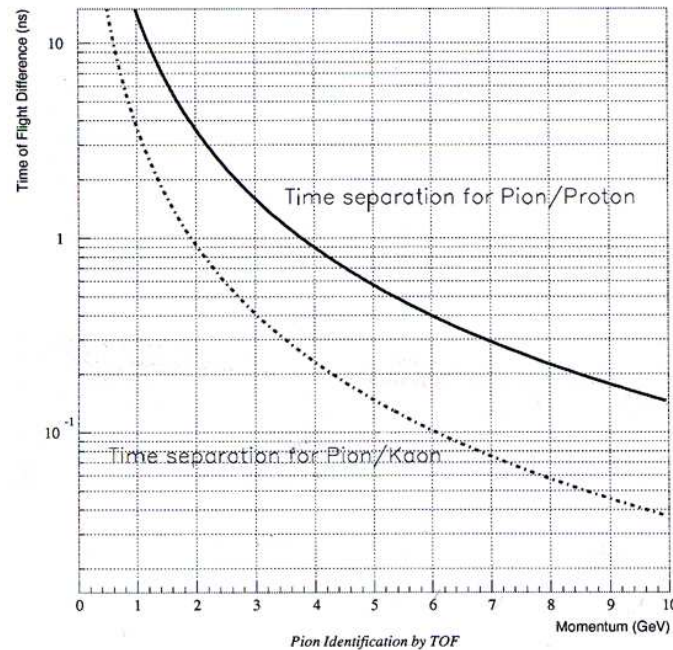


FIG. 2.18 – Différence de temps de vol en fonction de la quantité de mouvement.

2.2.4.5 Le Tcherenkov

Aux énergies plus hautes que 3-4 GeV, c'est le détecteur à effet Tcherenkov qui permet d'identifier les pions. Il faut rappeler que la mesure de la section efficace de production des pions est un des objectifs majeurs de l'expérience.

Ce détecteur repose sur la capacité à détecter les photons émis par une particule allant plus vite que la vitesse de la lumière dans un milieu (gaz). Pour chaque type de particule, il existe un seuil au dessus duquel il y a émission de photons (figure 2.20). Le nombre de photons produits dépend également de la vitesse selon $N_{photons} \propto (1 - \frac{1}{\beta^2 n^2})$ quand la vitesse β est au-dessus du seuil.

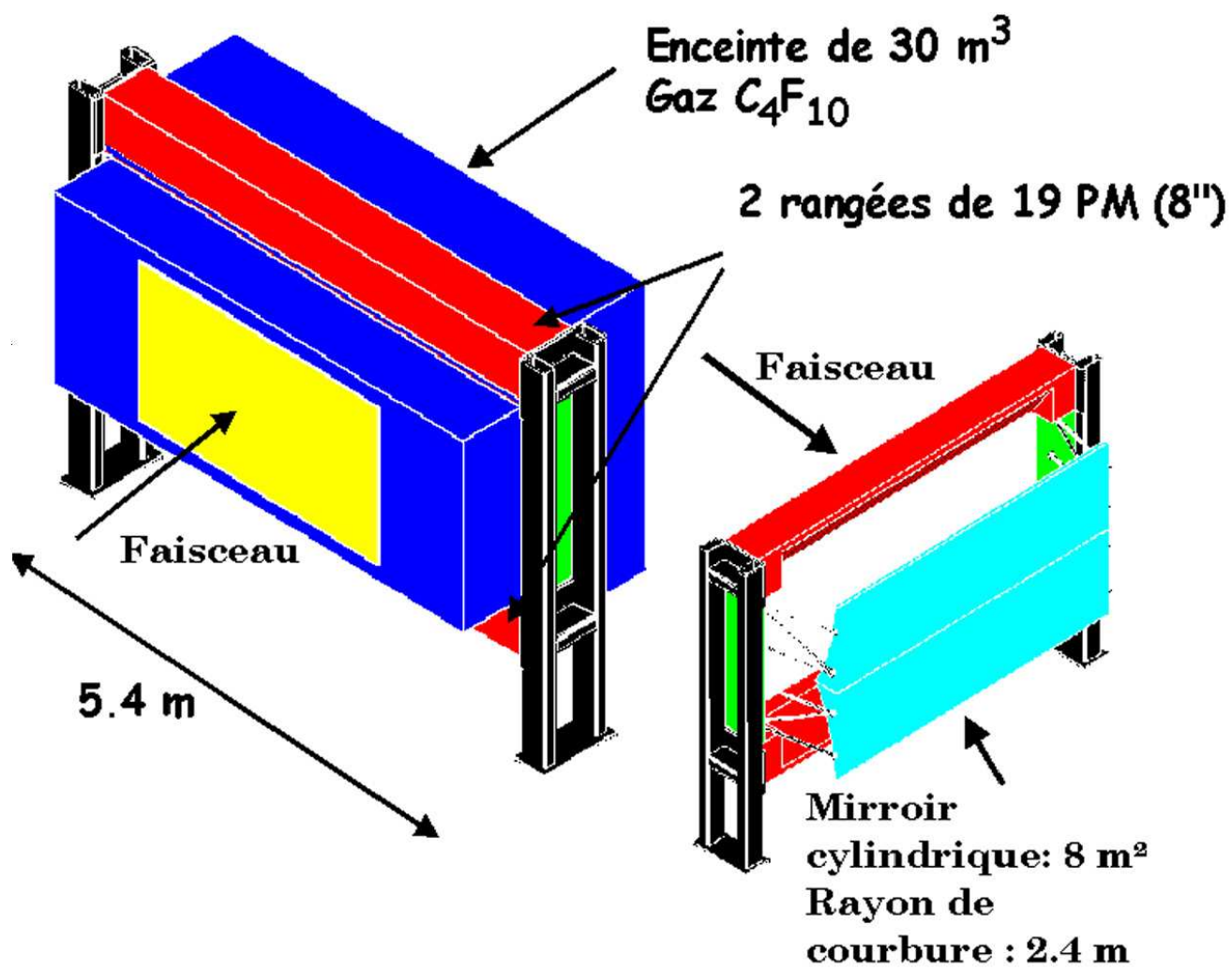
Le Tcherenkov de HARP est une enceinte parallélépipédique de 30 m³ de gaz C₄F₁₀ équipée de miroirs cylindriques (figure 2.19), qui réfléchissent les photons Tcherenkov vers 2 rangées de photomultiplicateurs équipés de cônes de Winston.

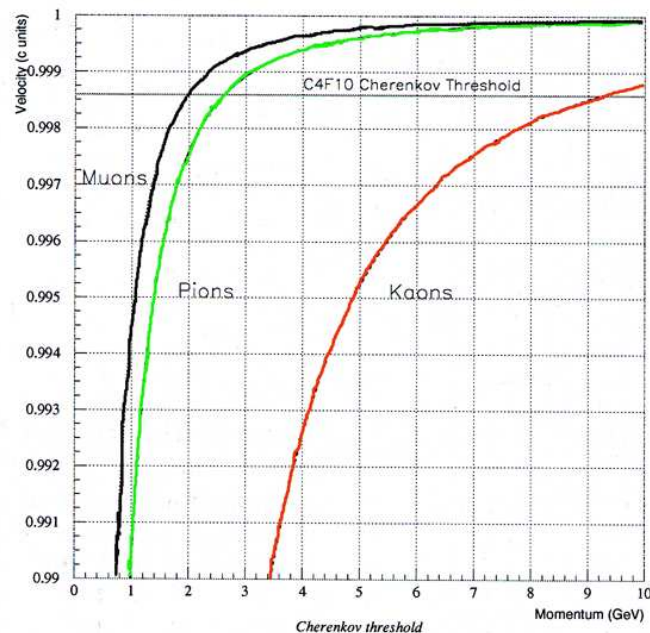
On voit donc sur la figure 2.20 que parmi les hadrons, seuls les pions donneront de la lumière entre 2 et 9 GeV et seront identifiés en corrélant la mesure de la quantité de mouvement et de la direction par les chambres à dérive, avec la lumière Tcherenkov.

On peut voir sur une simulation réalisée pour la proposition d'expérience [22] (figure 2.21) comment les divers détecteurs permettent une bonne séparation π /proton sur l'ensemble de l'espace des phases p_L (quantité de mouvement longitudinal) vs p_T (quantité de mouvement transverse). Les carrés représentent la densité d'événements simulés.

2.2.4.6 Les identificateurs d'électrons et de muons

Le dernier sous-détecteur est un identificateur de muons situé derrière un bloc de fer de 32 cm d'épaisseur. Son rôle est surtout d'identifier les muons pour les exclure de l'analyse. Pour séparer les électrons des pions et pour identifier les photons et donc remon-

FIG. 2.19 – *Le Tcherenkov de HARP.*

FIG. 2.20 – *Beta vs quantité de mouvement.*

ter aux π_0 , un calorimètre électromagnétique a été installé juste avant l'identificateur de muons : il s'agit d'un calorimètre de type spaghetti (sandwich plomb-fibres scintillantes), récupéré de l'expérience CHORUS [23]) constitué d'une partie électromagnétique de 5 longueurs d'interaction (X_0) et une partie hadronique de $10 X_0$. Les fibres verticales, sont groupées et lues aux 2 bouts par des photomultiplicateurs dont les signaux sont lus par des QDC ("charge" to digital converter). Entre le TOF et le calorimètre, une plaque de fer de 2 cm d'épaisseur permet aux photons de commencer leur gerbe, dont la position est repérée dans 3 chambres à dérive qui servent donc de détecteur de pied de gerbe avant le calorimètre.

2.2.5 Programme de mesure

Les tableaux 2.4 et 2.5 montrent le nombre d'événements enregistrés pour chaque énergie incidente et chaque cible considérée (les quantités de mouvement positives (négatives) concernent les particules incidentes de charge positive (négative)). Ces données ont été prises au cours des années 2001 et 2002. Pour les données prises avec les répliques des cibles utilisées par les expériences K2K et MiniBooNE, l'énergie nominale de leur faisceau a été utilisée : 12,9 GeV/c (K2K) et 8 GeV/c (MiniBooNE : tableau 2.6).

En été 2002, des prises de données ont été effectuées avec les cibles cryogéniques dont les évènements sont répertoriés dans le tableau 2.7 .

Je vais maintenant consacrer tout un chapitre au détecteur dont je me suis le plus occupé durant ma thèse, les chambres à dérive.

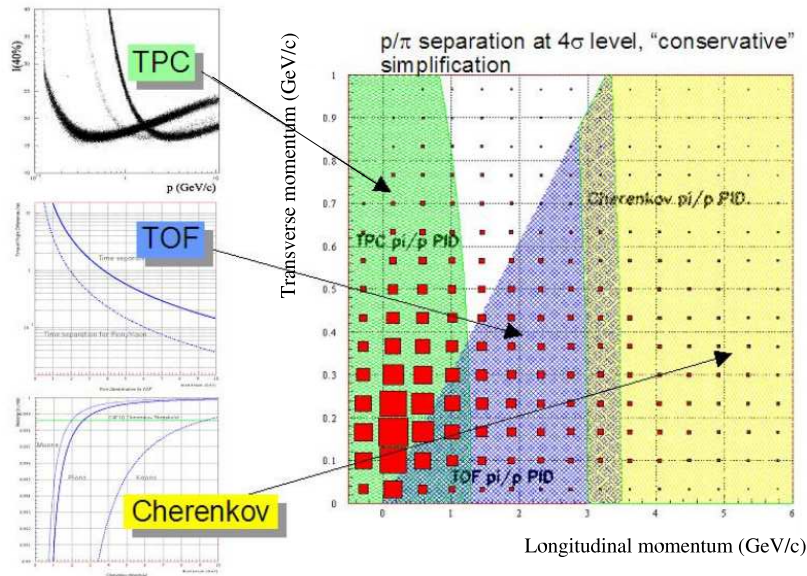


FIG. 2.21 – Distribution de pions suivant le plan $P_{\text{transverse}}-P_{\text{longitudinale}}$ pour un faisceau de 15 GeV/c de proton sur une cible mince de Beryllium. On peut voir la complémentarité des sous-détecteurs (TPC, Tcherenkov, et TOF) pour couvrir la majorité de l'espace des phases pour distinguer les pions et les protons.

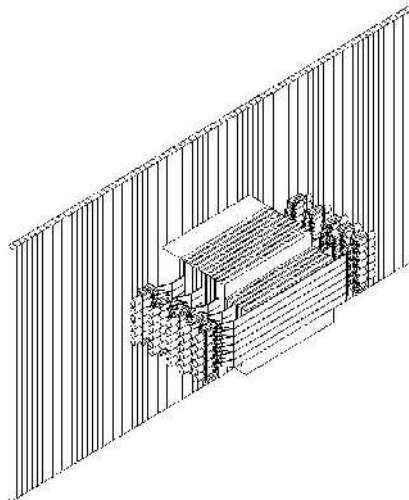


FIG. 2.22 – Le calorimètre électromagnétique et l'identificateur de muons.

TAB. 2.4 – liste des événements enregistrés en 2001 pour chaque configuration. (en millions d'évènements).

Cible	+3 GeV/c	+5 GeV/c	+8 GeV/c	+12 GeV/c	+15 GeV/c
Be mince	2,2	2,3	—	5,9	1,8
C mince	2,4	—	—	—	2,7
C épaisse	—	—	—	—	—
Al mince	2,4	2,1	—	3,8	2,3
Cu mince	3,6	1,5	—	2,0	3,0
Cu épaisse	—	—	—	—	—
Sn mince	1,3	—	—	—	2,2
Ta mince	4,9	1,5	—	1,6	2,6
Ta épaisse	1,1	—	—	—	—
Pb mince	4,4	1,5	—	2,5	2,3
Pb épaisse	1,0	—	—	—	—
vide	1,0	1,2	—	1,4	0,1
Cible	-3 GeV/c	-5 GeV/c	-8 GeV/c	-12 GeV/c	-15 GeV/c
Be mince	0,2	0,76	—	1,2	1,3
C mince	—	—	—	0,86	1,1
Al mince	0,27	0,79	—	0,88	0,90
Cu mince	—	—	—	1,11	0,7
Sn mince	—	—	—	—	1,2
Ta mince	0,82	0,80	—	0,6	2,1
Pb mince	0,94	0,34	—	0,7	1,4
vide	—	0,86	—	—	0,06
Cibles spéciales					
K2K at 12,9GeV /c	mince : 1,5, interm. : 2,2, épaisse : 1,4				
MiniBoone Be at 8, GeV /c	3,2				

TAB. 2.5 – Liste des évènements enregistrés en 2002 pour chaque configuration. (en million d'évènements)

Cible	+3 GeV/c	+5 GeV/c	+8 GeV/c	+12 GeV/c	+15 GeV/c
Be mince	1,54	1,95	2,23	0,63	0,91
vide	0,13	0,10	0,14	0,18	0,15
Be épaisse	1,16	1,09	1,09	1,81	0,74
C mince	1,63	3,37	2,27	2,15	0,81
vide	0,37	-	0,32	0,23	0,14
C épaisse	1,11	1,12	1,01	0,66	0,96
Al mince	1,93	2,07	2,22	0,73	0,71
vide	0,12	0,20	0,16	0,14	0,14
Al épaisse	1,34	1,29	1,27	1,15	0,70
Cu mince	1,16	2,49	2,95	0,85	0,65
vide	0,16	0,14	0,14	0,14	0,15
Cu épaisse	1,12	1,58	1,40	1,93	0,83
Sn mince	1,87	3,25	3,15	2,05	0,63
vide	0,21	0,14	0,17	0,14	0,15
Ta mince	2,62	2,38	2,33	1,04	0,62
vide	0,15	0,10	0,20	0,15	0,08
Ta épaisse	1,01	1,33	1,49	1,56	0,61
Pb mince	2,49	3,41	3,32	0,72	0,64
vide	0,10	0,10	0,16	0,08	0,09
Pb épaisse	0,76	1,26	0,98	1,12	1,10
Cible	-3 GeV/c	-5 GeV/c	-8 GeV/c	-12 GeV/c	-15 GeV/c
Be mince	2,33	1,39	1,76	1,21	0,85
vide	0,17	-	0,05	0,06	0,04
C mince	2,26	1,81	1,63	0,74	0,62
vide	0,24	0,10	0,07	-	0,05
Al mince	2,17	1,11	1,38	0,75	0,62
vide	-	-	0,7	0,09	0,05
Cu mince	3,20	1,23	2,28	0,85	1,06
vide	0,08	0,24	-	-	0,03
Sn mince	2,08	1,82	1,58	1,35	0,27
vide	-	-	0,07	0,08	-
Ta mince	1,66	1,60	1,38	1,03	0,40
vide	0,25	-	0,07	0,08	-
Pb mince	1,54	2,33	1,65	1,92	1,08
vide	0,11	-	0,10	0,06	0,04
Cibles spéciales					
K2K at 12,9GeV /c		mince : 3,4 , medium : 3,1, replica : 3,67			
MiniBoone Be at 8,9 GeV /c		mince : 7,3 , interm. : 5,17, épaisse : 2,84 replica : 4,05			

TAB. 2.6 – Liste des évènements enregistrés pour la cible de MiniBooNE en 2002 (en millions d'évènements) .

Cible	GeV/c	Syst. de déclenchement	Nombre d'évènements
Be mince	8,9	mince>8	7,31
Vide	8,9	mince>8	0,51
Be 50%	8,9	beam 14	5,17
Vide	8,9	beam14	0,36
Be 100%	8,9	beam 14	0,7
sans TPC	8,9	beam 14	2,3
TPC	8,9	beam 14	0,6
TPC	8,9	épaisse	2,0
sans TPC	8,9	épaisse	0,46

TAB. 2.7 – Liste des évènements enregistrés sur les cibles cryogéniques. (en millions d'évènements)

Target	+3 GeV/c	+5 GeV/c	+8 GeV/c	+12 GeV/c	+15 GeV/c
N-7	2,9	1,95	0,75	1,4	1,4
vide	0,41	0,20	0,15	0,22	0,23
O-8	1,63	3,37	2,27	2,15	0,81
vide	-	-	-	-	0,38
D-1	3,51	2,50	2,15	1,84	2,8
vide	1,35	1,14	0,90	0,50	0,43
H-1	5,28	3,24	2,13	2,10	0,88
vide	.39	-	-	0,49	-
Target	-3 GeV/c	-5 GeV/c	-8 GeV/c	-12 GeV/c	-15 GeV/c
N-7	0,62	0,44	0,61	0,39	0,69
vide	0,14	0,15	0,15	0,07	0,10
O-8	1,48	1,11	0,64	0,41	0,91
vide	-	-	-	-	0,37
D-1	0,55	0,49	.43	0,75	0,49
vide	0,04	0,25	0,24	0,30	0,26
H-1	.871	.697	.578	.613	.603
vide	0,37	-	-	-	-

Chapitre 3

Les Chambres à dérive

Les chambres à dérive de l'expérience HARP ont été récupérées de l'expérience NOMAD [24][1] (elles ont été conçues et fabriquées par le CEA de Saclay). Ce sont des chambres d'une bonne précision ($150\text{ }\mu\text{m}$ dans la direction de dérive en moyenne dans NOMAD) et de grande taille ($3\text{ m} \times 3\text{ m}$). La plupart des membres de l'équipe ayant en charge ces chambres les connaissaient déjà pour avoir fait partie de la collaboration NOMAD. Ces chambres présentent un certain nombre d'inconvénients :

- L'inconvénient principal est que ces chambres ne sont pas transparentes. Nous verrons qu'elles contiennent beaucoup de matière.
- Dans NOMAD, ces chambres avaient des fils horizontaux, parallèles au champ magnétique permettant des mesures de précision dans la direction perpendiculaire. Le champ magnétique de HARP étant vertical, il a fallu modifier le système de suspension mécanique des chambres pour avoir des fils verticaux et mesurer précisément les déflexions horizontales (X) dues au champ magnétique du spectromètre.
- Les normes de sécurité du CERN étant devenues plus strictes, le mélange gazeux autrefois utilisé dans NOMAD (argon-éthane) est devenu proscrit et il a donc fallu trouver et utiliser un nouveau mélange moins efficace. Le système de mélange et de distribution de gaz a donc dû être entièrement refait.
- les TDC (convertisseur temps numérique) devaient être dans la norme VME et non plus dans la norme FASTBUS dans un souci d'harmonisation avec les matériels utilisés par les autres sous-détecteurs. Ce qui veut dire que l'acquisition des chambres a dû être entièrement réécrite.

L'objectif en utilisant ces chambres est d'en tirer la meilleure résolution spatiale et cela afin d'obtenir une très bonne résolution en impulsion. En effet plus on connaîtra la position des traces avec précision de part et d'autre du dipôle et plus la détermination du rayon de courbure de la trajectoire sera précise. La conséquence immédiate est une estimation plus précise de la vitesse de la particule. Pour cela, il faut contrôler le mieux possible chacun des paramètres du fonctionnement de ces chambres.

Ce chapitre illustrera la stratégie adoptée quant à l'utilisation de ces chambres et leur fonctionnement. La procédure d'alignement des fils qui permet d'optimiser la résolution spatiale des chambres sera présentée dans le chapitre 4.

La disposition des chambres.

Les chambres à dérive sont l'unique détecteur de trace du spectromètre magnétique pour les particules dites "allant vers l'avant". La stratégie est de déterminer l'énergie d'une particule à l'aide du rayon de courbure obtenue lors de la déviation d'une particule chargée par le dipôle (figure 3.1).

Pour la partie placée avant le dipôle, nous avons un module qui doit assurer le repérage des traces sortant de la TPC. Après le dipôle, nous avons 4 modules qui sont chargés de couvrir une région d'acceptance plus importante. Les modules 3 et 4 sont disposés côte à côte à cet effet et le module 5 est placé devant ces derniers afin de couvrir la région centrale.

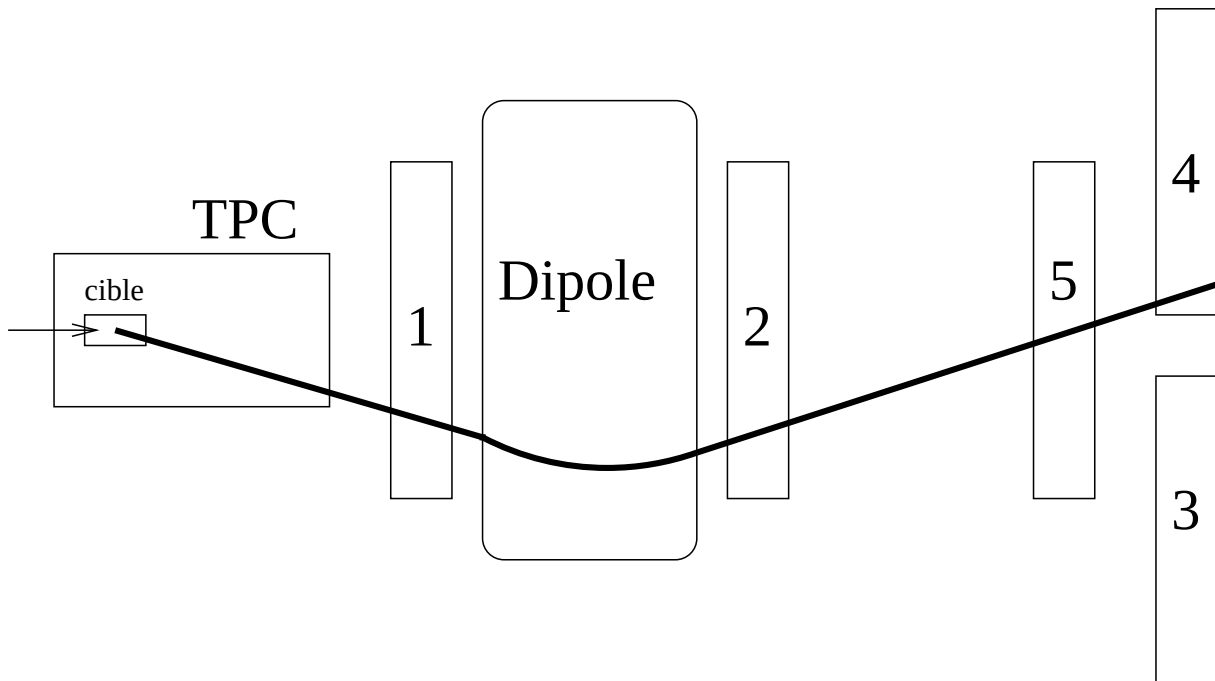


FIG. 3.1 – *Disposition des chambres.* Le module 1 doit s'efforcer de repérer une trace en provenance de la TPC. Les autres modules sont disposés de manière à augmenter l'acceptance après le dipôle.

Description des chambres.

Les chambres à dérive sont au nombre de 23 regroupées en 5 modules de 4 chambres. Les 3 dernières chambres restantes servent de détecteur de pied de gerbe (voir partie 2.2.4.6) et n'étaient pas initialement prévues : elles ont été rajoutées tardivement, mais sans modification mécanique.

Une chambre a la composition suivante (la structure d'une chambre est représentée sur la figure 3.2 :

- 3 volumes de dérive remplis du mélange gazeux qui sera ionisé au passage d'une particule chargée. Ils représentent notre milieu sensible, et contiennent les fils sensibles verticaux.
- Ces volumes sont séparés par 4 panneaux d'aramide en structure de nid d'abeille renforcés par 2 fines couches de kevlar-epoxy. Cette structure permet aux chambres d'être autoporteuses et d'éviter les lourds cadres métalliques classiquement utilisés pour supporter la tension des fils.

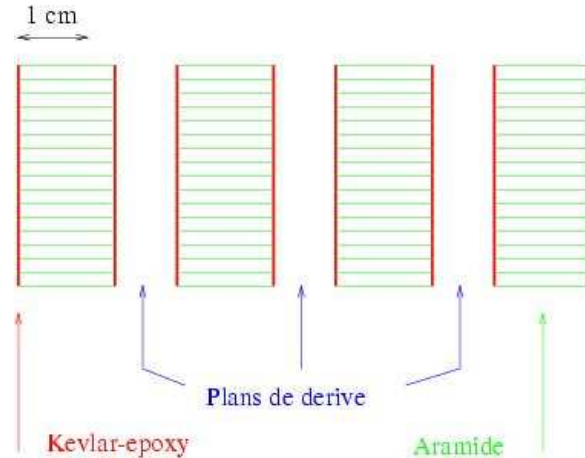


FIG. 3.2 – Structure d'une chambre à dérive.

3.1 Description et principe d'une cellule de dérive

Structure d'une cellule

Chaque plan de dérive contient une succession de fils d'anode et de cathode alternés, délimitant des cellules de dérive. Dans chaque chambre, l'orientation des fils par rapport à la verticale est de $+5$, 0 , et -5 degrés respectivement pour les 3 plans successifs. Les angles d'inclinaisons sont faibles ce qui est un inconvénient pour la résolution en Y , néanmoins les conditions de NOMAD recommandaient d'avoir les amplificateurs d'un même côté d'une chambre (figure 3.3). Ces orientations différentes permettent une vue "stéréoscopique" et donc de pouvoir effectuer une reconstruction en 3 dimensions de la trajectoire des particules. Les plans à fils inclinés possèdent 41 fils de lecture contre 44 pour le plan central. Les dimensions d'une cellule de dérive (figure 3.5) sont de 8 mm d'épaisseur de gaz et 32 mm de longueur de dérive maximale. Elles sont délimitées par 2 cathodes ou fils de champs (en Cu-Be avec un diamètre de $100\ \mu$) qui sont mis au potentiel de -2900 V . Au centre d'une cellule, on a une anode ou fil de lecture (en tungstène doré d'un diamètre de $20\ \mu$) porté à $+1300\text{ V}$.

La conséquence immédiate de l'orientation des fils est l'existence de zones mortes (figure 3.4) pour les plans avec des fils inclinés dans les régions les plus excentrées. Dans les faits, ces zones mortes n'ont pas de conséquences importantes. La taille des chambres est surdimensionnée par rapport à la zone utile effective de l'expérience. Bien sûr le problème pour les modules (3 et 4) qui sont côte à côte se pose mais il est compensé par le module 5 couvrant les régions mortes de celles-ci.

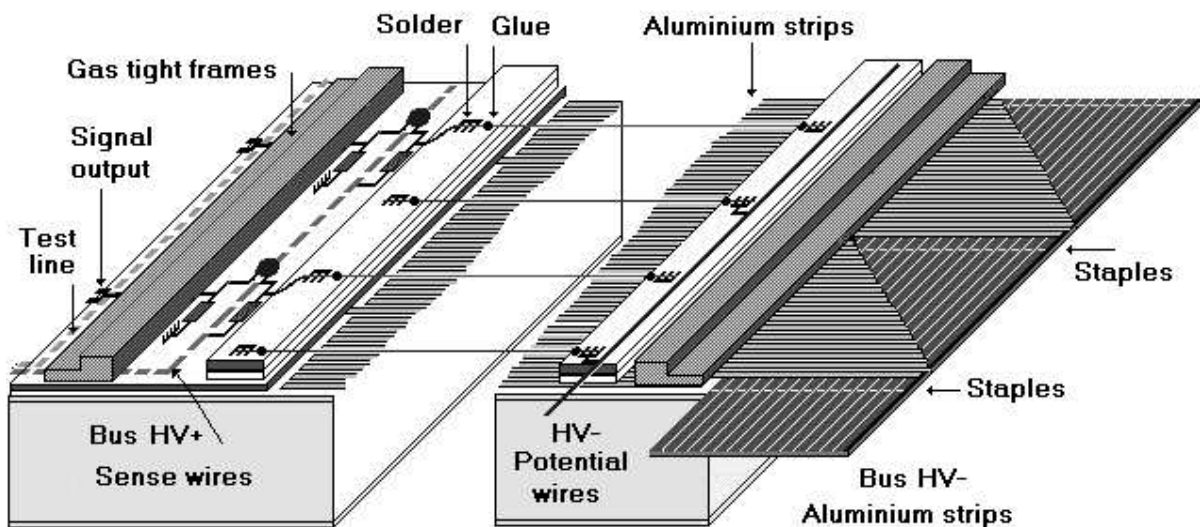


FIG. 3.3 – Coupe d'un plan X : l'électronique de lecture et la distribution du gaz sont situées à gauche (haut des chambres dans la configuration de HARP).

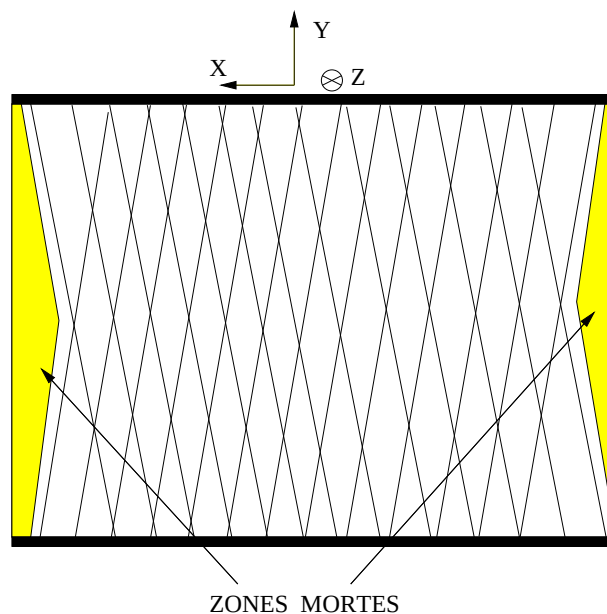


FIG. 3.4 – Zones mortes résultant de l'orientation des fils à -5 et $+5$ degrés..

Ionisation du gaz

Lorsqu'une particule chargée traverse un plan de dérive, elle interagit avec le gaz électromagnétiquement : elle ionise le gaz créant le long de sa trajectoire de multiples paires ions-électrons. Les électrons primaires vont dériver vers l'anode et vont également à leur tour créer des paires ions-électrons. En moyenne, 30 paires primaires pour 100 paires totales sont produites par centimètre.

Dérive des électrons

Les électrons dérivent vers l'anode grâce à un champ électrostatique. Le temps de dérive permet de remonter à la distance et donc à la position du passage de la particule dans la cellule de dérive. Il est important que la vitesse de dérive soit constante et donc que le champ soit uniforme. Le champ est assuré ici par de fines bandes (strips) d'aluminium placées de part et d'autre de la cellule et maintenues à des potentiels décroissants de -2900 V (face au fil de champ) à 0 V (face au fil de lecture).

On obtient ainsi un champ uniforme et constant d'environ 800 V/cm dans tout le volume de la cellule, sauf au voisinage immédiat du fil d'anode où le champ devient cylindrique.

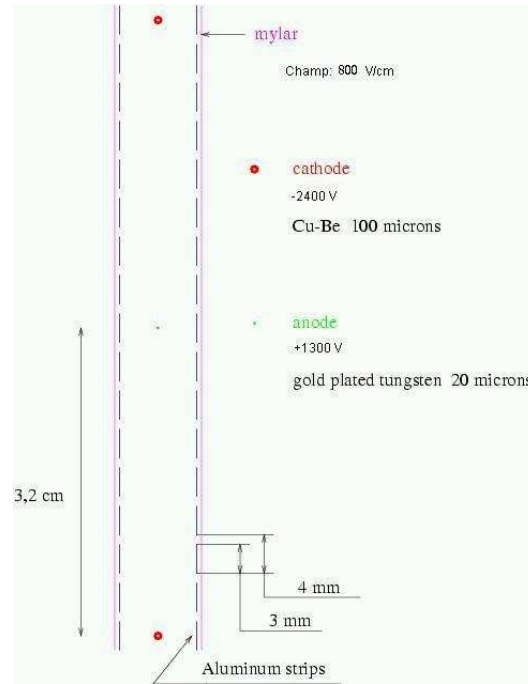


FIG. 3.5 – Cellule d'une chambre à dérive.

Avalanche et formation du signal

En s'approchant de l'anode, les électrons rencontrent un potentiel en $1/r$ et subissent une accélération violente qui leur permet d'arracher d'autres électrons pour produire un phénomène d'avalanche. Le gain de ce processus peut atteindre 10^6 : il dépend essentiellement du gaz, de la haute tension, et du diamètre du fil.

Le temps caractéristique de l'avalanche est très court mais c'est le signal induit sur l'anode par le mouvement de charges positives (plus lentes que les électrons) dans le champ électrostatique qui constitue le signal enregistrable. Ensuite ce dernier se propage sur le fil de lecture (anode) à la vitesse de 26 cm/ns pour atteindre l'électronique de lecture.

Choix du gaz

Le choix du gaz est essentiel pour un bon fonctionnement des chambres. Il s'agit en général du mélange d'un gaz rare et d'un ou plusieurs gaz polyatomiques. Le faible potentiel d'ionisation des gaz rares (11,6 eV pour l'argon) permet des gains très importants dans

l’avalanche, sans recourir à des tensions trop élevées. Par contre la production de photons UV dans la désexcitation des molécules d’argon peut conduire à de nouvelles avalanches parasites, s’ils ne sont pas absorbés : c’est le rôle du “quencher”. L’autre condition essentielle du mélange gazeux, c’est de présenter un plateau dans la distribution de la vitesse en fonction du champ électrostatique : en choisissant comme point de fonctionnement une valeur de champ sur ce plateau, on garantit une vitesse de dérive constante même en présence de petites inhomogénéités du champ.

Dans NOMAD, le mélange Ar-éthane (60 % - 40 %) avait l’avantage de présenter un très large plateau. Mais étant hautement inflammable, il est maintenant interdit d’utilisation au CERN. Par manque de temps, nous n’avons pas pu tester différents mélanges alternatifs et nous nous sommes donc rabattus sur un mélange Ar, CO_2 , CH_4 (90 %, 9 %, 1 %) assez classique dans les chambres à dérive. La figure 3.6 montre qu’il possède un “petit” plateau vers 700-800 V/cm et c’est cette valeur de fonctionnement que nous avons adoptée.

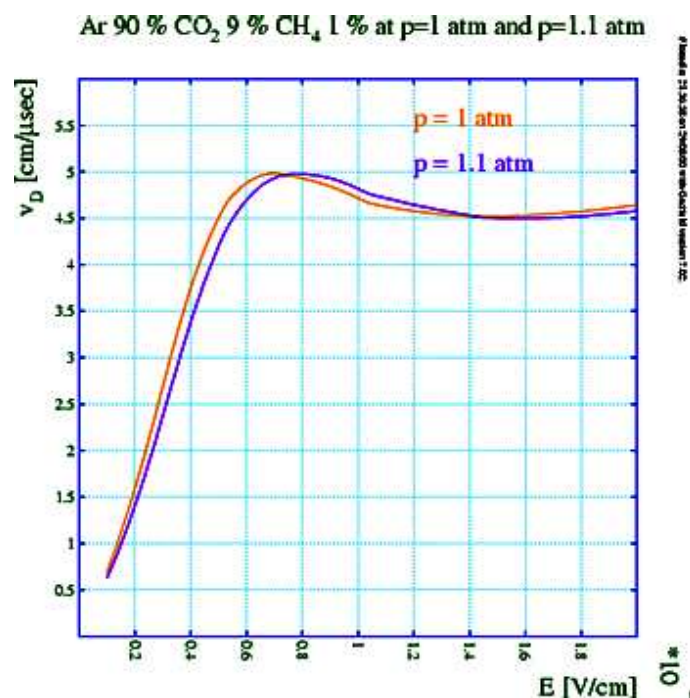


FIG. 3.6 – Variation de la vitesse de dérive en fonction du champ électrostatique pour un mélange Ar 90% - CO_2 9% - CH_4 1%

L’électronique de lecture

Un signal venant d’un fil atteint un amplificateur à travers un circuit RC qui isole le signal de la haute tension du fil de lecture. Le signal amplifié est envoyé vers un discriminateur qui transmet le signal au TDC (Time To Digital Converter). Celui-ci enregistre le temps du signal.

Le TDC enregistre de manière continue toutes les informations lui arrivant dans une fenêtre de $1 \mu s$ (donc beaucoup plus que les temps de dérive les plus importants) à partir d’un temps de référence donné par le système de déclenchement. La largeur de cette fenêtre est donnée par le programme d’acquisition. Elle est fixée en vue de limiter le temps

mort effectif lors du transfert de ses données vers l' "event builder"¹ pour les stocker sur disque. Ces TDC sont les modèles v767 fabriqués par la société CAEN [25]. Ils ont une résolution de 0,78125 ns mais la "physique" de la dérive des électrons ne permet pas de dissocier 2 signaux distants de moins d'une quinzaine de nanosecondes.

3.2 Les panneaux

Les chambres ont été conçues pour servir à la fois de cible aux neutrinos et de détecteur de traces dans l'expérience NOMAD.

Pour réaliser ce compromis, les panneaux qui séparent les plans de dérive ont les particularités suivantes :

- une faible longueur de radiation ($0,02 X_0$) par un choix de matériau à faible Z.
- Une bonne rigidité permettant de supporter la tension des fils.
- Une bonne planéité afin de maintenir un écart constant entre les différents plans de dérive.
- Une bonne étanchéité limitant les fuites de gaz

Un panneau est constitué d'une structure en nid d'abeille en Aramide, encadrée de 2 couches de Kevlar-Epoxy de $400 \mu\text{m}$ d'épaisseur pour assurer la rigidité. Les peaux de Kevlar représentent 75 % de la masse : on voit en effet sur la figure 3.7 que les interactions de neutrinos avaient lieu préférentiellement dans ces peaux. Cette quantité de matière est source de diffusion multiple et limite la résolution en impulsion malgré la bonne précision spatiale de ces chambres.

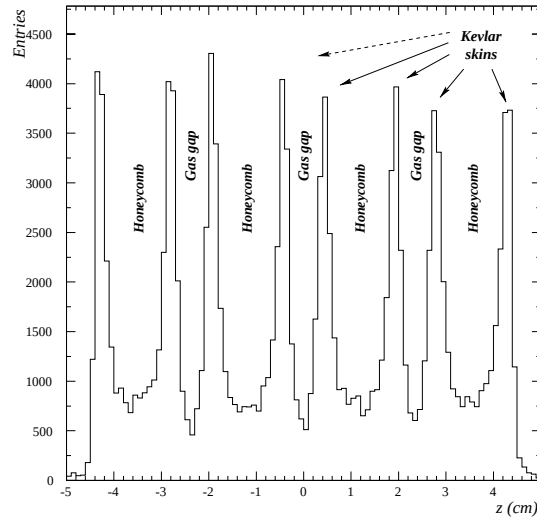


FIG. 3.7 – Distribution en Z des vertex d'interactions de neutrinos ramenés à une chambre (réalisée dans NOMAD[24]).

¹PC chargé de l'enregistrement d'un événement avec toutes les informations de tous les sous-détecteurs activés

3.3 La position des fils

Les fils, de champ et de lecture, sont placés dans le plan médian du volume gazeux. Ils sont soudés et collés aux 2 extrémités. Les fils de lecture ont une tension mécanique de 53g. Pour éviter que les fils ne s'écartent trop de leur position nominale sous l'effet du champ électrostatique, ils sont également fixés en 3 autres points intermédiaires sur des barres isolantes horizontales de 4 mm d'épaisseur (placées à $y = +75$ cm, 0 cm et -75 cm). Chaque fil peut être représenté sous la forme de 4 segments de droites dont les positions précises devront être obtenues par la procédure d'alignement (voir chapitre 4).

3.4 Conclusion

Les chambres à dérive de NOMAD représentent le meilleur compromis étant données les contraintes qui nous sont imposées :

- disponibilité : L'expérience a commencé à être opérationnelle à la prise de données physiques près de 15 mois après la création de la collaboration en février 2000. Le CERN nous avait attribué du temps de faisceau en 2001 et 2002. La rigueur au niveau des délais et le financement modeste de HARP ont conduit à l'utilisation de ces chambres à dérive. Elles n'avaient pas été utilisées depuis la fin de NOMAD mais le grand nombre de ces chambres nous permettait de sélectionner les plus fonctionnelles.
- coût : la réutilisation de ces chambres représente une économie. Les seuls coûts occasionnés ont porté sur les modifications mécaniques nécessaires à leur suspension et sur la construction des supports.
- performance : bien que fonctionnant dans des conditions différentes de celles de NOMAD, nous verrons dans le chapitre dévolu aux performances des chambres à dérive, qu'elles remplissent tout à fait leur rôle.

Au niveau des inconvénients, nous pouvons noter les points suivants :

- Conception : le problème premier de ces chambres réside dans leur conception. Elles sont sensées être suffisamment massives pour constituer une cible aux neutrinos de NOMAD. Or les mesures de précision visées par HARP, auraient réclamé des chambres aussi transparentes que possible pour ne pas perturber les trajectoires mesurées.
- Disposition : La disposition des modules de manière dispersée, change la façon de reconstituer les traces : on doit ici se baser sur la reconstruction de segments qu'il faut associer à travers la matière présente dans les autres détecteurs, alors que dans NOMAD [24], les modules se juxtaposaient permettant d'obtenir une trajectoire très fournie en points de mesure.

Dans l'ensemble, nous verrons dans les chapitres suivants que ces chambres à dérive ont rempli parfaitement leur rôle.

Troisième partie

Développement et environnement logiciel

Chapitre 4

Développement logiciel... vers l'analyse physique

Pour instruire la physique de cette expérience, il est nécessaire de passer par une phase de reconstruction qui, à partir des données des sous-détecteurs, permet de remonter à la trajectoire, l'impulsion et la nature des particules produites.

Dans ce chapitre, nous dévrons comment les données sont traitées afin de reconstruire des traces au sein des chambres à dérive. Une simulation détaillée de l'expérience est nécessaire non seulement pour évaluer les performances du programme de reconstruction mais également pour calculer l'acceptance qui permet une évaluation correcte de sections efficaces.

Ce chapitre qui traite des logiciels hors acquisitions (off-line¹) est construit suivant trois axes :

- une présentation de l'architecture informatique de la collaboration HARP et de son cadre général qui fournit des outils de développement d'algorithmes.
- la conception d'un programme de simulation réaliste, qui doit simuler non seulement la physique proprement dite mais également la réponse des détecteurs. De ce point de vue c'est un outil indispensable pour le développement des programmes de reconstruction des traces ou d'identification des particules. Je parlerai plus particulièrement de la digitisation des chambres à dérive que j'ai développée et qui simule leur réponse physique. Elle sera essentielle pour mesurer l'efficacité de la reconstruction des chambres.
- la reconstruction des traces dans les chambres à dérive à partir des coups laissés par le passage d'une particule chargée.

4.1 L'architecture informatique de la collaboration

En physique des particules, chaque expérience représente une association de différentes activités informatiques qu'il faut organiser. Ces activités peuvent être d'ordre très divers : un cadre général, que ce soit on-line ou off-line, et des fonctions spécifiques pour

¹Le terme off-line désigne tous les éléments fonctionnant en dehors de la prise de données en temps réel.

chaque sous-détecteur. Par exemple dans le cadre général du programme d'acquisition de l'expérience, qui coordonne et synchronise les tâches de lecture et d'écriture des données en temps réel, chaque sous-détecteur fournit les fonctions de lecture de ses données spécifiques (les TDC v767 CAEN pour les chambres à dérive) ; de même le programme de simulation est bâti autour d'un cadre général (qui fournit la description géométrique et matérielle du détecteur, les processus d'interaction particule-matière et la propagation des particules) et d'une gestion spécifique à chaque sous-détecteur du "signal" laissé par les particules et de son codage au plus près de la réalité des instruments. Ces exemples s'étendent également au programme de surveillance en ligne (monitoring) et au programme de reconstruction. Tous ces différents algorithmes sont organisés pour chaque sous-détecteur dans un cadre défini en vue d'un développement cohérent.

Pour la partie *Software off-line* à laquelle je me restreins ici, la collaboration HARP a décidé d'utiliser après maintes réflexions, de nouveaux logiciels soutenus par le CERN jusqu'à présent. Ceux-ci sont interfacés et codés en langage C++.

Les logiciels utilisés dans le secteur off-line sont les suivants :

- GAUDI [26] : cadre général du secteur off-line
- Objectivity[27] : Base de données
- GEANT 4 [28] : Simulation de la physique et réponse des détecteurs.
- ROOT[29] : monitoring des évènements.

4.1.1 GAUDI

Le cadre général choisi par la collaboration est celui développé pour l'expérience LHCb : GAUDI [26]. Sous une apparence simple dans le fonctionnement, il s'agit d'un cadre d'une certaine complexité dans son installation et sa maintenance. Avant tout, il me semble bon de redéfinir ce qu'est un cadre et quelle doit être sa fonction : un cadre est un squelette possédant des applications et gérant les relations entre ces applications. Les développeurs devront structurer leurs codes personnalisés autour de ces applications.

GAUDI fournit ainsi un certain nombre d'outils et de services pour son utilisation et pour le développement logiciel (voir figure 4.1) :

Services pour utilisateur :

- un service de sélection des événements (EventSelector) qui permet de gérer les appels à la base de données.
- un service d'interactivité et commande (jobOptions Service) qui est l'interface principale entre un utilisateur et le cadre. Par l'intermédiaire d'un fichier d'options, on exécute une liste ordonnée des algorithmes que l'on désire avec la possibilité d'initialiser les paramètres de chacun d'entre eux. Pour qu'un algorithme soit reconnu et donc sous le contrôle du cadre GAUDI, il doit posséder 3 méthodes clé : *initialize*, *execute* et *finalize* ; le contenu de ces méthodes est du ressort du développeur.

Les méthodes *initialize* et *finalize* ne sont appelées qu'une seule fois lors du traitement d'un ensemble d'événements (run). La méthode *execute* est appelée à chaque événement traité. Tous les algorithmes de HARP sont donc structurés par l'appel de ces trois méthodes. Sans entrer dans les détails du polymorphisme en code C++, il s'agit de rendre utilisable à notre besoin particulier des algorithmes généraux. Ceci est une tech-

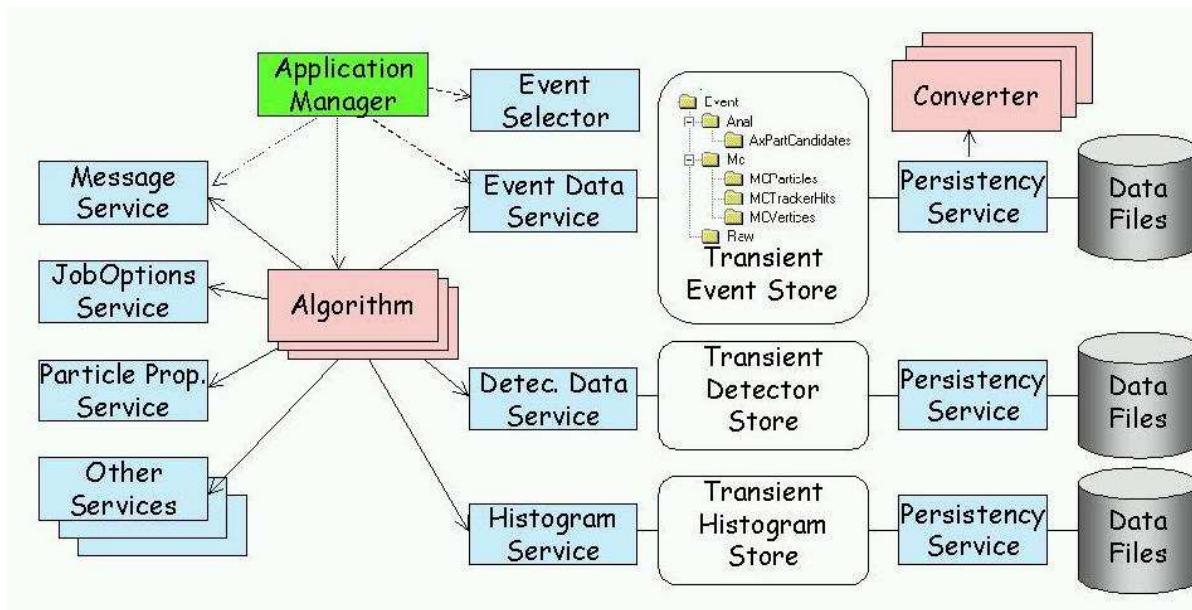


FIG. 4.1 – Architecture du cadre donné par GAUDI.

nique très courante lorsque l'on utilise un produit déjà existant.

Services pour développeur :

- un service de messagerie est disponible et d'une grande utilité lorsque l'on cherche à insérer différents niveaux d'alerte dans un code.
- GAUDI assure également l'accès aux données sur les détecteurs (géométrie, calibration, ...) à travers un service de gestion des paramètres de détecteurs (Detector Data Service). On verra dans la partie consacrée à la géométrie que cela permet de décrire chaque élément du détecteur avec un langage bien défini et unique.
- La manipulation d'objets ou d'informations entre différents algorithmes est accessible par un système de gestion des données. La figure 4.2 nous montre un schéma de fonctionnement de l'échange de données entre algorithmes : les données semblent passer d'un algorithme à l'autre, alors qu'en fait pendant le temps de traitement d'un événement, on utilise une zone de transition (Transient Event Data Store) qui permet d'avoir un accès libre aux données lues dans la base de données (Persistency Area) mais aussi aux résultats de traitements de ces mêmes données par les algorithmes utilisés.

Pour l'analyse physique, un service de gestion des histogrammes et des objets d'analyse NTuples est disponible reprenant la plupart des commandes de l'environnement d'analyse HBOOK [30] adaptées aux C++. Je rappelle que les Ntuples observent la même architecture que les bases de données : chaque événement d'un Ntuple contient des informations alphanumériques et est identifié à partir d'une clé unique (exemple : le numéro de l'événement). Les commandes non encore disponibles en C++ peuvent utiliser les routines fortran à l'aide de méthodes d'interface fortran/C++ de la CERNLIB.

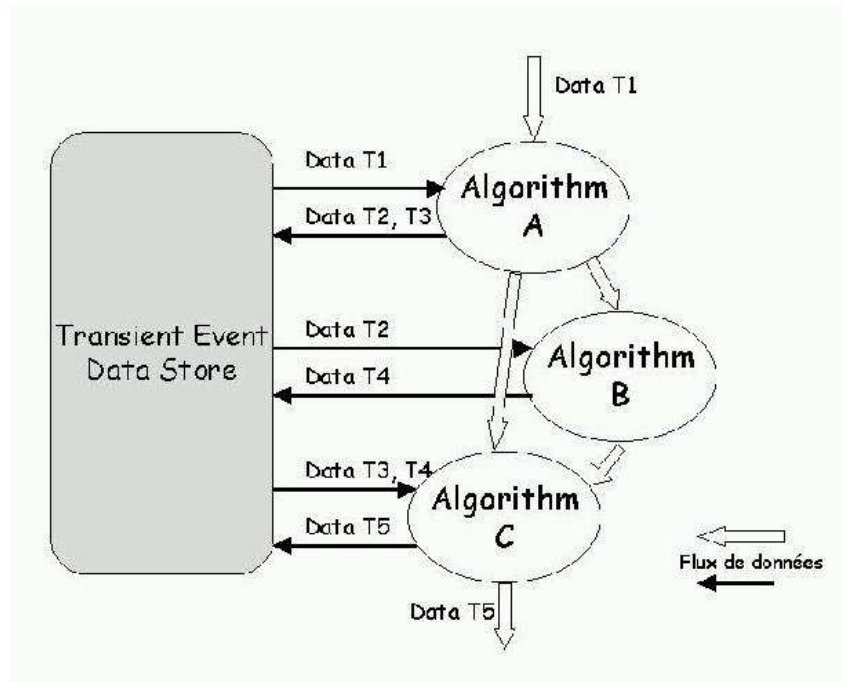


FIG. 4.2 – Gestion du transfert de données entre les différents algorithmes lancés sous GAUDI. Les flèches noires représentent les flux de données réels, les flèches blanches, les flux de données apparents pour l'utilisateur.

4.1.2 Objectivity

Objectivity [27] est le système de base de données choisi par la collaboration. Les motivations de ce choix étaient le soutien² du CERN, et une certaine expertise au sein de la collaboration.

Son niveau d'intervention dans la configuration de gestion de GAUDI (figure 4.1) concerne le stockage des données (fichiers aussi bien de données brutes que de données de calibration et de géométrie) mais également leur conversion dans la zone tampon ("Transient Store") afin de les mettre à disposition des algorithmes GAUDI.

L'implication d'Objectivity dans l'organigramme GAUDI ne le rend pas capital. Le traitement des données et leurs manipulations sont découplés de leur stockage sous Objectivity puisqu'ils se font dans la zone tampon qui est la source de l'information au niveau des logiciels d'analyse. Ainsi il devrait être possible de changer de base de données de manière invisible (dans un monde parfait) pour un utilisateur.

4.2 Simulation

La Simulation est un outil important par sa capacité à prévoir et à tester. J'ai personnellement été très impliqué dans l'ensemble de la simulation des chambres à dérive. Une très bonne simulation de la réponse du détecteur est requise en vue de tester les algorithmes de reconstruction des traces dans ces chambres à dérive. Bien sûr il s'agit de se rapprocher le plus possible de la réalité sur différents points : modèle adéquat

²pour information, à l'heure actuelle, le soutien du CERN n'est plus et la collaboration HARP se prépare à migrer vers une base de données ORACLE

de dérive des électrons, prise en compte de toutes les composantes du temps mesuré, bruit de fond électronique, efficacité, plans défectueux. En se rapprochant des conditions expérimentales, on justifie et crédibilise une comparaison des données simulées et réelles.

4.2.1 GEANT 4

La simulation est basée sur le programme GEANT 4 qui est une évolution de GEANT 3 (codé en fortran) vers le langage C++[31]. Les prémisses de ce programme ont commencé en 1994 avec la volonté de réadapter le code dans un cadre Orienté-Objet choisi par les grandes expériences LHC. Il faut préciser que GEANT 4 n'en est encore qu'en phase de développement ce qui rend son utilisation complexe du fait de bogues potentiels mais également de par son manque de documentation. Donc la grande difficulté pour moi fut de trouver des gens compétents avec qui interagir, quelques bonnes âmes de la collaboration GEANT 4, mais également de me former au développement sous ce type d'interface.

GEANT 4 est un cadre pour la simulation en physique des particules. De la même manière que GAUDI précédemment, ce produit est composé de classes abstraites généralistes et le rôle du développeur est justement de les rendre concrètes, et donc d'introduire les informations dont GEANT 4 a besoin pour réaliser une simulation personnalisée.

“physique” lors du passage d’une particule dans la matière.

GEANT 4 est capable suivant la particule concernée de la faire évoluer selon une liste de processus physiques très complète.

- désintégration : le “choix” se fait selon les modes connus pour chaque particule, et selon les probabilités associées.
- Interaction électromagnétique : effet photoélectrique, diffusion Compton, conversion gamma, diffusion multiple.
- interaction hadronique : interaction hadron-lepton et hadron-hadron, fragmentation, fission.

Les hadrons étant des objets complexes composés de partons, et l'interaction hadronique étant moins bien connue que l'interaction électromagnétique, plusieurs générateurs co-existent dans GEANT et peuvent être activés sur choix de l'utilisateur. Il est à noter que c'est exactement le secteur où les résultats de HARP auront une influence importante grâce à la grande statistique accumulée.

Fonctionnement global de GEANT 4

D'un point de vue technique, GEANT 4 propage toute particule à travers le détecteur par *pas*. Le pas est défini par le point de départ, la direction initiale et la longueur. C'est dans le calcul de cette longueur de pas que GEANT met en compétition les contraintes géométriques et les processus physiques. Par définition, un pas est toujours contenu dans un même volume géométrique et doit donc être interrompu à la rencontre d'une frontière entre 2 volumes ou milieux. A l'intérieur d'un même milieu, la longueur du pas est déterminée par celui des processus physiques possibles qui l'aura emporté : un tirage aléatoire d'une longueur (d'interaction ou de désintégration) est fait pour chacun

de ces processus concurrents, et c'est la longueur la plus courte qui est sélectionnée par GEANT. Le pas est alors complètement déterminé, mais également, toute particule nouvelle éventuellement issue du processus à l'œuvre, est stockée (avec son point de départ et sa direction initiale) dans une pile, pour un suivi ultérieur par GEANT. Un seuil de propagation permet à l'utilisateur d'indiquer à GEANT 4, les limites en énergie au-delà desquelles une particule sera propagée ou non.

Une étape essentielle est de définir les volumes “sensibles” qui entraîneront une réponse du ou des détecteurs. Il s'agit d'attribuer à ces objets des règles de fonctionnement que le développeur peut personnaliser pour produire un signal lors du passage d'une particule. Ces règles sont formulées par des algorithmes de type SD pour “Sensitive Detector”. Ceux-ci représentent la réponse caractéristique de chaque sous-détecteur et sont la concrétisation d'un algorithme abstrait de GEANT 4 appelé *G4VSensitiveDetector*. On associe ainsi à des volumes particuliers de la géométrie d'un détecteur, les zones de détection et de formation des signaux électroniques dus au passage d'une particule.

Au niveau de l'interactivité, il est possible d'initialiser une simulation suivant différents paramètres :

- nature, position, direction (constante ou aléatoire dans un angle solide donné), énergie du faisceau primaire
- dimensions horizontale et verticale du faisceau
- Nombre d'événements

Remarque : Tous ces paramètres peuvent être initialisés dans le cadre de GEANT 4, cependant pour mieux contrôler les étapes logiques entre les algorithmes de GAUDI et de GEANT 4, il était préférable d'utiliser au maximum les services de gestion de GAUDI pour la phase d'initialisation de la simulation.

4.2.2 La géométrie des chambres à dérive.

Il y a plusieurs manières de coder la géométrie d'un détecteur : on peut utiliser une méthode lourde où il s'agit d'écrire directement le code de chacun des éléments du détecteur à la GEANT 4. Ou bien on utilise un programme générique, une interface qui permet de lire une liste de fichiers de géométrie standardisés, ce qui est plus simple dans la réalisation. On a ainsi une lecture commune de la géométrie de tous les sous-détecteurs, lecture que l'on peut généraliser à tous les programmes de l'expérience, en particulier la reconstruction.

HarpDD

Le programme HarpDD (DD signifiant Detector Description) est l'interface qui permet de gérer la géométrie globale de l'expérience. Il possède deux composantes principales :

- La description des sous-détecteurs en tant qu'objets physiques.
- La localisation des zones de mesures (volumes sensibles) et les outils pour leur calibration.

L'architecture des sous-détecteurs est décrite au travers de fichiers ASCII suivant des bases de géométries rudimentaires qui sont utilisées dans GEANT 4 (BOX, CONE, TUBE, POLYGON, ...). GEANT 4 possède des règles dans le domaine géométrique très précises

et le format générique des fichiers de géométrie doit recueillir toutes les informations nécessaires à la réalisation d'une géométrie à la GEANT 4. Celle-ci se fait suivant une architecture en arbre en partant d'un espace général à trois dimensions dans lequel on dispose des éléments de plus en plus fins par inscrustation les uns dans les autres obtenant le degré de réalisme choisi. Bien sûr il faut savoir que trop de détails peut amener moins de performances en terme de traitement des informations.

La description géométrique utilise la nomenclature décrite en annexe D.

Au sein du cadre GAUDI, la procédure de lecture des fichiers de géométrie se fait via un service de type *DetectorServiceSvc* (figure 4.1) dévolue aux données relatives aux détecteurs. Ce service appelé *HarpDDSvc* fait appel à une unité de contrôle *HgDetUnitMgr* (figure 4.3) qui lit chaque fichier de géométrie en analysant les étiquettes (:ELEM, :MATE, :VOLU, ... ; Cf. annexe D) trouvées pour connaître le type de données.

La construction des matériaux proprement dite est déléguée à l'algorithme *HgMaterialFactory*. L'ordre d'appel des fichiers de géométrie est important : en effet, on ne peut pas construire un objet si l'on ne connaît par sa matière et on ne peut placer ce que l'on n'a pas construit. C'est pourquoi on appelle d'abord la construction des matériaux, puis la construction des éléments de volume et enfin leur disposition.

DetRep

DetRep (Detector Representation) est le programme qui joue le rôle d'interface entre HarpDD et GEANT 4. Ce service permet de créer un accès à la composition géométrique du détecteur dans le format GEANT 4. Dans les faits, il utilise un algorithme *HsG4VolumeMgr* qui obtient les éléments géométriques via *HgDetUnitMgr* et les matériaux via *HsMaterialMgr*. La figure 4.4 montre le résultat obtenu par l'interface graphique de GEANT 4 lorsque l'on a voulu afficher les chambres à dérive en détail. Les volumes principaux des chambres sont des tranches verticales, perpendiculaires au faisceau et donc à la direction générale de la plupart des particules produites vers l'avant dans les interactions. Cette géométrie simple facilite le calcul des pas dans GEANT et donc la vitesse d'exécution. On a inclu dans la géométrie les volumes importants mais également ceux de petites tailles qui représentent une surdensité de matière dans les panneaux, tels que les renforts, susceptibles de perturber la trajectoire des particules. On peut constater que les fils n'apparaissent pas car c'est tout à fait inutile : cela demanderait à GEANT 4 beaucoup de données à gérer ce qui affaiblirait beaucoup la vitesse d'exécution, de plus d'un point de vue matériel, ils ne représentent pas une quantité de matière significative.

4.3 La simulation dans les chambres à dérive

La Simulation est importante pour 2 raisons principalement : d'une part, elle permet de tester la qualité des générateurs en comparant des données simulées avec des données réelles. En effet, les générateurs hadroniques actuels ont été ajustés sur des données obtenues dans les années 70, et il sera intéressant de pouvoir les comparer avec les données de l'expérience HARP. D'autre part, la simulation est essentielle pour le développement

et le test des algorithmes de reconstruction.

4.3.1 La digitisation des chambres à dérive.

Pour chaque détecteur, il en va de la responsabilité de l'utilisateur de coder sa réponse dans GEANT 4 : c'est la *digitisation*. J'étais responsable de la digitisation des chambres

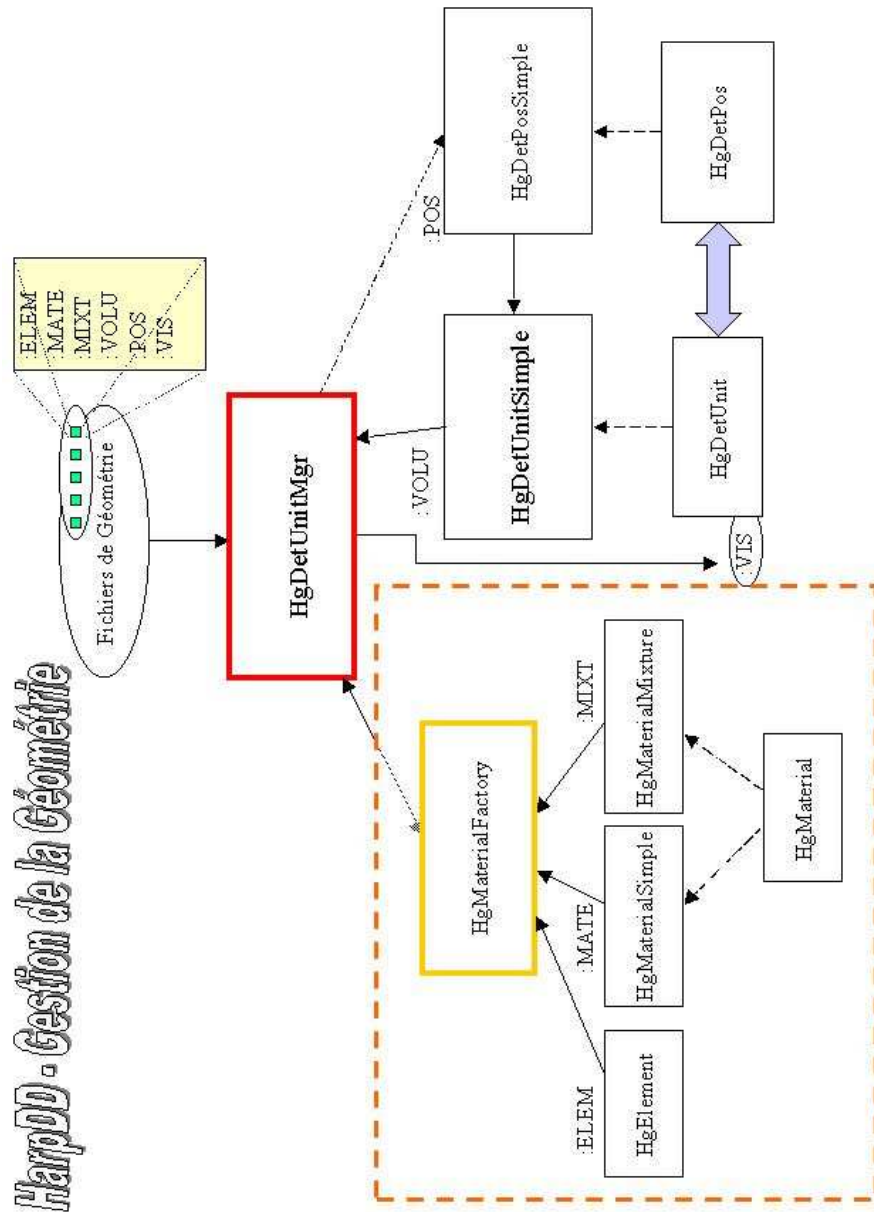


FIG. 4.3 – Schéma de l'algorithme permettant la lecture générique des fichiers de géométries dans le programme HarpDD. Les algorithmes entourés d'un cadre épais sont des singletons de contrôle.

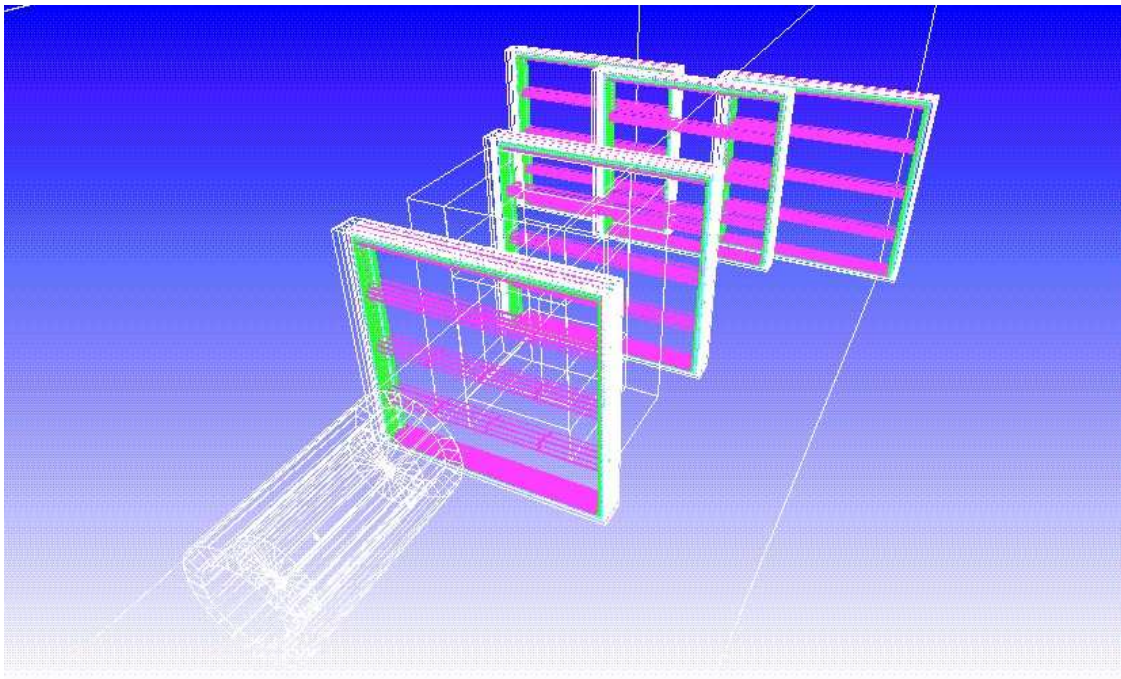


FIG. 4.4 – *Visualisation 3D de la géométrie des chambres à dérive*

à dérive. Je l'ai basée suivant un modèle particulier (figure 4.5) qui semble être ce qui correspond au mieux à la trajectoire des électrons secondaires, étant donnée la structure des lignes de champ de ces chambres.

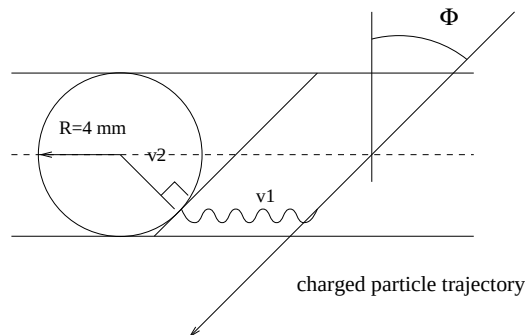


FIG. 4.5 – *Schéma utilisé pour la relation temps-distance : modèle de dérive de l'électron le plus rapide.*

Lorsqu'une particule traverse le milieu dit "sensible" (qui est notre milieu gazeux), l'électron déclencheur du signal est celui qui atteint le premier l'anode. Le modèle consiste à considérer que les électrons suivent une trajectoire rectiligne dans la région des lignes de champs plus ou moins parallèles et radiale lorsqu'ils subissent l'influence du potentiel des fils de lecture. La frontière entre les deux régimes est dans notre cas, un cercle de 8 mm de diamètre (c'est-à-dire l'épaisseur de gaz).

Le problème est un peu plus complexe, comme on le verra plus loin, lorsque GEANT 4 propage une particule en plusieurs pas (ou "steps") à travers le volume gazeux, du fait

de la diffusion de la particule incidente sur des constituants gazeux.

La digitisation repose sur une stratégie en deux temps : la création du hit dans un premier temps, qui représente l'objet "idéal" symbolisant l'endroit de passage de la particule dans un volume sensible ; puis vient la digitisation de ces hits ce qui nous donne des "digits". Ils représentent la traduction aussi proche que possible des instruments réels de cet objet idéal. Pour cela, on utilise le modèle précédent, mais on doit aussi tenir compte de la résolution de la chambre touchée, de son bruit, de son efficacité, de la propagation du signal sur le fil et dans l'électronique de lecture.

On positionne arbitrairement un hit à l'extrémité aval d'un step. Les attributs d'un hit sont :

- sa position.
- les coefficients directeurs de la trajectoire.
- la longueur du step auquel il appartient.

Il va s'agir ensuite de procéder à la digitisation de ces hits afin d'en extraire un temps, qui est la quantité mesurée des chambres à dérive.

On définit le temps de dérive comme le temps que met le premier électron pour arriver sur l'anode. Il est appelé "électron déclencheur". Ce temps est ensuite retraduit en distance (par convention la distance du fil à l'impact de la trace au centre du gap gazeux) par la relation temps-distance de la reconstruction. Pour la digitisation, il a fallu considérer tous les cas de figure que l'on peut rencontrer :

- Le cas le plus fréquent est celui où GEANT 4 propage une particule à travers le gaz d'une même cellule en un seul step. Celui-ci peut soit être à l'extérieur du cercle (figure 4.6.a), soit passer à l'intérieur (figure 4.6.b). Pour une trace interne au cercle, on ne considère que la composante radiale de la dérive.
- Pour le cas où la particule interagit par diffusion sur une molécule de gaz avec changement notable de direction (dans environ 10 % des cas), on se retrouve avec 2 steps dans une même cellule. On considère alors que l'on doit former un hit pour chaque step. Pour la digitisation, il s'agit donc de déterminer pour chaque step, l'électron déclencheur. On voit que pour un step si deux cas de position pour l'électron déclencheur (figure 4.7) sont possibles : (a) l'électron est à une extrémité du step - (b) l'électron est à une position calculée à partir de la tangente au cercle parallèle à la trajectoire. Enfin, et pour tenir compte du fait que les TDC ne peuvent séparer deux signaux à moins de 15 ns, on ne produira qu'un seul digit si tel est le cas.
- Il peut également arriver qu'une particule traverse le gap en étant à cheval sur deux cellules, voire plusieurs si la trajectoire est suffisamment verticale. Dans ce cas, le step GEANT 4 est subdivisé au cours de la digitisation en plusieurs steps, un par cellule. Les traces étant majoritairement proches de l'horizontale, on obtient principalement des cas à deux cellules.

La simulation a été mise au point en vue d'assurer un test efficace du programme de reconstruction. On a donc inclus sous forme d'options, que l'on peut activer à volonté, la prise en compte du temps de propagation du signal le long des fils, de la résolution en position des chambres, de leur efficacité... En effet ce qui différencie la création des hits de la création des digits est que le hit représente la trajectoire réelle d'une particule simulée

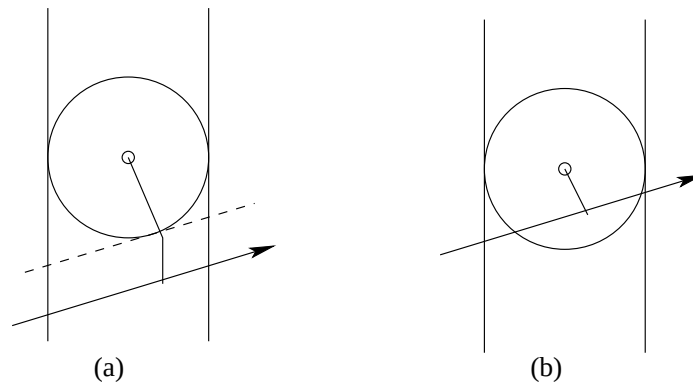


FIG. 4.6 – Cas d'un step

alors que la reconstruction se fera sur des données qui contiennent des incertitudes : la possibilité d'inclure ou non ces incertitudes est essentielle pour les tests d'algorithmes.

Nous pouvons voir (figure 4.9) que le programme de simulation permet d'obtenir une distribution des temps de dérive plate correspondant à ce à quoi l'on s'attendait.

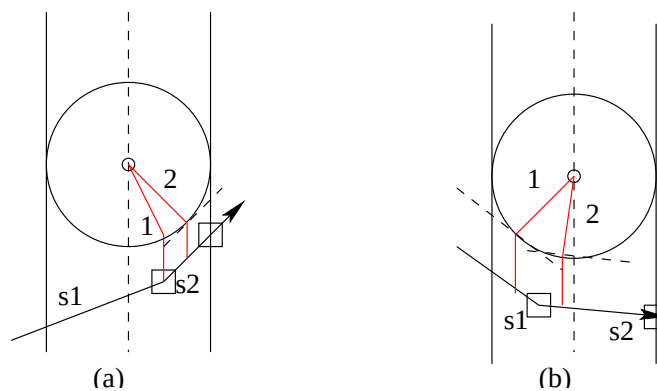


FIG. 4.7 – Cas d'un multistep

Schéma de fonctionnement sous GAUDI et GEANT 4

Le principe de fonctionnement du programme de simulation de la réponse des chambres à dérive suit un certain nombre d'étapes (figure 4.8).

La création des hits (coup idéal) est établie par un algorithme sous contrôle de GAUDI qui va assigner aux volumes sensibles les règles de fonctionnement (voir partie GEANT 4). Dans le cas des chambres à dérive, il utilise un module qui contient le processus de création des hits.

GAUDI se charge de stocker les hits enregistrés dans des containers utilisable par d'autres algorithmes. La digitisation utilise les objets hits comme entrée pour construire des objets Digit qui représente le coup réaliste c'est-à-dire avec des effets d'incertitudes quant à la réponse du détecteur. Dans le cas de données Monte-Carlo, on établie un lien entre le digit et le hit afin d'avoir un point de comparaison quant à la crédibilité de l'algorithme de reconstruction développé.

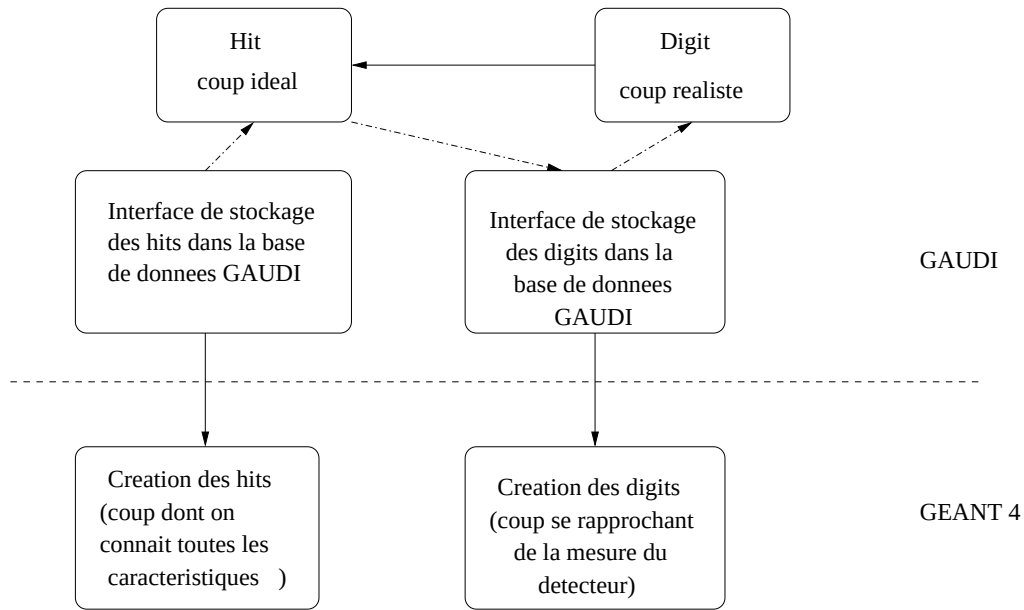


FIG. 4.8 – Schéma de fonctionnement de la Simulation.

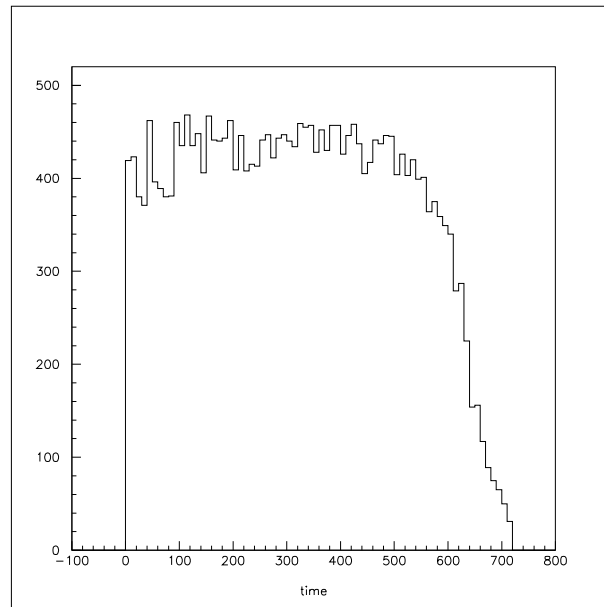


FIG. 4.9 – Distribution des temps de dérive (en ns) obtenus avec le module de digitisation.

4.3.2 Les paramètres optionnels.

Comme je l'ai dit auparavant, un des objectifs du programme de simulation est de mettre à disposition des outils nécessaires à la réalisation de tests pour juger des performances de l'algorithme de reconstruction. Au niveau des chambres à dérive, on va rechercher à faire le cheminement inverse d'une reconstruction.

Lorsque GEANT 4 propage une particule dans les milieux sensibles des chambres, on a des positions de particules parfaites que l'on va dégrader, pour simuler une résolution, puis convertir en temps. Pour la résolution, 3 choix sont possibles :

- Pas de résolution appliquée, ce qui revient à conserver des traces parfaites.
- une résolution constante que l'utilisateur peut choisir.

- une résolution générée par une fonction qui prend en compte l'information de la distance de dérive et de l'angle de la trace.

Le temps intéressant tel qu'il a été décrit dans la partie concernant la digitisation, est le temps de dérive. Cependant le temps enregistré lors de prises de données réelles, est une somme de plusieurs temps (voir plus loin la partie définissant le temps dans la reconstruction) que l'on peut choisir ou non de prendre en compte à diverses étapes de l'algorithme. Les temps à rajouter dans le cadre de la simulation sont le temps de vol de la particule depuis l'interaction et le temps de propagation du signal le long du fil.

La prise en compte d'un bruit de fond électronique est également réalisable par l'introduction de coups parasites autour de la trace.

4.4 Reconstruction des traces dans les chambres à dérive.

4.4.1 Avant-propos

Un des choix délicats de cette expérience est celui de l'algorithme de reconstruction puisque la précision des mesures dépend de l'efficacité de la reconstruction. Si délicat d'ailleurs que celui-ci n'est pas encore arrêté, plusieurs algorithmes étant développés en parallèle. Il existe toutefois un algorithme officiel basé sur des concepts développés par NOMAD. À l'heure actuelle, tous ces algorithmes sont dans un état de développement encore loin d'être finalisé ce qui amènera des résultats préliminaires quant à la recherche des performances de ces chambres à dérive. Mais après une brève description du programme officiel, je parlerai plutôt de l'algorithme en développement de l'équipe des chambres à dérive du LPNHE et DUBNA.

Les chambres à dérive ont pour réponse au passage d'une particule chargée, un signal provenant de la dérive des électrons vers le fil d'anode. Ce signal est amplifié, transporté et codé sous forme d'un temps dans un TDC.

L'objectif étant la reconstruction tridimensionnelle de la trajectoire de la particule, il faut traduire ce temps en distance (de dérive). La conversion du temps en une distance suit en fait le cheminement inverse de la simulation (où l'on cherchait à convertir une position en temps). Pour cela on utilise le modèle présenté dans la partie Simulation (figure 4.5).

4.4.2 Le temps dans les chambres à dérive.

Le temps que nous fournissent les TDCs des chambres à dérive résulte d'une somme de différents temps comme le montre la figure 4.10 :

$$t_{TDC} = t_{interaction} + t_{vol} + t_{dérive} + t_{fil} + t_{cable} \quad (4.1)$$

Le signal de déclenchement de l'expérience est distribué sur chaque TDC, sur un canal libre (le canal 109) (figure 4.11) :

$$t_{Canal109} = t_{interaction} + t_{trigger} \quad (4.2)$$

où $t_{interaction}$ est le temps au moment de l'interaction de la particule primaire avec la cible.

t_{vol} est le temps de vol de la particule de la cible jusqu'à un plan de dérive.

t_{derive} est le temps de dérive des électrons dans une cellule de dérive jusqu'au fil d'anode.

t_{fil} est le temps de propagation du signal le long du fil jusqu'au préamplificateur.

t_{cable} est le temps de parcours du signal partant des préamplificateurs vers les TDC.

$t_{trigger}$ est le temps de formation du trigger.

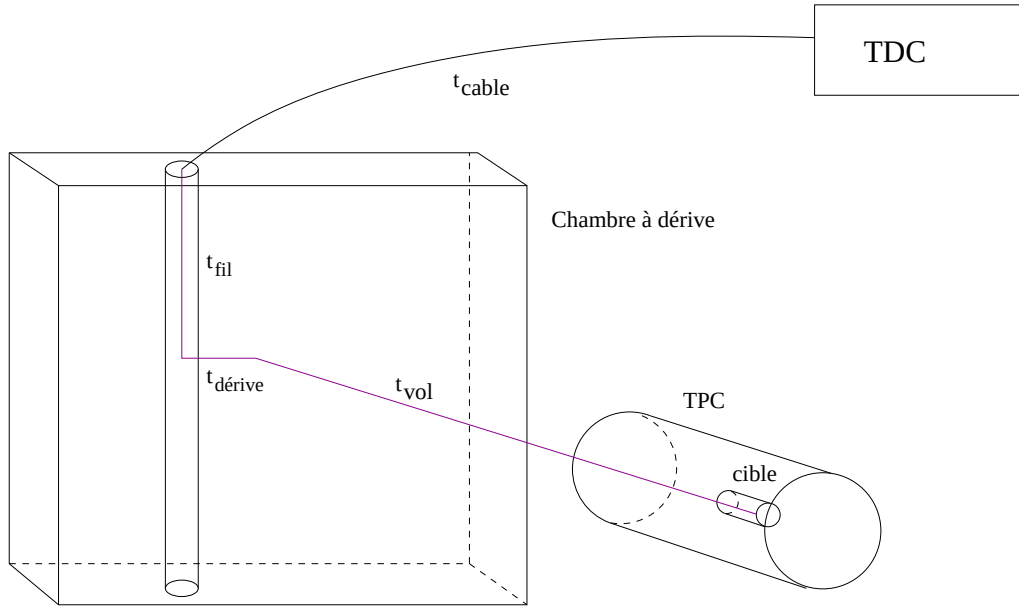


FIG. 4.10 – Décomposition du temps TDC.

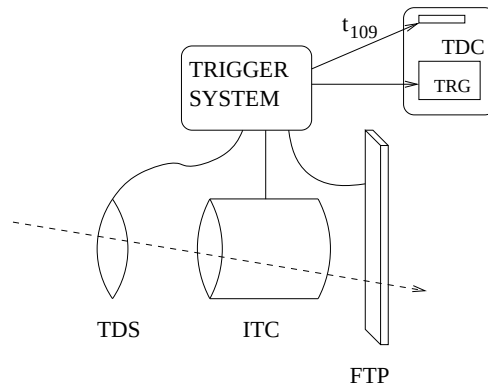


FIG. 4.11 – Composition du temps obtenu dans le canal 109.

En faisant la différence Δt , on obtient une durée (on continuera à l'appeler temps dans la suite) indépendante du temps de l'interaction.

$$\Delta t = t_{TDC} - t_{Canal109} = t_{vol} + t_{derive} + t_{fil} + t_{cable} - t_{trigger}.$$

Cette durée ou temps est finalement la donnée à traiter par l'algorithme de reconstruction.

t_{vol} et t_{fil} dépendent de la position des coups le long du fil. Cette position est affinée au cours de la reconstruction.

$t_{trigger} - t_{cable}$ est une constante de chaque fil. On fait l'hypothèse que le trigger n'a pas de dispersion trop grande en fonction des différents types de déclenchement (ceci est vrai à 1 ns près).

Etant armé de ces informations, on peut remonter alors à l'information souhaitée : t_{derive} .

4.4.3 L'algorithme de reconstruction

Un choix stratégique.

Rappelons d'abord les conditions de NOMAD pour lesquelles l'algorithme à base de triplets a été développé : l'expérience recherchait des oscillations de neutrinos $\nu_\mu \leftrightarrow \nu_\tau$ en séparant cinématiquement les événements CC (courant chargé) de ν_μ , de ν_e et d'éventuels ν_τ . D'où la nécessité d'une reconstruction précise des particules chargées produites lors d'une interaction neutrino. NOMAD comptait 49 chambres à dérive installées dans l'entrefer d'un aimant dont le champ horizontal était parallèle aux fils. La reconstruction des traces était basée d'abord sur la recherche de points géométriques tridimensionnels que l'on appelle "triplet". Un triplet est un ensemble de 3 hits dans les 3 plans U, Y et V d'une même chambre. Les relations entre les coordonnées dans les repères de fils (U,Y,V) et celles dans le repère absolu (X,Y) sont :

$$U = Y \cos 5^\circ - X \sin 5^\circ$$

$$V = Y \cos 5^\circ + X \sin 5^\circ$$

La construction de ces triplets reposait donc sur un critère qui permettait de juger si trois hits donnés pouvaient appartenir à une même trace en minimisant la quantité :

$$\Delta = U + V - 2.Y \cos 5^\circ$$

La recherche de trace consistait ensuite à construire l'hélice passant par 3 triplets et à collecter tous les hits dans une "route" autour de cette hélice. La collection des hits, l'extension de la trace et l'ajustement des paramètres de l'hélice étaient basés sur l'utilisation de la technique du filtre de Kalman, bien adapté à cette géométrie et à la présence importante de diffusion multiple entre les points de mesure, due à l'épaisseur des parois des chambres.

Cependant dans HARP, la configuration est très différente car les chambres à dérive sont regroupées par module de 4, modules qui sont séparés les uns des autres par un aimant et d'autres sous-détecteurs. Les trajectoires des particules de part et d'autre de l'aimant sont donc des droites (si l'on néglige la diffusion sur les matériaux traversés). Les modules ne nous offrent que des fragments de droites qui peuvent contenir au mieux

12 points de mesures. Une stratégie fondée sur les triplets à l'intérieur d'un même module peut paraître risquée car cela nous donne au maximum 4 triplets. L'inefficacité des chambres, plus importantes dans HARP que dans NOMAD du fait du changement de gaz, peut conduire à une inefficacité de reconstruction plus importante. Une optimisation de cet algorithme doit donc être trouvée si l'on tient à l'appliquer à HARP. C'est la voie suivie par le groupe "Reconstruction" de la collaboration.

Une alternative a été proposée par nos collègues de Dubna, basée sur la création de segments par projection. C'est celle que nous avons adoptée pour sa pertinence dans le traitement de l'alignement, et c'est celle que je vais décrire ci-dessous. Mais ni l'une ni l'autre des 2 méthodes n'est à l'heure actuelle suffisamment avancée pour être quantifiée en termes d'efficacité de reconstruction globale et de résolution de l'impulsion.

Avec l'aide des définitions ci-dessous, on peut expliciter une efficacité asymptotique, dans le cas des segments, comme étant le rapport du nombre de cas où l'on a au moins 2 segments 2D et au moins 2 hits sur 4 possibles dans la troisième projection sur le nombre de passage effectif de particule Monte-Carlo. On appelle également "efficacité des chambres", la probabilité de création d'un coup dans un plan de gaz après le passage d'une particule chargée.

On peut comparer les efficacités asymptotiques de ces algorithmes compte tenu de la répartition des hits dans un module pour une trace donnée pour différentes efficacités moyennes de chambres.

La figure 4.12 nous montre clairement que la stratégie de traitements des hits (décrit ci-dessous) par type de plan apporte un flexibilité plus importante que par la gestion de triplets.

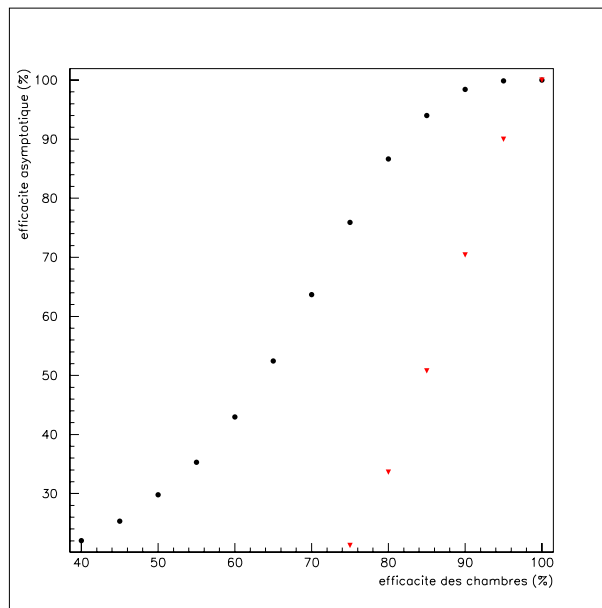


FIG. 4.12 – *Efficacité asymptotique d'algorithmes de reconstruction en fonction de l'efficacité des chambres à dérive : algorithme à base de triplets (triangles), et algorithme à base de segments (ronds).*

Quelques notions et définitions

L'algorithme de reconstruction utilisé repose avant tout sur une programmation orientée-objet dans laquelle l'on a défini les objets suivants :

- un hit : cela représente un coup dans un plan sensible de la chambre à dérive. Un hit est caractérisé par son temps et un numéro de canal TDC, qui une fois décodé peut être traduit en numéros de fil, plan, chambre et module.
- un segment 2D : ce type de segment représente une association de hits (dans un module) appartenant à un même type de plan : U,X ou V (c'est-à-dire un plan dont les fils sont à -5 ,0 et +5 degrés.). Par module, un segment 2D peut donc avoir au maximum 4 hits.
- un segment 3D : il s'agit d'un segment dans l'espace tridimensionnel créé à partir de 3 segments 2D de type U,X et V. Pour avoir une trace, il suffit ensuite d'associer les segments des différents modules.

Traitement des hits.

Dans un premier temps, un programme nous permet de récupérer, depuis la base de données, les hits enregistrés lors des prises de données : il s'agit du convertisseur d'objet dont le nom est *ObjectCnvAlg*, qui nous procure une liste de hits convertis dans le format *NdcEventHit* et la rend disponible dans la "Transient store". La conversion inclut plusieurs étapes, dont le décodage des mots des TDC pour les traduire en temps, numéro de fil, plan, chambre et module.

Ensuite on récupère la liste d'objets *NdcEventHit* que le programme *CreateNdcHitFromEventHit* (cf figure 4.13) va étudier en vue d'effectuer un premier traitement des données. Le traitement principal consiste à se rapprocher de la valeur du temps de dérive $t_{dérive}$. Les temps que l'on peut immédiatement retrancher sont le t_0 global qui est le temps de trigger (canal 109 sur chaque TDC) retenu comme temps initial de l'événement, mais également le temps spécifique au fil considéré que l'on avait représenté précédemment par t_{cable} (temps de transit du signal dans l'électronique et dans les câbles).

On obtient donc, comme indiqué au §4.2.2, $t_{dérive} + t_{vol} + t_{fil}$. À ce stade, ne connaissant pas la position de la trace le long du fil touché, on ne peut pas évaluer t_{vol} et t_{fil} et on fait donc l'hypothèse de départ que le hit est issu du milieu du fil. Le temps de dérive que l'on en déduit n'est qu'approché, et sera affiné lorsque la reconstruction aura permis de construire un segment 3D.

Avec la liste de hits obtenue, nous passons à la seconde étape de l'algorithme en utilisant l'algorithme de création de segments (*CreateNdcSegmentsAlg*) dans lequel nous recherchons à construire des traces.

Recherche de segments de plan (ou segments 2D)

Pour effectuer ce type de recherche, on essaie d'associer en un segment de droite les hits de 4 plans de même nature dans un même module. Cette association se fait en effec-

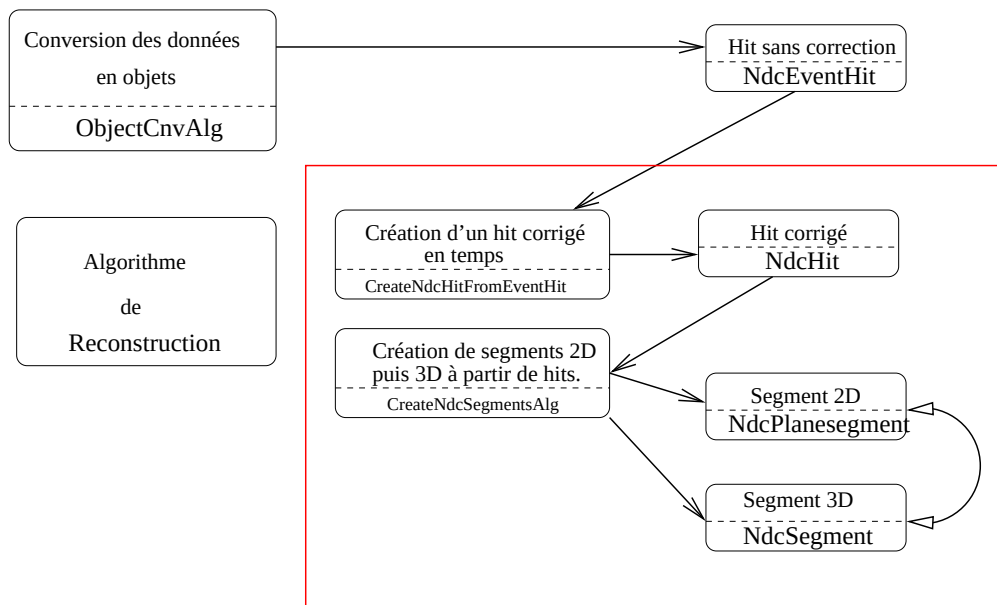


FIG. 4.13 – Schéma fonctionnel de l'algorithme de reconstruction.

tuant d'abord une présélection. Chaque hit est défini par un temps que l'on a converti en distance. Il reste cependant à savoir de quel côté du fil la particule est passée. L'association de N hits introduit 2^N possibilités de segments (voir figure 4.14 pour 4 hits). Cette présélection consiste à prendre des hits de plans 1 et 4 (les plans extrêmes) qui établiront un couloir d'une largeur donnée (5 mm par défaut). On cherche alors des hits dans les plans intermédiaires 2 et 3, qui seraient susceptibles de faire un segment avec les hits 1 et 4.

Avec 4 hits, nous avons donc 16 combinaisons possibles de signes. Lorsque nous avons une combinaison qui est acceptée par le couloir (≥ 3 hits), on procède à l'ajustement des paramètres de la droite et on calcule son χ^2 . Si le χ^2 est supérieur au χ^2_{MAX} , on rejette la combinaison. Un segment de projection est représenté par ses 4 hits associés et ses résultats d'ajustements des paramètres d'une droite arrangés par ordre croissant de χ^2 . On s'attend à ce que le χ^2 le plus faible corresponde à la combinaison la plus fiable pour la détection d'une trace. Nous verrons plus loin, quelle est l'efficacité de détection de la meilleure combinaison de l'algorithme lors des tests avec des données Monte-Carlo. Il est très fréquent que des hits se retrouvent impliqués dans la construction de plusieurs segments 2D. La combinaison ayant le plus faible χ^2 est retenue pour lever la dégénérescence.

Il arrive également que du fait de l'inefficacité d'un ou plusieurs plans d'un module, nous ne puissions associer 4 hits pour obtenir un segment 2D. C'est pourquoi après avoir tenté de produire des segments à 4 hits, nous essayons sur le même principe, de trouver des combinaisons de 3 hits pour lesquelles il nous est possible de faire des ajustements. Bien sûr on ne considère que des hits qui n'ont pas été employés par un segment à base de 4 hits.

Lorsque la construction des segments est terminée pour un événement, on se retrouve souvent avec 1 ou 2 ajustements en moyenne pour une même trace Monte-Carlo dans un module. Une sélection s'effectue sur le principe qu'un hit ne peut appartenir à plus d'un segment.

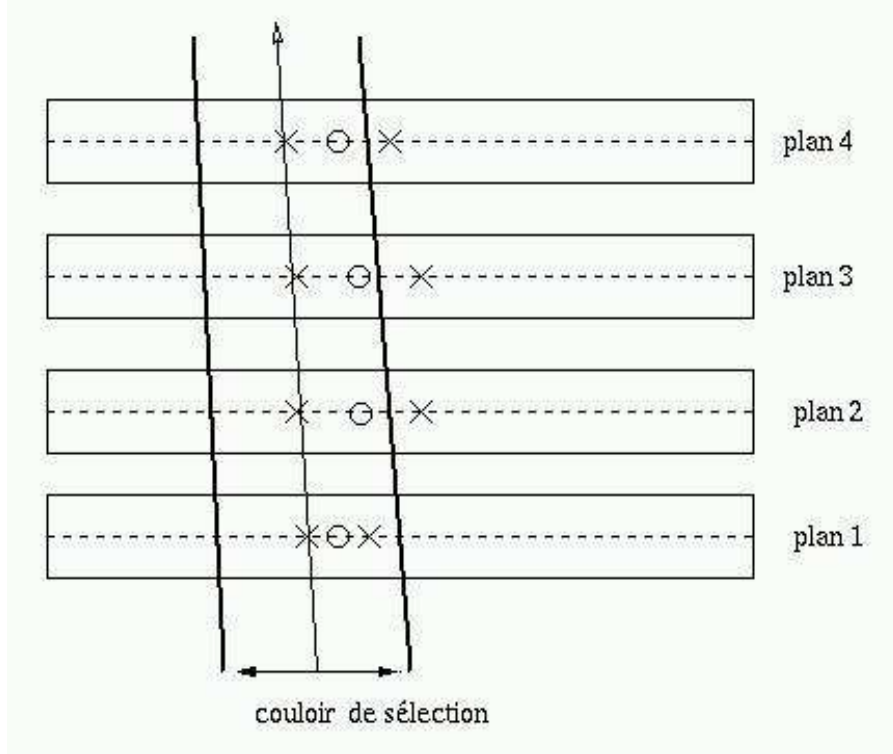


FIG. 4.14 – Création d'un segment en considérant tous les signes possibles : Les fils sont représentés par des ronds et les positions possibles d'un hit par des croix. Pour les plans 2 et 3, il n'existe qu'une possibilité pour les hits dans le couloir de sélection.

Recherche de segments

Les segments de projections sont comme leur nom l'indique des segments dans un espace à 2 dimensions. Le but maintenant est d'utiliser les segments bidimensionnels des projections U, X, et V afin de reconstruire des segments tridimensionnels. L'association des segments est soumise à une coupure de sorte que la somme des χ^2 de chaque segment de plan utilisé soit inférieure à la limite suivante :

$$\chi_{MAX}^2 = \chi_{dof}^2 \times (N_{hits} - 4)$$

où χ_{dof}^2 est le χ^2 par degré de liberté que l'on a fixé à 10.

Pour augmenter l'efficacité de reconstruction, on a aussi inclus la possibilité de construire un segment 3D à partir de 2 segments 2D et de relâcher partiellement la contrainte et la redondance induite par un 3^{ème} segment en recherchant des hits dans la troisième projection pour affiner l'ajustement des paramètres.

Ajustement des paramètres

La reconstruction d'un segment 3D est un processus établi sur 2 itérations. Lors de

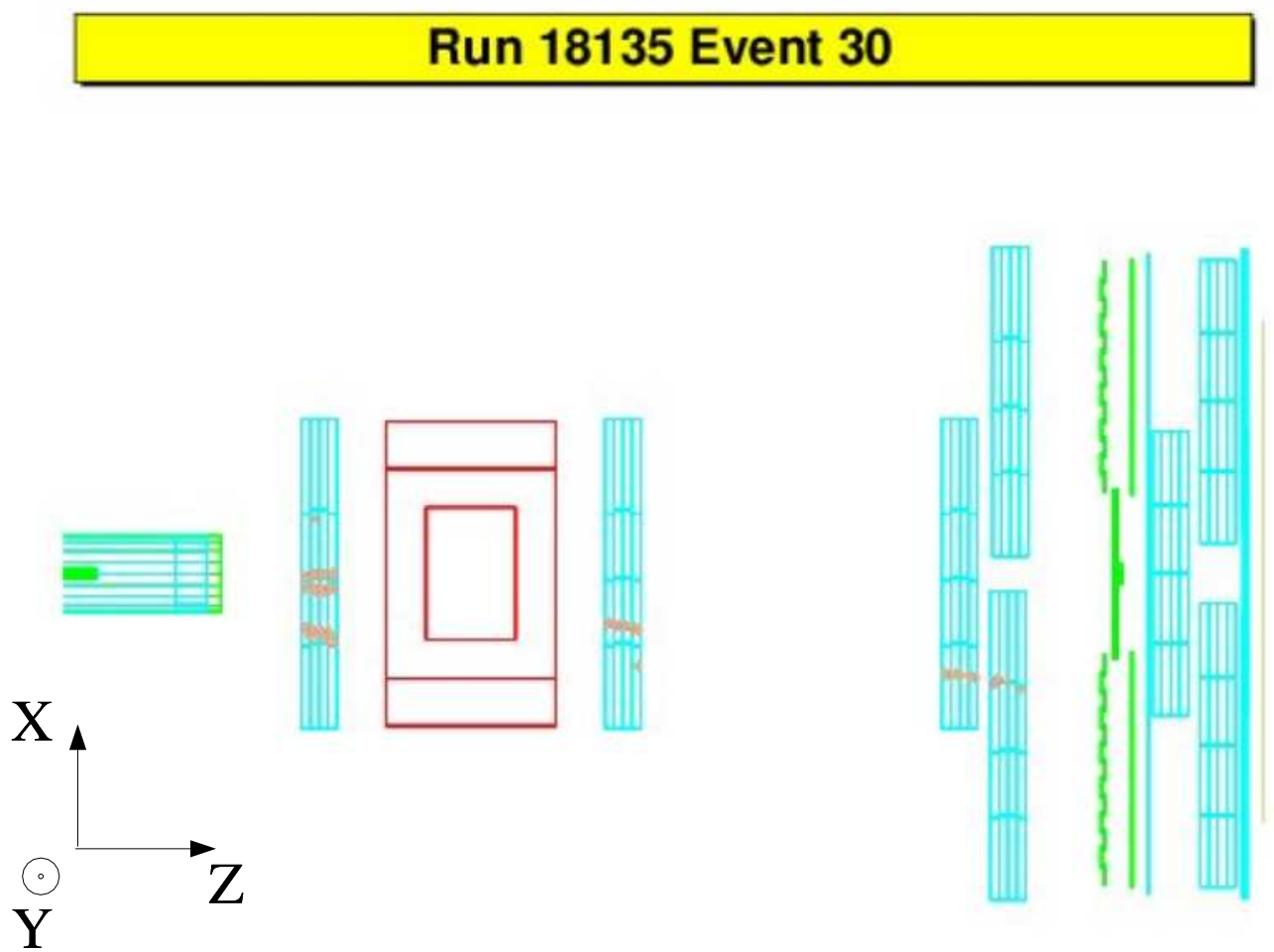


FIG. 4.15 – Hits laissés dans les chambres à dérive lors d'un événement (vue de dessus).

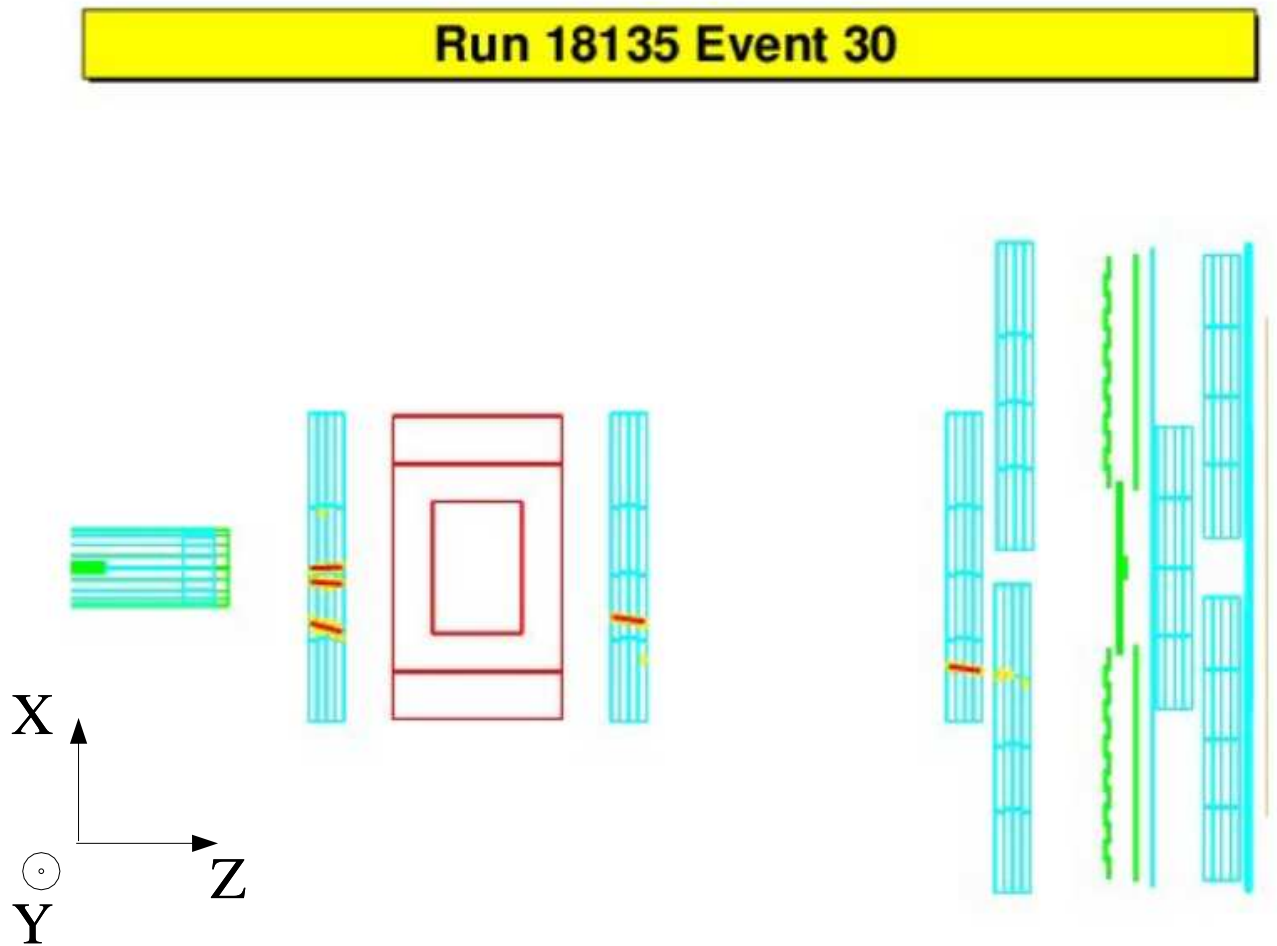


FIG. 4.16 – *Evènement reconstruit : des segments ont été produit en analysant les informations données par des hits (vue de dessus).*

la première itération, les hits ne possèdent pas d'information sur l'angle de la trace (dans le plan de dérive des électrons) et sur la coordonnée Y (position le long du fil). On se souvient que le calcul de la distance de dérive à partir du temps dépend de ces 2 paramètres. Le segment obtenu lors de la première itération possède une première estimation de l'angle et de la position en Y. La 2^{ème} itération consiste alors à recalculer les propriétés des segments à partir des distances de dérive des hits remises à jour : on obtient alors une véritable information 3D du passage d'une particule. Sur les figures 4.17 et 4.18, on peut se rendre compte que la prise en considération de la coordonnée en Y se reflète essentiellement dans une approche plus réaliste quant à la structure des fils.

Les figures 4.15 (position des hits dans un événement) et 4.16 (segments obtenus avec ce même événement) nous montre à première vue que l'on arrive à obtenir un nombre de segments assez important, excepté pour le module de droite où la disposition des hits n'a pas pu aboutir à un segment.

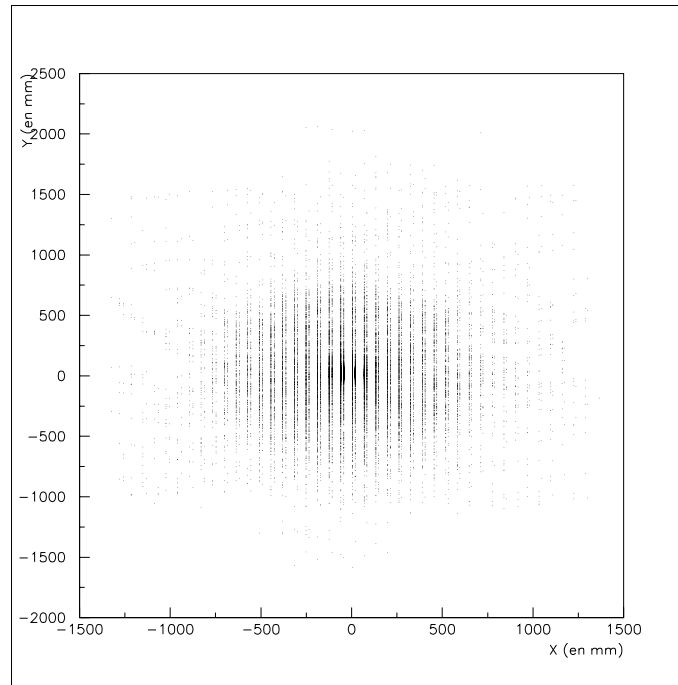


FIG. 4.17 – *Reconstruction de la position des fils lorsque l'information en Y n'est pas incluse (dimensions en mm).*

4.4.4 Performances de l'algorithme de reconstruction.

L'algorithme de reconstruction des segments 3D, bien que n'étant pas encore abouti et optimisé, est opérationnel et nous allons effectuer une première évaluation de ses performances. Pour cela, je vais l'utiliser avec des données Monte-Carlo produites dans différentes configurations. Nous commencerons avec un cas simple basé sur un faisceau de muons couvrant une grande partie de l'acceptance des chambres à dérive. Puis ensuite sur une situation plus réaliste avec des interactions de protons sur cible.

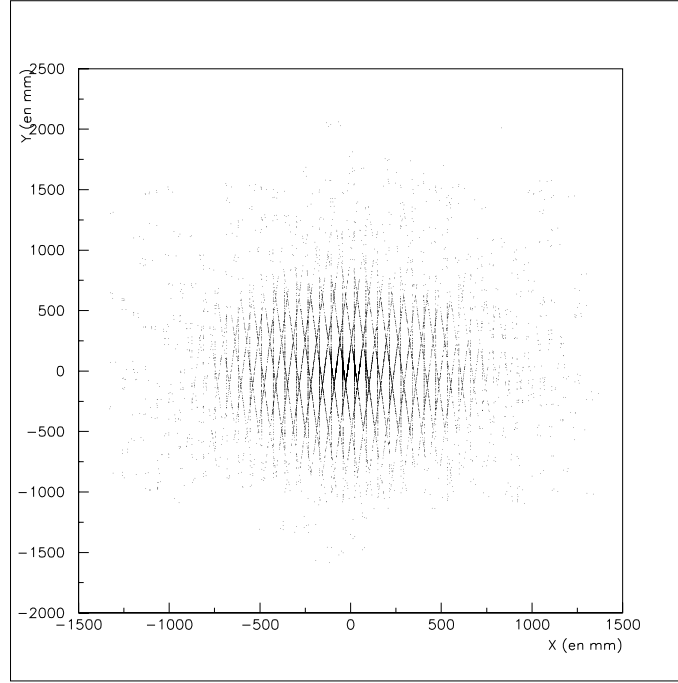


FIG. 4.18 – *Reconstruction de la position des fils lorsque l'on inclut l'information en Y (dimensions en mm).*

4.4.4.1 Conversion des temps en distances.

J'ai simulé le passage de muons dans les chambres et obtenu pour chaque coup des distances et des temps de dérive Monte-Carlo. La figure 4.20 nous montre par une droite de type $y=x$ que la relation pour transformer le temps en distance est satisfaisante. On peut ainsi retrouver une différence entre les 2 valeurs conforme avec ce qui était attendu (exemple : cas d'une résolution simulée fixe de $\sigma_{MC} = 0,5$ mm, figure 4.19).

4.4.4.2 Recherche des segments.

Un premier indicateur de fiabilité est la mesure de la *pureté* d'une trace. La pureté est la proportion de hits d'un segment produits par une particule. La particule de référence étant celle qui a le plus de hits dans un segment. Ceci arrive notamment lorsque des traces sont très proche ou que des électrons arrachés à la matière produisent des coups dans la même région. Dans le cas des muons, on s'attend à observer des traces de grande pureté (figure 4.21) : la simulation n'inclue pas ici de bruit électronique, et les seuls hits "parasites" peuvent être les quelques électrons δ . Le nombre de hits par segment évolue naturellement suivant l'efficacité des chambres (figures 4.21, 4.22, 4.23).

4.4.4.3 Efficacité de reconstruction

L'efficacité de reconstruction est définie comme étant le rapport du nombre de segments reconstruits sur le nombre de traces laissées par le passage de particules dans les modules.

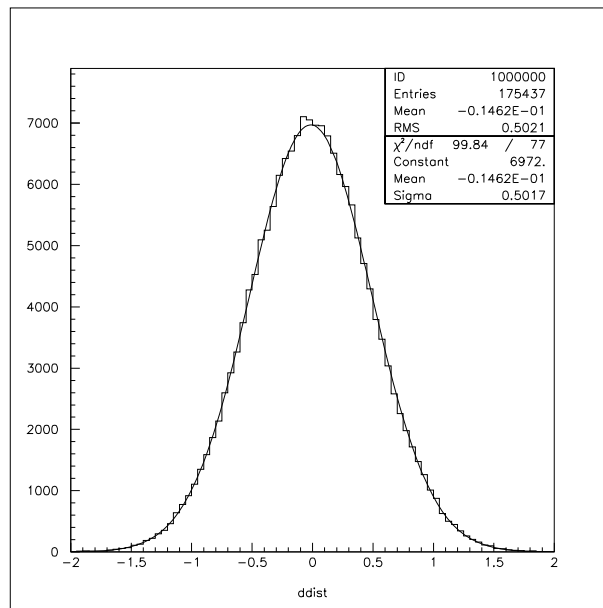


FIG. 4.19 – Différence (en mm) de la distance de dérive reconstruite et simulée. Événements simulés en considérant une résolution en position dans le plan perpendiculaire au fil de $500 \mu\text{m}$.

Identification de segments restructuribles

Fondamentalement, il n'existe pas dans le programme d'objets ou entité représentant la trace Monte-Carlo parfaite. On procède à l'identification des traces Monte-Carlo en triant les numéros d'identification (un par particule). On considère que l'on peut avoir un segment reconstruit lorsque l'on a au moins 8 hits par numéro d'identification. Pour accroître l'assurance d'obtention d'une trace restructurable, on vérifie que ces hits sont distribués correctement dans les différents plans (rejetant des accumulations de hits dans un plan). Par exemple, il est possible d'avoir 8 hits pour une particule donnée avec toutefois 2 hits dans un même plan : dans ce cas seulement 7 plans sont soumis à une reconstruction ce qui est insuffisant.

Mesure de l'efficacité de reconstruction

On compare ces traces Monte-Carlo avec les segments que l'on a obtenu par la reconstruction en comptant les hits Monte-Carlo en commun. La comparaison n'a de sens que si l'on prend en compte des segments qui ont une grande pureté. En considérant qu'une pureté de 80 % est suffisante, on établit une courbe d'efficacité en mesurant pour des efficacités simulées des chambres, les efficacités de reconstruction obtenues.

Résultats

La figure 4.24 nous montre donc que la courbe d'efficacité s'éloigne rapidement de la courbe asymptotique avec la chute de l'efficacité des chambres. En effet si l'on considérait une efficacité moyenne pour les chambres à dérive entre 80 et 90 %, on a alors une effica-

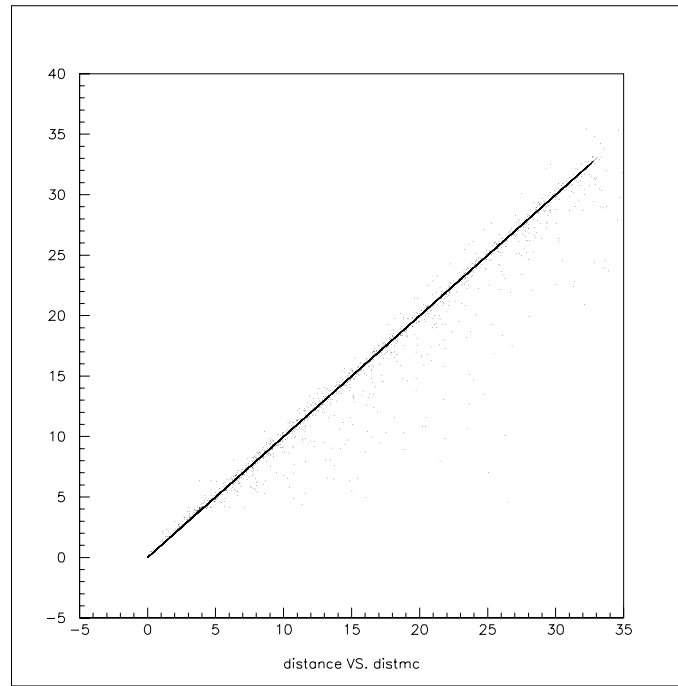


FIG. 4.20 – *Distance traduite du temps simulé vs distance Monte-Carlo simulée.*

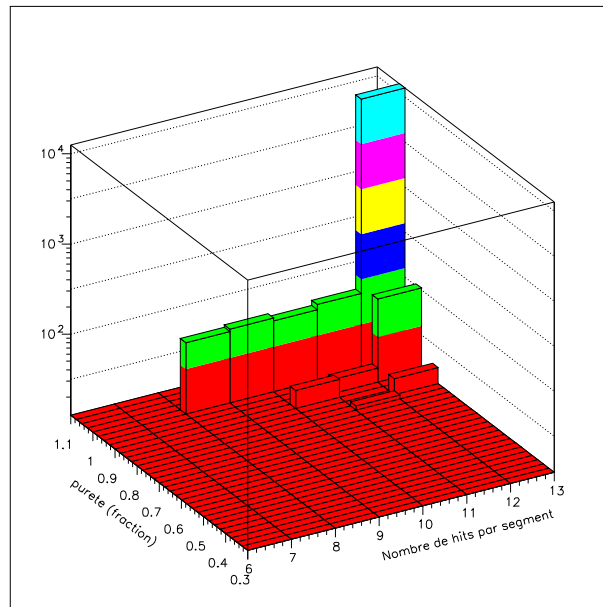


FIG. 4.21 – *Pureté des segments et nombre de hits par segment pour une efficacité de 100 % des chambres à dérivate.*

cité de reconstruction variant entre 75 et 95 %, ce qui souligne cette dégradation rapide. Une amélioration de l'algorithme significative peut être envisagée étant donnée la plage de 10 % restante entre reconstruction effective et asymptotique pour une efficacité des chambres de 80 %.

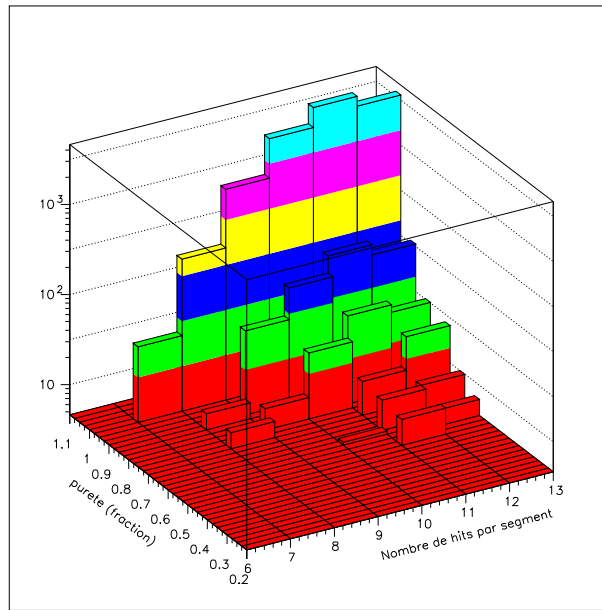


FIG. 4.22 – *Pureté des segments et nombre de hits par segment pour une efficacité de 90 % des chambres à dérive.*

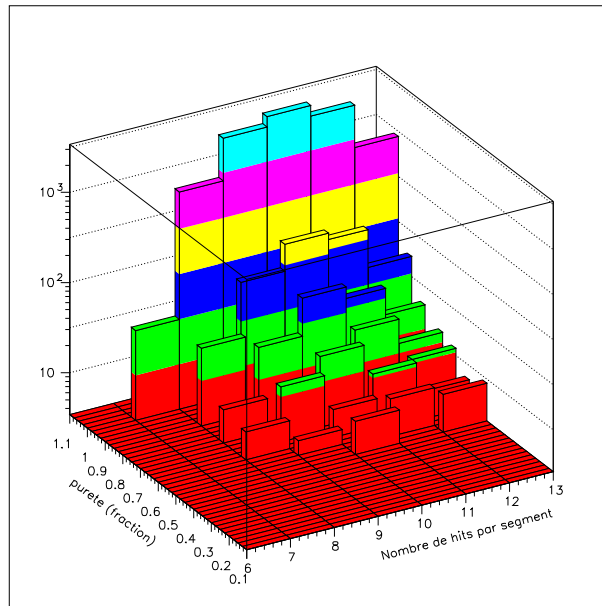


FIG. 4.23 – *Pureté des segments et nombre de hits par segment pour une efficacité de 80 % des chambres à dérive.*

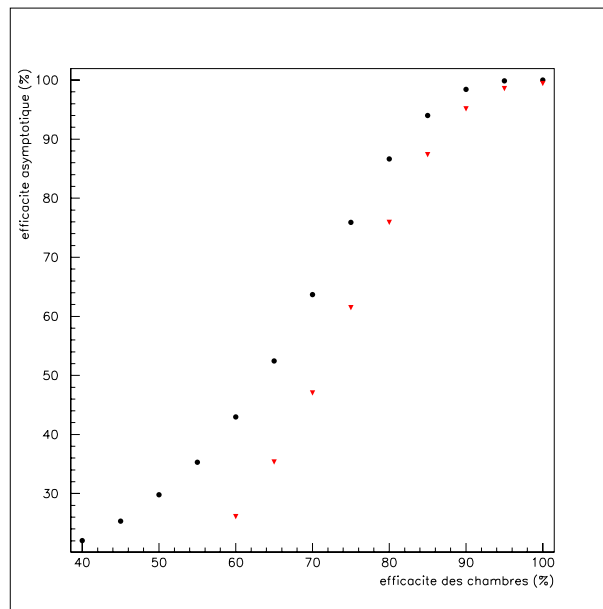


FIG. 4.24 – *Efficacité de reconstruction en fonction de l'efficacité des chambres avec une résolution dans l'axe de dérive de 0,5 mm (triangle) - L'efficacité asymptotique de reconstruction est représentée avec des ronds.*

Quatrième partie

Performances des chambres à dérive

Chapitre 5

Alignement

Avant de pouvoir tester sur les données physiques la qualité des algorithmes de reconstruction décrits au chapitre précédent, une étape importante pour notre groupe a été de procéder à l'alignement des chambres à dérive. Cet alignement est la procédure qui permet de connaître de manière précise la position de chaque fil, élément nécessaire pour une bonne interprétation des données. Elle permet aussi de déterminer pour chaque fil le temps de transit du signal dans les câbles et l'électronique (t_0), qu'on a appelé t_{cable} dans le chapitre précédent et pour chaque plan, les corrections à la vitesse de dérive utilisée dans la relation temps-distance.

5.1 Pourquoi effectuer un alignement ?

Un hit est le signal laissé par une particule traversant une zone sensible. Concrètement cela se traduit par un numéro de fil et un temps. Ce temps, une fois traduit en distance, représente par convention la position, par rapport au fil, de la particule dans le plan médian. La précision de la reconstruction des impacts des particules dans les chambres à dérive conditionne la qualité de la reconstruction des traces et leur résolution en impulsion. Cela est très important pour la qualité des résultats de HARP.

Deux éléments contribuent donc à une connaissance précise du point d'impact de la particule :

- la position du fil : dans le plan médian du volume gazeux : le tissage des fils étant une opération manuelle, la précision de leur positionnement selon la coordonnée X n'est pas absolue ; de plus lors du fonctionnement des chambres, les fils sont soumis à des champs électrostatiques ce qui peut les déformer. La procédure d'alignement a pour but d'obtenir les positions réelles de ces fils dans la direction de dérive des électrons lors du fonctionnement du détecteur.
- La relation temps-distance qui permet de calculer la distance de dérive. Elle se base sur un modèle qui décrit la dérive des électrons (cf figure 4.5).

5.2 Quelques notions

Pour comprendre la procédure de l'alignement, explicitons un certain nombre de notions :

Résidu : Il représente, dans un plan donné, l'écart entre le point de mesure (m) (position du hit utilisé par la trace dans ce plan) et la position de la trace reconstruite (e) :

$$r = m - e \quad (5.1)$$

C'est en effet l'étude des distributions des résidus qui va permettre d'évaluer les corrections aux différents paramètres des chambres à dérive.

La position du hit est donnée par :

$$m = u_0 + s \cdot d \quad (5.2)$$

où u_0 est la position du fil selon une direction orthogonale aux fils de ce plan (repère "tilté") ;

s est le signe, positif (vers les X positifs) si l'on se situe à gauche et négatif (vers les X négatifs) à droite du fil ;

d est la distance de dérive.

t_0 : il s'agit du temps de transit du signal entre le préamplificateur et le TDC. Il dépend donc principalement de la longueur des câbles et peut varier d'un fil à l'autre. L'autre composante minoritaire est le temps de passage dans les circuits électroniques. Schématiquement, le t_0 sera évalué pour chaque fil comme le temps à soustraire au temps du hit pour qu'une particule qui "touche" le fil ait un temps de dérive nul.

5.3 Procédure d'alignement

Cette procédure consiste en 2 étapes principales :

- Une première étape qui permet dans chaque module de connaître la position des fils, leurs t_0 et les corrections à la vitesse de dérive. Cette première étape, du fait de l'algorithme des segments utilisé, est insensible à un déplacement global du module.
- Une deuxième étape qui permet de repérer les modules les uns par rapport aux autres et par rapport aux autres sous-détecteurs de l'expérience.

5.3.1 Position du problème.

Les conditions expérimentales de HARP paraissent simples à première vue cependant les points de mesures que sont les chambres à dérive, sont dispersés ce qui pose problème. Mettre en œuvre l'alignement des chambres à dérive signifie parvenir à connaître la position des 5 points de colle de chaque fil (la forme d'un fil est approximé par une ligne brisée de 4 segments, cf §3.4). A partir de fils initialement rectilignes, nous devons obtenir les corrections sur la position en X des points de colle par l'intermédiaire de distributions

des résidus. Il faut donc beaucoup de statistique le long de chaque fil.

Il était prévu d'utiliser une procédure du même type que celle de l'expérience NOMAD en utilisant des muons dits "hors temps". Ces muons proviennent du halo du faisceau et ont l'intérêt d'arroser de manière suffisamment uniforme les chambres à dérive pour toucher tous les fils. Mais le trigger de l'expérience n'a pas permis de les sélectionner.

Nous avons donc dû utiliser les données de physique de notre expériences pour effectuer cet alignement. Le problème est que la plupart des traces sont dirigées vers l'avant : il est donc évident que les fils excentrés ne seront guère touchés (voir figure 5.1). Le problème est accentué pour les modules situés après le dipôle magnétique : les parties extrêmes des fils sont situées dans "l'ombre" du fer de l'aimant comme on peut le voir sur la figure 5.2 pour le module 2 et sur la figure 5.3 pour le module 5 qui est situé à environ 3,4 m en arrière. Néanmoins on peut considérer que les zones les plus touchées sont les zones "utiles" des chambres, donc celles qu'il faut corriger.

Nous avons choisi d'utiliser, pour la procédure d'alignement, les données d'interactions des protons de 12.9 GeV/c sur une cible épaisse (1 λ d'Aluminium, cible de K2K)

Nous avons par ailleurs enregistré des événements dits "cosmiques" tout au long de l'expérience et l'échantillon des plus horizontaux d'entre eux représente un complément utile à cette étude.

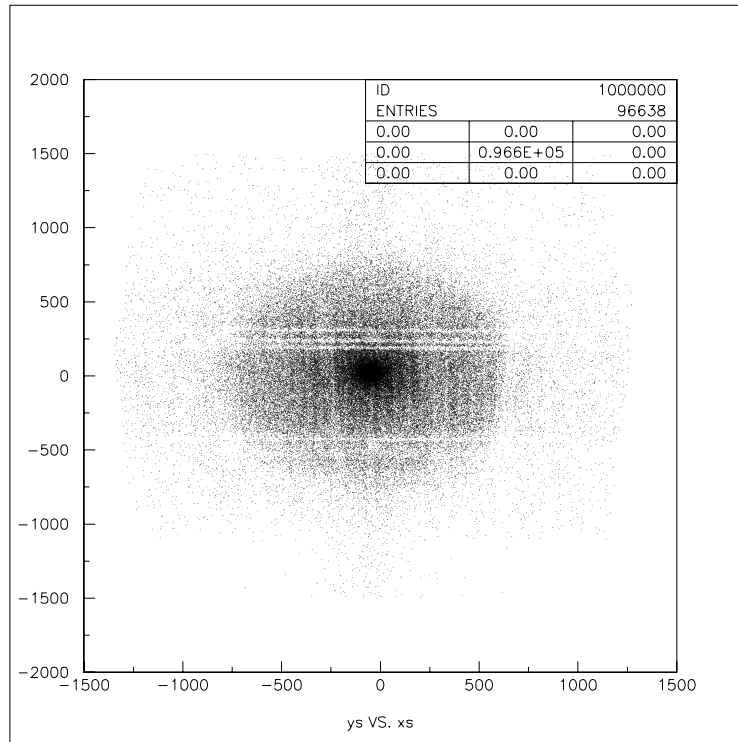


FIG. 5.1 – *Distribution des coups en x,y pour le module 1 des chambres à dérive. La distribution est "conique" avec le maximum dans l'axe du faisceau. Les traits horizontaux sont dûs aux inefficacités causées par les barrettes de supports des fils.*

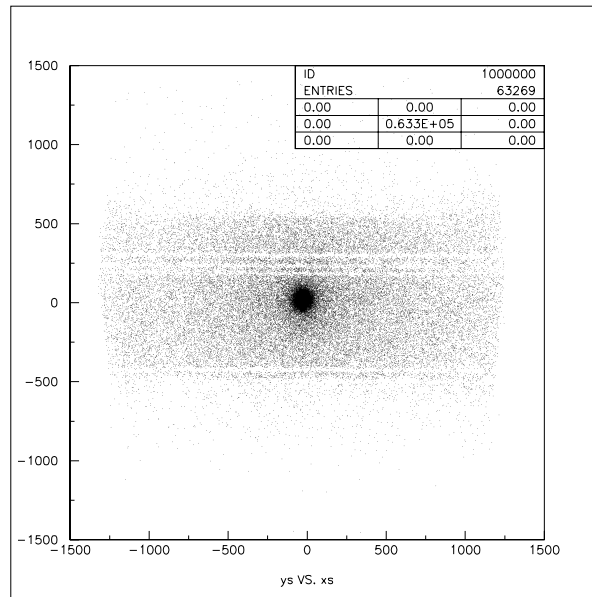


FIG. 5.2 – *Distribution des coups en x,y pour le module 2 des chambres à dérivation. Une coupure au niveau des parties supérieures et inférieures est à noter : il s'agit de l'influence du dipôle magnétique qui diminue l'acceptance des chambres à dérivation en aval de l'aimant.*

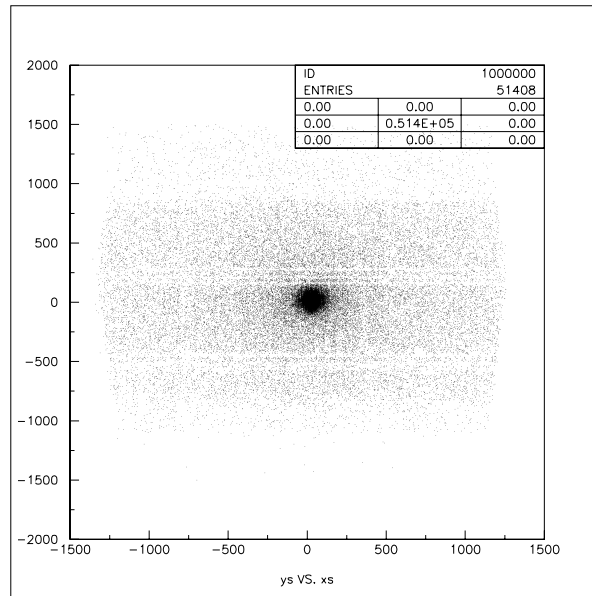


FIG. 5.3 – *Distribution des coups en x,y pour le module 5 des chambres à dérivation. L'effet d'acceptance du dipôle magnétique est encore visible.*

5.3.2 Alignement interne

5.3.2.1 Effet des différents paramètres sur les distributions des résidus.

Le résidu représente, dans un plan donné, l'écart de position entre une trace reconstruite et le hit associé. On s'attend à ce que sa distribution soit une gaussienne centrée en zéro et de largeur égale à la résolution spatiale des chambres à dérive (figure 5.4).

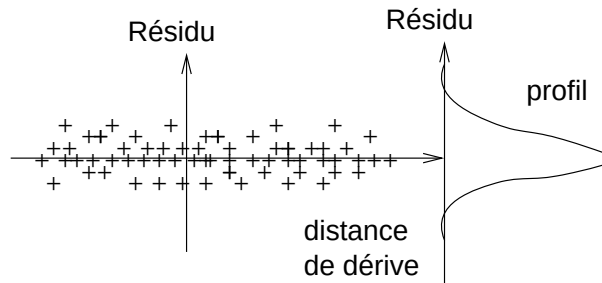


FIG. 5.4 – Distribution des résidus en fonction de la distance de dérive dans le cas idéal.

Effet de la position des fils de lecture.

La distribution moyenne des résidus peut s'écrire :

$$\langle r \rangle = \langle u_0 - e \rangle + \langle s \times d \rangle \quad (5.3)$$

Le terme $\langle s \times d \rangle$ est en moyenne nul, la cellule étant touchée uniformément de chaque côté de son fil. La moyenne n'est donc pas sensible, à la distance de dérive. Une erreur systématique sur la position du fil U_0 va se traduire par un décalage de la valeur moyenne des résidus (figure 5.5).

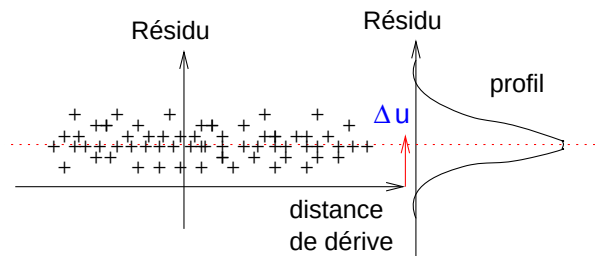


FIG. 5.5 – Distribution des résidus en fonction de la distance de dérive dans le cas où le fil est décalé de Δu par rapport à la position attendue.

Effet du temps propre des fils.

La valeur moyenne des résidus signés peut s'écrire :

$$\langle sr \rangle = \langle su_0 - se \rangle + \langle d \rangle \quad (5.4)$$

Ici pour la même raison, c'est le terme $\langle s(u_0 - e) \rangle$ qui est en moyenne nul. Par contre une mauvaise estimation du t_0 va se traduire par un offset systématique du temps de dérive et donc de la distance de dérive. La distribution des résidus signés est décalée, et la distribution des résidus comporte 2 pics symétriques (figure 5.6).

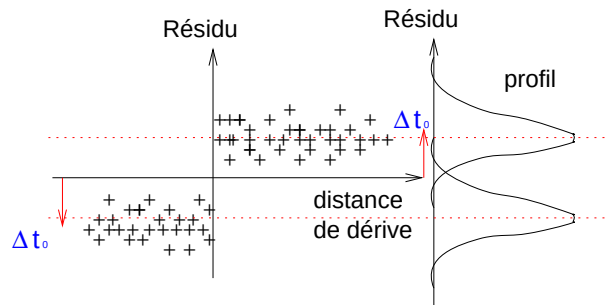


FIG. 5.6 – Distribution des résidus en fonction de la distance de dérive dans le cas où le temps propre du fil est erroné d'une quantité Δt_0 . On obtient dans ce cas une structure des résidus suivant deux pics.

Effet de la vitesse de dérive.

On peut voir sur l'expression de la moyenne des résidus signés, que si la vitesse de dérive, qui sert à calculer la distance de dérive (schématiquement $d = v_d \times t_{dérive}$), est différente de la vitesse de dérive réelle, la distribution de cette moyenne sera linéaire en fonction du temps de dérive. Sa pente est l'écart entre les 2 vitesses de dérive (figure 5.7). Dans la réalité, d'une part la relation temps-distance globale est plus compliquée qu'une simple relation linéaire (cf §4.3.1), et d'autre part, il faut tenir compte d'inhomogénéités du champ électrostatique qui entraînent des distorsions locales de la vitesse de dérive. On envisagera donc plutôt des corrections linéaires par morceau.

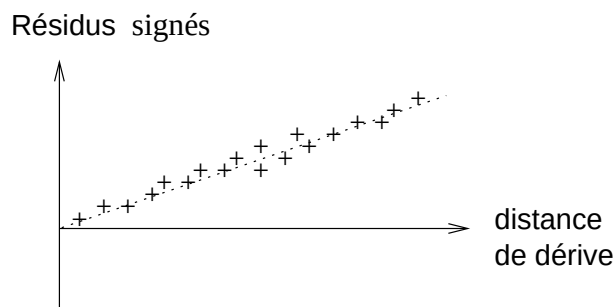


FIG. 5.7 – Distribution des résidus signés en fonction du temps de dérive lorsque la vitesse de dérive utilisée dans la reconstruction est différente de la vitesse de dérive réelle. La pente de la distribution est la différence des vitesses.

5.3.2.2 Comment corrige-t-on concrètement tous ces effets ?

La correction des paramètres que l'on a vus précédemment repose sur un processus itératif. Un grand nombre de traces reconstruites est nécessaire afin d'avoir une statistique suffisamment importante sur chaque fil pour en extraire des corrections significatives.

Après chaque itération, on corrige les paramètres en vérifiant la convergence des termes correctifs et le rétrécissement de la distribution des résidus.

Ajustement des t_0

Si l'ajustement du t_0 de chaque fil est mal pris en compte, l'effet est spectaculaire sur la distribution des résidus (figure 5.9), avec cette structure à 2 pics. La correction passe donc par le calcul, pour chaque fil, de la valeur moyenne du résidu signé. C'est ce qui est fait pour chaque plan de chambre à dérive en produisant un histogramme du profil des résidus signés en fonction du numéro de fil. Trois ou quatre itérations en moyenne sont nécessaires pour converger (figure 5.11).

La distribution des écarts entre les valeurs initiales et finales des t_0 est montrée (figure 5.8) : sa valeur moyenne 5 ns, confirme la nécessité de bien prendre en compte ce paramètre dans la reconstruction (cela correspond à une erreur sur la distance de $v_d \times 5$ ns = 250 μ m).

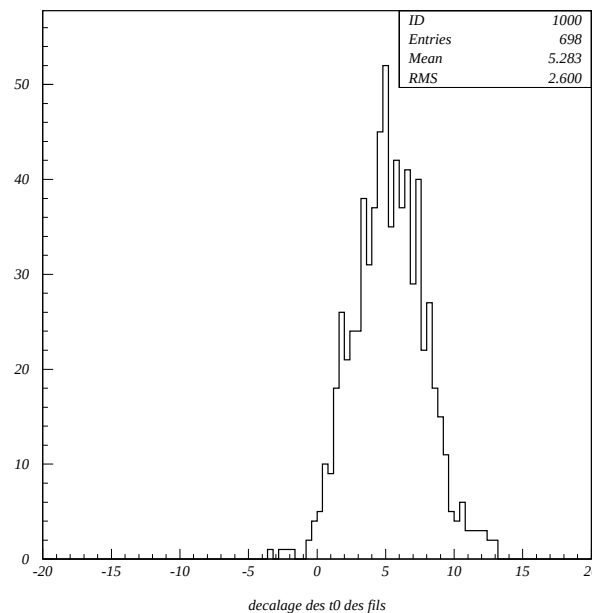


FIG. 5.8 – *Distribution des différences de t_0 avant et après correction (en ns).*

5.3.2.3 Corrections sur la position des fils.

On a vu que l'on pouvait représenter chaque fil sous la forme de 4 segments de droite. L'alignement doit donc préciser la position de chacun de ces points et non pas seulement des extrémités. Cela nous fait 5 points à ajuster pour chaque fil soit 12900 points pour l'ensemble des chambres à dérive.

Le principe est de partir d'un fil idéal rectiligne, dont la position est déduite des mesures des géomètres, pour arriver à un profil plus réaliste du fil en corrigeant la position de ces 5 points (figure 5.10).

Pour trouver le décalage en position de chaque point, on utilise la distribution des résidus le long du fil, dont la moyenne représente l'erreur sur la position du fil. Dans la

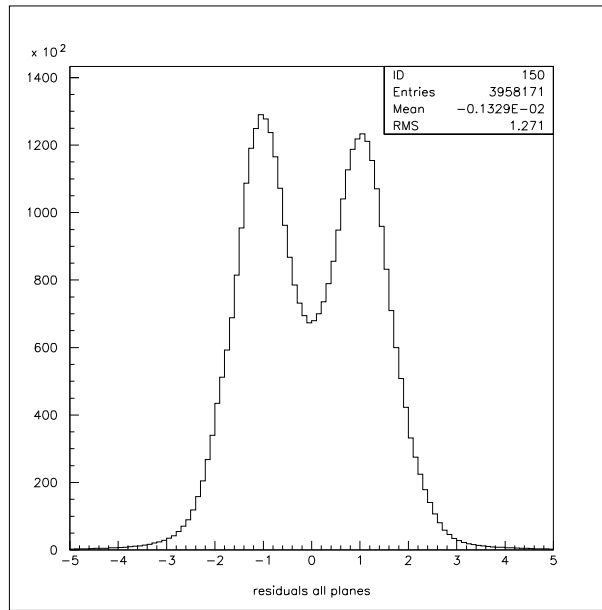


FIG. 5.9 – Distribution des résidus pour tous les plans avant toute correction.

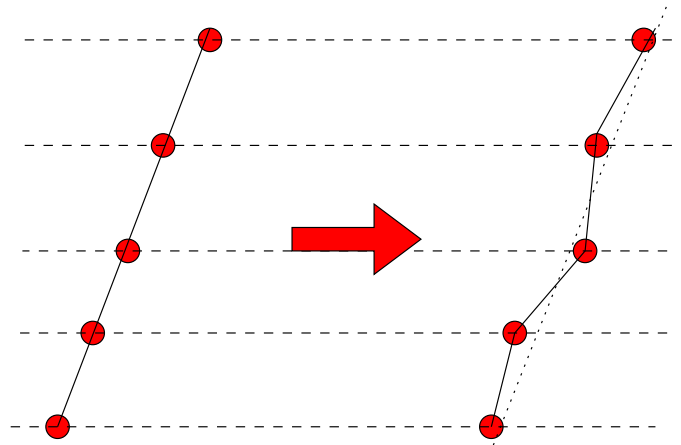


FIG. 5.10 – Ajustement de l'allure d'un fil avant (à gauche) et après la correction (à droite).

pratique, on réalise sur les données des résidus de chaque fil un ajustement de 4 segments de droite jointifs.

Pour éviter les biais dus à une trop faible statistique, on a imposé au moins 500 coups sur chaque fil traité. De plus, pour tenir compte de la coupure d'acceptance due à l'aimant en particulier pour le module 2, on a considéré comme fixes les points extrêmes quand la statistique était insuffisante. Là encore 3 à 4 itérations sont nécessaires pour converger. La figure 5.12 montre la distribution des corrections en X sur les positions des 5 points fixes de chaque fil. On voit là encore que des écarts non négligeables par rapport aux positions théoriques sont observés.

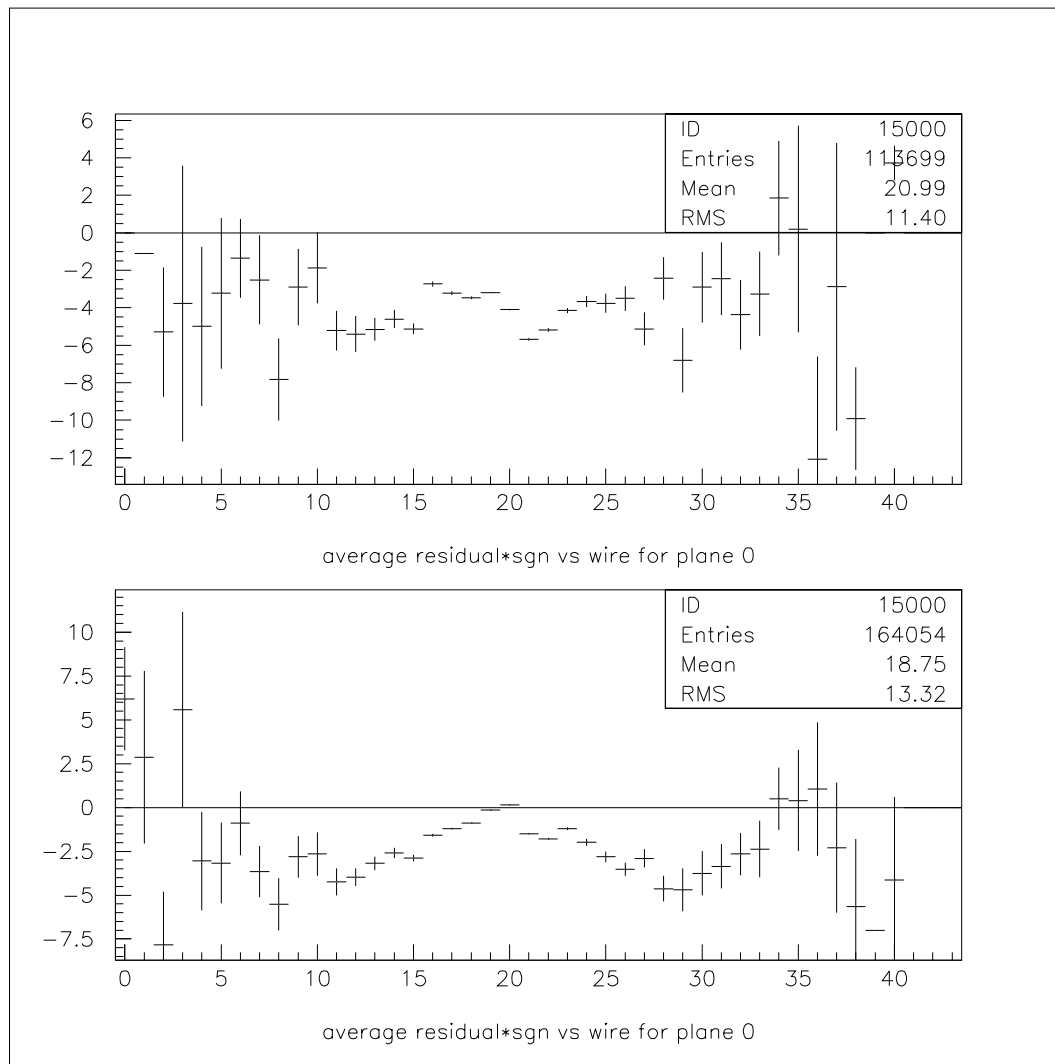


FIG. 5.11 – Profil des résidus signés divisés par la vitesse de dérive en fonction du numéro de fil pour un plan particulier et pour 2 itérations successives.

5.3.2.4 Une relation temps-distance affinée pour chaque plan.

L'effet des inhomogénéités du champ électrique dans les cellules d'un même plan ¹ peut être visualisé sur la figure 5.13 qui montre l'histogramme des résidus signés en fonction du temps de dérive. On constate qu'au lieu d'une droite unique, dont la pente indiquerait la correction à la vitesse de dérive, l'histogramme a plutôt la forme d'une droite brisée où les corrections à la vitesse de dérive apparaissent différentes par tranche de temps de dérive.

On a pu définir pour l'ensemble des 69 plans de dérive 8 zones de temps communes (tableau 5.1), dans lesquelles les distributions de résidus signés peuvent être ajustées par des segments de droite.

Ces corrections par morceaux à la vitesse de dérive sont ensuite intégrées au calcul de la distance de dérive. La convergence est atteinte lorsque d'une itération sur l'autre, la distribution de ces résidus signés en fonction du temps est plate (figure 5.17). La figure

¹le champ électrique a la même forme dans toutes les cellules d'un même plan car il est distribué à partir d'un même pont diviseur.

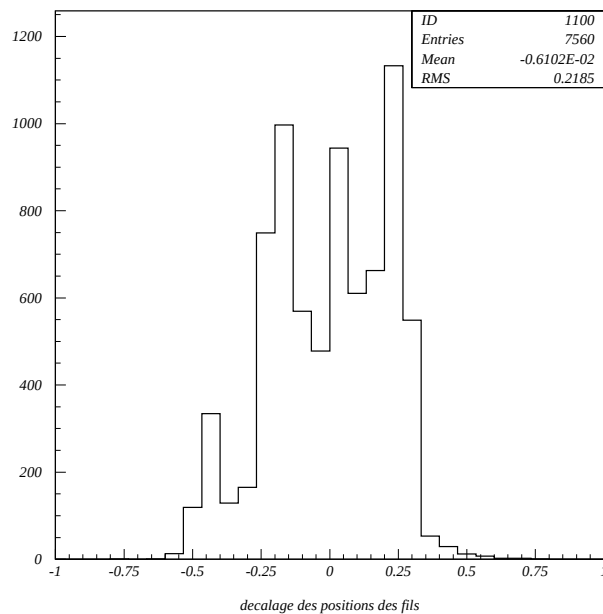


FIG. 5.12 – *Distribution des corrections en X sur les positions des points de colle des fils (en cm).*

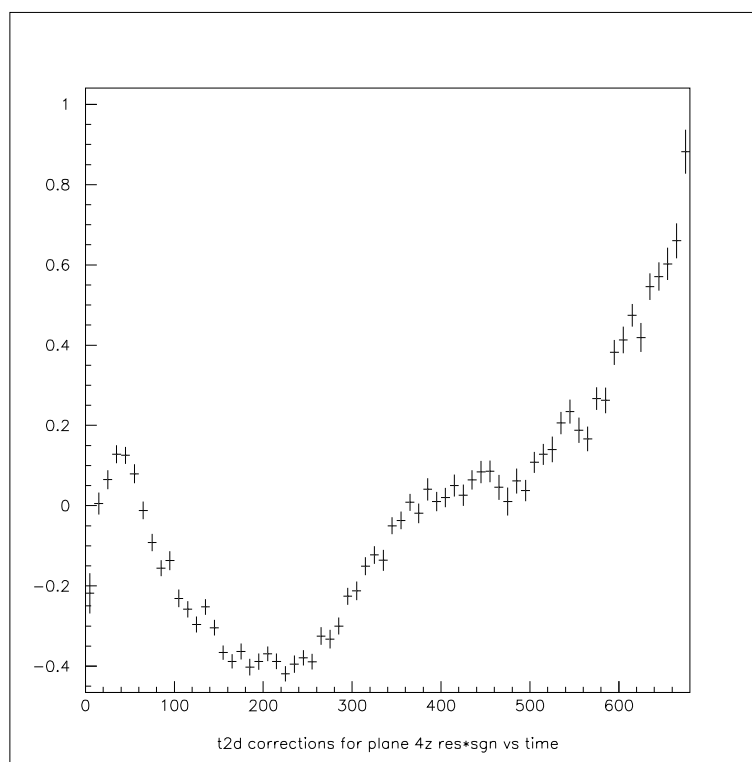


FIG. 5.13 – *Evolution des résidus signés en fonction du temps de dérive dans un plan (abscisse : en ns, ordonnée : en mm).*

5.14 montre la distribution des corrections à la vitesse de dérive sur l'ensemble des plans. Globalement il s'agit d'une correction de l'ordre de 2-3 %.

T1	T2	T3	T4	T5	T6	T7	T8
35	110	210	400	510	580	640	680

TAB. 5.1 – Bornes délimitant les 8 régions de dérive homogène (en ns)

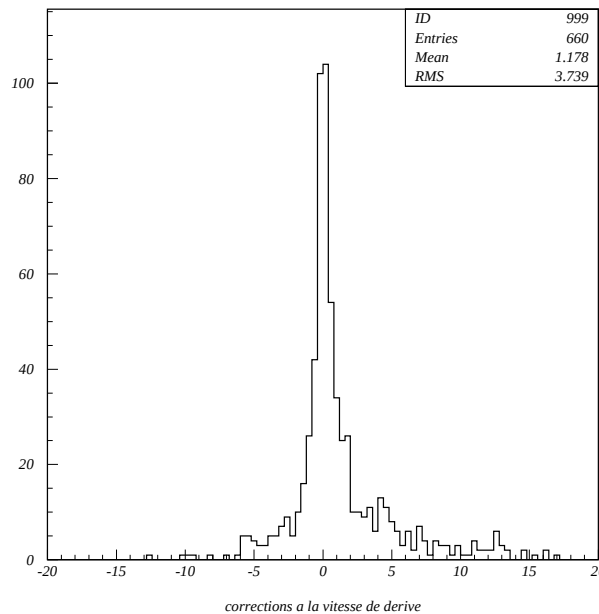


FIG. 5.14 – Correction à la vitesse de dérive (en pourcentage de la vitesse initiale).

5.3.2.5 Les résultats

On peut dans un premier temps voir l'évolution de la distribution des résidus après plusieurs itérations en prenant en compte la correction des t_0 (figure 5.15).

La structure à 2 bosses converge assez rapidement vers une structure à un pic. A partir de là, nous avons un point de départ satisfaisant pour corriger la position des fils.

Après 3 itérations sur 400000 événements, nous obtenons une distribution globale des résidus d'une largeur inférieure à $500 \mu\text{m}$ (figure 5.16). Les queues non gaussiennes sont dues au fait que dans cet histogramme sont cumulées des traces à différents angles et à toutes les distances de dérive. Or la résolution des chambres varie en fonction de l'angle et de la distance de dérive.

Trois autres itérations corrigeant la vitesse de dérive permettent de réduire encore la largeur de la distribution des résidus à moins de $350 \mu\text{m}$, précision amplement suffisante pour l'expérience HARP, puisque c'est l'incertitude due à la diffusion multiple qui dominera.

L'ensemble de la procédure est assez lourd puisqu'il a totalisé plus de 15 jours de CPU sur un PC Pentium III de 800 MHz.

5.3.2.6 Utilisation des fichiers de calibration.

L'ensemble des corrections précédentes est regroupé dans des fichiers de calibration intégrés à la base de données dédiée à la calibration des différents sous-détecteurs. Comme

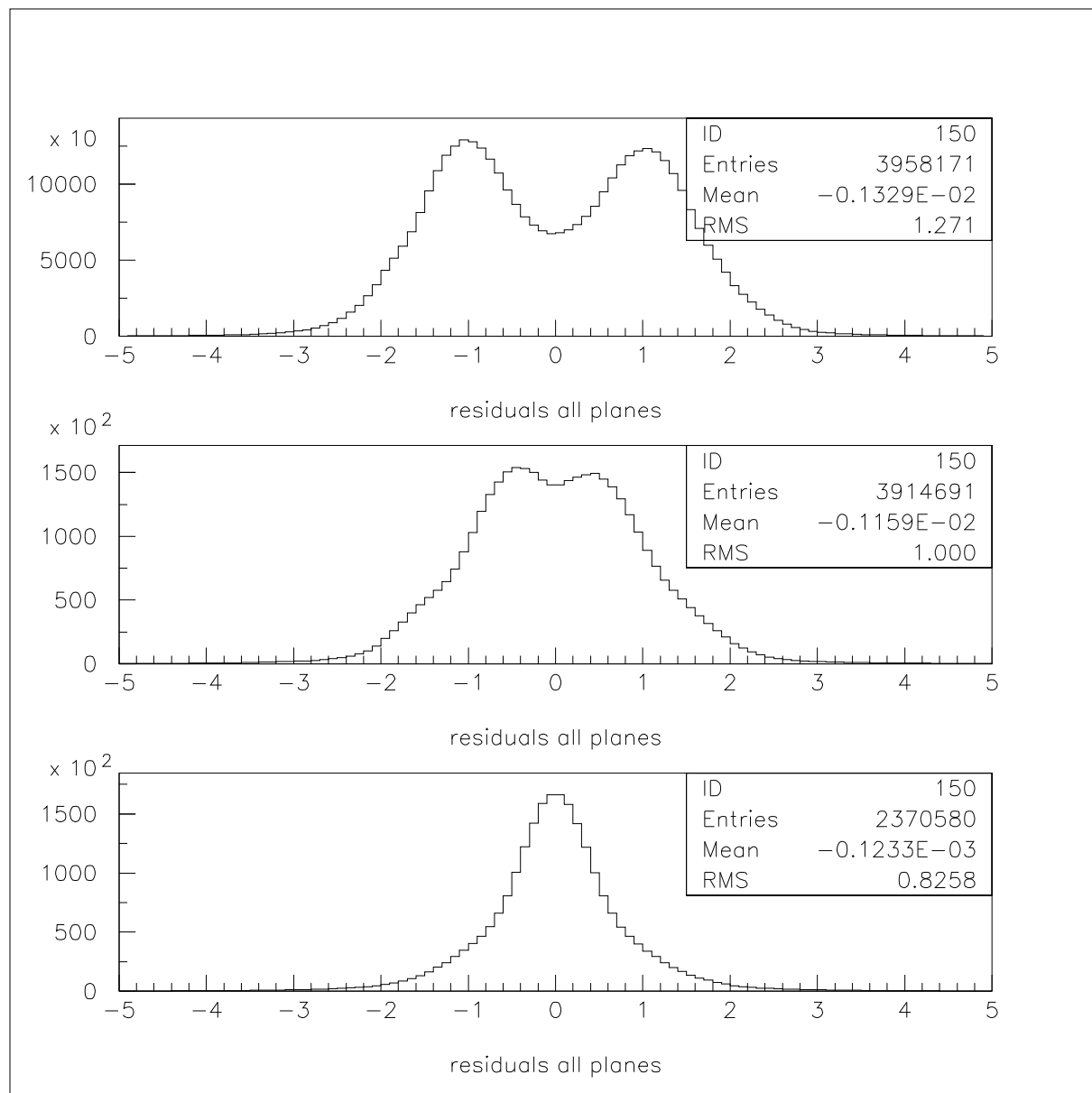
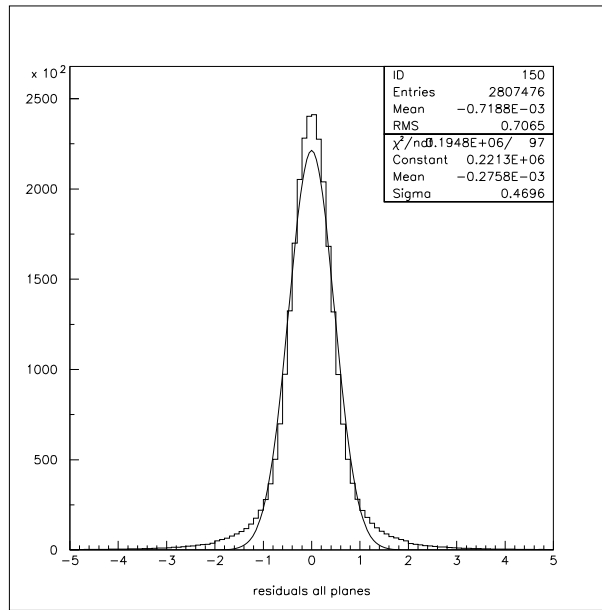


FIG. 5.15 – Evolution de la distribution des résidus (en mm) lorsque l'on corrige le temps t_0 de chaque fil à chaque itération. L'ordre des itérations va du haut vers le bas.

les conditions de prise de données peuvent varier (flux de gaz, température, pression, fuites de gaz) et influencer les performances des sous-détecteurs, on doit vérifier la validité de ces constantes au cours du temps et éventuellement en produire de nouvelles. Actuellement nous disposons d'un jeu de constantes pour 2001 et d'un autre pour 2002.

Une procédure est en cours pour reconstruire quelques millions de traces par jour de prise de données, et surveiller l'évolution de la largeur de la distribution des résidus. Si nécessaire, de nouveaux alignements seront réalisés.

FIG. 5.16 – *Distribution des résidus après correction de la position des fils.*

5.3.3 Alignement inter-modulaire

La procédure précédente repose sur une reconstruction autonome dans chaque module de chambres à dérive. Ce qui signifie qu'elle est insensible à la position absolue des modules : seules les positions relatives des fils dans le repère du module ont été déterminées par cette procédure. En l'absence d'un programme de reconstruction global des traces chargées dans HARP, pour aligner les modules les uns par rapport aux autres, nous avons utilisé des données d'interaction sans champ magnétique dans l'aimant, de façon à disposer de traces droites composées de segments extrapolables d'un module à l'autre.

5.3.3.1 Extrapolation de segments pour l'ajustement inter-modulaire

Méthode

Pour chaque segment d'un module, on effectue des extrapolations vers les autres modules. On calcule la différence de position et d'angle entre ces extrapolations et la position des segments reconstruits dans ces autres modules. Ces différences sont sensibles à la fois à un décalage en position des modules mais également à un décalage angulaire (figure 5.20).

On peut ainsi déduire des extrapolations croisées (figure 5.19) les différents décalages en position et angulaire de chacun des modules par rapport à un module de référence (nous avons choisi le module 2 arbitrairement). La méthode est explicitée dans l'annexe A. Reste enfin à aligner cet ensemble avec le reste de l'expérience : l'ensemble des chambres de faisceaux (MWPC) étant le sous-détecteur de référence, nous avons utilisé un run sans cible et sans champ magnétique pour repérer les écarts de position et d'angle suivant la même méthode.

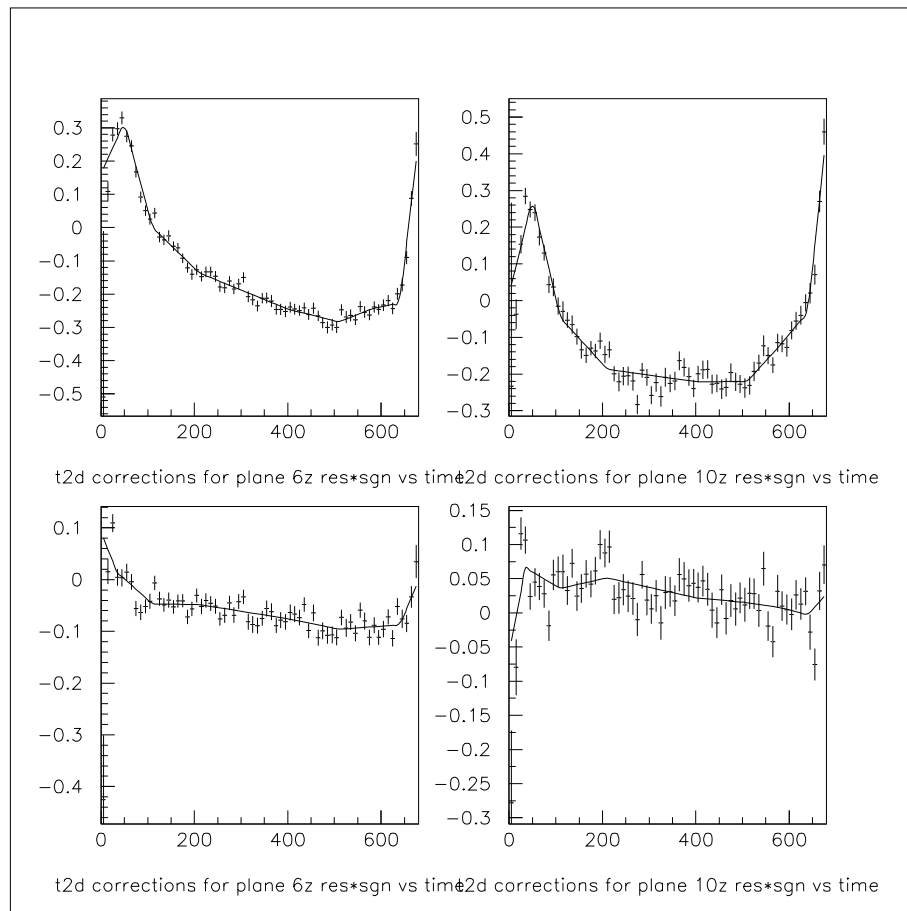


FIG. 5.17 – En haut les résidus signés en fonction du temps de dérive pour 2 plans différents avant correction : les écarts peuvent atteindre 1 mm d’amplitude. En bas, la même chose après correction : les écarts ont été diminués jusqu’à 200 μm d’amplitude.

5.3.4 Quelques remarques de conclusion.

Les informations obtenues lors de ces 2 phases de l’alignement sont stockées dans 2 bases de données. Le premier alignement (intrinsèque) fournit des informations qui évoluent dans le temps s’adaptant ainsi aux conditions de prise de données. En revanche l’alignement inter-modulaire est figé car il dépend uniquement de la position géométrique des modules.

Ces 2 étapes de l’alignement sont nécessaires à l’évolution même du programme de reconstruction afin de rendre plus efficace la phase d’assemblage des segments en une trace permettant la mesure de l’impulsion.

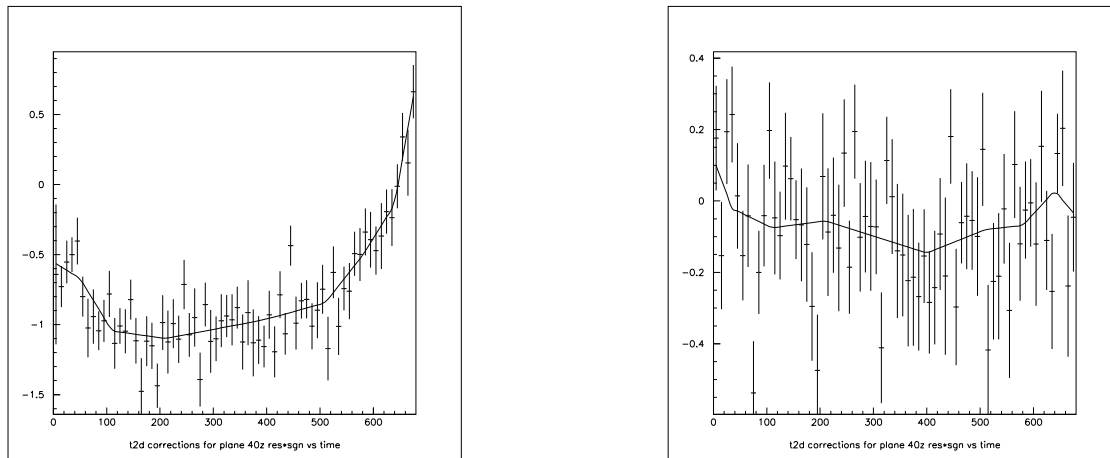


FIG. 5.18 – résidus signés en fonction du temps de dérive dans un plan des modules excentrés. Avant la correction (à gauche) et après la correction (à droite)

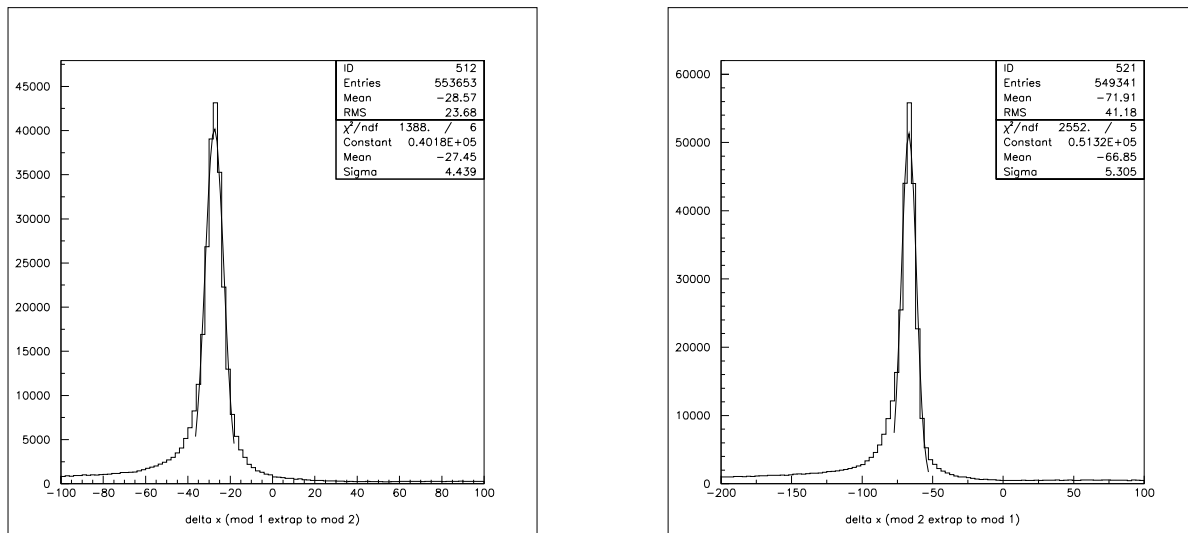


FIG. 5.19 – Différence de position (en mm) entre l'extrapolation d'un segment du module 1 vers le module 2 et le segment reconstruit dans le module 2, à gauche, et l'inverse à droite.

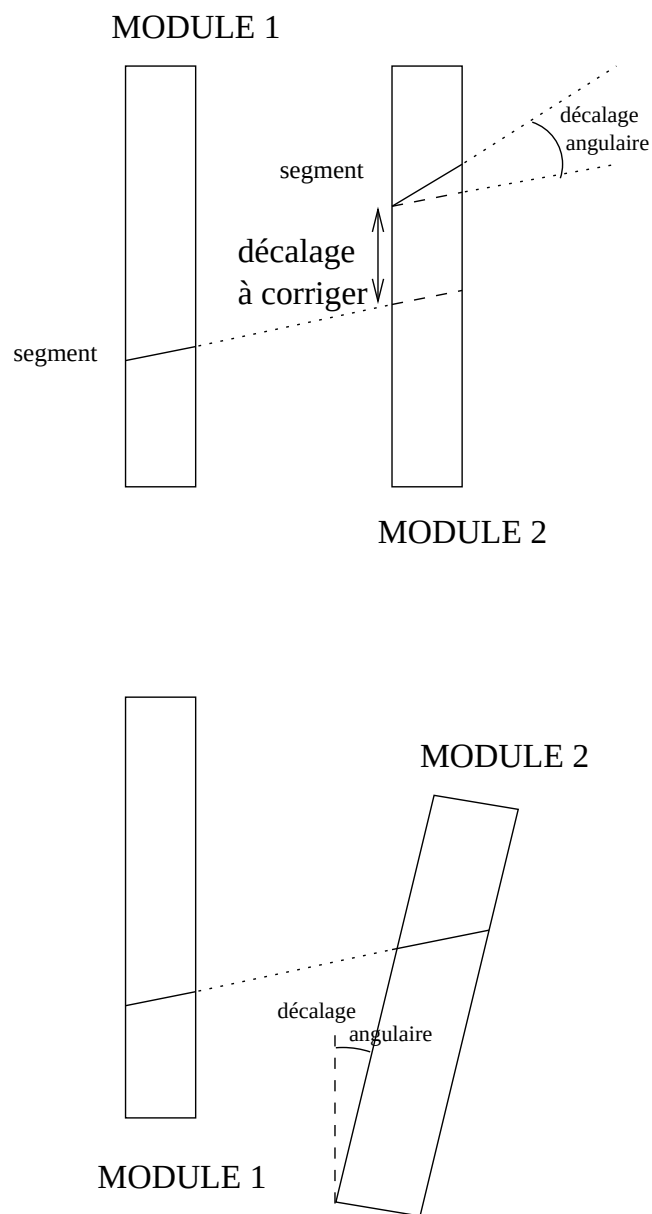


FIG. 5.20 – Effet d'un décalage de translation latérale et de rotation entre 2 modules pour une même trace : les deux segments ne correspondent pas (schéma du haut). La correction d'un module par rapport à l'autre nous permet d'établir la correspondance d'une seule et même trace (schéma du bas).

Chapitre 6

Efficacité des chambres.

Dans le chapitre 4 consacré au programme de reconstruction, nous avons étudié l'efficacité de reconstruction de segments en fonction de l'efficacité de hits des différents plans de chambres. Dans NOMAD, cette efficacité dépassait 90 %.

Nous allons nous intéresser ici à une mesure de l'efficacité des chambres à dérive dans HARP.

6.1 Etude de l'efficacité avec les rayons cosmiques

En l'absence d'un programme complet de reconstruction, les rayons cosmiques sont actuellement le meilleur moyen d'évaluer l'efficacité des chambres. En effet d'une part, à la différence des particules provenant du faisceau, les rayons cosmiques "arrosent" toute la surface des chambres (à l'acceptance du trigger près; voir plus loin). D'autre part, à la fin de la prise de données physiques de HARP (début octobre 2002), une prise de données exclusive de trigger cosmique a été effectuée et a permis d'accumuler près de 2 millions d'événements (dont plus de la moitié est en fait due au halo de muons du PS qui continuait à alimenter d'autres expériences).

6.1.1 Principe de l'étude.

Nous allons chercher à estimer l'efficacité des chambres en considérant que tous les plans doivent avoir au moins un coup au passage d'une particule chargée. L'absence de coup dans un plan est alors jugée comme une inefficacité.

La stratégie adoptée est alors simple. Le système de déclenchement sur les rayons cosmiques utilise une coïncidence entre le FTP, le TOF et le Cosmic Wall (CW). Pour simplifier l'étude, nous avons sélectionné les événements ne contenant a priori qu'une seule trace en demandant les événements pour qui il n'y a qu'un seul coup dans le CW.

Pour s'assurer qu'une trace traverse bien le détecteur, nous demandons qu'il y ait dans l'événement au moins 2 segments compatibles en position et angle dans 2 modules distincts. Nous extrapolons cette trace aux autres modules (figure 6.2), et nous y recherchons, plan par plan, les hits laissés par le passage de la particule. Cette méthode a l'avantage de décorréler la mesure de l'efficacité de hits de la capacité à reconstruire un segment dans

le module étudié.

La “zone de recherche” pour les hits s’étend de part et d’autre de la position extrapolée, à 3σ dans la direction de dérive, σ étant la largeur de la distribution de l’extrapolation correspondante (cf figure 5.19). Cela signifie qu’on peut être conduit à rechercher un hit dans une cellule voisine de celle “touchée” par l’extrapolation (figure 6.1). On mesure ainsi une efficacité globale pour chaque plan sensible : c’est un effet a priori de la qualité du gaz et du champ électrique qui sont les facteurs déterminants de l’efficacité, qualités communes pour tous les fils d’un même plan. Les résultats sont montrés sur la figure 6.3. Les fils réputés “morts” n’ont pas été exclus de ces chiffres.

On constate une efficacité moyenne de l’ordre de 80 % pour les modules “principaux” (1, 2 et 5) c’est-à-dire une perte d’environ 10 % par rapport à NOMAD, attribuable au changement de gaz. Le module 3 notablement moins efficace, avait été repéré “malade” lors de la prise de données : la distribution de gaz a été incriminée, sans qu’une solution ait pu être trouvée.

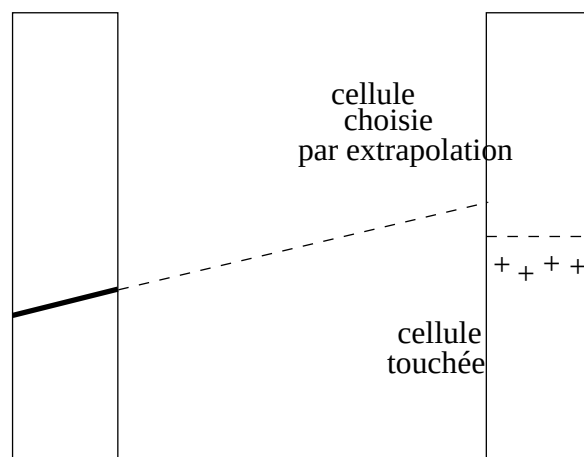


FIG. 6.1 – *Élargissement des recherches de coups dans les cellules adjacentes de la cellule sélectionnée par l’extrapolation .*

6.1.2 Influence de l’efficacité des chambres sur l’efficacité de reconstruction.

Lors de tests effectuées en vue d’étudier sur Monte-Carlo l’efficacité de reconstruction de notre algorithme de segments, nous avons pu montrer (au chapitre 4) une courbe en fonction de l’efficacité de hits. Cependant nous avons utilisé une efficacité uniforme pour tous les plans d’un module. Il semble intéressant d’utiliser les efficacités effectives des chambres au sein du code de la simulation pour se rendre compte de l’influence introduite sur l’efficacité de reconstruction.

Le tableau 6.1 contient les efficacités de reconstruction résultantes pour chaque module. Ces chiffres montrent que l’algorithme des segments actuel doit encore être amélioré compte tenu des inefficacités des chambres et surtout que la reconstruction des traces ne devra pas se contenter de rechercher des associations de segments, reconstruits de façon autonome mais devra aussi chercher à associer des hits “isolés” et des segments. Ce

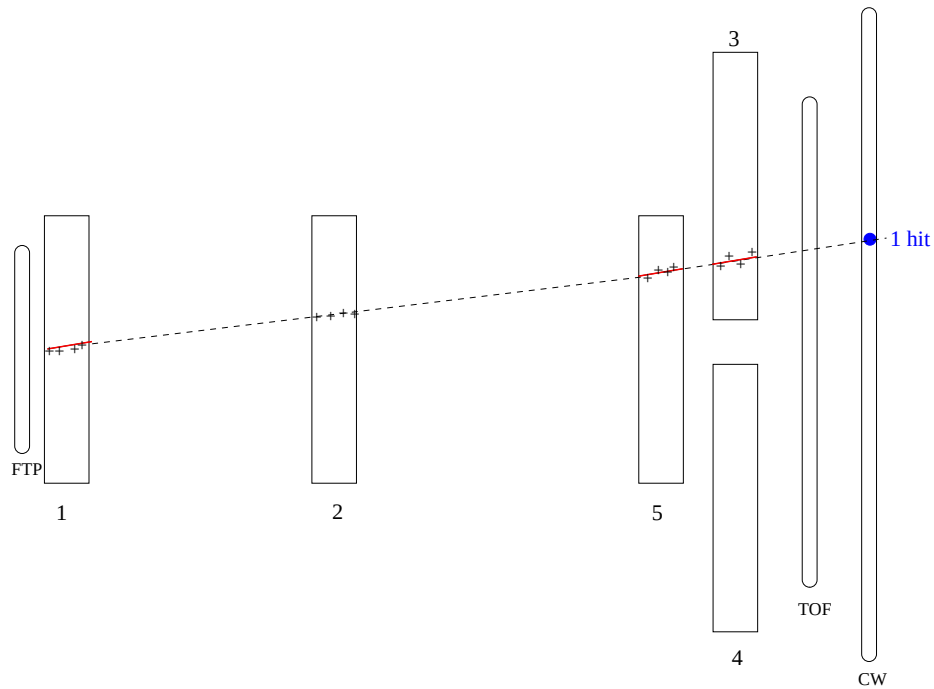


FIG. 6.2 – Recherche de coups appartenant à une trace par extrapolation de segments. Sur la figure (vue de dessus), on note la présence de segments reconstruits dans les modules 1, 3 et 5. Ayant identifié une trace, nous pouvons rechercher des hits appartenant à cette trace dans le module 2.

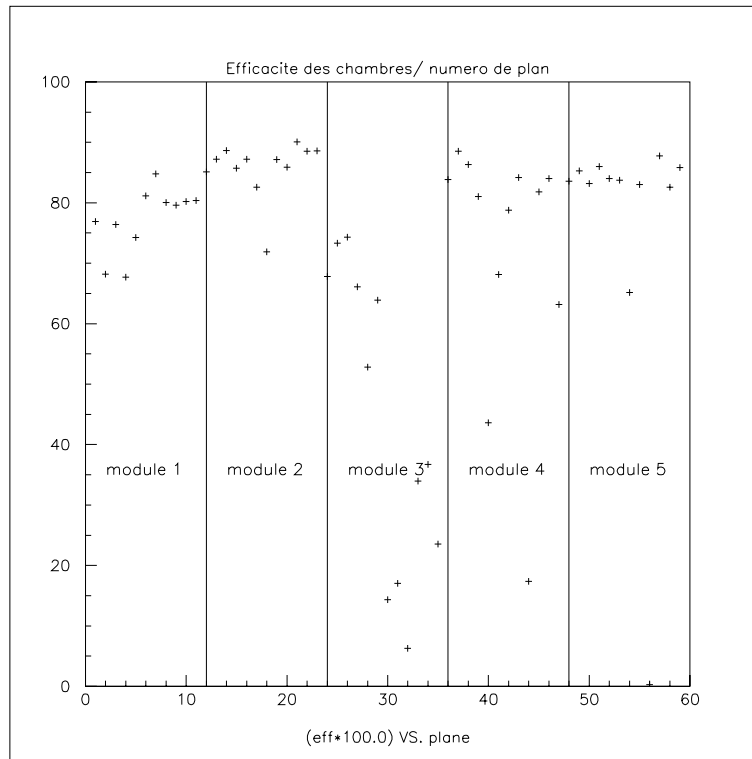


FIG. 6.3 – Efficacité des chambres à dérive pour chaque numéro de plan (résultat obtenu par l'analyse des rayons cosmiques).

# module	Efficacité de reconstruction	segments MC reconstruits	segments MC TOTAUX
1	79,35 %	3145	3963
2	92,13 %	3534	3836
3	4,23 %	21	496
4	54,24 %	326	601
5	69,68 %	2376	3410
TOTAL :		9402	12306

TAB. 6.1 – Efficacité de reconstruction par module pour 4000 μ MC en considérant une efficacité par plan réaliste.

développement est en cours sur la méthode du filtre de Kalman.

6.1.3 Efficacité de reconstruction à partir des données.

Avec les données cosmiques, il semble tout à fait possible d'évaluer l'efficacité de reconstruction des segments connaissant l'endroit où passe une trace pour chaque événement. Donc cette fois-ci, il s'agit de repérer l'existence d'un segment dans un module dans lequel on est sûr qu'une particule est passée. On définit toujours une trace par 2 segments provenant d'autres modules, segments compatibles en position et en orientation.

Le tableau 6.2 nous montre les résultats préliminaires quant à l'efficacité de reconstruction issue de données.

# module	Efficacité de reconstruction	Nombre de traces
1	77,16 %	552 985
2	87,72 %	456 603
3	2,39 %	104 186
4	22,51 %	150 029
5	72,70 %	561 100
TOTAL :		1 824 903

TAB. 6.2 – Efficacité de reconstruction obtenues à partir de l'analyse des données cosmiques.

Pour les modules 1,2 et 5 (modules centraux), les valeurs se comparent assez bien avec l'évaluation précédente sur les données Monte-Carlo.

6.1.4 Conclusion.

Nous avons vu dans ce chapitre, une méthode basée sur l'utilisation des rayons cosmiques, nous permettant d'obtenir une première évaluation de l'efficacité des chambres à dérive, puis dans le même temps une mesure de l'efficacité de reconstruction des segments (en se limitant à la reconstruction au sein d'un module). En l'état, il nous paraît évident

que les performances de ces chambres rendent l'actuel algorithme de reconstruction insuffisant pour une bonne exploitation des données prises par l'expérience. Des développements sont en cours à Dubna pour y remédier.

Chapitre 7

Conclusion

Dans ce manuscrit, nous avons pu constater que la physique hadronique reprend un intérêt de manière paradoxale. En effet elle semble désormais une étape incontournable mieux avancer dans la physique du neutrino qui s'apprête à franchir une nouvelle étape dans son histoire avec le projet d'usine à neutrino.

L'expérience HARP cherche à contribuer dans cette voie en améliorant de manière significative les mesures de sections efficaces lors de collisions proton-cible aux basses énergies. Elle vient compléter les résultats d'autres expériences menées en parallèle et s'efforce de faire mieux que les premières expériences passées. HARP s'y est attelé avec un calendrier très serré en récupérant des appareillages tels que les chambres à dérive, mais également en réalisant la construction de nouveaux éléments. On retient la conception complète de la TPC et du Cherenkov en très peu de temps (environ 14 mois).

Cette thèse s'est concentrée sur la mise en œuvre et l'exploitation des chambres à dérive. Celles-ci ont demandé des modifications "hardware" plus coûteuses que prévu : modification mécanique, matériel électronique nouveau, et changement de gaz (qui s'est révélé moins efficace).

Je me suis impliqué dans la conception d'une simulation détaillée des chambres à dérive dans le cadre de GEANT 4 (en mode de conception Orienté-Objet), nécessaire au développement des algorithmes de reconstruction. Le choix d'un algorithme préliminaire reposant sur la création de segments a permis d'effectuer une calibration des chambres à dérive par un alignement intrinsèque aux chambres ajustant ainsi la position des fils, leur t_0 et amenant à des corrections à la vitesse de dérive par plan. Un alignement global entre les modules a ensuite pu être effectué en vue d'insérer une cohérence géométrique à l'obtention de traces complètes.

La reconstruction par segment a permis d'effectuer une première évaluation des performances de ces chambres, en mesurant les efficacités de chaque plan au niveau des hits. Une étude Monte-Carlo a pu mettre en évidence l'efficacité de l'algorithme, et un recouplement des données et du Monte-Carlo a pu fournir des résultats soulignant la nécessité, malgré des signes encourageants, d'améliorer encore cet algorithme.

On est ainsi limité par un programme de reconstruction encore incomplet, cherchant maintenant à améliorer son association de segments par le filtre de Kalman, d'une part mais également sa recherche visant à collecter le mieux possible des électrons isolés.

Avant d'accéder à l'analyse physique, il reste à compléter le programme de reconstruction en améliorant les reconstructions intrinsèques de chaque sous-détecteur, et la mise en commun de leurs informations respectives dans une base de données bien définie. Une première sélection des événements sous la forme de NTuples est actuellement en cours.

En espérant que cette expérience apportera des satisfactions à la communauté des neutrinos quant à la mise en évidence concrète d'oscillations présente ou future, je tiens à vous remercier d'avoir lu cette thèse jusqu'au bout.

V

Annexes

Annexe A

Corrections pour l'alignement inter-modulaire

A.1 Position du problème

La résolution de ce problème relève du jeu de ce qui est et de ce qui apparaît. En effet nous avons vu dans le chapitre 5 qu'il existe une incohérence entre les extrapolations de 2 segments situés dans 2 modules différents et appartenant à une même trace en raison d'une rotation de module imprévue. Ces écarts ne correspondent pas car ils ne représentent pas concrètement des écarts de décalage linéaire.

Comme on peut le voir sur la figure A.1, nous avons dans le cas d'une extrapolation d'un segment du module 1 vers le module 2, une mesure de la valeur algébrique $d1$ dans le repère intrinsèque du module 1. Par contre, la position du segment du module 2 est définie dans le repère intrinsèque du module 2. Donc si l'on cherche à mesurer cette différence de position en l'état actuel des choses, on n'obtient pas de valeurs valables :

$(d1 - d2)$ n'a de sens que si la rotation est minime. En sélectionnant des traces de faibles incidences, nos observables nous renseignent sur le déplacement Δ recherché.

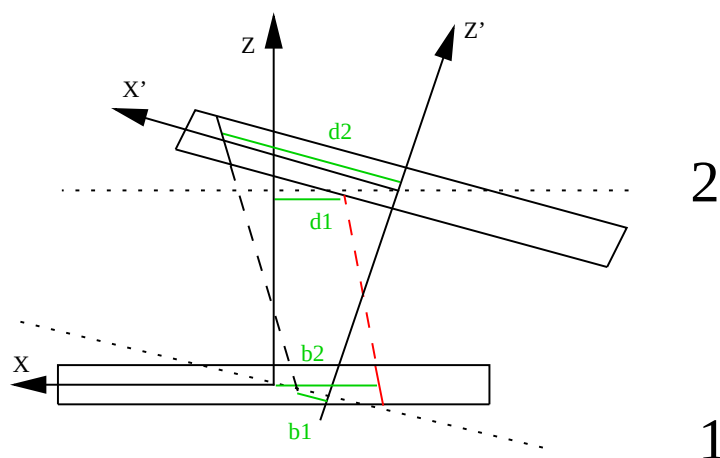


FIG. A.1 – La rotation d'un module introduit des valeurs d'écart non physiques pour les traces à incidence non nulle.

Il en est de même pour le cas inverse où l'on extrapole le segment du module 2 vers le module 1 pour le comparer avec la position du segment du module 1 en faisant $b_1 - b_2$.

A.2 Calcul des paramètres

Les véritables paramètres que l'on recherche (cf figure A.2) sont le décalage Δ et l'angle de rotation θ . L'angle est simple à obtenir car il correspond à la différence d'angle intrinsèque de chaque segment. Ceci est valable si l'on pense à des modules bien parallèles ; il suffit d'orienter les valeurs intrinsèques dans la direction théorique globale Z de HARP.

Effectuer la différence $d_1 - d_2$ revient à faire $x_B - O''C$.

Connaissant les coordonnées de

$O''(z, \Delta)$

$B(z, x_A + z \cdot \tan \alpha)$

On peut obtenir les coordonnées de C en résolvant un système de 2 équations à 2 inconnus :

$$x_c = x_a + z_c \tan \alpha \quad (\text{A.1})$$

$$\frac{z_c - z}{x_c - \Delta} = -\tan \theta \quad (\text{A.2})$$

ce qui nous donne :

$$x_c = \frac{x_a + z \cdot \tan \alpha + \Delta \tan \alpha \tan \theta}{1 + \tan \alpha \tan \theta} \quad (\text{A.3})$$

$$z_c = z - \frac{x_a + z \cdot \tan \alpha - \Delta}{1 + \tan \alpha \tan \theta} \cdot \tan \theta \quad (\text{A.4})$$

Armés de ces éléments, développons l'expression calculée à partir de nos observables :

$$x_B - O''C = x_A + z \cdot \tan \alpha - \frac{x_A - \Delta + z \cdot \tan \alpha}{(1 + \tan \alpha \tan \theta) \cdot \cos \theta} \quad (\text{A.5})$$

Ce qui nous donne pour des traces à incidence nulle :

$$x_B - O''C = x_A - \frac{x_A - \Delta}{\cos \theta} \quad (\text{A.6})$$

On connaît déjà θ par la différence des orientations des segments des 2 modules et on accède à la valeur estimée de l'écart Δ .

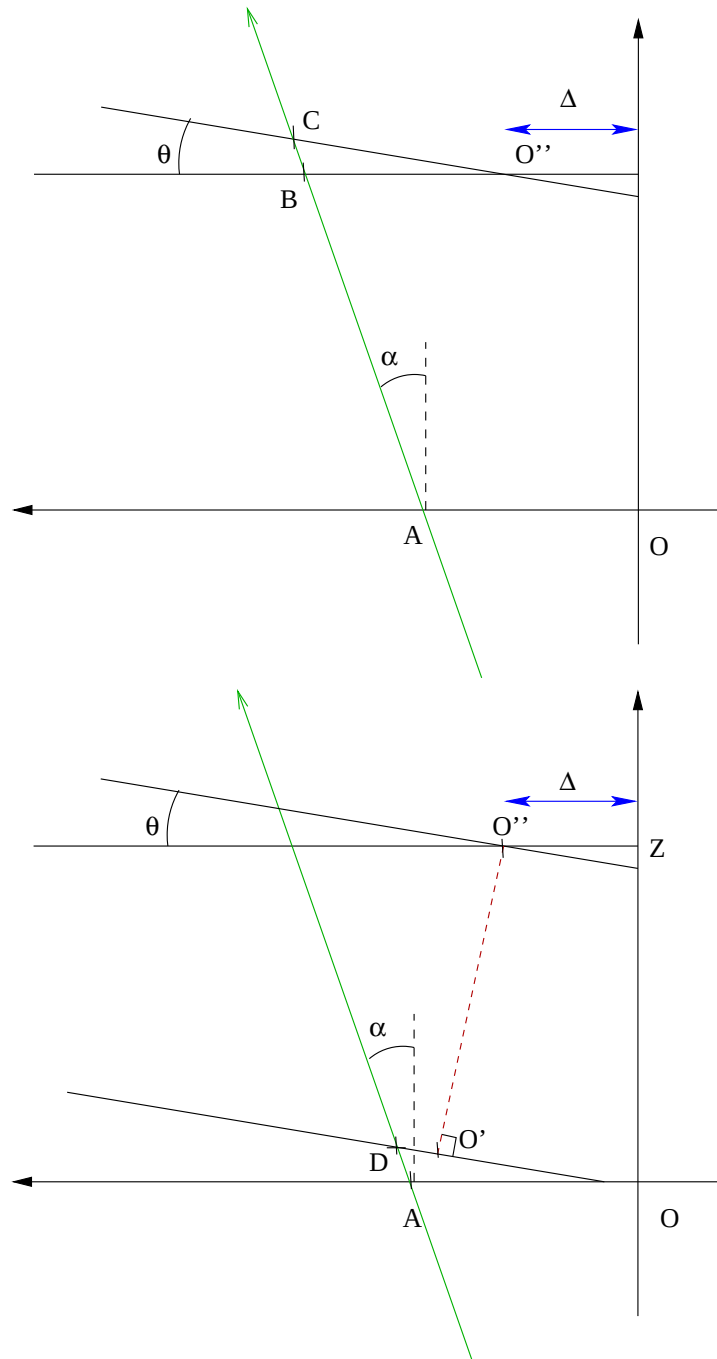


FIG. A.2 – Schémas permettant le calcul de la correction Δ et θ . En haut, une extrapolation du module 1 vers le module 2; en bas, une extrapolation du module 2 vers le module 1

Annexe B

Petit historique sur la production hadronique ($p+A \rightarrow \pi, K, \dots$)

La mesure de sections efficaces de production hadronique n'est en soi pas une nouveauté. Cependant nous allons voir à travers un petit historique que cette question reste toujours d'actualité.

B.1 Activités des années 60-70

À l'époque où l'on cherchait à connaître davantage les propriétés hadroniques, la communauté scientifique visait à effectuer des mesures de sections efficaces de production. Le catalyseur a été la physique du neutrino qui cherchait à progresser dans l'obtention de faisceaux de neutrinos mais également à mieux estimer les flux des neutrinos atmosphériques.

Dès le milieu des années 60 et 70, des chercheurs ont voulu compléter le spectre de connaissances sur ces sections efficaces. Le problème majeur réside dans le nombre important de points de mesure que cela entraîne. Pour avoir une “cartographie” complète des produits résultants de l'interaction proton-cible, il faut une statistique importante pour chaque partie du spectre en énergie des protons incidents et pour chaque type de cibles. Dans la réalité, la contribution statistique aux incertitudes ne furent pas importantes (mis à part aux très basses énergies) ; c'est lorsque l'on examine les erreurs systématiques que les difficultés apparaissent.

Le problème qui figurait comme dénominateur commun à toutes ces expériences était le problème de l'acceptance. On était souvent amené à effectuer les mesures avec un spectromètre disposé sur un bras ayant une faible couverture angulaire (souvent proche de l'axe 0 degré de l'expérience). Dans cette configuration, on peut souligner une expérience étudiant les interactions $p+\text{Be}$ pour des protons de 12,5 GeV/c. menée à Argonne (Illinois) dans [32] avec une couverture allant de 2 à 16 degrés. Les données pour des angles autres, devaient provenir d'extrapolations accompagnées d'incertitudes.

Il en fut de même pour les expériences qui ont suivi : dans [33], on étudia les interactions de protons de 19,2 GeV/c sur des cibles (Be, B_4C , Al, Cu, Pb) pour étudier leur influence, mais également sur une cible d'hydrogène pour les interactions $p+p$. Une extension de cette

dernière expérience a été réalisée dans [34], où l'on a utilisé des protons de 24 GeV/c sur un lot de cibles fines de même nature.

J'ai présenté ci-dessus les expériences majeures dans l'étude p+A. Le tableau B.1 donne un aperçu de toutes les expériences menées durant cette période.

$P_{\text{incidente}}$ (GeV/c)	Auteurs	référence	Cibles
8,65-11,8-18,8-23,1	Dekkers et al.	PR 137B p962 (1965)	p, Be, Al
10, 20, 30	Baker et al.	PRL 7 p101 (1961)	Be, Al
12,3	Marmer et al.	PR 179 p1294 (1964)	Be, Cu
12,3	Marmer et al.	PR D3 p1089 (1971)	Be
12,4	Cho et al.	ANL-HEP 7109 ¹ (1971)	Be
12,5	Lundy et al.	PRL 14 p504 (1966)	Be
12,5	Asbury et al.	PR 178 p2086 (1969)	Be, Al
19,2	Allaby et al.	PL 30B p549 (1969)	d
19,2	Allaby et al.	CERN report 70-12 (1970)	B_4C , Be, Al, Cu, Pb
19,3 - 21,5	Belletini et al.	NP 79 p609 (1966)	Li, Be, C, Al, Cu, Pb
24	Amaldi et al.	NP B39 p39 (1972)	d
25	Cocconi et al.	PRL 5 p19 (1960)	Al
26,6	Amaldi et al.	NC 30 p973 (1963)	C
30	Schwarzshild et al.	PR 129 p854 (1962)	C, Al, Be, Fe
30, 33	Fitch et al.	PR 126 p1849 (1962)	Be, Al
20,1 - 43,1 - 70	Bushnin et al.	PL 29B p48 (1969)	Al
43, 52, 70	Bino et al.	PL 30B p506 (1969)	Al
35, 70	Antipov et al.	PL 34B p164 (1971)	Al

TAB. B.1 – Listes des expériences de 1961 à 1972 ayant étudié l'interaction p+A..

Ensuite l'évolution en énergies des faisceaux de protons (avec l'apparation de nouveaux accélérateurs) a amené l'avènement de nouvelles expériences dans la region des hautes énergies.

B.2 NA20 et SPY/NA56

Ces deux expériences ont étudié les interactions proton-Béryllium a des énergies incidentes beaucoup plus importantes (de l'ordre de 400 GeV/c) apportant ainsi de nouvelles données dans cette gamme d'énergie incidente.

L'expérience NA20 [35] s'est déroulée au CERN avec l'objectif de mesurer la réaction p+Be pour une impulsion incidente de 400 GeV/c. Trois différentes épaisseurs de cibles ont été utilisées : 100, 300 et 500 mm. Les mesures sur les particules secondaires chargées ont été effectuées dans 4 gammes d'énergie (60, 120, 200, et 300 GeV/c) pour 2 impulsions transverses (0 et 0,5 GeV/c). Le principal atout de cette expérience fut sa capacité d'identification des particules grâce à un CEDAR (un compteur Tcherenkov différentiel) de 6 m de long rempli d'hélium pouvant être mis à une pression entre 10 et 14 bar, suivant

la sensibilité à la particule désirée. Cette sensibilité est essentielle sachant la difficulté de la séparation π -K aux hautes énergies.

SPY (pour Secondary Particle Yield) a repris l'idée de NA20 en apportant des améliorations significatives quant à la sensibilité en énergie et quant aux corrections apportées aux flux mesurés en tenant compte des désintégrations des particules étranges. Cependant cette expérience a été réalisée avec la motivation d'avoir une meilleure connaissance du flux de particules secondaires émises par la cible de NOMAD et de CHORUS. SPY a toutefois élargi son étude en envoyant un faisceau de proton de 450 GeV/c sur des cibles de beryllium de différentes épaisseurs. Le spectre mesuré en énergie des particules secondaires a été de 7 à 135 GeV/c accédant ainsi pour la première fois aux impulsions inférieures à 60 GeV/c pour cette énergie incidente. On voit notamment sur la figure B.1, la compatibilité des résultats de [35]. SPY s'est intéressée à la forme de la cible en cherchant à tester la cible utilisée dans NOMAD et CHORUS, qui avait été prévu spécifiquement pour limiter les ré-interactions, mais également pour assurer un refroidissement de la cible plus efficace. Cette cible était constitué de 6 morceaux de 100 mm d'épaisseur séparés par 90 mm d'air pour un diamètre de 3 mm. Ainsi une particule sortant d'une cible avec un angle d'au moins 2 degrés ne rencontrait pas de cibles.

La compréhension de ce type d'interaction intéresse également le secteur de la physique nucléaire qui cherche à comprendre les résultats obtenus lors de l'interaction noyau-noyau. Les expériences présentées ci-dessous rejoignent ainsi les expérience aux basses énergies en effectuant une similitude vers l'énergie par nucléon des noyaux.

B.3 E-802

E-802 a été menée au BNL Tandem-Alternating Gradient Synchrotron (AGS) cherchant à étudier l'interaction à 14,6 GeV/c d'ions ^{16}O et ^{28}Si sur des cibles de Be, Al, Cu, et Au. Le spectromètre utilisé pour les réactions p+A étaient exactement le même que celui pour les interactions ion-A. Il s'agissait d'un bras couvrant une acceptation angulaire de 5 à 58 degrés par rapport à l'axe du faisceau. Ce système avait été utilisé dans [32] dont j'avais parlé au début.

Les résultats de cette expérience sont présentés dans [36]. Dans cette article, il est indiqué qu'un accord satisfaisant avec [32] pour les sections efficaces absolues. Ceci dit, il reste difficile pour une tierce personne de vérifier cette accord. La principale source d'erreur provient de la normalisation de la section efficace (on l'estime entre 10 et 15 %).

B.4 E-910

Cette expérience a recherché en 1996 à combler la méconnaissance (manque de mesures ou de statistique) de données de section efficaces de production hadronique aux basses énergies. L'objectif a été d'étudier la production d'étrangeté mais également

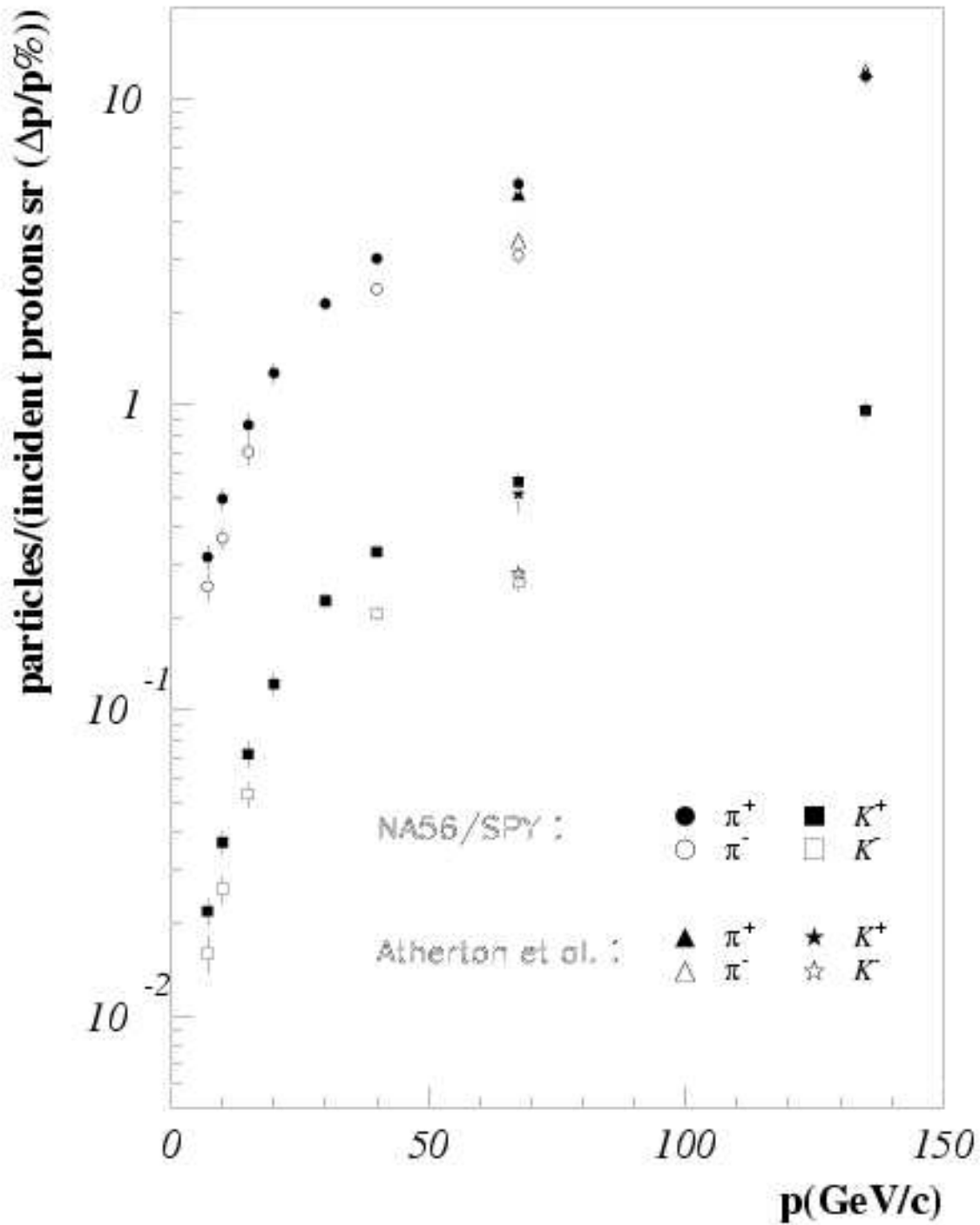


FIG. B.1 – Taux de production des pions et des kaons pour une cible de Be de 100 mm. Les mesures de Atherton et al. ont été rajoutées à titre de comparaison (avec un changement d'échelle tenant compte de la différence d'énergie pour les protons).

les productions de résonnances afin de comprendre l'augmentation importante de particules étranges dans les collisions complexes ion-noyau. Celle-ci travaille dans une gamme d'énergie allant d'environ 12 à 18 GeV/c (au regard de [37]) sur des cibles de Be, Cu et Au. Celle-ci ne s'intéressait qu'aux très basses énergies pour les particules secondaires (< 2 GeV/c). Lorsque l'on regarde l'allure du détecteur de E910 (figure B.2), on peut

reconnaitre une grande part de similitude entre E910 et HARP. cette dernière pourrait paraître comme étant une amélioration de conception pour augmenter l'acceptance.

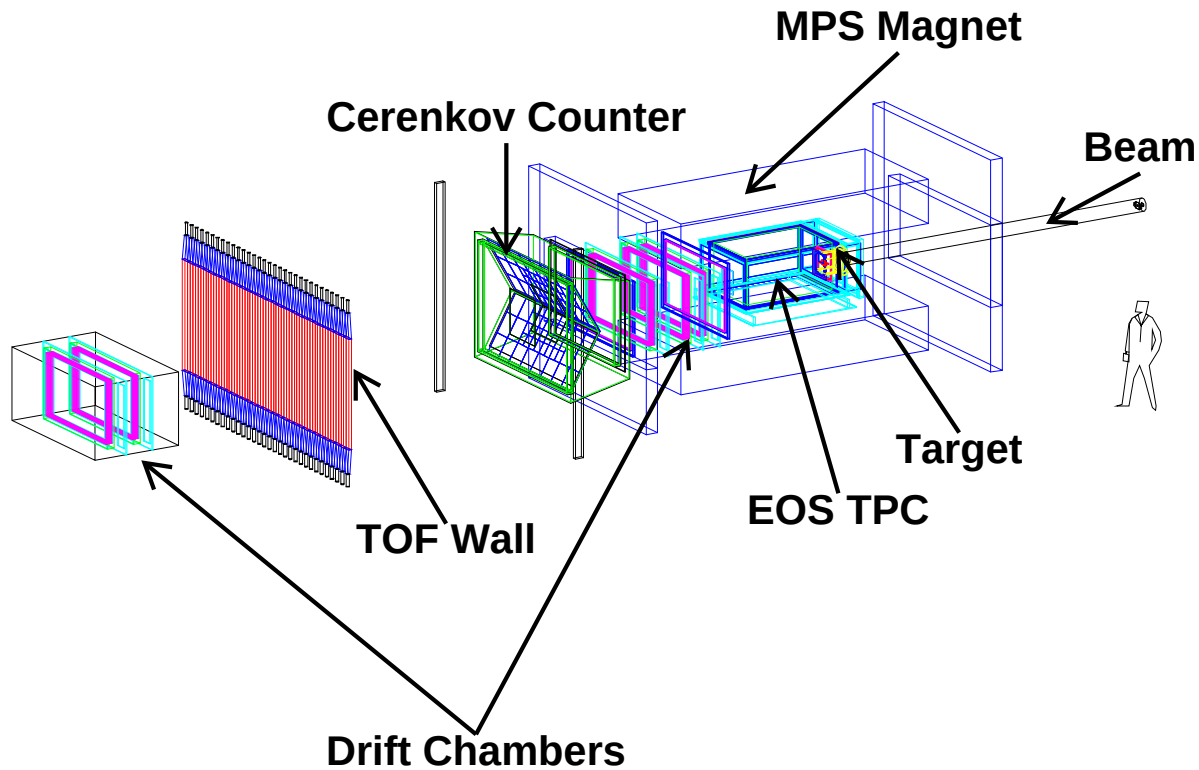


FIG. B.2 – Le détecteur de l'expérience E910.

La différence majeure se situe au niveau de la TPC avec la cible qui se trouve à l'extérieur de la TPC. On peut noter l'usage d'un champ magnétique dipolaire au niveau de la TPC exclusivement, celle-ci étant de forme parallélépipédique. La première conséquence concerne l'acceptance. Il faut reconnaître qu'un grand pas en ce sens avait été fait, amenant une plus grande sensibilité aux traces aux angles larges qui sont nombreuses à basse énergie incidente.

Actuellement ses résultats concernant les mesures de sections efficaces sont relativement limités (voir [37]). Cependant il serait bon de comparer leur résultats avec les premiers résultats de HARP.

Conclusion

Après ce petit tour d'horizon, on a pu constater 3 phases dans l'étude des interactions p+A :

- Dans les années 60-70, des expériences ont effectué les premières mesures de sections efficaces de production hadronique. Le domaine d'étude était les basses énergies (limitation des faisceaux de protons de l'époque). Un grand problème au niveau de l'acceptance est noter (électronique moins précise que maintenant pour l'identification des particules secondaires).
- Dans les années 80, les faisceaux de protons ont augmenté en énergie et de nouvelles

mesures ont été effectuées. A la fin des années 90, on réalise le besoin de mieux connaître les taux de production des cibles dans les expériences neutrinos.

- La physique nucléaire cherche à mieux comprendre les réactions élémentaires $p+A$ pour analyser les résultats ion-A. Ces réactions concernent les basses énergies.

Ainsi au début du 21^e siècle, un retour aux sources a lieu vers les basses énergies et se révèle être indispensable pour améliorer de manière significative des résultats concernant aussi bien la physique du neutrino que la physique nucléaire.

Annexe C

Les membres de la collaboration HARP

M.G. Catanesi, M.T. Muciaccia, E. Radicioni, S. Simone
Università degli Studi e Sezione INFN, Bari, Italy

C. Gößling, A. Grossheim,
Institut für Physik, Universität Dortmund, Germany

S. Bunyatov, G. Chelkov, A. Chukanov, D. Dedovitch, M. Gostkin, A. Guskov,
D. Khartchenko, O. Klimov, S. Kotov, A. Krasnoperov, Z. Kroumchtein, D. Kustov,
D. Naumov, Y. Nefedov, B. Popov, I. Potrap, V. Serdiouk, V. Tereshchenko,
A. Zhemchugov

Joint Institute for Nuclear Research, JINR Dubna, Russia

P. Arce, A. Cervera-Villanueva, F. Dydak, S. Giani, P. Gorbounov, A. Grant,
V. Ivanchenko, J.-C. Legrand, L. Linssen, J. Panman, I. Papadopoulos, J. Pasternak,
E. Tcherniaev, R. Veenhof, C. Wiebusch, J. Wotschack, P. Zucchelli
CERN, Geneva, Switzerland

R. Edgecock, M. Ellis, S. Robbins, F.J.P. Soler
Rutherford Appleton Laboratory, Chilton, Didcot, UK

E. Di Capua
Università degli Studi e Sezione INFN, Ferrara, Italy

A. Blondel, S. Borghi, M. Campanelli, S. Gilardoni, M.C. Morone, G. Prior
Section de Physique, Université de Genève, Switzerland

V. Ableev, C. Cavion, U. Gastaldi, M. Placentino
Laboratori Nazionali di Legnaro dell' INFN, Legnaro, Italy

J.S. Graulich, G. Grégoire
Institut de Physique Nucléaire, UCL, Louvain-la-Neuve, Belgium

M. Bonesini, M. Calvi, A. De Min, F. Ferri, M. Paganoni, F. Paleari
Università degli Studi e Sezione INFN, Milano, Italy

V. Chechin, V. Ermilova, V. Grichine, N. Polukhina, N. Starkov
**P. N. Lebedev Institute of Physics (FIAN), Russian Academy of Sciences,
Moscow, Russia**

S. Gninenko, M. Kirsanov, Yu. Musienko, A. Poljarush, A. Toropin
Institute for Nuclear Research, Moscow, Russia

V. Palladino
Università “Federico II” e Sezione INFN, Napoli, Italy

G. Barr, A. De Santo, C. Pattison, K. Zuber
Nuclear and Astrophysics Laboratory, University of Oxford, UK

M. Baldo Ceolin, G. Barichello, F. Bobisut, D. Gibin, A. Guglielmi, M. Laveder,
M. Mezzetto, M. Vascon
Università degli Studi e Sezione INFN, Padova, Italy

J. Dumarchez, S. Troquereau, F. Vannucci
LPNHE, Universités de Paris VI et VII, Paris, France

V. Ammosov, V. Gapienko, V. Koreshev, A. Semak, Yu. Sviridov, E. Usenko, V. Zaets
Institute for High Energy Physics, Protvino, Russia

U. Dore
Università “La Sapienza” e Sezione INFN Roma I, Roma, Italy

D. Orestano, F. Pastore, A. Tonazzo, L. Tortora
Università degli Studi e Sezione INFN Roma III, Roma, Italy

C. Booth, C. Buttar, P. Hodgson, L. Howlett
Dept. of Physics, University of Sheffield, UK

M. Bogomilov, M. Chizhov, D. Kolev, R. Tsenov
Faculty of Physics, St. Kliment Ohridski University, Sofia, Bulgaria

G. Maneva, S. Piperov, S. Stoykova, P. Temnikov
**Institute for Nuclear Research and Nuclear Energy, Academy of Sciences,
Sofia, Bulgaria**

M. Apollonio, P. Chimenti, G. Giannini, G. Santin
Università degli Studi e Sezione INFN, Trieste, Italy

J. Burguet-Castell, J.J. Gómez-Cadenas, A. Tornero, G. Vidal-Sitjes
Dept. de Física Atómica y Nuclear, Univ. de Valencia, Spain

Annexe D

Nomenclature pour décrire la géométrie

pour la nature des matériaux

- :ELEM pour l'intégration des éléments chimiques élémentaires.

Exemple :

légende : Nom Symbole Z A

:ELEM "Hydrogen" "H" 1. 1.00794

- :MATE pour décrire les matériaux de composition simple.

Exemple :

légende : Z A densité

:MATE "Aluminium" 13. 26.98 2.7

- :MIXT pour les matériaux composés d'éléments de type MATE.

Exemple :

légende : Nom - densité - nombre de composants

nom du composant - fraction ou poids

:MIXT "Honeycomb wall" 0.770 4

"Carbon" 0.722

"Hydrogen" 0.045

"Nitrogen" 0.096

"Oxygen" 0.137

pour les volumes

- :VOLU pour la conception des volumes en spécifiant leur nature et leurs dimensions.

Exemple :

légende : Nom - type - matière - nombre de dimensions - dimensions (ex : X,Y,Z)

:VOLU "MODULE" "BOX" "Air" 3 1642.0 1505.6 178.0

- :POS pour orienter et positionner un volume vis-à-vis d'un volume parent.

Exemple :

légende : Nom - No de copie - Objet parent - rotation - flag - coordonnées (X,Y,Z)

:POS "MODULE" 1 "ExperimentalHall" "RM0" :ONLY 0.0 0.0 2506.7

- :COLOR pour l'affectation de couleur à un volume.

Exemple :

légende : Nom de l'objet - code couleur

:COLOUR "DCGas1" 0. 1. 0.

- :VIS pour la visualisation lors de l'usage d'une interface graphique.

Exemple :

légende : Nom de l'objet - mode de visualisation (ON/OFF)

:VIS "DCGas1" ON

Bibliographie

- [1] P. Astier et al. *Final NOMAD results on $\nu/\mu \rightarrow \nu/\tau$ and $\nu/e \rightarrow \nu/\tau$ oscillations including a new search for ν/τ appearance using hadronic tau decays*, Nucl. Phys. B611 :3–39, 2001.
- [2] <http://nomadinfo.cern.ch>.
- [3] Gordon A. McGregor. First results from the sudbury neutrino observatory. 2002, nucl-ex/0205006.
- [4] Q. R. Ahmad et al. *Measurement of day and night neutrino energy spectra at SNO and constraints on neutrino mixing parameters*, Phys. Rev. Lett. 89 :011302, 2002.
- [5] K. Eguchi et al. First results from kamland : Evidence for reactor anti- neutrino disappearance. *Phys. Rev. Lett.*, 90 :021802, 2003.
- [6] *The Liquid Scintillator Neutrino Detector and LAMPF Neutrino Source*, Nuclear Instruments and Methods A 388 :149–172, 1997.
- [7] C. Athanassopoulos et al. Evidence for $\nu/\mu \rightarrow \nu/e$ neutrino oscillations from lsnd. *Phys. Rev. Lett.*, 81 :1774–1777, 1998.
- [8] Alessandro Strumia. Interpreting the lsnd anomaly : Sterile neutrinos or cpt- violation or...? *Phys. Lett.*, B539 :91–101, 2002.
- [9] MiniBooNE Collaboration. Miniboone experiment proposal. 1997.
- [10] R. Edgecock. *The MICE experiment*, ICFA Beam Dyn. Newslett. 29 :43–57, 2002.
- [11] Y. Obayashi. Atmospheric neutrino results from super-kamiokande experiment. *AIP Conf. Proc.*, 549 :767–771, 2002.
- [12] Vernon D. Barger et al. Oscillation measurements with upgraded conventional neutrino beams. 2001, hep-ph/0103052.
- [13] A. Blondel. Muon polarisation in the neutrino factory. *Nucl. Instrum. Meth.*, A451 :131–137, 2000.
- [14] Ralph Engel, T. K. Gaisser, and Todor Stanev. Pion production in proton collisions with light nuclei : Implications for atmospheric neutrinos. *Phys. Lett.*, B472 :113–118, 2000.
- [15] J. Pasternak, S. Schwenke, and J. Collot. Pion production at low energies. *Nucl. Instrum. Meth.*, A472 :557–560, 2000.
- [16] L. Scothmer P. Gobounov. Mwpc : the layout and geometry, harp note mwpc-2. May 2001.
- [17] Status report of the harp experiment, cern-spsc/2002-013 spsc/m 681. 25 March 2002.

- [18] <http://www.shef.ac.uk/phys/research/hep/harp/targets.html>.
- [19] I. Lehrs. *Nucl. Instrum. methods* 217, 1983, page 43.
- [20] B. Lasiuk et C. A. Whitten Relativistic Heavy Ion Group. Use of krypton 83 as a calibration source for the star tpc yale; <http://star.physics.yale.edu/users/lasiuk/docs.ps/kr-83.ps>.
- [21] Status report to the spsc, presentation de lucie linssen pour le spsc le 31 octobre 2000.
- [22] Harp Collaboration. Harp experiment proposal. 1999.
- [23] P. Annis et al. The chorus scintillating fiber tracker and opto-electronic readout system. *Nucl. Instrum. Meth.*, A412 :19–37, 1998.
- [24] M. Anfreville et al. The drift chambers of the nomad experiment. *Nucl. Instrum. Meth.*, A481 :339–364, 2002.
- [25] <http://www.caen.it/nuclear/product.php?mod=v767>.
- [26] <http://lhcb-comp.web.cern.ch/lhcb-comp/components/html/gaudimain.html>.
- [27] <http://www.objectivity.com>.
- [28] <http://wwwinfo.cern.ch/asd/geant4/geant4.html>.
- [29] <http://root.cern.ch>.
- [30] R. Brun and D. Lienart. Hbook user guide : Cern computer center program library long writeup : Version 4. CERN-Y250.
- [31] B. STROUSTRUP. Le langage c++, 3ème édition, éditions campuspress.
- [32] R.A Lundy et al. Pions and kaons production cross sections for 12.5 gev protons on be, phys. rev lett. 14 (1965) p.504.
- [33] James V. et al. Allaby. *High-energy particle spectra from proton interactions at 19.2 GeV/c*. CERN-70-12.
- [34] T. Eichten et al. *Particle Production in proton interactions in Nuclei at 24 GeV/c*, *Nucl. Phys. B* 44 (1972) p. 333.
- [35] H. W. Atherton et al. Precise measurements of particle production by 400-gev/c protons on beryllium targets. CERN-80-07.
- [36] T. Abbott et al. *Measurement of particle production in proton induced reactions at 14.6-GeV/c*, Phys. Rev. D45 :3906–3920, 1992.
- [37] I. Chemakin et al. *Inclusive soft pion production from 12.3 GeV/c and 17.5 GeV/c protons on Be, Cu and Au.*, Phys. Rev, C65 :024904, 2002.

Remerciements

Je tiens à remercier tous ceux qui ont contribué à cette aventure.

Bien sûr, dans un premier temps, je tiens à remercier toutes les personnes de mon groupe (ceci dit c'est très rapide), Jacques et François (pour Paris) et Boris, Alexei, Yura, et Slava (pour Dubna) avec j'ai eu un très grand plaisir à travailler et à partager de grands moments.

Je tiens à remercier le directeur du LPNHE pour m'avoir accueilli au sein de son laboratoire.

L'équipe informatique du LPNHE qui m'a beaucoup apporté lors de nos nombreuses conversations. Mon savoir-faire en informatique, c'est en partie à vous que je le dois.

Je remercie tous les thésards, stagiaires et autres catégories : Olivier, Pierre, Cyril, Arnaud, Saïd, Guillaume, Muriel, Xavier, Christophe, Marc, Samir pour tous les bons moments que l'on a passés ensemble. Ce qui me restera de ce laboratoire, c'est la grande communication entre tous les jeunes et cette amitié qui s'est développée.

Très grand MERCI à toute ma famille, que je n'ai pas eu le temps de beaucoup voir durant cette période. Mais vos encouragements ont été essentiels.

A Vladimir, qui est mon ami depuis déjà bien longtemps (oui mais on ne va pas s'arrêter là).

Et bien sûr, le plus grand des mercis à Nadia, qui me supporte depuis le début de cette aventure.

.

Résumé

La physique du neutrino a connu de grandes avancées avec la mise en évidence des oscillations (neutrinos solaires et atmosphériques). Pour étudier en détail la phénoménologie de ces neutrinos massifs, de nouvelles machines - les usines à neutrinos seront nécessaires, dont la conception dépend de certains développements ou expériences préalables.

L'expérience HARP se situe dans ce cadre et cherche à mesurer de façon détaillée et précise les sections efficaces de production hadronique (en particulier les pions) dans les interactions proton-noyau pour de faibles impulsions incidentes (2 à 15 GeV/c).

L'ensemble expérimental vise à une acceptance de 4π et cette thèse s'est intéressée au spectromètre vers l'avant et plus précisément aux chambres à dérive qui permettent la reconstruction des trajectoires des particules chargées les plus rapides. La simulation complète de leur fonctionnement, nécessaire aux calculs d'acceptance et aux développements des programmes de reconstruction, est présentée ici dans le cadre général de GEANT 4, adapté par la collaboration. Le problème délicat de l'alignement des chambres est traité ensuite, avant d'extraire une première évaluation de leurs performances : efficacité par plan et efficacité de reconstruction. Les développements sur le programme de reconstruction global permettront ensuite une analyse physique des données de HARP.

Discipline : Physique des particules

Mots-clés : Production Hadronique, HARP, neutrino, Chambres à dérive, simulation, reconstruction.

Adresse du laboratoire :

LPNHE
4, place Jussieu
Tour 33 - Rez de chaussée
75252 PARIS Cedex 05 - FRANCE
